



UNIVERSITÀ
DEGLI STUDI
FIRENZE

GPU programming basics

Prof. Marco Bertini



CUDA: atomic operations, privatization, algorithms

Atomic operations

- The basic atomic operation in hardware is something like a read-modify-write operation performed by a single hardware instruction on a memory location address
 - Read the old value, calculate a new value, and write the new value to the location
- The hardware ensures that no other threads can perform another read-modify-write operation on the same location until the current atomic operation is complete
 - Any other threads that attempt to perform an atomic operation on the same location will typically be held in a queue
 - All threads perform their atomic operations serially on the same location

Atomic Operations in CUDA

- Performed by calling functions that are translated into single instructions (a.k.a. intrinsic functions or intrinsics)
- Operation on one 32-bit or 64-bit word residing in global or shared memory.
- Atomic functions can only be used in device functions
- Atomic add, sub, inc, dec, min, max, exch (exchange), CAS (compare and swap), and, or, xor

Some examples

- Atomic Add
- `int atomicAdd(int* address, int val);`
 - reads the 32-bit word old from the location pointed to by address in global or shared memory, computes (old + val), and stores the result back to memory at the same address. The function returns old.
- Unsigned 32-bit integer atomic add
- `unsigned int atomicAdd(unsigned int* address, unsigned int val);`
- Unsigned 64-bit integer atomic add
- `unsigned long long int atomicAdd(unsigned long long int* address, unsigned long long int val);`
- Single-precision floating-point atomic add (capability > 2.0)
- `float atomicAdd(float* address, float val);`
- Double precision floating-point atomic add (capability > 6.0)
- `double atomicAdd(double* address, double val);`

atomicCAS

- `int atomicCAS(int* address, int compare, int val);`
- `unsigned int atomicCAS(unsigned int* address,
 unsigned int compare,
 unsigned int val);`
- `unsigned long long int atomicCAS(unsigned long long
 int* address,
 unsigned long long int
 compare,
 unsigned long long int
 val);`
- reads the 32-bit or 64-bit word `old` located at the address `address` in global or shared memory, computes `(old == compare ? val : old)`, and stores the result back to memory at the same address. These three operations are performed in one atomic transaction. The function returns `old` (Compare And Swap).

atomicCAS

- Note that any atomic operation can be implemented based on `atomicCAS()` (Compare And Swap). For example, `atomicAdd()` for double-precision floating-point numbers is not available on devices with compute capability lower than 6.0 but it can be implemented as follows:

```
#if __CUDA_ARCH__ < 600
__device__ double atomicAdd(double* address, double val)
{
    unsigned long long int* address_as_ull =
        (unsigned long long int*) address;

    unsigned long long int old = *address_as_ull;
    unsigned long long int assumed;
    do {
        assumed = old;
        old = atomicCAS(address_as_ull, assumed,
            __double_as_longlong(val + __longlong_as_double(assumed)));
        // Note: uses integer comparison to avoid hang in case
        //       of NaN (since NaN != NaN)
    } while (assumed != old);
    return __longlong_as_double(old);
}
#endif
```

atomicCAS

- Note that any atomic operation can be implemented based on `atomicCAS()` (Compare And Swap). For example, `atomicAdd()` for double-precision floating-point numbers is not available on devices with compute capability lower than 6.0 but it can be implemented as follows:

```
#if __CUDA_ARCH__ < 600
__device__ double atomicAdd(double* address, double val)
{
    unsigned long long int* address_as_ull =
        (unsigned long long int*) address;

    unsigned long long int old = *address_as_ull;
    unsigned long long int assumed;
    do {
        assumed = old;
        old = atomicCAS(address_as_ull, assumed,
            __double_as_longlong(val + __longlong_as_double(assumed)));
        // Note: uses integer comparison to avoid hang in case
        //       of NaN (since NaN != NaN)
    } while (assumed != old);
    return __longlong_as_double(old);
}
#endif
```

Reinterpret the bits in the 64-bit signed integer value as a double-precision floating point value.

Atomic operations and caches

- Atomic operations serialize simultaneous updates to a location, thus to improve performance the serialization should be as fast as possible
 - changing locations in global memory is slow: e.g. with an access latency of 200cycles and 1 Ghz clock the throuput of atomics is $\frac{1}{400} \text{ (atomics/clock)} * 1 \text{ G (clocks/secs)} = 2.5\text{M atomics/secs}$ (vs. the Gflops of modern GPUs)
- Cache memories are the primary tool for reducing memory access latency (e.g. 10s of cycle vs. 100s of cycles).
- Recent GPUs allow atomic operation to be performed in the last level cache, which is shared among all SMs.

Privatization

- The latency for accessing memory can be dramatically reduced by placing data in the shared memory.
Shared memory is private to each SM and has very short access latency (a few cycles); this directly translates into increase throughput of atomic operations.
- The problem is that due to the private nature of shared memory, the updates by threads in one thread block is no longer visible to threads in other blocks.

Privatization

- The idea of **privatization** is to replicate highly contended data structures into private copies so that each thread (or each subset of threads) can access a private copy.
The benefit is that the private copies can be accessed with much less contention and often at much lower latency.
- These private copies can dramatically increase the throughput for updating the data structures. The downside is that the private copies need to be merged into the original data structure after the computation completes. One must carefully balance between the level of contention and the merging cost.

Example: histogram computation

```
__global__ void histogram_kernel(const char *input, unsigned int *bins,
                                unsigned int num_elements,
                                unsigned int num_bins) {
    unsigned int tid = blockIdx.x * blockDim.x + threadIdx.x;
    // Privatized bins
    extern __shared__ unsigned int bins_s[];
    for (unsigned int binIdx = threadIdx.x; binIdx < num_bins;
         binIdx += blockDim.x) {
        bins_s[binIdx] = 0;
    }
    __syncthreads();

    // Histogram
    for (unsigned int i = tid; i < num_elements; i += blockDim.x * gridDim.x) {
        atomicAdd(&(bins_s[(unsigned int)input[i]]), 1);
    }
    __syncthreads();

    // Commit to global memory
    for (unsigned int binIdx = threadIdx.x; binIdx < num_bins;
         binIdx += blockDim.x) {
        atomicAdd(&(bins[binIdx]), bins_s[binIdx]);
    }
}
```


Example: histogram computation

```
__global__ void histogram_kernel(const char *input, unsigned int *bins,  
                                unsigned int num_elements,  
                                unsigned int num_bins) {  
    unsigned int tid = blockIdx.x * blockDim.x + threadIdx.x;  
    // Privatized bins  
    extern __shared__ unsigned int bins_s[];  
    for (unsigned int binIdx = threadIdx.x; binIdx < num_bins;
```

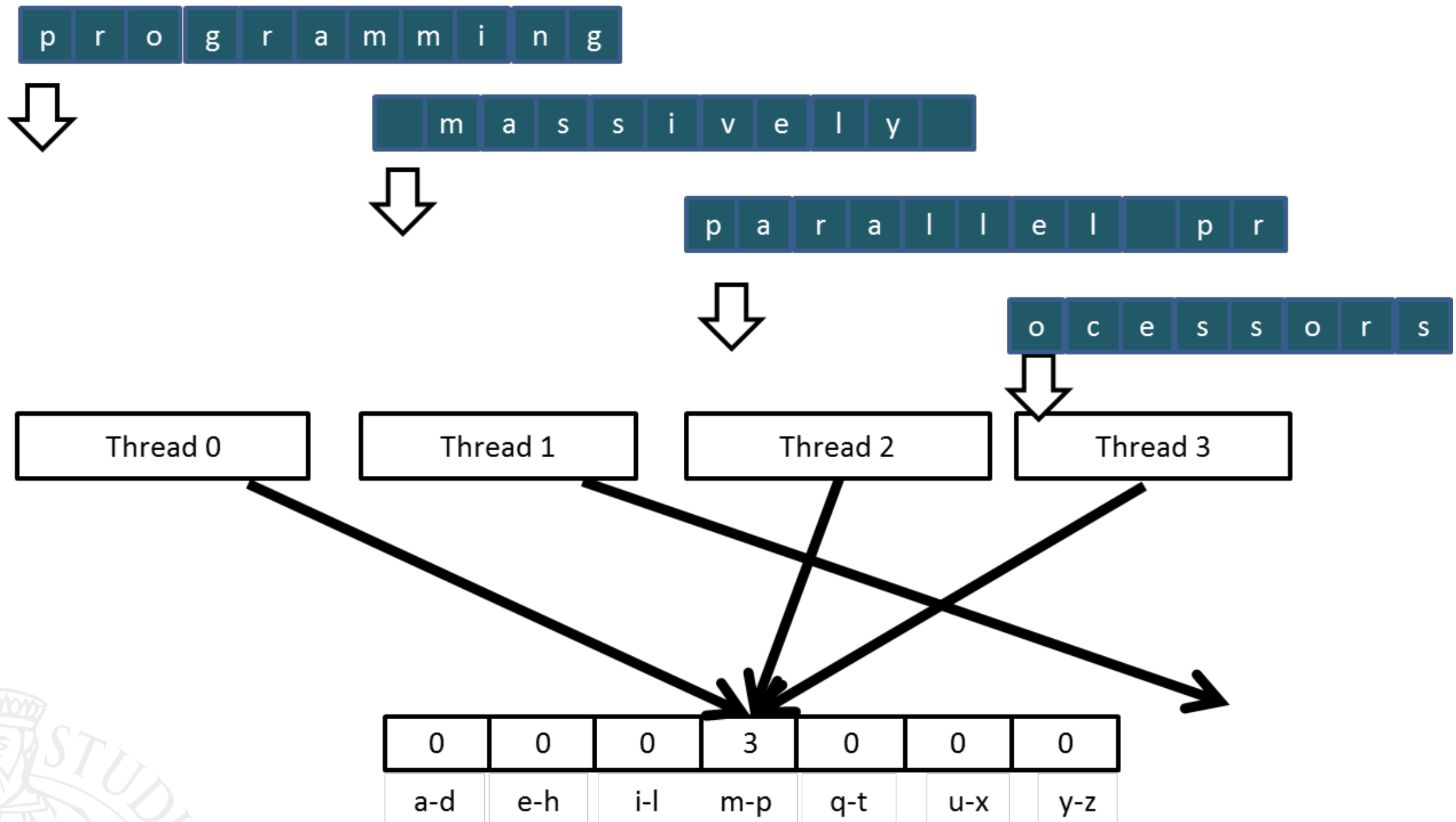
Dynamically allocated shared memory. To allocate it dynamically invoke the kernel with:

```
dim3 blockDim(256), gridDim(30);  
histogram_kernel<<<gridDim, blockDim,  
                num_bins * sizeof(unsigned int)>>>  
    (input, bins, num_elements, num_bins);  
}  
__syncthreads();  
  
// Commit to global memory  
for (unsigned int binIdx = threadIdx.x; binIdx < num_bins;  
     binIdx += blockDim.x) {  
    atomicAdd(&(bins[binIdx]), bins_s[binIdx]);  
}  
}
```

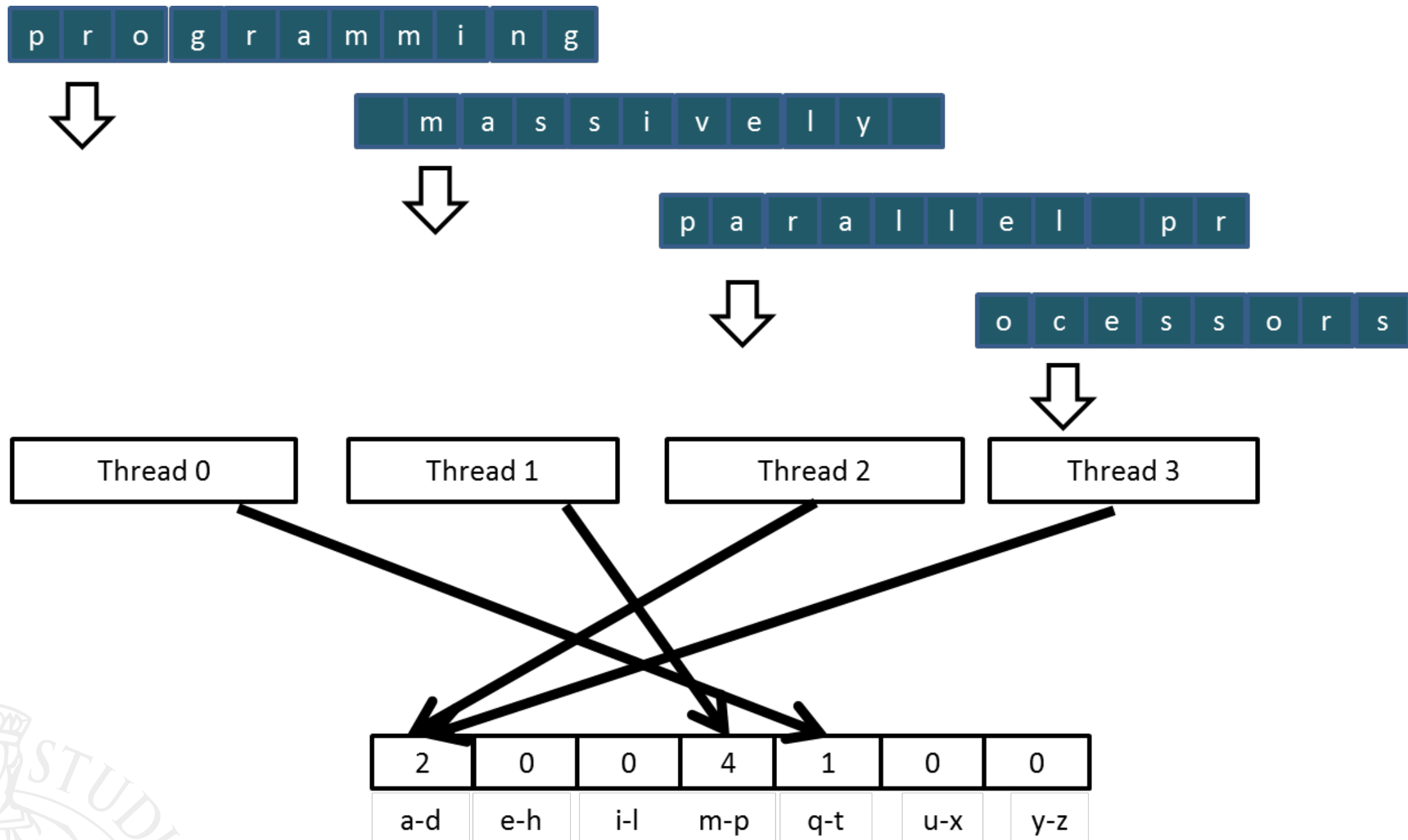
Improving memory access

- A simple parallel histogram algorithm partitions the input into sections
- Each section is given to a thread, that iterates through it
- This makes sense in CPU code, where we have few threads, each of which can efficiently use the cache lines when accessing memory
- This access is less convenient in GPUs

Sectioned Partitioning (Iteration #1)

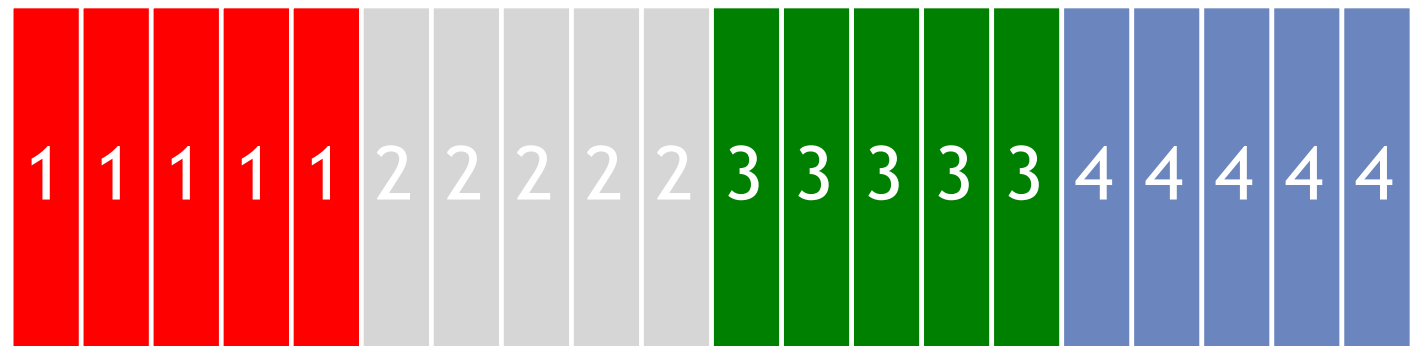


Sectioned Partitioning (Iteration #2)

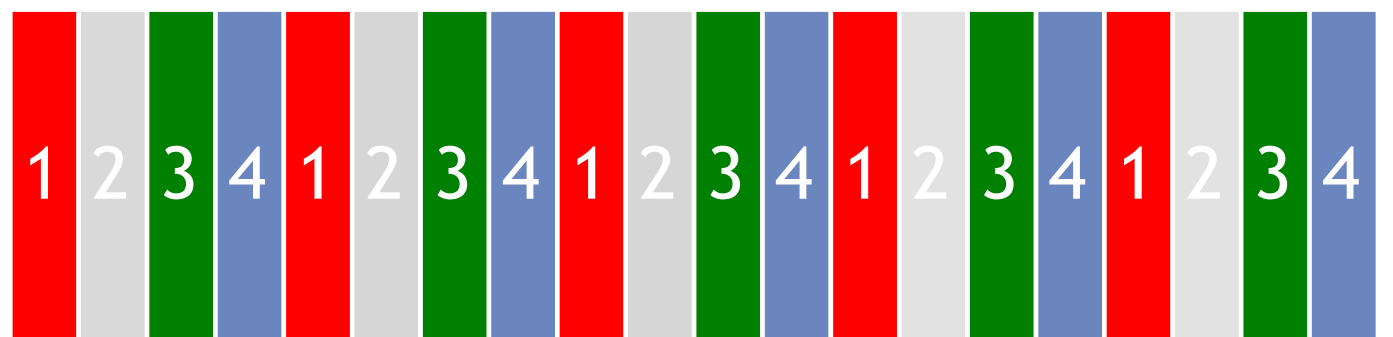


Input Partitioning Affects Memory Access Efficiency

- Sectioned partitioning results in poor memory access efficiency
 - Adjacent threads do not access adjacent memory locations
 - Accesses are not coalesced
 - DRAM bandwidth is poorly utilized

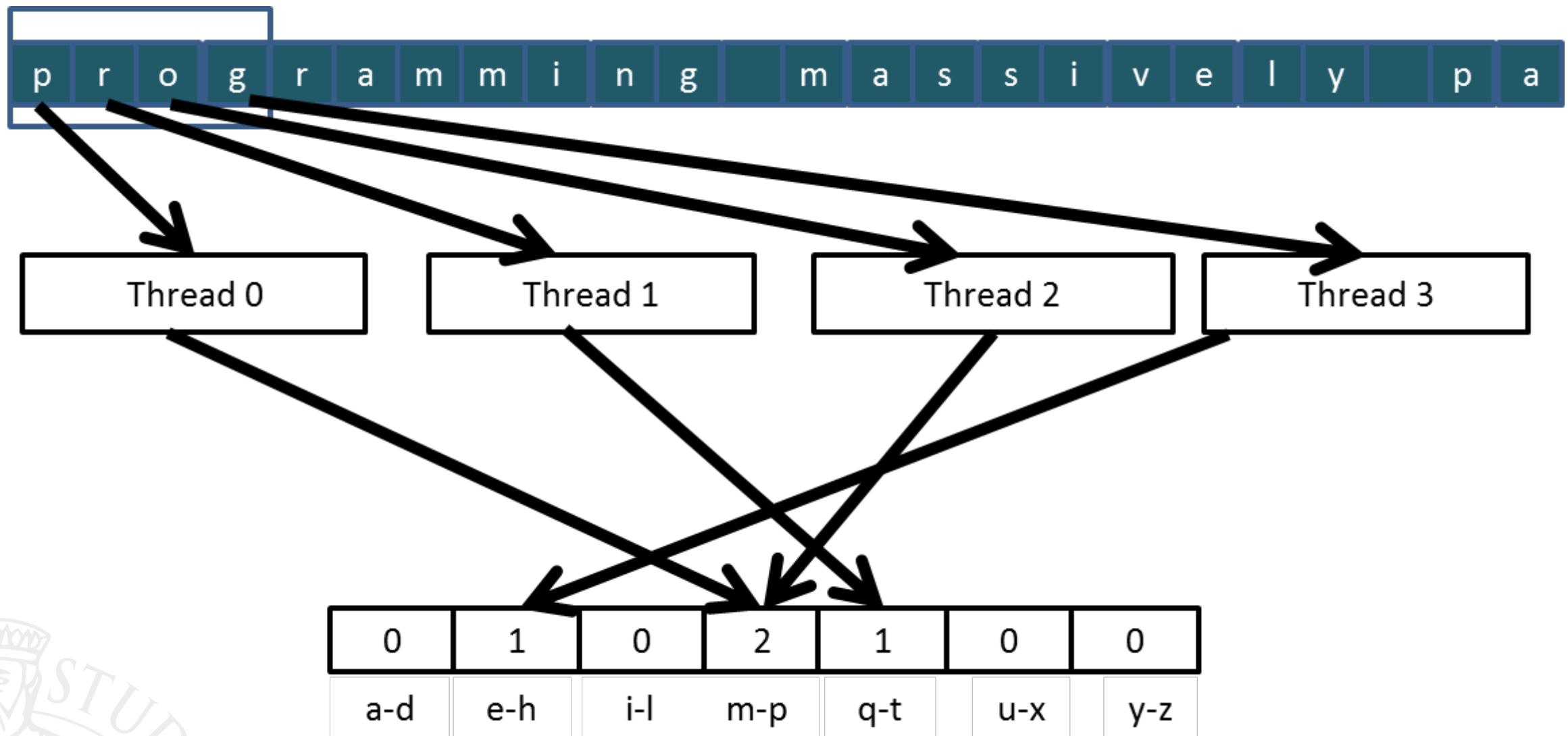


- Change to interleaved partitioning
 - All threads process a contiguous section of elements
 - They all move to the next section and repeat
 - The memory accesses are coalesced

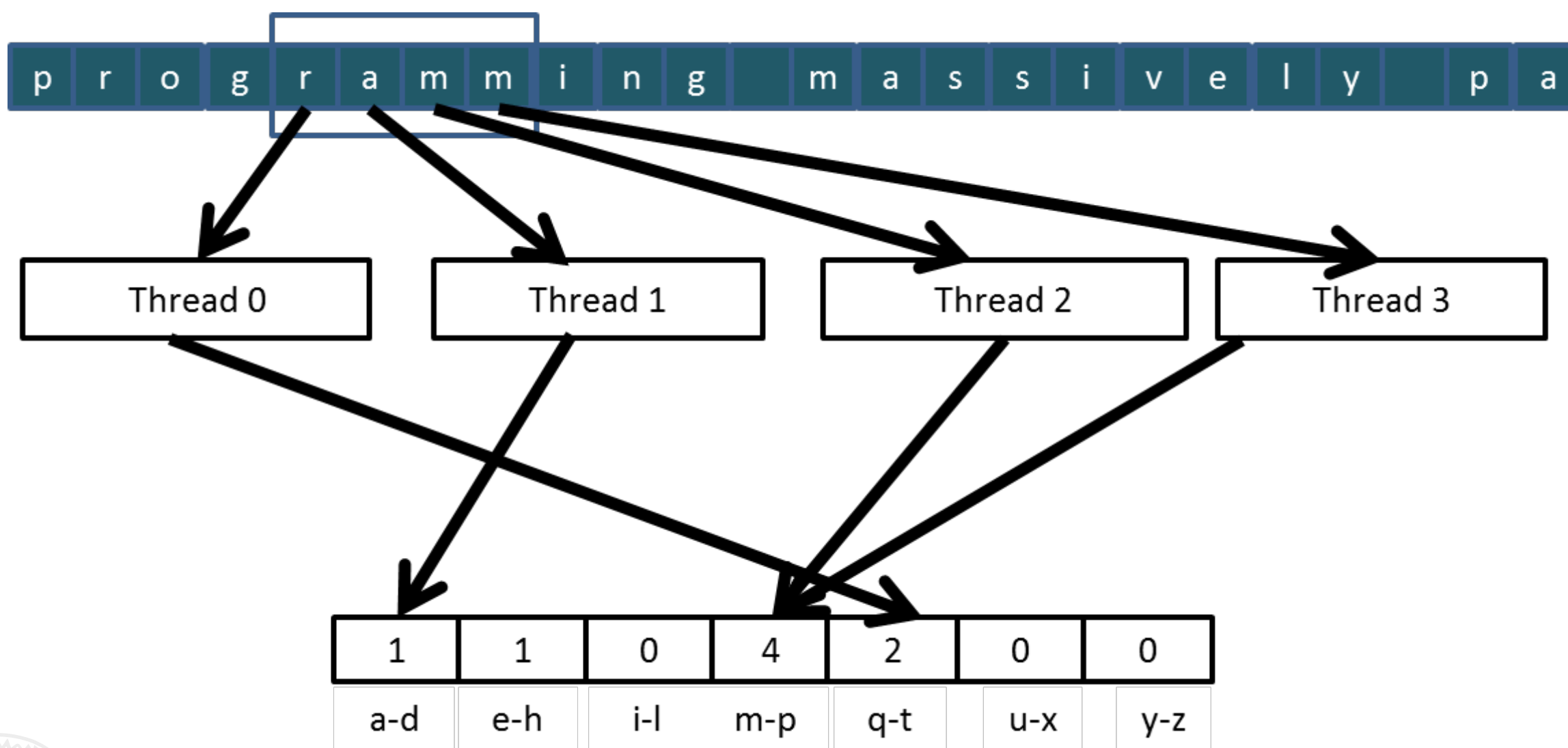


Interleaved Partitioning of Input

- For coalescing and better memory access performance



Interleaved Partitioning (Iteration 2)



A stride algorithm

```
__global__ void histo_kernel(unsigned char *buffer,  
                             long size, unsigned int *histo)  
{  
  
    int i = threadIdx.x + blockIdx.x * blockDim.x;  
  
    // stride is total number of threads  
    int stride = blockDim.x * gridDim.x;  
  
    // All threads handle blockDim.x * gridDim.x  
    // consecutive elements  
    while (i < size) {  
        atomicAdd( &(histo[buffer[i]]), 1);  
        i += stride;  
    }  
}
```


A stride algorithm

Calculates a **stride** value, which is the total number threads launched during kernel invocation (**blockDim.x*gridDim.x**). In the first iteration of the **while** loop, each thread index the input buffer using its global thread index: Thread 0 accesses element 0, Thread 1 accesses element 1, etc. Thus, all threads jointly process the first blockDim.x*gridDim.x elements of the input buffer.

```
int i = threadIdx.x + blockIdx.x * blockDim.x;
```

```
// stride is total number of threads
```

```
int stride = blockDim.x * gridDim.x;
```

```
// All threads handle blockDim.x * gridDim.x
```

```
// consecutive elements
```

```
while (i < size) {
```

```
    atomicAdd( &(histo[buffer[i]]), 1);
```

```
    i += stride;
```

```
}
```

```
}
```

A stride algorithm

Calculates a **stride** value, which is the total number threads launched during kernel invocation (**blockDim.x*gridDim.x**). In the first iteration of the **while** loop, each thread index the input buffer using its global thread index: Thread 0 accesses element 0, Thread 1 accesses element 1, etc. Thus, all threads jointly process the first blockDim.x*gridDim.x elements of the input buffer.

```
int i = threadIdx.x + blockIdx.x * blockDim.x;
```

```
// stride is total number of threads
```

```
int stride = blockDim.x * gridDim.x;
```

```
// All threads handle blockDim.x * gridDim.x
```

```
// consecutive elements
```

```
while (i < size) {
```

```
    atomicAdd( &(histo[buffer[i]]), 1);
```

```
    i += stride;
```

```
}
```

The while loop controls the iterations for each thread. When the index of a thread exceeds the valid range of the input buffer (**i** is greater than or equal to **size**), the thread has completed processing its partition and will exit the loop.



CUDA: parallel patterns - convolution

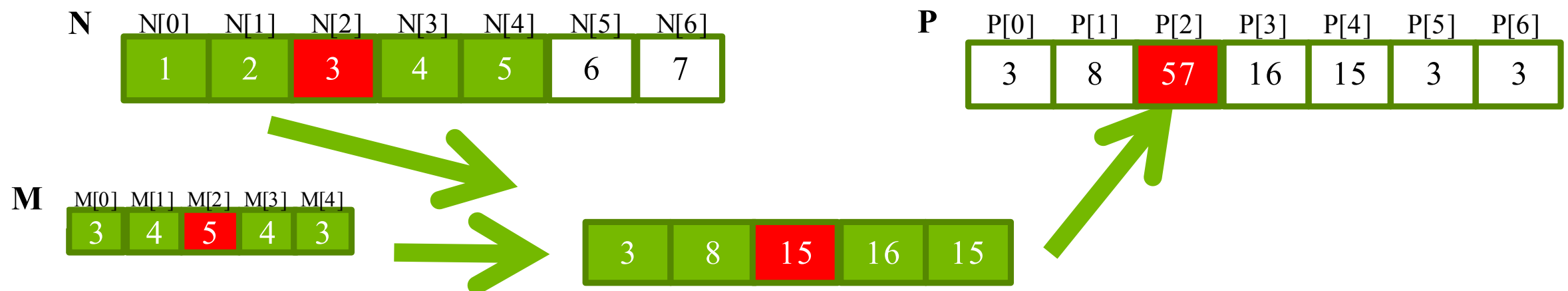
Convolution (stencil)

- An array operation where each output data element is a weighted sum of a collection of neighboring input elements
- The weights used in the weighted sum calculation are defined by an input mask array, commonly referred to as the convolution kernel
- We will refer to these mask arrays as **convolution masks** to avoid confusion.
- The value pattern of the mask array elements defines the type of filtering done

Convolution (stencil)

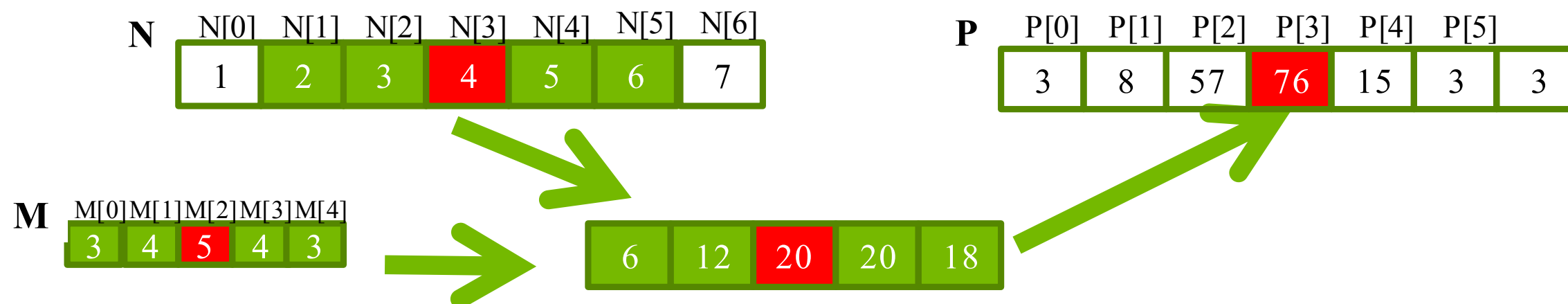
- An array operation where each output data element is a weighted sum of a collection of neighboring input elements
- The weights used in the weighted sum calculation are defined by an input mask array, commonly referred to as the convolution kernel
- We will refer to these mask arrays as **convolution masks** to avoid confusion.
- Often performed as a filter that transforms signal or pixel values into more desirable values.

1D Convolution Example

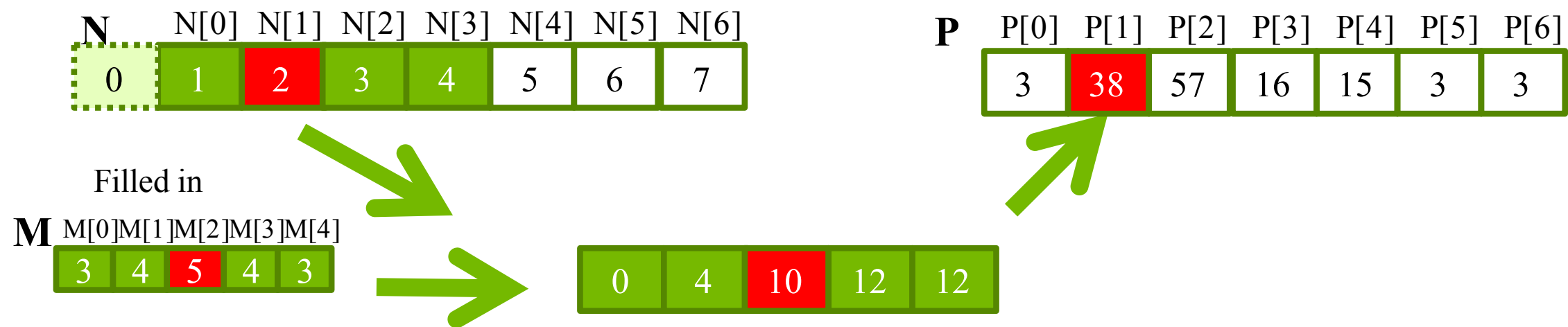


- Commonly used for audio processing
- Mask size is usually an odd number of elements for symmetry (5 in this example)
- The figure shows calculation of $P[2]$
- $$P[2] = N[0]*M[0] + N[1]*M[1] + N[2]*M[2] + N[3]*M[3] + N[4]*M[4]$$

Example: calculation of $P[3]$



Convolution Boundary Condition



- Calculation of output elements near the boundaries (beginning and end) of the array need to deal with “ghost” elements
- Different policies (0, replicates of boundary values, use of symmetrical values, etc.)

A 1D Convolution Kernel with Boundary Condition Handling

```
__global__ void convolution_1D_basic_kernel(float *N, float *M,
    float *P, int Mask_Width, int Width) {
    int i = blockIdx.x*blockDim.x + threadIdx.x;

    float Pvalue = 0;
    int N_start_point = i - (Mask_Width/2);

    for (int j = 0; j < Mask_Width; j++) {
        if (N_start_point + j >= 0 &&
            N_start_point + j < Width) {
            Pvalue += N[N_start_point + j]*M[j];
        }
    }

    P[i] = Pvalue;
}
```

- This kernel forces all elements outside the valid input range to 0

A 1D Convolution Kernel with Boundary Condition Handling

```
__global__ void convolution_1D_basic_kernel(float *N, float *M,
    float *P, int Mask_Width, int Width) {
    int i = blockIdx.x*blockDim.x + threadIdx.x;

    float Pvalue = 0;
    int N_start_point = i - (Mask_Width/2);

    for (int j = 0; j < Mask_Width; j++) {
        if (N_start_point + j >= 0 &&
            N_start_point + j < Width) {
            Pvalue += N[N_start_point + j]*M[j];
        }
    }

    P[i] = Pvalue;
}
```

- This kernel forces all elements outside the valid input range to 0

A 1D Convolution Kernel with Boundary Condition Handling

```
__global__ void convolution_1D_basic_kernel(float *N, float *M,
    float *P, int Mask_Width, int Width) {
    int i = blockIdx.x*blockDim.x + threadIdx.x;

    float Pvalue = 0;
    int N_start_point = i - (Mask_Width/2);

    for (int j = 0; j < Mask_Width; j++) {
        if (N_start_point + j >= 0 &&
            N_start_point + j < Width) {
            Pvalue += N[N_start_point + j]*M[j];
        }
    }

    P[i] = Pvalue;
}
```

Source of bad performance: 1 floating-point operation
per global memory access

- This kernel forces all elements outside the valid input range to 0

2D Convolution

N

1	2	3	4	5	6	7
2	3	4	5	6	7	8
3	4	5	6	7	8	9
4	5	6	7	8	5	6
5	6	7	8	5	6	7
6	7	8	9	0	1	2
7	8	9	0	1	2	3

P

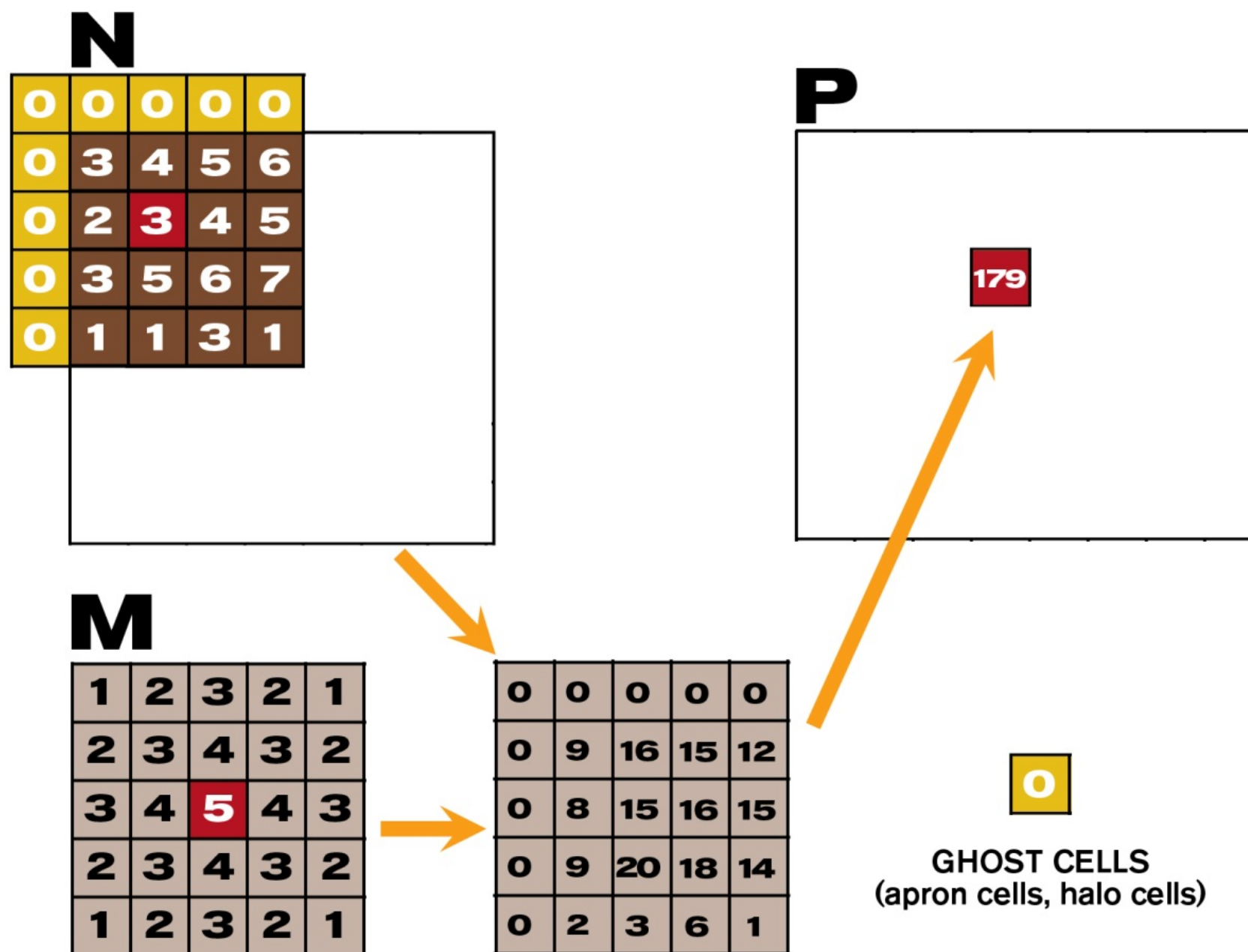
1	2	3	4	5			
2	3	4	5	6			
3	4	321	6	7			
4	5	6	7	8			
5	6	7	8	5			

M

1	2	3	2	1
2	3	4	3	2
3	4	5	4	3
2	3	4	3	2
1	2	3	2	1

1	4	9	8	5
4	9	16	15	12
4	16	25	24	21
8	15	24	21	16
5	12	21	16	5

2D Convolution – Ghost Cells



__global__

```
void convolution_2D_basic_kernel(unsigned char * in, unsigned char * mask, unsigned char * out,
    int maskwidth, int w, int h) {
    int Col = blockIdx.x * blockDim.x + threadIdx.x;
    int Row = blockIdx.y * blockDim.y + threadIdx.y;

    if (Col < w && Row < h) {
        int pixVal = 0;

        N_start_col = Col - (maskwidth/2);
        N_start_row = Row - (maskwidth/2);

        // Get the of the surrounding box
        for(int j = 0; j < maskwidth; ++j) {
            for(int k = 0; k < maskwidth; ++k) {

                int curRow = N_Start_row + j;
                int curCol = N_start_col + k;
                // Verify we have a valid image pixel
                if(curRow > -1 && curRow < h && curCol > -1 && curCol < w) {
                    pixVal += in[curRow * w + curCol] * mask[j*maskwidth+k];
                }
            }
        }

        // Write our new pixel value out
        out[Row * w + Col] = (unsigned char)(pixVal);
    }
}
```

__global__

```
void convolution_2D_basic_kernel(unsigned char * in, unsigned char * mask, unsigned char * out,  
    int maskwidth, int w, int h) {
```

```
    int Col = blockIdx.x * blockDim.x + threadIdx.x;  
    int Row = blockIdx.y * blockDim.y + threadIdx.y;
```

```
    if (Col < w && Row < h) {  
        int pixVal = 0;
```

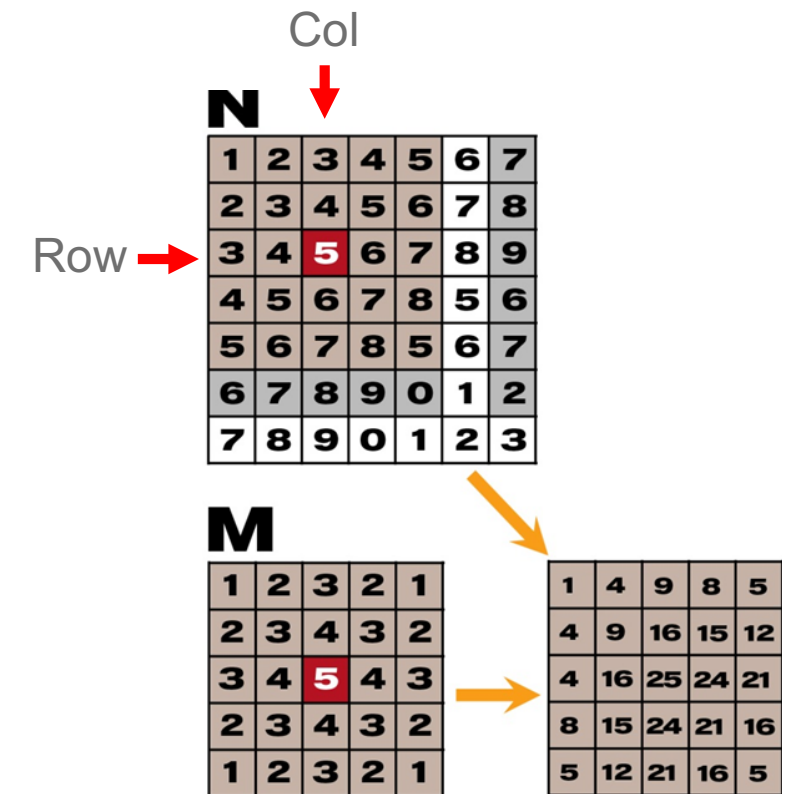
```
        N_start_col = Col - (maskwidth/2);  
        N_start_row = Row - (maskwidth/2);
```

```
        // Get the of the surrounding box  
        for(int j = 0; j < maskwidth; ++j) {  
            for(int k = 0; k < maskwidth; ++k) {
```

```
                int curRow = N_Start_row + j;  
                int curCol = N_start_col + k;  
                // Verify we have a valid image pixel  
                if(curRow > -1 && curRow < h && curCol > -1 && curCol < w) {  
                    pixVal += in[curRow * w + curCol] * mask[j*maskwidth+k];  
                }  
            }  
        }
```

```
        // Write our new pixel value out  
        out[Row * w + Col] = (unsigned char)(pixVal);
```

```
    }  
}
```



__global__

```
void convolution_2D_basic_kernel(unsigned char * in, unsigned char * mask, unsigned char * out,  
    int maskwidth, int w, int h) {
```

```
    int Col = blockIdx.x * blockDim.x + threadIdx.x;  
    int Row = blockIdx.y * blockDim.y + threadIdx.y;
```

```
    if (Col < w && Row < h) {  
        int pixVal = 0;
```

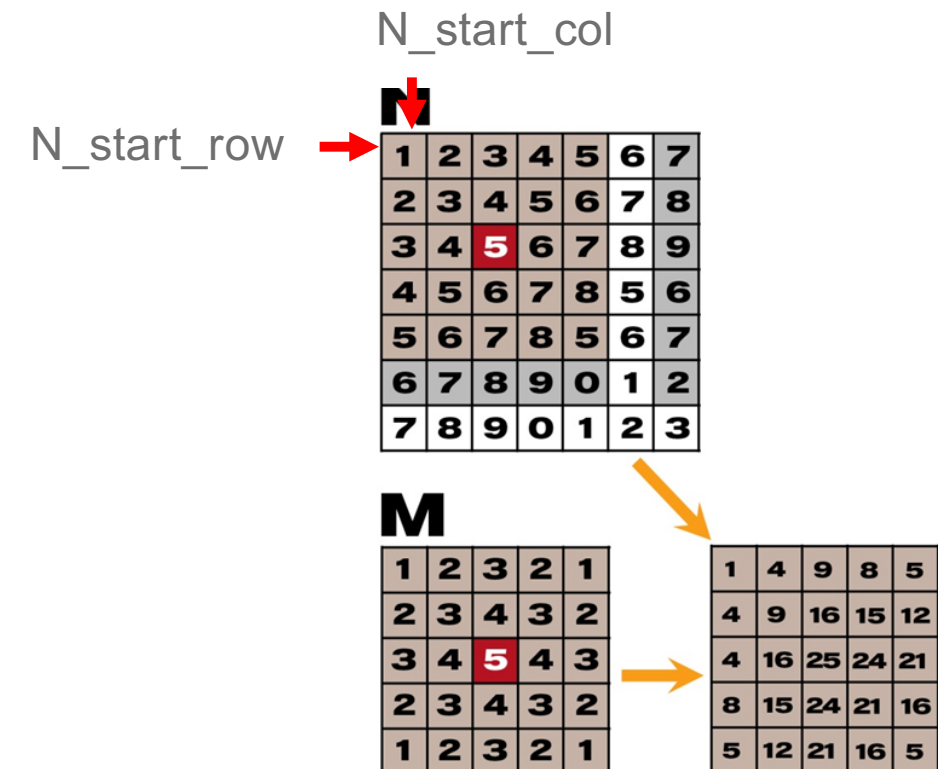
```
        N_start_col = Col - (maskwidth/2);  
        N_start_row = Row - (maskwidth/2);
```

```
        // Get the of the surrounding box  
        for(int j = 0; j < maskwidth; ++j) {  
            for(int k = 0; k < maskwidth; ++k) {
```

```
                int curRow = N_Start_row + j;  
                int curCol = N_start_col + k;  
                // Verify we have a valid image pixel  
                if(curRow > -1 && curRow < h && curCol > -1 && curCol < w) {  
                    pixVal += in[curRow * w + curCol] * mask[j*maskwidth+k];  
                }  
            }  
        }
```

```
        // Write our new pixel value out  
        out[Row * w + Col] = (unsigned char)(pixVal);
```

```
    }  
}
```



__global__

```
void convolution_2D_basic_kernel(unsigned char * in, unsigned char * mask, unsigned char * out,
    int maskwidth, int w, int h) {
    int Col = blockIdx.x * blockDim.x + threadIdx.x;
    int Row = blockIdx.y * blockDim.y + threadIdx.y;

    if (Col < w && Row < h) {
        int pixVal = 0;

        N_start_col = Col - (maskwidth/2);
        N_start_row = Row - (maskwidth/2);

        // Get the of the surrounding box
        for(int j = 0; j < maskwidth; ++j) {
            for(int k = 0; k < maskwidth; ++k) {

                int curRow = N_Start_row + j;
                int curCol = N_start_col + k;
                // Verify we have a valid image pixel
                if(curRow > -1 && curRow < h && curCol > -1 && curCol < w) {
                    pixVal += in[curRow * w + curCol] * mask[j*maskwidth+k];
                }
            }
        }

        // Write our new pixel value out
        out[Row * w + Col] = (unsigned char)(pixVal);
    }
}
```

__global__

```
void convolution_2D_basic_kernel(unsigned char * in, unsigned char * mask, unsigned char * out,
    int maskwidth, int w, int h) {
    int Col = blockIdx.x * blockDim.x + threadIdx.x;
    int Row = blockIdx.y * blockDim.y + threadIdx.y;

    if (Col < w && Row < h) {
        int pixVal = 0;

        N_start_col = Col - (maskwidth/2);
        N_start_row = Row - (maskwidth/2);

        // Get the of the surrounding box
        for(int j = 0; j < maskwidth; ++j) {
            for(int k = 0; k < maskwidth; ++k) {

                int curRow = N_Start_row + j;
                int curCol = N_start_col + k;
                // Verify we have a valid image pixel
                if(curRow > -1 && curRow < h && curCol > -1 && curCol < w) {
                    pixVal += in[curRow * w + curCol] * mask[j*maskwidth+k];
                }
            }
        }

        // Write our new pixel value out
        out[Row * w + Col] = (unsigned char)(pixVal);
    }
}
```

__global__

```
void convolution_2D_basic_kernel(unsigned char * in, unsigned char * mask, unsigned char * out,
    int maskwidth, int w, int h) {
    int Col = blockIdx.x * blockDim.x + threadIdx.x;
    int Row = blockIdx.y * blockDim.y + threadIdx.y;

    if (Col < w && Row < h) {
        int pixVal = 0;

        N_start_col = Col - (maskwidth/2);
        N_start_row = Row - (maskwidth/2);

        // Get the of the surrounding box
        for(int j = 0; j < maskwidth; ++j) {
            for(int k = 0; k < maskwidth; ++k) {

                int curRow = N_Start_row + j;
                int curCol = N_start_col + k;
                // Verify we have a valid image pixel
                if(curRow > -1 && curRow < h && curCol > -1 && curCol < w) {
                    pixVal += in[curRow * w + curCol] * mask[j*maskwidth+k];
                }
            }
        }

        // Write our new pixel value out
        out[Row * w + Col] = (unsigned char)(pixVal);
    }
}
```

Source of bad performance: 1 floating-point operation
per global memory access

Improving convolution kernel

- Use tiling for the N array element
- Use constant memory for the M mask
 - it's typically small and is not changed
 - can be read by all threads of the grid

```
#define MAX_MASK_WIDTH 10  
__constant__ float M[MAX_MASK_WIDTH];
```

```
cudaMemcpyToSymbol(M, M_h, Mask_Width*sizeof(float));
```

Improving convolution kernel

- Use tiling for the N array element
- Use constant memory for the M mask
 - it's typically small and is not changed
 - can be read by all threads of the grid

```
#define MAX_MASK_WIDTH 10  
__constant__ float M[MAX_MASK_WIDTH];
```

 global variable

```
cudaMemcpyToSymbol(M, M_h, Mask_Width*sizeof(float));
```

Convolution with constant memory

```
__global__ void convolution_1D_basic_kernel(float *N, float *P,  
                                           int Mask_Width, int Width) {  
  
    int i = blockIdx.x*blockDim.x + threadIdx.x;  
  
    float Pvalue = 0;  
  
    int N_start_point = i - (Mask_Width/2);  
  
    for (int j = 0; j < Mask_Width; j++) {  
        if (N_start_point + j >= 0 && N_start_point + j < Width) {  
            Pvalue += N[N_start_point + j]*M[j];  
        }  
    }  
  
    P[i] = Pvalue;  
}
```

Convolution with constant memory

```
__global__ void convolution_1D_basic_kernel(float *N, float *P,  
                                           int Mask_Width, int Width) {  
  
    int i = blockIdx.x*blockDim.x + threadIdx.x;  
  
    float Pvalue = 0;  
  
    int N_start_point = i - (Mask_Width/2);  
  
    for (int j = 0; j < Mask_Width; j++) {  
        if (N_start_point + j >= 0 && N_start_point + j < Width) {  
            Pvalue += N[N_start_point + j]*M[j];  
        }  
    }  
  
    P[i] = Pvalue;  
}
```

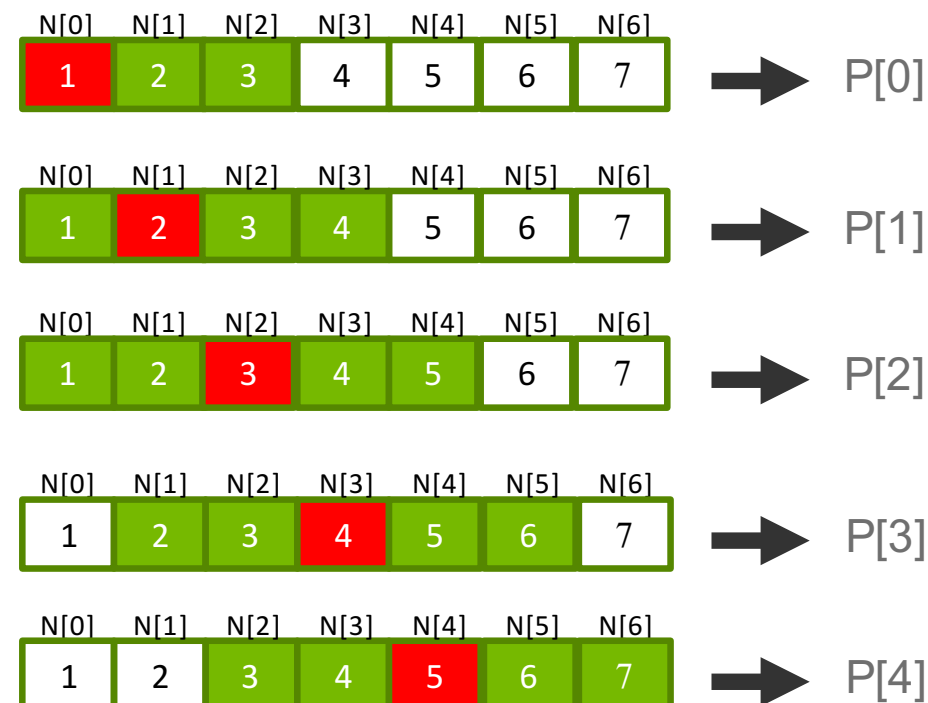
2 floating-point operations per global memory access (N)



CUDA: parallel patterns - convolution & tiling

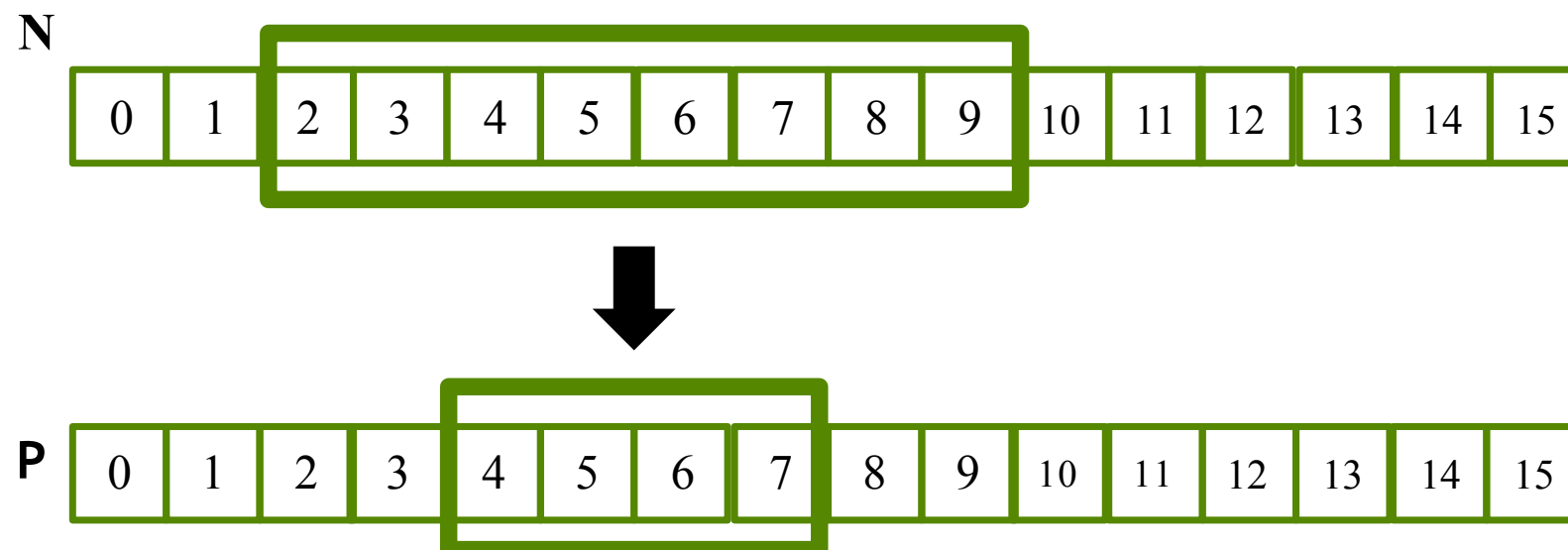
Tiling & convolution

- Calculation of adjacent output elements involve shared input elements
 - E.g., $N[2]$ is used in calculation of $P[0]$, $P[1]$, $P[2]$. $P[3]$ and $P[5]$ assuming a 1D convolution Mask_Width of width 5
- We can load all the input elements required by all threads in a block into the shared memory to reduce global memory accesses

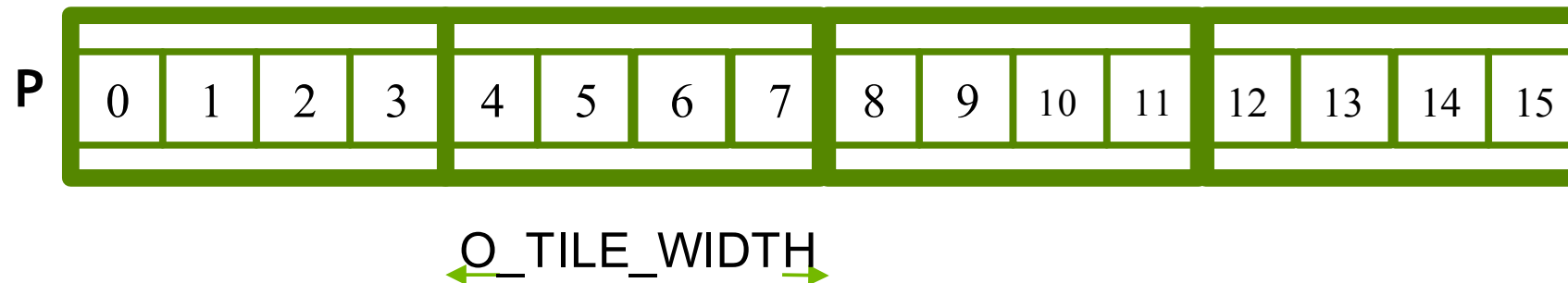


Input Data Needs

- Assume that we want to have each block to calculate T output elements
 - $T + \text{Mask_Width} - 1$ input elements are needed to calculate T output elements
 - $T + \text{Mask_Width} - 1$ is usually not a multiple of T , except for small T values
 - T is usually significantly larger than Mask_Width

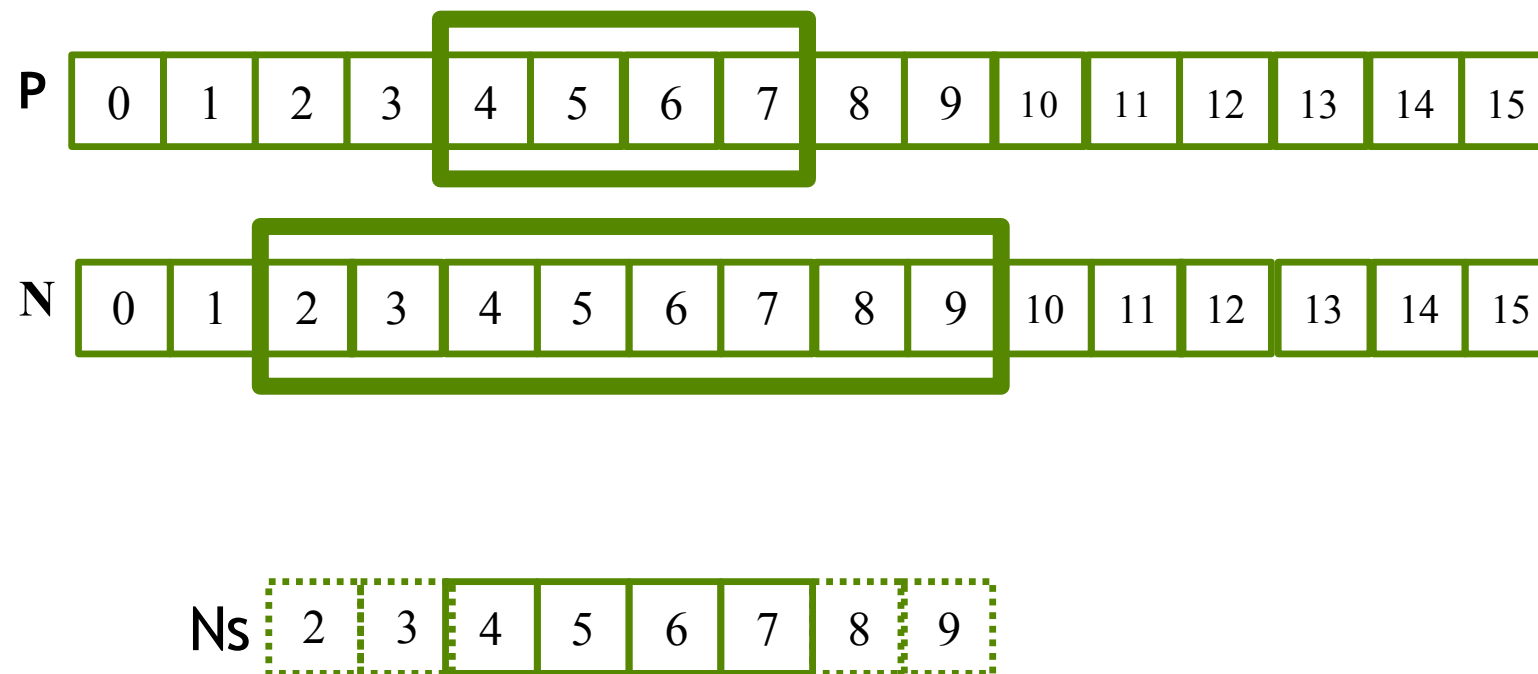


Definition – output tile



- Each thread block calculates an output tile
- Each output tile width is O_TILE_WIDTH
- For each thread:
 - O_TILE_WIDTH is 4 in this example

Definition - Input Tiles

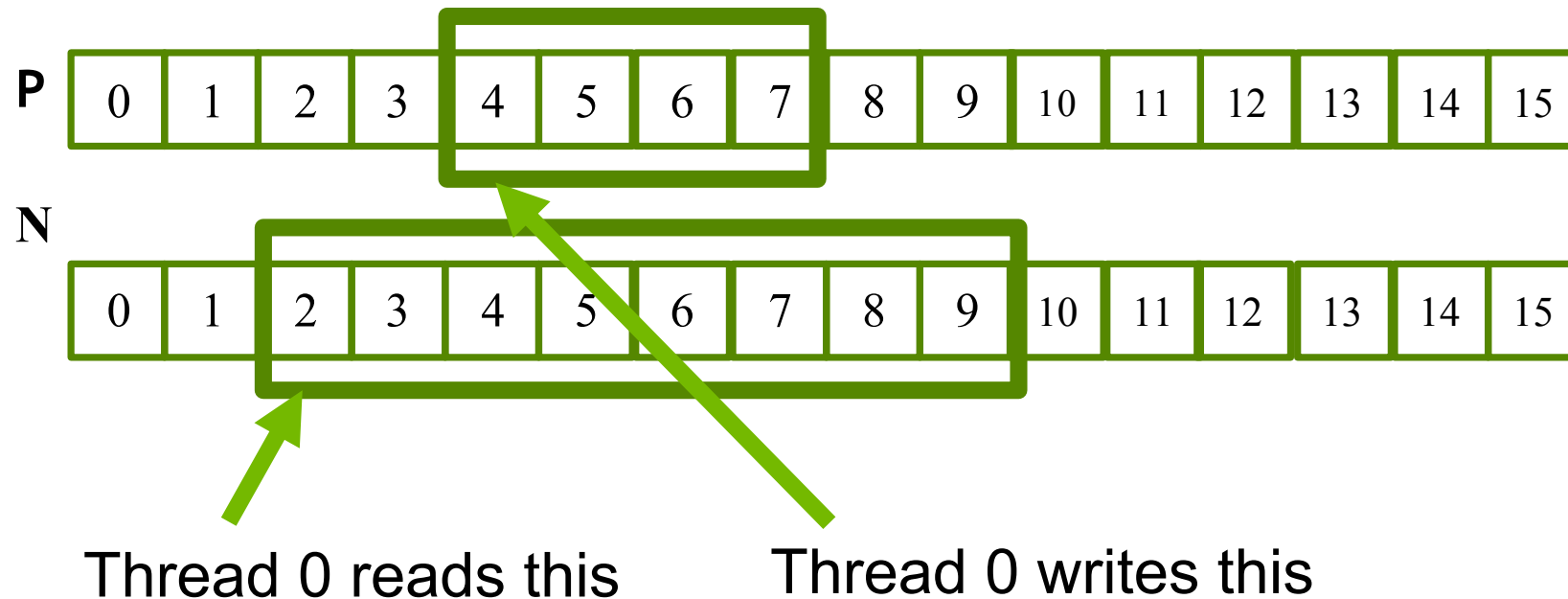


- Each input tile has all values needed to calculate the corresponding output tile.

Two Design Options

- Design 1: The size of each thread block matches the size of an output tile
 - All threads participate in calculating output elements
 - `blockDim.x` would be 4 in our example
 - Some threads need to load more than one input element into the shared memory
- Design 2: The size of each thread block matches the size of an input tile
 - Some threads will not participate in calculating output elements
 - `blockDim.x` would be 8 in our example
 - Each thread loads one input element into the shared memory

Thread to Input and Output Data Mapping



- For each thread:
 - $\text{Index}_i = \text{index}_o - n$
 - where n is $\text{Mask_Width} / 2$
 - n is 2 in this example

Loading input tiles

All threads participate:

```
float output = 0.0f;
```

```
if((index_i >= 0) && (index_i < Width)) {
```

```
    Ns[tx] = N[index_i];
```

```
} else {
```

```
    Ns[tx] = 0.0f;
```

```
}
```


Calculating output

- Some threads do not participate: Only Threads 0 through 0_TILE_WIDTH-1 participate in calculation of output.
- `index_o = blockIdx.x*0_TILE_WIDTH + threadIdx.x`

```
if (threadIdx.x < 0_TILE_WIDTH){  
    output = 0.0f;  
    for(j = 0; j < Mask_Width; j++) {  
        output += M[j] * Ns[j+threadIdx.x];  
    }  
    P[index_o] = output;  
}
```


Setting Block Size

```
#define O_TILE_WIDTH 1020
```

```
#define BLOCK_WIDTH (O_TILE_WIDTH + 4)
```

```
dim3 dimBlock(BLOCK_WIDTH,1, 1);
```

```
dim3 dimGrid((Width-1)/O_TILE_WIDTH+1, 1, 1)
```

- The Mask_Width is 5 in this example
- In general, block width should be
 - output tile width + (mask width-1)

Shared Memory Data Reuse

N_ds

Mask_Width is 5



Element 2 is used by thread 4 (1X)

Element 3 is used by threads 4, 5 (2X)

Element 4 is used by threads 4, 5, 6 (3X)

Element 5 is used by threads 4, 5, 6, 7 (4X)

Element 6 is used by threads 4, 5, 6, 7 (4X)

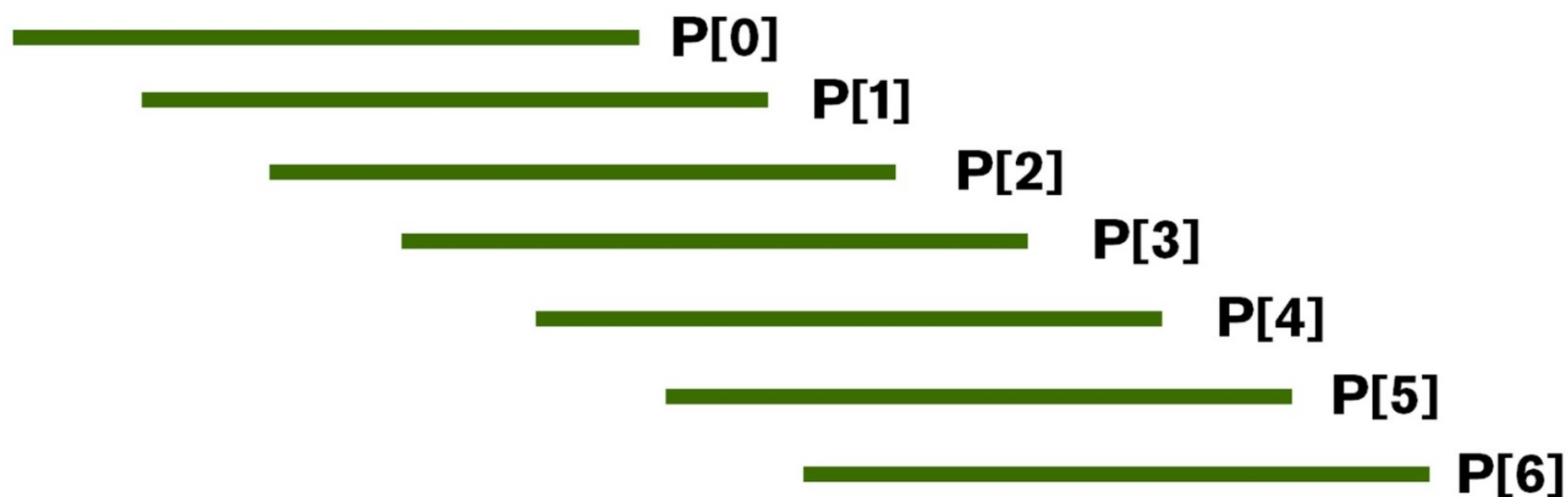
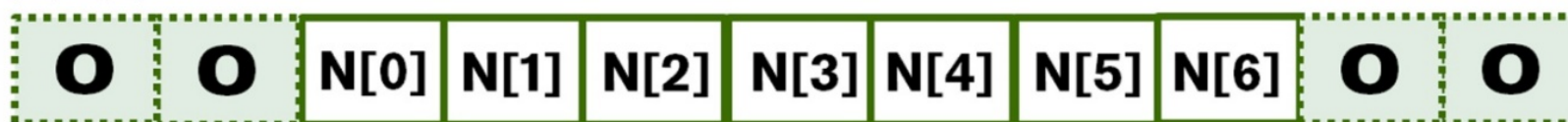
Element 7 is used by threads 5, 6, 7 (3X)

Element 8 is used by threads 6, 7 (2X)

Element 9 is used by thread 7 (1X)

Ghost/Halo cells

N





UNIVERSITÀ
DEGLI STUDI
FIRENZE

Evaluating tiling



An 8-element Convolution Tile

N_ds



P

Mask_Width is 5



- For Mask_Width=5, we load $8+5-1=12$ elements (12 memory loads)

Evaluating reuse

- Each output P element uses 5 N elements:

P[8] uses N[6], N[7], N[8], N[9], N[10]

P[9] uses N[7], N[8], N[9], N[10], N[11]

P[10] use N[8], N[9], N[10], N[11], N[12]

...

P[14] uses N[12], N[13], N[14], N[15], N[16]

P[15] uses N[13], N[14], N[15], N[16], N[17]

N_ds



Mask_Width is 5



Evaluating reuse

- Each output P element uses 5 N elements:

P[8] uses N[6], N[7], N[8], N[9], N[10]

P[9] uses N[7], N[8], N[9], N[10], N[11]

P[10] use N[8], N[9], N[10], N[11], N[12]

$(8+5-1)=12$ elements loaded

8*5 global memory accesses replaced by shared memory accesses

This gives a bandwidth reduction of $40/12=3.3$

N_ds



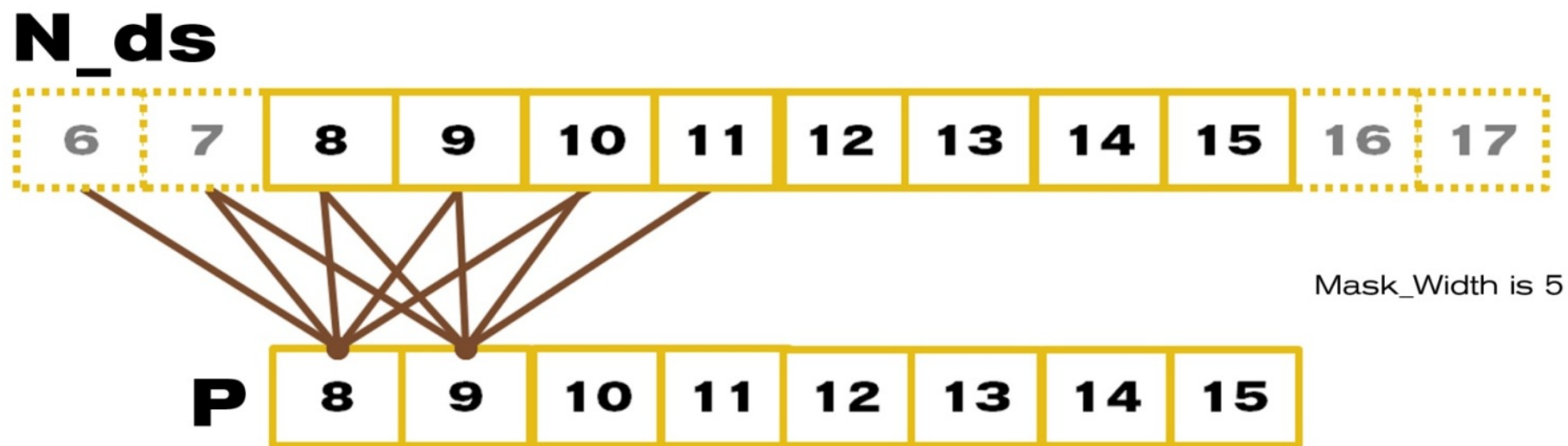
Mask_Width is 5



General 1D tiled convolution

- $O_TILE_WIDTH + MASK_WIDTH - 1$ elements loaded for each input tile
- $O_TILE_WIDTH * MASK_WIDTH$ global memory accesses replaced by shared memory accesses
- This gives a reduction factor of
$$\frac{(O_TILE_WIDTH * MASK_WIDTH)}{(O_TILE_WIDTH + MASK_WIDTH - 1)}$$
- This ignores ghost elements in edge tiles.

Another Way to Look at Reuse



- N[6] is used by P[8] (1X)
- N[7] is used by P[8], P[9] (2X)
- N[8] is used by P[8], P[9], P[10] (3X)
- N[9] is used by P[8], P[9], P[10], P[11] (4X)
- N[10] is used by P[8], P[9], P[10], P[11], P[12] (5X)
- ... (5X)
- N[14] is used by P[12], P[13], P[14], P[15] (4X)
- N[15] is used by P[13], P[14], P[15] (3X)

Another Way to Look at Reuse

- The total number of global memory accesses to the $(8+5-1)=12$ elements of N that is replaced by shared memory accesses is:

$$1+2+3+4+5*(8-5+1)+4+3+2+1 = 10+20+10 = 40$$

- So the reduction is $40/12 = 3.3$



General 1D tiling

- The total number of global memory accesses to the input tile can be calculated as
- $1 + 2 + \dots + \text{MASK_WIDTH} - 1 + \text{MASK_WIDTH} * (\text{O_TILE_WIDTH} - \text{MASK_WIDTH} + 1) + \text{MASK_WIDTH} - 1 + \dots + 2 + 1$
- $= \text{MASK_WIDTH} * (\text{MASK_WIDTH} - 1) + \text{MASK_WIDTH} * (\text{O_TILE_WIDTH} - \text{MASK_WIDTH} + 1)$
- $= \text{MASK_WIDTH} * \text{O_TILE_WIDTH}$
- For a total of $\text{O_TILE_WIDTH} + \text{MASK_WIDTH} - 1$ input tile elements

Examples of Bandwidth Reduction for 1D

- The reduction ratio is:
- $\text{MASK_WIDTH} * \text{O_TILE_WIDTH} / (\text{O_TILE_WIDTH} + \text{MASK_WIDTH} - 1)$

O_TILE_WIDTH	16	32	64	128	256
MASK_WIDTH= 5	4.0	4.4	4.7	4.9	4.9
MASK_WIDTH = 9	6.0	7.2	8.0	8.5	8.7



2D convolution tiles

- $(O_TILE_WIDTH + MASK_WIDTH - 1)^2$ input elements need to be loaded into shared memory
- The calculation of each output element needs to access $MASK_WIDTH^2$ input elements
- $O_TILE_WIDTH^2 * MASK_WIDTH^2$ global memory accesses are converted into shared memory accesses
- The reduction ratio is
 - $O_TILE_WIDTH^2 * MASK_WIDTH^2 / (O_TILE_WIDTH + MASK_WIDTH - 1)^2$

Bandwidth Reduction for 2D

- The reduction ratio is:

$$\frac{O_TILE_WIDTH^2 * MASK_WIDTH^2}{(O_TILE_WIDTH + MASK_WIDTH - 1)^2}$$

O_TILE_WIDTH	8	16	32	64
MASK_WIDTH = 5	11.1	16	19.7	22.1
MASK_WIDTH = 9	20.3	36	51.8	64

- Tile size has significant effect on of the memory bandwidth reduction ratio.
- This often argues for larger shared memory size

Credits

- These slides report material from:
 - NVIDIA GPU Teaching Kit

Books

- Programming Massively Parallel Processors: A Hands-on Approach, D. B. Kirk and W-M. W. Hwu, Morgan Kaufman - Chapt. 8