



ISCAS

中国科学院软件研究所学术年会
暨重点实验室科技活动周

2025 第十届

学术论文

The Seeds of the FUTURE Sprout from History: Fuzzing for Unveiling Vulnerabilities in Prospective Deep Learning Libraries

李志远, 吴敬征, 凌祥, 罗天悦, 芮志清, 武延军
lizhiyuan2021, jingzheng08, lingxiang@iscas.ac.cn

ACM International Conference on Software Engineering (ICSE 2025)

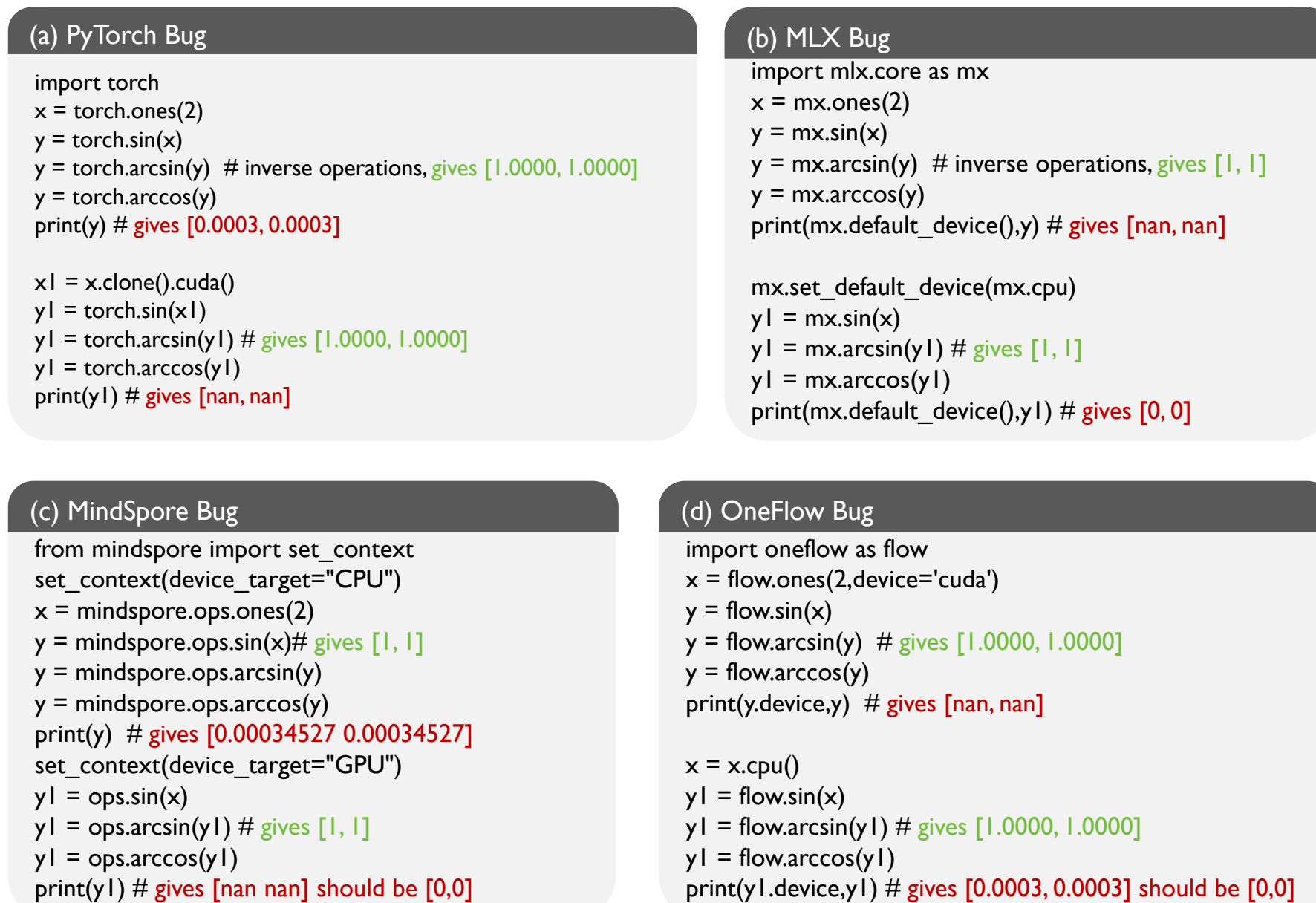
ACM SIGSOFT Distinguished Paper Award

背景 Background

近年来，深度学习技术依赖的人工智能框架安全性日益受到关注。作为检测人工智能框架安全性的关键技术，当前的模糊测试方法存在适用对象局限性、缺乏针对复杂API序列的测试以及大模型训练知识滞后等问题，难以满足开源人工智能框架的安全性检测需求。

动机 Motivation

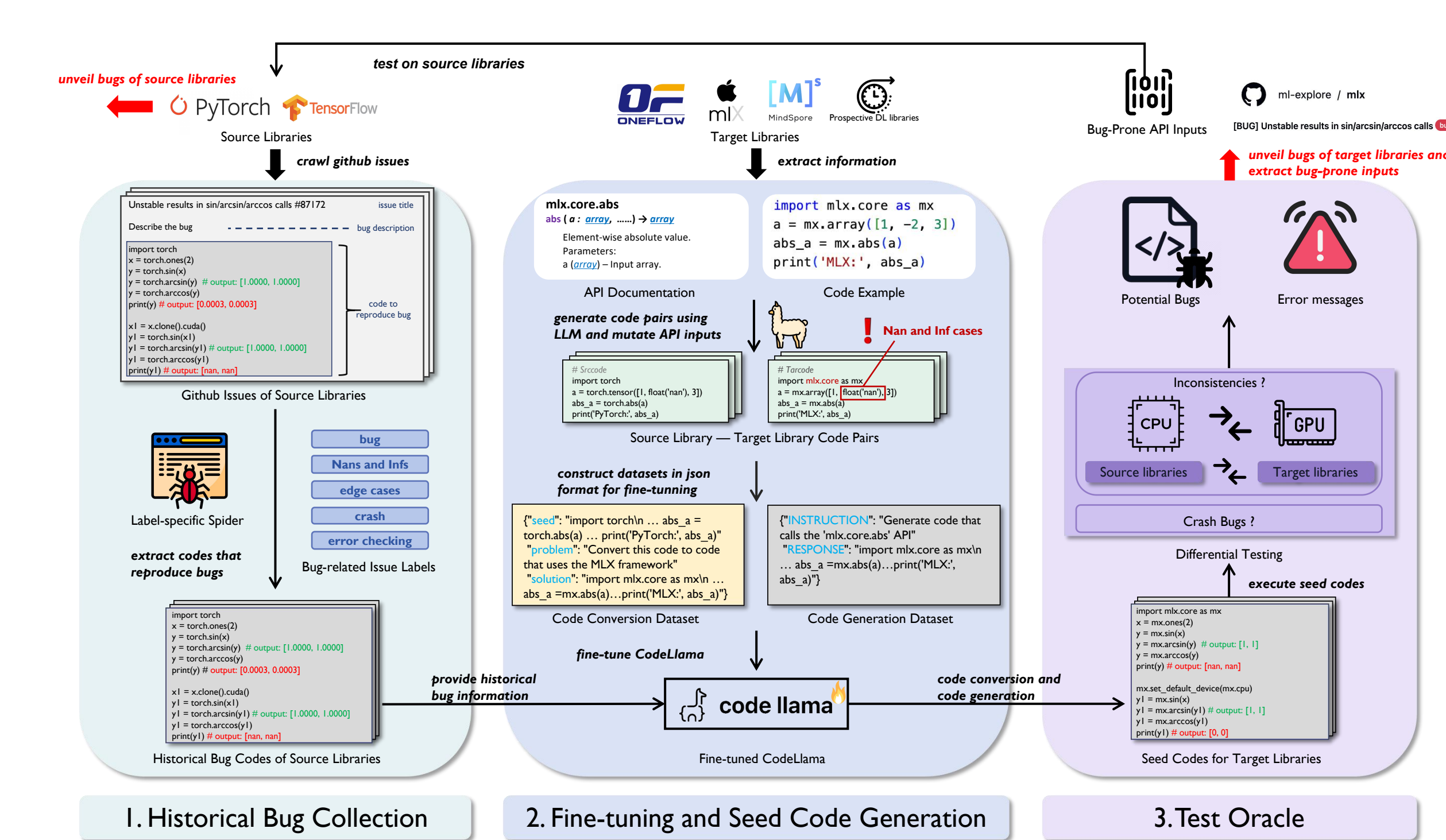
PyTorch和TensorFlow等传统人工智能框架包含大量历史漏洞信息，这些历史漏洞大多来自真实场景，覆盖了复杂的API序列。在苹果MLX、华为MindSpore以及一流OneFlow等新型框架的文档中明确指出这些新型框架的设计大多受传统框架启发。因此，通过研究团队进行实证研究表明，传统框架的漏洞可能在新型框架上仍然存在或以类似形式存在，故而可以用于新型框架的测试用例生成。



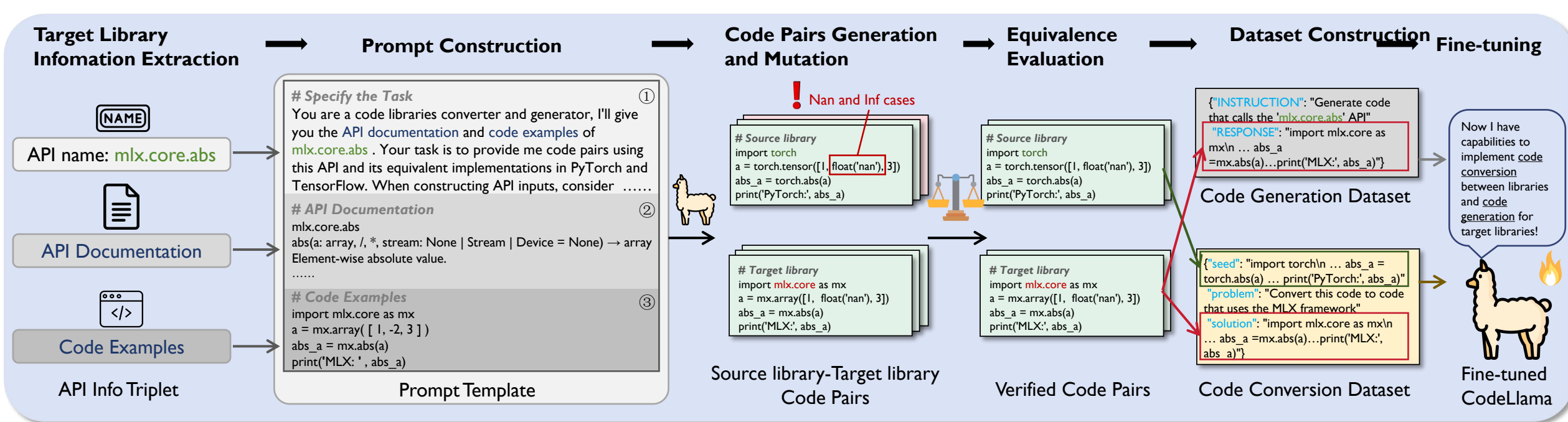
PyTorch某漏洞在MLX, MindSpore, OneFlow仍存在

方法 Methodology

针对上述需求和动机，研究团队创新性地提出了面向开源人工智能框架的模糊测试方法FUTURE。该方法利用新型开源AI框架的文档说明与代码示例构建微调数据集，然后通过LoRA技术实现对大语言模型的微调，并利用微调后模型构建框架间的代码映射关系，将PyTorch、TensorFlow等传统开源框架中的大量历史漏洞转换为新型开源框架的测试用例，通过差分测试运行测试用例并捕捉其中的异常行为，从而实现对开源人工智能框架安全性的有效检测。



FUTURE框架图



数据集构建与大模型微调过程

实验结果 Experiment Results

FUTURE在开源人工智能框架MindSpore、MLX、OneFlow、PyTorch进行了测试。结果显示，FUTURE累计发现新缺陷155个，其中10个缺陷获得CVE编号和4个缺陷获得CNVD编号。此外，该方法还发现了20个可优化问题与12个文档问题。该论文获得ICSE 2025杰出论文奖Distinguished Paper Award。



GitHub仓库



论文链接



This paper is supported by
Open Source Map