

Negative Binomial Variational Autoencoders for Overdispersed Latent Modeling

Yixuan Zhang^{1*}, Jinhao Sheng^{2*}, Wenxin Zhang³, Quyu Kong⁴, Feng Zhou^{5†}

¹Southeast University, ²China Medical University Shenyang,

³University of Chinese Academy of Science, ⁴Alibaba Cloud,

⁵ Center for Applied Statistics and School of Statistics, Renmin University of China

zhixuan@hotmail.com, jhsheng@cmu.edu.cn, feng.zhou@ruc.edu.cn

Abstract

Although artificial neural networks are often described as brain-inspired, their representations typically rely on continuous activations, such as the continuous latent variables in variational autoencoders (VAEs), which limits their biological plausibility compared to the discrete spike-based signaling in real neurons. Extensions like the Poisson VAE introduce discrete count-based latents, but their equal mean-variance assumption fails to capture overdispersion in neural spikes, leading to less expressive and informative representations. To address this, we propose NegBio-VAE, a negative-binomial latent-variable model with a dispersion parameter for flexible spike count modeling. NegBio-VAE preserves interpretability while improving representation quality and training feasibility via novel KL estimation and reparameterization. Experiments on four datasets demonstrate that NegBio-VAE consistently achieves superior reconstruction and generation performance compared to competing single-layer VAE baselines, and yields robust, informative latent representations for downstream tasks. Extensive ablation studies are performed to verify the model's robustness w.r.t. various components. Our code is available at <https://github.com/co234/NegBio-VAE>.

1. Introduction

Although artificial neural networks (ANNs) have historically been described as brain-inspired, their design choices are primarily driven by computational considerations rather than strict biological fidelity [1, 44]. A key distinction lies in how information is represented: while biological neurons communicate through sequences of action potentials (spike trains) [33], most machine learning models adopt continuous activations. This contrast has motivated a line

of work that investigates discrete, spike like representations as a pathway toward enriching the expressiveness of generative models [2, 16, 30]. From this perspective, studying count-based representations is not only biologically inspired but also methodologically valuable for expanding the modeling capacity of deep generative frameworks [15, 48].

Among these frameworks, the variational autoencoder (VAE) [23] is a powerful generative model grounded in Bayesian inference that learns structured latent representations of data, and is often described as brain-inspired due to its similarity to how the brain encodes sensory information [31, 42, 45]. While VAEs have achieved broad success, they typically employ continuous latent variables, in contrast to the discrete spike counts encoded by the brain. To bridge this gap, recent works have proposed extensions such as categorical or Poisson VAEs [18, 46, 49], which introduce discrete latent variables that not only offer greater biological plausibility but also enhance the capacity to model categorical or count structures in latent variables.

The main improvement presented in this paper builds on the Poisson VAE ($\mathcal{P}$ -VAE) [46], which encodes data as discrete spike counts drawn from a Poisson distribution. While the Poisson model provides a natural starting point, it imposes a restrictive assumption: the mean and variance of the discrete spike counts must be equal. In practice, however, neural spike trains often exhibit overdispersion, where the variance of the spike counts significantly exceeds the mean [32, 41, 43]. This has been linked to neurobiological sources such as trial-to-trial gain variability and network-level fluctuations [41]. While underdispersion can arise in neurons with refractory periods [4], overdispersion is the more prevalent and consequential deviation from Poisson statistics across cortical recordings [17, 40]. This coupling of the mean and variance limits the flexibility of the latent space, leading to underestimated uncertainty and reduced representational expressiveness.

To address this limitation, we adopt the negative binomial (NB) distribution [39], a two-parameter generalization of the Poisson distribution that introduces a dispersion pa-

*Equal contribution.

†Corresponding author.

parameter, allowing the variance to exceed the mean. This flexibility allows modeling of overdispersed spike counts, enabling latent representations that better capture the heterogeneous variability. Building on this idea, we propose NegBio-VAE (see Fig. 1), a principled extension of the VAE framework that preserves count-based representations while more accurately reflecting their statistical variability. While this formulation greatly enhances representational flexibility, it also introduces two challenges: (1) computing the KL divergence between NB distributions, and (2) performing reparameterized sampling. We address both with efficient approximations that make NegBio-VAE practically trainable. Empirically, NegBio-VAE demonstrates superior reconstruction quality, stronger generative performance, and more informative latent representations for downstream tasks.

Our main contributions are summarized as follows: (1) We propose NegBio-VAE, which introduces a dispersion parameter to model overdispersed latent spike counts and improve the flexibility of latent representations. (2) We develop efficient training strategies with two KL estimators (Monte Carlo and Dispersion Sharing) and two differentiable reparameterizations (Gumbel–Softmax and Continuous-time Simulation) for stable optimization. (3) Experiments on four benchmark datasets show that NegBio-VAE outperforms strong baselines in reconstruction and generation while learning more informative latent representations for downstream tasks.

2. Related Works

Brain-like ANNs, emerging at the intersection of neuroscience and machine learning, aim to mirror the brain’s functionality and structure. Related works can be categorized into two types: spiking neural networks (SNNs) and brain-like generative models. SNNs [6, 11, 15, 27, 54], like biological neurons, use discrete spikes for communication instead of continuous activations as in traditional ANNs. A notable model is the leaky integrate-and-fire (LIF) model, which simulates the temporal dynamics of spike generation. The second category includes generative models that learn data representations similar to how brain processes sensory information. Key works in this area include brain-like VAEs [19, 46, 51], GANs [12, 24, 38], and diffusion models [5, 20, 28]. Our work extends the $\mathcal{P}$ -VAE [46] by incorporating a NB distribution to better capture overdispersion in latent spike counts, enabling richer and more flexible variability in the latent representations.

Discrete VAEs are typically categorized into two types: discrete representations and discrete data. In VAEs with discrete representations, the variables capture the underlying discrete structure of the data. Most current works on discrete-representation VAEs use categorical distributions for the latent variables [10, 14, 18, 49]. Other works em-

ploy Bernoulli [19, 37] or Poisson distributions [46, 52]. These methods have achieved significant success in speech synthesis and image generation. The second category focuses on VAEs for discrete data, such as text, categorical, or count data. These models reconstruct discrete data, making them suitable for tasks like natural language processing and structured prediction [35, 53]. While [53] uses the NB distribution to model count data while keeping the latent variables continuous, our work extends NB modeling to discrete latent variables in a VAE.

3. Preliminaries

This section reviews VAE and $\mathcal{P}$ -VAE, first covering the standard VAE framework and then its adaptation to model latent spike counts with a Poisson distribution.

3.1. Variational Autoencoder

VAE [22] is a probabilistic generative model defining a joint distribution $p(\mathbf{x}, \mathbf{z})$ over data $\mathbf{x}$ and latent variables $\mathbf{z}$. Samples are generated by $\mathbf{z} \sim p(\mathbf{z})$ and decoded via $p_\theta(\mathbf{x} | \mathbf{z})$, while inference uses an approximate posterior $q_\phi(\mathbf{z} | \mathbf{x})$. Model parameters are learned by maximizing the evidence lower bound (ELBO), a tractable surrogate of $\log p(\mathbf{x})$ that balances reconstruction and latent regularization:

$$\mathcal{L}_{\text{VAE}} = \mathbb{E}_{q_\phi(\mathbf{z}|\mathbf{x})}[\log p_\theta(\mathbf{x} | \mathbf{z})] - \mathcal{D}_{\text{KL}}[q_\phi(\mathbf{z} | \mathbf{x}) || p(\mathbf{z})].$$

The first term enforces faithful reconstruction, while the second term regularizes the latent space. VAEs enable gradient-based optimization via the reparameterization trick [22], which introduces differentiable sampling between the encoder and decoder. Standard implementations assume an isotropic Gaussian prior $p(\mathbf{z})$, simplifying computation but limiting expressiveness.

3.2. Poisson VAE

To better mimic biological neuron activity, the $\mathcal{P}$ -VAE [46] was proposed to model spike counts as discrete latent variables. Specifically, it uses the Poisson distribution to represent the spike counts of K neurons, with the latent variable $\mathbf{z} \in \mathbb{Z}_0^{+K}$. The prior and variational posterior are defined as:

$$\begin{aligned} \text{Prior: } p(\mathbf{z}) &= \text{Poi}(\mathbf{z}; \mathbf{r}), \\ \text{Posterior: } q(\mathbf{z} | \mathbf{x}) &= \text{Poi}(\mathbf{z}; \mathbf{r} \odot \delta_r(\mathbf{x})), \end{aligned}$$

where both the prior Poisson and the posterior Poisson are factorized, i.e., $\text{Poi}(\mathbf{z}) = \prod_{i=1}^K \text{Poi}(z_i)$. Here, $\mathbf{r} \in \mathbb{R}^{+K}$ denotes the prior firing rates, and $\mathbf{r} \odot \delta_r(\mathbf{x})$ gives the posterior firing rates, with $\odot$ denoting element-wise multiplication. The encoder output $\delta_r(\mathbf{x}) \in \mathbb{R}^{+K}$ modulates the ratio of posterior to prior firing rates based on the input. In contrast to standard VAEs where latent variables are continuous and typically drawn from a Gaussian, the $\mathcal{P}$ -VAE models $\mathbf{z}$

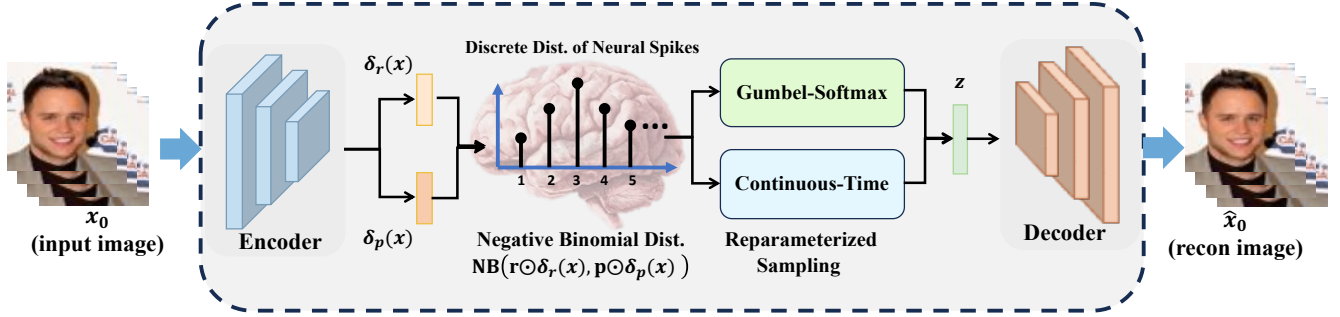


Figure 1. Overview of the proposed NegBio-VAE framework. The data are encoded as discrete spike counts drawn from a negative binomial distribution, whose variance exceeds the mean, enabling the model to capture overdispersed latent structures.

as a vector of discrete spike counts, which better resembles neural firing behavior. The objective of $\mathcal{P}$ -VAE is given by:

$$\mathcal{L}_{\mathcal{P}\text{-VAE}} = \mathbb{E}_{\text{Poi}(\mathbf{z}; \mathbf{r} \odot \delta_r(\mathbf{x}))} [\log p_\theta(\mathbf{x} | \mathbf{z})] + \sum_{i=1}^K r_i g(\delta_{r_i}), \quad (1)$$

where $g(a) = 1 - a + a \log a$ corresponds to the KL divergence between two Poisson distributions.

4. Methodology

A key limitation of the Poisson distribution is its restrictive assumption that the mean and variance of spike counts are equal. This assumption fails to capture the overdispersion frequently observed in neural spike train. To address this, we propose the NegBio-VAE, which applies a more flexible NB distribution. As a two-parameter generalization of the Poisson, the NB distribution introduces a dispersion parameter that allows the variance to exceed the mean. This makes it more suitable for modeling overdispersed spike counts. The NB distribution has been widely applied in various fields, such as spiking neuron models [34], RNA sequence analysis [9], and language modeling [55].

We begin by defining the prior and posterior distributions over the latent spike counts $\mathbf{z} \in \mathbb{Z}_0^{+K}$ as $p(\mathbf{z}) = \text{NB}(\mathbf{z}; \mathbf{r}, \mathbf{p})$ and $q(\mathbf{z} | \mathbf{x}) = \text{NB}(\mathbf{z}; \mathbf{r} \odot \delta_r(\mathbf{x}), \mathbf{p} \odot \delta_p(\mathbf{x}))$, respectively. Similar to $\mathcal{P}$ -VAE, both the prior and posterior NB distribution are factorized, i.e., $\text{NB}(\mathbf{z}) = \prod_{i=1}^K \text{NB}(z_i)$, $\delta_r(\mathbf{x})$ and $\delta_p(\mathbf{x})$ are outputs of the encoder, which captures the ratio of the posterior parameters to the prior parameters. With this setup, the ELBO of NegBio-VAE becomes:

$$\mathcal{L} = \mathbb{E}_{\text{NB}(\mathbf{z}; \mathbf{r} \odot \delta_r(\mathbf{x}), \mathbf{p} \odot \delta_p(\mathbf{x}))} [\log p_\theta(\mathbf{x} | \mathbf{z})] - \mathcal{D}_{\text{KL}}[\text{NB}(\mathbf{z}; \mathbf{r} \odot \delta_r(\mathbf{x}), \mathbf{p} \odot \delta_p(\mathbf{x})) || \text{NB}(\mathbf{z}; \mathbf{r}, \mathbf{p})]. \quad (2)$$

While this formulation enables greater flexibility, it also introduces two key technical challenges during the training of NegBio-VAE: (1) The second term in Eq. (2) requires calculating the KL divergence between two NB distributions;

(2) The first term in Eq. (2) requires reparameterized sampling from the NB distribution. We address each of these issues in the following sections.

4.1. KL Divergence between NB Distributions

In both vanilla VAE and $\mathcal{P}$ -VAE, the KL term is tractable due to closed-form solutions for Gaussian and Poisson distributions. However, no such form exists for the KL divergence between two NB distributions, which poses the first challenge for training NegBio-VAE. To address this, we propose two strategies: a **Monte Carlo** method for direct approximation, and a **dispersion sharing** technique that simplifies the KL divergence by partially tying posterior parameters to the prior.

(1) Monte Carlo. Optimizing the ELBO for NegBio-VAE is challenging due to the lack of an analytical form for the KL divergence between two NB distributions, preventing direct computation of the KL term. We address this using Monte Carlo estimation. Specifically, using $\mathcal{D}_{\text{KL}}[q(\mathbf{z}) || p(\mathbf{z})] = \mathbb{E}_q(\mathbf{z}) [\log q(\mathbf{z}) - \log p(\mathbf{z})]$, the KL divergence is approximated by sampling from the variational posterior and averaging the log-density difference between posterior and prior.

Substituting this expression into the ELBO in Eq. (2) we obtain the following objective:

$$\mathcal{L} = \mathbb{E}_{\text{NB}(\mathbf{z}; \mathbf{r} \odot \delta_r(\mathbf{x}), \mathbf{p} \odot \delta_p(\mathbf{x}))} [\log p_\theta(\mathbf{x} | \mathbf{z}) - \log \text{NB}(\mathbf{z}; \mathbf{r} \odot \delta_r(\mathbf{x}), \mathbf{p} \odot \delta_p(\mathbf{x})) + \log \text{NB}(\mathbf{z}; \mathbf{r}, \mathbf{p})].$$

Clearly, as long as we can implement reparameterized sampling from the NB distribution, we can use the above objective function to train NegBio-VAE.

(2) Dispersion Sharing. Although the KL divergence between two general NB distributions, $\text{NB}(z; r_1, p_1)$ and $\text{NB}(z; r_2, p_2)$, does not have an analytical solution, a tractable analytical form exists when the dispersion parameters are shared, i.e., $r_1 = r_2 = r$.

Based on this observation, we propose an alternative strategy for computing the KL term in the NegBio-VAE by constraining the prior $\text{NB}(\mathbf{z}; \mathbf{r}, \mathbf{p})$ and the posterior

$\text{NB}(\mathbf{z}; \mathbf{r} \odot \delta_r(\mathbf{x}), \mathbf{p} \odot \delta_p(\mathbf{x}))$ to share the same dispersion parameter, i.e., setting $\delta_r(\mathbf{x})$ to be $\mathbf{1}$. Then, the KL term in Eq. (2) admits a closed-form solution:

$$\mathcal{D}_{\text{KL}}[\text{NB}(\mathbf{z}; \mathbf{r}, \mathbf{p} \odot \delta_p(\mathbf{x})) || \text{NB}(\mathbf{z}; \mathbf{r}, \mathbf{p})] = \sum_{i=1}^K r_i g(p_i, \delta_{p_i}), \quad (3)$$

where $g(a, b)$ is defined as: $g(a, b) = \log b + \frac{1-ab}{ab} \log \left(\frac{1-ab}{1-a} \right)$, with $a \in (0, 1)$ and $b > 0$. The complete derivation can be found in appendix. Then, the final NegBio-VAE objective becomes:

$$\mathcal{L} = \mathbb{E}_{\text{NB}(\mathbf{z}; \mathbf{r}, \mathbf{p} \odot \delta_p(\mathbf{x}))} [\log p_\theta(\mathbf{x} | \mathbf{z})] + \sum_{i=1}^K r_i g(p_i, \delta_{p_i}). \quad (4)$$

Importantly, sharing the same dispersion parameter between the prior and posterior does not imply that they have identical means or variances. For the NB distribution, the mean is given by $r(1-p)/p$ and the variance by $r(1-p)/p^2$. Thus, even when r is the same for both, different p still allows the posterior to capture different distributional properties from the prior.

Both methods have advantages and limitations. The Monte Carlo method makes no assumptions about the variational posterior but may yield higher-variance gradient estimates. The dispersion sharing method instead assumes a shared dispersion parameter, enabling analytic KL computation. Although analytic KL does not guarantee lower gradient variance, it simplifies optimization and often improves training stability in practice while preserving the ability to capture overdispersion. We compare the performance of both strategies, and the results are presented in Sec. 5.

4.2. Reparameterized Sampling for NB Distribution

The second challenge in training NegBio-VAE lies in the sampling process. The expectation term in the ELBO requires reparameterized sampling from the NB distribution to allow efficient gradient-based optimization. Reparameterizing discrete distributions is more challenging compared to continuous ones, but it can be achieved through suitable relaxation techniques. In this section, we describe how to apply reparameterization to the NB distribution by leveraging a key property: the NB distribution can be represented as a continuous mixture of Poisson distributions, where the mixing weight being a Gamma distribution:

$$\text{NB}(z; r, p) = \int_0^\infty \text{Poi}(z | \lambda) \text{Gamma}(\lambda; r, \frac{p}{1-p}) d\lambda. \quad (5)$$

This implies that a sample from $\text{NB}(z; r, p)$ can be obtained by first sampling $\lambda \sim \text{Gamma}(r, \frac{p}{1-p})$, followed by sampling $z \sim \text{Poi}(\lambda)$.

The first step, sampling from the Gamma distribution, is straightforward to reparameterize via implicit reparameterization gradients [13]. In practice, PyTorch’s `Gamma.rsample()` function supports gradient propagation, as it uses the Marsaglia-Tsang algorithm in its underlying implementation and ensures differentiability through implicit gradient computation. The second step, sampling from the Poisson distribution, is more challenging, as it lacks a standard reparameterizable form. To address this, we adopt approximate relaxation techniques such as **Gumbel-Softmax Relaxation** [18] and **Continuous-Time Simulation** [46]. Both methods rely on a temperature parameter to transform “hard” counts into “soft” counts, thereby enabling differentiability. For implementation details, please see appendix.

(1) Gumbel-Softmax Relaxation. To enable differentiable sampling from a Poisson distribution, we adopt a relaxation-based strategy that treats the Poisson as a categorical distribution over a truncated support $\{0, 1, \dots, Z_{\max}\}$. By using the Gumbel-Softmax trick [18], we construct a soft approximation of the discrete counts:

$$\tilde{z} = \sum_{z=0}^{Z_{\max}} z \cdot \text{softmax} \left(\frac{\log \text{Poi}(z) + \epsilon_z}{\tau} \right),$$

where $\epsilon_z \sim \text{Gumbel}(0, 1)$ is an i.i.d. Gumbel noise and $\tau > 0$ is a temperature controlling the degree of relaxation. As $\tau \rightarrow 0$, the soft sample $\tilde{z}$ converges to the Poisson distribution.

(2) Continuous-Time Simulation. Following Vafaii et al. [46], we adopt the continuous-time simulation method, which leverages the connection between the Poisson distribution and the Poisson process. It models a Poisson-distributed count as the number of events occurring within the interval $[0, 1]$, where inter-arrival times follow an exponential distribution with rate λ . The soft count is computed by simulating the inter-arrival times and accumulating a temperature-smoothed approximation of the total event count:

$$\tilde{z} = \sum_{n=1}^M \sigma \left(\frac{1 - S_n}{\tau} \right),$$

where $S_n = \sum_{i=1}^n s_i$, $1 \leq n \leq M$, $\{s_i\}_{i=1}^M \sim \text{Exponential}(\lambda)$ and $\sigma(\cdot)$ is the sigmoid function, $\tau > 0$ is a temperature, and $\tau \rightarrow 0$ converges to the Poisson distribution. This approach enables differentiable Poisson sampling through reparameterizable exponential sampling.

Both Gumbel-Softmax and continuous-time relaxations are used for the Poisson step in the NB reparameterization. Theoretically, both approaches are valid. Empirically, under the same temperature, we find that the continuous-time relaxation tends to produce smoother count samples, whereas Gumbel-Softmax yields sharper ones. A detailed comparison of the two methods is provided in Sec. 5.

Table 1. Reconstruction and generation performance results on four benchmark datasets. The best and second-best results are marked in **bold** and underlined, respectively.

Dataset	Model	Reconstruction		Generation		
		MSE ↓	SSIM ↑	FID@5k ↓	FID@10k ↓	KID ↓
MNIST	$\mathcal{G}$ -VAE	0.0377	0.6790	152.5109	152.8226	0.1788 $\pm$ 0.0115
	$\mathcal{L}$ -VAE	0.0377	0.7124	132.7655	131.7514	0.1484 $\pm$ 0.0103
	$\mathcal{C}$ -VAE	0.0222	0.7712	135.4452	133.4826	0.1140 $\pm$ 0.0132
	$\mathcal{P}$ -VAE	<u>0.0125</u>	<u>0.8581</u>	105.3678	104.1416	0.1250 $\pm$ 0.0019
	NegBio-VAE _{MC-G}	0.0156	0.8487	79.6727	78.3802	0.0892 $\pm$ 0.0106
	NegBio-VAE _{MC-C}	0.0123	0.8661	<u>84.3853</u>	<u>83.0010</u>	<u>0.0906</u> $\pm$ 0.0111
	NegBio-VAE _{DS-G}	0.0168	0.7960	<u>87.6456</u>	87.4101	<u>0.1000</u> $\pm$ 0.0123
	NegBio-VAE _{DS-C}	<u>0.0125</u>	0.8554	106.4104	105.4089	0.1167 $\pm$ 0.0094
Fashion-MNIST	$\mathcal{G}$ -VAE	0.1417	0.1731	179.8126	179.2981	0.1828 $\pm$ 0.0106
	$\mathcal{L}$ -VAE	0.1274	0.2085	181.4542	179.5956	0.1847 $\pm$ 0.0112
	$\mathcal{C}$ -VAE	0.0238	0.6390	195.3205	193.0972	0.1835 $\pm$ 0.0219
	$\mathcal{P}$ -VAE	<u>0.0145</u>	<u>0.7387</u>	145.9776	146.0128	0.1667 $\pm$ 0.0133
	NegBio-VAE _{MC-G}	0.0180	0.7132	127.5248	125.9497	0.1468 $\pm$ 0.0130
	NegBio-VAE _{MC-C}	0.0152	0.7331	148.9795	147.7799	0.1688 $\pm$ 0.0128
	NegBio-VAE _{DS-G}	0.0186	0.6773	<u>133.0601</u>	<u>132.8822</u>	<u>0.1517</u> $\pm$ 0.0149
	NegBio-VAE _{DS-C}	0.0144	0.7406	155.5468	154.1402	0.1763 $\pm$ 0.0124
CIFAR _{16×16}	$\mathcal{G}$ -VAE	0.1027	0.4495	72.0683	69.7067	0.0607 $\pm$ 0.0074
	$\mathcal{L}$ -VAE	0.0807	0.5079	91.1614	89.9475	0.0857 $\pm$ 0.0096
	$\mathcal{C}$ -VAE	0.0664	0.4755	89.4235	88.7412	0.0463 $\pm$ 0.0105
	$\mathcal{P}$ -VAE	0.0357	<u>0.6791</u>	60.3653	59.1037	0.0582 $\pm$ 0.0098
	NegBio-VAE _{MC-G}	0.0470	0.6337	40.2788	39.8336	0.0348 $\pm$ 0.0065
	NegBio-VAE _{MC-C}	0.0456	0.6429	67.2898	65.6569	0.0727 $\pm$ 0.0096
	NegBio-VAE _{DS-G}	0.0388	0.6328	<u>41.7768</u>	<u>41.1260</u>	<u>0.0452</u> $\pm$ 0.0080
	NegBio-VAE _{DS-C}	0.0189	0.8089	64.9939	63.6688	0.0634 $\pm$ 0.0086
CelebA _{64×64}	$\mathcal{G}$ -VAE	0.4011	0.1772	195.1377	194.0974	0.2758 $\pm$ 0.0192
	$\mathcal{L}$ -VAE	0.3375	0.2161	199.9303	198.8191	0.2655 $\pm$ 0.0117
	$\mathcal{C}$ -VAE	0.0774	0.4662	166.2762	165.7814	0.1648 $\pm$ 0.0139
	$\mathcal{P}$ -VAE	0.0343	0.6354	<u>88.2312</u>	<u>87.8107</u>	0.0985 $\pm$ 0.0088
	NegBio-VAE _{MC-G}	0.0451	0.5922	89.7370	88.4573	0.1052 $\pm$ 0.0098
	NegBio-VAE _{MC-C}	0.0373	0.6165	104.3009	103.9739	0.1165 $\pm$ 0.0084
	NegBio-VAE _{DS-G}	0.0447	0.5982	84.2972	83.6357	<u>0.0992</u> $\pm$ 0.0098
	NegBio-VAE _{DS-C}	<u>0.0341</u>	<u>0.6329</u>	92.8698	91.3648	0.1069 $\pm$ 0.0098

5. Experiments

In this section, we compare NegBio-VAE with several well-known VAE variants on four standard benchmark datasets. These experiments are designed to evaluate the effectiveness of our model in terms of reconstruction quality, generative performance, and the expressiveness of the learned latent representations for downstream tasks.

5.1. Experimental Setup

This section introduces the datasets, baselines, metrics, and implementation details.

5.1.1. Datasets, Baselines and Metrics

We assess NegBio-VAE on four widely-used benchmark datasets: **MNIST** [8, 26], **Fashion-MNIST** [50],

CIFAR_{16×16} [25] and **CelebA-64** [29]. The model is compared with representative VAEs using either continuous or discrete latents. Continuous baselines include Gaussian VAE ($\mathcal{G}$ -VAE) [23], Laplace VAE ($\mathcal{L}$ -VAE), while discrete baselines include categorical VAE ($\mathcal{C}$ -VAE) [18] and Poisson VAE ($\mathcal{P}$ -VAE) [46]. We further examine four NegBio-VAE variants: **NegBio-VAE_{MC-G}**, **NegBio-VAE_{MC-C}**, **NegBio-VAE_{DS-G}**, and **NegBio-VAE_{DS-C}**, where **MC** denotes Monte Carlo, **DS** denotes dispersion sharing, and **G** and **C** indicate Gumbel-Softmax and continuous-time reparameterization. It is worth noting that we do not compare against certain strong baselines such as Nouveau VAE (NVAE) [47] and Very Deep VAE [7] as these models are built upon hierarchical latent structures, making a direct comparison with our single-layer NegBio-VAE

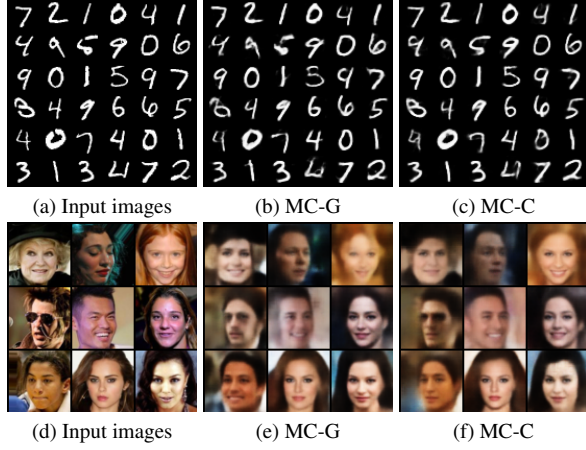


Figure 2. Visual reconstruction results on the MNIST (top) and CelebA-64 (bottom) datasets using the MC-series variants of NegBio-VAE.

unfair. Model performance is evaluated from two perspectives: reconstruction and generation. For reconstruction, mean squared error (MSE) and structural similarity index (SSIM) measure fidelity and structural preservation after latent compression and decoding. For generation, Fréchet Inception Distance (FID) and Kernel Inception Distance (KID) quantify the discrepancy between generated and real data distributions, reflecting sample quality and diversity.

5.1.2. Implementation.

The encoder $NB(\mathbf{z}; \mathbf{r} \odot \delta_r(\mathbf{x}), \mathbf{p} \odot \delta_p(\mathbf{x}))$ is implemented as a neural network that takes $\mathbf{x}$ as input and outputs $\delta_p(\mathbf{x})$, optionally $\delta_r(\mathbf{x})$. The decoder $p_\theta(\mathbf{x} | \mathbf{z})$ is modeled as a Gaussian distribution: $p_\theta(\mathbf{x} | \mathbf{z}) = \mathcal{N}(\mathbf{x}; f_\theta(\mathbf{z}), \sigma^2 \mathbf{I})$, where σ^2 is a hyperparameter. This yields the reconstruction term: $\log p_\theta(\mathbf{x} | \mathbf{z}) = -\frac{1}{2\sigma^2} \|\mathbf{x} - f_\theta(\mathbf{z})\|_2^2 + \text{const}$, which is equivalent to applying a coefficient $\beta = 2\sigma^2$ to the KL term in the ELBO, thereby balancing the trade-off between reconstruction and prior regularization. Unless otherwise specified, all VAEs use convolutional encoders and decoders, with the latent dimensionality fixed at 256.

5.2. Reconstruction

We first evaluate the reconstruction capability of the proposed method (Tab. 1). NegBio-VAE consistently achieves performance comparable to or better than existing single-layer VAE baselines across all datasets. Notably, the MC-C and DS-C variants attain the lowest MSE and highest SSIM on MNIST, Fashion-MNIST, and CIFAR_{16×16}, demonstrating their ability to effectively preserve both structural information and fine-grained image details. On more complex datasets like CelebA-64, NegBio-VAE exhibits slightly higher reconstruction errors, likely due to the stronger regularization introduced by its biologically inspired priors. However, this also yields a more structured latent representation

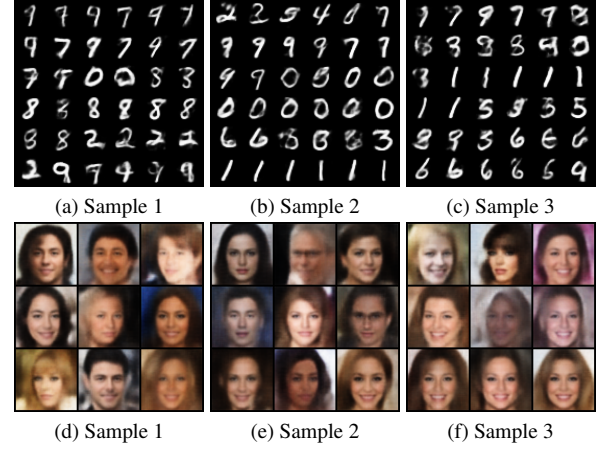


Figure 3. Samples randomly generated from the MNIST (top) and CelebA-64 (bottom) datasets using NegBio-VAE_{MC-G}.

tation (Sec. 5.4). Visual reconstruction results are presented in Fig. 2, with additional results provided in Appendix C.

5.3. Generation

The generative performance of NegBio-VAE is assessed by sampling from the latent space. As shown in Tab. 1, NegBio-VAE significantly outperforms traditional VAEs, with consistently lower FID and KID scores. The MC-G variant shows the largest advantage, attaining the best results across nearly all datasets. For example, reducing FID to 39.8 on CIFAR_{16×16}. These results indicate that the NB latent representation enhances flexibility, capturing richer and more diverse generative patterns. On CelebA-64, NegBio-VAE achieves the lowest FID and nearly the lowest KID compared to all baselines. Although NVAE is originally multi-layer, a single-layer version is used for fairness; even under this constraint, NegBio-VAE surpasses all baselines and could be extended to multi-layer latents to further improve performance. Visual results are shown in Fig. 3, with additional results in Appendix D.

5.4. Latent Analysis

We further evaluate latent representations on downstream tasks using two settings: fragmentation prediction, testing robustness under randomized labels, and few-shot learning, assessing classification with limited samples. All experiments are repeated 10 times, and we report mean performance to support robust comparisons across models. NVAE is excluded because even as a single-layer model, its latents are spatially structured feature maps rather than dense vectors, making them less compatible with standard tasks like classification or clustering.

5.4.1. Fragmentation Prediction

To evaluate robustness and discriminative power of the learned latent representations, we follow the setup in Vafai

Table 2. Evaluation of latent representations on MNIST for the fragmentation prediction task. Higher accuracy indicates more structured and generalizable latent representations. The best and second-best results are marked in **bold** and underlined, respectively.

Latent Dim	Model	Acc $\uparrow$ (N=200)	Acc $\uparrow$ (N=1000)	Acc $\uparrow$ (N=5000)	Acc $\uparrow$ (Shat. Dim.)
100	$\mathcal{G}$ -VAE	0.790 $\pm$ 0.0070	0.914 $\pm$ 0.0020	0.958 $\pm$ 0.0020	0.890 $\pm$ 0.0050
	$\mathcal{L}$ -VAE	<u>0.798</u> $\pm$ 0.0090	<u>0.912</u> $\pm$ 0.0020	0.958 $\pm$ 0.0020	<u>0.892</u> $\pm$ 0.0070
	$\mathcal{C}$ -VAE	0.783 $\pm$ 0.0070	0.896 $\pm$ 0.0030	0.941 $\pm$ 0.0040	0.886 $\pm$ 0.0070
	$\mathcal{P}$ -VAE	0.736 $\pm$ 0.0110	0.888 $\pm$ 0.0020	0.947 $\pm$ 0.0030	0.862 $\pm$ 0.0070
	NegBio-VAE	0.811 $\pm$ 0.0050	<u>0.912</u> $\pm$ 0.0010	<u>0.955</u> $\pm$ 0.0030	0.898 $\pm$ 0.0060

Table 3. Evaluation of latent representations on MNIST and CIFAR for the few-shot learning task. Higher accuracy indicates more structured and generalizable latent representations. The best and second-best results are marked in **bold** and underlined, respectively.

Model	Logistic Regression				k NN			
	Acc $\uparrow$ (1-shot)	Acc $\uparrow$ (5-shot)	Acc $\uparrow$ (10-shot)	Acc $\uparrow$ (20-shot)	Acc $\uparrow$ (1-shot)	Acc $\uparrow$ (5-shot)	Acc $\uparrow$ (10-shot)	Acc $\uparrow$ (20-shot)
MNIST								
$\mathcal{G}$ -VAE	0.409 $\pm$ 0.024	0.664 $\pm$ 0.022	0.736 $\pm$ 0.012	0.788 $\pm$ 0.010	0.228 $\pm$ 0.022	0.527 $\pm$ 0.015	0.653 $\pm$ 0.024	0.756 $\pm$ 0.008
$\mathcal{L}$ -VAE	0.411 $\pm$ 0.024	0.666 $\pm$ 0.025	0.742 $\pm$ 0.012	0.794 $\pm$ 0.012	0.230 $\pm$ 0.036	0.534 $\pm$ 0.014	0.654 $\pm$ 0.026	0.760 $\pm$ 0.010
$\mathcal{C}$ -VAE	<u>0.443</u> $\pm$ 0.034	0.683 $\pm$ 0.030	0.755 $\pm$ 0.011	0.807 $\pm$ 0.012	0.283 $\pm$ 0.032	0.593 $\pm$ 0.018	0.714 $\pm$ 0.013	0.791 $\pm$ 0.011
$\mathcal{P}$ -VAE	0.403 $\pm$ 0.031	<u>0.685</u> $\pm$ 0.030	<u>0.760</u> $\pm$ 0.015	<u>0.838</u> $\pm$ 0.013	0.224 $\pm$ 0.023	0.498 $\pm$ 0.020	0.629 $\pm$ 0.010	0.720 $\pm$ 0.013
NegBio-VAE	0.447 $\pm$ 0.031	0.715 $\pm$ 0.027	0.790 $\pm$ 0.011	0.865 $\pm$ 0.011	<u>0.273</u> $\pm$ 0.020	<u>0.591</u> $\pm$ 0.016	<u>0.710</u> $\pm$ 0.011	<u>0.786</u> $\pm$ 0.011
CIFAR								
$\mathcal{G}$ -VAE	0.142 $\pm$ 0.013	<u>0.206</u> $\pm$ 0.016	0.217 $\pm$ 0.014	0.238 $\pm$ 0.008	0.125 $\pm$ 0.015	0.144 $\pm$ 0.016	0.162 $\pm$ 0.011	0.182 $\pm$ 0.007
$\mathcal{L}$ -VAE	0.138 $\pm$ 0.015	0.202 $\pm$ 0.016	0.213 $\pm$ 0.014	0.235 $\pm$ 0.007	0.124 $\pm$ 0.012	0.134 $\pm$ 0.014	0.151 $\pm$ 0.010	0.174 $\pm$ 0.007
$\mathcal{C}$ -VAE	<u>0.158</u> $\pm$ 0.025	0.190 $\pm$ 0.018	0.223 $\pm$ 0.011	0.240 $\pm$ 0.013	<u>0.131</u> $\pm$ 0.018	<u>0.176</u> $\pm$ 0.015	<u>0.194</u> $\pm$ 0.010	<u>0.216</u> $\pm$ 0.009
$\mathcal{P}$ -VAE	0.154 $\pm$ 0.020	0.203 $\pm$ 0.016	<u>0.244</u> $\pm$ 0.013	<u>0.261</u> $\pm$ 0.012	0.120 $\pm$ 0.012	0.173 $\pm$ 0.015	0.188 $\pm$ 0.013	0.205 $\pm$ 0.010
NegBio-VAE	0.167 $\pm$ 0.023	0.221 $\pm$ 0.016	0.255 $\pm$ 0.011	0.266 $\pm$ 0.010	0.133 $\pm$ 0.024	0.192 $\pm$ 0.014	0.207 $\pm$ 0.012	0.233 $\pm$ 0.012

et al. [46], using MNIST with a fixed latent dimensionality of 100 and convolutional encoder-decoders for all models. We randomly split the test set into two sets of 5,000 samples each and train logistic regression classifiers using $N = 200, 1000, 5000$ labeled samples from one set. We then report accuracy on the other set (Tab. 2). For NegBio-VAE, we use the variant of DS-C. We also assess latent space structure via empirical shattering dimensionality [3, 21, 36, 46], defined as the average binary classification accuracy over disjoint class partitions using linear classifiers. As shown in Tab. 2, all models exhibit improved performance with increasing training samples. NegBio-VAE consistently ranks first or second, achieving 0.811 accuracy at $N = 200$ and 0.898 under the shattering metric. This demonstrates that NB latents enhance separability and robustness even under severe label perturbations, whereas conventional models like $\mathcal{C}$ -VAE are less resilient.

5.4.2. Few-shot Learning

We further evaluate the adaptability of the learned representations in low-data regimes through few-shot learning on MNIST and CIFAR_{16×16}. For each dataset, we use a k -shot setup with $k \in \{1, 5, 10, 20\}$, sampling k labeled samples per class for training and evaluating on the test set. Two lightweight classifiers—logistic regression and

k -nearest neighbors (k NN)—are used to assess the effectiveness of the latent representations. As shown in Tab. 3, NegBio-VAE consistently ranks first or second across all k values. On MNIST, it achieves the highest accuracy with logistic regression, with the gap widening as labeled samples increase, e.g., 0.865 at 20-shot v.s. 0.838 for the best baseline. On CIFAR_{16×16}, NegBio-VAE similarly maintains superior performance, e.g., 0.266 at 20-shot v.s. 0.261 for the best baseline. These results show that NegBio-VAE learns more discriminative and transferable representations, enabling accurate classification with minimal supervision and robust transfer across datasets of varying complexity.

5.5. Ablation Studies

To further analyze NegBio-VAE, we perform ablation studies on MNIST, investigating the impact of the encoder-decoder design, β scaling, and the number of Monte Carlo samples during training.

5.5.1. Encoder-Decoder Architectures

We compare encoder-decoder architectures using linear, multilayer perceptron (MLP), and convolutional networks. Results for linear encoders are shown in Fig. 4a, with further analyses in Appendix E. With a linear encoder, decoder choice strongly affects performance: MLP decoders achieve

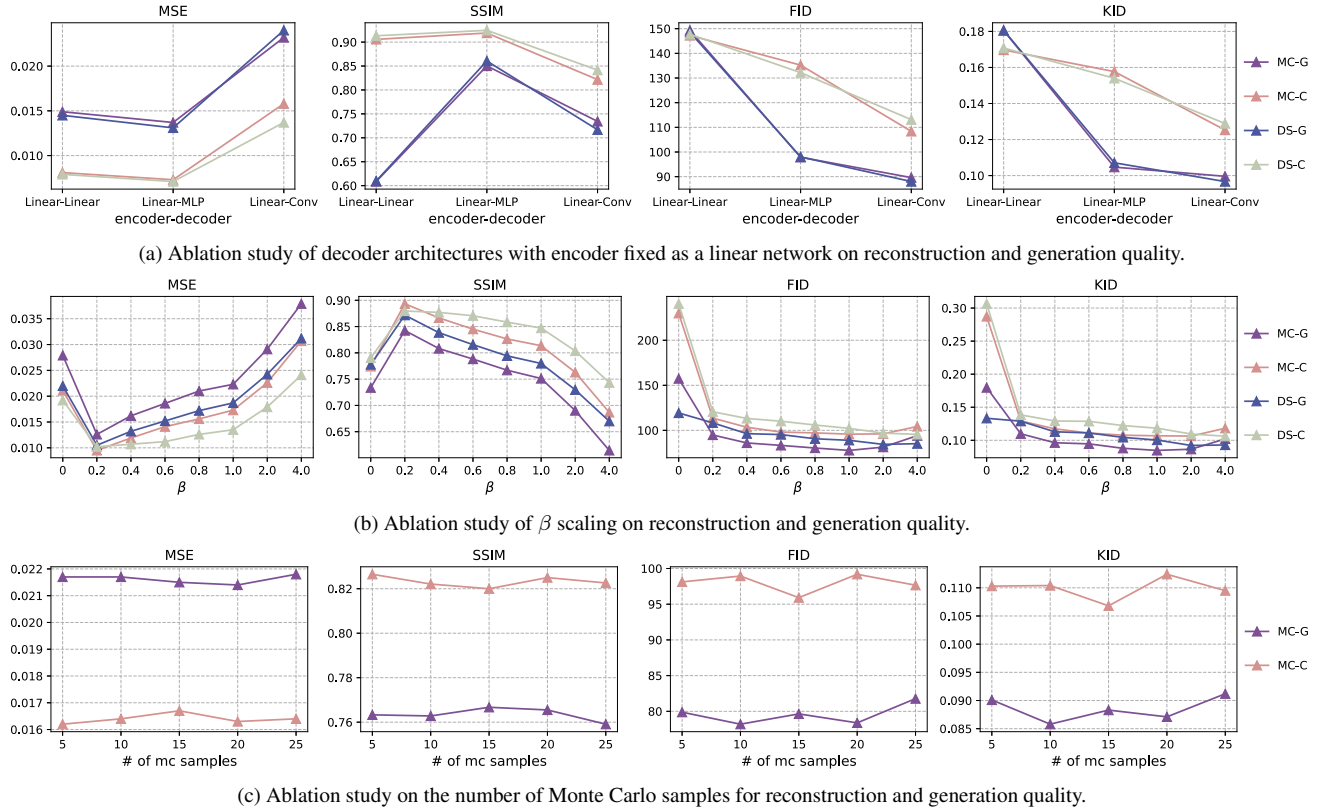


Figure 4. Ablation studies on decoder design, regularization strength, and the number of Monte Carlo samples in NegBio-VAE.

the lowest MSE and highest SSIM, convolutional decoders yield the best FID and KID, and linear decoders perform worst. These results highlight that enhancing decoder capacity, via nonlinear or convolutional architectures, is crucial for both reconstruction accuracy and generative quality.

5.5.2. Effect of β Scaling

We investigate the effect of the scaling factor β on NegBio-VAE performance (Fig. 4b). Small β values (0.2–0.4) yield the best reconstruction (lowest MSE, highest SSIM), especially for MC-C and DS-C variants. As β increases, FID improves and peaks around $\beta = 1.0$, indicating better generative fidelity. Beyond $\beta \geq 2.0$, reconstruction degrades and generative gains diminish. Overall, smaller β favors reconstruction, larger β favors generation, and intermediate values (≈ 0.6 – 1.0) provide the best trade-off.

5.5.3. Effect of Number of MC Samples

We examine the impact of the number of MC samples on model performance (Fig. 4c). As the number of MC samples increases from 5 to 25, all metrics remain relatively stable for both MC-G and MC-C variants, indicating that the proposed model is robust to the sampling variance. Specifically, MC-C achieves lower MSE and higher SSIM (better reconstruction), while MC-G attains lower FID and KID

(better generation). These results demonstrate that while increasing the number of MC samples provides only marginal gains, the model achieves a favorable balance between reconstruction and generation quality even with a few samples, confirming the effectiveness of our sampling strategy.

6. Conclusions

In this work, we presented NegBio-VAE, a generative model leveraging the NB distribution to capture overdispersed latent variables. By introducing a dispersion parameter, it extends beyond standard Poisson assumptions with minimal modification. Despite its simplicity, NegBio-VAE improves reconstruction and generation quality across benchmark datasets and outperforms existing VAE baselines in fidelity, generative quality, and the utility of latent representations for downstream tasks. While NegBio-VAE introduces greater flexibility in modeling overdispersed spike counts, the design choices, such as KL estimation and reparameterization, affect training and representations, a deeper theoretical understanding of these trade-offs remains open. Future work will explore adaptive reparameterization strategies based on data characteristics, and extend the framework to hierarchical latent structures similar to NVAE to further enhance model expressiveness.

Acknowledgments

This work was supported by the NSFC Projects (Nos. 62506069, 62576346), the MOE Project of Key Research Institute of Humanities and Social Sciences (22JJD110001), the fundamental research funds for the central universities, and the research funds of Renmin University of China (24XNKJ13), and Beijing Advanced Innovation Center for Future Blockchain and Privacy Computing.

References

- [1] Michael A Arbib. *The handbook of brain theory and neural networks*. MIT press, 2003. 1
- [2] Wyeth Bair and Christof Koch. Temporal precision of spike trains in extrastriate cortex of the behaving macaque monkey. *Neural computation*, 8(6):1185–1202, 1996. 1
- [3] Silvia Bernardi, Marcus K. Benna, Mattia Rigotti, Jérôme Munuera, Stefano Fusi, and C. Daniel Salzman. The geometry of abstraction in the hippocampus and prefrontal cortex. *Cell*, 183(4):954–967.e21, 2020. 7
- [4] Michael J. Berry, David K. Warland, and Markus Meister. The structure and precision of retinal spike trains. *Proceedings of the National Academy of Sciences*, 94(10):5411–5416, 1997. 1
- [5] Jiahang Cao, Ziqing Wang, Hanzhong Guo, Hao Cheng, Qiang Zhang, and Renjing Xu. Spiking denoising diffusion probabilistic models. In *Proceedings of the IEEE/CVF winter conference on applications of computer vision*, pages 4912–4921, 2024. 2
- [6] Xiang Cheng, Yunzhe Hao, Jiaming Xu, and Bo Xu. Lissn: Improving spiking neural networks with lateral interactions for robust object recognition. In *IJCAI*, pages 1519–1525. Yokohama, 2020. 2
- [7] Rewon Child. Very deep {vae}s generalize autoregressive models and can outperform them on images. In *International Conference on Learning Representations*, 2021. 5
- [8] Li Deng. The mnist database of handwritten digit images for machine learning research. *IEEE Signal Processing Magazine*, 29(6):141–142, 2012. 5
- [9] Yanming Di, Daniel W Schafer, Jason S Cumbie, and Jeff H Chang. The nbp negative binomial model for assessing differential gene expression from rna-seq. *Statistical applications in genetics and molecular biology*, 10(1), 2011. 3
- [10] Emilien Dupont. Learning disentangled joint continuous and discrete representations. *Advances in neural information processing systems*, 31, 2018. 2
- [11] Wei Fang, Zhaofer Yu, Yanqi Chen, Timothée Masquelier, Tiejun Huang, and Yonghong Tian. Incorporating learnable membrane time constant to enhance learning of spiking neural networks. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 2661–2671, 2021. 2
- [12] Linghao Feng, Dongcheng Zhao, and Yi Zeng. Spiking generative adversarial network with attention scoring decoding. *Neural Networks*, 178:106423, 2024. 2
- [13] Mikhail Figurnov, Shakir Mohamed, and Andriy Mnih. Implicit reparameterization gradients. In *Advances in Neural Information Processing Systems*. Curran Associates, Inc., 2018. 4
- [14] Vincent Fortuin, Matthias Hüser, Francesco Locatello, Heiko Strathmann, and Gunnar Rätsch. Som-vae: Interpretable discrete representation learning on time series. In *International Conference on Learning Representations*, 2019. 2
- [15] Samanwoy Ghosh-Dastidar and Hojjat Adeli. Spiking neural networks. *International journal of neural systems*, 19(04): 295–308, 2009. 1, 2
- [16] Tim Gollisch and Markus Meister. Rapid neural coding in the retina with relative spike latencies. *science*, 319(5866): 1108–1111, 2008. 1
- [17] Robbe L. T. Goris, J. Anthony Movshon, and Eero P. Simoncelli. Partitioning neuronal variability. *Nature Neuroscience*, 17:858–865, 2014. 1
- [18] Eric Jang, Shixiang Gu, and Ben Poole. Categorical reparameterization with gumbel-softmax. In *International Conference on Learning Representations*, 2017. 1, 2, 4, 5
- [19] Hiromichi Kamata, Yusuke Mukuta, and Tatsuya Harada. Fully spiking variational autoencoder. In *Proceedings of the AAAI conference on artificial intelligence*, pages 7059–7067, 2022. 2
- [20] Jaivardhan Kapoor, Auguste Schulz, Julius Vetter, Felix Pei, Richard Gao, and Jakob H Macke. Latent diffusion for neural spiking data. *Advances in Neural Information Processing Systems*, 37:118119–118154, 2024. 2
- [21] Matthew T. Kaufman, Marcus K. Benna, Mattia Rigotti, Fabio Stefanini, Stefano Fusi, and Anne K. Churchland. The implications of categorical and category-free mixed selectivity on representational geometries. *Current Opinion in Neurobiology*, 77:102644, 2022. 7
- [22] Diederik P Kingma and Max Welling. Auto-encoding variational bayes, 2014. 2
- [23] Diederik P Kingma, Max Welling, et al. Auto-encoding variational bayes, 2013. 1, 5
- [24] Vineet Kotariya and Udayan Ganguly. Spiking-gan: A spiking generative adversarial network using time-to-first-spike coding. In *2022 International Joint Conference on Neural Networks (IJCNN)*, pages 1–7. IEEE, 2022. 2
- [25] Alex Krizhevsky. Learning multiple layers of features from tiny images. Technical report, 2009. 5
- [26] Yann LeCun, Corinna Cortes, and CJ Burges. Mnist handwritten digit database. *ATT Labs [Online]*. Available: <http://yann.lecun.com/exdb/mnist>, 2, 2010. 5
- [27] Yuhang Li, Yufei Guo, Shanghang Zhang, Shikuang Deng, Yongqing Hai, and Shi Gu. Differentiable spike: Rethinking gradient-descent for training spiking neural networks. *Advances in neural information processing systems*, 34:23426–23439, 2021. 2
- [28] Mingxuan Liu, Jie Gan, Rui Wen, Tao Li, Yongli Chen, and Hong Chen. Spiking-diffusion: Vector quantized discrete diffusion model with spiking neural networks. In *2024 5th International Conference on Computer, Big Data and Artificial Intelligence (ICCBDAI)*, pages 627–631. IEEE, 2024. 2
- [29] Ziwei Liu, Ping Luo, Xiaogang Wang, and Xiaoou Tang. Deep learning face attributes in the wild. In *2015 IEEE In-*

- ternational Conference on Computer Vision (ICCV), pages 3730–3738, 2015. [5](#)
- [30] Zachary F Mainen and Terrence J Sejnowski. Reliability of spike timing in neocortical neurons. *Science*, 268(5216):1503–1506, 1995. [1](#)
- [31] Joseph Marino. Predictive coding, variational autoencoders, and biological connections. *Neural Computation*, 34(1):1–44, 2022. [1](#)
- [32] Dina Moshitch and Israel Nelken. Using tweedie distributions for fitting spike count data. *Journal of neuroscience methods*, 225:13–28, 2014. [1](#)
- [33] Donald H Perkel, George L Gerstein, and George P Moore. Neuronal spike trains and stochastic point processes: II. simultaneous spike trains. *Biophysical journal*, 7(4):419–440, 1967. [1](#)
- [34] Jonathan Pillow and James Scott. Fully bayesian inference for neural models with negative-binomial spiking. *Advances in neural information processing systems*, 25, 2012. [3](#)
- [35] Daniil Polykovskiy and Dmitry Vetrov. Deterministic decoding for discrete data in variational autoencoders. In *International conference on artificial intelligence and statistics*, pages 3046–3056. PMLR, 2020. [2](#)
- [36] Mattia Rigotti, Omri Barak, Melissa R. Warden, Xiao-Jing Wang, Nathaniel D. Daw, Earl K. Miller, and Stefano Fusi. The importance of mixed selectivity in complex cognitive tasks. *Nature*, 497:585–590, 2013. [7](#)
- [37] Jason Tyler Rolfe. Discrete variational autoencoders. In *International Conference on Learning Representations*, 2017. [2](#)
- [38] Bleema Rosenfeld, Osvaldo Simeone, and Bipin Rajendran. Spiking generative adversarial networks with a neural network discriminator: Local training, bayesian models, and continual meta-learning. *IEEE Transactions on Computers*, 71(11):2778–2791, 2022. [2](#)
- [39] GJS Ross and DA Preece. The negative binomial distribution. *Journal of the Royal Statistical Society: Series D (The Statistician)*, 34(3):323–335, 1985. [1](#)
- [40] Michael N. Shadlen and William T. Newsome. The variable discharge of cortical neurons: Implications for connectivity, computation, and information coding. *Journal of Neuroscience*, 18(10):3870–3896, 1998. [1](#)
- [41] Ian H Stevenson. Flexible models for spike count data with both over-and under-dispersion. *Journal of computational neuroscience*, 41:29–43, 2016. [1](#)
- [42] Katherine R Storrs, Barton L Anderson, and Roland W Fleming. Unsupervised learning predicts human perception and misperception of gloss. *Nature Human Behaviour*, 5(10):1402–1417, 2021. [1](#)
- [43] Wahiba Taouali, Giacomo Benvenuti, Pascal Wallisch, Frédéric Chavane, and Laurent U Perrinet. Testing the odds of inherent vs. observed overdispersion in neural spike counts. *Journal of neurophysiology*, 115(1):434–444, 2016. [1](#)
- [44] Amirhossein Tavanaei, Masoud Ghodrati, Saeed Reza Kheradpisheh, Timothée Masquelier, and Anthony Maida. Deep learning in spiking neural networks. *Neural networks*, 111:47–63, 2019. [1](#)
- [45] Hadi Vafaii, Jacob Yates, and Daniel Butts. Hierarchical vaes provide a normative account of motion processing in the primate brain. *Advances in Neural Information Processing Systems*, 36:46152–46190, 2023. [1](#)
- [46] Hadi Vafaii, Dekel Galor, and Jacob Yates. Poisson variational autoencoder. *Advances in Neural Information Processing Systems*, 37:44871–44906, 2024. [1](#), [2](#), [4](#), [5](#), [7](#), [3](#)
- [47] Arash Vahdat and Jan Kautz. Nvae: A deep hierarchical variational autoencoder. In *Advances in Neural Information Processing Systems*, pages 19667–19679. Curran Associates, Inc., 2020. [5](#)
- [48] Gido M Van de Ven, Hava T Siegelmann, and Andreas S Tolias. Brain-inspired replay for continual learning with artificial neural networks. *Nature communications*, 11(1):4069, 2020. [1](#)
- [49] Aaron Van Den Oord, Oriol Vinyals, et al. Neural discrete representation learning. *Advances in neural information processing systems*, 30, 2017. [1](#), [2](#)
- [50] Han Xiao, Kashif Rasul, and Roland Vollgraf. Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms. *arXiv preprint arXiv:1708.07747*, 2017. [5](#)
- [51] Srishti Yadav, Anshul Pundhir, Tanish Goyal, Balasubramanian Raman, and Sanjeev Kumar. Differentially private spiking variational autoencoder. In *International Conference on Pattern Recognition*, pages 96–112. Springer, 2025. [2](#)
- [52] Qiugang Zhan, Ran Tao, Xiurui Xie, Guisong Liu, Malu Zhang, Huajin Tang, and Yang Yang. Esvae: An efficient spiking variational autoencoder with reparameterizable poisson spiking sampling. *arXiv preprint arXiv:2310.14839*, 2023. [2](#)
- [53] He Zhao, Piyush Rai, Lan Du, Wray Buntine, Dinh Phung, and Mingyuan Zhou. Variational autoencoders for sparse and overdispersed discrete data. In *International conference on artificial intelligence and statistics*, pages 1684–1694. PMLR, 2020. [2](#)
- [54] Hanle Zheng, Yujie Wu, Lei Deng, Yifan Hu, and Guoqi Li. Going deeper with directly-trained larger spiking neural networks. In *Proceedings of the AAAI conference on artificial intelligence*, pages 11062–11070, 2021. [2](#)
- [55] Mingyuan Zhou and Lawrence Carin. Negative binomial process count and mixture modeling. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 37(2):307–320, 2013. [3](#)