



LLM-based Side-Channel Hardening

Advisor: **Hannes Weissteiner**

Motivation

Side channel attacks leak information about secrets using unintended channels, such as timing, power consumption, or memory access patterns. Writing code that is secure against side-channel attacks (constant-time code) is difficult and error-prone. We want to explore whether large language models (LLMs) can be useful to help developers eliminate side channels in their code.

The goal of the project is to explore if and how well LLMs can help to automatically harden code against side-channel attacks, while keeping the functionality of the original code.

Goals and Tasks

-  Review prior work on side-channel attacks and side-channel resistant coding techniques.
-  Explore how LLMs can be prompted to refactor vulnerable code into side-channel resistant code.
-  Implement a prototype that evaluates the effectiveness of different LLMs and prompts against small code snippets with known side-channel vulnerabilities.



Literature

- > [P. Kocher](#)
Timing Attacks on Implementations of Diffie-Hellman, RSA, DSS, and Other Systems
[CRYPTO](#)
- > [Intel](#)
Guidelines for Mitigating Timing Side Channels Against Cryptographic Implementations
[2019](#)

Courses & Deliverables

- ☒ **Introduction to Scientific Working**
Short report on background
Short presentation
- ☒ **Bachelor Project**
Project code and documentation
- ☒ **Bachelor's Thesis**
Project code
Thesis
Final presentation

Recommended if you're studying

- ☒ CS
- ☒ ICE
- ☒ SEM

Prerequisites

- > Programming experience (Python, C/C++)
- > Basic knowledge and interest in LLMs as coding assistants
- > Interest in system security and side-channel attacks

Advisor Contact