

Aprendizado Profundo

Autoencoders e Variational Autoencoders

Lucas Silveira Kupssinskü

Machine Learning Theory and Applications Laboratory
PUCRS

agosto de 2026

Visão Geral

1. Autoencoders
2. Variational Autoencoders
3. Treinamento e ELBO
4. Reparametrization Trick
5. VAE em ação
6. Panorama

Visão Geral

1. **Autoencoders**
2. Variational Autoencoders
3. Treinamento e ELBO
4. Reparametrization Trick
5. VAE em ação
6. Panorama

Autoencoders

Um *autoencoder* é uma rede neural treinada na tarefa de **reconstrução da própria entrada**.

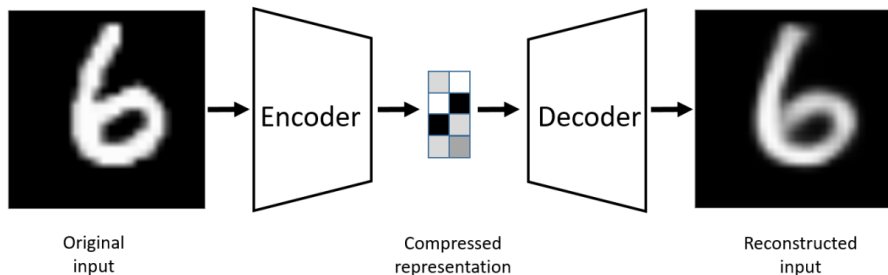
- Uma das primeiras formas de aprendizado **auto-supervisionado** / não supervisionado.
- A ideia é criar uma **representação útil** das instâncias de interesse na etapa intermediária (gargalo).

Objetivo

$$\operatorname{argmin}_{\theta, \phi} J(\mathbf{X}, f(g(\mathbf{X}, \phi), \theta))$$

onde g é o *encoder* (parâmetros ϕ) e f é o *decoder* (parâmetros θ).

Autoencoders – Pipeline



- *Encoder* g reduz a entrada a uma representação compacta (o **latente**).
- *Decoder* f tenta reconstruir a entrada original a partir do latente.

Autoencoders – Regularização

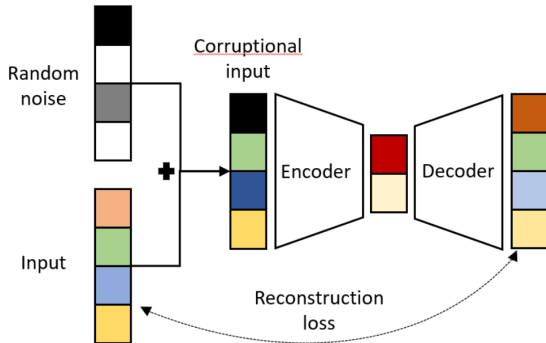
- É importante que o *autoencoder* tenha **regularização** para evitar soluções triviais.
- Se **não houver um gargalo**, o *autoencoder* pode simplesmente copiar a imagem.
- Se houver **unidades suficientes** no vetor latente, é possível indexar as imagens originais (memorização).

Formas comuns de regularização

- **Bottleneck**: dimensão latente $\ll$ dimensão da entrada.
- **Ruído** na entrada (denoising AE).
- **Esparsidade** na ativação do latente.

Denoising Autoencoders¹

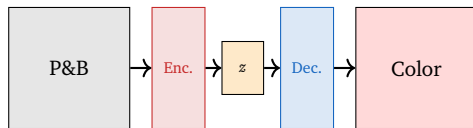
- Alguns *autoencoders* são treinados com o objetivo de **remover ruído** das instâncias de entrada.
- Tipicamente o *noise* segue uma:
 - distribuição **normal** (ruído aditivo);
 - distribuição de **Bernoulli** (ruído multiplicativo que invalida algumas *features*).
- A reconstrução é avaliada contra a entrada original (sem ruído).



¹Pascal Vincent et al. "Extracting and Composing Robust Features with Denoising Autoencoders". Em: [ICML](#). 2008, pp. 1096–1103.

Aplicação: Colorização

- **Entrada:** imagem P&B (1 canal). **Saída:** a mesma imagem em cores (3 canais RGB).
- *Encoder* comprime a cena em uma representação semântica; *decoder* gera as cores plausíveis.
- Perda: MSE no espaço de cor (RGB, ou *Lab* comparando apenas os canais *a*, *b*).
- Variação comum: *U-Net* convolucional com *skip connections* para preservar bordas.



Esquema: 1 canal $\rightarrow$ representação $\rightarrow$ 3 canais.

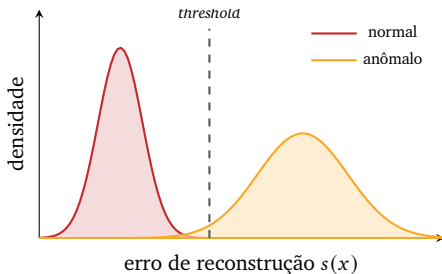
Problema **mal-posto**: várias colorações são plausíveis. MSE básico tende a cores desaturadas (média entre opções); variantes probabilísticas (VAE, GAN, difusão) geram cores mais vividas.

Aplicação: Detecção de Anomalias

- Instâncias **normais** reconstroem bem (erro baixo); **anomalias** reconstroem mal (erro alto).
- Score de anomalia:

$$s(x) = \|x - f(g(x, \phi), \theta)\|^2$$

- Um *threshold* sobre $s(x)$ separa as duas classes.
- Casos típicos: fraude em transações, defeitos industriais (*visual inspection*), intrusão em redes, monitoramento de equipamentos.



Por que funciona?

O AE só aprendeu a distribuição dos dados de treino, ele não tem como reconstruir o que nunca viu.

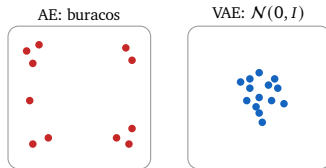
Visão Geral

1. Autoencoders
- 2. Variational Autoencoders**
3. Treinamento e ELBO
4. Reparametrization Trick
5. VAE em ação
6. Panorama

Por que o *autoencoder* comum não gera?

Para **gerar** bastaria amostrar um código z e decodificá-lo. No AE determinístico isso **falha**:

- O encoder envia cada x a um ponto z , mas **não controla a distribuição** dos códigos.
- O latente fica com **buracos** (regiões sem dado) e aglomerados irregulares.
- Sem uma $p(z)$ conhecida de onde amostrar, um z aleatório cai num buraco $\rightarrow$ decoder gera **lixo**.



A ideia do VAE

Forçar o encoder a mapear os dados a uma distribuição **contínua e conhecida** (gaussiana): então $z \sim \mathcal{N}(0, I)$ decodificado gera dados novos plausíveis.

Variational Autoencoders²

- São modelos geradores **probabilísticos** cujo objetivo é aprender uma distribuição $p(x)$ sobre os dados.
- É um **modelo não linear de variável latente**.
 - Depois do treinamento podemos usar apenas o decoder (amostragem) ou apenas o encoder (inferência).

Panorama

O VAE é a **fusão** de duas ideias:

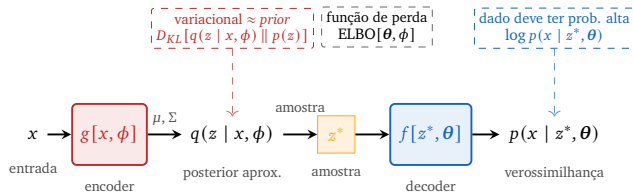
- *Autoencoder*: arquitetura encoder/decoder.
- **Inferência variacional**: aproximar a verossimilhança intratável via um limite inferior (ELBO).

Próximo slide: arquitetura completa, para servir de mapa ao longo da matemática.

²Diederik P. Kingma e Max Welling. “Auto-Encoding Variational Bayes”. Em: [ICLR](#). 2014.

Arquitetura VAE – Panorama

- **Encoder** $g[x, \phi]$ recebe x e produz $\mu(x), \Sigma(x)$.
- Definimos $q(z | x, \phi) = \mathcal{N}(\mu, \Sigma)$.
- Amostramos $z^* \sim q(z | x, \phi)$.
- **Decoder** $f[z^*, \theta]$ gera os parâmetros de $p(x | z^*, \theta)$.
- Perda: $-\text{ELBO} = \text{reconstrução} + \text{KL}$.



Adaptado de Prince (2023), Fig. 17.9.

Use como mapa

Toda a matemática a seguir refere-se a esses mesmos blocos.

Modelos de variáveis latentes

- Abordagem **indireta** para modelar uma distribuição $p(x)$ sobre uma variável multidimensional x .
- Marginalização de uma probabilidade conjunta $p(x, z)$:

$$p(x) = \int p(x, z) dz = \int p(x | z) p(z) dz$$

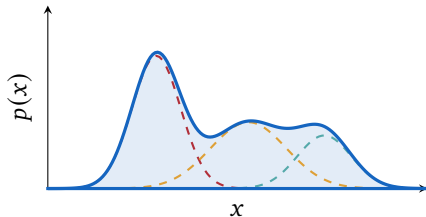
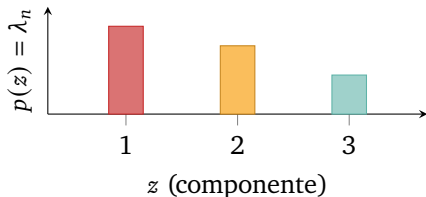
Motivação

Distribuições relativamente simples $p(x | z)$ e $p(z)$ podem, por marginalização, definir distribuições $p(x)$ **complexas**.

Exemplo: Mistura de Gaussianas 1D

- z é uma variável latente **discreta**, com distribuição a priori categórica:
 $p(z = n) = \lambda_n$.
- A verossimilhança segue uma normal: $p(x | z = n) = \mathcal{N}(\mu_n, \sigma_n^2)$.
- A distribuição marginal:

$$p(x) = \sum_n \lambda_n \mathcal{N}(\mu_n, \sigma_n^2).$$



O prior discreto $p(z)$ pesa cada componente, e a marginal $p(x)$ é a soma ponderada das gaussianas (tracejadas).

Modelos não lineares de variáveis latentes

Tanto z quanto x são **contínuos** e **multivariados**.

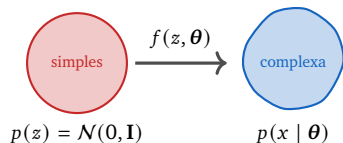
- Prior normal padrão: $p(z) = \mathcal{N}_z(0, \mathbf{I})$.
- Verossimilhança gaussiana com média não linear $f(z, \theta)$ e covariância esférica:

$$p(x | z, \theta) = \mathcal{N}_x\left(f(z, \theta), \sigma^2 \mathbf{I}\right).$$

- Marginal:

$$p(x | \theta) = \int p(x | z, \theta) p(z) dz.$$

- Integrando gaussiano, mas f não linear $\Rightarrow$ integral **sem forma fechada**.



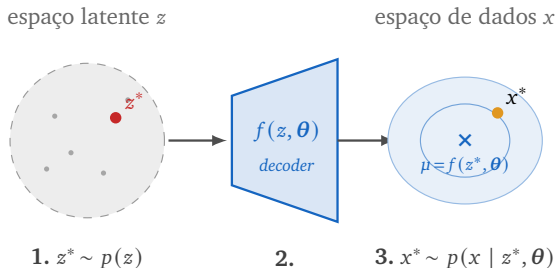
Função f não linear transforma uma gaussiana simples em uma distribuição arbitrariamente complexa.

Gerando novas instâncias x^*

Usamos **amostragem ancestral**:

1. Amostrar um latente z^* de $p(z)$ (*prior*).
2. Passar z^* pela rede $f(z^*, \theta)$ para obter a média da verossimilhança $p(x | z^*, \theta)$.
3. Amostrar x^* de $p(x | z^*, \theta)$.

A rede $f(z, \theta)$ é exatamente o *decoder* da arquitetura.



Amostramos z^* do *prior*, decodificamos com $f(z, \theta)$ e amostramos x^* da verossimilhança.

Na prática, costuma-se usar a média $f(z^*, \theta)$ diretamente, sem somar o ruído $\sigma^2 \mathbf{I}$.

Visão Geral

1. Autoencoders
2. Variational Autoencoders
- 3. Treinamento e ELBO**
4. Reparametrization Trick
5. VAE em ação
6. Panorama

Treinamento – O problema

Maximizar a log-verossimilhança direta é **inviável**:

$$\hat{\theta} = \operatorname{argmax}_{\theta} \left[\sum_{i=1}^N \int \mathcal{N}_{x^{(i)}} \left(f(z, \theta), \sigma^2 \mathbf{I} \right) \mathcal{N}_z(0, \mathbf{I}) dz \right]$$

- Integral **sem forma fechada**, pois f é não linear.
- Monte Carlo direto $p(x) \approx \frac{1}{K} \sum_k p(x | z_k)$ com $z_k \sim p(z)$ também **falha**: para um x específico, z amostrado do prior quase nunca explica x , então $p(x | z)$ é minúsculo na enorme maioria das amostras.

Estratégia

Maximizamos um **limite inferior** da log-verossimilhança, usando uma proposta inteligente para amostrar z .

Por que precisamos de um *encoder*?

Ideia: ao invés de amostrar z do prior, sortear z de uma distribuição que concentra massa onde z explica x .

- Essa distribuição ideal é a posterior $p(z | x)$, mas ela também é intratável.
- **Solução:** aproximamos com uma proposta parametrizada

$$q(z | x, \phi) \approx p(z | x),$$

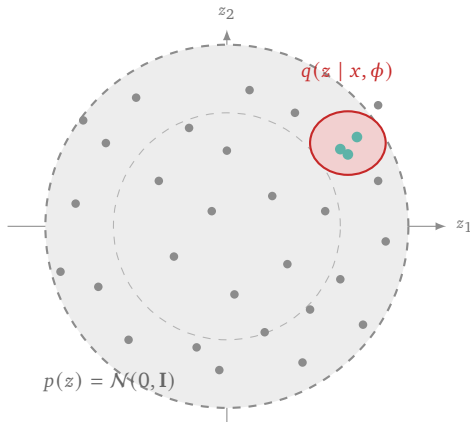
implementada por uma **segunda rede** com parâmetros ϕ .

- Essa rede é o *encoder* do VAE: para cada x , produz uma distribuição sobre z .

O que ganhamos

Com q no lugar de $p(z)$, podemos derivar um **limite inferior** (ELBO) que é apertado quando q aproxima bem a posterior, vamos ver isso a seguir.

Amostrar do *prior* é ineficiente



- Cinza: amostras do *prior* $p(z)$. A massa fica espalhada por todo o espaço latente.
- Verde: as raras amostras que caem onde z realmente explica x , isto é, onde $p(x | z)$ é alto.
- Acertar essa região por sorteio do *prior* exige um número **enorme** de amostras: $p(z)$ é pouco informativo sobre um x específico.

Solução

A proposta $q(z | x, \phi)$ já **concentra** a massa nessa região.

Desigualdade de Jensen

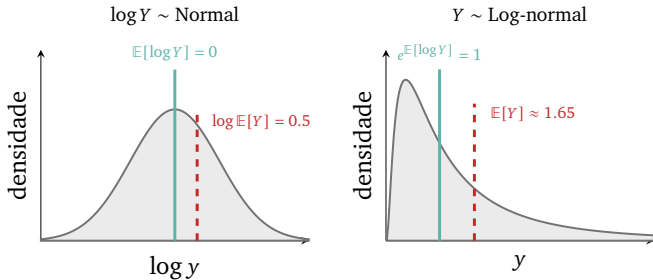
Para qualquer função **côncava** $g(\cdot)$:

$$g(\mathbb{E}[x]) \geq \mathbb{E}[g(x)]$$

Caso da função $\log$ (com $\log Y \sim \mathcal{N}(\mu, \sigma^2)$):

$$\log(\mathbb{E}[Y]) \geq \mathbb{E}[\log Y]$$

Aqui $\mu = 0$, $\sigma = 1$, então a folga $\sigma^2/2 = 0.5$.



A cauda da log-normal puxa a média aritmética $\mathbb{E}[Y]$ à direita da geométrica $e^{\mathbb{E}[\log Y]}$; logo $\log \mathbb{E}[Y] \geq \mathbb{E}[\log Y]$.

De onde vem o ELBO?

O nome **ELBO** (*Evidence Lower BOund*) é herdado de **inferência variacional**.

- **Evidência:** na linguagem bayesiana, $p(x | \theta) = \int p(x, z | \theta) dz$ é a verossimilhança marginal, a “evidência” que os dados dão sobre o modelo.
- **Limite inferior:** como a integral é intratável, construímos uma função $\mathcal{L}[\theta, \phi]$ tal que $\log p(x | \theta) \geq \mathcal{L}[\theta, \phi]$, e maximizamos $\mathcal{L}$ no lugar de $\log p(x)$.

Ingrediente que falta

Importance sampling: para qualquer q com suporte sobre z ,
 $p(x | \theta) = \mathbb{E}_{z \sim q}[p(x, z | \theta)/q(z)]$. O próximo slide combina esse passo com Jensen para obter o limite inferior.

Limite Inferior – Derivação

Reescrevemos a integral como uma esperança sob q e aplicamos Jensen na concavidade do log:

$$\begin{aligned}\log p(x \mid \theta) &= \log \left[\int p(x, z \mid \theta) dz \right] && \text{marginal} \\ &= \log \left[\int q(z) \frac{p(x, z \mid \theta)}{q(z)} dz \right] && \text{multiplica e divide por } q(z) \\ &= \log \left(\mathbb{E}_{z \sim q} \left[\frac{p(x, z \mid \theta)}{q(z)} \right] \right) && \text{integral ponderada por } q = \text{esperança} \\ &\geq \mathbb{E}_{z \sim q} \left[\log \frac{p(x, z \mid \theta)}{q(z)} \right] && \text{Jensen (log côncavo)} \\ &= \int q(z) \log \left(\frac{p(x, z \mid \theta)}{q(z)} \right) dz && \text{de volta à forma integral}\end{aligned}$$

ELBO

O $q(z)$ externo é a medida da esperança, permanece como peso; apenas o $q(z)$ do denominador entra no log. Esse limite é o *evidence lower bound*; na prática q é parametrizado: $q(z \mid x, \phi)$.

Limite Inferior – Conexão com a posterior

A partir daqui escrevemos q como $q(z | x, \phi)$ (o encoder condiciona em x). Aplicando a regra do produto $p(x, z | \theta) = p(z | x, \theta) p(x | \theta)$ no ELBO:

$\text{ELBO}[\theta, \phi] = \mathbb{E}_q \left[\log \frac{p(x, z \theta)}{q(z x, \phi)} \right]$	definição
$= \mathbb{E}_q \left[\log \frac{p(z x, \theta) p(x \theta)}{q(z x, \phi)} \right]$	regra do produto
$= \mathbb{E}_q [\log p(x \theta)] + \mathbb{E}_q \left[\log \frac{p(z x, \theta)}{q(z x, \phi)} \right]$	separa o log
$= \log p(x \theta) - \mathbb{E}_q \left[\log \frac{q(z x, \phi)}{p(z x, \theta)} \right]$	$\log p(x)$ sai da $\mathbb{E}$; troca o sinal
$= \log p(x \theta) - D_{KL}[q(z x, \phi) \ p(z x, \theta)]$	definição de KL

Rearranjando

$\log p(x | \theta) = \text{ELBO}[\theta, \phi] + D_{KL}[q(z | x, \phi) \| p(z | x, \theta)]$, a identidade exata do próximo slide.

ELBO – A identidade mestre

Da derivação anterior, chega-se à **identidade exata**:

$$\log p(x | \theta) = \underbrace{\int q(z | x, \phi) \log \frac{p(x, z | \theta)}{q(z | x, \phi)} dz}_{\text{ELBO}[\theta, \phi]} + \underbrace{D_{KL}[q(z | x, \phi) \| p(z | x, \theta)]}_{\geq 0}$$

- O gap entre $\log p(x)$ e o ELBO é **exatamente** a KL entre q e a posterior verdadeira.
- Bound **apertado** $\Leftrightarrow q(z | x, \phi) = p(z | x, \theta)$.
- Maximizar o ELBO faz **duas coisas ao mesmo tempo**: sobe $\log p(x)$ (aprende o modelo) e a aperta q contra a posterior (aprende o encoder).

ELBO – Reorganizando

$$\text{ELBO}[\boldsymbol{\theta}, \phi] = \int q(z | x, \phi) \log \left(\frac{p(x, z | \boldsymbol{\theta})}{q(z | x, \phi)} \right) dz$$

Aplicando $p(x, z) = p(x | z) p(z)$ e separando a razão:

$$\begin{aligned} &= \int q(z | x, \phi) \log \left(\frac{p(x | z, \boldsymbol{\theta}) p(z)}{q(z | x, \phi)} \right) dz \\ &= \int q(z | x, \phi) \log p(x | z, \boldsymbol{\theta}) dz - D_{KL}[q(z | x, \phi) \parallel p(z)] \end{aligned}$$

ELBO – Interpretação

$$\text{ELBO}[\theta, \phi] = \underbrace{\int q(z | x, \phi) \log p(x | z, \theta) dz}_{\text{Reconstruction Loss}} \underbrace{- D_{KL}[q(z | x, \phi) || p(z)]}_{\text{Regularização do latente}}$$

- **Primeiro termo:** log-verossimilhança esperada da reconstrução, sob z amostrados do encoder. Também chamado de *reconstruction loss*.
- **Segundo termo:** avalia quão perto a distribuição auxiliar $q(z | x, \phi)$ está do prior $p(z)$.

Sinal

O ELBO é **maximizado** (é um limite inferior da log-verossimilhança); a função de custo treinada é $-\text{ELBO}$, minimizada.

Termo de reconstrução = MSE (caso gaussiano)

Para $p(x | z, \theta) = \mathcal{N}(f(z, \theta), \sigma^2 \mathbf{I})$:

$$\begin{aligned}\log p(x | z, \theta) &= -\frac{D_x}{2} \log(2\pi\sigma^2) - \frac{1}{2\sigma^2} \|x - f(z, \theta)\|^2 \\ &= -\frac{1}{2\sigma^2} \|x - f(z, \theta)\|^2 + \text{const}\end{aligned}$$

Ponte com o AE clássico

Maximizar $\log p(x | z, \theta)$ é **exatamente** minimizar o erro quadrático $\|x - f(z, \theta)\|^2$. O ELBO se decompõe em

$$-\text{ELBO} = \underbrace{\frac{1}{2\sigma^2} \mathbb{E}_q[\|x - f(z, \theta)\|^2]}_{\text{reconstrução (MSE)}} + \underbrace{D_{KL}[q(z | x, \phi) \| p(z)]}_{\text{regularização}} + \text{const.}$$

KL entre gaussianas (1/2) – Partindo da definição

A KL entre duas distribuições quaisquer é:

$$D_{KL}[q \parallel p] = \int q(z) \log\left(\frac{q(z)}{p(z)}\right) dz = \mathbb{E}_q[\log q(z)] - \mathbb{E}_q[\log p(z)]$$

KL entre gaussianas (1/2) – Partindo da definição

A KL entre duas distribuições quaisquer é:

$$D_{KL}[q \parallel p] = \int q(z) \log\left(\frac{q(z)}{p(z)}\right) dz = \mathbb{E}_q[\log q(z)] - \mathbb{E}_q[\log p(z)]$$

Para $q(z|x, \phi) = \mathcal{N}(\mu, \Sigma)$ e $p(z) = \mathcal{N}(\mathbf{0}, \mathbf{I})$, o log de cada densidade gaussiana é uma forma quadrática:

$$\log p(z) = -\frac{D_z}{2} \log(2\pi) - \frac{1}{2} z^\top z$$

$$\log q(z) = -\frac{D_z}{2} \log(2\pi) - \frac{1}{2} \log |\Sigma| - \frac{1}{2} (z - \mu)^\top \Sigma^{-1} (z - \mu)$$

Os termos $-\frac{D_z}{2} \log(2\pi)$ cancelam. No próximo slide usamos a identidade $\mathbb{E}_q[zz^\top] = \Sigma + \mu\mu^\top$ (definição de covariância: $\Sigma = \mathbb{E}_q[zz^\top] - \mu\mu^\top$).

KL entre gaussianas (2/2) – Forma fechada

Tomando esperanças sob q (usando $\mathbb{E}_q[z] = \mu$ e $\mathbb{E}_q[zz^\top] = \Sigma + \mu\mu^\top$):

$$\mathbb{E}_q\left[-\frac{1}{2}z^\top z\right] = -\frac{1}{2}(\text{Tr}[\Sigma] + \mu^\top \mu)$$

$$\mathbb{E}_q\left[-\frac{1}{2}(z - \mu)^\top \Sigma^{-1}(z - \mu)\right] = -\frac{1}{2} D_z$$

Substituindo na definição da KL:

$$D_{KL}[\mathcal{N}(\mu, \Sigma) \parallel \mathcal{N}(\mathbf{0}, \mathbf{I})] = \frac{1}{2} \left(\underbrace{\text{Tr}[\Sigma]}_{\text{dispersão}} + \underbrace{\mu^\top \mu}_{\text{desvio da origem}} - \underbrace{D_z}_{\text{norm}} - \underbrace{\log \det \Sigma}_{\text{volume}} \right)$$

Exemplo Numérico: Calculando o Termo KL

Imagine um VAE treinado em MNIST com latente 2D. Para a imagem de um “7” específico, o encoder produz uma distribuição $q(z | x, \phi) = \mathcal{N}(\mu, \Sigma)$ sobre o latente, compacta o suficiente para “codificar 7”, mas com alguma incerteza para suavizar o espaço. O termo KL mede quão longe essa distribuição está do prior $\mathcal{N}(\mathbf{0}, \mathbf{I})$.

Encoder produz $\mu = (0.5, -0.8)$ e $\Sigma = \text{diag}(0.25, 1.5)$ (latente 2D, $D_z = 2$).

Calculando cada termo:

$$\text{Tr}[\Sigma] = 0.25 + 1.5 = 1.75$$

$$\mu^\top \mu = 0.25 + 0.64 = 0.89$$

$$\log \det \Sigma = \log(0.25) + \log(1.5)$$

$$\approx -1.386 + 0.405$$

$$= -0.981$$

$$D_z = 2$$

Substituindo

$$\begin{aligned} D_{KL} &= \frac{1}{2} (1.75 + 0.89 - 2 + 0.981) \\ &= \frac{1}{2} (1.621) \approx \mathbf{0.81} \end{aligned}$$

Os dois termos que dominam:

$\mu^\top \mu = 0.89$ (desvio da origem) e

$-\log \det \Sigma = +0.98$ (variância pequena

$\sigma^2 = 0.25$ é penalizada para evitar colapso).

VAE – Aproximação por *Monte Carlo*

O primeiro termo do ELBO ainda é uma integral intratável. Porém, como é uma **esperança**, podemos aproximá-lo por *Monte Carlo*:

$$\int q(z | x, \phi) \log p(x | z, \theta) dz \approx \log p(x | z^*, \theta), \quad z^* \sim q(z | x, \phi)$$

ELBO aproximado

$$\text{ELBO}[\theta, \phi] \approx \log p(x | z^*, \theta) - D_{KL}[q(z | x, \phi) \| p(z)]$$

Visão Geral

1. Autoencoders
2. Variational Autoencoders
3. Treinamento e ELBO
- 4. Reparametrization Trick**
5. VAE em ação
6. Panorama

Reparametrization Trick

- Como temos um **termo estocástico** no meio da rede neural (a amostragem $z^* \sim q(z \mid x, \phi)$), é inviável computar *backpropagation* diretamente.
- **Solução:** reparametrizar a saída do *encoder* de modo a tirar a amostragem do caminho do gradiente:

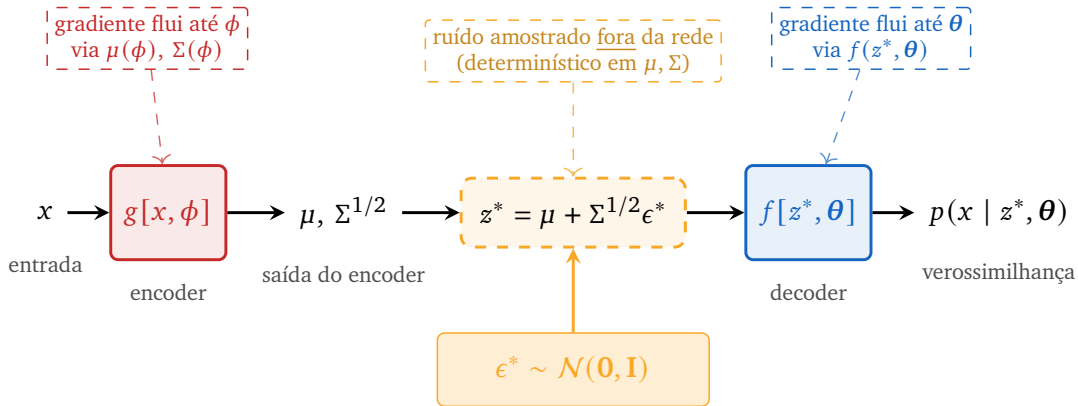
$$z^* = \mu + \Sigma^{1/2} \epsilon^*, \quad \epsilon^* \sim \mathcal{N}(0, \mathbf{I})$$

Para Σ diagonal (caso usual), $\Sigma^{1/2} \epsilon^* = \sigma \odot \epsilon^*$ (forma elemento a elemento).

Ideia chave

O ruído ϵ^* é amostrado fora da rede; μ e Σ permanecem no caminho diferenciável.

Reparametrization Trick – Diagrama



Adaptado de Prince (2023), Fig. 17.11.

Visão Geral

1. Autoencoders
2. Variational Autoencoders
3. Treinamento e ELBO
4. Reparametrization Trick
- 5. VAE em ação**
6. Panorama

VAE em ação – Amostras e interpolação



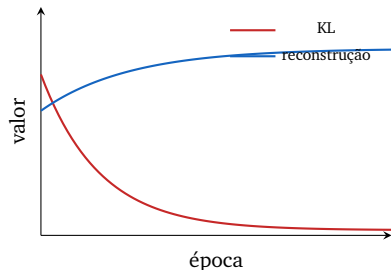
Amostragem: $z \sim p(z)$, depois $x = f(z, \theta)$. Resultado reconhecível mas borrado (limite inferior não é apertado).



Interpolação: transição suave entre dois dígitos confirma que o latente é contínuo e semanticamente organizado.

Patologia: *posterior collapse*

- Quando o decoder f é **poderoso demais**, ele consegue modelar x ignorando z .
- O termo KL pressiona $q(z | x, \phi) \rightarrow p(z)$: encoder colapsa para o prior. Latente vira inútil.
- Sintomas durante o treino:
 - $D_{KL} \rightarrow 0$;
 - $\mu(x) \rightarrow 0, \Sigma(x) \rightarrow \mathbf{I}$ para todo x .
- Mitigações:
 - **KL annealing**: rampa do peso do KL de 0 a 1;
 - **Free bits**: troca D_{KL} por $\sum_i \max(\lambda, D_{KL,i})$. Cada dimensão tem λ nats “grátis” antes do KL pressionar;
 - reduzir capacidade do decoder.



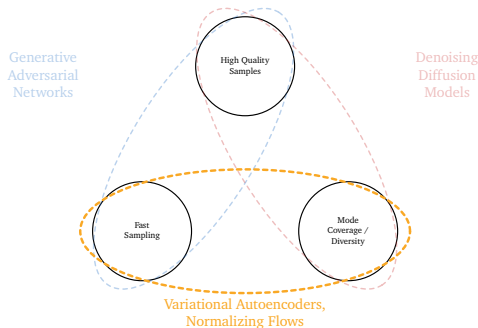
KL cai a zero enquanto a reconstrução estagna: o decoder ignora z .

Visão Geral

1. Autoencoders
2. Variational Autoencoders
3. Treinamento e ELBO
4. Reparametrization Trick
5. VAE em ação
- 6. Panorama**

Onde VAEs ficam no trilema?

- **Amostragem rápida:** ✓ (1 forward do decoder).
- **Cobertura / diversidade:** ✓ (maximização do ELBO cobre toda a distribuição).
- **Qualidade visual:** × (imagens tipicamente borradas; o limite inferior não é apertado).



Posicionamento

VAEs ocupam o canto velocidade + cobertura do trilema, ao lado dos *Normalizing Flows*. Para qualidade, combine VAEs com modelos difusores, base da arquitetura *Latent Diffusion*.

Resumo

- *Autoencoders*: encoder g e decoder f reconstruindo a entrada; latente útil para representação.
- *Denoising AE*: entrada corrompida; rede aprende a remover ruído.
- Aplicação de AE: detecção de anomalias via erro de reconstrução.
- *VAE*: *autoencoder* probabilístico, latente é uma distribuição $q(z | x, \phi)$.
- Identidade mestre: $\log p(x) = \text{ELBO} + D_{KL}[q||p(z | x)]$; bound apertado quando q aproxima a posterior.
- $\text{ELBO} = \text{reconstrução} - \text{KL}[q(z | x, \phi)||p(z)]$ (logo $-\text{ELBO} = \text{reconstrução} + \text{KL}$); para $p(x | z)$ gaussiana, reconstrução = MSE.
- *Reparametrization trick* permite backprop através da amostragem.
- *Posterior collapse*: $\text{KL} \rightarrow 0$, latente vira inútil; mitigado com KL annealing ou free bits.

Leituras Recomendadas

- Capítulo 17: Simon J. D. Prince. Understanding Deep Learning. MIT Press, 2023
- Diederik P. Kingma e Max Welling. “Auto-Encoding Variational Bayes”. Em: ICLR. 2014
- Pascal Vincent et al. “Extracting and Composing Robust Features with Denoising Autoencoders”. Em: ICML. 2008, pp. 1096–1103
- Irina Higgins et al. “ β -VAE: Learning Basic Visual Concepts with a Constrained Variational Framework”. Em: ICLR. 2017

Referências I

- [1] Irina Higgins et al. “ β -VAE: Learning Basic Visual Concepts with a Constrained Variational Framework”. Em: [ICLR](#). 2017.
- [2] Diederik P. Kingma e Max Welling. “Auto-Encoding Variational Bayes”. Em: [ICLR](#). 2014.
- [3] Simon J. D. Prince. [Understanding Deep Learning](#). MIT Press, 2023.
- [4] Pascal Vincent et al. “Extracting and Composing Robust Features with Denoising Autoencoders”. Em: [ICML](#). 2008, pp. 1096–1103.