

Este formulário reúne as principais expressões das unidades 4 e 5 e resultados matemáticos de apoio. Notação: $\mathcal{N}(\mu, \Sigma)$ denota a gaussiana de média μ e covariância Σ ; $\odot$ é o produto elemento a elemento; $\sigma(\cdot)$ é a função sigmoide.

Matemática Básica

Sigmoide e derivada:

$$\sigma(a) = \frac{1}{1 + e^{-a}}, \quad \sigma'(a) = \sigma(a)(1 - \sigma(a)).$$

Softmax (sobre logits $\mathbf{z} \in \mathbb{R}^K$, com temperatura τ ; $\tau = 1$ recupera a forma padrão):

$$\text{softmax}(\mathbf{z}/\tau)_i = \frac{e^{z_i/\tau}}{\sum_{j=1}^K e^{z_j/\tau}}.$$

Logaritmos:

$$\ln(ab) = \ln a + \ln b, \quad \ln \frac{a}{b} = \ln a - \ln b, \quad \ln a^b = b \ln a.$$

Divergência de Kullback-Leibler (caso contínuo):

$$D_{KL}[P \parallel Q] = \int P(x) \ln \frac{P(x)}{Q(x)} dx.$$

Traço (soma da diagonal principal) e determinante de matriz diagonal:

$$\text{Tr}[\text{diag}(d_1, \dots, d_n)] = \sum_{i=1}^n d_i, \quad \det[\text{diag}(d_1, \dots, d_n)] = \prod_{i=1}^n d_i, \quad \log \det = \sum_{i=1}^n \log d_i.$$

Determinante de matriz 2×2 geral:

$$\det \begin{bmatrix} a & b \\ c & d \end{bmatrix} = ad - bc.$$

Atenção e Transformers

Projeções (entrada X com um *token* por linha; d é a dimensão das projeções):

$$Q = XW_q, \quad K = XW_k, \quad V = XW_v.$$

Scaled dot-product attention:

$$\text{attention}(Q, K, V) = \text{softmax}\left(\frac{QK^\top}{\sqrt{d}}\right)V.$$

Máscara causal (aplicada aos *scores* antes do *softmax*; posições com $-\infty$ recebem peso 0):

$$\text{score}_{ij} \leftarrow -\infty \quad \text{se } j > i.$$

Multi-head attention (H cabeças em paralelo, cada uma com dimensão d_{model}/H ; concatenação seguida da projeção de saída W_O):

$$\text{MultiHead}(Q, K, V) = \text{Concat}(\text{head}_1, \dots, \text{head}_H) W_O, \quad \text{head}_i = \text{attention}(QW_Q^i, KW_K^i, VW_V^i).$$

Complexidade da *self-attention* no comprimento de sequência n : $O(n^2d)$.

Layer normalization e conexão residual (média e desvio padrão sobre a dimensão das *features*; γ e β treináveis):

$$\text{LayerNorm}(x) = \gamma \frac{x - \mu}{\sigma + \epsilon} + \beta, \quad \text{output} = \text{LayerNorm}(x + \text{Subblock}(x)).$$

Positional encoding sinusoidal (posição i , dimensões $2k$ e $2k+1$; somado ao *embedding*):

$$\text{PE}(i, 2k) = \sin\left(\frac{i}{10000^{2k/d}}\right), \quad \text{PE}(i, 2k+1) = \cos\left(\frac{i}{10000^{2k/d}}\right).$$

Tokenização

Score de merge do WordPiece (reconstrução popularizada pelo Hugging Face; o BPE, em contraste, mescla o par adjacente de maior $\text{freq}(xy)$):

$$\text{score}(x, y) = \frac{\text{freq}(xy)}{\text{freq}(x) \cdot \text{freq}(y)}.$$

Modelo de linguagem unigrama (segmentação $x = (x_1, \dots, x_n)$; probabilidades iniciais por contagem de substrings, com $N = \sum_{t'} \text{count}(t')$):

$$P(x) = \prod_{i=1}^n P(x_i), \quad P(t) = \frac{\text{count}(t)}{N}.$$

Viterbi (melhor segmentação por programação dinâmica; $\delta(j)$ é o log da melhor segmentação do prefixo com j caracteres, $s_{i:j}$ é o trecho entre as posições i e j , V é o vocabulário):

$$\delta(0) = 0, \quad \delta(j) = \max_{\substack{i < j \\ s_{i:j} \in V}} [\delta(i) + \ln P(s_{i:j})].$$

Log-likelihood do corpus e critério de poda do Unigram (a perda Δ_t mede quanto $\mathcal{L}$ cai ao remover o token t e re-segmentar; podam-se os tokens de menor Δ_t):

$$\mathcal{L} = \sum_{w \in \text{corpus}} \text{freq}(w) \ln P(\text{melhor seg. de } w), \quad \Delta_t = \mathcal{L} - \mathcal{L}_{\text{sem } t}.$$

Tabela de embedding (lookup da linha i da matriz E ; $|V|$ é o tamanho do vocabulário):

$$\mathbf{e}_i = E[i], \quad E \in \mathbb{R}^{|V| \times d}, \quad \#\text{params}_{\text{embedding}} = |V| \times d.$$

Treinamento e Inferência de LLMs

Loss de modelagem de linguagem (entropia cruzada autorregressiva sobre m *tokens*):

$$\mathcal{L} = - \sum_{t=1}^m \log P(y_t | y_{<t}).$$

Lei de escala de Chinchilla (volume ótimo de dados D^* , em *tokens*, para um modelo de N parâmetros):

$$D^* \approx 20 N.$$

LoRA (matriz congelada $W_0 \in \mathbb{R}^{d \times k}$; apenas $B \in \mathbb{R}^{d \times r}$ e $A \in \mathbb{R}^{r \times k}$ são treináveis):

$$h = W_0 x + \frac{\alpha}{r} B A x, \quad \#\text{params treináveis} = r(d + k).$$

Decodificação especulativa (rascunho q , alvo p ; *token* proposto $\tilde{x}$):

$$P(\text{aceitar } \tilde{x}) = \min\left(1, \frac{p(\tilde{x})}{q(\tilde{x})}\right); \quad \text{se rejeitado, reamostrar de } (p - q)_+ = \max(0, p - q) \text{ normalizada.}$$

Alinhamento por Preferências

Reward model (perda de Bradley-Terry sobre o par preferido y_w e preterido y_l):

$$\mathcal{L}_{RM} = -\log \sigma(r_\phi(x, y_w) - r_\phi(x, y_l)).$$

RLHF com PPO (objetivo da política com penalização KL em relação ao modelo de referência):

$$J(\theta) = \mathbb{E}_{x, y \sim \pi_\theta} [r_\phi(x, y) - \beta \text{KL}(\pi_\theta \parallel \pi_{\text{ref}})].$$

DPO (otimiza diretamente sobre os pares de preferência, sem *reward model* explícito):

$$\mathcal{L}_{DPO} = -\log \sigma\left(\beta \left[\log \frac{\pi_\theta(y_w | x)}{\pi_{\text{ref}}(y_w | x)} - \log \frac{\pi_\theta(y_l | x)}{\pi_{\text{ref}}(y_l | x)} \right]\right).$$

GANs

Loss original do gerador (saturante), com discriminador $f[\cdot, \phi]$ e gerador $g[\cdot, \theta]$:

$$L[\theta] = \sum_j \ln \left[1 - \sigma\left(f[g[z^{(j)}, \theta], \phi]\right) \right].$$

Loss não-saturante do gerador:

$$L_{\text{ns}}[\theta] = -\sum_j \ln \sigma\left(f[g[z^{(j)}, \theta], \phi]\right).$$

Discriminador ótimo (P é a densidade dos dados reais e P^* a dos gerados):

$$D^*(x) = \frac{P(x)}{P(x) + P^*(x)}.$$

Divergência de Jensen-Shannon e sua decomposição (com $M = \frac{1}{2}(P + P^*)$):

$$D_{JS}[P^* \parallel P] = \underbrace{\frac{1}{2} D_{KL}[P^* \parallel M]}_{\text{qualidade}} + \underbrace{\frac{1}{2} D_{KL}[P \parallel M]}_{\text{cobertura}}, \quad 0 \leq D_{JS} \leq \ln 2.$$

WGAN-GP (*gradient penalty* adicionada à *loss* do crítico, em pontos interpolados $\hat{x}$):

$$\lambda \left(\|\nabla_{\hat{x}} f[\hat{x}, \phi]\|_2 - 1 \right)^2.$$

Normalizing Flows

Mudança de variáveis (bijeção $z = f(x)$, $\dim \mathbf{z} = \dim \mathbf{x}$):

$$p_x(\mathbf{x}) = p_z(f(\mathbf{x})) |\det J_f(\mathbf{x})|, \quad \log p_x(\mathbf{x}) = \log p_z(f(\mathbf{x})) + \log |\det J_f(\mathbf{x})|.$$

Linear flow: $f(\mathbf{x}) = \theta \mathbf{x} + \mathbf{b}$.

Affine coupling (particiona a entrada em $\mathbf{x}^A, \mathbf{x}^B$; a *coupling network* produz $(\mathbf{s}, \mathbf{t}) = \text{NN}(\mathbf{x}^A)$):

$$\begin{aligned} \text{forward: } \mathbf{y}^A &= \mathbf{x}^A, & \mathbf{y}^B &= \mathbf{x}^B \odot \exp(\mathbf{s}) + \mathbf{t}, \\ \text{inversa: } \mathbf{x}^A &= \mathbf{y}^A, & \mathbf{x}^B &= (\mathbf{y}^B - \mathbf{t}) \odot \exp(-\mathbf{s}), & (\mathbf{s}, \mathbf{t}) &= \text{NN}(\mathbf{y}^A). \end{aligned}$$

Log-determinante da camada *affine coupling*:

$$\log|\det J| = \sum_i s_i.$$

Autoencoders Variacionais (VAE)

KL em forma fechada entre o posterior $\mathcal{N}(\mu, \Sigma)$ e o *prior* $\mathcal{N}(\mathbf{0}, \mathbf{I})$, com latente de dimensão D_z :

$$D_{KL}[\mathcal{N}(\mu, \Sigma) \parallel \mathcal{N}(\mathbf{0}, \mathbf{I})] = \frac{1}{2} (\text{Tr}[\Sigma] + \mu^\top \mu - D_z - \log \det \Sigma).$$

Reparametrization trick (com $\epsilon^* \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$; para Σ diagonal, $\Sigma^{1/2}\epsilon^* = \sigma \odot \epsilon^*$):

$$\mathbf{z}^* = \mu + \Sigma^{1/2} \epsilon^*.$$

Termo de reconstrução (*decoder* gaussiano $p(x | z, \theta) = \mathcal{N}_x(f[z, \theta], \mathbf{I})$):

$$\frac{1}{2} \|x - f[z^*, \theta]\|^2.$$

Objetivo (ELBO; a otimização minimiza $-\text{ELBO}$):

$$-\text{ELBO} \approx \underbrace{\frac{1}{2} \|x - f[z^*, \theta]\|^2}_{\text{reconstrução}} + \underbrace{D_{KL}[q(z | x, \phi) \parallel p(z)]}_{\text{regularização}}.$$

Modelos Difusores

Noise schedule (produto acumulado a partir de $\beta = (\beta_1, \dots, \beta_T)$):

$$\alpha_t = \prod_{s=1}^t (1 - \beta_s).$$

Processo forward (um passo, cadeia de Markov gaussiana, $\epsilon_{t-1} \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$):

$$\mathbf{z}_t = \sqrt{1 - \beta_t} \mathbf{z}_{t-1} + \sqrt{\beta_t} \epsilon_{t-1}.$$

Kernel de difusão (amostra direta do passo t a partir de x e $\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$):

$$\mathbf{z}_t = \sqrt{\alpha_t} x + \sqrt{1 - \alpha_t} \epsilon.$$

Objetivo de treino (a U-Net $g_t[\cdot, \theta_t]$ estima o ruído):

$$L_t = \|g_t[\mathbf{z}_t, \theta_t] - \epsilon\|^2.$$

Amostragem reversa (passo de *denoising*, com $\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ e σ_t do *schedule*):

$$\mathbf{z}_{t-1} = \frac{1}{\sqrt{1 - \beta_t}} \left(\mathbf{z}_t - \frac{\beta_t}{\sqrt{1 - \alpha_t}} g_t[\mathbf{z}_t, \theta_t] \right) + \sigma_t \epsilon.$$

Classifier-free guidance (predição condicionada $g_t[\mathbf{z}_t, c]$ e sem condição $g_t[\mathbf{z}_t, \emptyset]$; peso $w \geq 0$):

$$\tilde{g}_t = (1 + w) g_t[\mathbf{z}_t, c] - w g_t[\mathbf{z}_t, \emptyset].$$