

Towards Quantum Image Generation on Single Qubit using Quantum Information Bottleneck

Yijie Zhu, Vaneet Aggarwal, Debanjan Konar, Yuri Pashkin, Plamen Angelov and Richard Jiang

Abstract—Amidst the rapidly evolving landscape of information technology, the convergence of quantum computing and machine learning—referred to as quantum machine learning—offers promising potential to enhance classical algorithms. However, significant challenges remain in both hardware and software implementation during the Noisy Intermediate-Scale Quantum (NISQ) era, including imperfect qubits, architectural constraints, and high noise levels. In response to these obstacles, this research introduces a novel solution: Quantum Convolutional Variational Autoencoders (QCVAE), designed to operate with only a single qubit. This innovative approach efficiently utilizes a single qubit to manage large-scale data, making it particularly well-suited for quantum computers with limited resources. Simulation results demonstrate the robustness of QCVAE in handling image data, and its deployment on a real quantum computer showcases the model’s practical viability. Additionally, the proposed approach leverages the information bottleneck principle to optimize quantum embeddings, effectively mitigating the impact of prevalent quantum noise. By addressing these core challenges, QCVAE presents a compelling solution for advancing quantum computing applications within the constraints of current NISQ technology.

Impact Statement—This research introduces Quantum Convolutional Variational Autoencoders (QCVAE), a novel architecture bridging quantum computing with core challenges in artificial intelligence. Its primary contribution to AI lies in its ability to efficiently manage large-scale data, such as images, in resource-constrained quantum environments, fundamentally enhances the trade-off between model complexity and computational efficiency in generative AI systems. By leveraging quantum principles, QCVAE enhances classical AI capabilities, offering a novel solution to complex generative tasks while reducing the need for vast computational resources. The model incorporates a quantum information bottleneck, improving robustness and accuracy in noisy environments—a persistent issue in AI systems. This advancement directly contributes to the scalability of quantum generative AI models, enabling higher efficiency, greater adaptability, and better performance in tasks requiring advanced feature extraction and data generation. QCVAE represents a step forward in integrating quantum computing into AI, optimizing resource utilization without compromising the depth and complexity of AI models, and offering new pathways for developing more robust, generalizable, and efficient machine learning systems.

Index Terms—Quantum Machine Learning, Quantum Information Bottleneck, Quantum Neural Networks, Qubits.

Manuscript received October 14, 2024; revised September 1, 2025.

Yijie Zhu, Plamen Angelov and Richard Jiang are with LIRA Center, Lancaster University, Lancaster, LA1 4YW, United Kingdom (e-mail: r.jiang2@lancaster.ac.uk).

Yuri Pashkin is with Quantum Technology Centre, Lancaster University, Lancaster, LA1 4YW, United Kingdom.

Vaneet Aggarwal and Debanjan Konar are with Purdue Quantum Science and Engineering Institute (PQSEI) and School of Industrial Eng., Purdue University, West Lafayette IN 47907, USA.

This work was supported in part by the UK EPSRC under Grant EP/P009727/2, and the Leverhulme Trust under Grant RF-2019-492.

I. INTRODUCTION

The fusion of quantum computing and machine learning offers transformative efficiency gains and adversarial robustness [1, 2], yet confronts three NISQ-era constraints: (1) hardware limitations [3], (2) error susceptibility from incoherent and noise [4], and (3) barren plateaus in variational circuits [5]. While QML achieves high phase classification accuracy [6], these constraints delay quantum utility - reliable outperformance of classical methods [7] - despite hybrid encoding advances [8, 9].

Quantum machine learning now spans semiconductor design [10], biometrics [11, 12], and graph classification [13], achieving higher efficiency over classical counterparts [14]. However, NISQ devices’ gate errors [4] and restricted connectivity [3] create three barriers: qubit scarcity, error propagation, and variational algorithm instabilities [5, 15], limiting practical quantum supremacy [16]. Quantum convolutional networks [17] exemplify these challenges: MNIST processing [18] demands 784 qubits under naive encoding - exceeding IBM’s 127-qubit Eagle QPU and near-term 433-qubit Osprey. Even 1000-qubit prototypes [19] suffer fidelity loss from connectivity constraints, forcing most quantum generative studies onto sub-10-qubit tasks with synthetic data.

Pérez-Salinas et al. [20] established single-qubit neural networks through data-reuploading, effectively emulating multi-layer perceptrons while circumventing parameter bloat. This approach reduces quantum resource demands and introduces feature disentanglement through controlled state evolution on the Bloch sphere. Easom-McCaldin et al. [21] demonstrated that single-qubit systems could replicate CNNs’ hierarchical feature learning by treating sequential data patches as temporal convolution windows, achieving comparable classification accuracy with fewer parameters. The expressivity of single-qubit architectures, validated through quantum circuit capacity analysis [22], reveals their potential for generative adversarial networks and variational autoencoders.

Information bottleneck theory has emerged as a fundamental principle for optimizing representation learning across diverse computational paradigms. Recent advances demonstrate its effectiveness in neuromorphic computing, including surrogate gradient learning for spike-based systems [23], nonlinear information bottleneck applications in spiking neural networks [24], and robust event-based processing [25, 26]. These developments highlight the versatility of information-theoretic optimization principles across both classical and bio-inspired architectures. The extension of information bottleneck theory to quantum computing contexts [27, 28] represents a natu-

ral evolution, where quantum information bottleneck (QIB) principles can address the unique challenges of classical-to-quantum data encoding and noisy quantum environments. Our work contributes to this growing field by demonstrating the first application of QIB to single-qubit generative modeling, bridging information theory with resource-constrained quantum computation.

Building on this landscape, we introduce quantum convolutional variational autoencoders (QCVAE) on a single qubit, offering significant QML advancements. Our QCVAE achieves a universal quantum architecture [3, 20, 29] capable of processing large-scale data efficiently on NISQ devices. By utilizing a single qubit for quantum convolutional layers, we optimize feature extraction and integrate quantum transposed-convolution for data generation in resource-constrained environments. The innovation lies in quantum information bottleneck integration to enhance stability and mitigate quantum noise. Our single-qubit convolution method enables efficient encoding/decoding of large-scale data, including high-resolution images, using only one qubit, mitigating quantum decoherence, limited resources, and noise challenges in the NISQ era.

Experiments on both PennyLane Simulator and IBM Quantum Computer confirm that QCVAE successfully performs image generation using one qubit, demonstrating effectiveness for quantum generative tasks. Our model consistently outperforms classical baselines and existing quantum approaches in ideal and noisy environments, highlighting strong representational capacity of single-qubit architectures. By introducing a quantum autoencoder tailored for limited quantum resources while optimizing generative modeling, our work contributes significant advancement to quantum machine learning and AI-driven data generation.

This paper is organized as follows: Section 1 is the Introduction, Section 2 introduces the model and innovative methods proposed in this work, Section 3 is the experimental setup and results, Section 4 is the discussion, and Section 5 is the conclusion.

II. OUR PROPOSED QUANTUM CONVOLUTIONAL AUTOENCODER

Our novel approach replaces the traditional Euclidean space-based transformation matrices, commonly used in classical deep learning, with unitary matrices tailored for quantum computing. This adaptation significantly reduces the number of parameters in quantum neural networks while preserving expressive power. The workflow of our QCVAE is illustrated in Fig. 1: First, input image data is encoded using a quantum convolution operation based on the single-qubit method. The latent space representation is then extracted, followed by variational sampling for QCVAE. Next, the quantum-transposed convolution operation reconstructs the data during decoding. Finally, the variational quantum information bottleneck and its corresponding loss function optimize the model parameters. We conducted experiments on image denoising and generation, with results demonstrating that our model outperforms existing approaches in both tasks.

A. Encoder Based on Single-Qubit Convolution Method

Various quantum encoding methods exist, including basic, amplitude, and angle encoding. Basic encoding initializes the qubit's quantum state to its binary string equivalent, similar to classical computers. Amplitude encoding transforms data into corresponding superposition state probability magnitudes, while angle encoding employs quantum rotation gates for encoding. Despite their applications, these methods have certain drawbacks such as high qubit costs and challenges in quantum computer implementation. Hence, they may not be the most efficient for minimizing qubit usage. In contrast, single-qubit encoding, introduced in [20, 30], offers a strategy to encode a classical data vector into a unique Hilbert space. This is achieved through a series of single operations acting on each input data dimension, applied to a single qubit. A single qubit's Hilbert space is 2-dimensional, but its state can be repeatedly manipulated via parameterized quantum circuits (PQCs). By iteratively applying trainable unitary operations $U(\theta)$, the qubit evolves through a trajectory in Hilbert space, effectively encoding sequential or hierarchical features [31].

For an N -dimensional input, we partition the data into K blocks. Each block is encoded into the same qubit via distinct $U(\theta_i)$, simulating temporal entanglement across classical data block. Although the two-dimensional Hilbert space of a single quantum bit is limited, it can simulate a high-dimensional feature space through dynamic parameterization. For example, data with an input dimension of N is re-uploaded k times, which is equivalent to operating in a 2^k -dimensional virtual space [32]. This mimics multi-qubit entanglement in resource-constrained settings. We can use this encoding method to repeatedly reuse a single quantum bit, which can save quantum resources and improve the overall efficiency of the model compared to previous work. Based on the original design, we modified it to make the single qubit method more suitable for our convolution method. We expanded each unit to include three quantum rotation gates to fit the 3×3 convolution kernel and the unitary operation of single-qubit encoding can be expressed as

$$U = e^{-\alpha} R_Z(\beta) R_Y(\gamma) R_Z(\delta), \quad (1)$$

where $\alpha \in \mathbb{R}$ is the global phase factor, Euler angles $\beta, \gamma, \delta \in \mathbb{R}$ that define the extent of each rotation (R) around the Z , Y and Z axes, respectively. Within this method of encoding, these Euler angles are parameterized further and defined as

$$\beta = \theta_i + x_i \cdot \phi_i, \quad (2)$$

$$\gamma = \theta_{i+1} + x_{i+1} \cdot \phi_{i+1}, \quad (3)$$

$$\delta = \theta_{i+2} + x_{i+2} \cdot \phi_{i+2}, \quad (4)$$

where θ_i and ϕ_i are trainable weight parameters assigned to x_i , the value of the input vector x at dimension i . Therefore, the extent of rotation β, γ, δ is with respect to the weighted value of the input. By combining three rotation gates into a unitary operation, we obtain

$$U(\vec{\omega}) = e^{i\vec{\omega} \cdot \vec{\sigma}}, \quad (5)$$

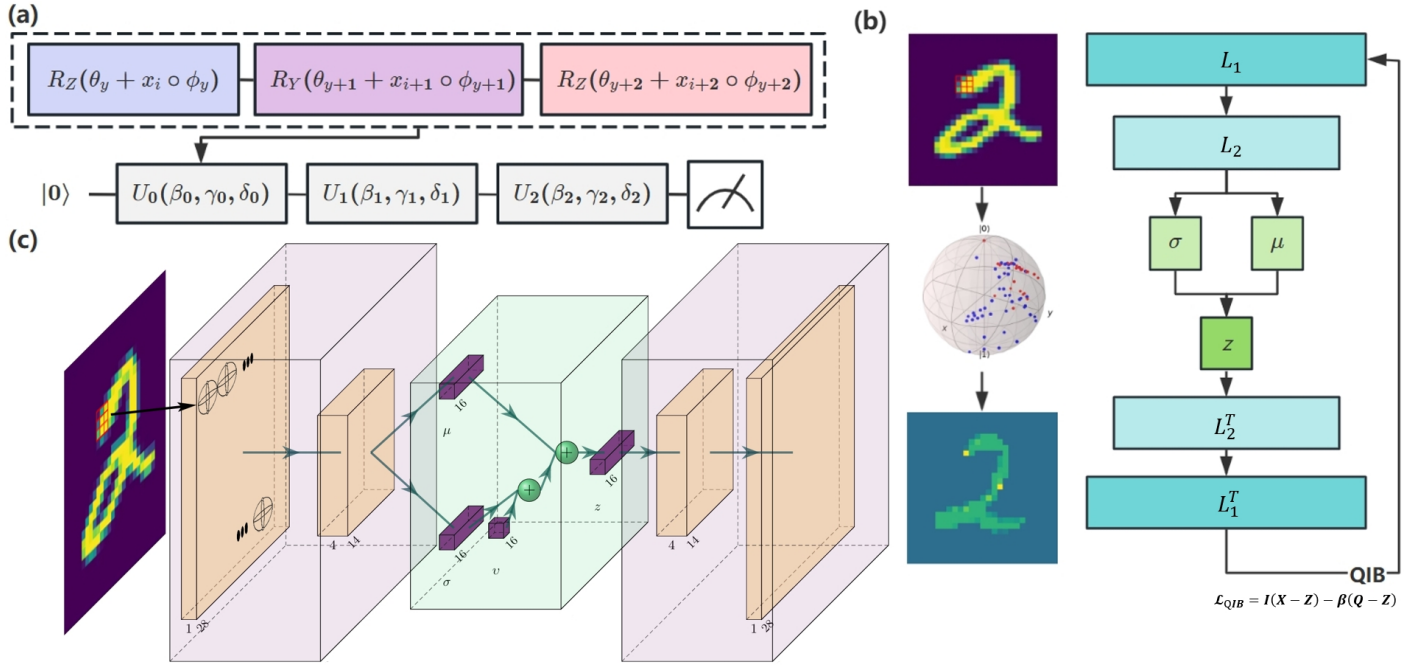


Fig. 1: Architecture and operational principles of the Quantum Convolutional Variational Autoencoder (QCVAE) for image processing. The input image undergoes a workflow: (1) quantum encoding via parameterized unitary operations $U(\theta)$, (2) encoder and decoder process based on quantum convolution, and (3) QIB-based gradient optimization. The quantum convolutional kernel utilizes a 3×3 grid (9 units), where each unit is encoded onto qubits through quantum rotation gates R_y, R_z . Triplets of rotation gates form a trainable unitary operation $U(\theta)$, serving as the fundamental quantum computing unit. The encoder ($L_1 \rightarrow L_n$) compresses image features into a latent representation z , optimized by the Quantum Information Bottleneck (QIB) loss L_{QIB} (defined in Eq. 15) to maximize task-relevant information while discarding redundancies. The decoder ($L_1^T \rightarrow L_n^T$) reconstructs images via inverse unitary operations (quantum transposed-convolutional layers L^T). Bloch sphere trajectories visualize qubit states in Hilbert space, with colors mapping to feature distributions of distinct inputs. Compared to classical VAEs, QCVAE enhances parameter efficiency and robustness in tasks like denoising and super-resolution by dynamically tuning $U(\theta)$ to adapt to complex image patterns. (c) is the model structure presented in the MNIST dataset.

where $\vec{\omega} = (\omega(\beta), \omega(\gamma), \omega(\delta))$, and, respectively:

$$\omega(\beta) = c \left(\sqrt{1 - \cos^2 c} \right)^{-1} \sin \left(\frac{\gamma - \delta}{2} \right) \sin \left(\frac{\beta}{2} \right), \quad (6)$$

$$\omega(\gamma) = c \left(\sqrt{1 - \cos^2 c} \right)^{-1} \cos \left(\frac{\gamma - \delta}{2} \right) \sin \left(\frac{\beta}{2} \right), \quad (7)$$

$$\omega(\delta) = c \left(\sqrt{1 - \cos^2 c} \right)^{-1} \sin \left(\frac{\gamma + \delta}{2} \right) \cos \left(\frac{\beta}{2} \right), \quad (8)$$

where $\cos c = \cos \left(\frac{\gamma + \delta}{2} \right) \cos \left(\frac{\beta}{2} \right)$. The single-qubit encoding method can be employed to encode up to three input dimensions per unitary operation. The input vector is thus cycled through in order to encode three-dimensional values until the entire input has been encoded. This method can be flexibly implemented on quantum circuits that process data of different structures and can increase the data capacity per qubit.

Quantum circuits form the core of quantum computing, enabling the realization of diverse functions through the synergy of qubits and quantum gates. In this investigation, our focus lies on the single-qubit approach, restricting the qubit number to one. While theoretically scalable to multiple qubits,

entanglement is established between them using the CNOT gate. The quantum convolution layer processes data embedded in the quantum circuit, treating combinations of R_X, R_Y , and R_Z gates as unitary operations as shown in Fig 1. The number of unitary operations on a single qubit is dictated by the input data size (quantum convolution kernel size), with each operation encoding three inputs through six parameters. For example, a 3×3 data size necessitates three unitary operations on a single qubit. This single-qubit method is scalable according to the input data size and applicable across various applications. The quantum system can be described by the following wave function:

$$|\psi(x)\rangle = \prod_{i=1}^L R_Z(\theta_{i_1} + x \cdot \phi_{i_1}) \cdot R_Y(\theta_{i_2} + x \cdot \phi_{i_2}) \cdot R_Z(\theta_{i_3} + x \cdot \phi_{i_3}) |0\rangle \quad (9)$$

For classic models, there is Universal Approximation Theorem (UAT) [33] to support its approximation capabilities. Similarly, for quantum models, UAT can be used to demonstrate approximation capabilities. According to [34], a quantum analogue can be constructed on the basis of UAT. Let f and g be a pair of functions, with $f \in \mathbb{R}^m \rightarrow [0, 1]$ and

$\varrho \in \mathbb{R}^m \rightarrow [0, 2\pi)$, there is:

$$\left| f(\vec{x}) e^{i\varrho(\vec{x})} - \left\langle 1 \prod_{i=1}^N U\left(\vec{x}, \vec{\theta}_i, \vec{\phi}_i\right) | 0 \right\rangle \right| < \epsilon \quad (10)$$

Where $\epsilon > 0$. Based on this quantum UAT, it can be considered that the Single-Qubit method is able to approximate the functions.

B. Decoder Based on Quantum Transpose Convolution

Within our QCAE, the decoding phase resembles that of classical CAE, and the quantum-transposed convolution is implemented through the single qubit method. The setting for the quantum circuit here is the same as that for the encoder. Following multiple convolution operations on the image, the feature map's size continues to decrease. For the autoencoder, it becomes essential to restore the image to its original size before further operations. While numerous methods exist for this purpose, early approaches are predominantly designed based on prior knowledge, often yielding suboptimal results in various scenarios. Transposed convolution, employed as an interpolation method in neural networks, distinguishes itself from conventional techniques. Unlike pre-defined interpolation methods, transposed convolution incorporates learnable parameters akin to standard convolution, enabling the acquisition of an optimal upsampling method through network learning.

The standard convolution operation establishes a many-to-one mapping relationship. In contrast, the transposed convolution aims to achieve a reverse operation. Considering the convolution equation $Y = XW$, we can construct a matrix V for W in a way that convolution is equivalent to matrix multiplication $Y' = VX'$, where Y' and X' represent the vectors of Y and X , respectively. Transposed convolution is thus equivalent to $Y' = V^T X'$. If the convolution transforms the input from dimensions (h, w) to (h', w') , the transposed convolution with the same parameters transforms it back from (h', w') to (h, w) . To provide a concrete example, when the padding is p and the stride is s , the first step involves inserting $s - 1$ rows or columns between existing ones. Subsequently, the input is padded with $k - p - 1$ (where k is the convolution kernel window). The kernel matrix is then flipped, and a conventional convolution operation is performed. If the input height (or width) is n , the kernel size is k , the padding is p , and the stride is s , the matrix size of the transposed convolution is given by $n' = sn + k - 2p - s$.

C. Quantum Information Bottleneck

In our quantum autoencoder, there exists a phase involving the transformation of classical data into quantum data via quantum circuits, akin to the compression of classical data. During this procedure, we applied the quantum extension of the information bottleneck theory, leveraging QIB to optimize our data embedding for the most effective results. Quantum analogs of algorithms associated with information bottlenecks have been explored [27, 28]. In many instances, particularly in real-world scenarios, the data at our disposal is classical

which is grounded in Euclidean space. To encode this data into the Hilbert space fundamental to quantum computing, specific transformations are necessary. This process inevitably introduces information loss. However, for representation learning, not all information is equally crucial. This implies that we can reframe the conversion of classical data into quantum data as a mechanism for eliminating redundant information, facilitated by the quantum information bottleneck.

While quantum information bottleneck principles have been explored in various quantum machine learning contexts, our implementation presents several key distinctions. Unlike previous QIB work that primarily focused on general quantum state compression or classification tasks, our approach specifically addresses the classical-to-quantum encoding challenge in generative models. Most significantly, we are the first to integrate QIB with single-qubit convolutional variational autoencoders, creating a unique synergy between resource-efficient quantum computing and information-theoretic optimization.

In classical contexts, the information bottleneck (IB) principle serves the purpose of identifying the optimal compressed representation Z of the input X , retaining all pertinent information essential for modeling the target Y . The degree of compression is measured through the classical mutual information $I(X : Z)$, with $I(Y : Z)$ quantifying the residual information between Y and Z . Achieving superior outcomes necessitates a substantial value for $I(Y : Z)$, while effective compression requires minimizing $I(X : Z)$. For general α and subsystems A and B , the α -Rényi mutual information [35] is defined as:

$$I_\alpha(A : B) = \frac{\alpha}{\alpha - 1} \log_2 \text{Tr} \sqrt[\alpha]{\text{Tr}_A \left(\rho_A^{(1-\alpha)/2} \rho_{AB}^\alpha \rho_A^{(1-\alpha)/2} \right)}, \quad (11)$$

For $\alpha \rightarrow 1$, one recovers the quantum mutual information $I_1(A : B) \equiv I(A : B)$.

There is some connection between the information bottleneck and Kullback–Leibler (KL) divergence here. We can regard KL divergence as an information bottleneck theory as an example of Loss function. Looking back at VAE, we want to maximize the probability value of generating true data and minimize the KL divergence of the true and estimated posterior distributions. KL divergence can be thought of as $I(Y : Z)$ in the information bottleneck. The formulas for the variational autoencoder and KL divergence are shown in the supplementary. The information bottleneck principle strikes a balance between accuracy and compression through the minimization of the Lagrangian $L_{IB} = I(X : Z) - \beta I(Y : Z)$, where β serves as a parameter enabling exploration of various regimes, allowing a preference for either accuracy or compression. For $\beta = 0$, minimizing L_{IB} prioritizes optimal compression with no concern for output results. Conversely, as $\beta \rightarrow \infty$, minimizing L_{IB} aims for optimal results without compression. In our case, we apply the IB principle to the distinct problem of determining the optimal embedding.

Our focus lies in the scenario where X and Z represent the classical input and representation spaces, and since our work is unsupervised learning, Q signifies the expression of output data for the quantum circuit in the decoder. Based on the

above-mentioned classical information bottleneck theory, we give the quantum information bottleneck theory. Subsequently, we define the quantum IB Lagrangian as follows:

$$L_{QIB} = I(X : Z) - \beta I(Q : Z), \quad (12)$$

where $I(A : B) = I_{\alpha \rightarrow 1}(A : B)$ is the quantum mutual information. Here, $I(X : Z)$ and $I(Q : Z)$ are quantum mutual information. Both $I(X : Z)$ and $I(Q : Z)$ can be formulated using Holevo's accessible information [36]. Our QIB implementation offers specific advantages for generative modeling: (1) it reframes inevitable information loss during classical-to-quantum encoding as selective preservation of task-relevant features, (2) it optimizes latent representations specifically for image reconstruction and generation tasks, and (3) it provides enhanced robustness against quantum noise by focusing learning on information-preserving features rather than noise-sensitive redundancies. Effective generalization occurs when $I(X : Z)$ is minimized, whereas achieving low error is feasible when $I(Q : Z)$ is maximized. Assuming $I_{\alpha}(X : Z)$ delineates a spectrum of generalization bounds corresponding to various loss functions, minimizing the generalization error aligns with minimizing $I(X : Z)$. The optimal embedding for a given β value is then derived as the minimum of $\rho(x)L_{IB}$. The explicit form of the quantum IB Lagrangian is obtained from the definition:

$$\begin{aligned} L_{QIB} = & (1 - \beta) S \left[\sum_x P(x) \rho(x) \right] \\ & - \sum_x P(x) S[\rho(x)] + \beta \sum_q P(q) S[\rho_q] \\ & + \sum_x \tilde{\lambda}_x \text{Tr}[\rho(x)] + \eta, \end{aligned} \quad (13)$$

where $S[\rho(x)]$ is the von Neumann entropy of quantum state ρ , the $\tilde{\lambda}_x$ are Lagrange multipliers to force correct normalization, and η contains all the terms that are independent of the embedding. The optimal embedding corresponds to a minimum of L_{IB} , which satisfies $\partial L_{IB} / \partial \rho(z) = 0$. When $\rho(x) = \rho(q)$, $p = 1$, otherwise $p = 0$. By explicit computation, we find that the above condition defines a recursive equation for the optimal embedding:

$$\tilde{\lambda}_z \rho(z) = e^{(1-\beta)\log p + \beta \sum_q P(q|z)\log \rho_q} \rho(q), \quad (14)$$

where $\rho = \sum_q P(q) \rho(q)$ and λ_z is directly related to $\tilde{\lambda}_z$ and is needed to enforce normalization. Alternatively, by restricting to pure state embeddings $\rho(x) = |\psi(x)\rangle \langle \psi(x)|$, we get

$$\tilde{\lambda}_z |\psi(x)\rangle = e^{(1-\beta)\log p + \beta \sum_q P(q|z)\log \rho_q} |\psi(x)\rangle \quad (15)$$

For $\beta = 0$, we get a constant embedding, while for large β , the optimal embedding for a given x is iteratively obtained from one of the eigenvectors of $\sum_q P(q|z)\log \rho_q$ with the largest eigenvalue, or a mixture of them. Since the state in single-qubit circuit in Eq.(9) is pure, $S(\rho(x)) = 0$ in Eq.(12). In order to train the embedding, we approximate the average

over the distribution $P(q, x)$, with empirical averages over the elements of the training set τ . According to Eq.(12), we have:

$$L_{QIB}^{\tau} = (1 - \beta) s(P_t) + \beta \sum_q \frac{T_q}{T} s(P_q) \quad (16)$$

where constant terms have been neglected, and by explicit computation, the purities are

$$P_t = \frac{T + 2 \sum_{x < q}^{\tau} F(\rho(x), \rho(q))^2}{T^2} \quad (17)$$

$$P_q = \frac{T_q + 2 \sum_{x < q}^{\tau_q} F(\rho(x), \rho(q))^2}{T_q^2} \quad (18)$$

where the ordering $x < q$ refers to the index of the inputs in the training set.

III. EXPERIMENTS

We opt for the MNIST dataset [18], the MNIST-Fashion dataset [37] the CIFAR-100 dataset [38], and the CelebA dataset [39] as our experimental datasets. MNIST comprises 60,000 training and 10,000 testing examples of handwritten digits (0-9). Fashion-MNIST contains 60,000 training and 10,000 testing examples of 28×28 grayscale clothing images across 10 categories. CIFAR-100 features 100 classes with 600 color images (32×32) each: 500 for training, 100 for testing, with fine and coarse labels. CelebA contains 202,599 celebrity face images representing 10,177 identities.

We established various tasks including image denoising and variational image generation models. Alongside classic models DAE, VAE and CAE, we compared classical methods with QIB to validate quantum algorithm advantages. Three quantum models were selected for comparison: Quantum Autoencoder (QAE) [40], Quantum Generative Adversarial Network (QGAN) [41], and Quantum Circuit AutoEncoder (varQCAE) [42]. Experiments utilized Fréchet Inception Distance (FID) [43], Structural Similarity Index Measure (SSIM) [44], PSNR [45] and Learned Perceptual Image Patch Similarity (LPIPS) [46] as benchmarks. Due to current quantum resource limitations, we cannot use annealing algorithms to optimize β . Therefore, we extracted small datasets from each dataset, testing with initial value 0 and step value 0.5, obtaining β values empirically using FID as the primary indicator (specific results in supplementary). Model parameters: learning rate 0.01, epochs 50, batch size 32, train/test split 8:2. Model sizes were adjusted according to data dimensions. For MNIST/Fashion datasets: one quantum convolution layer with 3×3 kernel, 4 output channels, 2 strides, 16 latent dimensions, and 18 single-qubit weights. QIB's $\beta = 1.5$. For CIFAR-100: structure $(3, 32, 32) \rightarrow (8, 16, 16) \rightarrow (16, 8, 8)$, 64 latent dimensions, $\beta = 2.0$. For CelebA: structure $(3, 128, 128) \rightarrow (8, 64, 64) \rightarrow (16, 32, 32) \rightarrow (32, 16, 16) \rightarrow (64, 8, 8)$, 128 latent dimensions, $\beta = 2.5$. Detailed structural figures in supplementary.

Our method follows a structured process: First, normalize datasets. Image data undergoes single-qubit convolution embedding, encoding input into latent space. Depending on

TABLE I: Denoising performance in ideal and depolarizing noisy quantum environments. Bold: best, underlined: second best.

Category	Method	Ideal Environment				Noisy Environment			
		FID	SSIM	PSNR	LPIPS	FID	SSIM	PSNR	LPIPS
MNIST									
Classical	DAE	25.2±1.8	0.85±0.02	19.6±0.5	0.15±0.01	-	-	-	-
	DCAE	<u>23.5±0.7</u>	<u>0.89±0.01</u>	<u>20.3±0.4</u>	0.12±0.01	-	-	-	-
Quantum	QAE	25.1±1.8	0.85±0.02	19.7±0.5	0.14±0.01	25.1±1.9	0.83±0.03	18.2±0.5	0.15±0.01
	varQCAE	25.5±1.4	0.87±0.01	19.5±0.4	0.13±0.01	25.6±1.7	0.86±0.02	18.5±0.4	0.14±0.01
	Our model	23.7±0.9	0.87±0.02	20.2±0.4	<u>0.10±0.01</u>	<u>23.9±1.0</u>	0.85±0.02	<u>19.1±0.3</u>	<u>0.12±0.01</u>
	Our model-QIB	23.1±0.8	0.90±0.01	20.4±0.3	0.09±0.01	23.6±1.0	0.88±0.02	19.4±0.3	0.11±0.01
Fashion-MNIST									
Classical	DAE	61.4±2.1	0.66±0.02	18.1±0.4	0.20±0.02	-	-	-	-
	DCAE	<u>60.1±1.2</u>	<u>0.70±0.01</u>	18.8±0.3	0.18±0.01	-	-	-	-
Quantum	QAE	61.1±1.9	0.68±0.02	18.3±0.4	0.19±0.01	61.8±1.0	0.65±0.04	17.3±0.4	0.20±0.02
	varQCAE	61.8±1.8	0.69±0.01	18.2±0.3	0.19±0.01	61.4±0.8	0.67±0.02	17.1±0.3	0.19±0.01
	Our model	60.6±1.0	0.69±0.02	<u>19.0±0.3</u>	<u>0.17±0.01</u>	<u>61.2±1.2</u>	0.69±0.02	<u>17.8±0.3</u>	<u>0.18±0.01</u>
	Our model-QIB	59.9±0.8	0.71±0.01	19.7±0.2	0.15±0.01	60.6±0.9	0.70±0.02	18.1±0.3	0.17±0.01

TABLE II: Image generation performance in ideal and depolarizing noisy quantum environments. Bold: best, underlined: second best.

Category	Method	Ideal Environment				Noisy Environment			
		FID	SSIM	PSNR	LPIPS	FID	SSIM	PSNR	LPIPS
MNIST									
Classical	VAE	23.8±0.6	0.90±0.01	20.6±0.4	0.12±0.01	-	-	-	-
	CVAE	24.7±0.9	<u>0.92±0.01</u>	20.2±0.5	0.13±0.01	-	-	-	-
Quantum	QGAN	26.4±1.3	0.93±0.01	19.9±0.6	0.14±0.01	26.6±1.9	0.90±0.03	18.6±0.6	0.15±0.02
	Our model	<u>23.1±0.5</u>	0.93±0.01	<u>21.2±0.3</u>	<u>0.11±0.01</u>	<u>23.3±0.7</u>	0.91±0.03	<u>20.3±0.4</u>	<u>0.12±0.01</u>
	Our model-QIB	22.8±0.6	0.93±0.01	21.6±0.3	0.10±0.01	22.9±0.7	0.91±0.02	20.7±0.3	0.11±0.01
Fashion-MNIST									
Classical	VAE	58.7±1.2	0.70±0.01	19.2±0.3	0.18±0.01	-	-	-	-
	CVAE	60.8±1.4	<u>0.72±0.01</u>	18.9±0.4	0.19±0.02	-	-	-	-
Quantum	QGAN	64.6±1.3	0.67±0.02	18.4±0.5	0.21±0.02	67.7±1.8	0.64±0.04	17.8±0.5	0.22±0.03
	Our model	<u>60.4±1.0</u>	0.73±0.01	<u>19.9±0.3</u>	<u>0.17±0.01</u>	<u>60.5±1.6</u>	0.71±0.02	<u>18.9±0.4</u>	<u>0.18±0.02</u>
	Our model-QIB	59.8±0.6	0.73±0.01	20.2±0.3	0.16±0.01	60.2±1.1	0.71±0.02	19.3±0.3	0.17±0.01

quantum convolution kernel size, appropriate unitary operations are applied to qubits. During measurement, expected values of target states $|0\rangle$ and $|1\rangle$ are obtained. In decoder phase, quantum transposed convolution operations reconstruct or generate data. Training involves assessing QIB impact and updating quantum rotation gate parameters. After iterations, optimal parameters are derived and evaluated using test data.

Experiments used PennyLane [47] with default "default.qubit" simulator and PyTorch [48]. The quantum computer was IBM's ibm_manila node [49] featuring 5 qubits, Quantum Volume 32, 2800 circuit operations/second capability, and Falcon r5.11L processor with linear layout. Single-qubit gate (SX, X, RZ) average error: $3.093\text{e-}4$. Two-qubit gate (CNOT) average error: $8.256\text{e-}3$. Average readout error: $2.594\text{e-}2$. Coherence times: average $T_1 = 155.49\mu\text{s}$, $T_2 = 88.69\mu\text{s}$.

A. Experiments in the Simulated Quantum Environment

Ideal simulated environment. Tables I to IV present results for each model in an ideal quantum environment, with classic AE and VAE serving as baselines. Our QCVAE significantly

outperforms other quantum autoencoders. Previous quantum autoencoders encountered challenges such as convergence difficulties and local optima during training, attributed to substantial increases in quantum gates and qubits. Previous quantum models had not handled data as large as images in their experiments. While encoding was feasible for small datasets, efficiency declined for relatively large-sized data like images. The experimental results clearly demonstrate QIB's optimization of data embedding. Compared to the non-optimized model, QIB results improved, particularly enhancing generated image stability.

For denoising tasks, our model exhibits notable advantages over classic Autoencoder (AE), as shown in the tables in the denoising part. Being a quantum model, quantum circuits inherently possess unique noise-handling mechanisms. Training challenges observed in previous quantum autoencoders persist. However, unlike image generation, the denoising task shows less substantial FID differences. We attribute this to quantum mechanisms' advantages in mitigating noise. While QIB aids model optimization, its impact is less pronounced than in generation tasks, possibly because denoising relies less on

TABLE III: Denoising performance on CIFAR-100 and CelebA datasets (Ideal and Noisy). Bold: best, underlined: second best.

Category	Method	Ideal Environment				Noisy Environment			
		FID	SSIM	PSNR	LPIPS	FID	SSIM	PSNR	LPIPS
CIFAR-100									
Classical	DAE	167.7±11.6	0.42±0.06	17.9±0.3	0.25±0.03	-	-	-	-
	DCAE	163.7±9.7	0.45±0.03	18.3±0.3	0.23±0.02	-	-	-	-
Quantum	QAE	166.4±11.1	0.43±0.03	18.1±0.3	0.24±0.02	165.2±11.8	0.39±0.05	17.1±0.3	0.25±0.03
	varQCAE	166.3±10.4	<u>0.45±0.02</u>	<u>18.4±0.3</u>	0.22±0.02	164.5±11.4	0.41±0.04	17.3±0.3	0.24±0.02
	Our model	<u>163.4±9.4</u>	<u>0.45±0.02</u>	18.3±0.5	<u>0.20±0.02</u>	<u>164.1±9.5</u>	0.42±0.04	<u>17.5±0.3</u>	<u>0.23±0.02</u>
	Our model-QIB	161.6±8.6	0.46±0.02	18.8±0.4	0.18±0.01	162.7±8.5	0.43±0.03	17.8±0.3	0.21±0.01
CELEBA									
Classical	DAE	87.5±4.0	0.62±0.02	18.7±0.4	0.22±0.02	-	-	-	-
	DCAE	86.1±4.0	<u>0.68±0.02</u>	19.4±0.3	0.20±0.02	-	-	-	-
Quantum	QAE	87.7±3.6	0.65±0.03	18.8±0.3	0.21±0.02	87.7±4.3	0.62±0.05	17.8±0.3	0.22±0.02
	varQCAE	86.7±3.3	0.68±0.02	19.0±0.3	0.20±0.02	87.5±4.0	0.65±0.02	18.0±0.3	0.21±0.02
	Our model	<u>84.4±3.3</u>	0.68±0.03	<u>19.5±0.3</u>	<u>0.19±0.01</u>	<u>84.6±3.6</u>	0.66±0.04	<u>18.5±0.3</u>	<u>0.20±0.01</u>
	Our model-QIB	83.7±2.8	0.70±0.02	19.7±0.3	0.17±0.01	84.2±3.3	0.68±0.02	18.7±0.3	0.19±0.01

TABLE IV: Image generation performance on CIFAR-100 and CelebA datasets (Ideal and Noisy). Bold: best, underlined: second best.

Category	Method	Ideal Environment					Noisy Environment			
		FID	SSIM	PSNR	LPIPS		FID	SSIM	PSNR	LPIPS
CIFAR-100										
Classical	VAE	155.7±11.6	0.45±0.02	18.7±0.3	0.24±0.02	-	-	-	-	
	CVAE	161.3±10.4	0.50±0.03	18.4±0.4	0.25±0.03	-	-	-	-	
Quantum	QGAN	167.7±12.2	0.41±0.04	18.1±0.5	0.27±0.03	173.2±13.4	0.37±0.06	17.7±0.5	0.28±0.04	
	Our model	154.1±9.3	0.52±0.03	19.5±0.3	0.22±0.02	155.5±8.6	0.48±0.05	18.4±0.4	0.23±0.03	
	Our model-QIB	153.5±8.4	0.54±0.02	19.8±0.3	0.20±0.01	154.4±7.4	0.50±0.04	18.7±0.3	0.21±0.02	
CELEBA										
Classical	VAE	85.7±3.8	0.65±0.01	20.5±0.4	0.20±0.01	-	-	-	-	
	CVAE	86.6±4.0	0.72±0.05	20.2±0.5	0.21±0.02	-	-	-	-	
Quantum	QGAN	90.1±3.3	0.65±0.02	19.8±0.6	0.23±0.02	94.3±5.1	0.63±0.06	18.6±0.6	0.25±0.03	
	Our model	84.8±3.3	0.72±0.01	20.8±0.3	0.19±0.01	85.7±3.3	0.70±0.04	19.4±0.4	0.20±0.02	
	Our model-QIB	84.8±2.9	0.74±0.02	21.1±0.3	0.18±0.01	84.9±2.9	0.72±0.03	19.7±0.3	0.19±0.01	

latent space compared to generation.

Figure 2 shows QCVAE-generated image samples during training under ideal conditions on the MNIST dataset. The model generates high-quality images overall. For different classes, generation quality varies, such as for digits 2 and 7. Since these inputs have more diverse shapes than other classes, generating better images is relatively difficult, but QCVAE still demonstrates good generalization ability. Figure 3(a) contains samples for MNIST-Fashion, CIFAR-100 and CelebA datasets, including original, reconstructed and generated images. Overall, our model maintains relatively high image reconstruction and generation quality. Due to CIFAR dataset quality issues, the visual effect is worse than other datasets, but FID comparisons show better performance than other models.

Ideal simulated environment. Table I displays the FID results for the QCAE model within an ideal simulated quantum environment. This segment of the experiment took place in an ideal quantum simulation environment without any noise and based on Ideal quantum circuit.

Figure 2 shows QCVAE-generated image samples during

training under ideal conditions on the MNIST dataset. The model generates high-quality images overall. For different classes, generation quality varies, such as for digits 2 and 7. Since these inputs have more diverse shapes than other classes, generating better images is relatively difficult, but QCVAE still demonstrates good generalization ability. Figure 3(a) contains samples for MNIST-Fashion, CIFAR-100 and CelebA datasets, including original, reconstructed and generated images. Overall, our model maintains relatively high image reconstruction and generation quality. Due to CIFAR dataset quality issues, the visual effect is worse than other datasets, but FID comparisons show better performance than other models.

Noisy simulated environment. To assess QCAE robustness against noise on NISQ devices, we conducted training and testing in simulated noisy environments. Noise was characterized by depolarization error, introducing random Pauli operations (X, Y, Z) to qubits, quantified as

$$\rho(x) = (1 - \epsilon) U(x) |0\rangle \langle 0| U(x)^\dagger + \epsilon \frac{1}{2^{N_Q}}, \quad (19)$$

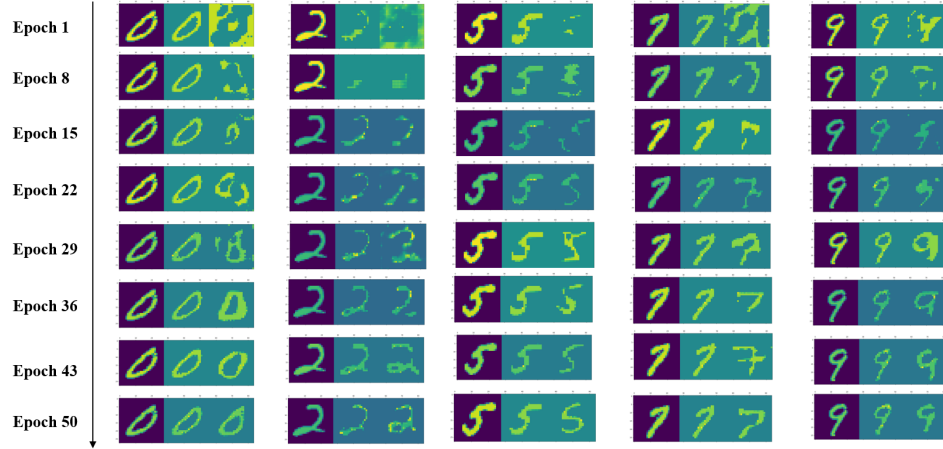


Fig. 2: Samples during training on the MNIST dataset. From top to bottom, the reconstructed image and generated image change with epoch as the training progresses. The left side of each set of images is the original image in the data set, which is also the image input to the model. In the middle is the reconstructed image based on the input image. On the right is a generated image based on random sampling.

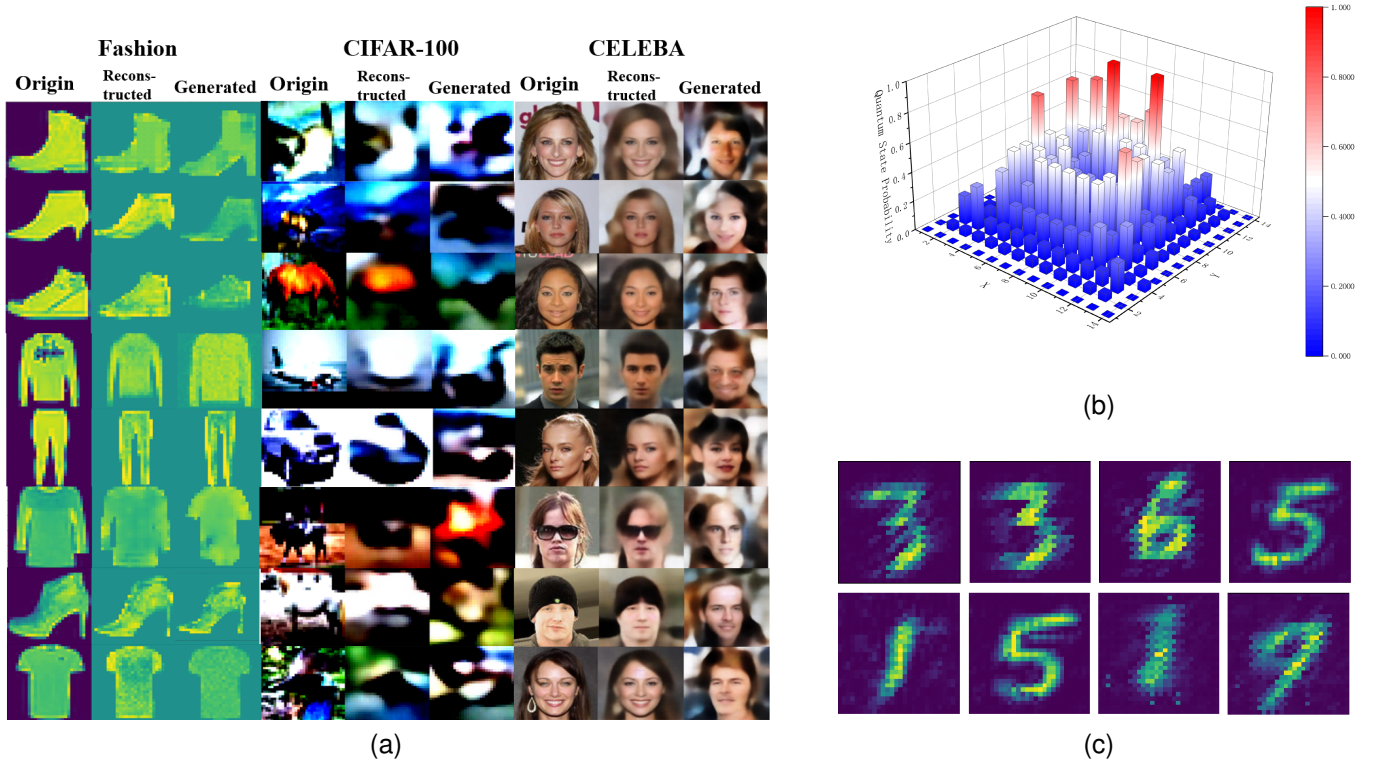


Fig. 3: Sample results and real quantum computer validation. (a) Generated samples across different datasets showing reconstruction and generation quality. For each tuple, from left to right are the original image, the reconstructed image and the generated image based on random z-space. Compared with the Fashion dataset and the CelebA dataset, the CIFAR-100 images are relatively inadequate in terms of the overall data quality (type, size, color, etc.), so the visual effect of the generated image is blurry. This can also be reflected in the FID results. (b) Feature representation on IBM quantum computer showing compressed representation of images. (c) MNIST images generated on actual quantum hardware demonstrating practical viability of the proposed QCVAE model.

$$E(\rho) = (1 - \varepsilon) \rho + \varepsilon \text{Tr}[\rho] \frac{I}{2^n}, \quad (20)$$

where ε is the depolarizing error parameter, n is the number of qubits, and I is the Pauli matrix. The depolarization error parameter was set to 0.1.

Our QCAE's FID results showcase robust error tolerance, attributed to its distinctive single-qubit implementation in tables in the noisy part. Particularly for denoising tasks with Gaussian noise applied directly to data, quantum simulator noise during training has minimal impact on results. Despite inherent NISQ-era functional limitations where unavoidable noise often causes notable performance degradation on real quantum computers, this experiment proves the QCAE model consistently maintains strong performance even with noise present.

B. Real Quantum Device Test Experiment

We assessed our single-qubit QCVAE model efficacy using the MNIST dataset on an actual IBM quantum computer. Model deployment in a real quantum environment occurred on an IBM online quantum computer through IBM Quantum Lab. However, the limited quantum volume of IBM NISQ systems imposed restrictions on available quantum gates. To overcome this limitation, we leveraged our single-qubit method, which combines parameters and encoding, reducing necessary quantum gates.

Figure 3(b,c) illustrates MNIST image features processed by the QCVAE model on a quantum computer, demonstrating the model's ability to generate high-quality images on quantum computers.

QGAN's [41] performance is not as strong as expected, despite incorporating an additional discriminator compared to autoencoders. While structurally similar to our model, QGAN's generator uses multiple quantum circuits to compensate for limited quantum resources and VQC training challenges. However, QGAN's quantum generator outputs image data row by row, limiting spatial correlation capture effectiveness. Due to VQC training constraints, original QGAN reduces image resolution to 8×8 before training. To match our experimental data size, we increased VQC instances based on 8 qubits, significantly increasing required quantum resources—an issue our single-qubit method avoids entirely. Moreover, QGAN's row-wise processing extracts local features less effectively than convolutional operations. Unlike QGAN, which relies on classical KL divergence-based loss functions, our approach leverages QIB to optimize latent space representations, preserving relevant information while reducing redundancy. Additionally, since QGAN's discriminator remains a classical neural network, it doesn't fully utilize quantum resources, lacking a dedicated VQC for adversarial learning. These problems are more significant on CIFAR-100 and CelebA datasets. A detailed comparison of QPU resources used in our model versus QGAN is presented in Table V.

IV. DISCUSSION

In the previous section, we introduced our experimental findings and conducted a comprehensive analysis. Our

TABLE V: Comparison of models running on quantum computers

	QGAN	Our model
Frequency(GHz)	5.552	4.963
Single-qubit gate AVG. error	4e-4	3.093e-4
Two-qubit gate AVG. error	1.5e-2	8.256e-3
Average readout error	6.55e-2	2.594e-2

methodology utilizes a single-qubit approach to build quantum convolutional layers, efficiently encoding features of image data while maximizing the use of a single qubit. This strategy shows potential for resource conservation, particularly in the NISQ era with its constraints on quantum computing resources. Considering the objectives of our research and the experimental context, our model can provide researchers with an efficient and resource-friendly research tool, enabling quantum computing to advance towards quantum utility.

The quantum information bottleneck becomes apparent in the optimization of our quantum model during training. Leveraging information bottlenecks in the transformation of classical data into quantum space allows us to turn information loss, initially considered a disadvantage, into an advantage for extracting useful information. Our experimental results further demonstrate that quantum information bottlenecks contribute to the stability of model training.

In the NISQ era, implementing quantum algorithms on noisy devices faces limitations due to inherent device noise. Nevertheless, our proposed model, employing the single-qubit method, exhibits robustness to noise in both simulated and actual quantum computing environments. While the quantum error correction algorithm is often employed to address errors caused by device noise, it requires additional resources. In contrast, the intrinsic robustness of our algorithm provides a means to conserve valuable resources in the NISQ era.

Efficiently utilizing the single-qubit method enables us to minimize the total qubit count, addressing the training difficulties associated with the “barren plateau” phenomenon in variational quantum circuits. Despite existing efforts [5, 50, 51] to mitigate the “barren plateau”, a comprehensive solution remains elusive. In the current realm of quantum machine learning utilizing variational quantum algorithms, the persistence of the “barren plateau” may endure. Nevertheless, our approach offers a practical means to navigate around this challenge.

We assessed the efficacy of our quantum circuit on an IBM quantum computer and observed its sustained capability. Notably, our circuit boasts fewer qubits and quantum volumes compared to conventional quantum circuits, suggesting a promising path for the integration of quantum machine learning in the future. At present, our model demonstrates the capability to address practical image-generation tasks on quantum computers.

This study employed a single-qubit model, utilizing only one qubit and forgoing potential advantages associated with multiple qubits. This deliberate choice renders the single-qubit method less susceptible to quantum decoherence, leveraging the inherent coherence of quantum systems. While future in-

vestigations could explore the potential benefits of employing multiple qubits, such as increased processing speed, these advantages may become achievable in the future. Notably, [22] showcased the general approximation power of the single-qubit method, while [52] utilized Fourier series to establish its approximation ability for any univariate function, albeit with uncertainty regarding its expressivity for multivariate functions. Our work takes a pragmatic approach, demonstrating the effectiveness of single-qubit-based quantum neural networks in handling complex data. We hope this practical case contributes to alleviating concerns in this domain.

While our single-qubit approach demonstrates clear advantages in resource efficiency, scaling to extremely large datasets presents both opportunities and challenges. Our experimental progression from MNIST to CelebA validates the theoretical scalability framework, showing consistent performance across scales. However, several scaling considerations merit discussion: (1) Computational Complexity: while quantum resource usage remains constant, classical preprocessing and longer quantum sequences may increase training time for ultra-high-resolution tasks, (2) Hardware Limitations: current NISQ device coherence times constrain maximum sequence lengths, potentially limiting scaling to extremely complex tasks, and (3) Feature Representation: very high-dimensional feature spaces may require hybrid classical-quantum architectures for optimal performance. In order to explore and study these issues, the following approaches are worth trying in the future: Combining classical convolutional preprocessing with quantum feature extraction for ultra-high-resolution images, future quantum hardware supporting multiple independent single-qubit processors for parallel feature extraction, breaking complex tasks into manageable sub-problems processed sequentially, and adapting quantum circuit depth to specific device capabilities. These approaches maintain the resource efficiency advantages while enabling scaling to more complex tasks.

V. CONCLUSION

In this work, we propose a quantum convolutional variational autoencoder model based on the single-qubit method and design a quantum information bottleneck optimization method for its implementation. Additionally, we demonstrate the application of varQCAE in various toy scenarios, such as quantum circuit denoising and image generation. Finally, we performed numerical evaluations and implemented the QCVAE application using PennyLane simulation platform and IBM Quantum Computer. This advancement opens pathways for practical applications, where quantum-enhanced generative models could enable efficient image reconstruction and augmentation under computational constraints. There is much potential for further progress, which includes the use of QCVAE for data generation and feature extraction. On the other hand, investigating practical applications in this work is also crucial. Our results in real quantum computer showing a potential application of quantum computing towards real-world applications.

REFERENCES

- [1] H.-Y. Huang, M. Broughton, M. Mohseni, R. Babbush, S. Boixo, H. Neven, and J. R. McClean, “Power of data in quantum machine learning,” *Nature communications*, vol. 12, no. 1, p. 2631, 2021.
- [2] M. T. West, S.-L. Tsang, J. S. Low, C. D. Hill, C. Leckie, L. C. Hollenberg, S. M. Erfani, and M. Usman, “Towards quantum enhanced adversarial robustness in machine learning,” *Nature Machine Intelligence*, vol. 5, no. 6, pp. 581–589, 2023.
- [3] J. Preskill, “Quantum computing in the nisq era and beyond,” *Quantum*, vol. 2, p. 79, 2018.
- [4] A. A. Clerk, M. H. Devoret, S. M. Girvin, F. Marquardt, and R. J. Schoelkopf, “Introduction to quantum noise, measurement, and amplification,” *Reviews of Modern Physics*, vol. 82, no. 2, p. 1155, 2010.
- [5] J. R. McClean, S. Boixo, V. N. Smelyanskiy, R. Babbush, and H. Neven, “Barren plateaus in quantum neural network training landscapes,” *Nature communications*, vol. 9, no. 1, p. 4812, 2018.
- [6] M. C. Caro, H.-Y. Huang, M. Cerezo, K. Sharma, A. Sornborger, L. Cincio, and P. J. Coles, “Generalization in quantum machine learning from few training data,” *Nature communications*, vol. 13, no. 1, p. 4919, 2022.
- [7] Y. Kim, A. Eddins, S. Anand, K. X. Wei, E. Van Den Berg, S. Rosenblatt, H. Nayfeh, Y. Wu, M. Zaletel, K. Temme *et al.*, “Evidence for the utility of quantum computing before fault tolerance,” *Nature*, vol. 618, no. 7965, pp. 500–505, 2023.
- [8] M. S. Benkner, Z. Lähner, V. Golyanik, C. Wunderlich, C. Theobalt, and M. Moeller, “Q-match: Iterative shape matching via quantum annealing,” in *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, October 2021, pp. 7586–7596.
- [9] E. Ovalle-Magallanes, J. G. Avina-Cervantes, I. Cruz-Aceves, and J. Ruiz-Pinales, “Hybrid classical-quantum convolutional neural network for stenosis detection in x-ray coronary angiography,” *Expert Systems with Applications*, vol. 189, p. 116112, 2022.
- [10] Y.-F. Yang and M. Sun, “Semiconductor defect detection by hybrid classical-quantum deep learning,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2022, pp. 2323–2332.
- [11] D. Konar, N. Sreekumar, R. Jiang, and V. Aggarwal, “Alz-qnet: A quantum regression network for studying alzheimer’s gene interactions,” *Computers in Biology and Medicine*, vol. 196, p. 110837, 2025.
- [12] D. Konar, S. Bhattacharyya, T. K. Gandhi, B. K. Panigrahi, and R. Jiang, “3-d quantum-inspired self-supervised tensor network for volumetric segmentation of medical images,” *IEEE Transactions on Neural Networks and Learning Systems*, vol. 35, no. 8, pp. 10312–10325, 2023.
- [13] Y. Zhu, R. Jiang, Q. Ni, and A. Bouridane, “Enable quantum graph neural networks on a single qubit with quantum walk,” *IEEE Transactions on Artificial Intelligence*, vol. 1, no. 01, pp. 1–10, 2025.

- [14] D. Silver, T. Patel, W. Cutler, A. Ranjan, H. Gandhi, and D. Tiwari, “Mosaik: Quantum generative adversarial networks for image generation on nisc computers,” in *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, October 2023, pp. 7030–7039.
- [15] M. Cerezo, A. Arrasmith, R. Babbush, S. C. Benjamin, S. Endo, K. Fujii, J. R. McClean, K. Mitarai, X. Yuan, L. Cincio *et al.*, “Variational quantum algorithms,” *Nature Reviews Physics*, vol. 3, no. 9, pp. 625–644, 2021.
- [16] J. Preskill, “Quantum computing and the entanglement frontier,” *arXiv preprint arXiv:1203.5813*, 2012.
- [17] I. Cong, S. Choi, and M. D. Lukin, “Quantum convolutional neural networks,” *Nature Physics*, vol. 15, no. 12, pp. 1273–1278, 2019.
- [18] L. Deng, “The mnist database of handwritten digit images for machine learning research [best of the web],” *IEEE signal processing magazine*, vol. 29, no. 6, pp. 141–142, 2012.
- [19] S. Krinner, S. Storz, P. Kurpiers, P. Magnard, J. Heinsoo, R. Keller, J. Luetolf, C. Eichler, and A. Wallraff, “Engineering cryogenic setups for 100-qubit scale superconducting circuit systems,” *EPJ Quantum Technology*, vol. 6, no. 1, p. 2, 2019.
- [20] A. Pérez-Salinas, A. Cervera-Lierta, E. Gil-Fuster, and J. I. Latorre, “Data re-uploading for a universal quantum classifier,” *Quantum*, vol. 4, p. 226, 2020.
- [21] P. Easom-McCaldin, A. Bouridane, A. Belatreche, R. Jiang, and S. Al-Maadeed, “Efficient quantum image classification using single qubit encoding,” *IEEE Transactions on Neural Networks and Learning Systems*, vol. 35, no. 2, pp. 1472–1486, 2022.
- [22] A. Pérez-Salinas, D. López-Núñez, A. García-Sáez, P. Forn-Díaz, and J. I. Latorre, “One qubit as a universal approximant,” *Phys. Rev. A*, vol. 104, p. 012405, Jul 2021. [Online]. Available: <https://link.aps.org/doi/10.1103/PhysRevA.104.012405>
- [23] S. Yang and B. Chen, “Effective surrogate gradient learning with high-order information bottleneck for spike-based machine intelligence,” *IEEE transactions on neural networks and learning systems*, 2023.
- [24] —, “Snib: improving spike-based machine learning using nonlinear information bottleneck,” *IEEE transactions on systems, man, and cybernetics: Systems*, vol. 53, no. 12, pp. 7852–7863, 2023.
- [25] S. Yang, B. Linares-Barranco, Y. Wu, and B. Chen, “Self-supervised high-order information bottleneck learning of spiking neural network for robust event-based optical flow estimation,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2024.
- [26] S. Yang, H. Wang, and B. Chen, “Sibols: robust and energy-efficient learning for spike-based machine intelligence in information bottleneck framework,” *IEEE Transactions on cognitive and developmental systems*, vol. 16, no. 5, pp. 1664–1676, 2023.
- [27] M. Hayashi and Y. Yang, “Efficient algorithms for quantum information bottleneck,” *Quantum*, vol. 7, p. 936, 2023.
- [28] L. Banchi, J. Pereira, and S. Pirandola, “Generalization in quantum machine learning: A quantum information standpoint,” *PRX Quantum*, vol. 2, p. 040321, Nov 2021. [Online]. Available: <https://link.aps.org/doi/10.1103/PRXQuantum.2.040321>
- [29] E. T. Campbell, B. M. Terhal, and C. Vuillot, “Roads towards fault-tolerant universal quantum computation,” *Nature*, vol. 549, no. 7671, pp. 172–179, 2017.
- [30] Y. Zhu, A. Bouridane, M. E. Celebi, D. Konar, P. Angelov, Q. Ni, and R. Jiang, “Quantum face recognition with multigate quantum convolutional neural network,” *IEEE Transactions on Artificial Intelligence*, vol. 5, no. 12, pp. 6330–6341, 2024.
- [31] Y. Du, M.-H. Hsieh, T. Liu, and D. Tao, “Expressive power of parametrized quantum circuits,” *Physical Review Research*, vol. 2, no. 3, p. 033125, 2020.
- [32] M. P. Cuéllar, “What we can do with one qubit in quantum machine learning: ten classical machine learning problems that can be solved with a single qubit,” *Quantum Machine Intelligence*, vol. 6, no. 2, p. 76, 2024.
- [33] K. Hornik, “Approximation capabilities of multilayer feedforward networks,” *Neural Networks*, vol. 4, no. 2, pp. 251–257, 1991. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/089360809190009T>
- [34] A. Pérez-Salinas, D. López-Núñez, A. García-Sáez, P. Forn-Díaz, and J. I. Latorre, “One qubit as a universal approximant,” *Physical Review A*, vol. 104, no. 1, p. 012405, 2021.
- [35] M. Berta, K. Seshadreesan, and M. Wilde, “Renyi generalizations of the conditional quantum mutual information,” *Journal of Mathematical Physics*, vol. 56, 03 2014.
- [36] A. S. Holevo, “Bounds for the quantity of information transmitted by a quantum communication channel,” *Problemy Peredachi Informatsii*, vol. 9, no. 3, pp. 3–11, 1973.
- [37] H. Xiao, K. Rasul, and R. Vollgraf. (2017) Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms.
- [38] A. Krizhevsky, G. Hinton *et al.*, “Learning multiple layers of features from tiny images,” 2009.
- [39] Z. Liu, P. Luo, X. Wang, and X. Tang, “Deep learning face attributes in the wild,” in *Proceedings of International Conference on Computer Vision (ICCV)*, December 2015.
- [40] Y. Du and D. Tao, “On exploring practical potentials of quantum auto-encoder with advantages,” *arXiv preprint arXiv:2106.15432*, 2021.
- [41] H.-L. Huang, Y. Du, M. Gong, Y. Zhao, Y. Wu, C. Wang, S. Li, F. Liang, J. Lin, Y. Xu *et al.*, “Experimental quantum generative adversarial networks for image generation,” *Physical Review Applied*, vol. 16, no. 2, p. 024051, 2021.
- [42] J. Wu, H. Fu, M. Zhu, W. Xie, and X.-Y. Li, “Quantum circuit autoencoder,” *arXiv preprint arXiv:2307.08446*, 2023.
- [43] M. Heusel, H. Ramsauer, T. Unterthiner, B. Nessler, and S. Hochreiter, “Gans trained by a two time-scale update rule converge to a local nash equilibrium,” *Advances in*

- neural information processing systems*, vol. 30, 2017.
- [44] Z. Wang, A. C. Bovik, H. R. Sheikh, and E. P. Simoncelli, "Image quality assessment: from error visibility to structural similarity," *IEEE transactions on image processing*, vol. 13, no. 4, pp. 600–612, 2004.
 - [45] A. Hore and D. Ziou, "Image quality metrics: Psnr vs. ssim," in *2010 20th international conference on pattern recognition*. IEEE, 2010, pp. 2366–2369.
 - [46] R. Zhang, P. Isola, A. A. Efros, E. Shechtman, and O. Wang, "The unreasonable effectiveness of deep features as a perceptual metric," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2018, pp. 586–595.
 - [47] V. Bergholm, J. Izaac, M. Schuld, C. Gogolin, S. Ahmed, V. Ajith, M. S. Alam, G. Alonso-Linaje, B. Akash-Narayanan, A. Asadi *et al.*, "Pennylane: Automatic differentiation of hybrid quantum-classical computations," *arXiv preprint arXiv:1811.04968*, 2018.
 - [48] A. Paszke, S. Gross, F. Massa, A. Lerer, J. Bradbury, G. Chanan, T. Killeen, Z. Lin, N. Gimelshein, L. Antiga *et al.*, "Pytorch: An imperative style, high-performance deep learning library," *Advances in neural information processing systems*, vol. 32, 2019.
 - [49] IBM, "Ibm quantum," <https://quantum-computing.ibm.com/>, 2023, [Online; accessed 1-March-2023].
 - [50] E. Grant, L. Wossnig, M. Ostaszewski, and M. Benedetti, "An initialization strategy for addressing barren plateaus in parametrized quantum circuits," *Quantum*, vol. 3, p. 214, 2019.
 - [51] J. Stokes, J. Izaac, N. Killoran, and G. Carleo, "Quantum natural gradient," *Quantum*, vol. 4, p. 269, 2020.
 - [52] Z. Yu, H. Yao, M. Li, and X. Wang, "Power and limitations of single-qubit native quantum neural networks," in *Advances in Neural Information Processing Systems*, S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh, Eds., vol. 35. Curran Associates, Inc., 2022, pp. 27 810–27 823. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2022/file/b250de41980b58d34d6aad3f4aedd4c-Paper-Conference.pdf