

Rate-Distortion Optimized Quantization: A Deep Learning Approach

Thuong Nguyen Canh, Motong Xu, and Byeungwoo Jeon
School of Electrical and Computer Engineering, Sungkyunkwan University, Korea
{ngcthuong, xumotong, bjeon}@skku.edu

Abstract—Rate-Distortion Optimized Quantization (RDOQ) is a very effective video encoding tool in MPEG HEVC/H.265. It sequentially decides the optimum quantized levels of each transform coefficient in a given transform block in the sense of the least rate-distortion cost. The best level of each coefficient is searched among multiple candidate quantized levels. Since the optimal quantized level of one coefficient is affected by its context (i.e., its location and previous quantized levels), RDOQ is difficult to parallelize. This paper is the first attempt to explore possibility of using deep learning in HEVC quantization. We set up a machine learning problem for RDOQ which predicts corresponding RDOQ quantized output upon receiving scalar quantization (SQ) result of a block as input. A residual learning framework is employed to predict the difference of SQ from RDOQ after further simplification that the residual becomes binary. By using a deep convolution neural network, the proposed deep learning based RDOQ (DL-RDOQ) is able to predict the optimal quantized levels without computing rate and distortion. Our experiments show potentially promising performance following RDOQ, especially at high bitrates.

Keywords—Quantization, HEVC, rate-distortion optimized quantization, deep learning, convolutional neural network.

I. INTRODUCTION

The recent video compression standard of ISO/IEC MPEG HEVC/ITU-T H.265, more popularly known as High Efficiency Video Coding (HEVC), has improved coding performance of ITU-T H.264/MPEG AVC (Advanced Video Coding) by about 50% [1]. Similarly, to its predecessors, it first generates residual signal by either inter or intra prediction, then quantizes the signal after transform. The statistical redundancy in the quantized signal is further reduced by entropy coding. However, HEVC has improved AVC in many aspects, for examples, by more flexible block partitioning, better prediction (through more intra prediction modes, advanced motion prediction, etc.), improved design of transform, optimal quantization by RDOQ, better adaptive arithmetic coding, and so on. As closed-circuit television becomes more popular for security application, its massive data storage requirement necessitates low bitrate-low complexity surveillance video system based on HEVC [18].

With the recent blooming of big data processing assisted by massive computation, deep learning (DL) has demonstrated its superior performance in many tasks from classification, recognition, low-level vision to image/video coding. The DL-based-video coding research can be grouped into (i) standard-compliant (requires no modification for standard) and; (ii) standard-non-compliant (requires modification of standard) ones. The first group encodes more effectively the existing coding standard with DL and it can be further divided to (1) fast methods for selecting block partition [2], intra prediction mode [3], etc.; (2) improving compression efficiency with

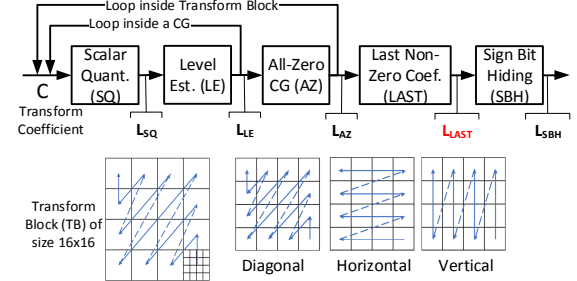


Fig. 1. Rate-Distortion Optimized Quantization in HEVC. Each coefficient group (CG) in the given transform block (TB) is scanned following a scan order (diagonal, horizontal, or vertical). Each coefficient in a CG is scanned also in the same order as CG.

better prediction [4], and pre/post filtering processing [5, 6]. The second group uses DL to perform end-to-end compression in different ways from standards with various approaches such as learned pixel generation network [7], joint up/down sampling [8], etc. In the first and second groups, DL has been applied to such problems already shown their outstanding performance such as super-resolution [9], denoising [10], etc.

In this paper, we apply the deep learning to HEVC encoding problem, more specifically, unlikely those DL-based approaches [2-8] so far, to quantization problem. Rate-Distortion Optimized Quantization (RDOQ) [11] is known to offer about 5% of bitrate reduction compared to the conventional scalar quantization plus dead-zone (SQDZ). It sequentially estimates the optimal quantization levels for each coefficient in a transform block (TB) in the sense of minimum rate and distortion. This is computationally very complex, so for more practical HEVC encoding, a computationally simpler method is implemented in HM [16]. However, even such a computationally simplified RDOQ still takes up about 12% of encoding time in HM [16]. As a result, most research for RDOQ is devoted to reducing its complexity by either even further simplification [12] or even skipping [12, 13] of the rate estimation. However, still less explored in implementing RDOQ is those problems dealing with its sequential nature deterring parallelization.

This is the first attempt to study deep learning in RDOQ. We propose to use deep convolutional neural network to predict RDOQ without actual estimation of rate and distortion, and it is named as Deep Learning-based RDOQ (DL-RDOQ). It is designed to predict optimal quantized TB as a whole, therefore, it will be friendly to parallelization when the deep learning framework itself supports parallelization.

The rest of this paper is organized as follows. Section II first introduces RDOQ and formulates deep learning approach for RDOQ. The proposed DL-RDOQ is addressed in Section III with the simplified residual learning framework. Finally, the paper is concluded in Section V.

¹ This work was supported by Institute for Information & Communications Technology Promotion (IITP) grant funded by the Korea government (MSIT) (No.2018-0-00348, Development of Intelligent Video Surveillance Technology to Solve Problem of Deteriorating Arrest Rate by Improving CCTV Constraint).

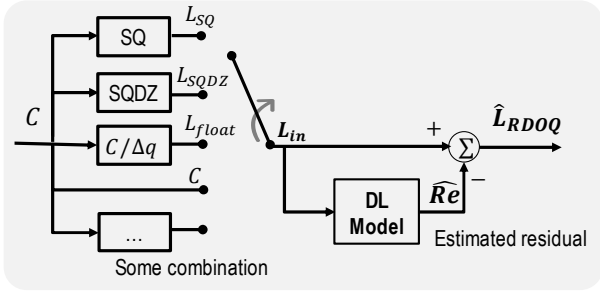


Fig. 2. The proposed residual learning approach for the rate-distortion optimized quantization.

II. RATE-DISTORTION OPTIMIZED QUANTIZATION

A. Optimal Rate-Distortion Optimized Quantization

RDOQ finds the optimal quantized level [11] of each transform coefficient by minimizing the rate (R) and distortion (D) cost in (1) of a given TB.

$$L_{RDOQ} = \arg \min_L J_\lambda(L) \text{ where } J_\lambda(L) = D + \lambda R. \quad (1)$$

RDOQ has practical implementation issues due to (i) huge number of candidates to search for achieving optimality; too excessive amount of computation to evaluate (ii) rate and (iii) distortion for those candidates.

B. Rate-Distortion Optimized Quantization in HEVC

Practical compromise in RDOQ [11] can be made by limiting the number of candidates [16] and/or simplifying the rate calculation via lookup tables [16]. The CABAC process in HEVC splits the transform coefficients of a TB (which can be of size 4×4 , 8×8 , 16×16 , or 32×32), denoted by C , to one or more coefficient groups (CG) of size 4×4 , and processes them in five internal steps shown in Fig. 1.

Firstly, scalar quantization (SQ) is executed as $l_{SQ} = \left\lfloor \frac{c}{\Delta_q} + \theta \right\rfloor$, $\theta = \frac{1}{2}$, where Δ_q denotes a quantization step size and $\lfloor \cdot \rfloor$ represents the floor operator. Following, the level estimation (LE) process selects the best quantized levels of a given CG from a candidate list consisting of l_{SQ} and $l_{SQ} - 1$. Value 0 is further considered as the third candidate if $l_{SQ} = 2$. The LE process is bypassed if $l_{SQ} = 0$. Thirdly, the detection process of All Zero (AZ) coding group decides whether to set all levels in the given CG to all zero based on rate-distortion cost. Fourthly, the last non-zero coefficient (LAST) process detects the best location for the last non-zero level. Finally, Sign Bit Hiding (SBH) process is used to hide a sign bit for the given CG.

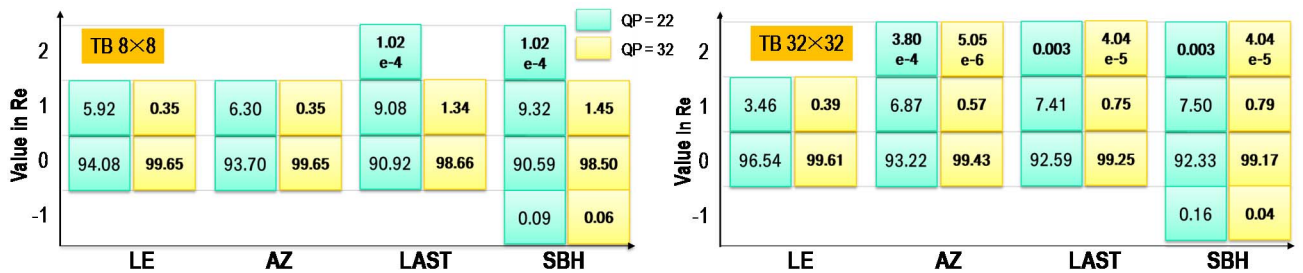


Fig. 3. Distribution (%) of the residual values in Re (vertical axis) at various RDOQ stages in Fig. 1 for Kimono sequence, Intra, TB size of 8×8 and 32×32 . This analysis tells that we can simplify the residual values by removing the case of value of 2 in Re as it rarely occurs.

Table 1. Possible input and output of DL-RDOQ. Multiple input and output can be used to improve the prediction performance.

Possible Input (L_{in})	Possible Output (L_{RDOQ})
Unquant. coeff: C	After level estimation: L_{LE}
Float coeff: C/Δ_q	After all zero coeff. group: L_{AZ}
Scalar quant.: L_{SQ}	After last nonzero coeff.: L_{LAST}
Deadzone quant.: L_{SQDZ}	After sign bit hiding: L_{SBH}

The capital letter denotes a 2D matrix representing a whole TB.

III. DEEP LEARNING-BASED RDOQ (DL-RDOQ)

A. Simplified RDOQ

This section addresses the problem of deep learning-based RDOQ, namely, DL-RDOQ, which predicts the optimal quantized levels of a whole TB without estimating rate and distortion. For practical reason, we investigate DL-RDOQ as a supervised learning problem with possible inputs and outputs given in Table 1. In addition, noting that the residual learning has better proven performance [10], we make DL-RDOQ implemented on HM following the residual learning scheme, that is, predicting the residual of a given TB, $Re \triangleq L_{SQ} - L_{out}$ as shown in Fig. 2.

To find out which RDOQ process is suitable for DL-RDOQ, we analyze the values in Re at various stages. As in Fig. 3, except in SBH, the residual signal assumes only three values of $\{0, 1, 2\}$ with a very small probability of residual value equal to 2, $\Pr\{re = 2\}$ where re is an element in the matrix Re . SBH can have additional -1 value. In addition, SBH is related to signs of coefficients in a CG, so a single miss prediction in SBH will cause a level l to be $-l$, which will significantly increase distortion. We, therefore, use the output after the LAST process L_{LAST} as the output of prediction and set value 2 to 1 for simplification. Now, element in the simplified Re only assumes values of 0 or 1.

B. Proposed Deep Learning-Based RDOQ (DL-RDOQ)

As deep learning has also drawn significant interest in video coding communities, this work studies DL-RDOQ. However, as the first research on DL-RDOQ, it is challenging to find a suitable network for RDOQ. It is because, DL is often applied to signals in the spatial domain like in image enhancement, while RDOQ deals with the DCT transformed signals with have less correlation and different characteristic.

1) Deep Convolution Neural Network

Convolutional neural network (CNN) is a class of deep learning methods which show high performance in many recognition tasks. CNN is well-known for its low complexity, transition invariance and weight sharing characteristic.

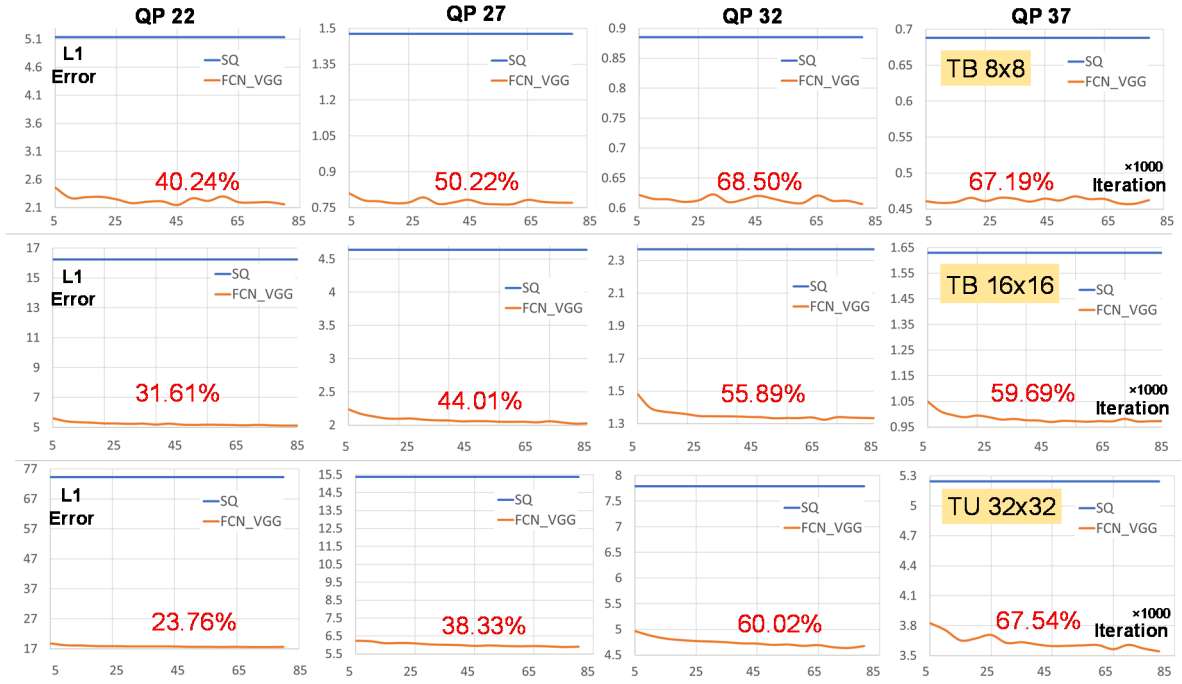


Fig. 4 Average L1 prediction error from RDOQ output in eq. (2) (vertical axis) over iterations (horizontal axis, $\times 1000$) respectively by DL-RDOQ implemented using FCN_VGG [14] and scalar quantization (SQ) for intra coded 1st frame of Kimono, BQTerrace, BasketballDrive with QPs (22~37). The percent number indicates error reduction ratio in (3).

The reason for using CNN to predict RDOQ is that local correlation exists in optimal quantization. Rate estimation is based on the context modeling for each level depending on the previous levels, its frequency location and quantized value. In addition, RDOQ is processed in CG units of size 4×4 . Therefore, CNN filter could learn to predict the optimal level of RDOQ utilizing the local/context characteristic. This work fully follows convolution network FCN_VGG [14] to validate the effectiveness of DL which only uses convolution and Relu activation layers. The DL-RDOQ can be parallelized using GPU thanks to the nature of CNN framework such as Caffe [15]. The training input/output pair is selected as L_{SQ} and $Re = L_{SQ} - L_{LAST}$.

2) Dataset Collection

To enable DL-RDOQ, training dataset is collected from HM 16.15 under coding configuration of Random Access (RA) with QP 22, 27, 32, 37. We collect the values of scalar quantization and RDOQ after each internal stage (LE, AZ, LAST, and SBH) together with information of the current TB (i.e., prediction mode, scan mode, CU size, TB size). The dataset is then grouped according to TB size (4×4 , 8×8 , 16×16 , 32×32). The dataset of each size is further divided into training, testing, and validation with proportion of 90:5:5.

3) Loss Function

Since the residual output is a 2D matrix having only 0 and 1 values as elements, we model the problem similar to semantic segmentation with only two labels, and employ the logistic loss function. In fact, the network structure FCN_VGG remains identical to the semantic segmentation [14]. It should be noted that the proposed network only mimics the RDOQ and does not utilize any rate information.

4) Training

Four networks corresponding to different TB sizes (4×4 , 8×8 , 16×16 , and 32×32) are implemented under Caffe framework [15] and trained with corresponding dataset of

different TB size. We set learning rate to 0.0001, momentum of 0.99, mini-batch size of 512 and total 80,000 iterations.

IV. EXPERIMENTAL RESULTS

A. Prediction Performance

To evaluate the prediction performance of RDOQ output, its average prediction error is computed as:

$$err = \frac{1}{m} \sum_{i=1}^m \|\hat{L}_{LAST} - L_{LAST}\|_1, \hat{L}_{LAST} = L_{SQ} - \widehat{Re} \quad (2)$$

where $\widehat{Re}$ denotes the predicted residual of a given TB, $\hat{L}_{LAST}$ denotes the RDOQ prediction, and m is the total number of TBs for evaluation. Since the residual only contain 0 or 1, L1 error is equivalent to the average number of different levels between DL-RDOQ and RDOQ.

The prediction results along iterations are shown in Fig. 5. The fact that the error is reduced much more by DL-RDOQ than by SQ over iterations clearly shows that the DL could predict the RDOQ output successfully. It is noteworthy to see the different degree of reduction with respect to different TB sizes, which certainly hints on further necessary investigation for designing different network for different TB size as future work.

To further evaluate the difference regarding TB sizes, we normalize the error by TB size, and name it as *error reduction ratio*. It represents the ratio of prediction errors between DL-RDOQ and SQ at the last iteration, and computed as,

$$error\ reduction\ ratio = \frac{err_{DL-RDOQ}}{err_{SQ}} * 100, \quad (3)$$

We observe that DL-RDOQ reduced L1 error by 50%~63% on average compared to SQ and has better error reduction ratio. We observe that the variation in residual characteristic has effects on the prediction results. That is, the smaller probability of $re = 1$, (i.e., large QPs, large TB size), the poorer prediction performance DL-RDOQ has.

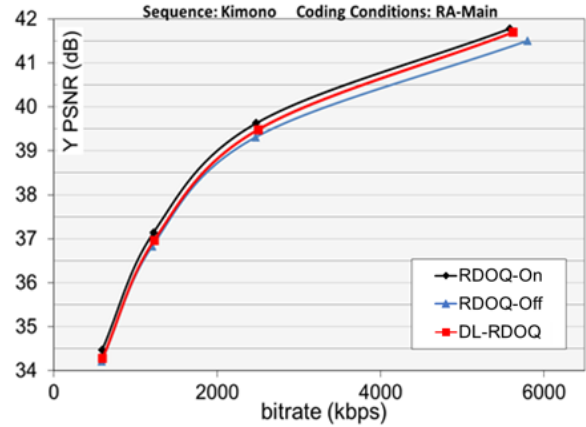
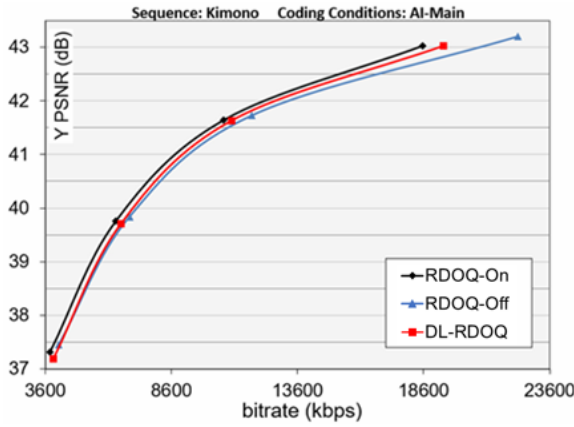


Fig. 6. Coding performance of DL-RDOQ implemented on HM 16.15 for Kimono (1920x1080, 50fps) at All Intra (left) and Random-Access (right).

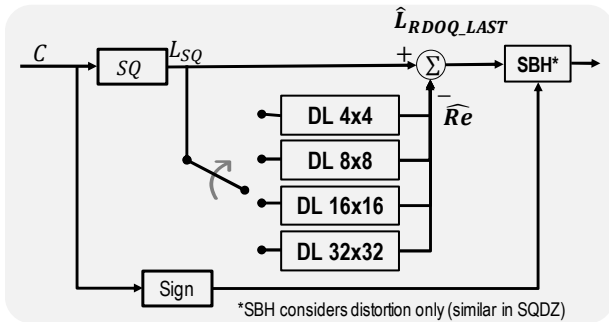


Fig. 5. The proposed residual DL-RDOQ implemented on HM 16.15.

B. Coding Performance

To test DL-RDOQ in HEVC encoding scenario, we implement the trained networks on top of the reference software HM 16.15 with Caffe C++ [15] interface. Four trained FCN_VGG networks are used to generate output when scalar quantized data is given as input. The DL-RDOQ prediction process does not consider the sign information. The corresponding TB size is provided to choose a proper residual DL-RDOQ network corresponding to the given size. The predicted residual signal is subtracted from the input L_{SQ} to obtain the RDOQ prediction after LAST. Sign information is used in SBH to deliver the final RDOQ as shown in Fig. 5.

We compare the coding performance of DL-RDOQ with RDOQ-On and RDOQ-Off with HM 16.15 [16]. SBH is on in both testing cases. For SBH, both rate and distortion are computed in the RDOQ-On case but distortion is computed only in the DL-RDOQ and RDOQ-Off cases in testing. It is because the rate is estimated in RDOQ-On but not in RDOQ-Off nor DL-RDOQ. The rate-distortion curves for all intra (AI) and random access (RA) are shown in Fig. 6 using 10 frames of the sequence Kimono (1920x1080, 50fps) for AI and 64 frames for RA under common test condition [17].

DL-RDOQ performs better than RDOQ-Off (or SQDZ) while fairly approximating the performance of RDOQ-On especially at high bit-rate. It is because, at high bit-rate, there is more residual value of 1 which leads to smaller prediction error of DL-RDOQ. As RDOQ-Off (SQDZ) is better than SQ, DL-RDOQ shows better performance than its initial input, SQ. On the other hand, this work only utilizes deep learning in DL-RDOQ as a black box solution. So, its performance boost can be archived more by further fine tuning such as adding more layers, customizing network structure so that to be able to better utilize knowledge on RDOQ.

V. CONCLUSION

This paper proposed a DL-based method to predict RDOQ without rate-distortion estimation via residual CNN network. The proposed method demonstrated that despite of using DL simply as a black box, we were able to predict RDOQ pretty well and produced much better performance than RDOQ-Off.

REFERENCES

- [1] G. Sullivan, and et al., "Overview of the High Efficiency Video Coding (HEVC) Standard," *IEEE Trans. Circ. Syst. Video Tech.*, vol. 22, no. 12, pp. 1649-1668, 2012.
- [2] M. Xu, and et al., "Reducing complexity of HEVC: A deep learning approach," *IEEE Trans. Image Process.*, vol. 27, no. 10, 2018.
- [3] L. Thorsten and O. Jorn, "Deep learning based intra prediction mode decision for HEVC," *Proc. of IEEE Picture Coding Symposium*, 2016.
- [4] J. Li, and et al., "Fully connected network-based intra prediction for image coding," *IEEE Trans. Image Process.*, vol. 27, no. 7, 2018.
- [5] R. Lin, and et al., "Deep CNN for Decompressed Video Enhancement," *Proc. of IEEE Data Compression Conference*, 2016.
- [6] R. Yang, M. Xu, Z. Wang, and T. Li, "Multi-Frame Quality Enhancement for Compressed Video," *arXiv:1803.04680*, 2018.
- [7] A. Oord, and et al., "Conditional image generation with PixelCNN decoders," *Inter. Conf. Neur. Info. Process. Sys.*, pp. 4797-4805, 2016.
- [8] F. Jiang, and et al., "An end-to-end compression framework based on convolutional neural networks," *IEEE Trans. Circ. Syst. Vid. Tech.*, 2017.
- [9] C. Dong, and et al. "Image super-resolution using deep convolutional networks," *arXiv: 501.00092*, Jul. 2015.
- [10] K. Zhang, and et al., "Beyond a Gaussian Denoiser: Residual learning of deep CNN for image denoising," *IEEE Trans. Image Process.*, vol. 26, no. 7, pp. 3142 – 3155, 2017.
- [11] M. Karczewicz, and et al., "Rate Distortion Optimized Quantization," document ITU-T SG16 Q.6, VCEG-AH21, 2008.
- [12] H. Lee, and et al., "Fast quantization method with simplified rate-distortion optimized quantization for an HEVC encoder," *IEEE Trans. Cir. and Sys. Video Tech.*, vol. 26, no. 1, pp. 106 – 116, 2016.
- [13] M. Xu, and et al., "Simplified rate-distortion optimized quantization for HEVC," *Proc. IEEE Inter. Sym. Broad. Mul. Sys. Broadcast.*, 2018.
- [14] E. Shelhamer et al., "Fully convolutional neural network for semantic segmentation," *IEEE Trans. Patt. Anal. Mach. Intell.*, vol. 39, no. 4, pp.640 – 651, 2016.
- [15] Y. Jia, and et al., "Caffe: Convolutional Architecture for Fast Feature Embedding," *ACM Inter. Conf. Multi Media*, pp. 674-678, 2014.
- [16] High Efficiency Video Coding Test Software 16.15, available at https://hevc.hhi.fraunhofer.de/svn_HEVCSoftware/tags/M-16.15.
- [17] F. Bossen, "Common HM Test Conditions and Software Reference Config," Joint Collaborative Team on Video Coding, JCTVC-L1100.
- [18] X. Zhang, and et al., "Optimizing the Hierarchical Prediction and Coding in HEVC for Surveillance and Conference Videos With Background Modeling," *IEEE Trans. Image Process.*, vol. 23, no. 10, pp. 4511 – 4526, 2014.