



Optimizing Deep Learning Accelerators for Real-Time Triggering in the Cherenkov Telescope Array: A Comparative Study of FPGA, GPU, and CPU Performance

Author:

Tamer VAN TEEFFELN

Supervisors:

Dr. Carlos ABELLAN

Dr. Iaroslava BEZSHYIKO

29/08/2024

Contents

	Page
1 Introduction	3
2 Methods	4
2.1 FPGA Inference Using Intel's Deep Learning Accelerator	4
2.2 FPGA Compilation and Optimization with Quartus	4
2.3 GPU and CPU Inference	5
2.4 Commands	5
2.4.1 Configuration	5
2.4.2 Model conversion	5
2.4.3 Compilation	6
2.4.4 Programming the FPGAs	7
2.4.5 Inference	7
2.4.6 Compiling Bitstream for Optimized Architectures	8
2.4.7 Power usage	9
2.5 FPGA Throughput, Latency and Area Estimate Scripts	9
2.6 GPU and CPU Scripts	13
3 Results and Discussion	20
3.1 Model dependence	20
3.2 Architecture dependence	23
3.3 Comparison with ResNet-50	26
3.4 Power usage and area estimates	27
4 Conclusion	29
Bibliography	31

1 Introduction

The Cherenkov Telescope Array (CTA) represents a groundbreaking international endeavor in the field of very high-energy gamma-ray astronomy, aiming to establish the world's most sophisticated ground-based observatory [1]. This ambitious project comprises an extensive network of Imaging Atmospheric Cherenkov Telescopes (IACTs) strategically distributed across multiple sites. The primary objective of these IACTs is to detect and analyze the faint Cherenkov light flashes generated when high-energy gamma rays interact with the Earth's atmosphere. Through this advanced detection system, the CTA seeks to provide unprecedented insights into some of the most energetic phenomena in the cosmos, including supernovae, black holes, and dark matter [1].

A critical component of the CTA's operational framework is its multi-tiered triggering system, which is responsible for discriminating valid Cherenkov events and initiating data recording processes. The triggering mechanism operates on two levels: local and global. At the local level, individual telescopes employ sophisticated algorithms to identify concentrations of Cherenkov light in their cameras. This is achieved either through thresholding pixel signals and detecting coincidences in adjacent pixels or by summing signals from pixel clusters and comparing the aggregate to a pre-determined threshold [2, 3]. Events that satisfy the local trigger criteria are temporarily stored, and a notification is transmitted to a global array trigger operating at the software level. The global trigger then analyzes potential coincidences across multiple telescopes, and upon detection, initiates the transmission of event information for further processing [4]. This intricate triggering scheme enables the CTA to operate with lower telescope trigger thresholds while mitigating the risk of readout system saturation [2–4].

Recent advancements in the field of deep learning have presented novel opportunities for enhancing the CTA's triggering system. Convolutional Neural Networks (CNNs) have demonstrated significant potential in improving event selection through more sophisticated and accurate data filtering techniques [5]. Deep learning models, such as those developed within the CTLearn framework, are trained on simulated event reconstructions and have been successfully employed to autonomously extract features from raw Cherenkov images, thereby improving the accuracy of event detection [5]. The integration of these deep learning techniques into the triggering process holds promise for reducing false positives and enhancing the precision of gamma-ray event detection [6].

Traditionally, Graphics Processing Units (GPUs) have been the preferred hardware platform for deep learning inference due to their capacity to efficiently handle the parallel computations required by neural networks [7]. Their architecture is particularly well-suited for the matrix operations that form the core of CNNs, making them ideal for the real-time processing demands of projects such as the CTA [8]. However, despite the widespread adoption of GPUs, Field-Programmable Gate Arrays (FPGAs) have emerged as a compelling alternative, offering advantages such as lower power consumption, customizable architecture, and potential for real-time processing [9]. FPGAs present a promising solution for deploying CNNs within the CTA's triggering system, particularly in resource-constrained environments where energy efficiency and low latency are paramount.

This study aims to compare the performance of CNNs implemented on FPGAs, GPUs, and CPUs within the context of the CTA's final triggering system. The performance metrics of interest are throughput—the rate at which data can be processed—and latency—the delay between input and processing output. By evaluating these architectures based on these key metrics, this research seeks to identify the most effective hardware platform for deploying deep learning models in the CTA's real-time, high-stakes environment.

2 Methods

This section outlines the tools and methodologies used to evaluate the performance of deep neural networks on FPGA and GPU platforms, specifically focusing on throughput and latency metrics. The corresponding code as well as commands are also given.

2.1 FPGA Inference Using Intel's Deep Learning Accelerator

For FPGA inference, Intel's deep learning accelerator (DLA) plugin for the FPGA AI Suite optimizer was used to deploy neural networks and measure performance. Two tools within this plugin were utilized:

- The `dla_compiler` tool was used to estimate throughput on various FPGA architectures using randomly generated data. This provided a theoretical performance benchmark for each architecture.
- The `dla_benchmark` tool was employed to measure real-time throughput on the FPGA architectures, also using randomly generated data, offering empirical validation of the compiler's estimates.

2.2 FPGA Compilation and Optimization with Quartus

Intel's Quartus software was employed for the bitstream compilation and optimization of different FPGA architectures using Intel's FPGA AI Suite. The architectures used and their corresponding FPGA boards were:

- **Arria 10 (A10):**
 - `A10_Generic.arch`
 - `A10_Performance.arch`
 - `A10_FP16_Generic.arch`
 - `A10_FP16_Performance.arch`
- **Agilex 7 (AGX7):**
 - `AGX7_Generic.arch`
 - `AGX7_Performance.arch`
 - `AGX7_FP16_Generic.arch`
 - `AGX7_FP16_Performance.arch`

It should be noted that the architectures without the FP16 specification indicate an FP11 precision instead and are thus expected to yield higher throughput. For the Arria 10 FPGA, both compilation and real-time inference were performed. The Agilex 7 FPGA, being in its testing phase and not yet available for industrial use, was only subjected to compilation for performance estimation.

Furthermore, optimization for each FPGA was performed in accordance to the maximum resources available. For Arria 10, the maximum available resources are 427,200 Adaptive Logic Modules (ALMs), 2,713 M20K memory blocks, and 1,518 DSP blocks. For the Agilex 7 FPGA, the AFG014 variant has 487,200 ALMs, 7,110 M20K blocks, and 4,510 DSP blocks, while the AGM039 variant offers 1,305,600 ALMs, 18,960 M20K blocks, and 12,300 DSP blocks. The assumed maximum core frequencies were set at 300 MHz for Arria 10 and 600 MHz for Agilex 7.

2.3 GPU and CPU Inference

Inference was conducted using an NVIDIA L40S GPU and an AMD EPYC 9254 CPU. To ensure the reliability of the results, latency was measured by averaging the processing time of over 10,000 samples, while throughput was calculated by averaging the rate at which these samples were processed.

2.4 Commands

2.4.1 Configuration

To execute the FPGA scripts effectively, it is necessary to configure the software environment by initializing the required tools and setting up the appropriate environment variables. The following commands must be run sequentially to ensure the proper configuration:

```
1 source /opt/openvino/openvino_2022.3.1_env/bin/activate
2 source /opt/intel/inteldevstack/init_env_user.sh
3 unset python_version
4 source /opt/intel/openvino_2022/setupvars.sh
5 source /opt/intel/fpga_ai_suite_2024.1/dla/bin/init_env.sh
6 source $HOME/coredla_work/coredla_work.sh
```

2.4.2 Model conversion

Before deploying the models on the FPGA, it is necessary to convert them to a format compatible with OpenVINO. The following Python script accomplishes this by converting the models and setting the appropriate tensor layouts:

```
1 import openvino as ov
2
3 dirname = "ctlearn_model"
4
5 # Specify that batch size is 1 because dynamic tensor sizes are not supported on
   FPGA
6 ov_model = ov.convert_model(dirname, input=[1,30,30,5])
7
8 input_name = list(ov_model.inputs[0].names)[0]
9 prep = ov.preprocess.PrePostProcessor(ov_model)
10
```

```

11 # Specify channel layout---batch (N) is 1, (H)eight is 30, (W)idth is 30, (C)
    hannels is 5
12 prep.input(input_name).tensor().set_layout(ov.Layout("NHWC"))
13
14 ov_model = prep.build()
15
16 # OpenVINO 2023.3 default behavior is to compress weights to fp16. We specify
    that we do not want that
17 ov.save_model(ov_model, 'ctlearn_model.xml', compress_to_fp16=False)

```

2.4.3 Compilation

To estimate the throughput and area usage of the FPGA, the `dla_compiler` tool is employed. This tool provides a comprehensive analysis of the performance and resource utilization based on the network model specified. The basic command to perform this estimation is as follows:

```

1 dla_compiler --march "${ARCH_PATH}/${arch}" --network-file "$network_file" \
2 --foutput-format=open_vino_hetero --o "$output_file" --batch-size=1 \
3 --fanalyze-area --fanalyze-performance $opt_params

```

In this command:

- `--march "${ARCH_PATH}/${arch}"` specifies the architecture file that defines the FPGA hardware configuration.
- `--network-file "$network_file"` indicates the path to the network model file that will be analyzed.
- `--foutput-format=open_vino_hetero` sets the output format to be compatible with OpenVINO for heterogeneous execution.
- `--o "$output_file"` defines the output file where the compiled model will be saved.
- `--batch-size=1` configures the batch size for the model, which is set to 1 in this case.
- `--fanalyze-area` enables the analysis of area utilization on the FPGA.
- `--fanalyze-performance` activates the performance analysis to estimate the throughput of the FPGA.
- `$opt_params` includes any additional optional parameters such as optimizations.

In scenarios where one desires to fully utilize the maximum resources available on the FPGA, optional parameters can be added to the command. These parameters specify the allocation of Adaptive Logic Modules (ALMs), M20K memory blocks, and DSP blocks, as well as additional optimization settings. The command with these optimizations is given below:

```

1 --gen-arch --mmax-resources=$ALMs,$M20Ks,$DSPs --gen-min-sb=2048 \
2 --mmax-resources-alm-util=75 --fassumed-fmax-core=300

```

Including these parameters not only optimizes the performance but also generates an optimized architecture that is tailored to the specific resource constraints and performance goals of the FPGA deployment.

2.4.4 Programming the FPGAs

In order to perform inference on the FPGA, it is necessary to program the bitstream onto the FPGA device first. This ensures that the FPGA is configured correctly for the specific design example. Below are the detailed steps to program the FPGA for different devices.

To program the DE10-Agilex-B2E2 board required for the PCIe-based design example for Agilex™ 7 devices, one must:

1. Build the runtime with the following commands:

```
1 cd $COREDLA_WORK/runtime
2 rm -rf build_Release
3 ./build_runtime.sh -de10_agilex
```

2. Program the bitstream onto the FPGA device with the following commands:

```
1 jtagdir=$COREDLA_WORK/runtime/build_Release/fpga_jtag_reprogram
2 bitsdir=$COREDLA_WORK/demo/bitstreams
3 $jtagdir/fpga_jtag_reprogram $bitsdir/AGX7_Performance.sof
4 curarch=$COREDLA_ROOT/example_architectures/AGX7_Performance.arch
```

To program the Intel® PAC with Arria® 10 GX FPGA hardware required for the PCIe-based design example for Arria® 10 devices, we use the following steps:

1. Build the runtime with the following commands:

```
1 cd $COREDLA_WORK/runtime
2 rm -rf build_Release
3 ./build_runtime.sh
```

2. Program the bitstream onto the FPGA device with the following commands:

```
1 fpgaconf -v $COREDLA_WORK/demo/bitstreams/A10_Performance.gbs
2 curarch=$COREDLA_ROOT/example_architectures/A10_Performance.arch
```

2.4.5 Inference

Once the FPGA has been programmed with the appropriate bitstream, inference can be performed using the `dla_benchmark` tool. This tool allows users to benchmark the performance of neural network models on the FPGA by specifying various parameters related to batch size, model path, and device architecture. The command to run inference is as follows:

```
1 $COREDLA_WORK/runtime/build_Release/dla_benchmark/dla_benchmark \
2 -b=1 -m "$model_path" -d=HETERO:FPGA,CPU -niter=8 \
3 -plugins_xml_file "$COREDLA_WORK/runtime/plugins.xml" \
4 -arch_file "${ARCH_PATH}/${arch}.arch" -api=async -perf_est -nireq=4 -bgr
```

In this command:

- `-b=1` specifies the batch size for inference.
- `-m "$model_path"` indicates the path to the model file to be used for inference.

- `-d=HETERO:FPGA,CPU` designates the heterogeneous execution across FPGA and CPU devices.
- `-niter=8` sets the number of iterations for the benchmark.
- `-plugins_xml_file "$COREDLA_WORK/runtime/plugins.xml"` specifies the plugins XML file required for the FPGA runtime.
- `-arch_file "$ARCH_PATH/$arch.arch"` points to the architecture file that matches the FPGA configuration.
- `-api=async` runs the benchmark in asynchronous mode for improved performance.
- `-perf_est` enables performance estimation during the benchmarking process.
- `-nireq=4` sets the number of inference requests to be processed simultaneously.
- `-bgr` specifies that the input data format is BGR (blue, green, red).

2.4.6 Compiling Bitstream for Optimized Architectures

When one is interested in compiling a bitstream after generating an optimized architecture, it is necessary to build the FPGA bitstream to deploy the PCIe-based design example on Arria 10 devices. This process ensures that the FPGA is configured according to the specified architecture. The following command is used to build the bitstream:

```

1 cd $COREDLA_WORK/demo
2 dla_build_example_design.py \
3 -a $COREDLA_ROOT/example_architectures/A10_Performance.arch \
4 --build -ed 1 \
5 --build-dir build_A10_Performance \
6 --output-dir $COREDLA_WORK/demo/my_bitstreams \
7 --n 1

```

In this command:

- `cd $COREDLA_WORK/demo` navigates to the demo directory where the bitstream will be built.
- `dla_build_example_design.py` is the script used to compile the design example.
- `-a $COREDLA_ROOT/example_architectures/A10_Performance.arch` specifies the path to the optimized architecture file generated for the Arria[®] 10 device.
- `--build -ed 1` initiates the build process for the specified design example.
- `--build-dir build_A10_Performance` sets the directory where the build files will be stored.
- `--output-dir $COREDLA_WORK/demo/my_bitstreams` designates the output directory where the compiled bitstream will be saved.
- `--n 1` ensures that a single instance of the architecture is built, which is crucial to avoid errors related to insufficient DSP resources.

It is important to include the `-n 1` option to build a single instance of the architecture. Without this option, one is likely to encounter an error due to insufficient DSP resources. If this error persists, it may be necessary to adjust the compilation process by lowering the `mmax-resources-alm-util` to a lower value (sometimes as low as 20 or 10) or by reducing the `mmax-resources` number for the DSPs during the generation of the optimized architecture.

2.4.7 Power usage

To monitor and evaluate the power consumption of the FPGA and GPU during inference, one must run the following commands: (1) For FPGA devices, the power usage can be checked by using the `fpgainfo` command. (2) For GPU devices, the power usage can be monitored using the `nvidia-smi` tool.

2.5 FPGA Throughput, Latency and Area Estimate Scripts

The following is a script that finds the throughput and area estimates of all architectures across all models including and excluding optimizations of Arria 10 and Agilex 7. The results are stored in a file titled “summary.csv”.

```

1 #!/bin/bash
2
3 # Define paths
4 BASE_PATH="/home/tvanteef/coredla_work/LST-AI-trigger-FPGA"
5 ARCH_PATH="/opt/intel/fpga_ai_suite_2024.1/dla/example_architectures"
6 MODELS=( "finaltrigger_cnn_8_16_fch_32_GlobalAvgPool_IB" "
7         testing_finaltrigger_cnn_8_16_fch_64_16" "
8         finaltrigger_cnn_8_16_fch_32_GlobalMaxPool" "
9         testing_finaltrigger_cnn_16_32_fch_128_32" "
10        finaltrigger_cnn_8_16_fch_32_GlobalMaxPool_noNSB" "
11        testing_finaltrigger_cnn_32_32_64_fch_256_64")
12
13 # Architecture arrays
14 A10_ARCHS=( "A10_FP16_Generic.arch" "A10_FP16_Performance.arch" "A10_Generic.arch"
15             " " "A10_Performance.arch")
16 AGX7_ARCHS=( "AGX7_FP16_Generic.arch" "AGX7_FP16_Performance.arch" "AGX7_Generic.
17             arch" "AGX7_Performance.arch")
18
19 # Optimized commands for A10
20 A10_OPTIMIZED="--gen-arch --mmax-resources=427200,2713,1518 --gen-min-sb=2048 --
21               mmax-resources-alm-util=75 --fassumed-fmax-core=300"
22
23 # Optimized commands for AGX7
24 AGX7_OPTIMIZED_1="--gen-arch --mmax-resources=487200,7110,4510 --gen-min-sb=2048
25                 --mmax-resources-alm-util=75 --fassumed-fmax-core=600"
26 AGX7_OPTIMIZED_2="--gen-arch --mmax-resources=1305600,18960,12300 --gen-min-sb
27                 =2048 --mmax-resources-alm-util=75 --fassumed-fmax-core=600"
28
29 # Output CSV file
30 CSV_FILE="$(pwd)/summary.csv"
31 echo "Model,Architecture,Optimization,Final Throughput (fps),ALMs,ALUTs,
32       Registers,DSPs,M20Ks,Memory ALMs" > $CSV_FILE
33
34 # Function to compile models
35 compile_model() {
36     local model=$1
37     local arch=$2
38     local opt_params=$3
39     local opt_label=$4
40
41     local network_file="${BASE_PATH}/${model}/ctlearn_model/Conversion/
42         ctlearn_model.xml"

```

```

31 local output_dir="${BASE_PATH}/Performance/${model}/${arch}"
32 mkdir -p $output_dir
33
34 local output_file="${output_dir}/ctlearn_model_${opt_label}.out.bin"
35 local log_file="${output_dir}/output_${opt_label}.txt"
36
37 dla_compiler --march "${ARCH_PATH}/${arch}" --network-file "$network_file" --
    foutput-format=open_vino_hetero --o "$output_file" --batch-size=1 --
    fanalyze-area --fanalyze-performance $opt_params > $log_file
38
39 # Parse and store the results
40 parse_and_store_results $log_file $model $arch $opt_label
41 }
42
43 # Function to parse log file and store results
44 parse_and_store_results() {
45     local log_file=$1
46     local model=$2
47     local arch=$3
48     local opt_label=$4
49
50     local throughput=$(grep -m 1 "FINAL THROUGHPUT =" $log_file | awk '{print $4}'
51 )
52     local ALMs=$(grep -m 1 "ALMs:" $log_file | awk '{print $2}')
53     local ALUTs=$(grep -m 1 "ALUTs:" $log_file | awk '{print $2}')
54     local Registers=$(grep -m 1 "Registers:" $log_file | awk '{print $2}')
55     local DSPs=$(grep -m 1 "DSPs:" $log_file | awk '{print $2}')
56     local M20Ks=$(grep -m 1 "M20Ks:" $log_file | awk '{print $2}')
57     local Memory_ALMs=$(grep -m 1 "Memory ALMs:" $log_file | awk '{print $3}')
58
59     echo "$model,$arch,$opt_label,$throughput,$ALMs,$ALUTs,$Registers,$DSPs,$M20Ks
    , $Memory_ALMs" >> $CSV_FILE
60 }
61
62 # Iterate over models and architectures
63 for model in "${MODELS[@]}; do
64     # A10 architectures
65     for arch in "${A10_ARCHS[@]}; do
66         compile_model $model $arch "" "regular"
67         compile_model $model $arch "$A10_OPTIMIZED" "optimized"
68     done
69
70     # AGX7 architectures
71     for arch in "${AGX7_ARCHS[@]}; do
72         compile_model $model $arch "" "regular"
73         compile_model $model $arch "$AGX7_OPTIMIZED_1" "optimized1"
74         compile_model $model $arch "$AGX7_OPTIMIZED_2" "optimized2"
75     done
76 done
77 echo "Compilation completed. Summary stored in $CSV_FILE."

```

Listing 1: Throughput and Area Estimates of Arria 10 and Agilex 7 for all models and architectures.

The next script does the same as the above but for the ResNet-50 model for comparison.

```

1 #!/bin/bash
2

```

```

3 # Define paths
4 ARCH_PATH="/opt/intel/fpga_ai_suite_2024.1/dla/example_architectures"
5 OUTPUT_BASE_PATH="." # Use current directory for output
6 MODEL="resnet50_model" # Use only ResNet-50 model
7
8 # Architecture arrays
9 A10_ARCHS=("A10_FP16_Generic.arch" "A10_FP16_Performance.arch" "A10_Generic.arch"
10 "A10_Performance.arch")
11
12 AGX7_ARCHS=("AGX7_FP16_Generic.arch" "AGX7_FP16_Performance.arch" "AGX7_Generic.
13 arch" "AGX7_Performance.arch")
14
15 # Optimized commands for A10
16 A10_OPTIMIZED="--gen-arch --mmax-resources=427200,2713,1518 --gen-min-sb=2048 --
17 mmax-resources-alm-util=75 --fassumed-fmax-core=300"
18
19 # Optimized commands for AGX7
20 AGX7_OPTIMIZED_1="--gen-arch --mmax-resources=487200,7110,4510 --gen-min-sb=2048
21 --mmax-resources-alm-util=75 --fassumed-fmax-core=600"
22
23 AGX7_OPTIMIZED_2="--gen-arch --mmax-resources=1305600,18960,12300 --gen-min-sb
24 =2048 --mmax-resources-alm-util=75 --fassumed-fmax-core=600"
25
26 # Output CSV file
27 CSV_FILE="$(pwd)/resnet_summary.csv"
28 echo "Model,Architecture,Optimization,Final Throughput (fps),ALMs,ALUTs,
29 Registers,DSPs,M20Ks,Memory ALMs" > $CSV_FILE
30
31 # Function to compile the model
32 compile_model() {
33     local model=$1
34     local arch=$2
35     local opt_params=$3
36     local opt_label=$4
37     local network_file="/home/tvanteef/model_performance/graph.xml" # Path to
38     ResNet-50 model
39     local output_dir="${OUTPUT_BASE_PATH}/${model}/${arch}"
40     mkdir -p $output_dir
41
42     local output_file="${output_dir}/resnet50_${opt_label}.out.bin"
43     local log_file="${output_dir}/output_resnet_${opt_label}.txt"
44
45     # Run the compilation and capture output
46     dla_compiler --march "${ARCH_PATH}/${arch}" --network-file "$network_file" --
47     foutput-format=open_vino_hetero --o "$output_file" --batch-size=1 --
48     fanalyze-area --fanalyze-performance $opt_params > $log_file
49
50     # Extract data from the log file
51     local throughput=$(grep -m 1 "FINAL THROUGHPUT =" $log_file | awk '{print $4}'
52 )
53     local ALMs=$(grep -m 1 "ALMs:" $log_file | awk '{print $2}')
54     local ALUTs=$(grep -m 1 "ALUTs:" $log_file | awk '{print $2}')
55     local Registers=$(grep -m 1 "Registers:" $log_file | awk '{print $2}')
56     local DSPs=$(grep -m 1 "DSPs:" $log_file | awk '{print $2}')
57     local M20Ks=$(grep -m 1 "M20Ks:" $log_file | awk '{print $2}')
58     local Memory_ALMs=$(grep -m 1 "Memory ALMs:" $log_file | awk '{print $3}')
59
60     # Append data to CSV file
61     echo "$model,$arch,$opt_label,$throughput,$ALMs,$ALUTs,$Registers,$DSPs,$M20Ks

```

```

    , $Memory_ALMs" >> $CSV_FILE
50 }
51
52 # Create the Output directory
53 mkdir -p "$OUTPUT_BASE_PATH"
54
55 # A10 architectures
56 for arch in "${A10_ARCHS[@]}; do
57     compile_model $MODEL $arch "" "regular"
58     compile_model $MODEL $arch "$A10_OPTIMIZED" "optimized"
59 done
60
61 # AGX7 architectures
62 for arch in "${AGX7_ARCHS[@]}; do
63     compile_model $MODEL $arch "" "regular"
64     compile_model $MODEL $arch "$AGX7_OPTIMIZED_1" "optimized1"
65     compile_model $MODEL $arch "$AGX7_OPTIMIZED_2" "optimized2"
66 done
67
68 echo "Compilation completed. Results stored in $CSV_FILE."

```

Listing 2: Throughput and Area Estimates of Arria 10 and Agilex 7 for all architectures on ResNet-50.

The following script evaluates the realtime throughput and latency performance of the Barvinok model across different architectures and saves them to a file. This is done using `dla_benchmark`.

```

1 #!/bin/bash
2
3 # Define paths
4 BASE_PATH="/home/tvanteef/coredla_work/LST-AI-trigger-FPGA"
5 COREDLA_WORK="/home/tvanteef/coredla_work"
6 ARCH_PATH="/opt/intel/fpga_ai_suite_2024.1/dla/example_architectures"
7 MODEL="finaltrigger_cnn_8_16_fch_32_GlobalAvgPool_IB"
8 A10_ARCHS=("A10_FP16_Generic" "A10_FP16_Performance" "A10_Generic" "
9     A10_Performance")
10 CSV_FILE="${BASE_PATH}/FPGA_benchmark_A10_smallmodel_Performance_results.csv"
11
12 # Function to run dla_benchmark and parse throughput and latency
13 run_and_parse_benchmark() {
14     local model=$1
15     local arch=$2
16     local model_path="${BASE_PATH}/${model}/ctlearn_model/Conversion/ctlearn_model
17         .xml"
18     local log_file="${BASE_PATH}/Performance/${model}/${arch}/benchmark_output.txt
19         "
20     mkdir -p "$(dirname "$log_file")"
21
22     fpgaconf -v $COREDLA_WORK/demo/bitstreams/${arch}.gbs
23     curarch=$COREDLA_ROOT/example_architectures/${arch}.arch
24
25     $COREDLA_WORK/runtime/build_Release/dla_benchmark/dla_benchmark -b=1 -m "
26         $model_path" -d=HETERO:FPGA,CPU -niter=8 -plugins_xml_file "$COREDLA_WORK/
27         runtime/plugins.xml" -arch_file "${ARCH_PATH}/${arch}.arch" -api=async -
28         perf_est -nireq=4 -bgr > $log_file
29
30     # Parse the first instance of throughput and latency from the log file
31     local throughput=$(grep -m 1 'FINAL THROUGHPUT =' $log_file | awk '{print $4}'
32         )

```

```

26 local latency=$(grep -m 1 'latency' $log_file | awk '{print $2}')
27 echo "$model,$arch,$throughput,$latency" >> $CSV_FILE
28 }
29
30 # Initialize CSV file
31 echo "Model,Architecture,Throughput (fps),Latency (ms)" > $CSV_FILE
32
33 # Iterate over A10 architectures
34 for arch in "${A10_ARCHS[@]}; do
35     run_and_parse_benchmark $MODEL $arch
36 done
37
38 echo "Benchmarking and parsing completed. Results saved to $CSV_FILE."

```

Listing 3: Realtime throughput and latency estimates for the Barvinok model across the Arria 10 architectures.

2.6 GPU and CPU Scripts

The following Python code calculates the throughput and latency estimates of the GPU through statistical means. One should set the sample size to a large value (100,000) get an accurate throughput or power reading.

```

1 import tensorflow as tf
2 import numpy as np
3 import time
4 import statistics
5 import os
6 import csv
7 import logging as log
8 import sys
9
10 # Setup logging
11 log.basicConfig(format='[ %(levelname)s ] %(message)s', level=log.INFO, stream=
    sys.stdout)
12
13 # Limit TensorFlow to only use GPU2
14 gpus = tf.config.list_physical_devices('GPU')
15 if gpus and len(gpus) > 2:
16     try:
17         # Set only GPU2 (index 2) as visible
18         tf.config.set_visible_devices(gpus[2], 'GPU')
19         # Optionally, configure GPU memory growth
20         tf.config.experimental.set_memory_growth(gpus[2], True)
21         log.info("Using GPU2 only")
22     except RuntimeError as e:
23         # Visible devices must be set before GPUs have been initialized
24         log.error(e)
25 else:
26     log.error("GPU2 is not available")
27
28 # Define the input shape (30, 30, 5)
29 input_shape = (30, 30, 5)
30
31 # Generate random data according to the specified input shape

```

```

32 num_samples = 500000 # You can vary this number as needed
33 data = np.random.random((num_samples, *input_shape))
34
35 def eval_throughput(model, data):
36     # Measure throughput
37     start_time = time.time()
38     model.predict(data)
39     end_time = time.time()
40     run_time = end_time - start_time
41     num_samples = data.shape[0]
42     throughput = num_samples / run_time
43
44     return throughput
45
46 def eval_latency(model, data, seconds_to_run=10, niter=10):
47     latencies = []
48     start_time = time.time()
49     time_point_to_finish = start_time + seconds_to_run
50     sample_index = 0
51
52     while time.time() < time_point_to_finish or len(latencies) < niter:
53         sample = data[sample_index:sample_index+1]
54         sample_start_time = time.time()
55         model.predict(sample)
56         sample_end_time = time.time()
57
58         latencies.append((sample_end_time - sample_start_time) * 1000) #
59             Convert to ms
60
61         sample_index = (sample_index + 1) % len(data)
62
63     median_latency = statistics.median(latencies)
64     average_latency = sum(latencies) / len(latencies)
65     min_latency = min(latencies)
66     max_latency = max(latencies)
67
68     return median_latency, average_latency, min_latency, max_latency
69
70 def main():
71     models = [
72         'finaltrigger_cnn_8_16_fch_32_GlobalMaxPool',
73         'finaltrigger_cnn_8_16_fch_32_GlobalMaxPool_noNSB',
74         'finaltrigger_cnn_8_16_fch_32_GlobalAvgPool_IB',
75         'testing_finaltrigger_cnn_16_32_fch_128_32',
76         'testing_finaltrigger_cnn_32_32_64_fch_256_64',
77         'testing_finaltrigger_cnn_8_16_fch_64_16'
78     ]
79
80     base_path = '/home/hep/tvanteef/LST-AI-trigger-FPGA/'
81     results = []
82
83     for model_name in models:
84         model_path = os.path.join(base_path, model_name, 'ctlearn_model')
85         if not os.path.exists(model_path):
86             log.warning(f'Model file not found: {model_path}')
87             continue

```

```

88     log.info(f'Benchmarking model: {model_name}')
89
90     # Load the model
91     model = tf.keras.models.load_model(model_path)
92
93     # Evaluate throughput
94     throughput = eval_throughput(model, data)
95
96     # Evaluate latency
97     median_latency, avg_latency, min_latency, max_latency = eval_latency(
98         model, data)
99
100    results.append({
101        'model': model_name,
102        'throughput_fps': throughput,
103        'median_latency_ms': median_latency,
104        'average_latency_ms': avg_latency,
105        'min_latency_ms': min_latency,
106        'max_latency_ms': max_latency
107    })
108
109    # Save results to CSV
110    output_file = 'benchmark_GPU_TensorFlow_results.csv'
111    with open(output_file, 'w', newline='') as csvfile:
112        fieldnames = ['model', 'throughput_fps', 'median_latency_ms', '
113            average_latency_ms', 'min_latency_ms', 'max_latency_ms']
114        writer = csv.DictWriter(csvfile, fieldnames=fieldnames)
115        writer.writeheader()
116        writer.writerows(results)
117
118    log.info(f'Benchmark results saved to {output_file}')
119
120    if __name__ == '__main__':
121        main()

```

Listing 4: Throughput and latency estimates of the GPU.

The analogue of the above code for the CPU can be found below.

```

1 import tensorflow as tf
2 import numpy as np
3 import time
4 import statistics
5 import os
6 import csv
7 import logging as log
8 import sys
9
10 # Setup logging
11 log.basicConfig(format='[ %(levelname)s ] %(message)s', level=log.INFO, stream=
12     sys.stdout)
13
14 # Load the data
15 data = np.load('/home/hep/tvanteef/input_triggerpatch_waveforms_run73.npy')
16
17 # Ensure data has the correct shape
18 if data.shape[-1] > 5:
19     data = data[..., :5]

```

```

19 elif data.shape[-1] < 5:
20     raise ValueError("Data must have at least 5 channels for the models")
21
22 # Resize images to 30x30 if necessary
23 data_resized = tf.image.resize(data, [30, 30])
24
25 def eval_throughput(model, data):
26     # Measure throughput
27     start_time = time.time()
28     model.predict(data)
29     end_time = time.time()
30     run_time = end_time - start_time
31     num_samples = data.shape[0]
32     throughput = num_samples / run_time
33
34     return throughput
35
36 def eval_latency(model, data, seconds_to_run=10, niter=10):
37     latencies = []
38     start_time = time.time()
39     time_point_to_finish = start_time + seconds_to_run
40     sample_index = 0
41
42     while time.time() < time_point_to_finish or len(latencies) < niter:
43         sample = data[sample_index:sample_index+1]
44         sample_start_time = time.time()
45         model.predict(sample)
46         sample_end_time = time.time()
47
48         latencies.append((sample_end_time - sample_start_time) * 1000) #
49             Convert to ms
50
51         sample_index = (sample_index + 1) % len(data)
52
53     median_latency = statistics.median(latencies)
54     average_latency = sum(latencies) / len(latencies)
55     min_latency = min(latencies)
56     max_latency = max(latencies)
57
58     return median_latency, average_latency, min_latency, max_latency
59
60 def main():
61     models = [
62         'finaltrigger_cnn_8_16_fch_32_GlobalMaxPool',
63         'finaltrigger_cnn_8_16_fch_32_GlobalMaxPool_noNSB',
64         'finaltrigger_cnn_8_16_fch_32_GlobalAvgPool_IB',
65         'testing_finaltrigger_cnn_16_32_fch_128_32',
66         'testing_finaltrigger_cnn_32_32_64_fch_256_64',
67         'testing_finaltrigger_cnn_8_16_fch_64_16'
68     ]
69
70     base_path = '/home/hep/tvanteeef/LST-AI-trigger-FPGA/'
71     results = []
72
73     for model_name in models:
74         model_path = os.path.join(base_path, model_name, 'ctlearn_model')
75         if not os.path.exists(model_path):

```



```

75         log.warning(f'Model file not found: {model_path}')
76         continue
77
78     log.info(f'Benchmarking model: {model_name}')
79
80     # Load the model
81     model = tf.keras.models.load_model(model_path)
82
83     # Ensure that the model runs on CPU
84     with tf.device('/CPU:0'):
85         # Evaluate throughput
86         throughput = eval_throughput(model, data_resized)
87
88         # Evaluate latency
89         median_latency, avg_latency, min_latency, max_latency = eval_latency
90             (model, data_resized)
91
92     results.append({
93         'model': model_name,
94         'throughput_fps': throughput,
95         'median_latency_ms': median_latency,
96         'average_latency_ms': avg_latency,
97         'min_latency_ms': min_latency,
98         'max_latency_ms': max_latency
99     })
100
101     # Save results to CSV
102     output_file = 'benchmark_CPU_TensorFlow_results.csv'
103     with open(output_file, 'w', newline='') as csvfile:
104         fieldnames = ['model', 'throughput_fps', 'median_latency_ms', '
105             average_latency_ms', 'min_latency_ms', 'max_latency_ms']
106         writer = csv.DictWriter(csvfile, fieldnames=fieldnames)
107         writer.writeheader()
108         writer.writerows(results)
109
110     log.info(f'Benchmark results saved to {output_file}')
111
112 if __name__ == '__main__':
113     main()

```

Listing 5: Throughput and latency estimates of the CPU.

The code below runs the ResNet-50 model on the GPU. Once again, a higher sample number yields a more accurate throughput reading.

```

1 import tensorflow as tf
2 import numpy as np
3 import time
4 import statistics
5 import csv
6
7 # Limit TensorFlow to only use GPU2
8 gpus = tf.config.list_physical_devices('GPU')
9 if gpus and len(gpus) > 2:
10     try:
11         # Set only GPU2 (index 2) as visible
12         tf.config.set_visible_devices(gpus[2], 'GPU')

```

```

13         # Optionally, configure GPU memory growth
14         tf.config.experimental.set_memory_growth(gpus[2], True)
15         print("Using GPU2 only")
16     except RuntimeError as e:
17         # Visible devices must be set before GPUs have been initialized
18         print(e)
19 else:
20     print("GPU2 is not available or does not exist. Using CPU.")
21
22 # Define the input shape (224, 224, 3) as required for ResNet-50
23 input_shape = (224, 224, 3)
24
25 # Generate random data according to the specified input shape
26 num_samples = 50000 # You can vary this number as needed
27 data = np.random.random((num_samples, *input_shape))
28
29 # Preprocess the images to match ResNet-50 input requirements
30 data_preprocessed = tf.keras.applications.resnet50.preprocess_input(data)
31
32 # Load ResNet-50 model from keras applications
33 model = tf.keras.applications.ResNet50(
34     include_top=True,
35     weights='imagenet',
36     input_tensor=None,
37     input_shape=input_shape,
38     pooling=None,
39     classes=1000
40 )
41
42 def eval_throughput(model, data):
43     # Measure throughput
44     start_time = time.time()
45     model.predict(data)
46     end_time = time.time()
47     run_time = end_time - start_time
48     num_samples = data.shape[0]
49     throughput = num_samples / run_time
50
51     return throughput
52
53 def eval_latency(model, data, seconds_to_run=10, niter=10):
54     latencies = []
55     start_time = time.time()
56     time_point_to_finish = start_time + seconds_to_run
57     sample_index = 0
58
59     while time.time() < time_point_to_finish or len(latencies) < niter:
60         sample = data[sample_index:sample_index+1]
61         sample_start_time = time.time()
62         model.predict(sample)
63         sample_end_time = time.time()
64
65         latencies.append((sample_end_time - sample_start_time) * 1000) #
66             Convert to ms
67
68         sample_index = (sample_index + 1) % len(data)

```

```

69     median_latency = statistics.median(latencies)
70     average_latency = sum(latencies) / len(latencies)
71     min_latency = min(latencies)
72     max_latency = max(latencies)
73
74     return median_latency, average_latency, min_latency, max_latency
75
76 # Function to reset TensorFlow device placement
77 def reset_device():
78     tf.keras.backend.clear_session()
79     if gpus and len(gpus) > 2:
80         tf.config.set_visible_devices(gpus[2], 'GPU') # Use GPU2
81     else:
82         tf.config.set_visible_devices([], 'GPU')
83
84 # Evaluate on GPU2 if available
85 if gpus and len(gpus) > 2:
86     reset_device()
87     with tf.device('/GPU:2'):
88         print("\nEvaluating on GPU2...")
89         model_gpu = tf.keras.applications.ResNet50(
90             include_top=True,
91             weights='imagenet',
92             input_tensor=None,
93             input_shape=input_shape,
94             pooling=None,
95             classes=1000
96         )
97         throughput_gpu = eval_throughput(model_gpu, data_preprocessed)
98         median_latency_gpu, avg_latency_gpu, min_latency_gpu, max_latency_gpu =
99             eval_latency(model_gpu, data_preprocessed)
100         print(f"Throughput on GPU2: {throughput_gpu:.2f} samples/sec")
101         print(f"Median Latency on GPU2: {median_latency_gpu:.4f} ms/sample")
102         print(f"Average Latency on GPU2: {avg_latency_gpu:.4f} ms/sample")
103         print(f"Min Latency on GPU2: {min_latency_gpu:.4f} ms/sample")
104         print(f"Max Latency on GPU2: {max_latency_gpu:.4f} ms/sample")
105 else:
106     throughput_gpu, median_latency_gpu, avg_latency_gpu, min_latency_gpu,
107     max_latency_gpu = None, None, None, None, None
108
109 # Save the results to a file
110 output_file = 'benchmark_ResNet50_results.csv'
111 results = []
112
113 if throughput_gpu is not None:
114     results.append({
115         'device': 'GPU2',
116         'throughput_samples_sec': throughput_gpu,
117         'median_latency_ms': median_latency_gpu,
118         'average_latency_ms': avg_latency_gpu,
119         'min_latency_ms': min_latency_gpu,
120         'max_latency_ms': max_latency_gpu
121     })
122
123 # Write results to a CSV file
124 with open(output_file, 'w', newline='') as csvfile:
125     fieldnames = ['device', 'throughput_samples_sec', 'median_latency_ms', '

```

```

124         'average_latency_ms', 'min_latency_ms', 'max_latency_ms']
125     writer = csv.DictWriter(csvfile, fieldnames=fieldnames)
126     writer.writeheader()
127     writer.writerows(results)
128 print(f'\nBenchmark results saved to {output_file}')
```

Listing 6: Throughput and latency estimates for the ResNet-50 model on the GPU.

3 Results and Discussion

The neural networks evaluated in this study, listed in order of their number of parameters, are as follows: Barvinok 1 (2,179 parameters), Verba (52,867 parameters), Barvinok 2 (2,179 parameters), Malva (210,435 parameters), Barvinok 3 (2,146 parameters), and Topolia (193,571 parameters).

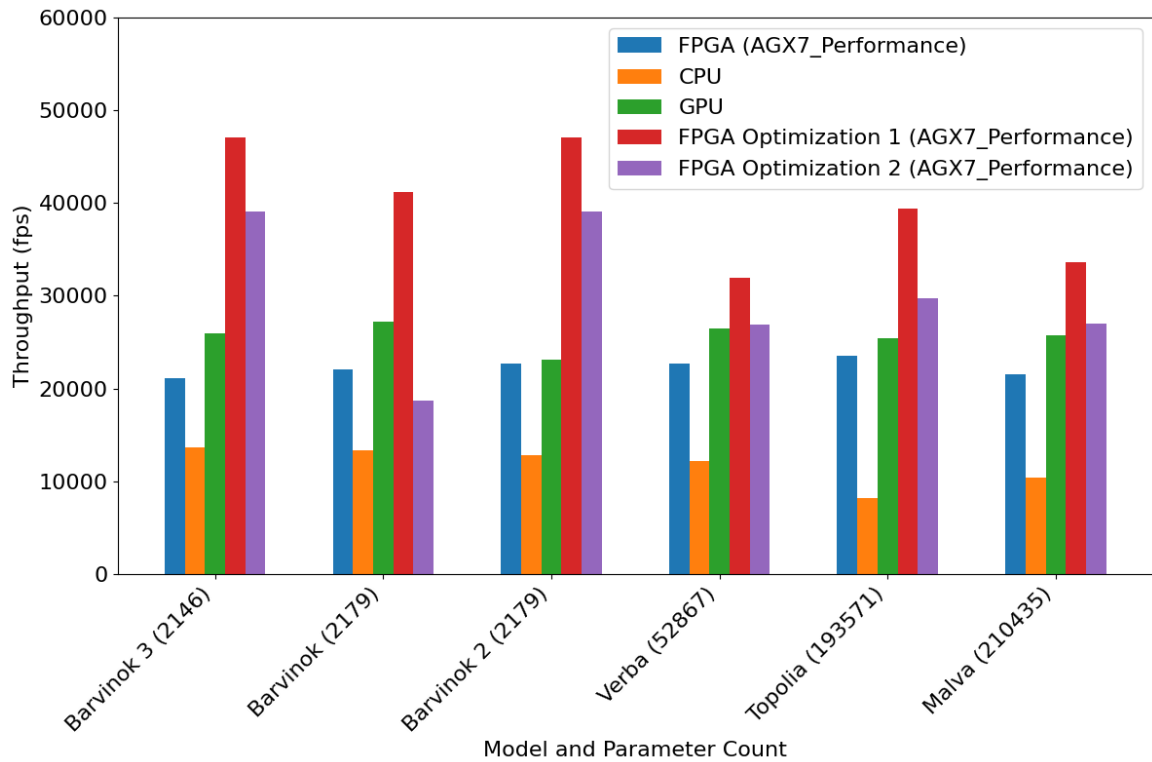
3.1 Model dependence

In this section, we analyze the dependence of throughput and latency on different neural network models across the evaluated hardware platforms, including CPUs, GPUs, and FPGAs. The results are illustrated in Figures 1 and 2.

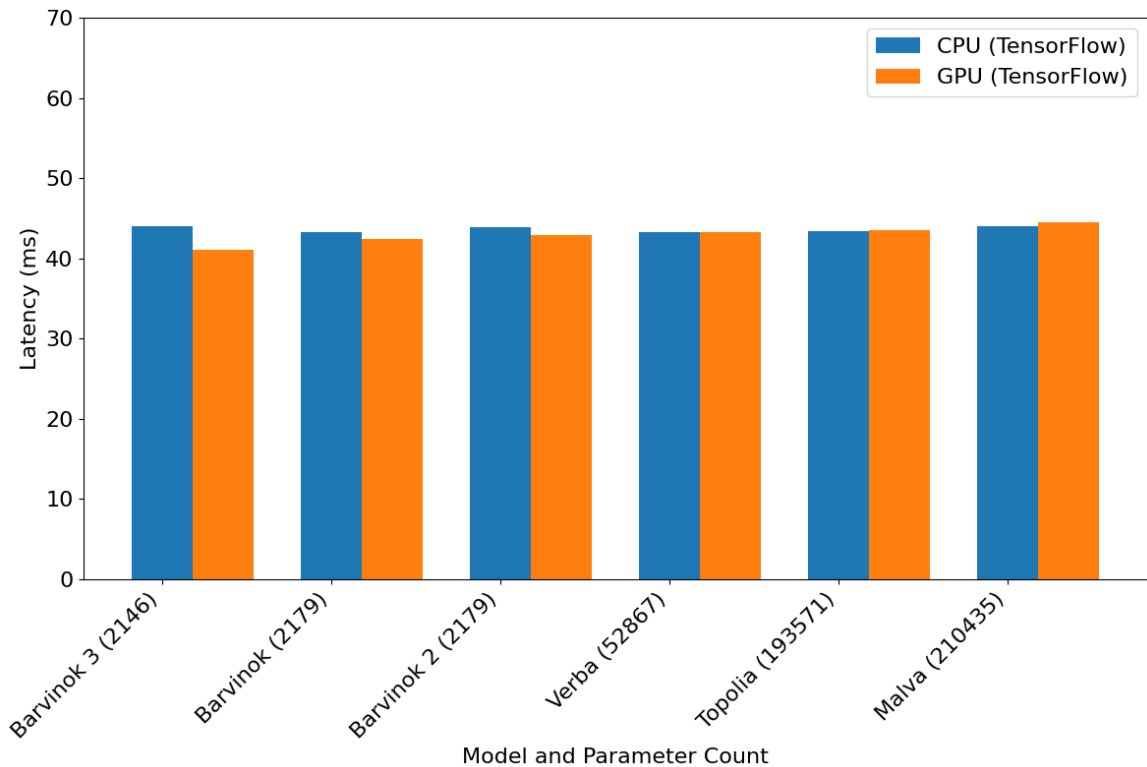
The results in Figure 1(a) indicate that the **Agilex 7 FPGA** exhibits the highest throughput among all platforms, particularly when optimized. This is especially evident in the significant performance gains observed with the Agilex 7 when compared to the GPU. Notably, Optimization 1 outperforms Optimization 2, despite Optimization 2 having access to a greater number of resources. This suggests that the larger configuration of the AGM039 variant may encounter scaling inefficiencies that could reduce performance. As a result, utilizing multiple cores or more sophisticated optimization strategies may be necessary to fully exploit the available resources in larger FPGA configurations.

The **Arria 10 FPGA**, as shown in Figure 2(a), demonstrates competitive performance, especially when optimized. In its unoptimized state, the Arria 10's throughput is on par with that of the GPU. However, after optimization, the Arria 10 is capable of matching or even surpassing GPU performance for certain models.

Latency measurements, depicted in Figures 1(b) and 2(b), reveal a significant advantage for the **Arria 10 FPGA**, which consistently outperforms both the CPU and GPU in terms of processing speed. In contrast, the latency figures for the CPU and GPU are quite similar.

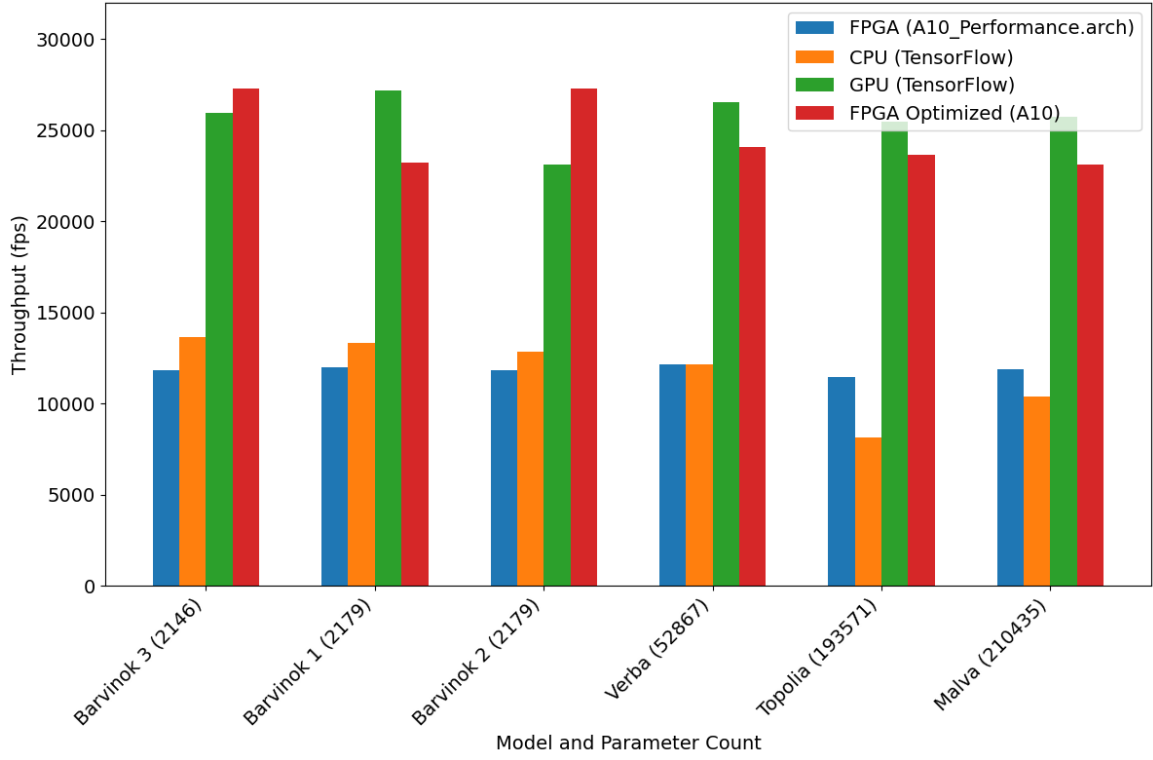


(a) Throughput vs. Model

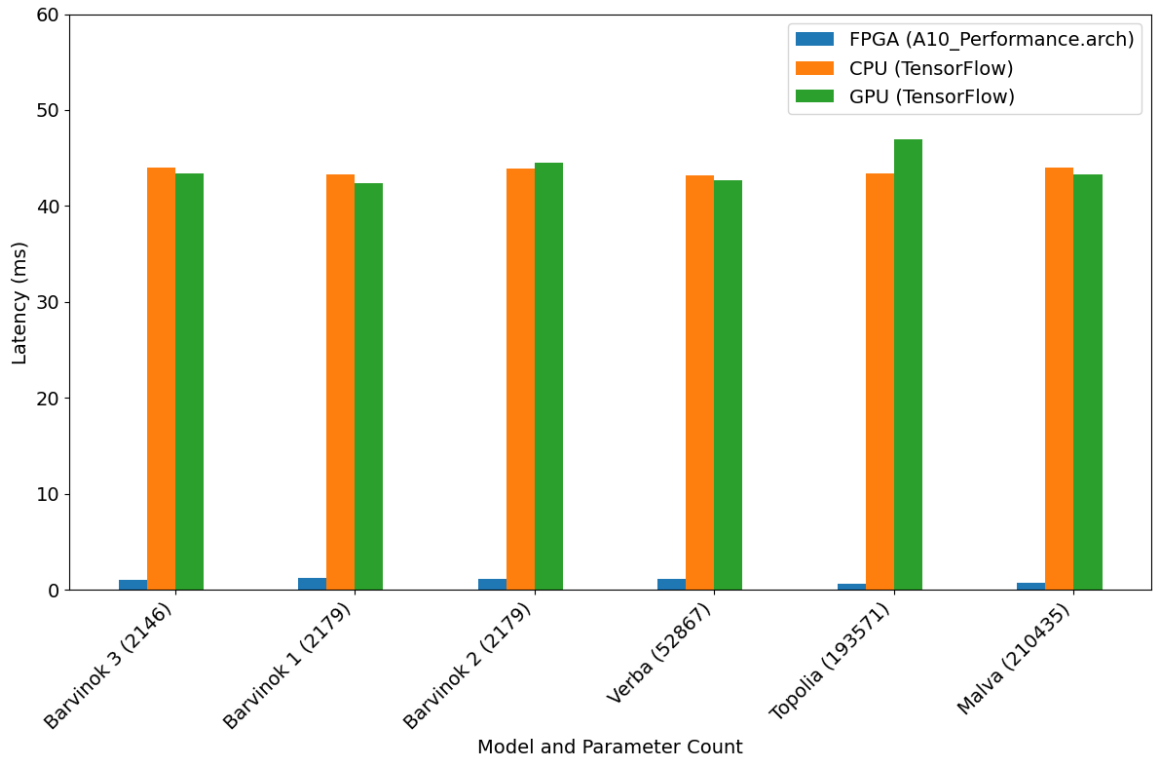


(b) Latency vs. Model

Figure 1: Performance comparison of the evaluated neural networks on the CPU, GPU, and Agilex 7 FPGA. (a) Throughput vs. Model and (b) Latency vs. Model. Optimization 1 corresponds to optimizing the AFG014 version of Agilex 7, while Optimization 2 corresponds to optimizing the AGM039 version. The FPGA values were computed using `dla_compiler` while the GPU and CPU values were using `tensorflow`.



(a) Throughput vs. Model



(b) Latency vs. Model

Figure 2: Performance comparison of the evaluated neural networks on the CPU, GPU, and Agilx 7 FPGA. (a) Throughput vs. Model and (b) Latency vs. Model. The FPGA values were computed using `dla_compiler` while the GPU and CPU values were using `tensorflow`. The latency calculation for Arria 10 was performed using `dla_benchmark` seeing that `dla_compiler` does not provide a latency estimate.

A comparison between the compiler estimates from `dla_compiler` and the real-time results from `dla_benchmark`, executed directly on the FPGA, reveals consistent similarities across the models. This consistency is illustrated in Figure 3, where the estimation was performed using the A10_Performance architecture.

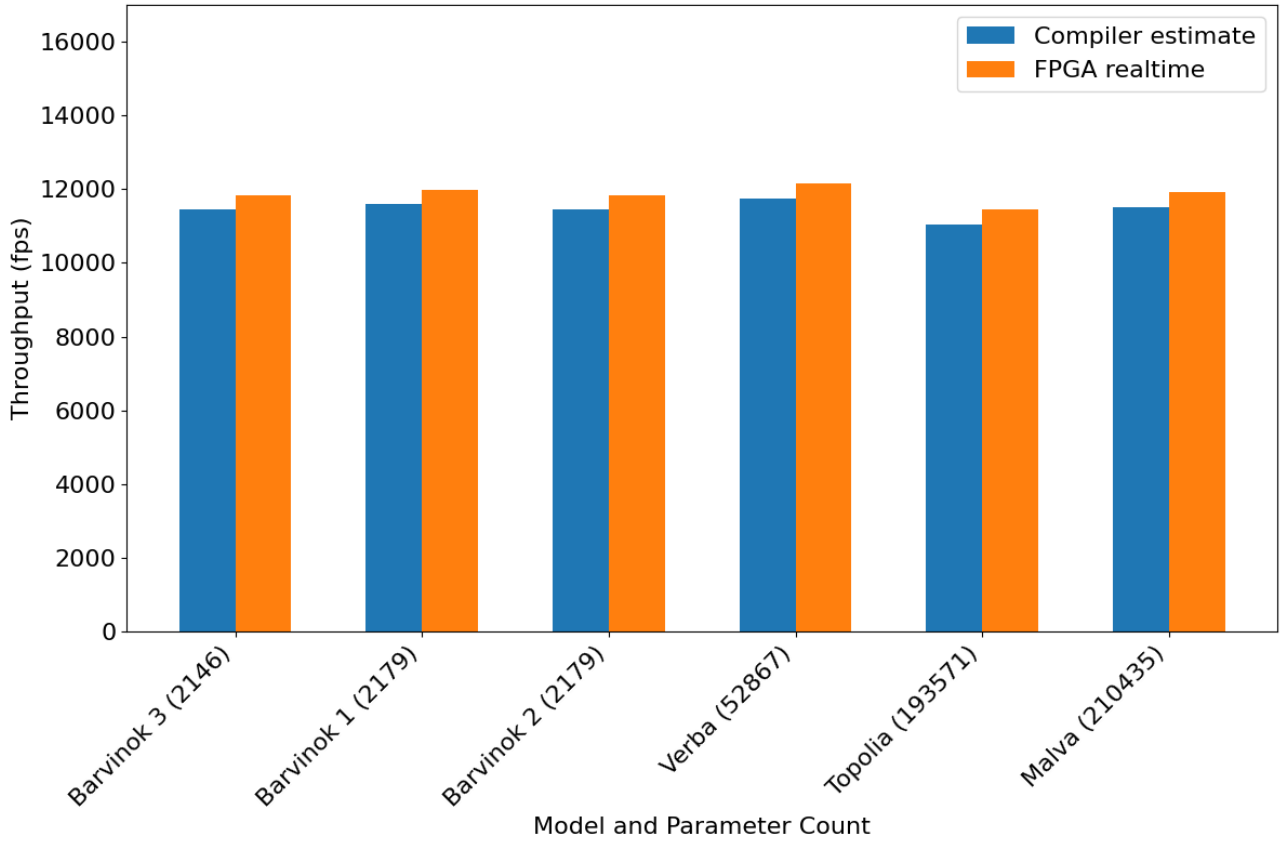
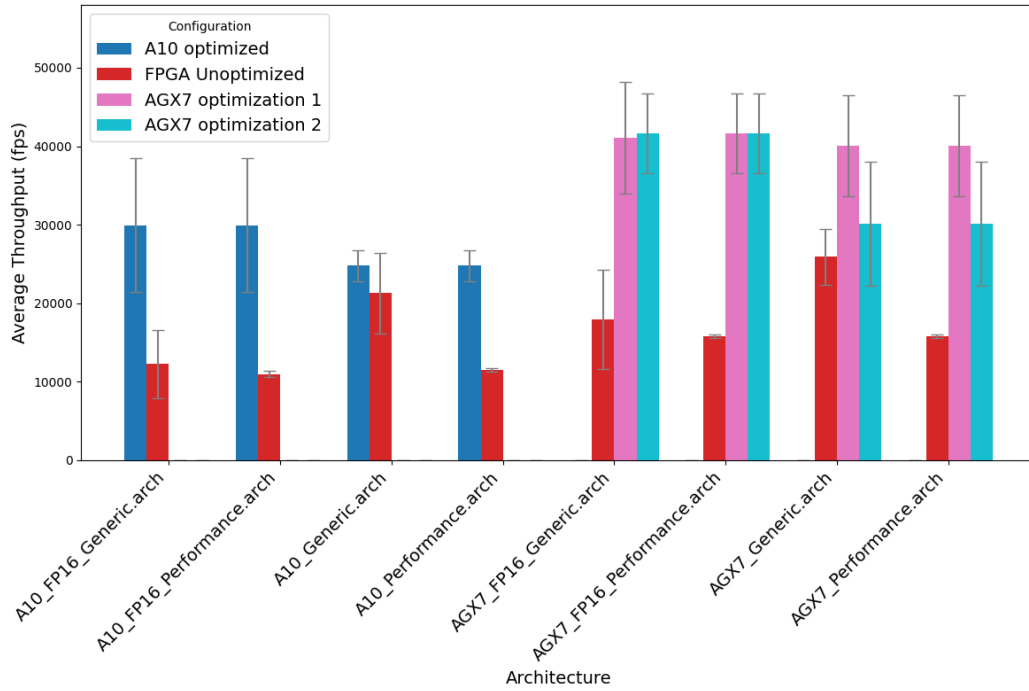


Figure 3: Throughput performance as measured by the compiler with `dla_compiler` and directly run on the Arria 10 FPGA with `dla_benchmark`. The estimates are given for the A10_Performance architecture.

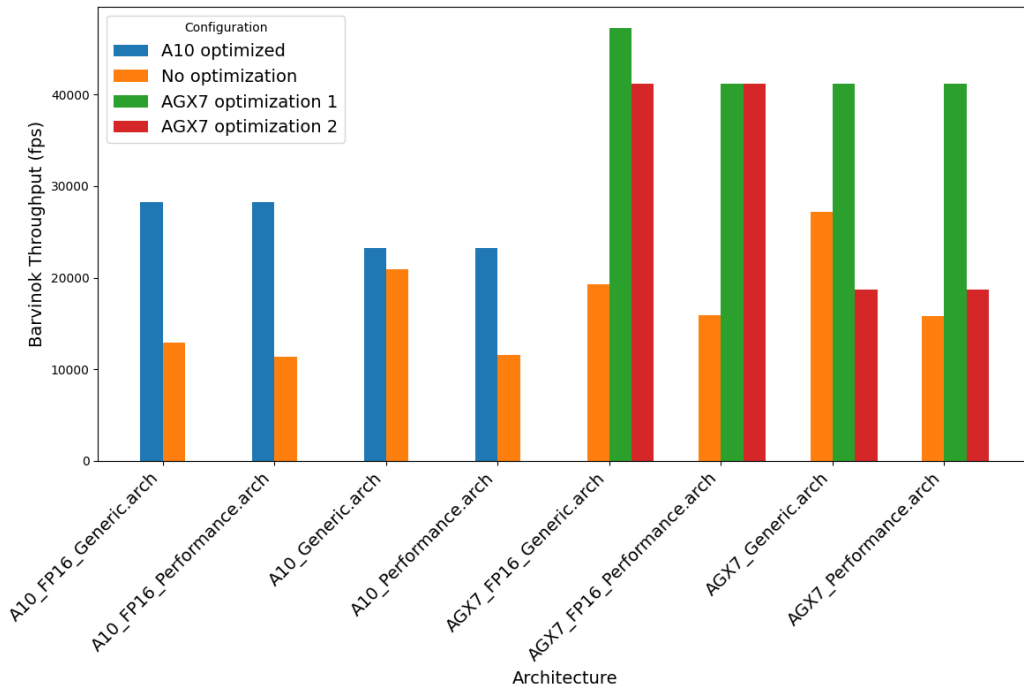
3.2 Architecture dependence

Figures 4(a) and 4(b) illustrate the throughput performance of different CNN models as estimated by `dla_compiler` for the Arria 10 and Agilex 7 FPGAs across various architectures.

A striking difference is observed between the FP16 Generic architecture and the regular Generic architecture. The FP16 Generic architecture consistently outperforms its regular counterpart, a trend that is less pronounced when comparing the Performance architectures. Moreover, the optimizations applied to the FPGA architectures yield substantial performance improvements, with an average 2x increase in throughput, as clearly seen in both figures.



(a) The mean throughput performance across all CNN models as estimated by `dla_compiler` for the Arria 10 and Agilex 7 FPGAs. Error bars indicate the standard deviation in throughput from the mean across all CNN models for a given architecture.



(b) The throughput performance for the Barvinok model as estimated by `dla_compiler` for the Arria 10 and Agilex 7 FPGAs.

Figure 4: Throughput performance comparisons for the Arria 10 and Agilex 7 FPGAs: (a) Mean throughput across all CNN models, and (b) Throughput for the Barvinok model. “FPGA Unoptimized” refers to the respective FPGA architecture, either A10 or AGX7, based on the prefix.

In general, the throughput varies significantly across the different CNN models, as evidenced by the large error bars in Figure 4(a). The error bars represent the standard deviation in throughput from the mean across all CNN models. This variability indicates that the performance of FPGAs is highly model-dependent, with some models achieving much higher throughput than others. However, it is notable that the unoptimized Performance architectures exhibit more consistent throughput across models, with relatively smaller error bars.

Figure 5 presents the throughput performance for the Barvinok model, including optimizations as measured by the compiler with `dla_compiler` and directly run on the Arria 10 FPGA using `dla_benchmark`. The optimized throughputs, as seen for both the Generic and Performance architectures, are identical. This consistency is expected as the FPGA is designed to optimize the architectures uniquely.

However, an interesting observation is that the optimized throughput for the Generic architecture is lower than the throughput for the unoptimized architecture. This could indicate that the optimization process is specifically tailored to the Barvinok model, which might inadvertently reduce the effectiveness of the generic optimization when applied more broadly. Further investigation, including compiling and testing the corresponding bitstreams, is suggested to fully understand the cause of this discrepancy.

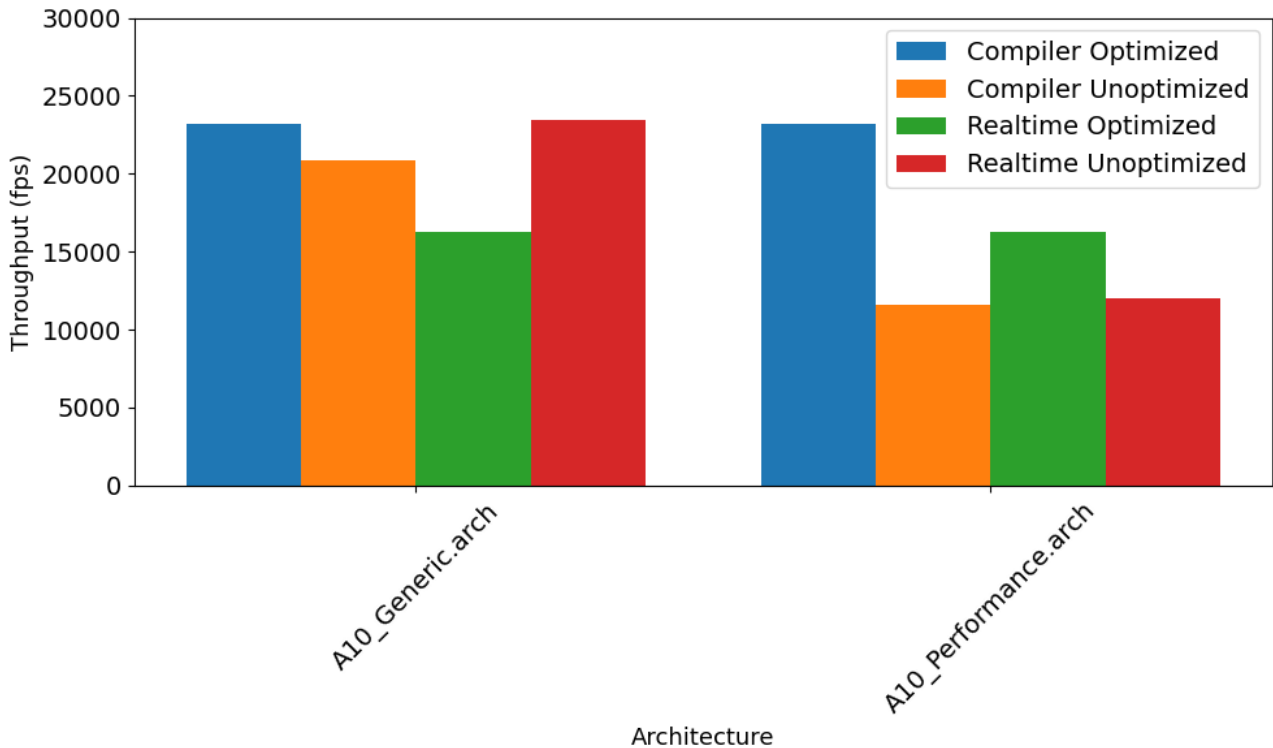


Figure 5: Throughput performance including optimizations as measured by the compiler with `dla_compiler` and directly run on the Arria 10 FPGA with `dla_benchmark`. The estimates are given for the Barvinok model.

3.3 Comparison with ResNet-50

Figures 6 and 7 illustrate the throughput performance for the ResNet-50 model with both the compiler-predicted and real-time performance across various architectures. Several key observations can be made from these graphs.

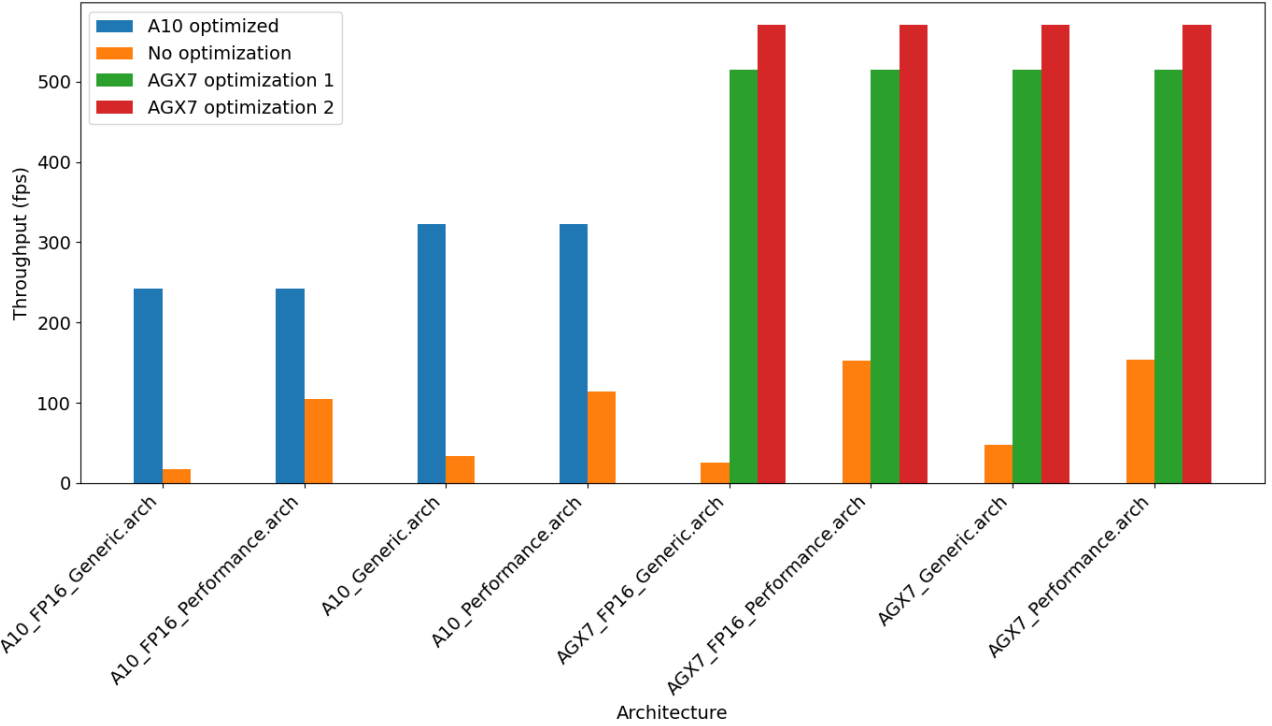


Figure 6: Compiler throughput performance as a function of architecture for the ResNet-50 model of 25.6M parameters for the Arria 10 and Agilex 7 FPGAs.

Firstly, it is evident that the compiler’s prediction for the expected performance boost is significantly higher for the ResNet-50 model compared to the previous CNN models (10x performance boost).

Secondly, an intriguing observation is the performance difference between the two Optimization 1 and Optimization 2. Unlike the previous CNN models where Optimization 1 typically outperformed Optimization 2, the ResNet-50 model shows a reversal of this trend and aligns better with expectations: with more resources, the net throughput should increase. Moreover, it is worth noting that the optimizations predicted by the compiler are consistent with those observed in real-time performance, as shown in the comparison between the compiler and real-time throughput in subfigure (b).

Thirdly, the observed consistency in performance improvements for the ResNet-50 model across both predicted and real-time data indicates that the previous inconsistencies noted with the Barvinok model are likely model-specific. The fact that the Barvinok model did not exhibit the same degree of performance consistency suggests that further testing and analysis are necessary to fully understand the underlying factors contributing to the observed discrepancies in performance gains.

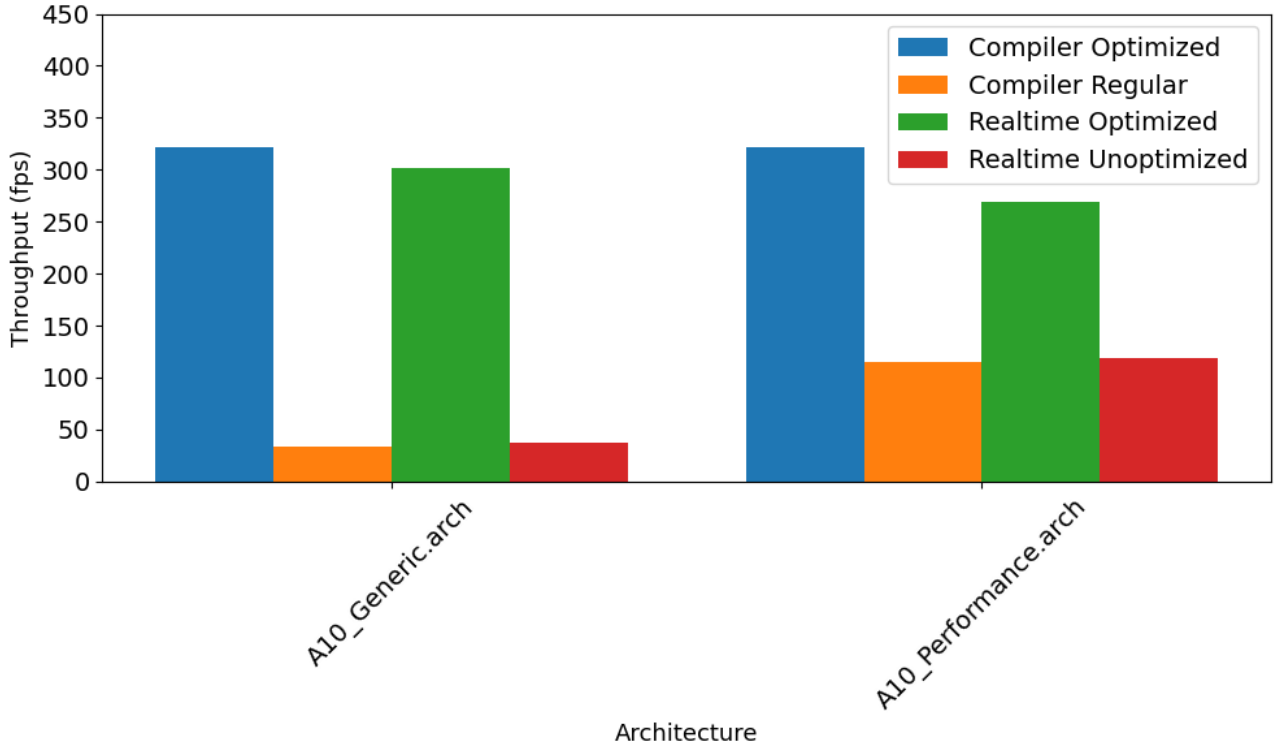


Figure 7: A comparison between the compiler and realtime throughput performance as a function of architecture for the ResNet-50 model of 25.6M parameters including optimizations.

Finally, the optimized throughput for ResNet-50, as shown in Figure 7, varies slightly between the two architectures. This difference is due to a technical adjustment during the Quartus bitstream compilation, where the estimated resource usage exceeded the available limit. To address this, the maximum resource bound was reduced to 1450 DSPs instead of the full 1518 DSPs, leading to a marginal reduction in performance.

Table 1: Throughput and Mean Latency of ResNet-50 on GPU2

Device	Throughput (samples/sec)	Mean Latency (ms)
GPU	788.84	48.56

Table 1 shows the throughput and latency performance of ResNet-50. As can be seen, the GPU has more than double the throughput of an FPGA.

3.4 Power usage and area estimates

The graph in Figure 8 clearly demonstrates the significant power efficiency of FPGA devices compared to GPUs. The power usage of the GPU is notably higher across all inference types, with ResNet-50 and Barvinok models consuming around 309 and 89 watts respectively, whereas the Arria 10 FPGA configurations consistently consume less than half of that (or less than 10 times for ResNet-50), typically ranging between 24.5 to 36 watts depending on the architecture.

Tables 2 and 3 presented below highlight the FPGA resource usage for both the Barvinok models and the ResNet-50 model across various architectures and optimization strategies. A key observation

is the disparity in performance and resource utilization between Optimization 1 and Optimization 2 across different models. For the CNN trigger models (Barvinok 2 and Barvinok 3), Optimization 1 outperforms Optimization 2 in throughput, despite Optimization 2 occupying a significantly larger area in terms of ALMs, M20Ks, and DSPs. Conversely, for the ResNet-50 model, Optimization 2 outperforms Optimization 1, as expected, given that Optimization 2 occupies more resources across all FPGA architectures. This correlation between resource allocation and performance aligns better with the typical behavior observed in FPGA optimizations.

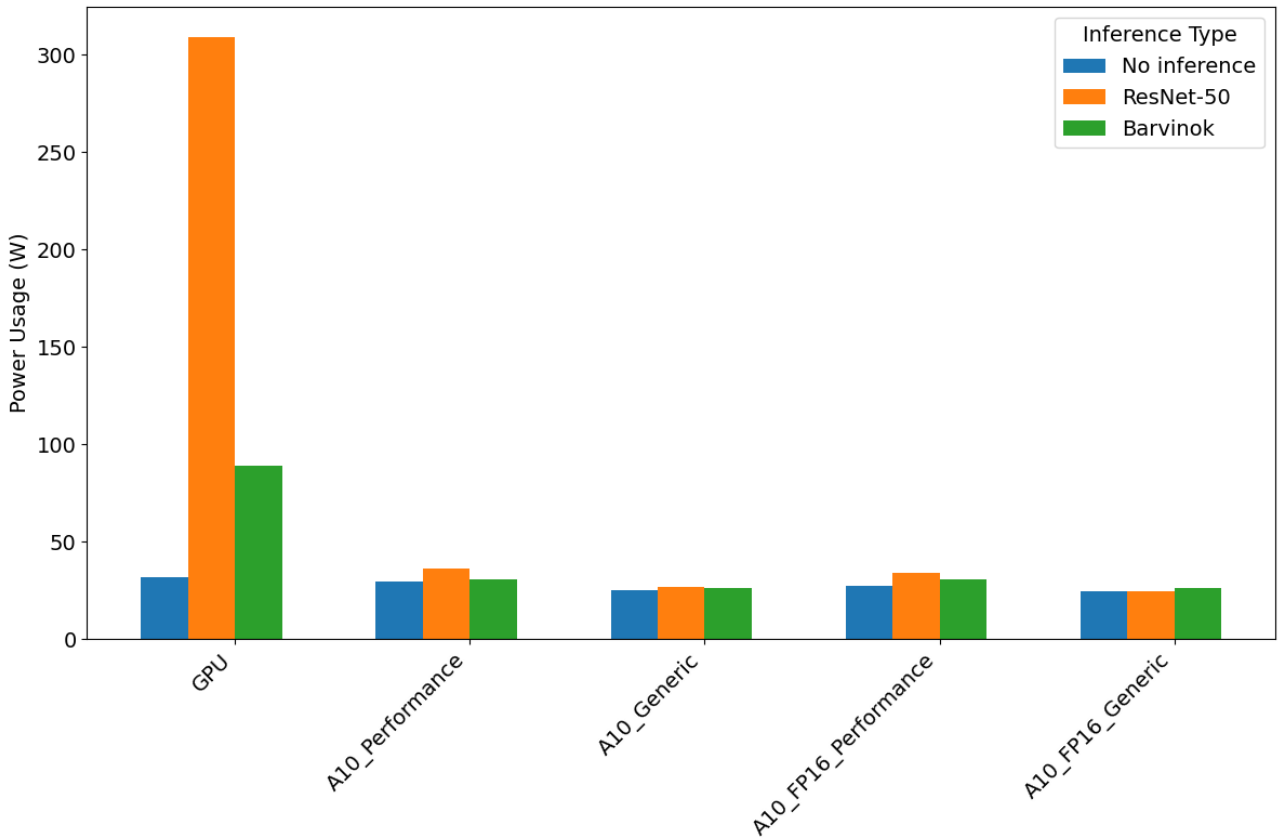


Figure 8: The power usage as a function of the device used: GPU and the Arria 10 FPGA. The maximum resources for

Additionally, it is important to note that for some A10 architectures used with the ResNet-50 model, the number of occupied DSPs exceeded the maximum specified limits. To address this, the maximum DSP allocation was reduced by the difference between the occupied area and the specified maximum, resulting in a final DSP count of 1518 during the execution of the `dla_compiler`.

Table 2: FPGA Resource Usage for Barvinok 2 and Barvinok 3 Models using the Performance architecture. The maximum specified resources are as follows. For Arria 10: 427,200 ALMs, 2,713 M20Ks, and 1,518 DSPs. For Agilex 7: Optimization 1 has 487,200 ALMs, 7,110 M20Ks, and 4,510 DSPs, while Optimization 2 offers 1,305,600 ALMs, 18,960 M20Ks, and 12,300 DSPs.

Optimization	Model	Occupied ALMs	Occupied M20Ks	Occupied DSPs
A10 Unoptimized	Barvinok 2	47259	953	650
A10 Optimized	Barvinok 2	47876	1545	362
A10 Unoptimized	Barvinok 3	47259	953	650
A10 Optimized	Barvinok 3	47876	1545	362
AGX7 Unoptimized	Barvinok 2	62451	1237	650
AGX7 Optimization 1	Barvinok 2	122489	1471	1226
AGX7 Optimization 2	Barvinok 2	175652	2486	2314
AGX7 Unoptimized	Barvinok 3	62451	1237	650
AGX7 Optimization 1	Barvinok 3	122489	1471	1226
AGX7 Optimization 2	Barvinok 3	175652	2486	2314

4 Conclusion

This study conducted a comprehensive evaluation of deep neural network performance on various hardware platforms, including FPGAs, GPUs, and CPUs, within the context of the Cherenkov Telescope Array’s real-time triggering system. The primary metrics of interest—throughput and latency—were assessed across different neural network models and hardware architectures.

Key findings from this research include the superior performance of the Agilex 7 FPGA, particularly when optimized, which outperformed both the Arria 10 FPGA and the GPU with double the throughput. However, the study also revealed that resource allocation does not always directly correlate with performance improvements, as demonstrated by the inconsistency between Optimization 1 and Optimization 2 across different models. Specifically, for the CNN trigger models, Optimization 1 (the AFG014 board) outperformed Optimization 2 (the AGM039 board) despite using fewer resources. Conversely, for the ResNet-50 model, Optimization 2, which utilized more resources, delivered better performance, aligning with expected outcomes.

The study also revealed general consistency between compiler-estimated and real-time throughput values, thereby validating the compiler’s predictions. However, discrepancies were noted between optimized compiler estimates and actual real-time throughput, particularly for the Barvinok model. The ResNet-50 model showed more consistent performance improvements across both predicted and real-time data. This suggests that the inconsistencies observed with the Barvinok model may be model-specific and require further investigation.

FPGAs showed remarkable power efficiency. The power consumption of FPGA configurations was consistently less than half that of the GPU across all inference types, making FPGAs a highly attractive option for the energy-constrained CTA project.

Future research directions include:

Table 3: FPGA area estimate and resource usage for ResNet-50 on various architectures of Arria 10. The maximum specified resources for Arria 10 are: 427,200 ALMs, 2,713 M20Ks, and 1,518 DSPs. For Agilex 7: Optimization 1 has 487,200 ALMs, 7,110 M20Ks, and 4,510 DSPs, while Optimization 2 offers 1,305,600 ALMs, 18,960 M20Ks, and 12,300 DSPs.

Optimization	Architecture	Occupied ALMs	Occupied M20Ks	Occupied DSPs
Unoptimized	A10_FP16_Generic.arch	24457	507	186
Optimized	A10_FP16_Generic.arch	165515	2670	1518
Unoptimized	A10_FP16_Performance.arch	67165	1489	1162
Optimized	A10_FP16_Performance.arch	167981	2670	1518
Unoptimized	A10_Generic.arch	26517	627	202
Optimized	A10_Generic.arch	134696	2649	1518
Unoptimized	A10_Performance.arch	47259	953	650
Optimized	A10_Performance.arch	136094	2649	1518
Unoptimized	AGX7_FP16_Generic.arch	30633	508	186
Optimized1	AGX7_FP16_Generic.arch	268240	6829	4362
Optimized2	AGX7_FP16_Generic.arch	268898	11949	4362
Unoptimized	AGX7_FP16_Performance.arch	90186	1490	1162
Optimized1	AGX7_FP16_Performance.arch	268319	6827	4362
Optimized2	AGX7_FP16_Performance.arch	268977	11947	4362
Unoptimized	AGX7_Generic.arch	34987	769	202
Optimized1	AGX7_Generic.arch	158282	5324	2314
Optimized2	AGX7_Generic.arch	158589	9164	2314
Unoptimized	AGX7_Performance.arch	62451	1237	650
Optimized1	AGX7_Performance.arch	158292	5324	2314
Optimized2	AGX7_Performance.arch	158599	9164	2314

1. Investigating optimization processes to understand the performance disparity between ResNet-50 (10x gain) and trigger models (1.5x gain).
2. Exploring multi-core FPGA configurations to address scaling inefficiencies in AGX7 deployments.
3. Expanding the analysis to encompass a wider range of neural network models and hardware configurations.

In conclusion, while FPGAs, especially the Agilex 7, show great potential for CTA's real-time triggering, ongoing optimization of hardware-software integration remains crucial for maximizing their performance in resource-constrained, high-performance environments.

Bibliography

- [1] W. Hofmann and R. Zanin, “The cherenkov telescope array,” 2023.
- [2] U. Schwanke, M. Shayduk, K.-H. Sulanke, S. Vorobiov, and R. Wischnewski, “A versatile digital camera trigger for telescopes in the cherenkov telescope array,” *Nuclear Instruments and Methods in Physics Research Section A: Accelerators, Spectrometers, Detectors and Associated Equipment*, vol. 782, p. 92–103, May 2015.
- [3] “Target: A digitizing and trigger asic for the cherenkov telescope array,” in *AIP Conference Proceedings*, Author(s), 2017.
- [4] J. Zhang, R. Zhou, S. Zhang, J. Zhang, H. Xiong, G. Hu, Y. Li, Z. Liu, and C. Yang, “Data acquisition and trigger system for imaging atmospheric cherenkov telescopes of the lhaaso,” *Journal of Instrumentation*, vol. 15, p. T02004, feb 2020.
- [5] D. Nieto, A. Brill, Q. Feng, T. B. Humensky, B. Kim, T. Miener, R. Mukherjee, and J. Sevilla, “Ctlearn: Deep learning for gamma-ray astronomy,” 2019.
- [6] P. Grespan, M. Jacquemont, R. López-Coto, T. Miener, D. Nieto-Castaño, and T. Vuillaume, “Deep-learning-driven event reconstruction applied to simulated data from a single large-sized telescope of cta,” 2021.
- [7] A. Nechi, L. Groth, S. Mulhem, F. Merchant, R. Buchty, and M. Berekovic, “Fpga-based deep learning inference accelerators: Where are we standing?,” *ACM Transactions on Reconfigurable Technology and Systems*, vol. 16, 09 2023.
- [8] Z. Szadkowski and A. Szadkowska, “Fpga implementation of the trigger based on a fuzzy logic for a detection of cosmic rays in water cherenkov detectors,” pp. 431–435, 12 2019.
- [9] P. Reddy and K. Ramanaiah, “Field-programmable gate array implementation of efficient deep neural network architecture,” *International Journal of Electrical and Computer Engineering (IJECE)*, vol. 14, p. 3863, 08 2024.