

Large Language Models Can Write Good Snippets Fast

- LLMs, especially frontier reasoning models since~ GPT-5 and Claude Sonnet 4.5, are now strong programmers
 - **A narrow, well-posed request** reliably generates a few hundred lines of self-consistent code in ~any language
- However, AI performance **degrades with problem size and complexity**
 - **Try prompting an LLM to read a 5000-line file and produce a correct GPU ported code in one completion. It will fail.**
 - Coding **agents** (Cursor, Claude Code, Codex) that interleave LLM inference with tool calls are **rapidly improving**.
 - But even for the best agents, simple prompting (“Port this code to GPU. No mistakes.”) doesn’t let the user walk away until an entire code is ported correctly, optimally, and suitably for production.

