

# Worst-case memory latency and distributed DRAM banks across chips

## Towards tighter bounds using Deterministic Network Calculus

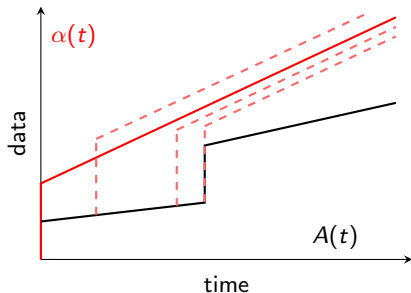
Edoardo Loni<sup>1</sup>   Raffaele Zippo<sup>1</sup>   Giovanni Stea<sup>1</sup>   Matteo Andreozzi<sup>2</sup>

<sup>1</sup>University of Pisa

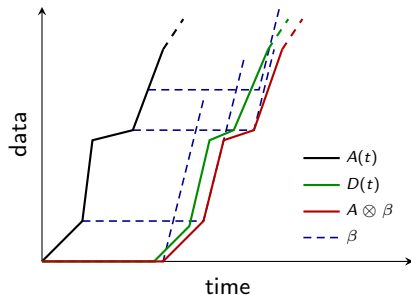
<sup>2</sup>ARM

Lund, RTSOPS 2026, 07-07-2026

# Context: Deterministic Network Calculus



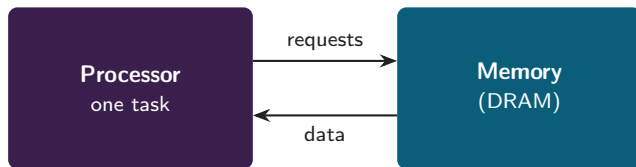
**Arrival curve**  $\alpha(t)$ : max traffic constraint;  
 $A(t) - A(s) \leq \alpha(t - s)$ .



**Service curve**  $\beta(t)$ : min service guaranteed;  
 $D \geq A \otimes \beta$ .

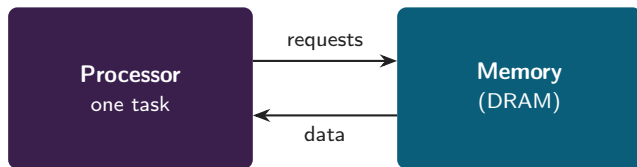
From  $\alpha$  and  $\beta$  we derive worst-case latency (WCD) and worst-case backlog (WCB) bounds.

# A simple model to start with



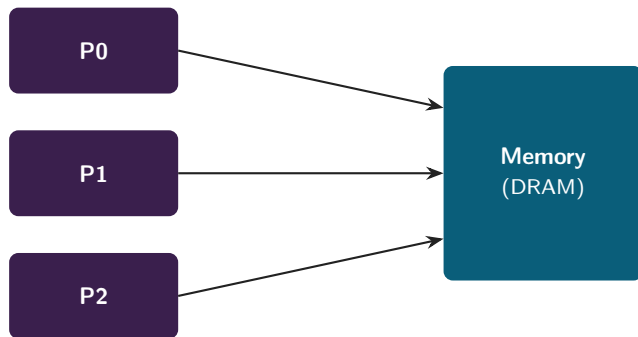
A simple computing system: one processor, running one task, that accesses the memory.

# A simple model to start with



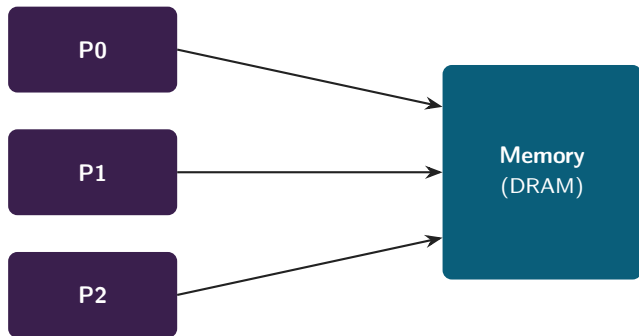
A simple computing system: one processor, running one task, that accesses the memory.  
Only need to derive  $\alpha$  for the requests and  $\beta$  for the memory.

## More processors, one memory



Concurrent processes compete for the same service.  
Can be addressed using *residual service curves*,  $\beta_2 = [\beta - (\alpha_0 + \alpha_1)]_{\uparrow}^+$ .

## More processors, one memory

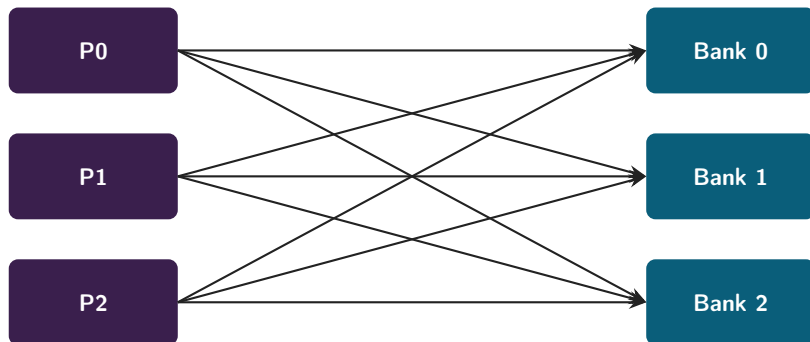


Concurrent processes compete for the same service.

Can be addressed using *residual service curves*,  $\beta_2 = [\beta - (\alpha_0 + \alpha_1)]_{\uparrow}^+$ .

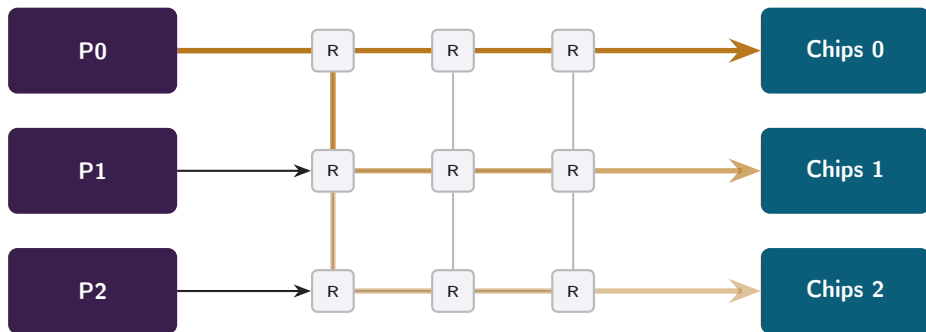
An **obvious bottleneck**: DRAM is multiple chips, banks etc. accessed in parallel.

## Requests spread across DRAM banks ...



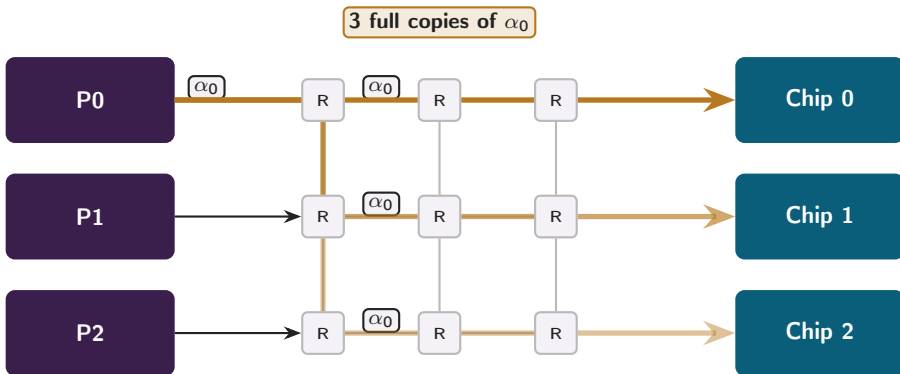
Each processor's traffic spreads to different banks.  
This leads to more parallelism and less contention **on average**.

...located on different chips on a Network-on-Chip



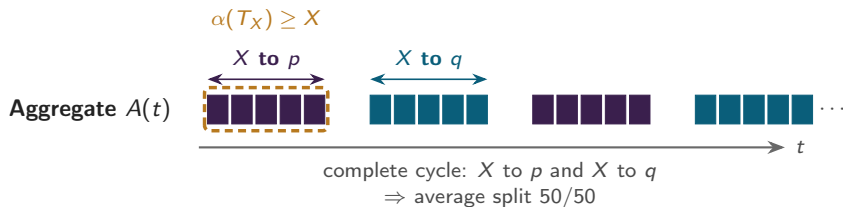
Traffic contention happens also on the NoC, and with other traffic going through it.  
Tight traffic modelling is critical to bound the interconnect latencies.

## State of the art: pessimistic splitting

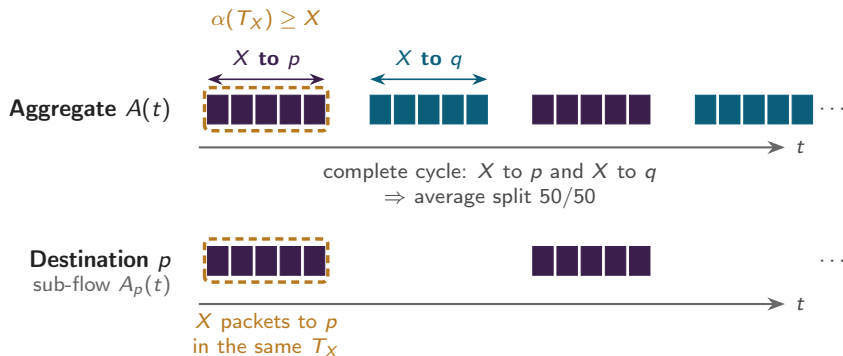


SOTA: **replicated in full** each processor's aggregate  $\alpha^i$  to every bank;  
pessimism compounds at every split point.

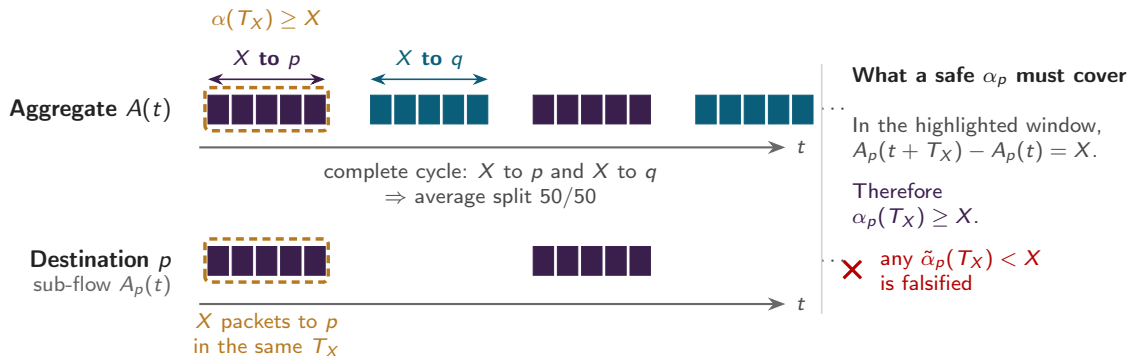
# State of the art: pessimistic splitting



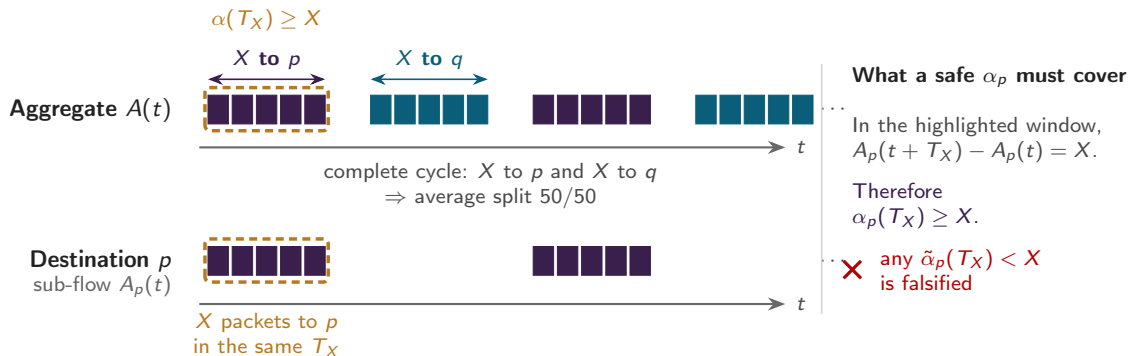
# State of the art: pessimistic splitting



# State of the art: pessimistic splitting

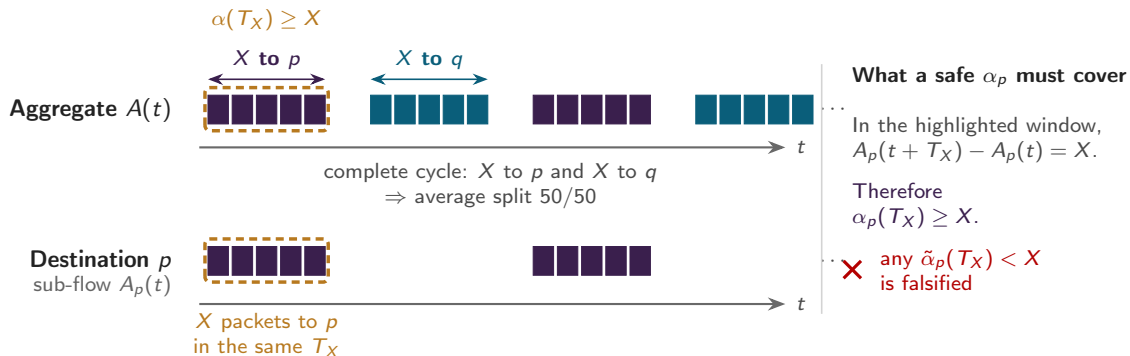


# State of the art: pessimistic splitting



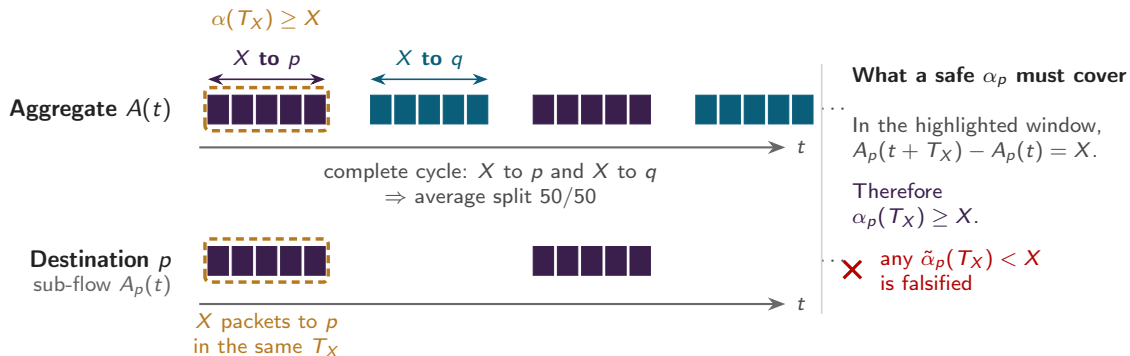
- Knowing the average ratio is 50/50 is **not enough**:  $\forall X$ , trace  $p^X q^X p^X \dots$  is admissible.

# State of the art: pessimistic splitting



- ▶ Knowing the average ratio is 50/50 is **not enough**:  $\forall X$ , trace  $p^X q^X p^X \dots$  is admissible.
- ▶ Since  $X$  is arbitrary, then  $\alpha_p$  must allow every burst allowed by  $\alpha$ .

# State of the art: pessimistic splitting



- ▶ Knowing the average ratio is 50/50 is **not enough**:  $\forall X$ , trace  $p^X q^X p^X \dots$  is admissible.
- ▶ Since  $X$  is arbitrary, then  $\alpha_p$  must allow every burst allowed by  $\alpha$ .
- ▶ Therefore the only safe  $\alpha_p$  is  $\alpha$ .

# A bottom-up approach

The problem is practical, but the underlying issue is theoretical.  
So we work on three levels

1. Which **hypotheses** does the formal model need to split an aggregate arrival curve?

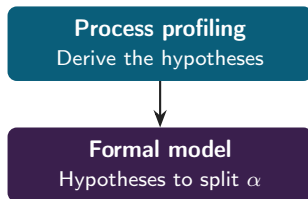
**Formal model**

Hypotheses to split  $\alpha$

# A bottom-up approach

The problem is practical, but the underlying issue is theoretical.  
So we work on three levels

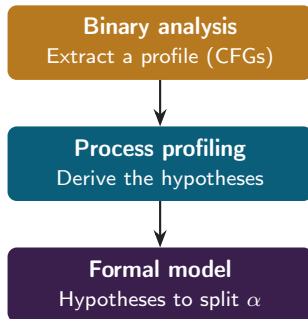
1. Which **hypotheses** does the formal model need to split an aggregate arrival curve?
2. Which **process profiling** is needed to derive these hypotheses?



# A bottom-up approach

The problem is practical, but the underlying issue is theoretical.  
So we work on three levels

1. Which **hypotheses** does the formal model need to split an aggregate arrival curve?
2. Which **process profiling** is needed to derive these hypotheses?
3. How do we **analyze** a given binary to extract such profile?



## Formal model - key idea

- ▶ Counter-example is an adversarial access pattern:  
arbitrary-sized burst in the short term hiding the long-term benefits

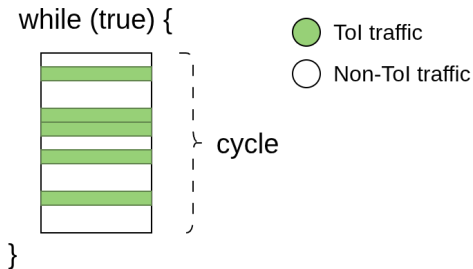
# Formal model - key idea

- ▶ Counter-example is an adversarial access pattern:  
arbitrary-sized burst in the short term hiding the long-term benefits
- ▶ To improve WC analysis, we need to make that pattern unfeasible

# Formal model - key idea

- ▶ Counter-example is an adversarial access pattern:  
arbitrary-sized burst in the short term hiding the long-term benefits
- ▶ To improve WC analysis, we need to make that pattern unfeasible
- ▶ We do so by adding constraints related to the **program loop**

# The cycle model

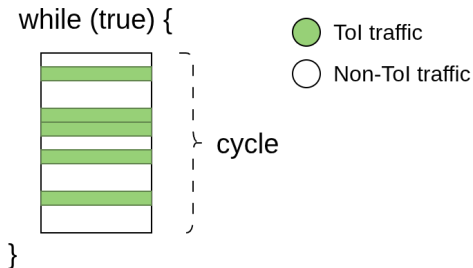


One cycle = traffic issued in one task iteration.

Tol stands for *Type of Interest*.

- Consider a (pseudo-)cyclic task: each iteration issues roughly the same *type* of traffic, i.e. predictable data structure accesses during each phase.

# The cycle model

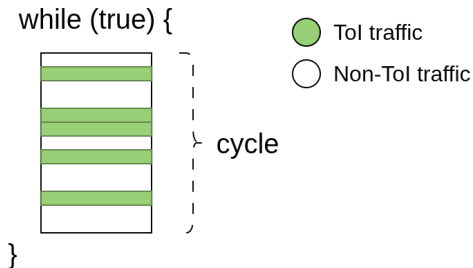


One cycle = traffic issued in one task iteration.

Tol stands for *Type of Interest*.

- ▶ Consider a (pseudo-)cyclic task: each iteration issues roughly the same *type* of traffic, i.e. predictable data structure accesses during each phase.
- ▶ Let  $N^{\min}$  be the minimum *aggregate* requests in a cycle
- ▶ Let  $N_{\text{toi}}^{\max}$  be the maximum amount of Tol requests in a cycle

# The cycle model

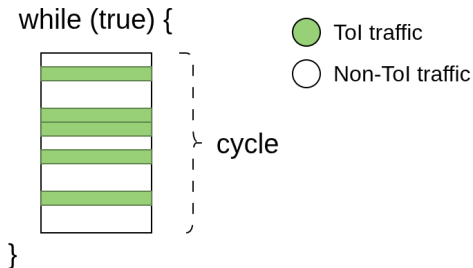


One cycle = traffic issued in one task iteration.

Tol stands for *Type of Interest*.

- ▶ Consider a (pseudo-)cyclic task: each iteration issues roughly the same *type* of traffic, i.e. predictable data structure accesses during each phase.
- ▶ Let  $N^{\min}$  be the minimum *aggregate* requests in a cycle
- ▶ Let  $N_{\text{toi}}^{\max}$  be the maximum amount of Tol requests in a cycle
- ▶ If  $N_{\text{toi}}^{\max} < N^{\min}$ , there is a minimum gap between each burst of Tol traffic, and arbitrarily long bursts are not possible

# The cycle model



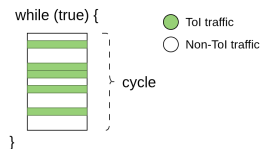
One cycle = traffic issued in one task iteration.

Tol stands for *Type of Interest*.

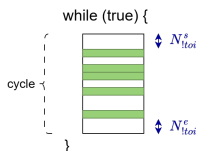
- ▶ Consider a (pseudo-)cyclic task: each iteration issues roughly the same *type* of traffic, i.e. predictable data structure accesses during each phase.
- ▶ Let  $N^{\min}$  be the minimum *aggregate* requests in a cycle
- ▶ Let  $N_{\text{toi}}^{\max}$  be the maximum amount of Tol requests in a cycle
- ▶ If  $N_{\text{toi}}^{\max} < N^{\min}$ , there is a minimum gap between each burst of Tol traffic, and arbitrarily long bursts are not possible
- ▶ Hence, this system admits  $\alpha_{\text{toi}} < \alpha$ .

# Progressively complex models

We iterated on the idea, finding improved versions which take more **structural information** into account



Basic

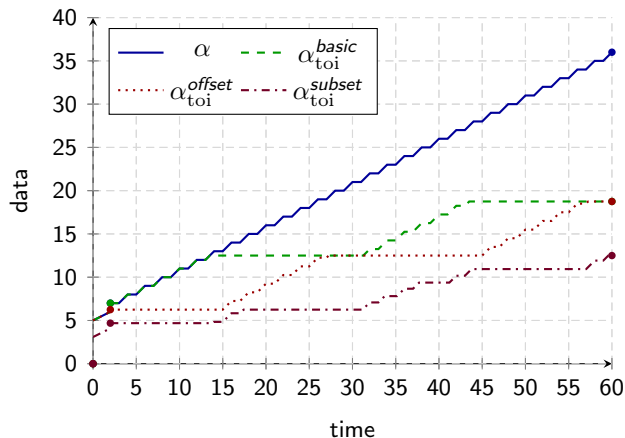


Offset-aware



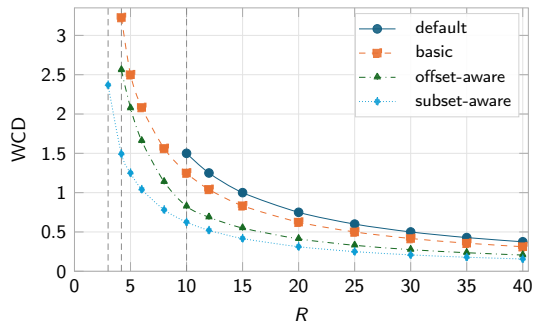
Subset-aware

# Progressively complex models



$\alpha$  (blue), and  $\alpha_{toi}$  from basic (green), offset-aware (red), subset-aware (purple).

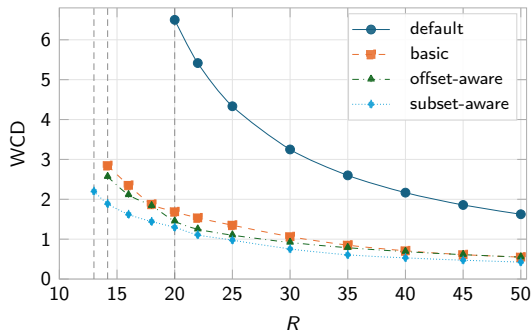
# Improvements in example scenarios



Scenario A: Tol traffic served by a rate-latency node.

► All three methods **widen the stability region**.

► Residual-service case: subset-aware WCD up to  $\sim 30\%$  smaller than the basic method's.



Scenario B: Tol is the competing flow, so our flow gets better residual service.

## Towards the real scenario

- ▶ The above work finds what the mathematical model *needs* to do proper splitting, remove pessimism and improve bounds.

## Towards the real scenario

- ▶ The above work finds what the mathematical model *needs* to do proper splitting, remove pessimism and improve bounds.
- ▶ The technique is general, and applies whenever there is some demultiplexing

# Towards the real scenario

- ▶ The above work finds what the mathematical model *needs* to do proper splitting, remove pessimism and improve bounds.
- ▶ The technique is general, and applies whenever there is some demultiplexing
  - ▶ memory requests to different DRAM banks/chip/etc.
  - ▶ memory requests that hit the cache or not

# Towards the real scenario

- ▶ The above work finds what the mathematical model *needs* to do proper splitting, remove pessimism and improve bounds.
- ▶ The technique is general, and applies whenever there is some demultiplexing
  - ▶ memory requests to different DRAM banks/chip/etc.
  - ▶ memory requests that hit the cache or not
- ▶ What's next: how do we actually derive this information from processes?

# Towards the real scenario

- ▶ The above work finds what the mathematical model *needs* to do proper splitting, remove pessimism and improve bounds.
- ▶ The technique is general, and applies whenever there is some **demultiplexing**
  - ▶ memory requests to different DRAM banks/chip/etc.
  - ▶ memory requests that hit the cache or not
- ▶ What's next: how do we actually derive this information from processes?
- ▶ **Annotated control-flow graphs** with address ranges and structural annotations on accesses of each program phase

# Towards the real scenario

- ▶ The above work finds what the mathematical model *needs* to do proper splitting, remove pessimism and improve bounds.
- ▶ The technique is general, and applies whenever there is some **demultiplexing**
  - ▶ memory requests to different DRAM banks/chip/etc.
  - ▶ memory requests that hit the cache or not
- ▶ What's next: how do we actually derive this information from processes?
- ▶ **Annotated control-flow graphs** with address ranges and structural annotations on accesses of each program phase
- ▶ Deriving these CFGs from binaries

# Takeaways

- ▶ Worst-case memory latencies has to consider many processors and memory elements
- ▶ While load-balancing is easy to model *on average*, its *worst-case* introduces lots of pessimism
- ▶ SOTA modelling forces *duplication* of arrival curves at each demultiplexing point
- ▶ We investigated the *structural information* needed to avoid this pessimism
- ▶ Next step is deriving these parameters out of tasks, with the current direction being *annotated CFGs*

Thank you!