

Video Compression With CNN-Based Postprocessing

Fan Zhang , Di Ma , Chen Feng, and David R. Bull, *University of Bristol, BS8 1TL, Bristol, U.K.*

In recent years, video compression techniques have been significantly challenged by the rapidly increased demands associated with high quality and immersive video content. Among various compression tools, postprocessing can be applied on reconstructed video content to mitigate visible compression artefacts and to enhance the overall perceptual quality. Inspired by advances in deep learning, we propose a new convolutional neural network based postprocessing approach, which has been integrated with two state-of-the-art coding standards, versatile video coding (VVC) and AOMedia Video (AV1). The results show consistent coding gains on all tested sequences at various spatial resolutions, with average bit rate savings of 4.0% and 5.8% against original VVC and AV1, respectively (based on the assessment of peak signal-to-noise ratio). This network has also been trained with perceptually inspired loss functions, which have further improved the reconstruction quality based on perceptual quality assessment (VMAF), with average coding gains of 13.9% over VVC and 10.5% against AV1.

The importance of video compression has come to the fore over the past decade driven by the tension between the huge quantities of video content consumed every day and the bandwidth available to transmit it. This challenge has been addressed through the creation of new video coding standards, the latest activity being by the joint video exploration team (JVET), who published the first version of H.266/versatile video coding (VVC)¹ in 2020. Compared to its predecessor, H.265/high-efficiency video coding (HEVC), VVC has achieved up to 50% performance improvement through the adoption of numerous sophisticated coding tools, in particular with improved support for formats with high spatial resolutions, high dynamic range and spherical content. Alongside VVC, the Alliance for Open Media (AOMedia) also published its first video coding format, AOMedia Video (AV1) in 2018, which has been reported to offer comparable coding performance to VVC.²

Machine learning, especially based on deep convolutional neural networks (CNNs), has been increasingly applied in the context of video compression and has achieved promising results both when used in

conjunction with conventional coding algorithms and in the form of new end-to-end architectures. In addition to conventional normative coding tools, deep learning can also be employed at the video decoder as a postprocessing stage to further reduce noticeable artefacts and enhance the visual quality of compressed content. For the state-of-the-art coding standards, such as VVC and AV1, most existing learning-based postprocessing approaches can only offer evident improvement for less-efficient coding configurations (e.g., intracoding), and the employed networks are normally trained to minimize average absolute pixel distortions rather than explicitly to improve the perceptual quality.

Based on the Generative Adversarial Network (GAN) paradigm, we propose a novel CNN-based postprocessing approach with an extension that achieves an improved perceptual reconstruction quality. This approach has been evaluated on the VVC Test Model (VTM) 4.0.1 and on AV1 libaom 1.0.0, with results showing consistent improvement on standard JVET test sequences for different QP values based on different quality measurements. We have further analyzed the computational complexities of different CNN structure variants and correlated them with overall coding gains.

This article is a comprehensive extension of our previous work,³ which solely focused on the peak

1070-986X © 2021 IEEE

Digital Object Identifier 10.1109/MMUL.2021.3052437

Date of publication 18 January 2021; date of current version 6 December 2021.

signal-to-noise ratio (PSNR) driven optimization of VVC compressed content. The primary differences are summarized below.

- The CNN model used for postprocessing has been extensively upgraded with a new GAN-based perceptual training strategy, which can significantly improve the perceptual quality of the final reconstructed content compared to the paper by Zhang *et al.*³
- This CNN-based postprocessing approach has been also trained and evaluated on AV1 compressed results (alongside VVC in the paper by Zhang *et al.*³) and achieved similar coding gains.

In the remainder of the article, we first survey the prior work in the field of deep video compression and, in particular, describe learning-based postprocessing approaches. Second, we present our proposed CNN-based postprocessing approach, describing the network architecture and how it was trained and evaluated. We then summarize and discuss the experimental results, highlighting the performance improvements over standard video codecs. Finally, we conclude the article and outline possible future work.

PRIOR WORK

Deep Video Compression

In the past few years, machine learning, in particular deep neural networks, has been increasingly applied to image and video compression, demonstrating significant potential compared to conventional coding methods. Learning-based video coding algorithms can be classified into two primary groups: new end-to-end network architectures and those that enhance individual conventional coding tools.

Machine learning has been employed to refine existing coding tools within a conventional coding framework, including intraprediction and interprediction, transformation, quantization, and in-loop filtering.⁴ New coding tools have also been developed with the support of neural networks, such as CNN-based spatial resolution and bit depth adaptation.⁵ Moreover, the classic hybrid video coding framework has been challenged by new deep network architectures, which enable end-to-end training and optimization.^{6,7} This latter approach often employs a general rate distortion framework with nonlinear transforms, which are based on convolutional filters and nonlinear activation functions. Although these solutions show great promise as an alternative to conventional codecs, their performance cannot still compete with

the latest standardized video codecs, including VVC and AV1.

CNN-Based Postprocessing

Compression processes often introduce various visible artefacts such as blocking mismatches, banding, and blurring, especially when large quantization steps are employed. These unpleasant distortions can be mitigated by filtering the reconstructed frames. When this enhancement process is performed outside of the encoding loop (generally after decoding), it is referred to as postprocessing.

In standardized codecs, filters have also been designed for use within the encoding loop to reduce compression artefacts. VVC employs three different types in-loop filters including deblocking filters (DBF), sample adaptive offset (SAO), and adaptive loop filters (ALF).¹ In AV1, in addition to DBFs, there are two other filtering operations: constrained directional enhancement filtering and loop restoration filtering.²

CNNs are now also playing an important role in image restoration (including super-resolution), and these approaches can also be employed for postprocessing of compressed video content to improve the overall reconstruction quality. Although various researchers have implemented CNN-based postprocessing approaches in the context of HEVC and VVC,^{8–11} most of these can only achieve coding gains for All Intraconfigurations, which offer lower coding efficiency compared to Random Access configurations based on hierarchical B frame structures. In addition, all the employed CNN models in these approaches were trained to optimize a simple loss function based on pixel distortions (ℓ_1 or ℓ_2 loss), which can lead to oversmoothed reconstruction results.

PROPOSED ALGORITHM

Figure 1 shows a high-level coding workflow with a CNN-based postprocessing module. This section focuses on the structure of the employed CNN architecture, and the details of network training and evaluation.

CNN Architecture

The CNN architecture used in this work is illustrated in Figure 2.

The *generator network* is a modified version based on the generator (SRResNet) of SRGAN,¹² and this has been previously employed by the authors in video compression systems based on spatial resolution and bit depth resampling.⁵ This network takes a 96×96 RGB (compressed) image block as input and produces

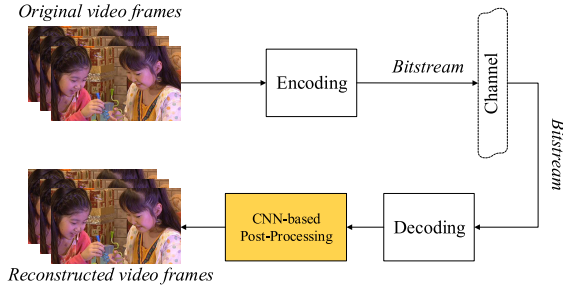


FIGURE 1. Typical coding workflow with a CNN-based post-processing module.

an image block with the same format, targeting to its corresponding original (uncompressed) counterpart. [We have used the same input/output formats as in SRResNet.]

Residual blocks (RB) form the basic unit in this network, which contains two convolutional layers and a parametric rectified linear unit (PReLU) activation function in between them. A skip connection is used between the input of each RB and the output of its second convolutional layer. The number of RB is configurable and was set to 16 in this work.

The input of the network is connected to these successive RBs through a convolutional layer (also with a ReLU). Between the network output and the output of the last RB, there is also a convolution layer (output layer) followed by a Tanh activation function. An additional skip connection is used between the

output of the input layer and the output of the last RB. A long skip connection is also employed between the input the first RB and the output of the output layer to produce the final output.

The discriminator used in this work is similar to that in SRGAN,¹² which takes the output of the generator (fake) and compares to its corresponding original (real). This network consists of one input layer (with a Leaky ReLU), seven identical convolutional layers, and two dense layers. Each of the convolutional layers is followed by a batch normalization layer and a leaky ReLU activation function. After the second dense layer, a Sigmoid activation function is employed to output a probability to predict how much the quality of the *real* image block is perceptually better than the *fake* one.

The parameters used in each convolutional layer including kernel sizes, feature map numbers, and stride values are shown in Figure 2.

Training Database

The training material is essential for learning-based algorithms. We need to ensure that the training content is diverse and covers various texture types in order to achieve good model generalization and avoid potential overfitting problems. To train the employed network, we have selected 432 uncompressed video sequences from a publicly available training database BVI-DVC,¹³ which was designed specifically for deep video compression. All these sequences have the same frame rate of 60 frames per second, YCbCr 4:2:0

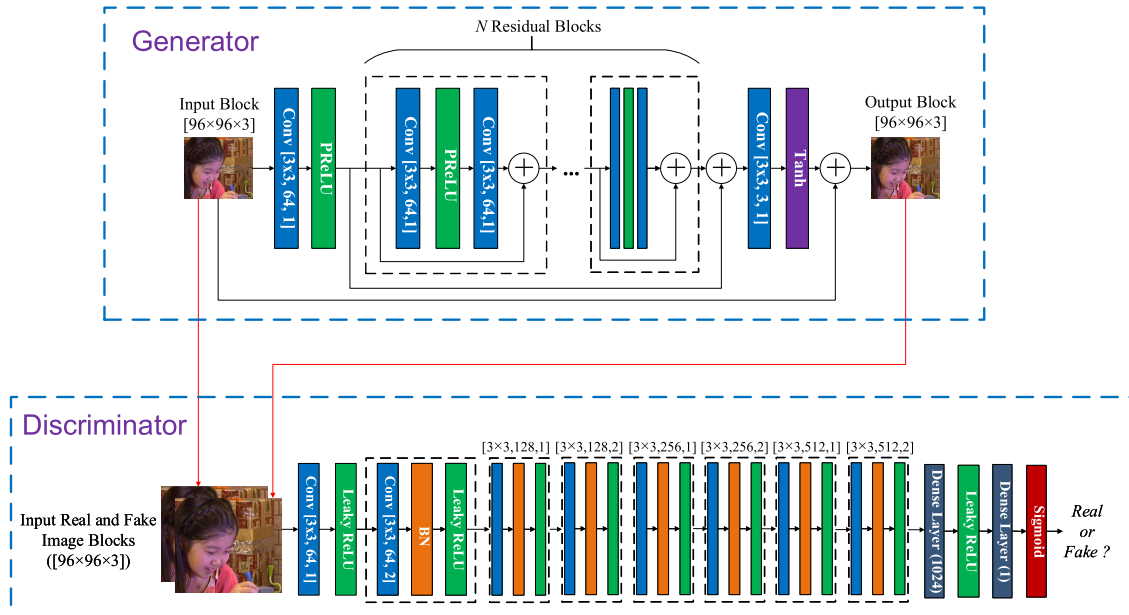


FIGURE 2. Employed GAN architecture, comprising generator, and discriminator stages.

format, and with four different spatial resolutions including 3 840×2 160, 1 920×1 080, 960×540, and 480×270. We have encoded these 432 original sequences using VVC VTM 4.0.1 and AV1 libaom 1.0.0 with the coding configurations summarized in Table 1.

For each codec, the reconstructed video frames for each QP value and their corresponding originals were randomly selected and segmented into 96×96 color image blocks (after converting to the RGB space from YCbCr 4:2:0). We have also rotated selected image blocks to further improve data diversity. As a result, for each video codec and QP group [Based on our preliminary results, we note that, through QP sub-grouping, additional coding gains (approximately 0.05dB based on PSNR) can be achieved compared to using a single CNN model.], there are over 1 00 000 image blocks pairs (compressed and original).

Training Strategy

We trained this network using two different methodologies: (i) only train the generator using ℓ_1 loss (mean absolute difference); (ii) jointly train both the generator and the discriminator based on perceptually inspired loss functions. The CNN models obtained by these two training methods are used to postprocess VVC and AV1 compressed content in the evaluation experiments, and their results are compared in the following section.

ℓ_1 loss is first employed to train the network (generator only) using the material generated for each QP group and codec. This results in a total number of nine CNN models for different evaluation scenarios:

$$\begin{cases} \text{CNN}_{\text{VVC}, \text{QP}22}, & \text{if } \text{QP}_{\text{eval}} \leq 24.5 \\ \text{CNN}_{\text{VVC}, \text{QP}27}, & \text{if } 24.5 < \text{QP}_{\text{eval}} \leq 29.5 \\ \text{CNN}_{\text{VVC}, \text{QP}32}, & \text{if } 29.5 < \text{QP}_{\text{eval}} \leq 34.5 \\ \text{CNN}_{\text{VVC}, \text{QP}37}, & \text{if } 34.5 < \text{QP}_{\text{eval}} \leq 39.5 \\ \text{CNN}_{\text{VVC}, \text{QP}42}, & \text{if } \text{QP}_{\text{eval}} > 39.5 \end{cases} \quad (1)$$

$$\begin{cases} \text{CNN}_{\text{AV1}, \text{QP}32}, & \text{if } \text{QP}_{\text{eval}} \leq 37.5 \\ \text{CNN}_{\text{AV1}, \text{QP}43}, & \text{if } 37.5 < \text{QP}_{\text{eval}} \leq 49 \\ \text{CNN}_{\text{AV1}, \text{QP}55}, & \text{if } 49 < \text{QP}_{\text{eval}} \leq 59 \\ \text{CNN}_{\text{AV1}, \text{QP}63}, & \text{if } \text{QP}_{\text{eval}} > 59 \end{cases} \quad (2)$$

Here, QP_{eval} represents the base QP value employed in the evaluation phase for the two different codecs, and $\text{CNN}_{\text{c}, \text{q}}$ is the CNN model trained for different codecs (VVC or AV1) and QP values.

Perceptual loss functions have been employed to train the whole GAN architecture following a two-stage training strategy. This was initially designed to

train the relativistic GAN for image generation,¹⁵ and has also been used to train the CNN models for spatial resolution and bit depth up-sampling.^{16,17} In the first stage, the generator is trained separately using the multiscale structural similarity index (MS-SSIM)¹⁸ as the loss function. The trained generator model is employed as the starting point when the generator and discriminator are trained together in the second phase.

The generator is trained using a combined loss function, $\mathcal{L}_{\text{gen}}$ in the second stage:

$$\mathcal{L}_{\text{gen}} = \mathcal{L}_{\text{SSIM}} + \alpha \cdot \mathcal{L}_{\ell_1} + \beta \cdot \mathcal{L}_{\text{G}}^a \quad (3)$$

in which $\mathcal{L}_{\text{SSIM}}$ stands for the SSIM¹⁹ loss (1-SSIM) between the generator output and the target, while $\mathcal{L}_{\ell_1}$ is the ℓ_1 loss between them. $\mathcal{L}_{\text{G}}^a$ is defined as the adversarial loss for the generator:

$$\begin{aligned} \mathcal{L}_{\text{G}}^a = & -E_r[\ln(1 - (\text{Sig}(O_d(I_r) - E_f[O_d(I_f)])))] \\ & - E_f[\ln(\text{Sig}(O_d(I_f) - E_r[O_d(I_r)]))]. \end{aligned} \quad (4)$$

Here, I_r and I_f are denoted as the *real* and *fake* image blocks, respectively. $E_r[\cdot]$ represents the mean operation for all the *real* (fake if I_r is replaced by I_f) image blocks, and $O_d(\cdot)$ is the output of the discriminator. “Sig” represents the Sigmoid function.

For the discriminator, the loss function $\mathcal{L}_D$ is given by the following:

$$\begin{aligned} \mathcal{L}_D = & -E_r[\ln(\text{Sig}(O_d(I_r) - E_f[O_d(I_f)]))] \\ & - E_f[\ln(1 - (\text{Sig}(O_d(I_f) - E_r[O_d(I_r)])))]]. \end{aligned} \quad (5)$$

The training configuration is summarized as follows. We implemented all networks based on the TensorFlow 1.8.0 framework, and set the learning rate and weight decay to 0.0001 and 0.1 (for every 100 epochs), respectively, for both training stages. The total number of training epochs are 200. We have used Adam optimization algorithm during the training with the hyperparameters of $\beta_1 = 0.9$ and $\beta_2 = 0.999$. The two weights α and β in (3) are set up to 0.025 and $5e-3$, respectively.

Evaluation Operation

In the evaluation stage, when we use the trained CNN models to enhance compressed video frames, each frame is segmented into 96×96 overlapping image blocks with an overlap of 4 pixels. These blocks are converted to RGB color space as the input of CNN models, and the output image blocks are then aggregated in the same way through simple blending to

TABLE 1. Coding configuration employed for VVC VTM and AV1 libaom.

Codec	Version	Configuration parameters
VVC VTM	4.0.1	Random Access configuration ¹⁴ . IntraPeriod=64, GOPSize=16, QP=22, 27, 32, 37, 42.
AV1 libaom	1.0.0-5ec3e8c (02/05/2020)	-i420 -psnr -usage=0 -verbose -cpu-used=0 -threads=0 -profile=0 -width=\$w -height=\$h -input-bit-depth=10 -bit-depth=10 -fps=\$fps/1001 -passes=1 -kf- max-dist=64 -kf-min-dist=64 -drop-frame=0 -static-thresh=0 -arnr- maxframes=7 -arnr-strength=5 -lag-in-frames=19 -aq-mode=0 -bias-pct=100 -minsection-pct=1 -maxsection-pct=10000 -auto-alt-ref=1 -min-q=0 -max-q=63 -max-gf-interval=16 -min-gf-interval=4 -frame-parallel=0 -color-primaries=bt709 -end-usage=q -sharpness=0 -undershoot-pct=100 -overshoot-pct=100 -tile- columns=0 -cq-level={32, 43, 55, 63} w/o -enable-fwd-kf=1

form the final video frame. Here, only the generator network is used in the evaluation (the discriminator is for training only).

RESULTS AND DISCUSSION

The proposed postprocessing approach has been utilized to enhance both VVC and AV1 compressed

TABLE 2. Compression performance of the proposed method benchmarked on the original VVC VTM 4.0.1. Negative BD-rate values indicate coding gains.

Method	VTM-PP (£1 loss)				VTM-PP (Perceptual loss)			
Metric	PSNR		VMAF		PSNR		VMAF	
QP Range	H-QPs	L-QPs	H-QPs	L-QPs	H-QPs	L-QPs	H-QPs	L-QPs
Class-Sequence	BD-rate	BD-rate	BD-rate	BD-rate	BD-rate	BD-rate	BD-rate	BD-rate
A1-Campfire	-3.3%	-2.3%	-5.6%	-4.6%	+0.2%	-0.4%	-10.4%	-10.7%
A1-FoodMarket4	-2.6%	-2.0%	-3.8%	-3.0%	-0.0%	+0.1%	-8.4%	-7.4%
A1-Tango2	-3.3%	-2.9%	-3.4%	-3.0%	-1.1%	-0.6%	-7.8%	-9.3%
A2-CatRobot1	-5.2%	-5.2%	-4.6%	-4.4%	-0.6%	-1.1%	-14.4%	-17.8%
A2-DaylightRoad2	-6.0%	-7.1%	-6.8%	-7.2%	-1.1%	-2.1%	-19.5%	-23.4%
A2-ParkRunning3	-0.8%	-0.4%	-2.3%	-0.2%	+2.1%	+2.4%	-11.7%	-11.1%
Class A	-3.5%	-3.3%	-4.4%	-3.7%	-0.1%	+0.3%	-10.1%	-13.3%
B-BQTerrace	-2.2%	-1.0%	-6.1%	-1.1%	+2.0%	+0.3%	-23.3%	-28.8%
B-BasketballDrive	-3.4%	-3.1%	-1.8%	+2.7%	-0.9%	-0.9%	-7.9%	-8.7%
B-Cactus	-3.4%	-3.0%	-5.1%	-4.4%	+0.2%	-0.2%	-15.8%	-17.1%
B-MarketPlace	-2.6%	-2.3%	-4.8%	-4.0%	+1.2%	+0.3%	-17.7%	-18.2%
B-RitualDance	-3.8%	-3.5%	-4.6%	-2.6%	-1.1%	-1.2%	-11.7%	-11.2%
Class B	-3.1%	-2.6%	-4.5%	-1.9%	+0.3%	-0.3%	-15.3%	-16.8%
C-BQMall	-5.6%	-5.6%	-4.7%	-6.8%	-2.1%	-2.5%	-14.3%	-13.6%
C-BasketballDrill	-3.9%	-3.6%	-3.8%	-2.8%	-1.4%	-1.6%	-12.3%	-11.7%
C-PartyScene	-4.1%	-4.3%	-5.9%	-4.1%	-0.4%	-1.4%	-16.1%	-13.8%
C-RaceHorses	-3.1%	-2.1%	-3.4%	+1.2%	+0.2%	-0.4%	-11.9%	-10.4%
Class C	-4.2%	-3.9%	-4.5%	-3.1%	-0.9%	-1.5%	-13.7%	-12.4%
D-BQSquare	-8.7%	-9.6%	-10.1%	-11.6%	-4.0%	-4.5%	-16.6%	-19.5%
D-BasketballPass	-6.1%	-5.6%	-5.4%	-4.0%	-3.0%	-2.8%	-9.8%	-8.1%
D-BlowingBubbles	-3.7%	-3.8%	-4.8%	-3.8%	-0.5%	-1.3%	-16.1%	-14.5%
D-RaceHorses	-4.8%	-4.2%	-5.2%	-1.0%	-1.6%	-1.9%	-11.8%	-10.5%
Class D	-5.8%	-5.8%	-6.4%	-5.1%	-2.3%	-2.6%	-13.6%	-13.2%
Overall	-4.0%	-3.8%	-4.9%	-3.4%	-0.6%	-1.2%	-13.5%	-14.2%
	BD-rate=-3.9%		BD-rate=-4.2%		BD-rate=-0.9%		BD-rate=-13.9%	

TABLE 3. Compression performance of the proposed method benchmarked on the original AV1 1.0.0. Negative BD-rate values indicate coding gains.

Method	AV1-PP (ℓ_1 loss)		AV1-PP (Perceptual loss)	
Metric	PSNR	VMAF	PSNR	VMAF
Class-Sequence	BD-rate	BD-rate	BD-rate	BD-rate
A1-Campfire	-5.0%	-7.3%	-1.1%	-11.8%
A1-FoodMarket4	-4.0%	-8.9%	-1.1%	-9.6%
A1-Tango2	-4.6%	-8.2%	-1.9%	-10.7%
A2-CatRobot1	-5.9%	-5.9%	-1.9%	-15.0%
A2-DaylightRoad2	-7.7%	-5.0%	-3.1%	-17.2%
A2-ParkRunning3	-1.9%	-2.2%	+0.4%	-11.3%
Class A	-4.9%	-6.3%	-1.5%	-12.6%
B-BQTerrace	-4.5%	+1.4%	-1.3%	-10.2%
B-BasketballDrive	-6.0%	-3.0%	-2.9%	-7.1%
B-Cactus	-3.8%	-3.3%	-0.7%	-11.9%
B-MarketPlace	-3.0%	-4.5%	-0.5%	-17.4%
B-RitualDance	-4.7%	-5.0%	-2.3%	-11.1%
Class B	-4.4%	-2.9%	-1.5%	-11.5%
C-BQMall	-6.3%	-2.4%	-2.9%	-9.4%
C-BasketballDrill	-6.8%	-2.4%	-3.1%	-9.9%
C-PartyScene	-7.8%	+2.3%	-3.3%	-9.1%
C-RaceHorses	-4.1%	-3.8%	-1.8%	-8.4%
Class C	-6.3%	-1.6%	-2.8%	-9.2%
D-BQSquare	-16.1%	+11.2%	-8.4%	-4.5%
D-BasketballPass	-7.0%	-3.8%	-3.9%	-8.2%
D-BlowingBubbles	-6.2%	+2.0%	-2.8%	-9.7%
D-RaceHorses	-5.6%	-3.0%	-3.0%	-7.7%
Class D	-8.7%	+1.6%	-4.5%	-7.5%
Overall	-5.8%	-2.7%	-2.4%	-10.5%

content. We have used all 19 test sequences from the JVET CTC standard dynamic range (SDR) testset. None of these sequences were included in the CNN training database.

VVC compressed content was generated using VTM 4.0.1 with the random access configuration. AV1 libaom 1.0.0 was used to produce AV1 content, with similar coding parameters to those for VVC. The employed configurations for both codecs are summarized in Table 1. The tested QP values are 22, 27, 32, 37, and 42 for VVC, and 32, 43, 55, and 63 for AV1.

The final video quality was evaluated using two assessment methods, PSNR and Video Multimethod Assessment Fusion (VMAF).²⁰ PSNR is the most commonly used quality metric in the video compression community, while VMAF is a machine learning based metric, which combines multiple existing quality metrics and a video feature using a support vector machine regression approach. Compared to PSNR,

VMAF has been reported to provide more accurate prediction of the perceptual quality. The performance of the proposed approach has been compared with two original codecs using the Bjntegaard delta rate measurements (BD-rate).²¹ For VVC, the compression performance is evaluated for both low (22–37) and high QP (27–42) ranges.

In order to further benchmark the performance of the proposed algorithm, another two state-of-the-art CNN-based postprocessing approaches, denoted as JVET-N0254¹⁰ and JVET-O0079,¹¹ are also compared here in the context of VVC (for low QP range only). Results of both are based on the RA configuration and were submitted to MPEG JVET meetings as VVC proposals.

Compression Performance

Tables 2 and 3 summarize the compression performance of the proposed method when it is applied to

TABLE 4. Comparison between the proposed method and two existing CNN-based PP approaches for VVC (low QP range).

Sequence (Class)	O0079 [11]	N0254 [10]	Proposed Method (ℓ_1)
	BD-rate (PSNR)	BD-rate (PSNR)	BD-rate (PSNR)
Class A (2160p)	-1.3%	-1.7%	-3.5%
Class B (1080p)	-1.5%	-1.1%	-3.1%
Class C (480p)	-3.3%	-1.4%	-4.2%
Class D (240p)	-5.0%	-1.4%	-5.8%
Overall	-2.6%	-1.4%	-4.0%

VVC and AV1 compressed content. For ℓ_1 trained CNNs, we note that the average bit-rate savings according to PSNR are 3.9% and 5.8% against the original VVC and AV1, respectively. If the perceptual quality metric VMAF is used to assess the video quality, the coding gains are 4.2% over VVC and 2.7% over AV1. When we use perceptual loss function trained models

for postprocessing, the coding gains appear much more significant based on the assessment of VMAF—13.9% and 10.5% over VVC and AV1, respectively. This is particularly evident for test sequences, such as ParkRunning3, BQTerrace, and BQSquare, which present a large number of sharp edges. We have also compared the proposed method with two JVET proposals, JVET-N0254¹⁰ and JVET-O0079,¹¹ in Table 4, where the ℓ_1 trained CNNs provides superior enhancement performance to these two works for all resolution classes according to PSNR.

Subjective Comparison

Figures 3 and 4 provide subjective comparisons between the reconstructed frames generated by the original VVC/AV1, ℓ_1 trained CNNs and perceptual loss function trained models. We can observe that the reconstructed blocks of the proposed method (for both ℓ_1 and perceptually trained models) exhibit fewer noticeable blocking artefacts compared to the anchor codecs. In addition, the CNN models trained using perceptual loss functions produce results with slightly

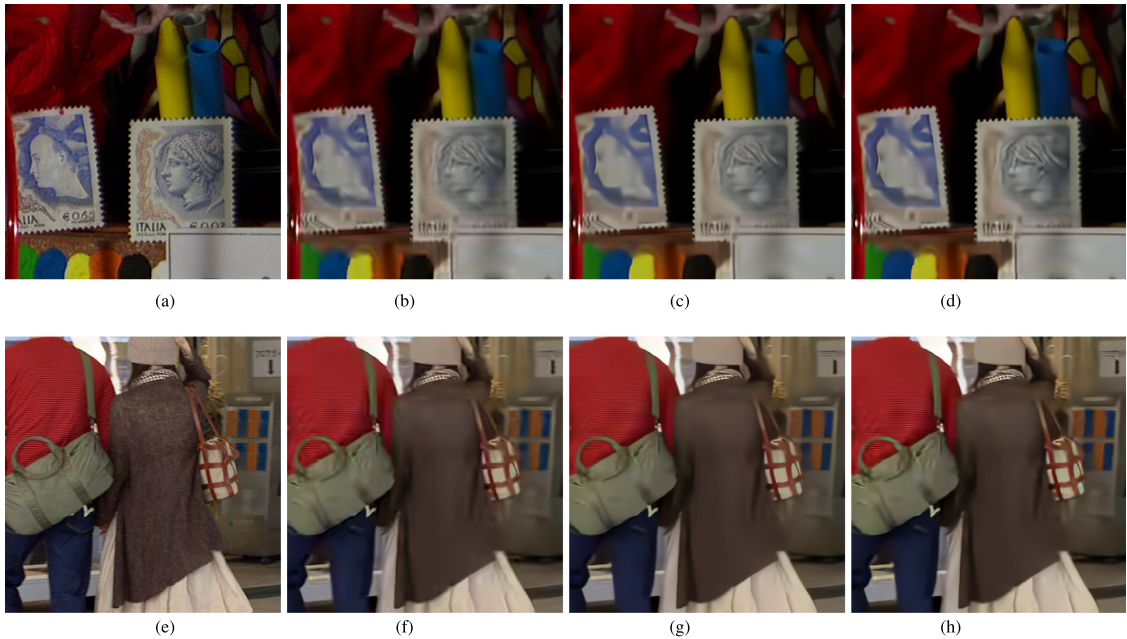


FIGURE 3. Example blocks of the reconstructed frames for the anchor VTM 4.0.1 and the proposed approach. These are from the 91st and 320th frames of *Cactus* and *BQMall* sequences, respectively. We can note that the results produced by the perceptual loss function trained models exhibit more textural detail and higher contrast than those generated by ℓ_1 trained networks (the difference is more evident within the regions with the stamps in *Cactus*, and the areas with the bags and the white skirt in *BQMall*). (a) Original. (b) VTM 4.0.1, QP = 42 (VMAF = 62.5). (c) VTM-PP: ℓ_1 , QP = 42 (VMAF = 64.0). (d) VTM-PP: GAN, QP = 42 (VMAF = 66.2). (e) Original. (f) VTM 4.0.1, QP = 42 (VMAF = 70.4). (g) VTM-PP: ℓ_1 , QP = 42 (VMAF = 72.2). (h) VTM-PP: GAN, QP = 42 (VMAF = 74.3).



FIGURE 4. Example blocks of the reconstructed frames for the anchor AV1 and the proposed approach. These are from the 91st and 320th frames of *Cactus* and *BQMall* sequences, respectively. The results produced by the perceptual loss function trained models exhibit more textural detail and higher contrast than those generated by ℓ_1 trained networks (the difference is more evident within the regions with the stamps in *Cactus*, and the areas with the bags and the white skirt in *BQMall*). (a) Original. (b) AV1, QP = 63 (VMAF = 75.2). (c) AV1-PP: ℓ_1 , QP = 63 (VMAF = 76.4). (d) AV1-PP: GAN, QP = 63 (VMAF = 78.1). (e) Original. (f) AV1, QP = 63 (VMAF = 83.7). (g) AV1-PP: ℓ_1 , QP = 63 (VMAF = 84.9). (h) AV1-PP: GAN, QP = 63 (VMAF = 86.2).

more textural detail and higher contrast than those generated by ℓ_1 trained networks.

Complexity Analysis

The relative computational complexity of the proposed method, benchmarked on the original VVC and AV1 codecs, is presented in Table 5. We have used a shared cluster computer at the University of Bristol, to execute all the computations. This computer contains multiple nodes with 2.4 GHz Inter CPUs, 138 GB RAM, and NVIDIA P100 GPU devices. We note that the average decoding complexity is 56.3 and 70.0 times compared to the original VVC VTM 4.0.1 and AV1,

respectively, due to the employment of CNN-based postprocessing at the decoder. This is for a configuration with 16 RB in the generator. This figure is slightly higher than those of two JVET proposals O0079¹¹ and N0254,¹⁰ in which the decoder complexities (inc. the CNN-based PP module) are 45.6 and 35.2 times that of the original VVC decoder (based on the whole JVET dataset).

Moreover, we have further investigated the relationship between the number of RB and compression performance. Figure 5 shows the coding gains (in terms of PSNR) and algorithm relative complexity using CNN models with different numbers of RB (N =

TABLE 5. Relative complexity of the proposed method benchmarked on original VVC and AV1 decoders.

Hose Codec	VVC	AV1
Class A (2160p)	18.6×	23.2×
Class B (1080p)	36.8×	38.3×
Class C (480p)	76.5×	83.3×
Class D (240p)	116.9×	166.7×
Average	56.3×	70.0×

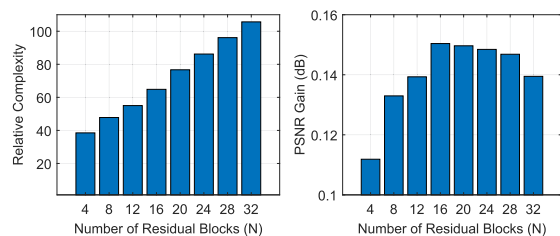


FIGURE 5. (Left) Relative complexity for different number of RB. (Right) PSNR gains for different number of RB.

4, 8, 12, 16, 20, 24, 28, and 32) to process VVC VTM QP 42 compressed content. Again the relative complexity here is benchmarked on the original VVC decoder. We observe that, when the number of RB (N) increases from 4 to 16, the PSNR gain relative to the original VVC content increases in a linear fashion. However, when the number of RB exceeds 20, the overall coding gain starts to decrease. This provides evidence that the residual block number in the proposed work is an optimal selection.

CONCLUSION

In this article, we presented a CNN-based postprocessing approach, which achieves evident and consistent coding gains over standardized video codecs, VVC and AV1. The employed CNN model was trained using both ℓ_1 and perceptually inspired methodologies. We would like to recommend future work focusing on computational complexity reduction and further improvement on the training methodology.

ACKNOWLEDGMENTS

This work was supported in part by the UK EPSRC under Grant EP/L016656/1 and Grant EP/M000885/1 and in part by the NVIDIA GPU Seeding Grants. We also appreciate the valuable advice provided by Dr. D. Mukherjee on AV1 coding configuration.

REFERENCES

1. Versatile video coding, ITU-T Std. ITU-T Rec. H. 266, 2020.
2. Y. Chen *et al.*, "An overview of coding tools in AV1: The first video codec from the alliance for open media," *APSIPA Trans. Signal Inf. Process.*, vol. 9, 2020.
3. F. Zhang, C. Feng, and D. R. Bull, "Enhancing VVC through CNN-based post-processing," in *Proc. IEEE Int. Conf. Multimedia Expo.*, 2020, pp. 1–6.
4. D. Liu, Y. Li, J. Lin, H. Li, and F. Wu, "Deep learning-based video coding: A review and a case study," *ACM Comput. Surveys*, vol. 53, no. 1, pp. 1–35, 2020.
5. F. Zhang, M. Afonso, and D. R. Bull, "ViSTRA2: Video coding using spatial resolution and effective bit depth adaptation," 2019, *arXiv:1911.02833*.
6. J. Ballé, V. Laparra, and E. P. Simoncelli, "End-to-end optimized image compression," in *Proc. Int. Conf. Learn. Representations*, 2017.
7. O. Rippel, S. Nair, C. Lew, S. Branson, A. G. Anderson, and L. Bourdev, "Learned video compression," in *Proc. IEEE Int. Conf. Comput. Vis.*, 2019, pp. 3454–3463.
8. L. Ma, Y. Tian, P. Xing, and T. Huang, "Residual-based post-processing for HEVC," *IEEE MultiMedia*, vol. 26, no. 4, pp. 67–79, Oct.–Dec. 2019.
9. H. Zhao, M. He, G. Teng, X. Shang, G. Wang, and Y. Feng, "A CNN-based post-processing algorithm for video coding efficiency improvement," *IEEE Access*, vol. 8, pp. 920–929, 2019.
10. Y. Wang, Z. Chen, Y. Li, L. Zhao, S. Liu, and X. Li, "Ce13: Dense residual convolutional neural network based in-loop filter (ce13-2.2 and ce13-2.3)," *JVET Meeting*, Geneva, Switzerland, Tech. Rep. *JVET-N0254*. ITU-T, ISO/IEC, 2019.
11. S. Wan *et al.*, "CE10: Integrated in-loop filter based on CNN (Tests 2.1, 2.2 and 2.3)," *JVET Meeting*, Geneva, Switzerland, Tech. no. *JVET-O0079*. ITU-T, ISO/IEC, 2019.
12. C. Ledig *et al.*, "Photo-realistic single image super-resolution using a generative adversarial network," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit.*, 2017, pp. 105–114.
13. D. Ma, F. Zhang, and D. R. Bull, "BVI-DVC: A training database for deep video compression," 2020, *arXiv:2003.13552*.
14. F. Bossen, J. Boyce, X. Li, V. Seregin, and K. Sühring, "JVET common test conditions and software reference configurations for SDR video," *JVET meeting*, Geneva, Switzerland, Tech. Rep. *JVET-M1001*. ITU-T and ISO/IEC, 2019.
15. A. Jolicoeur-Martineau, "The relativistic discriminator: A key element missing from standard GAN," 2018, *arXiv:1807.00734*.
16. D. Ma, F. Zhang, and D. R. Bull, "GAN-based effective bit depth adaptation for perceptual video compression," in *Proc. IEEE Int. Conf. Multimedia Expo.*, 2020, pp. 1–6.
17. D. Ma, M. F. Afonso, F. Zhang, and D. R. Bull, "Perceptually-inspired super-resolution of compressed videos," *Appl. Digit. Image Process. XLII*, vol. 11137, 2019, Art. no. 1113717.
18. Z. Wang, E. P. Simoncelli, and A. C. Bovik, "Multi-scale structural similarity for image quality assessment," in *Proc. Asilomar Conf. Signals, Syst. Comput.*, 2003, vol. 2, Art. no. 1398.
19. Z. Wang, A. Bovik, H. Sheikh, and E. Simoncelli, "Image quality assessment: From error visibility to structural similarity," *IEEE Trans. Image Process.*, vol. 13, no. 4, pp. 600–612, Apr. 2004.
20. Z. Li, A. Aaron, I. Katsavounidis, A. Moorthy, and M. Manohara, "Toward a practical perceptual video quality metric," *The Netflix Tech Blog*, 2016.
21. G. Bjøntegaard, "Calculation of average PSNR differences between RD-curves," *13th VCEG Meeting*, Austin, TX, USA, Tech. Rep. ITU-T VCEG-M33, Apr. 2001.

FAN ZHANG is currently a Research Fellow with the Visual Information Laboratory, University of Bristol, Bristol, U.K. His research interests include perceptual video compression, video quality assessment, and immersive video formats. He has authored and coauthored more than 50 academic papers and has contributed to two books on video compression. He received the Ph.D. degree from the University of Bristol in 2012.

DI MA is currently working toward the Ph.D. degree with the Visual Information Laboratory, University of Bristol, Bristol, U.K. His research interests include perceptual video compression, computer vision, machine learning, and radar signal processing. He received the M.Eng. degree (Hons) from the Tianjin Key Lab for Advanced Signal Processing, Civil Aviation University of China, Tianjin, China, in 2016.

CHEN FENG is currently working toward the Ph.D. degree with the Visual Information Laboratory, University of Bristol, Bristol, U.K. His research interests include machine learning for video coding. He received the B.Sc. degree from the University of Science and Technology Beijing, Beijing, China, in 2018, and the M.Sc. degree from the University of Bristol in 2019.

DAVID R. BULL is currently a Professor of signal processing with the University of Bristol, Bristol, U.K. He heads its Visual Information Laboratory and is the Director of the Bristol Vision Institute. He is also the Director of the UK Government's new MyWorld Strength in Places programme. He has authored and coauthored more than 600 papers and 4 books on image and video processing for streaming and broadcast applications. He is the Fellow of the IEEE.