

PAPER • OPEN ACCESS

Real-time semantic segmentation on FPGAs for autonomous vehicles with hls4ml

To cite this article: Nicolò Ghielmetti *et al* 2022 *Mach. Learn.: Sci. Technol.* **3** 045011

View the [article online](#) for updates and enhancements.

You may also like

- [Gradients should stay on path: better estimators of the reverse- and forward KL divergence for normalizing flows](#)
Lorenz Vaitl, Kim A Nicoli, Shinichi Nakajima et al.
- [‘Flux+Mutability’: a conditional generative approach to one-class classification and anomaly detection](#)
C Fanelli, J Giroux and Z Papandreou
- [Operationally meaningful representations of physical systems in neural networks](#)
Hendrik Poulsen Nautrup, Tony Metger, Raban Iten et al.



PAPER

OPEN ACCESS

RECEIVED
25 May 2022REVISED
27 August 2022ACCEPTED FOR PUBLICATION
21 October 2022PUBLISHED
4 November 2022

Original Content from
this work may be used
under the terms of the
[Creative Commons
Attribution 4.0 licence](#).

Any further distribution
of this work must
maintain attribution to
the author(s) and the title
of the work, journal
citation and DOI.



Real-time semantic segmentation on FPGAs for autonomous vehicles with hls4ml

Nicolò Ghielmetti^{1,8} , Vladimir Loncar^{1,9} , Maurizio Pierini¹ , Marcel Roed^{1,10}, Sioni Summers¹,
Thea Aarrestad² , Christoffer Petersson^{3,11}, Hampus Linander⁴, Jennifer Ngadiuba⁵ , Kelvin Lin^{6,12}
and Philip Harris^{7,*}

¹ European Organization for Nuclear Research (CERN), CH-1211 Geneva 23, Switzerland

² Institute for Particle Physics and Astrophysics, ETH Zürich, 8093 Zürich, Switzerland

³ Zenseact, Gothenburg 41756, Sweden

⁴ University of Gothenburg, Gothenburg 40530, Sweden

⁵ Fermi National Accelerator Laboratory, Batavia, IL 60510, United States of America

⁶ University of Washington, Seattle, WA 98195, United States of America

⁷ Massachusetts Institute of Technology, Cambridge, MA 02139, United States of America

⁸ Also at Politecnico di Milano, Italy.

⁹ Also at Institute of Physics Belgrade, Serbia.

¹⁰ Also at University of Oxford, United Kingdom.

¹¹ Also at Chalmers University of Technology, Sweden.

¹² Currently at Amazon, United States of America.

* Author to whom any correspondence should be addressed.

E-mail: pcharris@mit.edu

Keywords: FPGA, computer vision, deep learning, hls4ml, machine learning, autonomous vehicles, semantic segmentation

Abstract

In this paper, we investigate how field programmable gate arrays can serve as hardware accelerators for real-time semantic segmentation tasks relevant for autonomous driving. Considering compressed versions of the ENet convolutional neural network architecture, we demonstrate a fully-on-chip deployment with a latency of 4.9 ms per image, using less than 30% of the available resources on a Xilinx ZCU102 evaluation board. The latency is reduced to 3 ms per image when increasing the batch size to ten, corresponding to the use case where the autonomous vehicle receives inputs from multiple cameras simultaneously. We show, through aggressive filter reduction and heterogeneous quantization-aware training, and an optimized implementation of convolutional layers, that the power consumption and resource utilization can be significantly reduced while maintaining accuracy on the Cityscapes dataset.

1. Introduction

Deep Learning has strongly reshaped computer vision in the last decade, bringing the accuracy of image recognition applications to unprecedented levels. Improved pattern recognition capabilities have had a significant impact on the advancement of research in science and technology. Many of the challenges faced by future scientific experiments, such as the CERN High Luminosity Large Hadron Collider [1] or the Square Kilometer Array observatory [2], and technological challenges faced by, for example, the automotive industry, will require the capability of processing large amounts of data in real-time, often through edge computing devices with strict latency and power-consumption constraints. This requirement has generated interest in the development of energy-effective neural networks, resulting in efforts like tinyML [3], which aims to reduce power consumption as much as possible without negatively affecting the model accuracy.

Advances in deep learning for computer vision have had a crucial impact on the development of autonomous vehicles, enabling the vehicles to perceive their environment at ever-increasing levels of accuracy and detail. Deep neural networks are used for finding patterns and extracting relevant information from camera images, such as the precise location of the surrounding vehicles and pedestrians. In order for an autonomous vehicle to drive safely and efficiently, it must be able to react fast and make quick decisions.

This imposes strict latency requirements on the neural networks that are deployed to run inference on resource-limited embedded hardware in the vehicle.

In addition to algorithmic development, computer vision for autonomous vehicles has benefited from technological advances in parallel computing architecture [4]. The possibility of performing network training and inference on graphics processing units (GPUs) has made large and complex networks computationally affordable and testable on real-life problems. Due to their high efficiency, GPUs have become a common hardware choice in the automotive industry for on-vehicle deep learning inference.

Going beyond GPUs, as embedded computer vision-based systems are being developed and deployed in an emerging number of industries, including automotive, healthcare and surveillance, deep learning hardware accelerators are also utilizing field-programmable gate arrays (FPGAs) and application specific integrated circuits (ASICs). Such accelerators are also exploiting the possibility to parallelize the vast number of operations, thereby reducing latency and increasing throughput. In contrast to ASICs, for which the design cannot be modified upon fabrication, FPGAs are flexible and reconfigurable, and commonly used for prototyping and validating ASIC implementations. To fairly compare the performance and efficiency between GPUs and FPGAs is very difficult due to the strong dependence on the details of, for example, the implementation, open software support, data transfers and hardware design. See [5] for a recent survey on efficient semantic segmentation networks using GPUs.

In this paper, we investigate the possibility of exploiting FPGAs as a low-power, inference-optimized alternative to GPUs. By applying aggressive filter-reduction and quantization of the model bit precision at training time, and by introducing a highly optimized firmware implementation of convolutional layers, we achieve the compression required to fit semantic segmentation models on FPGAs. We do so by exploiting and improving the `hls4ml` library, which provides an automatic conversion of a given Deep Neural Network into C++ code, which is given as input to a high level synthesis (HLS) library. The HLS library then translates this into FPGA firmware, to be deployed on hardware. Originally developed for scientific applications in particle physics that require sub-microsecond latency [6–12], `hls4ml` has been successfully applied outside the domain of scientific research [13, 14], specifically in the context of tinyML applications [15].

Applying model compression at training time is crucial in order to minimize resource-consumption and maximize the model accuracy. To do so, we rely on quantization-aware training (QAT) through the QKeras [16] library, which has been interfaced to `hls4ml` in order to guarantee an end-to-end optimal training-to-inference workflow [13].

As a baseline, we start from the ENet [17] architecture, designed specifically to perform pixel-wise semantic segmentation for tasks requiring low latency operations. We modify the architecture, removing resource-consuming asymmetric convolutions, and dilated or strided convolutions. In addition, we apply filter ablation and quantization at training time. Finally, we optimize the implementation of convolutional layers in `hls4ml` in order to significantly reduce the resource consumption. With these steps, we obtain a good balance between resource utilization and accuracy, enabling us to deploy the whole network on a Xilinx ZCU102 evaluation board [18].

This paper is organized as follows: The baseline dataset and model are described in sections 2 and 3, respectively. The model compression and the specific optimization necessary to port the compressed model to the FPGA are described in sections 4 and 5. Conclusions are given in section 6.

2. Dataset

Our experiments are performed using the Cityscapes dataset [19], which involves 5000 traffic scene images collected in 50 different cities with varying road types and seasons. These images have fine-grained semantic segmentation annotations with pixel-level classification labels. We have limited ourselves to the four semantic classes Road, Car, Person and Background. According to the standard Cityscapes split, 2975 images are used for training, 500 for validation and 1525 for testing. We crop and resize the original images to have an input resolution of 240×152 pixels. As a pre-processing step, we normalize all pixel values (integer values in the range $[0, 255]$) to be in the $[0, 1]$ range by dividing each one by 256. In this way all inputs are smaller than one and can be represented by a fixed-point datatype using only 8 bits ($\log_2(256)$) (see section 4). An example image from the dataset is shown in figure 1, together with a visualization of its semantic segmentation mask.

For evaluation metrics we use two typical figures of merit for semantic segmentation:

- The model accuracy (Acc), defined as $\text{Acc} = \frac{TP+TN}{TP+TN+FP+FN}$, where TP , TN , FP , and FN are the fraction of true positives, true negatives, false positives, and false negatives, respectively.
- The mean of the class-wise Intersection over Union (mIoU), i.e. the average across classes of the Intersection-Over-Union (defined as $\text{IOU} = \frac{TP}{TP+FP+FN}$).



Figure 1. An downsampled image from the Cityscapes dataset (left) and the corresponding semantic segmentation target (right), in which the pixels belong to one of the classes {background (blue), road (teal), car (yellow), person (red)}.

Table 1. Model architecture parametrized by the number of filters in the bottlenecks f_i , with $i = 1, \dots, 5$.

Layer	Type	Output resolution
Initial	Downsample	$f_0 \times 120 \times 76$
3× bottleneck 1	Downsample	$f_1 \times 60 \times 38$
3× bottleneck 2	Downsample	$f_2 \times 30 \times 19$
3× bottleneck 3	Downsample	$f_3 \times 30 \times 19$
3× bottleneck 4	Upsample	$f_4 \times 60 \times 38$
3× bottleneck 5	Upsample	$f_5 \times 120 \times 76$
Final	Upsample	$4 \times 240 \times 152$

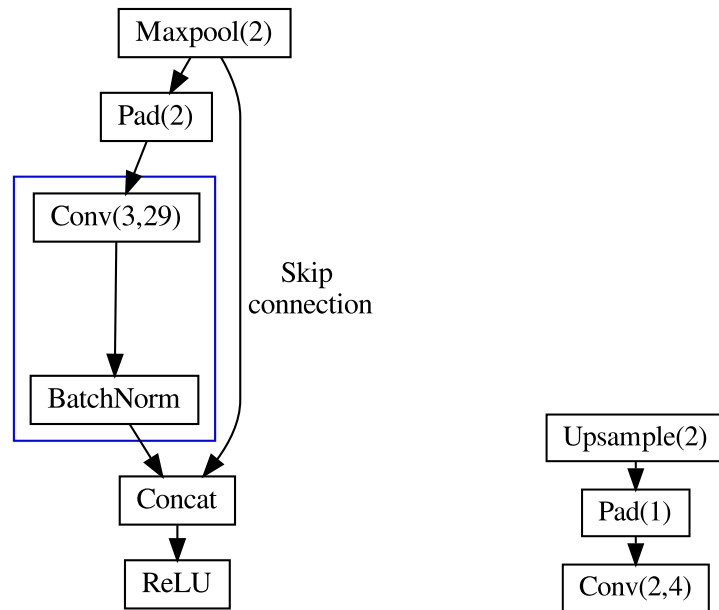
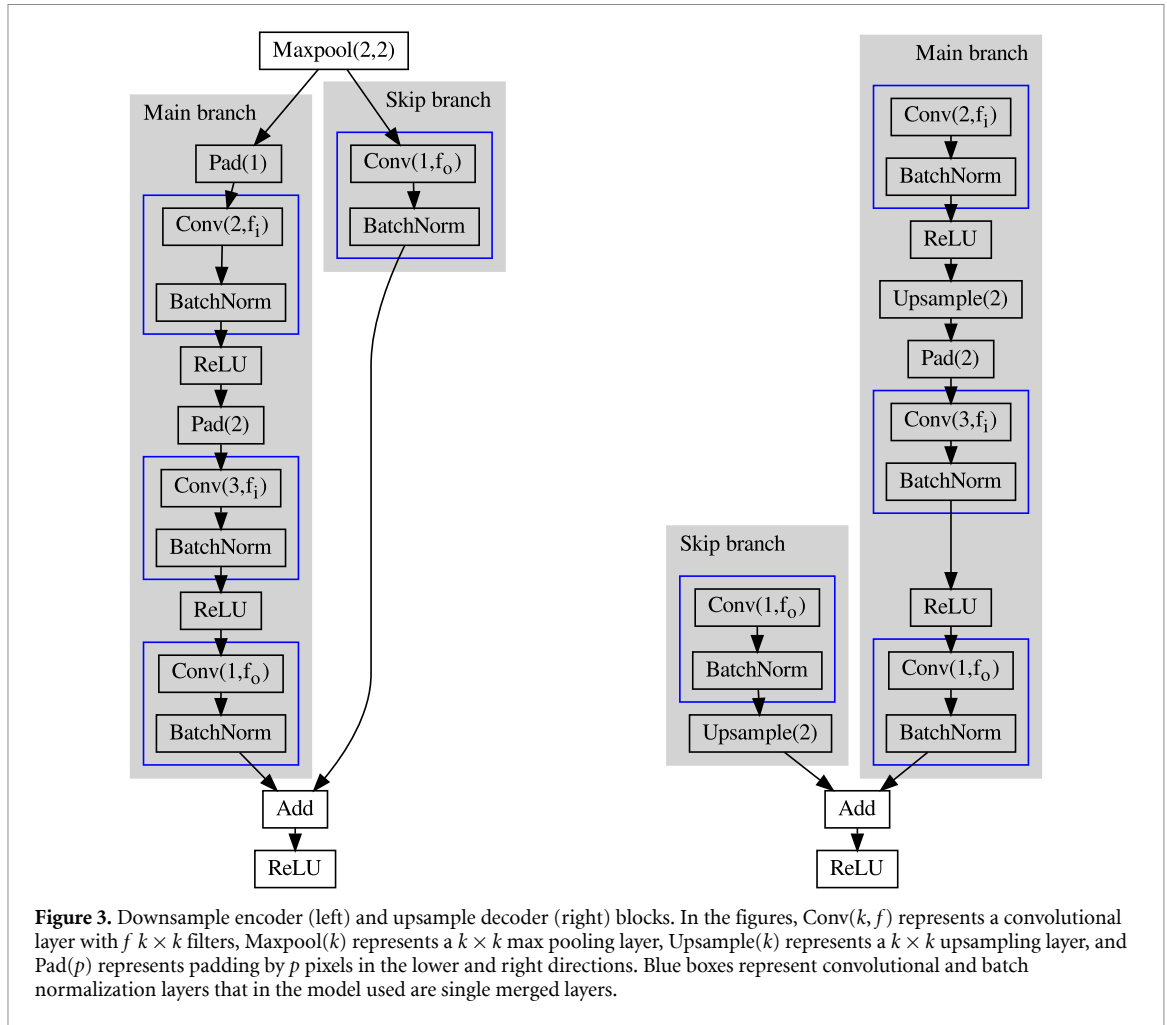


Figure 2. Initial (left) and final (right) block architecture. In the two diagrams, $\text{conv}(k, f)$ represents a convolutional layer with f $k \times k$ filters; $\text{maxpool}(k)$ and $\text{upsample}(k)$ represent a $k \times k$ max pooling or upsample layer, respectively; and $\text{pad}(p)$ represents padding by p pixels in the lower and right directions.

3. Baseline model

The architecture we use is inspired by a fully convolutional residual network called Efficient Neural Network (ENet) [17]. This network was designed for low latency and minimal resource usage. It is designed as a sequence of blocks, summarized in table 1. The initial block, shown in the left figure in figure 2, encodes the input into a $32 \times 120 \times 76$ tensor, which is then processed by a set of sequential blocks of bottlenecks. The first three blocks constitute the *downsampling encoder*, where each block consists of a series of layers as



summarized in the left diagram in figure 3. The final two blocks provide an *upsampling decoder*, as illustrated in the right diagram in figure 3. The final block is shown in the right diagram of figure 2.

Some differences from the original architecture in [17] is that we do not use asymmetric, dilated, or strided convolutions. To further reduce the resource usage, we use three bottlenecks per block instead of five, and we merge convolutional layers with batch normalization layers by rescaling convolutional filter weights with batch normalization parameters (implemented through a QConv2DBatchnorm layer). When we use QAT, this allows us to directly quantize the merged weights during the forward pass, rather than quantizing the batch normalization parameters and the convolutional filters separately. This merging of layers saves resources on the FPGA, since only the merged weights are used. Performing the merging already during training, ensures that the weights used during training and during inference are quantized the same way. The baseline ENet model is obtained fixing the six f hyperparameters of table 1 to (32, 64, 64, 64, 128, 48). This choice results in an architecture with 1.1×10^6 parameters, yielding a $\text{mIoU} = 63.2\%$ and an accuracy of 91.5%.

Note that, in a real world application, such as autonomous driving, this computer vision task of performing frame-by-frame semantic segmentation constitutes only a sub-task in the full software stack. For example, the outputs of this network will typically need to be transformed into 3d world coordinates and tracked over time before being sent to the planning and decision-making modules. Hence, the single-frame-based metrics, such as accuracy and mIoU , are primarily used to compare different semantic segmentation models, but they are insufficient for gauging the performance of the full autonomous driving stack in real world driving.

4. Model compression

We consider two compression techniques for the model at hand: filter-wise homogeneous pruning, obtained by reducing the number of filters on all the convolutional layers; and quantization, i.e. reducing the number

Table 2. Architecture reduction through internal filter ablation and corresponding performance. As a reference, the baseline architecture is reported on the first row. Highlighted in bold the three models considered further in this work.

Model name	f_i	f_1	f_2	f_3	f_4	f_5	Parameters	mIoU (%)	Accuracy (%)
Enet	32	64	64	64	128	48	1.1×10^6	63.2	91.5
Enet16	32	16	16	16	16	16	5×10^4	54.3	87.9
Enet12	32	12	12	12	12	12	3×10^4	52.0	86.8
Enet8	32	8	8	8	8	8	1.4×10^4	49.4	85.6
Enet6	32	6	6	6	6	6	9×10^3	45.9	84.0
Enet4	32	4	4	4	4	4	5×10^3	36.6	81.5

of bits allocated for the numerical representation of the network components and the output of each layer computation.

In addition, we use the AutoQKeras library [13], distributed with QKeras, to optimize the numerical representation of each component at training time as a hyperparameter. This is done using a mathematical model of the inference power consumption as a constraint in the loss function.

4.1. Filter multiplicity reduction

Normally, network pruning consists of zeroing specific network parameters that have little impact on the model performance. This could be done at training time or after training. In the case of convolutional layers, a generic pruning of the filter kernels would result in sparse kernels. It would then be difficult to take advantage of pruning during inference. To deal with this, filter ablation (i.e. the removal of an entire kernel) was introduced [20]. When filter ablation is applied, one usually applies a restructuring algorithm (e.g. Keras Surgeon [21]) to rebuild the model into the smaller-architecture model that one would use at inference. In this work, we take a simpler (and more drastic) approach: we treat the number of filters in the convolutional layers as a single hyperparameter, fixed across the entire network. We then reduce its value and repeat the training, looking for a good compromise between accuracy and resource requirements.

We repeat the procedure with different target filter multiplicities. The result of this procedure is summarized in table 2, where different pruning configurations are compared to the baseline Enet model.

Out of these models, we select two configurations that would be affordable on the FPGA at hand: a four-filters (Enet4) and an eight-filter (Enet8) configuration. As a reference for comparison, we also consider one version with 16 filters, Enet16, despite it being too large to be deployed on the FPGA in question. We then proceed by quantizing these models through QAT to further reduce the resource consumption.

4.2. Homogeneous quantization-aware training

Homogeneous QAT consists of repeating the model training while forcing the numerical representation of its weight and activation functions to a fixed $\langle T, I \rangle$ precision, where T is the total number of bits and I is the number of integer bits. This is done using the *straight-through estimator*, where quantization functions are applied to weights and activations during the forward pass of the training, but then assuming the quantization is the identity function in the backward pass, as the quantization function is not differentiable. The model training then converges to a minimum that might not be the absolute minimum of the full-precision training, but that would minimize the performance loss once quantization is applied. For a complete overview on quantization methods for neural networks, see [22].

In practice, we perform a homogeneous QAT replacing each layer of the model with its QKeras equivalent and exploiting the existing QKeras-to-hls4ml interface for FPGA deployment.

We study the impact of QAT for $T \in 2, 4, 8$ with $I = 0$, on the pruned models described above (Enet4, Enet8 and Enet16). The resulting performance is shown in table 3, where we label the three quantization configurations as Q2, Q4 and Q8, respectively.

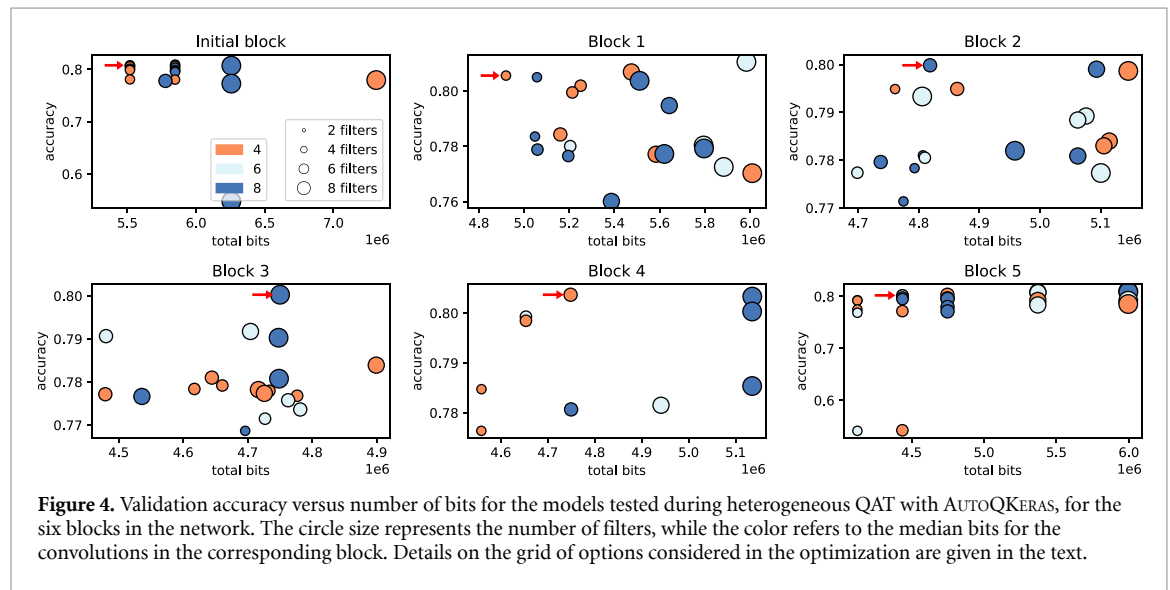
The resulting resource utilization for Enet4 and Enet8 falls within the regime of algorithms that we could deploy on the target FPGA. We observe similar drops in accuracy when going from full precision to Q8 and from Q4 to Q2, but little differences between the Q4 and Q8 models. In this respect, Q4 would offer a better compromise between accuracy and resources than Q8.

Out of these, the models with the highest accuracy and mIoU that would be feasible to fit on the FPGA, is the 8 filter model quantized to 8 bits (Enet8Q8) and the 8 filter model quantized to 4 bits (Enet8Q4).

The quantization of the model does not have to be homogeneous across layers. In fact, it has been demonstrated that a heterogeneous quantization is the best way to maintain high accuracy at low resource-cost [23]. We therefore define one final model with an optimized combination of quantizers.

Table 3. Homogeneously and heterogeneously quantized models with indicated bitwidth, filter architecture and number of parameters, together with their validation mean IOU trained with quantization aware training using QKeras. The corresponding values before quantization (from table 2) are also reported in the three first rows.

Model name	Quantization	f_i	f_1	f_2	f_3	f_4	f_5	Parameters	mIoU (%)	Accuracy (%)
Enet16	—	32	16	16	16	16	16	5×10^4	54.3	87.9
Enet8	—	32	8	8	8	8	8	1.4×10^4	49.4	85.6
Enet4	—	32	4	4	4	4	4	5×10^3	36.6	81.5
Enet16Q8	8	32	16	16	16	16	16	5×10^4	35.0	79.1
Enet8Q8	8	32	8	8	8	8	8	1.4×10^4	33.4	77.1
Enet4Q8	8	32	4	4	4	4	4	5×10^3	13.6	53.8
Enet16Q4	4	32	16	16	16	16	16	5×10^4	34.1	77.9
Enet8Q4	4	32	8	8	8	8	8	1.4×10^4	33.9	77.6
Enet4Q4	4	32	4	4	4	4	4	5×10^3	13.5	53.6
Enet16Q2	2	32	16	16	16	16	16	5×10^4	27.4	68.6
Enet8Q2	2	32	8	8	8	8	8	1.4×10^4	28.7	71.1
Enet4Q2	2	32	4	4	4	4	4	5×10^3	13.4	53.5
EnetHQ	Heterogeneous	8	2	4	8	4	3	5.3×10^3	36.8	81.1



4.3. Heterogeneous quantization aware training

Heterogeneous QAT consists in applying different quantization to different network components. For deep networks, one typically deals with the large number of possible configurations by using an optimization library. In our case, we use AUTOQKERAS [13]. In AUTOQKERAS, a hyperparameter search over individual layer quantization conditions and filter counts is performed. Since the model contains skip connections, the scan over number of filters needs to be handled with care. In particular, we use the block features of AUTOQKERAS to ensure that the filter count matches throughout a bottleneck, so that the tensor addition of the skip connection will have valid dimensions.

The search for best hyperparameters, including the choice of individual quantizers for kernels and activations, is carried out using a Bayesian strategy where the balance between accuracy and resource usage is controlled by targeting a metric derived from them both [13]. In our search we permit e.g. a 4% decrease in accuracy if the resource usage also is halved at the same time.

The hyperparameter scan is done sequentially over the blocks, i.e. the Bayesian search over quantization and filter count of the initial layer is performed first and is then frozen for the hyperparameter scan of the first bottleneck and so on. The rest of the model is kept in floating point until everything in the end is quantized.

Figure 4 shows the outcome of the heterogeneous QAT, in terms of validation accuracy and total number of bits for the six blocks in the network. The optimal configuration search is performed taking as a baseline the Enet4 model, scanning the kernel bits in $\{4, 8\}$ and fixing the number of kernels to four times a by-layer multiplicative chosen in $\{0.5, 0.75, 1.0, 1.25, 1.5, 1.75, 2.0\}$. The optimal configuration (EnetHQ) is obtained for $f_i = 8, f_1 = 2, f_2 = 4, f_3 = 8, f_4 = 4$, and $f_5 = 3$, resulting in 4.7×10^3 parameters, a mIoU = 36.8% and an accuracy of 81.1%. Out of all the quantized models, both homogeneous and heterogeneous, this is the one which performs the best.

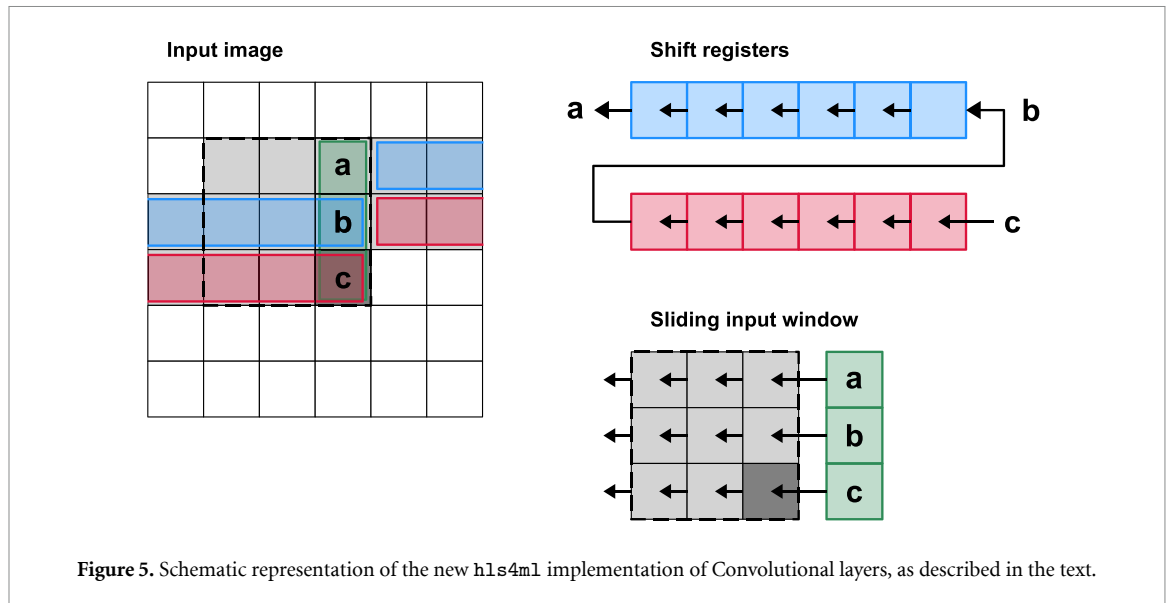


Table 4. Comparison of previous and proposed hls4ml implementation of the convolutional layer, in terms of relative reduction of resource utilization. The estimates are made targeting an xczu9eg-2ffvb1156 MPSoC device on a ZCU102 development kit.

Implementation	BRAM	DSP	FF	LUT
Encoded [14]	4752	5632	195 344	291 919
Line buffer	4064	5632	176 620	305 494
Improvement	−15%	0%	−1%	+5%

5. FPGA implementation, deployment and results

5.1. Resource-efficient convolution algorithm

The hls4ml library has an implementation of convolutional layers that is aimed at low-latency designs [14]. However, this implementation comes at the expense of high resource utilization. This is due to the number of times pixels of the input image are replicated to maintain the state of a sliding input window. For convolutional layers operating on wider images, like in our case, this overhead can be prohibitively large. In order to reduce the resource consumption of the convolutional layers of the model, we introduce a new algorithm that is more resource efficient.

The new implementation, dubbed ‘line buffer’, uses shift registers to keep track of previously seen pixels. The primary advantage of the line buffer implementation over the previous one is the reduction of the size of the buffer needed to store the replicated pixels. For an image of size $H \times W$, with a convolution kernel of size $K \times L$, the line buffer allocates $K - 1$ buffers (chain of shift registers) of depth W for the rows of the image, while the previous implementation allocates K^2 buffers of depth $K \times (W - K + 1)$ for the elements in the sliding input window.

The algorithm is illustrated on figure 5. Initially, each new pixel read from the input image stream is pushed into the shift register chain. If the shift register is full, the first element will be popped and it will be pushed into the next shift register in chain. The process is repeated for all $K - 1$ shift registers in the chain. The popped pixels are stacked with the input pixel into a column vector and are pushed as the rightmost column of the input window. The pixels popped from the leftmost column of the input window are not used further. In our implementation, the propagation of new pixels through the shift register chain and the insertion into the sliding input window are completed in a single clock cycle, making the implementation as efficient as the existing hls4ml implementation.

To compute the output from the populated sliding input window, we rely on the existing routines of hls4ml. We rely on a set of counters to keep track of input window state to know when to produce an output. The algorithm for maintaining the chain of shift registers and populating the sliding input window can be adapted for use in the pooling layers as well.

To compare the two implementations, we consider the resource utilization of an ENet bottleneck block consisting of 8 filters, implemented using either method. The results are summarized in table 4. We observe a substantial reduction in block random access memory (BRAM) usage, at the price of a small increase in look-up table (LUT) utilization.

Table 5. Effect of FIFO depth optimization on FPGA resource usage and model latency. The values in the table are taken from Vivado HLS estimates of resource usage. A comparison using physical resource usage is unfeasible since the model without optimization cannot be synthesized. The estimates are made targeting an xczu9eg-2ffvb1156 MPSoC device on a ZCU102 development kit.

Optimisation	BRAM	LUT	FF	DSP	Latency
No	7270	676 760	230 913	228	3.577 ms
Yes	1398	437 559	146 392	228	3.577 ms
Improvement	−81%	−35%	−37%	0%	0%

Table 6. Accuracy, mIoU, latency and resource utilization for the EnetHQ, Enet8Q4 and Enet8Q8 models. The latency is quoted for a batch size $b = 1$ and $b = 10$. Resources are expressed as a percentage of those available on the xczu9eg-2ffvb1156 MPSoC device on the ZCU102 development kit. The last row is a comparison to work presented in [24].

Model	Acc.	mIoU	Latency (ms)		BRAM	LUT	FF	DSP
			$b = 1$	$b = 10$				
EnetHQ	81.1%	36.8 %	4.9	30.6	224.5 (25%)	76 718 (30%)	87 059 (16%)	450 (18%)
Enet8Q4	77.6%	33.9 %	4.8	30.2	342.0 (37%)	166 741 (61%)	90 536 (16%)	0
Enet8Q8	77.1%	33.4 %	4.8	30.0	508.5 (56%)	126 458 (46%)	134 385 (25%)	1502 (60%)
ENet [24]	—	63.1%	30.38 (720) ^a	—	257	62 599	192 212	689

^a The former is without considering data transfer, pre- and post-processing. The number in parenthesis includes these additional overheads, averaged over 58 images, and is more comparable to the numbers we present.

5.2. FIFO depth optimization

With the dataflow compute architecture of hls4ml, layer compute units are connected with FIFOs, implemented as memories in the FPGA. These FIFOs contribute to the overall resource utilisation of the design. The read and write pattern of these FIFOs depends on the dataflow through the model, which is not predictable before the design has been scheduled by the HLS compiler, and is generally complex. With previous hls4ml releases, these memories have therefore been assigned a depth corresponding to the dimensions of the tensor in the model graph as a safety precaution.

To optimize this depth and thereby reduce resource consumption, we implemented an additional step in the compilation of the model to hardware. By using the clock-cycle accurate RTL simulation of the scheduled design, we can monitor the actual occupancy of each FIFO in the model when running the simulation over example images. This enables us to extract and set the correct size of the FIFOs, reducing memory usage compared to the baseline.

By applying this procedure, we observe a memory efficiency $\frac{\sum_l O_l}{\sum_l F_l} = 19.5\%$, where the index l runs across the layers, O_l is the observed occupancy for the l th layer, and F_l is the corresponding FIFO depth. The corresponding mean occupancy is found to be $\sum_l \frac{O_l}{F_l} = 4.7\%$.

We then resize every FIFO to its observed maximum occupancy and rerun the C-Synthesis, thereby saving FPGA resources and allowing larger models to fit on the FPGA. Table 5 shows the impact of such an optimization on the FPGA resources for one example model, Enet8Q8, demonstrating a significant reduction of resources, which are BRAM, LUT, flip-flop (FF), digital signal processor (DSP).

5.3. Results

The hardware we target is a Zynq UltraScale+ MPSoC device (xczu9eg-2ffvb1156) on a ZCU102 development kit, which targets automotive applications. After reducing the FPGA resource consumption through the methods described above, the highest accuracy models highlighted in table 3 are synthesized. These are the homogeneously quantized Enet8Q8 and Enet8Q4 models, as well as the heterogeneously quantized EnetHQ model. To find the lowest latency implementation, we run several attempts varying the reuse factor (RF) and the clock period. The RF indicates how many times a multiplier can be reused (zero for a fully parallel implementation). Lower RF leads to lower latency, but higher resource usage. We targeted reuse factors of 10, 20, 50, 100, and clock periods of 5, 7, 10 ns. For each model, we then chose the configuration yielding the lowest latency. For Enet8Q8, this is a target clock period of 7 ns and RF = 10. For Enet8Q4 and EnetHQ we use a clock period of 7 ns and RF = 6.

Inference performance of this model was measured on the ZCU102 target device. The final latency and resource utilization report is shown in table 6.

We measured the time taken by the accelerator to produce a prediction on batches of images, with batch sizes of $b = 1$ and $b = 10$. The same predictions have been executed 10^5 times, and the time average is taken as the latency. The single image latency (batch size of 1) is 4.8–4.9 ms for all three models. Exploiting the data

flow architecture, the latency to process images in a batch size of 10 is less than 10 times the latency observed for a batch size of 1. While in a real-world deployment of this model the latency to return the predictions of a single image is the most important metric, a system comprised of multiple cameras may be able to benefit from the speedup of batched processing by batching over the images captured simultaneously from different cameras. The model with the highest accuracy and lowest resource consumption is the heterogeneously quantized EnetHQ model. This model has an mIoU of 36.8% and uses less than 30% of the total resources.

Similar work on performing semantic segmentation on FPGAs include [24] and a comparison is given in table 6. Here, the original ENet model [17] is trained and evaluated on the Cityscapes dataset, and then deployed on a Xilinx Zynq 7035 FPGA using the Xilinx Vitis AI Deep Learning Processor Unit (DPU). There are some crucial differences between the approach taken here and that of [24]. In order to achieve the lowest possible latency, we implement a fully on-chip design with high layer parallelism. We optimize for latency, rather than frame rate, such that in a real-life application the vehicle response time could be minimized. Keeping up with the camera frame rate is a minimal requirement, but a latency lower than the frame interval can be utilized. In our approach, each layer is implemented as a separate module and data is streamed through the architecture layer by layer. Dedicated per-layer buffers ensure that just enough data is buffered in order to feed the next layer. This is highly efficient, but limits the number of layers that can be implemented on the FPGA. Consequently, in order to fit onto the FPGA in question, our model is smaller and achieves a lower mIoU. Jia *et al* [24] does not quote a latency, but a frame rate. A best-case latency is then computed as the inverse of this frame rate, which corresponds to 30.38 ms. However, this does not include any overhead latency like data transfer, pre- and post-processing. Including these, the average time per image increases to 720 ms.

6. Conclusions

In this paper, we demonstrate that we can perform semantic segmentation on a single FPGA on a Zynq MPSoC device using a compressed version of ENet. The network is compressed using automatic heterogeneous quantization at training time and a filter ablation procedure, and is then evaluated on the Cityscapes dataset. Inference is executed on hardware with a latency of 4.9 ms per image, utilizing 18% of the DSPs, 30% of the LUTs, 16% of the FFs and 25 % of the BRAMs. Processing the images in batches of ten results in a latency of 30 ms per batch, which is significantly faster than ten times the single-image batch inference latency. This is relevant when batching over images captured from different cameras simultaneously. By introducing an improved implementation of convolutional layers in `hls4ml`, we significantly reduce resource consumption, allowing for a fully-on-chip deployment of larger convolutional neural networks. This avoids latency overhead caused by data transfers between off-chip memory and FPGA processing elements, or between multiple devices. Also taking into account the favorable power-efficiency of FPGAs, we conclude that FPGAs offer highly interesting, low-power alternatives to GPUs for on-vehicle deep learning inference and other computer vision tasks requiring low-power and low-latency.

Data availability statement

All data that support the findings of this study are included within the article (and any supplementary files).

Acknowledgments

We acknowledge the Fast Machine Learning collective as an open community of multi-domain experts and collaborators. This community was important for the development of this project. M P, M R, S S and V L are supported by the European Research Council (ERC) under the European Union's Horizon 2020 research and innovation program (Grant Agreement No. 772369). M P and N G are supported by the European Research Council (ERC) under the European Union's Horizon 2020 research and innovation program (Grant Agreement No. 966696).

ORCID iDs

Nicolò Ghielmetti  <https://orcid.org/0000-0002-4660-9757>

Vladimir Loncar  <https://orcid.org/0000-0003-3651-0232>

Maurizio Pierini  <https://orcid.org/0000-0003-1939-4268>

Thea Aarrestad  <https://orcid.org/0000-0002-7671-243X>

Jennifer Ngadiuba  <https://orcid.org/0000-0002-0055-2935>

Philip Harris  <https://orcid.org/0000-0001-8189-3741>

References

- [1] Apollinari G et al 2017 *High-Luminosity Large Hadron Collider (HL-LHC): Technical Design Report V. 0.1* CERN yellow reports: monographs (Geneva: CERN) (<http://dx.doi.org/10.23731/CYRM-2017-004>)
- [2] Garrett M A et al 2010 (arXiv:1008.2871)
- [3] Banbury C R et al 2021 Benchmarking tinyml systems: challenges and direction (arXiv:2003.04821)
- [4] Raina R, Madhavan A and Ng A Y 2009 Large-scale deep unsupervised learning using graphics processors *Proc. 26th Annual Int. Conf. on Machine Learning* (ACM) pp 873–80
- [5] Holder C J and Shafique M 2022 On efficient real-time semantic segmentation: a survey (arXiv:2206.08605)
- [6] Duarte J et al 2018 Fast inference of deep neural networks in FPGAs for particle physics *J. Instrum.* **13** 07027
- [7] Summers S et al 2020 Fast inference of boosted decision trees in FPGAs for particle physics *J. Instrum.* **15** 05026
- [8] Loncar V et al 2021 Compressing deep neural networks on FPGAs to binary and ternary precision with hls4ml *Mach. Learn.: Sci. Technol.* **2** 015001
- [9] Iiyama Y et al 2021 Distance-weighted graph neural networks on fpgas for real-time particle reconstruction in high energy physics *Front. Big Data* **3** 598927
- [10] Heintz A et al 2020 Accelerated charged particle tracking with graph neural networks on FPGAs *3rd Machine Learning and the Physical Sciences Workshop at the 34th Annual Conf. on Neural Information Processing Systems* vol 12
- [11] Francescato S Giagu S, Riti F, Russo G, Sabetta L and Tortonesi F 2021 Model compression and simplification pipelines for fast deep neural network inference in FPGAs in HEP *Eur. Phys. J. C* **81** 969
Francescato S Giagu S, Riti F, Russo G, Sabetta L and Tortonesi F 2021 *Eur. Phys. J. C* **81** 1064 (erratum)
- [12] Sun C, Nakajima T, Mitsumori Y, Horii Y and Tomoto M 2022 Fast muon tracking with machine learning implemented in FPGA (arXiv:2202.04976)
- [13] Coelho C N Kuusela A, Li S, Zhuang H, Ngadiuba J, Aarrestad T K, Loncar V, Pierini M, Pol A A and Summers S 2021 Automatic heterogeneous quantization of deep neural networks for low-latency inference on the edge for particle detectors *Nat. Mach. Intell.* **3** 675–86
- [14] Aarrestad T et al 2021 Fast convolutional neural networks on FPGAs with hls4ml *Mach. Learn.: Sci. Technol.* **2** 045015
- [15] Fahim F et al 2021 hls4ml: an open-source codesign workflow to empower scientific low-power machine learning devices *TinyML Research Symp. 2021* vol 3
- [16] Coelho C et al 2019 Qkeras (available at: <https://github.com/google/qkeras>)
- [17] Paszke A, Chaurasia A, Kim S and Culurciello E 2016 Enet: a deep neural network architecture for real-time semantic segmentation (arXiv:1606.02147)
- [18] Xilinx ZCU102 evaluation board (available at: www.xilinx.com/products/boards-and-kits/ek-u1-zcu102-g.html#information)
- [19] Cordts M et al 2016 The cityscapes dataset for semantic urban scene understanding (arXiv:1604.01685)
- [20] Girshick R, Donahue J, Darrell T and Malik J 2014 Rich feature hierarchies for accurate object detection and semantic segmentation (arXiv:1311.2524)
- [21] Whetton B 2016 Keras surgeon (available at: <https://github.com/BenWhetton/keras-surgeon>)
- [22] Gholami A et al 2021 A survey of quantization methods for efficient neural network inference (arXiv:2103.13630)
- [23] Coelho C N, Kuusela A, Li S, Zhuang H, Ngadiuba J, Aarrestad T K, Loncar V, Pierini M, Pol A A and Summers S 2021 Automatic heterogeneous quantization of deep neural networks for low-latency inference on the edge for particle detectors *Nat. Mach. Intell.* **3** 675–86
- [24] Jia W, Cui J, Zheng X and Wu Q 2021 Design and implementation of real-time semantic segmentation network based on FPGA *2021 7th Int. Conf. on Computing and Artificial Intelligence, ICCAI* (New York: Association for Computing Machinery) pp 321–5