

6 REASONS LLMS AMPLIFY SOFTWARE RISK

Large Language Models (LLMs) are increasingly used to generate code. While they offer speed and efficiency, it's crucial to understand the inherent risks associated with LLM-generated code. This infographic highlights six key reasons why LLMs can amplify software risk, and why you can't simply trust the code they produce.

1 TRAINED ON INSECURE CODE

LLMs learn from public codebases that include flawed security patterns.

2 PRIORITIZES SYNTAX, NOT SECURITY

LLMs aim for code that runs, not code that's safe.

3 SECURITY DEPENDS ON THE PROMPT

Without explicit security instructions, LLMs won't enforce best practices.

4 FAILS TO ADAPT BY LANGUAGE

Secure coding varies by language, but LLMs apply patterns inconsistently.

5 NO REAL-TIME FEEDBACK LOOP

LLMs don't test or improve based on security outcomes.

6 FALSE SENSE OF SAFETY

Clean-looking code masks hidden vulnerabilities.