
Adaptive Density Estimation for Generative Models

Thomas Lucas (*)
Univ. Grenoble Alpes,
Inria, CNRS,
Grenoble INP, LJK
38000 Grenoble, France

Konstantin Shmelkov[†] (*)
Noah’s Ark Lab
Huawei

Cordelia Schmid
Univ. Grenoble Alpes,
Inria, CNRS,
Grenoble INP, LJK
38000 Grenoble, France

Karteek Alahari
Univ. Grenoble Alpes,
Inria, CNRS,
Grenoble INP, LJK
38000 Grenoble, France

Jakob Verbeek
Univ. Grenoble Alpes,
Inria, CNRS,
Grenoble INP, LJK
38000 Grenoble, France

Abstract

Unsupervised learning of generative models has seen tremendous progress over recent years, in particular due to generative adversarial networks (GANs), variational autoencoders, and flow-based models. GANs have dramatically improved sample quality, but suffer from two drawbacks: (i) they mode-drop, *i.e.*, do not cover the full support of the train data, and (ii) they do not allow for likelihood evaluations on held-out data. In contrast likelihood-based training encourages models to cover the full support of the train data, but yields poorer samples. These mutual shortcomings can in principle be addressed by training generative latent variable models in a hybrid adversarial-likelihood manner. However, we show that commonly made parametric assumptions create a conflict between them, making successful hybrid models non trivial. As a solution, we propose to use deep invertible transformations in the latent variable decoder. This approach allows for likelihood computations in image space, is more efficient than fully invertible models, and can take full advantage of adversarial training. We show that our model significantly improves over existing hybrid models: offering GAN-like samples, IS and FID scores that are competitive with fully adversarial models, and improved likelihood scores.

1 Introduction

Successful recent generative models of natural images can be divided into two broad families, which are trained in fundamentally different ways. The first is trained using likelihood-based criteria, which ensure that all training data points are well covered by the model. This category includes variational autoencoders (VAEs) [26, 27, 40, 41], autoregressive models such as PixelCNNs [47, 54], and flow-based models such as Real-NVP [9, 21, 25]. The second category is trained based on a signal that measures to what extent (statistics of) samples from the model can be distinguished from (statistics of) the training data, *i.e.*, based on the quality of samples drawn from the model. This is the case for generative adversarial networks (GANs) [2, 15, 23], and moment matching methods [29].

Despite tremendous recent progress, existing methods exhibit a number of drawbacks. Adversarially trained models such as GANs do not provide a density function, which poses a fundamental problem as it prevents assessment of how well the model fits held out and training data. Moreover, adversarial models typically do not allow to infer the latent variables that underlie observed images. Finally,

(*) The authors contributed equally. (†) Work done while at Inria.

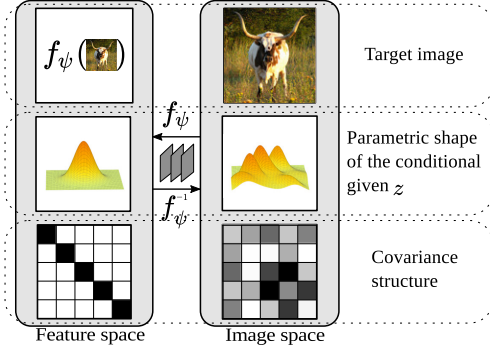


Figure 1: An invertible non-linear mapping f_ψ maps an image x to a vector $f_\psi(x)$ in feature space. f_ψ is trained to adapt to modelling assumptions made by a trained density p_θ in feature space. This induces full covariance structure and a non-Gaussian density.

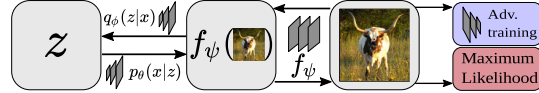


Figure 2: Variational inference is used to train a latent variable generative model in feature space. The invertible mapping f_ψ maps back to image space, where adversarial training can be performed together with MLE.

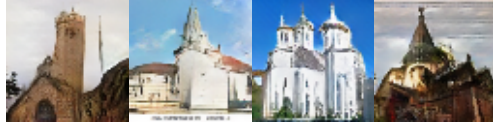


Figure 3: Our model yields compelling samples while the optimization of likelihood ensures coverage of all modes in the training support and thus sample diversity, here on LSUN churches (64×64).

adversarial models suffer from mode collapse [2], *i.e.*, they do not cover the full support of the training data. Likelihood-based model on the other hand are trained to put probability mass on all elements of the training set, but over-generalise and produce samples of substantially inferior quality as compared to adversarial models. The models with the best likelihood scores on held-out data are autoregressive models [36], which suffer from the additional problem that they are extremely inefficient to sample from [39], since images are generated pixel-by-pixel. The sampling inefficiency makes adversarial training of such models prohibitively expensive.

In order to overcome these shortcomings, we seek to design a model that (i) generates high-quality samples typical of adversarial models, (ii) provides a likelihood measure on the entire image space, and (iii) has a latent variable structure to allow for efficient sampling, that permits adversarial training. Additionally we show that, (iv) a successful hybrid adversarial-likelihood paradigm requires going beyond simplifying conditional independence assumptions commonly used with likelihood based latent variable models. These simplifying assumptions on the conditional distribution on data x given latents z , $p(x|z)$, include full independence across the dimensions of x and/or simple parametric forms such as Gaussian [26], or use fully invertible networks [9, 25]. These assumptions create a conflict between achieving high sample quality and high likelihood scores on held-out data. Autoregressive models, such as pixelCNNs [47, 54], do not make factorization assumptions, but are extremely inefficient to sample from. As a solution, we propose learning a non-linear invertible function f_ψ between the image space and an abstract feature space as illustrated in Figure 1. Training a model with full support in this feature space induces a model in the image space that does not make Gaussianity or independence assumptions in the conditional $p(x|z)$. Trained by MLE, f_ψ adapts to modelling assumptions made by p_θ so we refer to this approach as “adaptive density estimation”.

We experimentally validate our approach on the CIFAR-10 dataset with an ablation study. Our model significantly improves over existing hybrid models, producing GAN-like samples, and IS and FID scores that are competitive with fully adversarial models. At the same time, we obtain likelihoods on held-out data comparable to state-of-the-art likelihood-based methods which requires covering the full support of the dataset. We further confirm these observations with quantitative and qualitative experimental results on the STL-10, ImageNet and LSUN datasets.

2 Related work

Mode-collapse in GANs has received considerable attention, and stabilizing the training process as well as improved and bigger architectures have been shown to alleviate this issue [2, 18, 38]. Another line of work focuses on allowing the discriminator to access batch statistics of generated images, as pioneered by [23, 46], and further generalized by [30, 33]. This enables comparison of distributional statistics by the discriminator rather than only individual samples. Other approaches to encourage diversity among GAN samples include the use of maximum mean discrepancy [1], optimal transport [48], determinantal point processes [14] and bayesian formulations of adversarial training [43] that allow model parameters to be sampled. In contrast to our work, these models lack an inference network, and do not define an explicit density over the full data support.

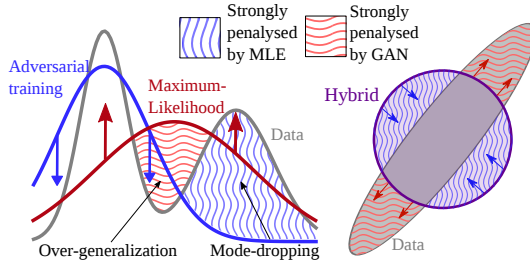


Figure 4: (Left) Maximum likelihood training pulls probability mass towards high-density regions of the data distribution, while adversarial training pushes mass out of low-density regions. (Right) Independence assumptions become a source of conflict in a joint training setting, making hybrid training non-trivial.

An other line of research has explored inference mechanisms for GANs. The discriminator of BiGAN [10] and ALI [12], given pairs (x, z) of images and latents, predict if z was encoded from a real image, or if x was decoded from a sampled z . In [53] the encoder and the discriminator are collapsed into one network that encodes both real images and generated samples, and tries to spread their posteriors apart. In [6] a symmetrized KL divergence is approximated in an adversarial setup, and uses reconstruction losses to improve the correspondence between reconstructed and target variables for x and z . Similarly, in [42] a discriminator is used to replace the KL divergence term in the variational lower bound used to train VAEs with the density ratio trick. In [34] the KL divergence term in a VAE is replaced with a discriminator that compares latent variables from the prior and the posterior in a more flexible manner. This regularization is more flexible than the standard KL divergence. The VAE-GAN model [28] and the model in [22] use the intermediate feature maps of a GAN discriminator and of a classifier respectively, as target space for a VAE. Unlike ours, these methods do not define a likelihood over the image space.

Likelihood-based models typically make modelling assumptions that conflict with adversarial training, these include strong factorization and/or Gaussianity. In our work we avoid these limitations by learning the shape of the conditional density on observed data given latents, $p(x|z)$, beyond fully factorized Gaussian models. As in our work, Flow-GAN [16] also builds on invertible transformations to construct a model that can be trained in a hybrid adversarial-MLE manner, see Figure 2. However, Flow-GAN does not use efficient non-invertible layers we introduce, and instead relies entirely on invertible layers. Other approaches combine autoregressive decoders with latent variable models to go beyond typical parametric assumptions in pixel space [7, 19, 32]. They, however, are not amenable to adversarial training due to the prohibitively slow sequential pixel sampling.

3 Preliminaries on MLE and adversarial training

Maximum-likelihood and over-generalization. The de-facto standard approach to train generative models is maximum-likelihood estimation. It maximizes the probability of data sampled from an unknown data distribution p^* under the model p_θ w.r.t. the model parameters θ . This is equivalent to minimizing the Kullback-Leibler (KL) divergence, $\mathcal{D}_{\text{KL}}(p^* || p_\theta)$, between p^* and p_θ . This yields models that tend to cover all the modes of the data, but put mass in spurious regions of the target space; a phenomenon known as “over-generalization” or “zero-avoiding” [4], and manifested by unrealistic samples in the context of generative image models, see Figure 4. Over-generalization is inherent to the optimization of the KL divergence oriented in this manner. Real images are sampled from p^* , and p_θ is explicitly optimized to cover all of them. The training procedure, however, does not sample from p_θ to evaluate the quality of such samples (ideally using the inaccessible $p^*(x)$ as a score). Therefore p_θ may put mass in spurious regions of the space without being heavily penalized. We refer to this kind of training procedure as “coverage-driven training” (CDT). This optimizes a loss of the form $\mathcal{L}_C(p_\theta) = \int_x p^*(x) s_c(x, p_\theta) dx$, where $s_c(x, p_\theta) = \ln p_\theta(x)$ evaluates how well a sample x is covered by the model. Any score s_c that verifies: $\mathcal{L}_C(p_\theta) = 0 \iff p_\theta = p^*$, is equivalent to the log-score, which forms a justification for MLE on which we focus.

Explicitly evaluating sample quality is redundant in the regime of unlimited model capacity and training data. Indeed, putting mass on spurious regions takes it away from the support of p^* , and thus reduces the likelihood of the training data. In practice, however, datasets and model capacity are finite, and models *must* put mass outside the finite training set in order to generalize. The maximum likelihood criterion, by construction, only measures *how much* mass goes off the training data, not *where* it goes. In classic MLE, generalization is controlled in two ways: (i) inductive bias, in the form of model architecture, controls *where* the off-dataset mass goes, and (ii) regularization controls to which extent this happens. An adversarial loss, by considering samples of the model p_θ , can provide

a second handle to evaluate and control where the off-dataset mass goes. In this sense, and in contrast to model architecture design, an adversarial loss provides a “trainable” form of inductive bias.

Adversarial models and mode collapse. Adversarially trained models produce samples of excellent quality. As mentioned, their main drawbacks are their tendency to “mode-drop”, and the lack of measure to assess mode-dropping, or their performance in general. The reasons for this are two-fold. First, defining a valid likelihood requires adding volume to the low-dimensional manifold learned by GANs to define a density under which training and test data have non-zero density. Second, computing the density of a data point under the defined probability distribution requires marginalizing out the latent variables, which is not trivial in the absence of an efficient inference mechanism. When a human expert subjectively evaluates the quality of generated images, samples from the model are compared to the expert’s implicit approximation of p^* . This type of objective may be formalized as $\mathcal{L}_Q(p_\theta) = \int_x p_\theta(x) s_q(x, p^*) dx$, and we refer to it as “quality-driven training” (QDT). To see that GANs [15] use this type of training, recall that the discriminator is trained with the loss $\mathcal{L}_{\text{GAN}} = \int_x p^*(x) \ln D(x) + p_\theta(x) \ln(1 - D(x)) dx$. It is easy to show that the optimal discriminator equals $D^*(x) = p^*(x)/(p^*(x) + p_\theta(x))$. Substituting the optimal discriminator, $\mathcal{L}_{\text{GAN}}$ equals the Jensen-Shannon divergence,

$$\mathcal{D}_{\text{JS}}(p^*||p_\theta) = \frac{1}{2} \mathcal{D}_{\text{KL}}(p^*||\frac{1}{2}(p_\theta + p^*)) + \frac{1}{2} \mathcal{D}_{\text{KL}}(p_\theta||\frac{1}{2}(p_\theta + p^*)), \quad (1)$$

up to additive and multiplicative constants [15]. This loss, approximated by the discriminator, is symmetric and contains two KL divergence terms. Note that $\mathcal{D}_{\text{KL}}(p^*||\frac{1}{2}(p_\theta + p^*))$ is an integral on p^* , so *coverage driven*. The term that approximates it in $\mathcal{L}_{\text{GAN}}$, *i.e.*, $\int_x p^*(x) \ln D(x)$, is however independent from the generative model, and disappears when differentiating. Therefore, it cannot be used to perform coverage-driven training, and the generator is trained to minimize $\ln(1 - D(G(z)))$ (or to maximize $\ln D(G(z))$), where $G(z)$ is the deterministic generator that maps latent variables z to the data space. Assuming $D = D^*$, this yields

$$\int_z p(z) \ln(1 - D^*(G(z))) dz = \int_x p_\theta(x) \ln \frac{p_\theta(x)}{p_\theta(x) + p^*(x)} dx = \mathcal{D}_{\text{KL}}(p_\theta||(\frac{1}{2}(p_\theta + p^*))), \quad (2)$$

which is a quality-driven criterion, favoring sample quality over support coverage.

4 Adaptive Density Estimation and hybrid adversarial-likelihood training

We present a hybrid training approach with MLE to cover the full support of the training data, and adversarial training as a trainable inductive bias mechanism to improve sample quality. Using both these criteria provides a richer training signal, but satisfying both criteria is more challenging than each in isolation for a given model complexity. In practice, model flexibility is limited by (i) the number of parameters, layers, and features in the model, and (ii) simplifying modeling assumptions, usually made for tractability. We show that these simplifying assumptions create a conflict between the two criteria, making successful joint training non trivial. We introduce Adaptive Density Estimation as a solution to reconcile them.

Latent variable generative models, defined as $p_\theta(x) = \int_z p_\theta(x|z)p(z) dz$, typically make simplifying assumptions on $p_\theta(x|z)$, such as full factorization and/or Gaussianity, see *e.g.* [11, 26, 31]. In particular, assuming full factorization of $p_\theta(x|z)$ implies that any correlations not captured by z are treated as independent per-pixel noise. This is a poor model for natural images, unless z captures each and every aspect of the image structure. Crucially, this hypothesis is problematic in the context of hybrid MLE-adversarial training. If p^* is too complex for $p_\theta(x|z)$ to fit it accurately enough, MLE will lead to a high variance in a factored (Gaussian) $p_\theta(x|z)$ as illustrated in Figure 4 (right).

This leads to unrealistic blurry samples, easily detected by an adversarial discriminator, which then does not provide a useful training signal. Conversely, adversarial training will try to avoid these poor samples by dropping modes of the training data, and driving the “noise” level to zero. This in turn is heavily penalized by maximum likelihood training, and leads to poor likelihoods on held-out data.

Adaptive density estimation. The point of view of regression hints at a possible solution. For instance, with isotropic Gaussian model densities with fixed variance, solving the optimization problem $\theta^* \in \arg \max_\theta \ln(p_\theta(x))$ is similar to solving $\min_\theta \|\mu_\theta(z) - x\|_2$, *i.e.*, ℓ_2 regression, where $\mu_\theta(z)$ is the mean of the decoder $p_\theta(x|z)$. The Euclidean distance in RGB space is known to

be a poor measure of similarity between images, non-robust to small translations or other basic transformations [35]. One can instead compute the Euclidean distance in a feature space, $\|f_\psi(x_1) - f_\psi(x_2)\|_2$, where f_ψ is chosen so that the distance is a better measure of similarity. A popular way to obtain f_ψ is to use a CNN that learns a non-linear image representation, that allows linear assessment of image similarity. This is the idea underlying GAN discriminators, the FID evaluation measure [20], the reconstruction losses of VAE-GAN [28] and classifier based perceptual losses as in [22].

Despite their flexibility, such similarity metrics are in general degenerate in the sense that they may discard information about the data point x . For instance, two different images x and y can collapse to the same points in feature space, i.e., $f_\psi(x) = f_\psi(y)$. This limits the use of similarity metrics in the context of generative modeling for two reasons: (i) it does not yield a valid measure of likelihood over inputs, and (ii) points generated in the feature space f_ψ cannot easily be mapped to images. To resolve this issue, we chose f_ψ to be a bijection. Given a model p_θ trained to model $f_\psi(x)$ in feature space, a density in image space is computed using the change of variable formula, which yields $p_{\theta,\psi}(x) = p_\theta(f_\psi(x)) |\det(\partial f_\psi(x)/\partial x^\top)|$. Image samples are obtained by sampling from p_θ in feature space, and mapping to the image space through f_ψ^{-1} . We refer to this construction as Adaptive Density Estimation. If p_θ provides efficient log-likelihood computations, the change of variable formula can be used to train f_ψ and p_θ together by maximum-likelihood, and if p_θ provides fast sampling adversarial training can be performed efficiently.

MLE with adaptive density estimation. To train a generative latent variable model $p_\theta(x)$ which permits efficient sampling, we rely on amortized variational inference. We use an inference network $q_\phi(z|x)$ to construct a variational evidence lower-bound (ELBO),

$$\mathcal{L}_{\text{ELBO}}^\psi(x, \theta, \phi) = \mathbb{E}_{q_\phi(z|x)} [\ln(p_\theta(f_\psi(x)|z))] - \mathcal{D}_{\text{KL}}(q_\phi(z|x)||p_\theta(z)) \leq \ln p_\theta(f_\psi(x)). \quad (3)$$

Using this lower bound together with the change of variable formula, the mapping to the similarity space f_ψ and the generative model p_θ can be trained jointly with the loss

$$\mathcal{L}_C(\theta, \phi, \psi) = \mathbb{E}_{x \sim p^*} \left[-\mathcal{L}_{\text{ELBO}}^\psi(x, \theta, \phi) - \ln \left| \det \frac{\partial f_\psi(x)}{\partial x^\top} \right| \right] \geq -\mathbb{E}_{x \sim p^*} [\ln p_{\theta,\psi}(x)]. \quad (4)$$

We use gradient descent to train f_ψ by optimizing $\mathcal{L}_C(\theta, \phi, \psi)$ w.r.t. ψ . The $\mathcal{L}_{\text{ELBO}}$ term encourages the mapping f_ψ to maximize the density of points in feature space under the model p_θ , so that f_ψ is trained to match modeling assumptions made in p_θ . Simultaneously, the log-determinant term encourages f_ψ to maximize the volume of data points in feature space. This guarantees that data points cannot be collapsed to a single point in the feature space. We use a factored Gaussian form of the conditional $p_\theta(\cdot|z)$ for tractability, but since f_ψ can arbitrarily reshape the corresponding conditional image space, it still avoids simplifying assumptions in the image space. Therefore, the (invertible) transformation f_ψ avoids the conflict between the MLE and adversarial training mechanisms, and can leverage both.

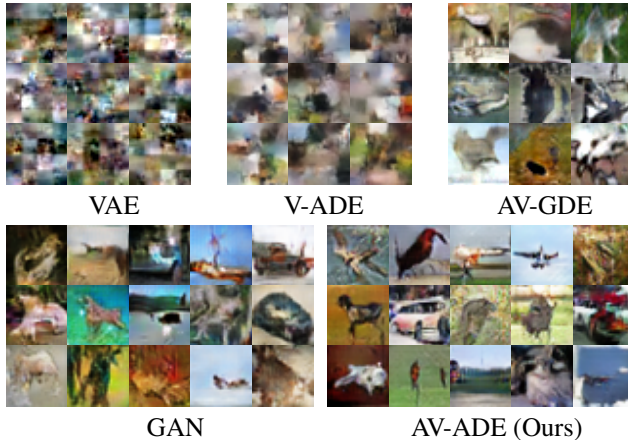
Adversarial training with adaptive density estimation. To sample the generative model, we sample latents from the prior, $z \sim p_\theta(z)$, which are then mapped to feature space through $\mu_\theta(z)$, and to image space through f_ψ^{-1} . We train our generator using the modified objective proposed by [51], combining both generator losses considered in [15], i.e. $\ln[(1 - D(G(z)))/D(G(z))]$, which yields:

$$\mathcal{L}_Q(p_{\theta,\psi}) = -\mathbb{E}_{p_\theta(z)} \left[\ln D(f_\psi^{-1}(\mu_\theta(z))) - \ln(1 - D(f_\psi^{-1}(\mu_\theta(z)))) \right]. \quad (5)$$

Assuming the discriminator D is trained to optimality at every step, it is easy to demonstrate that the generator is trained to optimize $\mathcal{D}_{\text{KL}}(p_{\theta,\psi}||p^*)$. The training procedure, written as an algorithm in Appendix H, alternates between (i) bringing $\mathcal{L}_Q(p_{\theta,\psi})$ closer to its optimal value $\mathcal{L}_Q^*(p_{\theta,\psi}) = \mathcal{D}_{\text{KL}}(p_{\theta,\psi}||p^*)$, and (ii) minimizing $\mathcal{L}_C(p_{\theta,\psi}) + \mathcal{L}_Q(p_{\theta,\psi})$. Assuming that the discriminator is trained to optimality at every step, the generative model is trained to minimize a bound on the symmetric sum of two KL divergences: $\mathcal{L}_C(p_{\theta,\psi}) + \mathcal{L}_Q^*(p_{\theta,\psi}) \geq \mathcal{D}_{\text{KL}}(p^*||p_{\theta,\psi}) + \mathcal{D}_{\text{KL}}(p_{\theta,\psi}||p^*) + \mathcal{H}(p^*)$, where the entropy of the data generating distribution, $\mathcal{H}(p^*)$, is an additive constant independent of the generative model $p_{\theta,\psi}$. In contrast, MLE and GANs optimize one of these divergences each.

5 Experimental evaluation

We present our evaluation protocol, followed by an ablation study to assess the importance of the components of our model (Section 5.1). We then show the quantitative and qualitative performance of



	f_ψ	Adv.	MLE	BPD ↓	IS ↑	FID ↓
GAN	×	✓	×	[7.0]	6.8	31.4
VAE	×	×	✓	4.4	2.0	171.0
V-ADE [†]	✓	×	✓	3.5	3.0	112.0
AV-GDE	×	✓	✓	4.4	5.1	58.6
AV-ADE [†]	✓	✓	✓	3.9	7.1	28.0

Table 1: Quantitative results. [†] : Parameter count decreased by 1.4% to compensate for f_ψ . [Square brackets] denote that the value is approximated, see Section 5.

Figure 5: Samples from GAN and VAE baselines, our V-ADE, AV-GDE and AV-ADE models, all trained on CIFAR-10.

our model, and compare it to the state of the art on the CIFAR-10 dataset in Section 5.2. We present additional results and comparisons on higher resolution datasets in Section 5.3.

Evaluation metrics. We evaluate our models with complementary metrics. To assess sample quality, we report the Fréchet inception distance (FID) [20] and the inception score (IS) [46], which are the de facto standard metrics to evaluate GANs [5, 55]. Although these metrics focus on sample quality, they are also sensitive to coverage, see Appendix D for details. To specifically evaluate the coverage of held-out data, we use the standard bits per dimension (BPD) metric, defined as the negative log-likelihood on held-out data, averaged across pixels and color channels [9].

Due to their degenerate low-dimensional support, GANs do not define a density in the image space, which prevents measuring BPD on them. To endow a GAN with a full support and a likelihood, we train an inference network “around it”, while keeping the weights of the GAN generator fixed. We also train an isotropic noise parameter σ . For both GANs and VAEs, we use this inference network to compute a lower bound to approximate the likelihood, *i.e.*, an upper bound on BPD. We evaluate all metrics using held-out data not used during training, which improves over common practice in the GAN literature, where training data is often used for evaluation.

5.1 Ablation study and comparison to VAE and GAN baselines

We conduct an ablation study on the CIFAR-10 dataset. ¹ Our GAN baseline uses the non-residual architecture of SNGAN [38], stable and quick to train, without spectral normalization. The same convolutional architecture is kept to build a VAE baseline. ² It produces the mean of a factorizing Gaussian distribution. To ensure a valid density model we add a trainable isotropic variance σ . We train the generator for coverage by optimizing $\mathcal{L}_Q(p_\theta)$, for quality by optimizing $\mathcal{L}_C(p_\theta)$, and for both by optimizing the sum $\mathcal{L}_Q(p_\theta) + \mathcal{L}_C(p_\theta)$. The model using Variational inference with Adaptive Density Estimation (ADE) is referred to as V-ADE. The addition of adversarial training is denoted AV-ADE, and hybrid training with a Gaussian decoder as AV-GDE. The bijective function f_ψ ³ increases the number of weights by approximately 1.4%, which we compensate for with a slight decrease in the width of the generator for fair comparison. ⁴ See Appendix B for details.

Experimental results in Table 1 confirm that the GAN baseline yields better sample quality (IS and FID) than the VAE baseline, *e.g.*, obtaining inception scores of 6.8 and 2.0, respectively. Conversely, VAE achieves better coverage, with a BPD of 4.4, compared to 7.0 for GAN. An identical generator trained for both quality and coverage, AV-GDE, obtains a sample quality that is in between that of the GAN and the VAE baselines, in line with the analysis in Section 4. Samples from the different models in Figure 5 confirm these quantitative results. Using f_ψ and training with $\mathcal{L}_C(p_\theta)$ only, denoted by V-ADE in the table, leads to improved sample quality with IS up from 2.0 to 3.0 and FID down from 171 to 112. Note that this quality is still below the GAN baseline and our AV-GDE model.

¹ We use the standard split of 50k/10k train/test images of 32×32 pixels. ² In the VAE model, some intermediate feature maps are treated as conditional latent variables, allowing for hierarchical top-down sampling (see Appendix B). Experimentally, we find that similar top-down sampling is not effective for the GAN model.

³ implemented as a small Real-NVP with 1 scale, 3 residual blocks, 2 layers per block. ⁴ This is too small to make a significant difference in experiments.

When f_ψ is used with coverage and quality driven training, AV-ADE, we obtain improved IS and FID scores over the GAN baseline, with IS up from 6.8 to 7.1, and FID down from 31.4 to 28.0. The examples shown in the figure confirm the high quality of the samples generated by our AV-ADE model. Our model also achieves a better BPD than the VAE baseline. These experiments demonstrate that our proposed bijective feature space substantially improves the compatibility of coverage and quality driven training. We obtain improvements over both VAE and GAN in terms of held-out likelihood, and improve VAE sample quality to, or beyond, that of GAN. We further evaluate our model using the recent precision and recall approach of [44] in the classification framework of [49] in Appendix E. Additional results showing the impact of the number of layers and scales in the bijective similarity mapping f_ψ (Appendix F), reconstructions qualitatively demonstrating the inference abilities of our AV-ADE model (Appendix G) are presented in the appendix.

5.2 Refinements and comparison to the state of the art

We now consider further refinements to our model, inspired by recent generative modeling literature. Four refinements are used: (i) adding residual connections to the discriminator [18] (rd), (ii) leveraging more accurate posterior approximations using inverse auto-regressive flow [27] (iaf); see Appendix B, (iii) training wider generators with twice as many channels (wg), and (iv) using a hierarchy of two scales to build f_ψ (s2); see Appendix F. Table 2 shows consistent improvements with these additions, in terms of BPD, IS, FID.

Refinements	BPD ↓	IS ↑	FID ↓
GAN	[7.0]	6.8	31.4
GAN (rd)	[6.9]	7.4	24.0
AV-ADE	3.9	7.1	28.0
AV-ADE (rd)	3.8	7.5	26.0
AV-ADE (wg, rd)	3.8	8.2	17.2
AV-ADE (iaf, rd)	3.7	8.1	18.6
AV-ADE (s2)	3.5	6.9	28.9

Table 2: Model refinements.

Table 3 compares our model to existing hybrid approaches and state-of-the-art generative models on CIFAR-10. In the category of hybrid models that optimize likelihood, denoted by Hybrid (L) in the table, FlowGAN(H) optimizes MLE and an adversarial loss, and FlowGAN(A) is trained adversarially. The AV-ADE model significantly outperforms these two variants both in terms of BPD, from 4.2 to between 3.5 and 3.8, and quality, *e.g.*, IS improves from 5.8 to 8.2. Compared to models that train an inference network adversarially, denoted by Hybrid (A), our model shows a substantial improvement in IS from 7.0 to 8.2. Note that these models do not allow likelihood evaluation, thus BPD values are absent.

Hybrid (L)	BPD ↓	IS ↑	FID ↓
AV-ADE (wg, rd)	3.8	8.2	17.2
AV-ADE (iaf, rd)	3.7	8.1	18.6
AV-ADE (S2)	3.5	6.9	28.9
FlowGAN(A) [16]	8.5	5.8	
FlowGAN(H) [16]	4.2	3.9	

Hybrid (A)	BPD ↓	IS ↑	FID ↓
AGE [53]		5.9	
ALI [12]		5.3	
SVAE [6]		6.8	
α -GAN [42]		6.8	
SVAE-r [6]		7.0	

Compared to adversarial models, which are not optimized for support coverage, AV-ADE obtains better FID (17.2 down from 21.7) and similar IS (8.2 for both) compared to SNGAN with residual connections and hinge-loss, despite training on 17% less data than GANs (test split removed). The improvement in FID is likely due to this measure being more sensitive to support coverage than IS. Compared to models optimized with MLE only, we obtain a BPD between 3.5 and 3.7, comparable to 3.5 for Real-NVP demonstrating a good coverage of the support of held-out data. We computed IS and FID scores for MLE based models using publicly released code, with provided parameters (denoted by $\dagger$ in the table) or trained ourselves (denoted by $\ddagger$). Despite being smaller (for reference Glow has 384 layers VS at most 10 for our deeper generator), our AV-ADE model generates better samples, *e.g.*, IS up from 5.5 to 8.2 (samples displayed in Figure 6), owing to quality driven training controlling where the off-dataset mass goes. Additional samples from our AV-ADE model and comparison to others models are given in Appendix A.

Adversarial	BPD ↓	IS ↑	FID ↓
mmd-GAN[1]		7.3	25.0
SNGan [38]		7.4	29.3
BatchGAN [33]		7.5	23.7
WGAN-GP [17]		7.9	
SNGAN _(R,H)		8.2	21.7

MLE	BPD ↓	IS ↑	FID ↓
Real-NVP [9]	3.5	4.5 $\dagger$	56.8 $\dagger$
VAE-IAF [27]	3.1	3.8 $\dagger$	73.5 $\dagger$
Pixcnn++ [47]	2.9	5.5	
Flow++ [21]	3.1		
Glow [25]	3.4	5.5$\ddagger$	46.8$\ddagger$

Table 3: Performance on CIFAR10, without labels. MLE and Hybrid (L) models discard the test split. $\dagger$: computed by us using provided weights. $\ddagger$: computed by us using provided code to (re)train models.

5.3 Results on additional datasets

To further validate our model we evaluate it on STL10 (48×48), ImageNet and LSUN (both 64×64). We use a wide generator to account for the higher resolution, without iaf, a single scale in f_ψ , and no residual blocks (see Section 5.2). The architecture and training hyper-parameters are not tuned, besides adding one layer at resolution 64×64 , which demonstrates the stability of our approach. On STL10, Table 4 shows that our AV-ADE improves inception score over SNGAN, from 9.1 up to



Figure 6: Samples from models trained on CIFAR-10. Our AV-ADE spills less mass on unrealistic samples, owing to adversarial training which controls where off-dataset mass goes.

9.4, and is second best in FID. Our likelihood performance, between 3.8 and 4.4, and close to that of Real-NVP at 3.7, demonstrates good coverage of the support of held-out data. On the ImageNet dataset, maintaining high sample quality, while covering the full support is challenging, due to its very diverse support. Our AV-ADE model obtains a sample quality behind that of MMD-GAN with IS/FID scores at 8.5/45.5 vs. 10.9/36.6. However, MMD-GAN is trained purely adversarially and does not provide support guarantees, unlike our approach.

Figure 7 shows samples from our generator trained on a single GPU with 11 Gb of memory on LSUN classes. It yields compelling samples compared to those of the Glow model, despite having more flexibility (over 500 VS 7 layers) showing that Glow spills more mass outside of the training support. Additional samples and other LSUN categories are presented in Appendix A.

STL-10 (48×48)	BPD ↓	IS ↑	FID ↓	ImageNet (64×64)	BPD ↓	IS ↑	FID ↓	LSUN	Real-NVP	Glow	Ours
AV-ADE (wg, wd)	4.4	9.4	44.3	AV-ADE (wg, wd)	4.90	8.5	45.5	Bedroom	(2.72/×)	(2.38/208.8 [†])	(3.91, 21.1)
AV-ADE (iaf, wd)	4.0	8.6	52.7	Real-NVP	3.98			Tower	(2.81/×)	(2.46/214.5 [†])	(3.95, 15.5)
AV-ADE (s2)	3.8	8.6	52.1	Glow	3.81			Church	(3.08/×)	(2.67/222.3 [†])	(4.3, 13.1)
Real-NVP	3.7[‡]	4.8 [‡]	103.2 [‡]	Flow++	3.69			Classroom	×	×	(4.6, 20.0)
BatchGAN		8.7	51	MMD-GAN		10.9	36.6	Restaurant	×	×	(4.7, 20.5)
SNGAN (Res-Hinge)		9.1	40.1								

Table 4: Results on the STL-10, ImageNet, and LSUN datasets. AV-ADE (wg, rd) is used for LSUN.



Figure 7: Samples from models trained on LSUN Churches (C) and bedrooms (B). Our AV-ADE model over-generalises less and produces more compelling samples. See Appendix A for more classes and samples.

6 Conclusion

We presented a generative model that leverages invertible network layers to relax the conditional pixel independence assumption commonly made in VAE decoders. It allows for efficient feed-forward sampling, and can be trained using a maximum likelihood criterion that ensures coverage of the data generating distribution, as well as an adversarial criterion that ensures high sample quality.

Acknowledgments

The authors would like to thank Corentin Tallec, Mathilde Caron, Adria Ruiz and Nikita Dvornik for useful feedback and discussions. Acknowledgments also go to our anonymous reviewers, who contributed valuable comments and remarks.

This work was supported in part by the grants ANR16-CE23-0006 “Deep in France”, LabEx PERSYVAL-Lab (ANR-11-LABX0025-01) as well as the Indo-French project EVEREST (no. 5302-1) funded by CEFIPRA and a grant from ANR (AVENUE project ANR-18-CE23-0011),

References

- [1] M. Arbel, D. J. Sutherland, M. Binkowski, and A. Gretton. On gradient regularizers for mmd gans. In *NeurIPS*, 2018.
- [2] M. Arjovsky, S. Chintala, and L. Bottou. Wasserstein generative adversarial networks. In *ICML*, 2017.
- [3] P. Bachman. An architecture for deep, hierarchical generative models. In *NIPS*, 2016.
- [4] C. Bishop. *Pattern recognition and machine learning*. Springer-Verlag, 2006.
- [5] A. Brock, J. Donahue, and K. Simonyan. Large scale GAN training for high fidelity natural image synthesis. In *ICLR*, 2019.
- [6] L. Chen, S. Dai, Y. Pu, E. Zhou, C. Li, Q. Su, C. Chen, and L. Carin. Symmetric variational autoencoder and connections to adversarial learning. In *AISTATS*, 2018.
- [7] X. Chen, D. Kingma, T. Salimans, Y. Duan, P. Dhariwal, J. Schulman, I. Sutskever, and P. Abbeel. Variational lossy autoencoder. In *ICLR*, 2017.
- [8] H. De Vries, F. Strub, J. Mary, H. Larochelle, O. Pietquin, and A. C. Courville. Modulating early visual processing by language. In *NIPS*, 2017.
- [9] L. Dinh, J. Sohl-Dickstein, and S. Bengio. Density estimation using real NVP. In *ICLR*, 2017.
- [10] J. Donahue, P. Krähenbühl, and T. Darrell. Adversarial feature learning. In *ICLR*, 2017.
- [11] G. Dorta, S. Vicente, L. Agapito, N. Campbell, and I. Simpson. Structured uncertainty prediction networks. In *CVPR*, 2018.
- [12] V. Dumoulin, I. Belghazi, B. Poole, A. Lamb, M. Arjovsky, O. Mastropietro, and A. Courville. Adversarially learned inference. In *ICLR*, 2017.
- [13] V. Dumoulin, J. Shlens, and M. Kudlur. A learned representation for artistic style. In *ICLR*, 2017.
- [14] M. Elfeki, C. Couprie, M. Riviere, and M. Elhoseiny. GDPP: Learning diverse generations using determinantal point processes. In *arxiv.org/pdf/1812.00068*, 2018.
- [15] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio. Generative adversarial nets. In *NIPS*, 2014.
- [16] A. Grover, M. Dhar, and S. Ermon. Flow-gan: Combining maximum likelihood and adversarial learning in generative models. In *AAAI Press*, 2018.
- [17] I. Gulrajani, F. Ahmed, M. Arjovsky, V. Dumoulin, and A. Courville. Improved training of Wasserstein GANs. In *NIPS*, 2017.
- [18] I. Gulrajani, F. Ahmed, M. Arjovsky, V. Dumoulin, and A. C. Courville. Improved training of Wasserstein GANs. In *NIPS*, 2017.
- [19] I. Gulrajani, K. Kumar, F. Ahmed, A. A. Taiga, F. Visin, D. Vazquez, and A. Courville. PixelVAE: A latent variable model for natural images. In *ICLR*, 2017.

- [20] 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. In *NIPS*, 2017.
- [21] J. Ho, X. Chen, A. Srinivas, Y. Duan, and P. Abbeel. Flow++: Improving flow-based generative models with variational dequantization and architecture design. In *ICML*, 2019.
- [22] X. Hou, L. Shen, K. Sun, and G. Qiu. Deep feature consistent variational autoencoder. 2017.
- [23] T. Karras, T. Aila, and S. L. abd J. Lehtinen. Progressive growing of GANs for improved quality, stability, and variation. In *ICLR*, 2018.
- [24] D. Kingma and J. Ba. Adam: A method for stochastic optimization. In *ICLR*, 2015.
- [25] D. P. Kingma and P. Dhariwal. Glow: Generative flow with invertible 1x1 convolutions. In *NeurIPS*, 2018.
- [26] D. P. Kingma and M. Welling. Auto-encoding variational Bayes. In *ICLR*, 2014.
- [27] D. P. Kingma, T. Salimans, R. Józefowicz, X. Chen, I. Sutskever, and M. Welling. Improving variational autoencoders with inverse autoregressive flow. In *NIPS*, 2016.
- [28] A. Larsen, S. Sønderby, H. Larochelle, and O. Winther. Autoencoding beyond pixels using a learned similarity metric. In *ICML*, 2016.
- [29] Y. Li, K. Swersky, and R. S. Zemel. Generative moment matching networks. In *ICML*, 2015.
- [30] Z. Lin, A. Khetan, G. Fanti, and S. Oh. PacGAN: The power of two samples in generative adversarial networks. In *NeurIPS*, 2018.
- [31] O. Litany, A. Bronstein, M. Bronstein, and A. Makadia. Deformable shape completion with graph convolutional autoencoders. In *CVPR*, 2018.
- [32] T. Lucas and J. Verbeek. Auxiliary guided autoregressive variational autoencoders. In *ECML*, 2018.
- [33] T. Lucas, C. Tallec, Y. Ollivier, and J. Verbeek. Mixed batches and symmetric discriminators for GAN training. In *ICML*, 2018.
- [34] A. Makhzani, J. Shlens, N. Jaitly, and I. Goodfellow. Adversarial autoencoders. In *ICLR*, 2016.
- [35] M. Mathieu, C. Couprie, and Y. LeCun. Deep multi-scale video prediction beyond mean square error. In *ICLR*, 2016.
- [36] J. Menick and N. Kalchbrenner. Generating high fidelity images with subscale pixel networks and multidimensional upscaling. In *ICLR*, 2019.
- [37] T. Miyato and M. Koyama. cGANs with projection discriminator. In *ICLR*, 2018.
- [38] T. Miyato, T. Kataoka, M. Koyama, and Y. Yoshida. Spectral normalization for generative adversarial networks. In *ICLR*, 2018.
- [39] P. Ramachandran, T. L. Paine, P. Khorrami, M. Babaeizadeh, S. Chang, Y. Zhang, M. A. Hasegawa-Johnson, R. H. Campbell, and T. S. Huang. Fast generation for convolutional autoregressive models. In *ICLR workshop*, 2017.
- [40] D. Rezende and S. Mohamed. Variational inference with normalizing flows. In *ICML*, 2015.
- [41] D. Rezende, S. Mohamed, and D. Wierstra. Stochastic backpropagation and approximate inference in deep generative models. In *ICML*, 2014.
- [42] M. Rosca, B. Lakshminarayanan, D. Warde-Farley, and S. Mohamed. Variational approaches for auto-encoding generative adversarial networks. *arXiv preprint arXiv:1706.04987*, 2017.
- [43] Y. Saatchi and A. G. Wilson. Bayesian gan. 2017.
- [44] M. S. M. Sajjadi, O. Bachem, M. Lucic, O. Bousquet, and S. Gelly. Assessing generative models via precision and recall. In *NeurIPS*, 2018.

- [45] T. Salimans and D. P. Kingma. Weight normalization: A simple reparameterization to accelerate training of deep neural networks. In *NIPS*, 2016.
- [46] T. Salimans, I. Goodfellow, W. Zaremba, V. Cheung, A. Radford, and X. Chen. Improved techniques for training GANs. In *NIPS*, 2016.
- [47] T. Salimans, A. Karpathy, X. Chen, and D. P. Kingma. PixelCNN++: Improving the PixelCNN with discretized logistic mixture likelihood and other modifications. In *ICLR*, 2017.
- [48] T. Salimans, H. Zhang, A. Radford, and D. Metaxas. Improving gans using optimal transport. In *ICLR*, 2018.
- [49] K. Shmelkov, C. Schmid, and K. Alahari. How good is my GAN? In *ECCV*, 2018.
- [50] C. K. Sønderby, T. Raiko, L. Maaløe, S. K. Sønderby, and O. Winther. Ladder variational autoencoders. In *NIPS*, 2016.
- [51] C. K. Sønderby, J. Caballero, L. Theis, W. Shi, and F. Huszár. Amortised MAP inference for image super-resolution. In *ICLR*, 2017.
- [52] H. Thanh-Tung, T. Tran, and S. Venkatesh. Improving generalization and stability of generative adversarial networks. In *ICLR*, 2019.
- [53] D. Ulyanov, A. Vedaldi, and V. Lempitsky. It takes (only) two: Adversarial generator-encoder networks. In *AAAI*, 2018.
- [54] A. van den Oord, N. Kalchbrenner, and K. Kavukcuoglu. Pixel recurrent neural networks. In *ICML*, 2016.
- [55] H. Zhang, I. Goodfellow, D. Metaxas, and A. Odena. Self-attention generative adversarial networks. *arXiv preprint arXiv:1805.08318*, 2018.