

Fractional Motion Estimation for Point Cloud Compression

Haoran Hong^{*}, Eduardo Pavez^{*}, Antonio Ortega^{*}, Ryosuke Watanabe[†],

Keisuke Nonaka[†]

^{*}University of Southern California, Los Angeles CA. 90089 USA

[†]KDDI Research, Inc., Japan

^{*}{haoranh, pavezcar, aortega}@usc.edu

[†]{ru-watanabe, ki-nonaka}@kddi-research.jp

Abstract

Motivated by the success of fractional pixel motion in video coding, we explore the design of motion estimation with fractional-voxel resolution for compression of color attributes of dynamic 3D point clouds. Our proposed block-based fractional-voxel motion estimation scheme takes into account the fundamental differences between point clouds and videos, i.e., the irregularity of the distribution of voxels within a frame and across frames. We show that motion compensation can benefit from the higher resolution reference and more accurate displacements provided by fractional precision. Our proposed scheme significantly outperforms comparable methods that only use integer motion. The proposed scheme can be combined with and add sizeable gains to state-of-the-art systems that use transforms such as Region Adaptive Graph Fourier Transform and Region Adaptive Haar Transform.

1 Introduction

Recent progress in 3D acquisition and reconstruction technology makes the capture of 3D scenes ubiquitous. In dynamic point clouds, each frame consists of a list of data points with 3D coordinates and RGB color values. Since point clouds in raw format would require a huge amount of bandwidth for transmission there has been a significant interest in point cloud compression techniques, which has led to MPEG standardization efforts [1], considering both video-based point cloud compression (V-PCC) and geometry-based point cloud compression (G-PCC) [1, 2].

Methods for inter-frame (temporal) prediction have been proposed to achieve efficient compression of dynamic point clouds. These methods can be grouped into three main categories. In *voxel-based* schemes [3], where a motion vector (MV) is estimated for each voxel, a few points in both the prediction and reference frames are selected as anchors to establish correspondence via spectral matching, leading to a set of sparse MVs. Then, using a smoothness constraint, a dense set of MVs can be obtained from the sparse set to provide motion for all remaining points. In *patch-based* techniques [4], motion estimation (ME) is considered as an unsupervised 3D point registration process wherein a MV is estimated by iterative closest point (ICP) [5] for each patch generated by K-means clustering. In this paper we focus on *block-based* methods, where frames to be predicted are partitioned into several non-overlapping

This work was funded in part by KDDI Research, Inc. and by the National Science Foundation (NSF CNS-1956190).

3-dimensional blocks of a given size. For each block, the best matching block in a reference frame is selected according to specific matching criteria, which can be based purely on geometry, e.g., an ICP-based approach that generates rigid transforms [6], or can use a combination of geometry and color attribute information [7]. Recent work has also focused on block-based motion search speedup, including both efficient search pattern design and search window reduction [8, 9, 10, 11].

Our work is motivated by the observation that ME with sub-pixel accuracy is an essential tool for modern video coding [12], while all the aforementioned ME methods for dynamic point clouds are based on integer-voxel displacements. There are two main reasons why a direct extension of video-based fractional ME to 3D contexts is not straightforward. First, point clouds are irregularly distributed within each frame, i.e., only those voxels that correspond to the surfaces of objects in the scene contain attribute information. Thus, while interpolation of attributes at new voxel locations can be based in conventional methods, we have the additional challenge of choosing only those new voxel locations that are consistent with object surfaces, even though those surfaces are not explicitly known. For example, we would like to avoid creating additional, fractional accuracy voxels *inside* an object. Second, voxels are inconsistent across frames, i.e., both the number of voxels and their distribution in space are different from frame to frame. Thus, since two matching blocks in consecutive frames will in general have a different number of voxels containing attribute information, we will need to develop alternatives to the one-to-one pixel (or sub-pixel) matching commonly used for conventional video.

In this paper, we focus on fractional-voxel motion estimation (FvME) under the assumption that integer-voxel MVs (IvMV) have already been obtained using an existing integer-voxel motion estimation (IvME) scheme [7, 8, 9, 10]. Specifically, in this paper we use precomputed IvMV from a public database [13]. In our approach, we start by creating fractional voxels between pairs of *neighboring* occupied integer voxels. Neighboring voxels are used to favor consistency with object surfaces, without requiring explicit estimation of the surfaces. Then, a higher resolution point cloud is obtained by interpolating attributes at each fractional voxel from the values at nearby integer voxels. FvME is implemented by searching fractional-voxel MVs (FvMV) around the positions given by IvMV and selecting the fractional displacement leading to the lowest motion compensation prediction error. Motion-compensated prediction is implemented by directly copying, as the attribute for a voxel in a block in the current frame, the attribute of the *nearest* voxel in the matched block in the reference frame. Our proposed FvME scheme leads to improved performance over transform-based approaches without inter or intra prediction and is also significantly better than temporal prediction methods based on the IvMV from the database [13].

2 Fractional-Voxel Motion Estimation and Compensation

2.1 Motivation

Real-world scenes and objects are captured by multiple calibrated and synchronized RGB or RGB-D cameras clusters from various viewing angles [14, 15]. After stitching

and voxelization, dynamic point clouds are generated on integer grids. Note that the 3D voxel coordinates are obtained as integer approximations to the “true” positions of the object in 3D space, while the optimal displacement between frames is unlikely to be exactly integer. Thus, a fractional voxel displacement can be better than an integer-one, so that higher resolution MVs have the potential to provide more accurate motion and hence more accurate MC prediction.

Furthermore, distortion due to lossy coding in previously reconstructed point cloud frames can lead to higher prediction errors, while camera noise, lighting change in the capture environment, object movements, etc., may also result in noisy color attributes and in imperfect matches during motion search [16]. Thus, as for conventional video, where it is well known that fractional motion compensation contributes to noise removal, the process of generating higher resolution point clouds and attributes at fractional voxel locations can contribute to denoising and lead to improvements in the quality of the reference frames.

2.2 Occupied fractional voxels

In this section, we define fractional voxels and describe our proposed method for interpolation. Based on the same design philosophy used for images and videos, fractional voxels are created at selected intermediate locations between voxels on the integer resolution grid. We define a fractional voxel of 1/2 resolution (1/2-voxel), as a voxel at the mid point between any two neighboring integer-voxels.

As noted in the introduction, not all integer-voxels are “occupied” in 3D space, and those that are occupied typically correspond to object surfaces. Thus, in our proposed method, new fractional voxels are created only at locations in 3D space that are (approximately) consistent with the surfaces implied the location of occupied integer voxels and attributes are interpolated only at these newly created fractional-voxels. We say that two integer voxels with coordinates v_j and v_k are neighbors if their distance is below a threshold ρ . Then, a fractional voxel is created only between neighbors v_j and v_k (assumed to be close enough so that they are likely to belong to the same surface) and the corresponding interpolated color attribute is computed as:

$$C(v_i) = \frac{1}{2} (C(v_j) + C(v_k)), \quad (1)$$

with $L(v_i) = \frac{1}{2} (L(v_j) + L(v_k))$ and $\text{dist}(v_j, v_k) \leq \rho, v_j, v_k \in V_i,$

where v_i is a voxel in the fractional-voxel set V_f with color signal $C(v_i)$, v_j and v_k are voxels in the integer-voxel set V_i with color signal $C(v_j)$ and $C(v_k)$, respectively. $L(\cdot)$ represents the coordinates of the voxel. ρ is the distance threshold and $\text{dist}(v_j, v_k)$ measures the Euclidean distance between the coordinates of v_j and v_k .

Note that different pairs of integer voxels may produce the same fractional voxel. Thus, to remove repeated fractional voxels after interpolation, attributes that belong to the same fractional voxel and are obtained by interpolation from different pairs of neighboring voxels are merged by averaging. Fig. 1b shows several examples of possible fractional-voxels locations, where we can see that interpolation based

on neighboring integer voxels tends to favor increasing the voxel resolution on the (implicit) surface where the voxels are located.

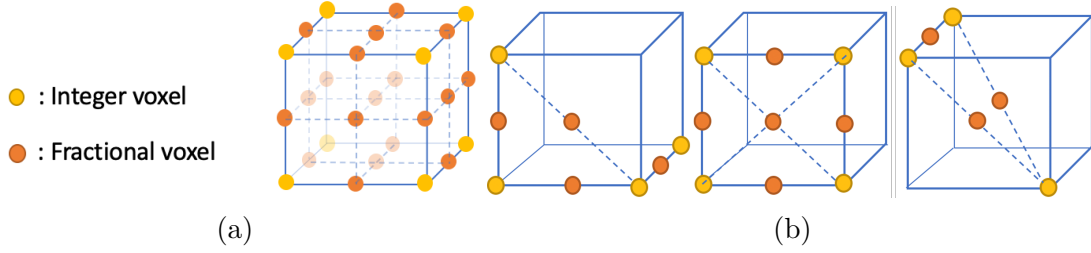


Figure 1: Integer and fractional voxels. Figure 1a depicts all possible candidate integer and 1/2-voxels positions. Figure 1b shows 3 examples of occupied integer voxel positions with corresponding fractional voxels obtained from neighboring integer voxels. Note that these interpolated fractional voxels are more likely to belong to the same surface as the neighboring integer voxels they were obtained from.

2.3 ME with fractional-voxel accuracy

Due to the inconsistency of voxel distributions in consecutive frames, it is difficult to establish exact one-to-one correspondences between the voxels in two matching blocks. To generalize MC prediction for fractional motion in 3D space, we start by super-resolving the reference frame as described in Section 2.2. As we can see from Fig. 2, the continuity among voxels and their corresponding attributes is significantly increased underlying surfaces, which provide better predictors when high resolution motion is available. The low pass filtering used for interpolation also contributes to attribute noise removal.

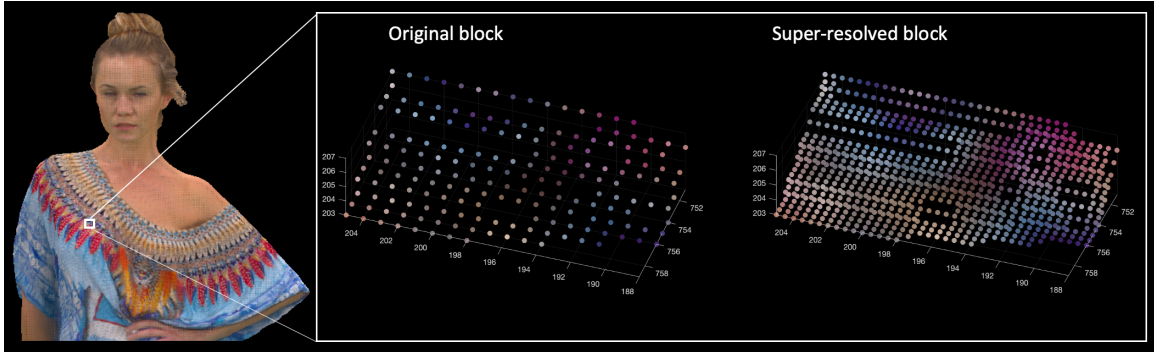


Figure 2: Comparison between the original and super-resolved reference block.

Next, we estimate MVs in fractional precision for MC. The entire ME process is a coarse-to-fine procedure, including IvME and FvME. Each estimated MV is obtained as the sum of an IvMV and a FvMV displacement. Assuming the IvMV MV_i is given, the optimal FvMV MV_f^{opt} is searched from a set of candidate fractional displacements. Since we super-resolve the reference frame in 1/2-voxel precision, each coordinate of

a fractional displacement MV_f can take values in $\{-\frac{1}{2}, 0, \frac{1}{2}\}$, resulting in 27 possible displacements. For a given fractional displacement MV_f , we predict each attribute in the current block from its nearest voxel in the translated super-resolved reference block, as depicted in Fig. 3. Then the displacement with the smallest prediction error is chosen, that is,

$$\begin{aligned}
MV_f^{opt} &= \arg \min_{MV_f} \sum_{v_i \in V(B_p)} E_{pred}(C(v_i), C(v_{j'})), \\
\text{s.t. } j' &= \arg \min_j (\text{dist}(v_i, v_j)), v_j \in V(B_{rMC}^s), \\
L_b(B_{rMC}^s) &= L_b(B_r^s) + MV, \\
MV &= MV_f + MV_i,
\end{aligned} \tag{2}$$

where B_r^s and B_{rMC}^s represent the super-resolved reference block before and after translation with MV , respectively, v_i and v_j are voxels with color signals $C(v_i)$ and $C(v_j)$ in blocks B_p and B_{rMC}^s , respectively. $E_{pred}(\cdot, \cdot)$ is the function for measuring the prediction error. $L_b(\cdot)$ represents the coordinates of the block while $V(\cdot)$ represents the set of voxels within the block.

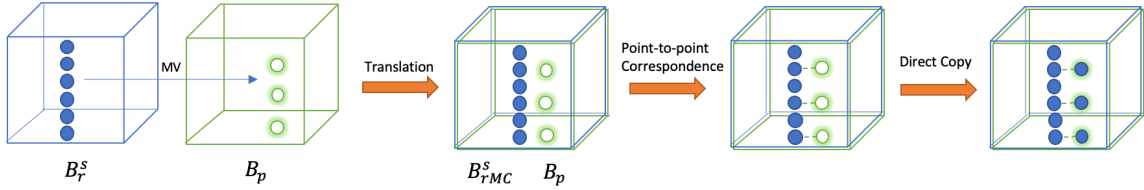


Figure 3: Motion-compensated prediction.

2.4 MC prediction with fractional-voxel accuracy

Finally, we apply MC prediction using the obtained MVs in fractional precision. Specifically, once the voxels in the reference block are translated using the integer motion vector MV_i , they are further shifted by the obtained optimal fractional displacement MV_f^{opt} , as shown in (2). Then, temporal correspondence are established from voxels in the predicted block B_p to their nearest neighbours in the translated super-resolved reference block B_{rMC}^s for motion-compensated prediction. The attribute of each voxel in the predicted block is predicted by copying the attribute of its corresponding voxel in the reference frame, that is,

$$\forall v_i \in B_p, C(v_i) = C(v_{j'}) \text{ s.t. } j' = \arg \min_j (\text{dist}(v_i, v_j)), v_j \in B_{rMC}^s. \tag{3}$$

3 Experiments

3.1 Dataset

In this section, we evaluate the proposed FvME scheme for compression of color attributes of the dataset of [15], which consists of four sequences: *longdress*, *redandblack*,

loot and *soldier*. Each sequence contains 300 frames.

Note that we assume IvMVs are given and are used to estimate FvMVs. Since IvMVs derived using different algorithms may lead to different FvMVs with disparate coding performance, we start from the publicly available 3D motion vector database [13]. The IvMVs in [13] are selected to minimize a hybrid distance metric, $\delta = \delta_g + 0.35\delta_c$, which combines δ_g , the average Euclidean distance between the voxels, and δ_c , the average color distance in Y-channel¹. We only consider motion for $16 \times 16 \times 16$ sized blocks. We implement a conventional inter-coding system where previously decoded frames are used as reference.

3.2 Experimental Settings

Following the MPEG Call for Proposal (CfP) for point cloud compression [17], we evaluate the proposed block-based FvME scheme (Proposed FvME) in groups of 32 frames, with the first frame coded in intra mode, and the rest coded using inter prediction. The threshold distance between integer voxels for interpolating fractional voxels is set to $\rho = \sqrt{3}$ in (1). Colors are transformed from RGB to YUV spaces, and each of the Y, U, and V channels are processed independently. When searching for the best candidate FvMV in (2), we use the squared distances to measure prediction errors. All blocks in the intra-coded frames undergo region adaptive graph Fourier transform (RA-GFT) [18] while, in the inter-coded frames, all blocks are motion-compensated. After MC prediction, the residues are transformed using the graph Fourier transform (GFT) [19]. To compute the GFT, a threshold graph is built inside every block wherein voxels are connected to each other if their Euclidean distance is less than or equal to $\sqrt{3}$. If after thresholding, the resulting graph for a block is not connected, a complete graph is built instead, which results in the transformed coefficients consisting of a single approximation coefficient (DC) and multiple detail coefficients (AC) for each block. The DC coefficients of all blocks are concatenated and encoded together. Then, the AC coefficients are coded block by block. This approach is equivalent to a single level RA-GFT [18]. For all transforms, we perform uniform quantization and entropy code the coefficients using the adaptive run-length Golomb-Rice algorithm (RLGR) [20]. As for FvMVs overheads, since there are 27 FvMVs in total, 8 bits are used to signal each FvMV. For IvMVs, we use 4 bits to signal the value and 1 bit to signal sign for each axis and therefore, 15 total bits are used to represent an IvMV. The overheads of FvMVs and IvMVs are entropy coded by Lempel–Ziv–Markov chain algorithm [21].

We considered the following schemes as baselines. IvME using the database motion (DM) for MC prediction and using DM with additional integer local refinement (DM+RF) for MC prediction. The local refinement uses different criteria that aims

¹Note that the resulting IvMVs aim to select matching blocks with similar geometry (δ_g) and color attributes (δ_c) but there is no guarantee that this metric, and in particular the relative weight between the distances (0.35) is the optimal choice in terms of coding efficiency. Thus, it will be shown in our motion compensation experiments, that these IvMVs can sometimes lead to performance below that of encoding methods that do not use motion compensation. In these cases performance can be improved by local refinement of the IvMVs from the database.

to minimize color errors only, instead of the hybrid errors used in the database. DM is refined by additional local search in integer precision to improve its matching accuracy over the original ones. The locally refined range for each axis is set to be $[-1, 1]$, which entirely encloses fractional positions searched in the proposed FvME scheme.

To evaluate the benefits of high resolution references and FvMVs, we propose two inter coding schemes which are using super-resolved reference blocks, with and without fractional motion vectors for compensated prediction. First, to evaluate the super-resolution method, we implement a scheme that considers IvME using integer local refined DM and super-resolved reference blocks for prediction, which is denoted by “DM+RF+SR”. The difference between DM+RF and DM+RF+SR is the resolution of the reference block. Then, to evaluate benefits of FvMVs, we implement a scheme that uses fractional resolution in both reference blocks and motion vectors, which is denoted by “proposed FvME”. For a fair comparison between inter coding schemes, all other test conditions are the same.

Additionally, to make our performance evaluation more complete, we include two state of the art (all intra) anchor solutions, namely, RA-GFT [18] and region adaptive Haar transform (RAHT) [22]. For RA-GFT, block size 16 is used. The residues are entropy coded by RLGR.

3.3 Evaluation Metrics

The evaluation metrics are the number of bits per voxel (bpv) and average peak signal-to-noise ratio over Y component (PSNR-Y),

$$PSNR_Y = -10 \log_{10} \left(\frac{1}{T} \sum_{t=1}^T \frac{\|Y_t - \hat{Y}_t\|_2^2}{255^2 N_t} \right), bpv = \frac{\sum_{t=1}^T b_t}{\sum_{t=1}^T N_t}, \quad (4)$$

where Y_t and $\hat{Y}_t$ represent original and reconstructed signals on the same voxels of t -th frame respectively, T is the total number of frames, b_t is the bits required to encode YUV components of t -th, including IvMVs and FvMVs overhead when necessary, and N_t is the total number of occupied voxels in t -th frame. The Bjontegaard-Delta [23] results for bitrate (BD-rate) are also reported.

3.4 Experimental Results and Analysis

Rate distortion (RD) curves are shown in Fig. 4. We first note that using only the original IvME from the database [13], results in sub-optimal performance compared to RAHT and RA-GFT. This is in part due to the criteria used in [13] to choose the optimal MV based on geometry and color information. After local refinement with integer precision, the performance of IvME (DM+RF) improves significantly with respect to IvME (DM) but it is still far from being competitive with other techniques. Further improvements have been shown to be achievable by using per block intra/inter mode decision [10].

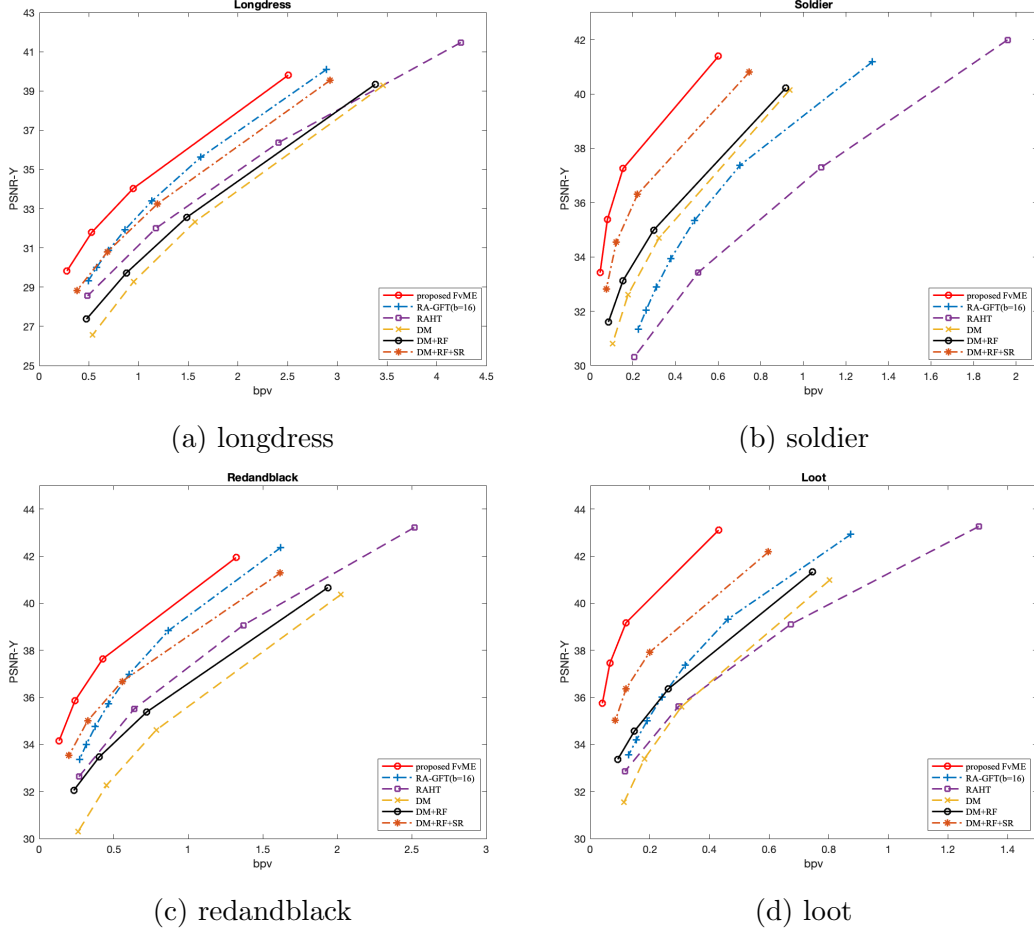


Figure 4: Rate distortion curves of 8iVFBv2 sequences.

anchors \ sequences	<i>longdress</i>	<i>soldier</i>	<i>redandblack</i>	<i>loot</i>
IvME(DM+RF)	−43.93%	−63.96%	−57.98%	−64.43%
RAHT	−39.94%	−81.91%	−51.10%	−73.69%
RA-GFT(b=16)	−16.19%	−72.46%	−24.37%	−61.44%

Table 1: The BD-rate performances of the proposed scheme over baselines

After the reference blocks are super-resolved, the performance of the proposed IvME (DM+RF+SR) is further improved with respect to DM+RF, even without increasing MV resolution. The DM+RF+SR scheme can be better than the intra schemes in some cases with the advantage of complexity lower than that of the proposed FvME. Finally, after we increase MV resolution to 1/2-voxel, further coding gains are obtained, outperforming intra coding baselines, RA-GFT and RAHT, with average gains of $2.8dB$ and $4.6dB$, respectively. The method is always better than DM+RF+SR but at the cost of higher complexity due to additional motion search. The results show that both interpolated fractional voxels and high resolution MVs lead to higher coding gain and outperform both inter coding with IvME and non predictive transform based schemes.

Table 1 summarizes the performance of the proposed method over IvME (DM+RF), RA-GFT, and RAGT in terms of BD-rate. The proposed FvME can achieve 57% average bitrate reduction over IvME (DM+RF). Compared with the prior arts, the proposed scheme can achieve 61% and 43% bitrate reduction on average over RAHT and RA-GFT respectively.

Compared with IvME (DM+RF), the proposed FvME scheme increases the number of voxels at most eightfold (because of super resolution), and requires evaluating additional 27 fractional displacements for ME. Therefore, the complexity of the proposed FvME is larger than the complexity of IvME (DM+RF) by a constant factor (independent of the point cloud size).

4 Conclusions

This paper describes a fractional-voxel motion estimation scheme tailored for attribute compression in 3D dynamic point clouds. Our scheme defines and identifies fractional voxels to be interpolated and provides a motion compensation prediction method by super-resolution and temporal correspondence. Extensive experiments show superior performance over the prior art. Our work reveals the benefits given by high resolution references and the further improvements given by fractional-voxel motion vectors for dynamic point clouds color compression.

5 References

- [1] Sebastian Schwarz, Marius Preda, Vittorio Baroncini, Madhukar Budagavi, Pablo Cesar, Philip A. Chou, Robert A. Cohen, Maja Krivokuća, Sébastien Lasserre, Zhu Li, Joan Llach, Khaled Mammou, Rufael Mekuria, Ohji Nakagami, Ernestasia Siahaan, Ali Tabatabai, Alexis M. Tourapis, and Vladyslav Zakharchenko, “Emerging MPEG standards for point cloud compression,” *IEEE Journal on Emerging and Selected Topics in Circuits and Systems*, vol. 9, no. 1, pp. 133–148, 2019.
- [2] D. Graziosi, O. Nakagami, S. Kuma, A. Zaghetto, T. Suzuki, and A. Tabatabai, “An overview of ongoing point cloud compression standardization activities: video-based (V-PCC) and geometry-based (G-PCC),” *APSIPA Transactions on Signal and Information Processing*, vol. 9, pp. e13, 2020.
- [3] Dorina Thanou, Philip A. Chou, and Pascal Frossard, “Graph-based compression of dynamic 3D point cloud sequences,” *IEEE Transactions on Image Processing*, vol. 25, no. 4, pp. 1765–1778, 2016.
- [4] Yiqun Xu, Wei Hu, Shanshe Wang, Xinfeng Zhang, Shiqi Wang, Siwei Ma, Zongming Guo, and Wen Gao, “Predictive generalized graph Fourier transform for attribute compression of dynamic point clouds,” *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 31, no. 5, pp. 1968–1982, 2021.
- [5] P.J. Besl and Neil D. McKay, “A method for registration of 3-d shapes,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 14, no. 2, pp. 239–256, 1992.
- [6] Rufael Mekuria, Kees Blom, and Pablo Cesar, “Design, implementation, and evaluation of a point cloud codec for tele-immersive video,” *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 27, no. 4, pp. 828–842, 2017.

- [7] Ricardo L. de Queiroz and Philip A. Chou, "Motion-compensated compression of dynamic voxelized point clouds," *IEEE Transactions on Image Processing*, vol. 26, no. 8, pp. 3886–3895, 2017.
- [8] Camilo Dorea and Ricardo L. de Queiroz, "Block-based motion estimation speedup for dynamic voxelized point clouds," in *2018 25th IEEE International Conference on Image Processing (ICIP)*, 2018, pp. 2964–2968.
- [9] Camilo Dorea, Edson M. Hung, and Ricardo L. de Queiroz, "Local texture and geometry descriptors for fast block-based motion estimation of dynamic voxelized point clouds," in *2019 IEEE International Conference on Image Processing (ICIP)*, 2019, pp. 3721–3725.
- [10] André L. Souto and Ricardo L. de Queiroz, "On predictive RAHT for dynamic point cloud coding," in *2020 IEEE International Conference on Image Processing (ICIP)*, 2020, pp. 2701–2705.
- [11] Cristiano Santos, Mateus Gonçalves, Guilherme Corrêa, and Marcelo Porto, "Block-based inter-frame prediction for dynamic point cloud compression," in *2021 IEEE International Conference on Image Processing (ICIP)*, 2021, pp. 3388–3392.
- [12] B. Girod, "Motion-compensating prediction with fractional-pel accuracy," *IEEE Transactions on Communications*, vol. 41, no. 4, pp. 604–612, 1993.
- [13] Ricardo L. de Queiroz Andre L. Souto and Camilo Dorea, "A 3d motion vector database for dynamic point clouds," 2020.
- [14] Kyungjin Lee, Juheon Yi, Youngki Lee, Sunghyun Choi, and Young Min Kim, "Groot: A real-time streaming system of high-fidelity volumetric videos," in *Proceedings of the 26th Annual International Conference on Mobile Computing and Networking*, New York, NY, USA, 2020, MobiCom '20, Association for Computing Machinery.
- [15] Taos Myers Eugene d'Eon, Bob Harrison and Philip A. Chou, "8i voxelized full bodies - a voxelized point cloud dataset," *ISO/IEC JTC1/SC29 Joint WG11/WG1 (MPEG/JPEG) input document WG11M40059/WG1M74006*, January 2017.
- [16] Zhao Wang, Shiqi Wang, Jian Zhang, and Siwei Ma, "Adaptive progressive motion vector resolution selection based on rate-distortion optimization," *IEEE Transactions on Image Processing*, vol. 26, no. 1, pp. 400–413, 2017.
- [17] "Call for proposals for point cloud compression v2," *MPEG 3DG and Requirements, ISO/IEC JTC1/SC29/WG11 MPEG2017/N16763*, 2017.
- [18] Eduardo Pavez, Benjamin Girault, Antonio Ortega, and Philip A. Chou, "Region adaptive graph Fourier transform for 3D point clouds," in *2020 IEEE International Conference on Image Processing (ICIP)*, 2020, pp. 2726–2730.
- [19] Cha Zhang, Dinei Florencio, and Charles Loop, "Point cloud attribute compression with graph transform," in *2014 IEEE International Conference on Image Processing (ICIP)*. IEEE, 2014, pp. 2066–2070.
- [20] H.S. Malvar, "Adaptive run-length/Golomb-Rice encoding of quantized generalized gaussian sources with unknown statistics," in *Data Compression Conference (DCC'06)*, 2006, pp. 23–32.
- [21] I. Pavlov, "7z format," <https://www.7-zip.org/7z.html>.
- [22] Ricardo L. de Queiroz and Philip A. Chou, "Compression of 3D point clouds using a region-adaptive hierarchical transform," *IEEE Transactions on Image Processing*, vol. 25, no. 8, pp. 3947–3956, 2016.
- [23] G. Bjøntegaard, "Calculation of average PSNR differences between rd curves," *ITU-T SG 16/Q.6 13th VCEG Meeting, Austin, Texas, USA, document VCEG-M33*, 2001.