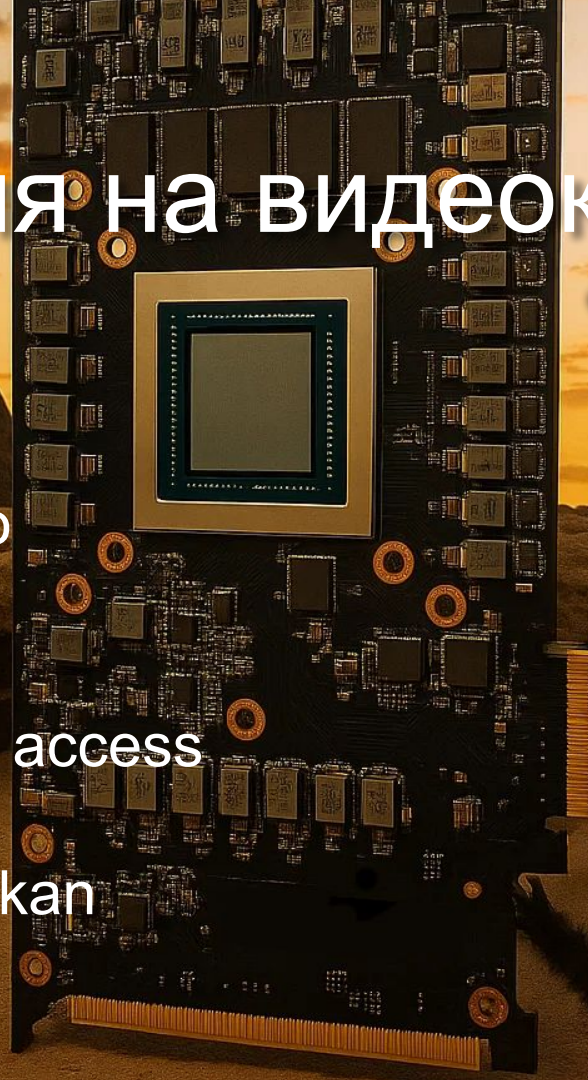


Вычисления на видеокартах

Лекция 2

- Архитектура GPU
- Модель массового параллелизма
- code divergence
- coalesced memory access
- Код кернала:
OpenCL CUDA Vulkan

 Vulkan NVIDIA
CUDA OpenCL™

План лекции

Большая часть из лекции “CS Space: Видеокарты: что они могут?”

- 1) **Вычисления:**
shared instruction pointer, code divergence
- 2) **Работа с памятью:**
hyper-threading, latency hiding, occupancy, registers pressure/spilling, coalesced memory access pattern, cache lines, local/shared memory
- 3) **Общая картина ЭВМ-архитектуры:** CPU - RAM - PCI-E - VRAM - GPU
- 4) **Модель вычислений массового параллелизма**
- 5) **Профилирование и оптимизация:** NVIDIA Nsight, SoL анализ

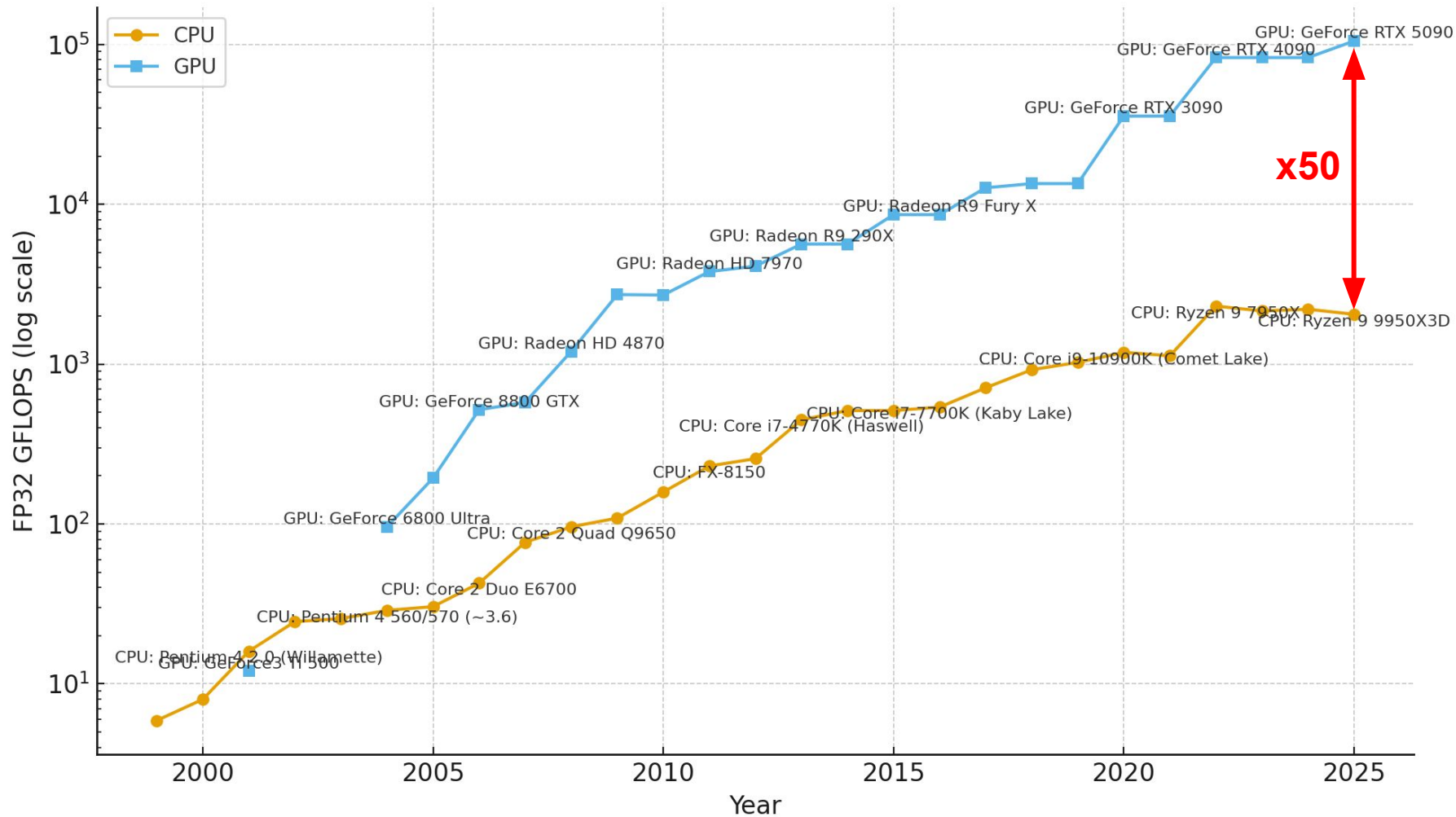
И обсудим как выглядит код на видеокарте:

- 6) **Пример кода:** A+B на C++ (OpenMP), OpenCL, CUDA, Vulkan (GLSL)

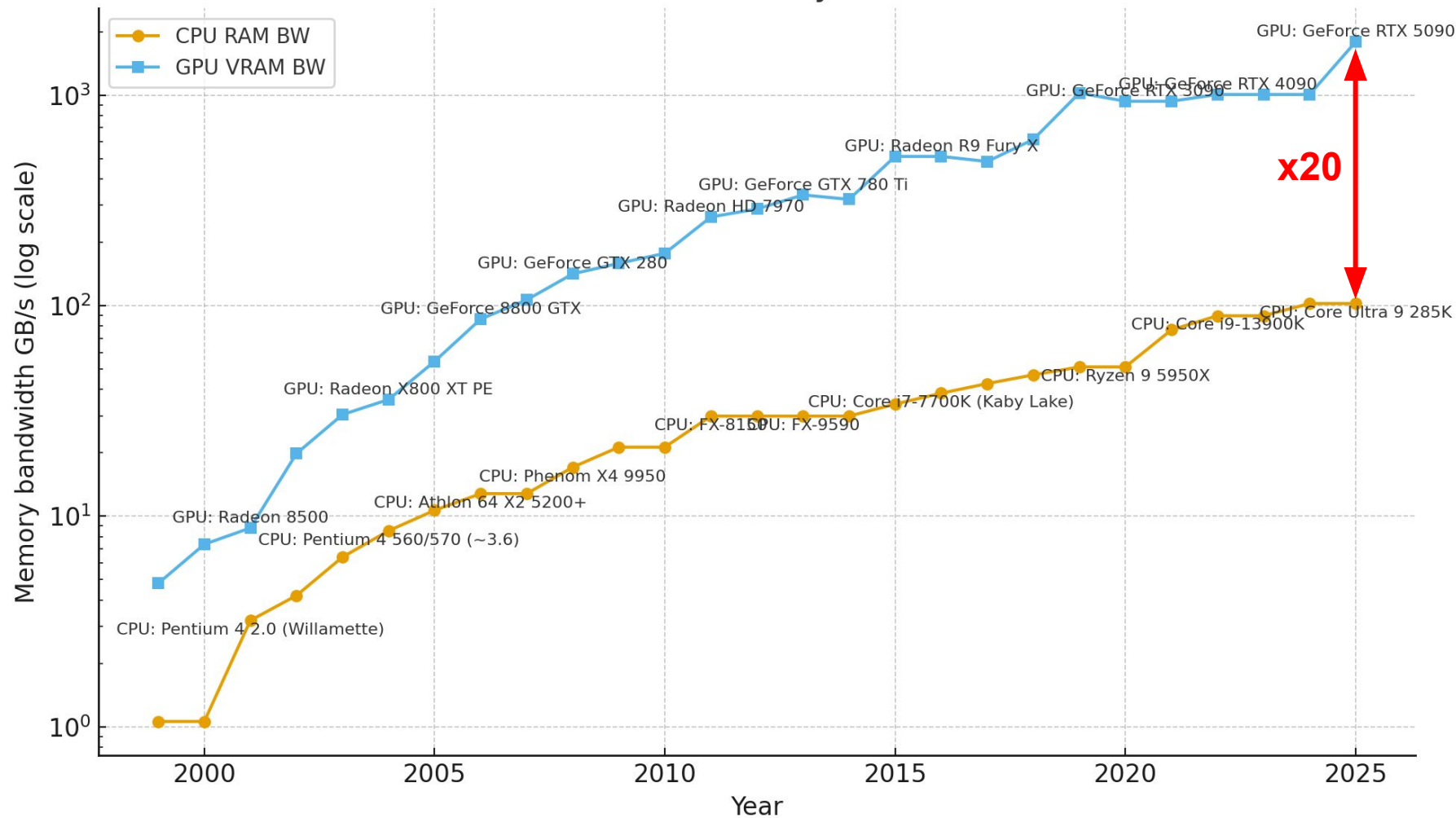
Глава 1: Вычисления

shared instruction pointer, code divergence

CPU vs GPU FP32 Peak

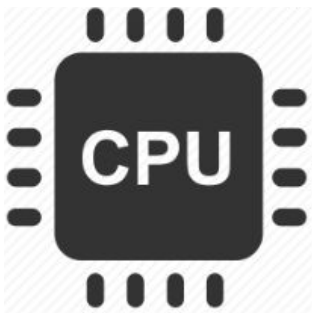


CPU vs GPU Memory Bandwidth



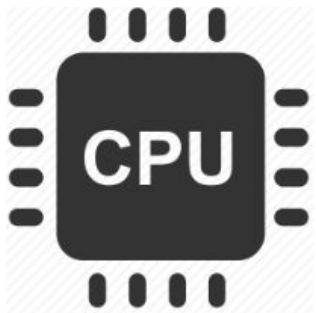
Архитектура

$100 \cdot 10^9$ FLOPS (**F**loating-point operations **p**er **s**econd)

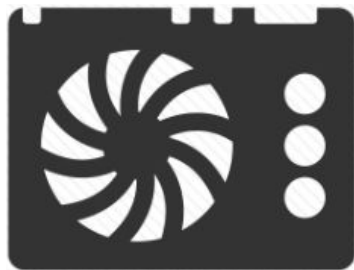


Архитектура

$100 \cdot 10^9$ FLOPS (**F**loating-point operations **p**er **s**econd)



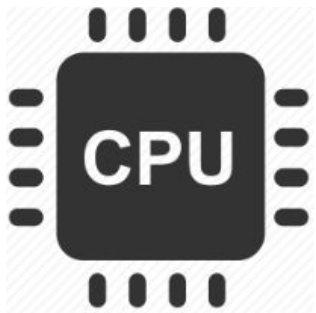
$100 \cdot 10^{12}$ FLOPS (**х1000 раз больше**)



GPU

Архитектура

$100 \cdot 10^9$ FLOPS



40 GB/s memory bandwidth

$100 \cdot 10^{12}$ FLOPS (**х1000 раз больше**)

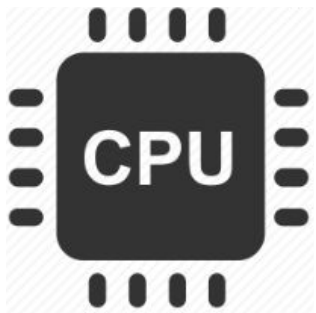


1000 GB/s (**х25 раз больше**)

GPU

Архитектура

$100 \cdot 10^9$ FLOPS

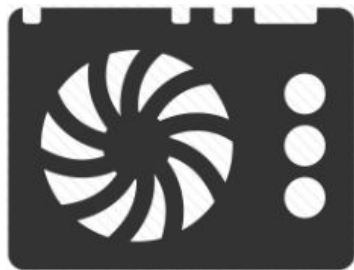


Мало ядер, но они **МОЩНЫЕ**

6 GHz



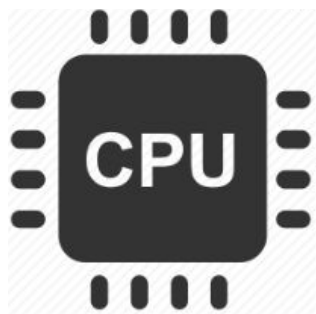
$100 \cdot 10^{12}$ FLOPS



GPU

Архитектура

$100 \cdot 10^9$ FLOPS



$100 \cdot 10^{12}$ FLOPS



GPU

Мало ядер, но они **МОЩНЫЕ**

6 GHz

ТЫСЯЧИ слабых ядер

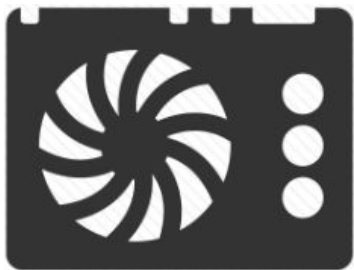
1 GHz



Архитектура

Как уместить ТЫСЯЧИ ядер?
На CPU ведь это почему-то невозможно!

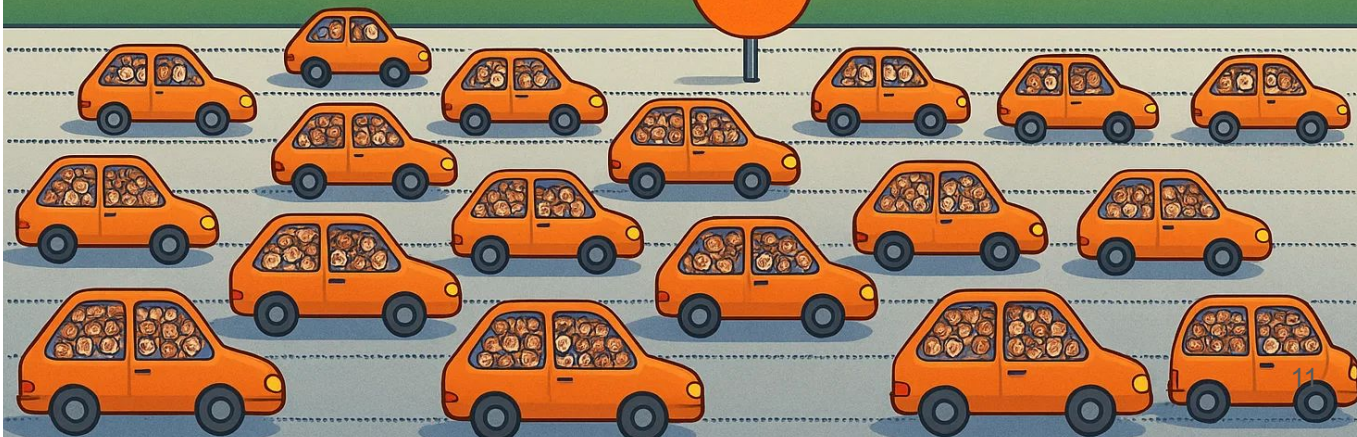
$100 \cdot 10^{12}$ FLOPS



GPU

ТЫСЯЧИ слабых ядер

1 GHz



Архитектура

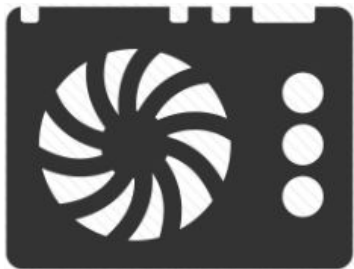


warp

32 слабых медленных
CUDA ядер

Как уместить ТЫСЯЧИ ядер?
На CPU ведь это почему-то невозможно!

$100 \cdot 10^{12}$ FLOPS



GPU



ТЫСЯЧИ слабых ядер

1 GHz

SM - Streaming Multiprocessor



Архитектура

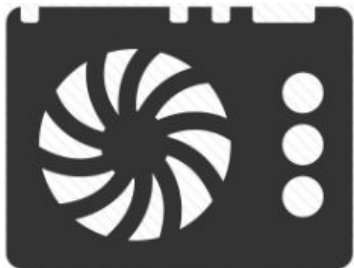


warp

32 слабых медленных
CUDA ядер

Как уместить ТЫСЯЧИ ядер?
На CPU ведь это почему-то невозможно!

$100 \cdot 10^{12}$ FLOPS



GPU



SM - Streaming Multiprocessor



RTX 3090: **10496 CUDA cores** = 82 SM · 4 warps · 32 ALUs



SM - Streaming Multiprocessor

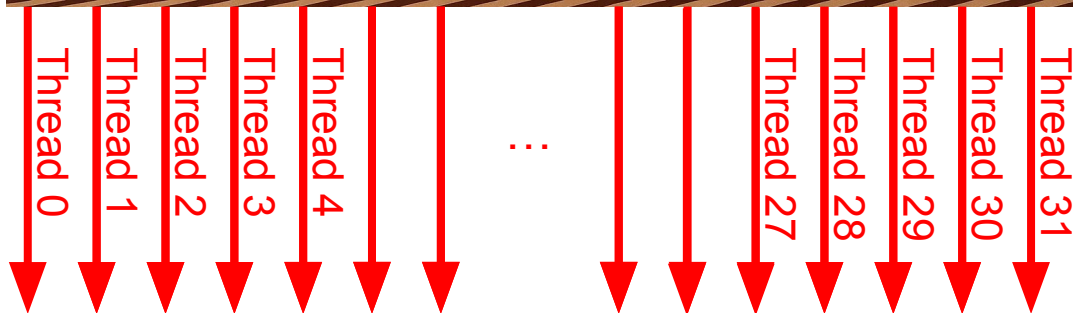


<https://www.techpowerup.com/review/nvidia-geforce-ampere-architecture-board-design-gaming-tech-software/3.html>

CPU ядро



GPU warp



Как уместить ТЫСЯЧИ ядер?
На CPU ведь это почему-то невозможно!

Архитектура

GPU warp



Thread 0

Thread 1

Thread 2

Thread 3

Thread 4

⋮

Thread 27

Thread 28

Thread 29

Thread 30

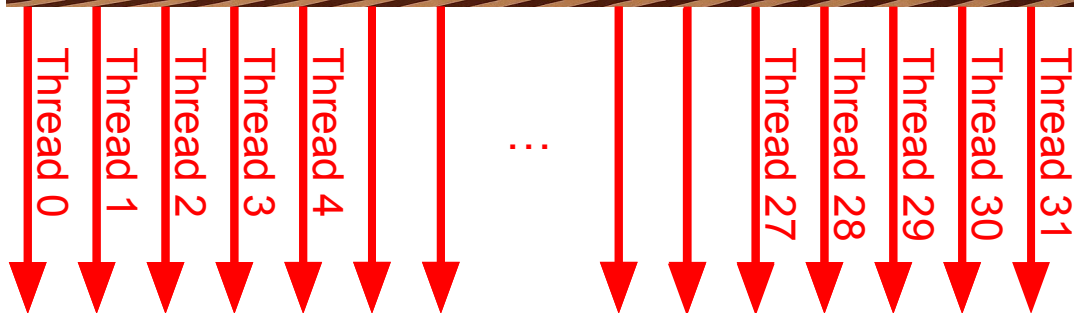
Thread 31

```
239 → int warpThreadID = threadIdx.x % 32;  
240 int result = 0;  
241 if (warpThreadID < 16) {  
242     result = dataA[warpThreadID];  
243 } else {  
244     result = dataB[warpThreadID];  
245 }
```

Один Instruction pointer
на все потоки warp-a!

Архитектура

GPU warp

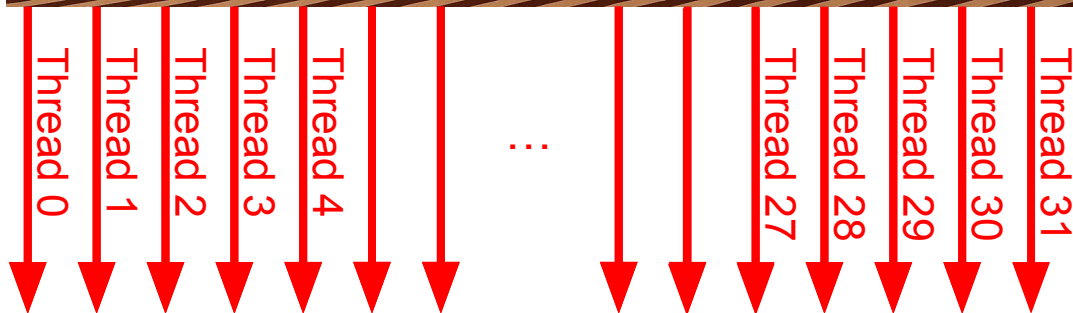


Один Instruction pointer
на все потоки warp-a!

```
239 int warpThreadID = threadIdx.x % 32;  
240 → int result = 0;  
241 if (warpThreadID < 16) {  
242     result = dataA[warpThreadID];  
243 } else {  
244     result = dataB[warpThreadID];  
245 }
```

Архитектура

GPU warp

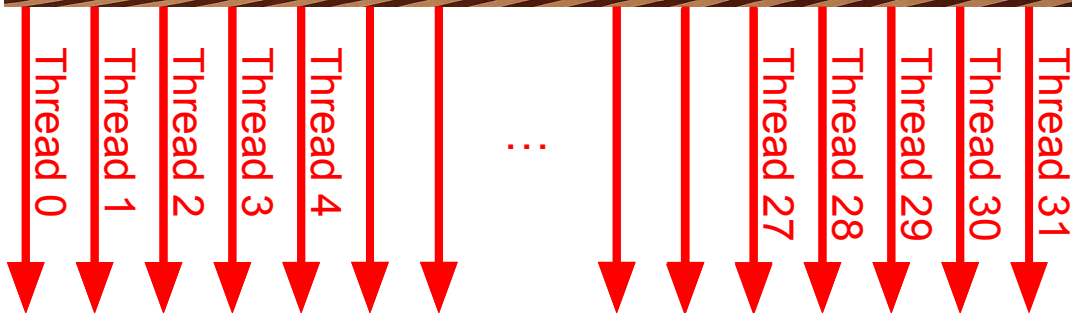


Один Instruction pointer
на все потоки warp-a!

```
239 int warpThreadID = threadIdx.x % 32;  
240 int result = 0;  
241 → if (warpThreadID < 16) {  
242     result = dataA[warpThreadID];  
243 } else {  
244     result = dataB[warpThreadID];  
245 }
```

Архитектура

GPU warp

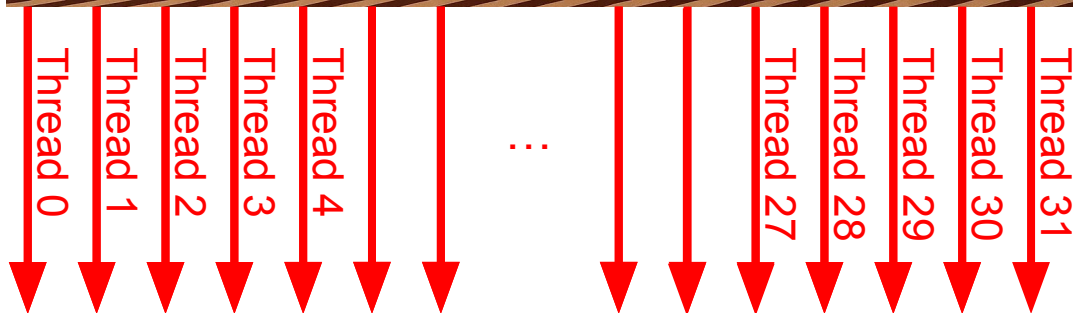


Один Instruction pointer
на все потоки warp-a!

```
239 int warpThreadID = threadIdx.x % 32;  
240 int result = 0;  
241 if (warpThreadID < 16) {  
242     result = dataA[warpThreadID];  
243 } else {  
244     result = dataB[warpThreadID];  
245 }
```

Архитектура

GPU warp

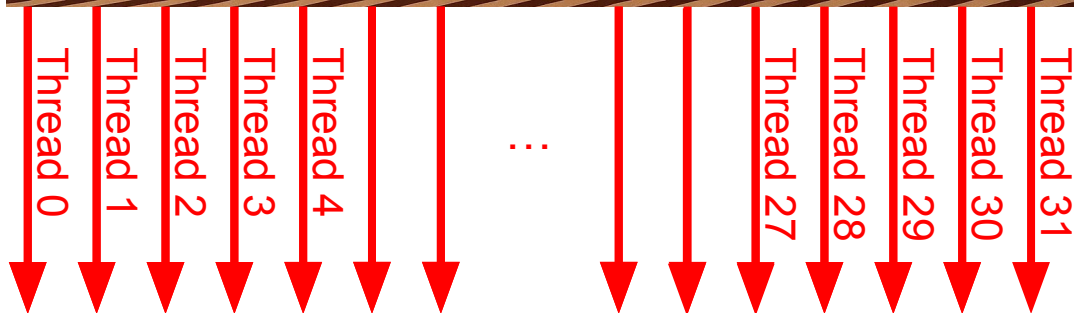


Один Instruction pointer
на все потоки warp-a!

```
239 int warpThreadID = threadIdx.x % 32;  
240 int result = 0;  
241 if (warpThreadID < 16) {  
242     result = dataA[warpThreadID];  
243 } else {  
244     result = dataB[warpThreadID];  
245 }
```

Архитектура

GPU warp

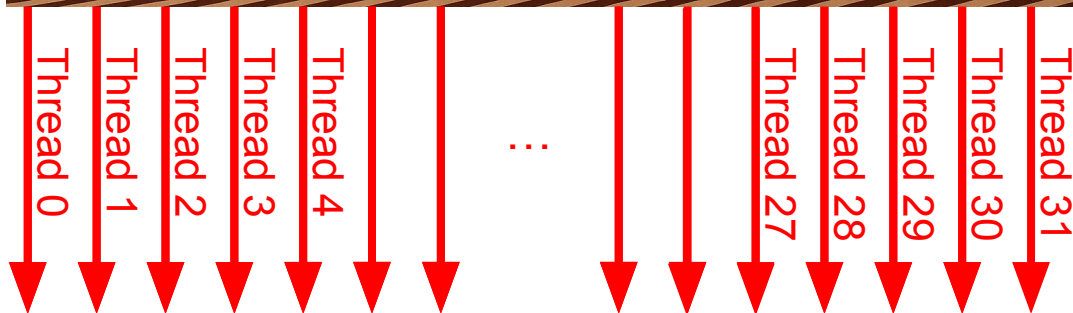


Один Instruction pointer
на все потоки warp-a!

```
239 int warpThreadID = threadIdx.x % 32;  
240 int result = 0;  
241 if (warpThreadID < 16) {  
242     result = dataA[warpThreadID];  
243 } else {  
244 → result = dataB[warpThreadID];  
245 }
```

Архитектура

GPU warp



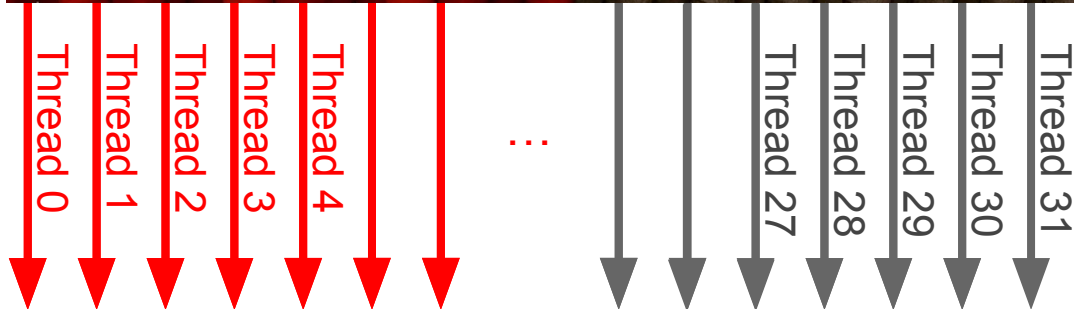
Один Instruction pointer
на все потоки warp-a!

```
239 int warpThreadID = threadIdx.x % 32;  
240 int result = 0;  
241 if (warpThreadID < 16) {  
242     result = dataA[warpThreadID];  
243 } else {  
244     result = dataB[warpThreadID];  
245 }
```

Выходит зайдём в обе ветки if-а?

Архитектура

GPU warp

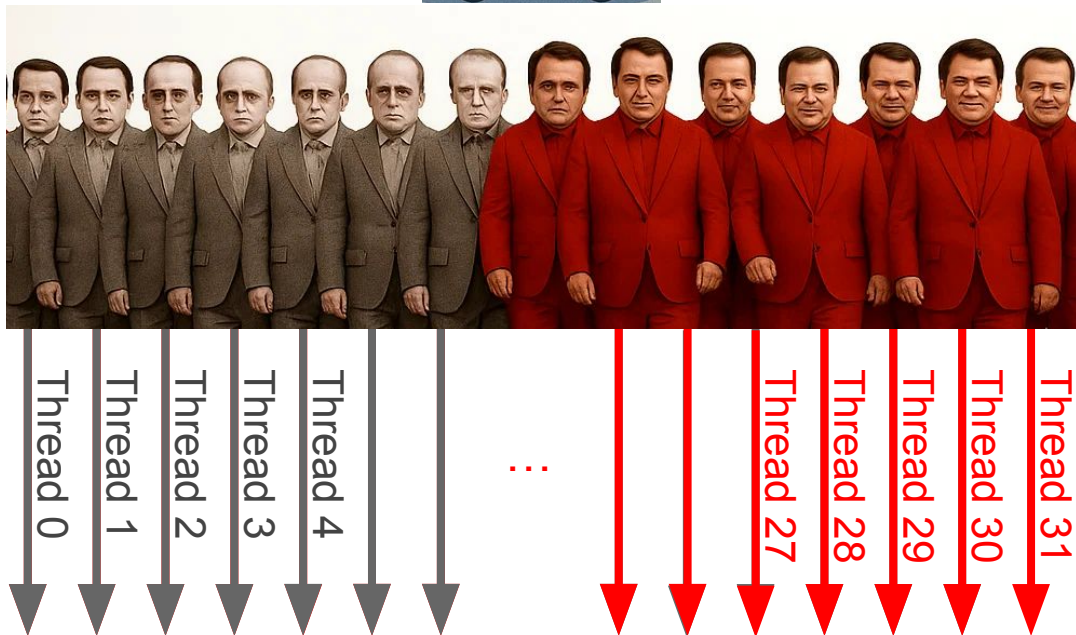
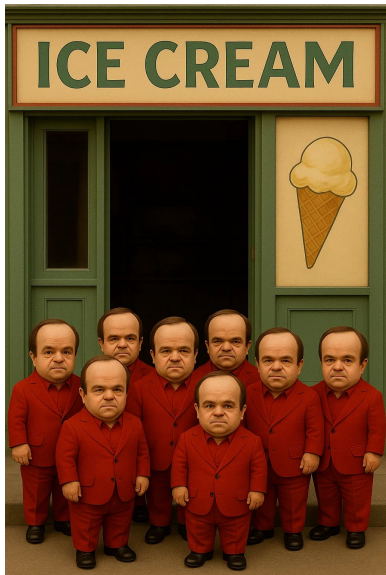


Один Instruction pointer
на все потоки warp-a!

```
239 int warpThreadID = threadIdx.x % 32;  
240 int result = 0;  
241 if (warpThreadID < 16) {  
242     result = dataA[warpThreadID]; // exec mask: [++++ ++++++ ++++++ ++++++ ---- ---- ---- ----]  
243 } else {  
244     result = dataB[warpThreadID]; // exec mask: [---- ---- ---- ---- ++++++ ++++++ ++++++ ++++++]  
245 }
```

Архитектура

GPU warp

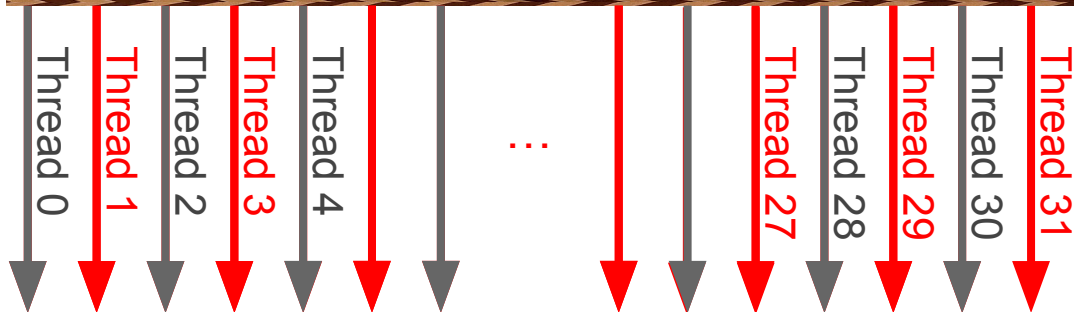


Один Instruction pointer
на все потоки warp-a!

```
239 int warpThreadID = threadIdx.x % 32;  
240 int result = 0;  
241 if (warpThreadID < 16) {  
242     result = dataA[warpThreadID]; // exec mask: [++++ ++++++ ++++++ ++++++ ---- ---- ---- ----]  
243 } else {  
244     result = dataB[warpThreadID]; // exec mask: [---- ---- ---- ---- ++++++ ++++++ ++++++ ++++++]  
245 }
```

Архитектура

GPU warp



Один Instruction pointer
на все потоки warp-a!

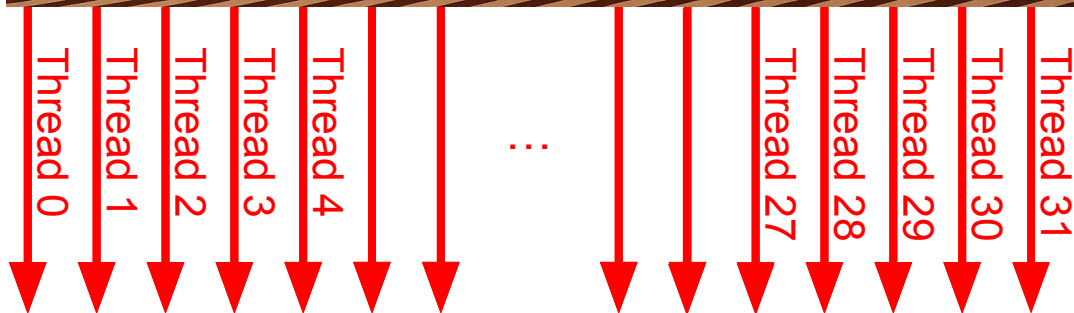
```
239 int warpThreadID = threadIdx.x % 32;  
240 int result = 0;  
241 if (warpThreadID % 2 == 0) {  
242     result = dataA[warpThreadID]; // exec mask: [+--+ +--+ +--+ +--+ +--+ +--+ +--+]  
243 } else {  
244     result = dataB[warpThreadID]; // exec mask: [-+++ -+++ -+++ -+++ -+++ -+++ -+++ -+++]  
245 }
```

Архитектура

GPU warp



```
241 if (predicate1) {  
242     if (predicate2) {  
243         if (predicate3) {  
244             // A1  
245         } else {  
246             // A2  
247         }  
248     } elif (predicate4) {  
249         // A3  
250     }  
251 } else {  
252     // A4  
253 }
```

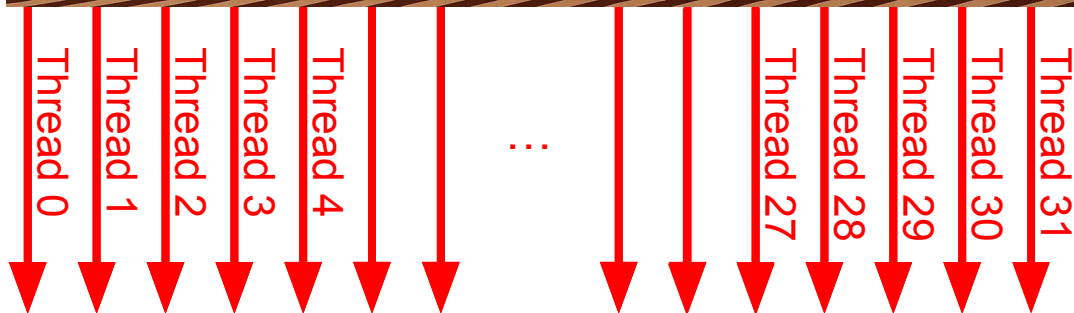


Один Instruction pointer
на все потоки warp-a!

Сколько “времени” будет работать warp?

Архитектура

GPU warp



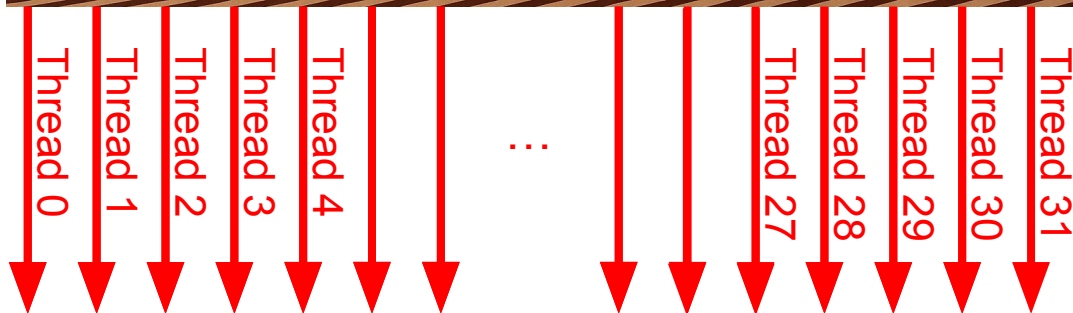
Один Instruction pointer
на все потоки warp-a!

```
241 if (predicate1) {  
242     if (predicate2) {  
243         if (predicate3) {  
244             // A1  
245         } else {  
246             // A2  
247         }  
248     } elif (predicate4) {  
249         // A3  
250     }  
251 } else {  
252     // A4  
253 }
```

Сколько “времени” будет работать warp?
 $A1 + A2 + A3 + A4$ при **code divergence**

Архитектура

GPU warp

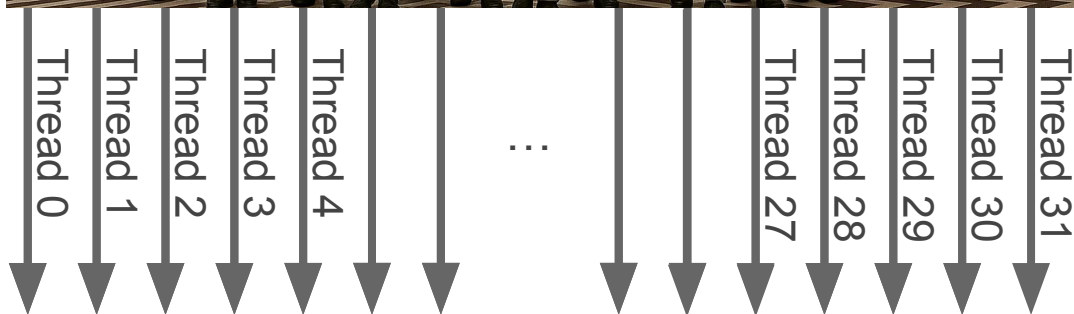
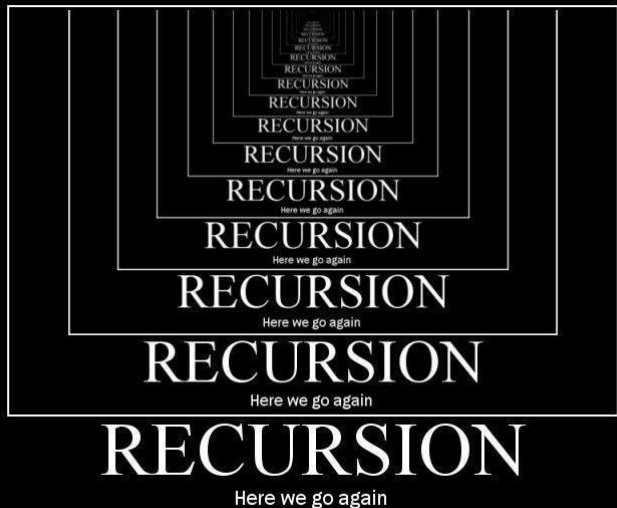


Один Instruction pointer
на все потоки warp-a!

Сколько “времени” будет работать достаточно глубокая рекурсия?

Архитектура

GPU warp

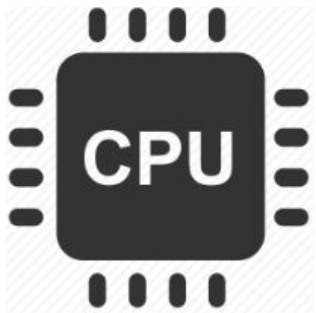


Один Instruction pointer
на все потоки warp-a!

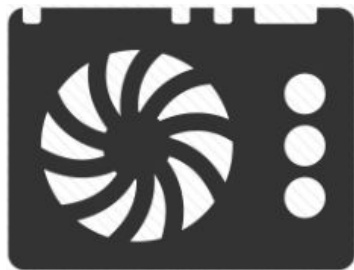
Глава 2: Работа с памятью

hyper-threading, latency hiding, occupancy, registers pressure/spilling,
coalesced memory access pattern, cache lines, local/shared memory, bank conflicts

Архитектура



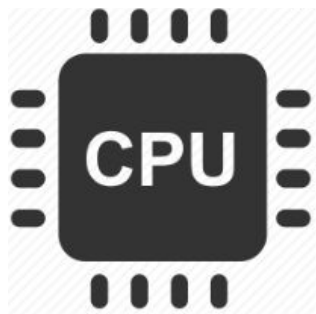
40
GB/s



1000
GB/s

GPU

Архитектура



40
GB/s

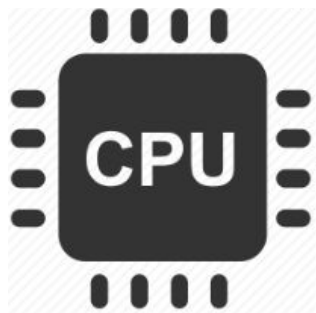
Память - малая пропускная способность
НИЗКАЯ LATENCY



1000
GB/s

GPU

Архитектура



40
GB/s



GPU

1000
GB/s

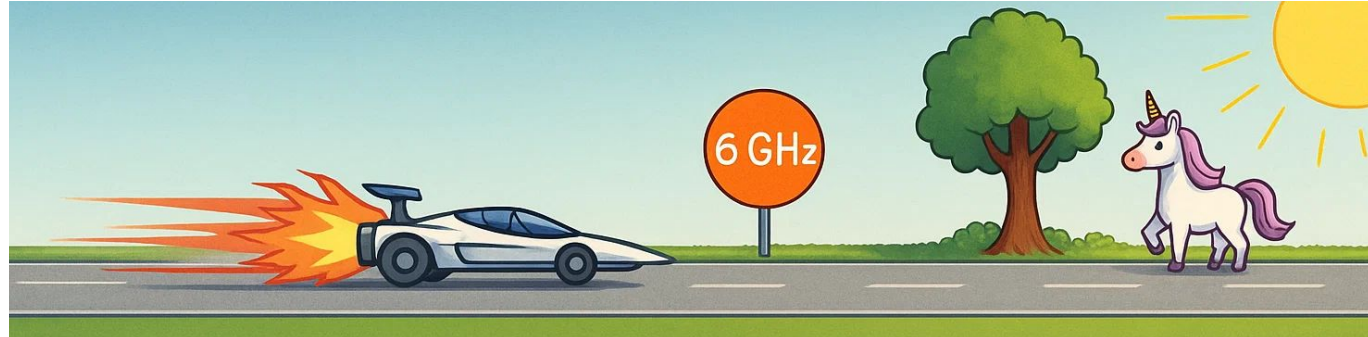
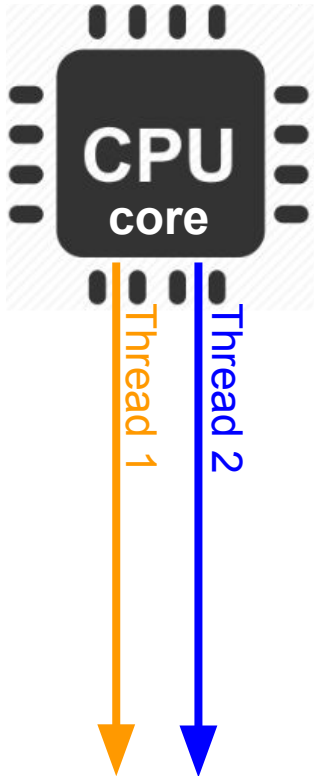
Память - малая пропускная способность
НИЗКАЯ LATENCY



Память - **ОГРОМНАЯ** пропускная способность
но большая latency



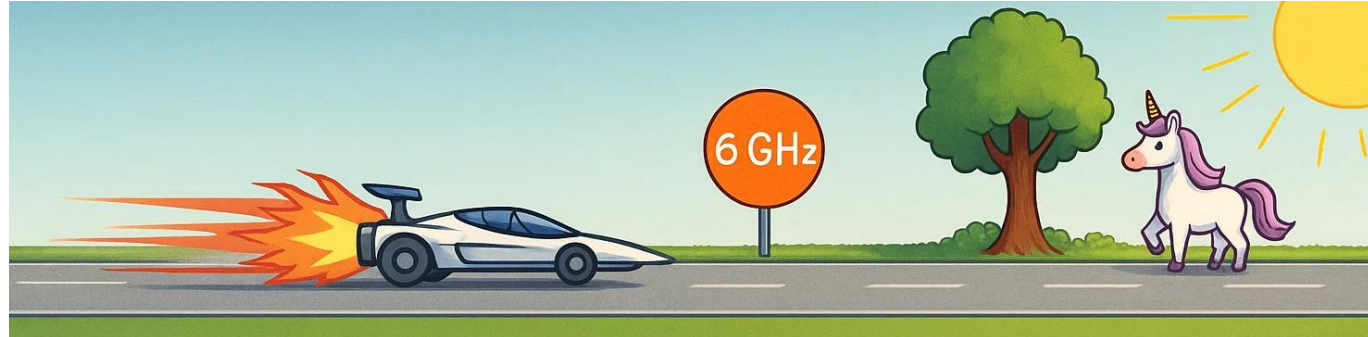
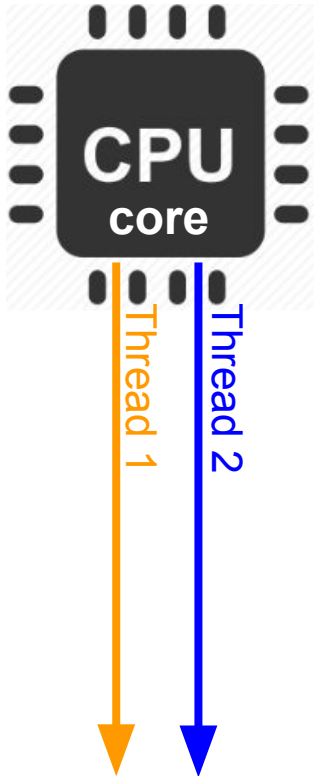
Архитектура CPU: многопоточность



Многопоточность благодаря переключению контекста!

Что нужно перещелкнуть в состоянии ЦПУ ядра для другого потока?

Архитектура CPU: многопоточность

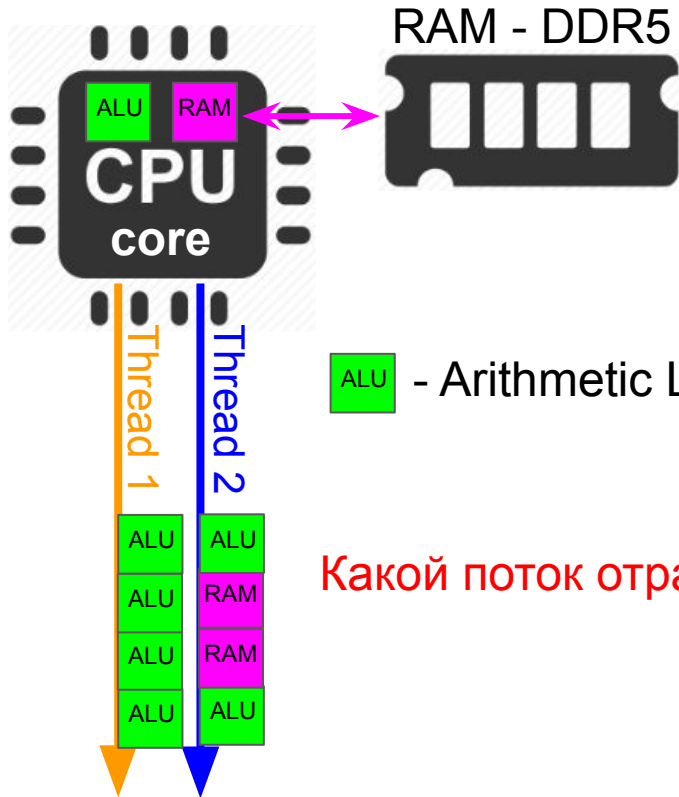


Многопоточность благодаря переключению контекста!

Context switch:

- Обновляет Instruction pointer (указатель на строку кода)
- Подгружает значения регистров процессора

Архитектура CPU: Hyper-Threading, SMT

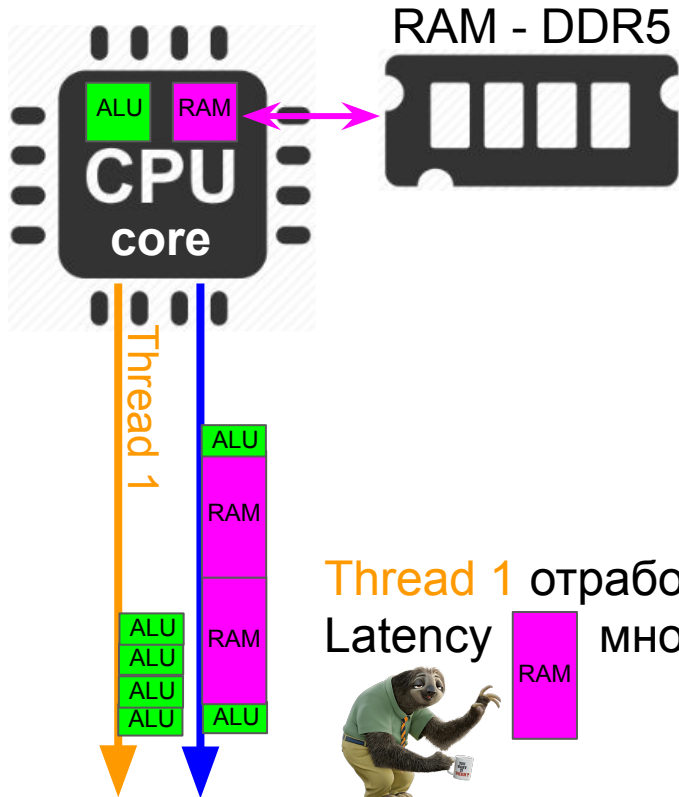


 - Arithmetic Logical Unit

Какой поток отработает быстрее?



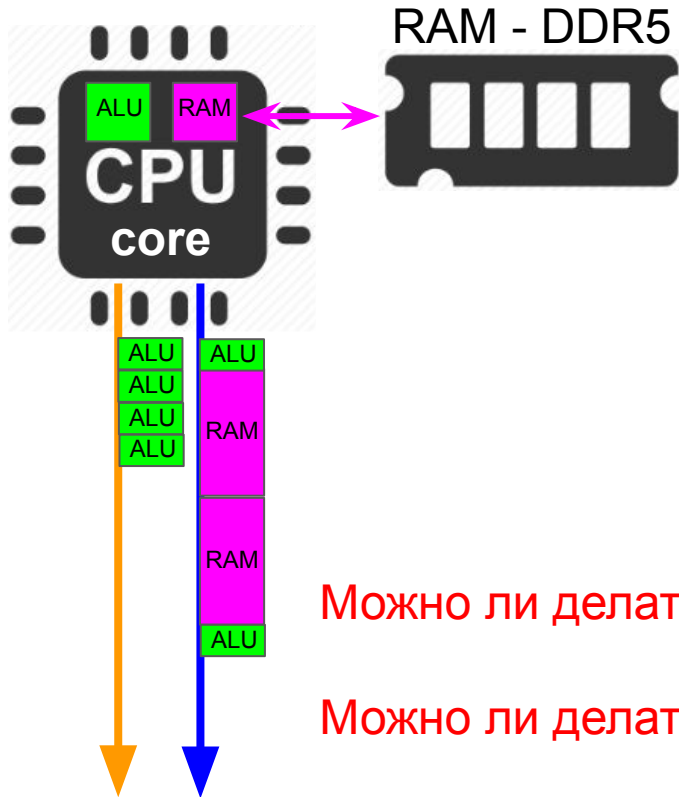
Архитектура CPU: Hyper-Threading, SMT



Thread 1 отработает быстрее!
Latency много больше чем у



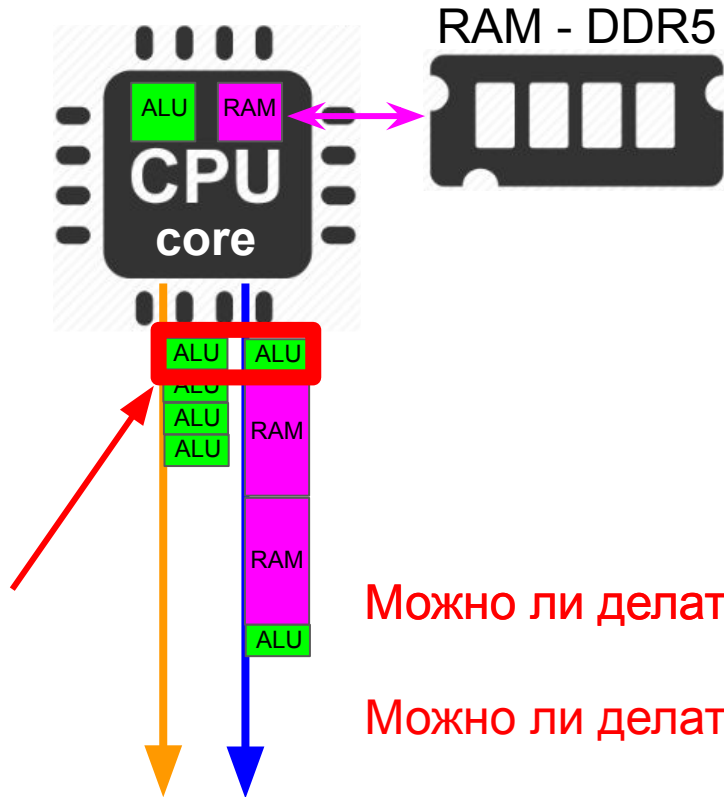
Архитектура CPU: Hyper-Threading, SMT



Можно ли делать **ALU** пока ждем **RAM** ?

Можно ли делать **ALU** пока ждем **ALU** ?

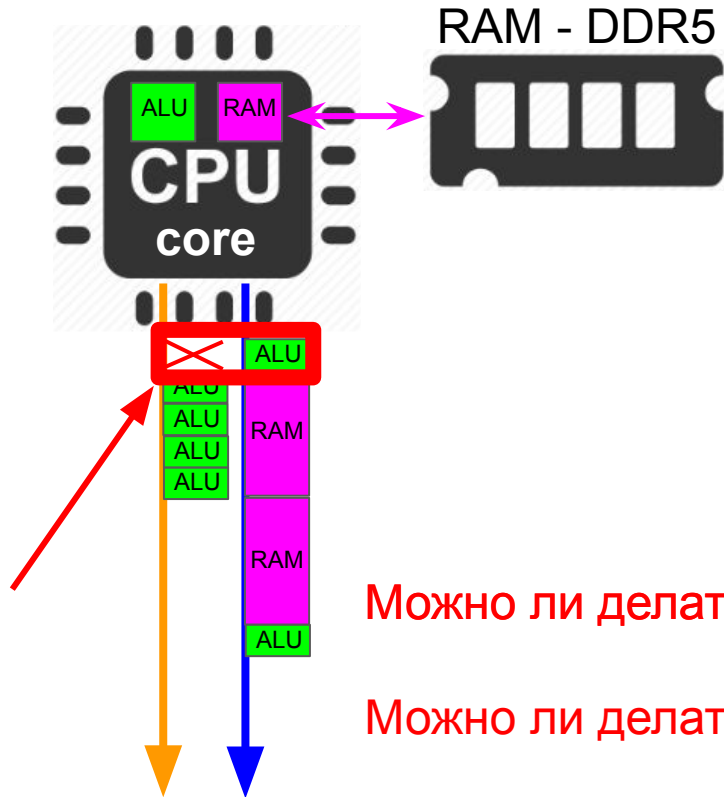
Архитектура CPU: Hyper-Threading, SMT



Можно ли делать **ALU** пока ждем **RAM** ?

Можно ли делать **ALU** пока ждем **ALU** ?

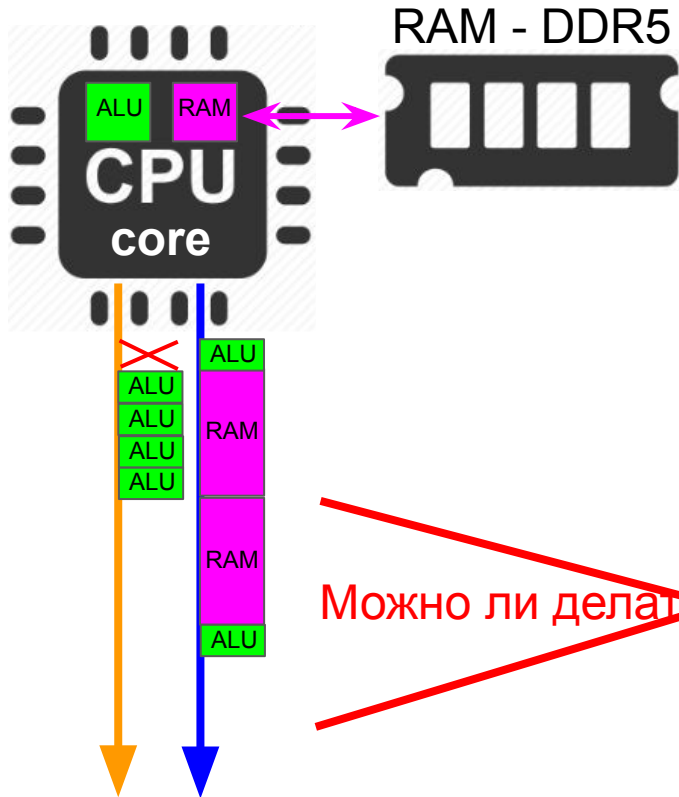
Архитектура CPU: Hyper-Threading, SMT



Можно ли делать **ALU** пока ждем **RAM** ?

Можно ли делать **ALU** пока ждем **ALU** ?

Архитектура CPU: Hyper-Threading, SMT

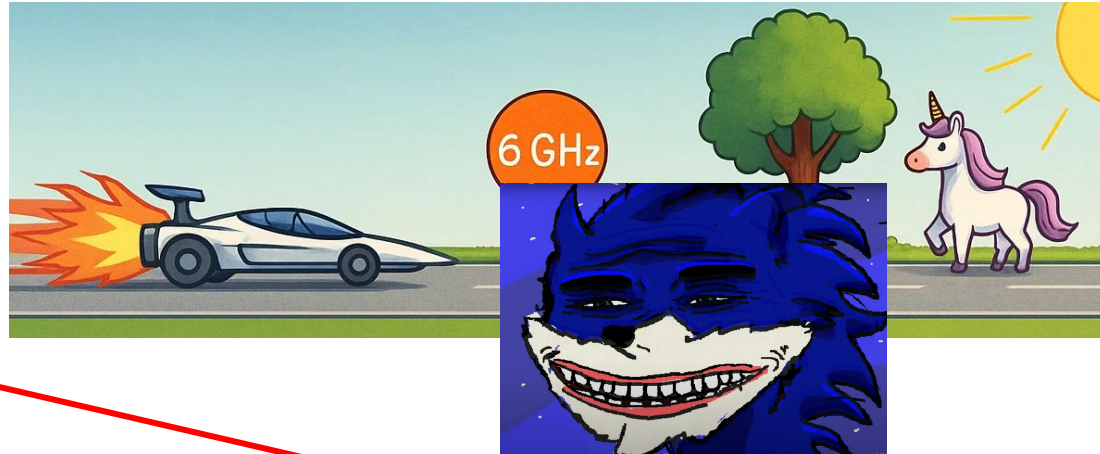
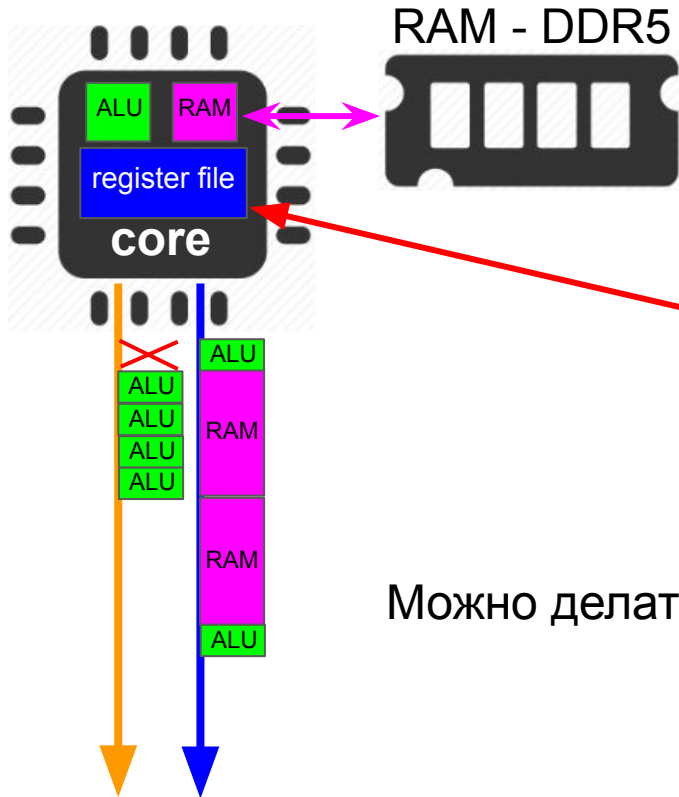


Можно ли делать ALU пока ждем RAM ?

Context switch - дорого!
Долго подгружает регистры!



Архитектура CPU: Hyper-Threading, SMT

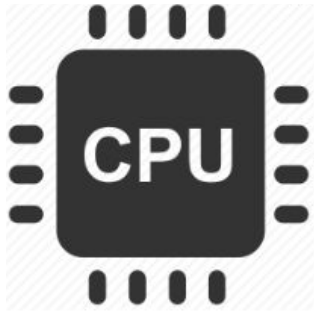


Давайте держать два множества регистров!

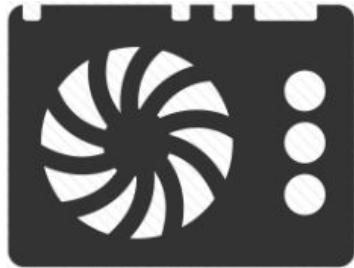
Можно делать **ALU** пока ждем **RAM** !

Это и есть SMT и HT!

Архитектура GPU



40
GB/s



1000
GB/s

Память - малая пропускная способность
НИЗКАЯ LATENCY



Память - **ОГРОМНАЯ** пропускная способность
но большая latency



Архитектура GPU



warp

32 слабых медленных
CUDA ядер

SM - Streaming Multiprocessor



Архитектура GPU



wavefront
64 threads



nvidia
warp
32 threads



warp

32 слабых медленных
CUDA ядер

SM - Streaming Multiprocessor



Архитектура GPU



wavefront

64 threads

**Compute Units
(CU)**



nVIDIA®

warp

32 threads

**Streaming Multiprocessor
(SM)**



warp

32 слабых медленных
CUDA ядер

SM - Streaming Multiprocessor



Архитектура GPU



warp

32 слабых медленных
CUDA ядер

SM - Streaming Multiprocessor



Архитектура GPU



warp

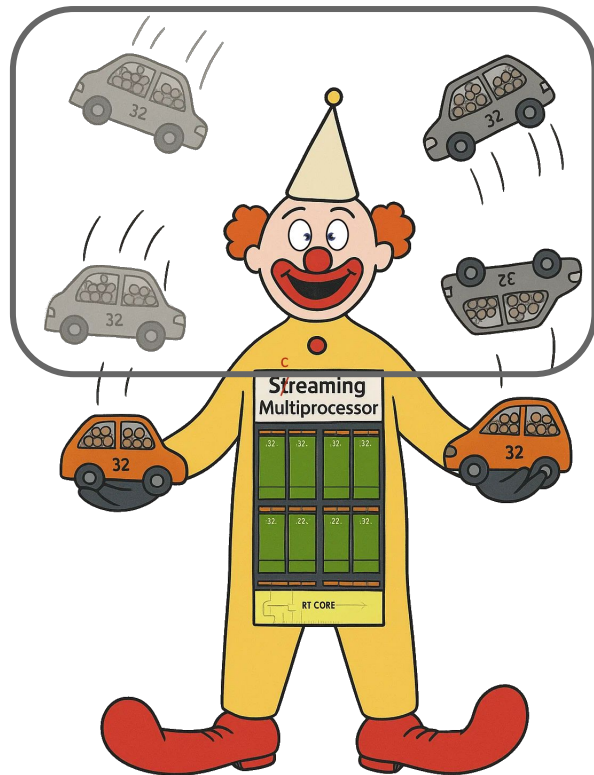
32 слабых медленных
CUDA ядер

Register File для быстрого
жонглирования warp-ами!
(быстрый "context switch")

SM - Streaming Multiprocessor



Архитектура GPU



warp

32 слабых медленных
CUDA ядер

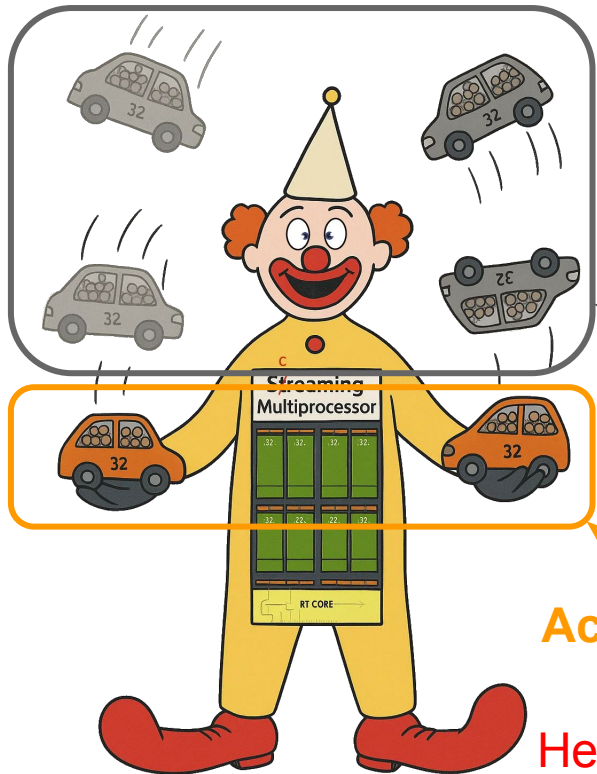
Register File для быстрого
жонглирования warp-ами!
(быстрый "context switch")

Inactive warps: ждут когда
придут данные из VRAM

SM - Streaming Multiprocessor



Архитектура GPU



warp

32 слабых медленных
CUDA ядер

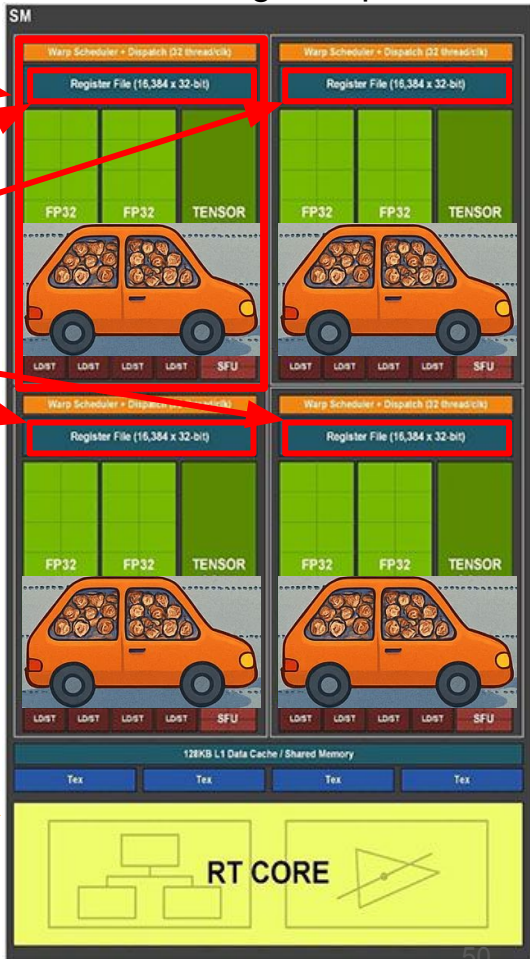
Register File для быстрого
жонглирования warp-ами!
(быстрый "context switch")

Inactive warps: ждут когда
придут данные из VRAM

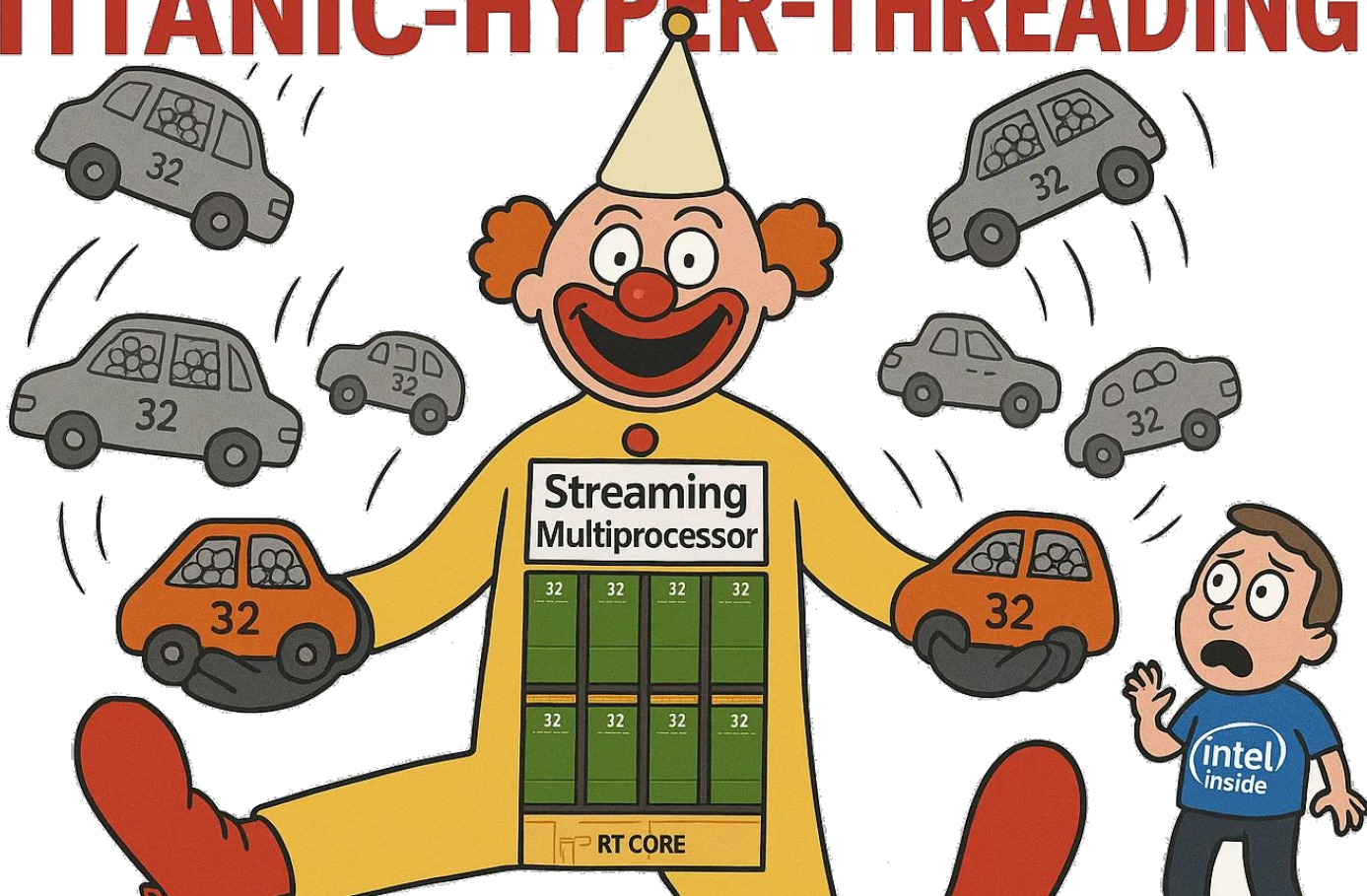
Active warps: данные в регистрах
Работаем-работаем! 🧐

Не напоминает ли вам это что-то?

SM - Streaming Multiprocessor



ULTRA SUPER MEGA GIGA TITANIC-HYPER-THREADING



Архитектура GPU



warp

32 слабых медленных
CUDA ядер

Register File для быстрого
жонглирования warp-ами!
(быстрый "context switch")

Occupancy = насколько
много доступно warp-ов для
жонглирования.

SM - Streaming Multiprocessor



Архитектура GPU



warp

32 слабых медленных
CUDA ядер

Register File для быстрого
жонглирования warp-ами!
(быстрый "context switch")

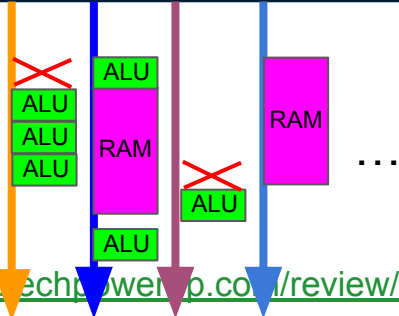
Осцирансу = насколько
много доступно warp-ов для
жонглирования.

**Если осцирансу высокая,
то что?**

SM - Streaming Multiprocessor



Архитектура GPU



warp

32 слабых медленных
CUDA ядер

Register File для быстрого
жонглирования warp-ами!
(быстрый "context switch")

Осцирансу = насколько
много доступно warp-ов для
жонглирования.

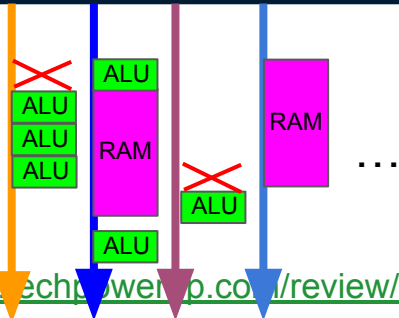
Если осцирансу высокая, то
хорошо скрывается latency
доступа к памяти т.к. всегда
найдется готовый warp!

SM - Streaming Multiprocessor



Архитектура GPU

А почему у SM может быть разное число warp-ов?



warp

32 слабых медленных
CUDA ядер

Register File для быстрого
жонглирования warp-ами!
(быстрый “context switch”)

Осцирансу = насколько
много доступно warp-ов для
жонглирования.

Если осцирансу высокая, то
хорошо скрывается latency
доступа к памяти т.к. всегда
найдется готовый warp!

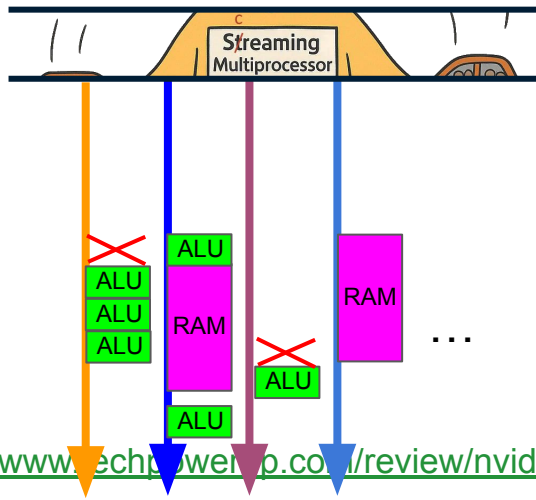
SM - Streaming Multiprocessor



Архитектура GPU

А почему у SM может быть
разное число warp-ов?

Registers Pressure!



32 слабых медленных
CUDA ядер

Register File для быстрого
жонглирования warp-ами!
(быстрый "context switch")

Осцирансу = насколько
много доступно warp-ов для
жонглирования.

Если осцирансу высокая, то
хорошо скрывается latency
доступа к памяти т.к. всегда
найдется готовый warp!



warp

SM - Streaming Multiprocessor

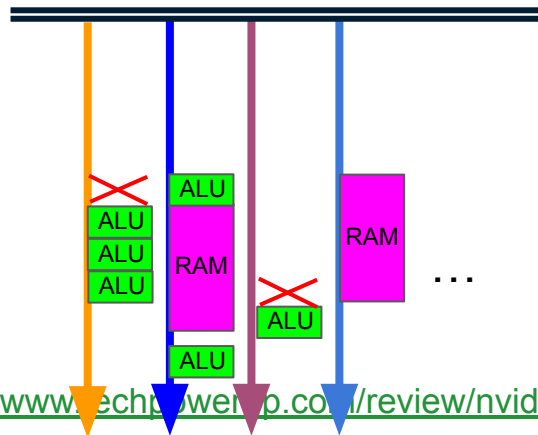


Архитектура GPU

А почему у SM может быть
разное число warp-ов?

Registers Pressure!

Вплоть до **Registers Spilling!**



Register File для быстрого
жонглирования warp-ами!
(быстрый “context switch”)

Occupancy = насколько
много доступно warp-ов для
жонглирования.

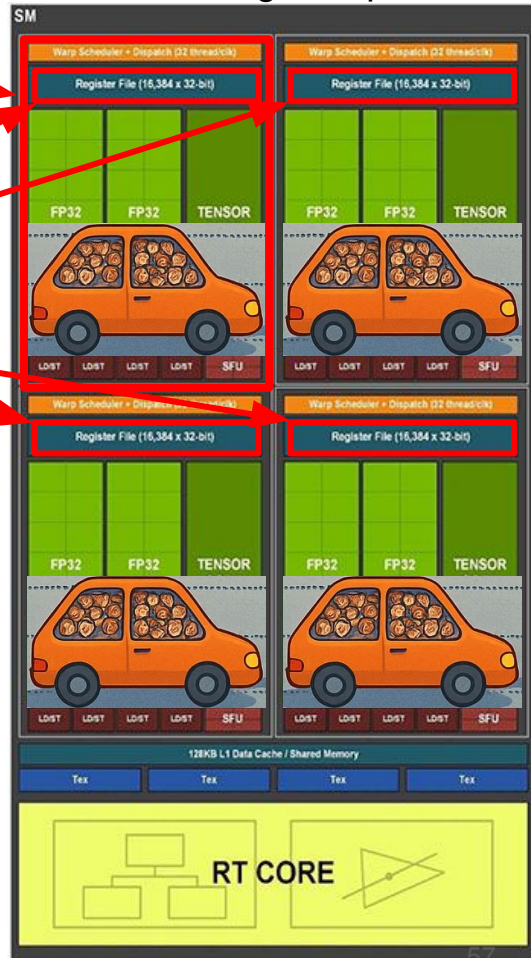
Если occupancy высокая, то
хорошо скрывается latency
доступа к памяти т.к. всегда
найдется готовый warp!



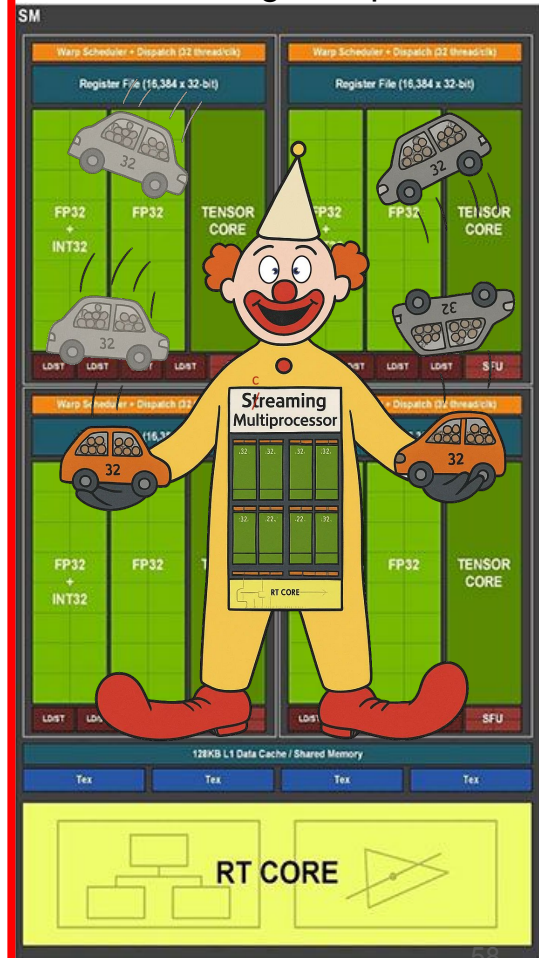
warp

32 слабых медленных
CUDA ядер

SM - Streaming Multiprocessor

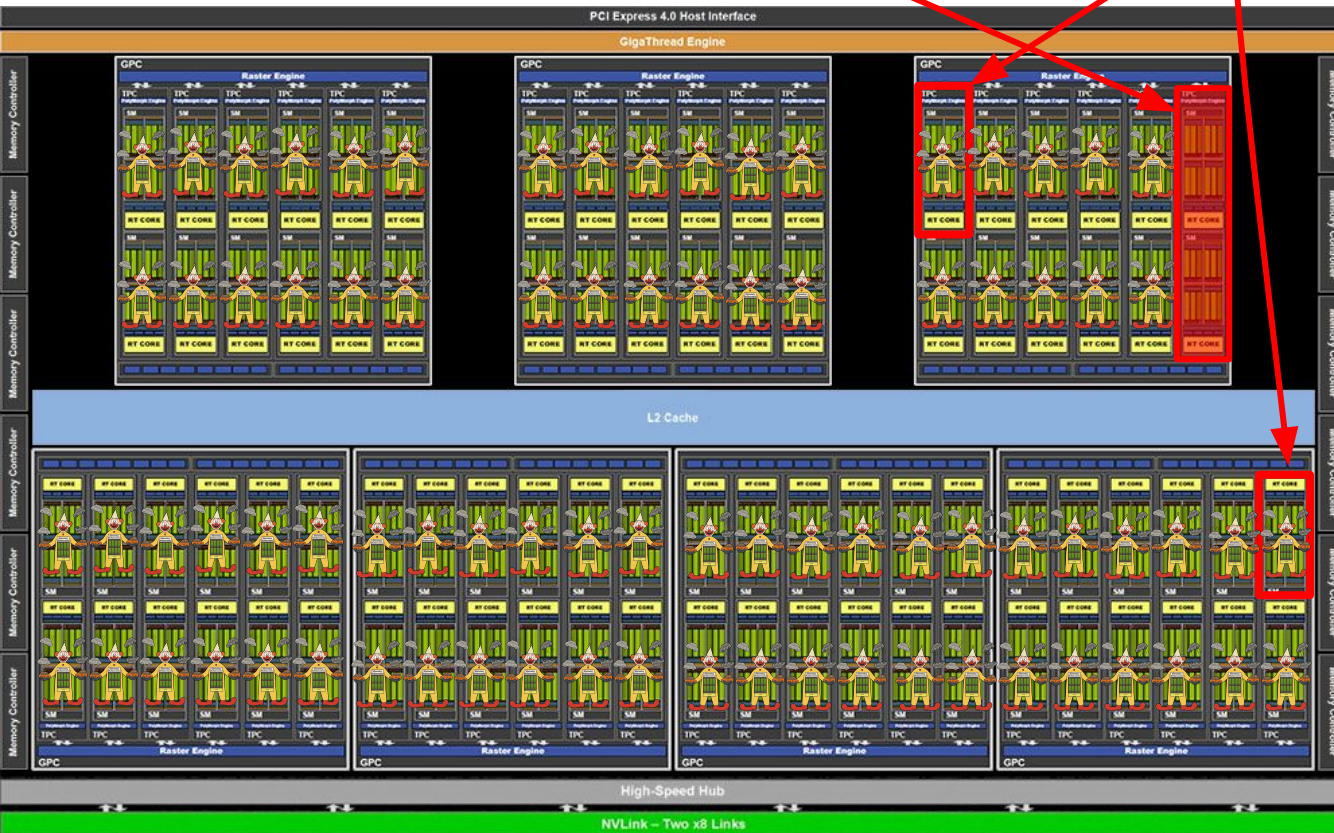


SM - Streaming Multiprocessor



RTX 3090: 10496 CUDA cores = 82 SM · 4 warps · 32 ALUs

Куда сбежали два клоуна?



SM - Streaming Multiprocessor

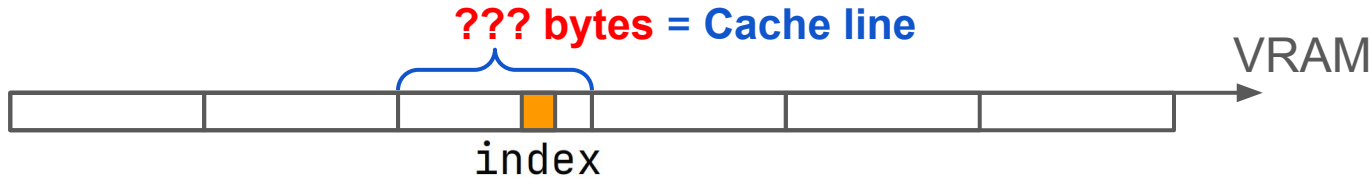


Архитектура VRAM (**coalesced** memory access pattern)



```
float value = data[index];
```

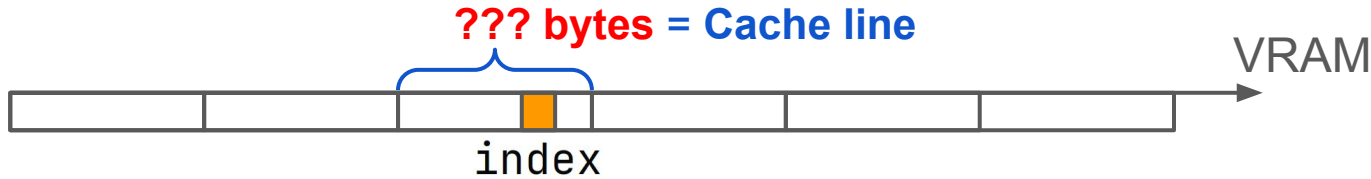
Архитектура VRAM (**coalesced** memory access pattern)



```
float value = data[index];
```

Какого размера вы бы решили
сделать **Cache line**?

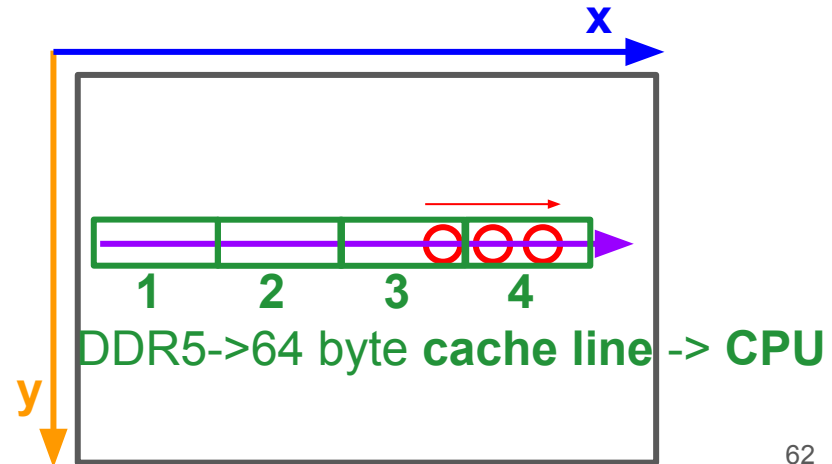
Архитектура VRAM (**coalesced** memory access pattern)



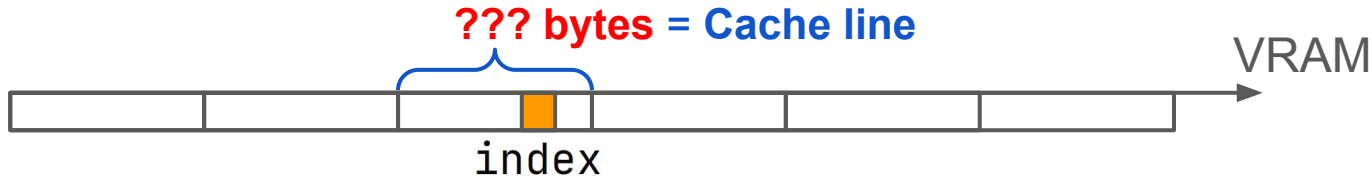
```
float value = data[index];
```

Какого размера вы бы решили сделать **Cache line**?

Напоминание про CPU:



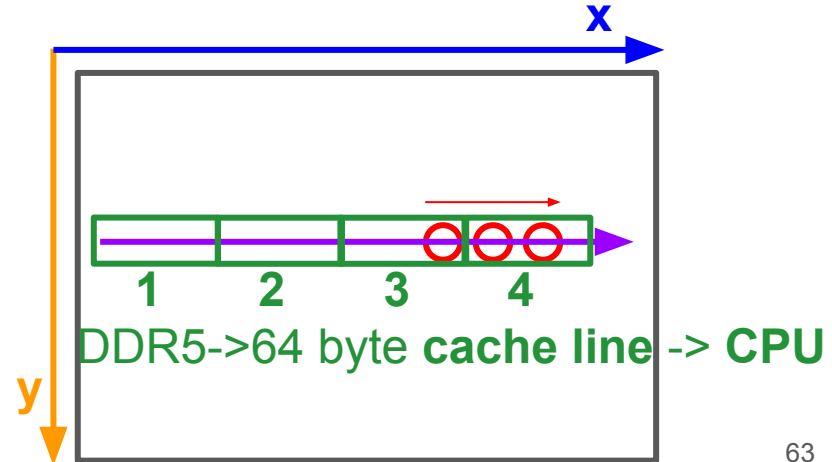
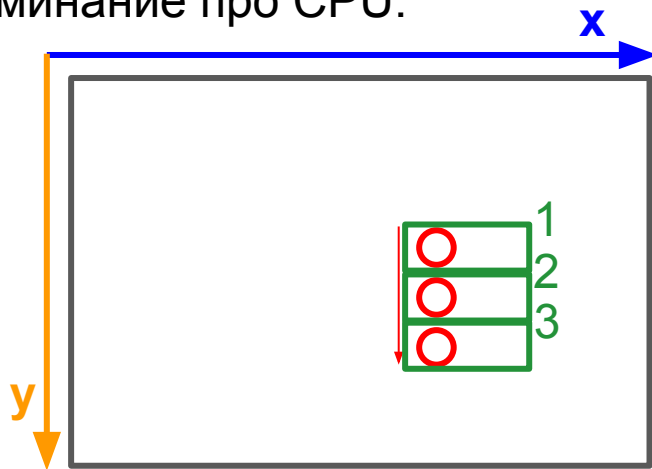
Архитектура VRAM (**coalesced** memory access pattern)



```
float value = data[index];
```

Какого размера вы бы решили сделать **Cache line**?

Напоминание про CPU:



Архитектура VRAM (**coalesced** memory access pattern)

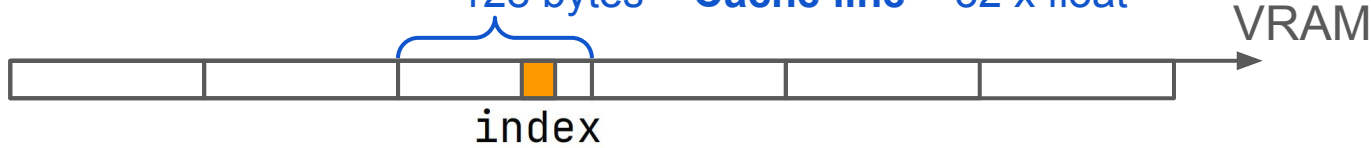
128 bytes = **Cache line** = 32 x float



```
float value = data[index];
```

Архитектура VRAM (**coalesced** memory access pattern)

128 bytes = **Cache line** = 32 x float



```
float value = data[index];
```



Архитектура VRAM (**coalesced** memory access pattern)

128 bytes = **Cache line** = 32 x float



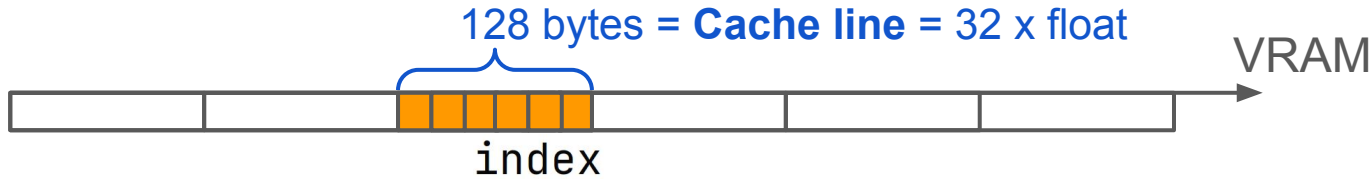
float value = data[index]; index | cache lines count

[1024+0; 1024+32) ???

Сколько **cache line**-ов чтобы покрыть заказ?
Сколько потребуется **транзакций** из VRAM?



Архитектура VRAM (**coalesced** memory access pattern)



<code>float value = data[index];</code>	index	cache lines count
---	-------	-------------------

<code>[1024+0; 1024+32)</code>	1 транзакция (coalesced)
--------------------------------	--------------------------------------



Архитектура VRAM (**coalesced** memory access pattern)

128 bytes = **Cache line** = 32 x float



float value = data[index]; index | cache lines count

[1024+0; 1024+32)

1 транзакция
(**coalesced**)

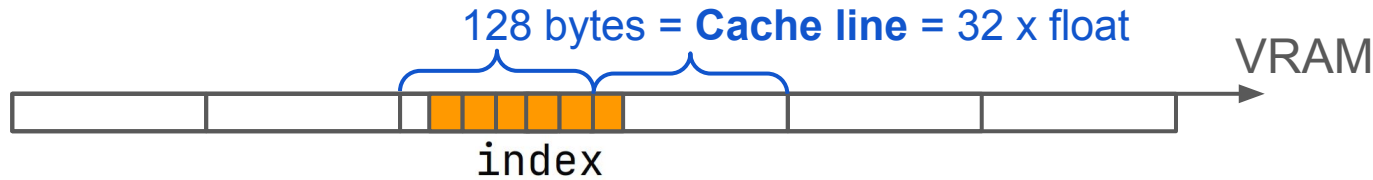
[1024+1; 1024+33)

???

Сколько **cache line**-ов чтобы покрыть заказ?
Сколько потребуется **транзакций** из VRAM?



Архитектура VRAM (**coalesced** memory access pattern)



`float value = data[index];` index | cache lines count

[1024+0; 1024+32)

1 транзакция
(**coalesced**)

[1024+1; 1024+33)

2 транзакции
(**coalesced**)

Насколько просела достигнутая
полезная пропускная способность VRAM?

Архитектура VRAM (**coalesced** memory access pattern)

128 bytes = **Cache line** = 32 x float



float value = data[index]; index | cache lines count

[1024+0; 1024+32)

1 транзакция
(**coalesced**)

[1024+1; 1024+33)

2 транзакции
(**coalesced**)

{1024 + i*32}

???

Сколько **cache line**-ов чтобы покрыть заказ?
Сколько потребуется **транзакций** из VRAM?



Архитектура VRAM (**coalesced** memory access pattern)

128 bytes = **Cache line** = 32 x float



float value = data[index]; index | cache lines count

[1024+0; 1024+32)

1 транзакция
(**coalesced**)

[1024+1; 1024+33)

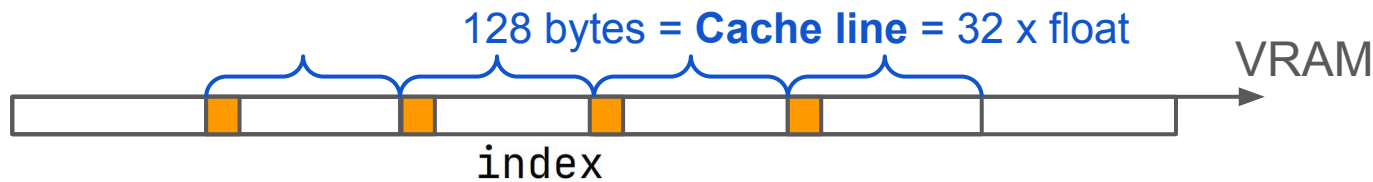
2 транзакции
(**coalesced**)

{1024 + i*32}

???

Сколько **cache line**-ов чтобы покрыть заказ?
Сколько потребуется **транзакций** из VRAM?

Архитектура VRAM (**coalesced** memory access pattern)



`float value = data[index];` index cache lines count



`[1024+0; 1024+32)`

1 транзакция
(**coalesced**)

`[1024+1; 1024+33)`

2 транзакции
(**coalesced**)

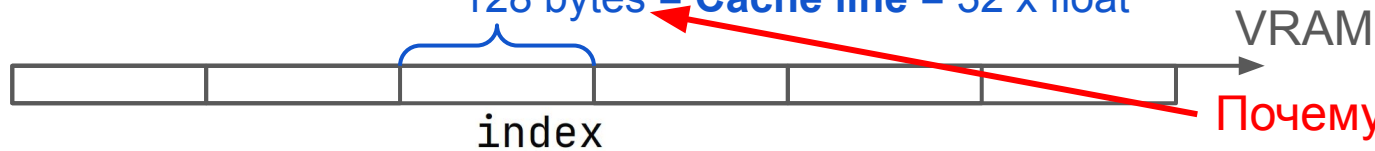
`{1024 + i*32}`

32 транзакции
(**uncoalesced**)

Насколько просела достигнутая
полезная пропускная способность VRAM?⁷²

Архитектура VRAM (**coalesced** memory access pattern)

128 bytes = **Cache line** = 32 x float



Почему именно 128 байт?

`float value = data[index];` index cache lines count

[1024+0; 1024+32)

1 транзакция
(**coalesced**)

[1024+1; 1024+33)

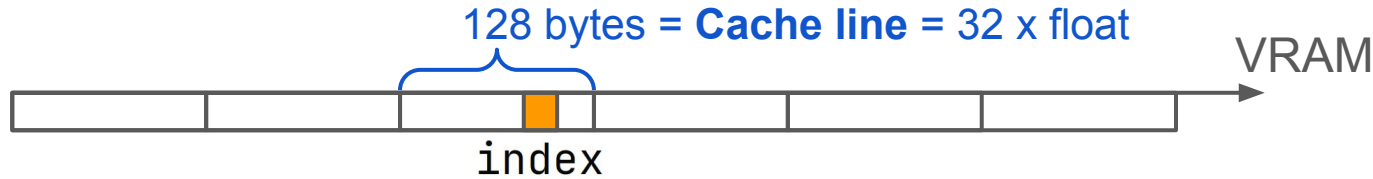
2 транзакции
(**coalesced**)

{1024 + i*32}

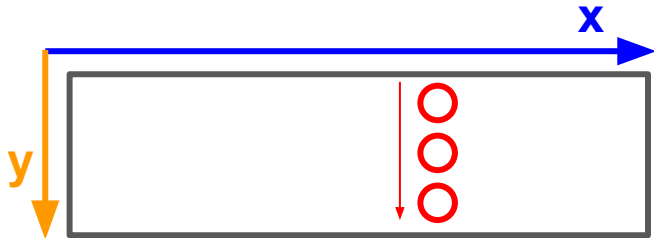
32 транзакции
(**uncoalesced**)



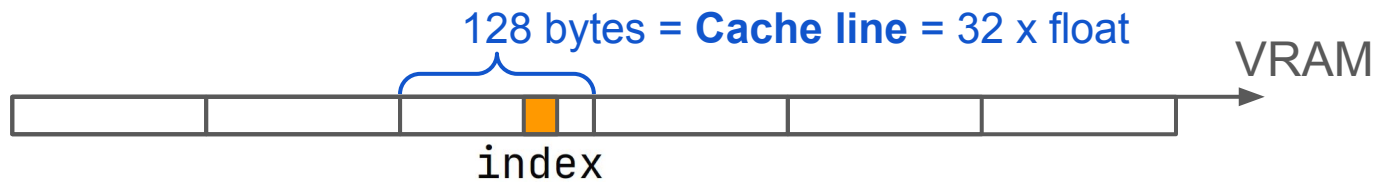
Архитектура VRAM (**coalesced** memory access pattern)



```
sum = 0;
for (int x = 0; x < width; ++x) {
    for (int y = 0; y < height; ++y) {
        sum += image[y * width + x];
    }
}
```

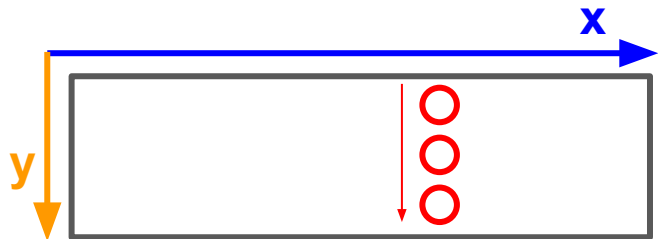


Архитектура VRAM (**coalesced** memory access pattern)

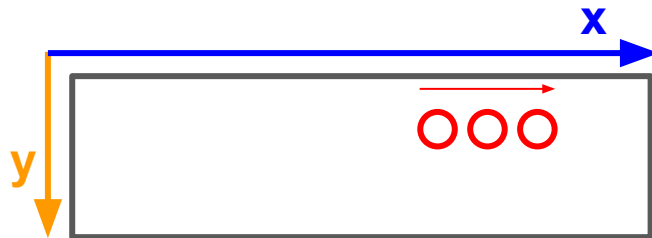


Но это однопоточный код, как в него добавить массовый параллелизм?

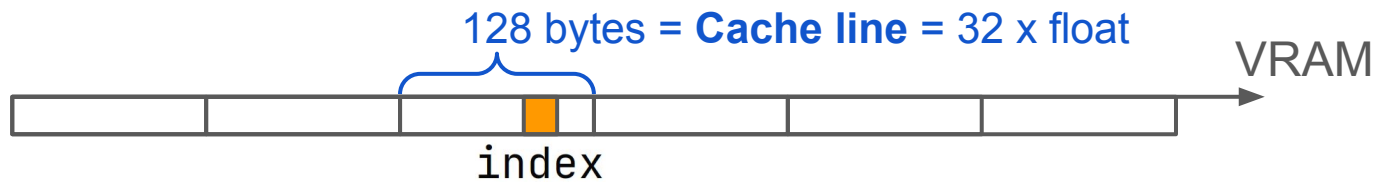
```
sum = 0;
for (int x = 0; x < width; ++x) {
    for (int y = 0; y < height; ++y) {
        sum += image[y * width + x];
    }
}
```



```
sum = 0;
for (int y = 0; y < height; ++y) {
    for (int x = 0; x < width; ++x) {
        sum += image[y * width + x];
    }
}
```



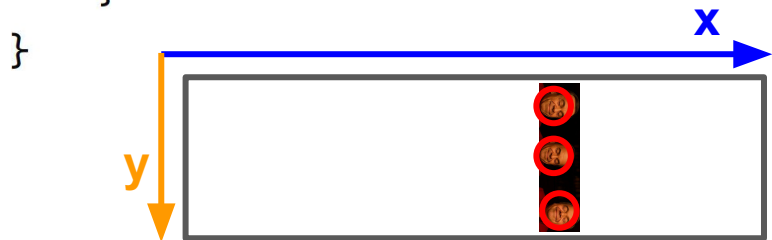
Архитектура VRAM (**coalesced** memory access pattern)



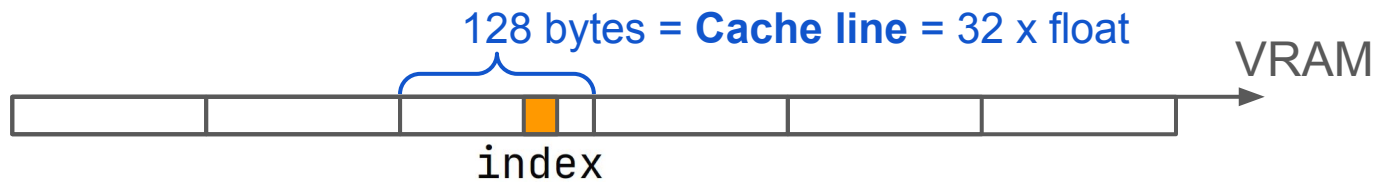
Но это однопоточный код, как в него добавить массовый параллелизм?

```
sum = 0;  
for (int x = 0; x < width; ++x) {  
    y =  
    sum += image[y * width + x];  
}
```

← Номер потока в *warp*



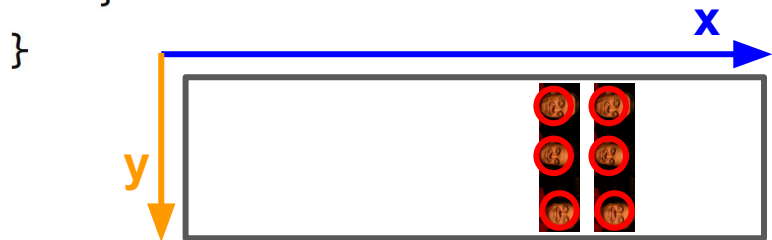
Архитектура VRAM (**coalesced** memory access pattern)



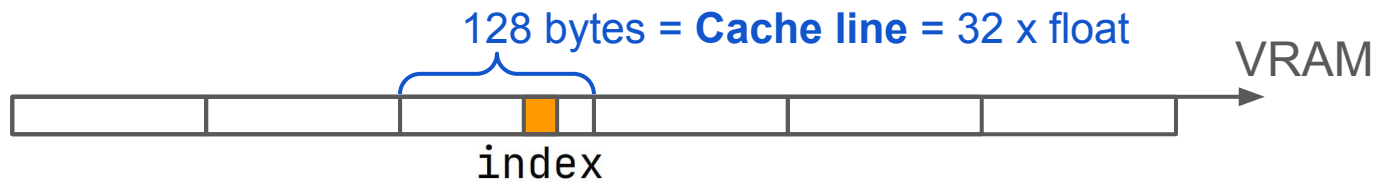
Но это однопоточный код, как в него добавить массовый параллелизм?

```
sum = 0;  
for (int x = 0; x < width; ++x) {  
    y =  
    sum += image[y * width + x];  
}
```

Номер потока в *warp*



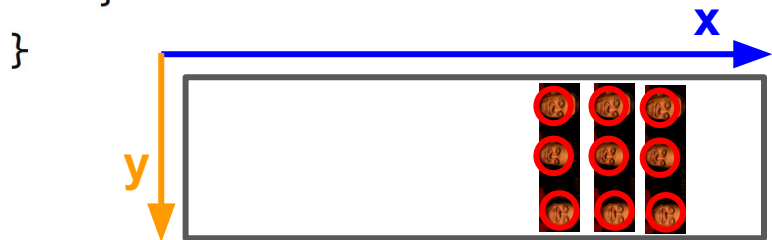
Архитектура VRAM (**coalesced** memory access pattern)



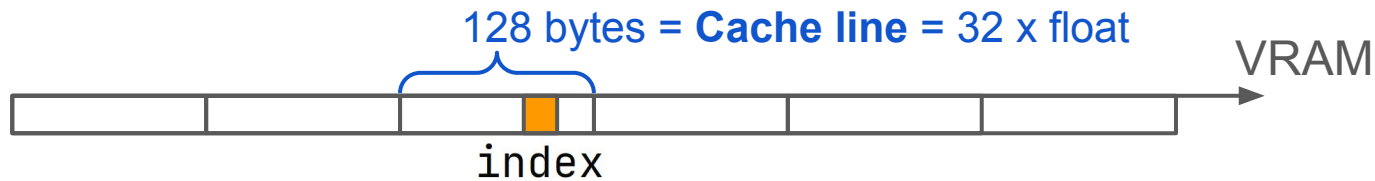
Но это однопоточный код, как в него добавить массовый параллелизм?

```
sum = 0;  
for (int x = 0; x < width; ++x) {  
    y =  
    sum += image[y * width + x];  
}
```

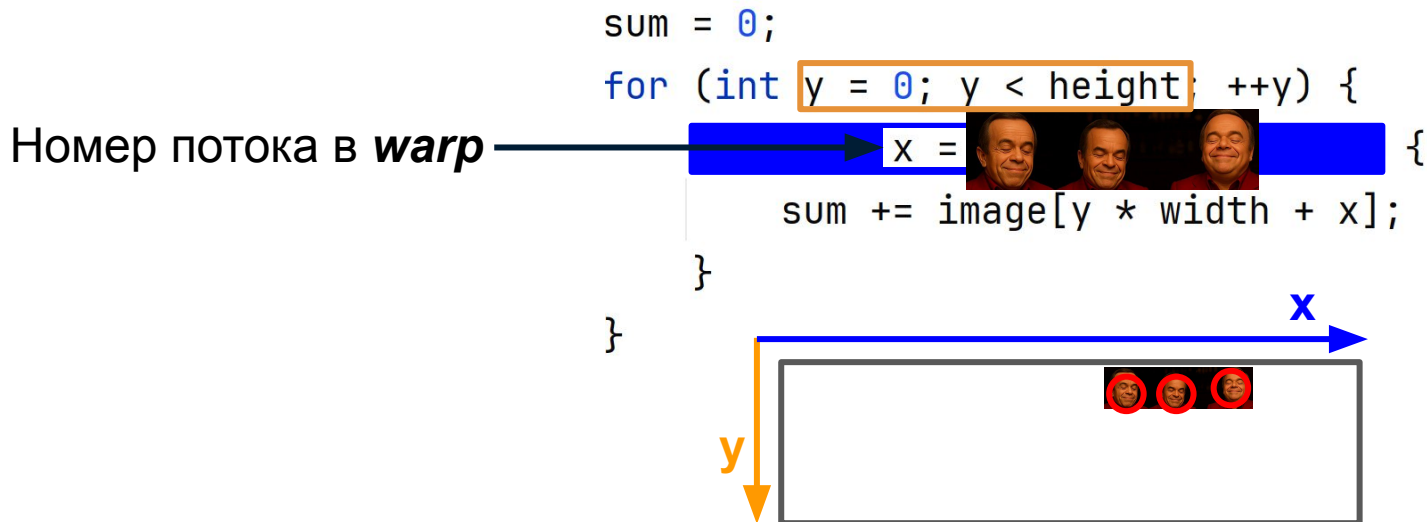
Номер потока в *warp*



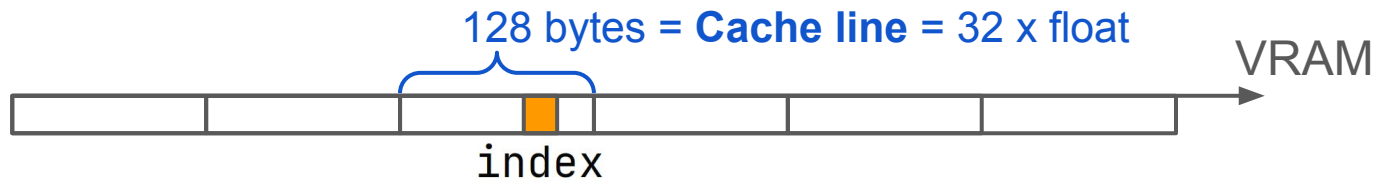
Архитектура VRAM (**coalesced** memory access pattern)



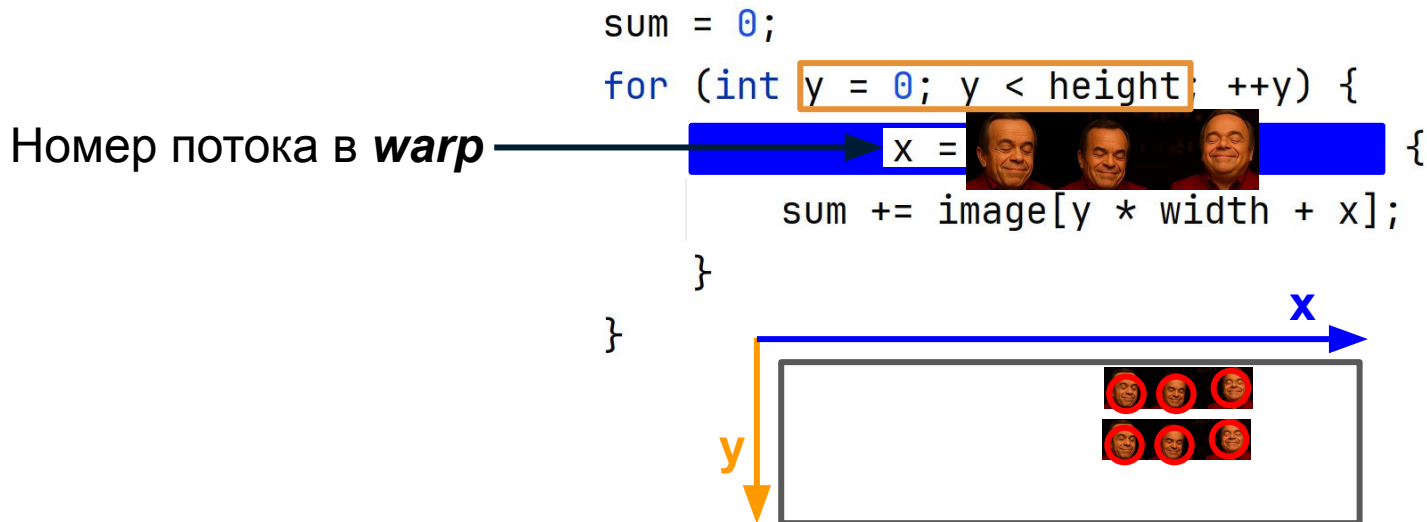
Но это однопоточный код, как в него добавить массовый параллелизм?



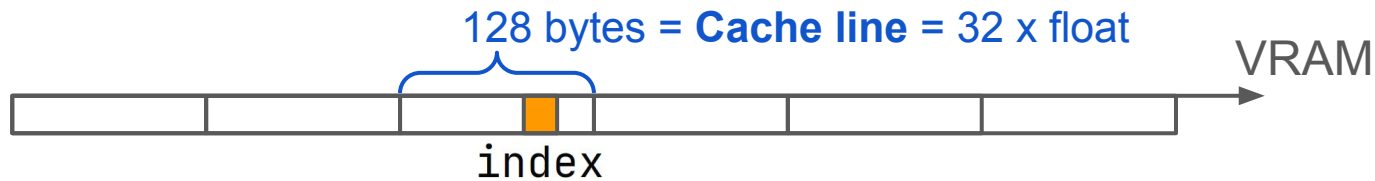
Архитектура VRAM (**coalesced** memory access pattern)



Но это однопоточный код, как в него добавить массовый параллелизм?

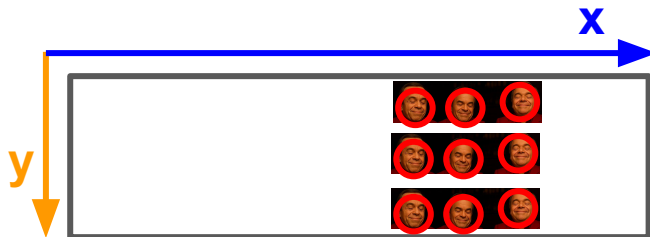


Архитектура VRAM (**coalesced** memory access pattern)

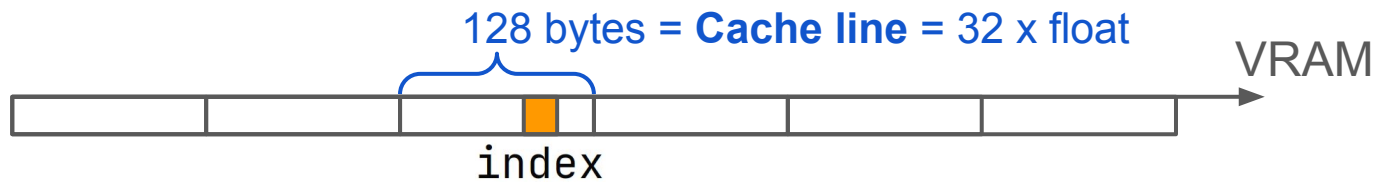


Но это однопоточный код, как в него добавить массовый параллелизм?

```
sum = 0;  
for (int y = 0; y < height; ++y) {  
    Номер потока в warp → x =  {  
        sum += image[y * width + x];  
    }  
}
```

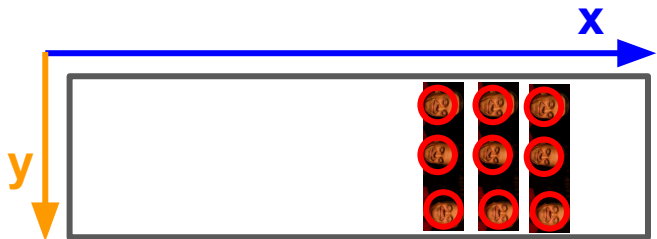


Архитектура VRAM (**coalesced** memory access pattern)

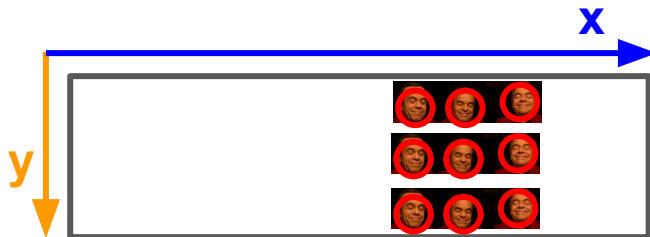


Какой код выбрали бы вы?

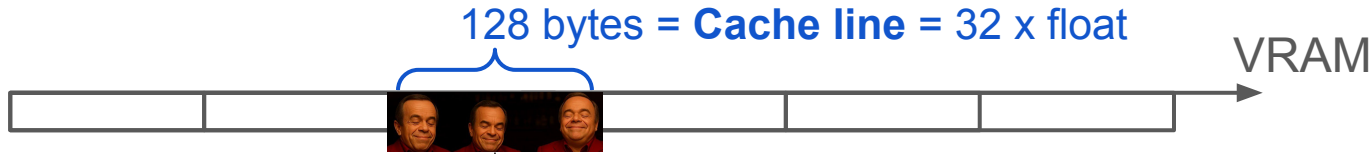
```
sum = 0;
for (int x = 0; x < width; ++x) {
    y =
    sum += image[y * width + x];
}
```



```
sum = 0;
for (int y = 0; y < height; ++y) {
    x =
    sum += image[y * width + x];
}
```



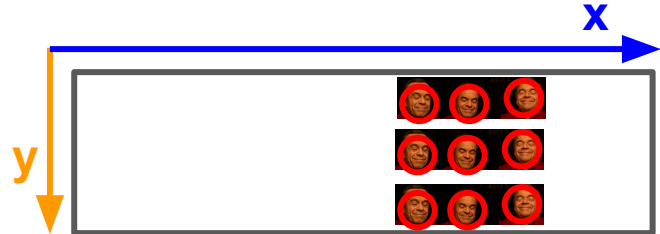
Архитектура VRAM (**coalesced** memory access pattern)



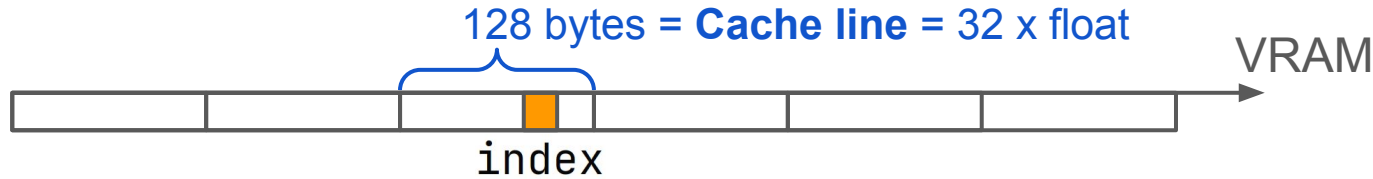
Какой код выбрали бы вы?

```
sum = 0;  
for (int y = 0; y < height; ++y) {  
    x =  {  
        sum += image[y * width + x];  
    }  
}
```

Они загрузят cache line
и распределят в рамках *warp*-а
байты, произойдет **broadcast**

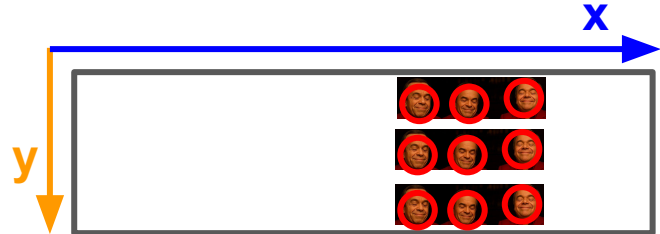


Архитектура VRAM (**coalesced** memory access pattern)

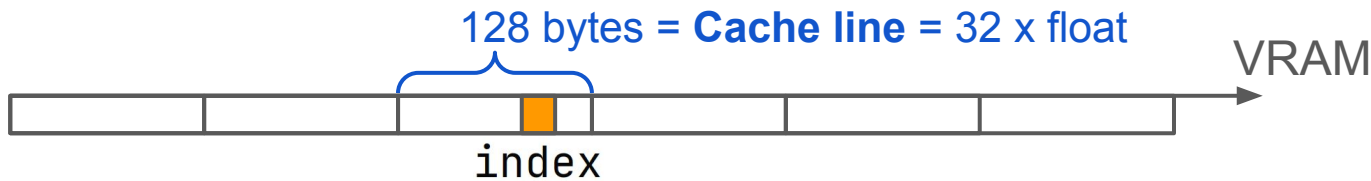


SIMT (thread)

```
sum = 0;  
for (int y = 0; y < height; ++y) {  
    x =  {  
        sum += image[y * width + x];  
    }  
}
```

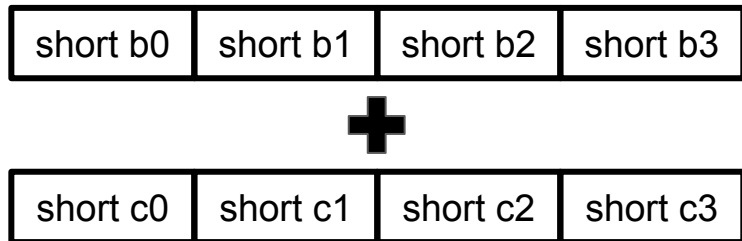


Архитектура VRAM (coalesced memory access pattern)



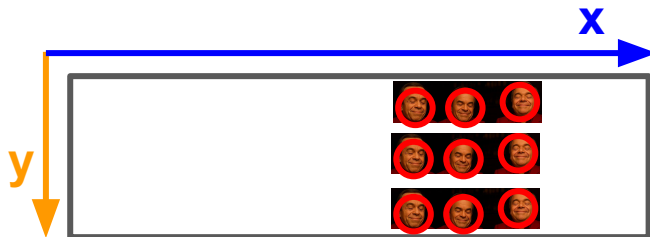
SIMD (data)

```
__m256i maskAll = _mm256_setzero_si256();
__m256i iters = _mm256_setzero_si256();
__m256 xs = _mm256_set1_ps(0.0f);
__m256 ys = _mm256_set1_ps(0.0f);
for (iteration = 0; iteration < MAX_ITERAT
    __m256 xsn = _mm256_add_ps(_mm256_sub
    __m256 ysn = _mm256_add_ps(_mm256_mul
    xs = _mm256_add_ps(_mm256_andnot_ps((
    ys = _mm256_add_ps(_mm256_andnot_ps((
maskAll = (__m256i) _mm256_or_ps(_mm256_cmp_ps(_mm256_add_ps(_mm256_mul_ps(xs, xs), _mm256_mul_ps(ys,
iters = _mm256_add_epi16(iters, _mm256_andnot_si256(maskAll, _mm256_set1_epi16(1)));
int mask = _mm256_movemask_epi8(maskAll);
if (mask == (int) 0xffffffff) {
    break;
}
```



SIMT (thread)

```
sum = 0;
for (int y = 0; y < height; ++y) {
    X =  {
        sum += image[y * width + x];
    }
}
```

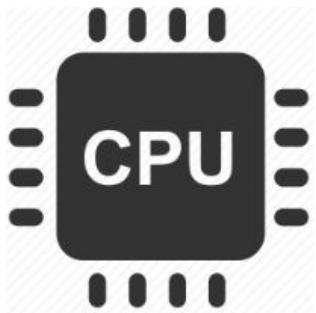


Глава 3: Общая картина ЭВМ-архитектуры

CPU - RAM - PCI-E - VRAM - GPU

Архитектура

$100 \cdot 10^9$ FLOPS



40 GB/s
low latency

RAM - DDR5

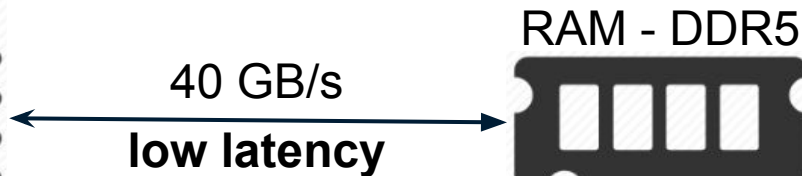
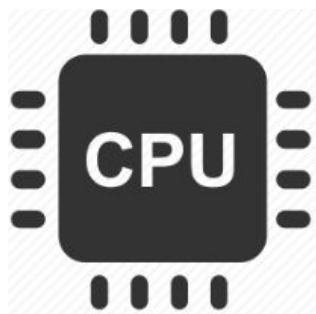


$100 \cdot 10^{12}$ FLOPS



Архитектура

$100 \cdot 10^9$ FLOPS



RAM - DDR5



$100 \cdot 10^{12}$ FLOPS

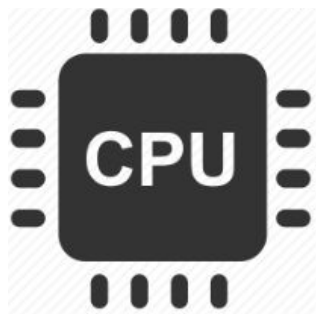


VRAM - GDDR6 или HBM (high bandwidth)

Архитектура

Где лучше исполнять OS?
(например реагировать на клики пользователя)
Какие есть метрики качества?

$100 \cdot 10^9$ FLOPS



40 GB/s
low latency

RAM - DDR5



$100 \cdot 10^{12}$ FLOPS



1000 GB/s
(x25 раз больше)

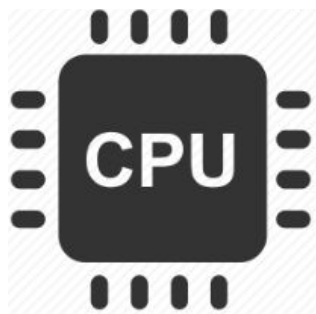


VRAM - GDDR6 или HBM (high bandwidth)

Архитектура

Где быстрее сложить два массива чисел?

$100 \cdot 10^9$ FLOPS



40 GB/s

RAM - DDR5



$100 \cdot 10^{12}$ FLOPS



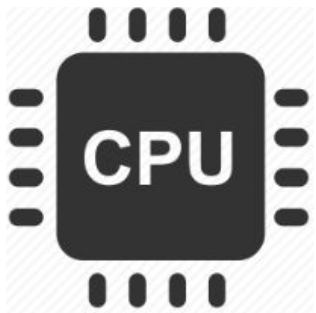
1000 GB/s



VRAM - GDDR6 или HBM (**high bandwidth**)

Архитектура

$100 \cdot 10^9$ FLOPS



Где быстрее сложить два массива чисел?
И во что мы упираемся - в память или в ALUs?
(ALUs = Arithmetic Logic Units)

40 GB/s

RAM - DDR5



$100 \cdot 10^{12}$ FLOPS



1000 GB/s

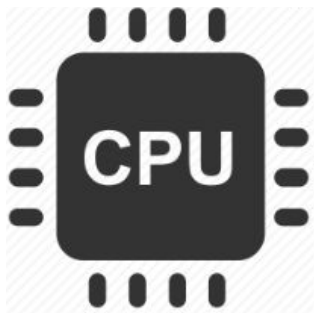


VRAM - GDDR6 или HBM (**high bandwidth**)

Архитектура

Где быстрее сложить два массива чисел?
И во что мы упираемся - в память или в ALUs?
А что если данные находятся в **RAM**?

$100 \cdot 10^9$ FLOPS



40 GB/s

RAM - DDR5



$100 \cdot 10^{12}$ FLOPS



1000 GB/s

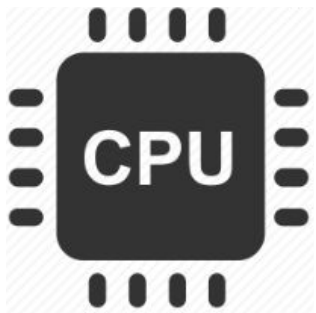


VRAM - GDDR6

PCI-E 5.0 x16
16 GB/s

Архитектура

$100 \cdot 10^9$ FLOPS



Где быстрее сложить два массива чисел?
И во что мы упираемся - в память или в ALUs?
А что если данные находятся в **RAM**?

40 GB/s

RAM - DDR5



$100 \cdot 10^{12}$ FLOPS



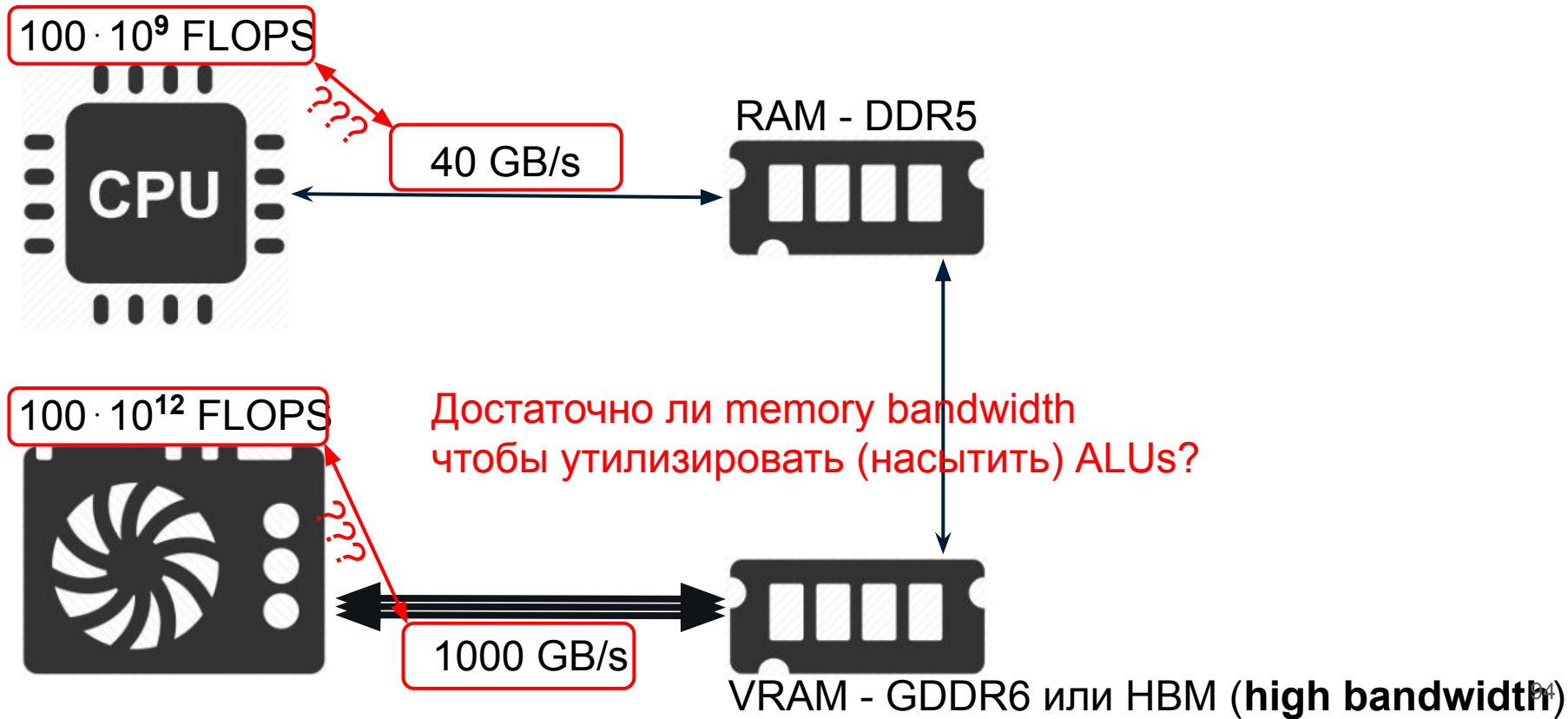
PCI-E 5.0 x16
16 GB/s

1000 GB/s

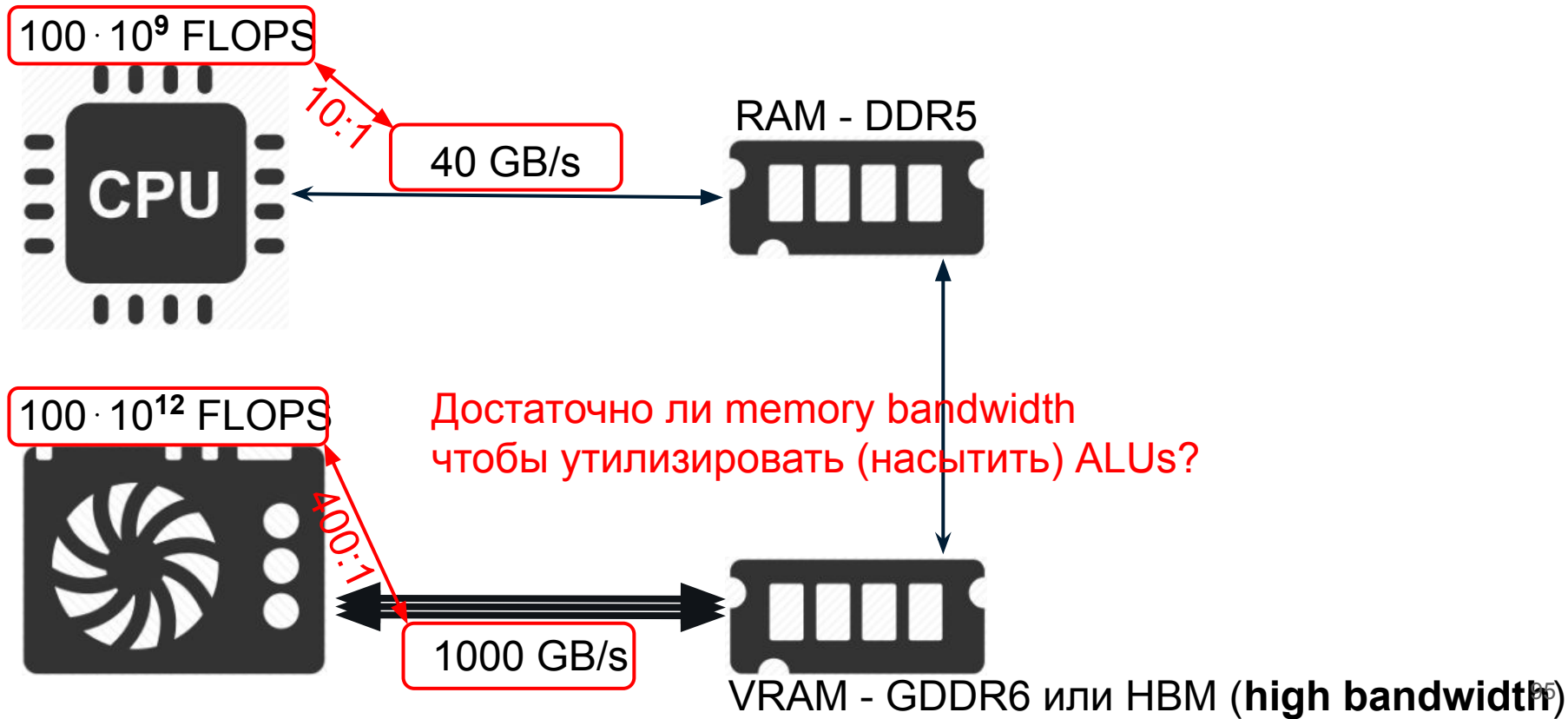


VRAM - GDDR6

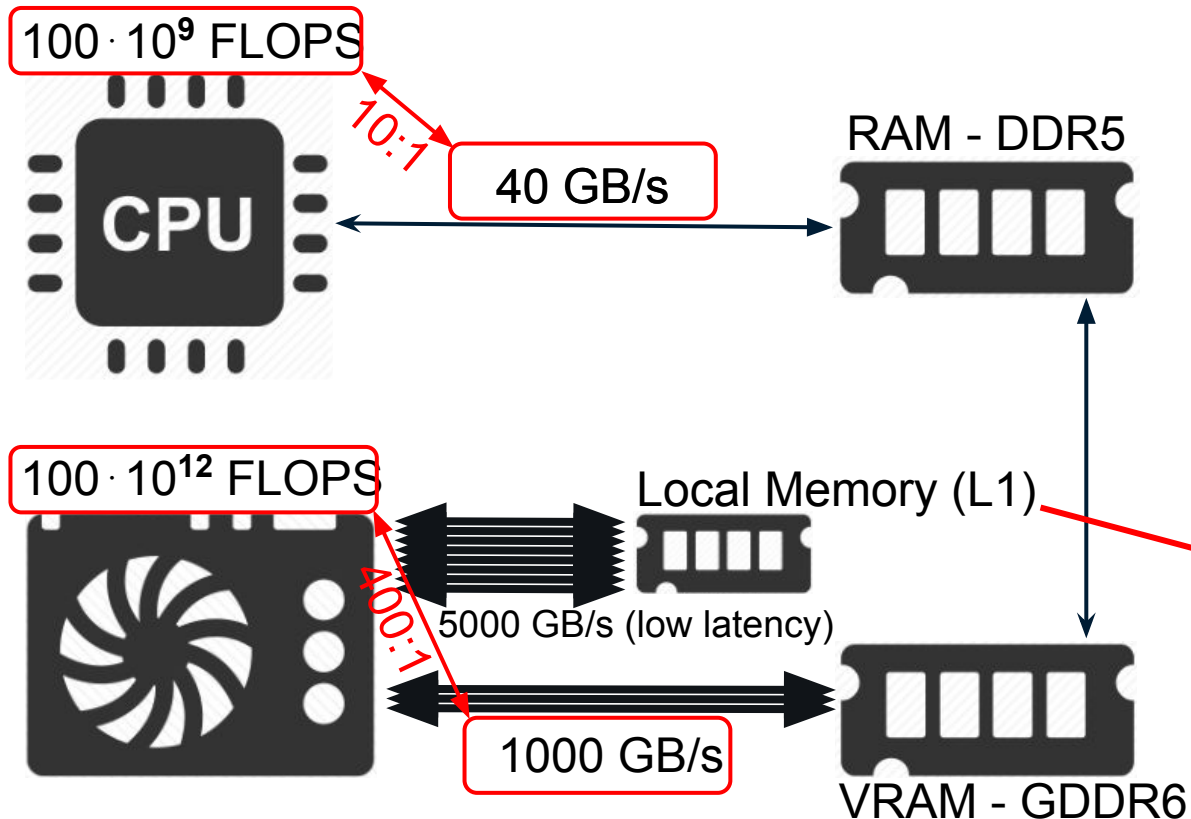
Архитектура



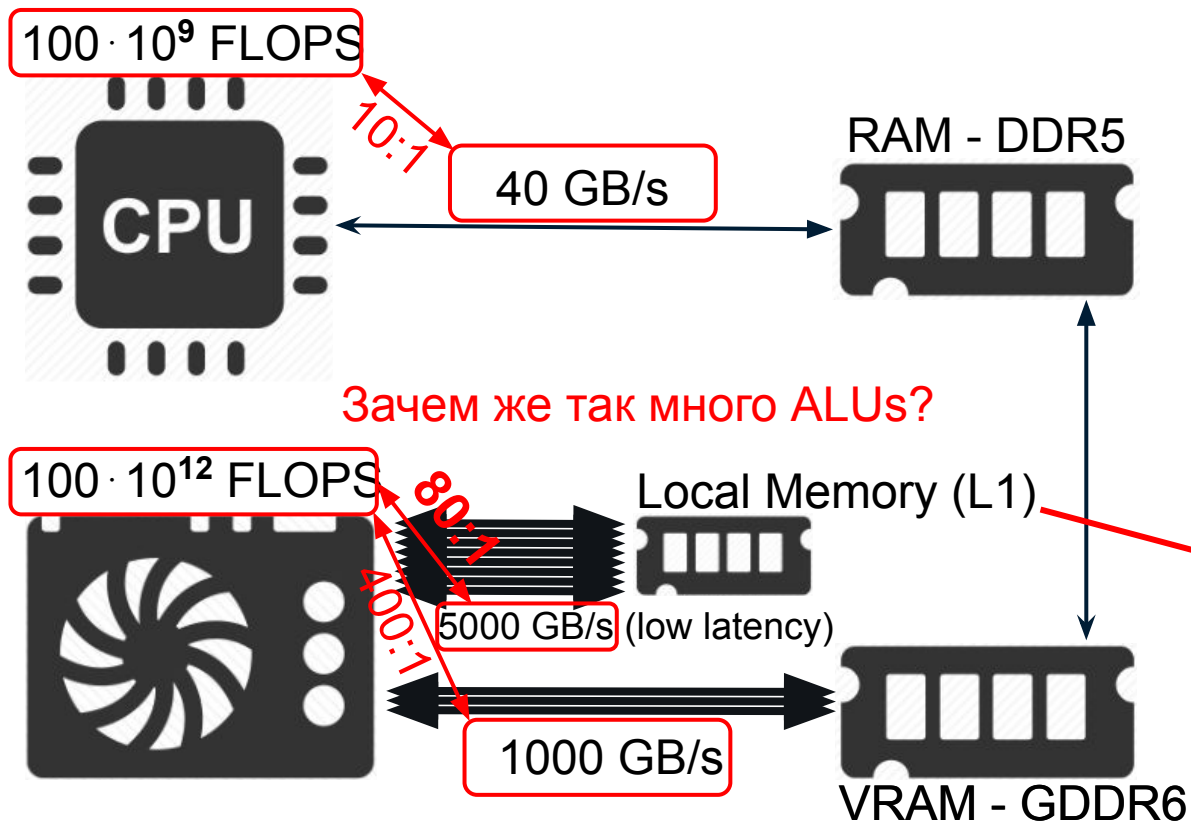
Архитектура



Архитектура

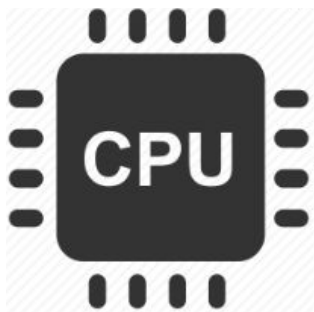


Архитектура



Архитектура

$100 \cdot 10^9$ FLOPS



40 GB/s



RAM - DDR5

Можно ли ее использовать
так же как VRAM?

$100 \cdot 10^{12}$ FLOPS



5000 GB/s (low latency)

Local Memory (L1)



1000 GB/s



VRAM - GDDR6



The diagram illustrates the internal architecture of the NVIDIA RTX 3090. It is divided into two main processing blocks, each containing a GPC (Graphics Processing Cluster) and a Raster Engine. Each GPC contains 6 RT Cores and 6 SMs (Streaming Multiprocessors). The Raster Engine contains 6 TPCs (Texture Processing Clusters). The diagram is color-coded: blue for L2 Cache, green for RT Cores, yellow for SMs, and red for TPCs. A red box highlights one of the RT Cores in the right GPC. The diagram is connected to a High-Speed Hub at the bottom, which is connected to a VLink - Two x8 Links. On the right side, there are three vertical bars labeled Memory Controller, Memory Controller, and Memory Controller.



5000 GB/s (low latency)



SM - Streaming Multiprocessor

The diagram illustrates the architecture of a Streaming Multiprocessor (SM). It is composed of four identical SM units arranged in a 2x2 grid. Each unit contains the following components:

- Warp Scheduler + Dispatch (32 threads/cb)**: Located at the top of each unit.
- Register File (16,384 x 32-bit)**: A large blue block below the scheduler.
- FP32 + INT32**: Two green blocks for floating-point and integer processing.
- TENSOR CORE**: A green block for tensor operations.
- LDST**: Four small red blocks for load/store operations.
- SFU**: A small red block for scalar function units.

Below the SM units is a **128KB L1 Data Cache / Shared Memory** block. At the bottom of the diagram are four **Tex** blocks, likely representing texture units. The entire architecture is enclosed in a red border.

$P_{\text{max}} = 1.6 \times 10^8$ Pa, $\sigma_{\text{max}} = 1.6 \times 10^{-7}$ s
 $P_{\text{max}} = 1.6 \times 10^8$ Pa, $\sigma_{\text{max}} = 1.6 \times 10^{-7}$ s

Age Group	Percentage
18-24	10%
25-34	15%
35-44	20%
45-54	25%
55-64	30%
65-74	35%
75-84	40%
85+	45%

FR32	FR32	TENSOR	FR32	FR32	TENSOR
------	------	--------	------	------	--------

INT32

Age Group	Percentage
18-24	10%
25-34	15%
35-44	20%
45-54	25%
55-64	30%
65-74	35%
75-84	40%
85+	45%

[illegible]

LDST LDST LDST LDST SFU LDST LDST LDST LDST SFU

Year	Number of cases	Number of deaths
1990	1,000	100
1991	1,000	100
1992	1,000	100
1993	1,000	100
1994	1,000	100
1995	1,000	100
1996	1,000	100
1997	1,000	100
1998	1,000	100
1999	1,000	100
2000	1,000	100
2001	1,000	100
2002	1,000	100
2003	1,000	100
2004	1,000	100
2005	1,000	100
2006	1,000	100
2007	1,000	100
2008	1,000	100
2009	1,000	100
2010	1,000	100
2011	1,000	100
2012	1,000	100
2013	1,000	100
2014	1,000	100
2015	1,000	100
2016	1,000	100
2017	1,000	100
2018	1,000	100
2019	1,000	100
2020	1,000	100

[illegible]

Age Group	Percentage
18-24	10%
25-34	15%
35-44	20%
45-54	25%
55-64	30%
65-74	35%
75-84	40%
85+	45%

Age Group	Percentage
18-24	10%
25-34	15%
35-44	20%
45-54	25%
55-64	30%
65-74	35%
75-84	40%
85+	45%

1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24	25	26	27	28	29	30	31	32	33	34	35	36	37	38	39	40	41	42	43	44	45	46	47	48	49	50	51	52	53	54	55	56	57	58	59	60	61	62	63	64	65	66	67	68	69	70	71	72	73	74	75	76	77	78	79	80	81	82	83	84	85	86	87	88	89	90	91	92	93	94	95	96	97	98	99	100
---	---	---	---	---	---	---	---	---	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	-----

[illegible]

FP32	FP32	TENSOR CORE	FP32	FP32	TENSOR CORE
------	------	-------------	------	------	-------------

+		CORE	+		CORE
INT32			INT32		

IN 102				IN 102			
--------	--	--	--	--------	--	--	--

Age Group	Percentage
18-24	10%
25-34	15%
35-44	20%
45-54	25%
55-64	30%
65-74	35%
75-84	40%
85+	45%

Age Group	Percentage
18-24	~12%
25-34	~28%
35-44	~18%
45-54	~15%
55-64	~10%
65-74	~8%
75-84	~5%
85+	~2%

1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24	25	26	27	28	29	30	31	32	33	34	35	36	37	38	39	40	41	42	43	44	45	46	47	48	49	50	51	52	53	54	55	56	57	58	59	60	61	62	63	64	65	66	67	68	69	70	71	72	73	74	75	76	77	78	79	80	81	82	83	84	85	86	87	88	89	90	91	92	93	94	95	96	97	98	99	100
---	---	---	---	---	---	---	---	---	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	----	-----

LD/ST	LD/ST	LD/ST	LD/ST	SFU	LD/ST	LD/ST	LD/ST	LD/ST	SFU
-------	-------	-------	-------	-----	-------	-------	-------	-------	-----

© 2006 The Authors
Journal compilation © 2006 Blackwell Publishing Ltd

128KB L1 Data Cache / Shared Memory

--	--

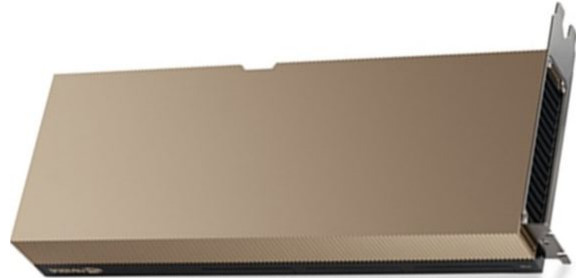
© 2004 Blackwell Publishing Ltd, *Journal of Internal Medicine* 255: 111–118

RT CORE

[illegible]

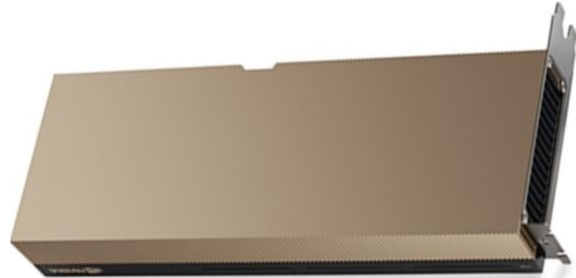
Архитектура GPU

- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)



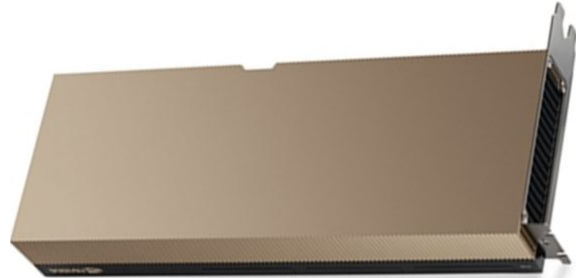
Архитектура GPU

- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра

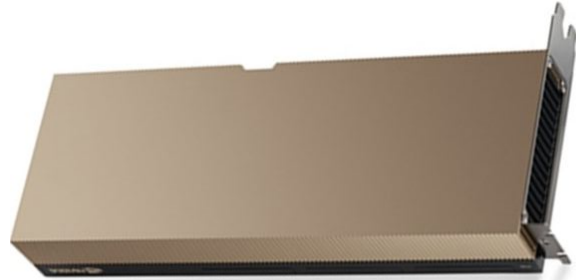


Архитектура GPU

- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра
 - но у warp-а общий **instruction pointer** (опасность **code divergence**)

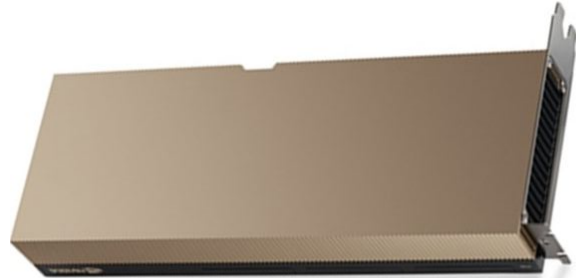


Архитектура GPU



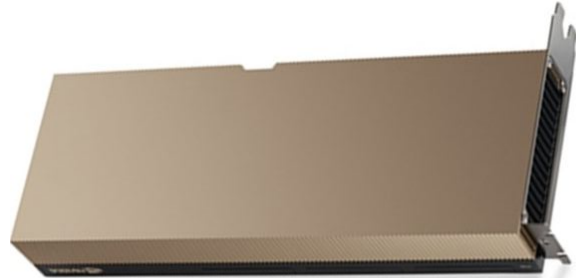
- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра
 - но у warp-а общий **instruction pointer** (опасность **code divergence**)
 - активен **warp** который не ждет данных из **VRAM** (при высокой **occupancy**)

Архитектура GPU



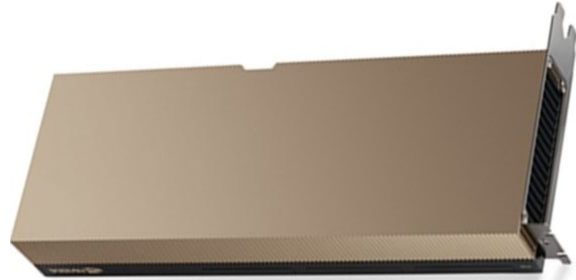
- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра
 - но у warp-а общий **instruction pointer** (опасность **code divergence**)
 - активен **warp** который не ждет данных из **VRAM** (при высокой **occupancy**)
- 2) Огромная пропускная способность **VRAM**:

Архитектура GPU



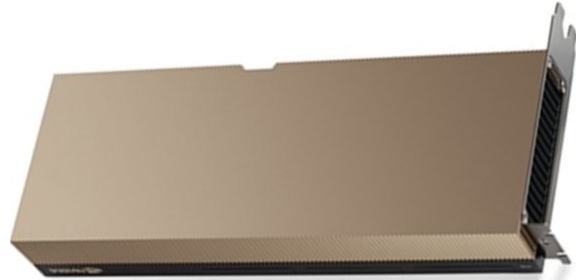
- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра
 - но у warp-а общий **instruction pointer** (опасность **code divergence**)
 - активен **warp** который не ждет данных из **VRAM** (при высокой **occupancy**)
- 2) Огромная пропускная способность **VRAM**:
 - но большая задержка (**latency**)

Архитектура GPU



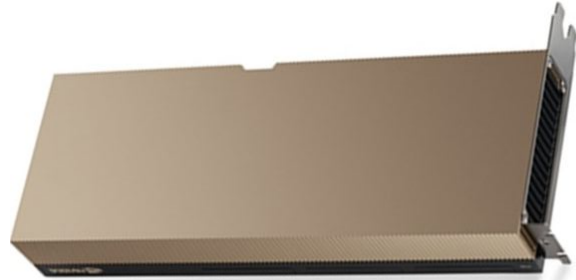
- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра
 - но у warp-а общий **instruction pointer** (опасность **code divergence**)
 - активен **warp** который не ждет данных из **VRAM** (при высокой **occupancy**)
- 2) Огромная пропускная способность **VRAM**:
 - но большая задержка (**latency**)
 - скрывается за счет **сверх-SMT**/Hyper-Threading (при высокой **occupancy**)

Архитектура GPU



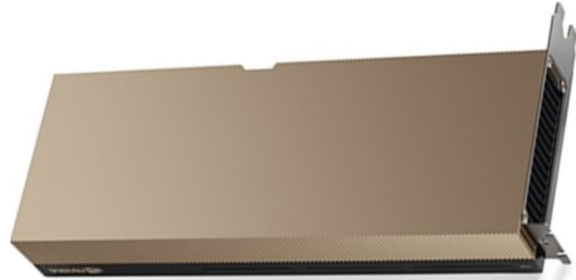
- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра
 - но у warp-а общий **instruction pointer** (опасность **code divergence**)
 - активен **warp** который не ждет данных из **VRAM** (при высокой **occupancy**)
- 2) Огромная пропускная способность **VRAM**:
 - но большая задержка (**latency**)
 - скрывается за счет **сверх-SMT**/Hyper-Threading (при высокой **occupancy**)
 - транзакции разбиты на 128-байтные **cache line**

Архитектура GPU



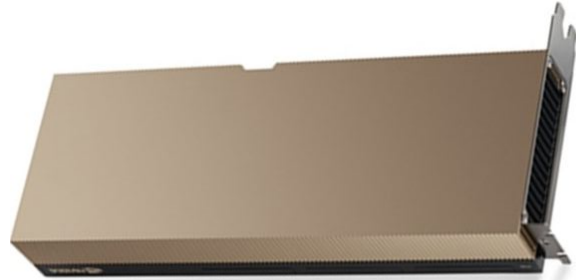
- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра
 - но у warp-а общий **instruction pointer** (опасность **code divergence**)
 - активен **warp** который не ждет данных из **VRAM** (при высокой **occupancy**)
- 2) Огромная пропускная способность **VRAM**:
 - но большая задержка (**latency**)
 - скрывается за счет **сверх-SMT**/Hyper-Threading (при высокой **occupancy**)
 - транзакции разбиты на 128-байтные **cache line**
 - но при **un-coalesced** memory access pattern - малая проп. способность

Архитектура GPU



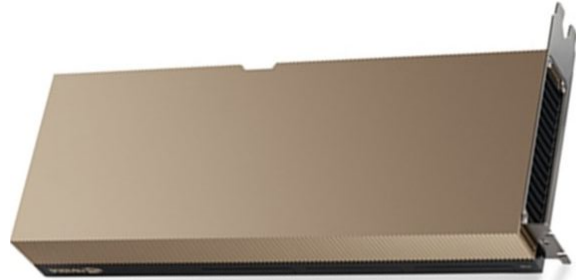
- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра
 - но у warp-а общий **instruction pointer** (опасность **code divergence**)
 - активен **warp** который не ждет данных из **VRAM** (при высокой **occupancy**)
- 2) Огромная пропускная способность **VRAM**:
 - но большая задержка (**latency**)
 - скрывается за счет **сверх-SMT**/Hyper-Threading (при высокой **occupancy**)
 - транзакции разбиты на 128-байтные **cache line**
 - но при **un-coalesced** memory access pattern - малая проп. способность
- 3) Локальная память (**local/shared memory**):

Архитектура GPU



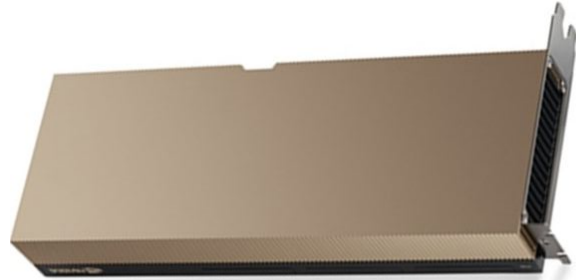
- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра
 - но у warp-а общий **instruction pointer** (опасность **code divergence**)
 - активен **warp** который не ждет данных из **VRAM** (при высокой **occupancy**)
- 2) Огромная пропускная способность **VRAM**:
 - но большая задержка (**latency**)
 - скрывается за счет **сверх-SMT**/Hyper-Threading (при высокой **occupancy**)
 - транзакции разбиты на 128-байтные **cache line**
 - но при **un-coalesced** memory access pattern - малая проп. способность
- 3) Локальная память (**local/shared memory**):
 - **титаническая** пропускная способность + низкая задержка

Архитектура GPU



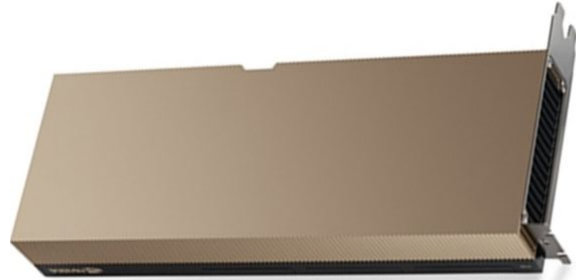
- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра
 - но у warp-а общий **instruction pointer** (опасность **code divergence**)
 - активен **warp** который не ждет данных из **VRAM** (при высокой **occupancy**)
- 2) Огромная пропускная способность **VRAM**:
 - но большая задержка (**latency**)
 - скрывается за счет **сверх-SMT**/Hyper-Threading (при высокой **occupancy**)
 - транзакции разбиты на 128-байтные **cache line**
 - но при **un-coalesced** memory access pattern - малая проп. способность
- 3) Локальная память (**local/shared memory**):
 - **титаническая** пропускная способность + низкая задержка
 - нет проблемы с coalesced-паттерном

Архитектура GPU



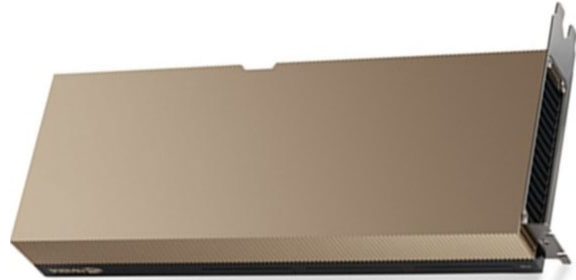
- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра
 - но у warp-а общий **instruction pointer** (опасность **code divergence**)
 - активен **warp** который не ждет данных из **VRAM** (при высокой **occupancy**)
- 2) Огромная пропускная способность **VRAM**:
 - но большая задержка (**latency**)
 - скрывается за счет **сверх-SMT**/Hyper-Threading (при высокой **occupancy**)
 - транзакции разбиты на 128-байтные **cache line**
 - но при **un-coalesced** memory access pattern - малая проп. способность
- 3) Локальная память (**local/shared memory**):
 - **титаническая** пропускная способность + низкая задержка
 - нет проблемы с coalesced-паттерном
 - **но?**

Архитектура GPU



- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра
 - но у warp-а общий **instruction pointer** (опасность **code divergence**)
 - активен **warp** который не ждет данных из **VRAM** (при высокой **occupancy**)
- 2) Огромная пропускная способность **VRAM**:
 - но большая задержка (**latency**)
 - скрывается за счет **сверх-SMT**/Hyper-Threading (при высокой **occupancy**)
 - транзакции разбиты на 128-байтные **cache line**
 - но при **un-coalesced** memory access pattern - малая проп. способность
- 3) Локальная память (**local/shared memory**):
 - **титаническая** пропускная способность + низкая задержка
 - нет проблемы с coalesced-паттерном
 - но **небольшая** + **локальная** для каждого **warp WorkGroup=Block**

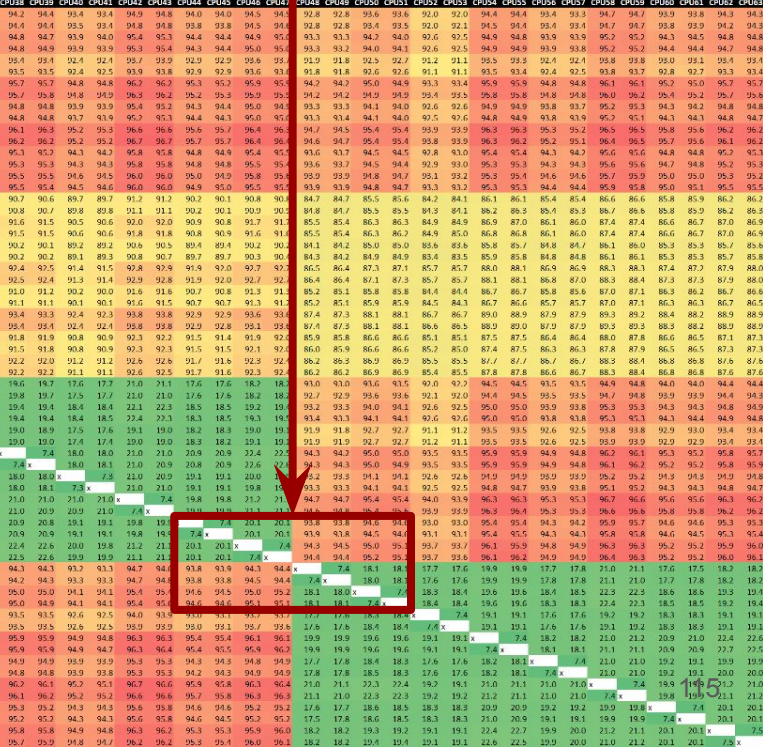
Архитектура GPU



- 1) Десятки тысяч ядер (много **FLOPs**):
 - сгруппированы по 32 в **warp**-ы (по 64 в **wavefront**-ы у AMD)
 - но слабые ядра
 - но у warp-а общий **instruction pointer** (опасность **code divergence**)
 - активен **warp** который не ждет данных из **VRAM** (при высокой **occupancy**)
- 2) Огромная пропускная способность **VRAM**:
 - но большая задержка (**latency**)
 - скрывается за счет **сверх-SMT**/Hyper-Threading (при высокой **occupancy**)
 - транзакции разбиты на 128-байтные **cache line**
 - но при **un-coalesced** memory access pattern - малая проп. способность
- 3) Локальная память (**local/shared memory**):
 - **титаническая** пропускная способность + низкая задержка
 - нет проблемы с coalesced-паттерном Позже мы обсудим **bank conflicts**
 - но **небольшая** + **локальная** для каждого **warp** **WorkGroup=Block**
И в этом секрет успеха! Рецепт к **масштабируемости!**



More Latency



64 ядра = 8 x Чипетов по 8 ядер

AMD Ryzen Threadripper 7980X Core-to-Core Latency

CPU0	17.2	18.0	18.1	17.5	17.5	19.6	19.1	27.6	17.5	20.9	20.9	17.6	17.6	18.1	17.1	87.94	87.1	90.2	90.2	88.8	88.7	91.1	91.1	89.8	89.8	92.0	91.9	90.6	90.5	90.7	90.7	92.9	91.1	93.5	93.5	92.0	92.0	94.4	94.4	93.5	93.4	94.8	94.8	93.8	93.8	94.4	94.4	92.8	92.8	93.5	93.5	92.0	92.1	94.5	94.4	93.9	94.7	94.7	93.8	93.9	94.2	94.3			
CPU1	18.0	18.1	17.7	7.1	18.2	18.3	19.4	18.5	18.1	22.4	20.0	18.4	18.4	19.2	1.1	89.8	90.1	90.6	90.7	89.3	88.3	91.5	91.5	90.2	90.3	92.4	92.5	91.1	91.0	91.2	91.1	93.6	93.5	93.9	93.9	92.5	92.5	94.8	94.7	93.9	94.0	95.4	95.3	94.4	94.4	94.9	94.9	93.3	93.3	94.7	94.0	90.6	92.5	94.9	94.8	93.9	93.9	95.2	95.2	94.3	94.5	94.8	94.8		
CPU2	18.0	18.1	17.7	7.1	18.2	18.3	19.4	18.5	18.1	22.2	22.2	18.5	18.5	19.8	1.1	89.8	90.1	90.6	90.7	89.3	89.3	91.6	91.5	90.2	90.3	92.4	92.5	91.0	91.0	91.3	91.1	93.4	93.5	93.9	93.9	92.5	92.5	94.8	94.9	93.9	93.9	95.4	95.4	94.4	94.4	95.0	95.0	93.3	93.3	94.0	94.1	92.6	92.5	94.9	94.9	93.8	93.8	95.2	95.2	94.4	94.7	94.8	94.8		
CPU3	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	88.6	88.5	89.3	89.3	87.8	87.8	90.1	90.1	88.8	88.8	91.0	91.0	91.5	91.5	91.0	91.0	91.5	91.5	92.6	92.6	92.5	92.5	91.1	91.1	93.5	93.5	92.5	92.5	93.8	93.8	92.9	92.9	93.6	93.6	91.8	91.8	92.6	92.6	91.1	91.1	93.5	93.4	92.4	92.5	93.8	93.7	92.8	92.7	93.3	93.3
CPU4	15.6	15.6	15.3	9.1	15.8	15.8	1.1	1.1	1.1	1.1	1.1	1.1	1.1	1.1	1.1	90.4	90.8	9.6	91.6	90.1	90.1	92.4	92.5	91.1	91.0	93.3	93.4	91.5	91.8	92.0	92.0	94.3	94.3	94.9	94.4	93.5	93.5	96.7	96.7	94.8	94.8	96.7	96.7	96.3	96.3	96.9	96.9	94.7	94.7	95.0	94.9	93.3	93.4	95.9	95.8	94.8	94.8	96.1	96.1	95.2	95.0	95.7	95.7		
CPU5	17.1	17.1	17.1	17.1	17.1	17.1	17.1	17.1	17.1	17.1	17.1	17.1	17.1	17.1	17.1	89.9	90.5	91.8	91.8	90.2	92.5	92.5	91.1	91.1	93.5	93.4	91.9	91.8	92.0	91.1	91.1	91.5	91.5	94.9	94.9	95.8	95.8	96.7	96.7	96.8	96.8	96.2	96.2	96.9	96.9	95.8	95.8	96.8	96.8	96.0	96.2	95.8	95.8	96.0	96.2	95.8	95.8	96.0	96.2	95.8	95.8				
CPU6	17.4	17.4	17.4	17.4	17.4	17.4	17.4	17.4	17.4	17.4	17.4	17.4	17.4	17.4	17.4	89.8	90.5	90.7	90.7	89.2	89.5	91.6	91.6	90.5	90.5	92.5	92.5	91.1	91.1	91.5	91.5	94.0	93.9	92.5	92.5	94.8	94.8	93.7	93.9	96.7	96.7	96.3	96.3	94.4	94.3	95.0	95.0	93.4	93.4	94.1	94.0	92.5	92.6	94.8	94.9	93.8	93.8	95.2	95.1	94.3	94.3	94.8	94.7	94.8	94.7
CPU7	21.0	20.9	22.4	22.2	22.0	19.0	21.0	21.0	21.0	21.0	21.0	21.0	21.0	21.0	21.0	91.3	91.1	92.1	90.6	92.8	92.9	91.5	91.5	91.2	93.8	93.8	92.2	92.2	92.5	92.3	94.8	94.8	95.2	95.2	93.8	93.8	96.1	96.3	95.2	95.3	96.6	96.6	96.9	96.9	96.4	96.4	96.7	96.7	94.7	94.5	96.4	96.4	93.9	93.9	96.3	96.3	95.2	95.2	96.5	96.5	95.8	95.8	95.6	95.6	
CPU8	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	90.6	93.8	93.8	91.5	91.5	91.4	91.4	93.9	93.9	94.4	94.4	93.8	93.8	95.3	95.3	94.3	94.3	95.8	95.8	94.8	94.8	95.4	95.4	93.6	93.7	94.5	94.5	92.9	93.0	95.3	95.3	94.4	94.4	95.6	95.6	94.8	94.8	95.2	95.2		
CPU9	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	90.6	93.8	93.8	91.5	91.5	91.4	91.4	93.9	93.9	94.4	94.4	93.8	93.8	95.3	95.3	94.3	94.3	95.8	95.8	94.8	94.8	95.4	95.4	93.6	93.7	94.5	94.5	92.9	93.0	95.3	95.3	94.4	94.4	95.6	95.6	94.8	94.8	95.2	95.2		
CPU10	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	90.6	93.8	93.8	91.5	91.5	91.4	91.4	93.9	93.9	94.4	94.4	93.8	93.8	95.3	95.3	94.3	94.3	95.8	95.8	94.8	94.8	95.4	95.4	93.6	93.7	94.5	94.5	92.9	93.0	95.3	95.3	94.4	94.4	95.6	95.6	94.8	94.8	95.2	95.2		
CPU11	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	90.6	93.8	93.8	91.5	91.5	91.4	91.4	93.9	93.9	94.4	94.4	93.8	93.8	95.3	95.3	94.3	94.3	95.8	95.8	94.8	94.8	95.4	95.4	93.6	93.7	94.5	94.5	92.9	93.0	95.3	95.3	94.4	94.4	95.6	95.6	94.8	94.8	95.2	95.2		
CPU12	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	90.6	93.8	93.8	91.5	91.5	91.4	91.4	93.9	93.9	94.4	94.4	93.8	93.8	95.3	95.3	94.3	94.3	95.8	95.8	94.8	94.8	95.4	95.4	93.6	93.7	94.5	94.5	92.9	93.0	95.3	95.3	94.4	94.4	95.6	95.6	94.8	94.8	95.2	95.2		
CPU13	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	90.6	93.8	93.8	91.5	91.5	91.4	91.4	93.9	93.9	94.4	94.4	93.8	93.8	95.3	95.3	94.3	94.3	95.8	95.8	94.8	94.8	95.4	95.4	93.6	93.7	94.5	94.5	92.9	93.0	95.3	95.3	94.4	94.4	95.6	95.6	94.8	94.8	95.2	95.2		
CPU14	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	90.6	93.8	93.8	91.5	91.5	91.4	91.4	93.9	93.9	94.4	94.4	93.8	93.8	95.3	95.3	94.3	94.3	95.8	95.8	94.8	94.8	95.4	95.4	93.6	93.7	94.5	94.5	92.9	93.0	95.3	95.3	94.4	94.4	95.6	95.6	94.8	94.8	95.2	95.2		
CPU15	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	90.6	93.8	93.8	91.5	91.5	91.4	91.4	93.9	93.9	94.4	94.4	93.8	93.8	95.3	95.3	94.3	94.3	95.8	95.8	94.8	94.8	95.4	95.4	93.6	93.7	94.5	94.5	92.9	93.0	95.3	95.3	94.4	94.4	95.6	95.6	94.8	94.8	95.2	95.2		
CPU16	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	90.6	93.8	93.8	91.5	91.5	91.4	91.4	93.9	93.9	94.4	94.4	93.8	93.8	95.3	95.3	94.3	94.3	95.8	95.8	94.8	94.8	95.4	95.4	93.6	93.7	94.5	94.5	92.9	93.0	95.3	95.3	94.4	94.4	95.6	95.6	94.8	94.8	95.2	95.2		
CPU17	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	90.6	93.8	93.8	91.5	91.5	91.4	91.4	93.9	93.9	94.4	94.4	93.8	93.8	95.3	95.3	94.3	94.3	95.8	95.8	94.8	94.8	95.4	95.4	93.6	93.7	94.5	94.5	92.9	93.0	95.3	95.3	94.4	94.4	95.6	95.6	94.8	94.8	95.2	95.2		
CPU18	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	90.6	93.8	93.8	91.5	91.5	91.4	91.4	93.9	93.9	94.4	94.4	93.8	93.8	95.3	95.3	94.3	94.3	95.8	95.8	94.8	94.8	95.4	95.4	93.6	93.7	94.5	94.5	92.9	93.0	95.3	95.3	94.4	94.4	95.6	95.6	94.8	94.8	95.2	95.2		
CPU19	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	90.6	93.8	93.8	91.5	91.5	91.4	91.4	93.9	93.9	94.4	94.4	93.8	93.8	95.3	95.3	94.3	94.3	95.8	95.8	94.8	94.8	95.4	95.4	93.6	93.7	94.5	94.5	92.9	93.0	95.3	95.3	94.4	94.4	95.6	95.6	94.8	94.8	95.2	95.2		
CPU20	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	90.6	93.8	93.8	91.5	91.5	91.4	91.4	93.9	93.9	94.4	94.4	93.8	93.8	95.3	95.3	94.3	94.3	95.8	95.8	94.8	94.8	95.4	95.4	93.6	93.7	94.5	94.5	92.9	93.0	95.3	95.3	94.4	94.4	95.6	95.6	94.8	94.8	95.2	95.2		
CPU21	17.6	17.6	17.6	18.2	18.2	18.2	18.2	17.6	17.6	18.2	18.2	17.6	17.6	18.2	18.2	91.4	91.4	91.7	91.7	89.5	89.6	92.0	92.0	90.6	9																																								

64 ядра = 8 x Чиплетов (по 8 ядер в чиплете)

64 ядра = 8 x Чиплетов (по 8 ядер в чиплете)

AMD Ryzen Threadripper 7980X Core-to-Core Latency

CPU0	180	180	175	175	175	177	174	170	168	176	160	16	CPU8	902	901	888	888	911	911	897	896	918	919	907	906	907	906	913	913	929	934	935	920	922	942	944	934	934	939	948	940	940	945	945	928	928	936	936	920	920	944	944	944	944	933	947	947	933	933	943	943																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																												
CPU1	180	180	180	180	180	180	180	180	180	180	180	18	CPU9	902	901	888	888	911	911	897	896	918	919	907	906	907	906	913	913	929	934	935	920	922	942	944	934	934	939	948	940	940	945	945	928	928	936	936	920	920	944	944	944	944	933	947	947	933	933	943	943																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																												
CPU2	180	181	181	181	181	181	181	181	181	181	181	18	CPU10	902	901	888	888	911	911	897	896	918	919	907	906	907	906	913	913	929	934	935	920	922	942	944	934	934	939	948	940	940	945	945	928	928	936	936	920	920	944	944	944	944	933	947	947	933	933	943	943																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																												
CPU3	180	181	171	171	171	172	172	172	172	172	172	18	CPU11	902	901	888	888	911	911	897	896	918	919	907	906	907	906	913	913	929	934	935	920	922	942	944	934	934	939	948	940	940	945	945	928	928	936	936	920	920	944	944	944	944	933	947	947	933	933	943	943																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																												
CPU4	175	175	182	182	182	182	182	182	182	182	182	18	CPU12	902	901	888	888	911	911	897	896	918	919	907	906	907	906	913	913	929	934	935	920	922	942	944	934	934	939	948	940	940	945	945	928	928	936	936	920	920	944	944	944	944	933	947	947	933	933	943	943																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																												
CPU5	176	176	183	183	183	183	183	183	183	183	183	18	CPU13	902	901	888	888	911	911	897	896	918	919	907	906	907	906	913	913	929	934	935	920	922	942	944	934	934	939	948	940	940	945	945	928	928	936	936	920	920	944	944	944	944	933	947	947	933	933	943	943																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																												
CPU6	176	176	183	183	183	183	183	183	183	183	183	18	CPU14	902	901	888	888	911	911	897	896	918	919	907	906	907	906	913	913	929	934	935	920	922	942	944	934	934	939	948	940	940	945	945	928	928	936	936	920	920	944	944	944	944	933	947	947	933	933	943	943																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																												
CPU7	176	176	183	183	183	183	183	183	183	183	183	18	CPU15	902	901	888	888	911	911	897	896	918	919	907	906	907	906	913	913	929	934	935	920	922	942	944	934	934	939	948	940	940	945	945	928	928	936	936	920	920	944	944	944	944	933	947	947	933	933	943	943																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																												
CPU8	176	176	183	183	183	183	183	183	183	183	183	18	CPU16	902	901	888	888	911	911	897	896	918	919	907	906	907	906	913	913	929	934	935	920	922	942	944	934	934	939	948	940	940	945	945	928	928	936	936	920	920	944	944	944	944	933	947	947	933	933	943	943																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																												
CPU9	176	176	183	183	183	183	183	183	183	183	183	18	CPU17	902	901	888	888	911	911	897	896	918	919	907	906	907	906	913	913	929	934	935	920	922	942	944	934	934	939	948	940	940	945	945	928	928	936	936	920	920	944	944	944	944	933	947	947	933	933	943	943																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																																												
CPU10	210	209	224	222	220	210	210	208	218	217	214	198	199	209	208	913	913	921	921	906	908	928	929	915	912	918	938	922	922	925	923	948	948	952	952	939	938	961	963	952	953	965	966	957	958	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	956	

64 ядра = 8 x Чиплетов (по 8 ядер

CPU19	CPU20	CPU21	CPU22	CPU23	CPU24	CPU25	CPU26	CPU27	CPU28	CPU29	CPU30	CPU31	CPU32	CPU33	CPU34	CPU35	CPU36	CPU37	CPU38	CPU39	CPU40	CPU41	CPU42
90.1	88.8	88.8	91.1	91.1	89.7	89.6	91.8	91.9	90.7	90.6	90.7	90.6	93.1	92.9	93.4	93.5	92.0	92.2	94.2	94.4	93.4	93.4	93.4

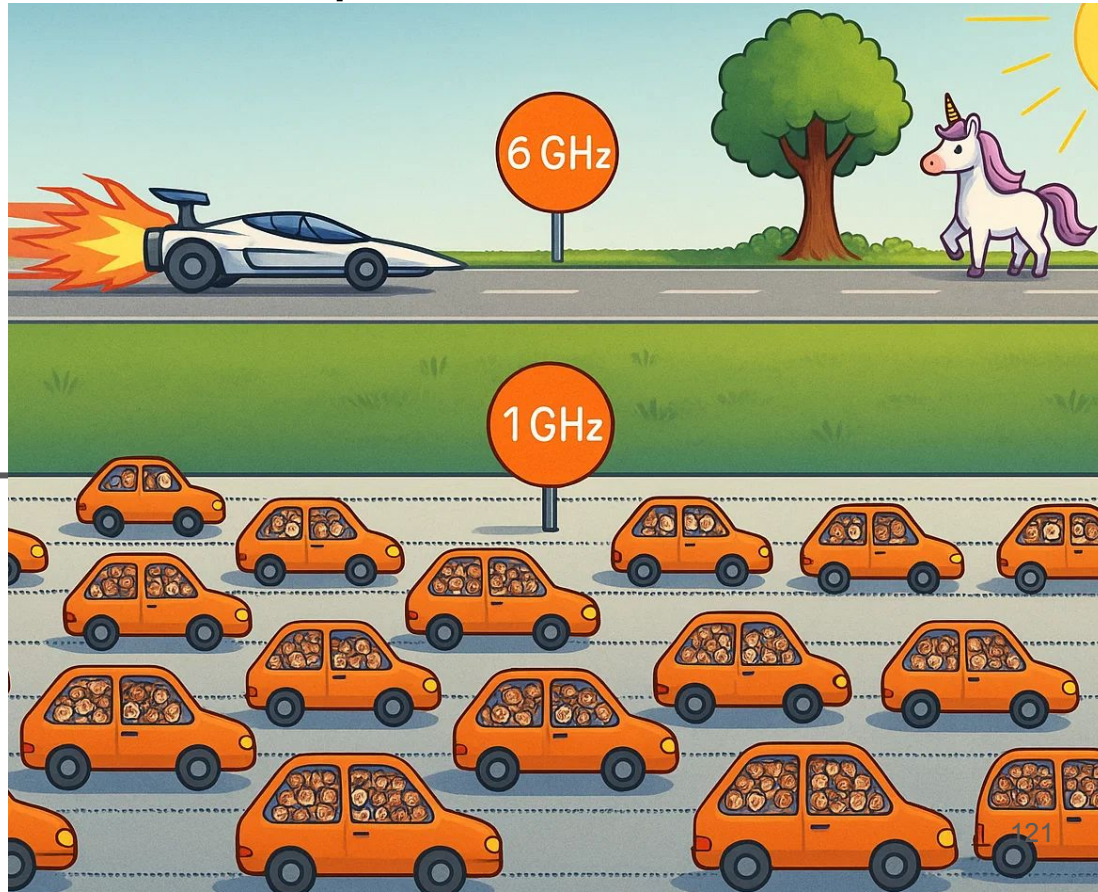
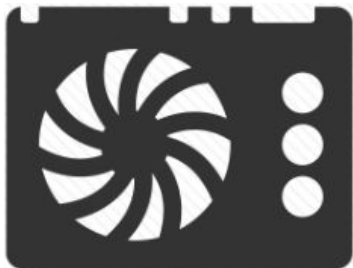
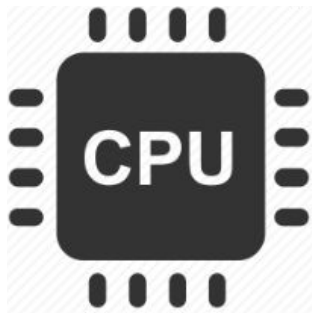
[illegible]

Перерыв!



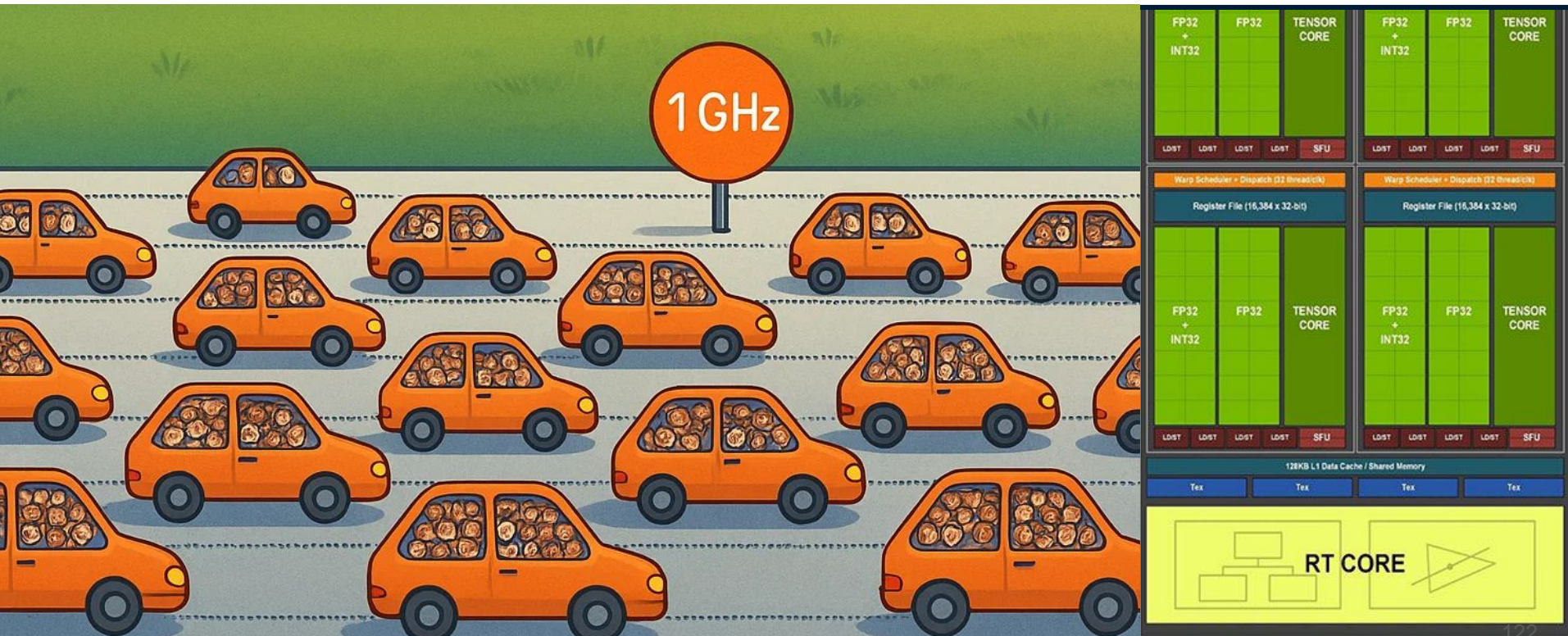
Глава 4: Модель вычислений массового параллелизма

Модель вычислений массового параллелизма



Модель вычислений массового параллелизма

SM - Streaming Multiprocessor



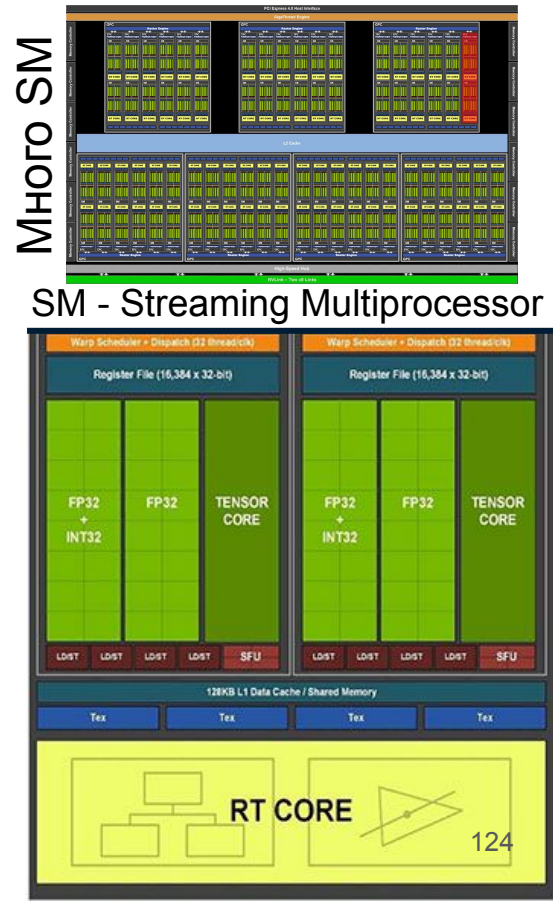
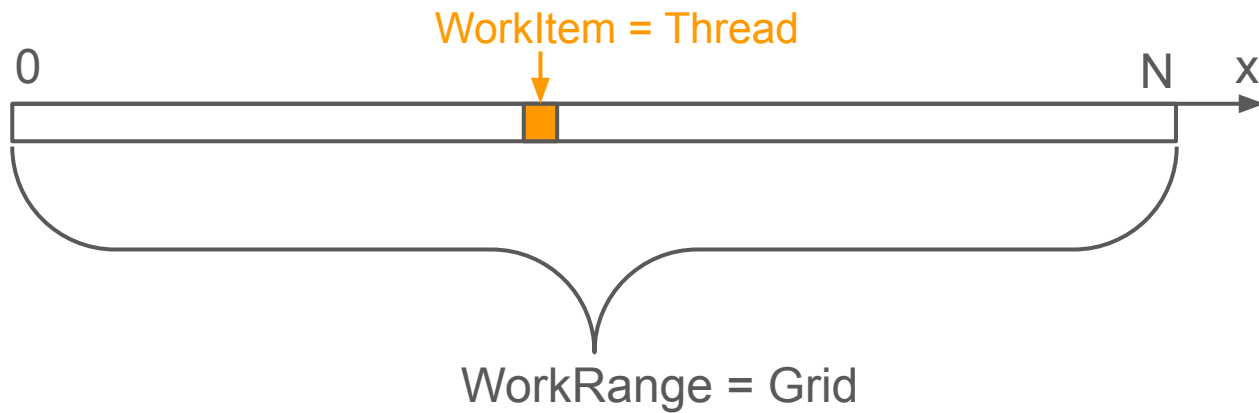
Модель вычислений массового параллелизма



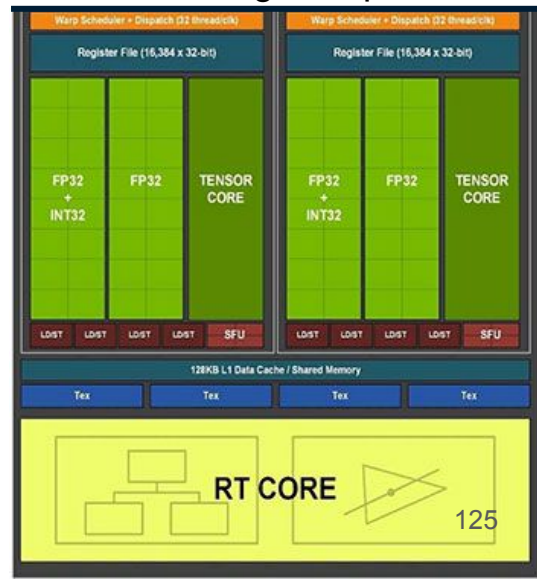
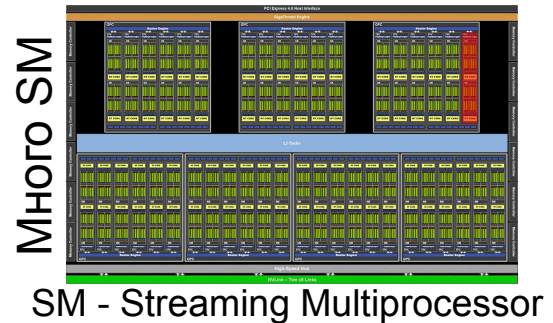
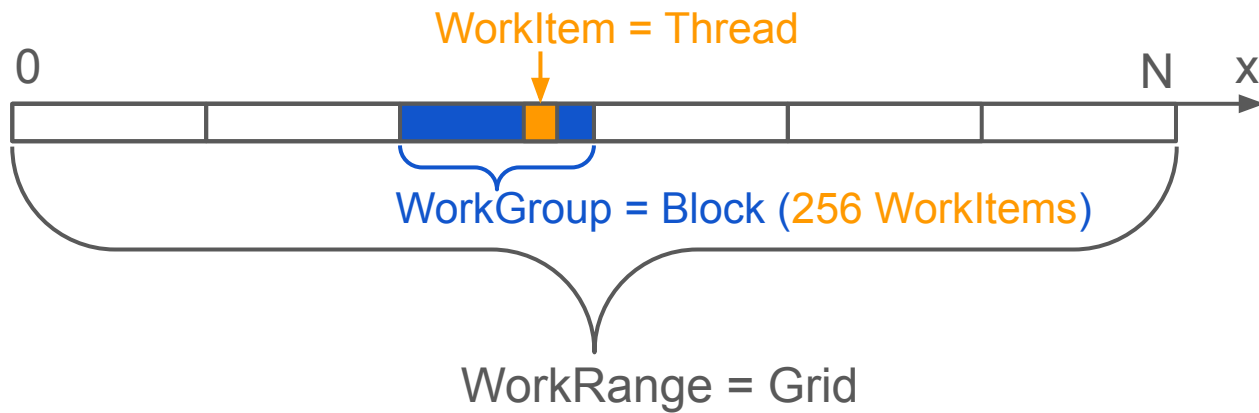
SM - Streaming Multiprocessor



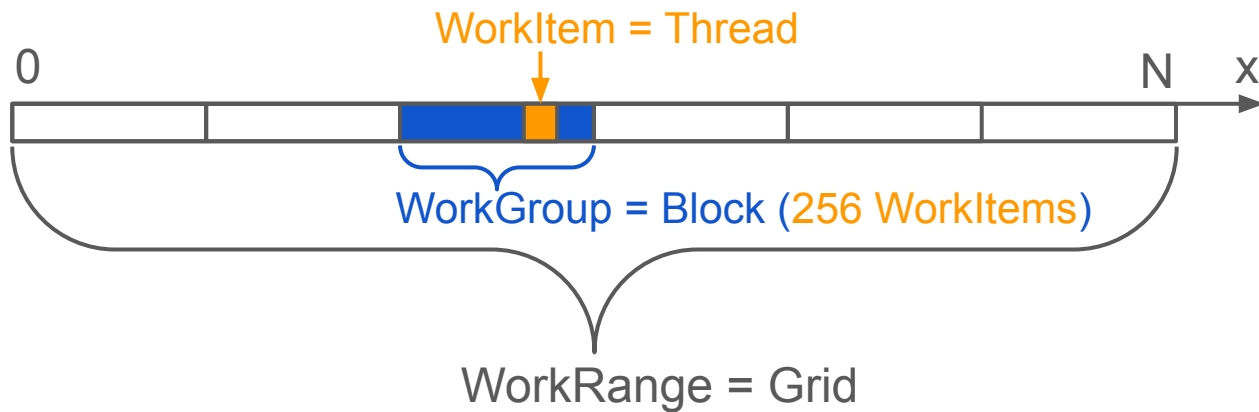
Модель вычислений массового параллелизма



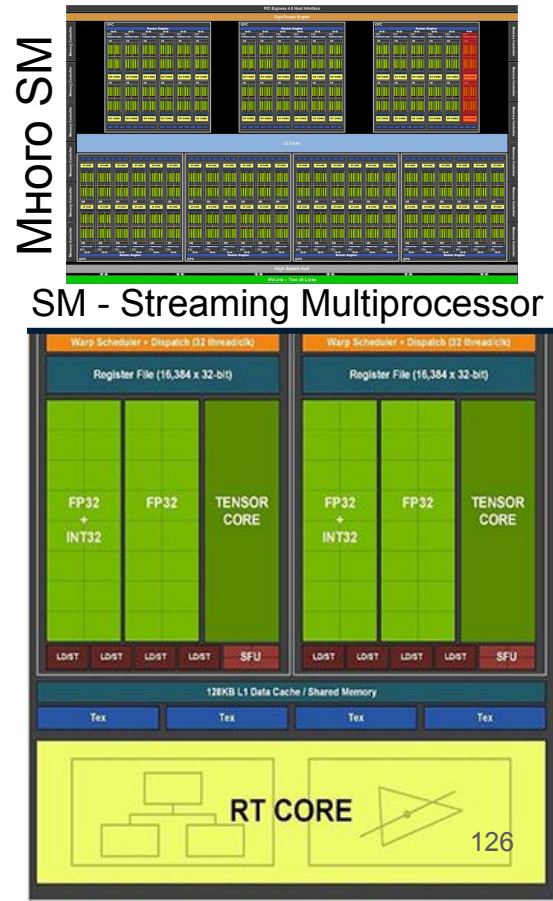
Модель вычислений массового параллелизма



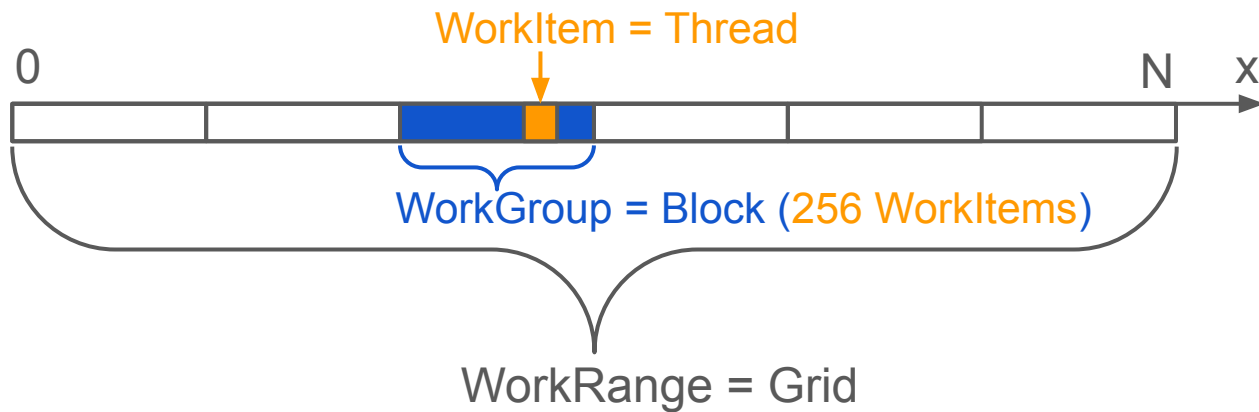
Модель вычислений массового параллелизма



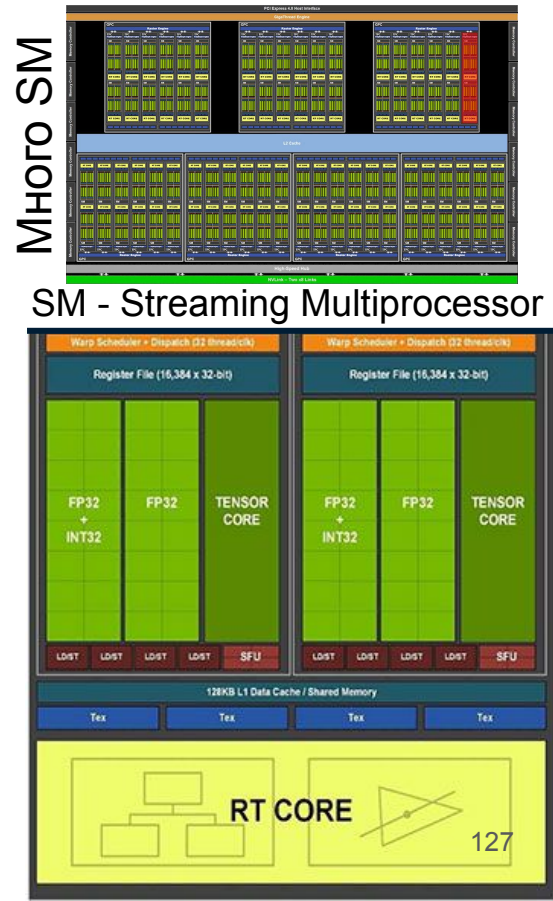
Но ведь в одном warp 32 потока!
(у **AMD**: 64 потока в wavefront)
Как так?



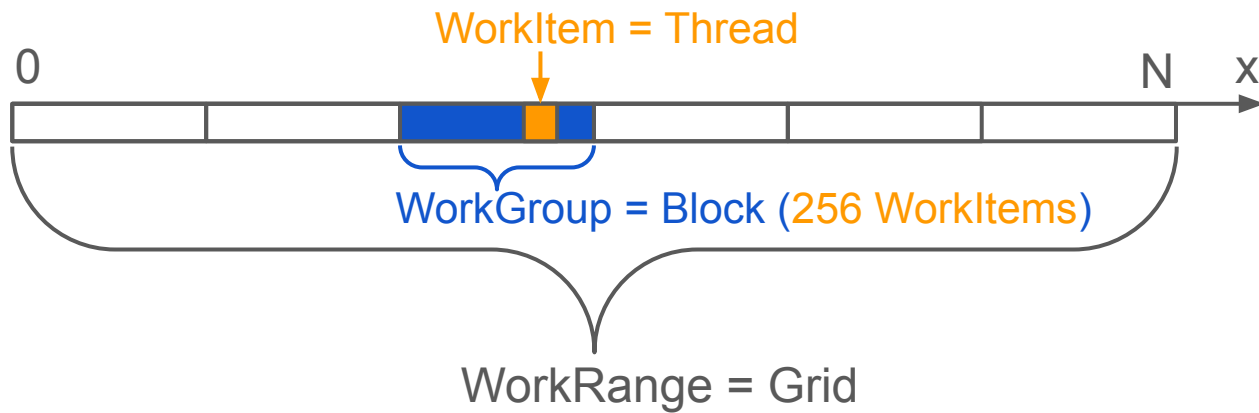
Модель вычислений массового параллелизма



WorkGroup = Block (256 WorkItems) = 8 x Warps (32 threads)

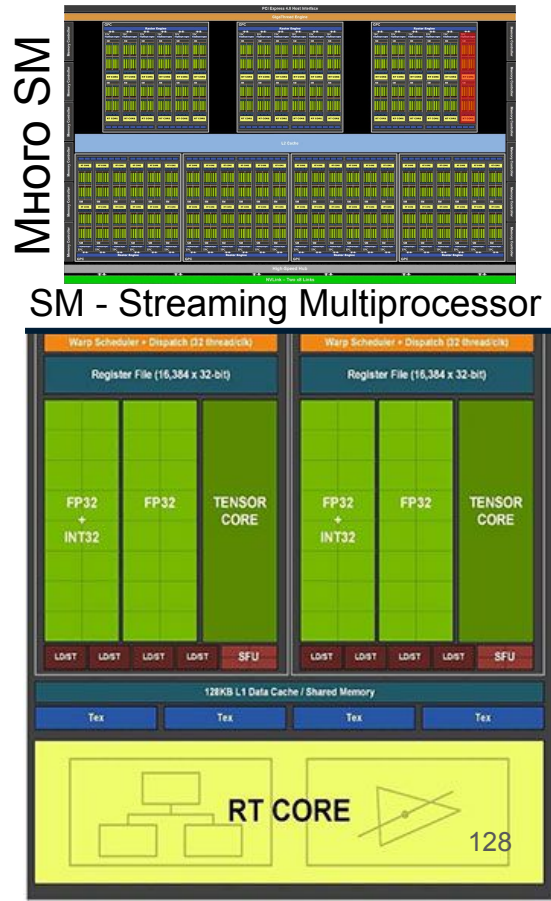


Модель вычислений массового параллелизма

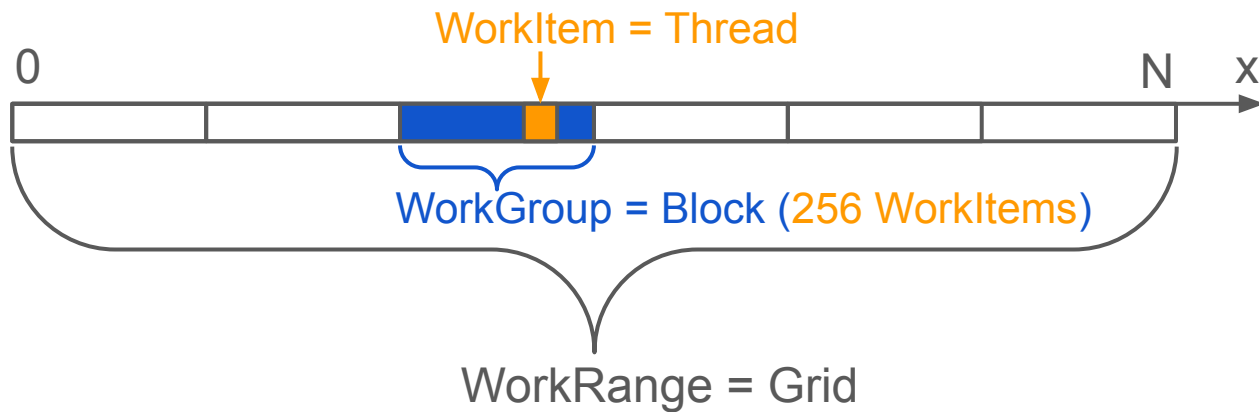


WorkGroup = Block (256 WorkItems) = 8 x Warps (32 threads)

Зачем нам контроль над **WorkGroups**?
Зачем на уровне API знать про **warps**?

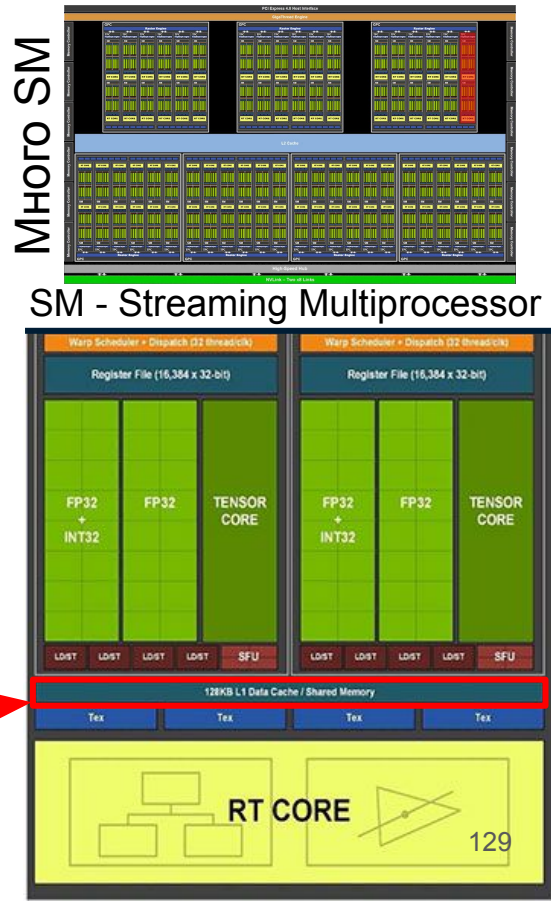


Модель вычислений массового параллелизма

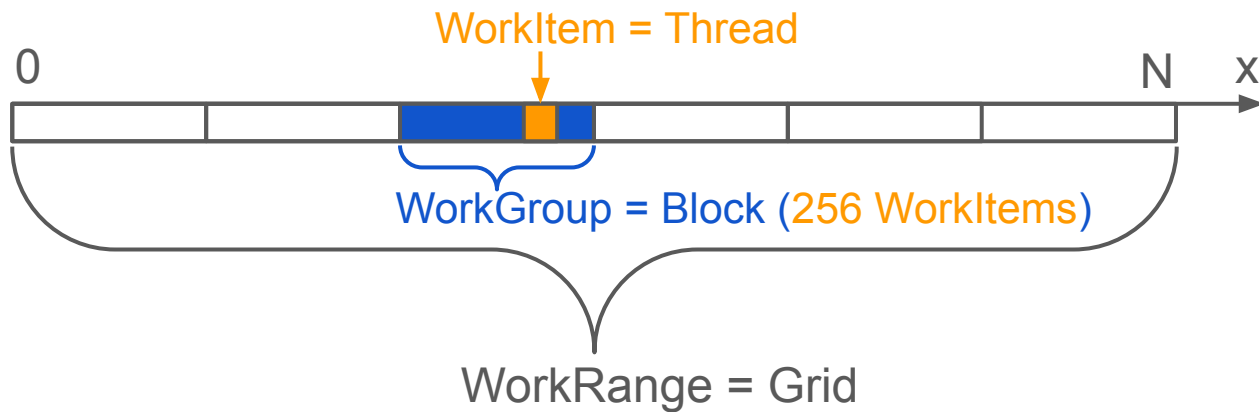


WorkGroup = Block (256 WorkItems) = 8 x Warps (32 threads)

- WorkItems коммуницируют через VRAM
- WorkItems в рамках одной WorkGroup общаются эффективнее - через **Shared/Local Memory (L1)**



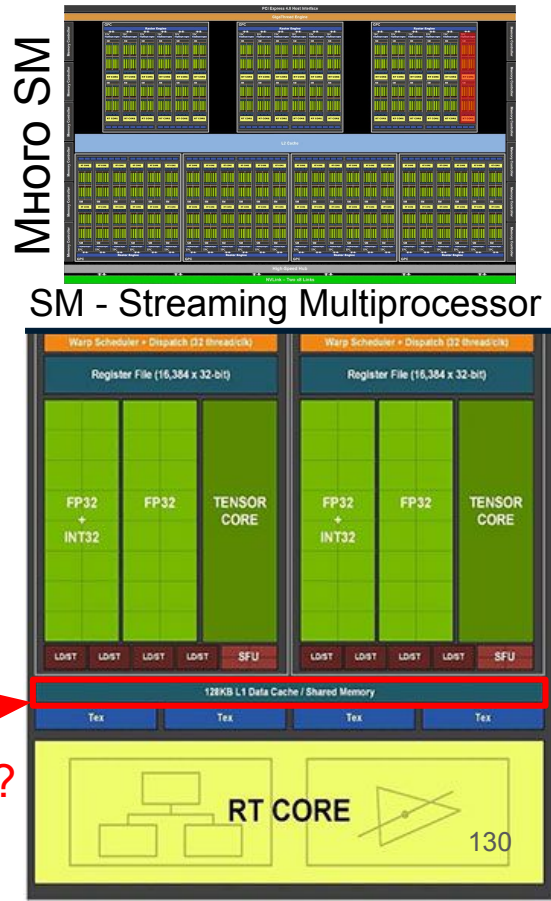
Модель вычислений массового параллелизма



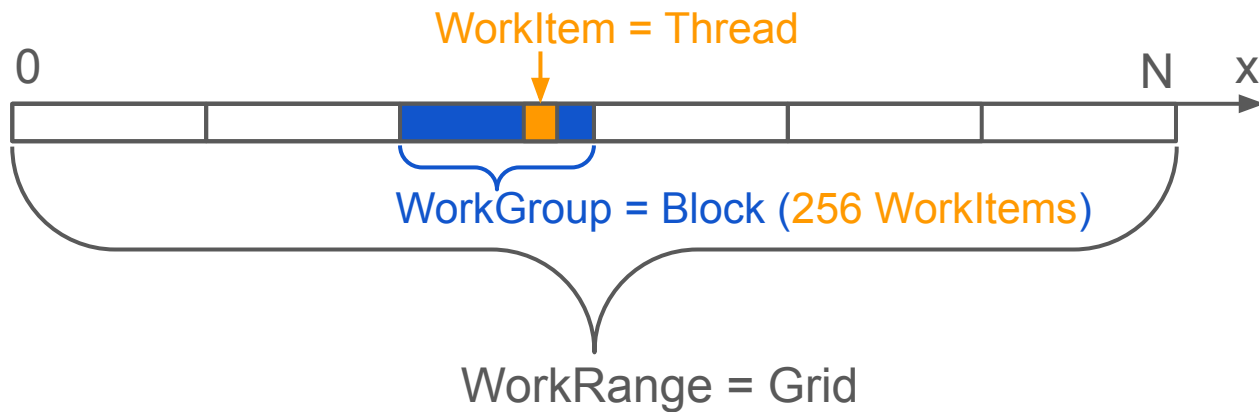
WorkGroup = Block (256 WorkItems) = 8 x Warps (32 threads)

- WorkItems коммуницируют через VRAM
- WorkItems в рамках одной WorkGroup общаются эффективнее - через **Shared/Local Memory (L1)**

Могут ли **warp**-ы одной **WorkGroup** исполняться на разных SM?



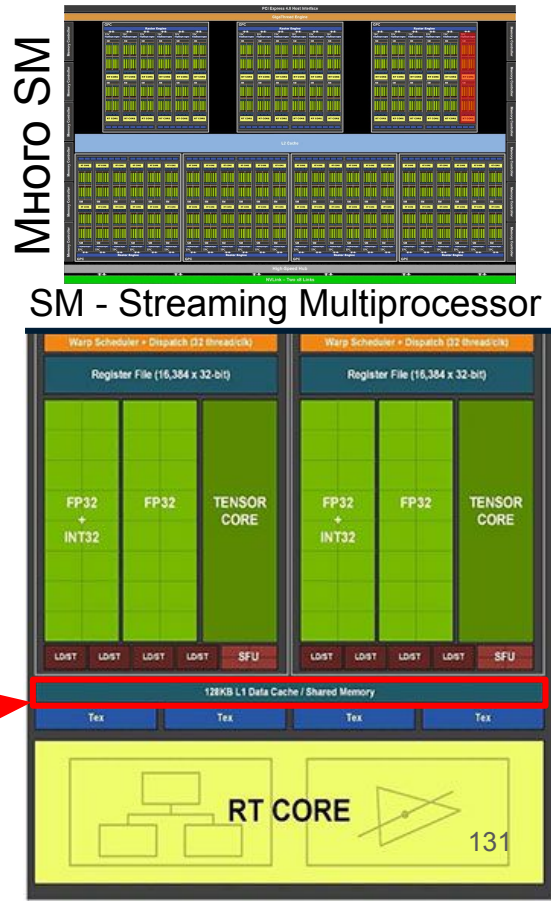
Модель вычислений массового параллелизма



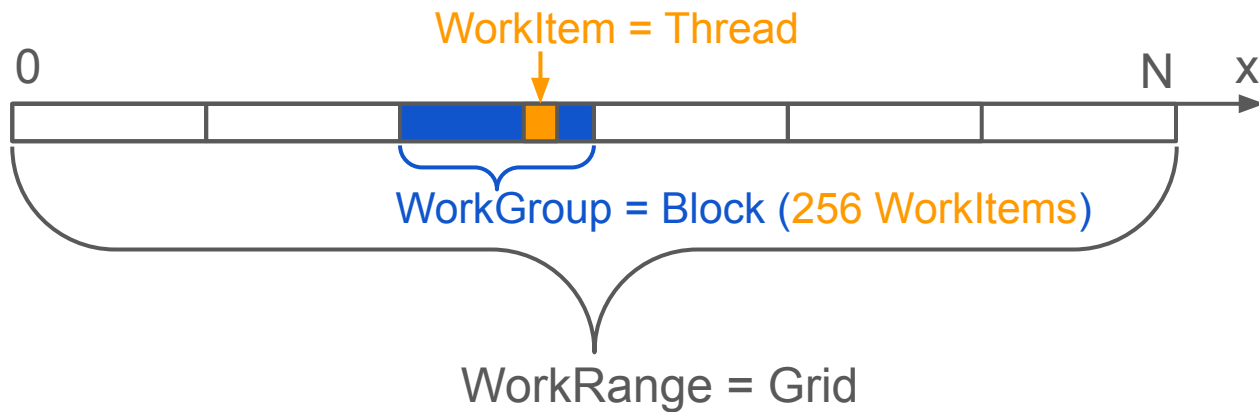
WorkGroup = Block (256 WorkItems) = 8 x Warps (32 threads)

- WorkItems коммуницируют через VRAM
- WorkItems в рамках одной WorkGroup общаются эффективнее - через **Shared/Local Memory (L1)**

А если один такой поток записал в **Local Memory** число - увидит ли другой поток этой **WorkGroup**?



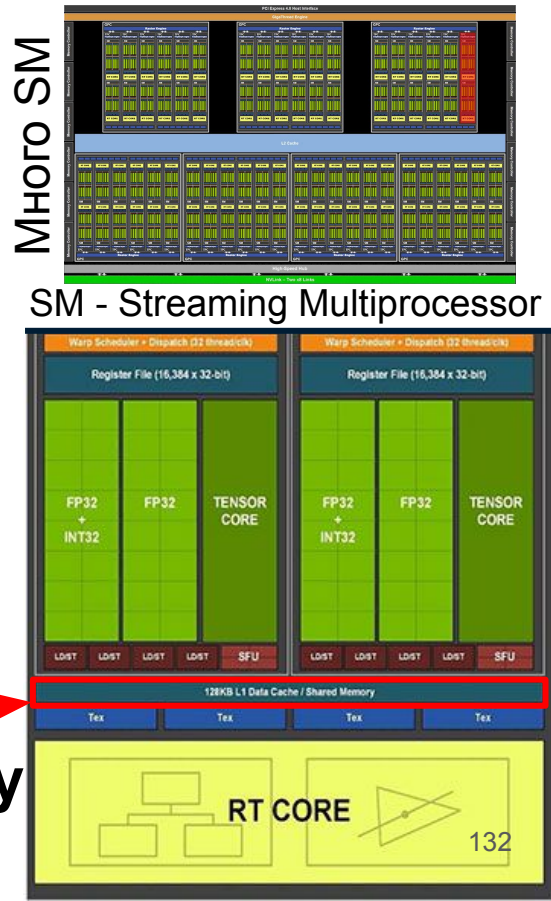
Модель вычислений массового параллелизма



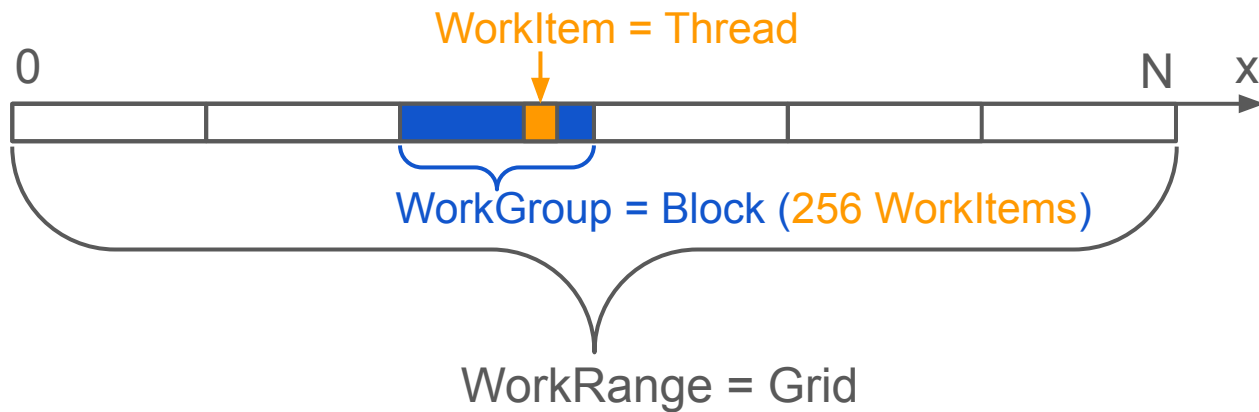
WorkGroup = Block (256 WorkItems) = 8 x Warps (32 threads)

- WorkItems коммуницируют через VRAM
- WorkItems в рамках одной WorkGroup общаются эффективнее - через **Shared/Local Memory (L1)**

А если один такой поток записал в **Local Memory** memory число - увидит ли другой поток этой **WorkGroup**? **barrier!**

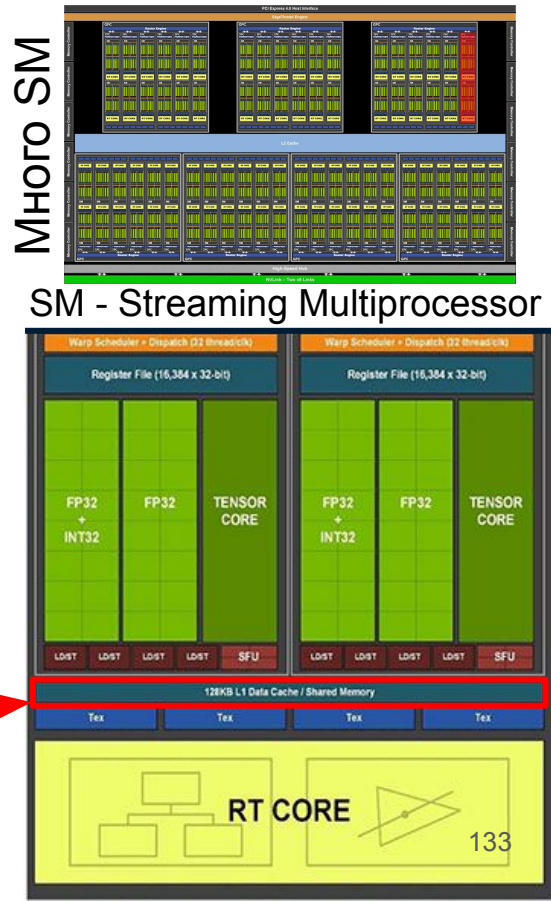


Модель вычислений массового параллелизма



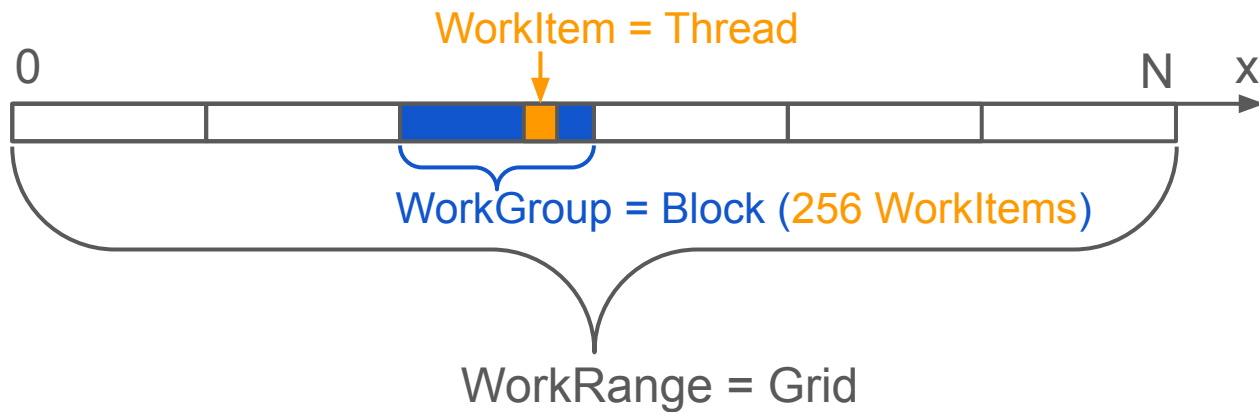
WorkGroup = Block (256 WorkItems) = 8 x Warps (32 threads)

- **WorkItems** коммуницируют через VRAM
- **WorkItems** в рамках одной **WorkGroup** общаются эффективнее - через **Shared/Local Memory (L1)**



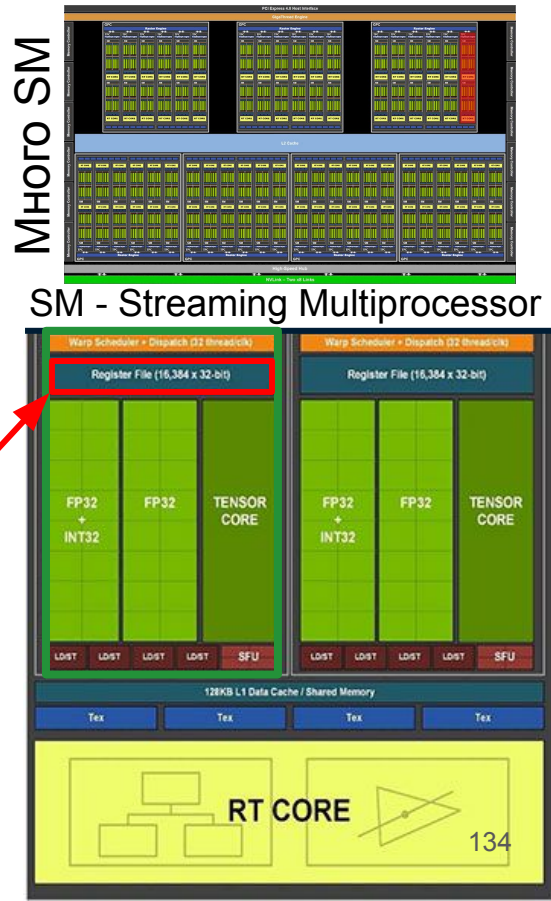
Могут ли потоки одного warp-а общаться еще эффективнее?

Модель вычислений массового параллелизма

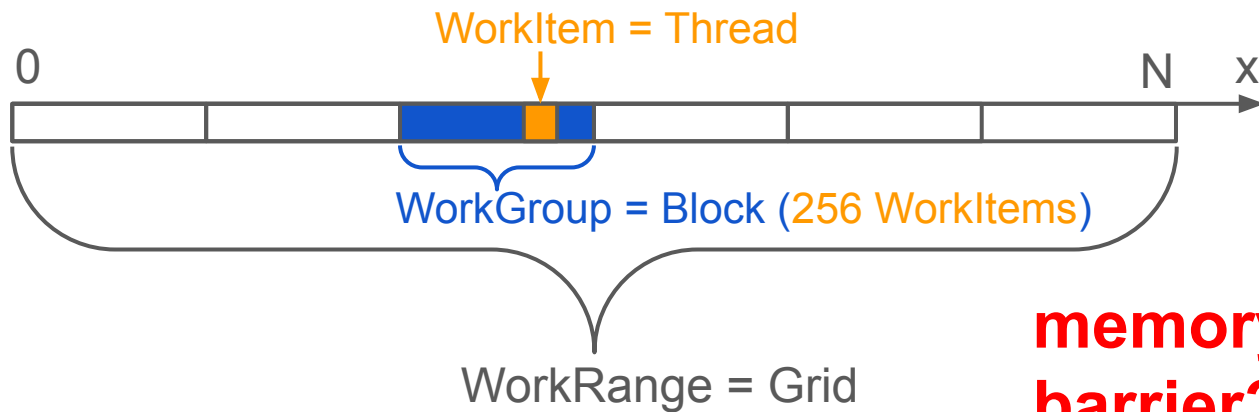


WorkGroup = Block (256 WorkItems) = 8 x Warps (32 threads)

- WorkItems коммуницируют через VRAM
- WorkItems в рамках одной WorkGroup общаются эффективнее - через **Shared/Local Memory (L1)**
- WorkItems в рамках одного warp-а могут подглядывать друг другу в **регистры** (shuffle instr.)



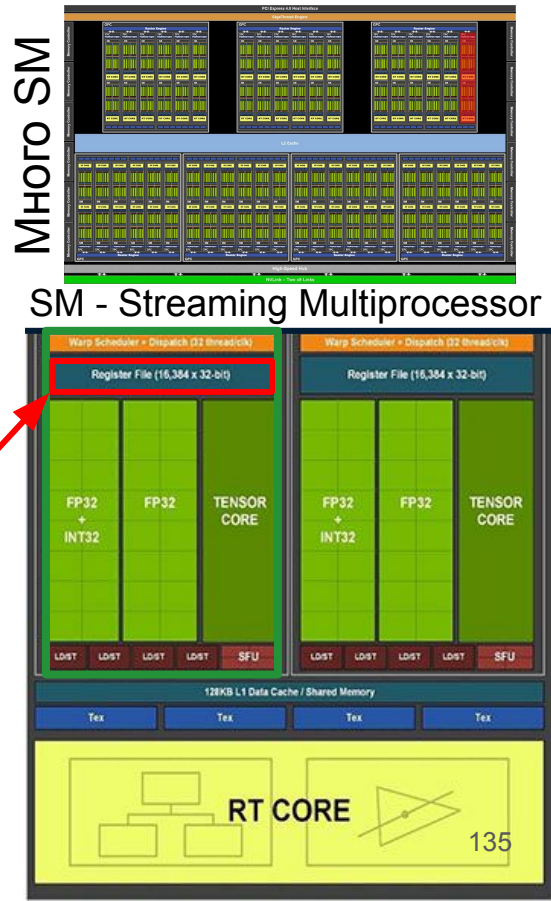
Модель вычислений массового параллелизма



WorkGroup = Block (256 WorkItems) = 8 x Warps (32 threads)

- WorkItems коммуницируют через VRAM
- WorkItems в рамках одной WorkGroup общаются эффективнее - через **Shared/Local Memory (L1)**
- WorkItems в рамках одного warp-а могут подглядывать друг другу в **регистры** (shuffle instr.)

memory barrier?

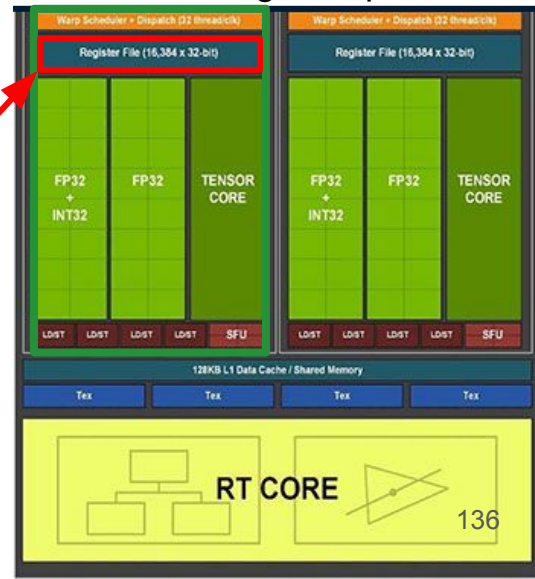


GLSL (Graphics Library Shading Language)

dFdx(...), dFdy(...) — return the partial derivative of an argument
with respect to x or y

- **WorkItems** коммуницируют через VRAM
- **WorkItems** в рамках одной **WorkGroup** общаются эффективнее - через **Shared/Local Memory (L1)**
- **WorkItems** в рамках одного warp-а могут подглядывать друг другу в **регистры** (**shuffle instr.**)

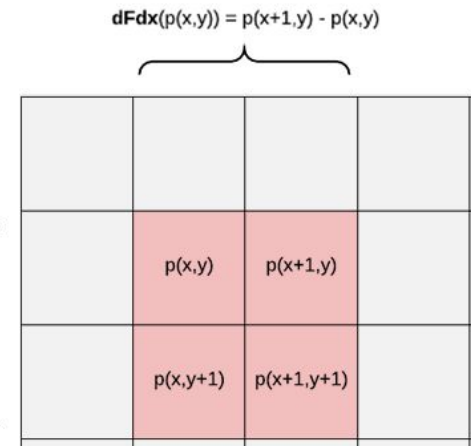
SM - Streaming Multiprocessor



GLSL (Graphics Library Shading Language)

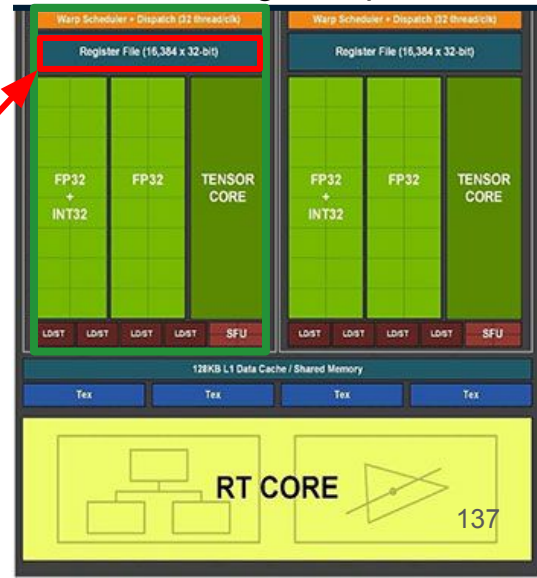
dFdx(...), dFdy(...) — return the partial derivative of an argument with respect to x or y

$$dFdy(p(x,y)) = p(x,y+1) - p(x,y)$$



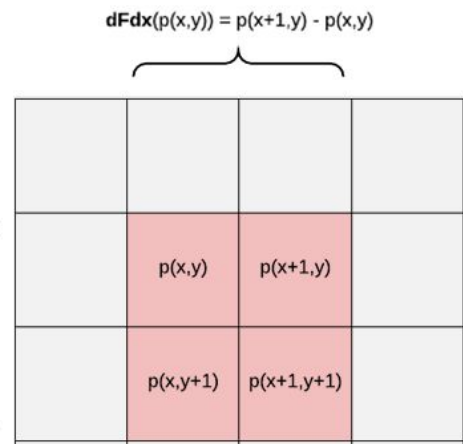
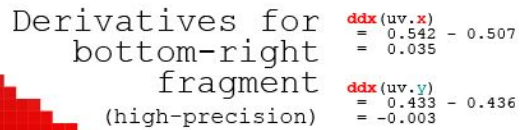
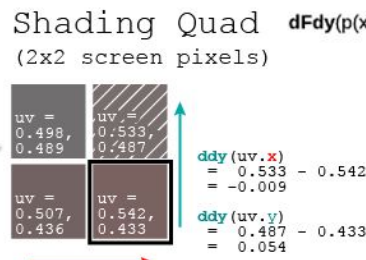
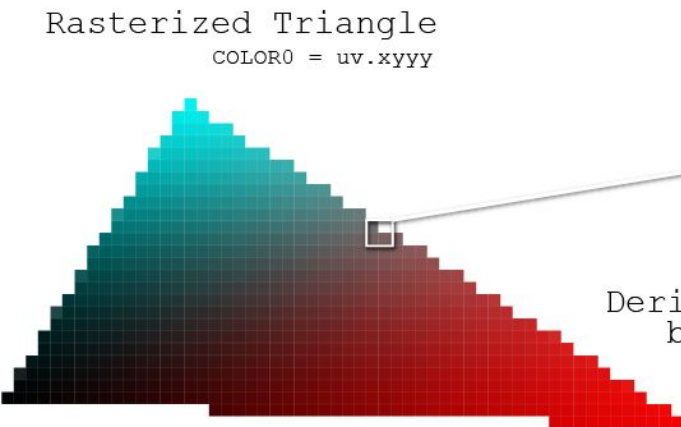
SM - Streaming Multiprocessor

- **WorkItems** коммуницируют через VRAM
- **WorkItems** в рамках одной **WorkGroup** общаются эффективнее - через **Shared/Local Memory (L1)**
- **WorkItems** в рамках одного warp-а могут подглядывать друг другу в **регистры** (**shuffle instr.**)

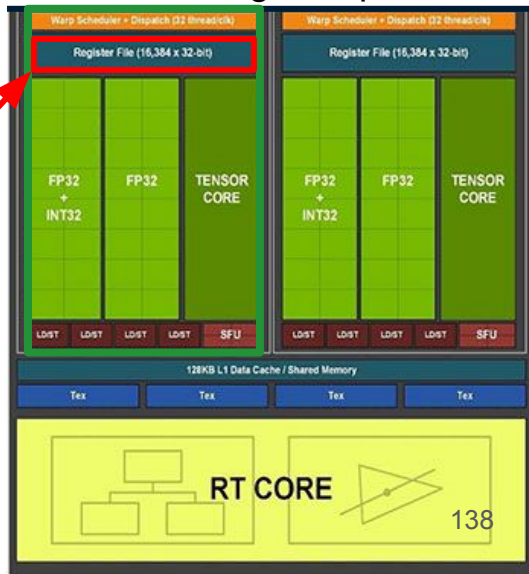


GLSL (Graphics Library Shading Language)

dFdx(...), dFdy(...) — return the partial derivative of an argument with respect to x or y

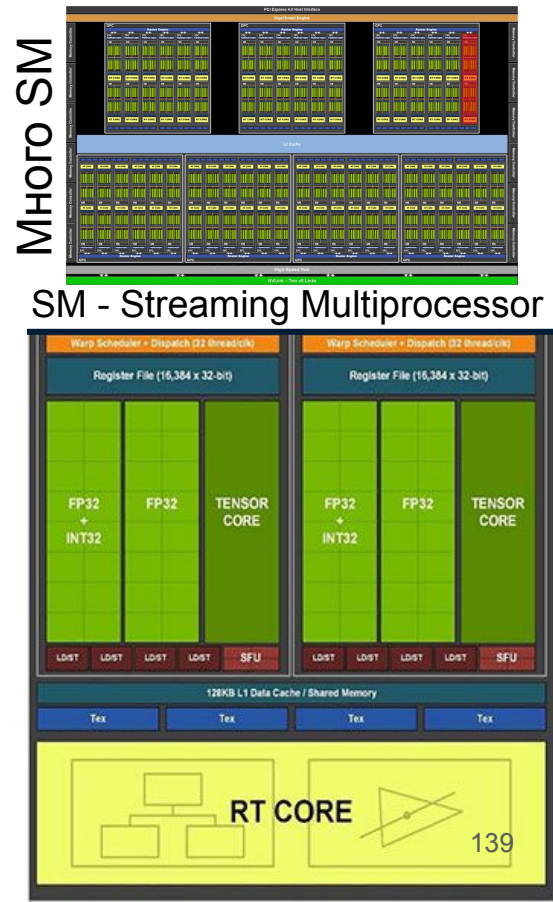
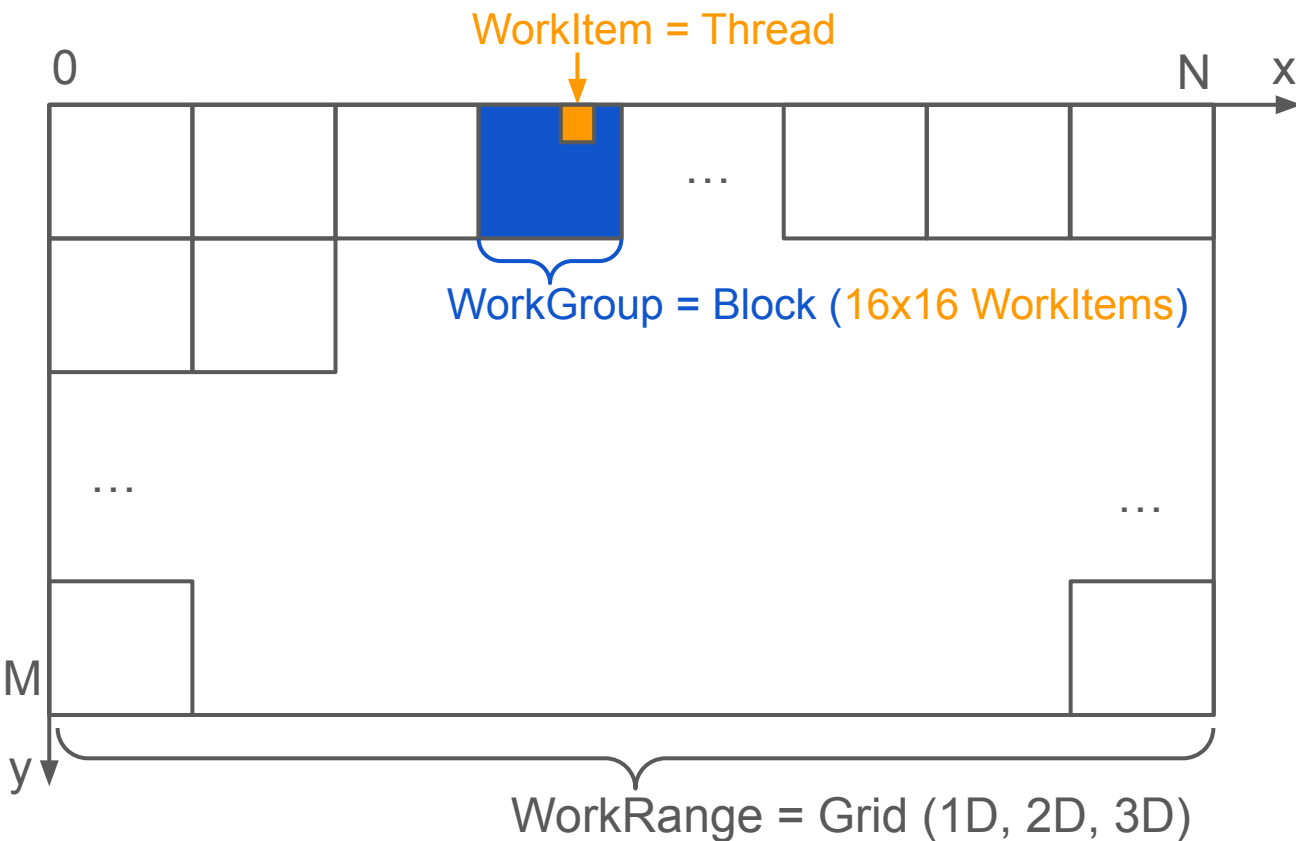


SM - Streaming Multiprocessor

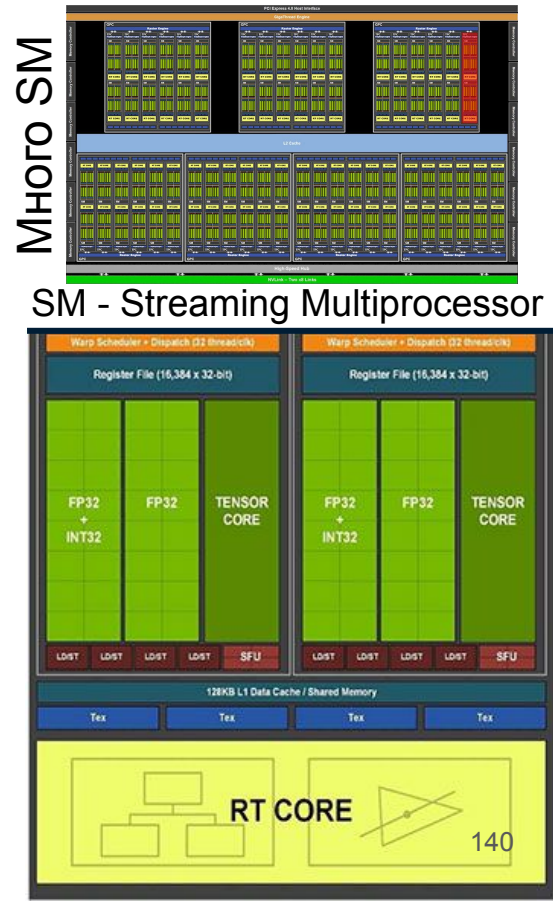
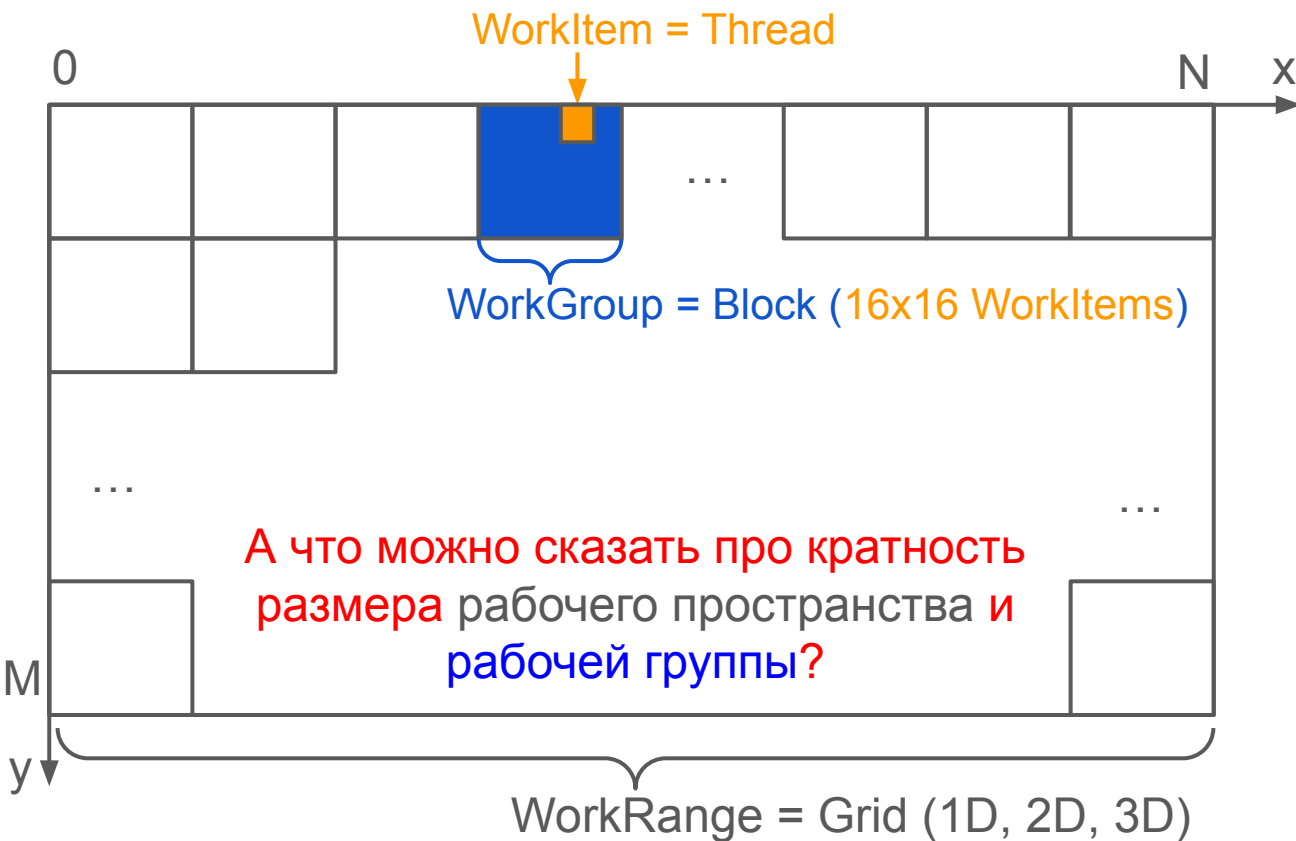


- **WorkItems** коммуницируют через VRAM
- **WorkItems** в рамках одной **WorkGroup** общаются эффективнее - через **Shared/Local Memory (L1)**
- **WorkItems** в рамках одного warp-а могут подглядывать друг другу в **регистры** (**shuffle instr.**)

Модель вычислений массового параллелизма



Модель вычислений массового параллелизма

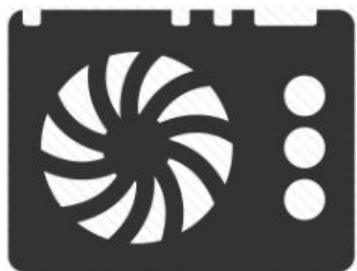


Глава 5: Профилирование и оптимизация

МНОГО вычислений? (**ALU**)

МНОГО читать+писать? (**VRAM**)

$100 \cdot 10^{12}$ FLOPS



1000 GB/s



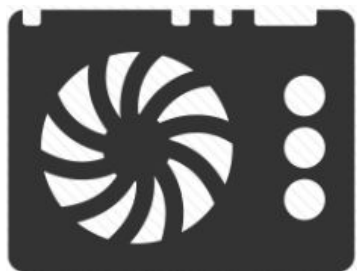
VRAM - GDDR6

Чего больше?
Как сравнить яблоки и
апельсины?

МНОГО вычислений? (**ALU**)

МНОГО читать+писать? (**VRAM**)

$100 \cdot 10^{12}$ FLOPS



1000 GB/s



VRAM - GDDR6

Чего больше?

МНОГО вычислений? (**ALU**)

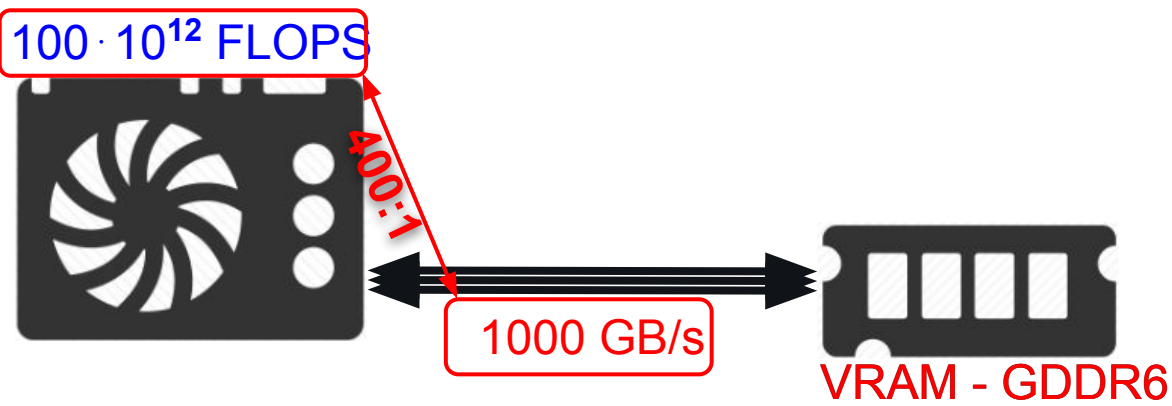
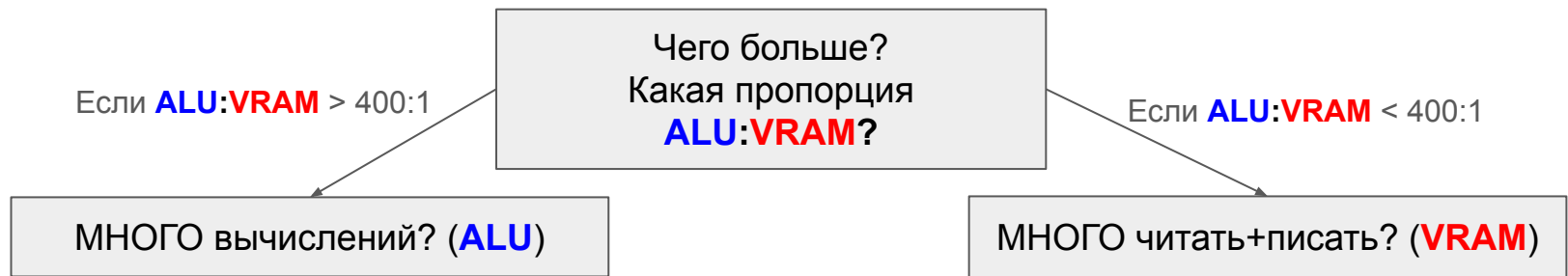
МНОГО читать+писать? (**VRAM**)

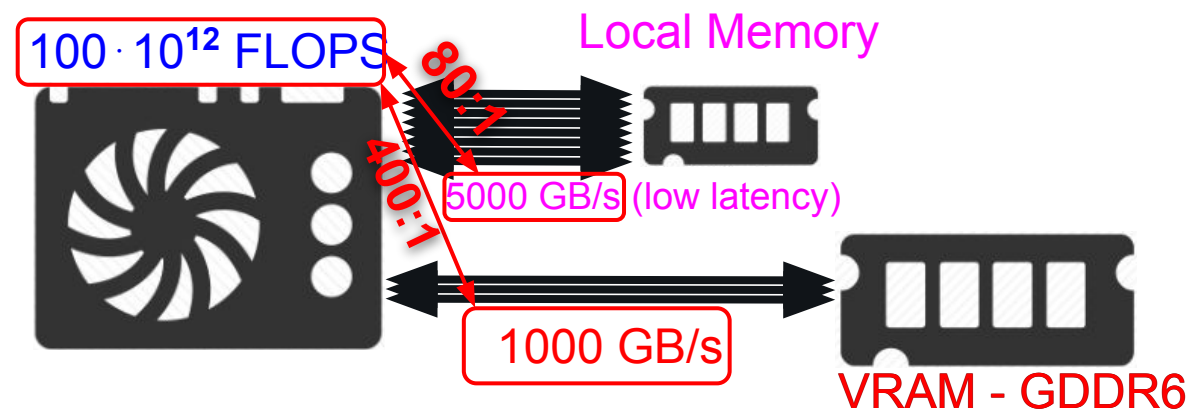
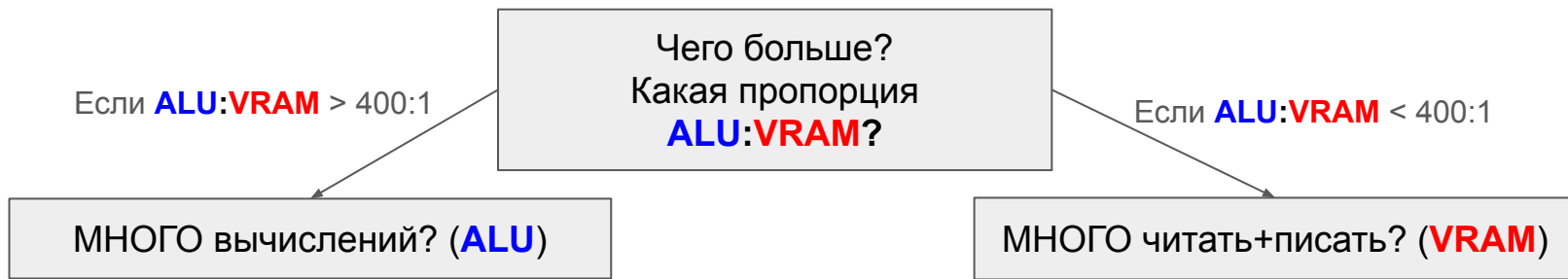
$100 \cdot 10^{12}$ FLOPS

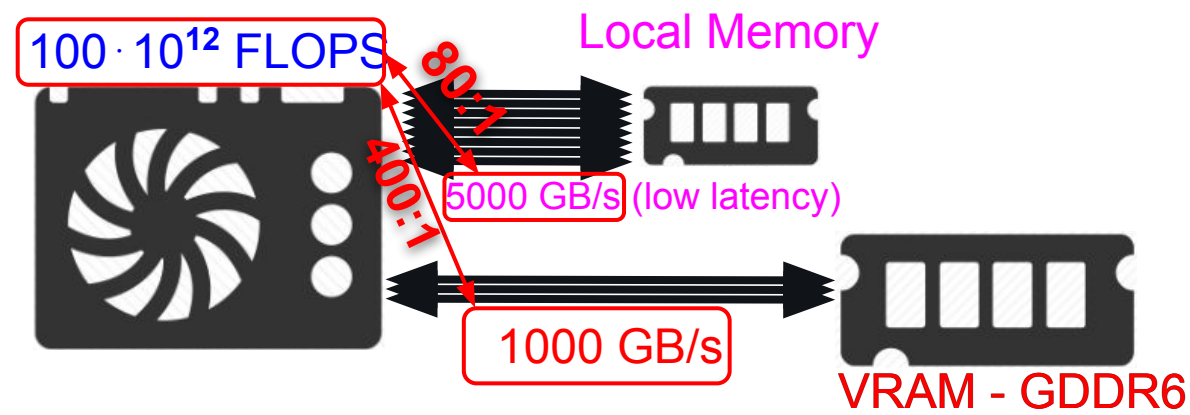
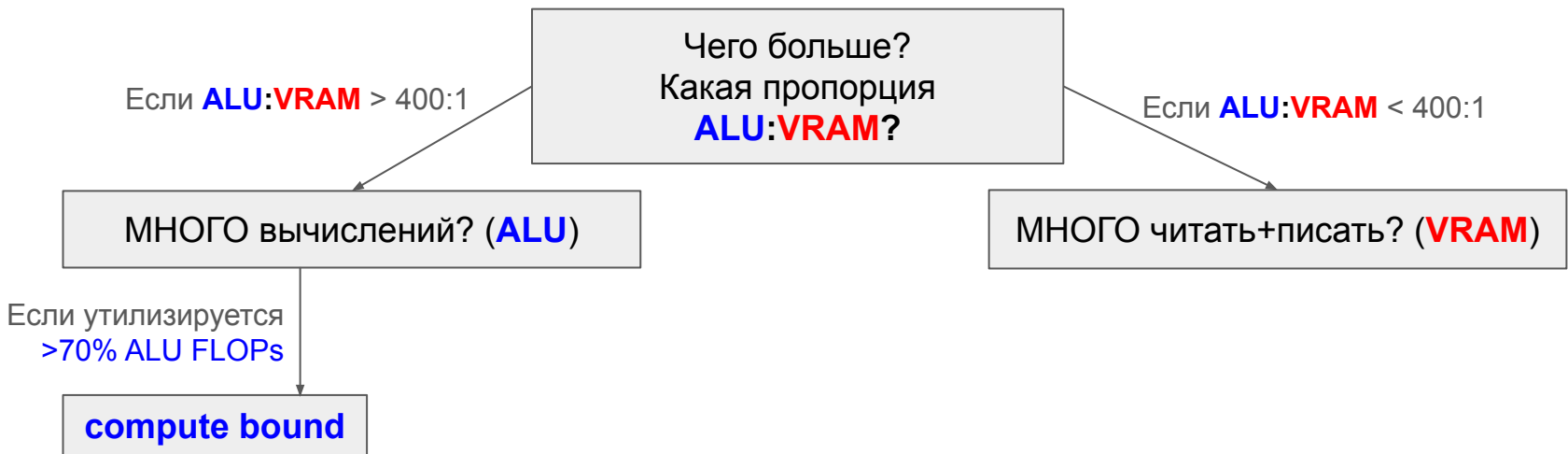
400:1

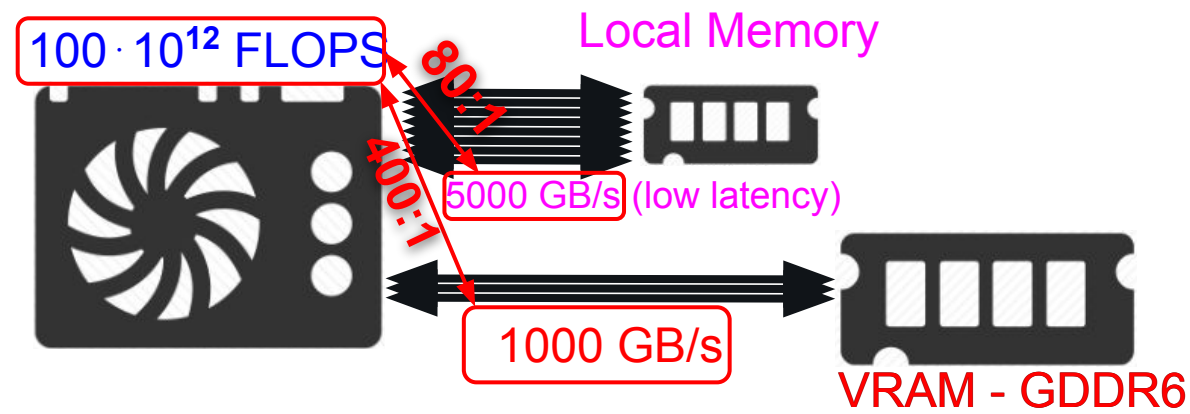
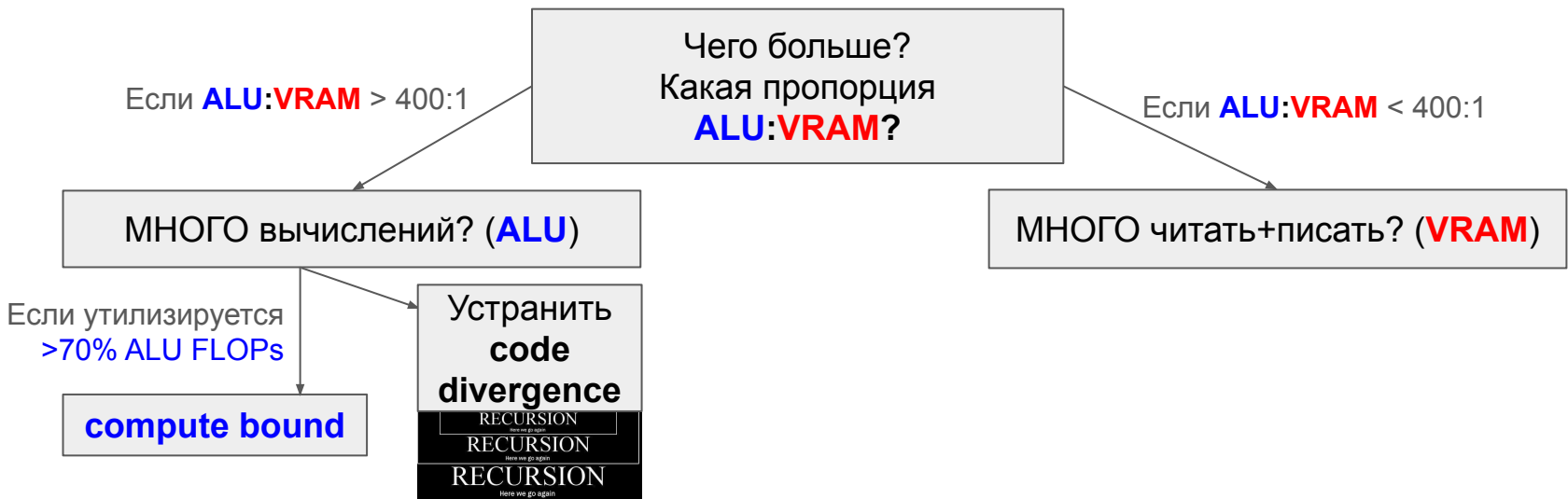
1000 GB/s

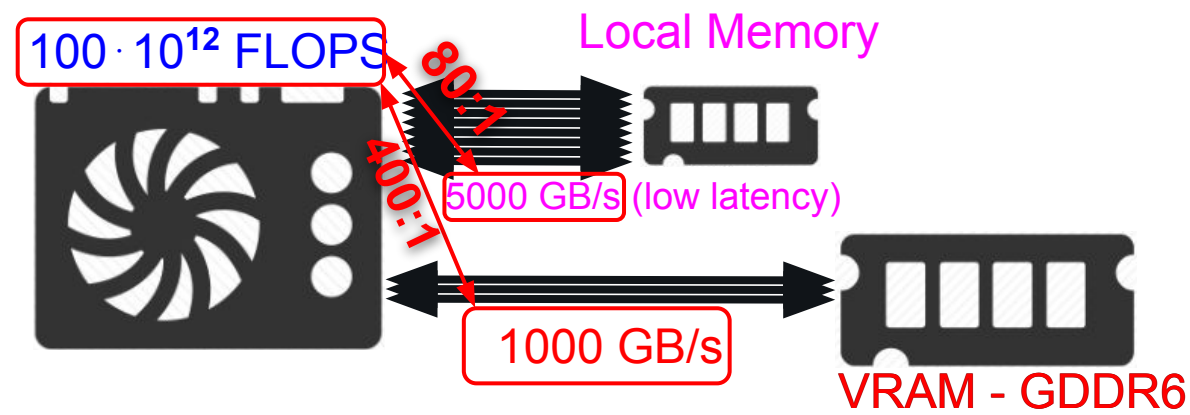
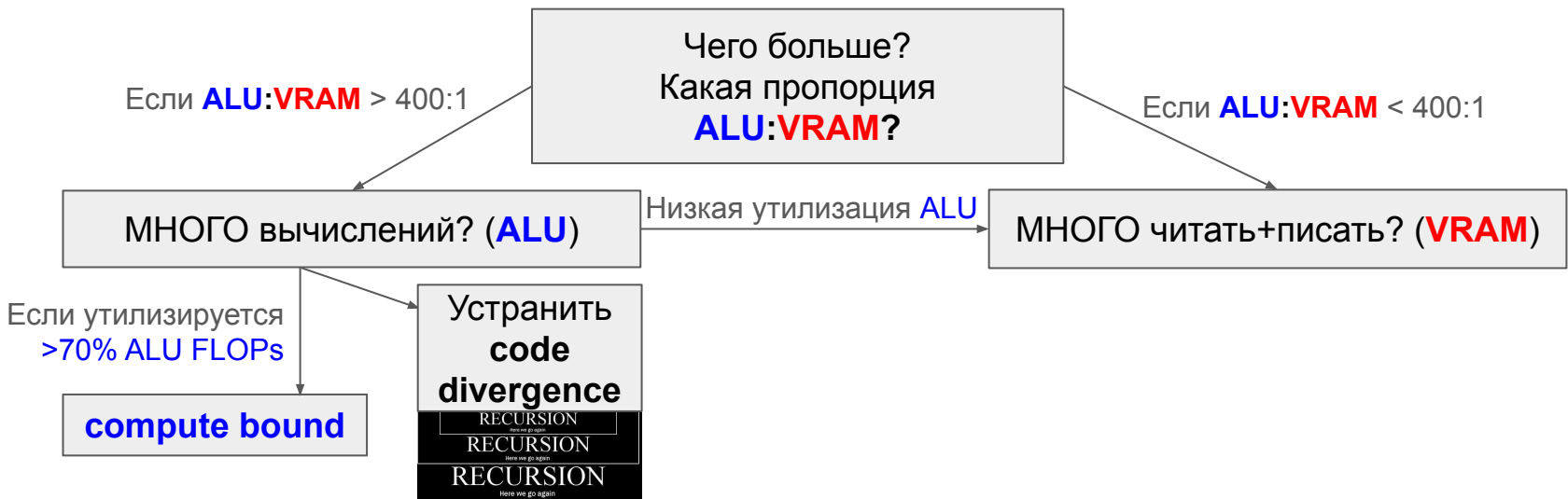
VRAM - GDDR6

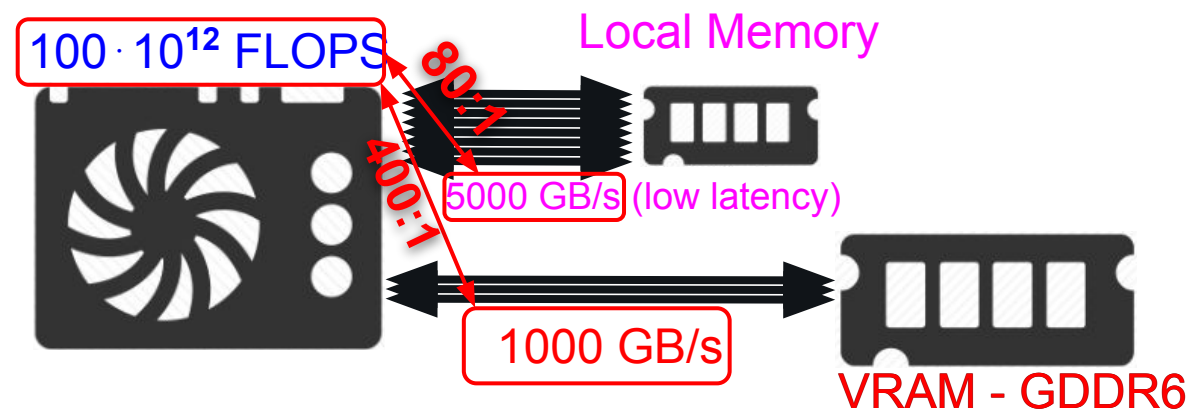


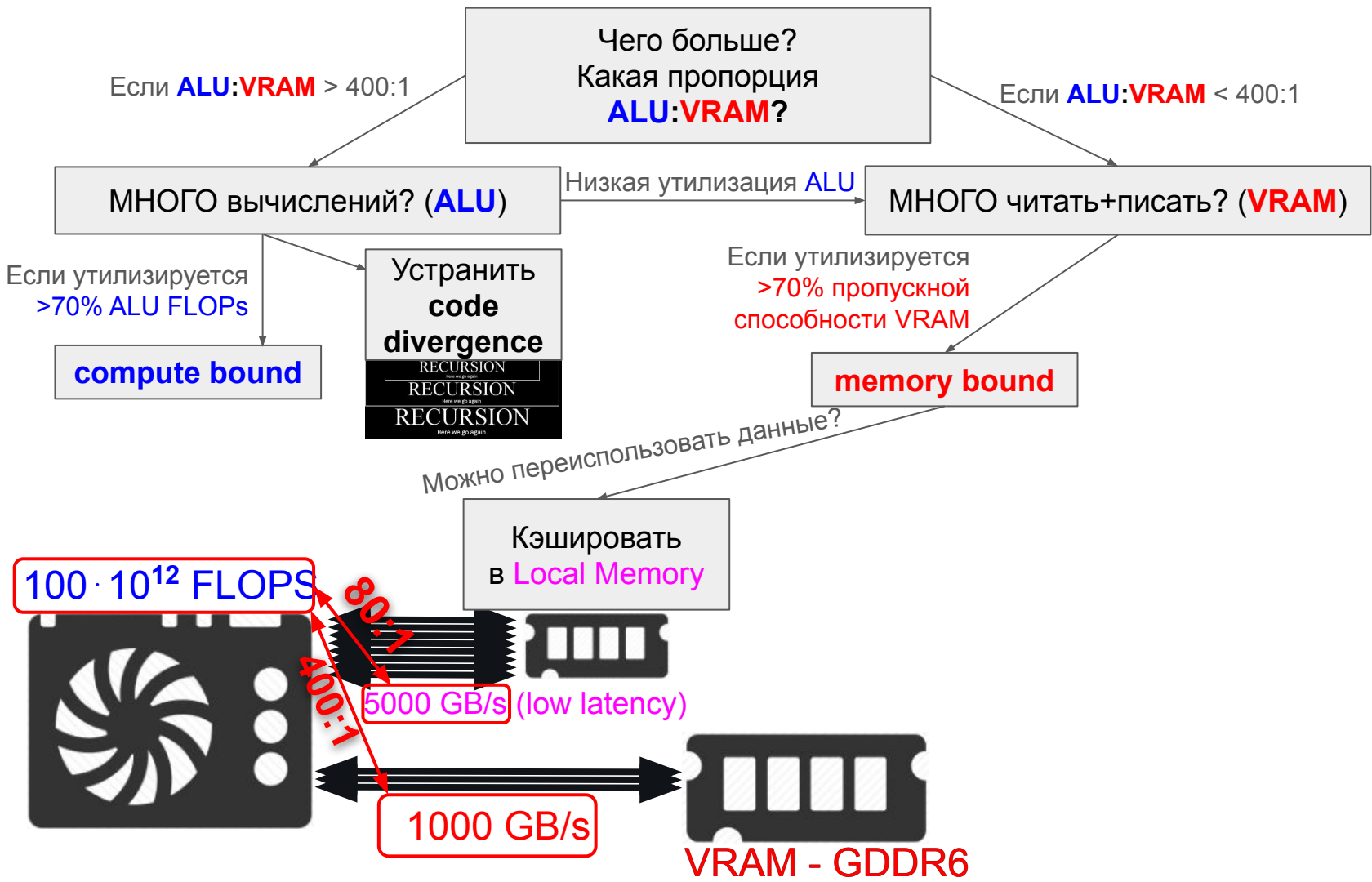


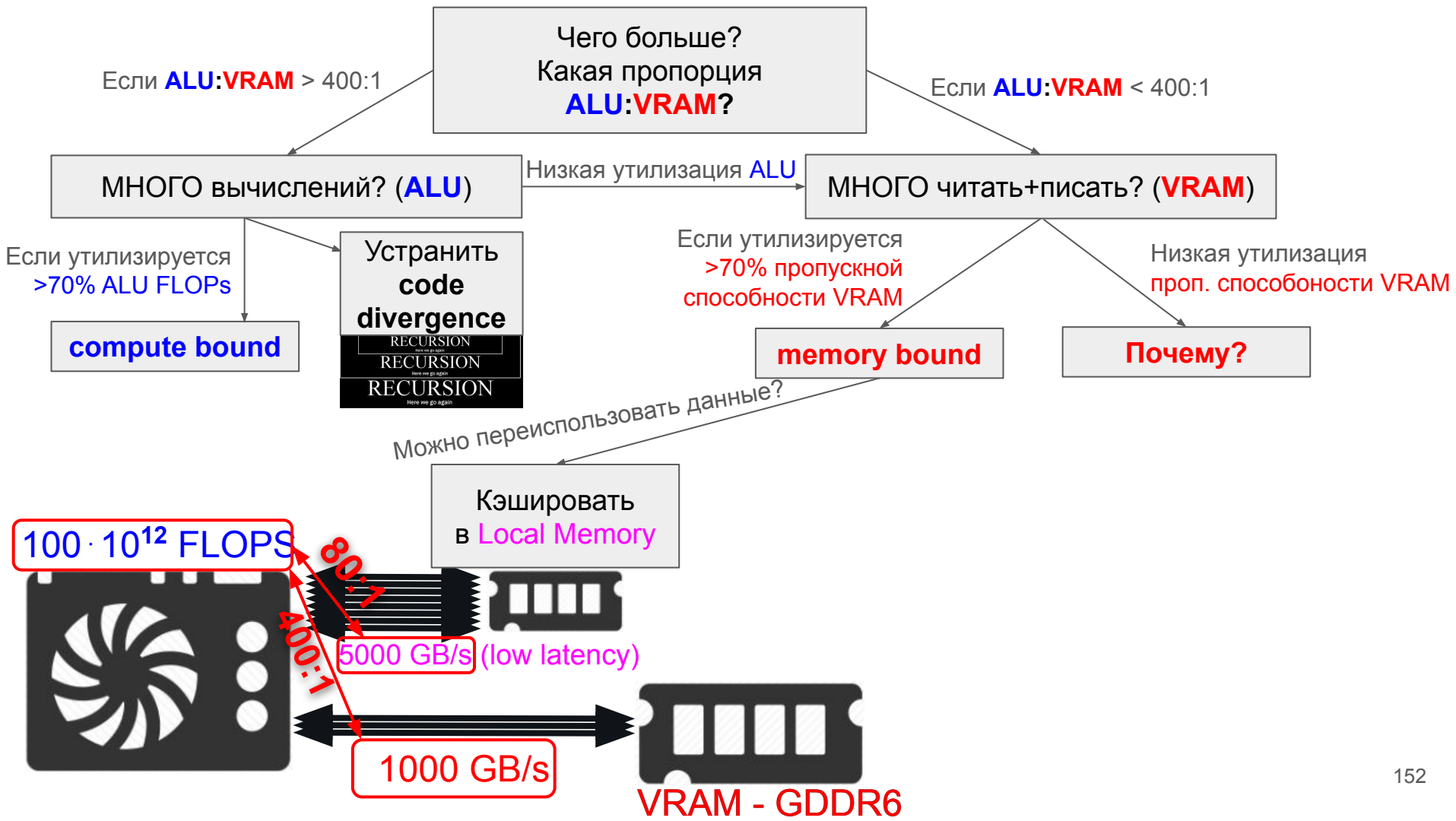


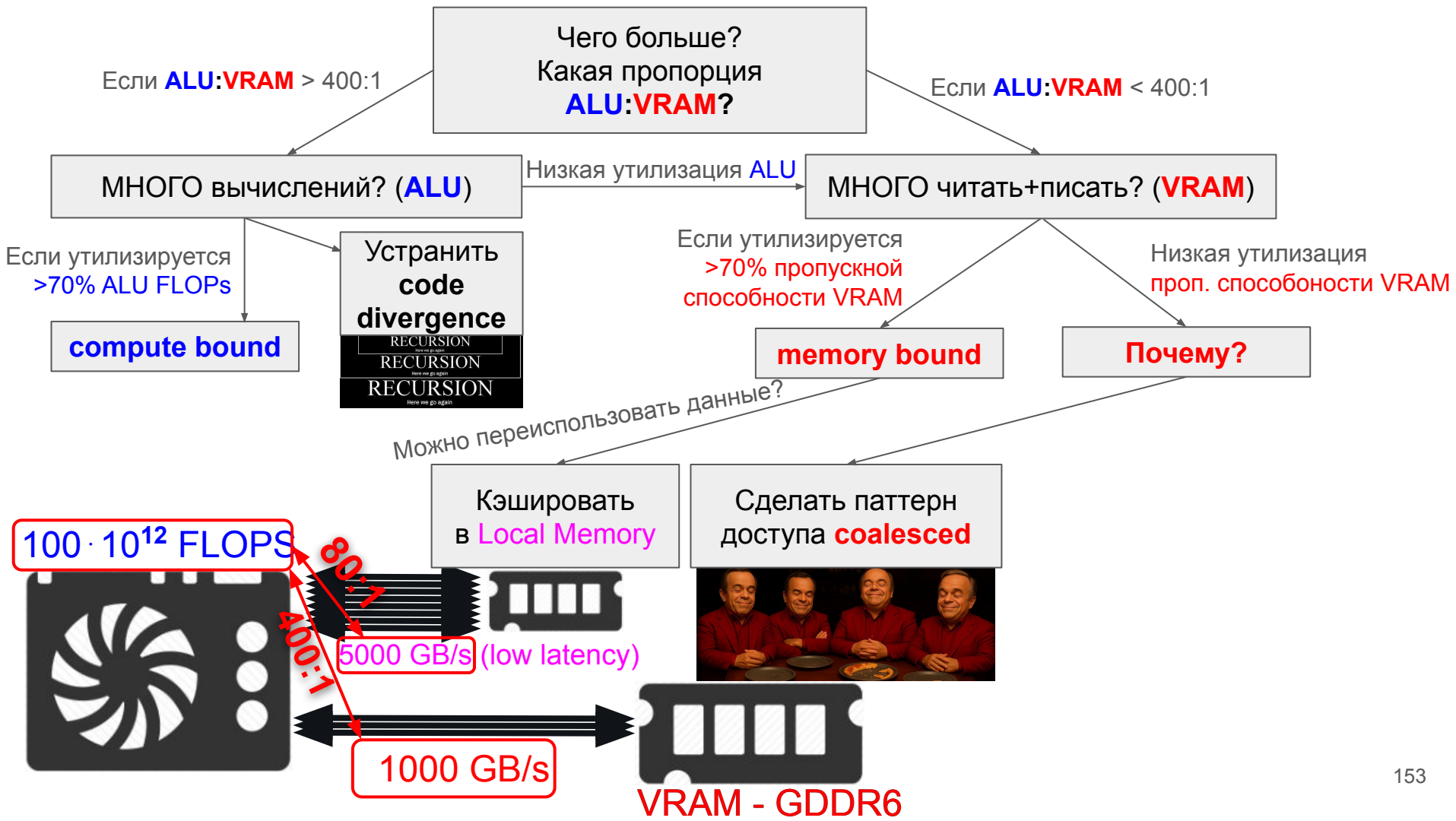


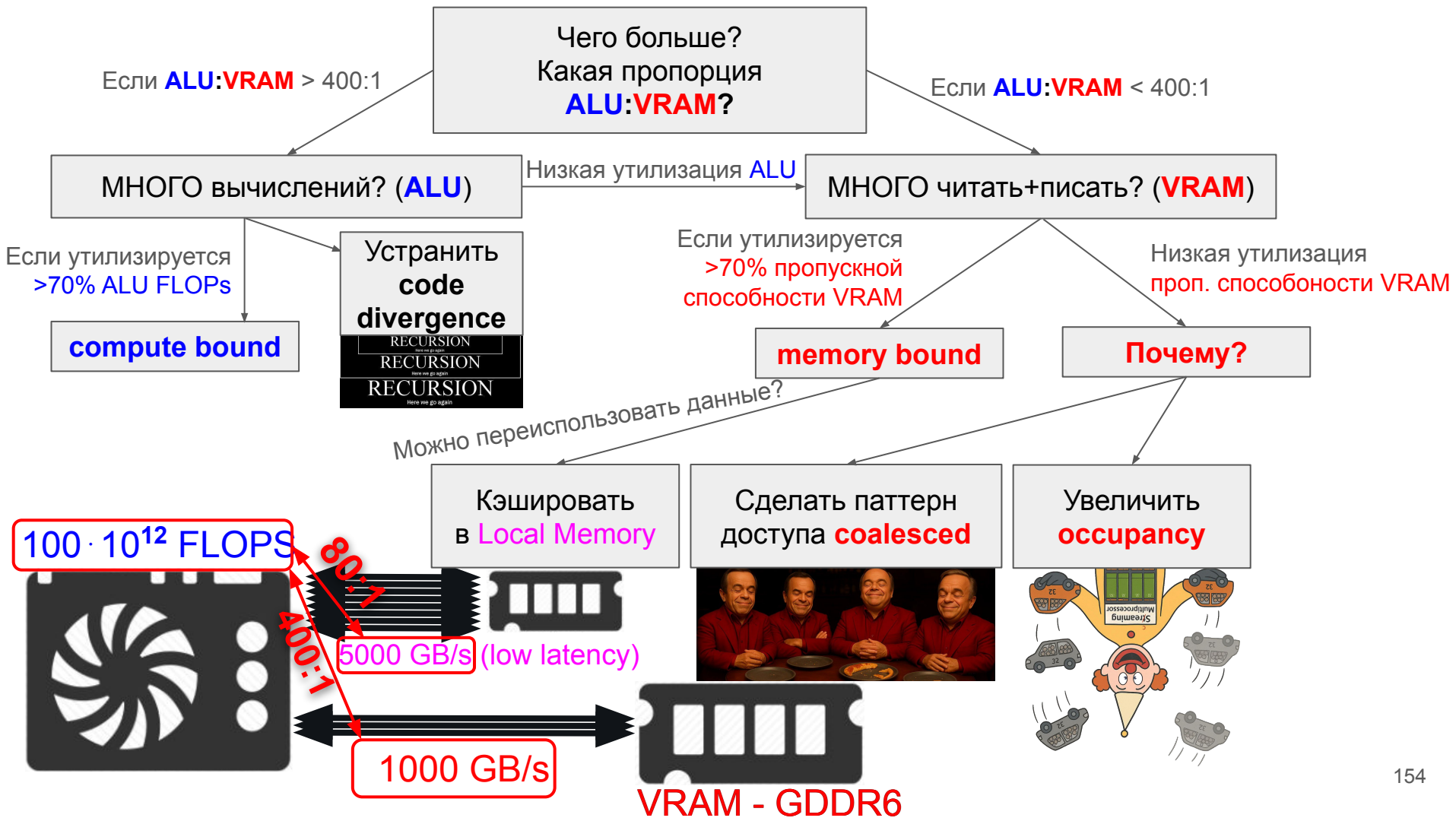


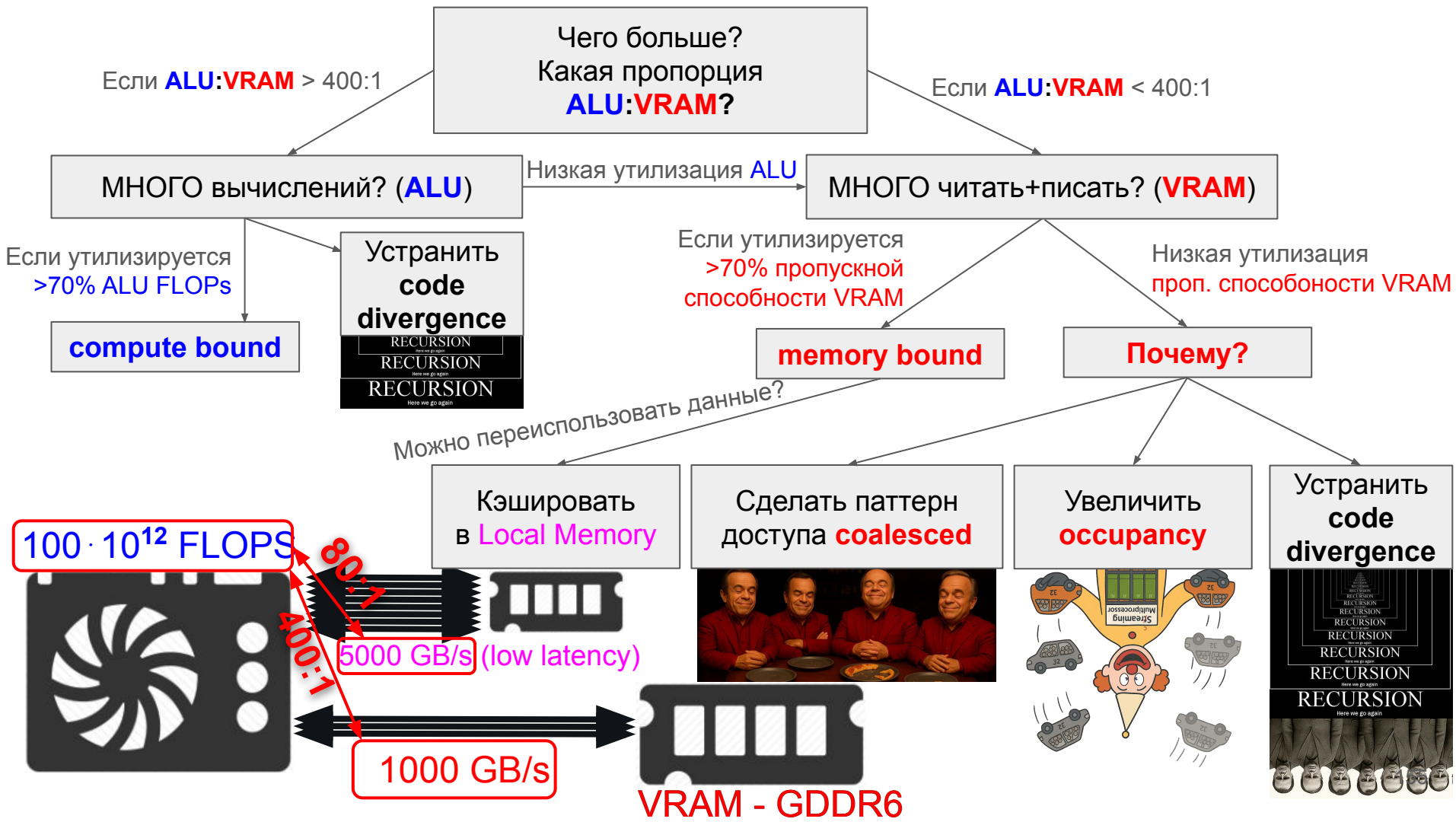


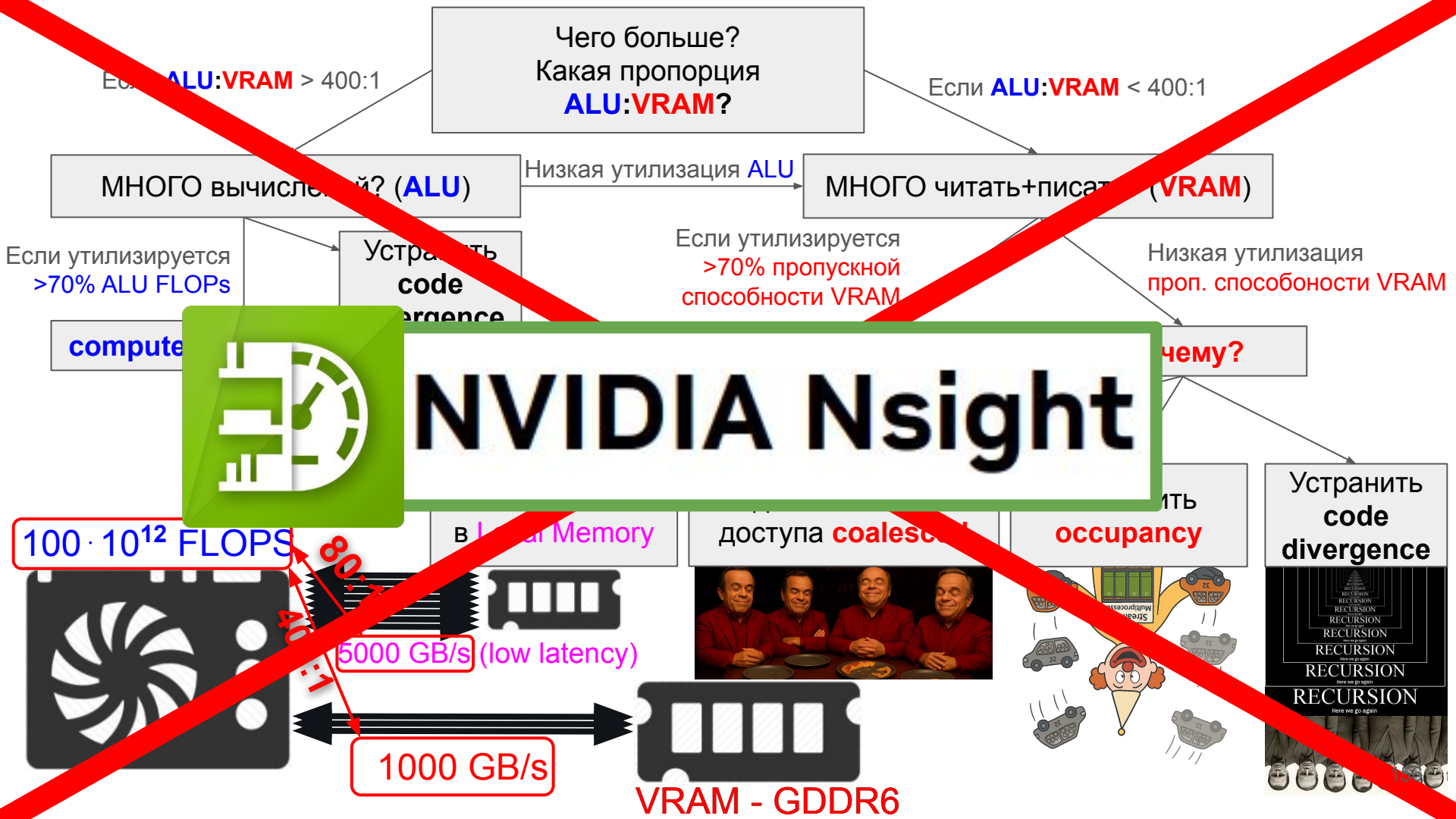






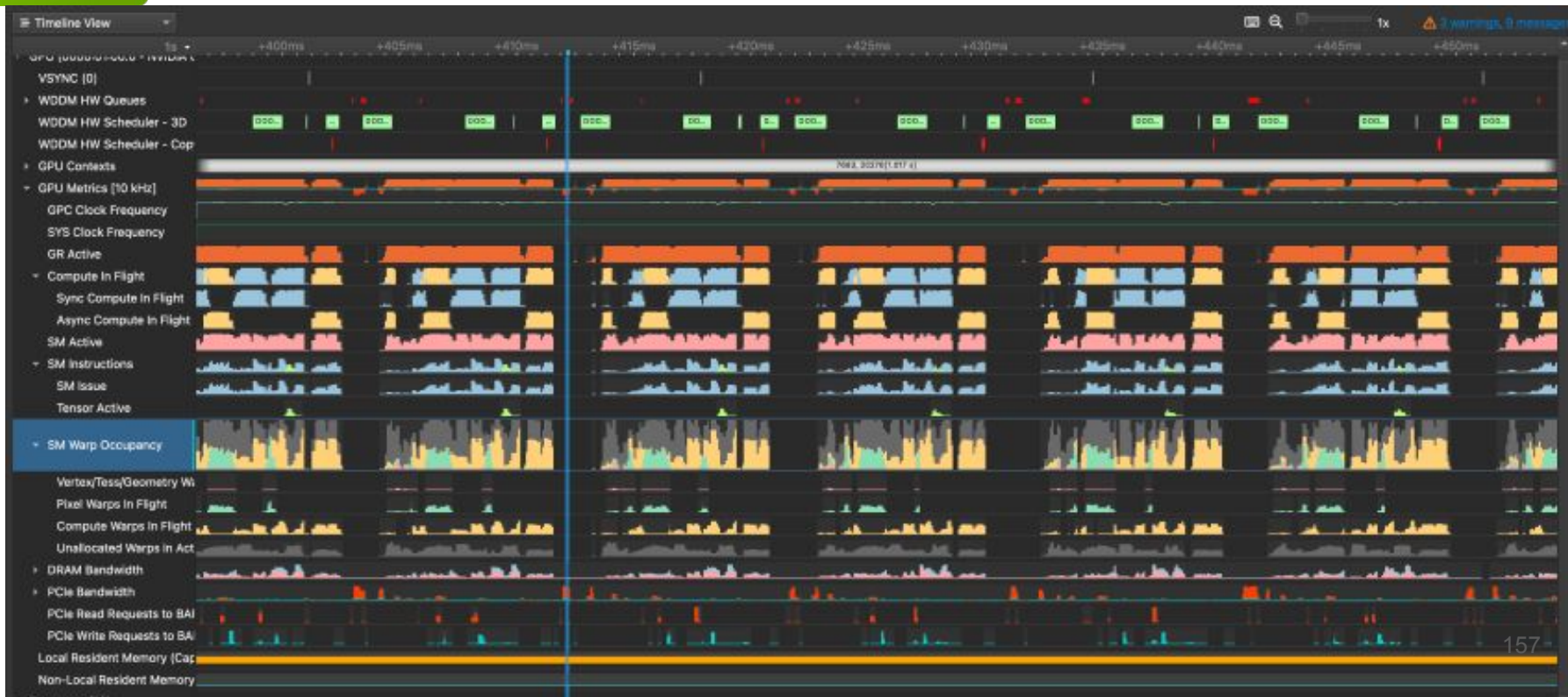








NVIDIA Nsight





Action Debug Profile Tools Window Help

[X] Disconnect [X] Terminate [X] Profile Kernel [X] Baselines [X] Metric Details

addConstDouble3.ncu-rep X addConstDouble3.ncu-rep X Untitled 1 X

Force [X] [Y] [Z]

Result	Time	GPU	SM Frequency	CC
double3 539 - addConstDouble3 (4096, 1, 1)x(256, 1, 1)	101.73 usecond	0 - NVIDIA TITAN V	1.17 cycle/nsecond	7.0

Source: uncoalescedGlobalAccesses.cu [X] [Y] [Z] Find...

Instructions Executed [X] [Y] [Z] [X] [Y] [Z]

Source

32 * global memory and generates an output array of double3 in global

33 */

34

35 #include <stdio.h>

36 #include <cuda_runtime_api.h>

37

38 #define BLOCK_SIZE 256

39

40 #define RUNTIME_API_CALL(apiFuncCall)

41 do {

42 cudaError_t _status = apiFuncCall;

43 if (_status != cudaSuccess) {

44 fprintf(stderr, "%s:%d: error: function %s failed with error

45 __FILE__, __LINE__, #apiFuncCall, cudaGetErrorStrr

46 exit(EXIT_FAILURE);

47 }

48 } while (0)

49

50 __global__ void addConstDouble3(int numElements, double3 *d_in, double3 *d_out)

51 {

52 int index = blockIdx.x * blockDim.x + threadIdx.x;

53 if (index < numElements)

54 {

55 double3 a = d_in[index];

56 a.x += k;

57 a.y += k;

58 a.z += k;

59 d_out[index] = a;

60 }

61 }

62 }

63 int main (int argc, char *argv[])

64 {

65 // Error code to check return values for CUDA calls

66 cudaError_t err = cudaSuccess;

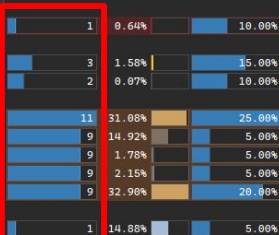
67 double constK = 10.0;

68

69 int kernelOption = 0;

70 if (argc > 1)

71 {

Live arp Stall Sampling
Registers (All Samples)Instructions
Executed

Report	Result	Time	GPU	SM Frequency	CC
addConstDouble	539 - addConstDouble (12288, 1, 1)x(256, 1, 1)	84.35 usecond	0 - NVIDIA TITAN V	1.16 cycle/nsecond	7.0

Source: uncoalescedGlobalAccesses.cu [X] [Y] [Z] Find...

Navigation: Instructions Executed [X] [Y] [Z] [X] [Y] [Z]

Source

32 * global memory and generates an output array of double3 in global

33 */

34

35 #include <stdio.h>

36 #include <cuda_runtime_api.h>

37

38 #define BLOCK_SIZE 256

39

40 #define RUNTIME_API_CALL(apiFuncCall)

41 do {

42 cudaError_t _status = apiFuncCall;

43 if (_status != cudaSuccess) {

44 fprintf(stderr, "%s:%d: error: function %s failed with error

45 __FILE__, __LINE__, #apiFuncCall, cudaGetErrorStrr

46 exit(EXIT_FAILURE);

47 }

48 } while (0)

49

50 __global__ void addConstDouble(int numElements, double *d_in, double *d_out)

51 {

52 int index = blockIdx.x * blockDim.x + threadIdx.x;

53 if (index < numElements)

54 {

55 d_out[index] = d_in[index] + k;

56 }

57 }

58

59 int main (int argc, char *argv[])

60 {

61 // Error code to check return values for CUDA calls

62 cudaError_t err = cudaSuccess;

63 double constK = 10.0;

64

65 int kernelOption = 0;

66 if (argc > 1)

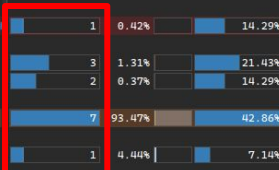
67 {

68 kernelOption = atoi(argv[1]);

69 }

70

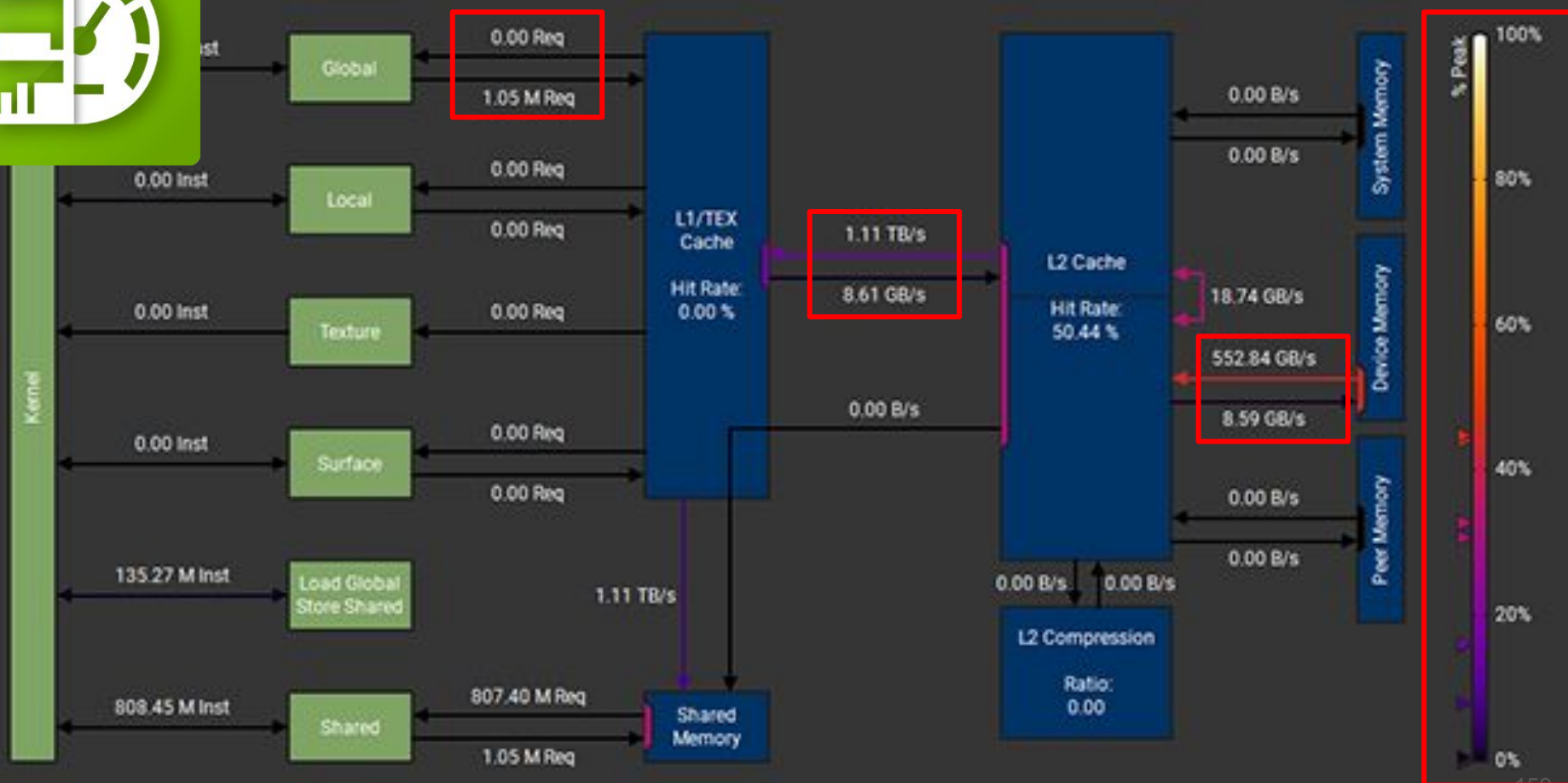
71 int numElements = 1024 * 1024;

Live arp Stall Sampling
Registers (All Samples)Instructions
Executed



Memory Chart

Show As: Throughput





2 - 566 - mergeRanksAndIndicesKernel



Clear Baselines

Apply Rules

Occupancy Calculator

Source Comparison

Copy as Image

	Time	Cycles	GPU	SM Frequency	Process	Attributes
mergeRanksAndIndicesKernel (64, 1, 1)x(256, 1, 1)	4.19 us	4,326	0 - NVIDIA RTX A4500	1.03 cycle/ns	[762349] CuMergeSort	
mergeSortSharedKernel (4096, 1, 1)x(512, 1, 1)	555.10 us	5,79,537	0 - NVIDIA RTX A4500	1.04 cycle/ns	[762349] CuMergeSort	
generateSampleRanksKernel (64, 1, 1)x(256, 1, 1)	26.72 us	27,984	0 - NVIDIA RTX A4500	1.05 cycle/ns	[762349] CuMergeSort	



GPU Speed Of Light Throughput

GPU Throughput Chart



High-level overview of the throughput for compute and memory resources of the GPU. For each unit, the throughput reports the achieved percentage of utilization with respect to the theoretical maximum. Breakdowns show the throughput for each individual sub-metric of Compute and Memory to clearly identify the highest contributor. High-level overview of the utilization for compute and memory resources of the GPU presented as a roofline chart.

Compute (SM) Throughput [%]	10.84 (-75.36%, z=-0.61)	Duration [us]	4.19 (-98.56%, z=-0.75)
Memory Throughput [%]	7.88 (-87.90%, z=-1.23)	Elapsed Cycles [cycle]	4,326 (-98.58%, z=-0.75)
L1/TEX Cache Throughput [%]	10.84 (-79.01%, z=-0.82)	SM Active Cycles [cycle]	2,192.52 (-99.26%, z=-0.74)
L2 Cache Throughput [%]	7.88 (-51.10%, z=-0.75)	SM Frequency [cycle/ns]	1.03 (-1.32%, z=-1.38)
DRAM Throughput [%]	5.53 (-83.26%, z=-1.08)	DRAM Frequency [cycle/ns]	7.55 (-0.53%, z=-1.39)

GPU Throughput



PM Sampling



Timeline view of PM metrics sampled periodically over the workload duration. Data is collected across multiple passes. Use this section to understand how workload behavior changes over its runtime.

Baselines

Difference Bars Green/Red

	Name	Launch	Report	Time	Cycles	Registers	GPU	SM Frequency	CC	Process
✓	mergeSort	560 - mergeSortSharedKernel (4096, ...	/home/prabhanshuc/Documents/me...	555.10 us	5,79,537	17	0 - NVIDIA RTX ...	1.04 cycle/ns	8.6	[762349] CuMergeSort
✓	generateSample	563 - generateSampleRanksKernel (6...	/home/prabhanshuc/Documents/me...	26.72 us	27,984	18	0 - NVIDIA RTX ...	1.05 cycle/ns	8.6	[762349] CuMergeSort
✓	mergeRanks	566 - mergeRanksAndIndicesKernel (...	/home/prabhanshuc/Documents/me...	4.19 us	4,326	16	0 - NVIDIA RTX ...	1.03 cycle/ns	8.6	[762349] CuMergeSort

Глава 6: Как выглядит код

OpenMP, OpenCL, CUDA, Vulkan (GLSL)

Пример: A + B

A

10	34	12	34	54	113	...	1
----	----	----	----	----	-----	-----	---

+

B

32	12	57	12	14	126	...	5
----	----	----	----	----	-----	-----	---

||

C

42	46	69	46	68	239	...	6
----	----	----	----	----	-----	-----	---

Пример: A + B

A

10	34	12	34	54	113	...	1
----	----	----	----	----	-----	-----	---

+

B

32	12	57	12	14	126	...	5
----	----	----	----	----	-----	-----	---

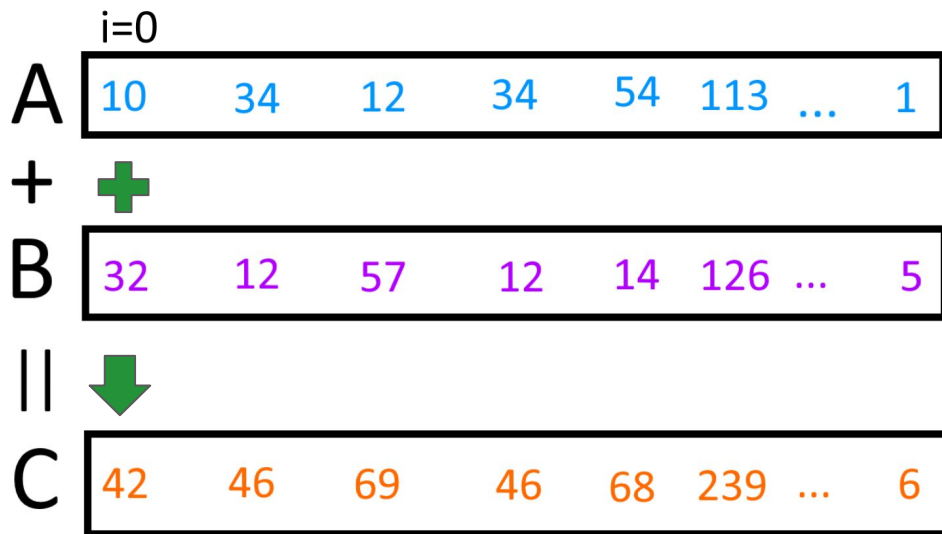
||

C

42	46	69	46	68	239	...	6
----	----	----	----	----	-----	-----	---

```
void solveCPU(int[] a, int[] b, int c[], int n) {  
    for (int i = 0; i < n; ++i) {  
        int sum = a[i] + b[i];  
        c[i] = sum;  
    }  
}
```

Пример: A + B



```
void solveCPU(int[] a, int[] b, int c[], int n) {  
    for (int i = 0; i < n; ++i) {  
        int sum = a[i] + b[i];  
        c[i] = sum;  
    }  
}
```

Пример: A + B

	i=0	i=1						
A	10	34	12	34	54	113	...	1
+	+	+						
B	32	12	57	12	14	126	...	5
	↓	↓						
C	42	46	69	46	68	239	...	6

```
void solveCPU(int[] a, int[] b, int c[], int n) {  
    for (int i = 0; i < n; ++i) {  
        int sum = a[i] + b[i];  
        c[i] = sum;  
    }  
}
```

Пример: A + B

	i=0	i=1	i=2				
A	10	34	12	34	54	113 ...	1
+	+	+	+				
B	32	12	57	12	14	126 ...	5
	↓	↓	↓				
C	42	46	69	46	68	239 ...	6

```
void solveCPU(int[] a, int[] b, int c[], int n) {  
    for (int i = 0; i < n; ++i) {  
        int sum = a[i] + b[i];  
        c[i] = sum;  
    }  
}
```

Пример: A + B

	i=0	i=1	i=2	i=3	i=4	i=5		i=N-1
A	10	34	12	34	54	113	...	1
+	+	+	+	+	+	+		+
B	32	12	57	12	14	126	...	5
	↓	↓	↓	↓	↓	↓		↓
C	42	46	69	46	68	239	...	6

```
void solveCPU(int[] a, int[] b, int c[], int n) {  
    for (int i = 0; i < n; ++i) {  
        int sum = a[i] + b[i];  
        c[i] = sum;  
    }  
}
```

**Какая
асимптотика?**

Пример: A + B

	i=0	i=1	i=2	i=3	i=4	i=5		i=N-1
A	10	34	12	34	54	113	...	1
+	+	+	+	+	+	+		+
B	32	12	57	12	14	126	...	5
	↓	↓	↓	↓	↓	↓		↓
C	42	46	69	46	68	239	...	6

```
void solveCPU(int[] a, int[] b, int c[], int n) {  
    for (int i = 0; i < n; ++i) {  
        int sum = a[i] + b[i];  
        c[i] = sum;  
    }  
}
```

O(N)

Пример: A + B

A	10	34	12	34	54	113	...	1
+								
B	32	12	57	12	14	126	...	5
C	42	46	69	46	68	239	...	6



NVIDIA RTX 4090 - 16 тысяч ядер!

Пример: A + B

A	10	34	12	34	54	113	...	1
+								
B	32	12	57	12	14	126	...	5
C	42	46	69	46	68	239	...	6



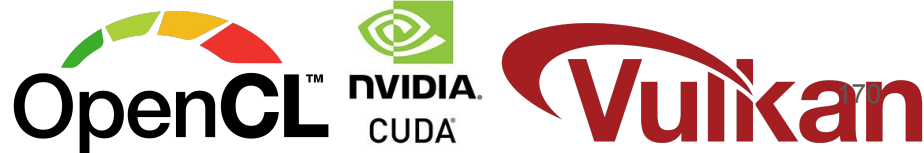
```
void solveCPU(int[] a, int[] b, int c[], int n) {  
    for (int i = 0; i < n; ++i) {  
        int sum = a[i] + b[i];  
        c[i] = sum;  
    }  
}
```

O(N)

```
void solveGPU(int[] a, int[] b, int c[], int n) {  
    const int i = get_global_id(0);  
    c[i] = b[i] + a[i];  
}
```

Супер многопоточно!

NVIDIA RTX 4090 - 16 тысяч ядер!



Пример: A + B

A	10	34	12	34	54	113	...	1
+								
B	32	12	57	12	14	126	...	5
C	42	46	69	46	68	239	...	6



```
void solveCPU(int[] a, int[] b, int c[], int n) {  
    for (int i = 0; i < n; ++i) {  
        int sum = a[i] + b[i];  
        c[i] = sum;  
    }  
}
```

O(N)

```
void solveGPU(int[] a, int[] b, int c[], int n) {  
    const int i = get_global_id(0);  
    c[i] = b[i] + a[i];  
}
```

Какая асимптотика?

NVIDIA RTX 4090 - 16 тысяч ядер!



Пример: A + B

	i=0	i=1	i=2	i=3	i=4	i=5	i=N-1
A	10	34	12	34	54	113 ...	1
+	+	+	+	+	+	+	+
B	32	12	57	12	14	126 ...	5
	↓	↓	↓	↓	↓	↓	↓
C	42	46	69	46	68	239 ...	6

```
void solveCPU(int[] a, int[] b, int c[], int n) {  
    for (int i = 0; i < n; ++i) {  
        int sum = a[i] + b[i];  
        c[i] = sum;  
    }  
}
```

O(N)

```
void solveGPU(int[] a, int[] b, int c[], int n) {  
    const int i = get_global_id(0);  
    c[i] = b[i] + a[i];  
}
```

O(N / 16384)

NVIDIA RTX 4090 - 16 тысяч ядер!



OpenCL™

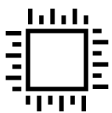
**NVIDIA
CUDA®**

Vulkan

Пример: A + B (N=100.000.000)

```
for (size_t i = 0; i < n; ++i) {  
    cs[i] = as[i] + bs[i];  
}
```

Пример: A + B (N=100.000.000)



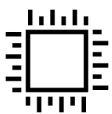
CPU - Intel 13700K

```
for (size_t i = 0; i < n; ++i) {  
    cs[i] = as[i] + bs[i];  
}
```

a + b median time: 0.068 sec (+-0.00481664)

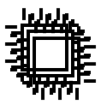
a + b median RAM bandwidth: 16.4351 GB/s

Пример: A + B (N=100.000.000)



CPU - Intel 13700K

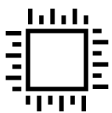
```
for (size_t i = 0; i < n; ++i) {  
    cs[i] = as[i] + bs[i];  
}  
  
a + b median time: 0.068 sec (+-0.00481664)  
a + b median RAM bandwidth: 16.4351 GB/s
```



CPU - Intel 13700K

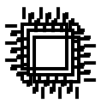
```
#pragma omp parallel for  
for (ptrdiff_t i = 0; i < n; ++i) {  
    cs[i] = as[i] + bs[i];  
}  
  
a + b median time: 0.047 sec (+-0.00153623)  
a + b median RAM bandwidth: 23.7784 GB/s
```

Пример: A + B (N=100.000.000)



CPU - Intel 13700K

```
for (size_t i = 0; i < n; ++i) {  
    cs[i] = as[i] + bs[i];  
}  
  
a + b median time: 0.068 sec (+-0.00481664)  
a + b median RAM bandwidth: 16.4351 GB/s
```



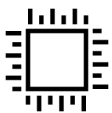
CPU - Intel 13700K

```
#pragma omp parallel for  
for (ptrdiff_t i = 0; i < n; ++i) {  
    cs[i] = as[i] + bs[i];  
}  
  
a + b median time: 0.047 sec (+-0.00153623)  
a + b median RAM bandwidth: 23.7784 GB/s
```

**Почему разница такая
незначительная?**

**Как проверить гипотезу?
Как спровоцировать
значительную разницу?**

Пример: **100*cos**(A + B) (N=100.000.000)

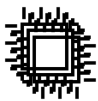


CPU - Intel 13700K

```
for (size_t i = 0; i < n; ++i) {  
    cs[i] = 100.0*cos( Left: as[i] + bs[i]);  
}
```

median time: 1.178 sec (+-0.0235009)

median RAM bandwidth: 0.948716 GB/s



CPU - Intel 13700K

```
#pragma omp parallel for
```

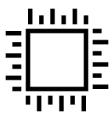
```
for (ptrdiff_t i = 0; i < n; ++i) {  
    cs[i] = 100*cos( Left: as[i] + bs[i]);  
}
```

median time: 0.109 sec (+-0.0121988)

median RAM bandwidth: 10.2531 GB/s

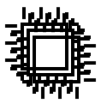
x10.9

Пример: A + B (N=100.000.000)



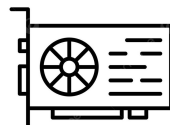
CPU - Intel 13700K

```
for (size_t i = 0; i < n; ++i) {  
    cs[i] = as[i] + bs[i];  
}  
a + b median time: 0.068 sec (+-0.00481664)  
a + b median RAM bandwidth: 16.4351 GB/s
```



CPU - Intel 13700K

```
#pragma omp parallel for  
for (ptrdiff_t i = 0; i < n; ++i) {  
    cs[i] = as[i] + bs[i];  
}  
a + b median time: 0.047 sec (+-0.00153623)  
a + b median RAM bandwidth: 23.7784 GB/s
```

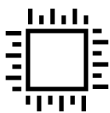


GPU - NVIDIA RTX 4090

```
__kernel void aplusb(__global const float* a,  
                    __global const float* b,  
                    __global float* c,  
                    unsigned int n)  
{  
    const unsigned int index = get_global_id( dimindx: 0);  
    if (index >= n)  
        return;  
  
    c[index] = a[index] + b[index];  
}
```

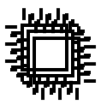
```
a + b median time: 0.002 sec (+-0.000632456)  
a + b median VRAM bandwidth: 558.794 GB/s
```

Пример: A + B (N=100.000.000)



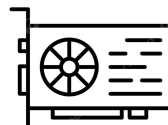
CPU - Intel 13700K

```
for (size_t i = 0; i < n; ++i) {  
    cs[i] = as[i] + bs[i];  
}  
  
a + b median time: 0.068 sec (+-0.00481664)  
a + b median RAM bandwidth: 16.4351 GB/s
```



CPU - Intel 13700K

```
#pragma omp parallel for  
for (ptrdiff_t i = 0; i < n; ++i) {  
    cs[i] = as[i] + bs[i];  
}  
  
a + b median time: 0.047 sec (+-0.00153623)  
a + b median RAM bandwidth: 23.7784 GB/s
```

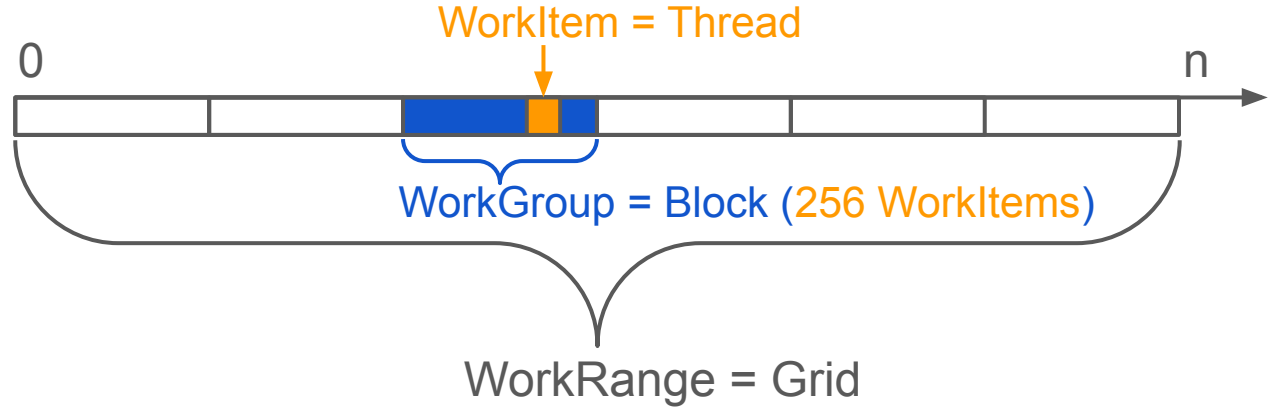


GPU - NVIDIA RTX 4090

```
__kernel void aplusb(__global const float* a,  
                    __global const float* b,  
                    __global float* c,  
                    unsigned int n)  
{  
  
    const unsigned int index = get_global_id( dimindx: 0 );  
    if (index >= n)  
        return;  
  
    c[index] = a[index] + b[index];  
}  
  
a + b median time: 0.002 sec (+-0.000632456)  
a + b median VRAM bandwidth: 558.794 GB/s
```

x23.5

Пример: $A + B$ ($N=100.000.000$)



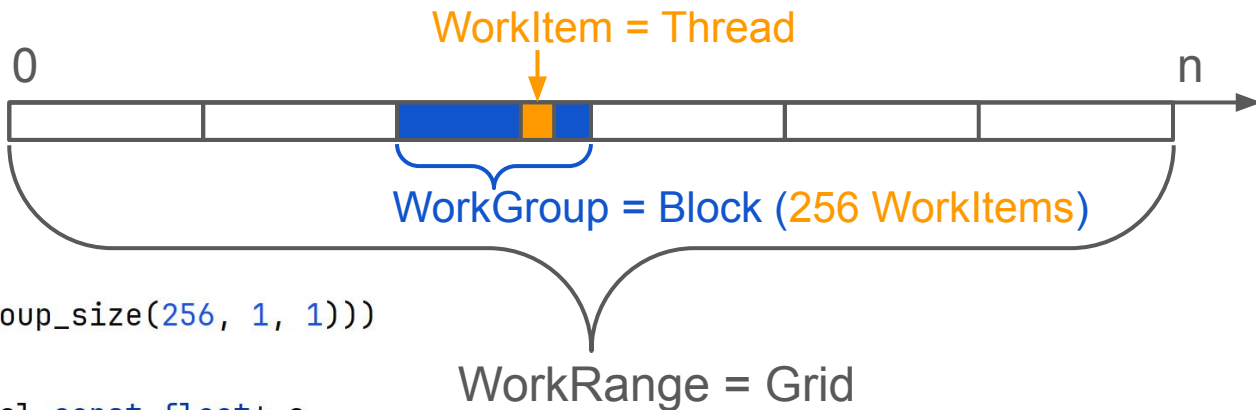
Пример: A + B (N=100.000.000)

```
1 #ifdef __CLION_IDE__
```

```
2 #include <libgpu/opencl/cl/clion_defines.cl>
```

```
3 #endif
```

```
4  
5 #line 5
```



```
10 __attribute__((reqd_work_group_size(256, 1, 1)))
```

```
12 __kernel void aplusb(__global const float* a,  
13                     __global const float* b,  
14                     __global float* c,  
15                     unsigned int n)
```

```
16 {
```

```
17     const unsigned int index = get_global_id(dimindx: 0);
```

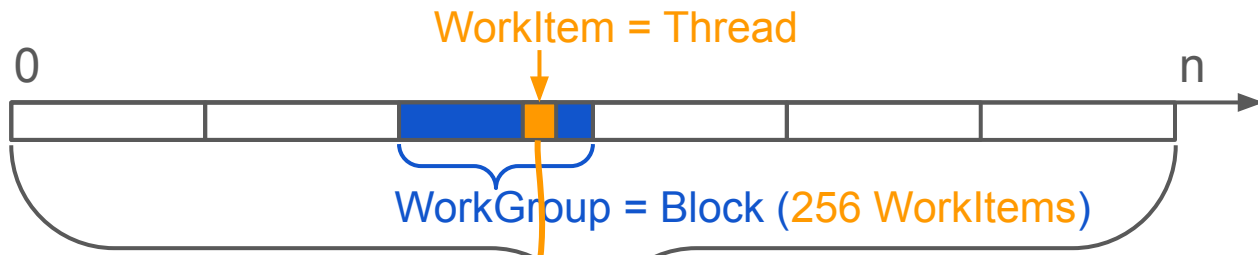
```
18     if (index >= n)
```

```
19         return;
```

```
20  
21     c[index] = a[index] + b[index];
```

```
22 }
```

Пример: A + B (N=100.000.000)



#line 5

```
__attribute__((reqd_work_group_size(256, 1, 1)))
```

```
__kernel void aplusb(__global const float* a,  
                    __global const float* b,  
                    __global float* c,  
                    unsigned int n)
```

```
{
```

```
    const unsigned int index = get_global_id(dimindx: 0);
```

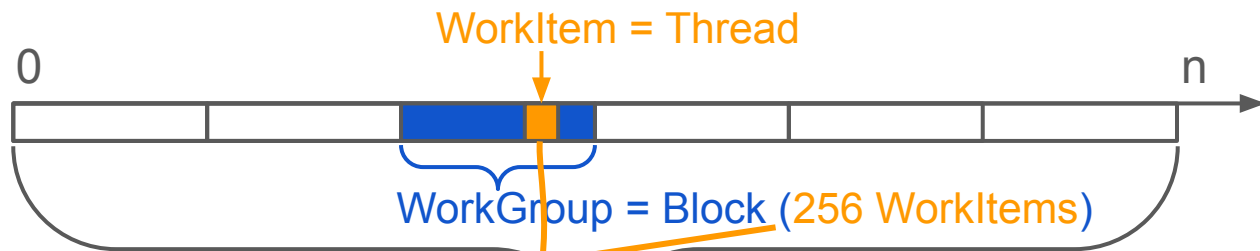
```
    if (index >= n)
```

```
        return;
```

```
    c[index] = a[index] + b[index];
```

```
}
```

Пример: A + B (N=100.000.000)



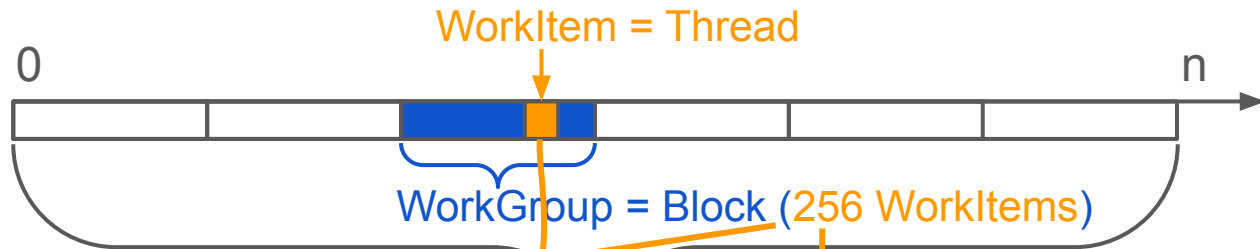
#line 5

```
__attribute__((reqd_work_group_size(256, 1, 1)))
```

```
__kernel void aplusb(__global const float* a,
                    __global const float* b,
                    __global float* c,
                    unsigned int n)
{
    const unsigned int index = get_global_id( dimindx: 0);
    if (index >= n)
        return;

    c[index] = a[index] + b[index];
}
```

Пример: A + B (N=100.000.000)



#line 5

```
__attribute__((reqd_work_group_size(256, 1, 1)))
```

```
__kernel void aplusb(__global const float* a,  
                    __global const float* b,  
                    __global float* c,  
                    unsigned int n)
```

```
{
```

```
    const unsigned int index = get_global_id( dimindx: 0);
```

```
    if (index >= n)
```

```
        return;
```

```
    c[index] = a[index] + b[index];
```

```
}
```

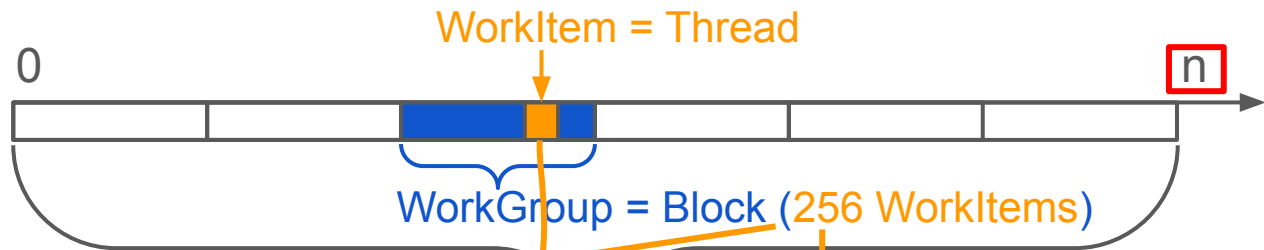
WorkRange = Grid

HOST-side (CPU, C++)

```
gpu::WorkSize workSize(GROUP_SIZE, n);
```

```
ocl_aplusb.exec(workSize, a_gpu, b_gpu, c_gpu, n);
```

Пример: A + B (N=100.000.000)



А где задается
размер рабочего
пространства?

WorkRange = Grid

```
1 #ifndef __CLION_IDE__
2 #include <libgpu/opencl/cl/clion_defines.cl>
3 #endif
```

#line 5

```
10 __attribute__((reqd_work_group_size(256, 1, 1)))
```

```
12 __kernel void aplusb(__global const float* a,
13                     __global const float* b,
14                     __global float* c,
15                     unsigned int n)
```

```
16 {
```

```
17     const unsigned int index = get_global_id( dimindx: 0);
```

```
18     if (index >= n)
```

```
19         return;
```

```
21     c[index] = a[index] + b[index];
```

```
22 }
```

HOST-side (CPU, C++)

```
gpu::WorkSize workSize(GROUP_SIZE, n);
```

```
ocl_aplusb.exec(workSize, a_gpu, b_gpu, c_gpu, n);
```

Пример: A + B (N=100.000.000)



```
1 #ifndef __CLION_IDE__
2 #include <libgpu/opencl/cl/clion_defines.cl>
3 #endif
```

#line 5

```
10 __attribute__((reqd_work_group_size(256, 1, 1)))
```

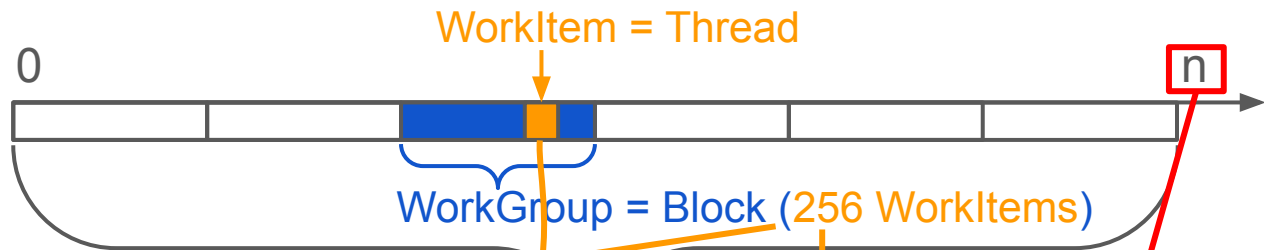
```
12 __kernel void aplusb(__global const float* a,
13                     __global const float* b,
14                     __global float* c,
15                     unsigned int n)
```

```
16 {
17     const unsigned int index = get_global_id( dimindx: 0);
```

```
18     if (index >= n)
19         return;
```

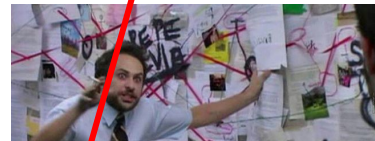
```
21     c[index] = a[index] + b[index];
```

```
22 }
```



WorkRange = Grid

А где задается
размер рабочего
пространства?



HOST-side (CPU, C++)

```
gpu::WorkSize workSize(GROUP_SIZE, n);
```

```
ocl_aplusb.exec(workSize, a_gpu, b_gpu, c_gpu, n);
```

Пример: A + B (N=100.000.000)



```
1 #ifndef __CLION_IDE__
2 #include <libgpu/opencl/cl/clion_defines.cl>
3 #endif
```

#line 5

```
10 __attribute__((reqd_work_group_size(256, 1, 1)))
```

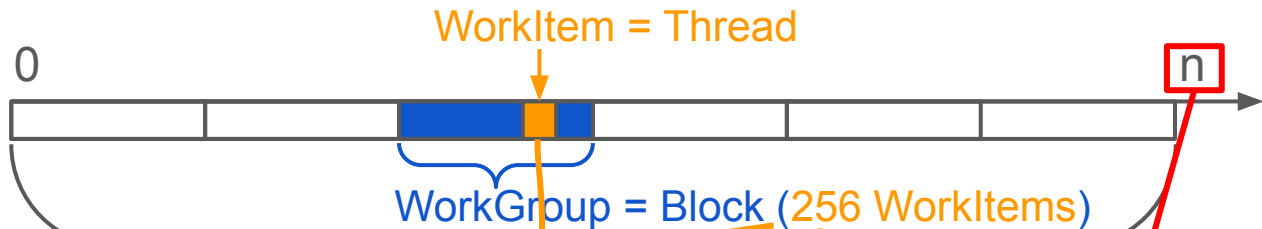
```
12 __kernel void aplusb(__global const float* a,
13                     __global const float* b,
14                     __global float* c,
15                     unsigned int n)
```

```
16 {
17     const unsigned int index = get_global_id( dimindx: 0);
```

```
18     if (index >= n)
19         return;
```

```
21     c[index] = a[index] + b[index];
```

```
22 }
```



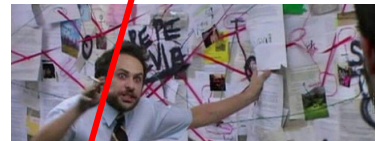
WorkRange = Grid

А где задается
размер рабочего
пространства?

HOST-side (CPU, C++)

```
gpu::WorkSize workSize(GROUP_SIZE, n);
```

```
ocl_aplusb.exec(workSize, a_gpu, b_gpu, c_gpu, n);
```



Пример: A + B (N=100.000.000)



```
1 #ifdef __CLION_IDE__
2 #include <libgpu/opencl/cl/clion_defines.cl>
3 #endif
```

#line 5

```
10 __attribute__((reqd_work_group_size(256, 1, 1)))
```

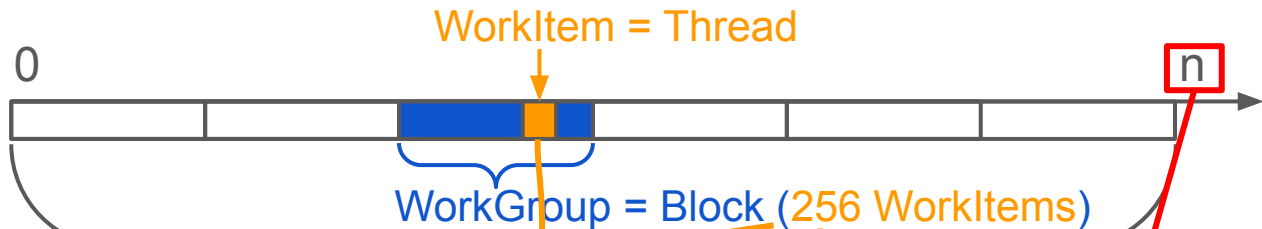
```
12 __kernel void aplusb(__global const float* a,
13                     __global const float* b,
14                     __global float* c,
15                     unsigned int n)
```

```
16 {
17     const unsigned int index = get_global_id( dimindx: 0);
```

```
18     if (index >= n)
19         return;
```

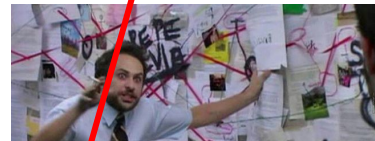
```
21     c[index] = a[index] + b[index];
```

```
22 }
```



WorkRange = Grid

А где задается
размер рабочего
пространства?



HOST-side (CPU, C++)

```
gpu::WorkSize workSize(GROUP_SIZE, n);
```

```
ocl_aplusb.exec(workSize, a_gpu, b_gpu, c_gpu, n);
```

Пример: A + B (N=100.000.000)

```
1 #ifndef __CLION_IDE__
```

```
2 #include <libgpu/opencl/cl/clion_defines.cl>
```

```
3 #endif
```

```
4  
5 #line 5
```



```
6  
7  
8  
9  
10 __attribute__((reqd_work_group_size(256, 1, 1)))
```

```
11  
12 __kernel void aplusb(__global const float* a,  
13                     __global const float* b,  
14                     __global float* c,  
15                     unsigned int n)
```

```
16 {
```

```
17     const unsigned int index = get_global_id(dimindx: 0);
```

```
18     if (index >= n)
```

```
19         return;
```

```
20  
21     c[index] = a[index] + b[index];
```

```
22 }
```



```
blockIdx.x * blockDim.x + threadIdx.x;
```

```
1 __global__ void aplusb(const float* a,  
2                       const float* b,  
3                       float* c,  
4                       unsigned int n)
```

```
5 {
```

```
6     const unsigned int index = blockIdx.x * blockDim.x + threadIdx.x;
```

```
7     if (index >= n)
```

```
8         return;
```

```
9  
10     c[index] = a[index] + b[index];
```

```
11 }
```

Пример: A + B (N=100.000.000)



```
1  __global__ void aplusb(const float* a,  
2                          const float* b,  
3                          float* c,  
4                          unsigned int n)  
5  {  
6      const unsigned int index = blockIdx.x * blockDim.x + threadIdx.x;  
7      if (index >= n)          WorkGroup = Block (256 WorkItems)  
8          return;  
9  
10     c[index] = a[index] + b[index];  
11 }
```

HOST-side (CPU, C++)

Пример: $A + B$ ($N=100.000.000$)

```
12 void cuda_aplusb(const gpu::WorkSize &workSize,  
13                 const float* a, const float* b, float* c, unsigned int n,  
14                 cudaStream_t stream)  
15 {  
16     const unsigned int blockSize = 256;  
17     const unsigned int gridSize = (n + blockSize - 1) / blockSize;  
18     aplusb<<<gridSize, blockSize, 0, stream>>>(a, b, c, n);  
19     CUDA_CHECK_KERNEL(stream);  
20 }  
21
```



```
1 __global__ void aplusb(const float* a,  
2                       const float* b,  
3                       float* c,  
4                       unsigned int n)  
5 {  
6     const unsigned int index = blockIdx.x * blockDim.x + threadIdx.x;  
7     if (index >= n)  
8         return;  
9  
10    c[index] = a[index] + b[index];  
11 }
```

WorkGroup = Block (256 WorkItems)

HOST-side (CPU, C++)

Пример: A + B (N=100.000.000)

```
12 void cuda_aplusb(const gpu::WorkSize &workSize,  
13                 const float* a, const float* b, float* c, unsigned int n,  
14                 cudaStream_t stream)  
15 {  
16     const unsigned int blockSize = 256;  
17     const unsigned int gridSize = (n + blockSize - 1) / blockSize;  
18     aplusb<<<gridSize, blockSize, 0, stream>>>(a, b, c, n);  
19     CUDA_CHECK_KERNEL(stream);  
20 }  
21
```



А где задается
размер рабочего
пространства?

```
1 __global__ void aplusb(const float* a,  
2                       const float* b,  
3                       float* c,  
4                       unsigned int n)  
5 {  
6     const unsigned int index = blockIdx.x * blockDim.x + threadIdx.x;  
7     if (index >= n) WorkGroup = Block (256 WorkItems)  
8         return;  
9  
10    c[index] = a[index] + b[index];  
11 }
```

HOST-side (CPU, C++)

Пример: $A + B$ ($N=100.000.000$)

```
12 void cuda_aplusb(const gpu::WorkSize &workSize,  
13                 const float* a, const float* b, float* c, unsigned int n,  
14                 cudaStream_t stream)  
15 {  
16     const unsigned int blockSize = 256;  
17     const unsigned int gridSize = (n + blockSize - 1) / blockSize;  
18     aplusb<<gridSize, blockSize, 0, stream>>>(a, b, c, n);  
19     CUDA_CHECK_KERNEL(stream);  
20 }  
21
```



А где задается
размер рабочего
пространства?

```
1 __global__ void aplusb(const float* a,  
2                       const float* b,  
3                       float* c,  
4                       unsigned int n)  
5 {  
6     const unsigned int index = blockIdx.x * blockDim.x + threadIdx.x;  
7     if (index >= n) WorkGroup = Block (256 WorkItems)  
8         return;  
9  
10    c[index] = a[index] + b[index];  
11 }
```

Пример: A + B (N=100.000.000)

GLSL (Graphics Library Shading Language)



```
#version 450
#include <libgpu/vulkan/vk/common.vk>
```

```
layout (std430, binding = 0) readonly buffer AsIn { uint as[]; };
layout (std430, binding = 1) readonly buffer BsIn { uint bs[]; };
layout (std430, binding = 2) writeonly buffer CsOut { uint cs[]; };
```

```
layout (push_constant) uniform PushConstants {
    uint n;
} params;
```

```
layout (local_size_x = 256) in;
```

```
void main()
```

```
{
    const uint index = gl_GlobalInvocationID.x;
```

```
    if (index >= params.n)
```

```
        return;
```

```
    cs[index] = as[index] + bs[index];
```

```
}
```

```
__kernel void aplusb_matrix(__global const uint* a,  
                             __global const uint* b,  
                             __global      uint* c,  
                             unsigned int width,  
                             unsigned int height)
```

```
{
```

```
    const unsigned int index = get_global_id(0);
```

```
    for (int k = 0; k < width; ++k) {
```

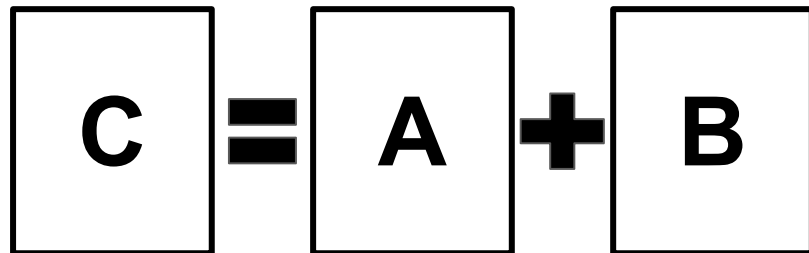
```
        unsigned int col = k;
```

```
        unsigned int row = index;
```

```
        c[row * width + col] = a[row * width + col]  
                                + b[row * width + col];
```

```
    }
```

```
}
```



```
__kernel void aplusb_matrix(__global const uint* a,  
                            __global const uint* b,  
                            __global      uint* c,  
                            unsigned int width,  
                            unsigned int height)
```

```
{
```

```
    const unsigned int index = get_global_id(0);
```

```
    for (int k = 0; k < width; ++k) {
```

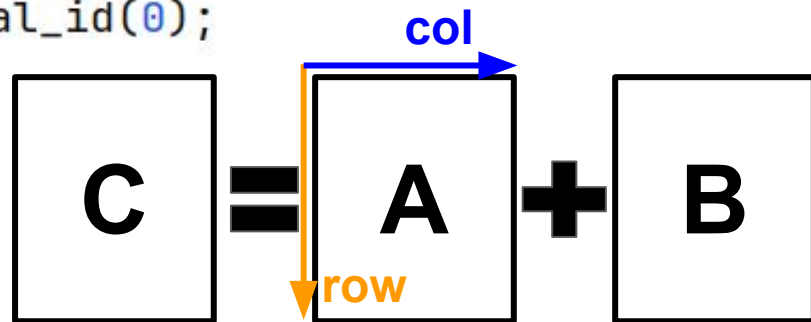
```
        unsigned int col = k;
```

```
        unsigned int row = index;
```

```
        c[row * width + col] = a[row * width + col]  
                                + b[row * width + col];
```

```
    }
```

```
}
```



```
__kernel void aplusb_matrix(__global const uint* a,  
                           __global const uint* b,  
                           __global      uint* c,  
                           unsigned int width,  
                           unsigned int height)
```

```
{
```

```
    const unsigned int index = get_global_id(0);
```

```
    for (int k = 0; k < width; ++k) {
```

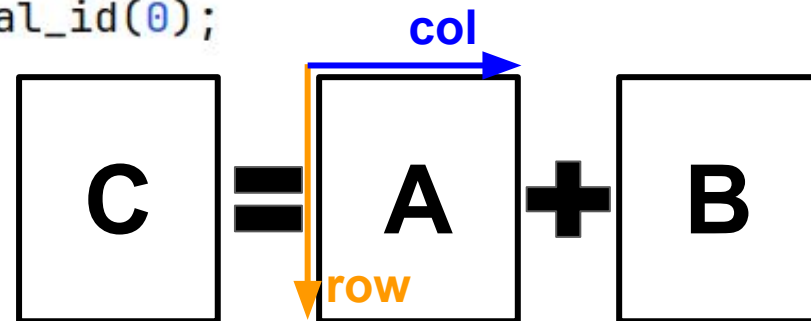
```
        unsigned int col = k;
```

```
        unsigned int row = index;
```

```
        c[row * width + col] = a[row * width + col]  
                               + b[row * width + col];
```

```
    }
```

```
}
```



coalesced memory access?

```
__kernel void aplusb_matrix(__global const uint* a,
                           __global const uint* b,
                           __global      uint* c,
                           unsigned int width,
                           unsigned int height)
```

```
{
```

```
    const unsigned int index = get_global_id(0);
```

```
    for (int k = 0; k < width; ++k) {
```

```
        unsigned int col = k;
```

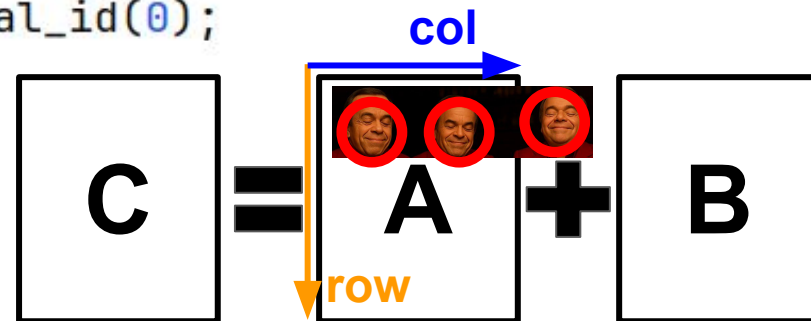
```
        unsigned int row = index;
```

```
        c[row * width + col] = a[row * width + col]
                               + b[row * width + col];
```

```
    }
```

```
}
```

Как выглядит *warp*?



coalesced memory access?

```
__kernel void aplusb_matrix(__global const uint* a,
                           __global const uint* b,
                           __global      uint* c,
                           unsigned int width,
                           unsigned int height)
```

```
{
```

```
    const unsigned int index = get_global_id(0);
```

```
    for (int k = 0; k < width; ++k) {
```

```
        unsigned int col = k;
```

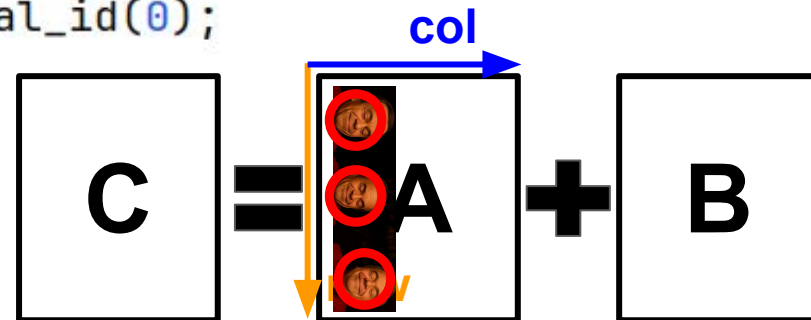
```
        unsigned int row = index;
```

```
        c[row * width + col] = a[row * width + col]
                               + b[row * width + col];
```

```
    }
```

```
}
```

Как выглядит *warp*?



coalesced memory access?

```
__kernel void aplusb_matrix(__global const uint* a,
                           __global const uint* b,
                           __global      uint* c,
                           unsigned int width,
                           unsigned int height)
```

HE coalesced memory access!!!

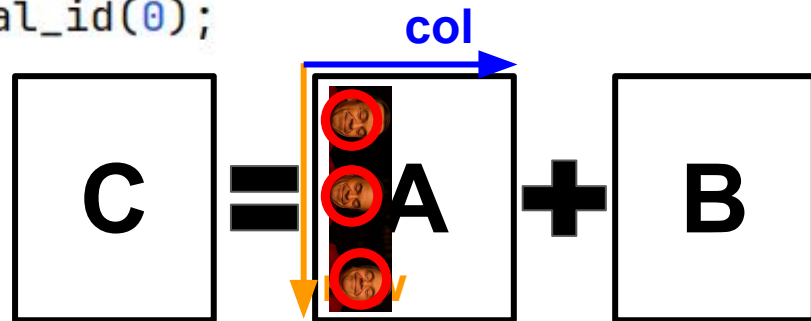
```
{
    const unsigned int index = get_global_id(0);
```

```
    for (int k = 0; k < width; ++k) {
```

```
        unsigned int col = k;
```

```
        unsigned int row = index;
```

```
        c[row * width + col] = a[row * width + col]
                               + b[row * width + col];
```



```
__kernel void aplusb_matrix(__global const uint* a,  
                            __global const uint* b,  
                            __global      uint* c,  
                            unsigned int width,  
                            unsigned int height)
```

```
{
```

```
    const unsigned int index = get_global_id(0);
```

```
    for (int k = 0; k < height; ++k) {
```

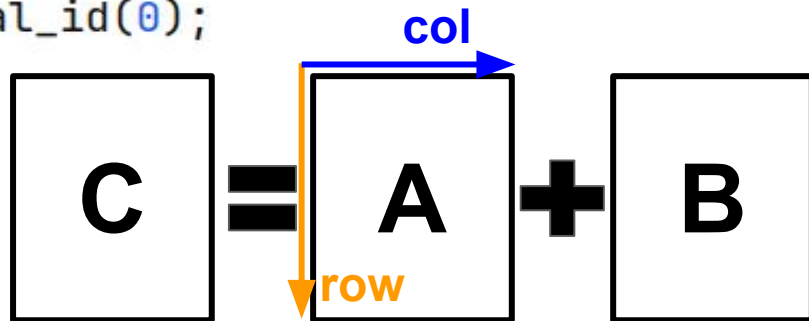
```
        unsigned int col = index;
```

```
        unsigned int row = k;
```

```
        c[row * width + col] = a[row * width + col]  
                                + b[row * width + col];
```

```
    }
```

```
}
```



```
__kernel void aplusb_matrix(__global const uint* a,  
                           __global const uint* b,  
                           __global      uint* c,  
                           unsigned int width,  
                           unsigned int height)
```

```
{
```

```
    const unsigned int index = get_global_id(0);
```

```
    for (int k = 0; k < height; ++k) {
```

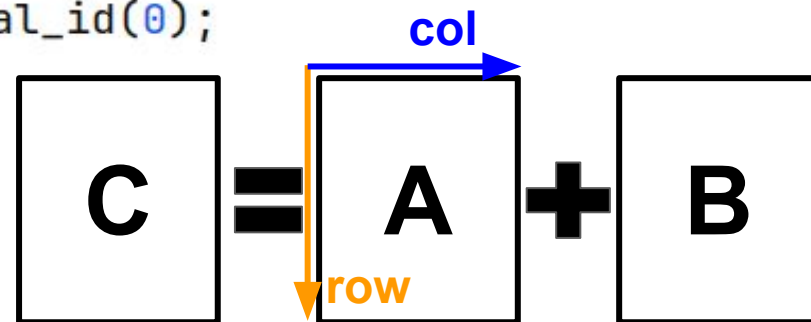
```
        unsigned int col = index;
```

```
        unsigned int row = k;
```

```
        c[row * width + col] = a[row * width + col]  
                               + b[row * width + col];
```

```
    }
```

```
}
```



coalesced memory access?

```
__kernel void aplusb_matrix(__global const uint* a,
                           __global const uint* b,
                           __global      uint* c,
                           unsigned int width,
                           unsigned int height)
```

```
{
```

```
    const unsigned int index = get_global_id(0);
```

```
    for (int k = 0; k < height; ++k) {
```

```
        unsigned int col = index;
```

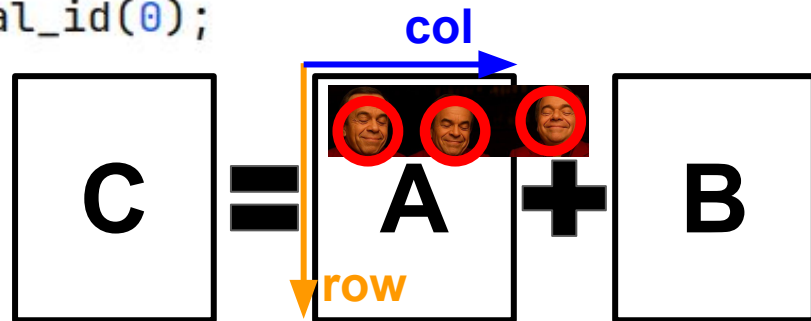
```
        unsigned int row = k;
```

```
        c[row * width + col] = a[row * width + col]
                               + b[row * width + col];
```

```
    }
```

```
}
```

Как выглядит *warp*?



coalesced memory access?

```
__kernel void aplusb_matrix(__global const uint* a,
                           __global const uint* b,
                           __global      uint* c,
                           unsigned int width,
                           unsigned int height)
```

```
{
```

```
    const unsigned int index = get_global_id(0);
```

```
    for (int k = 0; k < height; ++k) {
```

```
        unsigned int col = index;
```

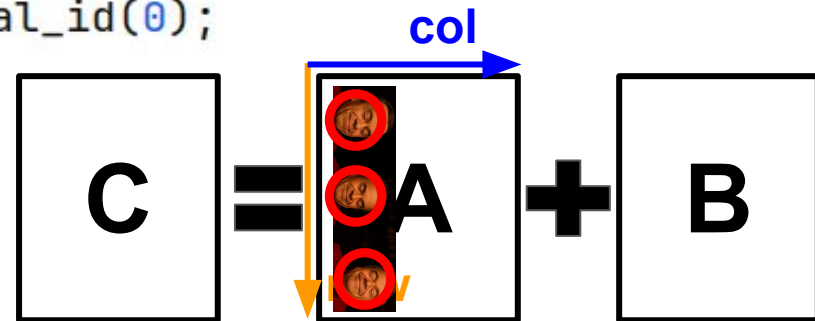
```
        unsigned int row = k;
```

```
        c[row * width + col] = a[row * width + col]
                               + b[row * width + col];
```

```
    }
```

```
}
```

Как выглядит *warp*?



coalesced memory access?

```
__kernel void aplusb_matrix(__global const uint* a,
                           __global const uint* b,
                           __global      uint* c,
                           unsigned int width,
                           unsigned int height)
```

```
{
```

coalesced memory access!!!

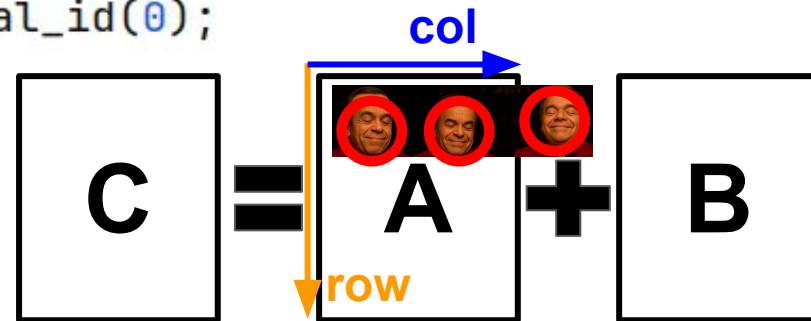
```
const unsigned int index = get_global_id(0);
```

```
for (int k = 0; k < height; ++k) {
```

```
    unsigned int col = index;
```

```
    unsigned int row = k;
```

```
    c[row * width + col] = a[row * width + col]
                          + b[row * width + col];
```



```
}
```



CS Space

Клуб технологий и науки



Vulkan



nVIDIA
CUDA

OpenCL™

Вопросы?



[@UnicornGlade](#)



[@PolarNick239](#)



polarnick239@gmail.com



Николай Полярный

Activate Ubuntu

Go to Settings to activate Ubuntu.