

# Autoencoders and Variational Autoencoders

From compressing inputs to sampling a distribution

---

Prof. Nipun Batra

ES 667 · Deep Learning · IIT Gandhinagar

# **Why reconstruction is not generation**

---

A classifier only learns the conditional  $p(y \mid x)$ . A **generative** model is more ambitious:

# The generative-modeling question

A classifier only learns the conditional  $p(y \mid x)$ . A **generative** model is more ambitious:

**learn  $p_\theta(x)$  well enough to generate plausible new  $x$**

# The generative-modeling question

A classifier only learns the conditional  $p(y \mid x)$ . A **generative** model is more ambitious:

learn  $p_\theta(x)$  well enough to **generate** plausible new  $x$

**Discriminative** —  $p(y \mid x)$ : given an image, name it. Never has to imagine an image.

**Generative** —  $p_\theta(x)$ : what does a **plausible image even look like**? Must model the data itself.

# The generative-modeling question

A classifier only learns the conditional  $p(y \mid x)$ . A **generative** model is more ambitious:

learn  $p_\theta(x)$  well enough to **generate** plausible new  $x$

**Discriminative** —  $p(y \mid x)$ : given an image, name it. Never has to imagine an image.

**Generative** —  $p_\theta(x)$ : what does a **plausible image even look like**? Must model the data itself.

today: build a generative model by **learning a latent space we can sample from**

Compress the input to a code  $z$ , then reconstruct it:

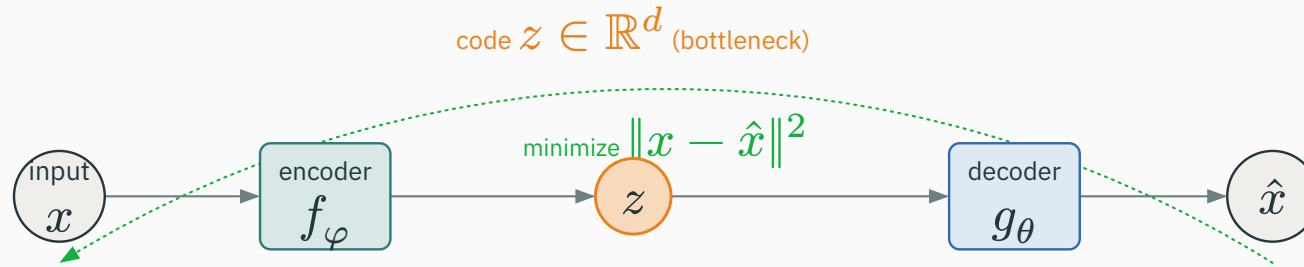
Compress the input to a code  $z$ , then reconstruct it:

$$x \xrightarrow{f_\varphi} z \xrightarrow{g_\theta} \hat{x}$$



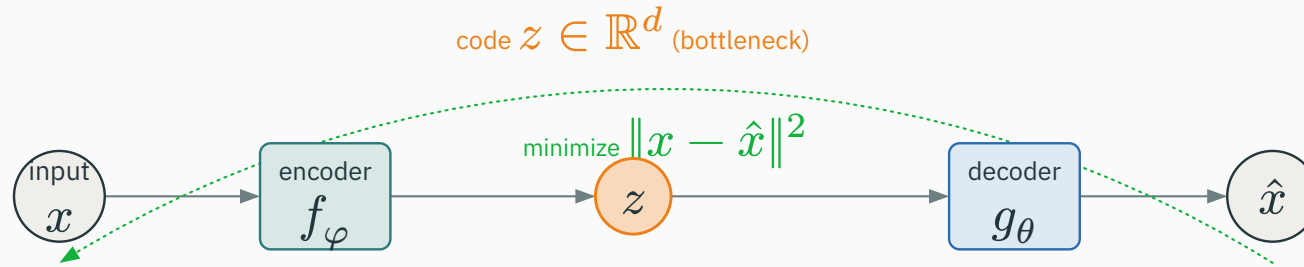
Compress the input to a code  $z$ , then reconstruct it:

$$x \xrightarrow{f_\varphi} z \xrightarrow{g_\theta} \hat{x}$$



Compress the input to a code  $z$ , then reconstruct it:

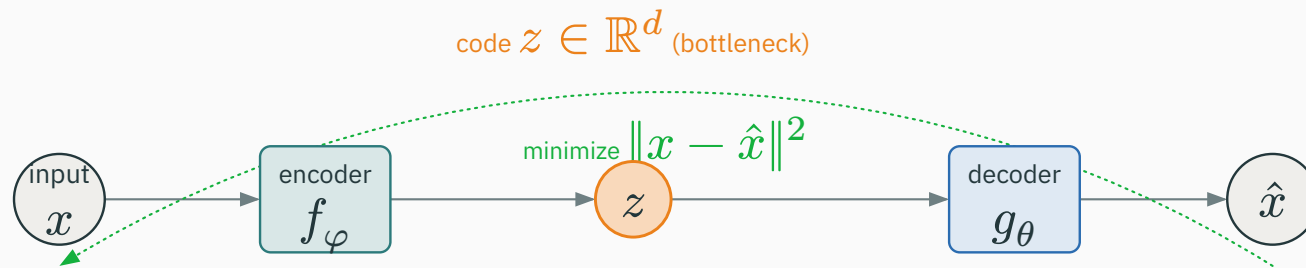
$$x \xrightarrow{f_\varphi} z \xrightarrow{g_\theta} \hat{x}$$



**encoder**  $f_\varphi$  maps  $x \rightarrow z$ ; **decoder**  $g_\theta$  maps  $z \rightarrow \hat{x}$ ; train to make  $\hat{x} \approx x$

Compress the input to a code  $z$ , then reconstruct it:

$$x \xrightarrow{f_\varphi} z \xrightarrow{g_\theta} \hat{x}$$



**encoder**  $f_\varphi$  maps  $x \rightarrow z$ ; **decoder**  $g_\theta$  maps  $z \rightarrow \hat{x}$ ; train to make  $\hat{x} \approx x$

(VAE convention:  $\theta$  is the **generative / decoder** side,  $\varphi$  the encoder — the reverse of the classifier's  $f_\theta \rightarrow g_\varphi$  in L18)

Why does a narrow code  $z$  force **learning** rather than **copying**?

Why does a narrow code  $z$  force **learning** rather than **copying**?

If  $z$  is constrained (few dimensions), the network cannot store every input coordinate. To reconstruct well it must keep the **salient structure** and discard the rest.

Why does a narrow code  $z$  force **learning** rather than **copying**?

If  $z$  is constrained (few dimensions), the network cannot store every input coordinate. To reconstruct well it must keep the **salient structure** and discard the rest.

a wide-open code learns the **identity** map and learns nothing useful

# The reconstruction objective

Train encoder and decoder together to minimize squared reconstruction error:

Train encoder and decoder together to minimize squared reconstruction error:

$$\mathcal{L}_{\text{AE}} = \sum_i \left\| x_i - g_{\theta}(f_{\varphi}(x_i)) \right\|_2^2$$



Train encoder and decoder together to minimize squared reconstruction error:

$$\mathcal{L}_{\text{AE}} = \sum_i \left\| x_i - g_{\theta}(f_{\varphi}(x_i)) \right\|_2^2$$

This is not an arbitrary choice. A **Gaussian-output** decoder gives

$$p_{\theta}(x \mid z) = \mathcal{N}(g_{\theta}(z), \sigma_x^2 I) \quad \Rightarrow \quad -\log p_{\theta}(x \mid z) \propto \|x - g_{\theta}(z)\|^2$$

Train encoder and decoder together to minimize squared reconstruction error:

$$\mathcal{L}_{\text{AE}} = \sum_i \left\| x_i - g_{\theta}(f_{\varphi}(x_i)) \right\|_2^2$$

This is not an arbitrary choice. A **Gaussian-output** decoder gives

$$p_{\theta}(x \mid z) = \mathcal{N}(g_{\theta}(z), \sigma_x^2 I) \quad \Rightarrow \quad -\log p_{\theta}(x \mid z) \propto \|x - g_{\theta}(z)\|^2$$

squared error = negative Gaussian log-likelihood — a first taste of the probabilistic view

Even without a probabilistic story, a trained autoencoder is useful:

Even without a probabilistic story, a trained autoencoder is useful:

- **denoising** — corrupt the input, reconstruct the clean signal;
- **compression** — a compact code  $z$  for storage or transmission;
- **visualization** — a 2-D or 3-D code to plot the data;
- **features** — reuse  $z$  (or the encoder) for a downstream task.

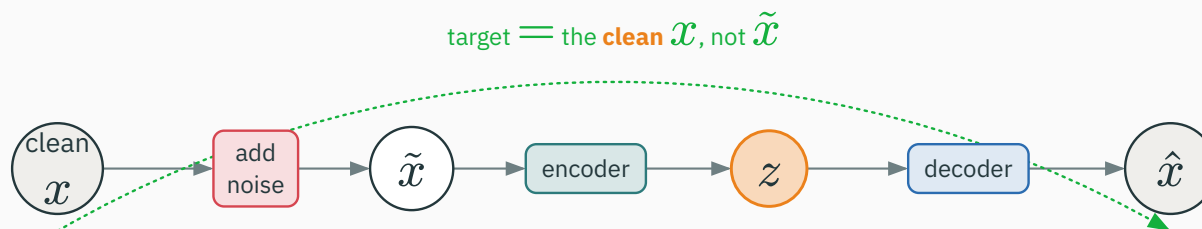
Even without a probabilistic story, a trained autoencoder is useful:

- **denoising** — corrupt the input, reconstruct the clean signal;
- **compression** — a compact code  $z$  for storage or transmission;
- **visualization** — a 2-D or 3-D code to plot the data;
- **features** — reuse  $z$  (or the encoder) for a downstream task.

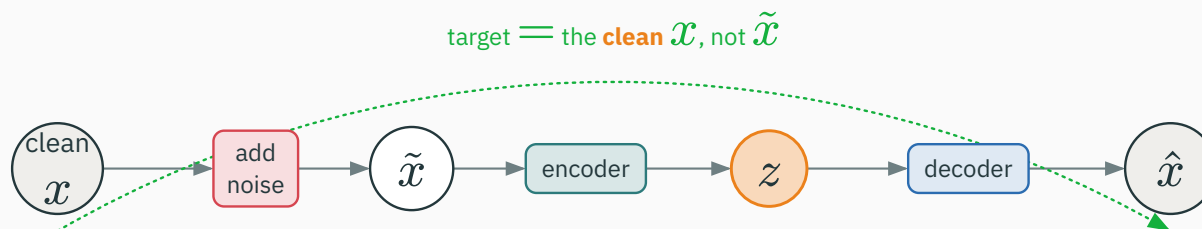
all of these use  $z$  as a **representation** — none of them require **sampling** new  $x$

Force the code to capture structure by making the task harder — corrupt, then restore:

Force the code to capture structure by making the task harder — corrupt, then restore:



Force the code to capture structure by making the task harder — corrupt, then restore:



the target is the **clean**  $x$ ; the encoder must learn what survives noise, not memorize pixels



So can we just feed a random  $z$  to the decoder and get a new image?

So can we just feed a random  $z$  to the decoder and get a new image?

**No.** An autoencoder never constrains **where** codes live. Encoded training points occupy **disconnected, irregular** regions. A random  $z$  lands **between** them and decodes to nonsense.

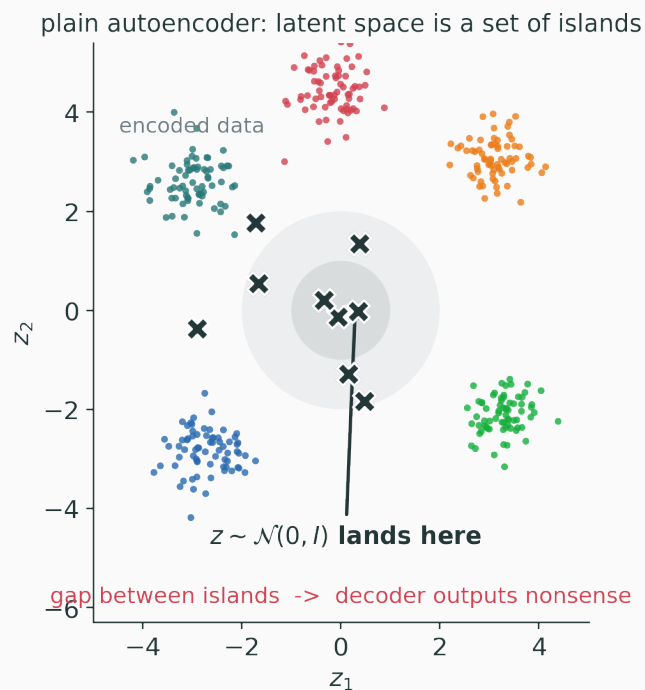
So can we just feed a random  $z$  to the decoder and get a new image?

**No.** An autoencoder never constrains **where** codes live. Encoded training points occupy **disconnected, irregular** regions. A random  $z$  lands **between** them and decodes to nonsense.

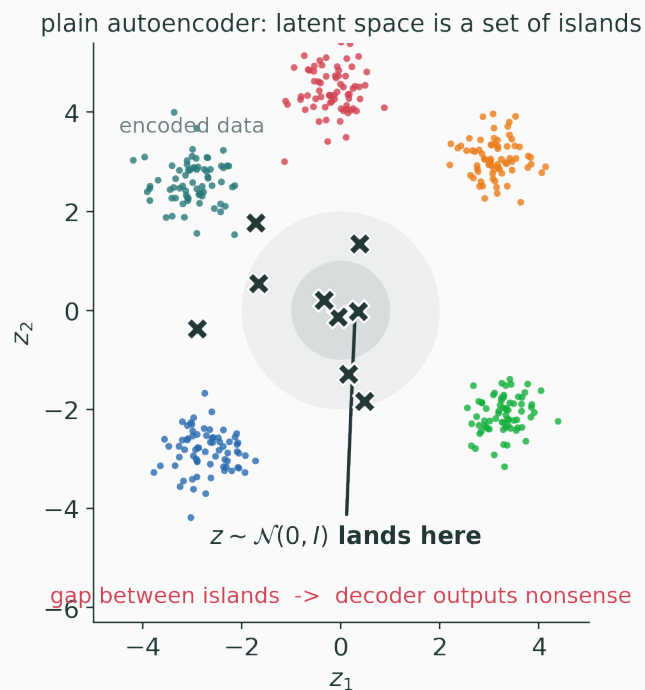
reconstruction says nothing about the **empty** parts of latent space

Encode the data — the codes form islands with wide empty gaps:

Encode the data — the codes form islands with wide empty gaps:



Encode the data — the codes form islands with wide empty gaps:



samples from  $\mathcal{N}(0, I)$  mostly land in the **gap between islands**  $\rightarrow$  the decoder was never trained there

For a latent space to be **generative**, sampling the prior must produce plausible data:

For a latent space to be **generative**, sampling the prior must produce plausible data:

$$z \sim p(z) \Rightarrow g_{\theta}(z) \text{ should be plausible}$$



For a latent space to be **generative**, sampling the prior must produce plausible data:

$$z \sim p(z) \Rightarrow g_{\theta}(z) \text{ should be plausible}$$

We need two things at once: (1) a **known** distribution  $p(z)$  we can sample, and (2) a decoder trained so that **every** likely  $z$  maps to something realistic. The VAE engineers exactly this.

## INTERACTIVE

↗ [nipunbatra.github.io/interactive-articles/vae-latent-explorer](https://nipunbatra.github.io/interactive-articles/vae-latent-explorer)

Move a point around a 2-D latent space and watch the decoded output. See the islands of a plain autoencoder, and how prior samples fall into the gaps between them.

# **The VAE probabilistic model**

---

A VAE first **defines** how data is generated from a simple latent cause:

A VAE first **defines** how data is generated from a simple latent cause:

$$z \sim p(z) = \mathcal{N}(0, I), \quad x \sim p_{\theta}(x \mid z)$$

A VAE first **defines** how data is generated from a simple latent cause:

$$z \sim p(z) = \mathcal{N}(0, I), \quad x \sim p_{\theta}(x \mid z)$$

1. sample a simple latent  $z$  from a fixed **standard Gaussian** prior;
2. **render** an observation  $x$  with a neural decoder  $p_{\theta}(x \mid z)$ .

A VAE first **defines** how data is generated from a simple latent cause:

$$z \sim p(z) = \mathcal{N}(0, I), \quad x \sim p_{\theta}(x \mid z)$$

1. sample a simple latent  $z$  from a fixed **standard Gaussian** prior;
2. **render** an observation  $x$  with a neural decoder  $p_{\theta}(x \mid z)$ .

the prior is **fixed and sampleable by construction** — that is what fixes the islands problem

# A decoder is a conditional distribution

The decoder does not output “an image” — it outputs the **parameters of a distribution** over  $x$ :



# A decoder is a conditional distribution

The decoder does not output “an image” — it outputs the **parameters of a distribution** over  $x$ :

$$p_{\theta}(x \mid z) = \mathcal{N}(g_{\theta}(z), \sigma_x^2 I) \quad (\text{Gaussian example})$$

# A decoder is a conditional distribution

The decoder does not output “an image” — it outputs the **parameters of a distribution** over  $x$ :

$$p_{\theta}(x \mid z) = \mathcal{N}(g_{\theta}(z), \sigma_x^2 I) \quad (\text{Gaussian example})$$

the familiar “reconstruction image” is just the conditional **mean**  $g_{\theta}(z)$

To train, we would like the **posterior** over the latent cause of each  $x$ :

To train, we would like the **posterior** over the latent cause of each  $x$ :

$$p_{\theta}(z \mid x) = \frac{p_{\theta}(x \mid z) p(z)}{p_{\theta}}(x)$$

To train, we would like the **posterior** over the latent cause of each  $x$ :

$$p_{\theta}(z \mid x) = \frac{p_{\theta}(x \mid z) p(z)}{p_{\theta}}(x)$$

$$p_{\theta}(x) = \int p_{\theta}(x \mid z) p(z) dz \quad (\text{intractable integral})$$

To train, we would like the **posterior** over the latent cause of each  $x$ :

$$p_{\theta}(z \mid x) = \frac{p_{\theta}(x \mid z) p(z)}{p_{\theta}}(x)$$

$$p_{\theta}(x) = \int p_{\theta}(x \mid z) p(z) dz \quad (\text{intractable integral})$$

The evidence  $p_{\theta}(x)$  integrates over **all**  $z$  through a nonlinear decoder. No closed form — so the posterior is out of reach too.

Sidestep the intractable posterior: let the **encoder** predict a tractable Gaussian approximation.

Sidestep the intractable posterior: let the **encoder** predict a tractable Gaussian approximation.

$$q_{\varphi}(z \mid x) = \mathcal{N}\left(\mu_{\varphi}(x), \text{diag}\left(\sigma_{\varphi}(x)^2\right)\right)$$



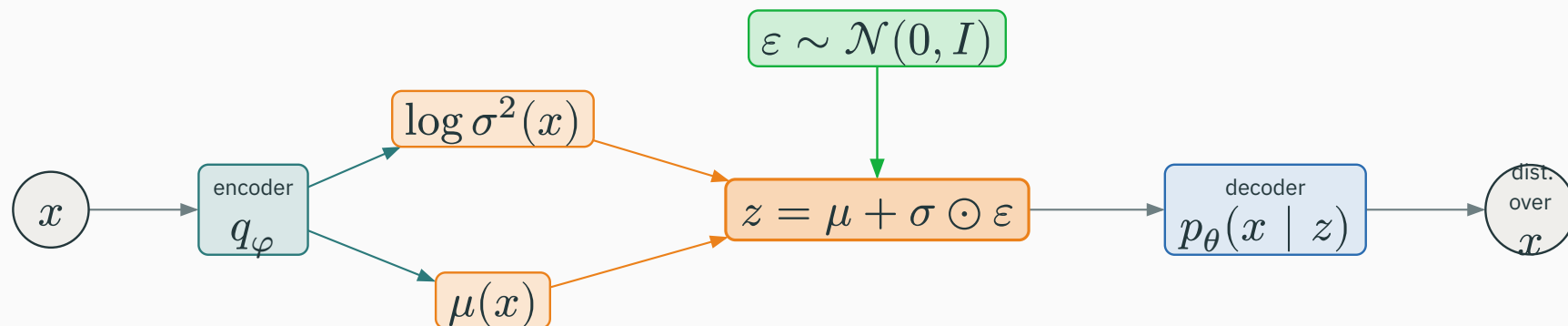
Sidestep the intractable posterior: let the **encoder** predict a tractable Gaussian approximation.

$$q_{\varphi}(z \mid x) = \mathcal{N}\left(\mu_{\varphi}(x), \text{diag}\left(\sigma_{\varphi}(x)^2\right)\right)$$

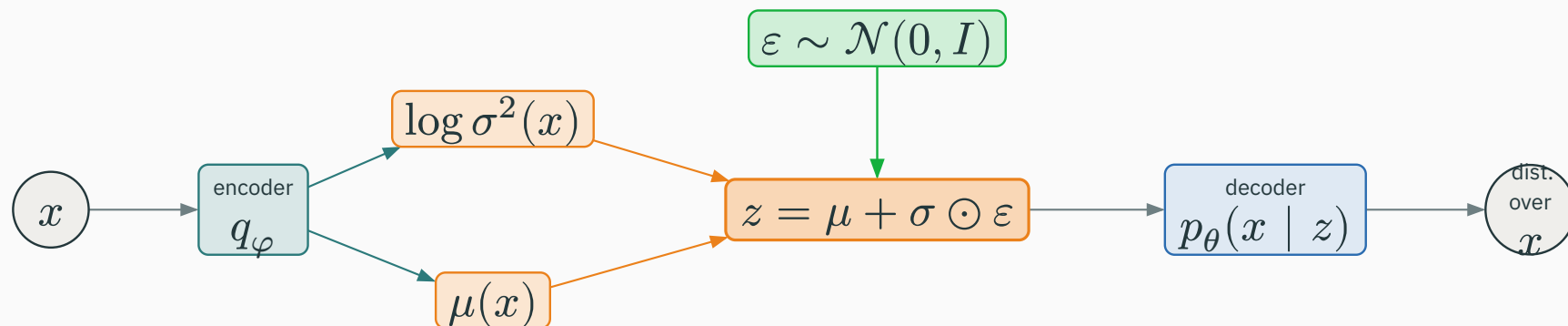
The encoder now outputs a **distribution**, not a single point: a mean  $\mu_{\varphi}(x)$  and a per-dimension variance  $\sigma_{\varphi}(x)^2$ . Each input maps to a small cloud in latent space.

The whole forward model — encoder predicts a distribution, we sample, the decoder renders:

The whole forward model — encoder predicts a distribution, we sample, the decoder renders:



The whole forward model — encoder predicts a distribution, we sample, the decoder renders:



$x \rightarrow \text{encoder} \rightarrow (\mu, \log \sigma^2) \rightarrow \text{sample } z \rightarrow \text{decoder} \rightarrow \text{distribution over } x$

We cannot maximize  $\log p_{\theta}(x)$  directly — but we can maximize a **lower bound** on it:

We cannot maximize  $\log p_\theta(x)$  directly — but we can maximize a **lower bound** on it:

$$\log p_\theta(x) \geq \underbrace{\mathbb{E}_{q_\varphi(z|x)}[\log p_\theta(x|z)]}_{\text{reconstruct}} - \underbrace{D_{\text{KL}}(q_\varphi(z|x) \parallel p(z))}_{\text{organize latent space}}$$

We cannot maximize  $\log p_\theta(x)$  directly — but we can maximize a **lower bound** on it:

$$\log p_\theta(x) \geq \underbrace{\mathbb{E}_{q_\varphi(z|x)}[\log p_\theta(x|z)]}_{\text{reconstruct}} - \underbrace{D_{\text{KL}}(q_\varphi(z|x) \parallel p(z))}_{\text{organize latent space}}$$

two terms, two jobs — this is the object we actually optimize

This bound is the **Evidence Lower BOund**:



This bound is the **Evidence Lower BOund**:

**maximize a tractable lower bound on the data log-likelihood**

This bound is the **Evidence Lower BOund**:

**maximize a tractable lower bound on the data log-likelihood**

Understand what each term **does** before any algebra: the first term **explains the data**, the second **keeps latent codes near the prior** so we can sample them later.

$$\mathbb{E}_{q_{\varphi}(z \mid x)}[\log p_{\theta}(x \mid z)]$$

$$\mathbb{E}_{q_{\varphi}(z \mid x)}[\log p_{\theta}(x \mid z)]$$

Read it as a recipe:

$$\mathbb{E}_{q_{\varphi}(z \mid x)}[\log p_{\theta}(x \mid z)]$$

Read it as a recipe:

1. take input  $x$ , encode it, and **sample** a plausible latent explanation  $z \sim q_{\varphi}(z \mid x)$ ;
2. require the decoder to **reconstruct**  $x$  from that  $z$ .

$$\mathbb{E}_{q_{\varphi}(z \mid x)}[\log p_{\theta}(x \mid z)]$$

Read it as a recipe:

1. take input  $x$ , encode it, and **sample** a plausible latent explanation  $z \sim q_{\varphi}(z \mid x)$ ;
2. require the decoder to **reconstruct**  $x$  from that  $z$ .

high when the sampled latent actually explains the observed input

$$D_{\text{KL}}(q_{\varphi}(z \mid x) \parallel p(z))$$

$$D_{\text{KL}}(q_{\varphi}(z \mid x) \parallel p(z))$$

Read it as a constraint:



$$D_{\text{KL}}(q_{\varphi}(z \mid x) \parallel p(z))$$

Read it as a constraint:

each example's encoded cloud  $q_{\varphi}(z \mid x)$  should stay **close to the shared prior**  $\mathcal{N}(0, I)$ . This is what pulls every code toward one common, sampleable region.

$$D_{\text{KL}}(q_{\varphi}(z \mid x) \parallel p(z))$$

Read it as a constraint:

each example's encoded cloud  $q_{\varphi}(z \mid x)$  should stay **close to the shared prior**  $\mathcal{N}(0, I)$ . This is what pulls every code toward one common, sampleable region.

without it, the encoder would scatter codes into islands again

# The Gaussian KL — show, do not derive

D DERIVATION

For a diagonal-Gaussian posterior against the standard-normal prior, the KL is **closed form**:

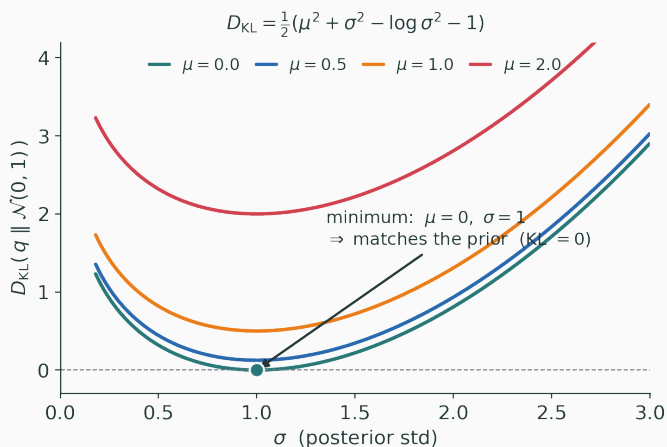
For a diagonal-Gaussian posterior against the standard-normal prior, the KL is **closed form**:

$$D_{\text{KL}}(\mathcal{N}(\mu, \sigma^2) \parallel \mathcal{N}(0, 1)) = \frac{1}{2} \sum_j (\mu_j^2 + \sigma_j^2 - \log \sigma_j^2 - 1)$$

# The Gaussian KL — show, do not derive

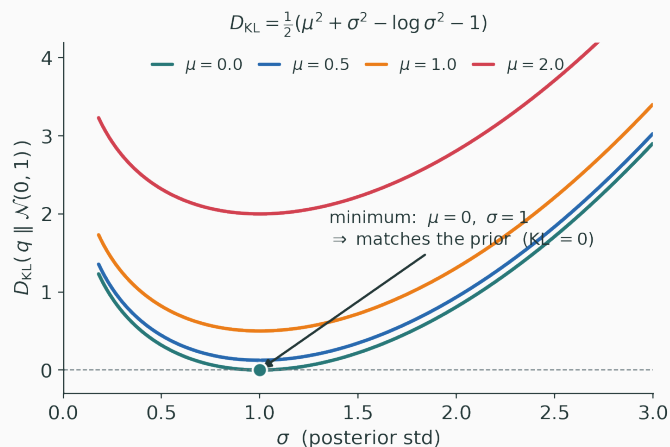
For a diagonal-Gaussian posterior against the standard-normal prior, the KL is **closed form**:

$$D_{\text{KL}}(\mathcal{N}(\mu, \sigma^2) \parallel \mathcal{N}(0, 1)) = \frac{1}{2} \sum_j (\mu_j^2 + \sigma_j^2 - \log \sigma_j^2 - 1)$$



For a diagonal-Gaussian posterior against the standard-normal prior, the KL is **closed form**:

$$D_{\text{KL}}(\mathcal{N}(\mu, \sigma^2) \parallel \mathcal{N}(0, 1)) = \frac{1}{2} \sum_j (\mu_j^2 + \sigma_j^2 - \log \sigma_j^2 - 1)$$



minimized at  $\mu = 0, \sigma = 1$  (KL = 0) — the KL pushes each code toward the prior

The two terms **pull in opposite directions** — the objective finds the balance:

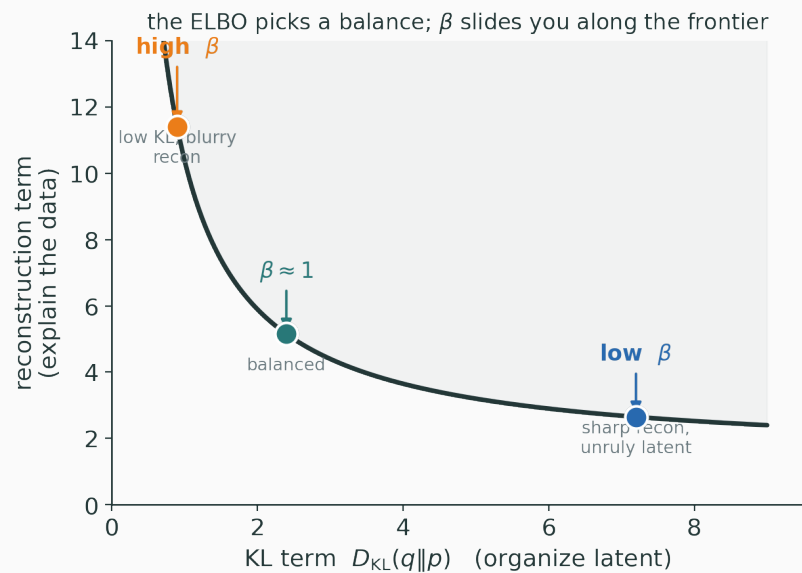
The two terms **pull in opposite directions** — the objective finds the balance:

**Reconstruction** wants every input to keep its **own detailed** code; **KL** wants **all** codes to collapse to **one** Gaussian cloud at the origin.



The two terms **pull in opposite directions** — the objective finds the balance:

**Reconstruction** wants every input to keep its **own detailed** code; **KL** wants **all** codes to collapse to **one** Gaussian cloud at the origin.



# Backpropagation through sampling

---

The ELBO needs a **sample**  $z \sim \mathcal{N}(\mu, \sigma^2)$  inside the forward pass. But:

The ELBO needs a **sample**  $z \sim \mathcal{N}(\mu, \sigma^2)$  inside the forward pass. But:

**Sampling is not differentiable.** A raw draw  $z \sim \mathcal{N}(\mu, \sigma^2)$  seems to block gradients from flowing back into  $\mu$  and  $\sigma$  — the encoder cannot learn.

The ELBO needs a **sample**  $z \sim \mathcal{N}(\mu, \sigma^2)$  inside the forward pass. But:

**Sampling is not differentiable.** A raw draw  $z \sim \mathcal{N}(\mu, \sigma^2)$  seems to block gradients from flowing back into  $\mu$  and  $\sigma$  — the encoder cannot learn.

we need the randomness **without** breaking the chain rule

Move the randomness to an **external** variable, then transform it deterministically:

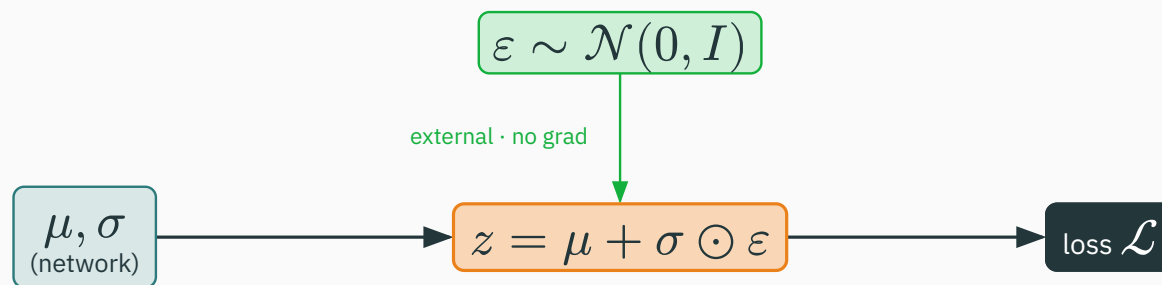
Move the randomness to an **external** variable, then transform it deterministically:

$$\varepsilon \sim \mathcal{N}(0, I), \quad z = \mu + \sigma \odot \varepsilon$$

# The reparameterization trick

Move the randomness to an **external** variable, then transform it deterministically:

$$\varepsilon \sim \mathcal{N}(0, I), \quad z = \mu + \sigma \odot \varepsilon$$



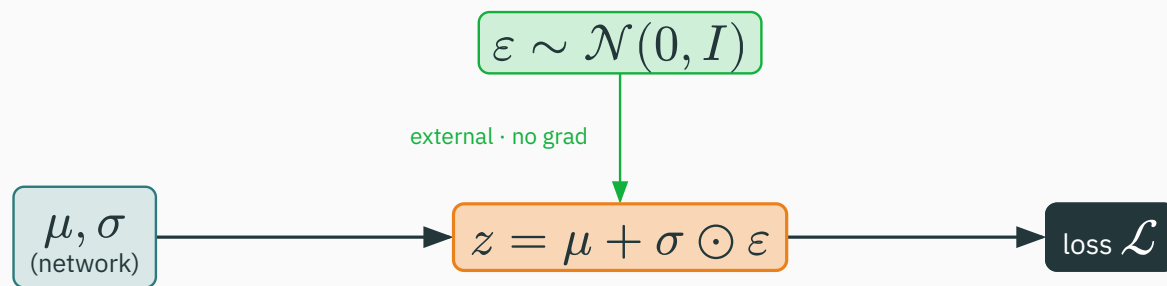
$\nabla$  flows back along this **deterministic** path



# The reparameterization trick

Move the randomness to an **external** variable, then transform it deterministically:

$$\varepsilon \sim \mathcal{N}(0, I), \quad z = \mu + \sigma \odot \varepsilon$$



$\nabla$  flows back along this **deterministic** path

$\varepsilon$  is sampled **outside** the network; the map  $\mu, \sigma \rightarrow z$  is a smooth, differentiable function

Take a 2-D latent with  $\mu = [1, -1]$ ,  $\sigma = [0.5, 0.2]$ , and a draw  $\varepsilon = [-0.4, 1.5]$ :

Take a 2-D latent with  $\mu = [1, -1]$ ,  $\sigma = [0.5, 0.2]$ , and a draw  $\varepsilon = [-0.4, 1.5]$ :

$$z = \mu + \sigma \odot \varepsilon = \begin{pmatrix} 1 + 0.5 \cdot (-0.4) \\ -1 + 0.2 \cdot 1.5 \end{pmatrix} = \begin{pmatrix} \mathbf{0.8} \\ \mathbf{-0.7} \end{pmatrix}$$

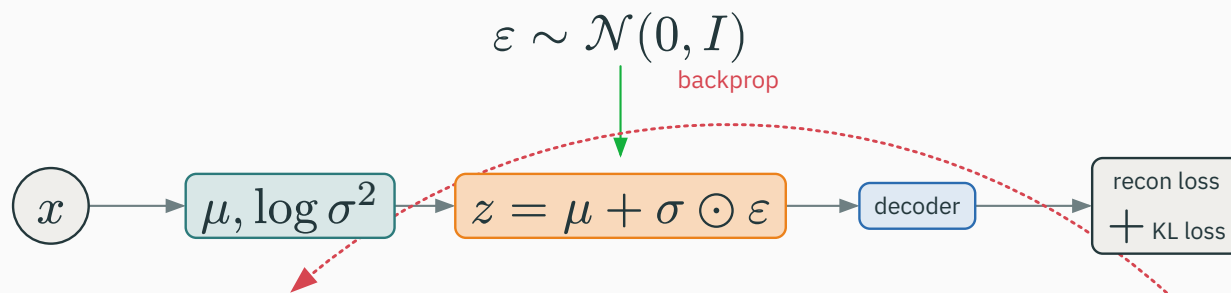
Take a 2-D latent with  $\mu = [1, -1]$ ,  $\sigma = [0.5, 0.2]$ , and a draw  $\varepsilon = [-0.4, 1.5]$ :

$$z = \mu + \sigma \odot \varepsilon = \begin{pmatrix} 1 + 0.5 \cdot (-0.4) \\ -1 + 0.2 \cdot 1.5 \end{pmatrix} = \begin{pmatrix} \mathbf{0.8} \\ \mathbf{-0.7} \end{pmatrix}$$

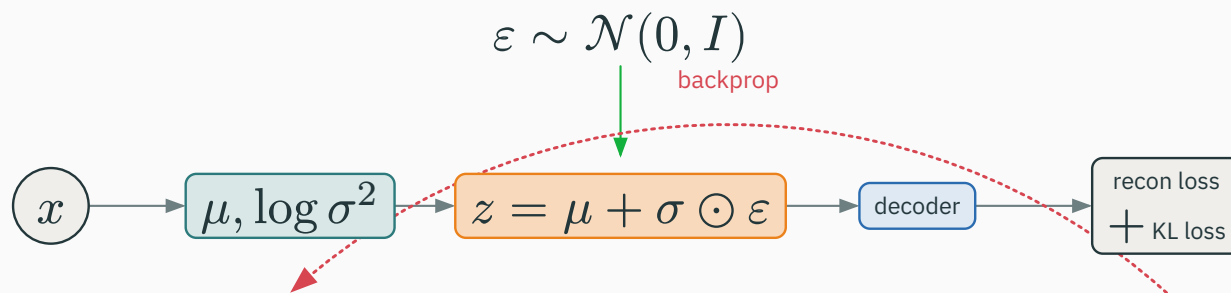
grow  $\sigma \rightarrow$  the same  $\varepsilon$  moves  $z$  **further** from  $\mu$ : more spread, a larger KL

Every step: encode, reparameterize, decode, score both terms, and backprop.

Every step: encode, reparameterize, decode, score both terms, and backprop.



Every step: encode, reparameterize, decode, score both terms, and backprop.



data **enters** the loop; gradients flow through  $z$  into the encoder **and** the decoder

Generation throws away the encoder — just sample the prior and decode:



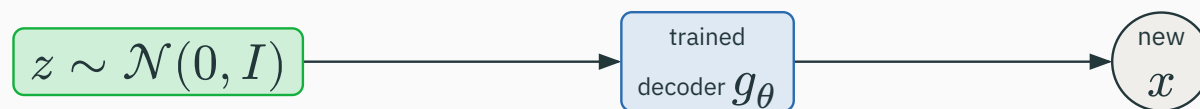
Generation throws away the encoder — just sample the prior and decode:

no encoder, no data at generation time



Generation throws away the encoder — just sample the prior and decode:

no encoder, no data at generation time



**training** begins with data; **generation** begins with  $z \sim \mathcal{N}(0, I)$  — nothing else

Walk a straight line between two codes and decode along the path:

Walk a straight line between two codes and decode along the path:

VAE latent: interpolation decodes to a gradual, valid morph



plain AE: the path crosses empty latent gaps -> invalid decodings



Walk a straight line between two codes and decode along the path:

VAE latent: interpolation decodes to a gradual, valid morph



plain AE: the path crosses empty latent gaps -> invalid decodings



a good latent space yields **gradual semantic change**; a bad one crosses gaps into invalid images

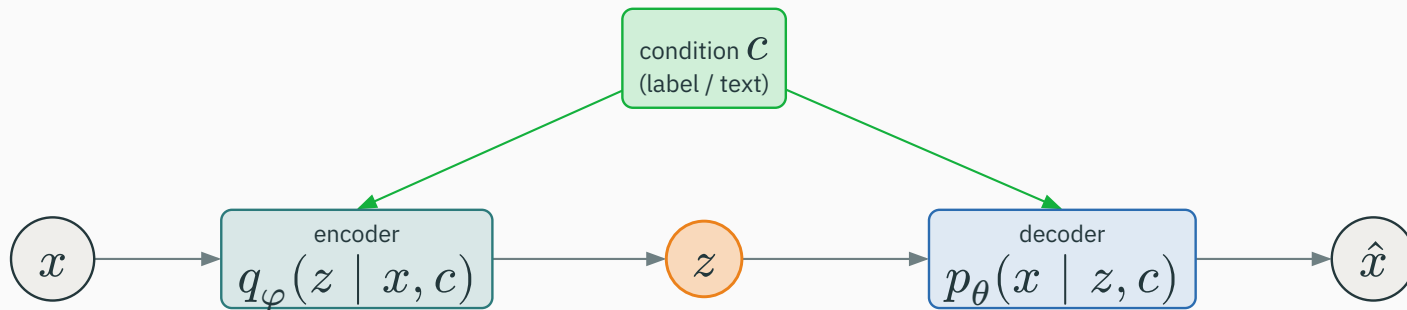
Feed a **condition**  $c$  (a class label, a text embedding) into both networks:

Feed a **condition**  $c$  (a class label, a text embedding) into both networks:

$$q_{\varphi}(z \mid x, c), \quad p_{\theta}(x \mid z, c)$$

Feed a **condition**  $c$  (a class label, a text embedding) into both networks:

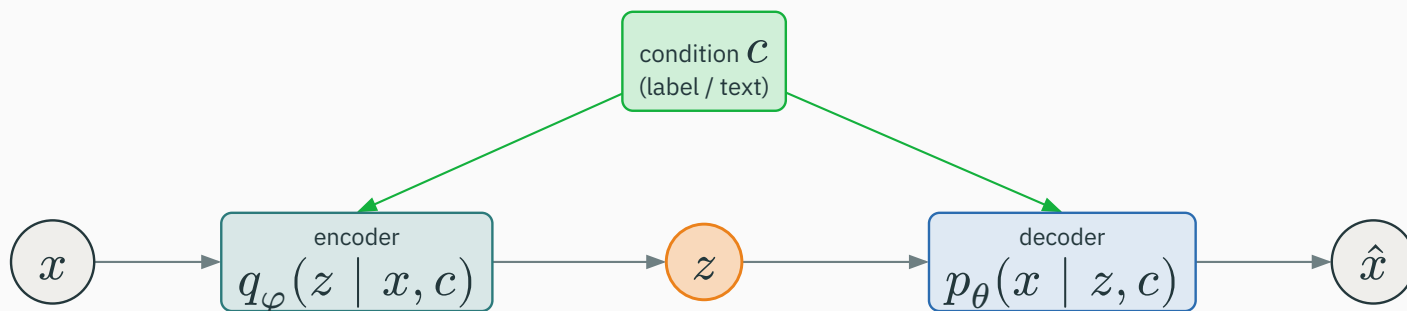
$$q_{\varphi}(z \mid x, c), \quad p_{\theta}(x \mid z, c)$$





Feed a **condition**  $c$  (a class label, a text embedding) into both networks:

$$q_{\varphi}(z \mid x, c), \quad p_{\theta}(x \mid z, c)$$



at generation time: **choose**  $c$ , sample  $z \sim p(z)$ , decode —  $c$  steers **what** is generated

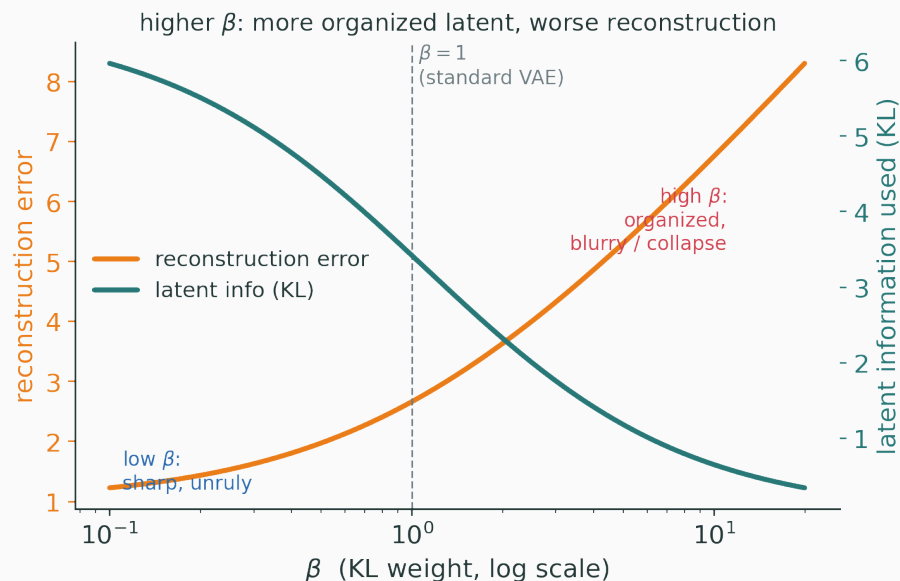
Put a **weight**  $\beta$  on the KL term and dial the balance explicitly:

Put a **weight**  $\beta$  on the KL term and dial the balance explicitly:

$$\mathcal{L} = -\mathbb{E}_{q[\log p_{\theta}(x | z)]} + \beta D_{\text{KL}}(q_{\varphi}(z | x) \parallel p(z))$$

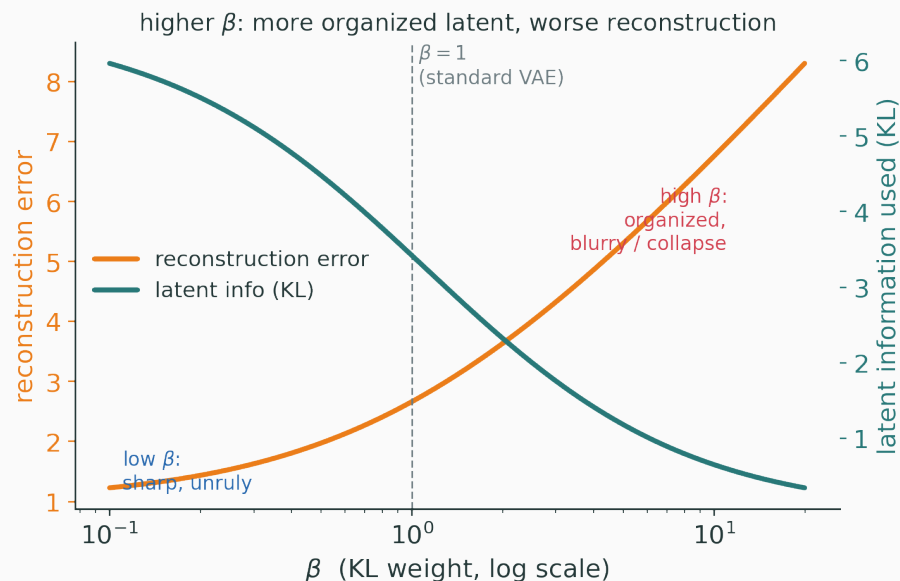
Put a **weight**  $\beta$  on the KL term and dial the balance explicitly:

$$\mathcal{L} = -\mathbb{E}_{q[\log p_{\theta}(x | z)]} + \beta D_{\text{KL}}(q_{\varphi}(z | x) \parallel p(z))$$



Put a **weight**  $\beta$  on the KL term and dial the balance explicitly:

$$\mathcal{L} = -\mathbb{E}_{q[\log p_{\theta}(x | z)]} + \beta D_{\text{KL}}(q_{\varphi}(z | x) \parallel p(z))$$



higher  $\beta \rightarrow$  more organized latent factors, but **softer / worse** reconstructions

# **Honest limitations and the bridge onward**

---

A recurring complaint: VAE samples look **soft**. The cause is the per-pixel likelihood.

A recurring complaint: VAE samples look **soft**. The cause is the per-pixel likelihood.

each sample is a plausible sharp output; a per-pixel likelihood rewards their average





A recurring complaint: VAE samples look **soft**. The cause is the per-pixel likelihood.

each sample is a plausible sharp output; a per-pixel likelihood rewards their average

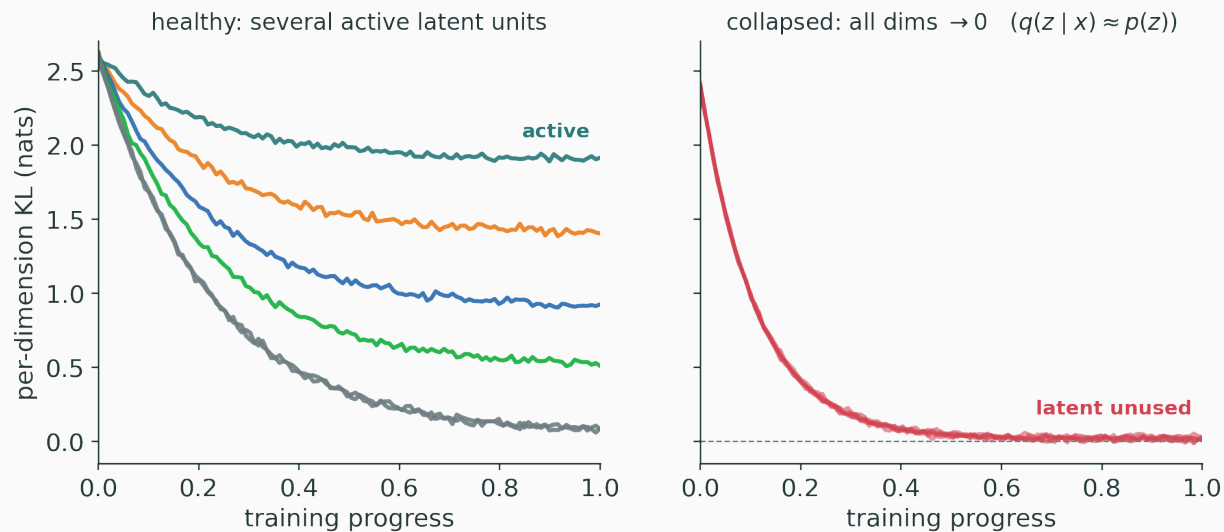


several sharp outputs are all plausible; a per-pixel Gaussian rewards their **average**, which is blurry

The opposite failure — the model **ignores** the latent entirely:

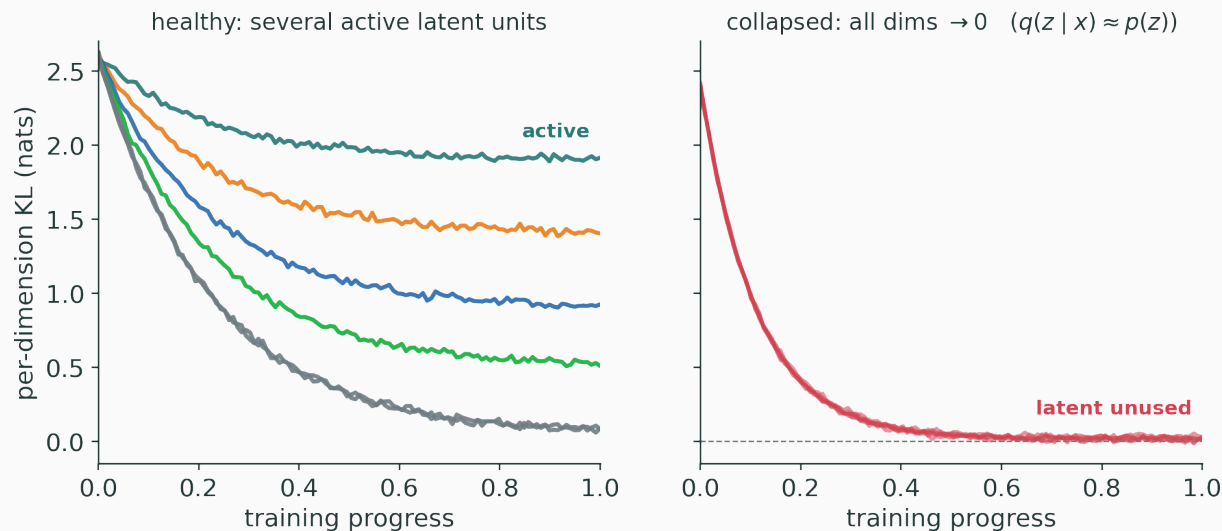


The opposite failure — the model **ignores** the latent entirely:





The opposite failure — the model **ignores** the latent entirely:



Two routes to the same failure: an over-powerful **autoregressive** decoder can predict  $x$  from its own context alone, and — more often for a plain feed-forward

decoder like ours — the KL term simply dominates. Either way  $q_\varphi(z \mid x) \rightarrow p(z)$  for every  $x$  and the latent goes **unused** — not magically regularized.

|                | <b>Autoencoder</b>   | <b>VAE</b>                   |
|----------------|----------------------|------------------------------|
| code           | one point            | a distribution               |
| sampling       | unreliable (islands) | defined by the prior         |
| objective      | reconstruction       | likelihood lower bound       |
| typical output | useful features      | smooth, often softer samples |



|                | <b>Autoencoder</b>   | <b>VAE</b>                   |
|----------------|----------------------|------------------------------|
| code           | one point            | a distribution               |
| sampling       | unreliable (islands) | defined by the prior         |
| objective      | reconstruction       | likelihood lower bound       |
| typical output | useful features      | smooth, often softer samples |

the KL term is the single change that turns a point-code into a **sampleable** latent space

Both use a latent  $z$ , but their **learning signals** differ radically:

Both use a latent  $z$ , but their **learning signals** differ radically:

**VAE** — an **explicit** probabilistic story; maximize a likelihood bound (reconstruction + KL).

**GAN** — **no** explicit likelihood; a generator learns by trying to **fool a discriminator**.

Both use a latent  $z$ , but their **learning signals** differ radically:

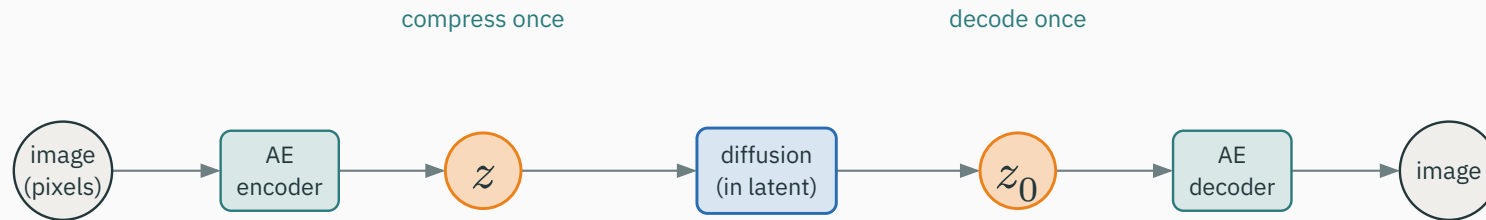
**VAE** — an **explicit** probabilistic story; maximize a likelihood bound (reconstruction + KL).

**GAN** — **no** explicit likelihood; a generator learns by trying to **fool a discriminator**.

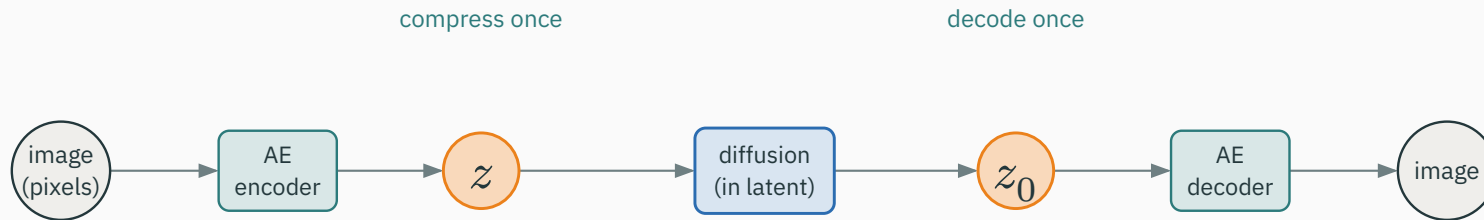
VAEs optimize a bound on  $p(x)$ ; GANs play a game — sharper samples, no density

Modern **latent diffusion** leans directly on the autoencoder:

Modern **latent diffusion** leans directly on the autoencoder:



Modern **latent diffusion** leans directly on the autoencoder:



compress pixels to a compact latent, run diffusion **there**, decode once — the AE is an active component

On a small VAE diagram, label the three steps and predict a failure:



On a small VAE diagram, label the three steps and predict a failure:

**Label** — (1) the **reconstruction** step, (2) the **KL** step, (3) the **sampling** step.

**Predict** — what happens to the latent space if we **remove the KL term** entirely?

On a small VAE diagram, label the three steps and predict a failure:

**Label** — (1) the **reconstruction** step, (2) the **KL** step, (3) the **sampling** step.

**Predict** — what happens to the latent space if we **remove the KL term** entirely?

no KL → the codes drift back into islands → prior samples decode to nonsense again

**VAEs make latent spaces **sampleable** by **regularizing** them**

**VAEs make latent spaces *sampleable* by *regularizing* them**

Next: **GANs** trade the explicit likelihood for an **adversarial game**

— giving up density estimation to chase **sharper** samples

# Deeper: latent-variable models

---

PCA finds the best **linear** low-dimensional reconstruction:

PCA finds the best **linear** low-dimensional reconstruction:

$$x \longrightarrow W^\top x \longrightarrow WW^\top x$$

PCA finds the best **linear** low-dimensional reconstruction:

$$x \longrightarrow W^{\top} x \longrightarrow W W^{\top} x$$

An autoencoder replaces the two linear maps with **neural networks**, so it can bend around **nonlinear** manifolds.



PCA finds the best **linear** low-dimensional reconstruction:

$$x \longrightarrow W^\top x \longrightarrow WW^\top x$$

An autoencoder replaces the two linear maps with **neural networks**, so it can bend around **nonlinear** manifolds.

an intuition link, **not** a claim that every autoencoder is PCA — the nonlinearity is the point

Track the dimensions as data flows through the encoder:

Track the dimensions as data flows through the encoder:

$$x \in \mathbb{R}^{H \times W \times C} \xrightarrow{\text{encoder}} z \in \mathbb{R}^d$$

Track the dimensions as data flows through the encoder:

$$x \in \mathbb{R}^{H \times W \times C} \xrightarrow{\text{encoder}} z \in \mathbb{R}^d$$

A concrete example — a  $32 \times 32$  grayscale image compressed to a 16-dimensional code:

$$x \in \mathbb{R}^{1,024} \longrightarrow z \in \mathbb{R}^{16} \quad (\text{a } 64 \times \text{squeeze})$$

Track the dimensions as data flows through the encoder:

$$x \in \mathbb{R}^{H \times W \times C} \xrightarrow{\text{encoder}} z \in \mathbb{R}^d$$

A concrete example — a  $32 \times 32$  grayscale image compressed to a 16-dimensional code:

$$x \in \mathbb{R}^{1,024} \longrightarrow z \in \mathbb{R}^{16} \quad (\text{a } 64 \times \text{squeeze})$$

what must the 16 numbers preserve so the decoder can rebuild all 1,024 pixels?

| Code dimension                  | Pressure on the model                    |
|---------------------------------|------------------------------------------|
| $d < \dim(x)$ (undercomplete)   | must compress salient structure          |
| $d \geq \dim(x)$ (overcomplete) | may learn an identity shortcut           |
| overcomplete + regularization   | can still learn sparse / robust features |

| Code dimension                  | Pressure on the model                    |
|---------------------------------|------------------------------------------|
| $d < \dim(x)$ (undercomplete)   | must compress salient structure          |
| $d \geq \dim(x)$ (overcomplete) | may learn an identity shortcut           |
| overcomplete + regularization   | can still learn sparse / robust features |

the bottleneck can be **dimensional**, **stochastic**, **architectural**, or **regularization**-based

Grounding the familiar loss in probability. If the decoder is Gaussian:



Grounding the familiar loss in probability. If the decoder is Gaussian:

$$p_{\theta}(x \mid z) = \mathcal{N}(g_{\theta}(z), \sigma_x^2 I)$$

Grounding the familiar loss in probability. If the decoder is Gaussian:

$$p_{\theta}(x \mid z) = \mathcal{N}(g_{\theta}(z), \sigma_x^2 I)$$

then the negative log-likelihood is **proportional to squared error**:

$$-\log p_{\theta}(x \mid z) = \frac{1}{2\sigma_x^2} \|x - g_{\theta}(z)\|^2 + \text{const}$$

Grounding the familiar loss in probability. If the decoder is Gaussian:

$$p_{\theta}(x \mid z) = \mathcal{N}(g_{\theta}(z), \sigma_x^2 I)$$

then the negative log-likelihood is **proportional to squared error**:

$$-\log p_{\theta}(x \mid z) = \frac{1}{2\sigma_x^2} \|x - g_{\theta}(z)\|^2 + \text{const}$$

MSE is a **modeling choice** (a Gaussian decoder), not a law of nature

# Data type determines the decoder output

| <b>Data</b>                  | <b>Decoder distribution</b> | <b>Reconstruction term</b> |
|------------------------------|-----------------------------|----------------------------|
| normalized continuous pixels | Gaussian                    | squared error              |
| binary image                 | Bernoulli                   | binary cross-entropy       |
| tokens                       | categorical                 | cross-entropy              |

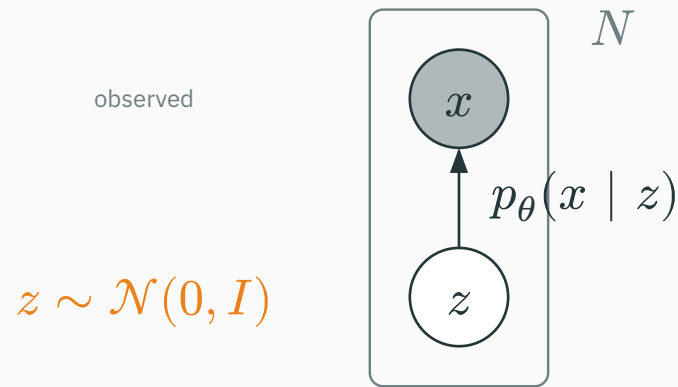
# Data type determines the decoder output

| Data                         | Decoder distribution | Reconstruction term  |
|------------------------------|----------------------|----------------------|
| normalized continuous pixels | Gaussian             | squared error        |
| binary image                 | Bernoulli            | binary cross-entropy |
| tokens                       | categorical          | cross-entropy        |

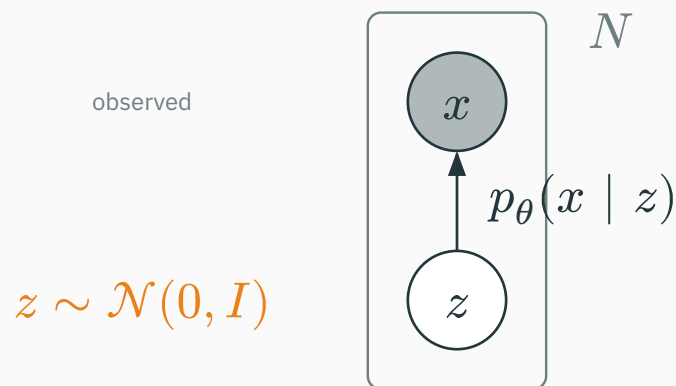
the decoder predicts **distribution parameters** — match the distribution to the data type

At generation time the arrow is simple: cause  $\rightarrow$  observation.

At generation time the arrow is simple: cause  $\rightarrow$  observation.



At generation time the arrow is simple: cause  $\rightarrow$  observation.



training **inverts** it: infer the hidden cause  $z$  that plausibly generated each observed  $x$



# Marginal likelihood is an integral

The quantity we truly want to maximize:

# Marginal likelihood is an integral

The quantity we truly want to maximize:

$$p_{\theta}(x) = \int p_{\theta}(x \mid z) p(z) dz$$

The quantity we truly want to maximize:

$$p_{\theta}(x) = \int p_{\theta}(x \mid z) p(z) dz$$

For a nonlinear neural decoder and high-dimensional  $z$ , this integral has **no closed form** and cannot be evaluated exactly — hence the lower bound.

Choose a tractable family  $q_\varphi(z \mid x)$  to **imitate** the true posterior  $p_\theta(z \mid x)$ :

Choose a tractable family  $q_{\varphi}(z \mid x)$  to **imitate** the true posterior  $p_{\theta}(z \mid x)$ :

**Amortized inference.** One encoder network  $\varphi$  handles **every** input — instead of solving a fresh optimization for each image, we **amortize** the cost into shared weights.

Choose a tractable family  $q_{\varphi}(z \mid x)$  to **imitate** the true posterior  $p_{\theta}(z \mid x)$ :

**Amortized inference.** One encoder network  $\varphi$  handles **every** input — instead of solving a fresh optimization for each image, we **amortize** the cost into shared weights.

classical VI optimizes per-datapoint; the VAE trains one network to predict  $q$  for all of them

The identity behind the bound's **name**:

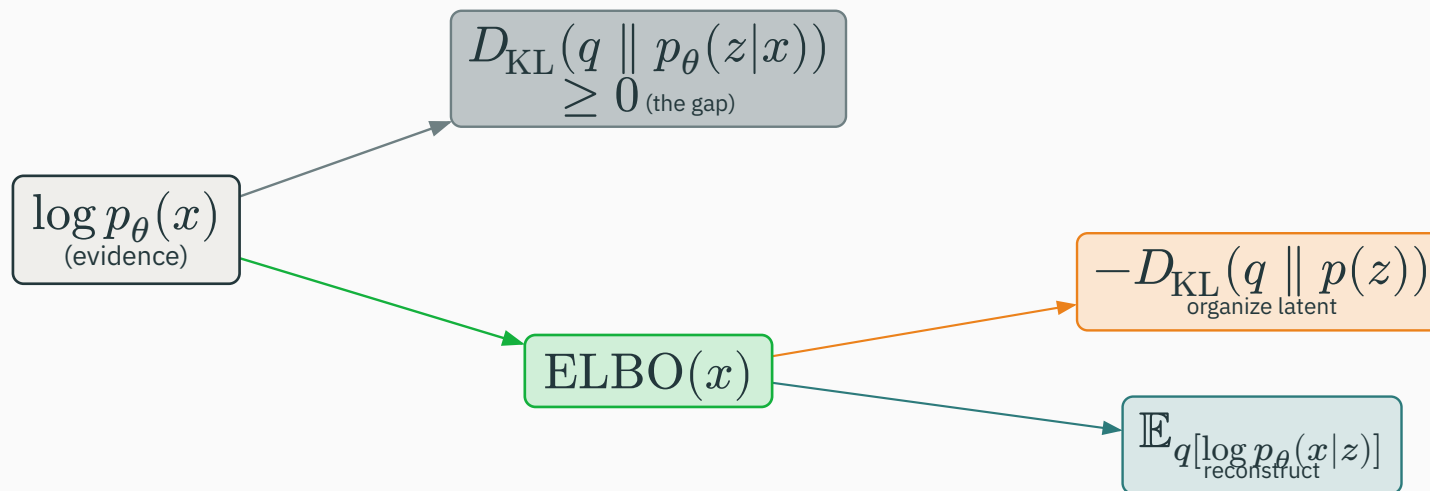
The identity behind the bound's **name**:

$$\log p_{\theta}(x) = \text{ELBO}(x) + D_{\text{KL}}(q_{\varphi}(z \mid x) \parallel p_{\theta}(z \mid x))$$



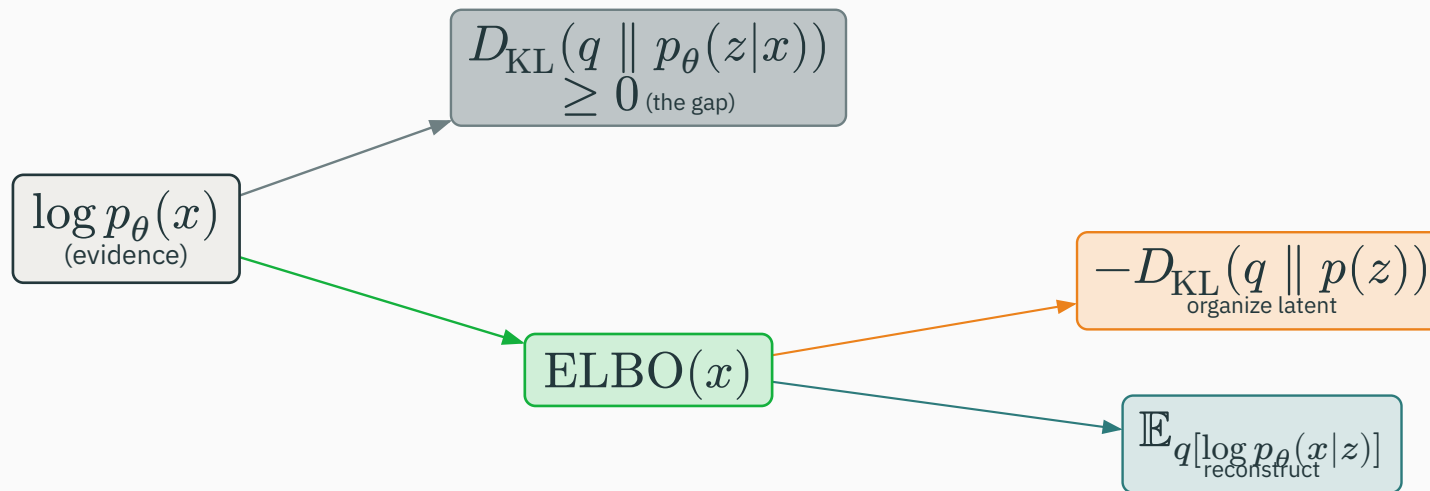
The identity behind the bound's **name**:

$$\log p_{\theta}(x) = \text{ELBO}(x) + D_{\text{KL}}(q_{\varphi}(z | x) \parallel p_{\theta}(z | x))$$



The identity behind the bound's **name**:

$$\log p_{\theta}(x) = \text{ELBO}(x) + D_{\text{KL}}(q_{\varphi}(z | x) \parallel p_{\theta}(z | x))$$



the KL gap is  $\geq 0$ , so  $\text{ELBO} \leq \log p_{\theta}(x)$  – a genuine **lower bound**

Start from the definition and substitute the joint:

Start from the definition and substitute the joint:

$$\text{ELBO}(x) = \mathbb{E}_{q_{\varphi}(z \mid x)} [\log p_{\theta}(x, z) - \log q_{\varphi}(z \mid x)]$$

Start from the definition and substitute the joint:

$$\text{ELBO}(x) = \mathbb{E}_{q_{\varphi}(z \mid x)} [\log p_{\theta}(x, z) - \log q_{\varphi}(z \mid x)]$$

$$p_{\theta}(x, z) = p_{\theta}(x \mid z) p(z)$$

Start from the definition and substitute the joint:

$$\text{ELBO}(x) = \mathbb{E}_{q_{\varphi}(z \mid x)} [\log p_{\theta}(x, z) - \log q_{\varphi}(z \mid x)]$$

$$p_{\theta}(x, z) = p_{\theta}(x \mid z) p(z)$$

substitute the joint, then split the logarithm — the next slide collects the pieces

# Rearrange into the two familiar terms

Grouping the split logarithm gives back the working objective:

# Rearrange into the two familiar terms

Grouping the split logarithm gives back the working objective:

$$\text{ELBO} = \underbrace{\mathbb{E}_{q[\log p_{\theta}(x | z)]}}_{\text{explain the data}} - \underbrace{D_{\text{KL}}(q_{\varphi}(z | x) \parallel p(z))}_{\text{organize the latent}}$$



# Rearrange into the two familiar terms

Grouping the split logarithm gives back the working objective:

$$\text{ELBO} = \underbrace{\mathbb{E}_{q[\log p_{\theta}(x | z)]}}_{\text{explain the data}} - \underbrace{D_{\text{KL}}(q_{\varphi}(z | x) \parallel p(z))}_{\text{organize the latent}}$$

same two terms as Part II — now **derived**, not asserted

# Which KL is which?

Two KL expressions appear — students routinely conflate them:

# Which KL is which?

Two KL expressions appear — students routinely conflate them:

$D_{\text{KL}}(q(z \mid x) \parallel p(z))$  regularizes each code toward the **prior** (in the loss)

# Which KL is which?

Two KL expressions appear — students routinely conflate them:

$D_{\text{KL}}(q(z \mid x) \parallel p(z))$  regularizes each code toward the **prior** (in the loss)

$D_{\text{KL}}(q(z \mid x) \parallel p(z \mid x))$  the **gap** between ELBO and true log-likelihood

# Which KL is which?

Two KL expressions appear — students routinely conflate them:

$D_{\text{KL}}(q(z \mid x) \parallel p(z))$  regularizes each code toward the **prior** (in the loss)

$D_{\text{KL}}(q(z \mid x) \parallel p(z \mid x))$  the **gap** between ELBO and true log-likelihood

one is a **term you optimize**; the other is **how loose the bound is**

We **maximize** the ELBO, equivalently **minimize** its negative:

We **maximize** the ELBO, equivalently **minimize** its negative:

$$\mathcal{L}_{\text{VAE}} = \underbrace{-\mathbb{E}_{q[\log p_{\theta}(x | z)]}}_{\text{“reconstruction loss”}} + \underbrace{D_{\text{KL}}(q_{\varphi}(z | x) \parallel p(z))}_{\text{“KL loss”}}$$

We **maximize** the ELBO, equivalently **minimize** its negative:

$$\mathcal{L}_{\text{VAE}} = \underbrace{-\mathbb{E}_{q[\log p_{\theta}(x | z)]}}_{\text{“reconstruction loss”}} + \underbrace{D_{\text{KL}}(q_{\varphi}(z | x) \parallel p(z))}_{\text{“KL loss”}}$$

in code the first term is literally the MSE or BCE; the second is the closed-form Gaussian KL



The expectation over  $q_{\varphi}(z \mid x)$  is approximated by **sampling** latent vectors:

The expectation over  $q_\varphi(z \mid x)$  is approximated by **sampling** latent vectors:

$$\mathbb{E}_{q[\cdot]} \approx \frac{1}{K} \sum_{k=1}^K [\cdot]_{z^{(k)}}, \quad z^{(k)} = \mu + \sigma \odot \varepsilon^{(k)}$$

The expectation over  $q_\varphi(z \mid x)$  is approximated by **sampling** latent vectors:

$$\mathbb{E}_{q[\cdot]} \approx \frac{1}{K} \sum_{k=1}^K [\cdot]_{z^{(k)}}, \quad z^{(k)} = \mu + \sigma \odot \varepsilon^{(k)}$$

**one** reparameterized sample per input ( $K = 1$ ) is usually enough for the standard objective

# Why output log-variance?

Networks predict  $\ell = \log \sigma^2$  rather than  $\sigma$  directly, then recover:

# Why output log-variance?

Networks predict  $\ell = \log \sigma^2$  rather than  $\sigma$  directly, then recover:

$$\ell = \log \sigma^2 \quad \longrightarrow \quad \sigma = \exp\left(\frac{1}{2} \ell\right)$$

# Why output log-variance?

Networks predict  $\ell = \log \sigma^2$  rather than  $\sigma$  directly, then recover:

$$\ell = \log \sigma^2 \quad \longrightarrow \quad \sigma = \exp\left(\frac{1}{2} \ell\right)$$

$\ell$  ranges over **all** of  $\mathbb{R}$  — safe for an unconstrained linear output — while  $\exp(\cdot)$  keeps the variance **strictly positive**. No clamping, no  $\log 0$ .

# Why output log-variance?

Networks predict  $\ell = \log \sigma^2$  rather than  $\sigma$  directly, then recover:

$$\ell = \log \sigma^2 \quad \longrightarrow \quad \sigma = \exp\left(\frac{1}{2} \ell\right)$$

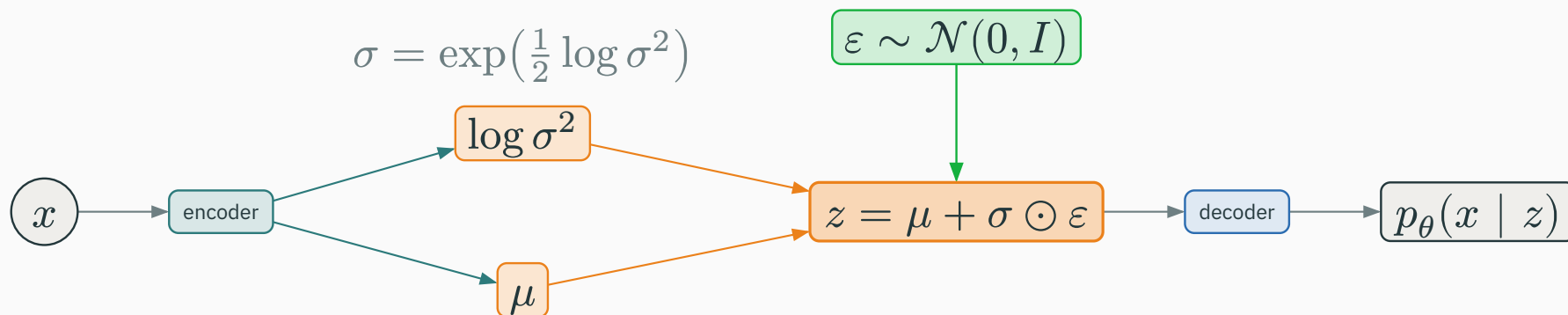
$\ell$  ranges over **all** of  $\mathbb{R}$  — safe for an unconstrained linear output — while  $\exp(\cdot)$  keeps the variance **strictly positive**. No clamping, no  $\log 0$ .

$$\text{e.g. } \ell = -1.386 \rightarrow \sigma = \exp(-0.693) = 0.5$$

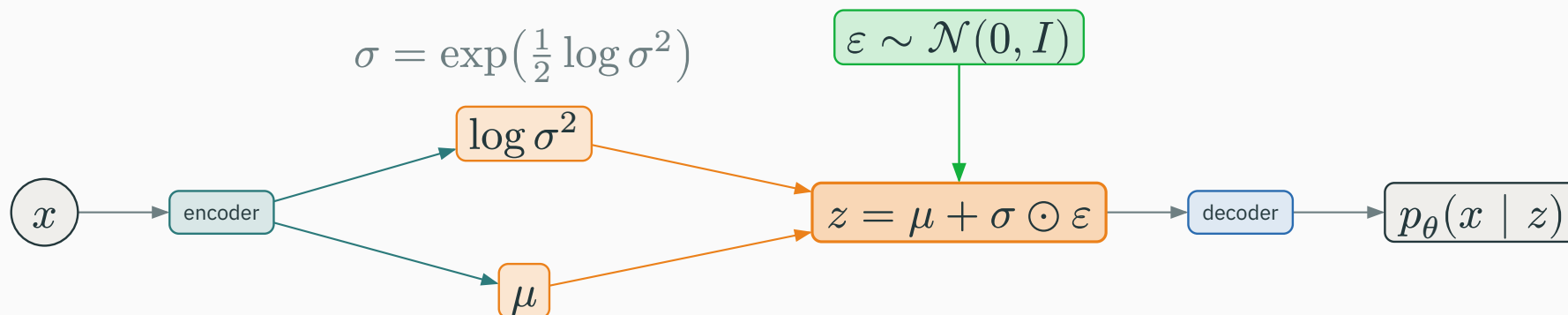
The full graph with the log-variance parameterization:



The full graph with the log-variance parameterization:



The full graph with the log-variance parameterization:



$\varepsilon$  is drawn independently; gradients flow through the **deterministic** transform  $\mu, \sigma \rightarrow z$

Take the same encoded point,  $\mu = [0.8, -0.7]$ ,  $\sigma = [0.5, 0.2]$ , with  $\sigma_x^2 = 1$  and  $\|x - \hat{x}\|^2 = 3.0$ :

Take the same encoded point,  $\mu = [0.8, -0.7]$ ,  $\sigma = [0.5, 0.2]$ , with  $\sigma_x^2 = 1$  and  $\|x - \hat{x}\|^2 = 3.0$ :

$$\text{KL} = \frac{1}{2}[(0.64 + 0.25 + 1.386 - 1) + (0.49 + 0.04 + 3.219 - 1)] = \mathbf{2.01}$$

Take the same encoded point,  $\mu = [0.8, -0.7]$ ,  $\sigma = [0.5, 0.2]$ , with  $\sigma_x^2 = 1$  and  $\|x - \hat{x}\|^2 = 3.0$ :

$$\text{KL} = \frac{1}{2}[(0.64 + 0.25 + 1.386 - 1) + (0.49 + 0.04 + 3.219 - 1)] = \mathbf{2.01}$$

$$\text{recon} = \frac{1}{2}\|x - \hat{x}\|^2 = 1.50, \quad \mathcal{L}_{\text{VAE}} = 1.50 + 2.01 = \mathbf{3.51}$$

Take the same encoded point,  $\mu = [0.8, -0.7]$ ,  $\sigma = [0.5, 0.2]$ , with  $\sigma_x^2 = 1$  and  $\|x - \hat{x}\|^2 = 3.0$ :

$$\text{KL} = \frac{1}{2}[(0.64 + 0.25 + 1.386 - 1) + (0.49 + 0.04 + 3.219 - 1)] = \mathbf{2.01}$$

$$\text{recon} = \frac{1}{2}\|x - \hat{x}\|^2 = 1.50, \quad \mathcal{L}_{\text{VAE}} = 1.50 + 2.01 = \mathbf{3.51}$$

both terms are in **nats**; the optimizer trades one against the other

| Operation            | Latent source          | Question answered                 |
|----------------------|------------------------|-----------------------------------|
| reconstruction       | $z \sim q(z \mid x)$   | can the model explain this input? |
| unconditional sample | $z \sim p(z)$          | what does the model generate?     |
| interpolation        | path between two codes | is the latent geometry smooth?    |

| Operation            | Latent source          | Question answered                 |
|----------------------|------------------------|-----------------------------------|
| reconstruction       | $z \sim q(z \mid x)$   | can the model explain this input? |
| unconditional sample | $z \sim p(z)$          | what does the model generate?     |
| interpolation        | path between two codes | is the latent geometry smooth?    |

where  $z$  comes from changes the meaning of the decoder's output



Two candidate encoders for the **same** data:

Two candidate encoders for the **same** data:

| Encoder                     | Reconstruction | KL  |
|-----------------------------|----------------|-----|
| highly individualized codes | 1              | 100 |
| codes near the prior        | 10             | 1   |

Two candidate encoders for the **same** data:

| Encoder                     | Reconstruction | KL  |
|-----------------------------|----------------|-----|
| highly individualized codes | 1              | 100 |
| codes near the prior        | 10             | 1   |

totals 101 vs 11: neither is “obviously best” — the **objective** picks the balance, not reconstruction alone

Encode a dataset into  $(z_1, z_2)$  and decode a **grid** of latent points:

# Visualize a 2-D latent space

Encode a dataset into  $(z_1, z_2)$  and decode a **grid** of latent points:



# Visualize a 2-D latent space

Encode a dataset into  $(z_1, z_2)$  and decode a **grid** of latent points:



a good map shows **clusters**, **smooth directions**, **empty regions**, and where **prior samples** land

Hold every coordinate fixed and move a **single** latent dimension:

Hold every coordinate fixed and move a **single** latent dimension:

$$z_j \leftarrow z_j + \delta, \quad (\text{decode several values of } \delta)$$



Hold every coordinate fixed and move a **single** latent dimension:

$$z_j \leftarrow z_j + \delta, \quad (\text{decode several values of } \delta)$$

Interpretable changes along  $z_j$  are **interesting evidence** about the learned representation — but **not proof** that the factors are perfectly disentangled.

For a label or text condition  $c$ , condition **both** networks:

For a label or text condition  $c$ , condition **both** networks:

$$q_{\varphi}(z \mid x, c), \quad p_{\theta}(x \mid z, c)$$

For a label or text condition  $c$ , condition **both** networks:

$$q_{\varphi}(z \mid x, c), \quad p_{\theta}(x \mid z, c)$$

**Controlled by  $c$**  — the attribute you set (the digit class, the caption).

**Still diverse through  $z$**  — the style, stroke, and variation **within** that attribute.

For a label or text condition  $c$ , condition **both** networks:

$$q_{\varphi}(z \mid x, c), \quad p_{\theta}(x \mid z, c)$$

**Controlled by  $c$**  — the attribute you set (the digit class, the caption).

**Still diverse through  $z$**  — the style, stroke, and variation **within** that attribute.

sample: choose  $c$ , draw  $z \sim p(z)$ , decode  $p_{\theta}(x \mid z, c)$

$$\mathcal{L} = -\mathbb{E}_{q[\log p_{\theta}(x | z)]} + \beta D_{\text{KL}}(q \parallel p)$$

$$\mathcal{L} = -\mathbb{E}_{q[\log p_{\theta}(x | z)]} + \beta D_{\text{KL}}(q \parallel p)$$

Increasing  $\beta$  strengthens the pressure toward a simple prior.

$$\mathcal{L} = -\mathbb{E}_{q[\log p_{\theta}(x | z)]} + \beta D_{\text{KL}}(q \parallel p)$$

Increasing  $\beta$  strengthens the pressure toward a simple prior.

It **can** encourage more separated factors on curated data — but “disentangled” is **not** a guaranteed, universal outcome, and reconstruction suffers.



# Posterior collapse, diagnose it

The symptom to watch for:

The symptom to watch for:

$$q_{\varphi}(z \mid x) \approx p(z) \quad \text{for every } x$$

The symptom to watch for:

$$q_{\varphi}(z \mid x) \approx p(z) \quad \text{for every } x$$

A powerful decoder predicts  $x$  well from context alone, so the KL term drives every posterior to the prior. The latent has become **unused** — inspect per-dimension KL, not just reconstruction.

Design choices, not a new module:

Design choices, not a new module:

- **weaken or schedule** the KL pressure early in training (KL annealing);
- **constrain** the decoder's capacity so it cannot ignore  $z$ ;
- **force** more information through the latent (e.g. a free-bits floor);
- **inspect** latent-use statistics, not only reconstructions.

Design choices, not a new module:

- **weaken or schedule** the KL pressure early in training (KL annealing);
- **constrain** the decoder's capacity so it cannot ignore  $z$ ;
- **force** more information through the latent (e.g. a free-bits floor);
- **inspect** latent-use statistics, not only reconstructions.

the fixes all **rebalance** how much the model is allowed to rely on the decoder alone

Some autoencoders encode into a **finite codebook** instead of a continuous Gaussian:

Some autoencoders encode into a **finite codebook** instead of a continuous Gaussian:

continuous VAE latent      versus      discrete codebook token



Some autoencoders encode into a **finite codebook** instead of a continuous Gaussian:

continuous VAE latent      versus      discrete codebook token

this previews **tokenized** visual representations — the basis of many image + video generators

Some autoencoders encode into a **finite codebook** instead of a continuous Gaussian:

continuous VAE latent      versus      discrete codebook token

this previews **tokenized** visual representations — the basis of many image + video generators

## INTERACTIVE

↗ [nipunbatra.github.io/interactive-articles/vq-codebook](https://nipunbatra.github.io/interactive-articles/vq-codebook)

A codebook widget: snap continuous encodings to their nearest discrete code vectors.

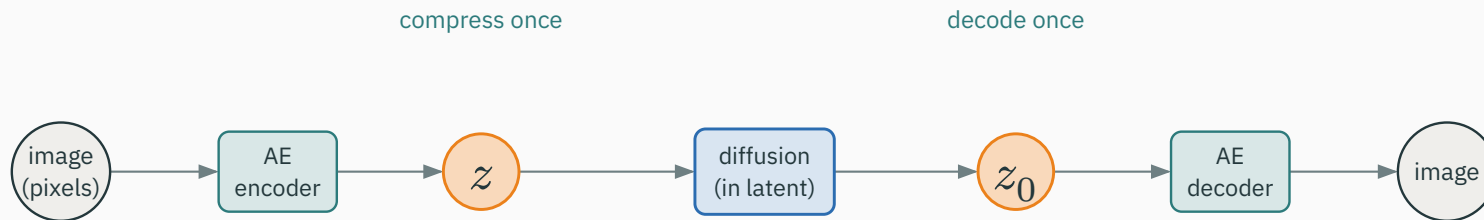
Latent diffusion uses the autoencoder to change **what space** the generator works in:

Latent diffusion uses the autoencoder to change **what space** the generator works in:

pixels  $\longrightarrow$  compact continuous latents

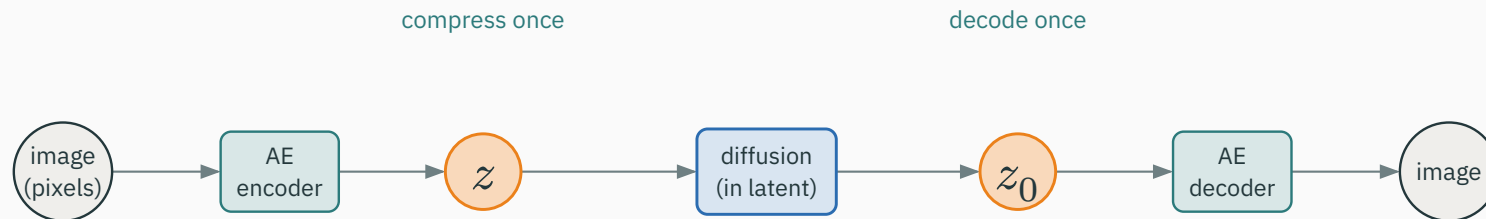
Latent diffusion uses the autoencoder to change **what space** the generator works in:

pixels  $\longrightarrow$  compact continuous latents



Latent diffusion uses the autoencoder to change **what space** the generator works in:

pixels  $\longrightarrow$  compact continuous latents



the VAE is not historical background — it is an **active systems component** of latent diffusion

# How should generative models be evaluated?

No single number suffices — use several lenses:

# How should generative models be evaluated?

No single number suffices — use several lenses:

- **sample quality and diversity** (are outputs realistic **and** varied?);
- **reconstruction / likelihood** metrics where appropriate;
- **downstream utility** of the learned features;
- **human evaluation** with a transparent protocol;
- **memorization and data-provenance** checks.



# How should generative models be evaluated?

No single number suffices — use several lenses:

- **sample quality and diversity** (are outputs realistic **and** varied?);
- **reconstruction / likelihood** metrics where appropriate;
- **downstream utility** of the learned features;
- **human evaluation** with a transparent protocol;
- **memorization and data-provenance** checks.

a model can score well on one lens and fail badly on another

| Misconception                          | Correction                                            |
|----------------------------------------|-------------------------------------------------------|
| “the encoder outputs one latent point” | it outputs a <b>distribution’s parameters</b>         |
| “KL makes samples better directly”     | it <b>organizes codes</b> to match the sampling prior |
| “blurry samples prove VAEs are wrong”  | they reflect likelihood / output trade-offs           |
| “interpolation proves understanding”   | it is a <b>diagnostic</b> , not proof                 |

Test yourself before the next lecture:

Test yourself before the next lecture:

1. Why is an ordinary autoencoder **not guaranteed** to generate from a random  $z$ ?

Test yourself before the next lecture:

1. Why is an ordinary autoencoder **not guaranteed** to generate from a random  $z$ ?
2. What does the KL term **compare** in a VAE?

Test yourself before the next lecture:

1. Why is an ordinary autoencoder **not guaranteed** to generate from a random  $z$ ?
2. What does the KL term **compare** in a VAE?
3. Why does  $z = \mu + \sigma \odot \varepsilon$  **permit** backpropagation?

Test yourself before the next lecture:

1. Why is an ordinary autoencoder **not guaranteed** to generate from a random  $z$ ?
2. What does the KL term **compare** in a VAE?
3. Why does  $z = \mu + \sigma \odot \varepsilon$  **permit** backpropagation?
4. What is **posterior collapse**, and how would you **detect** it?

**VAE**: explain data through a latent **likelihood**



**GAN**: learn realism through an adversarial **critic**



**VAE**: explain data through a latent **likelihood**



**GAN**: learn realism through an adversarial **critic**

the next model **removes** the explicit reconstruction target and introduces a **game**

A VAE turns a code into a **sampleable latent variable**: reconstruction + a KL-regularized prior = a generative model.

Next: **Generative Adversarial Networks** → diffusion → modern generative models.