

CS59200: Final Project Report

Vidush Singhal and Vedant Paranjape

Section 1: Summary

OpenMP [1] is a language agnostic application programming interface (API). It allows the programmer to annotate certain regions in the program using pragmas. These pragmas are directives that inform the compiler to generate certain type of code. For example, a *parallel for* pragma allows the programmer to parallelize trivial loops without dependencies.

As an example consider listing 1 that can execute the inner computation in parallel across 4 threads on a CPU.

Listing 1: OpenMP CPU Example

```
1 #include <omp.h>
2 #include <stdio.h>
3
4 #pragma omp parallel for num_threads(4)
5 for (i = 1; i < SIZE; i++) {
6     a[i] = b[i] + c[i];
7 }
```

Listing 2: OpenMP Offload Example

```
1 #include <omp.h>
2 #include <stdio.h>
3
4 #pragma omp target teams num_teams(3)
5 for (i = 0; i < 12; i++) {
6     a[i] = b[i] + c[i];
7 }
```

Each thread compute the i th element of the array a , using, the i th element of arrays b and c . This reduces the burden on the programmer by not worrying about target dependent instructions for parallelization for instance, the OpenMp runtime can handle efficient parallelism.

The goal with OpenMP offload is to "offload", that is, send computation to the device from the host machine. The programmer can use directives to annotate certain parts of the program that are computationally intensive. These parts can benefit from execution on the GPU device rather than the host machine. This can allow the user to speed up their application.

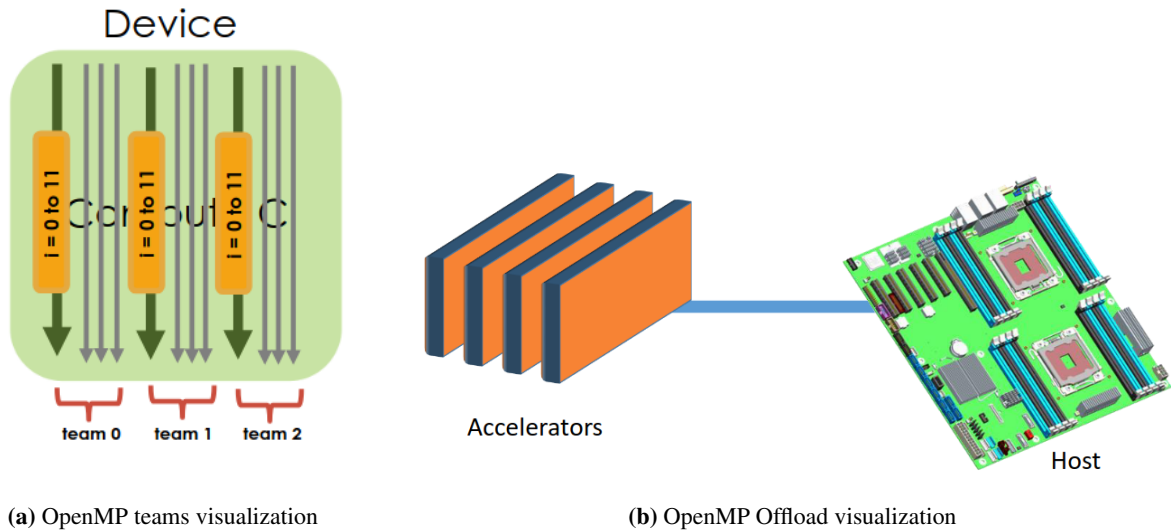
A compiler pragma like *pragma omp target* [2, 3] allows the programmer to automatically send the computation to devices like the GPU. The compiler will automatically generate code that will copy the data from the host to the device and vice-versa once the computation finishes. The device and the host have both have a copy of the same data and the host would not be able to access that data when it is executing on the device.

Additional pragmas like *teams* allows to form groups of threads for mapping the execution on GPUs, similar to warps in cuda.

For example, listing 2 shows an example where the loop executes in teams of 3 with each team having 4 threads. The visualization for such a case is shown in figure 1b

The following is a description of all the offload related OpenMp pragmas and their description.

- **omp declare target:** Specifies functions and variables that are created or mapped to a device.



- **omp declare variant:** Identifies a variant of a base procedure and specifies the context in which this variant is used.
- **omp dispatch:** Determines if a procedure variant is called for a given procedure.
- **omp distribute:** Specifies that the iterations of one or more loops should be distributed among the initial threads of all thread teams in a league.
- **omp interop:** Identifies a foreign runtime context and identifies runtime characteristics of that context, enabling interoperability with it.
- **omp requires:** Lists the features that an implementation must support so that the program compiles and runs correctly.
- **omp target enter data:** Specifies that variables are mapped to a device data environment.
- **omp target exit data:** Specifies that variables are unmapped from a device data environment.
- **omp teams:** Creates a league of thread teams inside a target region to execute a structured block in the initial thread of each team.

In this project, we want to see how efficient offloading C++17 parallel STL functions would be on the GPU [4]. STL [5] provides various algorithms and we can self implement them for testing performance then testing. We can also try to then automatically do it using a compiler like LLVM.

Section 2: Proof of concept and intended results

To establish a proof of concept, we decided to hand implement some of the standard algorithms and compare their performance in different configurations. For each example, we compare the performance of the GPU offloaded version and CPU OpenMP version and the CPU fully serial version. For this project we compared the performance of the following algorithms which we specifically handwrote using OpenMp pragmas.

- **copy:** This program is similar to STL copy such that it takes an input array and copied it to an output location, i.e, another array.

- **fill**: This program is similar to the STL fill function, such that it takes an input array and the value that should be used to fill the array. This means that every element in the array will be set to this value.
- **transform_reduce**: This program takes an array and a unary function to be applied to each array. It then does function application using the specified function for each array element and also reduces all the values of the array into a single value to be returned
- **for_each**: This program takes an array and then iterates over each element of the array and applies that function onto each element of the array.
- **any_of**: This program takes an array and then iterates over each element of the array to check if any one of the element matches the result of the function. If it does, it will set the result variable to true.
- **count**: This program takes an array and then counts each element of the array only if it is equal to the result of the function. It then returns the count as the result variable.

Section 3: Results

We were able to make LLVM generate code for GPU using OpenMP pragmas. The original plan was to use libc++ source to automatically transform to GPU code, but due to some upstream bugs in libc++ it was not possible to due in this time frame. So instead we implemented a custom version of libc++ using OpenMP pragmas. We wrote simple benchmarks which used these OpenMP pragma enabled libc++ functions to benchmark the efficiency of our implementation. Using the LLVM Compiler with Nvidia (nvptx) codegen backend enabled, we were able to automatically generate OpenMP code for Nvidia GPUs. For validation purposes we were also able to get it running on AMD GPUs. Our custom libc++ implementation supports `std::copy`, `std::fill`, `std::transform_reduce`, `std::for_each`, `std::any_of` and `std::count`. These implementation use OpenMP pragmas to parallelize code on GPUs and CPUs alike. Our implementation supports offloading to Nvidia, AMD GPUs and CPUs.

Section 3.1: Automatically generating GPU Code using LLVM

As mentioned in the previous section we were able to implement few algorithms, in this section we will explain in depth about `std::count()` function which we implemented. All the other algorithms are similar in terms of GPU code generation, but for the brevity of this report we stick to `std::count`. All the other benchmarks are uploaded to the github repository.

In the code listing 2, we show our implementation of `std::count` which has been modified to have a tiled loop and use OpenMP pragmas to speedup the code. The pragmas are device agnostic, and it can be deployed to CPU or GPUs alike. In the code listing 7 we use the `std::count` function defined below to implement a small benchmark to compare the performance of our implementation across CPUs and GPUs. We use the following commands to compile this for Nvidia GPU and AMD CPU. Clang has an option (`-fopenmp-targets=`) to specify which target to compile this code for, so if we specify Nvidia GPU it will generate code for it. The compile commands we used are shown below in the code listing. As mentioned before, we use `-fopenmp-targets` to specify `nvptx64-nvidia-cuda` to target Nvidia GPUs. LLVM is able to support a wide variety of targets as LLVM supports many compiler code gen backends like Nvidia's PTX, AMD's HIP, etc.

```

1 # Compile command for GPU
2 clang++ -O3 -fopenmp -fopenmp-targets=nvptx64-nvidia-cuda count_offload.cpp -o
   count_offload
3
4 # Compile command for CPU
5 clang++ -O3 -fopenmp count_offload.cpp -o count_omp

```

```

1 void count(long *array, int size, int *result, long (*func_ptr)(long)) {
2     #pragma omp target teams distribute parallel for map(tofrom: array[0:size])
3     for(int i = 0; i < (size/5000); i++) {
4         for(int j = 0; j < 5000; j++) {
5             if (array[i*5000 + j] == func_ptr(array[i*5000 + j])) {
6                 *result += *result + 1;
7             }
8         }
9     }
10 }

```

Figure 2: Figure showing implementation of std::count using OpenMP pragmas.

In the implementation of std::count, we tiled the for loop so that it can better perform on GPUs. It checks if each array element is equal to the result of the function passed, and if it is, it will increment the result. We added omp target distribute pragmas at the start of the for loop so that OpenMP is able to parallelize it. The benchmark simply defines a large array, fills it with dummy data and then invokes std::count with this data. We see significant performance gains when run on a GPU vs on a CPU.

Section 3.2: Offloading OpenMP code to GPUs using OpenMPOffload

To offload OpenMP code to GPUs, we need to build LLVM to support generating code for NVPTX. This is trivial to do, as LLVM support NVPTX as a code generation backend. The CMake command for configuring LLVM is depicted in code listing 3. In that we enable NVPTX backend by adding it to **DLLVM_TARGETS_TO_BUILD="NVPTX;X86"**.

We then use LLVM to compile our benchmark codes using the following commands:

```

1 # Compile command for GPU
2 clang++ -O3 -fopenmp -fopenmp-targets=nvptx64-nvidia-cuda count_offload.cpp -o
   count_offload

```

Section 3.3: Analysing Nvidia PTX code generated using LLVM OpenMP Offload

In this section we analyze the generated code to understand how LLVM generates code for the NVPTX backend and how it offloads the CUDA kernels to the GPU. We add the -S -emit-llvm command to dump the LLVM IR for the generated code. Similarly we also add -save-temps to dump all the intermediate files generated during compilation. The files created during compilation are shown in Figure 4. As you can see, there is a *.cubin file that is generated by LLVM, and there are other files as well (count_offload-openmp-nvptx64-nvidia-cuda.o) which contain the actual NVPTX assembly generated by LLVM. We have added a reduced snippet of the NVPTX assembly in code listing 9. The code has PTX assembly generated by the LLVM Compiler.

Now we explain the exact flow of compilation followed by LLVM to generate the PTX binary. We compile the std::count benchmark using the command given below. As a result the LLVM Compiler generates a CUDA kernel for the functions that use OpenMP pragma, so it internally converts the pragmas to CUDA kernel parameters. It will generate a wrapper function which will setup the CUDA kernel parameters like size, warps, threads, etc and then calls the actual implementation of the CUDA kernel. The wrapper function can be seen in code listing 6, on Line 7, it allocates a target kernel object and following that it sets the target kernel parameters. On Line 24, there is a call to __tgt_target_kernel() which is a generic OpenMP function to

```

1  cmake -G Ninja -DCMAKE_BUILD_TYPE=Release \
2      -DCLANG_INCLUDE_TESTS=ON \
3      -DCMAKE_INSTALL_PREFIX="/local/scratch/a/paranjav/llvm-project/install-
        release" \
4      -DCMAKE_CXX_STANDARD=17 \
5      -DCMAKE_C_COMPILER=gcc \
6      -DCMAKE_CXX_COMPILER=g++ \
7      -DLIBOMPTARGET_ENABLE_DEBUG=ON \
8      -DCMAKE_EXPORT_COMPILE_COMMANDS=ON \
9      -DCMAKE_CXX_LINK_FLAGS="-Wl,-rpath_, $LD_LIBRARY_PATH" \
10     -DLLVM_ENABLE_ASSERTIONS=ON \
11     -DLLVM_ENABLE_LIBCXX=ON \
12     -DLIBC_GPU_BUILD=OFF \
13     -DLLVM_LIBC_FULL_BUILD=ON \
14     -DOPENMP_ENABLE_LIBOMPTARGET=ON \
15     -DLIBOMPTARGET_NVPTX_ENABLE_BCLIB=true \
16     -DLLVM_TARGETS_TO_BUILD="NVPTX;X86" \
17     -DLLVM_ENABLE_PROJECTS="clang;lld" \
18     -DLLVM_ENABLE_RUNTIMES="openmp;offload;libcxxabi;libcxx;libunwind" \
19     -DLIBCXX_ENABLE_THREADS=ON \
20     /local/scratch/a/paranjav/llvm-project/llvm

```

Figure 3: CMake command

call a target kernel. The target kernel can be a GPU kernel or a CPU kernel, it entirely depends on the target specified. At this stage LLVM generates platform agnostic code.

The LLVM IR for the function called in `_tgt.target_kernel` is in code listing 8. The LLVM IR is of a nested loop, and thus it matches the `std::count` implementation. On Line 31, it calls `_kmprc_fork_call` which calls the NVPTX implementation of the function. The assembly for this function is listed in code listing 9, as you can see on Line 24, the actual PTX assembly implementation of the count function starts. This NVPTX assembly dump was generated from the `*.o` file using `cuobjdump` which is an disassembly utility provided by CUDA. The said assembly was generated by LLVM NVPTX backend.

To reiterate, the steps for compilation are as follows:

- Compile OpenMP code to OpenMP kernel LLVM IR.
- Compile the OpenMP kernel LLVM IR to the target backend specified during compilation.
- Generate CUDA wrapper code for the OpenMP LLVM IR.
- Compile the OpenMP kernel to NVPTX backend using LLVM Compiler.
- Link the files and generate a unified assembly.

```

1  clang++ -O3 -fopenmp -fopenmp-targets=nvptx64-nvidia-cuda count\_offload.cpp -o
        count\_offload -S -emit-llvm

```

Section 3.4: Running the binary generated by LLVM

Now that LLVM has generated offload binary for our `std::count` benchmark, we can simply run it and it gets offloaded to the Nvidia GPU. The output of the program is shown below.

```

paranjav@io:~/offload_stl_examples/llvm-ir$ ls -ltr
total 4656
-rw-r--r-- 1 paranjav ecnuser 1355 Dec 14 22:56 count_offload.cpp
-rw-r--r-- 1 paranjav ecnuser 987815 Dec 14 22:56 count_offload-host-x86_64-unknown-linux-gnu.ii
-rw-r--r-- 1 paranjav ecnuser 26796 Dec 14 22:56 count_offload-host-x86_64-unknown-linux-gnu.bc
-rw-r--r-- 1 paranjav ecnuser 1069147 Dec 14 22:56 count_offload-openmp-nvptx64-nvidia-cuda.ii
-rw-r--r-- 1 paranjav ecnuser 267212 Dec 14 22:56 count_offload-openmp-nvptx64-nvidia-cuda.bc
-rw-r--r-- 1 paranjav ecnuser 36074 Dec 14 22:56 count_offload-openmp-nvptx64-nvidia-cuda.s
-rw-r--r-- 1 paranjav ecnuser 20328 Dec 14 22:56 count_offload-openmp-nvptx64-nvidia-cuda.o
-rw-r--r-- 1 paranjav ecnuser 20472 Dec 14 22:56 count_offload-openmp-x86_64-unknown-linux-gnu.out
-rw-r--r-- 1 paranjav ecnuser 97621 Dec 14 22:56 count_offload-host-x86_64-unknown-linux-gnu.s
-rw-r--r-- 1 paranjav ecnuser 30528 Dec 14 22:56 count_offload-host-x86_64-unknown-linux-gnu.o
-rw-r--r-- 1 paranjav ecnuser 20328 Dec 14 22:56 count_offload-host-x86_64-unknown-linux-gnu-nvptx64-nvidia-cuda-sm_89.o
-rw-r--r-- 1 paranjav ecnuser 19720 Dec 14 22:56 count_offload.nvptx64.sm_89.img
-rw-r--r-- 1 paranjav ecnuser 20328 Dec 14 22:56 count_offload-host-x86_64-unknown-linux-gnu-nvptx64-nvidia-cuda-sm_89.cubin
-rw-r--r-- 1 paranjav ecnuser 22188 Dec 14 22:56 count_offload.openmp.image.wrapper.bc
-rw-r--r-- 1 paranjav ecnuser 22280 Dec 14 22:56 count_offload.openmp.image.wrapper.o
-rwxr-xr-x 1 paranjav ecnuser 39752 Dec 14 22:57 count_offload
-rw-r--r-- 1 paranjav ecnuser 699427 Dec 14 22:58 count_offload.asm
-rw-r--r-- 1 paranjav ecnuser 1166144 Dec 14 23:02 count_offload-openmp-nvptx64-nvidia-cuda.ll
-rw-r--r-- 1 paranjav ecnuser 169957 Dec 14 23:06 cuda-kernel.asm
paranjav@io:~/offload_stl_examples/llvm-ir$

```

Figure 4: OpenMP Offload file dump

```

1 ./count_offload
2 Execution time: 2.00091 seconds
3 Count: 0

```

We ran all our benchmarks on Nvidia GPU, CPU OpenMP offload and vanilla CPU runtime. The results are explained in the next section.

Section 3.5: Issues with upstream LLVM libc++

The initial plan was to offload LLVM’s libc++ implementation to Nvidia GPUs. We spent several days trying to get it working, but due to some incorrect CMake configurations, LLVM always generated code for Intel TBB. It was unable to change the backend to say Nvidia for the libc++ implementation. We have filed a upstream issue for the same as addressed in the Future works section. We will port our implementation to LLVM’s libc++ as soon as the upstream issue is resolved. But we believe that our implementation gives a proof of concept that standard C++ algorithms can be effortlessly offloaded to GPUs.

Section 4: Evaluation

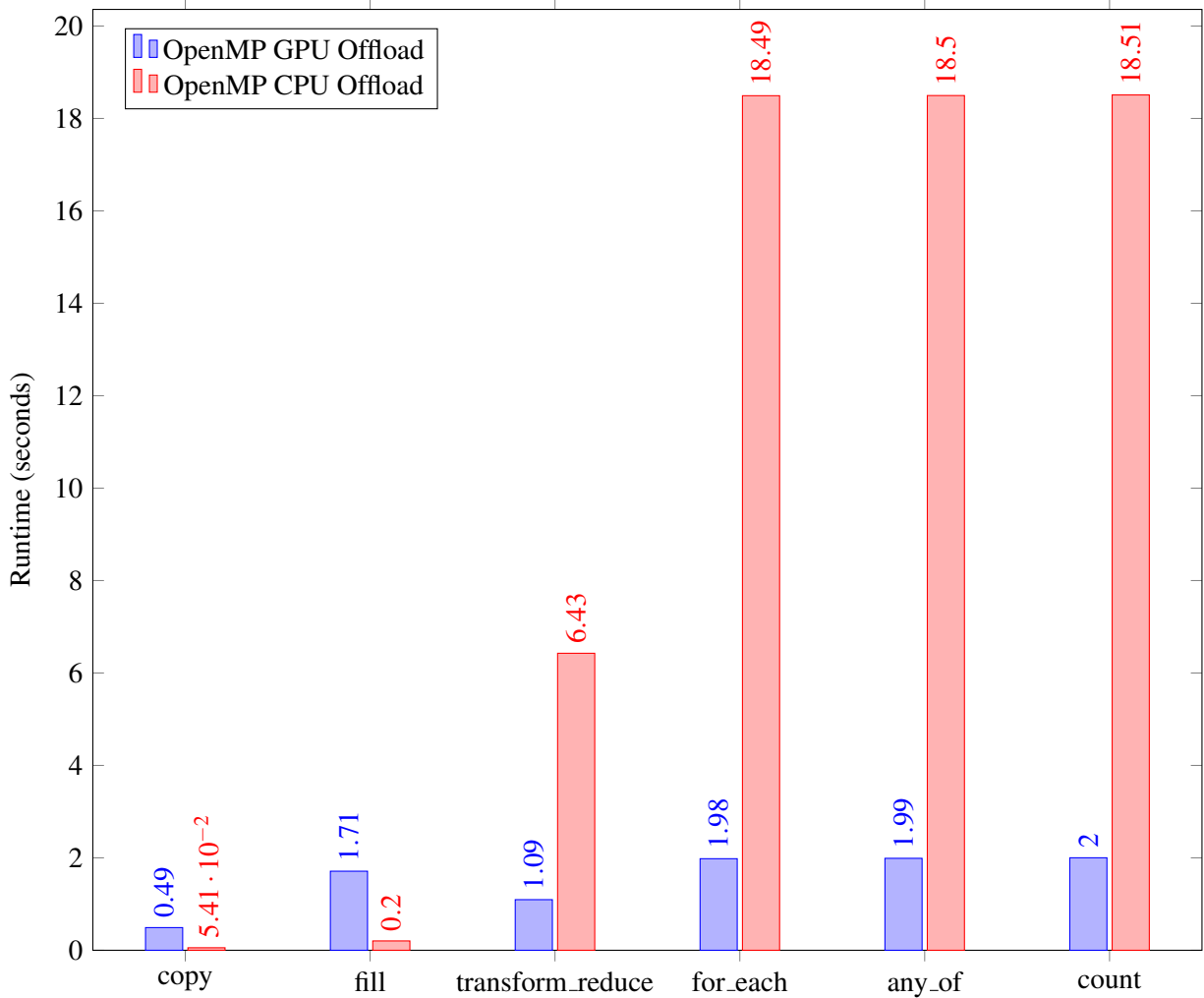
We performed our test on an AMD Ryzen Threadripper PRO 5965WX 24-Cores Processor with an NVIDIA GeForce RTX 4070 Ti GPU. The L1D and L1I cache size was 768K, the L2 cache size was 12M and the L3 cache size was 128M. The CPU max clock speed was 7.021 GHz.

For offloading support, we build LLVM using the CMake command shown in listing 3 command. Table 2 shows the results of our experiments. We now provide explanation of the results and why the offloading is slow for copy and fill but faster for transform_reduce and for_each. We also show the compile commands that we used for each program in listings 5. The programs used can be found in the linked public github repository https://github.com/vidsinghal/offload_stl_examples.

Offloading involves a significant overhead. The memory for the input and output arrays must be allocated on the device and the output arrays must be copied back from the device to the host. If the arrays are very large this may take significant time as well. The cost of the overhead is reduced is the computation on the arrays is computationally intensive. If the computation is large and parallel, the significant number of threads on GPUs can reduce the time it takes to do the computation thus reducing the overall time of the program. On the other hand, if the computation is lightweight, the overhead to copy the arguments will dominate thus not providing enough speedup.

This is exactly what we see in our results shown in Table 2 and Figure 1. Copy and Fill are lightweight traversals over the array. The computational intensity is low thus copying is an overkill. The offload versions are slower than the CPU OpenMP versions for this reason. Even the serial CPU version is slightly faster than the offloaded versions. However, for transform, any_of, count and foreach, since the computational intensity is high, we see significant speedups with the offloaded version than the OpenMP CPU and serial CPU version.

Table 1: Bar plot for Runtime of benchmarks



Section 4: Conclusion

We have shown through some simple hand crafted examples that offloading some STL function can be beneficial as it may speed up the computation by executing them on the GPU. However, that's only the case if the function is computationally intensive. For lightweight functions, there is a slowdown from offloading that function. We will work on the cost model that can infer which functions may be beneficial to offload. For the future work, we plan to offload the STL function automatically using the LLVM compiler. We are in process to do so and are dealing with compiler bugs atm. With this ability we can potentially speed up the user's application significantly automatically using OpenMP offload.

Table 2: Performance comparison of different handwritten STL functions.

Benchmark	OpenMP Offload	OpenMP CPU	CPU
copy	0.490246 (s)	0.0541481 (s)	0.256001 (s)
fill	1.71053 (s)	0.202087 (s)	1.2373 (s)
transform_reduce	1.09482 (s)	6.42661 (s)	148.325 (s)
for_each	1.9809 (s)	18.4916 (s)	600.848 (s)
any_of	1.98938 (s)	18.4952 (s)	620.799 (s)
count	2.00091 (s)	18.5082 (s)	600.873 (s)

```

1 [singhal2@linux -5: algorithm_tests]$ make all
2 clang++ -O3 -fopenmp --offload-arch=native \
3     for_each_offload.cpp -o for_each_offload
4 clang++ -O3 -fopenmp --offload-arch=native \
5     transform_reduce.cpp -o transform_reduce_offload
6 clang++ -O3 -fopenmp --offload-arch=native \
7     fill_offload.cpp -o fill_offload
8 clang++ -O3 -fopenmp --offload-arch=native \
9     copy_offload.cpp -o copy_offload
10 clang++ -O3 -fopenmp for_each_offload.cpp -o for_each_omp
11 clang++ -O3 -fopenmp transform_reduce.cpp \
12     -o transform_reduce_omp
13 clang++ -O3 -fopenmp fill_offload.cpp -o fill_omp
14 clang++ -O3 -fopenmp copy_offload.cpp -o copy_omp
15 clang++ -O3 for_each_offload.cpp -o for_each_cpu
16 clang++ -O3 transform_reduce.cpp -o transform_reduce_cpu
17 clang++ -O3 fill_offload.cpp -o fill_cpu
18 clang++ -O3 copy_offload.cpp -o copy_cpu

```

Figure 5: Figure showing compile commands that were used.

Section 5: Future Work

Now that we have proved that offloading some parallel STL algorithms to GPUs using OpenMP offloading can be beneficial, we will implement code to do this automatically. For this we will use LLVM, we will parse the C++ STL functions and write a custom OpenMP backend for these functions such that they can be automatically offloaded to the GPU. There is some ongoing work on this in the form of this PR: <https://github.com/llvm/llvm-project/pull/116869>. We tried to build the PR with the OpenMP backend for PSTL. However, it seems like there is a bug with their implementation and the build defaults to the TBB backend. As it can be seen in the PR, we posted a comment with the error and the possible explanation of the error and we are in progress to fix these errors. With the help of the contributors we should be able to get the initial proof of concept compiler working. However, as our results show, this is only beneficial for computationally intensive functions. Thus, we would probably add a static cost estimation of the function and only generate offloaded code for parallel STL algorithms that are computationally intensive, that is, costly according to our cost model.

References

- [1] Accessed: 2024-12-11. URL: <https://www.openmp.org/wp-content/uploads/2021-10-20-Webinar-OpenMP-Offload-Programming-Introduction.pdf>.
- [2] Accessed: 2024-12-11. URL: https://www.olcf.ornl.gov/wp-content/uploads/2021/08/ITOpenMP_Day1.pdf.
- [3] Accessed: 2024-12-11. URL: <https://discourse.llvm.org/t/gsoc-2024-offloading-libcxx/77238>.
- [4] Accessed: 2024-12-11. URL: <https://llvm.org/OpenProjects.html#offload-libcxx>.
- [5] Accessed: 2024-12-11. URL: <https://www.geeksforgeeks.org/the-c-standard-template-library-stl/>.

Appendix A: Code Listings

```
1 ; Function Attrs: nounwind uwtable
2 define dso_local void @_Z5countPliPiPF1lE(ptr noundef %array, i32 noundef %
   size, ptr noundef %result, ptr noundef %func_ptr) local_unnamed_addr #3 {
3 entry:
4   %.offload_baseptrs = alloca [4 x ptr], align 8
5   %.offload_ptrs = alloca [4 x ptr], align 8
6   %.offload_sizes = alloca [4 x i64], align 8
7   %kernel_args = alloca %struct.__tgt_kernel_arguments, align 8
8   %1 = getelementptr inbounds i8, ptr %.offload_baseptrs, i64 8
9   ...
10  ...
11  store i32 3, ptr %kernel_args, align 8
12  %9 = getelementptr inbounds nuw i8, ptr %kernel_args, i64 4
13  store ptr %.offload_ptrs, ptr %11, align 8
14  ...
15  ...
16  %16 = getelementptr inbounds nuw i8, ptr %kernel_args, i64 64
17  call void @llvm.memset.p0.i64(ptr noundef nonnull align 8 dereferenceable
   (36) %16, i8 0, i64 36, i1 false)
18  %17 = call i32 @__tgt_target_kernel(ptr nonnull @3, i64 -1, i32 0, i32 0,
   ptr nonnull @.__omp_offloading_fd04_8135d223__Z5countPliPiPF1lE_155.
   region_id, ptr nonnull %kernel_args)
19  %.not = icmp eq i32 %17, 0
20  br i1 %.not, label %omp_offload.cont, label %omp_offload.failed
21
22 omp_offload.failed:                                ; preds = %entry
23   call void (ptr, i32, ptr, ...) @__kmpc_fork_teams(ptr nonnull @3, i32 4, ptr
   nonnull @.__omp_offloading_fd04_8135d223__Z5countPliPiPF1lE_155.
   omp_outlined, i64 %size.casted.sroa.0.0.insert.ext, ptr %array, ptr %
   func_ptr, ptr %result)
24   br label %omp_offload.cont
25
26 omp_offload.cont:                                ; preds = %omp_offload.
   failed, %entry
27   ret void
28 }
```

Figure 6: Figure showing LLVM IR of the OpenMP offload function

```

1  #include <stdio.h>
2  #include <omp.h>
3  #include <chrono>
4  #include <iostream>
5  #include <cassert>
6
7  void count(long *array, int size, int *result, long (*operation)(long));
8  long complex(long i);
9  #pragma omp declare target indirect to(complex)
10
11 int main() {
12     int size = 2e8;
13     long *array = (long*) malloc(sizeof(long) * size);
14     int val = 10;
15
16     #pragma omp parallel for
17     for(int i = 0; i < size; i++){
18         array[i] = val;
19     }
20
21     long (*func_ptr)(long) = complex;
22     int result = 0;
23     auto start = std::chrono::high_resolution_clock::now();
24
25     count(array, size, &result, func_ptr);
26
27     auto end = std::chrono::high_resolution_clock::now();
28
29     std::chrono::duration<double> duration = end - start;
30     std::cout << "Execution time: " << duration.count() << " seconds" << std::
        endl;
31
32     std::cout << "Count: " << result << std::endl;
33 }
34
35 long complex(long i){
36     long j = i;
37     j = j * 1;
38     j = j + 1;
39     j = j * 10;
40     j = j / 4;
41     j = j + 1;
42
43     for(int k = 1; k < 10000; k++){
44         j = j * k;
45     }
46
47     return j;
48 }

```

Figure 7: Figure showing benchmark utilizing std::count using OpenMP pragmas.

```

1      ; Function Attrs: alwaysinline norecurse nounwind uwtable
2      define internal void @__omp_offloading_fd04_8135d223__Z5countPliPiPF11E_155.
      omp_outlined(ptr noalias nocapture noundef readonly %.global_tid., ptr
      noalias nocapture readnone %.bound_tid., i64 noundef %size, ptr noundef %
      array, ptr noundef %func_ptr, ptr noundef %result) #4 {
3      entry:
4          %.omp.comb.lb = alloca i32, align 4
5          %.omp.comb.ub = alloca i32, align 4
6          %.omp.stride = alloca i32, align 4
7          %.omp.is_last = alloca i32, align 4
8          %size.addr.sroa.0.0.extract.trunc = trunc i64 %size to i32
9          %cmp = icmp sgt i32 %size.addr.sroa.0.0.extract.trunc, 4999
10         br i1 %cmp, label %omp.precond.then, label %omp.precond.end
11
12         omp.precond.then:                                ; preds = %entry
13         %div7 = udiv i32 %size.addr.sroa.0.0.extract.trunc, 5000
14         %sub3 = add nsw i32 %div7, -1
15         call void @llvm.lifetime.start.p0(i64 4, ptr nonnull %.omp.comb.lb) #6
16         ...
17         ...
18         %0 = load i32, ptr %.global_tid., align 4, !tbaa !9
19         call void @__kmpc_for_static_init_4(ptr nonnull @1, i32 %0, i32 92, ptr
      nonnull %.omp.is_last, ptr nonnull %.omp.comb.lb, ptr nonnull %.omp.
      comb.ub, ptr nonnull %.omp.stride, i32 1, i32 1)
20         ...
21         br i1 %cmp6.not8, label %omp.loop.exit, label %omp.inner.for.body.lr.ph
22
23         omp.inner.for.body.lr.ph:                        ; preds = %omp.precond.then
24         %size.casted.sroa.0.0.insert.ext = and i64 %size, 2147483647
25         br label %omp.inner.for.body
26
27         omp.inner.for.body:                               ; preds = %omp.inner.for.
      body.lr.ph, %omp.inner.for.body
28         %3 = phi i32 [ %cond, %omp.inner.for.body.lr.ph ], [ %8, %omp.inner.for.
      body ]
29         %.omp.iv.09 = phi i32 [ %2, %omp.inner.for.body.lr.ph ], [ %add, %omp.
      inner.for.body ]
30         ...
31         call void (ptr, i32, ptr, ...) @__kmpc_fork_call(ptr nonnull @3, i32 6,
      ptr nonnull @__omp_offloading_fd04_8135d223__Z5countPliPiPF11E_155.
      omp_outlined.omp_outlined, i64 %5, i64 %6, i64 %size.casted.sroa.0.0.
      insert.ext, ptr %array, ptr %func_ptr, ptr %result)
32         %7 = load i32, ptr %.omp.stride, align 4, !tbaa !9
33         %add = add nsw i32 %7, %.omp.iv.09
34         ...
35
36         omp.loop.exit:                                   ; preds = %omp.inner.for.
      body, %omp.precond.then
37         call void @__kmpc_for_static_fini(ptr nonnull @1, i32 %0)
38         br label %omp.precond.end
39
40         ...
41         ...
42     }

```

Figure 8: Figure showing LLVM IR of the OpenMP for loop which does the computation

```

1  //
2  // Generated by LLVM NVPTX Back-End
3  //
4
5  .version 8.2
6  .target sm_89
7  .address_size 64
8  ...
9  ...
10 ;
11 .weak .global .align 4 .u32 __omp_rtl_assume_teams_oversubscription;
12 .weak .global .align 4 .u32 __omp_rtl_assume_threads_oversubscription;
13 .visible .global .align 8 .u64 __omp_offloading_fd04_8135d223_complex_18 =
    _Z7complex1;
14 ...
15 ().__unnamed_2), generic(
    __omp_offloading_fd04_8135d223__Z5countPliPiPFllE_l55_dynamic_environment)
    };
16 // _ZN4omp5state9TeamStateE has been demoted
17 // _ZN4omp5state12ThreadStatesE has been demoted
18 .weak .global .align 8 .b8 __omp_rtl_device_memory_pool[16];
19 .weak .global .align 8 .b8 __omp_rtl_device_memory_pool_tracker[32];
20 .weak .global .align 4 .u32 __omp_rtl_debug_kind;
21 .weak .global .align 4 .u32 __omp_rtl_assume_no_thread_state;
22 ...
23 ...
24 // .weak __omp_offloading_fd04_8135d223__Z5countPliPiPFllE_l55 // -- Begin
    function __omp_offloading_fd04_8135d223__Z5countPliPiPFllE_l55
25 .weak .entry __omp_offloading_fd04_8135d223__Z5countPliPiPFllE_l55(
26     .param .u64 __omp_offloading_fd04_8135d223__Z5countPliPiPFllE_l55_param_0,
27     .param .u64 __omp_offloading_fd04_8135d223__Z5countPliPiPFllE_l55_param_1,
28     .param .u64 __omp_offloading_fd04_8135d223__Z5countPliPiPFllE_l55_param_2,
29     .param .u64 __omp_offloading_fd04_8135d223__Z5countPliPiPFllE_l55_param_3,
30 )
31 .maxntid 128, 1, 1 //
    @__omp_offloading_fd04_8135d223__Z5countPliPiPFllE_l55
32 {
33     .reg .pred      %p<70>;
34     .reg .b32       %r<144>;
35     ...
36     ...
37     setp.ne.s32     %p1, %r1, 0;
38     @%p1 bra        $L__BB1_2;
39 // %bb.1: // %if.then.i.i1
40     mov.b32         %r46, 0;
41     st.shared.v2.u32 [_ZN4omp5state9TeamStateE], {%r46, %r46};
42     st.shared.v2.u32 [_ZN4omp5state9TeamStateE+8], {%r46, %r46};
43     mov.b32         %r47, 1;
44     ...
45     ...
46 $L__BB1_8: // %if.then.i.i.i.i.i9
47     mov.u64         %rd217, %rd5;
48     bra.uni         $L__BB1_9;
49 $L__BB1_12: // %if.then.i.i.i.i.i.i.i19
50     mov.u64         %rd218, %rd7;
51     bra.uni         $L__BB1_13;
52 // -- End function
53 }

```

Figure 9: Figure showing NVPTX for the benchmark generated using LLVM