

Objective

IMPROVING INTRA-WARP CACHE UTILIZATION ON GPUS THROUGH COMPILER-GUIDED EXPLOITATION OF TEXTURE HARDWARE

Method

Mechanism

YI-TING LAI

ADVISOR: YI-PING YOU

Note: We assume NVIDIA GPUs.



國立陽明交通大學
NATIONAL YANG MING CHIAO TUNG UNIVERSITY

INSTITUTE OF COMPUTER SCIENCE AND ENGINEERING
NATIONAL YANG MING CHIAO TUNG UNIVERSITY

sslolo
系統軟體實驗室

OUTLINE

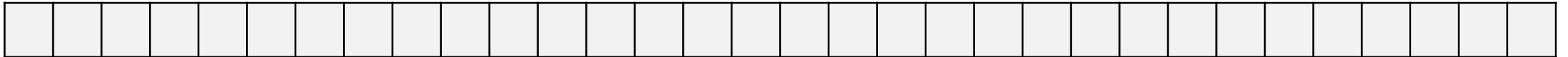
- What is Intra-Warp Cache Utilization?
- Why Is It Important?
- How Can Texture Hardware Help?
- How Can a Compiler Exploit It?
- How Effective Is the Exploitation?
- Limitations of the Mechanism
- Conclusion

WHAT IS INTRA-WARP CACHE UTILIZATION?

GPU EXECUTION HARDWARE: SIMT-CORE ARRAY

Analogous to SIMD lanes,
executing the same instruction in parallel.

SIMT CORE × 32

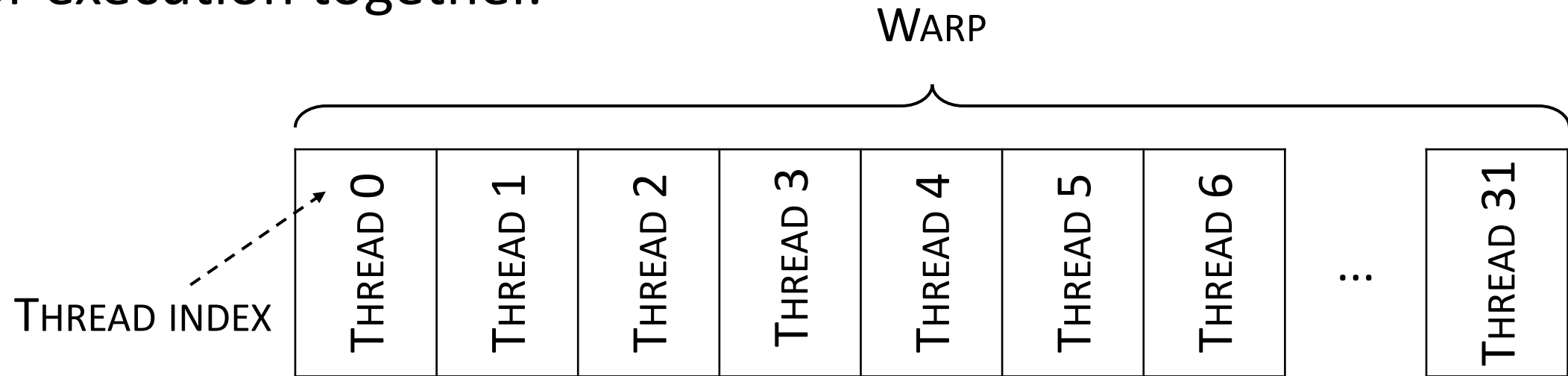


SIMD: single instruction, multiple data

SIMT: single instruction, multiple threads

GPU EXECUTION MODEL: THREAD AND WARP

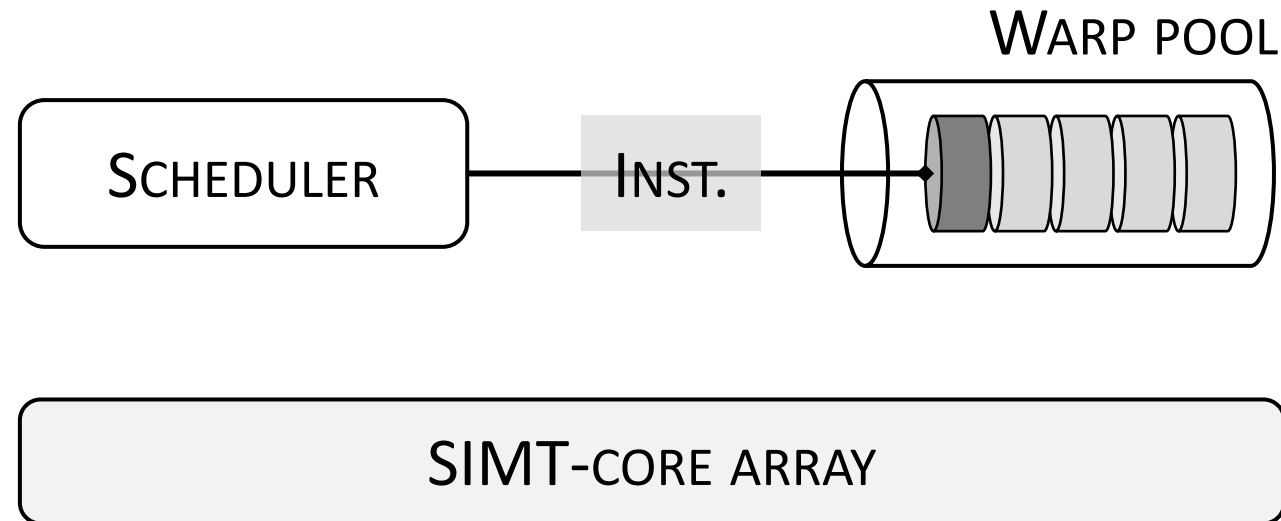
The SPMD programming model decomposes a workload into many *threads*. Every 32 threads form a *warp* and are scheduled for execution together.



SPMD: single program, multiple data

SCHEDULING UNIT: WARP

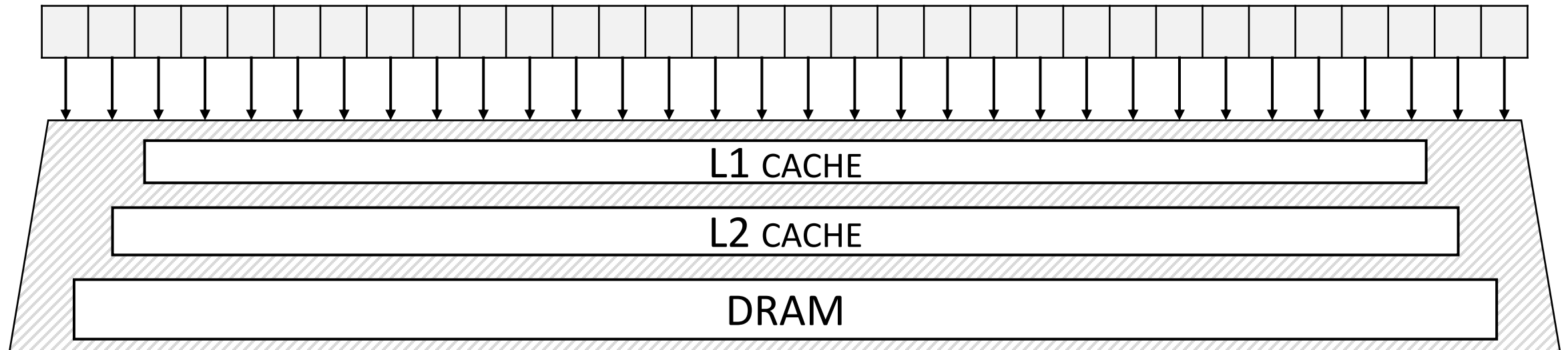
Each cycle, the scheduler selects a warp and issues its next instruction; the selected warp may differ from the previous cycle.



MULTI-ADDRESS MEMORY INSTRUCTION

A memory instruction requests up to 32 addresses—one from each thread—which are sent to the memory hierarchy.

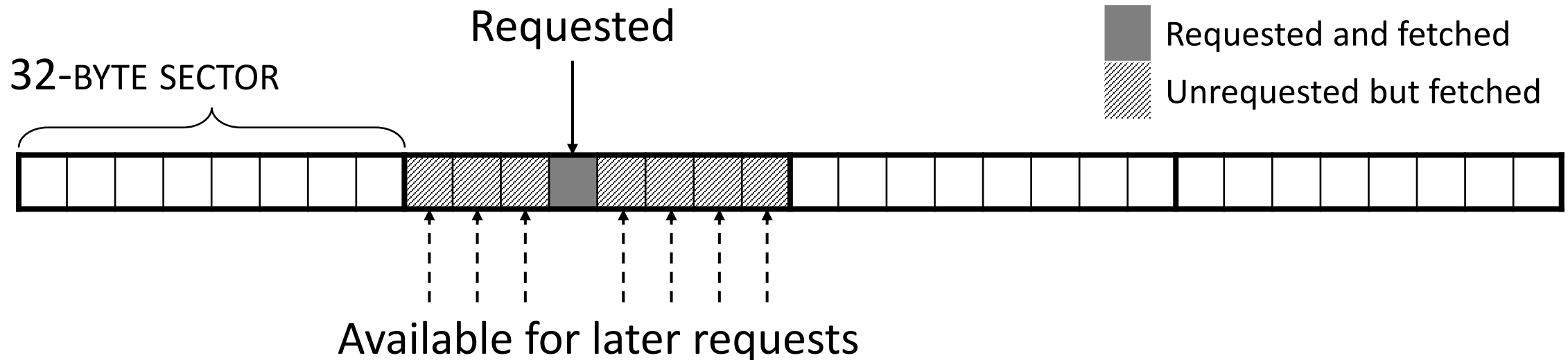
SIMT-CORE ARRAY



CACHE FETCH GRANULARITY

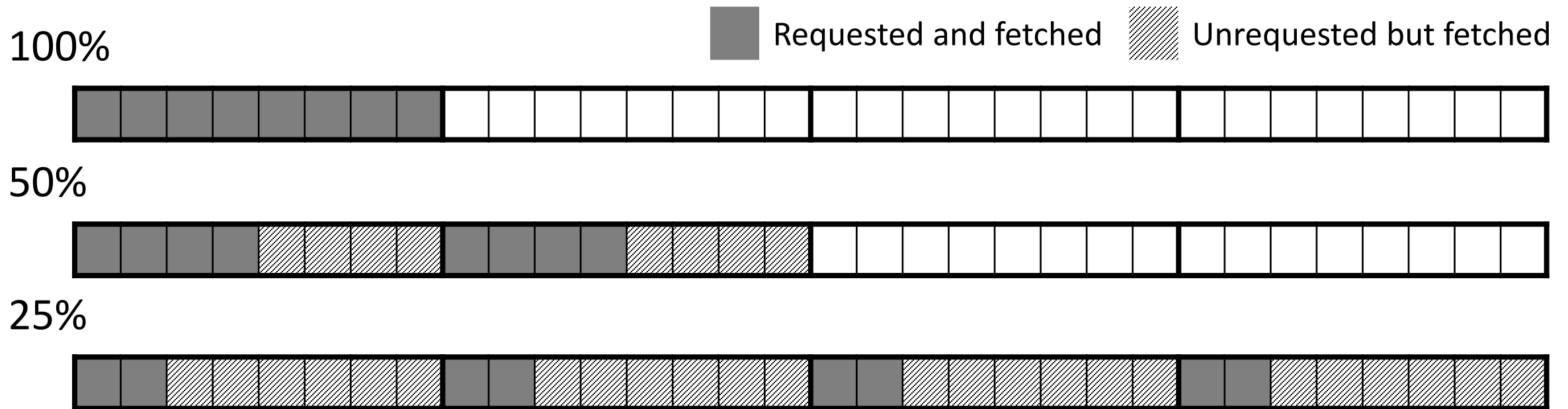
On a cache miss, data is fetched in 32-byte cache sectors, not necessarily limited to the requested data.

Unrequested-but-fetched data may later be requested by subsequent memory instructions as cache hits.



DEFINITION: INTRA-WARP CACHE UTILIZATION

The percentage of fetched cache-sector data that is actually requested by a memory instruction.



Note: This is actually cache-sector utilization; for the cache-line counterpart, see Chapter 3.

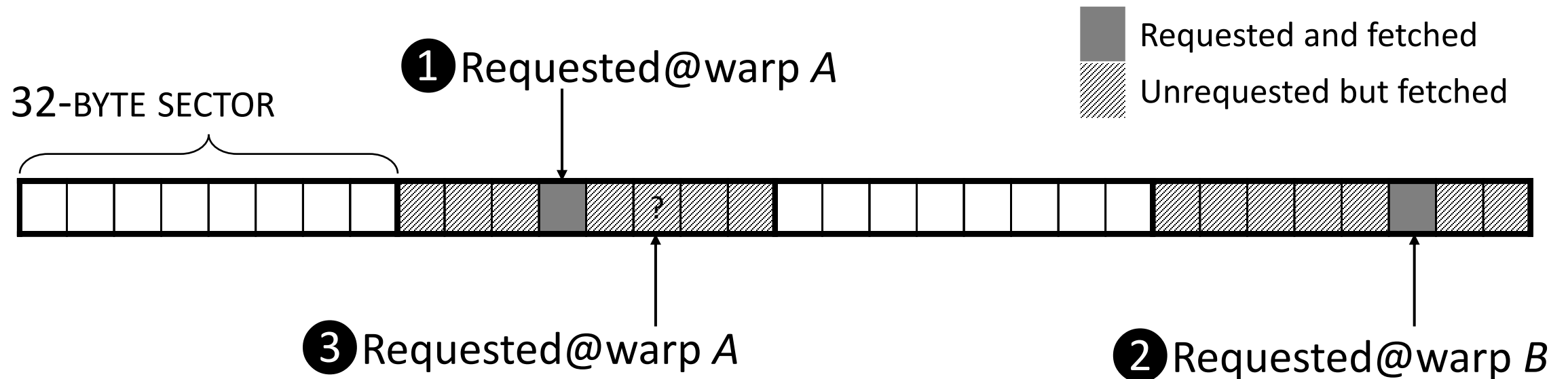
WHY IS INTRA-WARP CACHE UTILIZATION IMPORTANT?

INTER-WARP CACHE CONTENTION

Cache is shared across all warps in the pool.

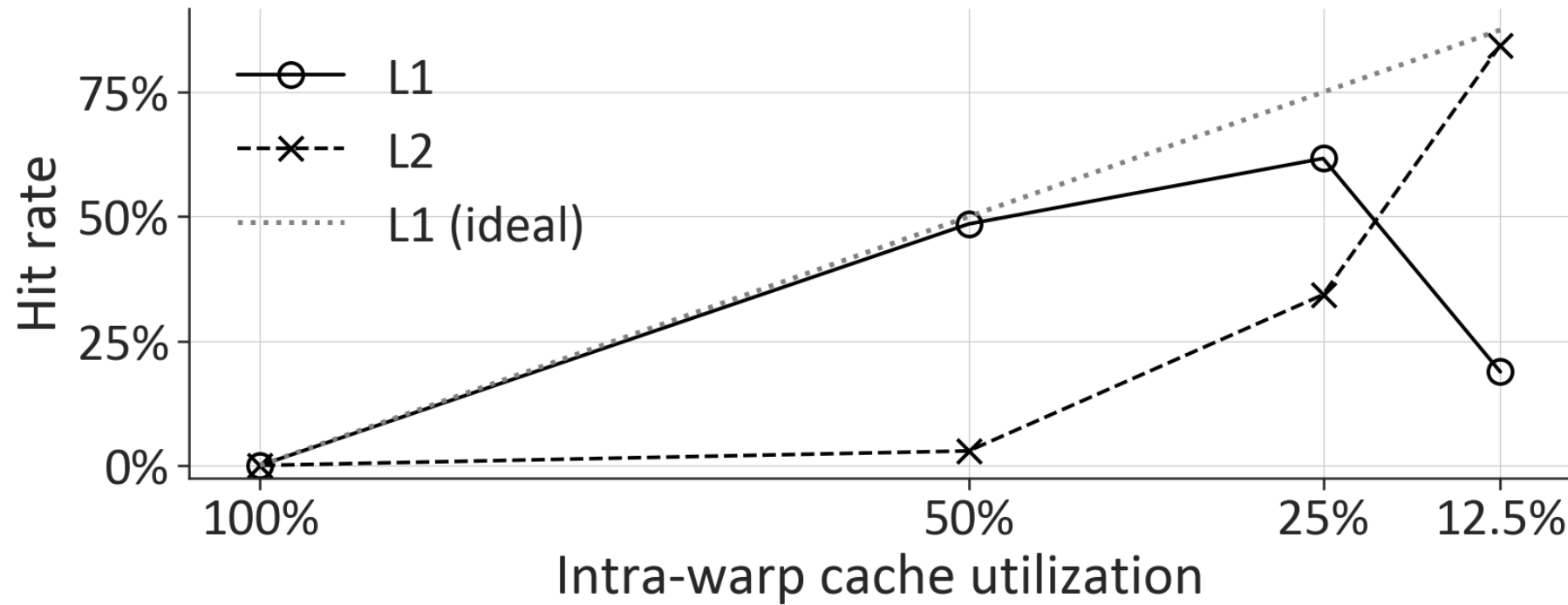
Other warps can evict data before it is requested later.

Note: Eviction granularity is a cache line.



KEY OBSERVATION I

Memory instructions should maximize utilization of fetched cache data rather than leaving it to subsequent memory instructions.



Scenario: streaming access; each element is accessed exactly once.

TWO WAYS TO IMPROVE INTRA-WARP CACHE UTILIZATION

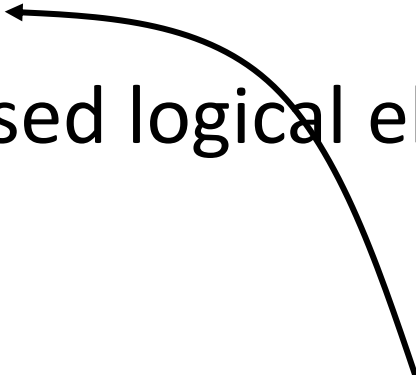
Access pattern:

make threads access nearby logical elements together.

Memory layout:

place jointly accessed logical elements at nearby memory addresses.

Our proposed direction



WHAT IS MEMORY LAYOUT?

MEMORY LAYOUT: $(x, y) \rightarrow i$

		LOGICAL			
		x			
		0	1	2	3
y	0	(0, 0)	(1, 0)	(2, 0)	(3, 0)
1	(0, 1)	(1, 1)	(2, 1)	(3, 1)	
2	(0, 2)	(1, 2)	(2, 2)	(3, 2)	
3	(0, 3)	(1, 3)	(2, 3)	(3, 3)	

		PHYSICAL					
		i					
		0	1	2	3	4	5
		????	????	????	????	????	

Determines the mapping between logical data and its physical placement in memory (in DRAM and within cache lines).

DEFAULT MEMORY LAYOUT: LINEAR

		LOGICAL			
		x			
		0	1	2	3
y	0	(0, 0)	(1, 0)	(2, 0)	(3, 0)
	1	(0, 1)	(1, 1)	(2, 1)	(3, 1)
	2	(0, 2)	(1, 2)	(2, 2)	(3, 2)
	3	(0, 3)	(1, 3)	(2, 3)	(3, 3)

		PHYSICAL					
		i					
		0	1	2	3	4	5
		(0, 0)	(1, 0)	(2, 0)	(3, 0)	(0, 1)	

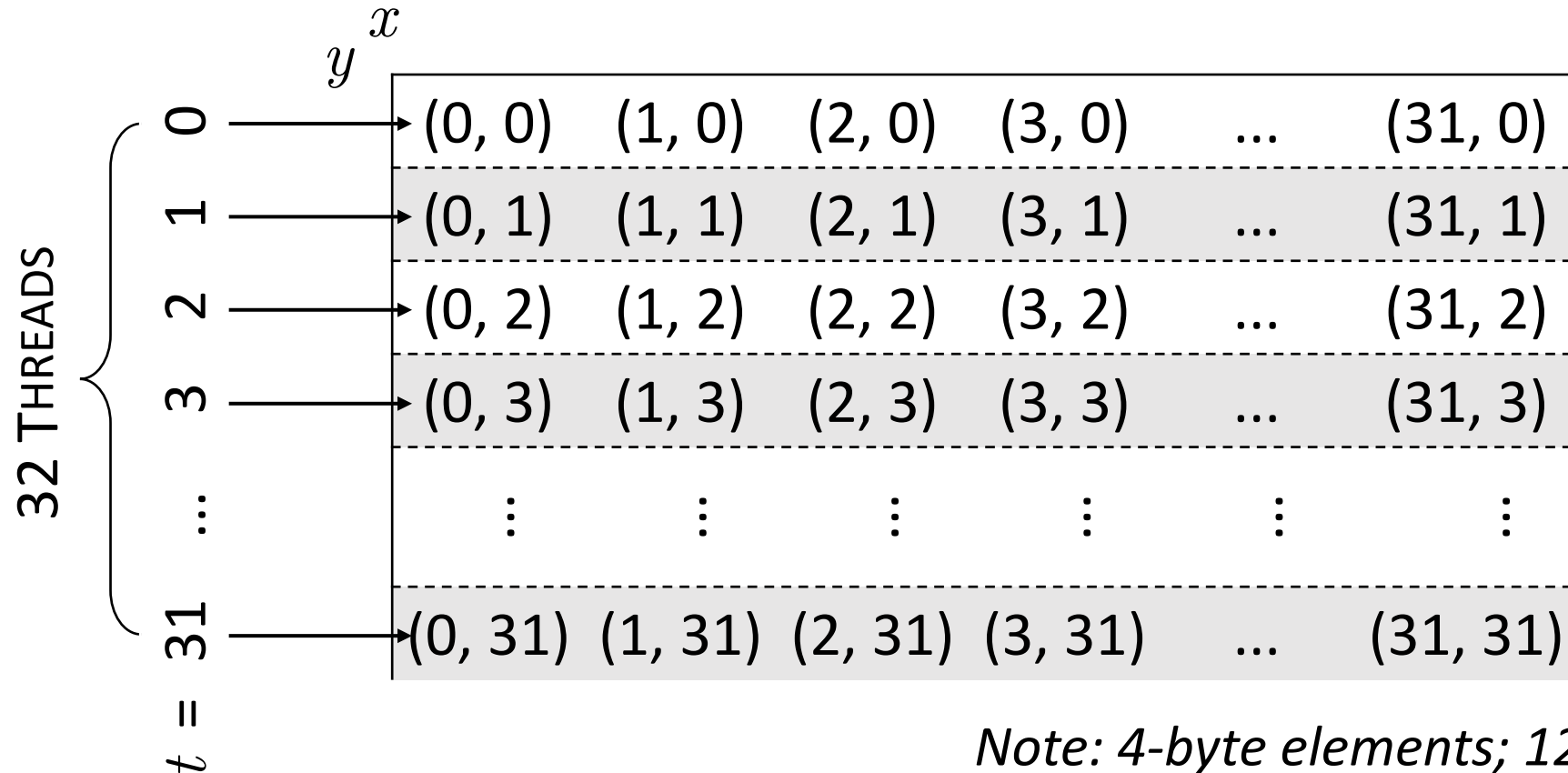
----- Row boundary

By default, data copied to GPU is stored in a linear layout, with consecutive data of each row placed consecutively in memory.

WHY CAN LINEAR LAYOUT LEAD TO LOW INTRA-WARP CACHE UTILIZATION?

WORST-CASE SCENARIO: CROSS-ROW ACCESS

Each thread _{t} in a warp accesses $(x, y + t)$.

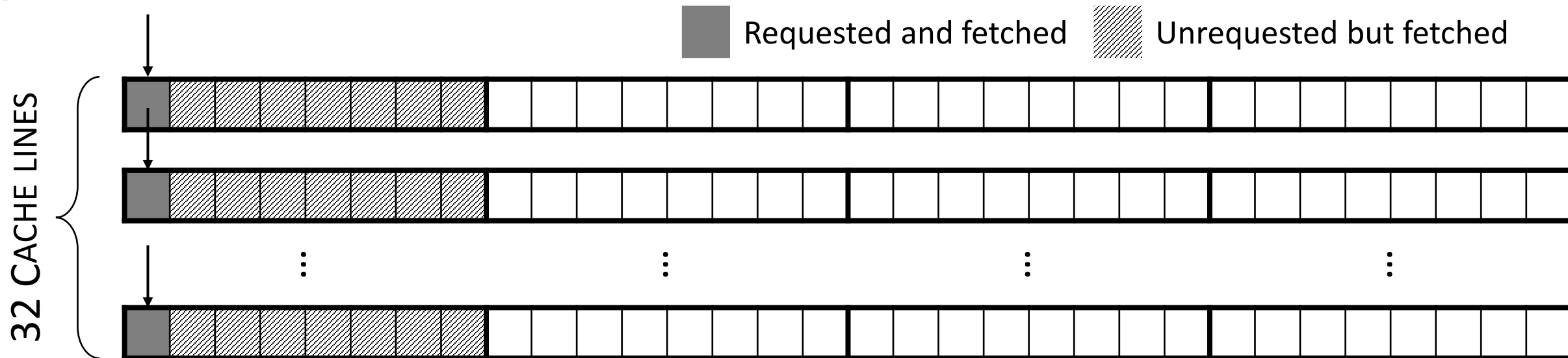


Note: 4-byte elements; 128-bytes logical row.

WORST-CASE SCENARIO: CROSS-ROW ACCESS

Each thread_{*t*} in a warp accesses $(x, y + t)$.

Because consecutive rows are separated by one cache line under the linear layout, the accesses map to 32 distinct cache sectors and lines, achieving 12.5% intra-warp cache utilization.



BLOCKED LAYOUT IMPROVES WORST-CASE UTILIZATION

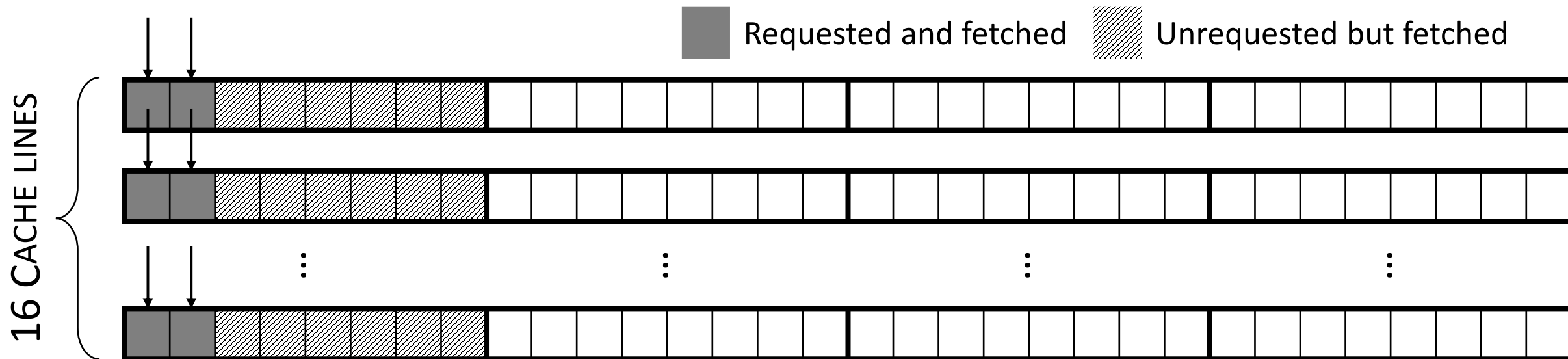
		LOGICAL			
		x			
		0	1	2	3
y	0	(0, 0)	(1, 0)	(2, 0)	(3, 0)
	1	(0, 1)	(1, 1)	(2, 1)	(3, 1)
	2	(0, 2)	(1, 2)	(2, 2)	(3, 2)
	3	(0, 3)	(1, 3)	(2, 3)	(3, 3)

PHYSICAL						
i	0	1	2	3	4	5
	(0, 0)	(1, 0)	(0, 1)	(1, 1)	(2, 0)	
-----			Row boundary			
—————			Block boundary			

When accesses may follow either logical dimension, a blocked layout can provide higher worst-case intra-warp cache utilization.

BLOCKED LAYOUT IMPROVES WORST-CASE UTILIZATION

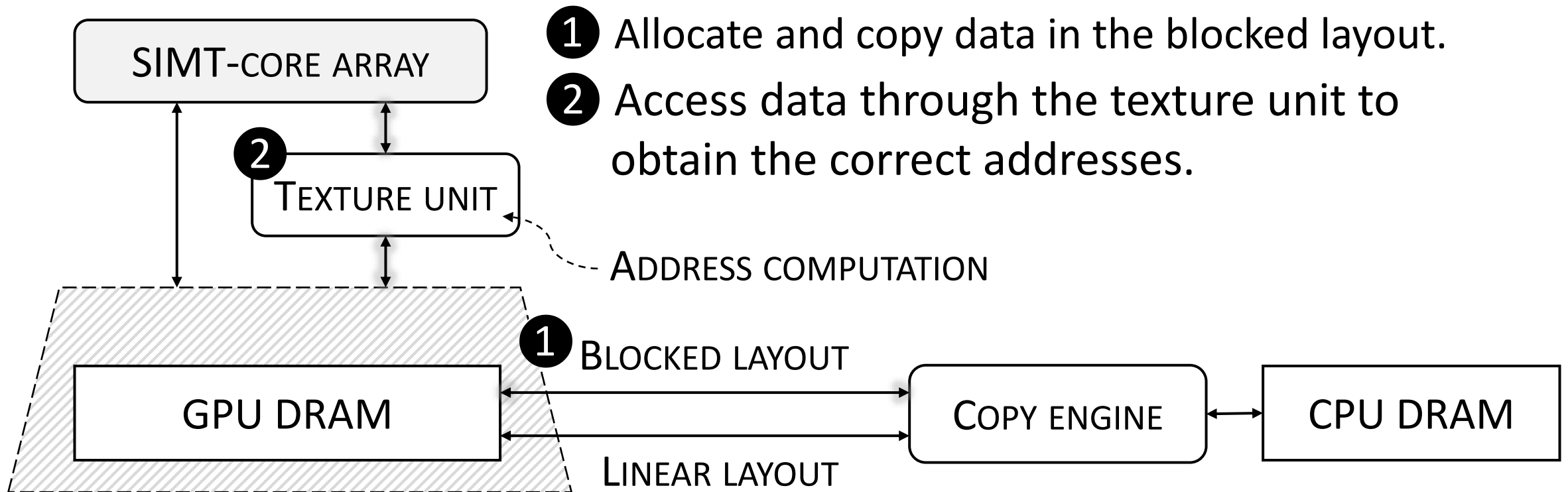
For the same cross-row access, a 2×2 blocked layout pairs data from adjacent rows within the same cache sector, doubling intra-warp cache utilization from 12.5% to 25%.



HOW CAN TEXTURE HARDWARE HELP?

KEY OBSERVATION II

Existing texture hardware already employs a blocked layout.



PROPOSED APPROACH

- Key observation I:
Memory instructions should maximize utilization of fetched cache data rather than leaving it to subsequent memory instructions.
 - Key observation II:
Existing texture hardware already employs a blocked layout.
- Use compiler analysis to identify data with low intra-warp cache utilization, then transparently allocate and copy it in the blocked layout and access it through the texture unit.

HOW IS THE APPROACH REALIZED?

GPU SUBMISSION WORKFLOW

Note: This is CUDA pseudocode.

```
/* execute on GPU */
```

```
@gpu_kernel
```

```
def kernel(gpu_data):
```

```
    v = access(gpu_data, x, y + thread_idx);
```

```
/* execute on CPU */
```

```
gpu_data = allocate_on_gpu(data_shape);
```

```
copy_to_gpu(cpu_data, gpu_data, data_shape);
```

```
launch_kernel(kernel, /* kernel argument */ gpu_data);
```

```
copy_from_gpu(gpu_data, cpu_data, data_shape);
```

*AUTO*TEX: THE COMPILER PIPELINE

Intra-warp access-
pattern analysis

Identify data with low intra-warp cache utilization from thread-index (t) use in its (x, y) access pattern.

```
@gpu_kernel
def kernel(gpu_data):
    v = access(gpu_data, x, y + thread_idx);
```

AUTO T EX: THE COMPILER PIPELINE

Intra-warp access-
pattern analysis



Kernel-argument
provenance analysis

Identify data with low intra-warp cache utilization from thread-index (t) use in its (x, y) access pattern.

Identify sharing kernels that must use the texture unit if their shared data is allocated in the blocked layout.

```
gpu_data = allocate_on_gpu(data_shape);
```

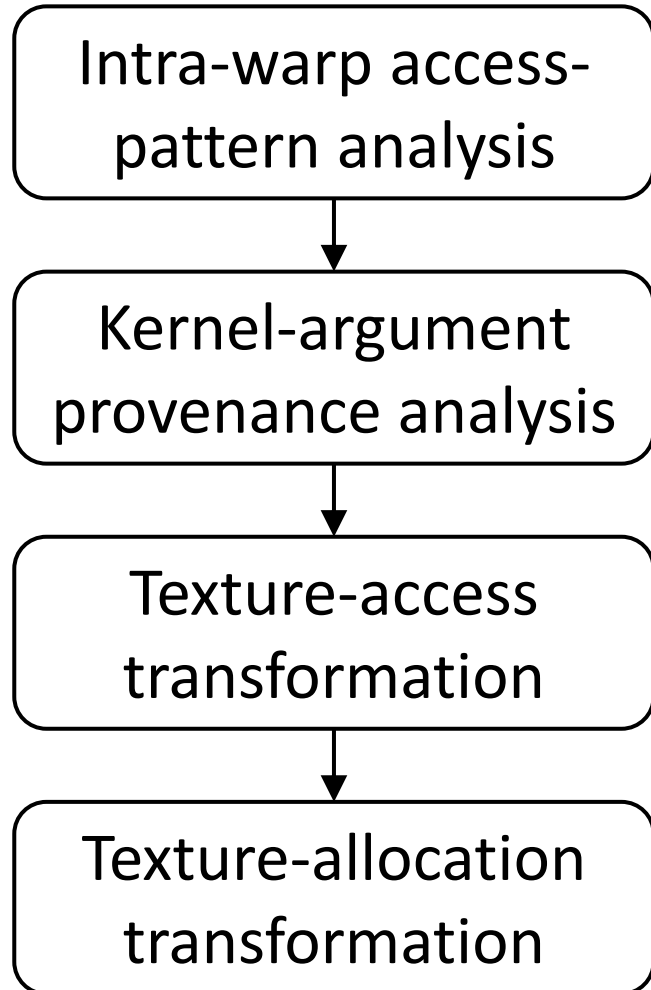
```
...
```

```
launch_kernel(kernel_1, gpu_data);
```

```
launch_kernel(kernel_2, gpu_data);
```

```
...
```

AUTO T EX: THE COMPILER PIPELINE



Identify data with low intra-warp cache utilization from thread-index (t) use in its (x, y) access pattern.

Identify sharing kernels that must use the texture unit if their shared data is allocated in the blocked layout.

Rewrite accesses to such data through the texture unit.
(2)

Allocate and copy such data using the blocked layout.
(1)

*Note: As an implementation detail,
2 is rewritten before 1.*

TRANSFORMED GPU SUBMISSION

Note: This is CUDA pseudocode.

```
/* execute on GPU */
```

```
@gpu_kernel
```

```
def kernel(gpu_data):
```

```
    v = tex_access(gpu_data, x, y + thread_idx);
```

```
/* execute on CPU */
```

```
gpu_data = allocate_on_gpu(data_shape, BLOCKED_LAYOUT);
```

```
copy_to_gpu(cpu_data, gpu_data, data_shape, BLOCKED_LAYOUT);
```

```
launch_kernel(kernel, /* kernel argument */ gpu_data);
```

```
copy_from_gpu(gpu_data, cpu_data, data_shape, BLOCKED_LAYOUT);
```

HOW EFFECTIVE IS THE APPROACH?

EXPERIMENTAL ENVIRONMENT

Note: Also evaluated on an RTX 3050 Laptop GPU; See Chapter 7.

Hardware specification

GPU	<u>Model</u>	<u>NVIDIA GeForce RTX 4080</u>
	Arch.	Ada Lovelace
	# SMs	76
	L1 cache	128 KiB
	L2 cache	64 MiB
	DRAM	15.6 GiB
CPU	Model	Intel® Core™ i7-12700KF
	DRAM	31.2 GiB

Software version

OS	Distribution	openSUSE Leap 16.0
	Linux kernel	6.12.0
Compiler	<u>LLVM</u>	<u>19.1.7</u>
CUDA	Driver	580.119.02
	Toolkit	12.9
	<u>Driver API</u>	<u>13.0</u>
	<u>Runtime API</u>	<u>12.9</u>

BENCHMARK SUITES AND TRANSFORMED KERNELS

	NPB-GPU	PolyBench-ACC	Rodinia	Total
# benchmarks	8	21	26	55
# transformed benchmarks	3	9	4	16
# kernels	106	47	68	221
# transformed kernels	25 (24%)	19 (40%)	8 (12%)	52
# low-utilization kernels ^a	16	10	6	32
# sharing-only kernels ^b	9	9	2	20

^aLow-utilization kernels: transformed kernels containing low-utilization accesses.

^bSharing-only kernels: no low-utilization accesses; transformed only because they share data with another low-utilization kernel.

Label numbers for transformed kernels		
1	IS/full_verify_gpu_kernel_1	19 SP/compute_rhs_gpu_kernel_1
2	IS/full_verify_gpu_kernel_2	20 SP/compute_rhs_gpu_kernel_2
3	LU/error_gpu_kernel	21 SP/rhs_norm_gpu_kernel_1
4	LU/jacld_blts_gpu_kernel	22 SP/txinvr_gpu_kernel
5	LU/jacu_buts_gpu_kernel	23 SP/x_solve_gpu_kernel
6	LU/pintgr_gpu_kernel_1	24 SP/y_solve_gpu_kernel
7	LU/pintgr_gpu_kernel_2	25 SP/z_solve_gpu_kernel
8	LU/pintgr_gpu_kernel_3	26 adi/adi_kernel1
9	LU/rhs_gpu_kernel_1	27 adi/adi_kernel3
10	LU/rhs_gpu_kernel_2	28 adi/adi_kernel4
11	LU/rhs_gpu_kernel_3	29 adi/adi_kernel6
12	LU/rhs_gpu_kernel_4	30 atax/atax_kernel1
13	LU/setbv_gpu_kernel_1	31 atax/atax_kernel2
14	LU/setbv_gpu_kernel_2	32 bicg/bicg_kernel1
15	LU/setbv_gpu_kernel_3	33 bicg/bicg_kernel2
16	LU/setiv_gpu_kernel	34 gemver/gemver_kernel1
17	LU/ssor_gpu_kernel_2	35 gemver/gemver_kernel2
18	SP/add_gpu_kernel	36 gemver/gemver_kernel3
		37 gesummv/gesummv_kernel
		38 gramschmidt/gramschmidt_kernel1
		39 gramschmidt/gramschmidt_kernel2
		40 gramschmidt/gramschmidt_kernel3
		41 mvt/mvt_kernel1
		42 mvt/mvt_kernel2
		43 syr2k/syr2k_kernel
		44 syrk/syrk_kernel
		45 gaussian/Fan1
		46 gaussian/Fan2
		47 huffman/pack2
		48 huffman/vlc_encode_kernel_sm64huff
		49 kmeans/invert_mapping
		50 kmeans/kmeansPoint
		51 leukocyte/GICOV_kernel
		52 leukocyte/dilate_kernel

EVALUATION METRICS

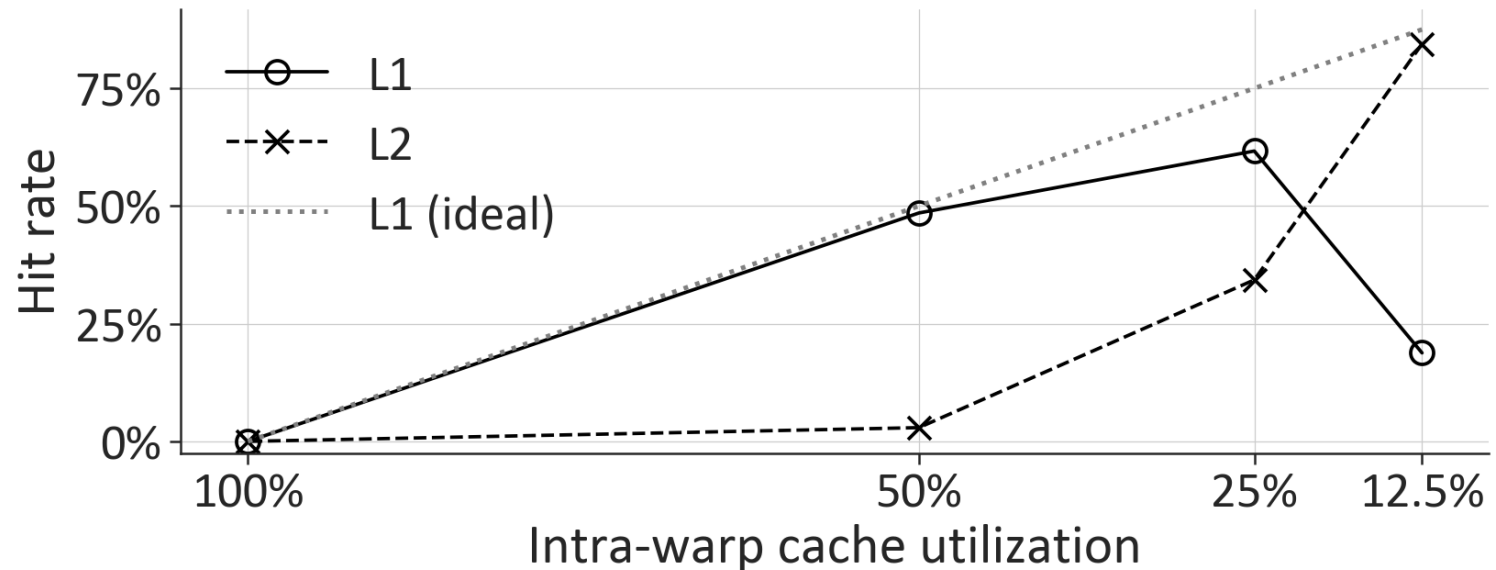
Hit rate alone is ambiguous: higher utilization can lower it by leaving less fetched data for later hits, while effective caching raises it toward the ideal. We therefore evaluate effectiveness directly through (1) fetched-sector reduction and (2) speedup.

Note:

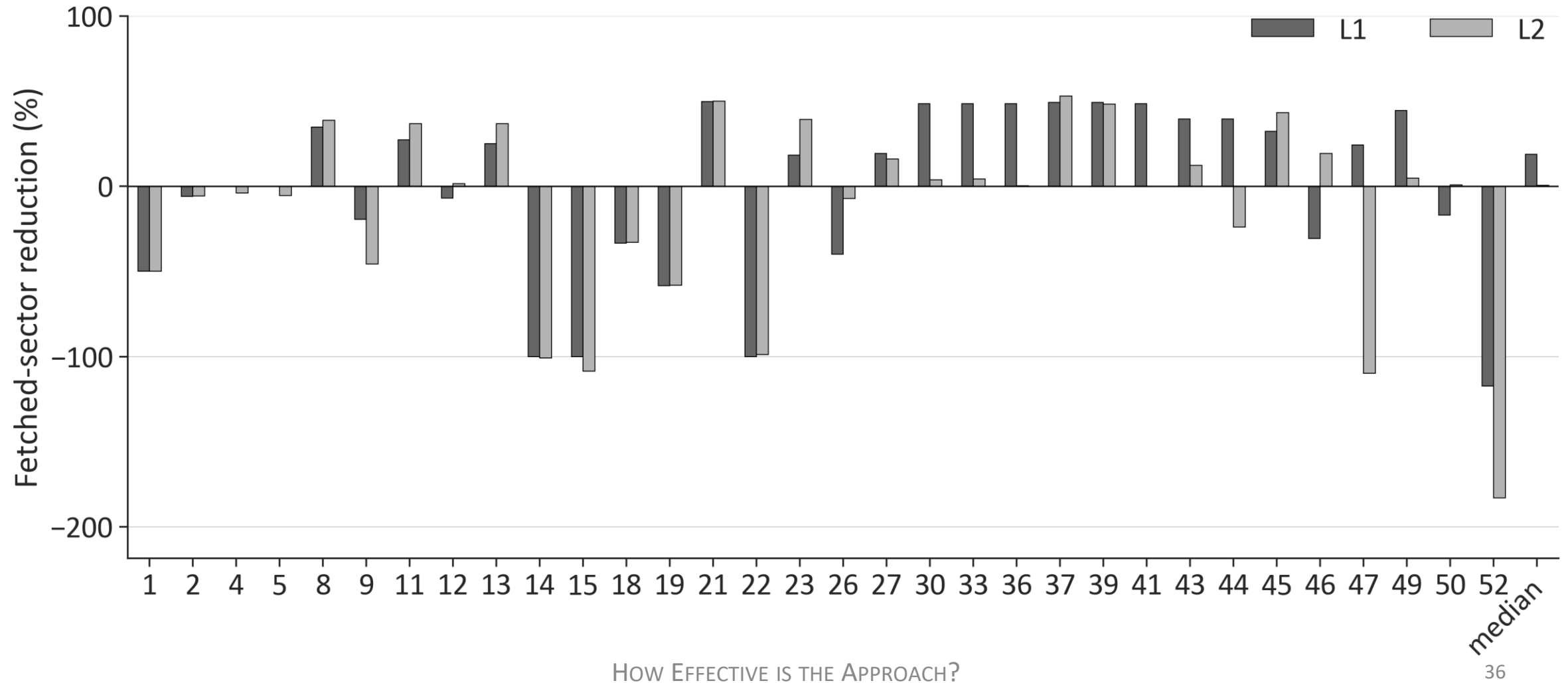
$$\text{Reduction} = \left(1 - \frac{\text{AutoTex}}{\text{Baseline}}\right) \times 100\%,$$

$$\text{Speedup} = \frac{\text{Baseline}}{\text{AutoTex'}}$$

where Baseline uses the LLVM O3 pipeline.

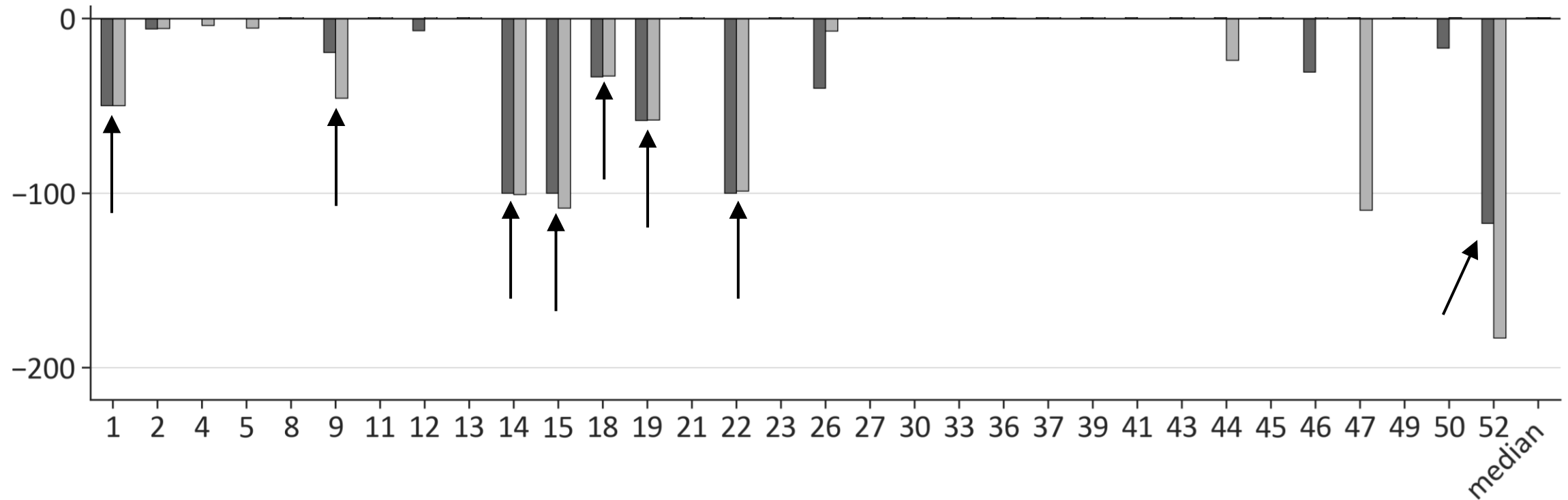


FETCHED SECTOR REDUCTION ON LOW-UTILIZATION KERNELS

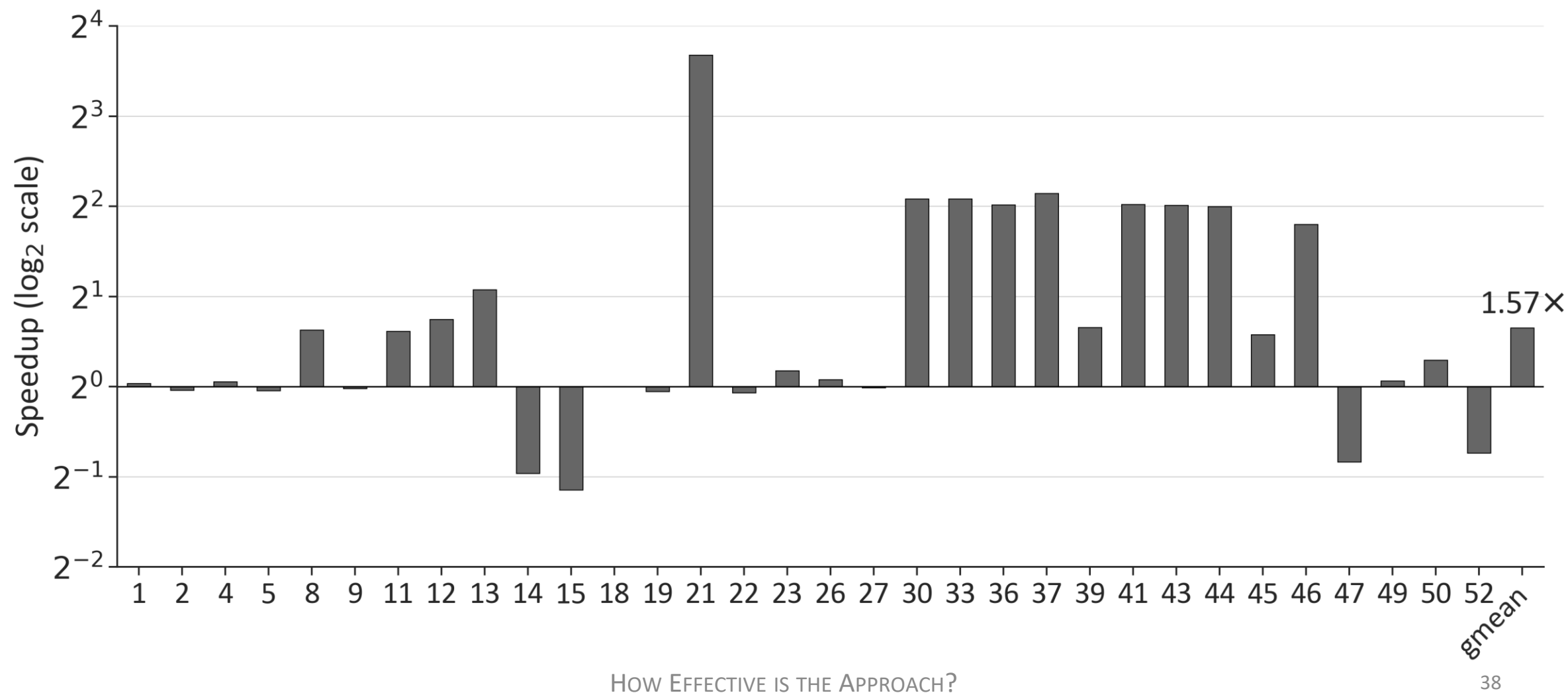


FALSE POSITIVES IN ACCESS-PATTERN ANALYSIS

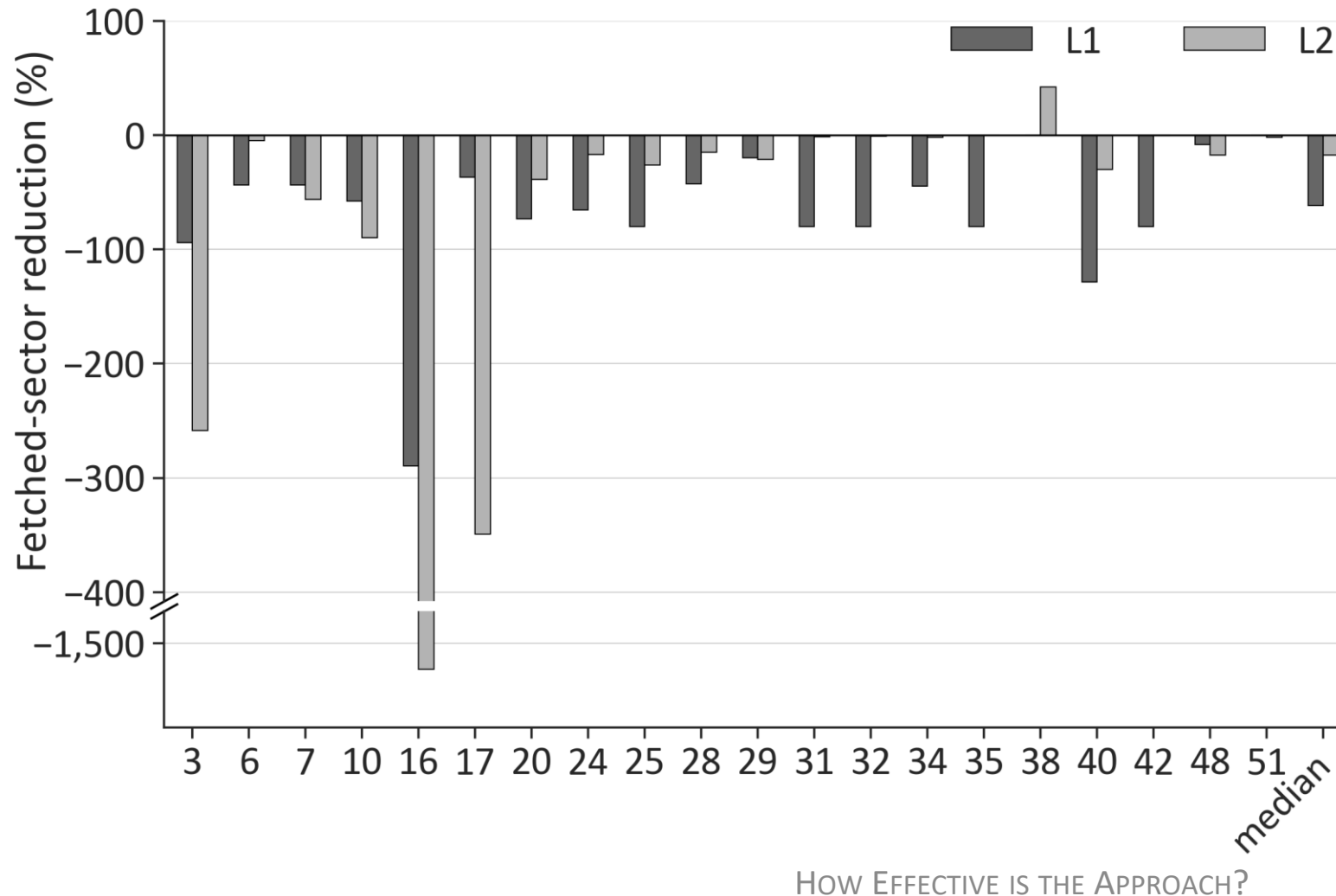
AutoTex currently treats arithmetic on thread indices as indicative of low utilization, e.g., $(x, y + t / r)$. Such false positives can increase fetched sectors.



SPEEDUP ON LOW-UTILIZATION KERNELS

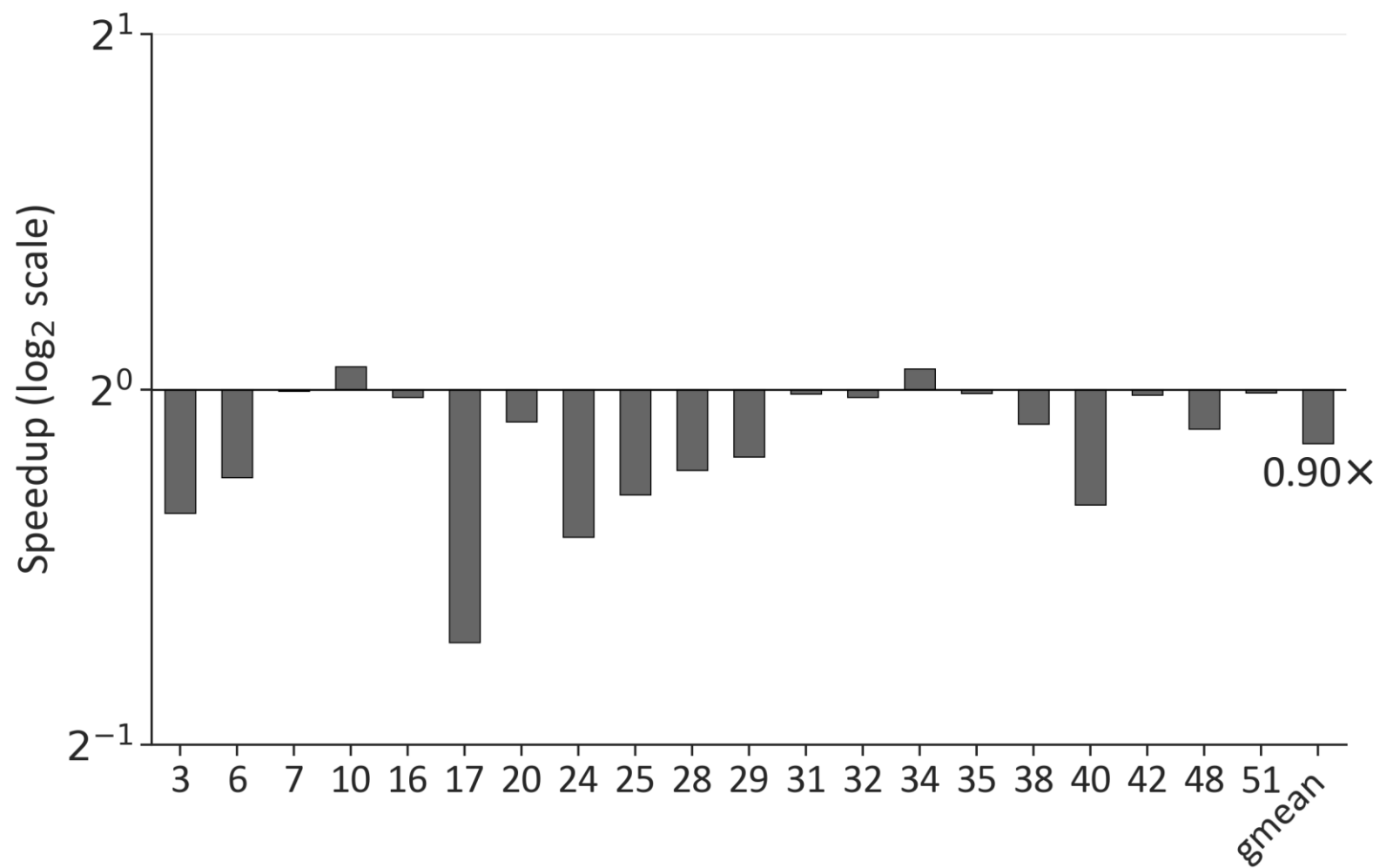


FETCHED-SECTOR REDUCTION ON SHARING-ONLY KERNELS



The shared data is already accessed with high intra-warp cache utilization, so the blocked layout can increase fetched sectors.

SPEEDUP ON SHARING-ONLY KERNELS



HOW EFFECTIVE IS THE APPROACH?

PER-KERNEL EVALUATION TAKEAWAYS

- Blocked layout reduces fetched sectors and improves performance, while false positives can cause mis-optimization.
- Sharing-only transformations can instead increase fetched sectors and degrade performance.

How do these opposing effects translate to overall benchmark speedup?

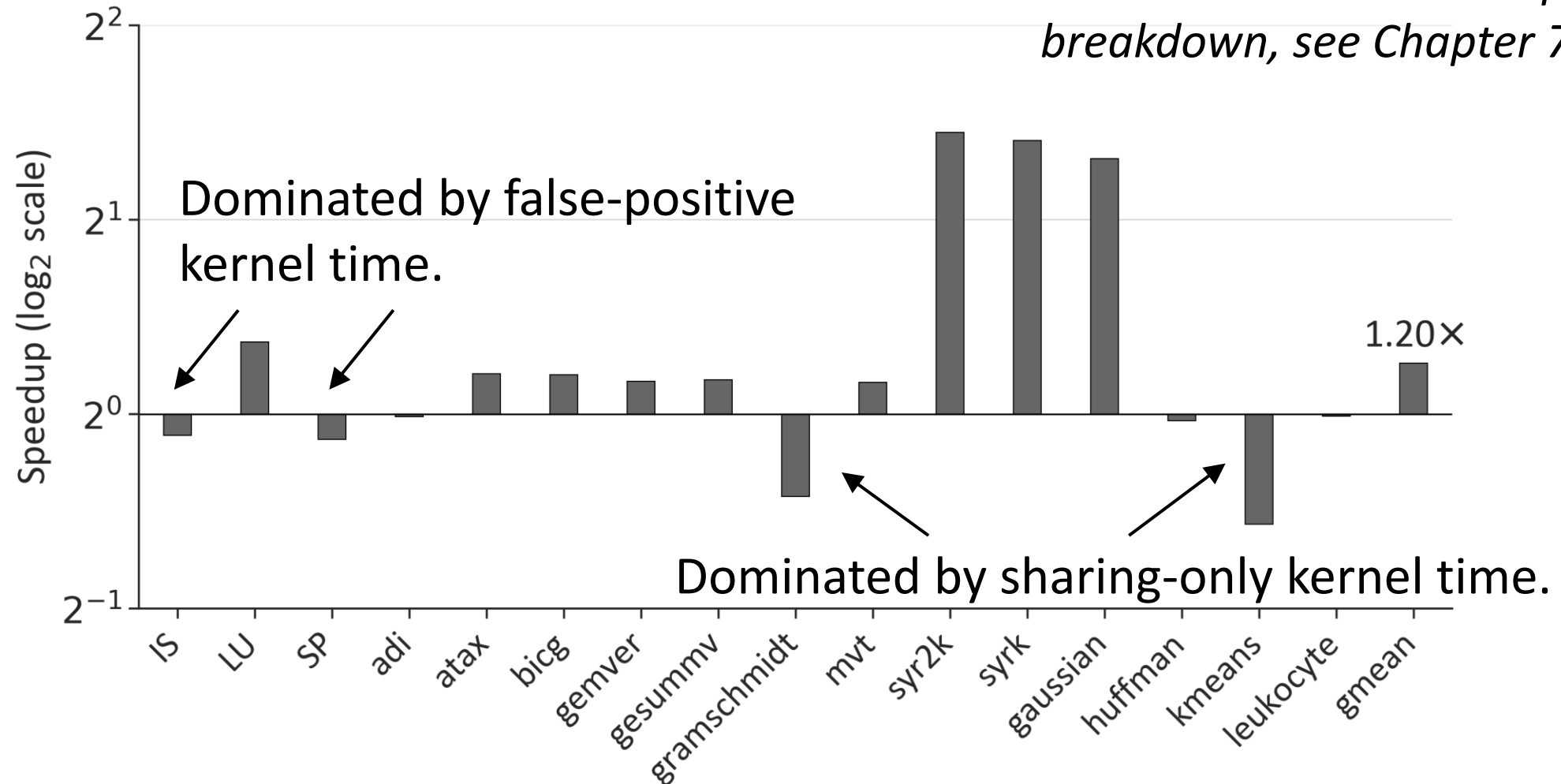
BENCHMARK RUNTIME COMPONENTS

In addition to kernel execution, benchmark time includes

- data allocation (`allocate_on_gpu`),
- data copy (`copy_to_gpu`, `copy_from_gpu`), and
- kernel-launch (`launch_kernel`) time.

OVERALL BENCHMARK SPEEDUP

Note: For a detailed component breakdown, see Chapter 7.



LIMITATIONS OF THE LAYOUT MECHANISM

- 2D blocked layouts: up to $131,072 \times 65,536$ 4-byte elements;
3D: up to 16,384 4-byte elements per dimension.
- Other specialized GPU hardware, particularly tensor cores and the Tensor Memory Accelerator, relies on linear layout.

CONCLUSION

- **Objective:** Improve intra-warp cache utilization on GPUs.
- **Mechanism:** Exploit the hardware-managed blocked layout provided by texture hardware.
- **Method:** Use compiler analysis and transformation to transparently apply the blocked layout to data with low intra-warp cache utilization and access it through the texture unit.
- **Result:** 1.57× geometric-mean speedup on low-utilization kernels and 1.20× on transformed benchmarks.
- **Contribution:** Texture hardware can be repurposed as a compiler-exploitable memory-layout optimization.