

PURPOSE
OSCAR Workshop



DATE
06/28/2026



Agile Compiler Infrastructure for HW/SW Co-Design

Agustin N. Coppari Hollmann

Kris Dong

Dima Nikiforov

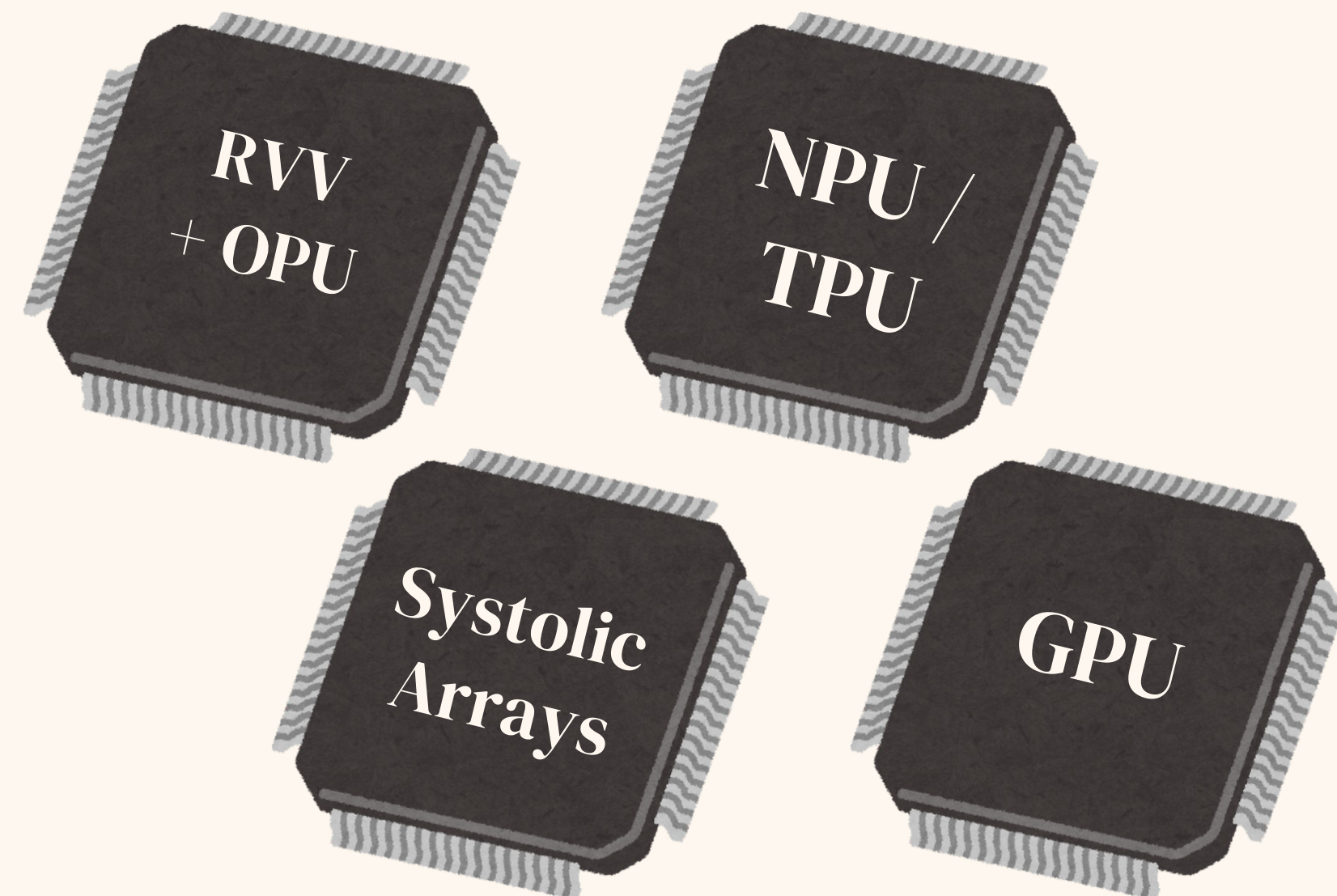
Ashvin Verma

Sophia Yakun Shao

PROJECT PART OF
SLICE

PRESENTED BY
Agustin N. Coppari Hollmann





Many Targets → Lots of opportunities

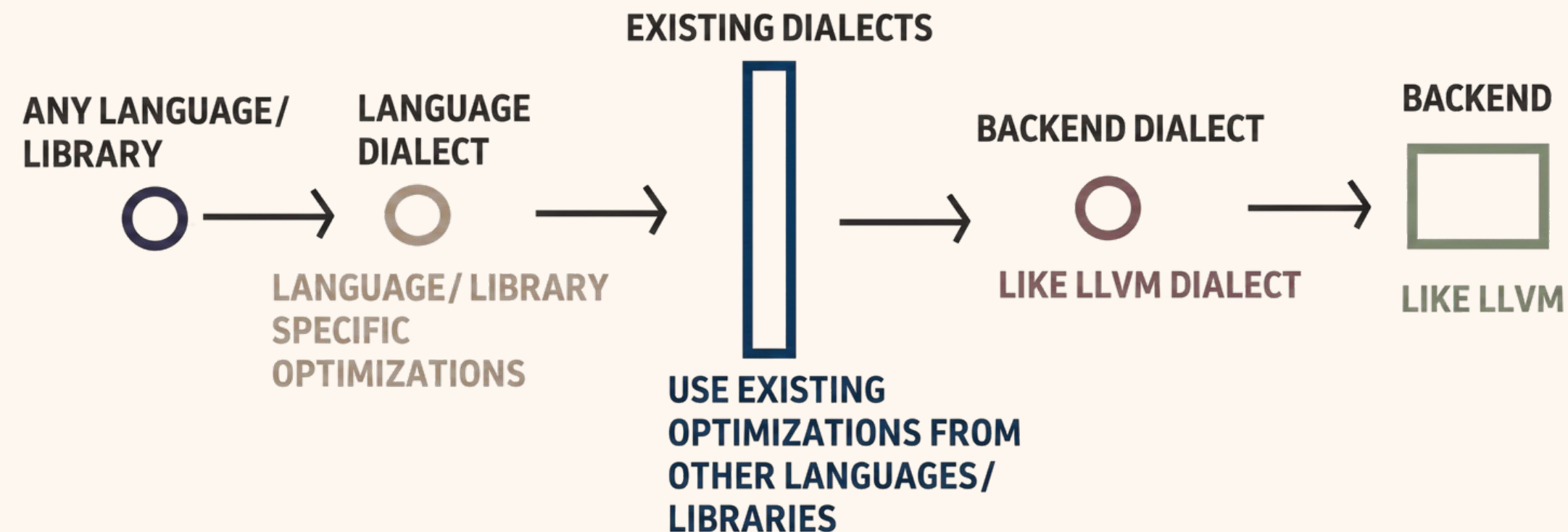
Compilers for AI



What is MLIR?

Multi-Level Intermediate Representation (MLIR) [1]

- **1 IR** for **many “Dialects”** (custom operations, types, and attributes)
- Supports **progressive lowering** with Dialects “Co-existing”
- (Primary use) Bridges **ML Frameworks to HW Targets**





Where is MLIR used?

PURPOSE
OSCAR Workshop

LICENSE.txt

[Release] CUDA Tile IR 13.1 - Initial open-source release

2 months ago

RELEASE.md

Add RELEASE.md (#5896)

last year

CODE_OF_CONDUCT.md

Update project stewards (#2380)

2 years ago

pytorch-hash.txt

Handle pytorch pybind11 changes and bump nightly pins (#...

4 months ago

README.md

[docs] Add 2025 AsiaLLVM talk about data-tiling from Hanh...

7 months ago

RELEASING.md

Update docs, experimental and samples (#19065)

2 years ago

configure_bazel.py

Adding iree-bazel-* wrapper tools for improved Bazel workfl...

3 weeks ago

README

Contributing

Apache-2.0 license

IREE: Intermediate Representation Execution Environment

IREE (Intermediate Representation Execution Environment, pronounced as "eerie") is an [MLIR](#)-based end-to-end compiler and runtime that lowers Machine Learning (ML) models to a unified IR that scales up to meet the needs of

C++ 48.7%

MLIR 22.1%





MLIR Compilers

PURPOSE
OSCAR Workshop



Triton

- Kernel based
- Optimized for GPU
- Tile-based abstraction
- Widely adopted for GPU
- CPU implementation much less mature (even less for RISC-V)



OpenXLA

- Macro-level optimization
- Optimized for TPU
- Plug-in your proprietary drivers / backend compiler as a black-box
- Not a code generator on its own for your ASIC



IREE [2]

- Code-gen + Kernels based
- Can generate bare-metal exec.
- Allows you to include your own pipelines and MLIR dialects
- High customization comes with a cost of maintaining it

Less ← Customizable / Research friendly → More



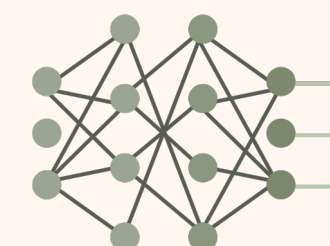
Your average experience with ML compilers

PURPOSE
OSCAR Workshop



- Sit between SW and HW
- Deal with new HW
- Deal with the newest models
- Does their best to do research
- Has to enable SW X on HW Y

Your Favorite
Architecture & ISA
E.g.: RISC-V



ML Model(s)



(Partial) Spec



Executorch

...



XLA



EXO



IREE

Need to change
significant infra...

Model not
supported...

Quantization not
supported...

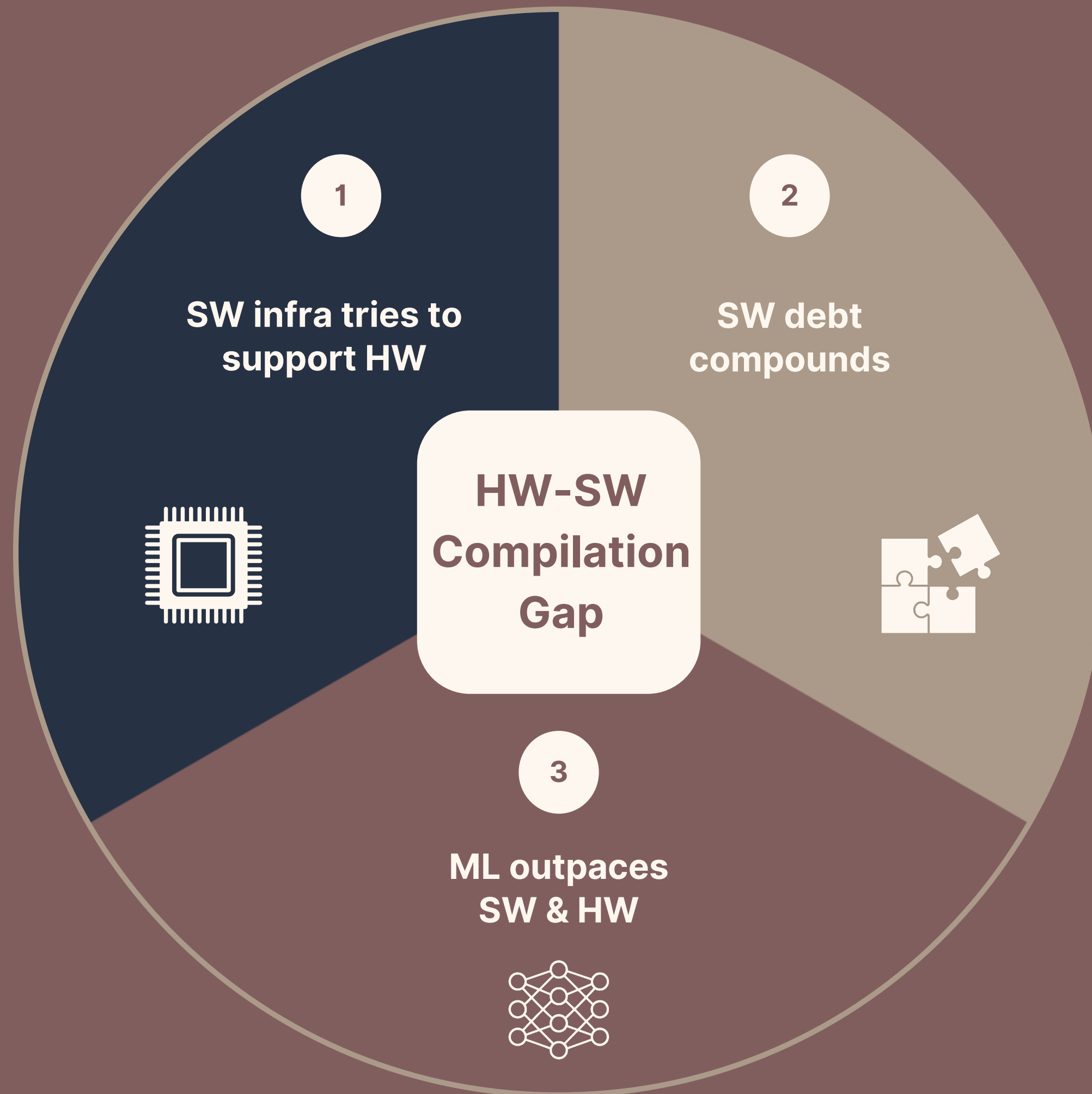
Iteration painfully
slow...

Introduce...
New IR

Compile...
Latest VLA

Introduce...
new target

Compile...
Latest Quantization



SW

/

HW

Models change quickly

Local kernel wins do not imply E2E wins

Targets are increasingly specialized

New ops / quantization schemes appear

Target bringup requires stack-wide changes

Each target exposes different SW interaction

Generated kernels overfit to shapes

Compiler development is slow

Simulation fidelity is expensive

Capture can erase structure

DSE needs faithful workload facts, not guesses

Developed for yesterday's algorithms

•

•

•

•

•

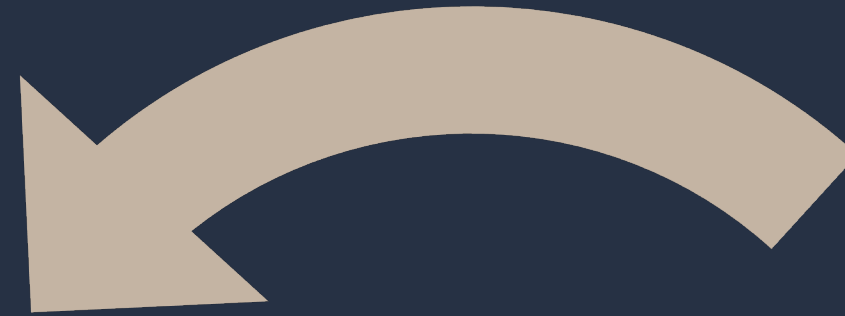
•

•

•

•

SW



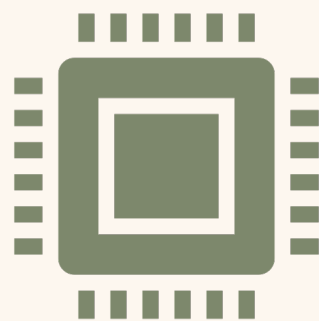
HW





Why Merlin and for who?

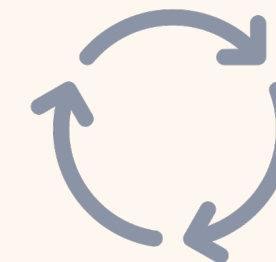
PURPOSE
OSCAR Workshop



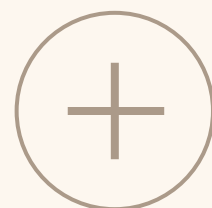
Expose multiple entry points for the compiler to evolve with the HW



Native integration with Zephyr RTOS multicore



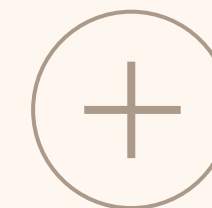
Build to interact with LLMs and agents from the ground-up



You need for your SW stack to no fall behind your HW

You need a runtime friendly for HW simulation

You want to iterate closer to the speed ML does



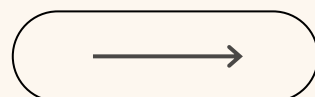
**HW
Co-design**



**Lightweight
MLIR compilation**



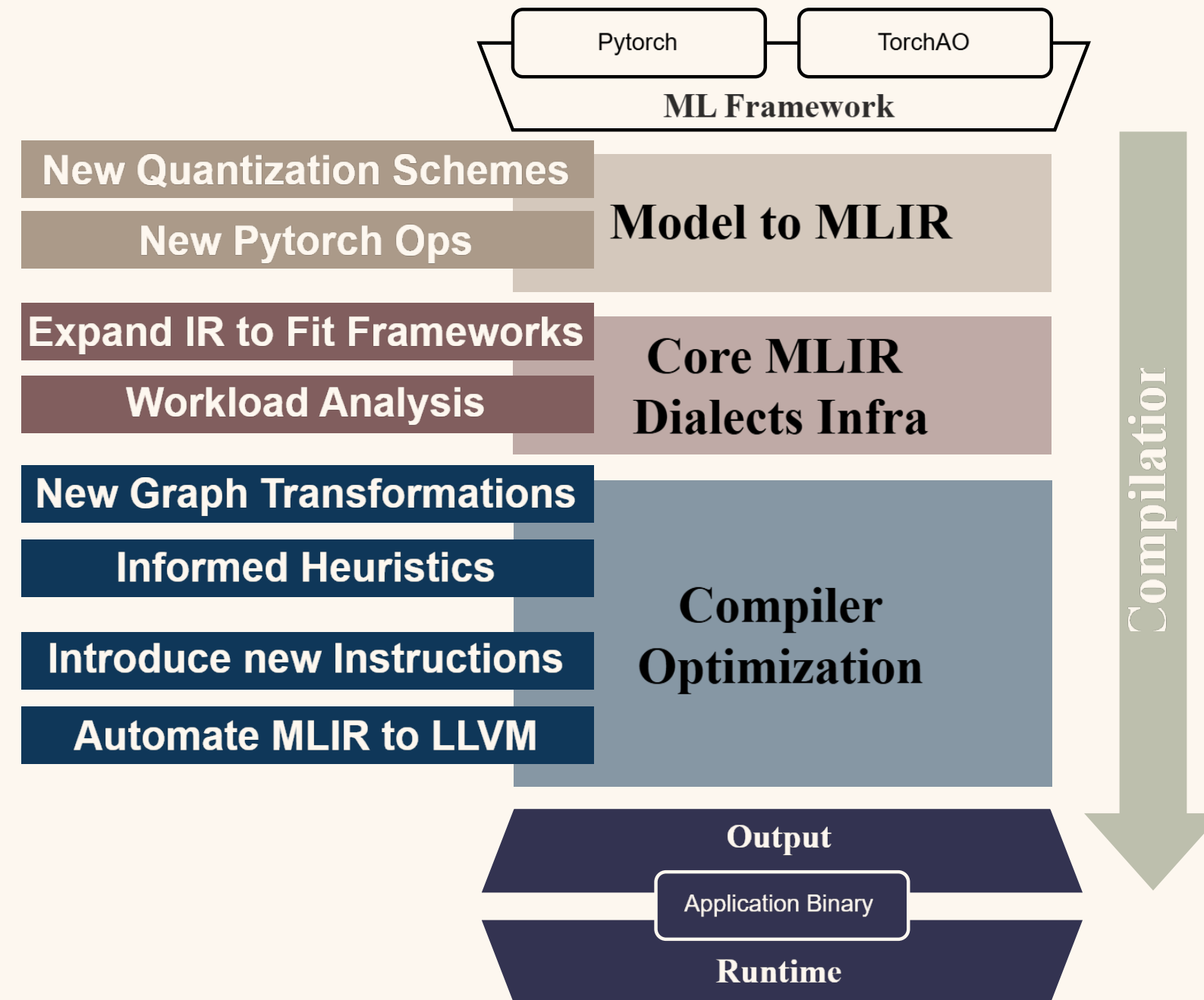
**Bring up
flexibility**





Full-stack agile development

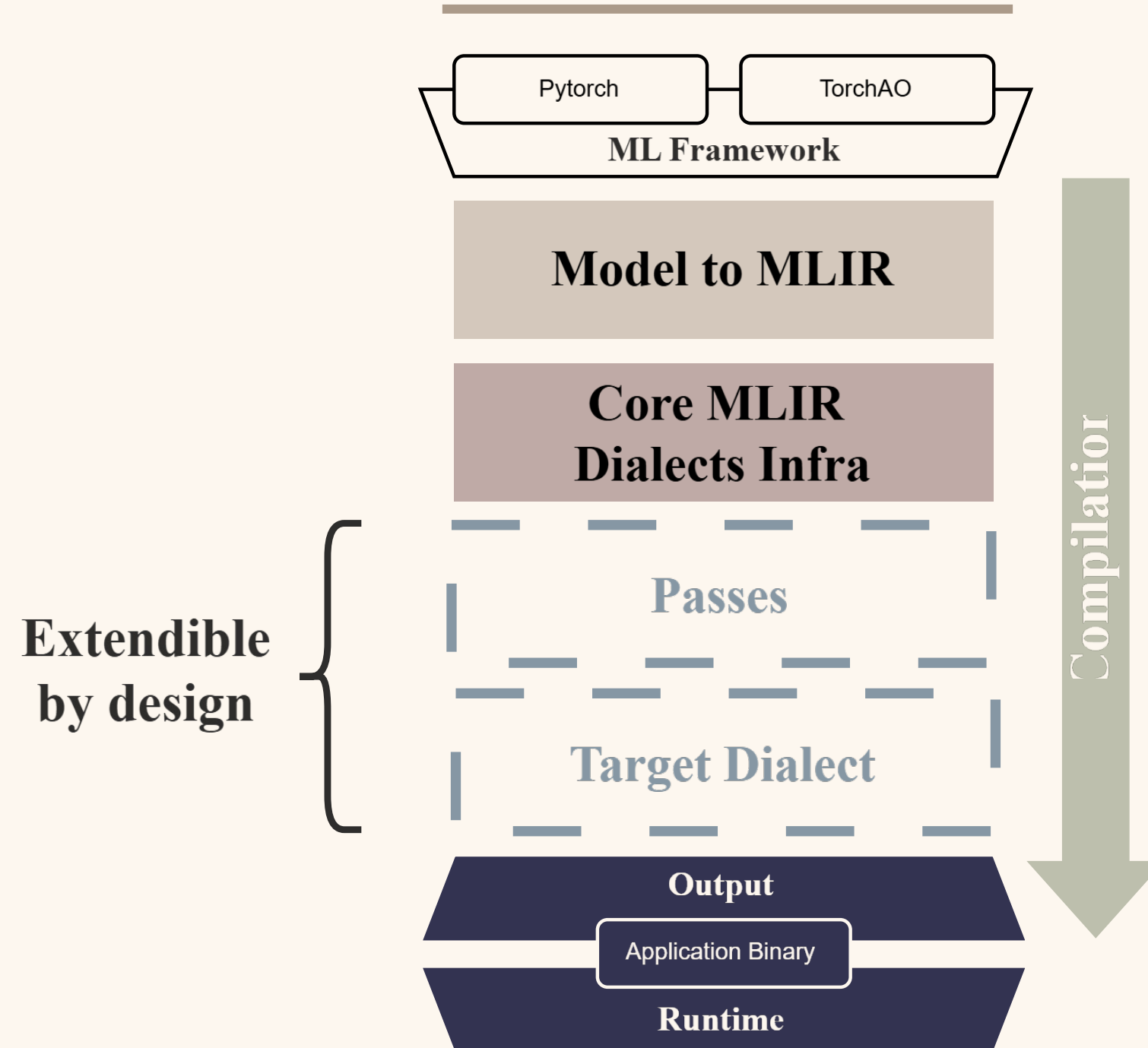
PURPOSE
Merlin ME-Commons





Full-stack agile development

PURPOSE
Merlin ME-Commons

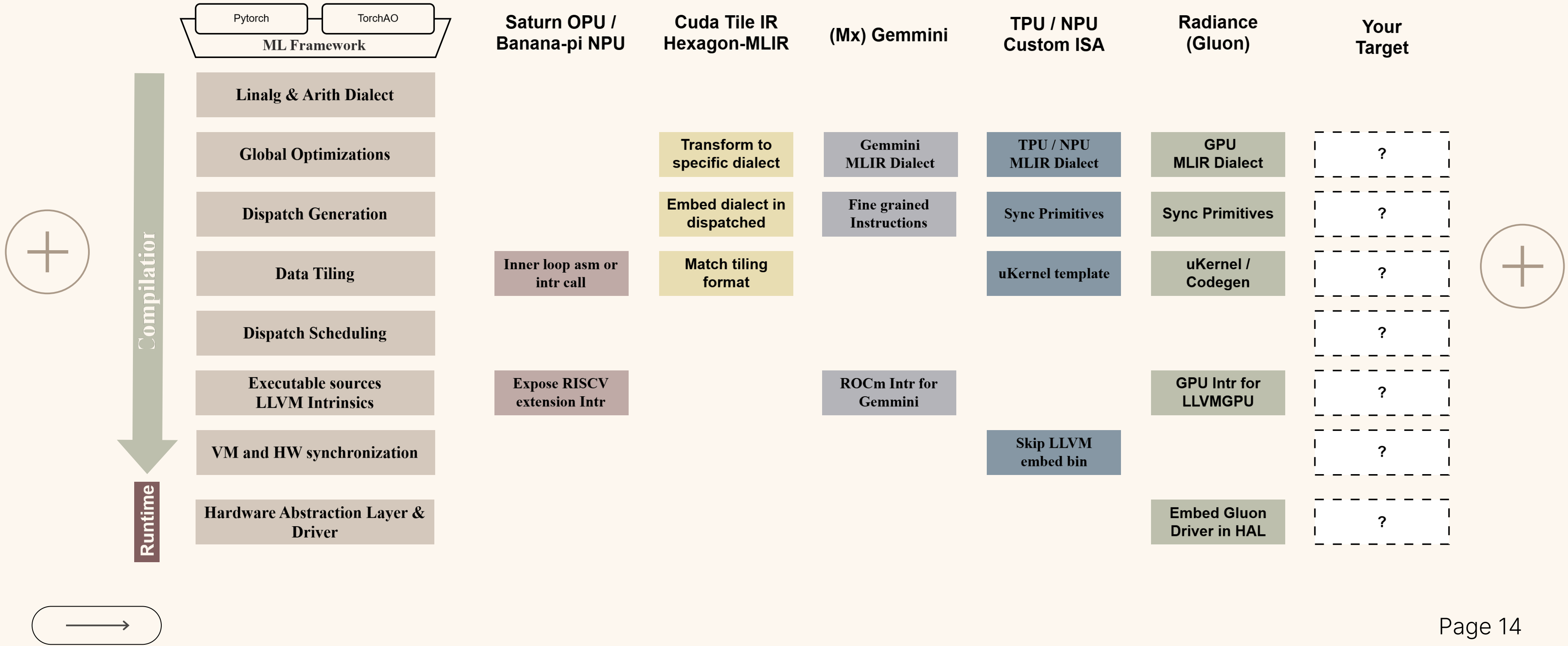


Bring up of
your target
in MLIR



Writing a compiler for your HW target

PURPOSE
OSCAR Workshop



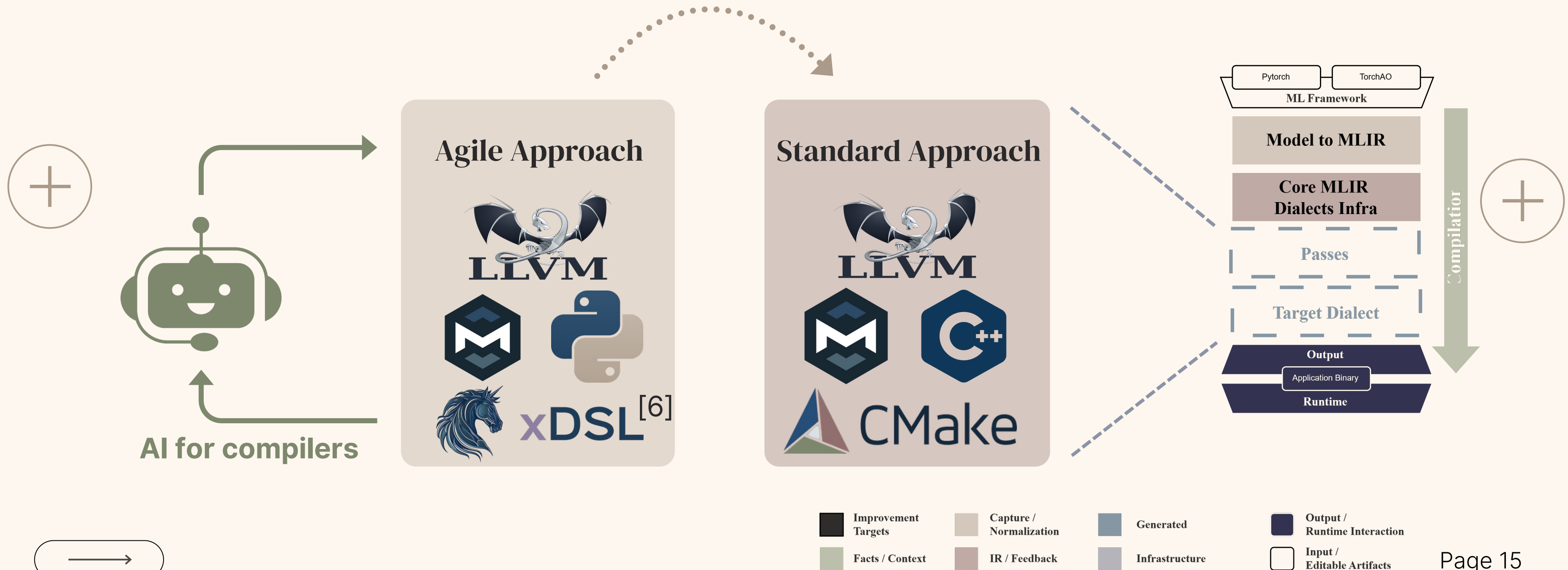


MLIR dialect target bring up

PURPOSE
OSCAR Workshop

We support multiple ways to integrate your dialect & passes

But which approach should you use?

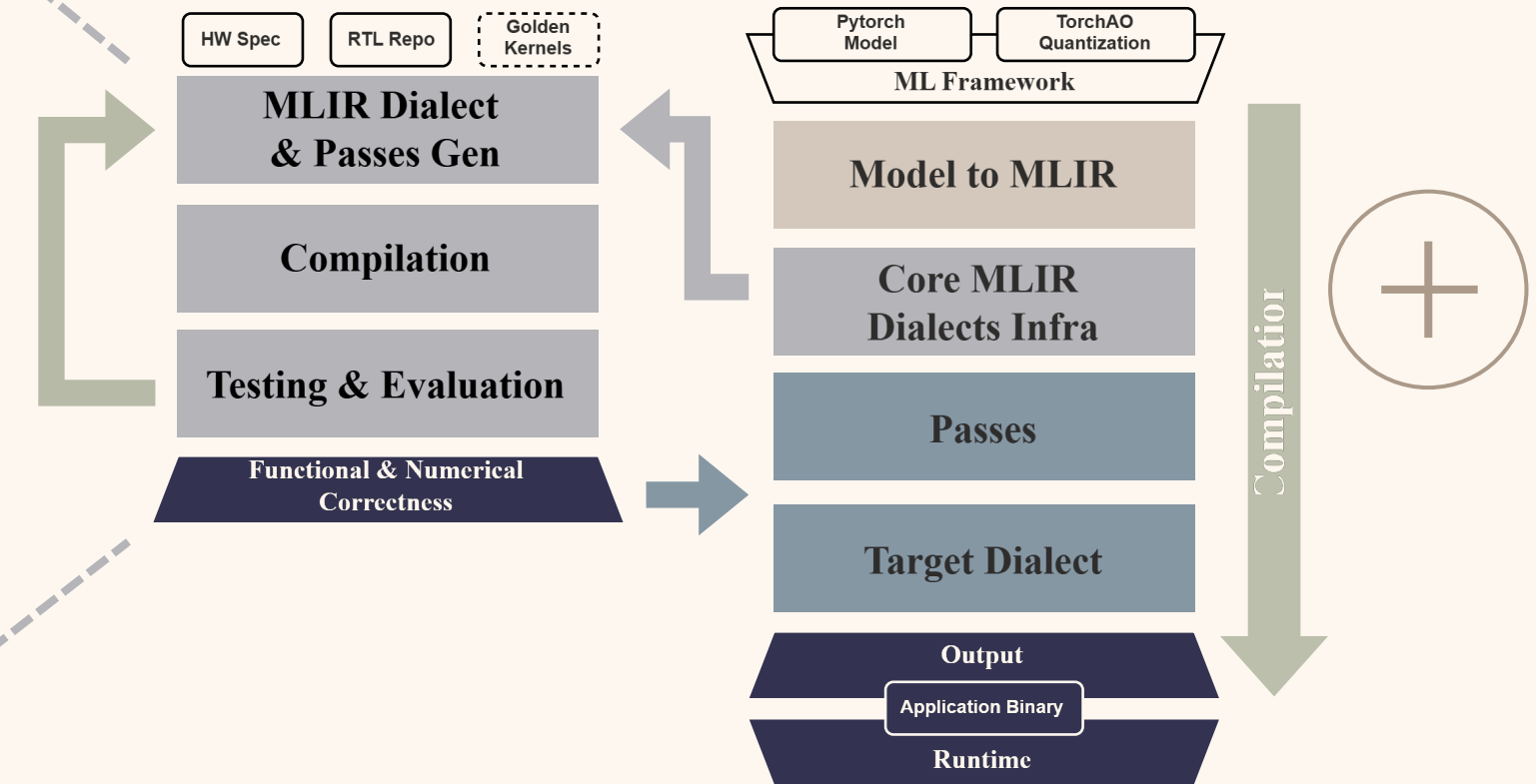
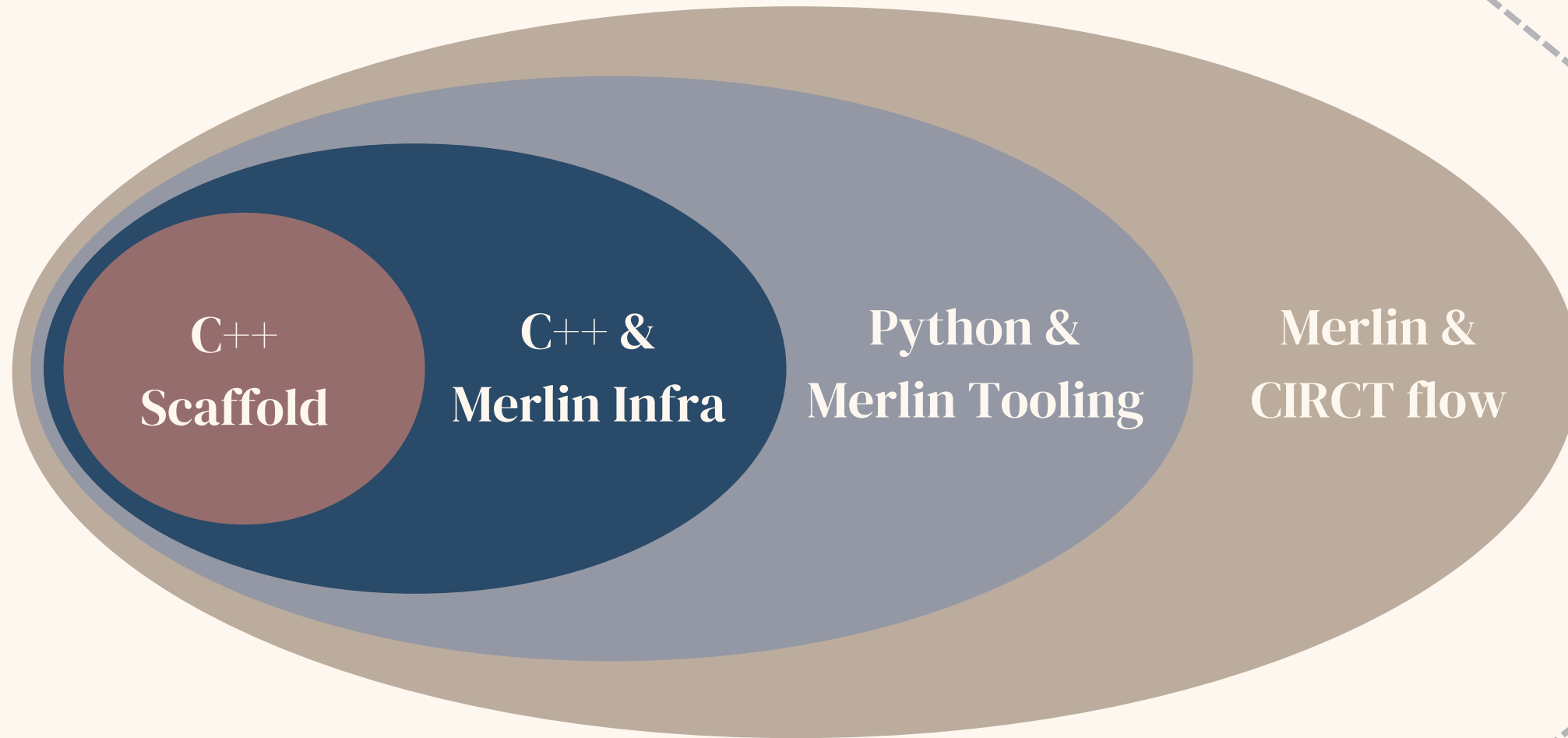




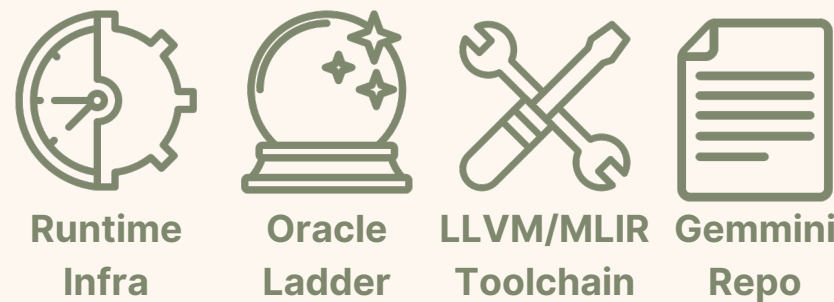
Agentic experiments:

[3] Gemmini case study

PURPOSE
OSCAR Workshop



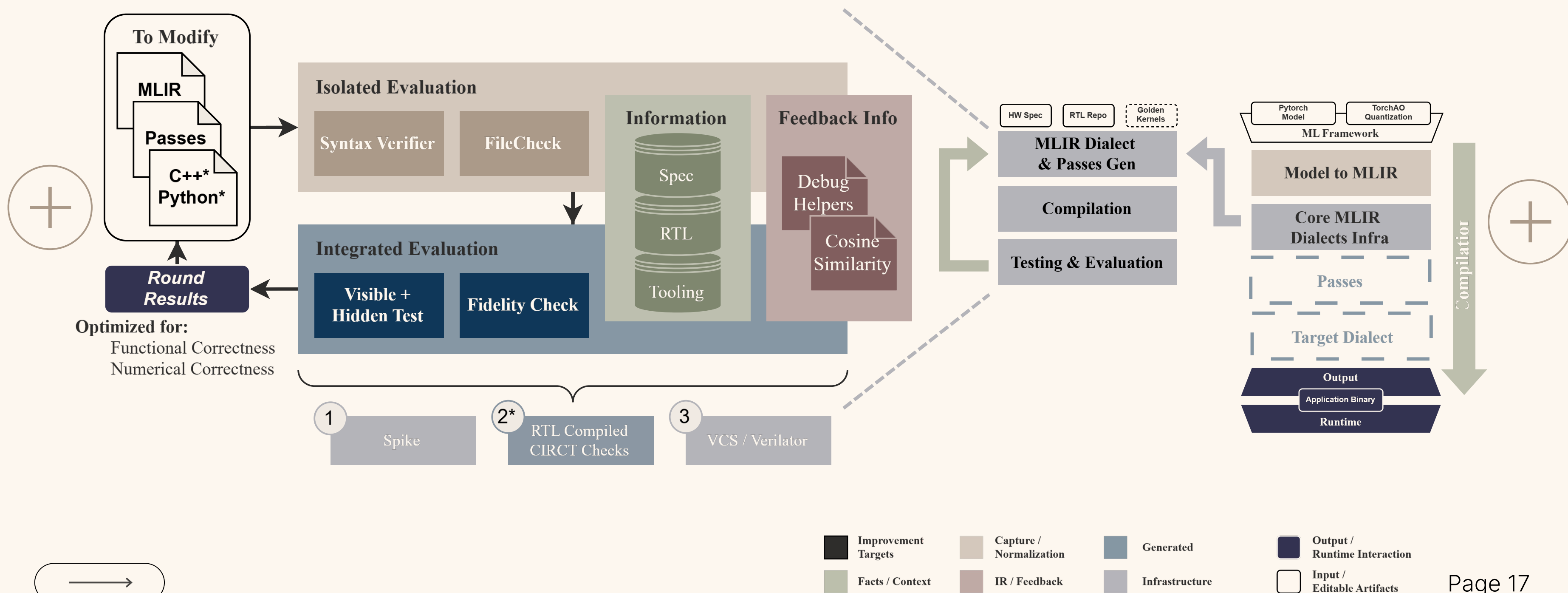
Task → Same
Model → Same
Grading → Same
Sandbox → Same





Agentic MLIR generation loop

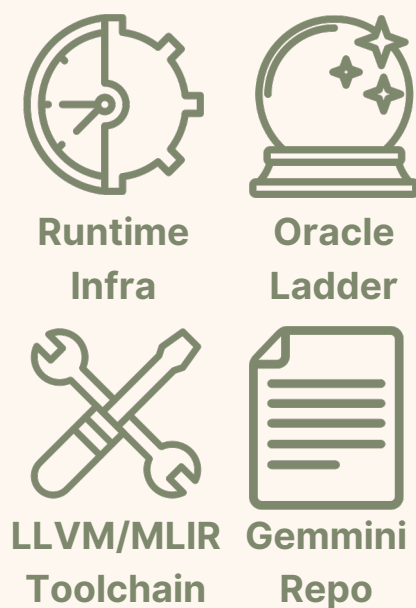
PURPOSE
OSCAR Workshop



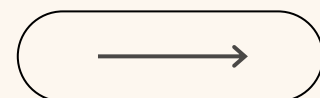
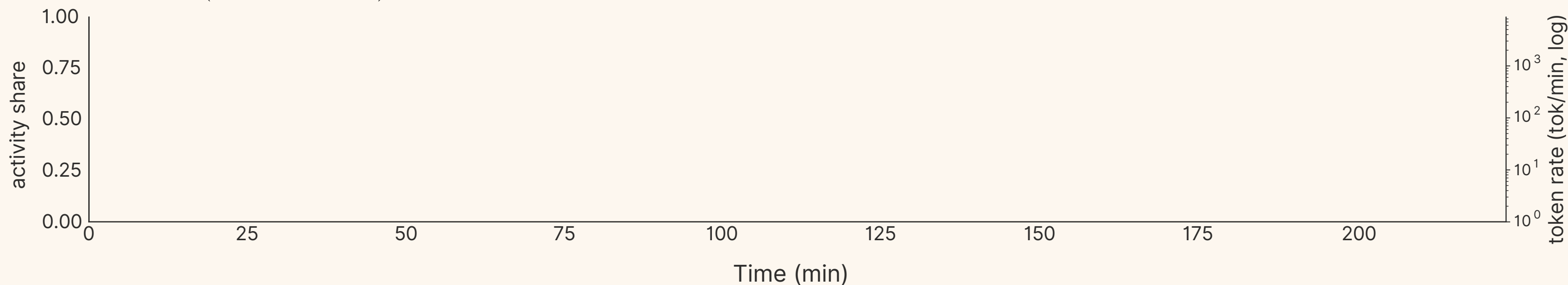


Experiment results: C++ Scaffold

PURPOSE
OSCAR Workshop



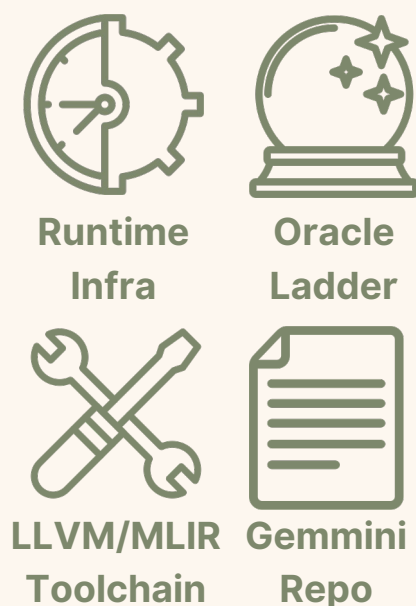
raw C++ (from scratch)



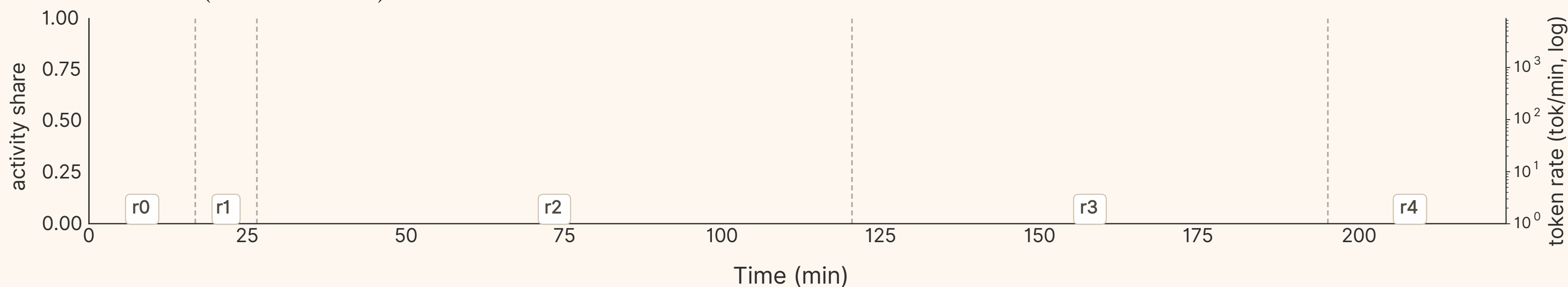


Experiment results: C++ Scaffold

PURPOSE
OSCAR Workshop



raw C++ (from scratch)

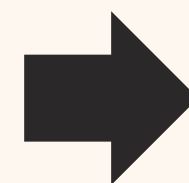
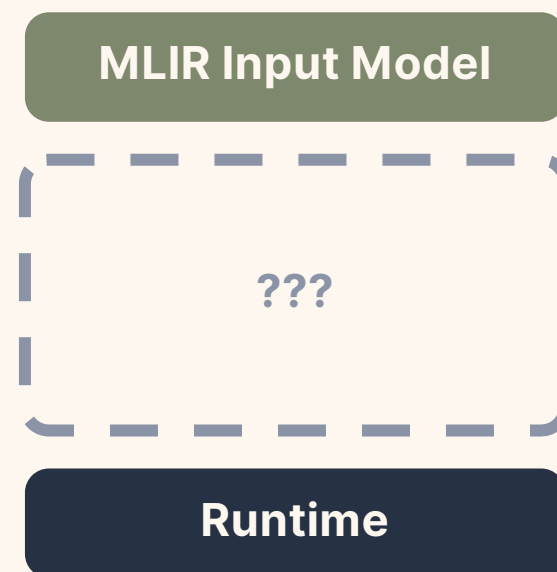
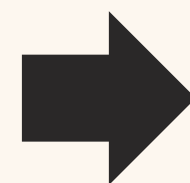
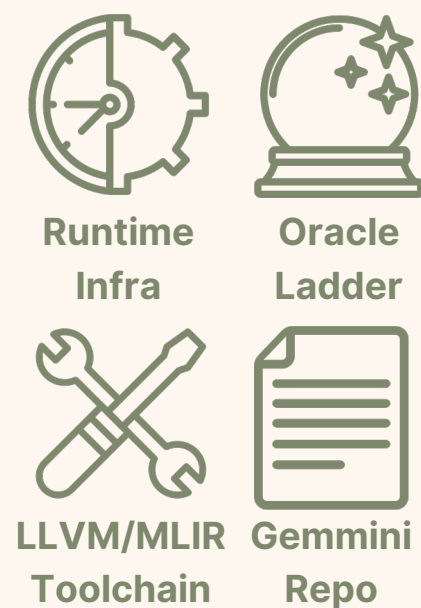


..... round

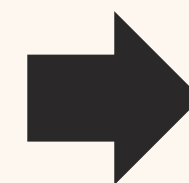


Experiment results: C++ Scaffold

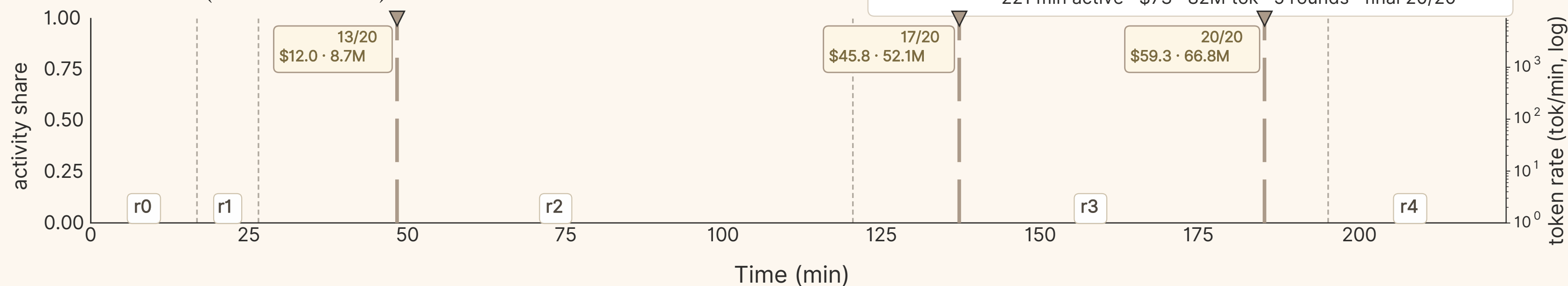
PURPOSE
OSCAR Workshop



*"Given X & runtime
contract build
working dialect"*



raw C++ (from scratch)



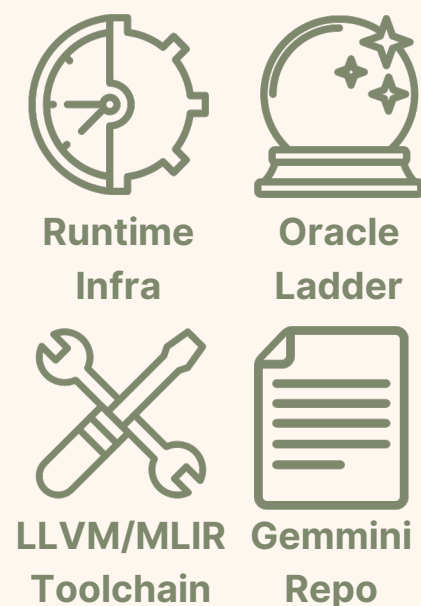
-- test-pass milestone round



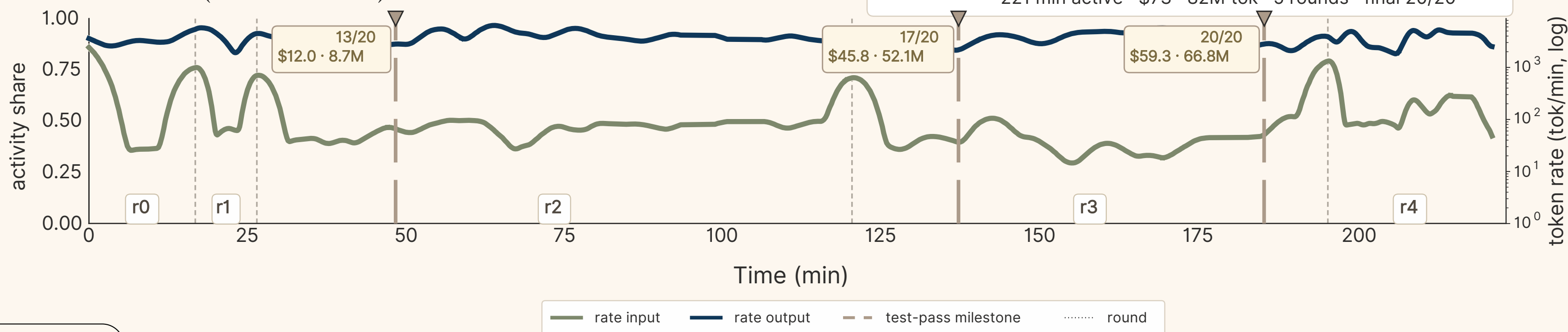


Experiment results: C++ Scaffold

PURPOSE
OSCAR Workshop



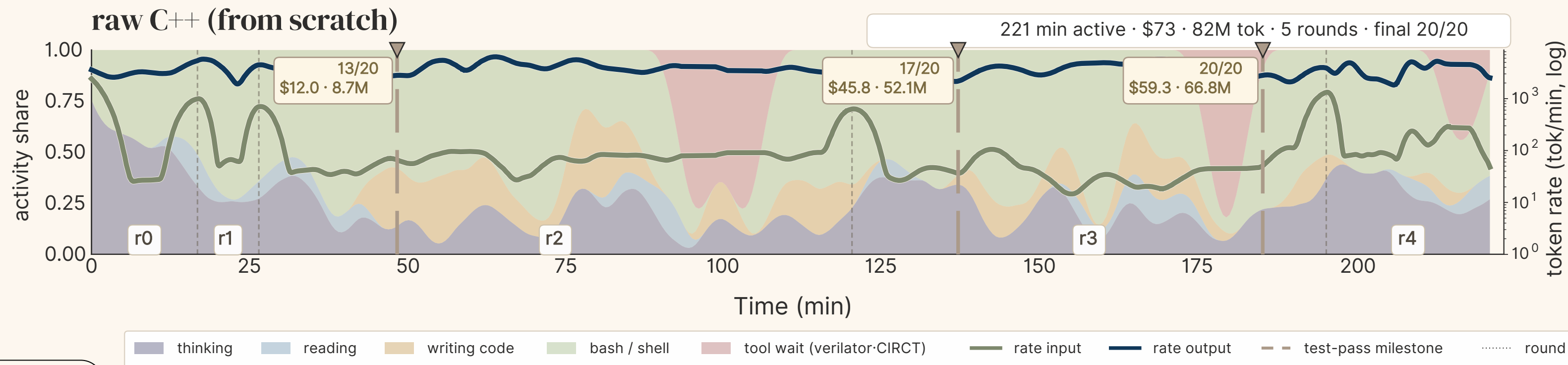
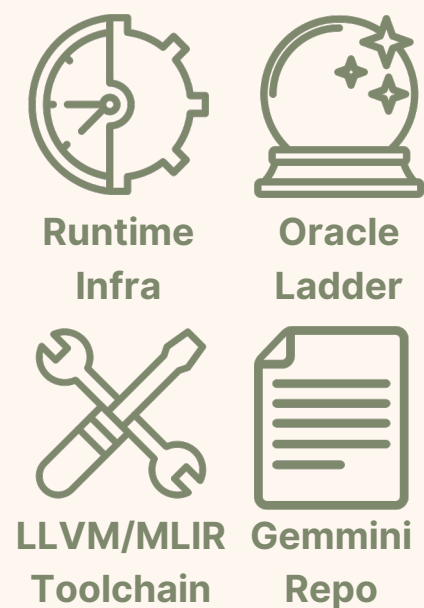
raw C++ (from scratch)





Experiment results: C++ Scaffold

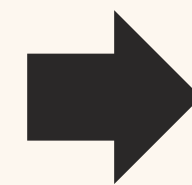
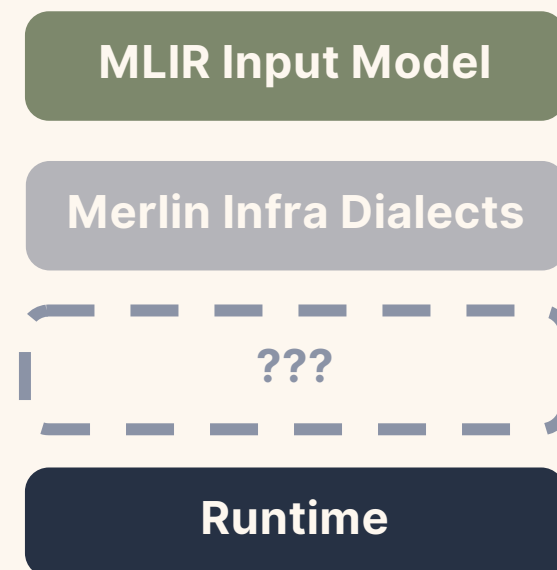
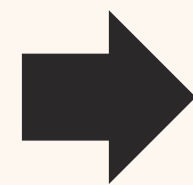
PURPOSE
OSCAR Workshop



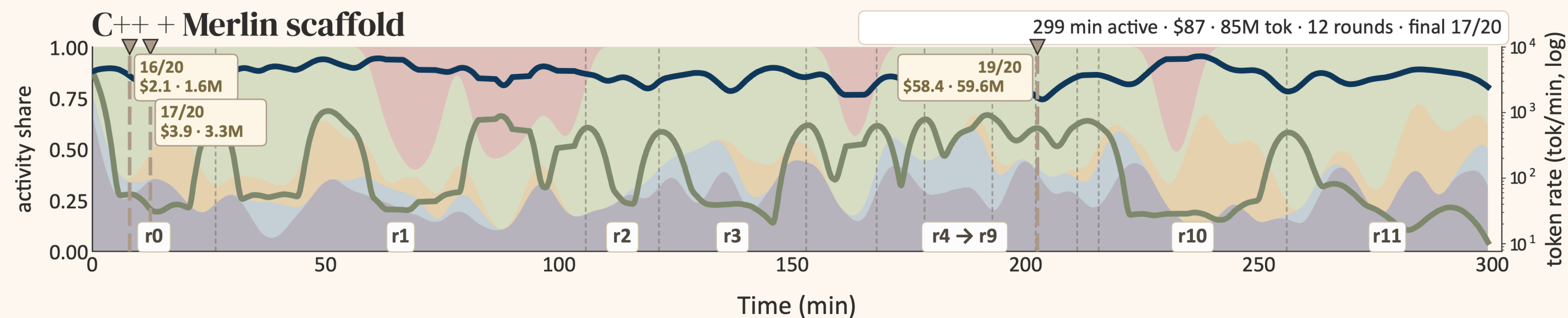
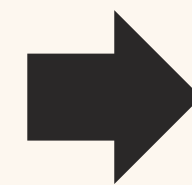


Experiment results: C++ & Merlin infra

PURPOSE
OSCAR Workshop



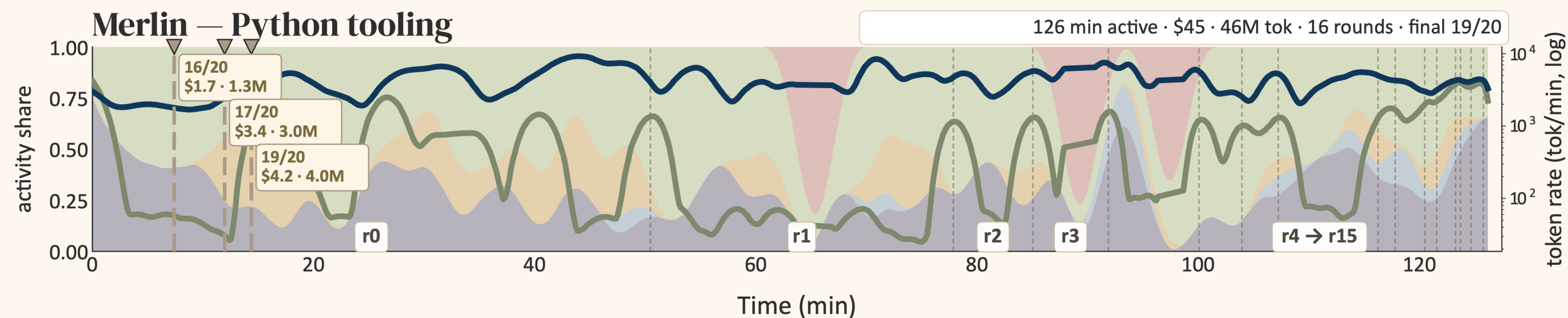
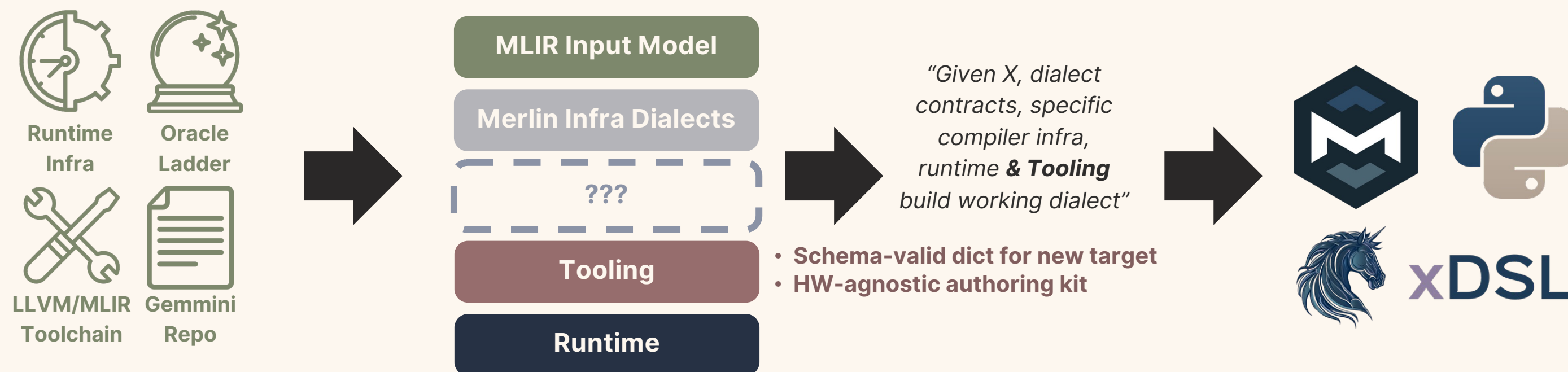
"Given X , **dialect contracts, specific compiler infra & runtime build working dialect**"



→ thinking reading writing code bash / shell tool wait (verilator-CIRCT) rate input rate output test-pass milestone round



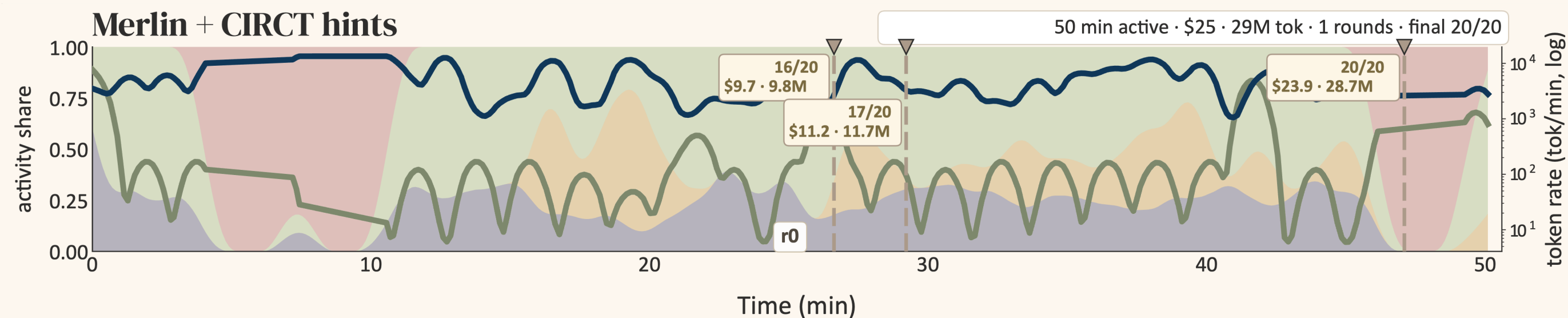
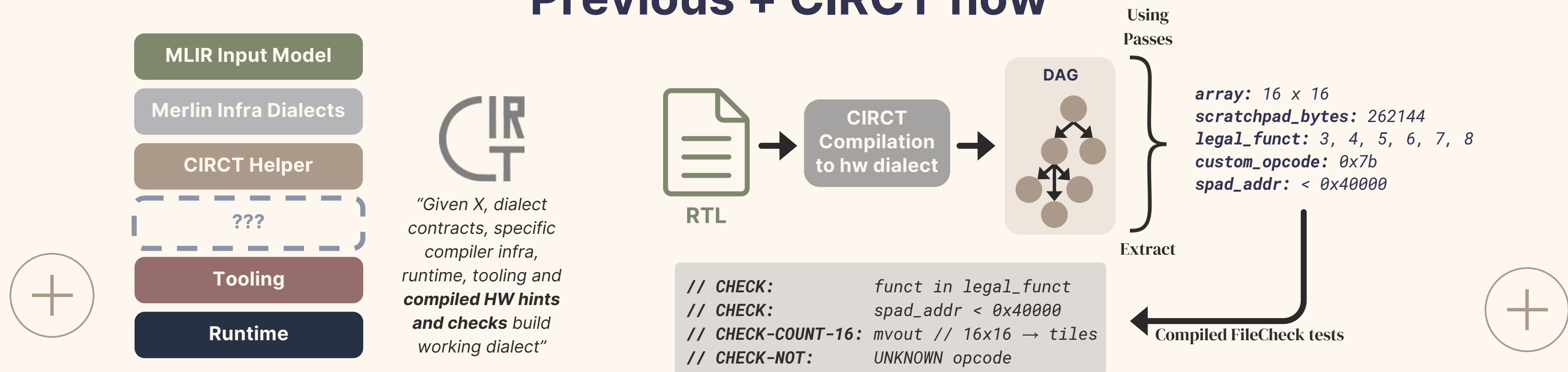
Experiment results: Python & Merlin Tooling





Experiment results: Previous + CIRCT flow

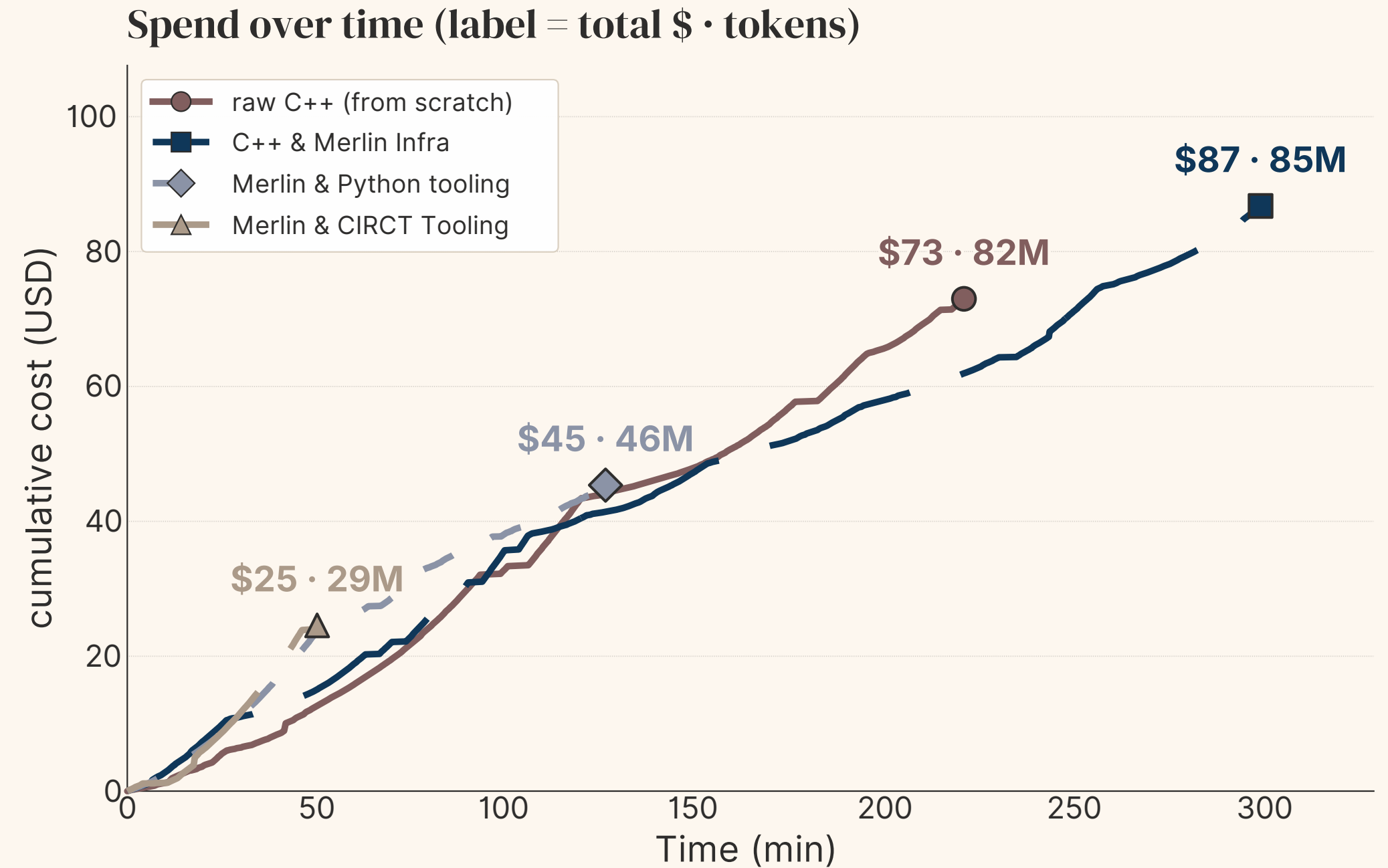
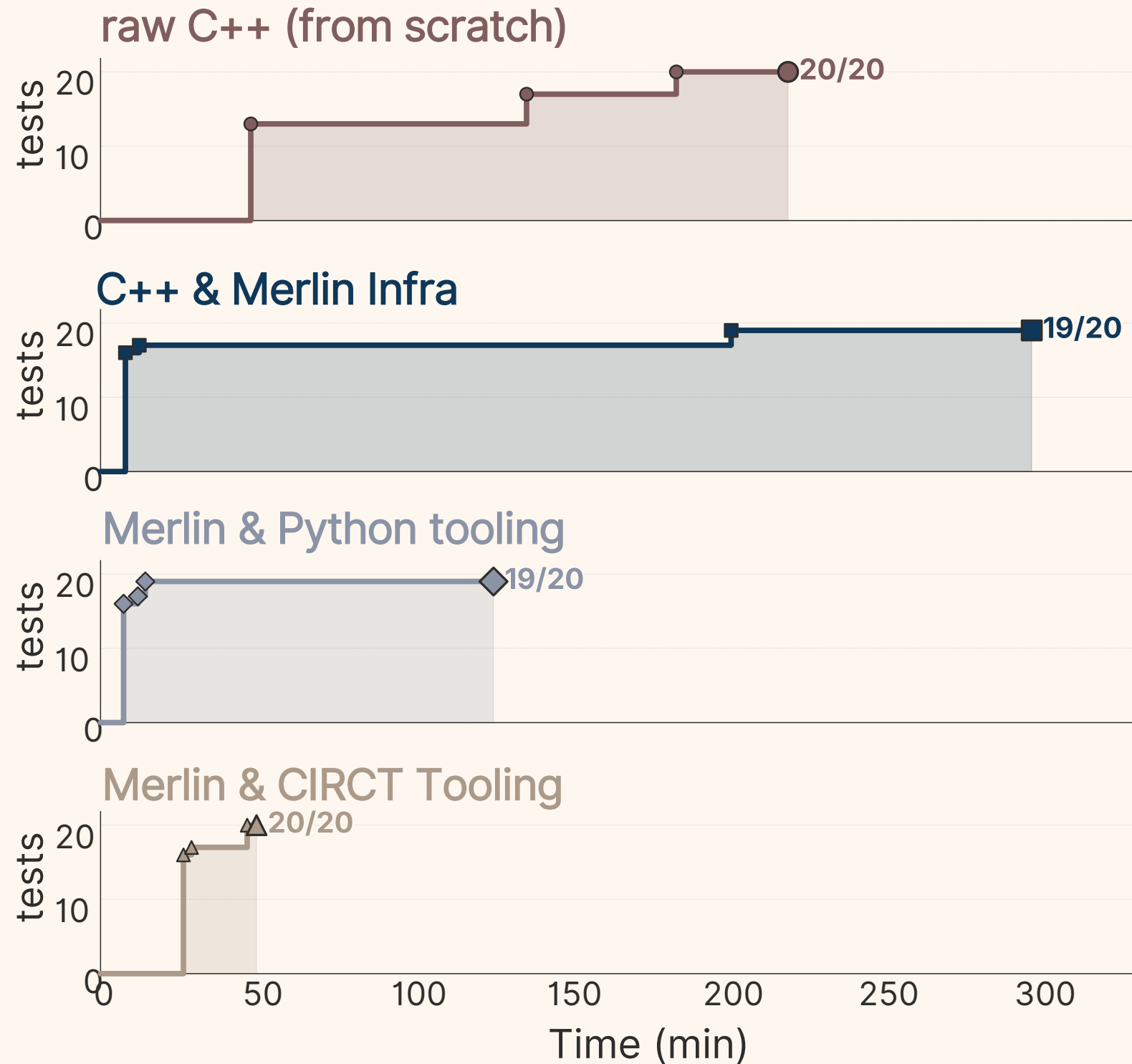
PURPOSE
OSCAR Workshop





Experiment results overview

PURPOSE
OSCAR Workshop

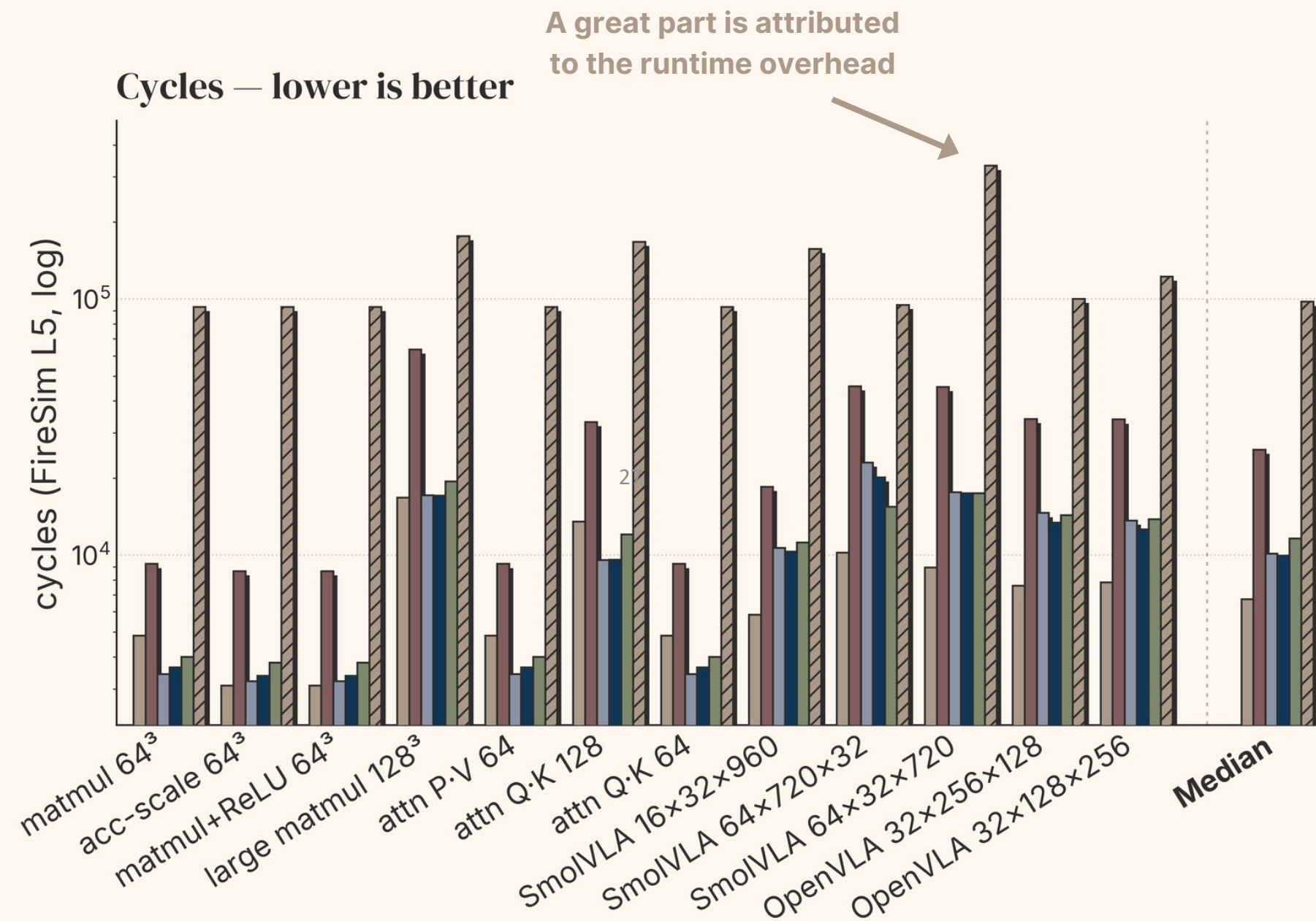




Experiment results: Gen vs. IREE & Golden Lib

PURPOSE
OSCAR Workshop

golden (C lib) raw C++ scaffold C++ (Merlin) Python-Merlin Merlin+CIRCT IREE dialect *

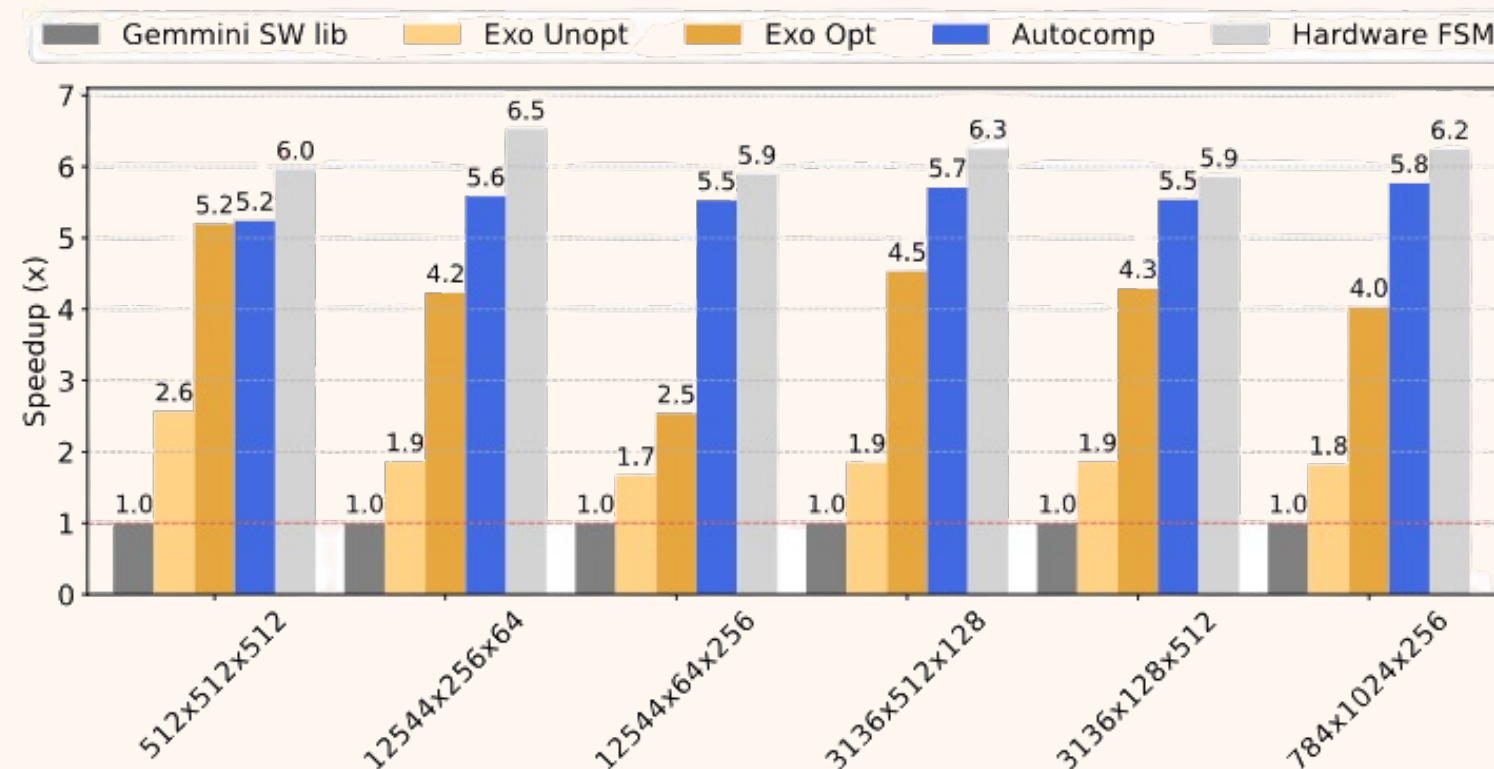


Optimizing your Compiler with Kernels



Generated kernels advance fast but overfit

PURPOSE
OSCAR Workshop

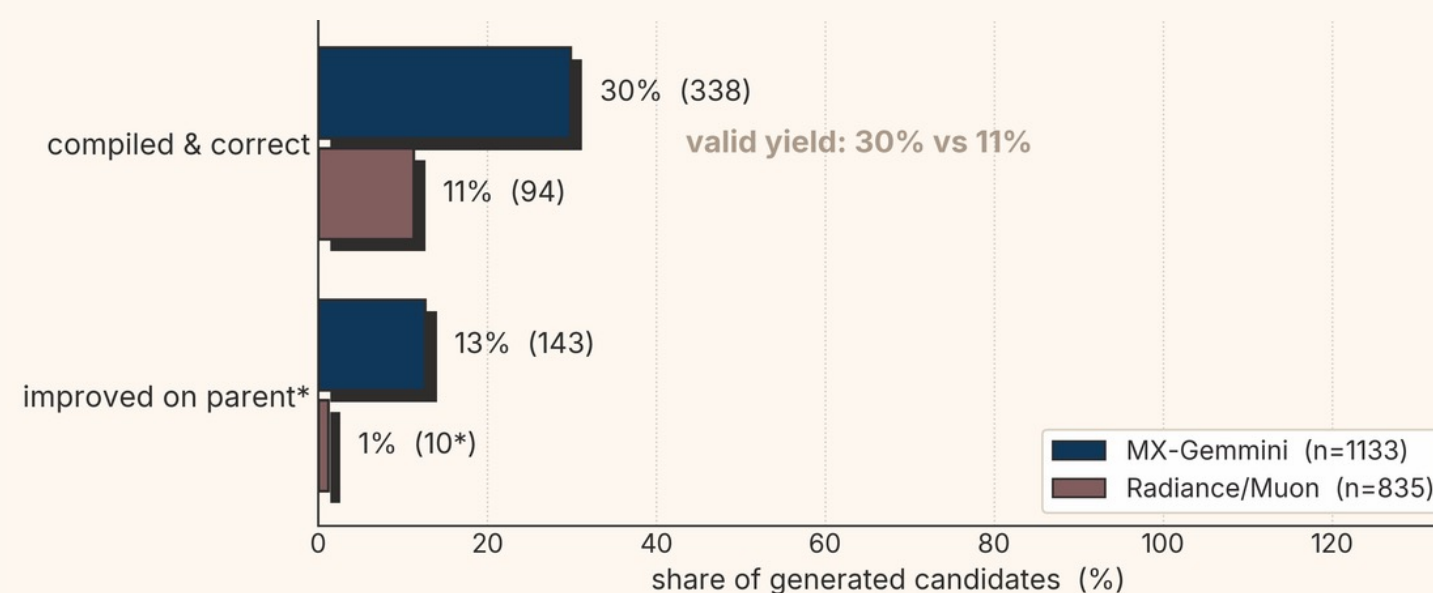


Great results in a short period of time!

Good iteration speed!

Case-by-case optimizations

Iteration in simulators or firesim costly

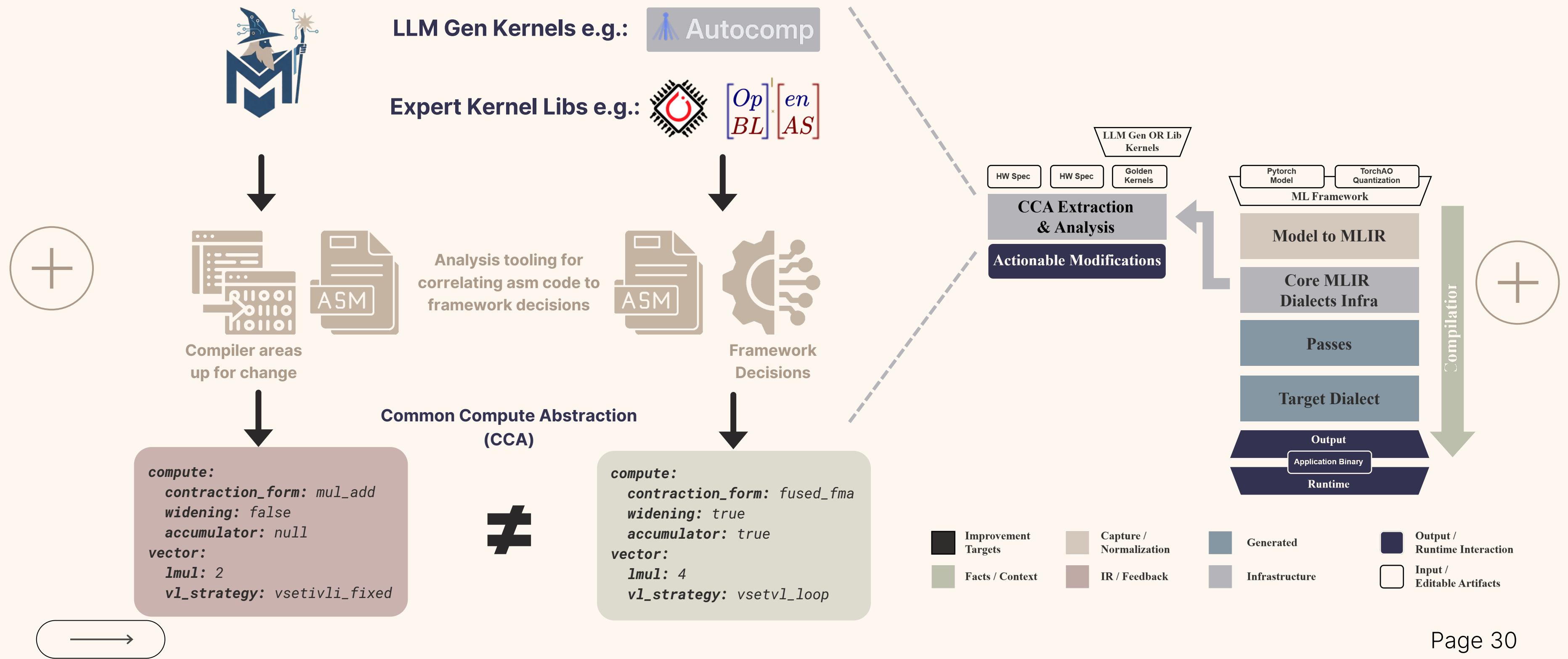


Not Kernels OR Compiler → BUT BEST OF BOTH



Extracting information from kernels

PURPOSE
Merlin ME-Commons





Integrating kernel learnings

PURPOSE
Merlin ME-Commons

Our CCA

≠

Expert CCA

What to improve?

Tiling /
Dataflow

Fusion /
Layout

Register
Residency

Instruction
Selection

Runtime /
Sync

Actions to take:

Pass

Heuristic

Flag

Knob

Beam Search

Fork Validity & Build

Did the CCA divergence close?

Functional
Correctness

Does the IR behave as intended?

Utilized HW
Capabilities

Real or fake speedup?

HW-in-the-loop
True Cost

What is the performance?

If fail at gate:
Candidate pruned

Changing a
heuristic like
MR

LLM Gen OR Lib
Kernels

HW Spec

HW Spec

Golden
Kernels

CCA Extraction
& Analysis

Actionable Modifications

Beam Search

Pytorch
Model

TorchAO
Quantization

Model to MLIR

Core MLIR
Dialects Infra

Passes

Target Dialect

Output

Application Binary

Runtime

Compiler

Improvement
Targets

Capture /
Normalization

Generated

Output /
Runtime Interaction

Facts / Context

IR / Feedback

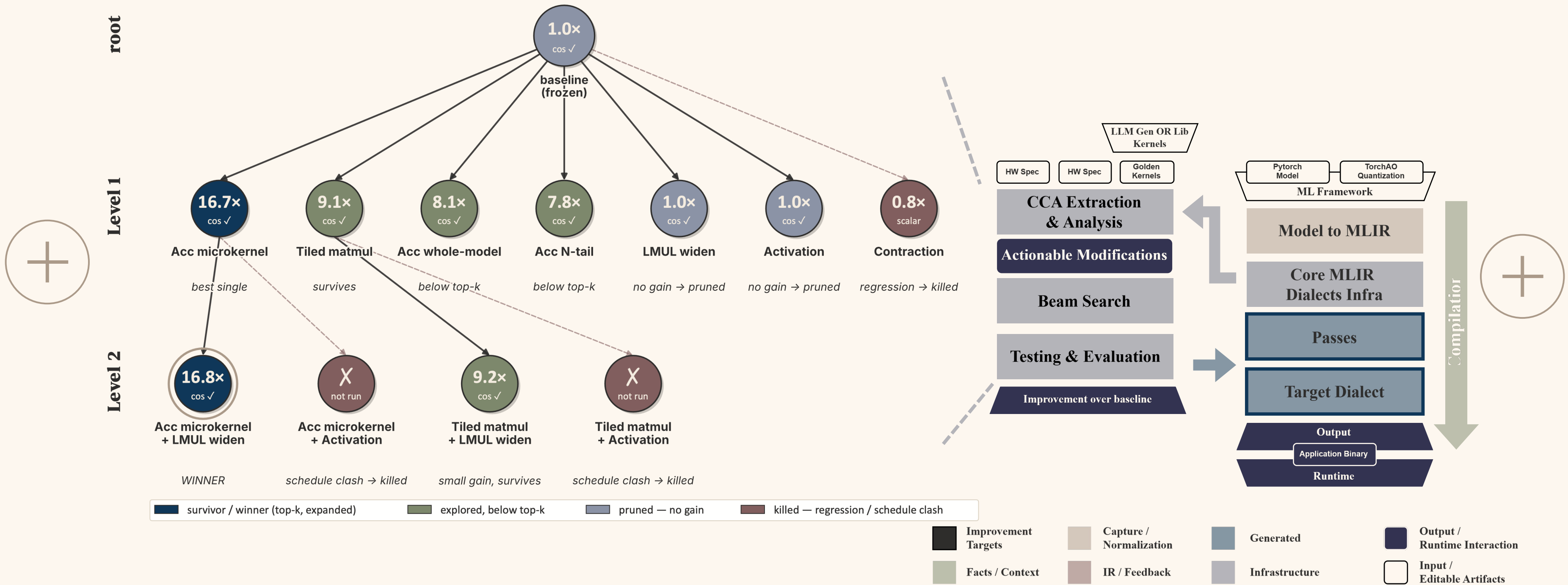
Infrastructure

Input /
Editable Artifacts



Beam search through each modification

PURPOSE
Merlin ME-Commons



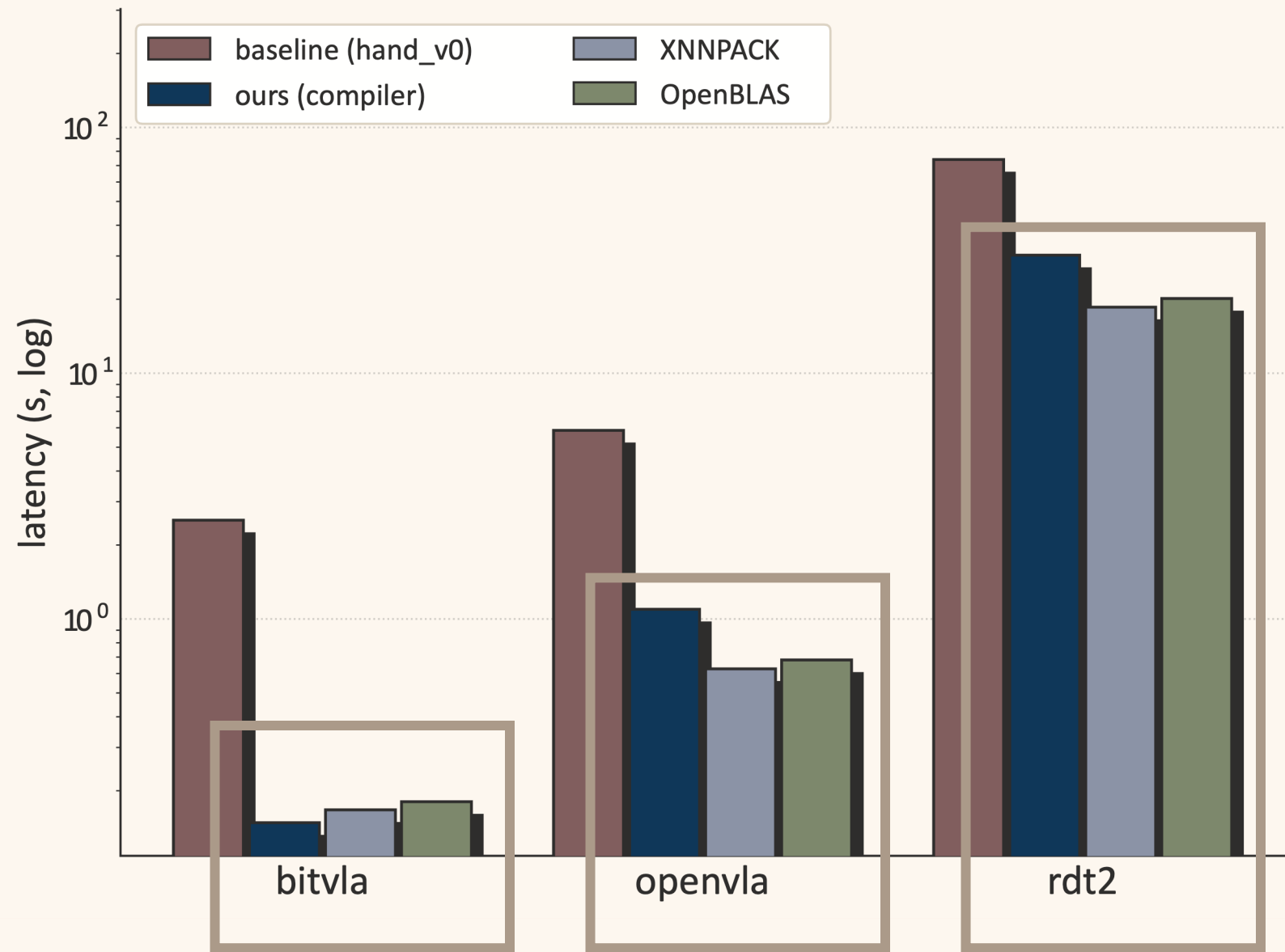


Experiment results

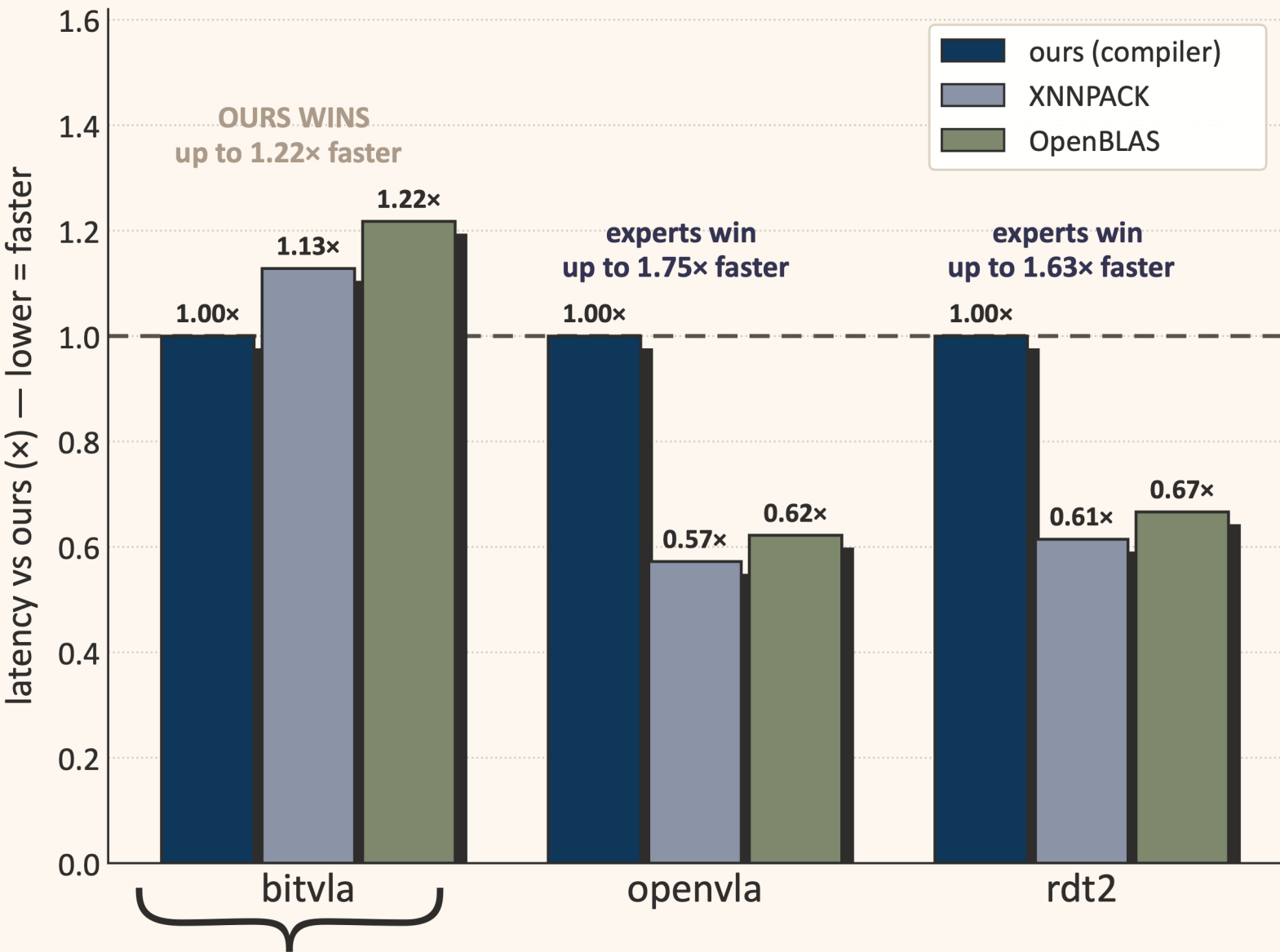
on Saturn RVV [4]

PURPOSE
OSCAR Workshop

All four backends



Head-to-head comparison



Partitioning of model
beyond XNNPACK lib
boundaries

We find ways to optimize
compute & requant for
1.58-bit weights

Optimized packing &
removal of dispatch
overhead



Extending our study cases

PURPOSE
OSCAR Workshop

1st BASELINE

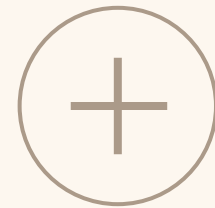


IREE

Currently all targets mentioned have working examples in our OOT IREE extension



All targets will be migrated through target gen experiments!



Spatial Array

e.g.: Gemmini

Working!

NPU

e.g.: Atlas NPU

TODO

Mix precision
Spatial Array

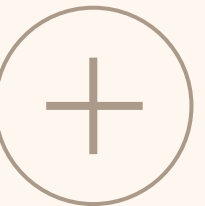
e.g.: Gemmini-MX

TODO

SIMT Core

e.g.: Radiance

In progress



[5]



RADIANCE

Muon Talk

@OSCAR:

1:30 p.m. Today!





Points left to address

**Agentic
MLIR Gen**

- + data points
- + HW targets

- **Compare to
TVM & EXO**

**Improvements
from CCA**

**Workload
Analysis Tool**

- **Test with
DSE tool**

PURPOSE
OSCAR Workshop





THANK YOU
FOR YOUR TIME!

References

- [1] C. Lattner, M. Amini, U. Bondhugula, A. Cohen, A. Davis, J. Pienaar, R. Riddle, T. Shpeisman, N. Vasilache, and O. Zinenko, “MLIR: A compiler infrastructure for the end of Moore’s law,” arXiv preprint arXiv:2002.11054, 2020.
- [2] The IREE Authors, “IREE,” 2019. [Online]. Available: <https://iree.dev/>. Repository: <https://github.com/iree-org/iree>
- [3] H. Genc, S. Kim, A. Amid, A. Haj-Ali, V. Iyer, P. Prakash, J. Zhao, D. Grubb, H. Liew, H. Mao, A. Ou, C. Schmidt, S. Steffl, J. Wright, I. Stoica, J. Ragan-Kelley, K. Asanovic, B. Nikolic, and Y. S. Shao, “Gemmini: Enabling systematic deep-learning architecture evaluation via full-stack integration,” in Proceedings of the 58th Annual Design Automation Conference (DAC), 2021.
- [4] J. Zhao, D. Grubb, M. Rusch, T. Wei, K. Anderson, B. Nikolic, and K. Asanovic, “The Saturn Microarchitecture Manual,” EECS Department, University of California, Berkeley, Tech. Rep. UCB/EECS-2024-215, Dec. 2024. [Online]. Available: <http://www2.eecs.berkeley.edu/Pubs/TechRpts/2024/EECS-2024-215.html>
- [5] H. Kim, R. R. Yan, S. Anand, J. You, M. Shi, C. W. Fletcher, and Y. S. Shao, “Radiance,” GitHub repository, 2026. [Online]. Available: <https://github.com/ucb-bar/radiance>
- [6] Fehr, M., Weber, M., Ulmann, C., Lopoukhine, A., Lücke, M.P., Degioanni, T., Vasiladiotis, C., Steuwer, M. and Grosser, T., 2025, March. xdsl: Sidekick compilation for ssa-based compilers. In *Proceedings of the 23rd ACM/IEEE International Symposium on Code Generation and Optimization* (pp. 179-192).