

# future of **AI-HW** & **GPU** in SoCET

April 30th, 2026

background

~~future~~ of **AI-HW** & **GPU**  
in SoCET + little bit about future

April 30th, 2026

**motivation**

**history & growth**

**future plans**

**who should join?**

**motivation**

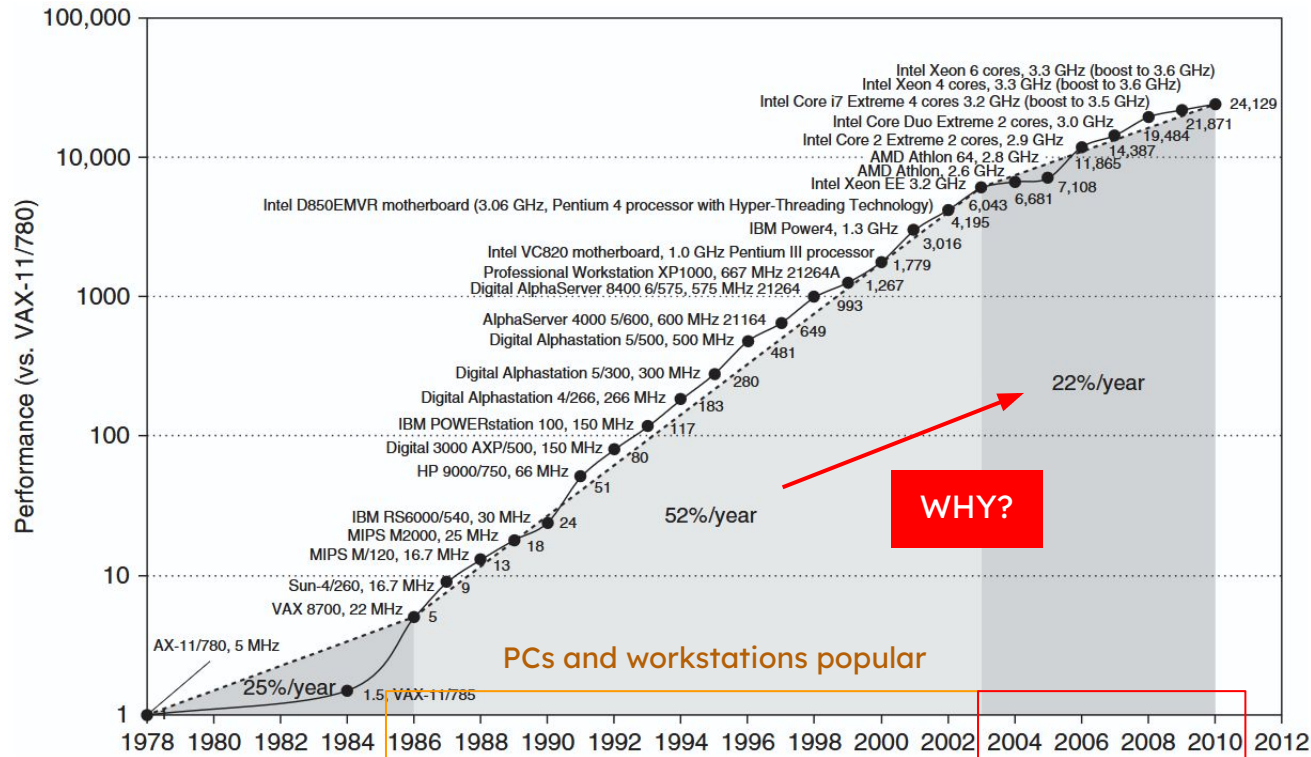
history & growth

future plans

who should join?

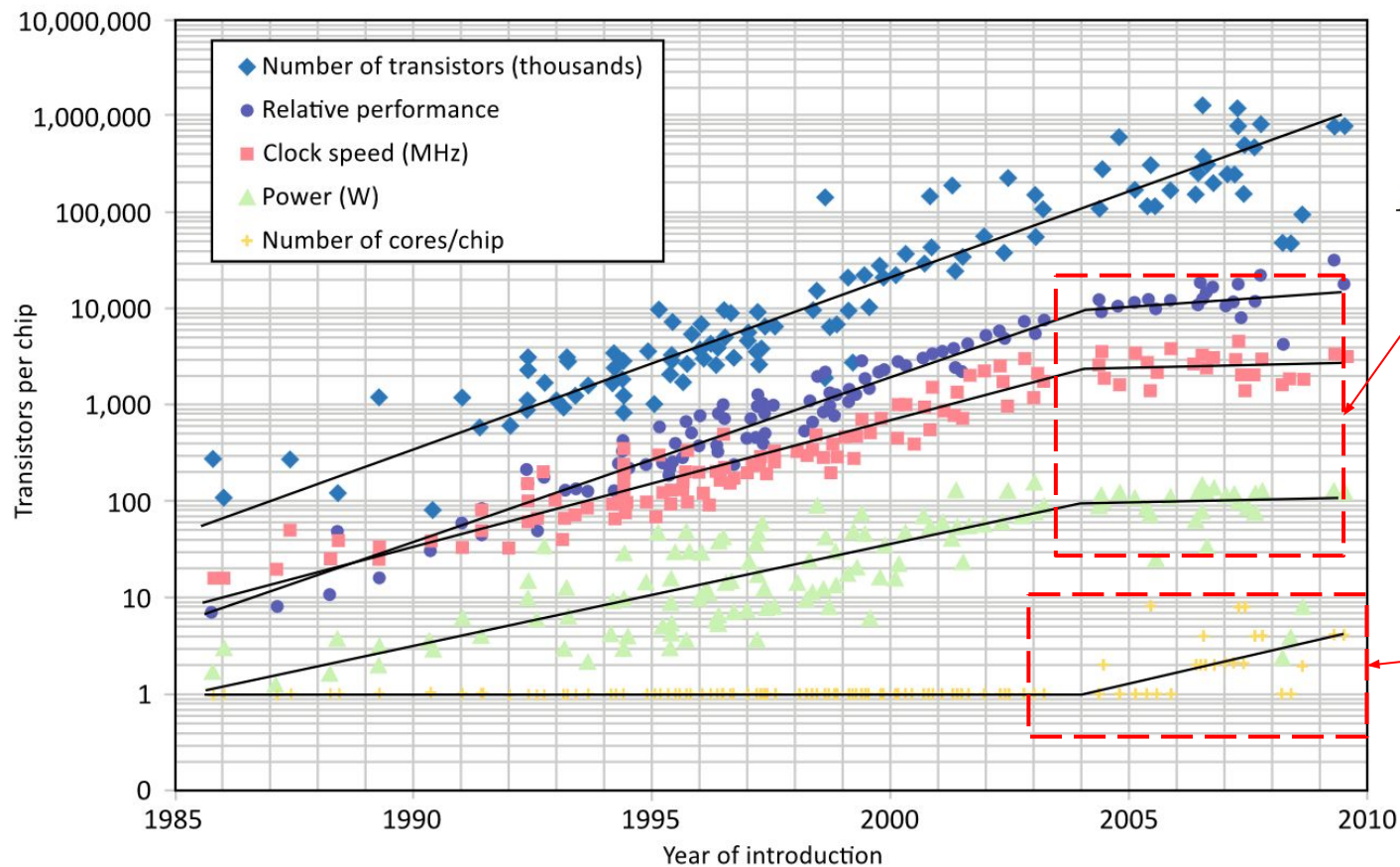
# key ideas:

- ❑ specializing silicon is important
- ❑ CPUs are *not enough* today
- ❑ programmability required
- ❑ transparency limited



?

?



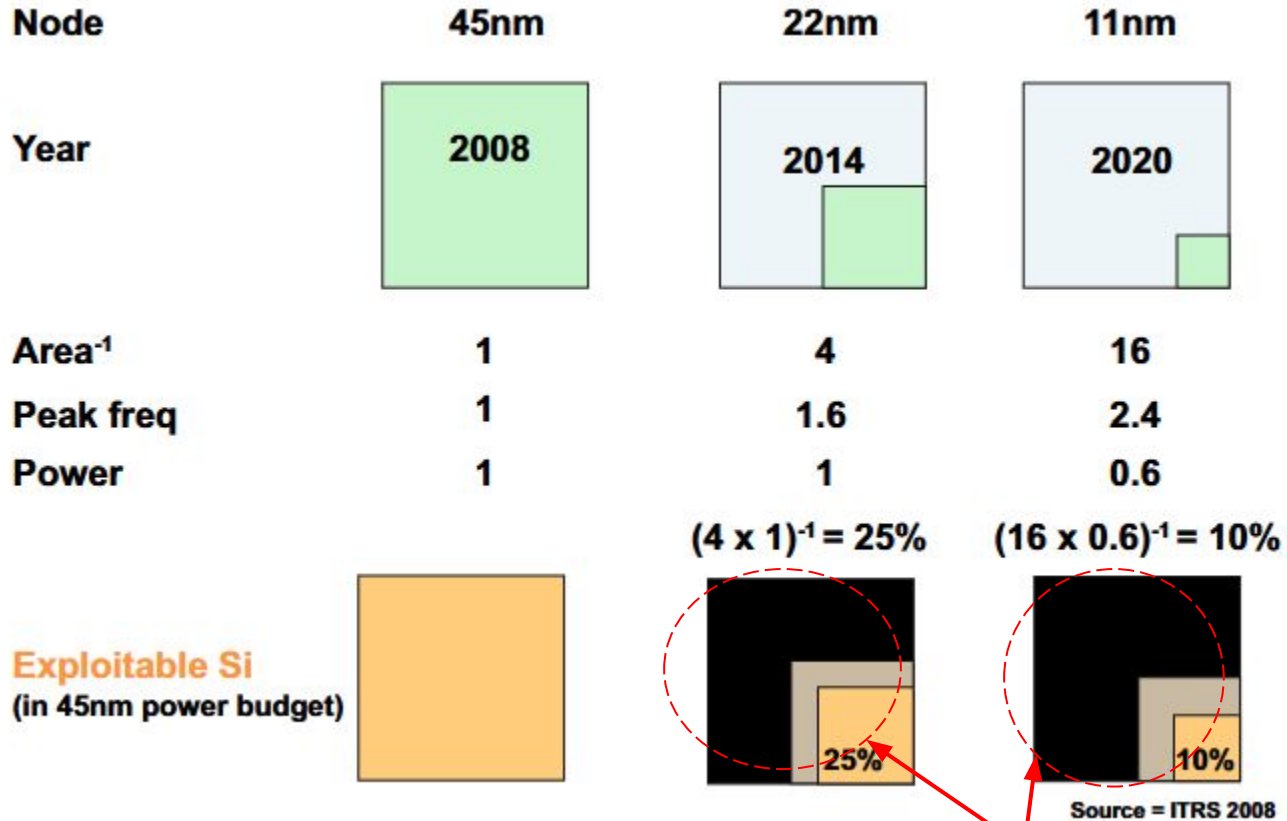
(“Dennard Scaling”)

$$P_{total} = \alpha C V_{dd}^2 f + I_{leak} V_{dd}$$

computer architects in the late 1990s had “free lunch” by scaling  
(1) voltage (2) frequency (etc.)

we moved to multi-core systems instead  
(more cores running at same  $V$  and  $f$ )

benefits of scaling were limited by the “*power wall*”



(2) complex hardware in CPUs could not be kept busy enough to justify its power and area cost.

(1) “dark silicon”  
(had to be switched off)

benefits of scaling were limited by the “*power wall*”

adding functionality to single core limited by “*utilization wall*”

Solution:  
Specialized domain-specific cores sitting alongside  
general-purpose core  
adding functionality to the

examples of “specialized” cores?



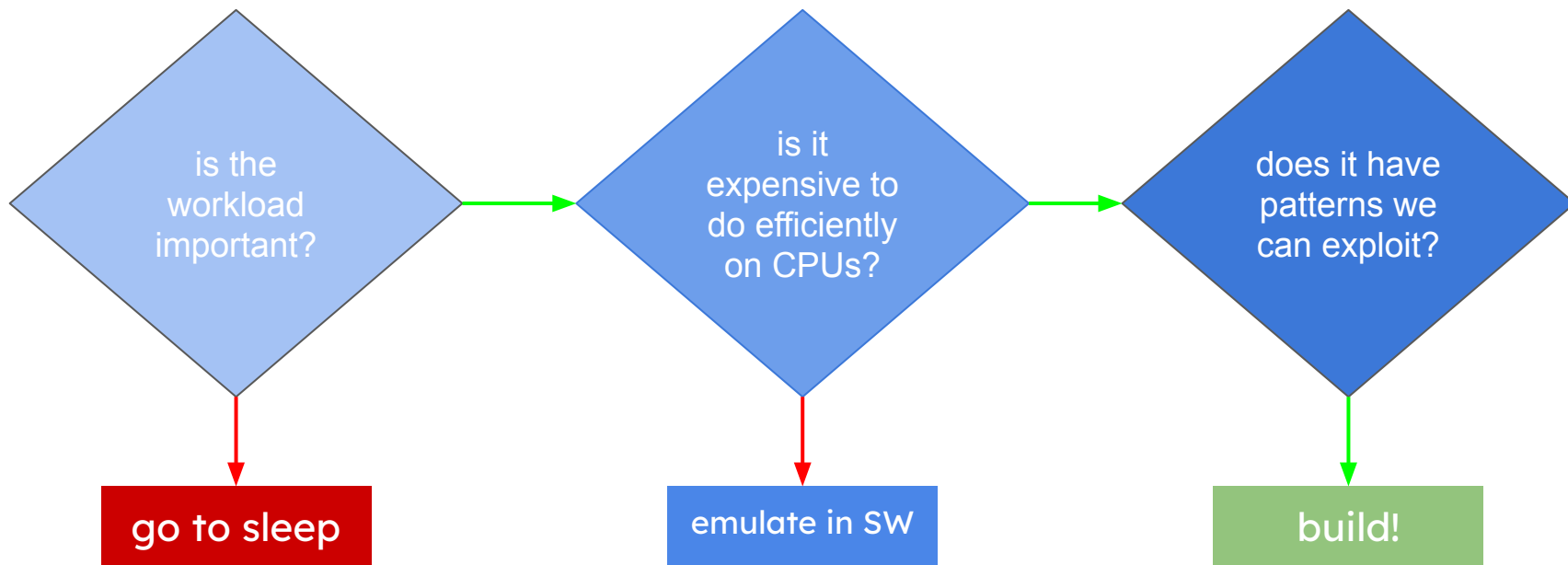
**TPU (Tensor Processing Unit)**



**GPU (Graphics Processing Unit)**

- + iPhone Neural Engine,
- + Pixel Mali GPU,
- + Snapdragon Hexagon NPU

# when do you build a specialized core?

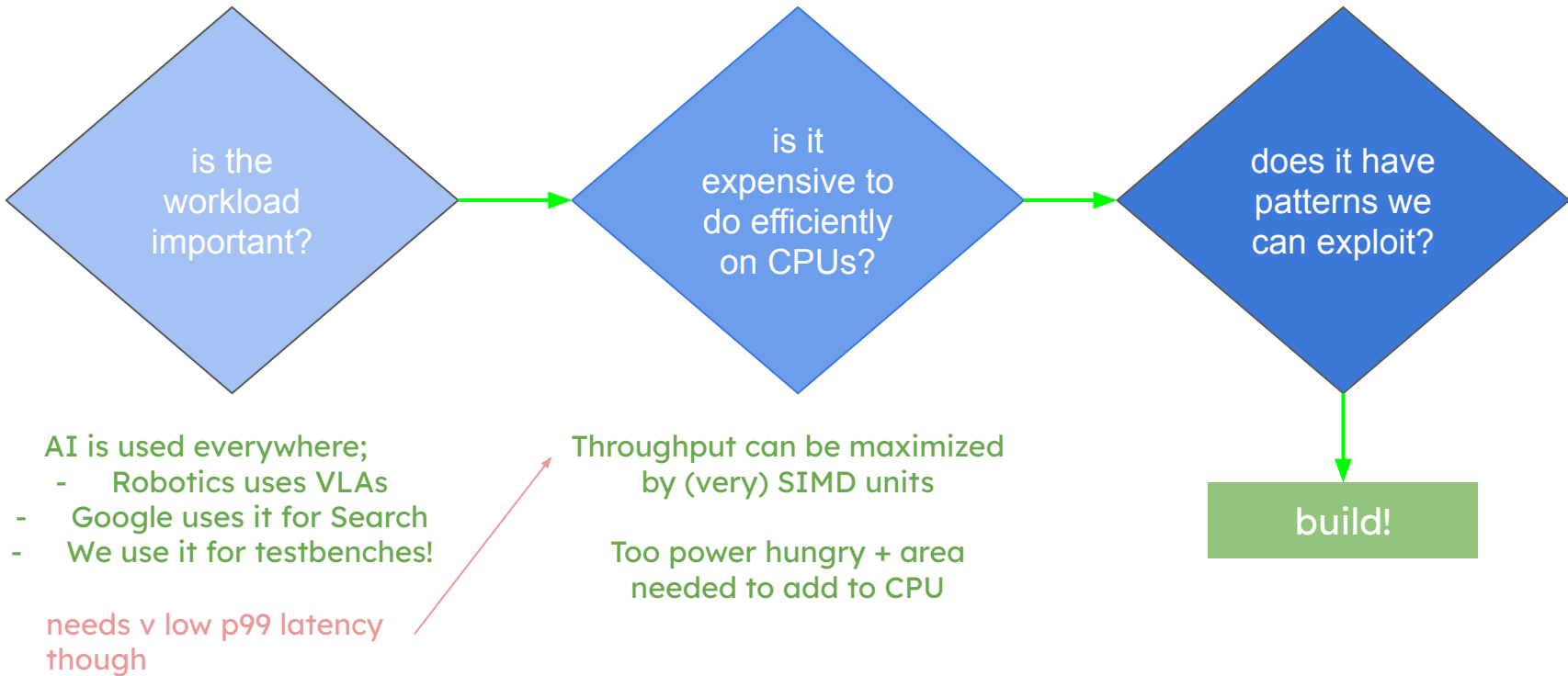


nerds build useless fun stuff;  
engineers don't

# why build TPUs (ML/AI)?

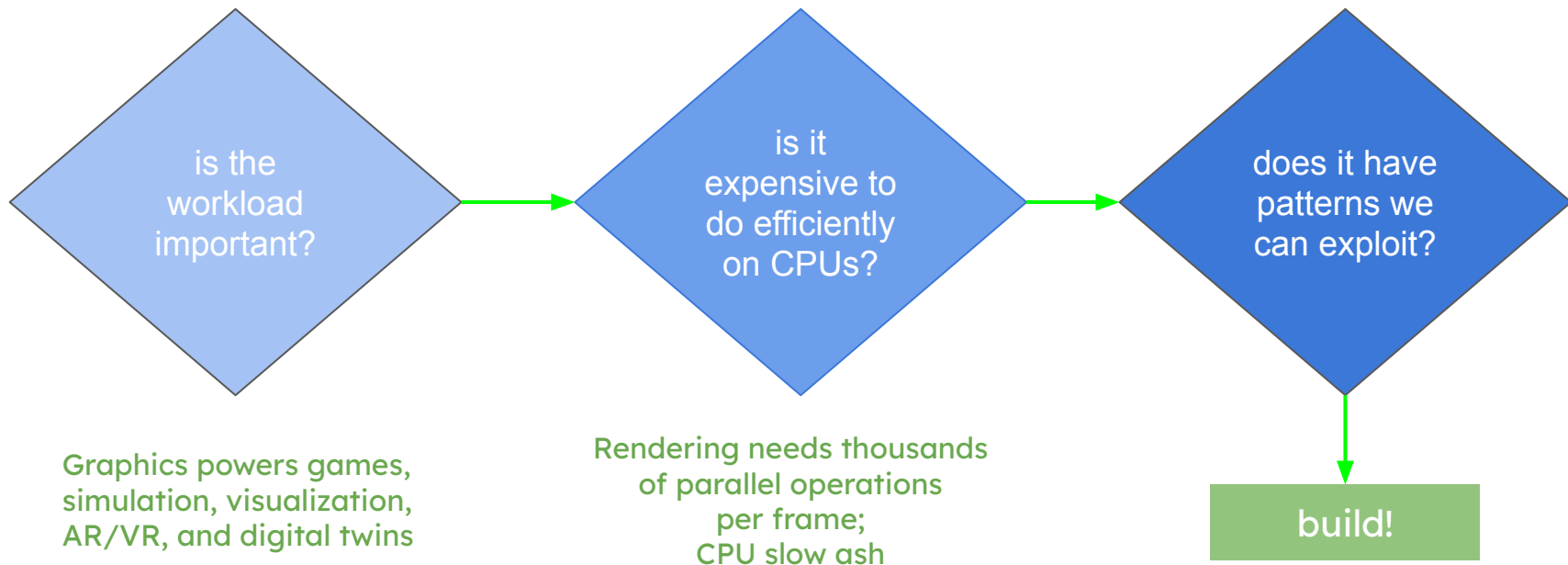
ML just FC/Conv layers;  
Those layers are repeated  
Tensor math;

Regular compute/mem  
patterns!



# why build GPUs (graphics)?

Vertices, pixels, shaders,  
Textures repeat;  
Independent math and  
streaming data;



# Nvidia & Google don't open source “secret-sauce” – why?

## VAGUELY Open Sourced!

programming model,  
compiler interface,  
performance guidance,  
high-level arch

**promote end-users**

**define ecosystem/moat**

**create the engineering  
trends shaping the  
market**

**competitive advantage**

**security**

**freedom to escape  
hardware  
backwards  
compatibility**

## NOT EXPOSED!

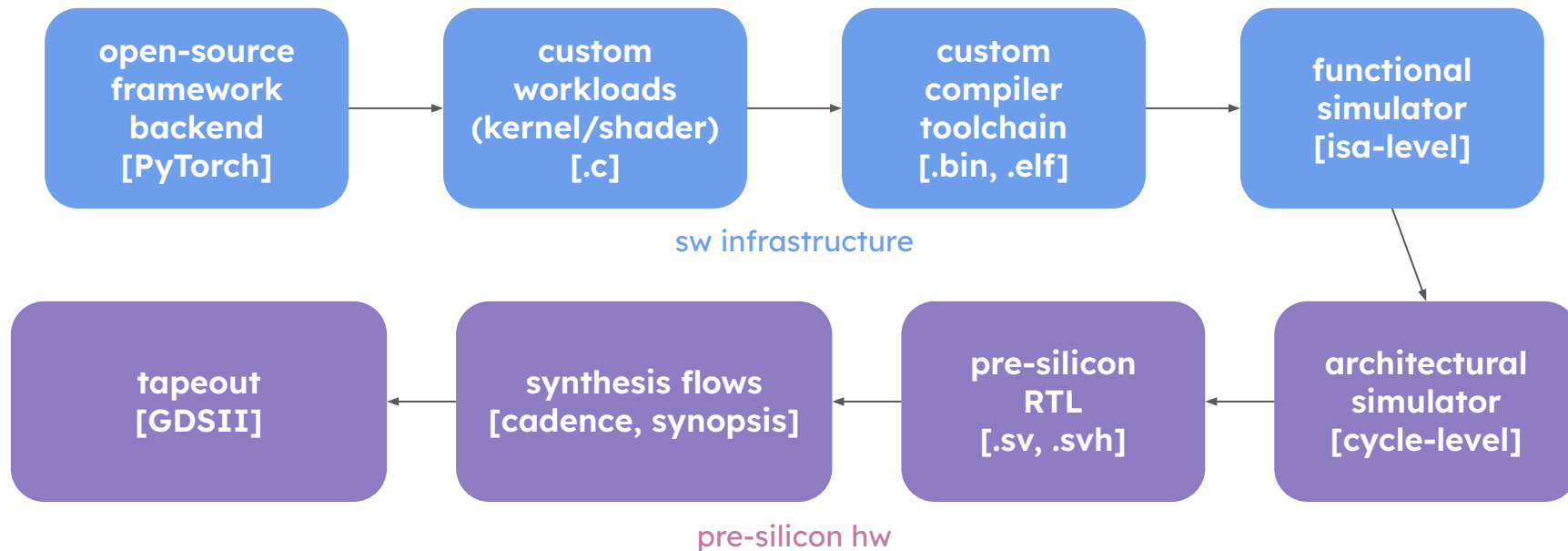
Exact core microarchitecture,  
Precise scheduler behavior,  
Exact cache replacement,  
Operand collector details,  
cache PnR information,  
etc.

# and so,

- ❑ industry trends showed potential for **DSAs**
- ❑ AI & Graphics are **perfect for custom silicon**
- ❑ Best industrial systems are *closed off*
- ❑ DSA startups are everywhere (Taalas!)
- ❑ academic tape-outs and FOSS toolchains **rare**

**that's where we come in!**

Build open-source custom *architectures* and *toolchains*;  
inspired by industrial systems, but designed for research



The goal is that **anyone** can understand the architecture, modify it, run workloads on it, and eventually help us tape it out!

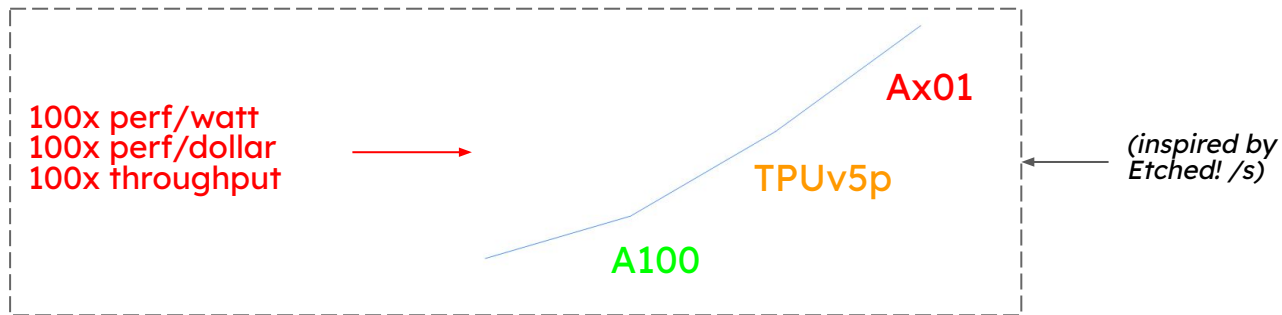
motivation

**history & growth**

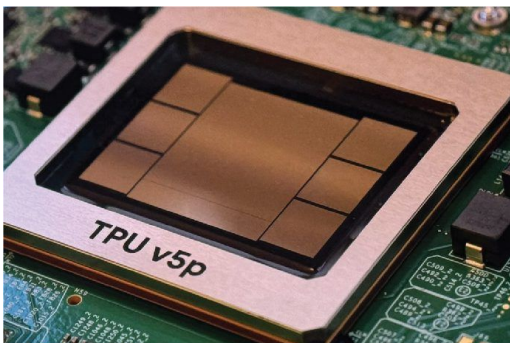
future plans

who should join?

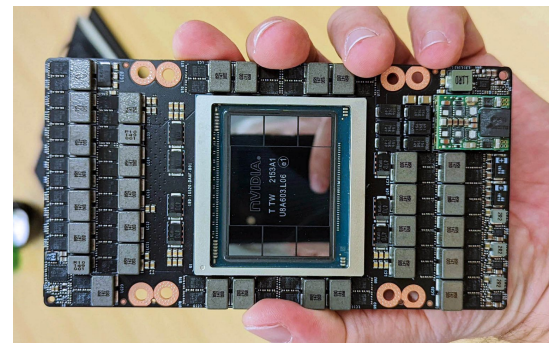
many many weeks ago,  
in the MSEE 180 boardroom,  
a new sub-team was formed!  
aptly named as AI Hardware  
Sooraj, Timmy and Malcolm were  
inspired by **Nvidia** and **Google** to  
build a custom ASIC 



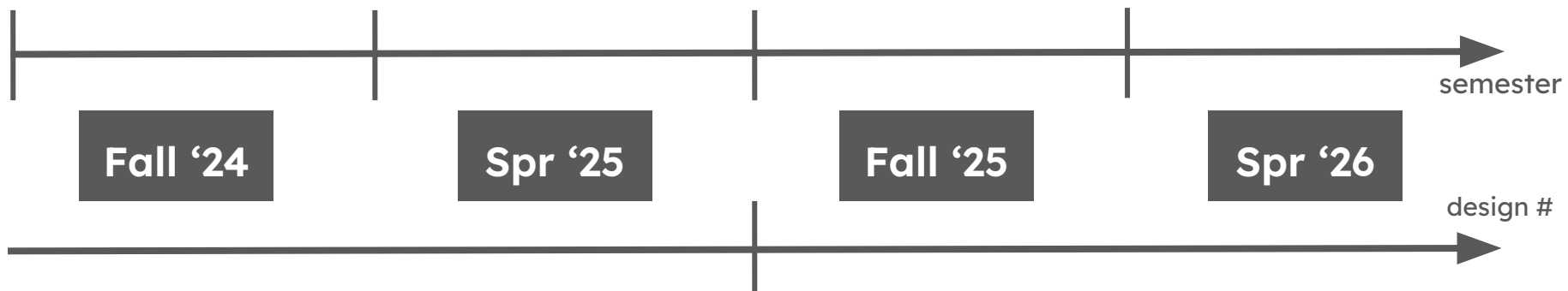
SoCET's Atalla Ax01!



Google's TPUv5p

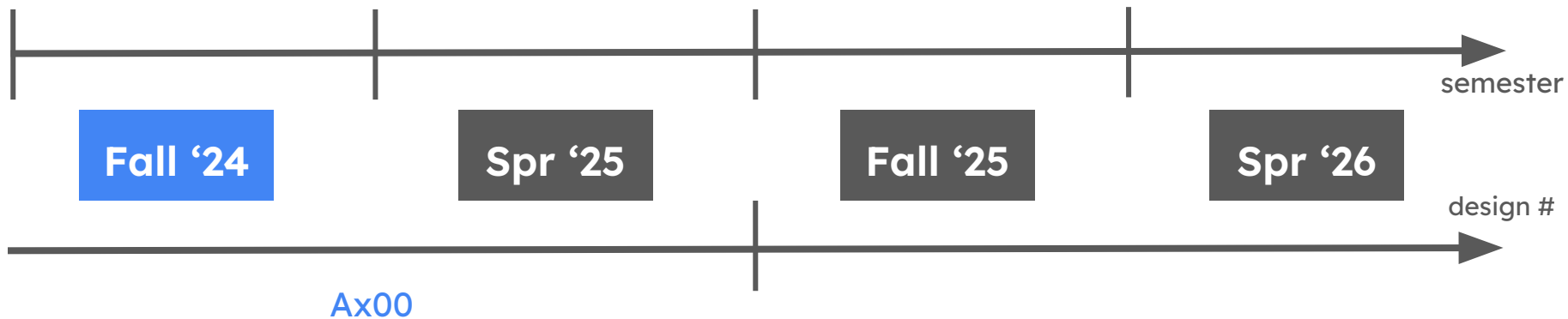


Nvidia's A100



Systolic Array, ISA,  
Memory, Scheduler

13 members  
4 teams



Systolic Array, ISA,  
Memory, Scheduler

13 members  
4 teams

23 members  
4 teams

semester

Fall '24

Spr '25

Fall '25

Spr '26

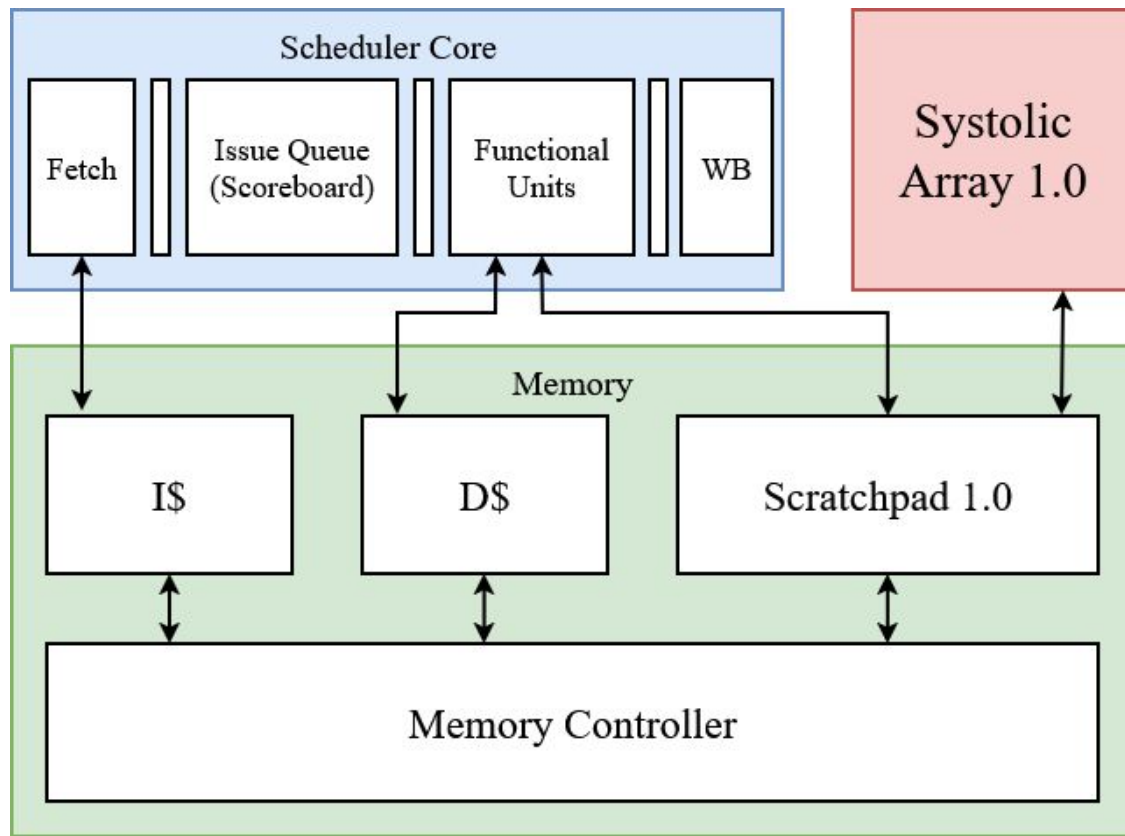
design #

Ax00

scoreboard  
O3

blocking  
caches

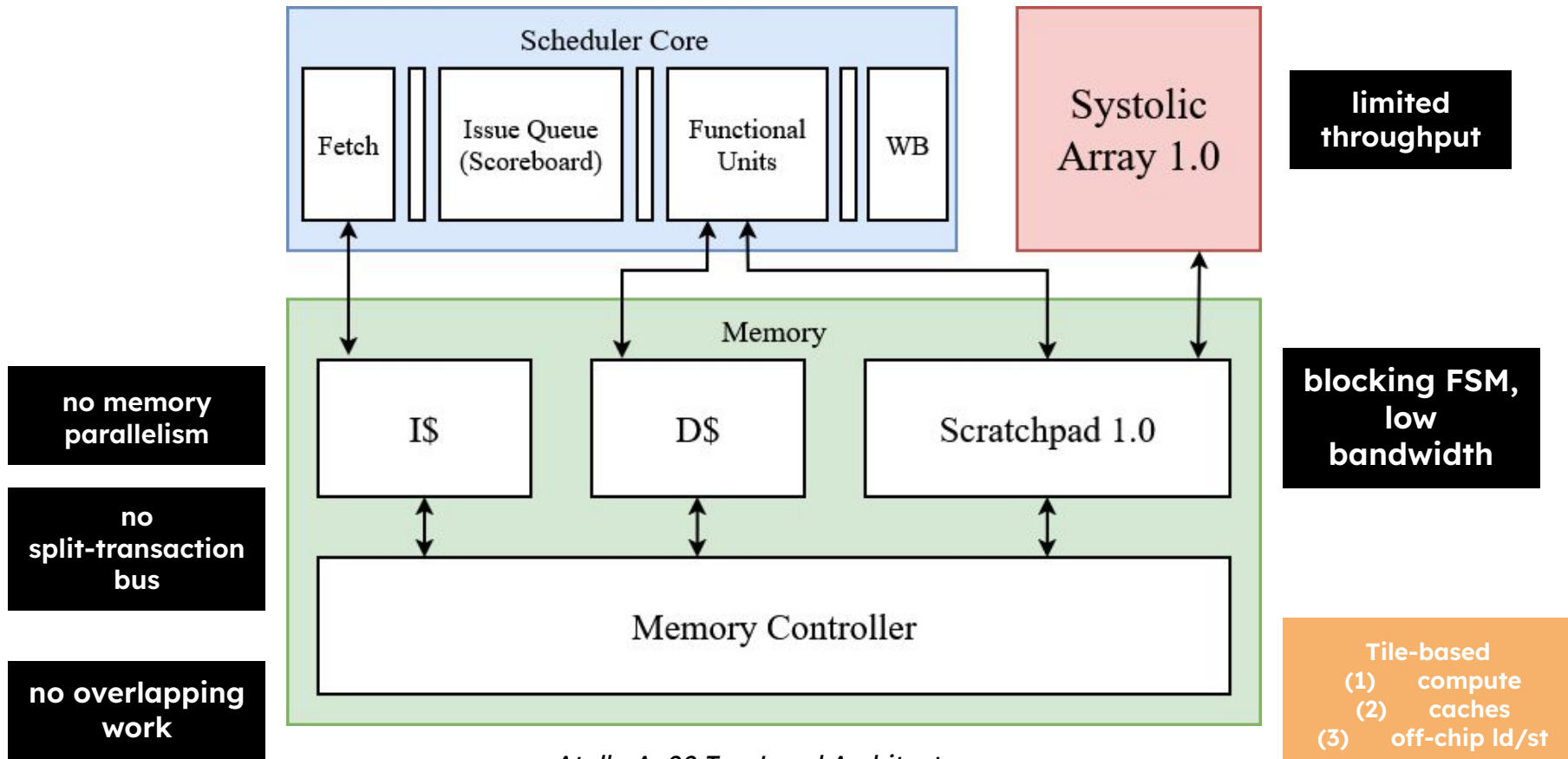
blocking  
ddr4-cntrl



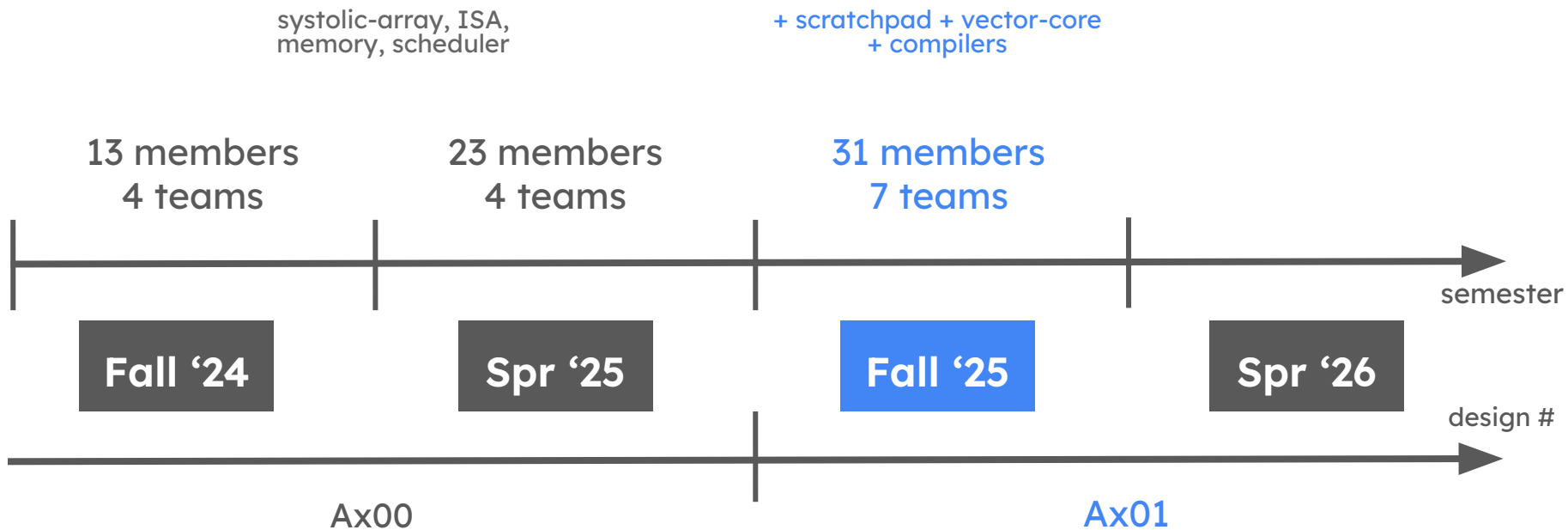
32x32 SA

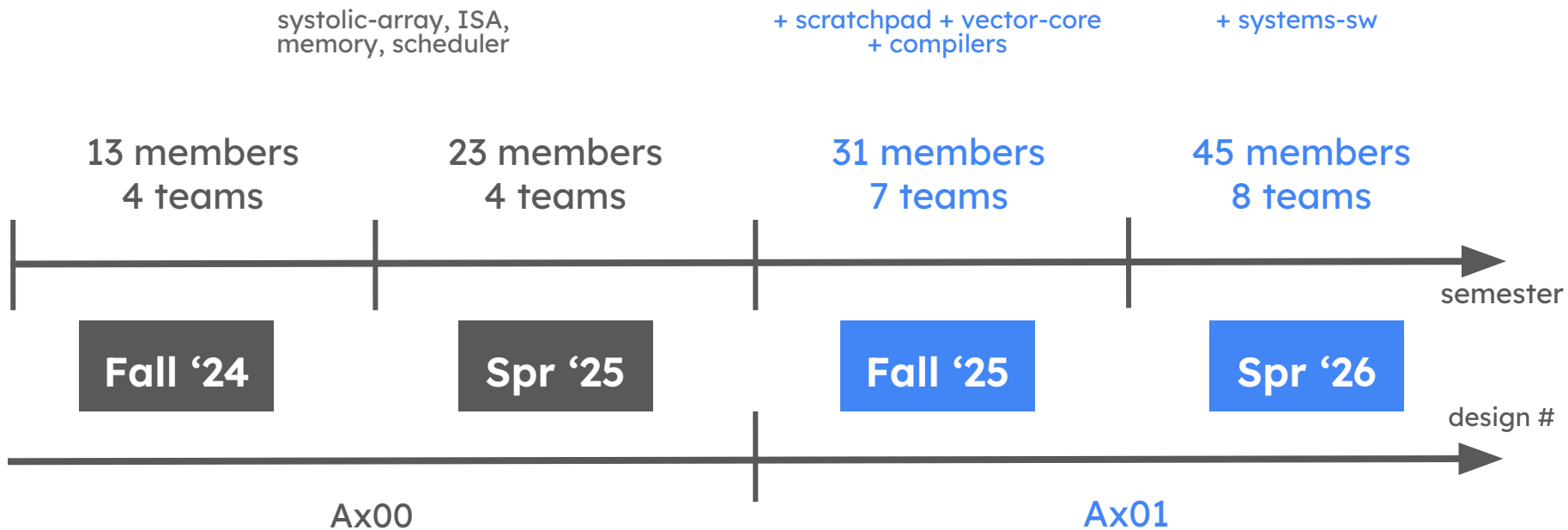
matrix  
registers

*Atalla Ax00 Top-Level Architecture*

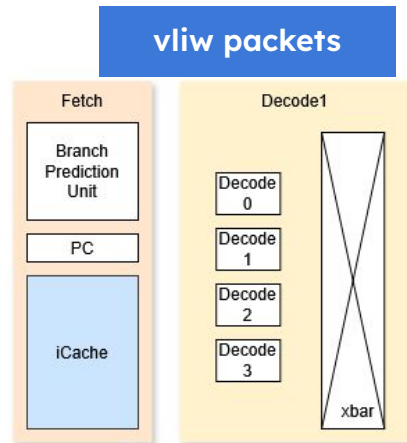


*Atalla Ax00 Top-Level Architecture*

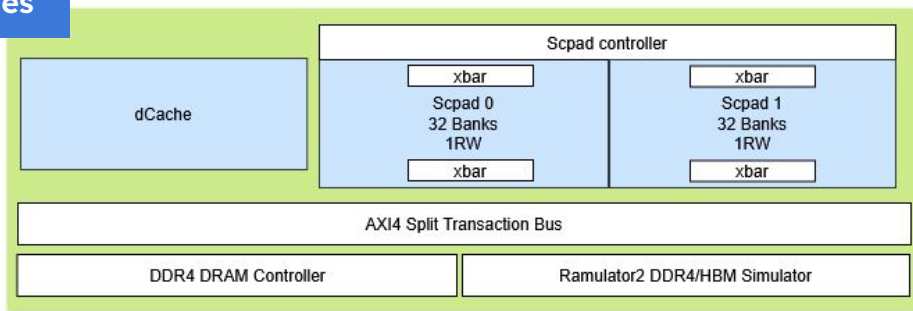








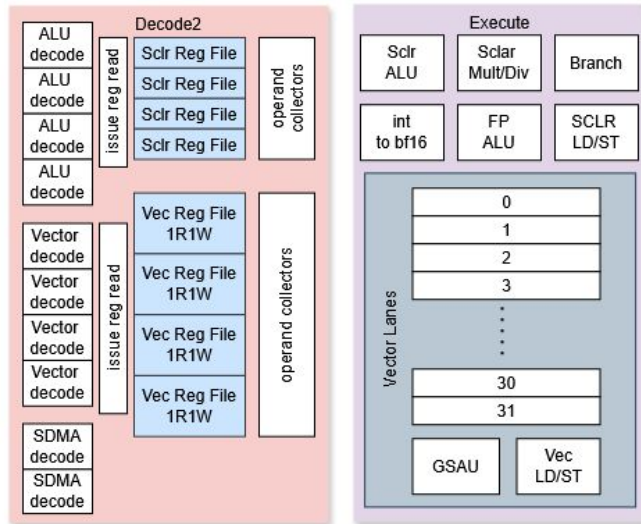
**non-blocking  
banked caches**



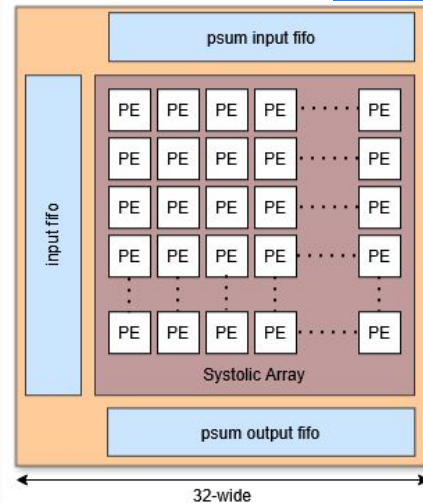
**plug-and-play HBM behaviour**

## Atalla Ax01 Top-Level Architecture

**vector-core!**



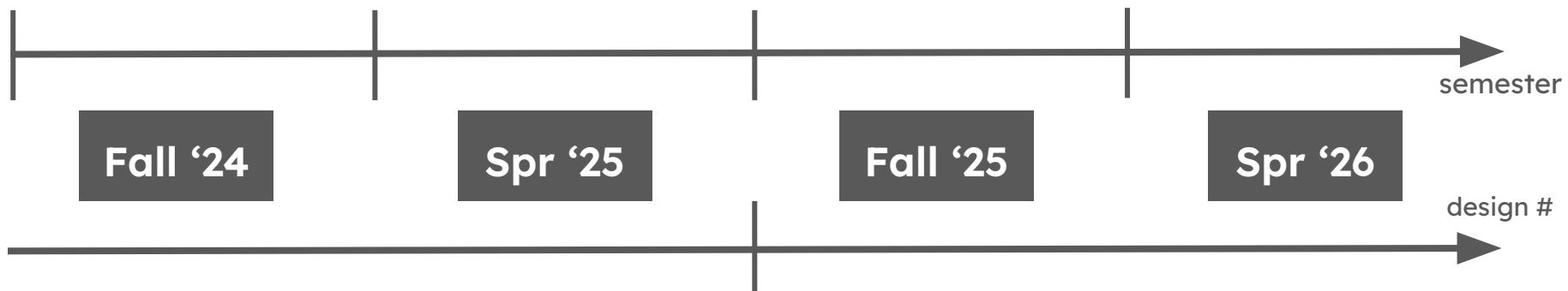
**32x32 SA**



**2 scpads,  
non-blocking access**

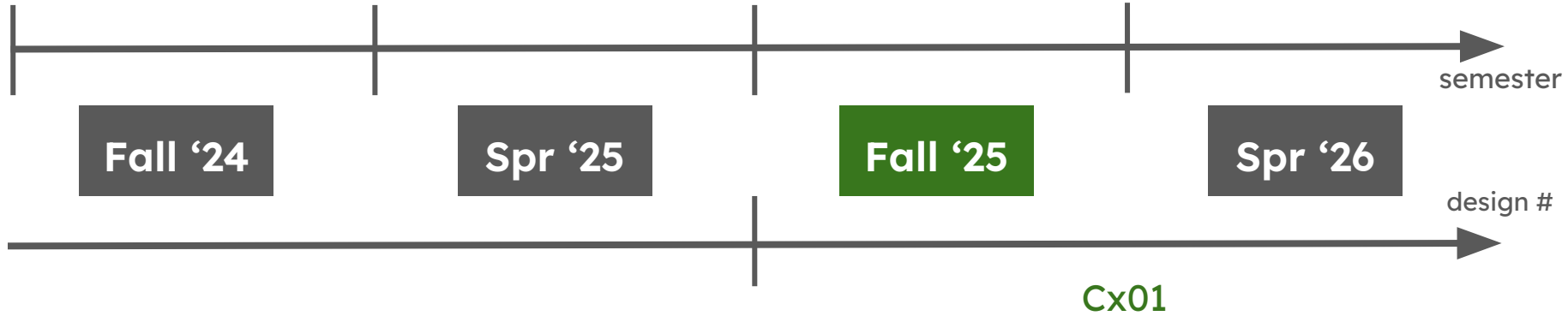
Vector-based  
(1) compute  
(2) caches

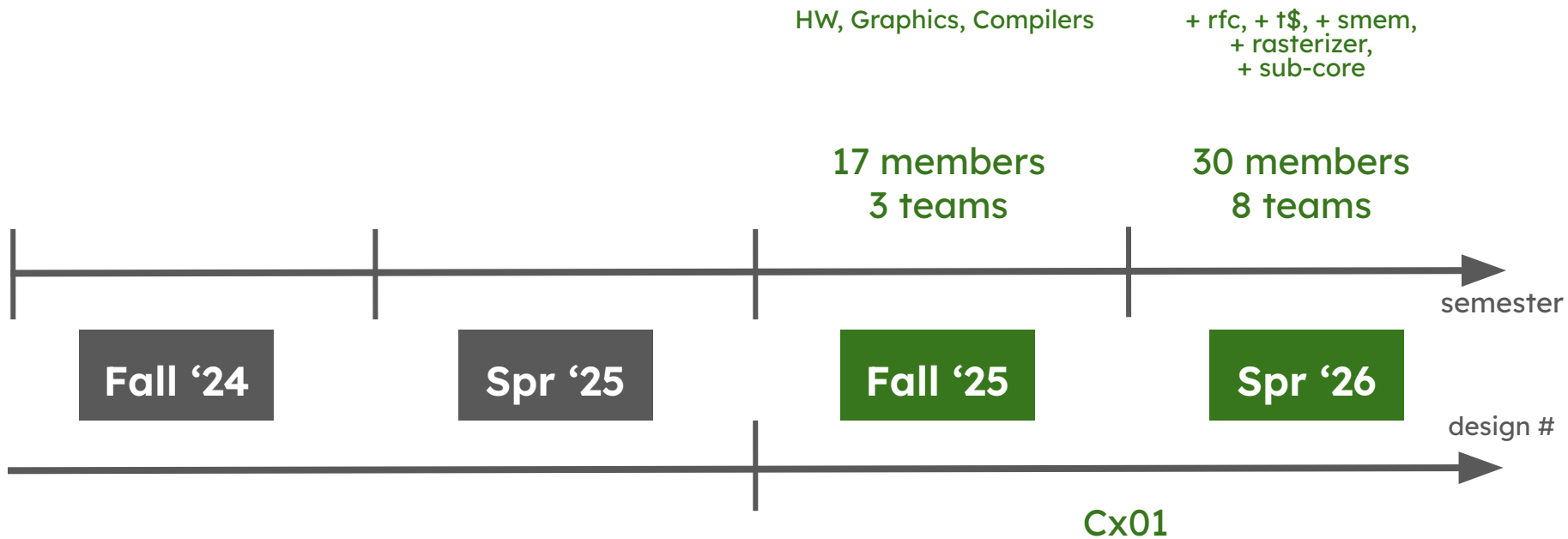
Tile-based  
(3) off-chip ld/st<sup>32</sup>



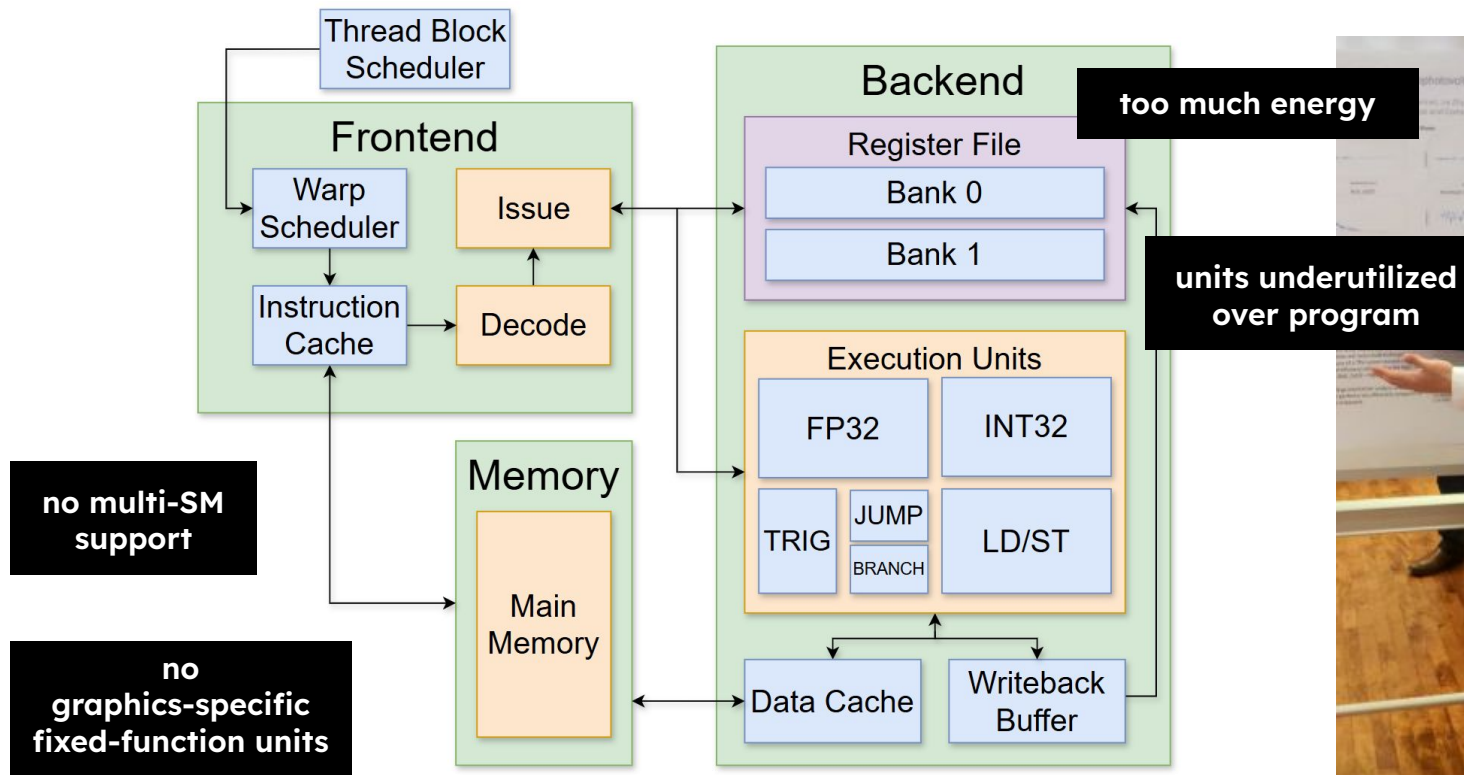
HW, Graphics, Compilers

17 members  
3 teams









Cardinal Cx01 Top-Level Architecture



motivation

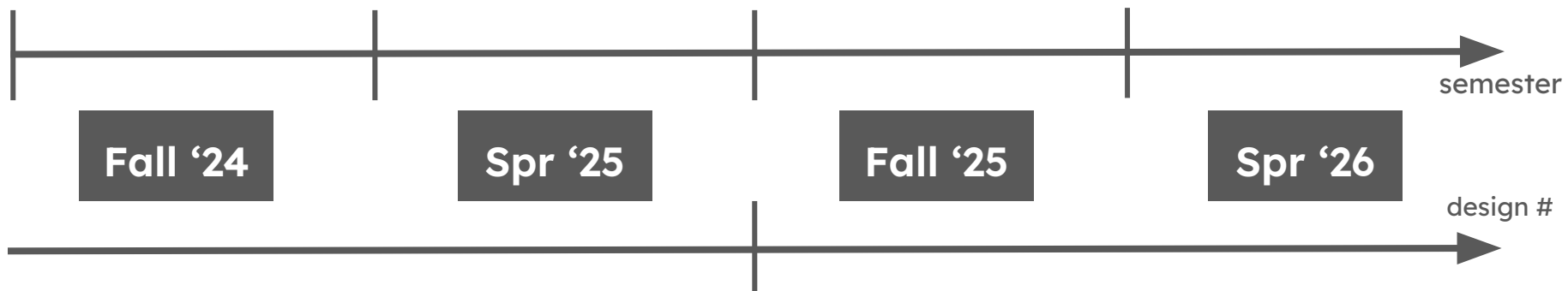
history & growth

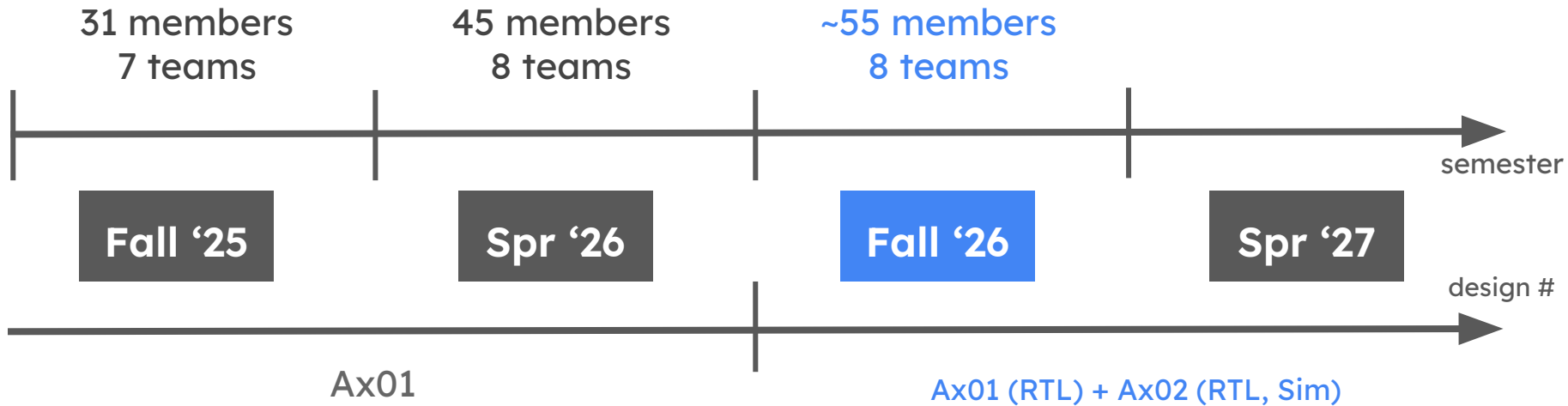
**future plans**

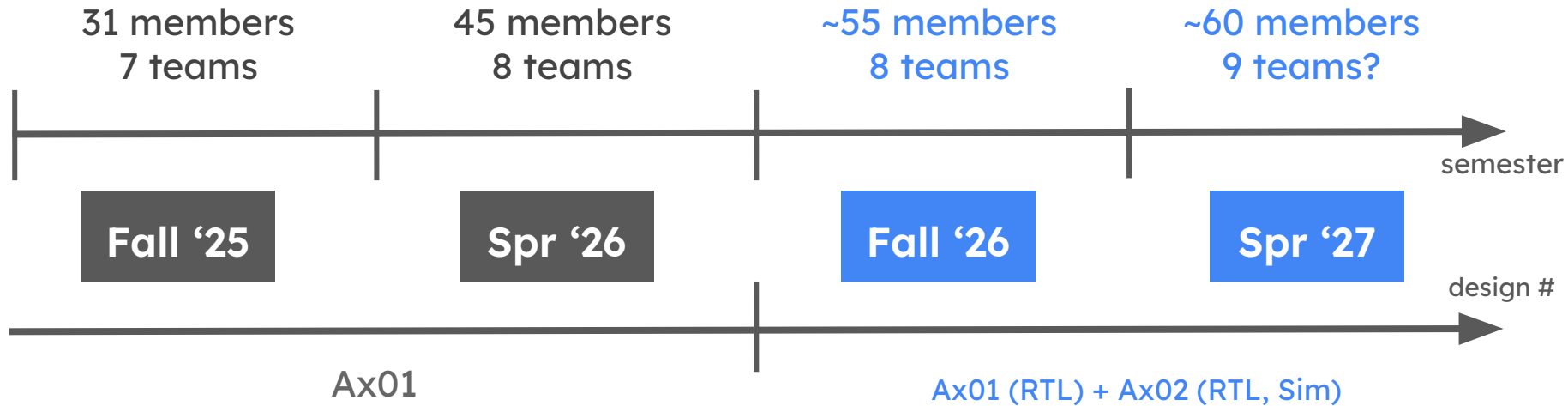
who should join?

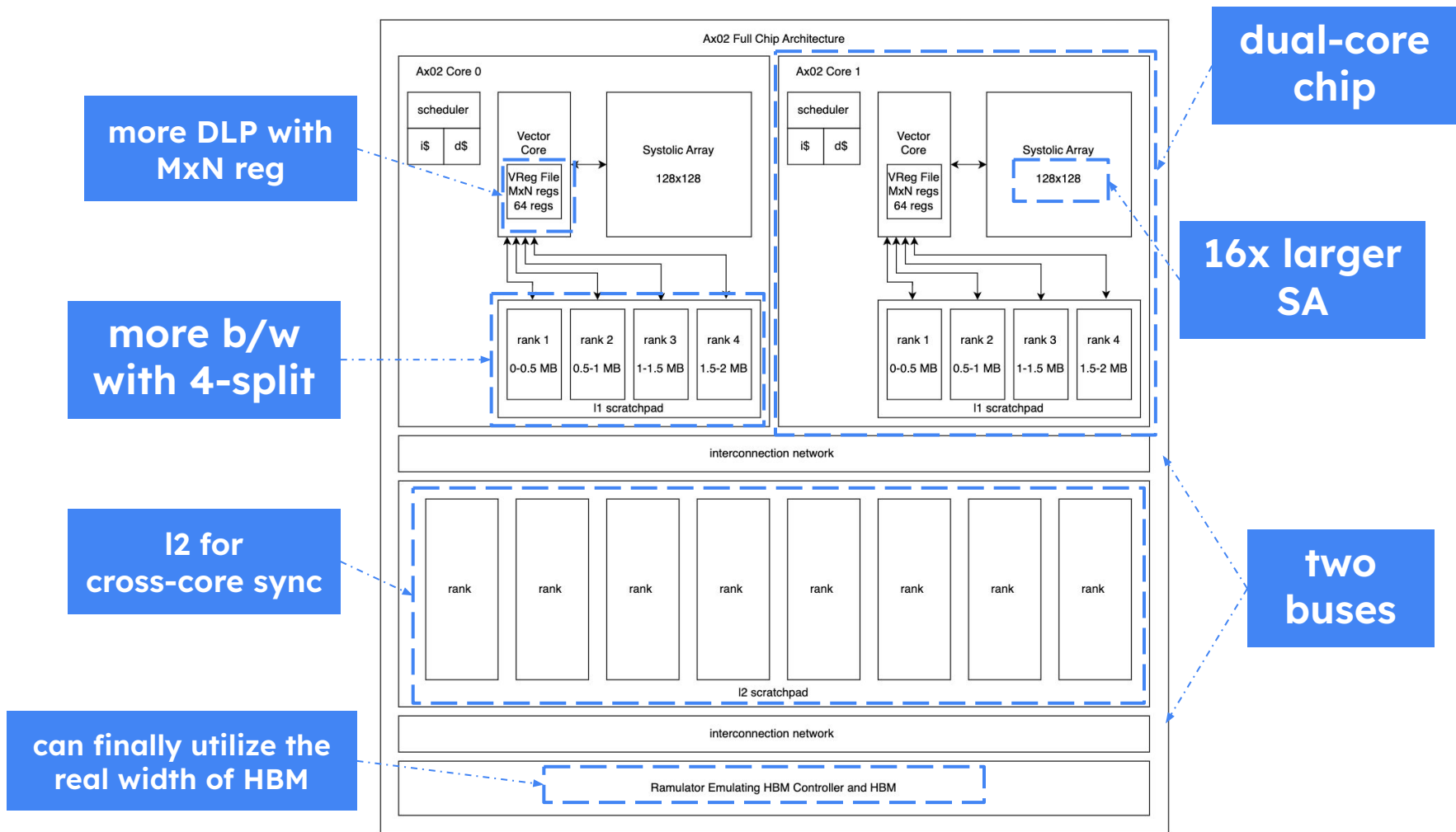
# **lessons learned:**

- ❑ architecture changes first in functional sim**
- ❑ RTL needs minimum 98% code coverage**
- ❑ every arch edit needs compiler and sw efforts (hw/sw co-design)**
- ❑ physical design imperative in design cycle**
- ❑ communication skills most important**









Atalla Ax02 Top-Level Architecture

- hierarchical tiling strategies
- better lowering/codegen
- kernel fusion (etc.)
- better memory management

- more custom kernel support
- identify more models and potential arch changes
- improved hw utilization

- optimizations in the backend
- best unroll-factor? (heuristics)
- better register-usage
- support for new arch changes

- support new arch changes
- better metric support
- multi-core behaviour?

open-source  
framework  
backend  
[PyTorch]

custom  
workloads  
(kernel/shader)  
[.c]

custom  
compiler  
toolchain  
[.bin, .elf]

functional  
simulator  
[isa-level]

pre-silicon hw

sw infrastructure

FPGA emulation  
[HAPS]

synthesis flows  
[cadence, synopsys]

pre-silicon  
RTL  
[.sv, .svh]

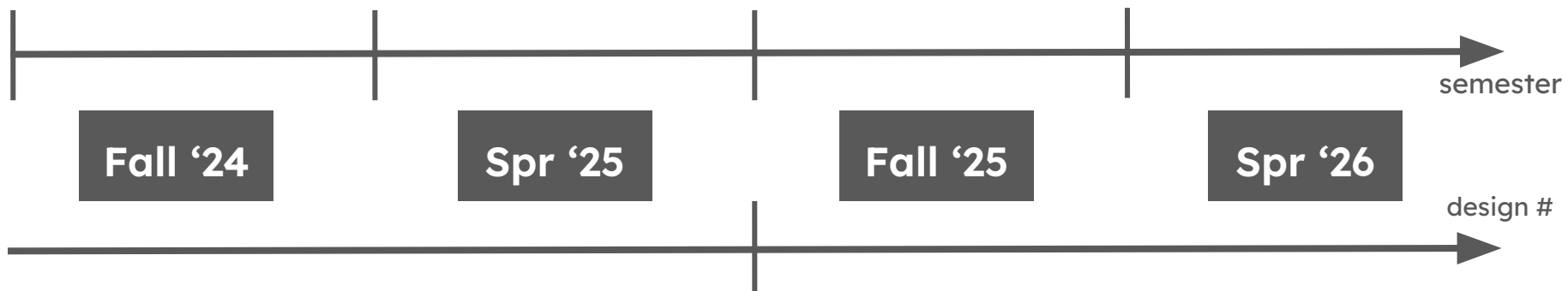
architectural  
simulator  
[cycle-level]

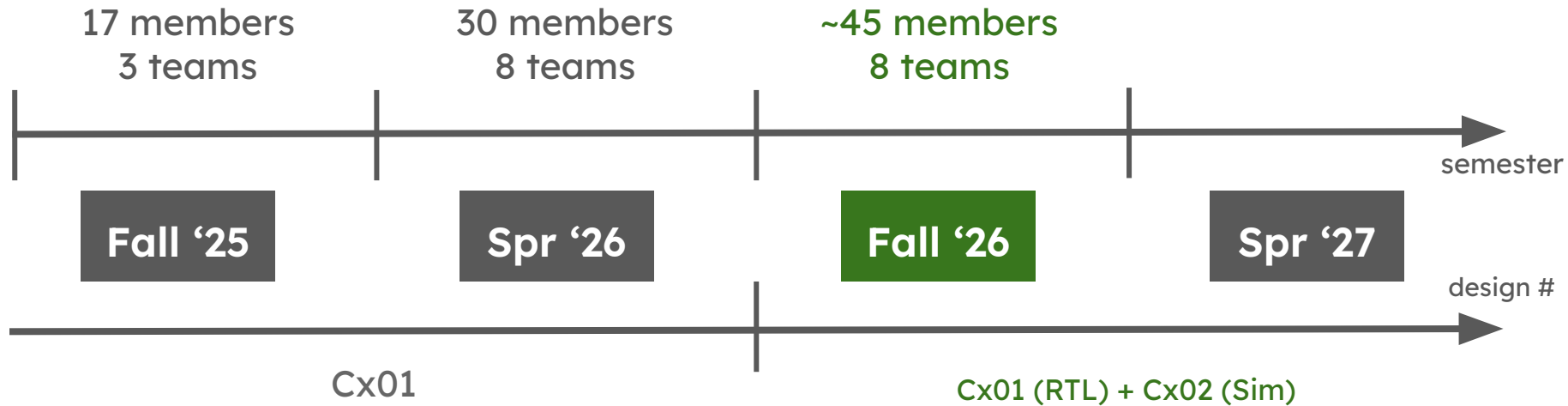
- port entire onto the HAPS
- run entire model flow
- accurately simulate memory hierarchy

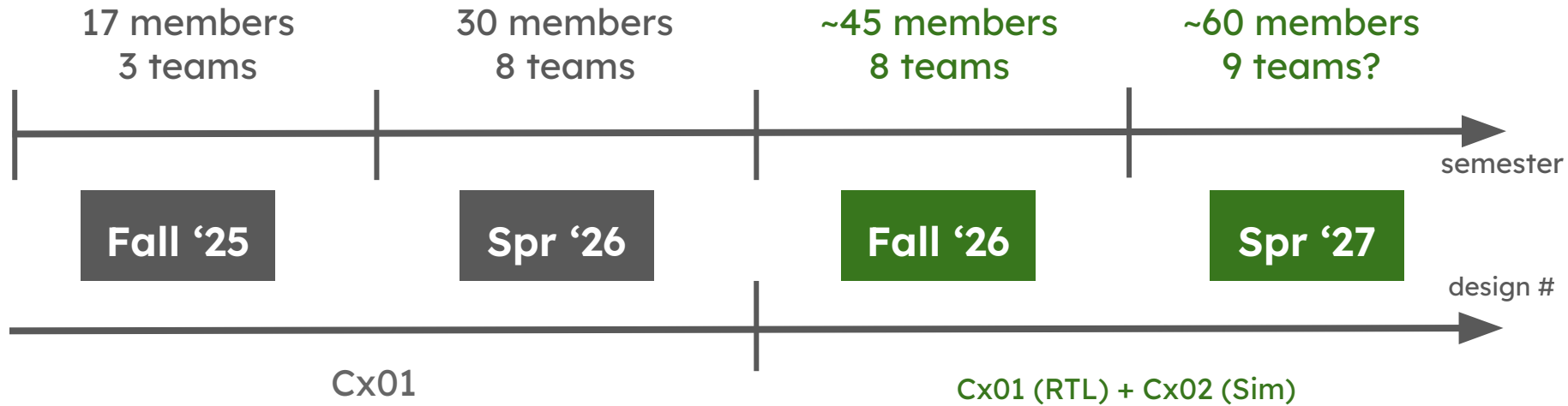
- improved cadence flows
- port to synopsys tools

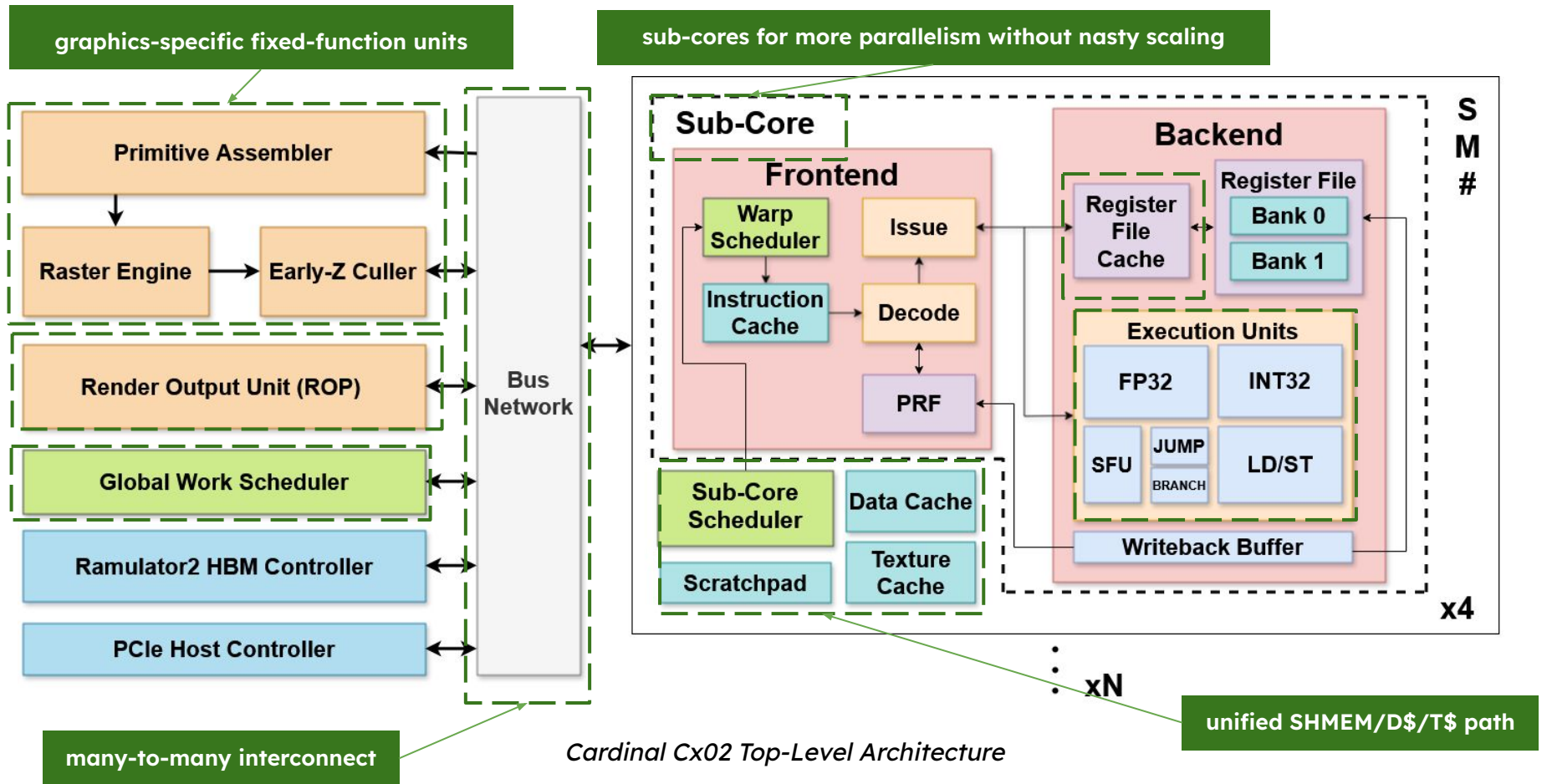
- complete RTL verification (lot!)
- better parameterization
- Ax02 arch changes implementations

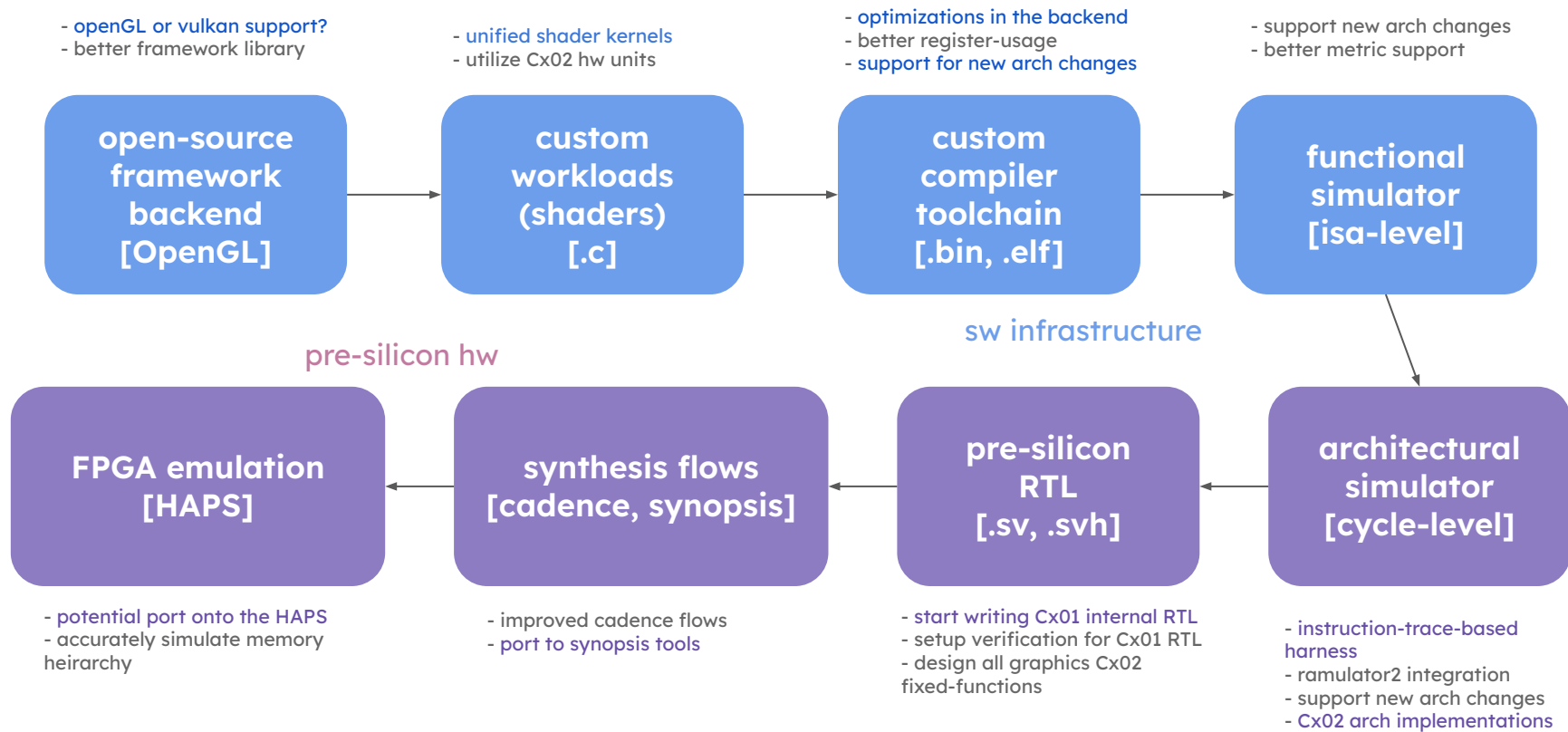
- instruction-trace-based harness
- ramulator2 integration
- support new arch changes
- better OOP composability
- Ax03 arch changes











motivation

history & growth

future plans

**who should join?**

# what you will gain!

## communication skills

- interviews/internships/fte
- teamwork
- how to present

## portfolio-building

- real results to showcase
- shows awareness of latest trends
- improves credibility

## technical skills

- RTL design + verification
- computer architecture
- software infrastructure

## mentorship

- how to think
- how to identify opportunities
- create your own projects

apply to join, if you are:

- ❑ excited about *comp.arch* & *software*
- ❑ willing to work across disciplines
- ❑ want to understand chip design flow
- ❑ capable of being consistent (!)
- ❑ good at communicating with others
- ❑ quick to onboard

**DON'T** apply to join, if you are:

- ❑ **low** motivation
- ❑ **not** willing to learn-by-doing
- ❑ have **busy** semesters
- ❑ **not** doing senior-design with us
- ❑ **not** wanting to take ownership

# who are we looking for? (any)

- ❑ good at **compilers** (ppci, llvm exp.)
- ❑ wonderful at **RTL & verification**
- ❑ capable in graphics/ml workloads
- ❑ proficient in python/**c++** (for sim)
- ❑ background in **physical design**
- ❑ **fpga emulation** (HAPS pref.) experience
- ❑ pytorch or openGL/vulkan background
- ❑ coursework in compiler/arch/DL/graphics

**questions?**