



HERMES

Accelerating Long-Latency Load Requests via Perceptron-Based Off-Chip Load Prediction

Rahul Bera, Konstantinos Kanellopoulos, Shankar Balachandran,
David Novo, Ataberk Olgun, Mohammad Sadrosadati, Onur Mutlu

The Key Problem

Long-latency **off-chip** load requests

The Key Problem

Long-latency **off-chip** load requests



Often **block instruction retirement** from
Reorder Buffer (ROB)

The Key Problem

Long-latency **off-chip** load requests



Often **block instruction retirement** from
Reorder Buffer (ROB)



Limits performance

Traditional Solutions



Traditional Solutions



1

Employ sophisticated **prefetchers**

Traditional Solutions



1

Employ sophisticated **prefetchers**

2

Increase size of on-chip caches

Key Observations



Observation 1

Nearly **50%** of the off-chip requests
in a no-prefetching system
still go to the main memory
even in presence of state-of-the-art prefetcher

Key Observations

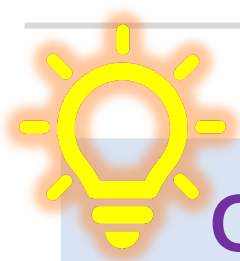


Observation 1

Nearly **50%** of the off-chip requests
in a no-prefetching system
still go to the main memory
even in presence of state-of-the-art prefetcher

70% of these off-chip loads **block the ROB**

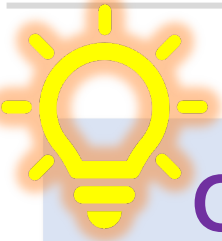
Key Observations



Observation 2

On-chip cache access latency significantly contributes to total off-chip load latency

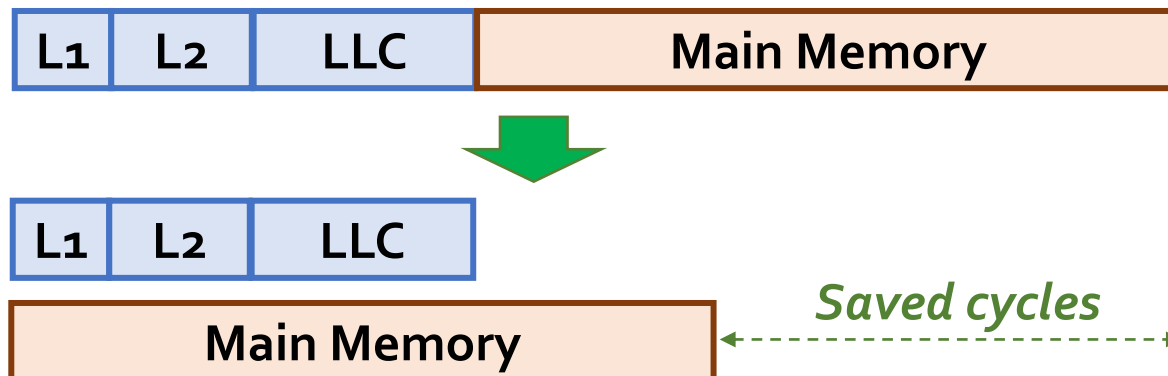
Key Observations



Observation 2

On-chip cache access latency significantly contributes to total off-chip load latency

40% of the stalls caused by an off-chip load can be reduced by **removing on-chip cache access latency from its critical path**



And Things are Getting Worse...

And Things are Getting Worse...

Surprisingly High Latency Discovered During Alder Lake Test With DDR5-6400 Memory

Jason R. Wilson · Sep 16, 2021, 02:26 PM EDT

393

Memory	90138 MB/s	88844 MB/s	77750 MB/s	92.5 ns
L1 Cache	3422.0 GB/s	1717.56GB/s	3422.2 GB/s	0.9 ns
L2 Cache	985.06 GB/s	820.36 GB/s	986.84 GB/s	2.6 ns
L3 Cache	512.61 GB/s	475.36 GB/s	493.62 GB/s	12.0 ns
CPU Type	10-Core			
CPU Stepping	A0			
CPU Clock	4596.6 MHz (original: 3200 MHz, overclock: 43%)			
CPU FSB	99.9 MHz (original: 100 MHz)			
CPU Multiplier	46x	North Bridge Clock		4296.9 MHz
Memory Bus	3198.8 MHz	DRAM:FSB Ratio		48:3
Memory Type	Dual Channel DDR5-6400 (40-40-40-85 CR2)			
Chipset	Z790			
Motherboard	ASUS ROG Strix Z790-A			

Trending Stories

PlayStation Call of Duty: Modern Warfare II Comparison Video That the Game is Being Hammered by Last-Gen Consoles

95 Active Readers

Diablo IV Gets Almost One Hour of Leaked Gameplay Footage

56 Active Readers

Hours Into The ETH Merge, NVIDIA GeForce & AMD Radeon Graphics Card Prices Hit Their Lowest Since 3090 Ti Drops Below \$1000

28 Active Readers

SpaceX Starship Might Not Fly Until Orbit This Year Suggests NASA

20 Active Readers

AMD Brings Ray Tracing Support to Radeon RX 6000 Series

And Things are Getting Worse...

Surprisingly High Latency Discovered During Alder Lake Test With DDR5-6400

Jason R. Wilson · Sep 16, 2021, 02:26 PM EDT

The Intel 12th Gen Core i9-12900K Review: Hybrid Performance Brings Hybrid Complexity

by [Dr. Ian Cutress](#) & [Andrei Frumusanu](#) on November 4, 2021 9:00 AM EST

474
Comments

+ Add A
Comment

Posted in [CPUs](#) [Intel](#) [DDR4](#) [DDR5](#) [Alder Lake](#) [Golden Cove](#) [Gracemont](#) [Windows 11](#) [Core i9-12900K](#)
[P-core](#) [E-Core](#) [Thread Director](#)

Memory	90138 MB/s	88844	P-core	E-Core	Thread Director
L1 Cache	3422.0 GB/s	1717.56GB/s	3422.2 GB/s	0.9 ns	
L2 Cache	985.06 GB/s	820.36 GB/s	986.84 GB/s	2.6 ns	
L3 Cache	512.61 GB/s	475.36 GB/s	493.62 GB/s	12.0 ns	
CPU Type	10-Core				
CPU Stepping	A0				
CPU Clock	4596.6 MHz (original: 3200 MHz, overclock: 43%)				
CPU FSB	99.9 MHz (original: 100 MHz)				
CPU Multiplier	46x	North Bridge Clock		4296.9 MHz	
Memory Bus	3198.8 MHz	DRAM:FSB Ratio		48:3	
Memory Type	Dual Channel DDR5-6400 (40-40-40-85 CR2)				
Chipset	Z790				
Motherboard	ASUS ROG Strix Z790-A				

Trend

PlayStation
Warfare
That the
Last-Gen

95

Diablo
Leaked

56

Hours
GeForce

Card P
3090 T

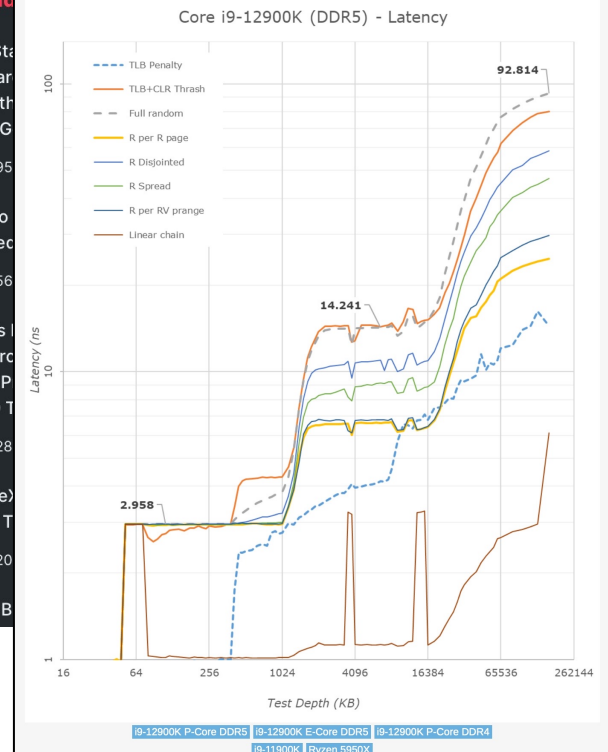
28

Space
Orbit T

20

AMD B

Cache Latencies and DDR5 vs DDR4



And Things are Getting Worse...

Surprisingly High Latency Test With DDR5-6400

Jason R. Wilson · Sep 16, 2021, 02:26 PM EDT

Memory

90138 MB/s

88844

L1 Cache

3422.0 GB/s

1717.56GB/s

3422.2 GB/s

0.9 ns

The Intel 12th Gen Core i9-12900K Review: Hybrid Performance Brings Hybrid Complexity

by [Dr. Ian Cutress](#) & [Andrei Frumusanu](#) on November 4, 2021 9:00 AM EST

Posted in [CPUs](#) [Intel](#) [DDR4](#) [DDR5](#) [Alder Lake](#) [Golden Cove](#) [Gracemont](#) [Windows 11](#) [Core i9-12900K](#)
[P-core](#) [E-Core](#) [Thread Director](#)

474
Comments

+ Add A
Comment



Posted by u/neeteshkurup 11 months ago

506

The L3 cache issue is more than latency related?



Discussion

AIDA64 Cache & Memory Benchmark

	Read	Write	Copy	Latency
Memory	94022 MB/s	94428 MB/s	95908 MB/s	80.4 ns
L1 Cache	8293.5 GB/s	4197.5 GB/s	8195.1 GB/s	0.9 ns
L2 Cache	3986.6 GB/s	3943.5 GB/s	4159.1 GB/s	2.8 ns
L3 Cache	515.07 GB/s	2145.7 GB/s	2051.0 GB/s	10.1 ns
CPU Type	32-Core AMD Ryzen Threadripper 3970X (Castle Peak, Socket TR4+)			
CPU Stepping	55P-B0			
CPU Clock	4350.8 MHz			
CPU FSB	100.0 MHz (original: 100 MHz)			
CPU Multiplier	43.5x	North Bridge Clock		

About Community



r/Amd

Welcome to /r/AMD — the subreddit for all things AMD; come talk about Ryzen, Threadripper, EPYC, RDNA3, rumors, show-off your build and more. /r/AMD is community run and does not represent AMD in any capacity unless specified.

1.5m
Members

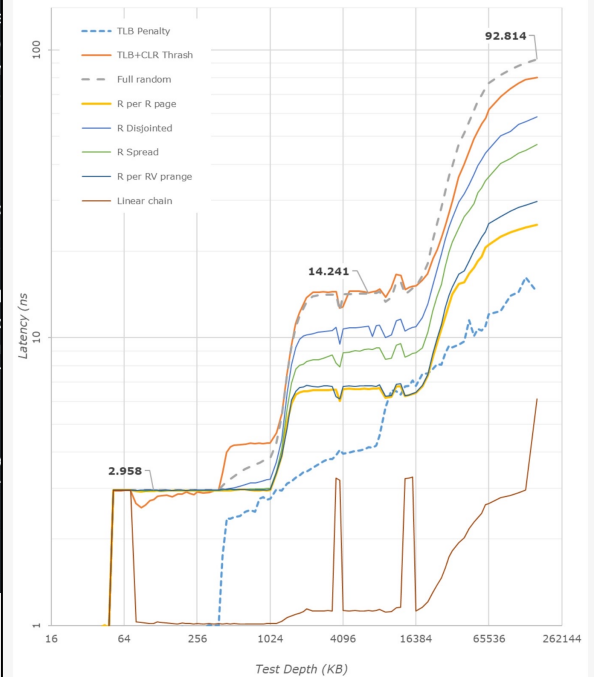
1.1k
Online

Created Jul 1, 2010

Join

Cache Latencies and DDR5 vs DDR4

Core i9-12900K (DDR5) - Latency



Our Goal

Improve processor performance
by **removing on-chip cache access latency**
from the **critical path of off-chip loads**



HERMES



HERMES



Predicts which load requests might go off-chip



HERMES

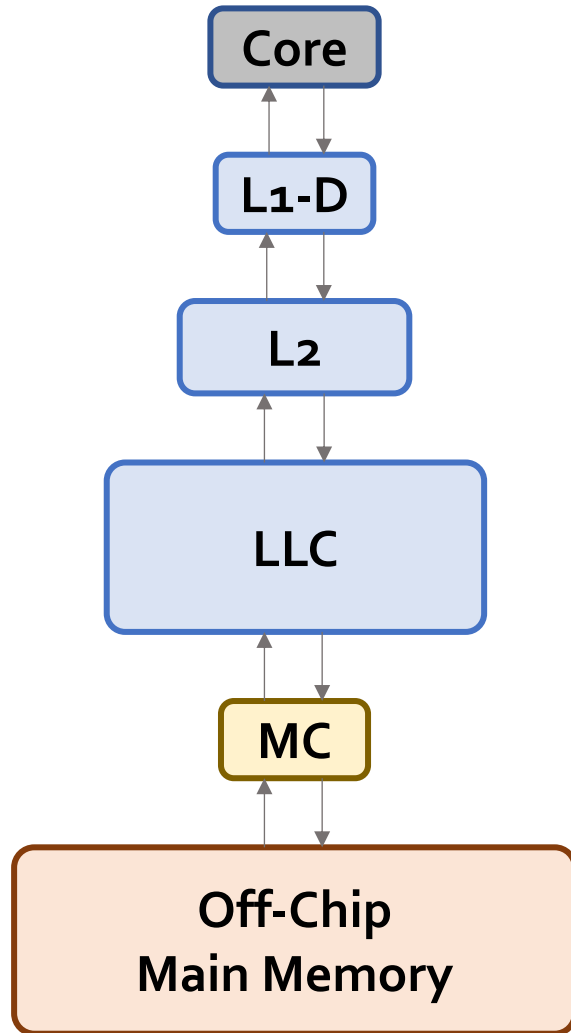


Predicts which load requests might go off-chip

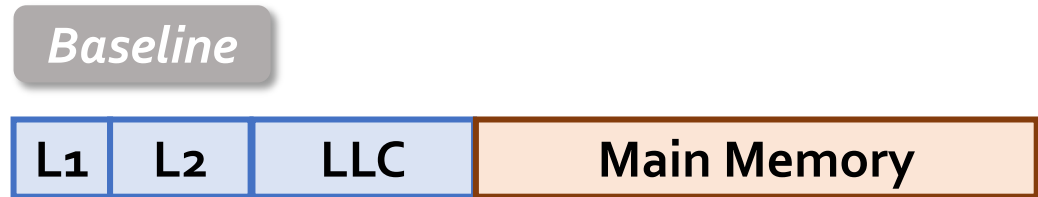
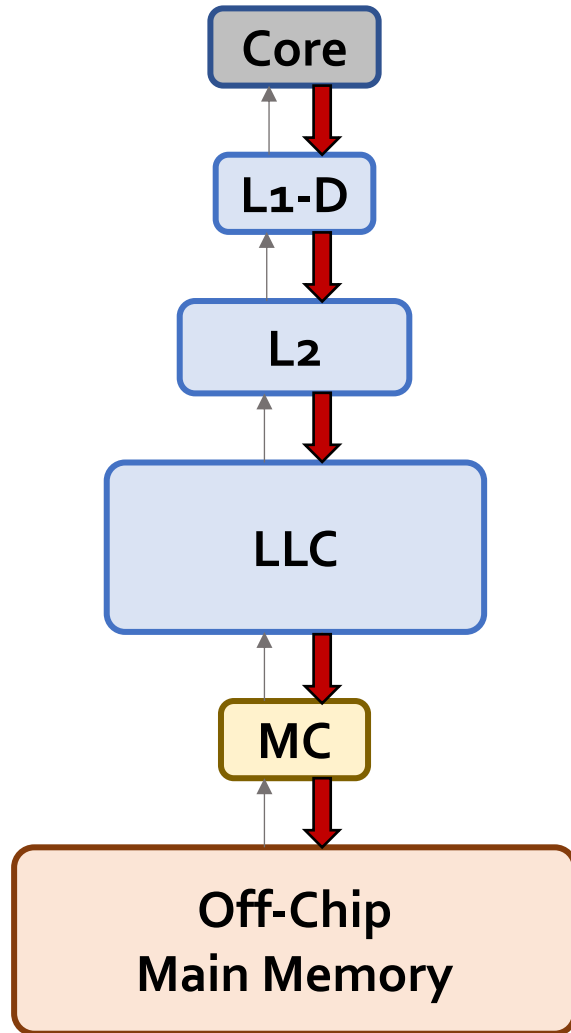


Starts **fetching** data **directly** from main memory while concurrently accessing the cache hierarchy

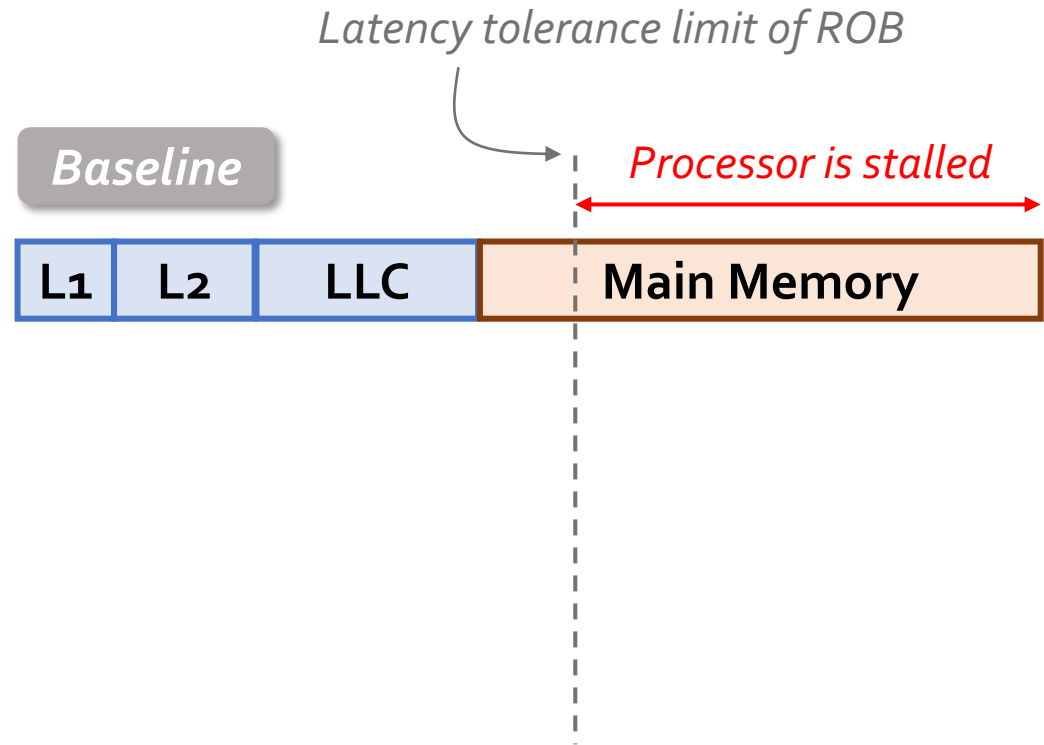
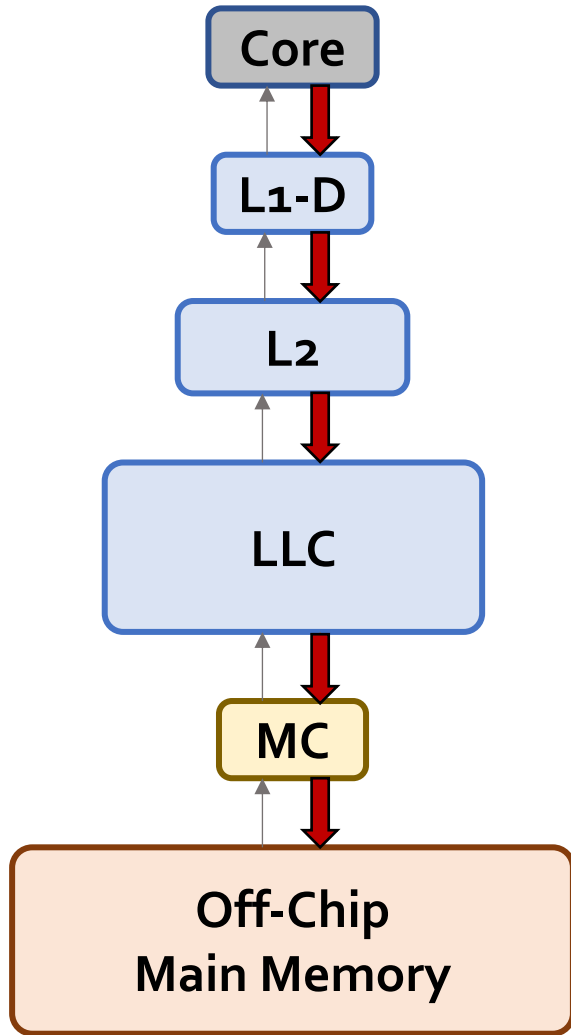
Hermes Overview



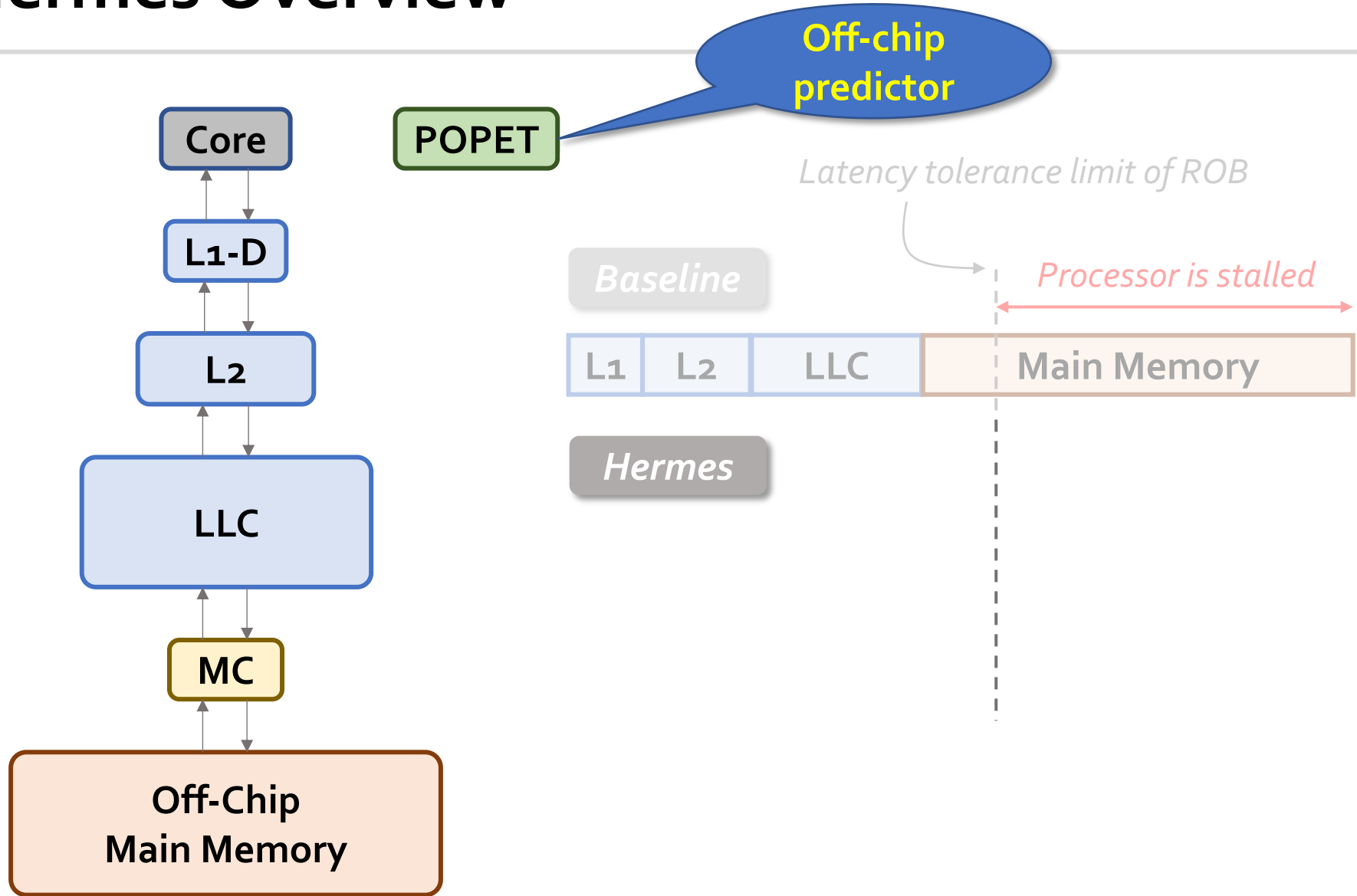
Hermes Overview



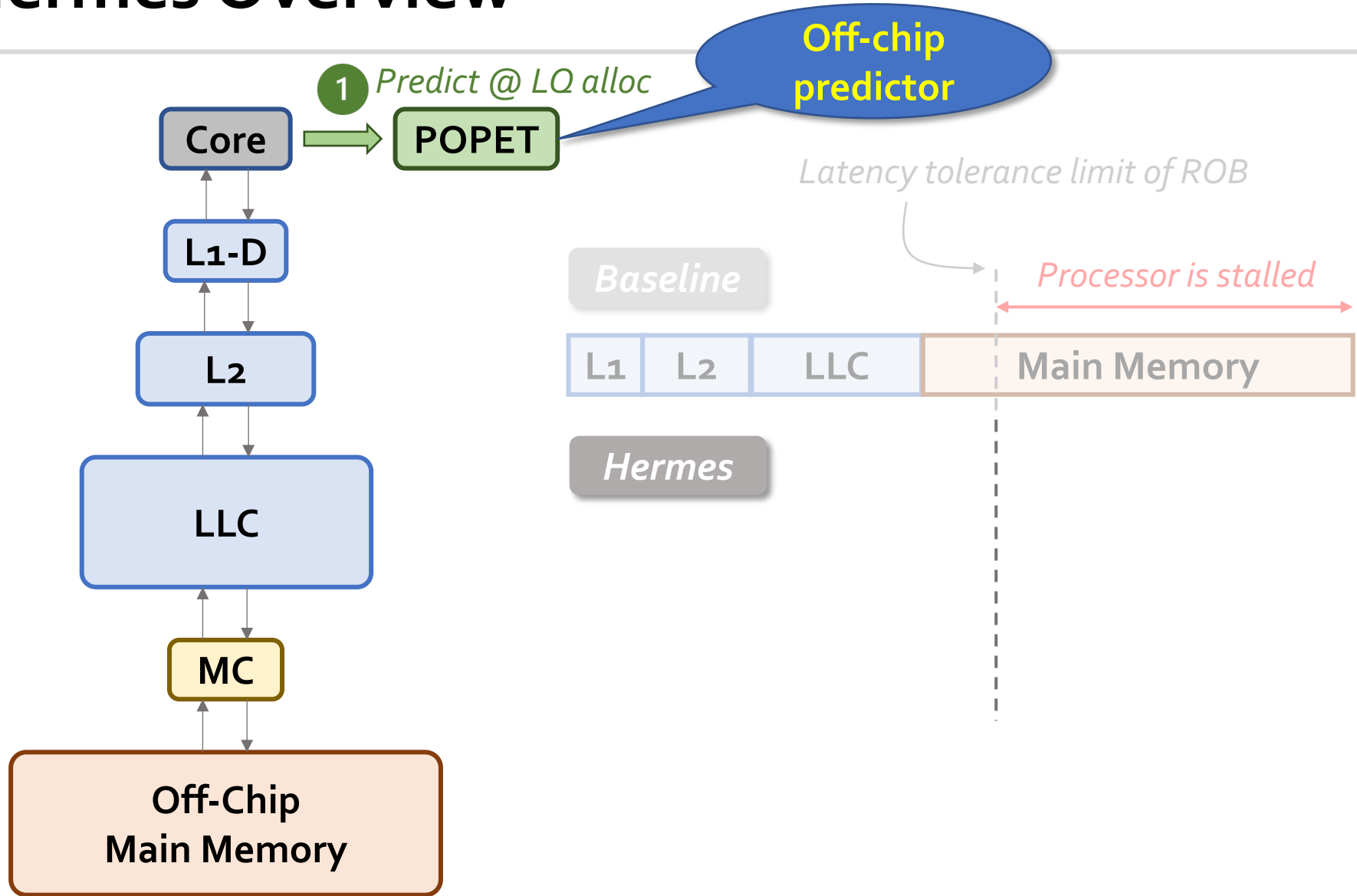
Hermes Overview



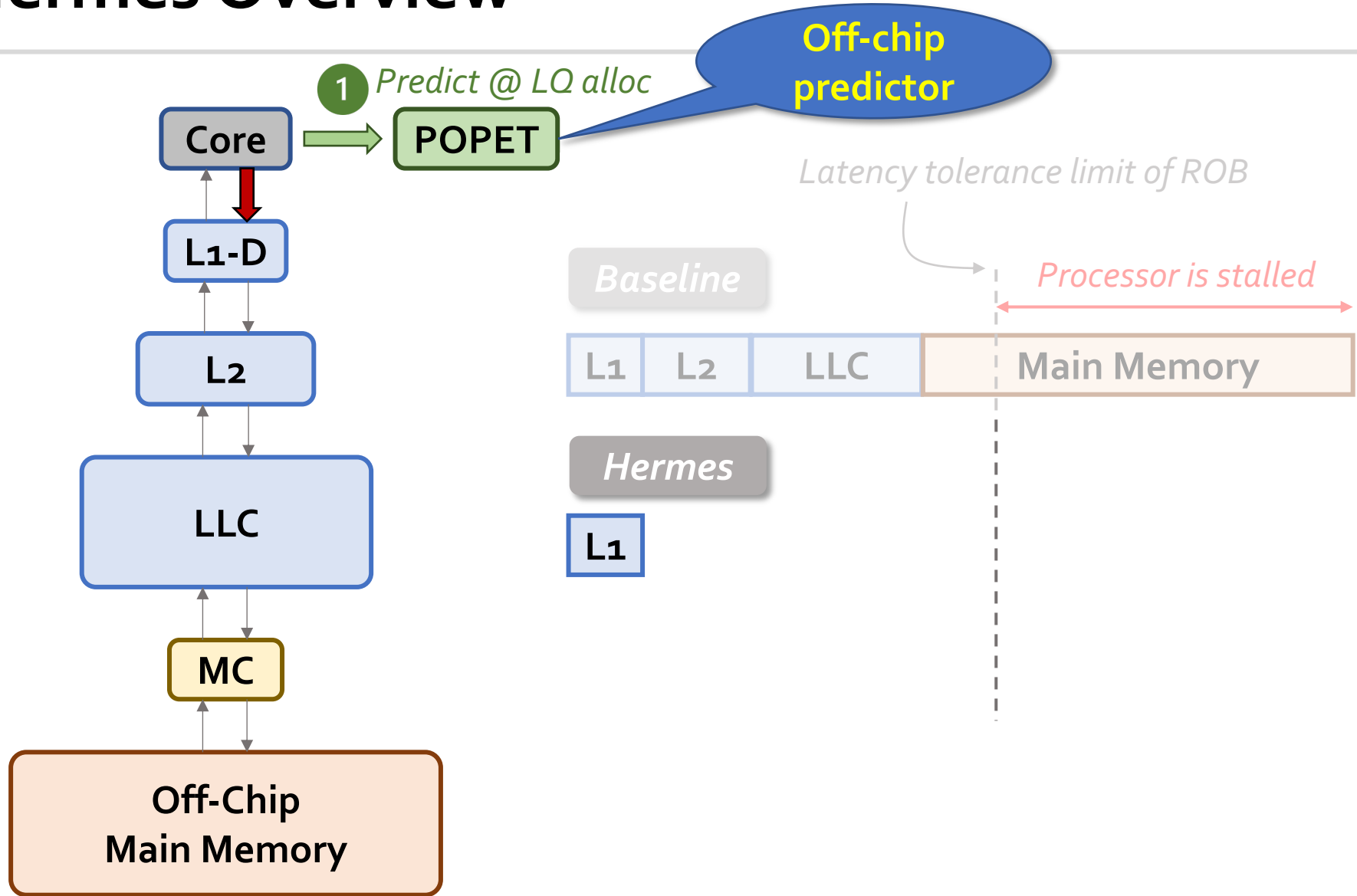
Hermes Overview



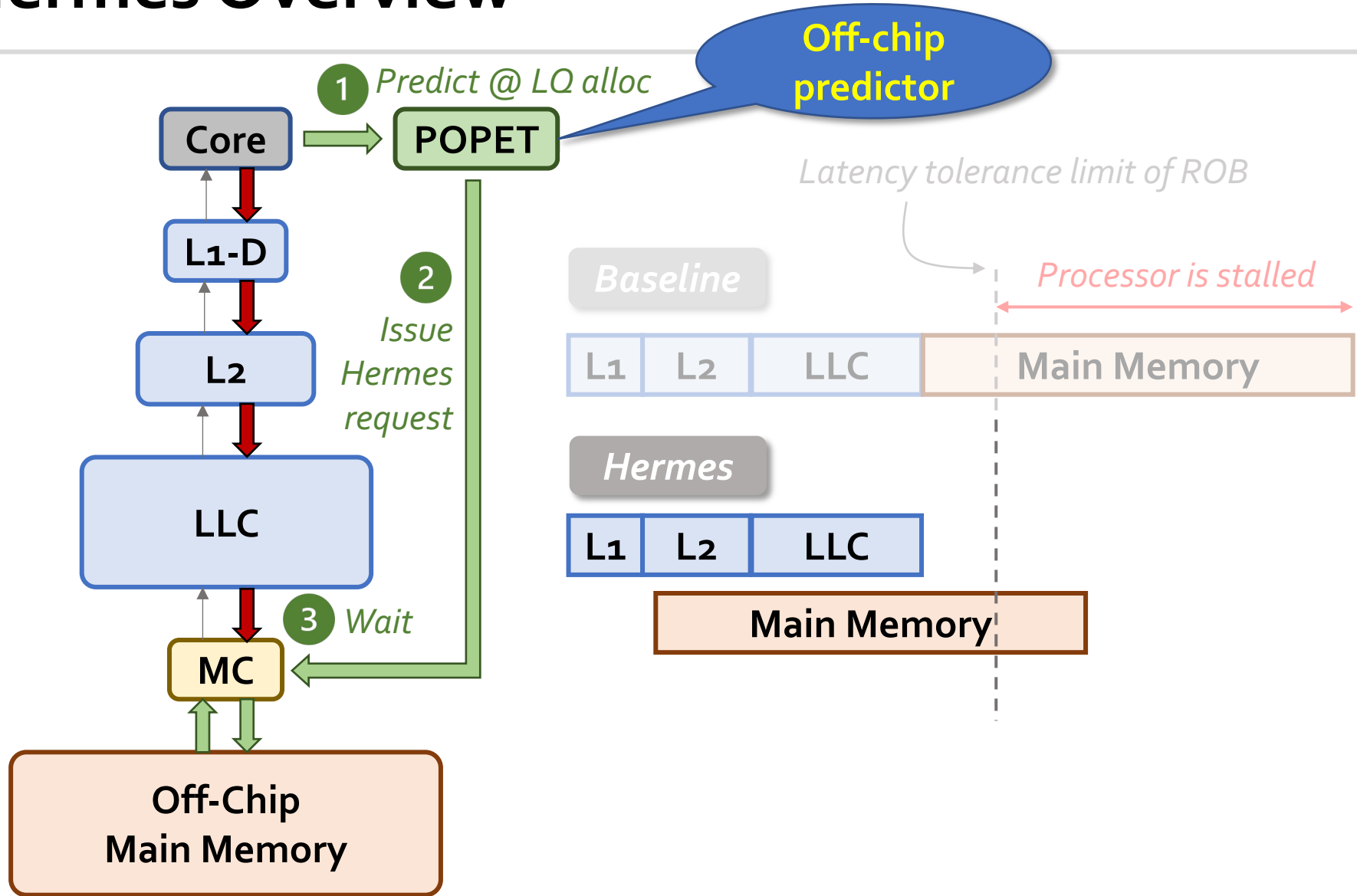
Hermes Overview



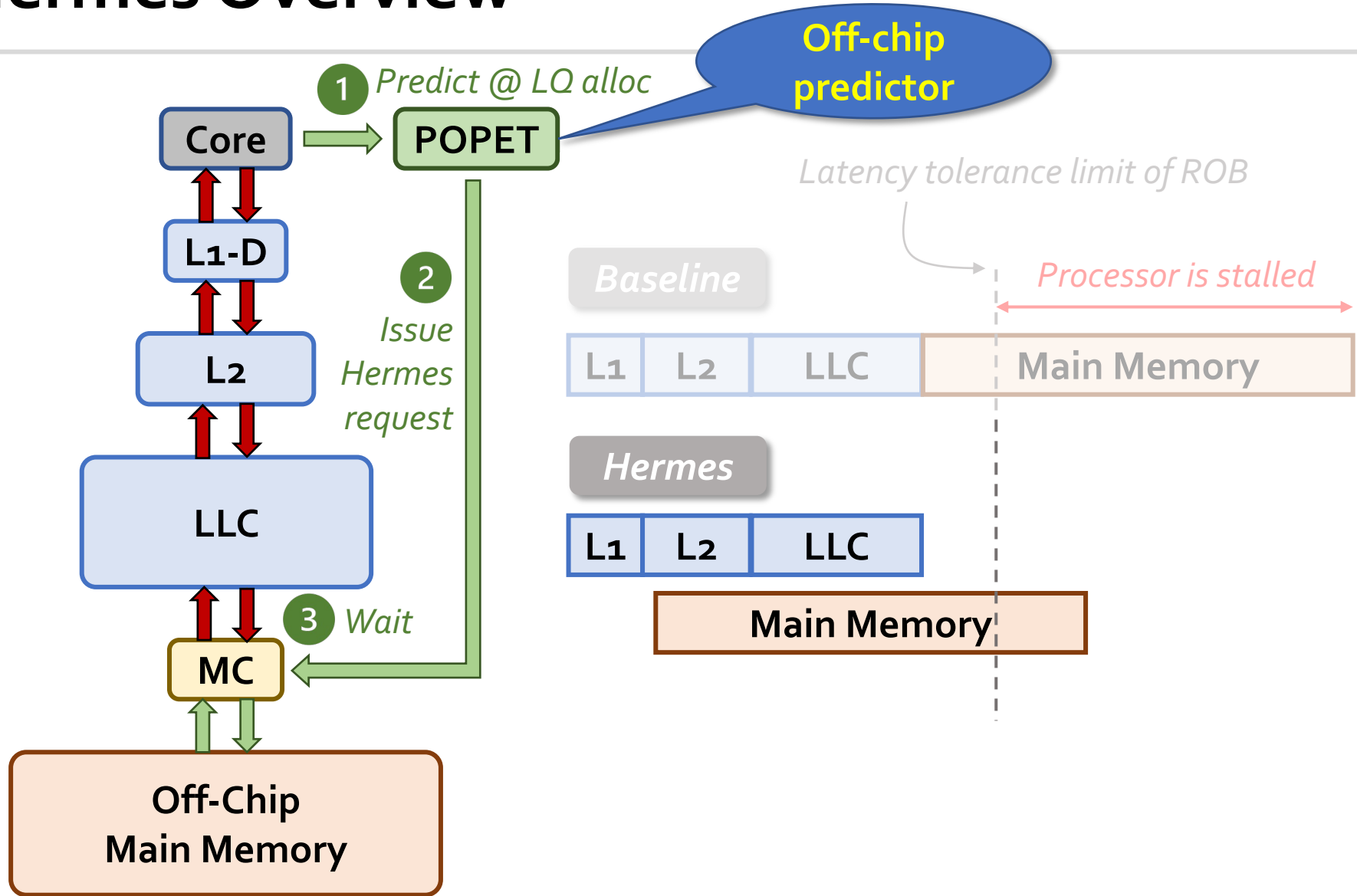
Hermes Overview



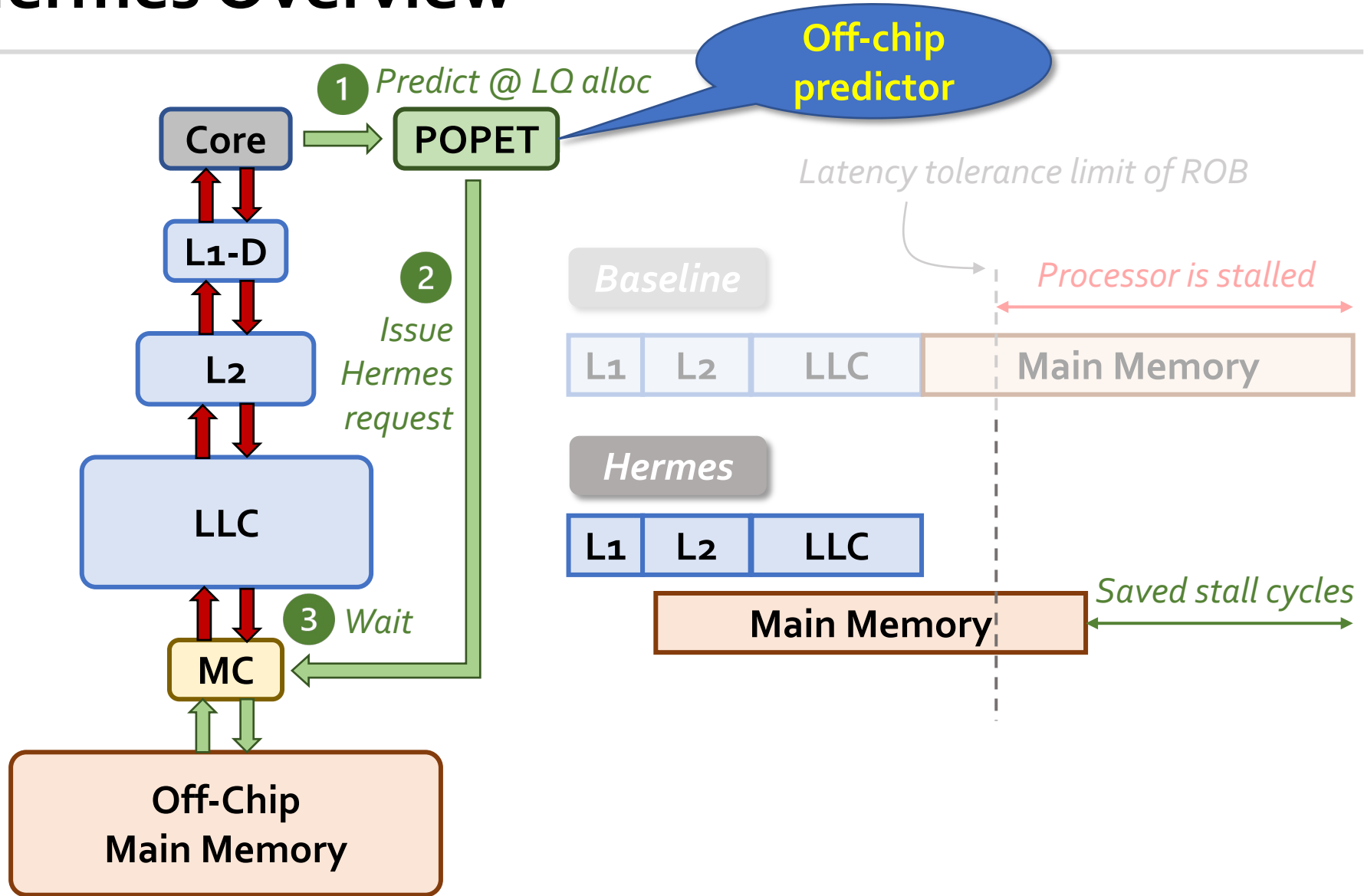
Hermes Overview



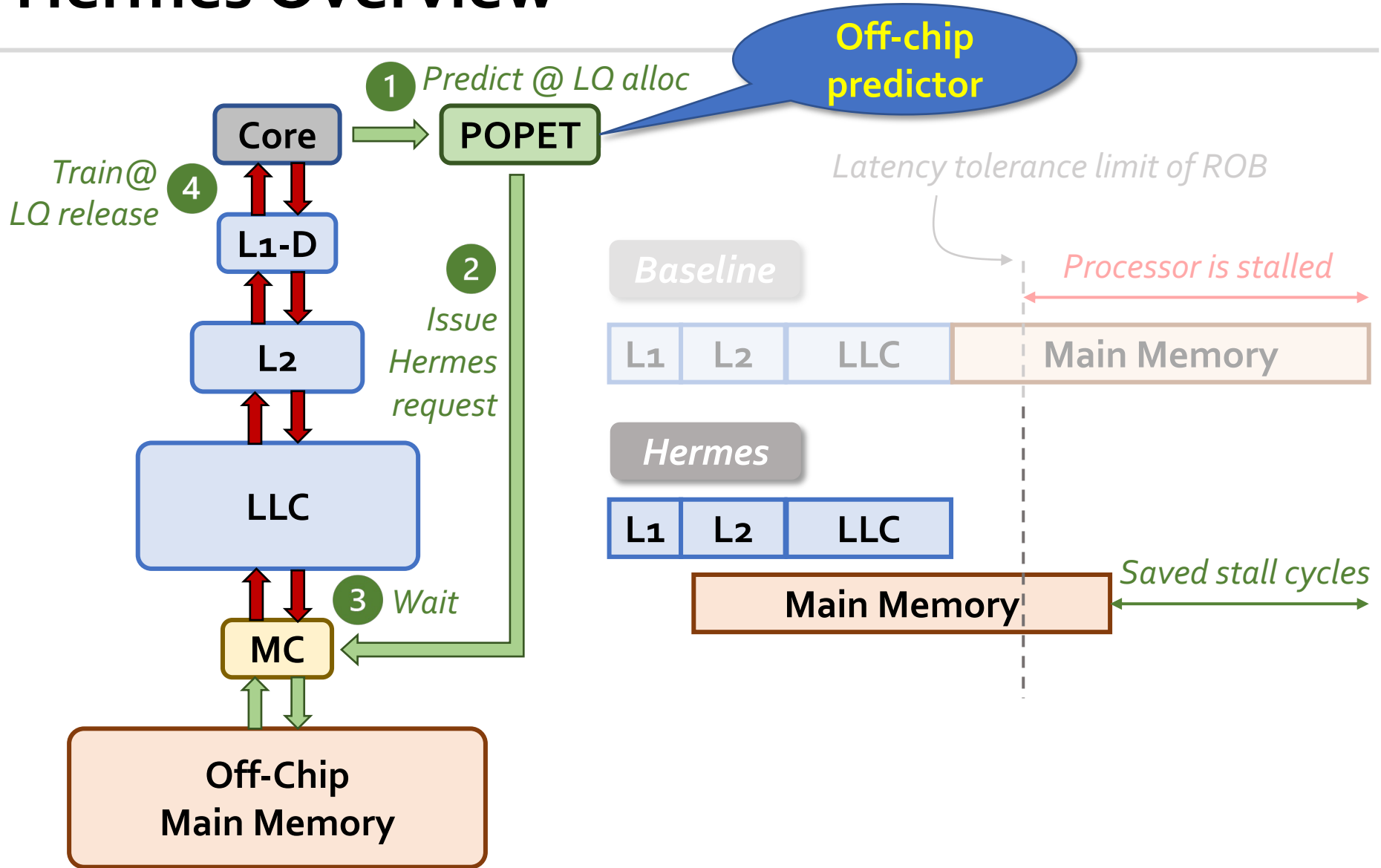
Hermes Overview



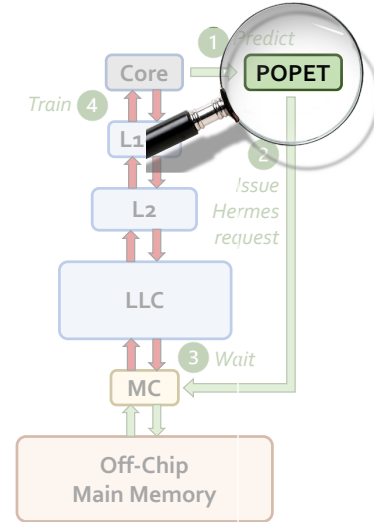
Hermes Overview



Hermes Overview



Designing The Off-Chip Predictor

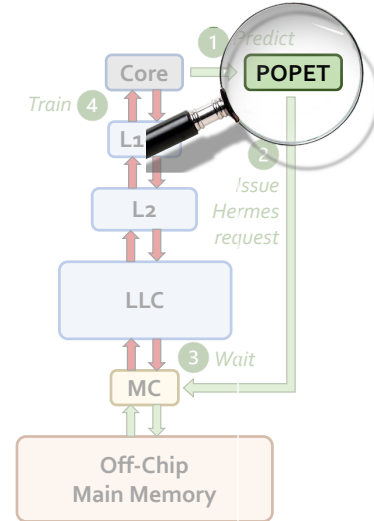
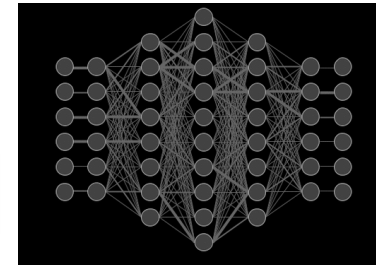


Designing The Off-Chip Predictor

Tracking cache contents

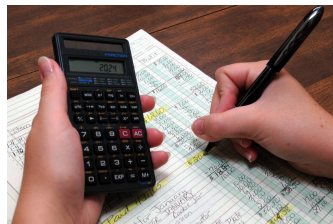


Learning from program behavior



Designing The Off-Chip Predictor

Tracking cache contents

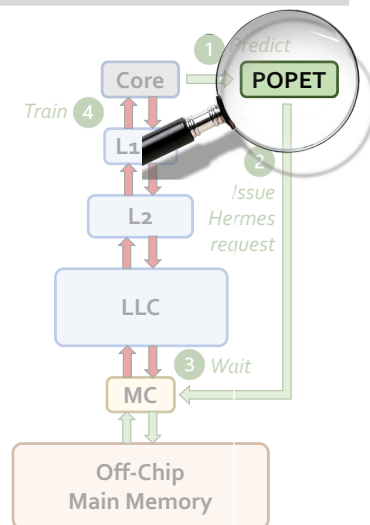


✗ Large metadata

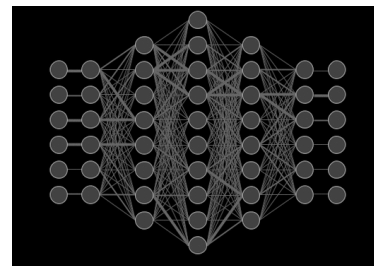
- Metadata size increases with cache hierarchy size

✗ Track every cache operations

- Gets tricky based on the cache hierarchy configuration (e.g., inclusivity, bypassing,...)

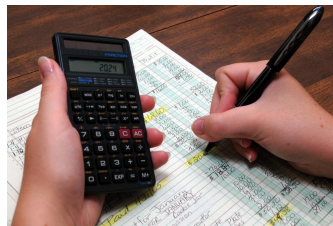


Learning from program behavior



Designing The Off-Chip Predictor

Tracking cache contents

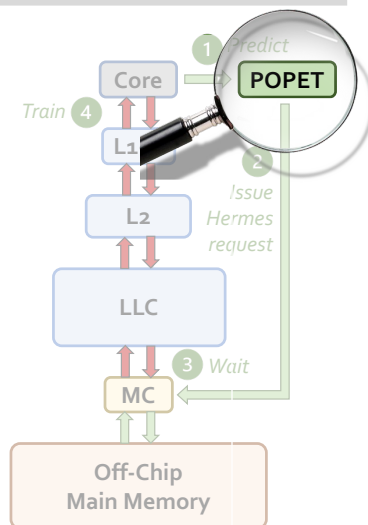


✗ Large metadata

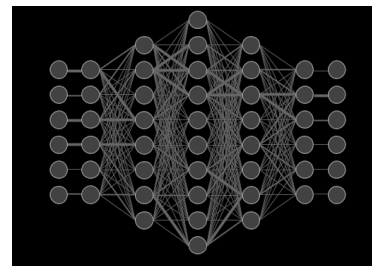
- Metadata size increases with cache hierarchy size

✗ Track every cache operations

- Gets tricky based on the cache hierarchy configuration (e.g., inclusivity, bypassing,...)



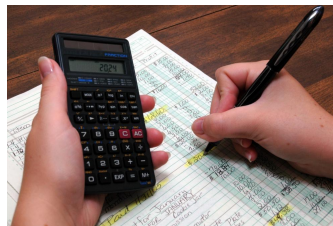
Learning from program behavior



Correlate different program features with off-chip loads

Designing The Off-Chip Predictor

Tracking cache contents

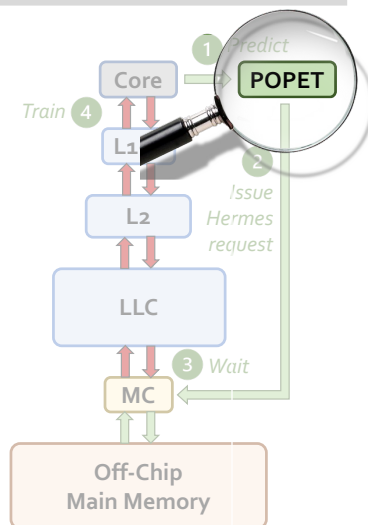


❌ Large metadata

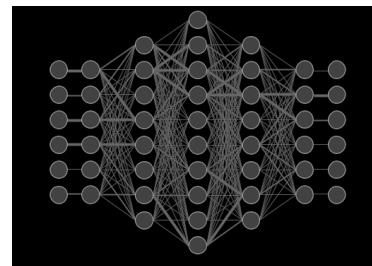
- Metadata size increases with cache hierarchy size

❌ Track every cache operations

- Gets tricky based on the cache hierarchy configuration (e.g., inclusivity, bypassing,...)



Learning from program behavior



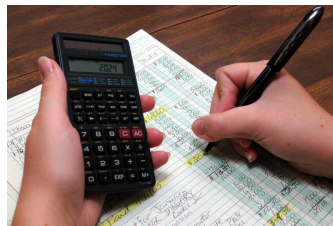
Correlate different program features with off-chip loads

✅ Lower storage overhead

✅ Lower design complexity

Designing The Off-Chip Predictor

Tracking cache contents

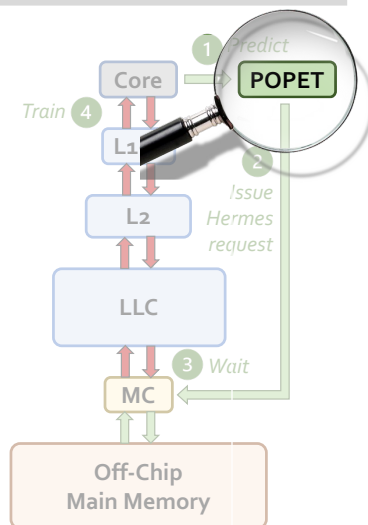


❌ Large metadata

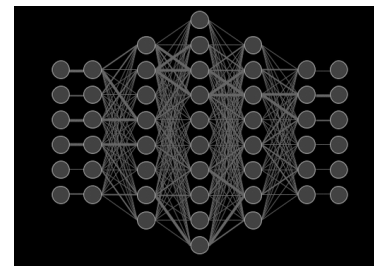
- Metadata size increases with cache hierarchy size

❌ Track every cache operations

- Gets tricky based on the cache hierarchy configuration (e.g., inclusivity, bypassing,...)



Learning from program behavior

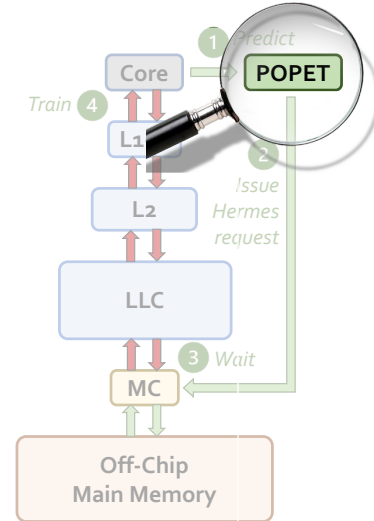


Correlate different program features with off-chip loads

- ✅ Lower storage overhead
- ✅ Lower design complexity

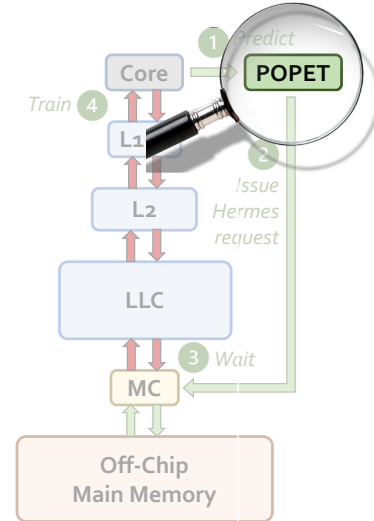
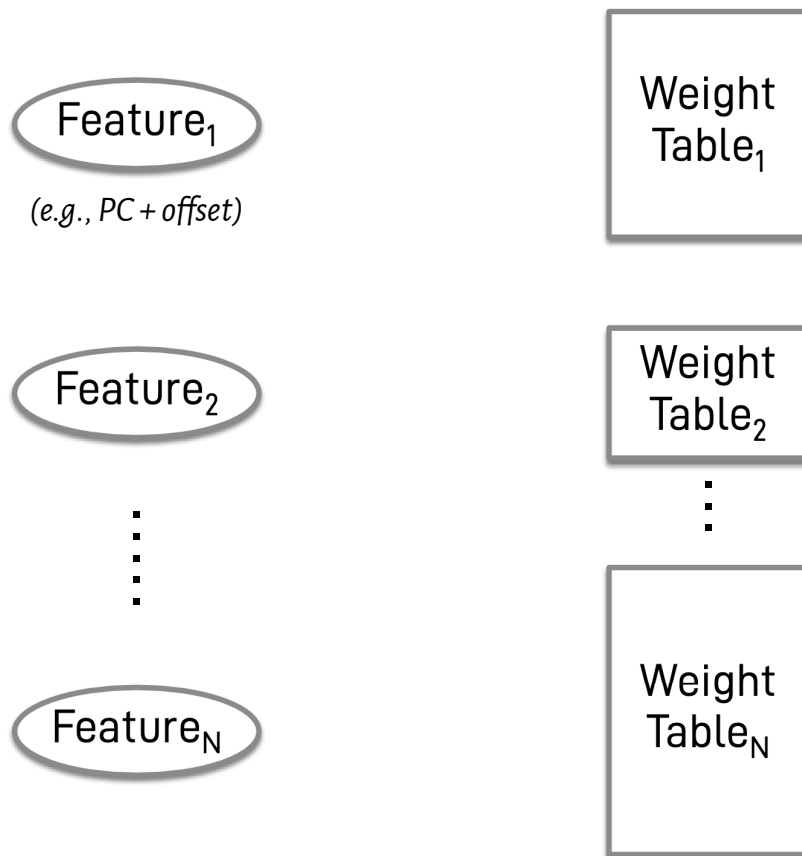
POPET: Perceptron-Based Off-Chip Predictor

- Multi-feature hashed perceptron model^[1]



POPET: Perceptron-Based Off-Chip Predictor

- Multi-feature hashed perceptron model^[1]
 - Each feature has its own *weight table*
 - Stores *correlation* between *feature value* and *off-chip prediction*



Predicting using POPET

- Using simple **table lookups**, **addition**, and **comparison**

Feature₁
(e.g., PC + offset)

Weight
Table₁

Feature₂

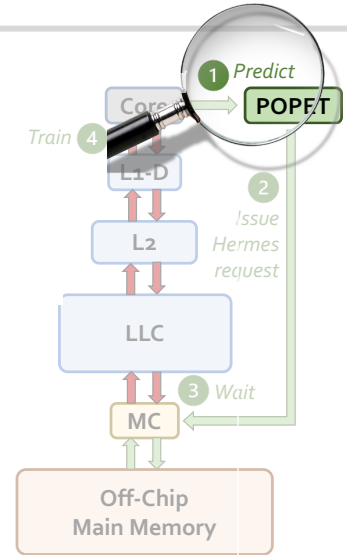
Weight
Table₂

⋮

⋮

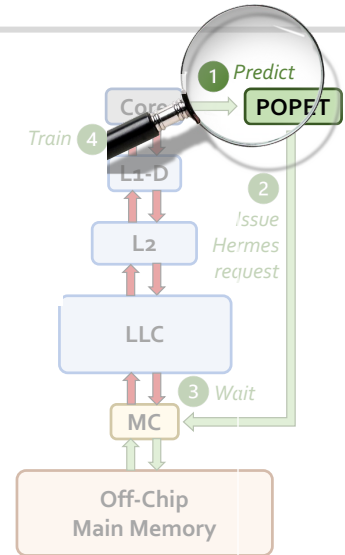
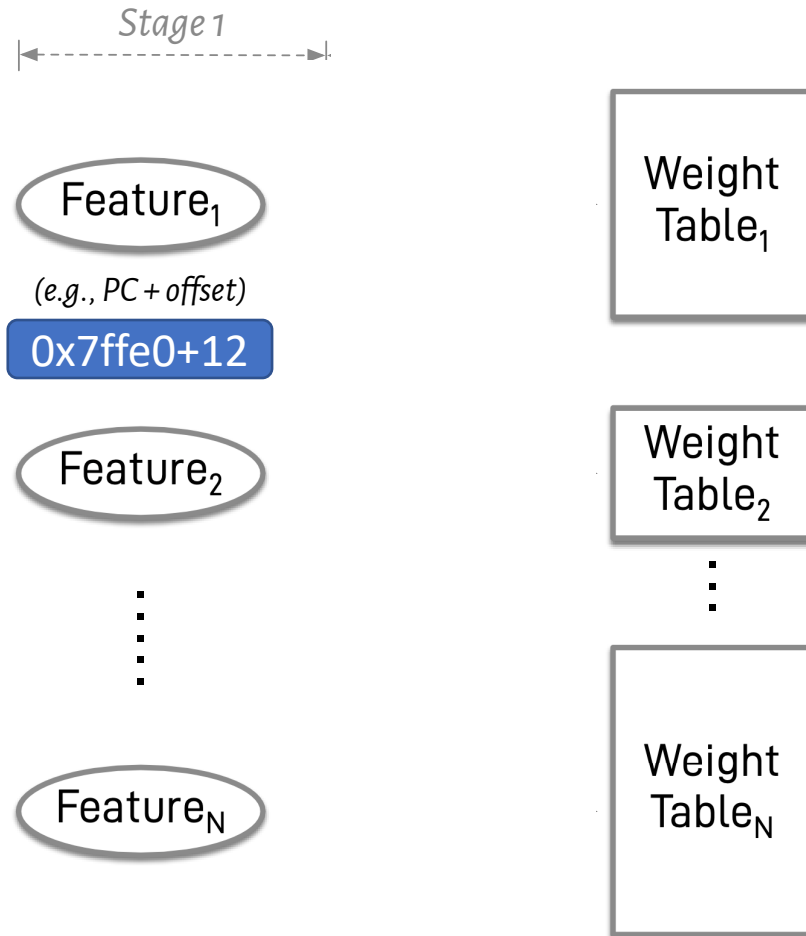
Feature_N

Weight
Table_N



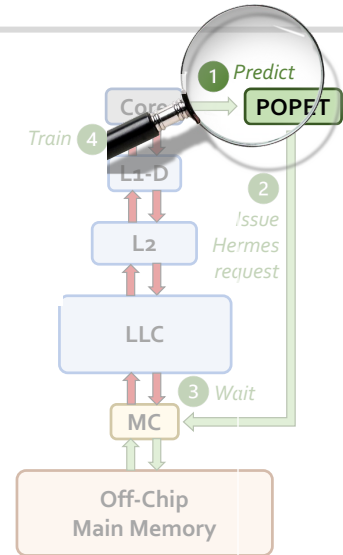
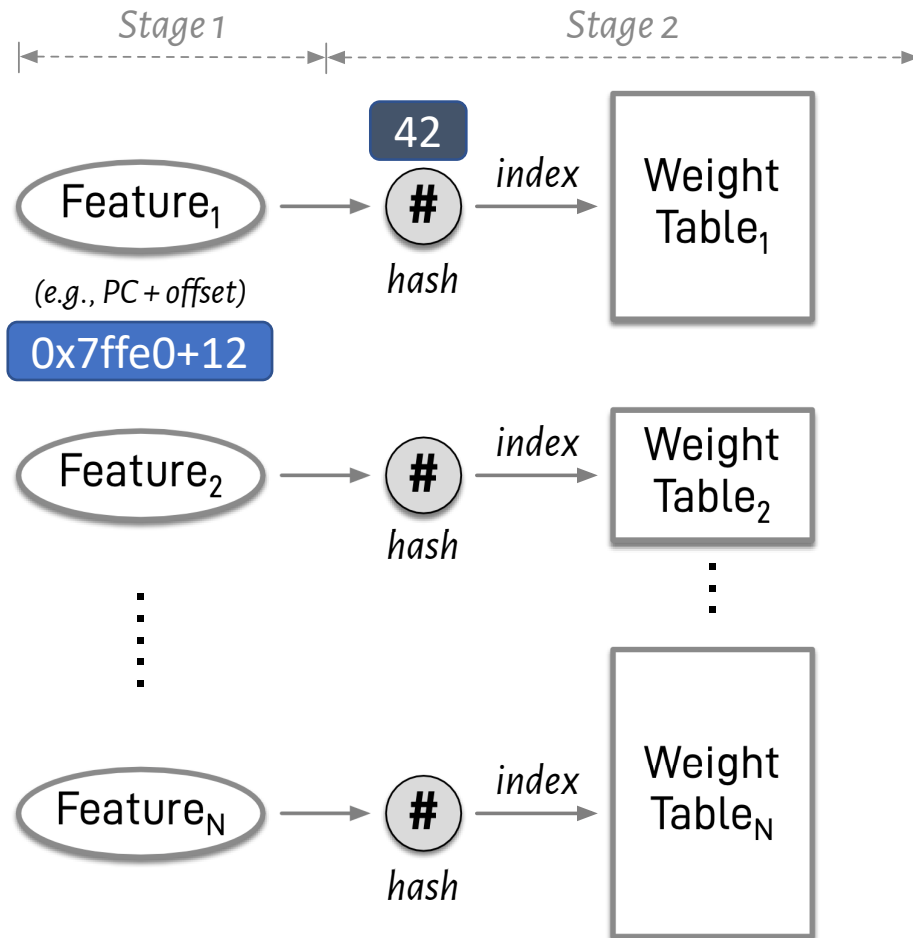
Predicting using POPET

- Using simple **table lookups**, **addition**, and **comparison**



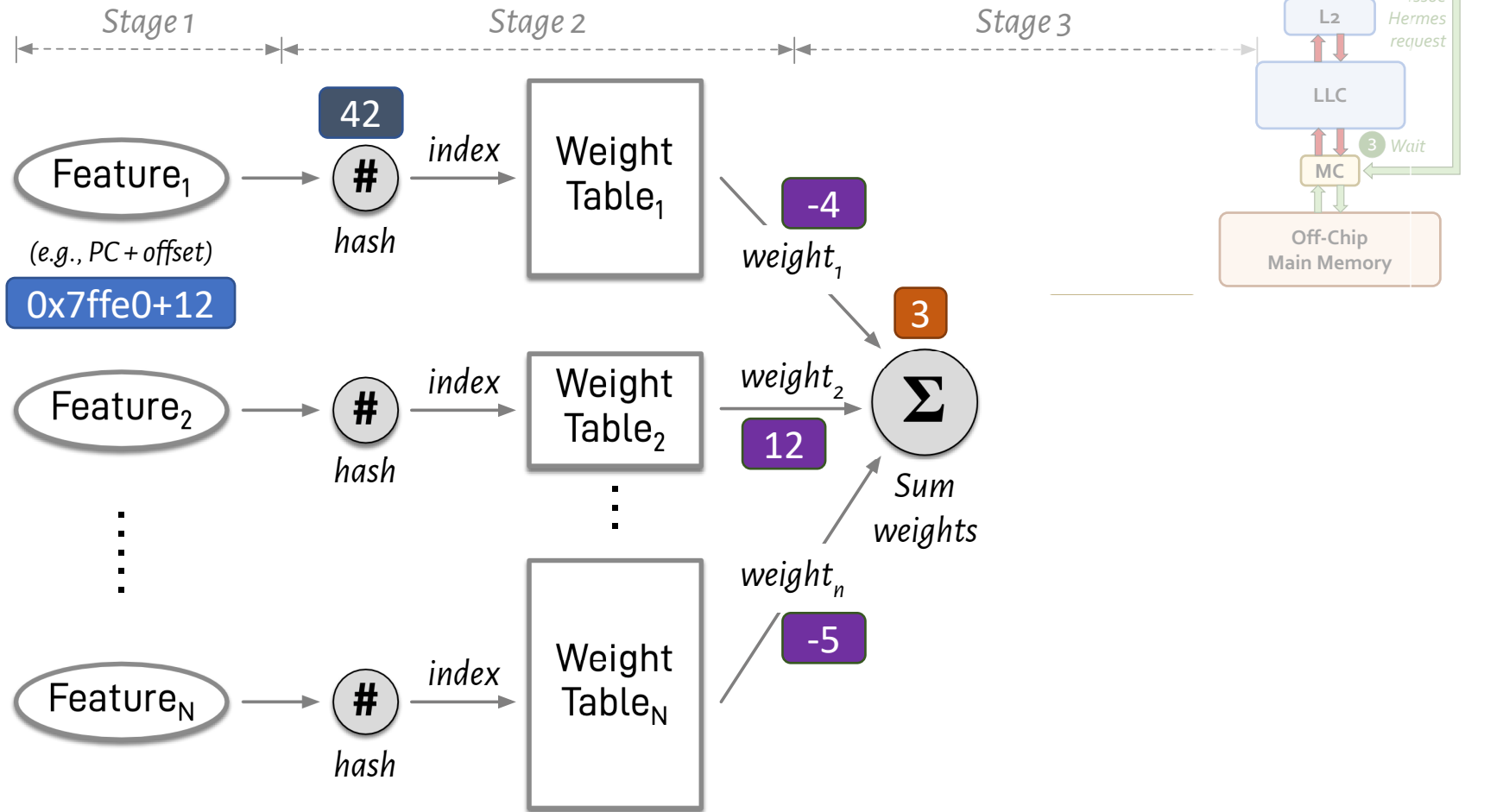
Predicting using POPET

- Using simple **table lookups**, **addition**, and **comparison**



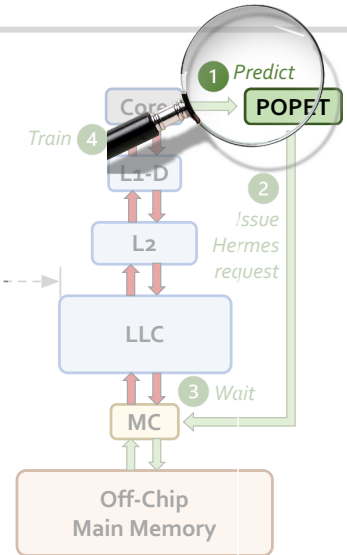
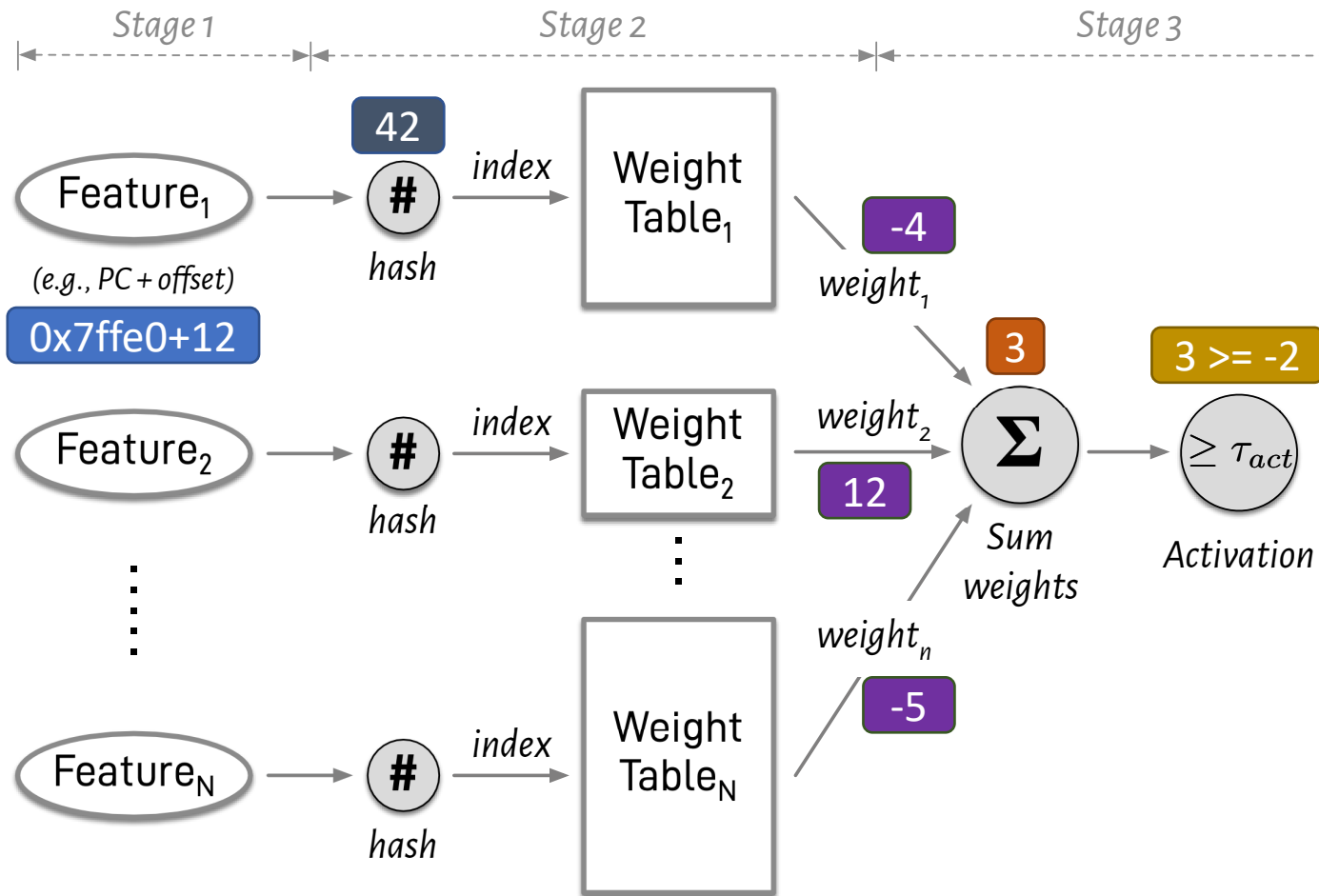
Predicting using POPET

- Using simple **table lookups**, **addition**, and **comparison**



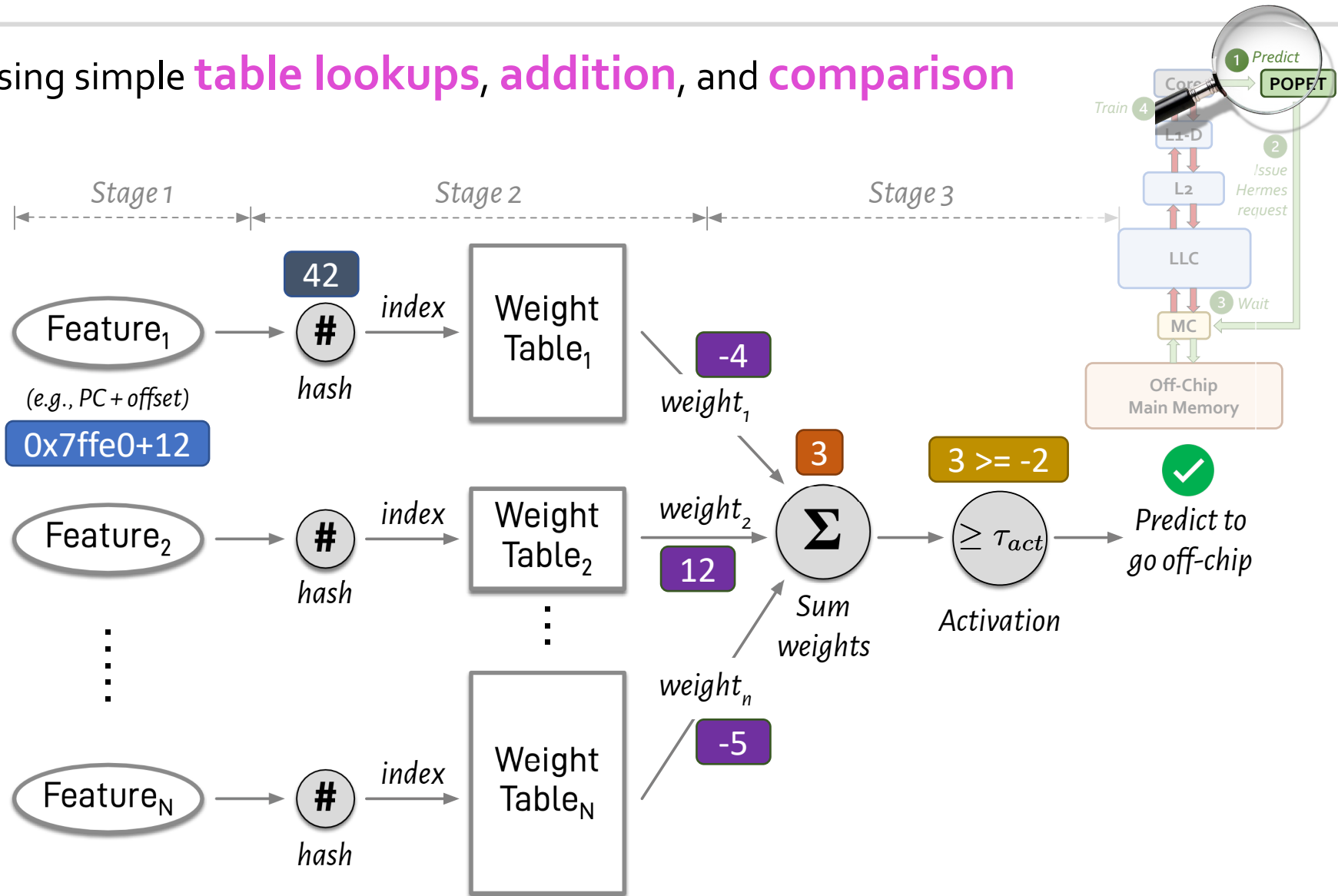
Predicting using POPET

- Using simple **table lookups**, **addition**, and **comparison**



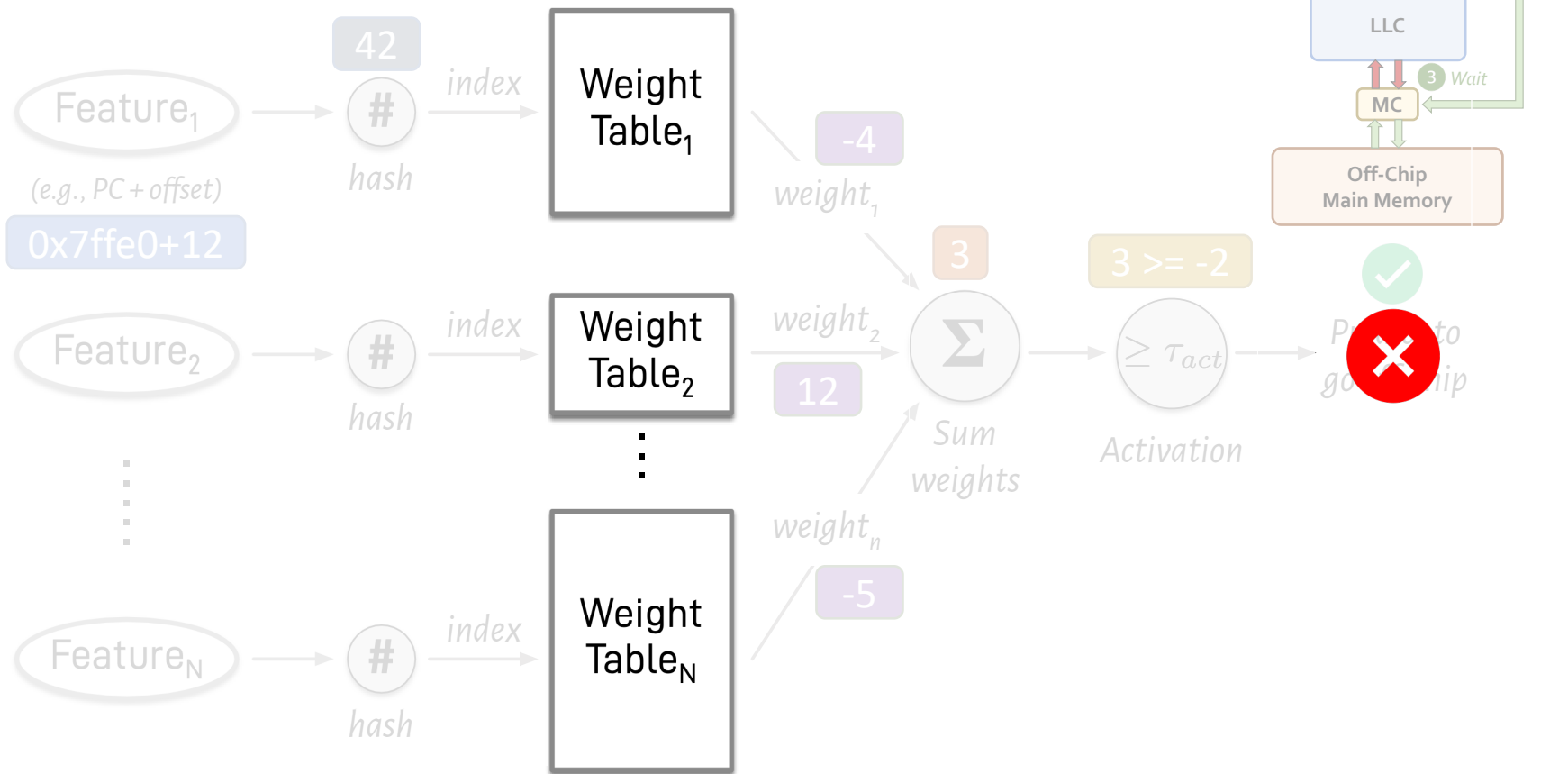
Predicting using POPET

- Using simple **table lookups**, **addition**, and **comparison**



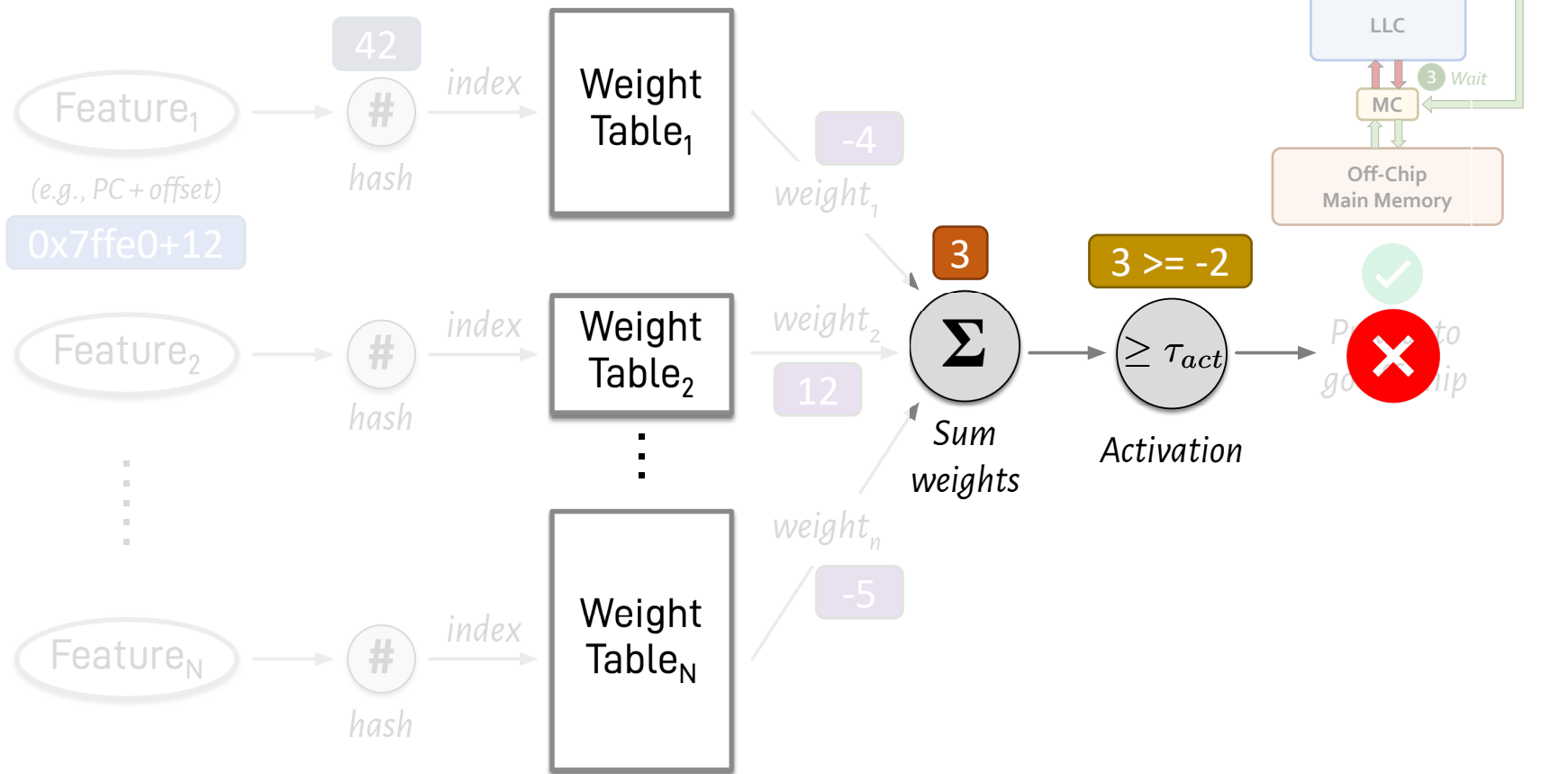
Training POPET

- Using simple **increment** or **decrement** of feature weights



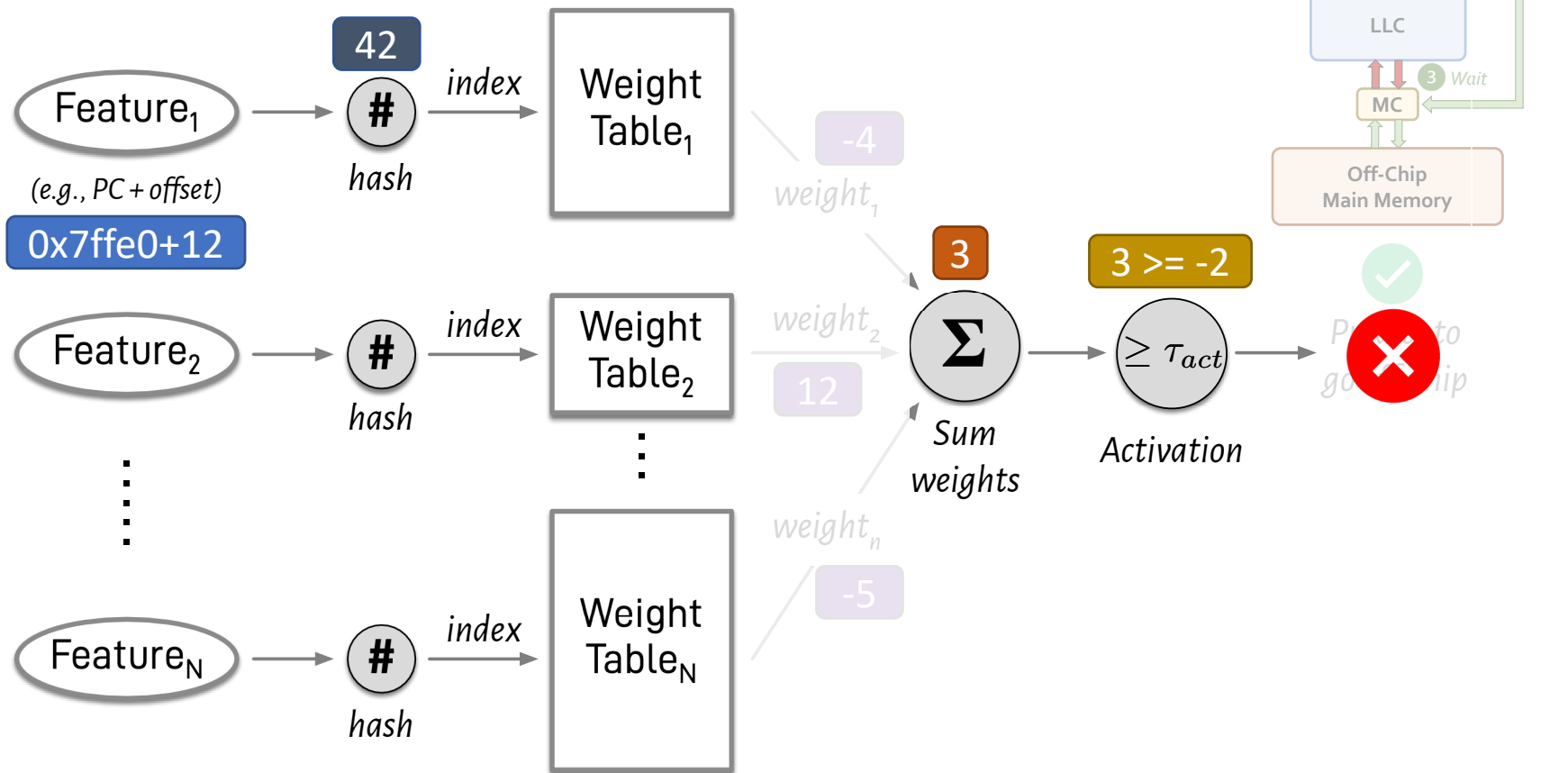
Training POPET

- Using simple **increment** or **decrement** of feature weights



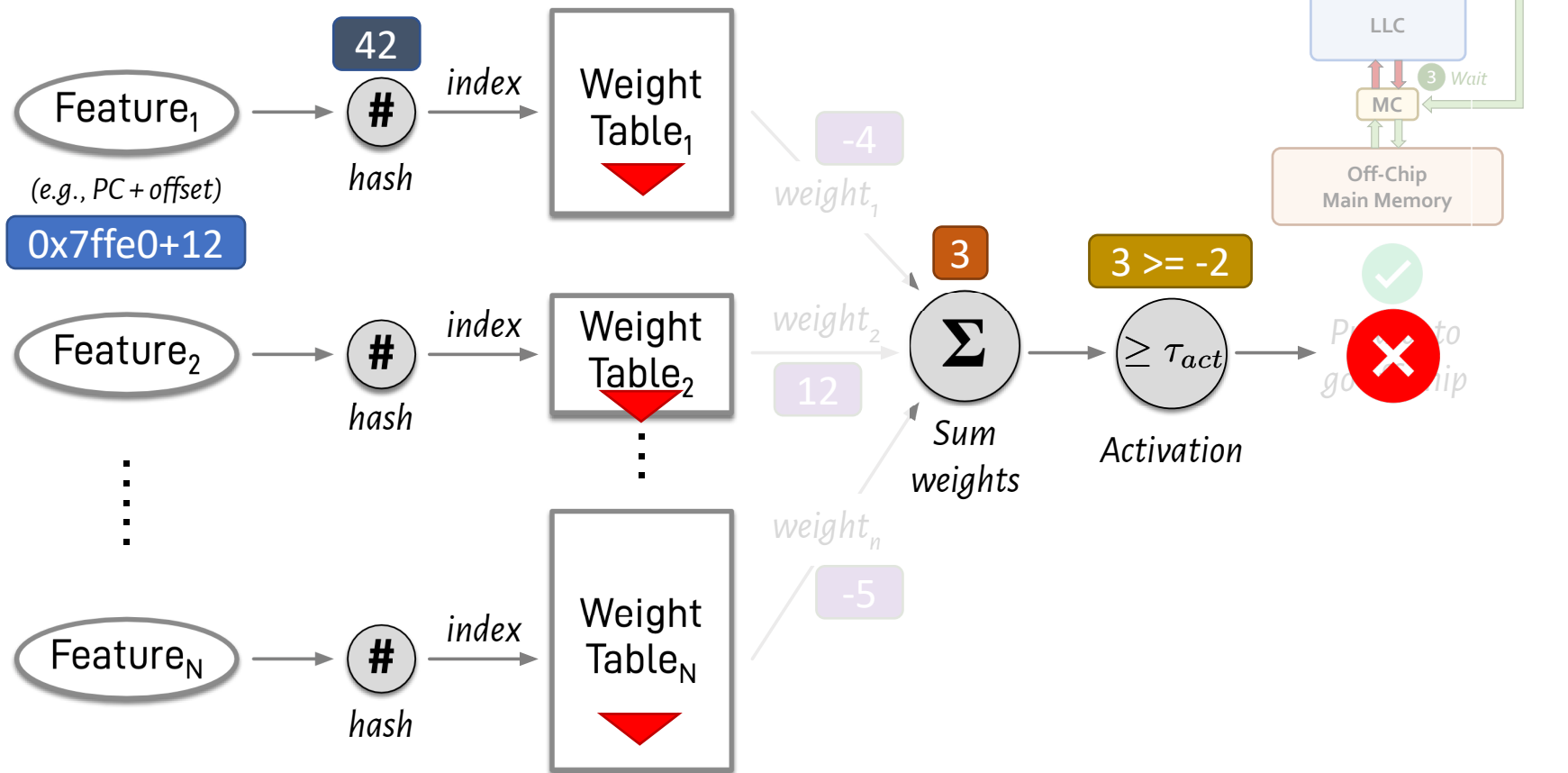
Training POPET

- Using simple **increment** or **decrement** of feature weights



Training POPET

- Using simple **increment** or **decrement** of feature weights



Evaluation

Simulation Methodology

- **ChampSim** trace driven simulator

Simulation Methodology

- **ChampSim** trace driven simulator
- **110 single-core** memory-intensive traces
 - SPEC CPU 2006 and 2017
 - PARSEC 2.1
 - Ligra
 - Real-world applications
- **220 eight-core** homogeneous and heterogeneous trace mixes

Simulation Methodology

- **ChampSim** trace driven simulator
- **110 single-core** memory-intensive traces
 - SPEC CPU 2006 and 2017
 - PARSEC 2.1
 - Ligra
 - Real-world applications
- **220 eight-core** homogeneous and heterogeneous trace mixes

LLC Prefetchers

- Pythia [Bera+, MICRO'22]
- Bingo [Bakshalipour+, HPCA'19]
- MLOP [Shakerinava+, 3rd Prefetching Championship'19]
- SPP + Perceptron filter [Bhatia+, ISCA'20]
- SMS [Somogyi+, ISCA'06]

Simulation Methodology

- **ChampSim** trace driven simulator
- **110 single-core** memory-intensive traces
 - SPEC CPU 2006 and 2017
 - PARSEC 2.1
 - Ligra
 - Real-world applications
- **220 eight-core** homogeneous and heterogeneous trace mixes

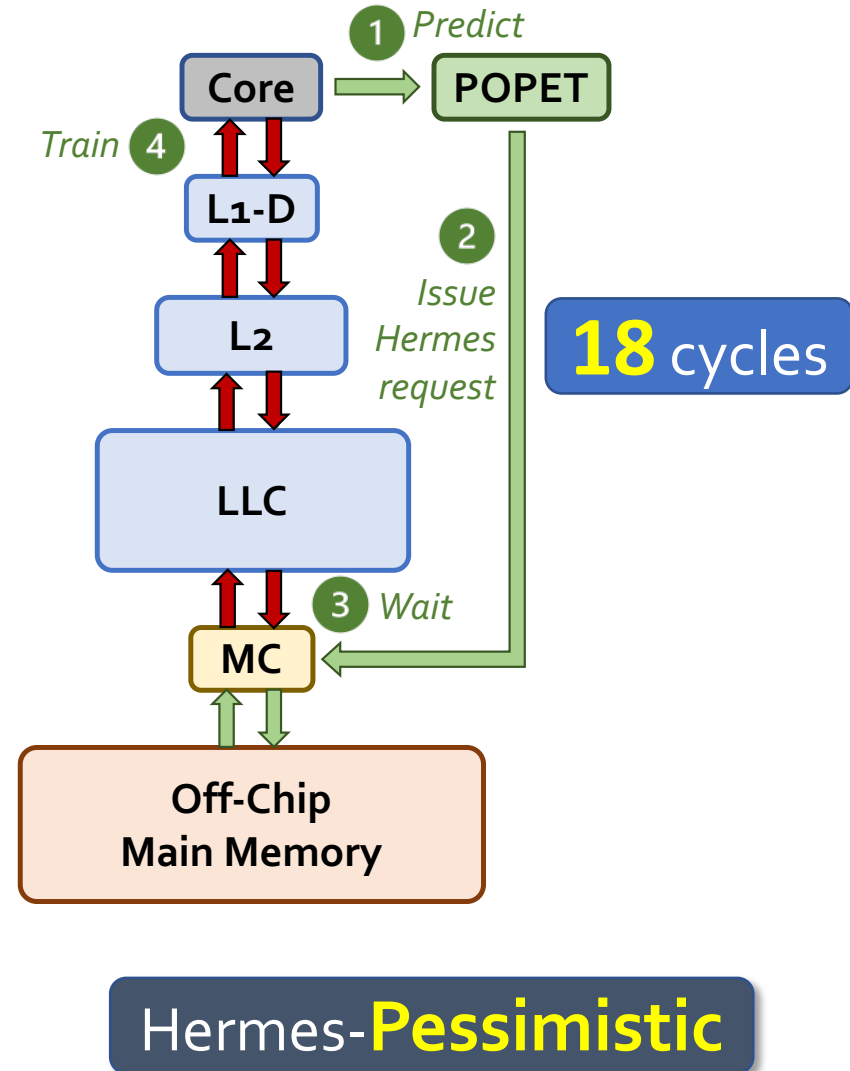
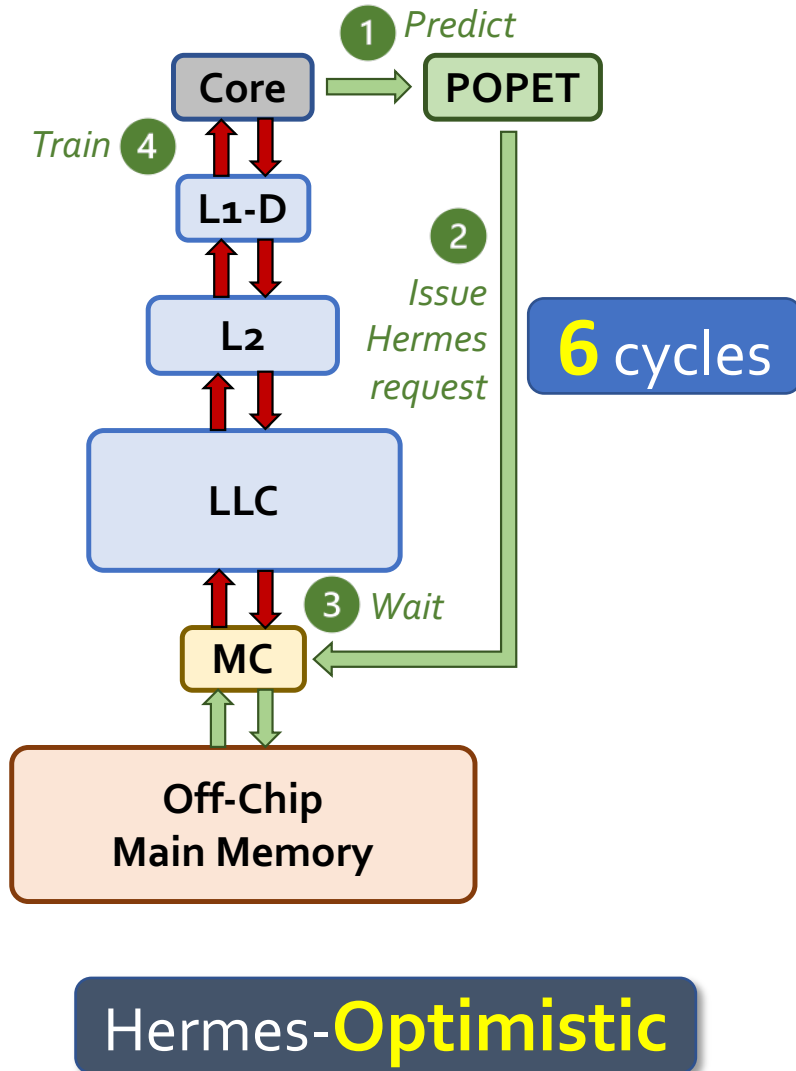
LLC Prefetchers

- Pythia [Bera+, MICRO'22]
- Bingo [Bakshalipour+, HPCA'19]
- MLOP [Shakerinava+, 3rd Prefetching Championship'19]
- SPP + Perceptron filter [Bhatia+, ISCA'20]
- SMS [Somogyi+, ISCA'06]

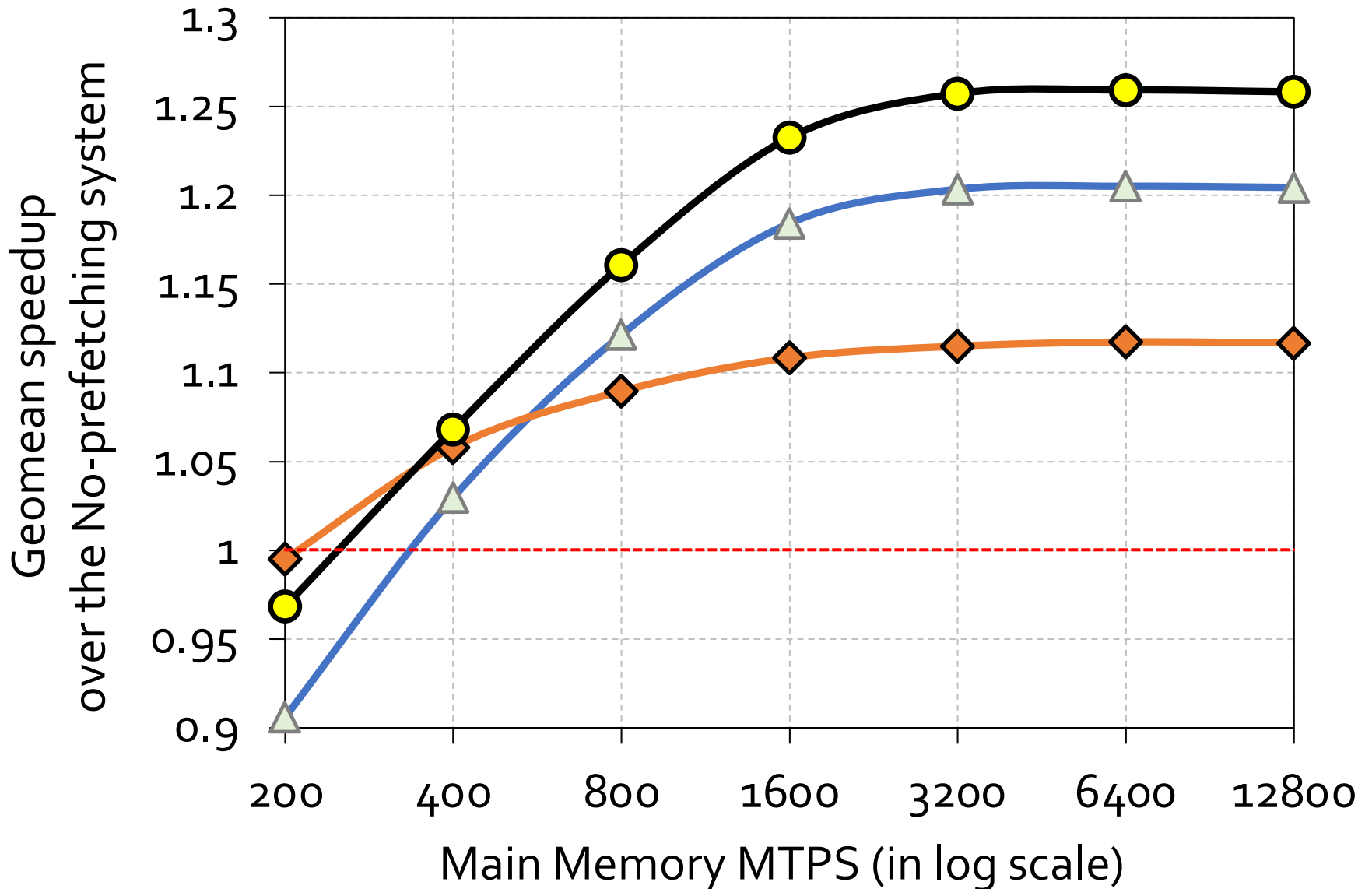
Off-Chip Predictors

- HMP [Yoaz+, ISCA'99]
- Address Tag-Tracking based Predictor (TTP)

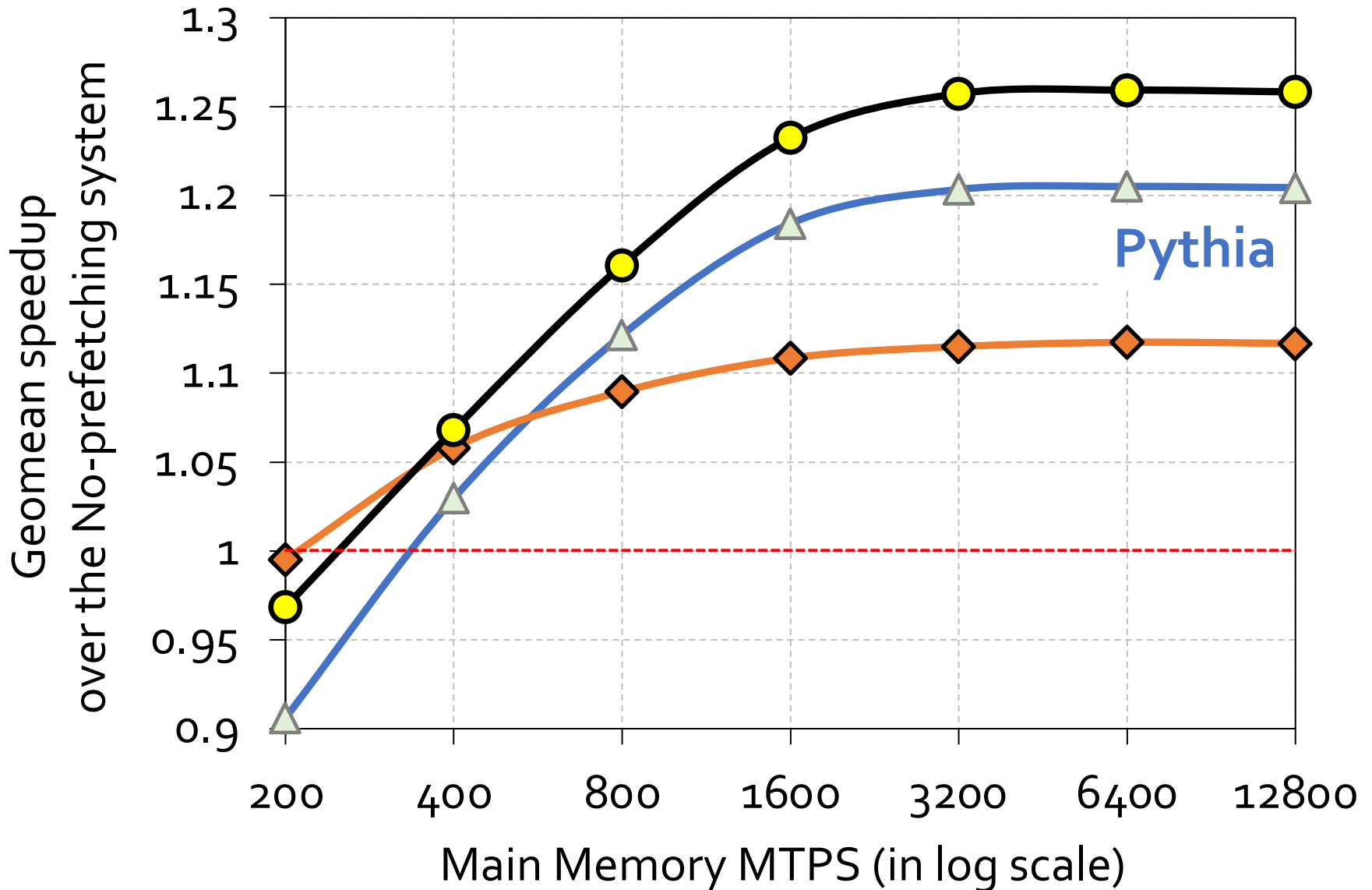
Simulation Methodology



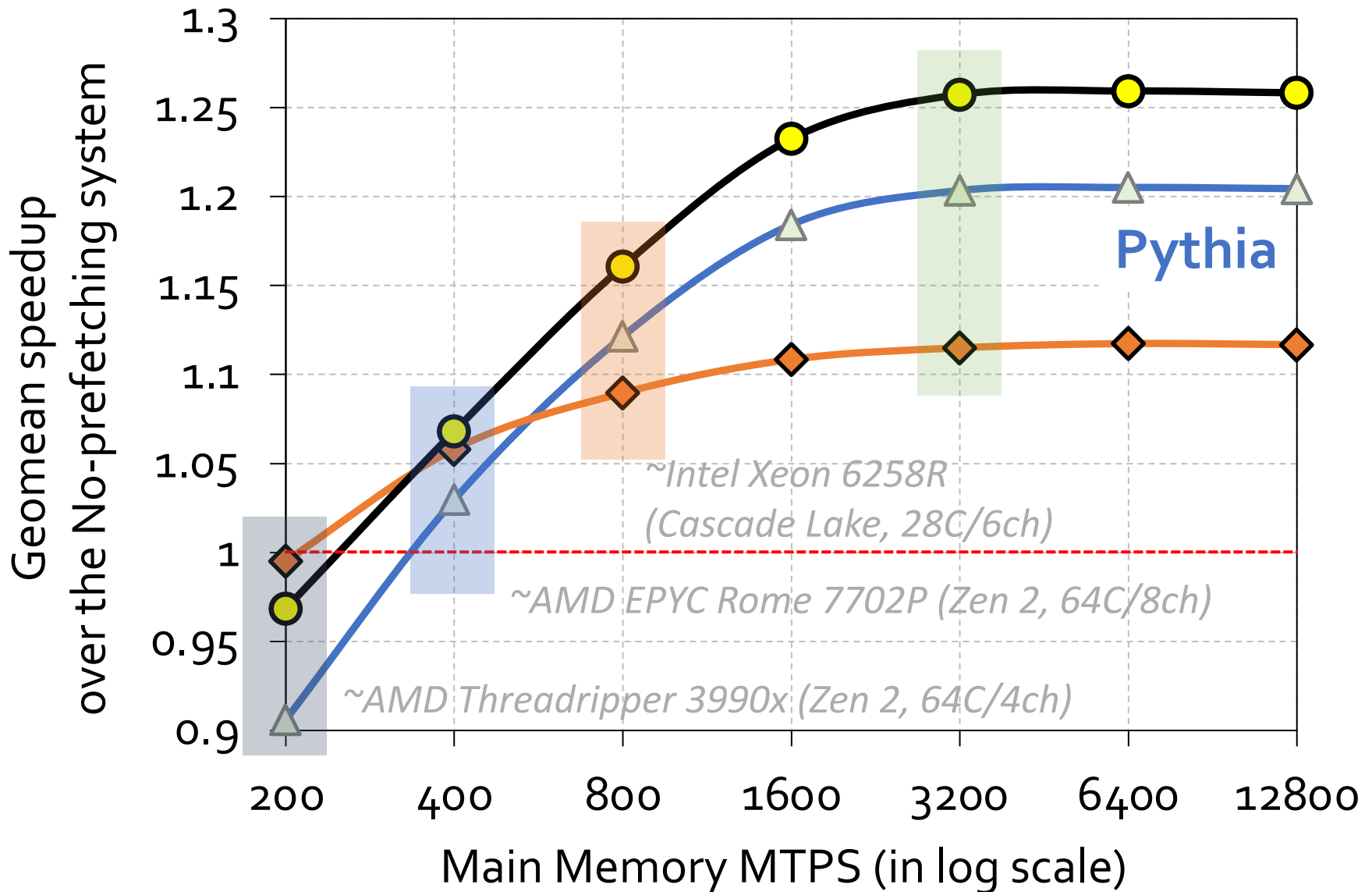
Performance with Varying Memory Bandwidth



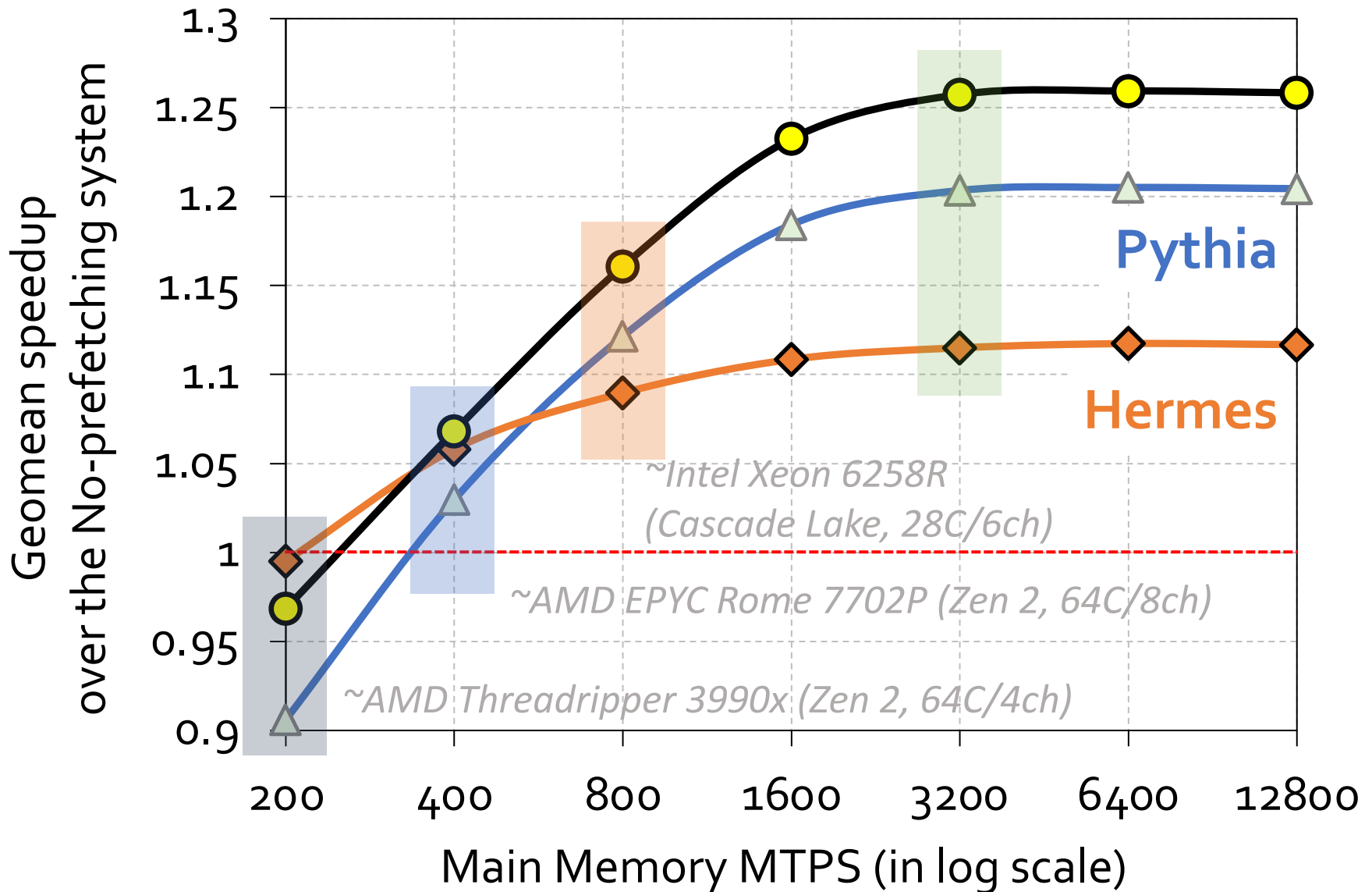
Performance with Varying Memory Bandwidth



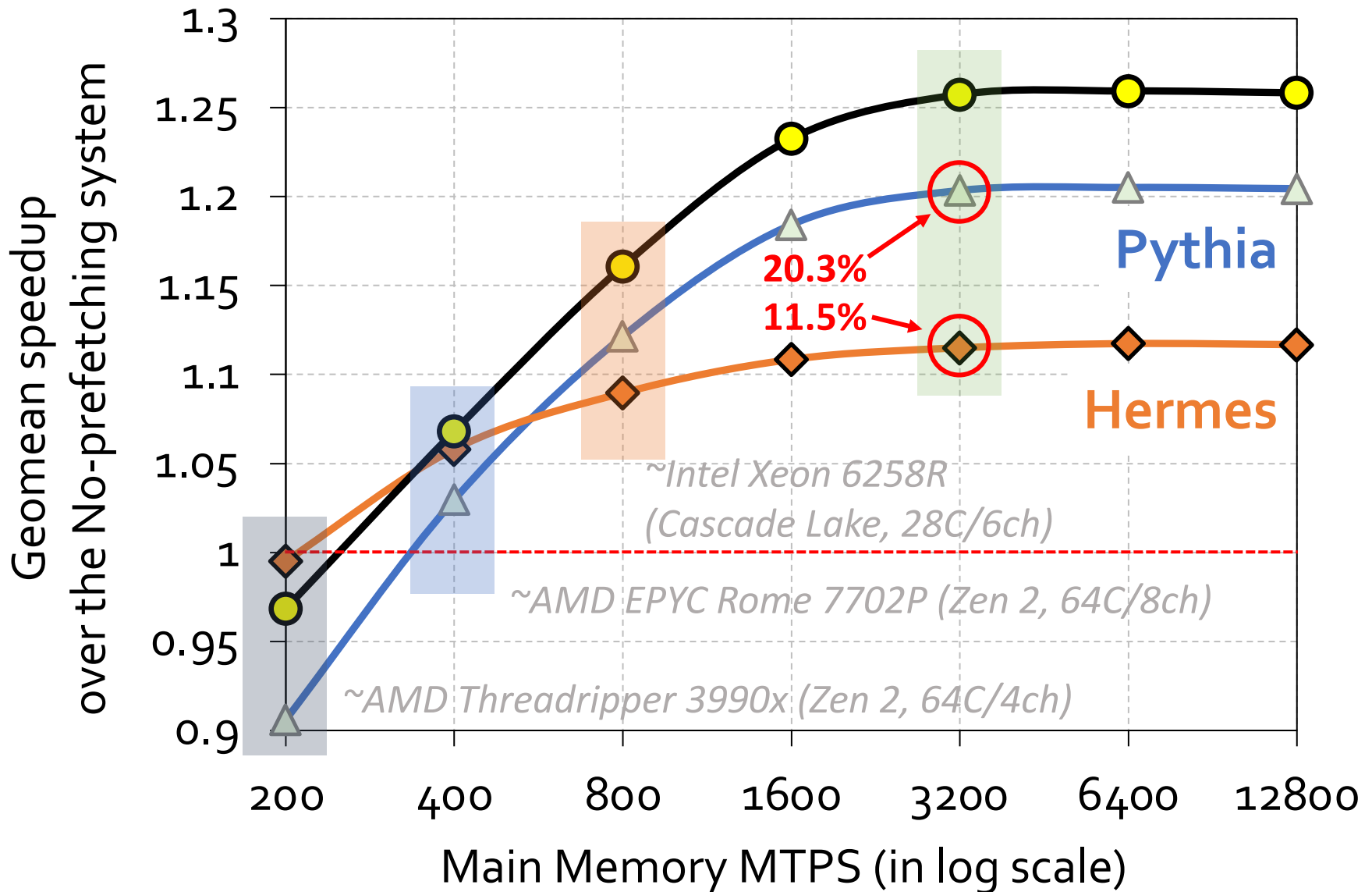
Performance with Varying Memory Bandwidth



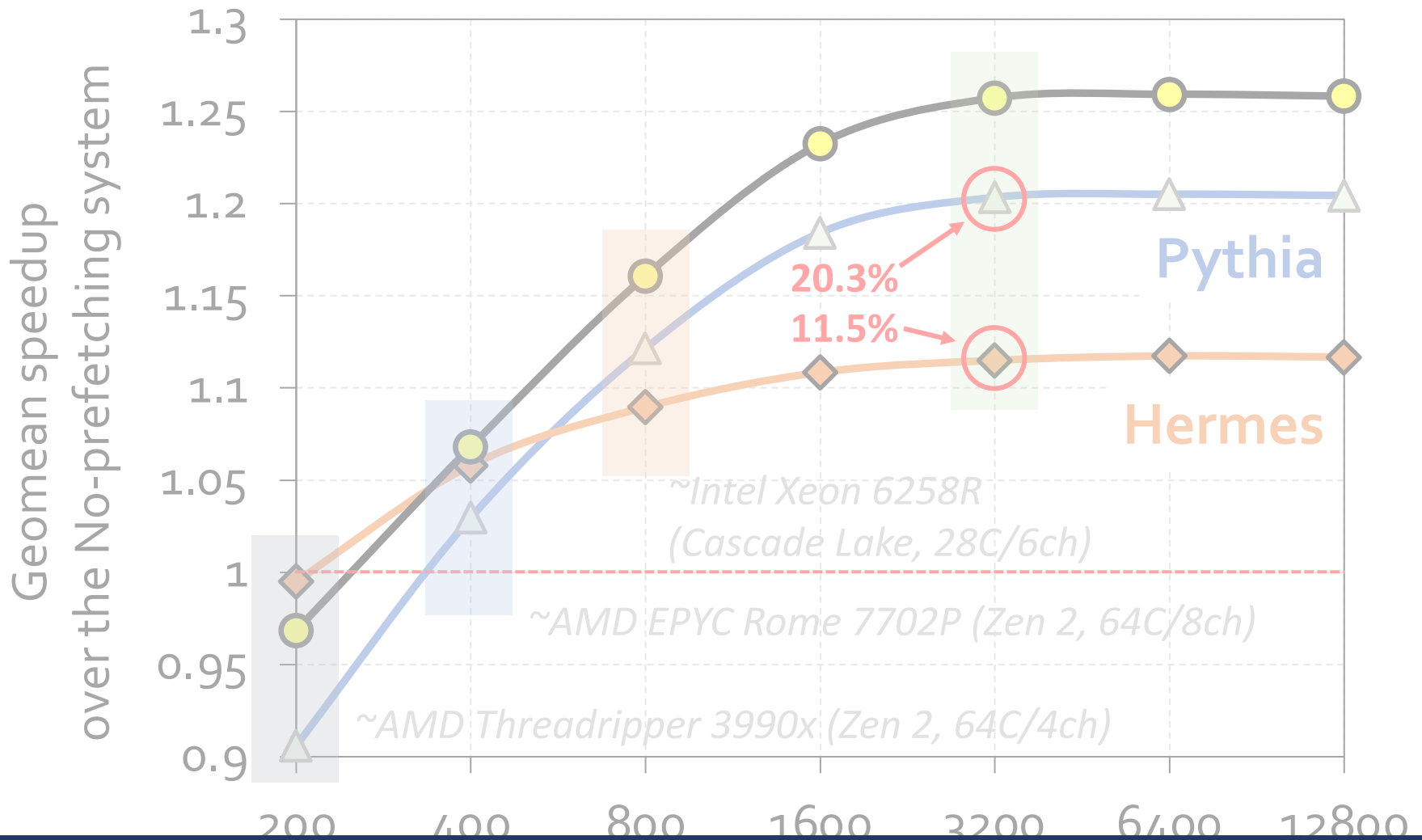
Performance with Varying Memory Bandwidth



Performance with Varying Memory Bandwidth

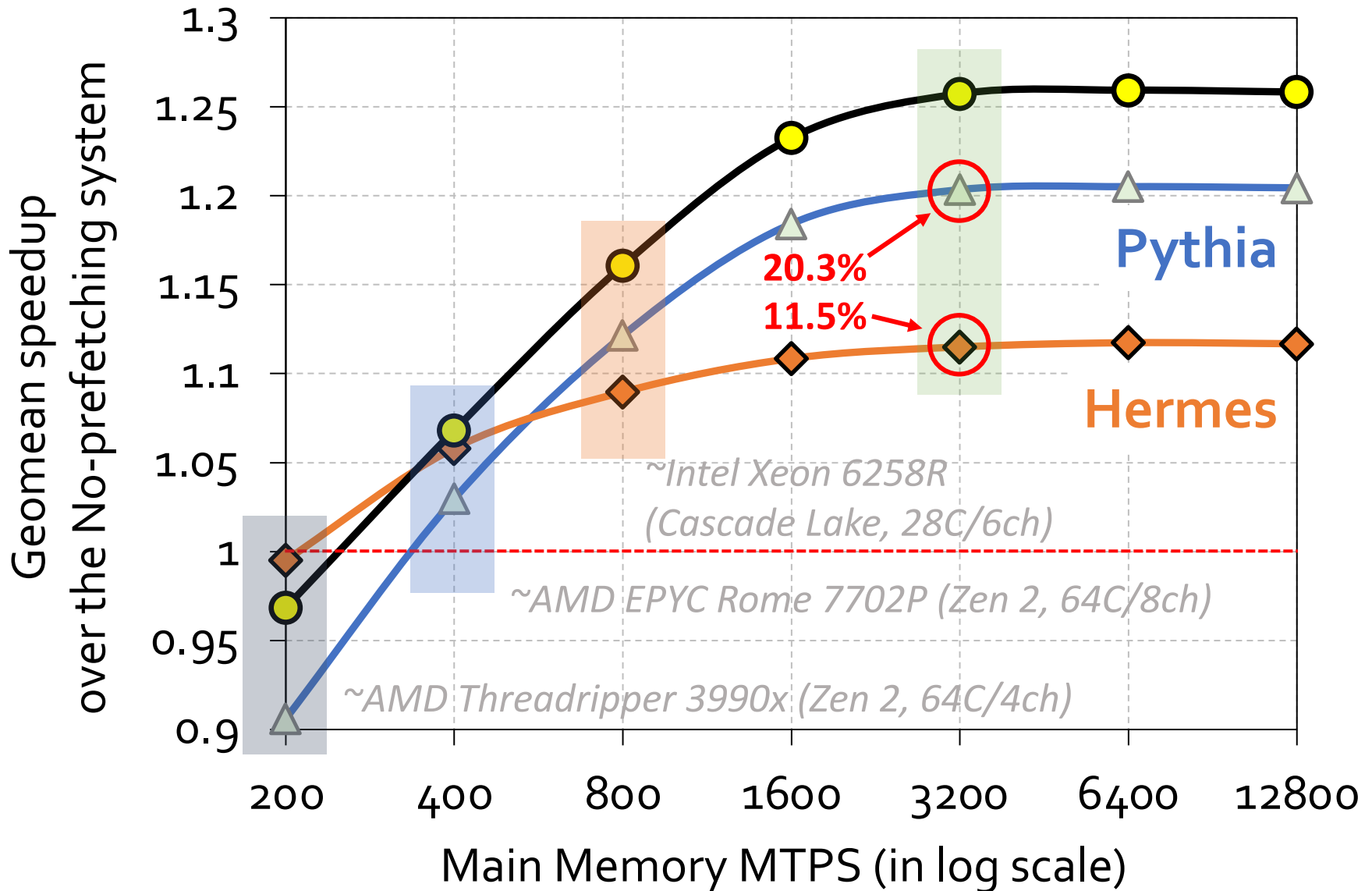


Performance with Varying Memory Bandwidth

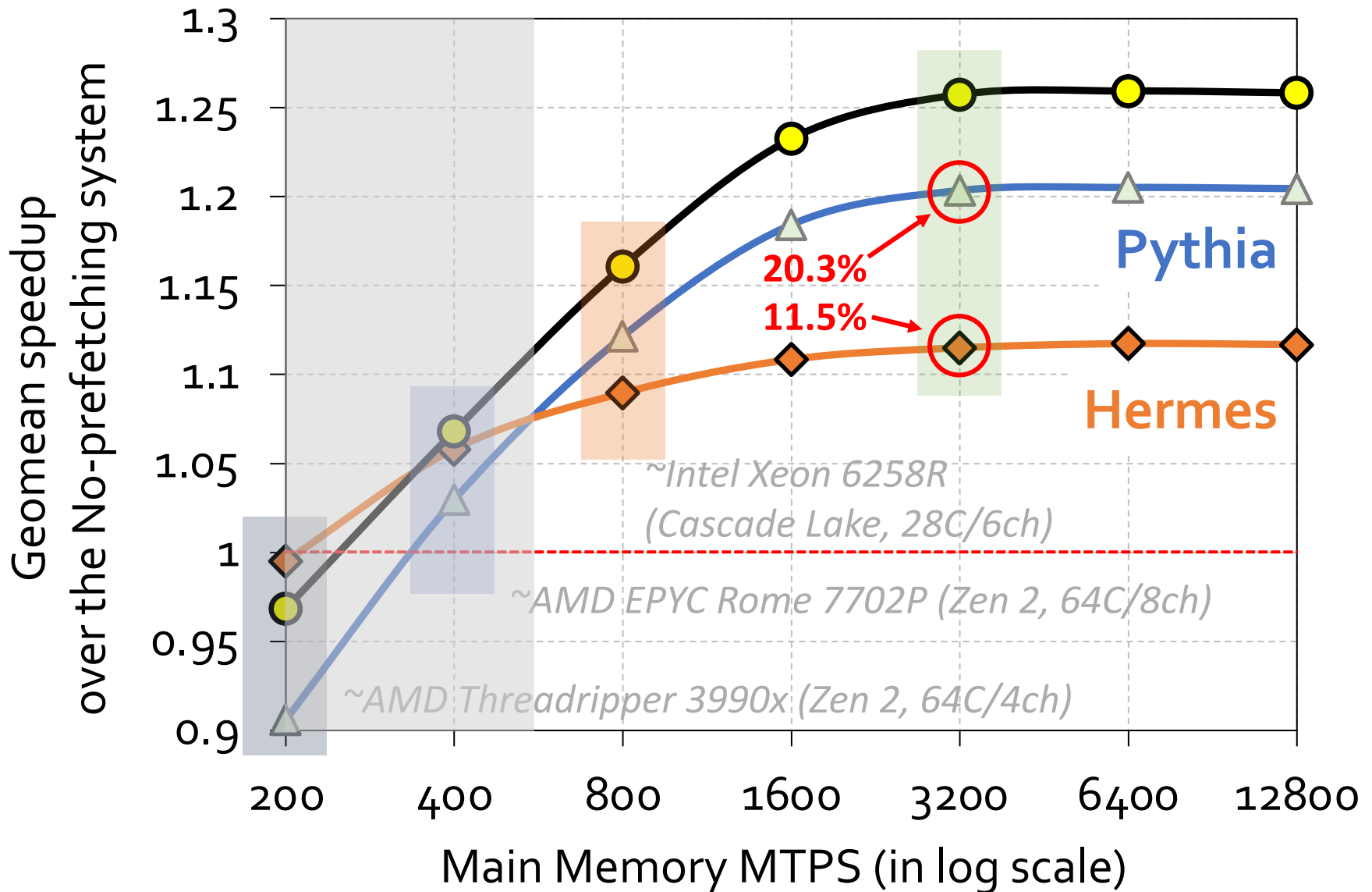


Hermes alone provides nearly **50%** performance of **Pythia**

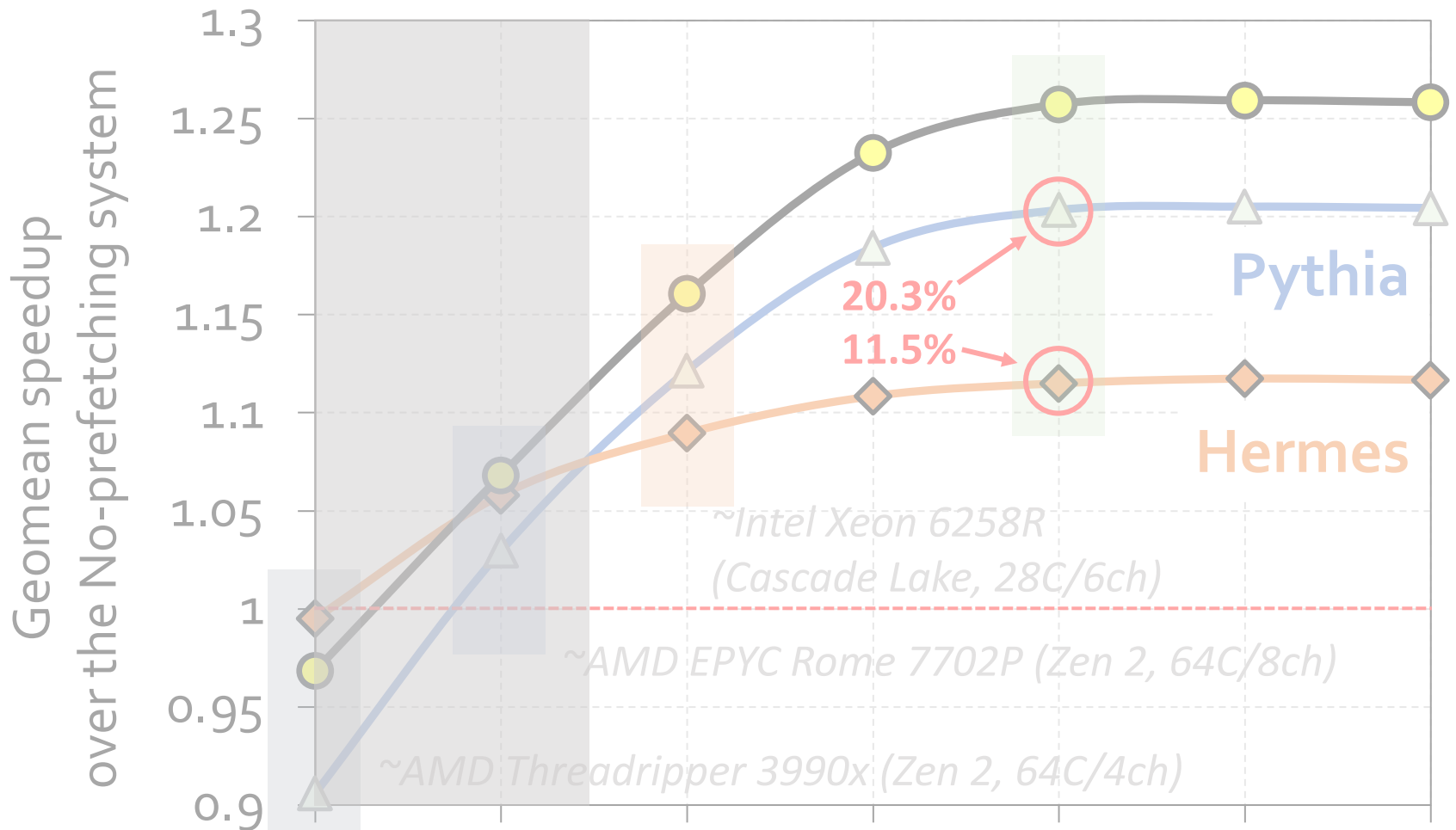
Performance with Varying Memory Bandwidth



Performance with Varying Memory Bandwidth

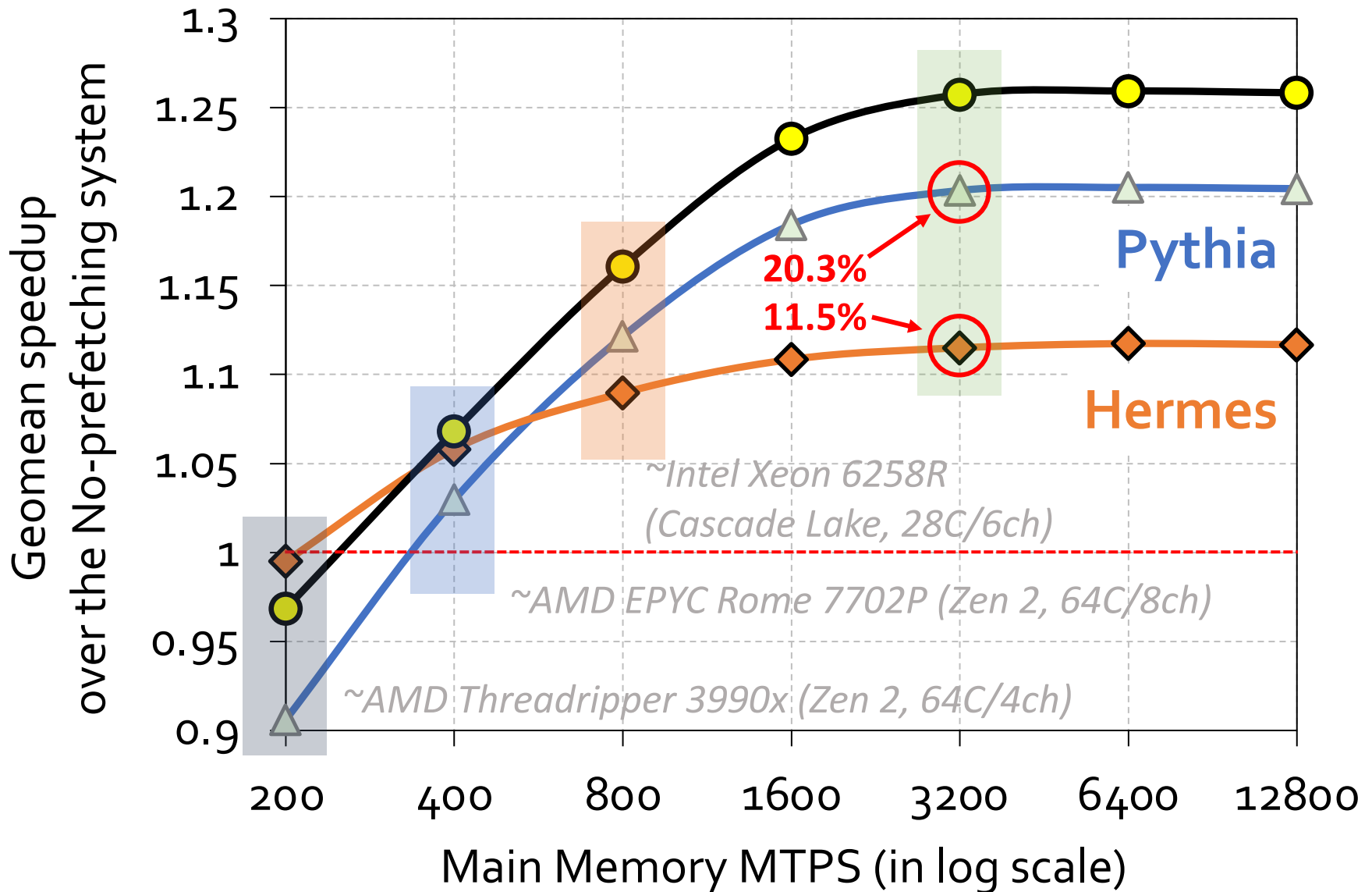


Performance with Varying Memory Bandwidth

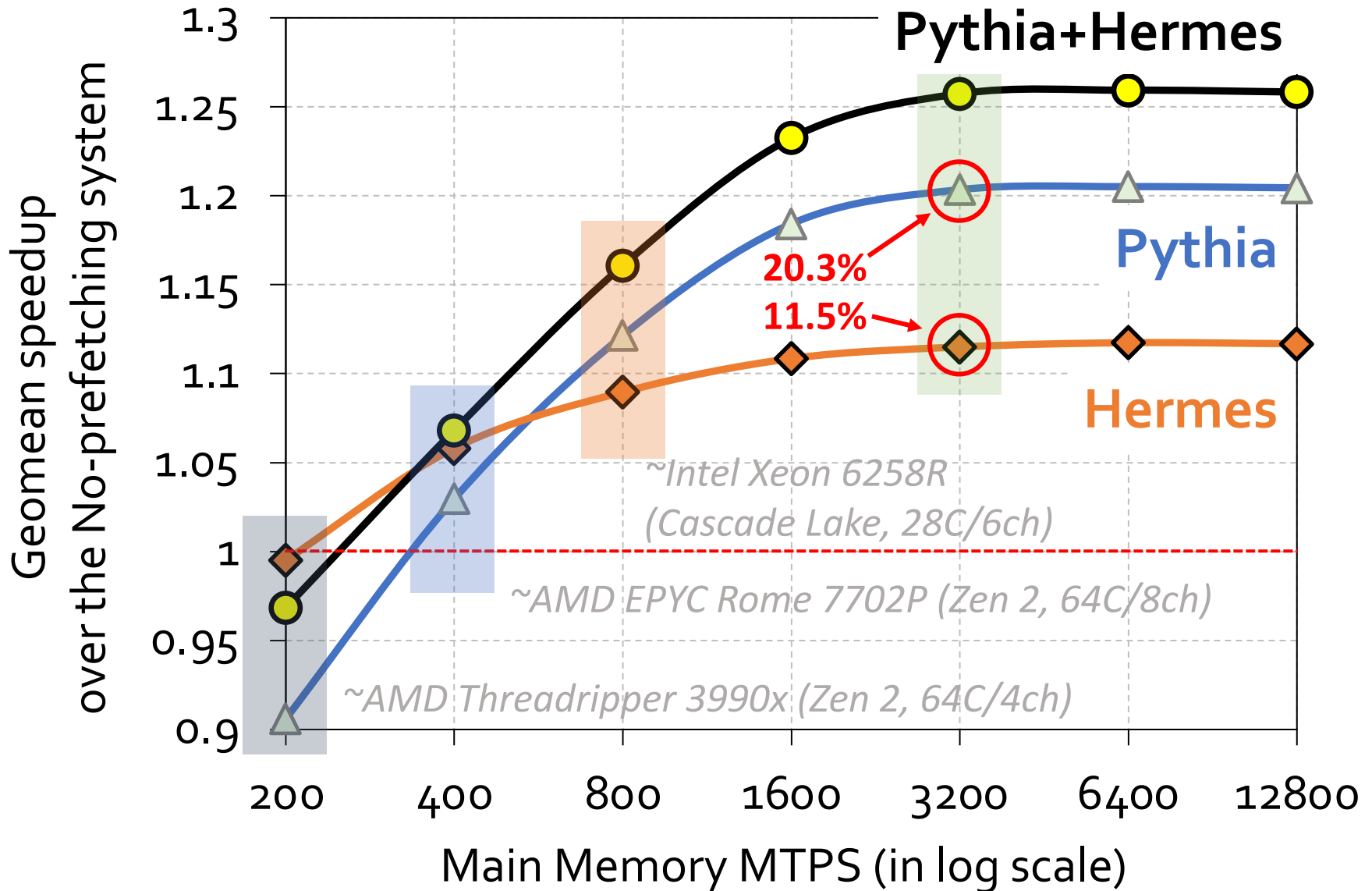


In **extremely bandwidth-constrained** configurations, Hermes alone outperforms Pythia

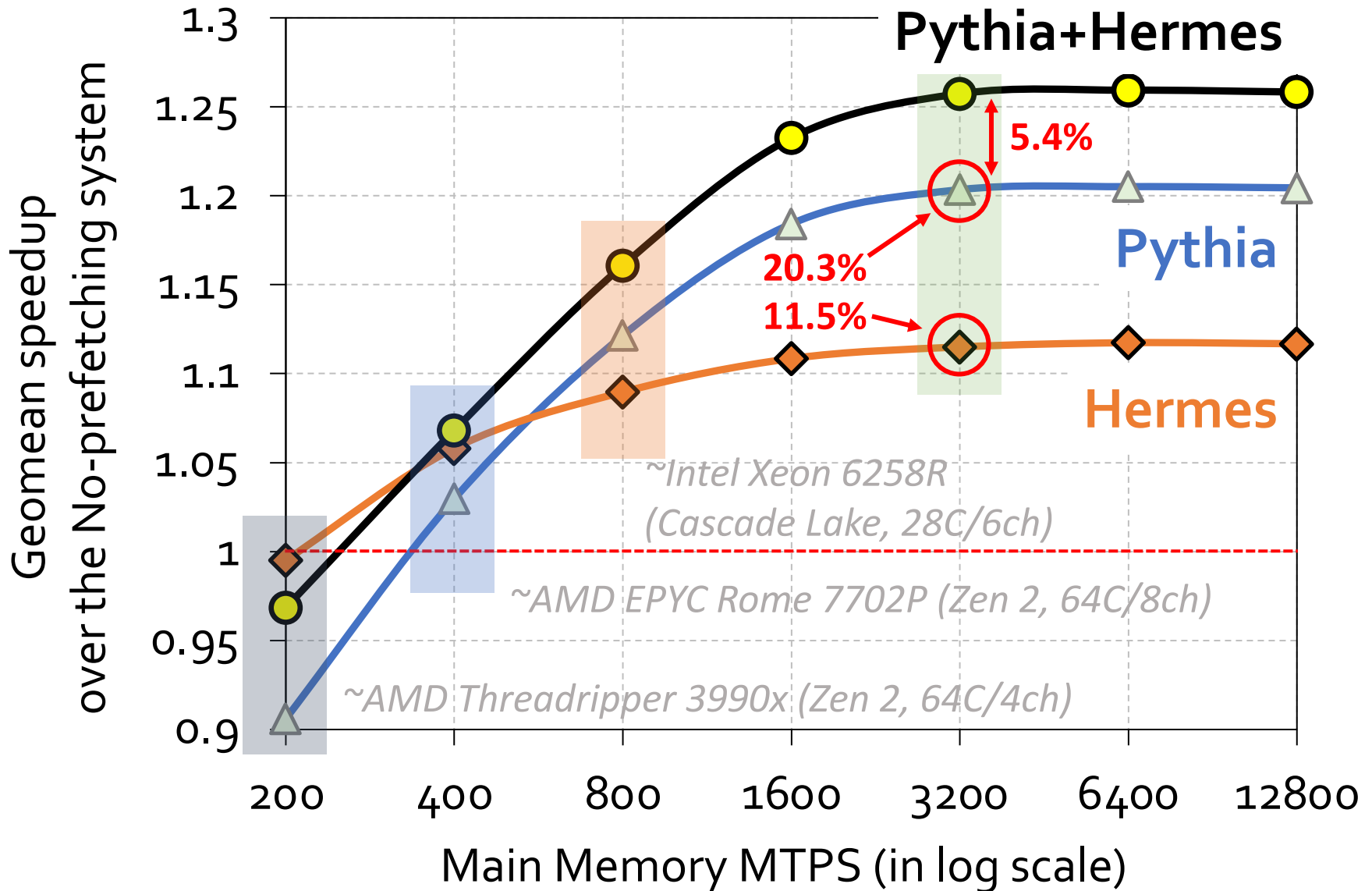
Performance with Varying Memory Bandwidth



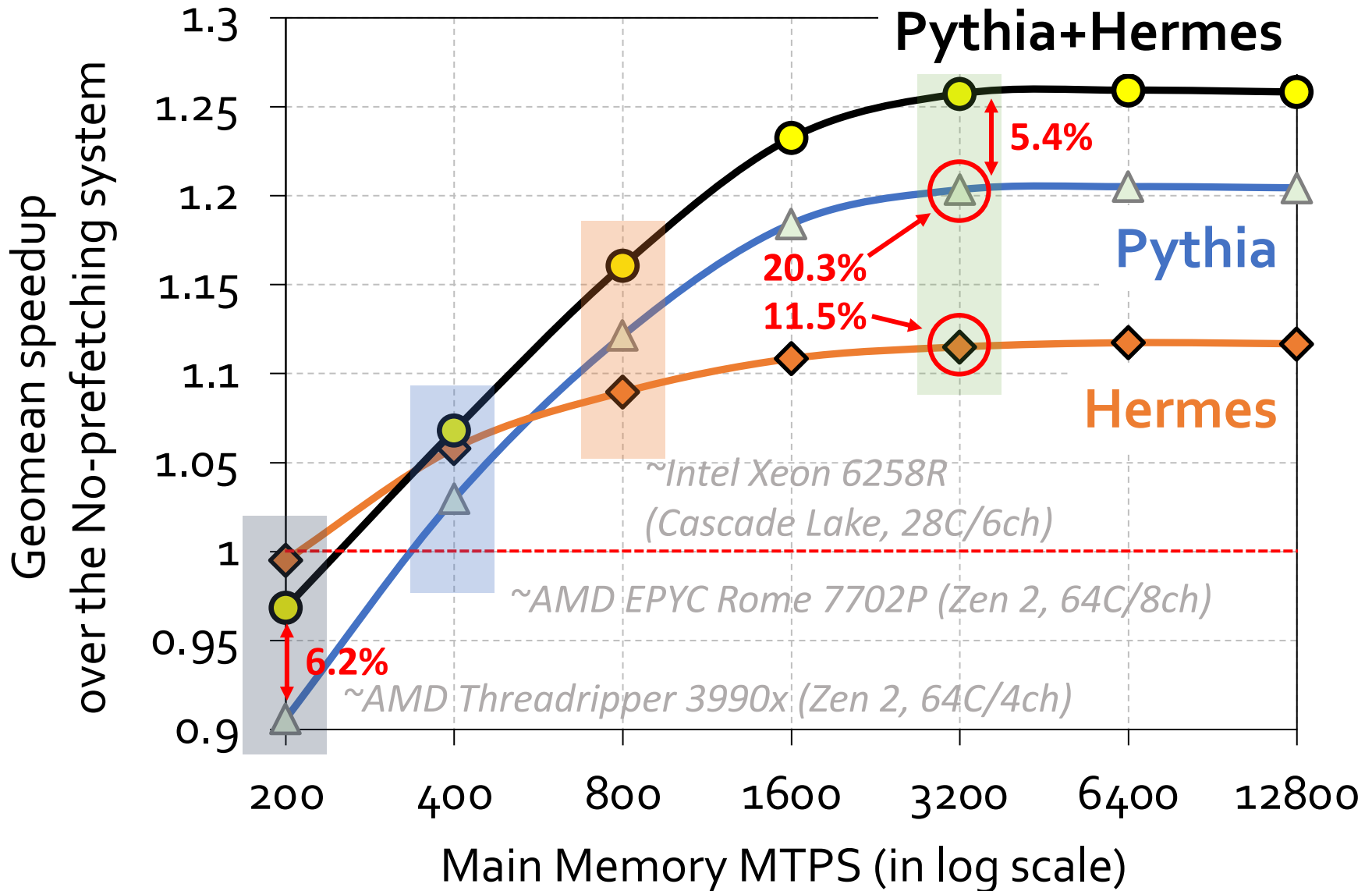
Performance with Varying Memory Bandwidth



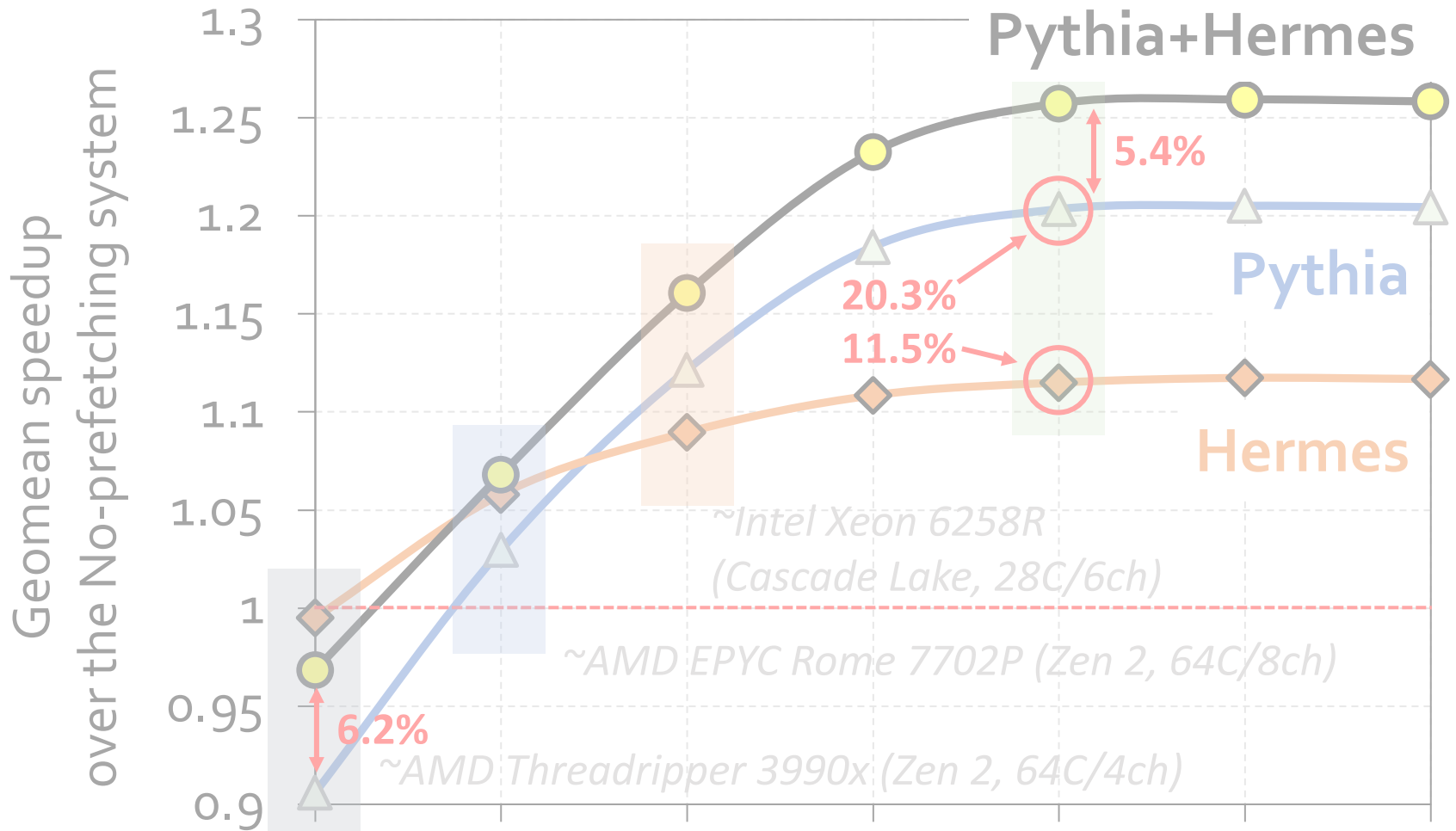
Performance with Varying Memory Bandwidth



Performance with Varying Memory Bandwidth

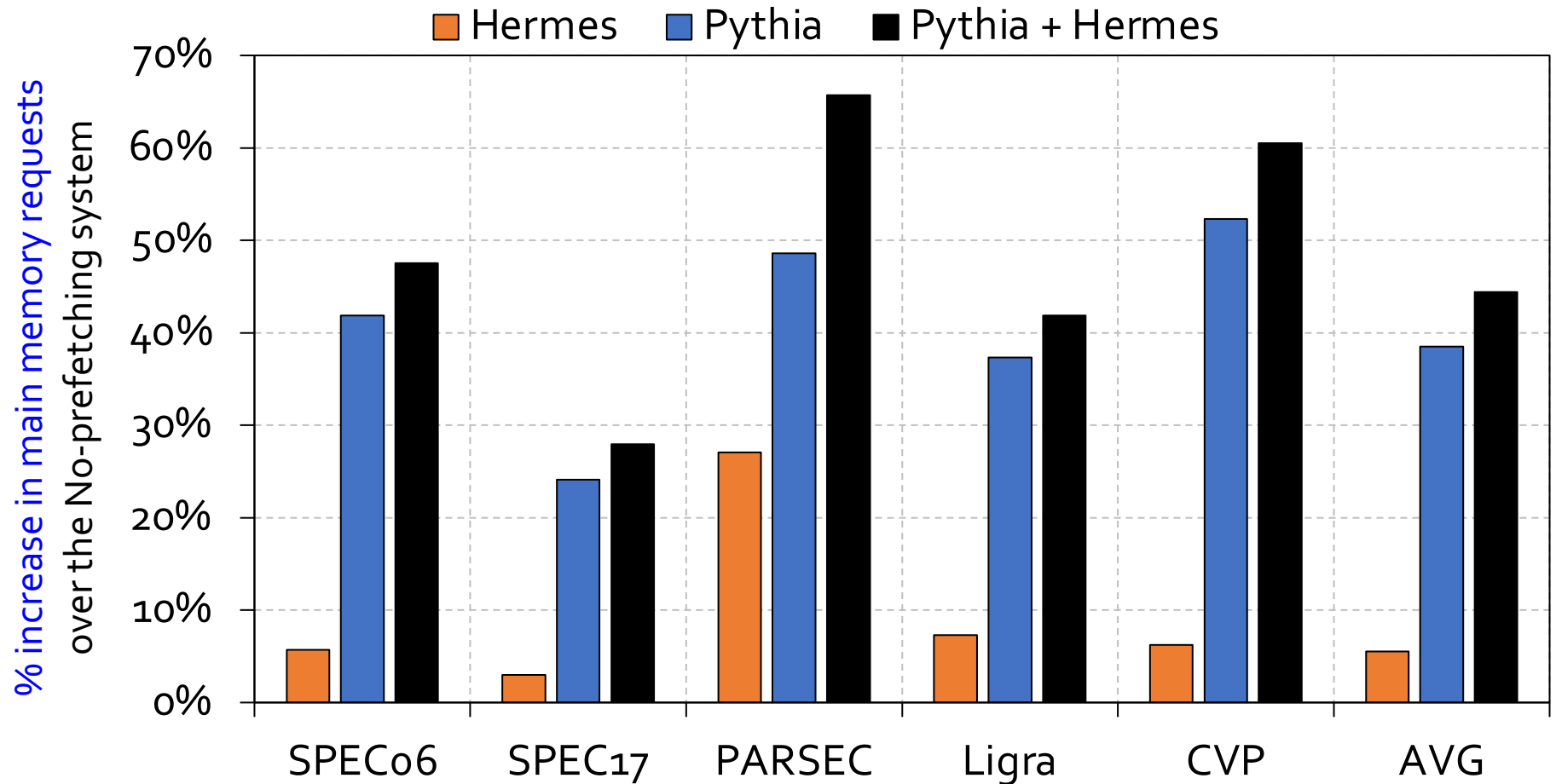


Performance with Varying Memory Bandwidth

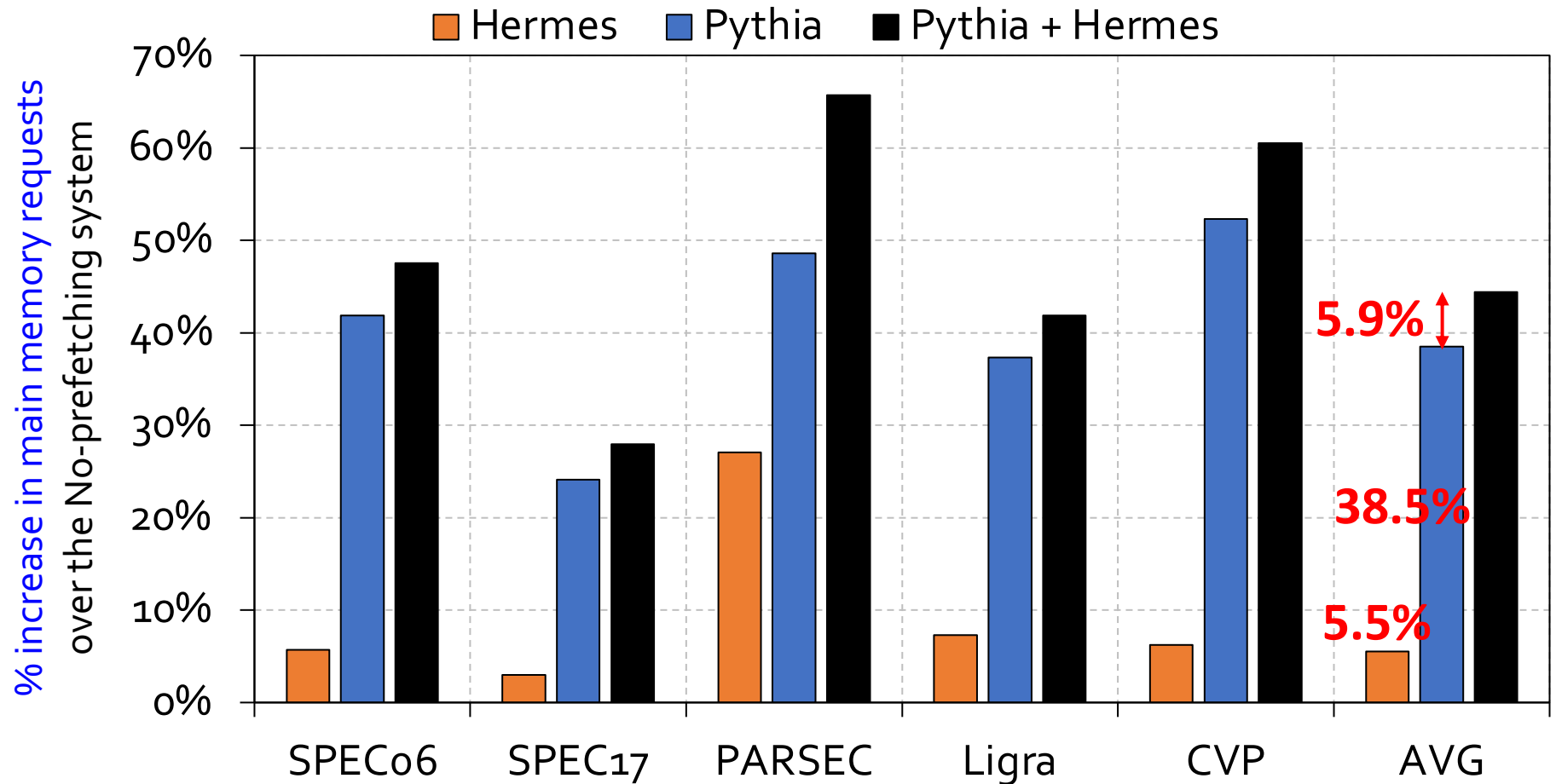


Hermes+Pythia outperforms Pythia
in a wide range of bandwidth configurations

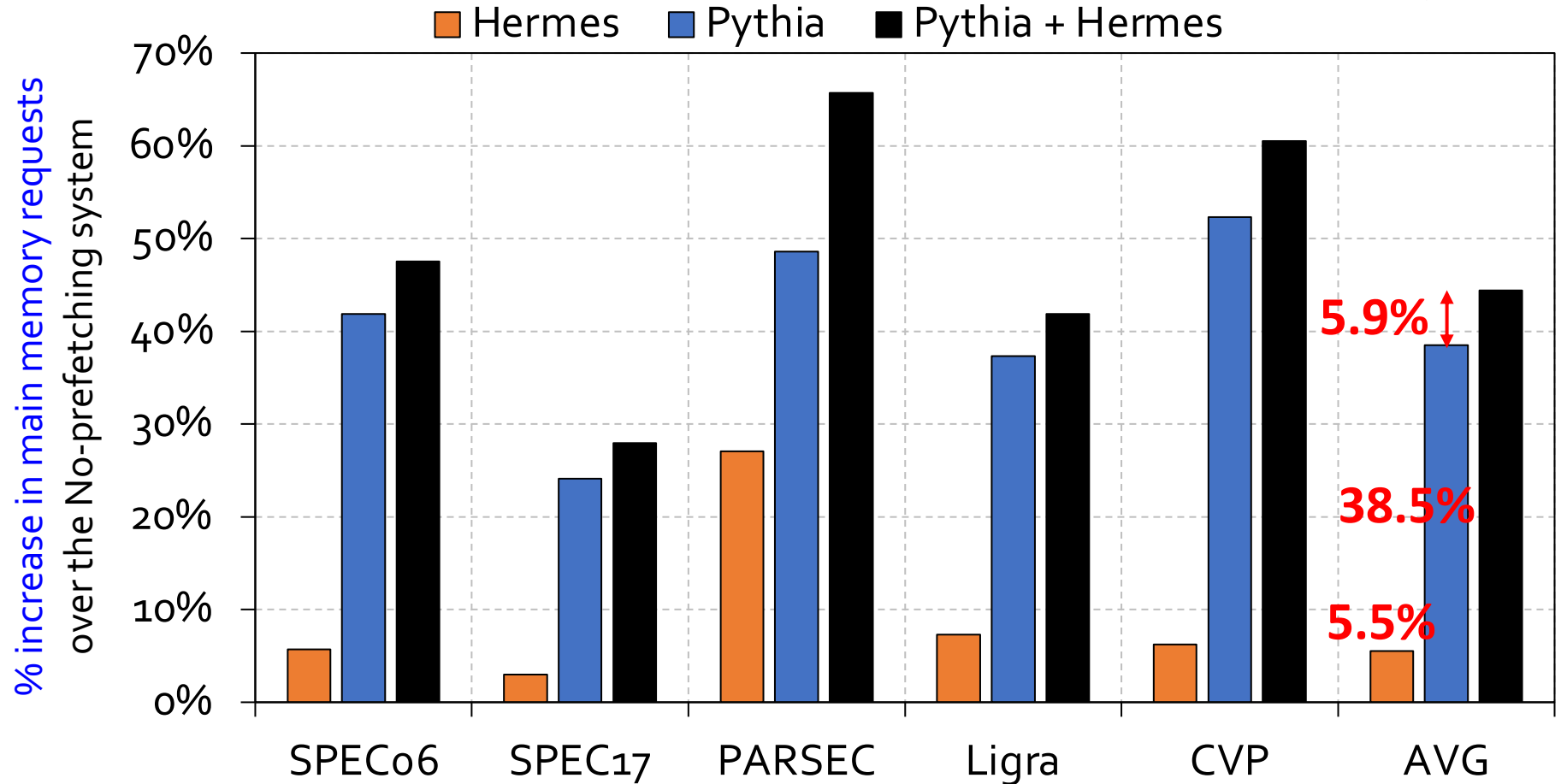
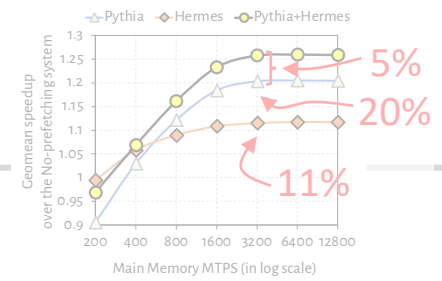
Bandwidth Efficiency



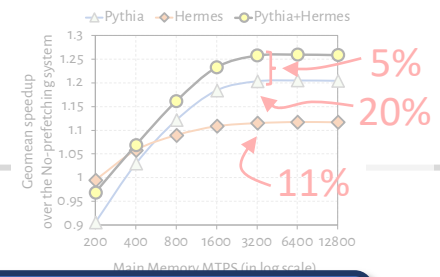
Bandwidth Efficiency



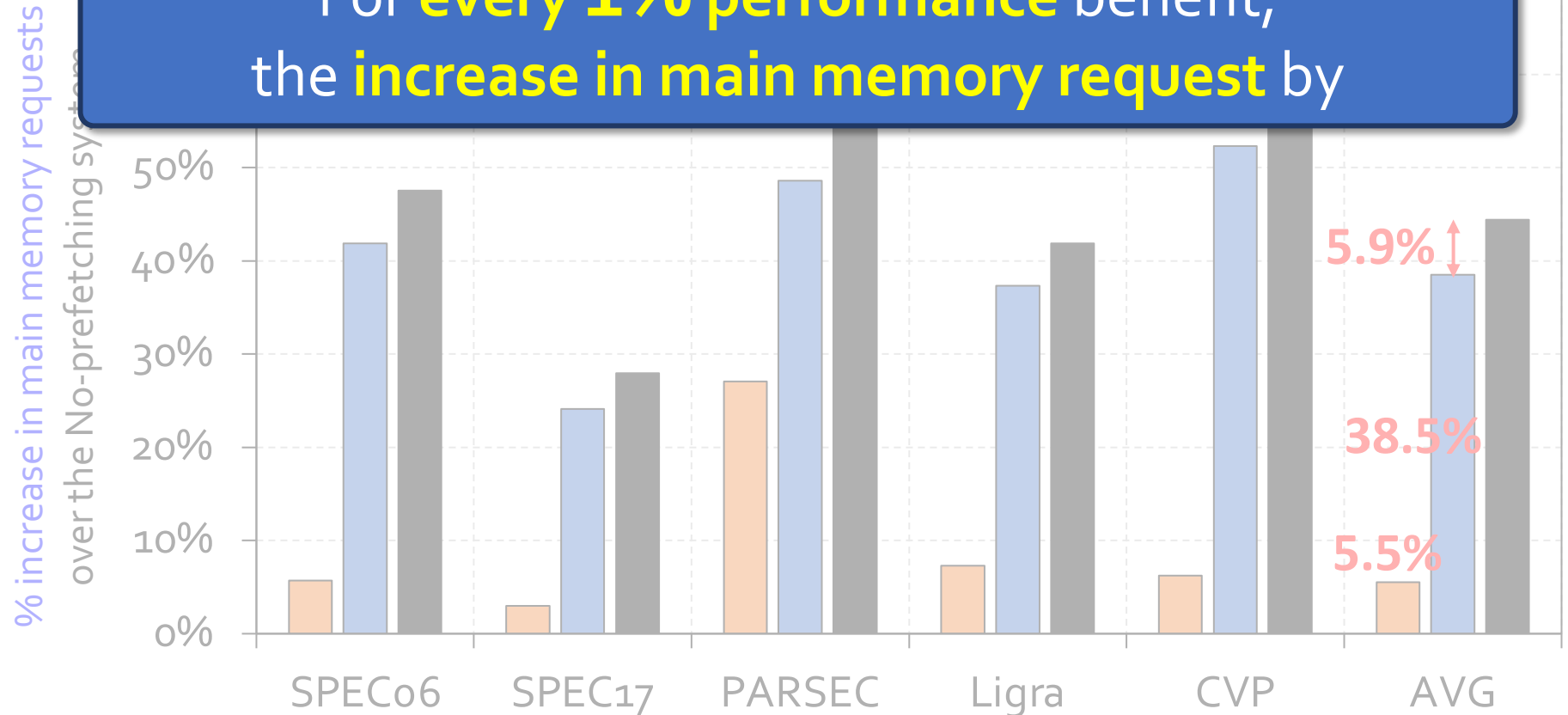
Bandwidth Efficiency



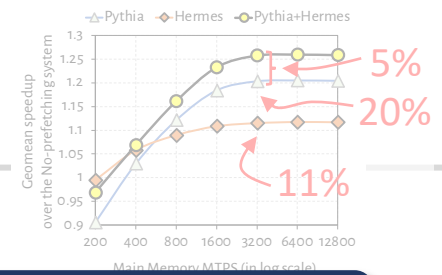
Bandwidth Efficiency



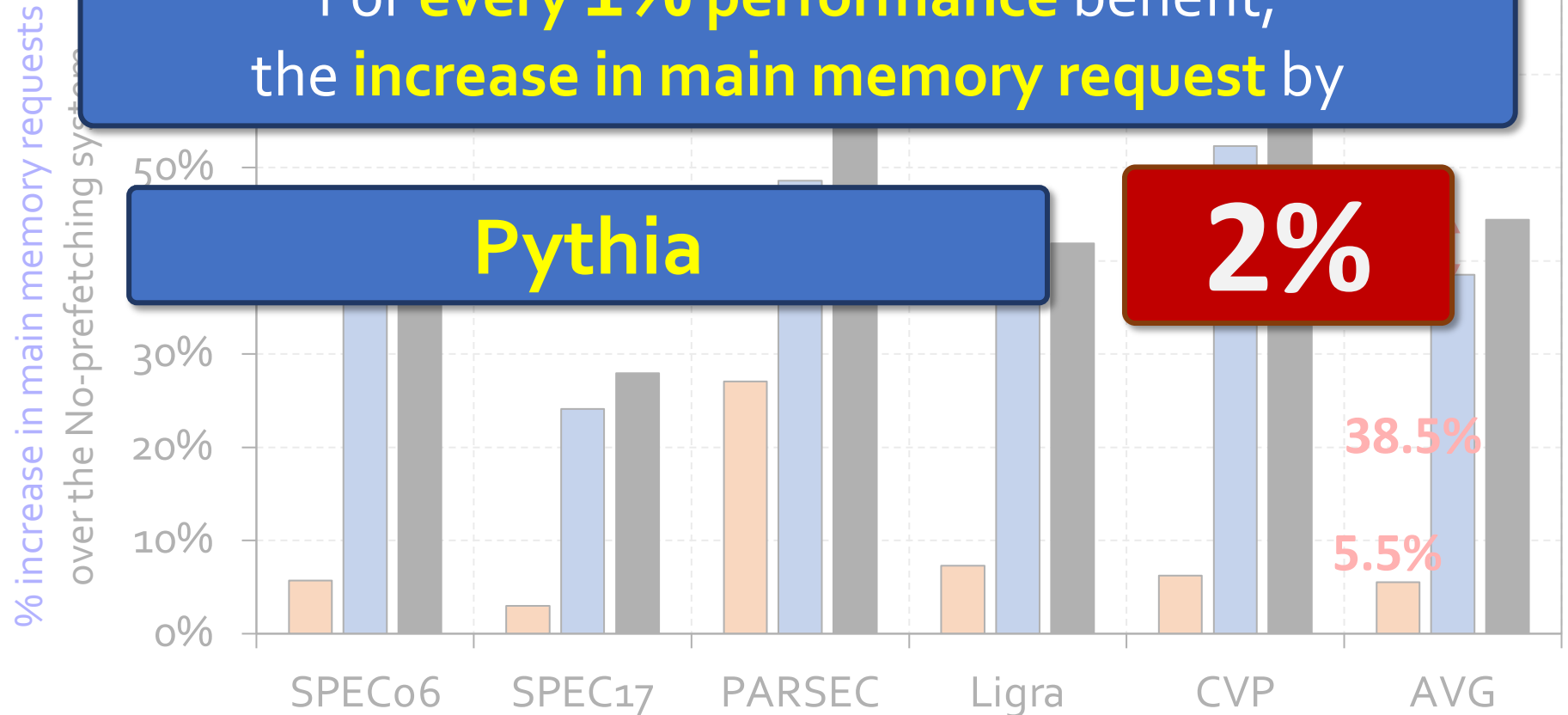
For **every 1%** performance benefit, the **increase in main memory request** by



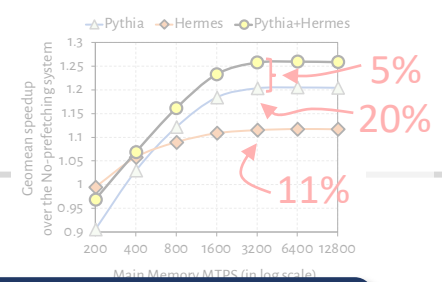
Bandwidth Efficiency



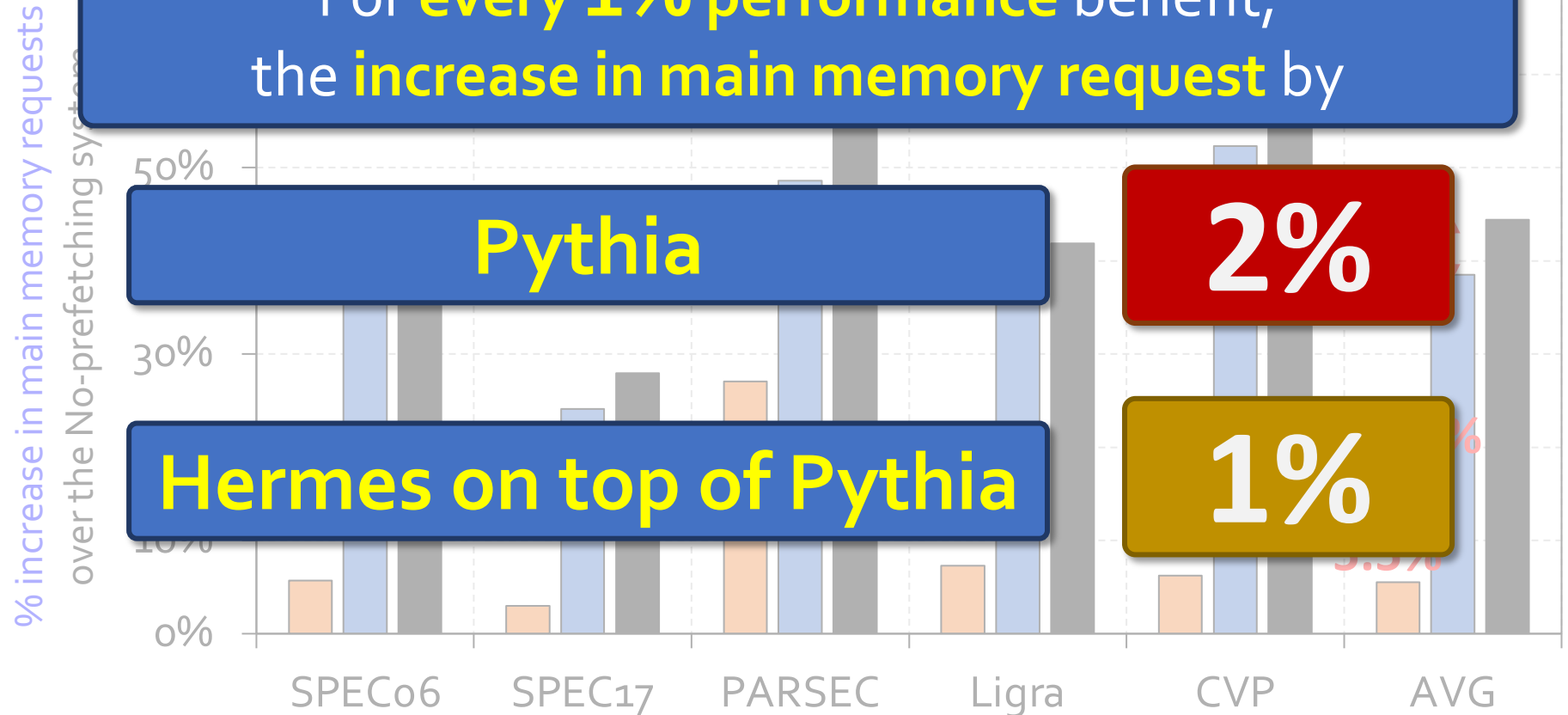
For **every 1%** performance benefit,
the **increase in main memory request** by



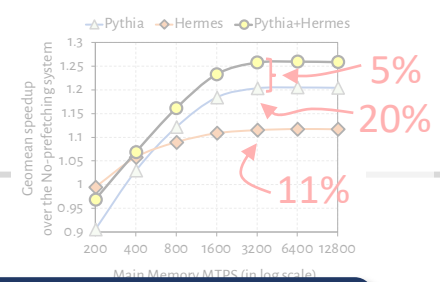
Bandwidth Efficiency



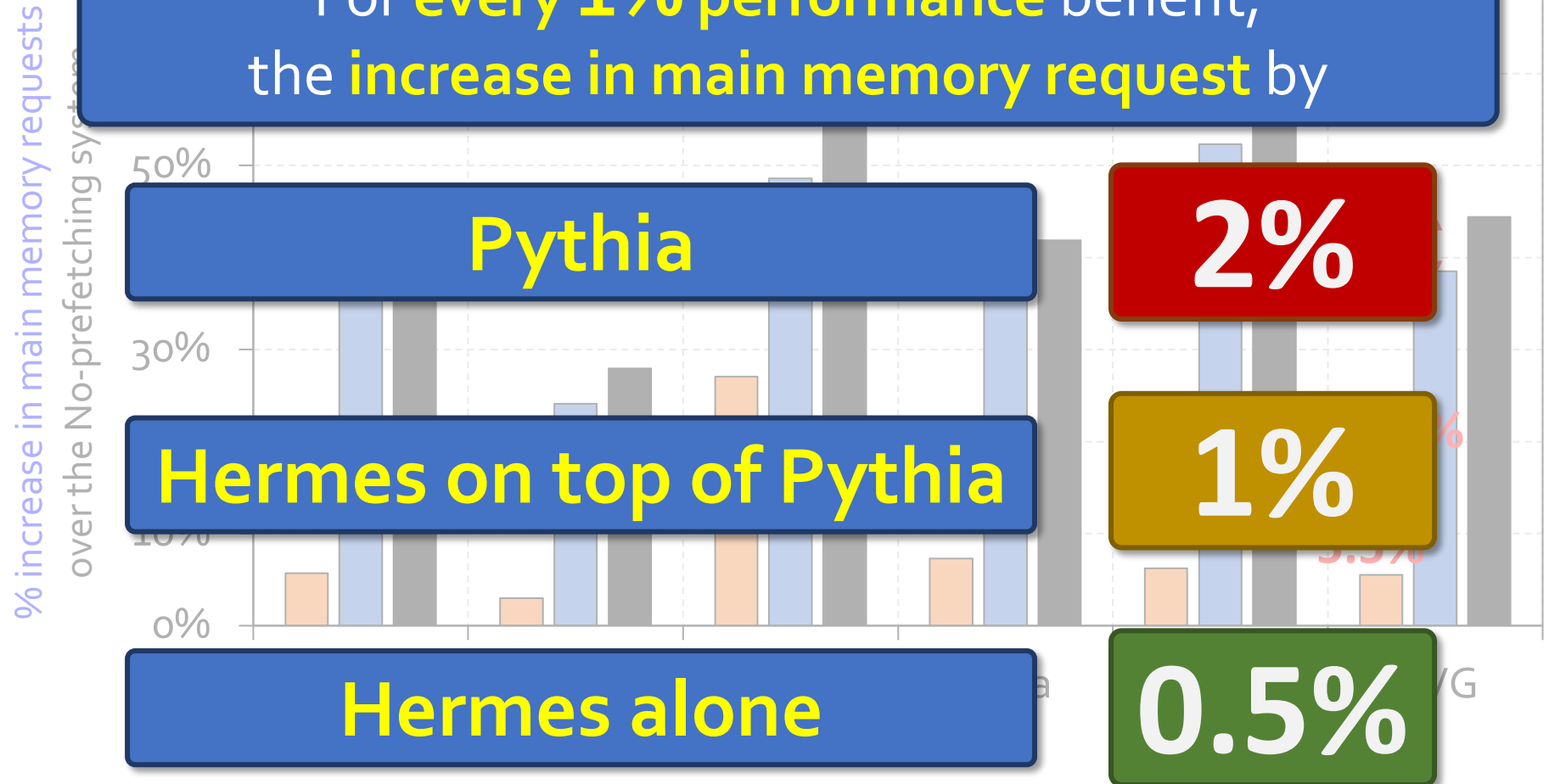
For **every 1%** performance benefit, the **increase in main memory request** by



Bandwidth Efficiency

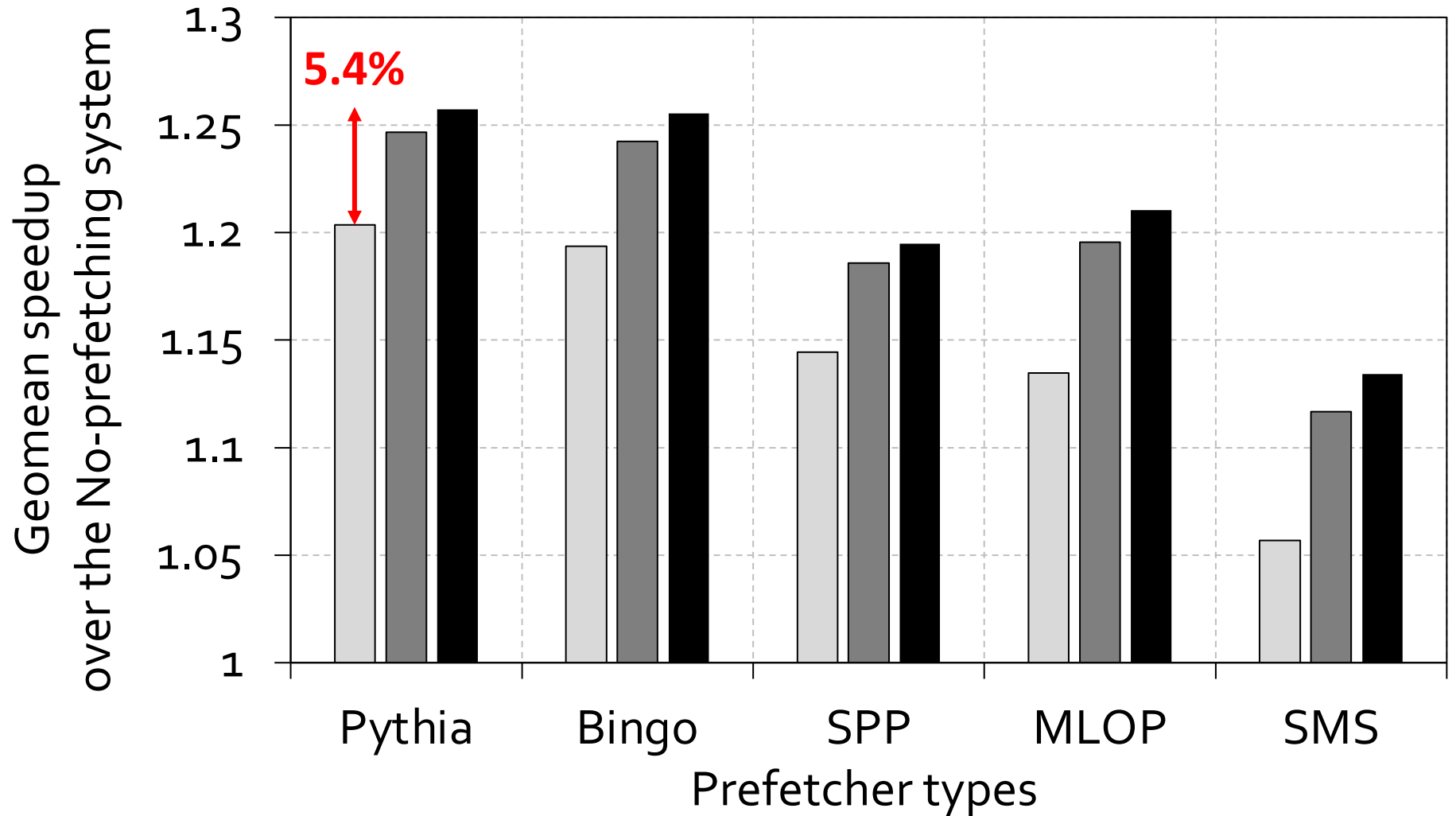


For **every 1%** performance benefit, the **increase in main memory request** by



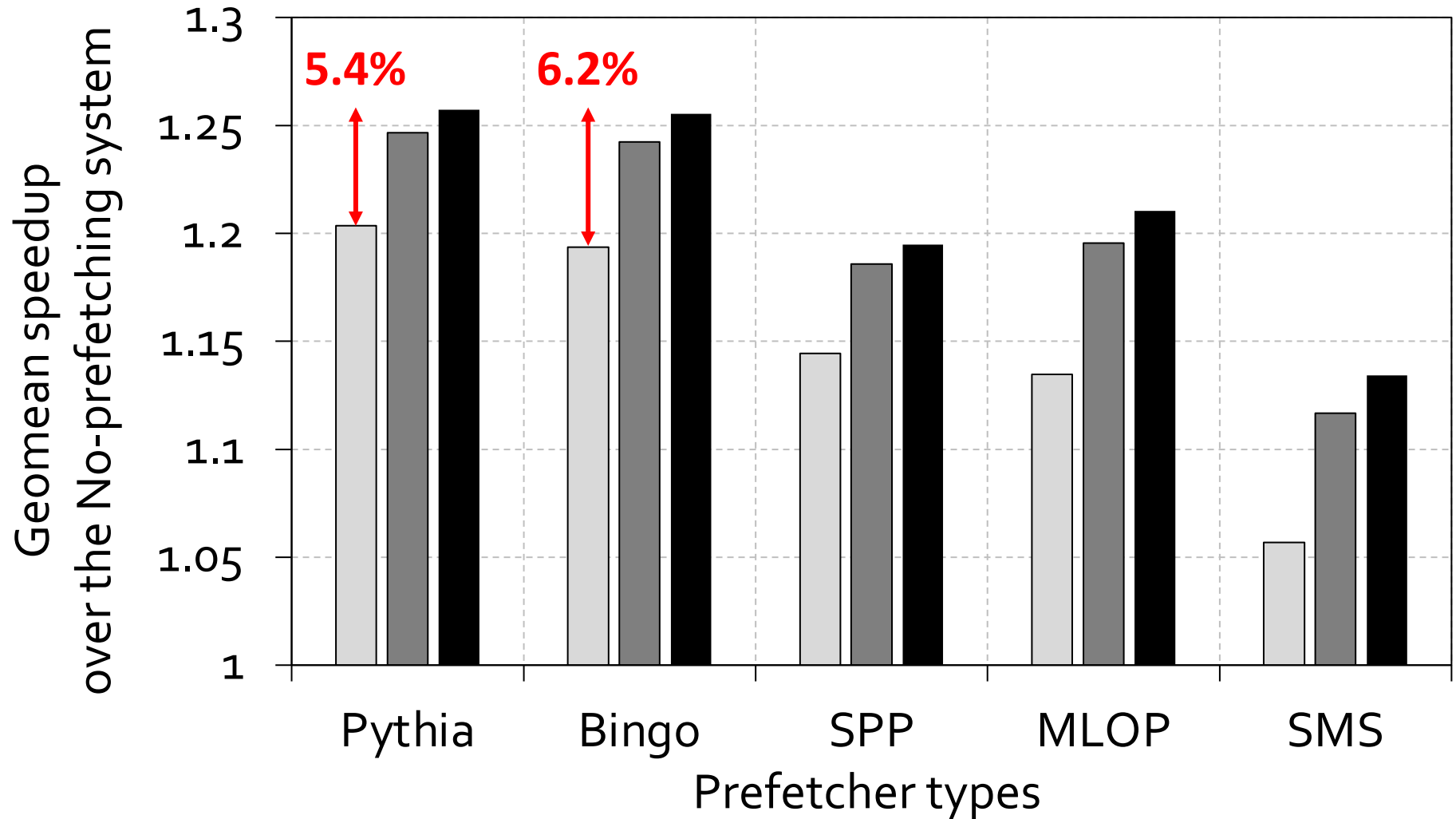
Performance with Varying Baseline Prefetcher

■ Prefetcher-only ■ Prefetcher + Hermes-P ■ Prefetcher + Hermes-O



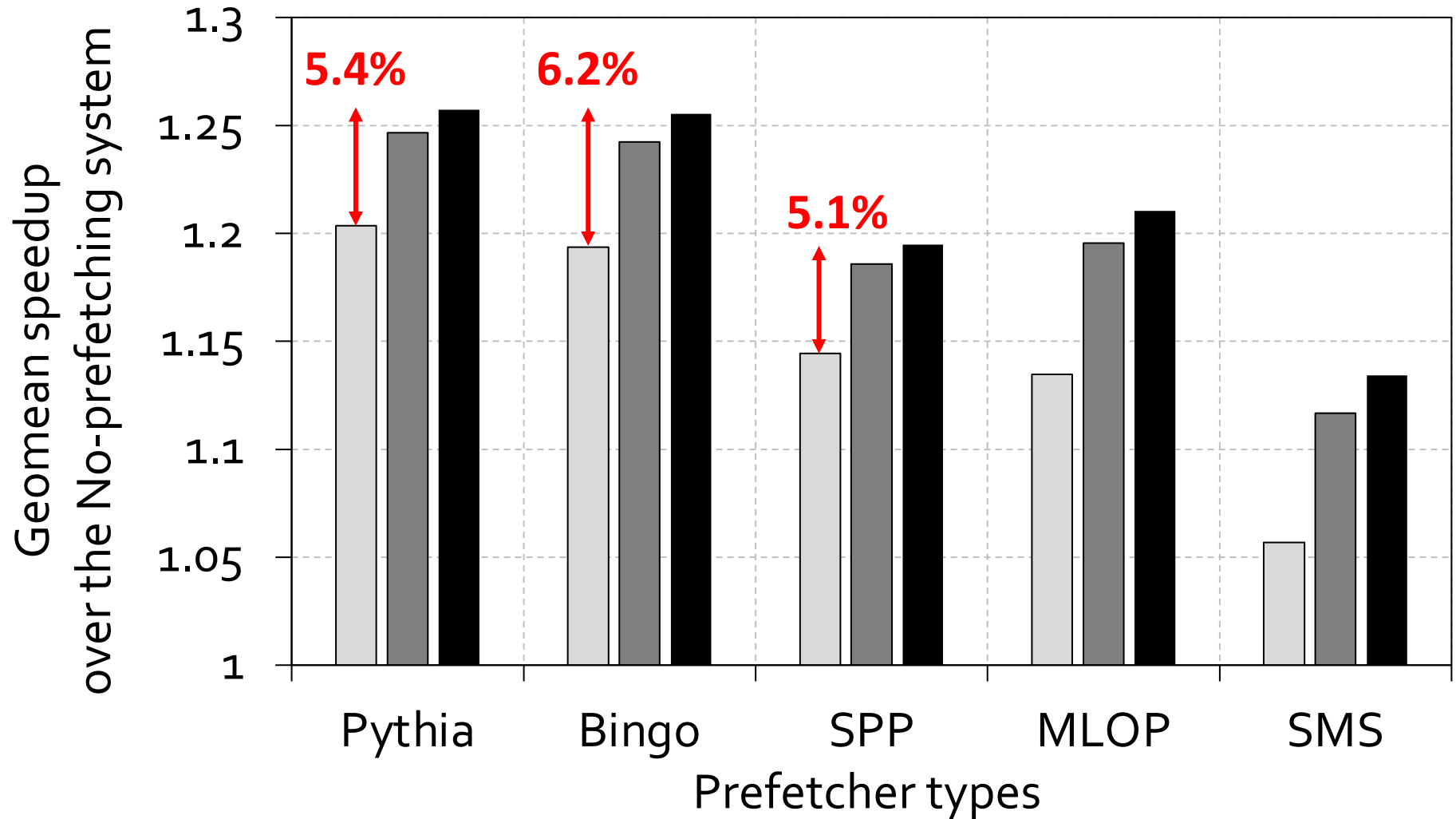
Performance with Varying Baseline Prefetcher

■ Prefetcher-only ■ Prefetcher + Hermes-P ■ Prefetcher + Hermes-O



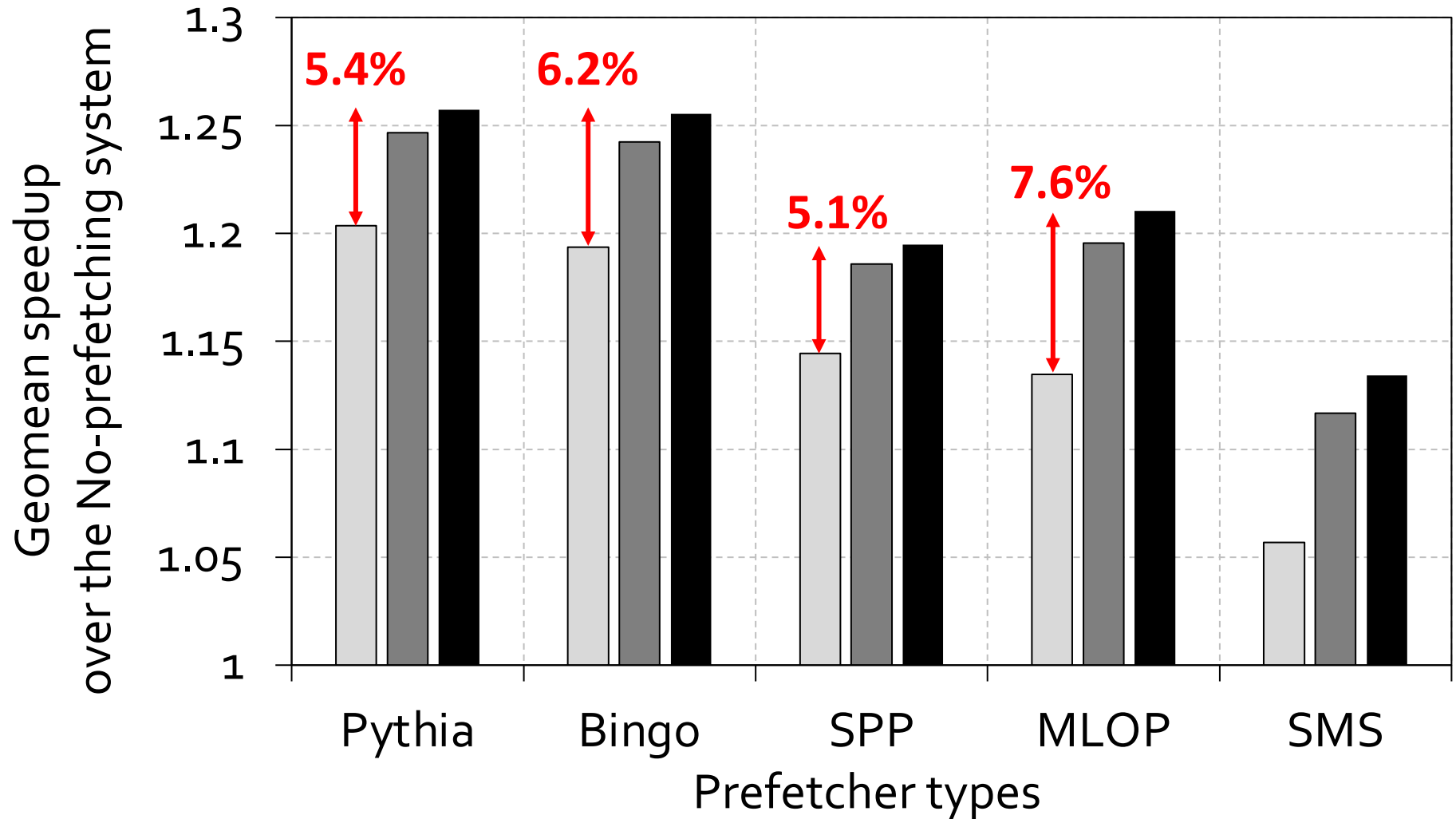
Performance with Varying Baseline Prefetcher

■ Prefetcher-only ■ Prefetcher + Hermes-P ■ Prefetcher + Hermes-O



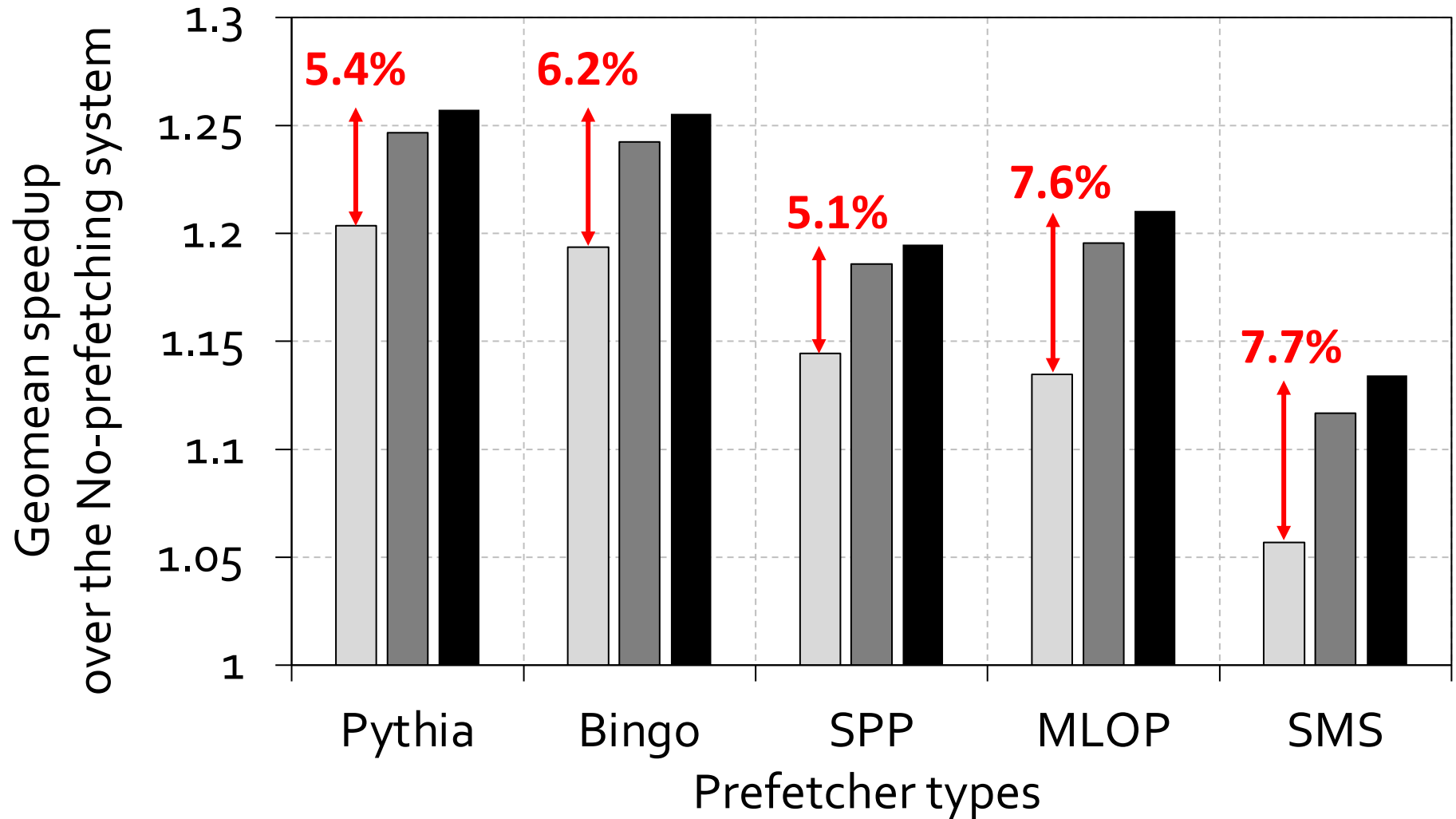
Performance with Varying Baseline Prefetcher

■ Prefetcher-only ■ Prefetcher + Hermes-P ■ Prefetcher + Hermes-O



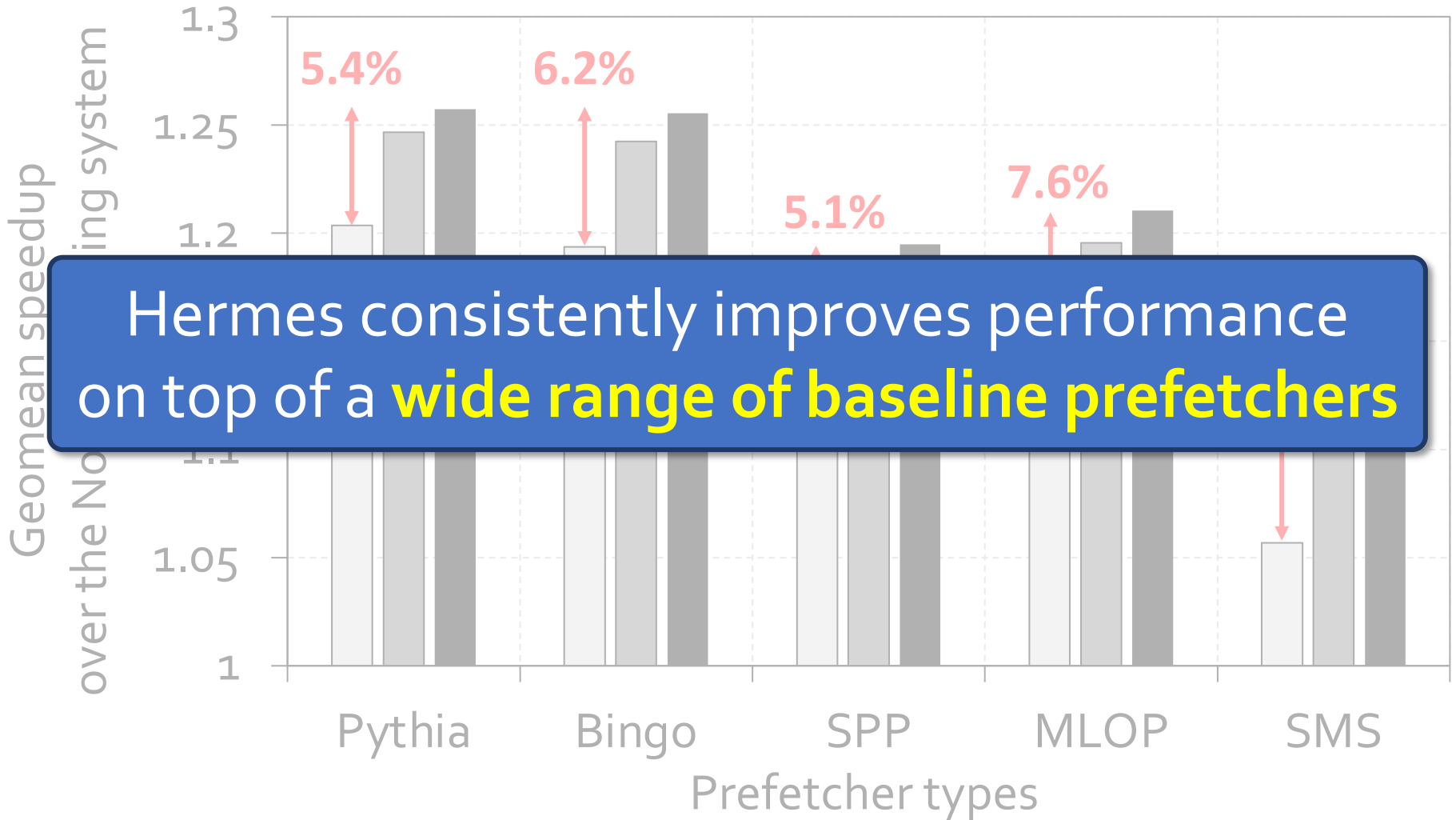
Performance with Varying Baseline Prefetcher

■ Prefetcher-only ■ Prefetcher + Hermes-P ■ Prefetcher + Hermes-O



Performance with Varying Baseline Prefetcher

□ Prefetcher-only ■ Prefetcher + Hermes-P ■ Prefetcher + Hermes-O



Overhead of Hermes

Overhead of Hermes



4 KB storage overhead

Overhead of Hermes



4 KB storage overhead



1.5% power overhead

On top of an Intel Alder Lake-like performance-core^[2] configuration

More in the Paper

- Wide range of **sensitivity studies** by varying
 - Cache hierarchy access latency
 - Hermes request issue latency
 - Activation threshold
- Comparison of POPET's **accuracy** and **coverage** against **HMP** and **TTP**
- Understanding **usefulness of each program feature**
- Hermes's effect on **stall cycle reduction**
- **Performance** analysis in **eight-core** system

More in the Paper

- Wide range of **sensitivity studies** by varying
 - **Cache hierarchy access latency**



Hermes: Accelerating Long-Latency Load Requests via Perceptron-Based Off-Chip Load Prediction

Rahul Bera¹ Konstantinos Kanellopoulos¹ Shankar Balachandran² David Novo³
Ataberk Olgun¹ Mohammad Sadrosadati¹ Onur Mutlu¹

¹ETH Zürich ²Intel Processor Architecture Research Lab ³LIRMM, Univ. Montpellier, CNRS

<https://arxiv.org/pdf/2209.00188.pdf>

- Hermes's effect on **stall cycle reduction**
- **Performance** analysis in **eight-core** system

To Summarize...

Summary

Summary

Hermes advocates **off-chip load prediction**,
a different form of speculation than
load address predictions by prefetchers

Summary

Hermes advocates **off-chip load prediction**,
a different form of speculation than
load address predictions by prefetchers

Off-chip load prediction can be applied **by itself**
or **combined with load address prediction**
to provide performance improvement

Summary

Hermes...

Summary

Hermes...



Identifies **74%** off-chip loads,

Summary

Hermes...



Identifies **74%** off-chip loads,



that misses a **4.5 MB** cache hierarchy
using only **4 KB** storage overhead,

Summary

Hermes...



Identifies **74%** off-chip loads,



that misses a **4.5 MB** cache hierarchy
using only **4 KB** storage overhead,



with **77%** accuracy,

Summary

Hermes...



Identifies **74%** off-chip loads,



that misses a **4.5 MB** cache hierarchy
using only **4 KB** storage overhead,



with **77%** accuracy,



and provides **5.4%** performance gain

Hermes is Open Sourced



 **CMU-SAFARI / Hermes** Public

Unwatch 3 Fork 1 Star 6

[Code](#) [Issues](#) [Pull requests](#) [Actions](#) [Projects](#) [Wiki](#) [Security](#) [Insights](#) [Settings](#)

main 1 branch 5 tags Go to file Add file Code

Rahul Bera Minor changes in experiment file		9eaeca0 5 days ago	21 commits
branch	Initial commit for MICRO'22 AE		2 months ago
config	Minor update		5 days ago
cvp_tracer	Initial commit for MICRO'22 AE		2 months ago
experiments	Minor changes in experiment file		5 days ago
inc	Upgraded TTP implementation. All experiment files and rollup script...		6 days ago
logo	Github theme-adapting logo		2 months ago
mcpat	Initial commit for MICRO'22 AE		2 months ago
prefetcher	Initial commit for MICRO'22 AE		2 months ago
replacement	Initial commit for MICRO'22 AE		2 months ago
scripts	Added support to launch multiple experiments in parallel when loc...		25 days ago
src	Minor update		5 days ago
tools	Initial commit for MICRO'22 AE		2 months ago
tracer	Initial commit for MICRO'22 AE		2 months ago
traces	Fixed the link mismatch in artifact_traces.csv		2 months ago
.gitignore	Added extra experiment file and an example python script to genera...		25 days ago
LICENSE	Updated License		2 months ago
LICENSE.ChampSim	Updated License		2 months ago

About

A speculative mechanism to accelerate long-latency off-chip load requests by removing on-chip cache access latency from their critical path, as described by MICRO 2022 paper by Bera et al. (<https://arxiv.org/pdf/2209.00188.pdf>)

arxiv.org/abs/2209.00188

[machine-learning](#) [cache](#) [perceptron](#) [computer-architecture](#) [microarchitecture](#) [perceptron-learning-algorithm](#) [prefetching](#)

[Readme](#) [MIT, Unknown licenses found](#) [6 stars](#) [3 watching](#) [1 fork](#)

Releases

[v1.2](#) Latest 6 days ago [+ 4 releases](#)

Packages

No packages published

Hermes is Open Sourced



All workload traces

CMU-SAFARI Hermes

<> Code

main

Rahul Bera Minor changes in experiment file 9eaeca0 5 days ago 21 commits

branch	Initial commit for MICRO'22 AE	2 months ago
config	Minor update	5 days ago
cvp_tracer	Initial commit for MICRO'22 AE	2 months ago
experiments	Minor changes in experiment file	5 days ago
inc	Upgraded TTP implementation. All experiment files and rollup script...	6 days ago
logo	Github theme-adapting logo	2 months ago
mcpat	Initial commit for MICRO'22 AE	2 months ago
prefetcher	Initial commit for MICRO'22 AE	2 months ago
replacement	Initial commit for MICRO'22 AE	2 months ago
scripts	Added support to launch multiple experiments in parallel when loc...	25 days ago
src	Minor update	5 days ago
tools	Initial commit for MICRO'22 AE	2 months ago
tracer	Initial commit for MICRO'22 AE	2 months ago
traces	Fixed the link mismatch in artifact_traces.csv	2 months ago
.gitignore	Added extra experiment file and an example python script to genera...	25 days ago
LICENSE	Updated License	2 months ago
LICENSE.ChampSim	Updated License	2 months ago

long-latency off-chip load requests by removing on-chip cache access latency from their critical path, as described by MICRO 2022 paper by Bera et al. (<https://arxiv.org/pdf/2209.00188.pdf>)

arxiv.org/abs/2209.00188

machine-learning cache perceptron
computer-architecture microarchitecture
perceptron-learning-algorithm prefetching

Readme
MIT, Unknown licenses found
6 stars
3 watching
1 fork

Releases 5

v1.2 Latest 6 days ago
+ 4 releases

Packages
No packages published

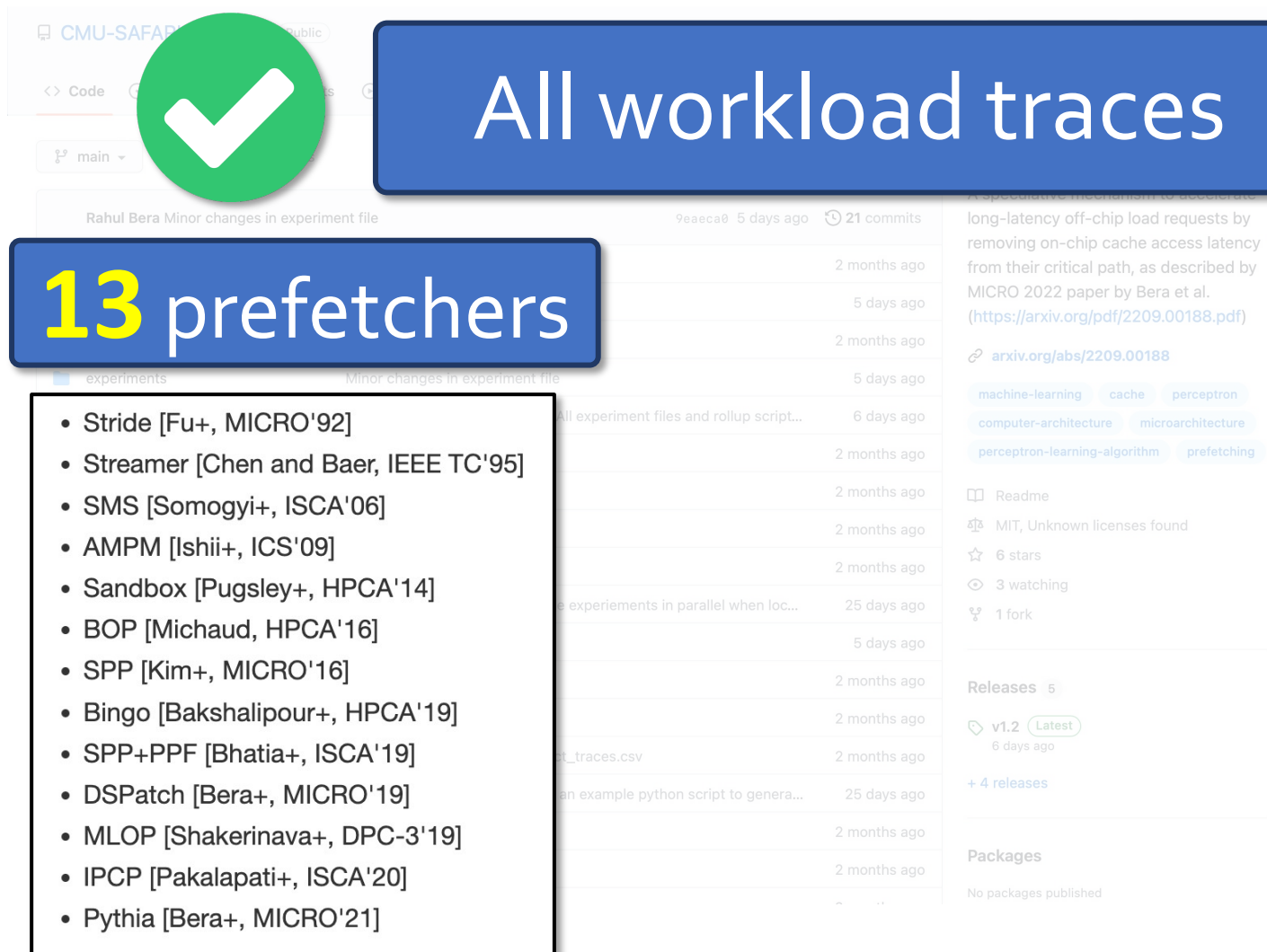
Hermes is Open Sourced



All workload traces

13 prefetchers

- Stride [Fu+, MICRO'92]
- Streamer [Chen and Baer, IEEE TC'95]
- SMS [Somogyi+, ISCA'06]
- AMPM [Ishii+, ICS'09]
- Sandbox [Pugsley+, HPCA'14]
- BOP [Michaud, HPCA'16]
- SPP [Kim+, MICRO'16]
- Bingo [Bakshalipour+, HPCA'19]
- SPP+PPF [Bhatia+, ISCA'19]
- DSPatch [Bera+, MICRO'19]
- MLOP [Shakerinava+, DPC-3'19]
- IPCP [Pakalapati+, ISCA'20]
- Pythia [Bera+, MICRO'21]



Hermes is Open Sourced



All workload traces

13 prefetchers

- Stride [Fu+, MICRO'92]
- Streamer [Chen and Baer, IEEE TC'95]
- SMS [Somogyi+, ISCA'06]
- AMPM [Ishii+, ICS'09]
- Sandbox [Pugsley+, HPCA'14]
- BOP [Michaud, HPCA'16]
- SPP [Kim+, MICRO'16]
- Bingo [Bakshalipour+, HPCA'19]
- SPP+PPF [Bhatia+, ISCA'19]
- DSPatch [Bera+, MICRO'19]
- MLOP [Shakerinava+, DPC-3'19]
- IPCP [Pakalapati+, ISCA'20]
- Pythia [Bera+, MICRO'21]

9 off-chip predictors

Predictor type	Description
Base	Always NO
Basic	Simple confidence counter-based threshold
Random	Random Hit-miss predictor with a given positive probability
HMP-Local	Hit-miss predictor [Yoaz+, ISCA'99] with local prediction
HMP-GShare	Hit-miss predictor with GShare prediction
HMP-GSkew	Hit-miss predictor with GSkew prediction
HMP-Ensemble	Hit-miss predictor with all three types combined
TTP	Tag-tracking based predictor
Perc	Perceptron-based OCP used in this paper

Easy To Define Your Own Off-Chip Predictor

- Just extend the **OffchipPredBase** class

```
8  class OffchipPredBase
9  {
10 public:
11     uint32_t cpu;
12     string type;
13     uint64_t seed;
14     uint8_t dram_bw; // current DRAM bandwidth bucket
15
16     OffchipPredBase(uint32_t _cpu, string _type, uint64_t _seed) : cpu(_cpu), type(_type), seed(_seed)
17     {
18         srand(seed);
19         dram_bw = 0;
20     }
21     ~OffchipPredBase() {}
22     void update_dram_bw(uint8_t _dram_bw) { dram_bw = _dram_bw; }
23
24     virtual void print_config();
25     virtual void dump_stats();
26     virtual void reset_stats();
27     virtual void train(ooo_model_instr *arch_instr, uint32_t data_index, LSQ_ENTRY *lq_entry);
28     virtual bool predict(ooo_model_instr *arch_instr, uint32_t data_index, LSQ_ENTRY *lq_entry);
29 };
30
31 #endif /* OFFCHIP_PRED_BASE_H */
32
```

Easy To Define Your Own Off-Chip Predictor

- Define your own `train()` and `predict()` functions

```
19 void OffchipPredBase::train(ooo_model_instr *arch_instr, uint32_t data_index, LSQ_ENTRY *lq_entry)
20 {
21     // nothing to train
22 }
23
24 bool OffchipPredBase::predict(ooo_model_instr *arch_instr, uint32_t data_index, LSQ_ENTRY *lq_entry)
25 {
26     // predict randomly
27     // return (rand() % 2) ? true : false;
28     return false;
29 }
```

Easy To Define Your Own Off-Chip Predictor

- Define your own **train()** and **predict()** functions

```
19 void OffchipPredBase::train(ooo_model_instr *arch_instr, uint32_t data_index, LSQ_ENTRY *lq_entry)
20 {
21     // nothing to train
22 }
23
24 bool OffchipPredBase::predict(ooo_model_instr *arch_instr, uint32_t data_index, LSQ_ENTRY *lq_entry)
25 {
26     // predict randomly
27     // return (rand() % 2) ? true : false;
28     return false;
29 }
```

- Get statistics like **precision** (aka accuracy) and **recall** (aka coverage) out of the box

```
Core_0_offchip_pred_true_pos 2358716
Core_0_offchip_pred_false_pos 276883
Core_0_offchip_pred_false_neg 132145
Core_0_offchip_pred_precision 89.49
Core_0_offchip_pred_recall 94.69
```

Off-Chip Prediction Can Further Enable...

Off-Chip Prediction Can Further Enable...

Prioritizing loads that are likely go off-chip
in **cache queues** and **on-chip network routing**

Off-Chip Prediction Can Further Enable...

Prioritizing loads that are likely go off-chip
in **cache queues** and **on-chip network routing**

Better instruction scheduling
of data-dependent instructions

Off-Chip Prediction Can Further Enable...

Prioritizing loads that are likely go off-chip
in **cache queues** and **on-chip network routing**

Better instruction scheduling
of data-dependent instructions

and many more...



HERMES

Accelerating Long-Latency Load Requests via Perceptron-Based Off-Chip Load Prediction

Rahul Bera, Konstantinos Kanellopoulos, Shankar Balachandran,
David Novo, Ataberk Olgun, Mohammad Sadrosadati, Onur Mutlu

