

Lista de Exercícios 2

Fundamentos da Inteligência Artificial Generativa

Conteúdos: Arquiteturas generativas, redes neurais e mecanismos de atenção — aprofundamento técnico e cálculos

Nome: _____ Data: ____/____/____

1. Explique a diferença entre um modelo de linguagem autoregressivo (como o usado no GPT), que gera texto palavra por palavra prevendo a próxima com base nas anteriores, e um modelo de difusão (usado em geradores de imagem como o Stable Diffusion), que parte de ruído aleatório e o refina progressivamente.

Fórmula(s): Questão teórica — modelo autoregressivo × modelo de difusão

2. Um neurônio de rede neural recebe três entradas: $x_1 = 0,3$, $x_2 = 0,7$ e $x_3 = 0,2$, com pesos $w_1 = 0,5$, $w_2 = 0,2$ e $w_3 = 0,9$, e bias $b = -0,1$. Calcule a soma ponderada z recebida por esse neurônio.

Fórmula(s): $z = x_1w_1 + x_2w_2 + x_3w_3 + b$

3. Utilizando o resultado da questão anterior ($z = 0,37$), calcule a saída do neurônio após passar pela função de ativação sigmoide $f(z) = 1/(1+e^{-z})$, sabendo que $e^{-0,37} \approx 0,691$.

Fórmula(s): $f(z) = 1/(1 + e^{-z})$

4. Uma GAN foi treinada por 1000 épocas (iterações completas de treinamento). Ao final, o discriminador classificou corretamente 520 de 1000 imagens (metade reais, metade geradas) apresentadas em um teste. Calcule a acurácia do discriminador e explique o que significaria uma acurácia próxima de 50% nesse contexto.

Fórmula(s): $\text{Acurácia} = \text{acertos}/\text{total}$

5. Explique o conceito de 'embedding' (ou vetor de representação) no contexto de modelos de linguagem, e por que palavras com significados semelhantes tendem a ter embeddings numericamente próximos.

Fórmula(s): Questão teórica — embeddings e proximidade semântica

6. Explique a diferença entre a arquitetura Transformer (usada em modelos como GPT e BERT) e as redes neurais recorrentes (RNNs), quanto à forma de processar sequências de texto e à capacidade de paralelização do treinamento.

Fórmula(s): Questão teórica — Transformer × RNN (paralelização e dependências de longo alcance)

7. Um modelo de IA Generativa de texto tem uma 'janela de contexto' de 8.000 tokens. Sabendo que, em média, 1 token corresponde a cerca de 0,75 palavras em português, calcule aproximadamente quantas palavras cabem nessa janela de contexto.

Fórmula(s): $\text{Palavras} \approx \text{tokens} \times 0,75$

8. Explique o que é 'temperatura' (temperature) em um modelo de geração de texto, e como ajustar esse parâmetro para valores mais altos ou mais baixos afeta a criatividade/aleatoriedade das respostas geradas.

Fórmula(s): Questão teórica — temperatura em geração de texto

9. Explique o processo de 'fine-tuning' (ajuste fino) de um modelo de IA Generativa pré-treinado, e cite uma vantagem desse processo em relação a treinar um modelo do zero para uma tarefa específica.

Fórmula(s): Questão teórica — fine-tuning de modelo pré-treinado

10. Explique o que são 'alucinações' (hallucinations) em modelos de IA Generativa de texto, e cite uma estratégia que pode ajudar a reduzir esse problema ao utilizar esses modelos.

Fórmula(s): Questão teórica — alucinações em modelos generativos
