

Este formulário reúne as principais expressões das unidades 1 a 3 e resultados matemáticos de apoio. Notação: $\odot$ é o produto elemento a elemento; $\sigma(\cdot)$ é a função sigmoide; $\varphi(\cdot)$ é uma função de ativação genérica; η é a taxa de aprendizado; $\nabla_{\theta} \mathcal{L}$ é o gradiente da função de custo em relação aos parâmetros.

Matemática Básica

Logaritmos e exponenciais:

$$\ln(ab) = \ln a + \ln b, \quad \ln \frac{a}{b} = \ln a - \ln b, \quad \ln a^b = b \ln a, \quad \log_2 x = \frac{\ln x}{\ln 2}, \quad e^{a+b} = e^a e^b.$$

Funções de Ativação e Derivadas

$$\begin{aligned} \sigma(a) &= \frac{1}{1 + e^{-a}}, & \sigma'(a) &= \sigma(a)(1 - \sigma(a)), & \sigma'(0) &= \frac{1}{4}. \\ \tanh(a) &= \frac{e^a - e^{-a}}{e^a + e^{-a}}, & \tanh'(a) &= 1 - \tanh^2(a), & \tanh'(0) &= 1. \\ \text{ReLU}(a) &= \max(0, a), & \text{ReLU}'(a) &= \begin{cases} 1, & a > 0, \\ 0, & a < 0. \end{cases} \end{aligned}$$

Softmax (converte *logits* $\mathbf{z} \in \mathbb{R}^K$ em probabilidades; com temperatura τ):

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

Perceptron e MLP

Perceptron de Rosenblatt (saída binária em $\{-1, +1\}$; primeira componente de X igual a 1, truque do bias):

$$\hat{y} = \varphi(X \theta^\top), \quad \varphi(z) = \text{sinl}(z) = \begin{cases} +1, & z \geq 0, \\ -1, & z < 0. \end{cases}$$

Algoritmo de aprendizado do Perceptron (atualização para uma instância mal classificada $X^{(i)}$):

$$\theta \leftarrow \theta + y^{(i)} X^{(i)}.$$

Camada densa (pré-ativação e ativação, forma vetorizada; $A^{(0)} = X$):

$$Z^{(l)} = A^{(l-1)} \theta^{(l)\top} + b^{(l)}, \quad A^{(l)} = \varphi(Z^{(l)}).$$

Funções de Custo

Erro quadrático (L_2 loss, MSE quando dividida por N ; por instância, $J = \frac{1}{2}(y - \hat{y})^2$):

$$\mathcal{L}(\theta) = \sum_{i=1}^N (y^{(i)} - \hat{y}^{(i)})^2.$$

Entropia cruzada binária (classificação binária):

$$\mathcal{L}_{\text{BCE}}(\boldsymbol{\theta}) = - \sum_{i=1}^N \left[y^{(i)} \log \hat{y}^{(i)} + (1 - y^{(i)}) \log(1 - \hat{y}^{(i)}) \right].$$

Entropia cruzada categórica (K classes, saída *softmax*; $y^{(i)}$ é o índice da classe correta):

$$\mathcal{L}(\boldsymbol{\theta}) = - \sum_{i=1}^N \log \text{softmax}(f(\mathbf{x}^{(i)}; \boldsymbol{\theta}))_{y^{(i)}}.$$

Backpropagation

Regra da cadeia por camada (rede de duas camadas com custo J):

$$\frac{\partial J}{\partial \theta^{(2)}} = \frac{\partial J}{\partial a^{(2)}} \frac{\partial a^{(2)}}{\partial z^{(2)}} \frac{\partial z^{(2)}}{\partial \theta^{(2)}}, \quad \frac{\partial J}{\partial \theta^{(1)}} = \frac{\partial J}{\partial a^{(2)}} \frac{\partial a^{(2)}}{\partial z^{(2)}} \frac{\partial z^{(2)}}{\partial a^{(1)}} \frac{\partial a^{(1)}}{\partial z^{(1)}} \frac{\partial z^{(1)}}{\partial \theta^{(1)}}.$$

Delta da camada (definição, erro na camada de saída L e propagação do erro; $a^{(l)} = \varphi'(z^{(l)})$):

$$\delta^{(l)} = \frac{\partial J}{\partial z^{(l)}}, \quad \delta^{(L)} = \frac{\partial J}{\partial a^{(L)}} \odot \frac{\partial a^{(L)}}{\partial z^{(L)}}, \quad \delta^{(l)} = (\delta^{(l+1)} \theta^{(l+1)}) \odot a'^{(l)}.$$

Gradiente dos parâmetros (delta da camada vezes a ativação da camada anterior; $a^{(0)} = X$):

$$\frac{\partial J}{\partial \theta^{(l)}} = \delta^{(l)\top} a^{(l-1)}.$$

Descida de Gradiente e Otimizadores

Descida de gradiente:

$$\boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \eta \nabla_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}_t).$$

Momentum (média móvel exponencial dos gradientes, $v_0 = 0$):

$$v_{t+1} = \beta v_t + (1 - \beta) \nabla_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}_t), \quad \boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \eta v_{t+1}.$$

Nesterov (gradiente avaliado no ponto deslocado pelo momentum):

$$v_{t+1} = \beta v_t + (1 - \beta) \nabla_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}_t - \eta \beta v_t), \quad \boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \eta v_{t+1}.$$

AdaGrad (acumulador G cresce monotonicamente; operações elemento a elemento):

$$G_{t+1} = G_t + (\nabla_{\boldsymbol{\theta}} \mathcal{L})^2, \quad \boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \frac{\eta}{\sqrt{G_{t+1}} + \epsilon} \nabla_{\boldsymbol{\theta}} \mathcal{L}.$$

RMSProp (média móvel exponencial no lugar da soma):

$$\mathbb{E}[g^2]_{t+1} = \rho \mathbb{E}[g^2]_t + (1 - \rho) (\nabla_{\boldsymbol{\theta}} \mathcal{L})^2, \quad \boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \frac{\eta}{\sqrt{\mathbb{E}[g^2]_{t+1}} + \epsilon} \nabla_{\boldsymbol{\theta}} \mathcal{L}.$$

Adam (momentos m e v com correção de viés; $m_0 = v_0 = 0$):

$$m_{t+1} = \beta_1 m_t + (1 - \beta_1) \nabla_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}_t), \quad v_{t+1} = \beta_2 v_t + (1 - \beta_2) (\nabla_{\boldsymbol{\theta}} \mathcal{L}(\boldsymbol{\theta}_t))^2, \\ \tilde{m}_{t+1} = \frac{m_{t+1}}{1 - \beta_1^{t+1}}, \quad \tilde{v}_{t+1} = \frac{v_{t+1}}{1 - \beta_2^{t+1}}, \quad \boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \eta \frac{\tilde{m}_{t+1}}{\sqrt{\tilde{v}_{t+1}} + \epsilon}.$$

Regularização, Normalização e Inicialização

Regularização L1 e L2 (penalidade Ω somada à função de custo; L2 corresponde ao *weight decay*):

$$\tilde{\mathcal{L}}(\theta) = \mathcal{L}(\theta) + \lambda \Omega(\theta), \quad \Omega_{L2} = \|\theta\|_2^2, \quad \Omega_{L1} = \|\theta\|_1.$$

Batch normalization (estatísticas do *mini-batch* B ; γ e β treináveis; em inferência usam-se as médias móveis μ_{mov} e σ_{mov}):

$$\mu_h = \frac{1}{|B|} \sum_{i \in B} h_i, \quad \sigma_h = \sqrt{\frac{1}{|B|} \sum_{i \in B} (h_i - \mu_h)^2}, \quad h_i \leftarrow \frac{h_i - \mu_h}{\sigma_h + \epsilon}, \quad h_i \leftarrow \gamma h_i + \beta.$$

Inicialização de pesos (gaussianas de média zero com variâncias distintas; $n^{(l-1)}$ é o *fan-in* da camada; Xavier para tanh e sigmoide, He para ReLU):

$$\text{Xavier: } \theta_{ij}^{(l)} \sim \mathcal{N}\left(0, \frac{1}{n^{(l-1)}}\right), \quad \text{He: } \theta_{ij}^{(l)} \sim \mathcal{N}\left(0, \frac{2}{n^{(l-1)}}\right).$$

Variante fan_avg (compromisso entre *forward* e *backward*, com $n^{(l)}$ o *fan-out*):

$$\text{Xavier: } \sigma_\theta^2 = \frac{2}{n^{(l-1)} + n^{(l)}}, \quad \text{He: } \sigma_\theta^2 = \frac{4}{n^{(l-1)} + n^{(l)}}.$$

Camadas Convolucionais

Dimensão espacial de saída (entrada $W \times W$, *kernel* K , *stride* S , *padding* P ; vale também para *pooling*):

$$O = \left\lfloor \frac{W - K + 2P}{S} \right\rfloor + 1.$$

Kernel efetivo com dilatação D (substitui K na fórmula acima):

$$K_{\text{ef}} = D(K - 1) + 1.$$

Conexão residual (bloco com atalho de identidade):

$$y = F(x) + x.$$

IoU (*intersection over union* entre duas *bounding boxes* A e B):

$$\text{IoU}(A, B) = \frac{|A \cap B|}{|A \cup B|}.$$

Redes Recorrentes

RNN vanilla (um passo de tempo):

$$z_t = x_t \theta_{ih}^T + h_{t-1} \theta_{hh}^T + b_h, \quad h_t = \tanh(z_t).$$

Gradiente retropropagado no tempo (fator multiplicado a cada passo; origem do desaparecimento e da explosão do gradiente):

$$\frac{\partial h_t}{\partial h_{t-1}} \propto \theta_{hh} (1 - \tanh^2(z_t)).$$

Gradient clipping por norma (limita o passo preservando a direção):

$$g \leftarrow c \frac{g}{\|g\|} \quad \text{se } \|g\| > c.$$

LSTM (portões f_t , i_t , o_t e candidato $\tilde{c}_t$; $[h_{t-1}, x_t]$ é a concatenação):

$$\begin{aligned} f_t &= \sigma(\theta_f [h_{t-1}, x_t] + b_f), & i_t &= \sigma(\theta_i [h_{t-1}, x_t] + b_i), & \tilde{c}_t &= \tanh(\theta_c [h_{t-1}, x_t] + b_c), \\ c_t &= f_t \odot c_{t-1} + i_t \odot \tilde{c}_t, & o_t &= \sigma(\theta_o [h_{t-1}, x_t] + b_o), & h_t &= o_t \odot \tanh(c_t). \end{aligned}$$

Caminho da memória (derivada que mitiga a dissipação do gradiente):

$$\frac{\partial c_t}{\partial c_{t-1}} = f_t.$$

GRU (portões de *reset* r_t e de *update* z_t):

$$\begin{aligned} r_t &= \sigma(\theta_r [h_{t-1}, x_t] + b_r), & z_t &= \sigma(\theta_z [h_{t-1}, x_t] + b_z), \\ \tilde{h}_t &= \tanh(\theta_h [r_t \odot h_{t-1}, x_t] + b_h), & h_t &= (1 - z_t) \odot h_{t-1} + z_t \odot \tilde{h}_t. \end{aligned}$$

Modelagem de Linguagem e Atenção

Probabilidade de uma sequência (fatoração autorregressiva de m *tokens*):

$$p(y_1, \dots, y_m) = \prod_{t=1}^m p(y_t \mid y_{<t}).$$

Perplexidade (a partir da *log-likelihood* em bits, com $\log_2$):

$$L(y_{1:m}) = \sum_{t=1}^m \log_2(p(y_t \mid y_{<t})), \quad \text{Perplexidade}(y_{1:m}) = 2^{-\frac{1}{m}L(y_{1:m})}, \quad \text{Perplexidade} \in [1, |V|].$$

Atenção em seq2seq (estado do decodificador h_t , estados do codificador $s_1, \dots, s_m$; scores de produto escalar, pesos via *softmax* e vetor de contexto):

$$\text{score}(h_t, s_k) = h_t^\top s_k, \quad a_k^{(t)} = \frac{\exp(\text{score}(h_t, s_k))}{\sum_{i=1}^m \exp(\text{score}(h_t, s_i))}, \quad c^{(t)} = \sum_{k=1}^m a_k^{(t)} s_k.$$