

Muitas vezes projetos de LLM com RAG ignoram o aspecto crucial da fragmentação de texto dentro de seu pipeline, o que impacta a qualidade das tarefas intensivas em conhecimento (Zhao, J., Meta-chunking: Learning efficient text segmentation via logical perception)

Técnicas de chunking inteligentes combinadas com o uso de RAG podem gerar embeddings mais ricos, potencialmente melhorando o desempenho de LLMs, mesmo aqueles com menor capacidade, ao focar o contexto nas informações mais críticas para a tarefa de codificação?

Faz sentido fragmentarmos o texto através de sumarização semântica em vez de divisão sequencial, reduzindo o número de tokens e aumentando a relevância dos chunks, e assim obter melhores resultados na geração de códigos mesmo com LLMs mais simples?

Nosso trabalho explora o pré-processamento focado em agrupar e resumir informações-chave de documentos-fonte criando tokens por contexto que são utilizados no RAG e no seu impacto na geração de códigos-fonte.

Como exemplo, carregamos a documentação da API dos Correios e, no contexto de engenharia de software, a IA pôde responder perguntas sobre a API, gerando código Python com base na documentação oficial.

Para comparar, criamos um pré-processamento alternativo com divisão sequencial e analisamos a diferença dos códigos-fonte e a quantidade de tokens gerados/ utilizados. Além disso, executamos cenários de testes com Llama e o Deep Seek, para verificar a diferença de código entre os dois LLMs.

