

2 **Supplementary Material**

1 **NOTATION**

3 Notation is given in Table S1.

2 **ADDITIONAL EXPERIMENTAL DETAILS**

4 Table S2 lists the point clouds in our dataset and the number of points in each point cloud. The point clouds
 5 in the top part of the table are point clouds derived from meshes created in Guo et al. (2019); Meka et al.
 6 (2020), while the point clouds in the bottom part of the table are MPEG point clouds d'Eon et al. (2017);
 7 Alliez et al. (2017). The point clouds are shown in Fig. 1 in the main text and in Fig. S1 below. Point cloud
 8 visualization is performed in Meshlab Cignoni et al. (2008).



Figure S1: Point clouds *longdress*, *loot*, *redandblack*, *soldier*, *mask*, *shiva*, and *klimt*.

9 We implement the LVAC framework in Python using Tensorflow. The entire point cloud constitutes one
 10 batch. All configurations are trained in about 25K steps using the Adam optimizer and a learning rate of
 11 0.01, with low bit rate configurations typically taking longer to converge. Each step takes 0.5-3.0 s on
 12 an NVIDIA P100 class GPU in eager mode with various debugging checks in place. Training loss as a
 13 function of steps for point cloud *rock* is shown in Fig. S2.

3 **COORDINATE BASED NETWORKS**

14 Figure S3 shows the RD performance of different networks: (left) $mlp(35 \times 256 \times 3)$, (middle) $mlp(35 \times 64 \times 3)$,
 15 and (right) $pa(3 \times 32 \times 3)$ for point cloud *rock*, along with baselines. We give the corresponding plots for other
 16 point clouds in Fig. S7, Fig. S8, Fig. S9, and Fig. S10. At higher bit rates, higher target levels perform
 17 better.

4 **GENERALIZATION**

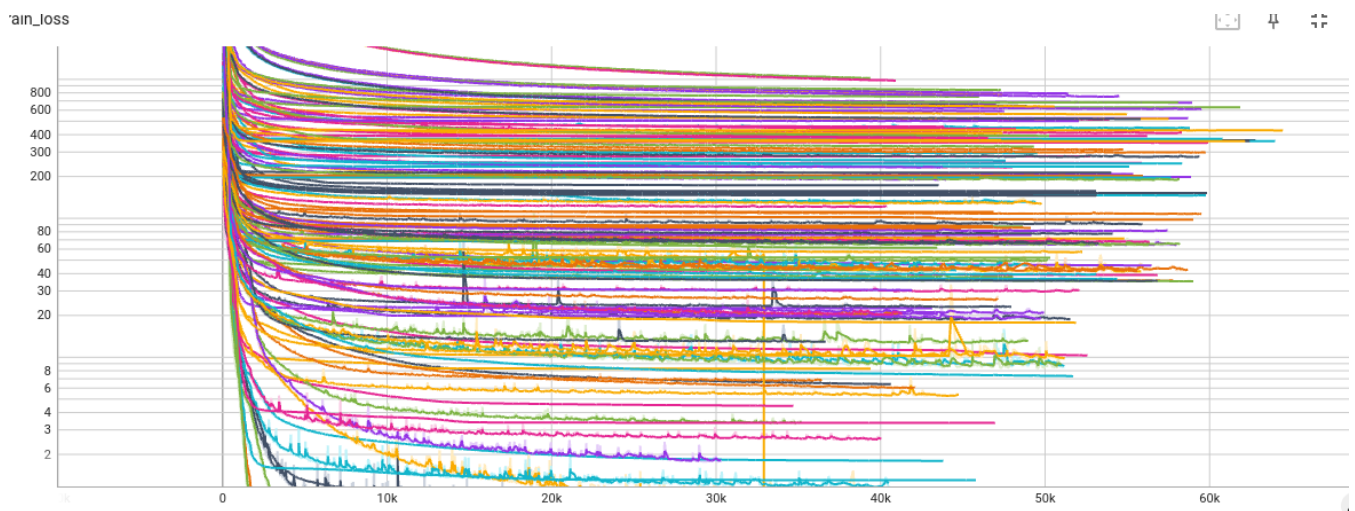
18 In the main text, we provide the RD performance on point cloud *rock* using CBNs pre-trained on point
 19 cloud *basketball*. The corresponding plots for other point clouds are given in Fig. S11 and Fig. S12.

symbol	description
$\mathbf{x}$	position vector in $\mathbb{R}^d$, typically $d = 3$ with $\mathbf{x} = (x, y, z)$
$\mathbf{y}$	attribute vector in $\mathbb{R}^r$, typically $\mathbf{y} = (R, G, B)$
$f : \mathbf{x} \mapsto \mathbf{y}$	volumetric function
$f_{\theta, \mathbf{Z}} : \mathbf{x} \mapsto \mathbf{y}$	volumetric function in family parameterized by $\theta, \mathbf{Z}$
$d(f, f_{\theta, \mathbf{Z}})$	distortion between two volumetric functions
$\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^{N_p}$	point cloud with N_p points or N_p occupied voxels in voxel grid
L	binary tree depth, or $\log_2$ of number of all voxels in voxel grid
$\ell = 0, 1, \dots, L$	level of binary tree between root ($\ell = 0$) and leaves ($\ell = L$)
$\mathbf{n}, \mathbf{n}_L$, and $\mathbf{n}_R$	positions of a block at level ℓ and its left and right children at level $\ell + 1$
$B, B_{\mathbf{n}}$, or $B_{\ell, \mathbf{n}}$	block at level ℓ and position (or <i>offset</i>) $\mathbf{n}$
$N_x \times N_y \times N_z$	dimensions of a block, in voxels
$w, w_{\mathbf{n}}$, or $w_{\ell, \mathbf{n}}$	weight of block (i.e., number of point cloud points $\mathbf{x}_i$ in block)
$\mathbf{z}, \mathbf{z}_{\mathbf{n}}$, or $\mathbf{z}_{\ell, \mathbf{n}}$	C -dimensional latent vector for block
$f_{\theta}(\mathbf{x} - \mathbf{n}; \mathbf{z}_{\mathbf{n}})$	continuous function representing attributes across block at offset $\mathbf{n}$
θ	parameters of the coordinate based network (CBN) representing f_{θ}
$\mathbf{Z}$	$N \times C$ matrix of latent vectors, at a target level having N blocks
$\mathbf{V}$	$N \times C$ matrix of transform coefficients of $\mathbf{Z}$
$\mathbf{T}_a$	$N \times N$ analysis transform for computing $\mathbf{V} = \mathbf{T}_a \mathbf{Z}$
$\mathbf{T}_s$	$N \times N$ synthesis transform for computing $\mathbf{Z} = \mathbf{T}_s \mathbf{V}$
$\mathbf{S}$	$N \times N$ diagonal matrix of spatial normalizing weights
Δ	$C \times C$ diagonal matrix of quantization stepsizes
$\tilde{\mathbf{V}} = \mathbf{S}^{-1} \mathbf{V}$	$N \times C$ matrix of spatially-normalized transform coefficients
$\mathbf{U} = \tilde{\mathbf{V}} \Delta^{-1}$	$N \times C$ matrix of stepsize-scaled spatially-normalized coefficients
$\hat{\mathbf{U}} = \lfloor \mathbf{U} \rfloor$	$N \times C$ matrix of integers for entropy coding
$\hat{\mathbf{V}} = \hat{\mathbf{U}} \Delta$	$N \times C$ matrix of quantized spatially-normalized transform coefficients
$\hat{\mathbf{V}} = \mathbf{S} \hat{\mathbf{V}}$	$N \times C$ matrix of quantized transform coefficients
$\hat{\mathbf{Z}} = \mathbf{T}_s \hat{\mathbf{V}}$	$N \times C$ matrix of quantized latent vectors
$\hat{\theta}$	quantized parameters reconstructed at decoder along with $\hat{\mathbf{Z}}$
$f_{\hat{\theta}, \hat{\mathbf{Z}}} : \mathbf{x} \mapsto \mathbf{y}$	reconstructed volumetric function $f_{\hat{\theta}, \hat{\mathbf{Z}}}(\mathbf{x}) = \sum_{\mathbf{n}} f_{\hat{\theta}}(\mathbf{x} - \mathbf{n}; \hat{\mathbf{z}}_{\mathbf{n}}) \mathbb{1}_{B_{\mathbf{n}}}(\mathbf{x})$
D or $D(\hat{\theta}, \hat{\mathbf{Z}})$	distortion $d(f, f_{\hat{\theta}, \hat{\mathbf{Z}}})$, typically $\sum_i \ \mathbf{y}_i - f_{\hat{\theta}, \hat{\mathbf{Z}}}(\mathbf{x}_i)\ ^2$
R or $R(\hat{\theta}, \hat{\mathbf{Z}})$	bit rate including bits both for $\hat{\theta}$ and $\hat{\mathbf{Z}}$
J or $J(\hat{\theta}, \hat{\mathbf{Z}})$	loss function, or Lagrangian $D + \lambda R$, with Lagrange multiplier $\lambda > 0$
<i>linear</i> (3x3)	linear model with 3 input and 3 output channels
<i>mlp</i> (35x256x3)	MLP with 35=3+C input and 3 output channels, and 256 hidden units
<i>mlp</i> (35x64x3)	MLP with 35=3+C input and 3 output channels, and 64 hidden units
<i>pa</i> (3x32x3)	PA network with 3 input and 3 output channels, and 32 modulation inputs
$l=21, 24, 27, 30$	target levels (in binary tree where CBNs are located)
<i>CBN</i>	Coordinate Based Network
<i>LVAC</i>	Learned Volumetric Attribute Coding (our framework)
<i>RAHT</i>	Region Adaptive Hierarchical Transform
<i>RLGR</i>	adaptive Run-Length Golomb-Rice entropy coder (Malvar, 2006)
<i>cbe</i>	Continuous Batched Entropy model
<i>G-PCC</i>	MPEG Geometry-based Point Cloud Codec
<i>TMC13</i>	Test Model 13 (software) for MPEG G-PCC
<i>Deep PCAC</i>	Deep Point Cloud Attribute Coder (Sheng et al., 2021)

Table S1. Notation

Table S2. Voxelized point clouds in our dataset.

Point Cloud	# points
<i>rock</i>	837434
<i>chair</i>	791416
<i>scooter</i>	959388
<i>juggling</i>	798441
<i>basketball</i>	868224
<i>basketball2</i>	948870
<i>jacket</i>	805882
<i>longdress</i> (aka <i>longdress_vox10_1300</i>)	857966
<i>loot</i> (aka <i>loot_vox10_1200</i>)	805285
<i>redandblack</i> (aka <i>redandblack_vox10_1550</i>)	757691
<i>soldier</i> (aka <i>soldier_vox10_0690</i>)	1089091
<i>mask</i> (aka <i>Egyptian_mask_vox12</i>)	272684
<i>shiva</i> (aka <i>Shiva_00035_vox12</i>)	1009132
<i>klimt</i> (aka <i>Staue_Klimt_vox12</i>)	499660

**Figure S2:** Training loss as a function of steps on point cloud *rock*.

5 SIDE INFORMATION

20 In the main text, we show the penalty required to transmit side information, for the entropy model and
 21 for the CBNs for point cloud *rock*. The corresponding plots for other point clouds are given in Fig. S14,
 22 Fig. S15, and Fig. S16.

6 NORMALIZATION

23 We show the bit rate improvement due to normalization for point cloud *rock* in the main text. We give the
 24 corresponding results for other point clouds in Fig. S17 and Tab. S3.

7 SUBJECTIVE QUALITY

25 Figure S4 shows further results on subjective quality.

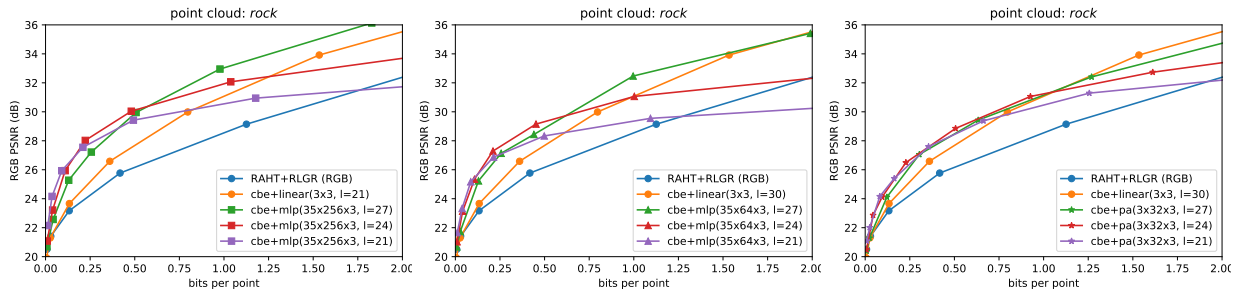


Figure S3: Coordinate Based Networks, by network. (left) $mlp(35 \times 256 \times 3)$. (middle) $mlp(35 \times 64 \times 3)$. (right) $pa(3 \times 32 \times 3)$.

8 ADDITIONAL POINT CLOUDS

Figure S5, Fig. S7, Fig. S9, Fig. S11, Fig. S12, Fig. S13, Fig. S14, Fig. S15, Fig. S16, Tab. S3, Fig. S17, and Fig. S18 show results for point clouds *chair*, *scooter*, *juggling*, *basketball*, *basketball2*, and *jacket* corresponding respectively to Fig. 5(a), Fig. 6, Fig. S3, Fig. 8, Fig. 9, Fig. 10, Table 1, Fig. 11, and Fig. 12 for *rock*. Figure S6, Fig. S8, and Fig. S10 show results for MPEG point clouds *longdress*, *loot*, *redandblack*, *soldier*, *mask*, *shiva*, and *klimt* corresponding respectively to Fig. 5(b), Fig. 6, and Fig. S3.

MPEG point cloud data from Fig. S8 and Fig. S10 are used in Tab. S4 to compare LVAC to Sparse PCAC (Wang and Ma, 2022), showing that LVAC and Sparse PCAC are roughly comparable. Sparse PCAC appeared while this work was in review, and Wang and Ma (2022) cite our LVAC pre-print, but we do our best to compare them here. Beware, however, that Figs. S8 and S10 — like Figs. 6 and S3 but unlike Figs. 8, 9, 10, 12, S14, S15, S16, and S18 — and hence Tab. S4, do not reflect the side information needed to completely specify the decoding. To assess this side information, please see Sec. 5.



Figure S4: Subjective quality of point cloud *rock*. Each row is a different bit rate. Original is shown in the main text. Zoom in to see differences.

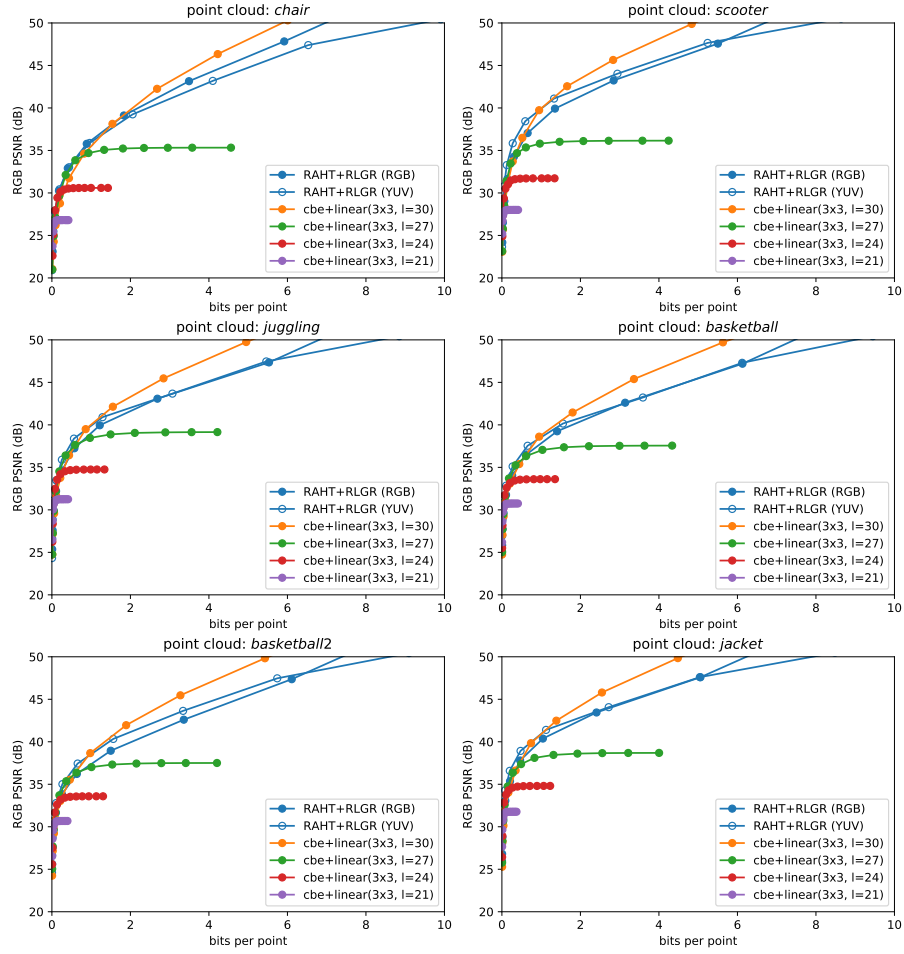


Figure S5: Baselines for six point clouds. *RAHT+RLGR (RGB)* and *(YUV)* are shown against 3×3 linear models at levels 30, 27, 24, and 21, which optimize the colorspace by minimizing $D + \lambda R$ using the *cbe* entropy model.

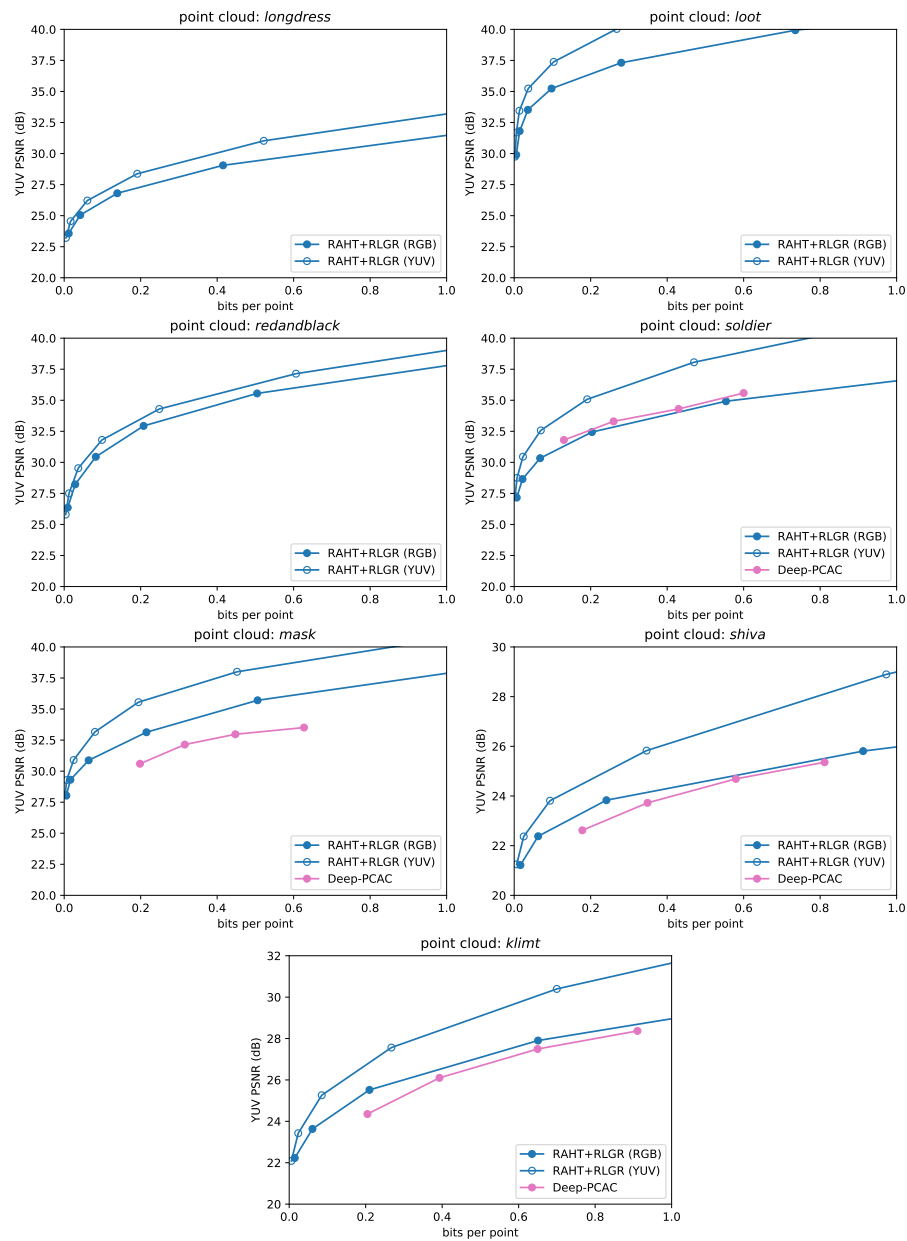


Figure S6: Baselines for MPEG point clouds: *RAHT+RLGR (RGB)*, *RAHT+RLGR (YUV)*, and (for some point clouds) Deep-PCAC Sheng et al. (2021). Baselines for other point clouds are shown in Fig. S5.

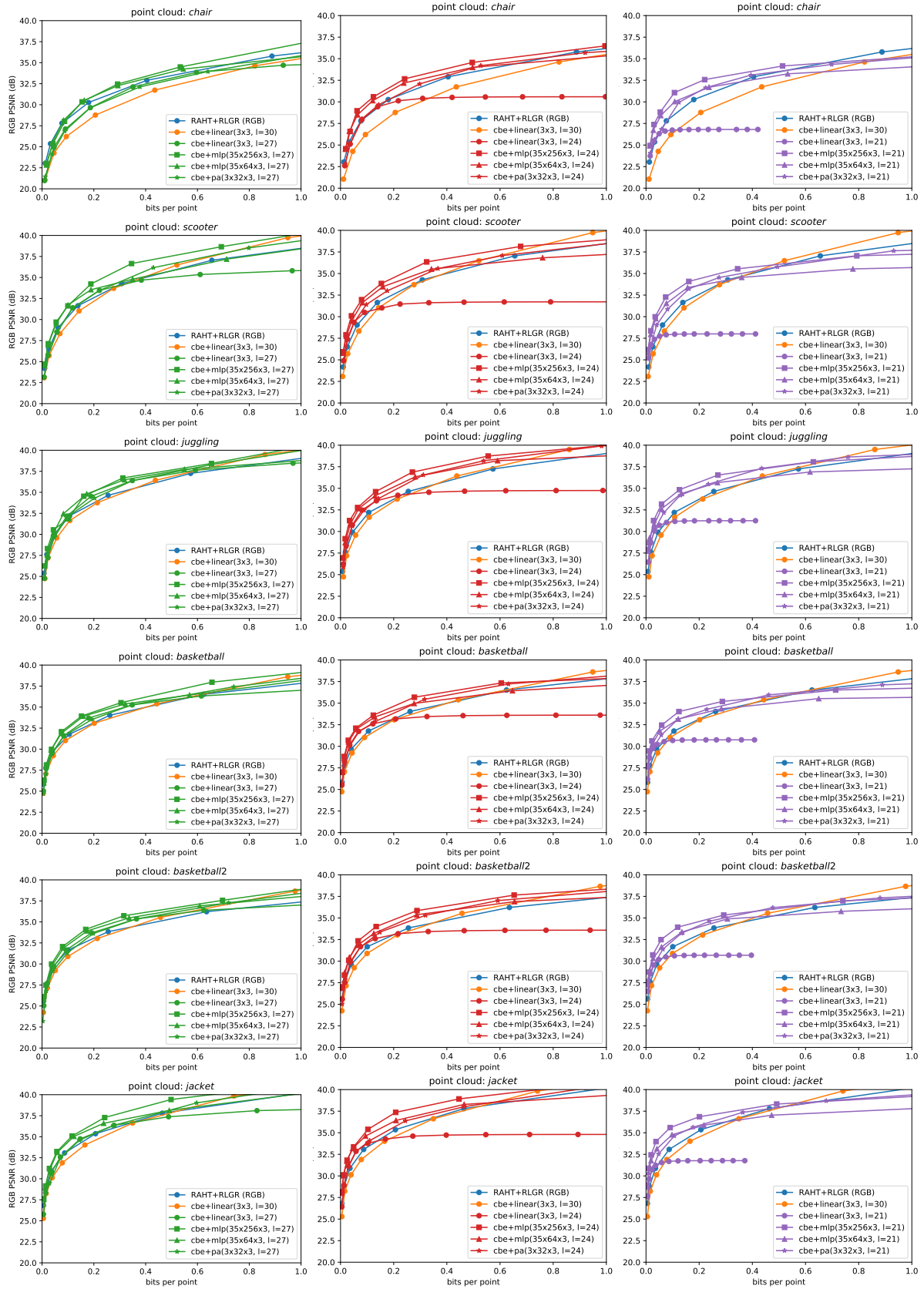


Figure S7: Coordinate Based Networks, by target level. Each row is a different point cloud. Left, middle, right columns each show $mlp(35 \times 256 \times 3)$, $mlp(35 \times 64 \times 3)$, and $pa(3 \times 32 \times 3)$ CBNs, along with baselines, at levels 27, 24, 21. See Fig. S8 for other point clouds.

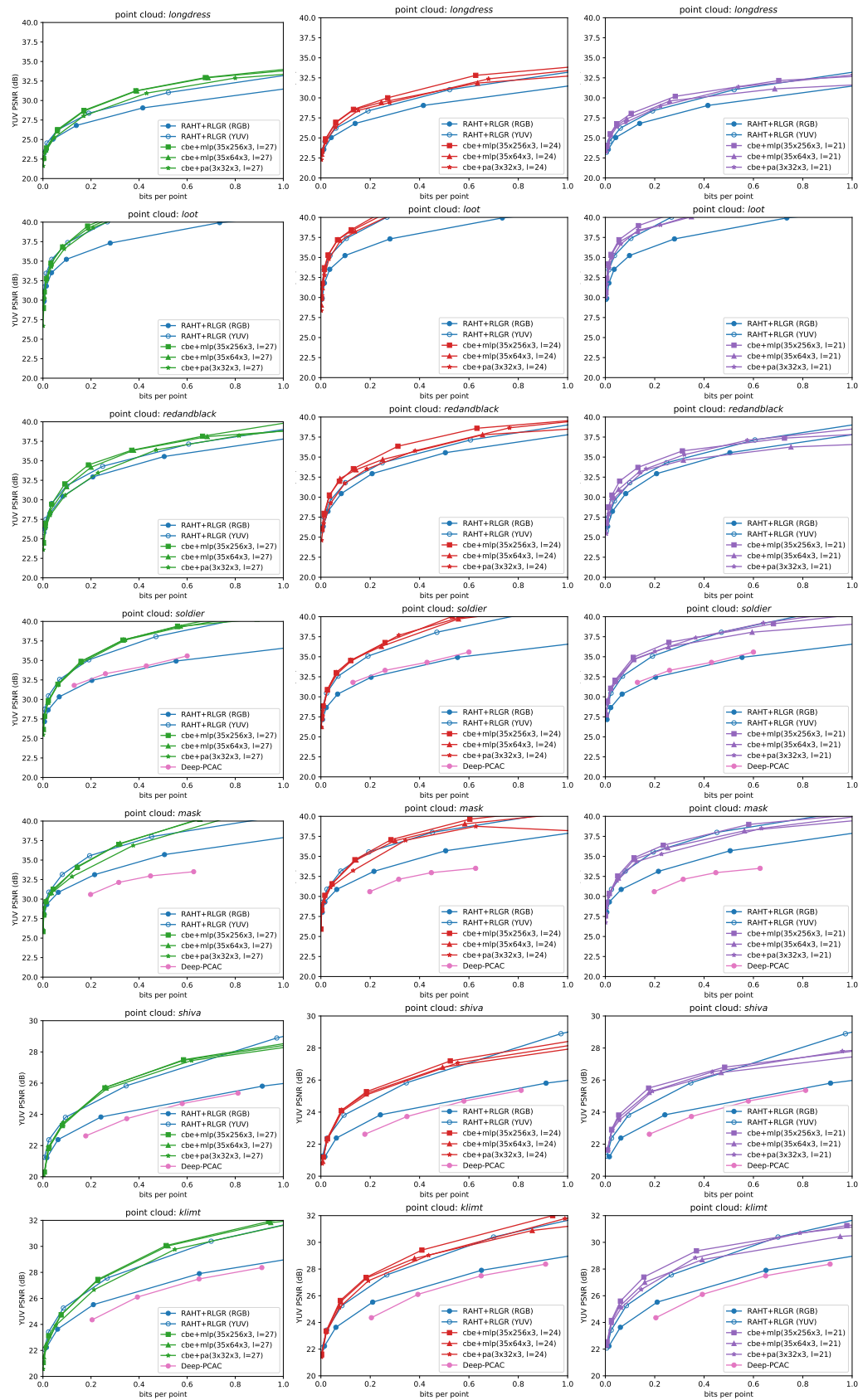


Figure S8: Coordinate Based Networks, by target level, for MPEG point clouds. Each row is a different point cloud. Left, middle, right columns each show $mlp(35 \times 256 \times 3)$, $mlp(35 \times 64 \times 3)$, and $pa(3 \times 32 \times 3)$ CBNs, along with baselines, at levels 27, 24, 21. See Fig. S7 for other point clouds.

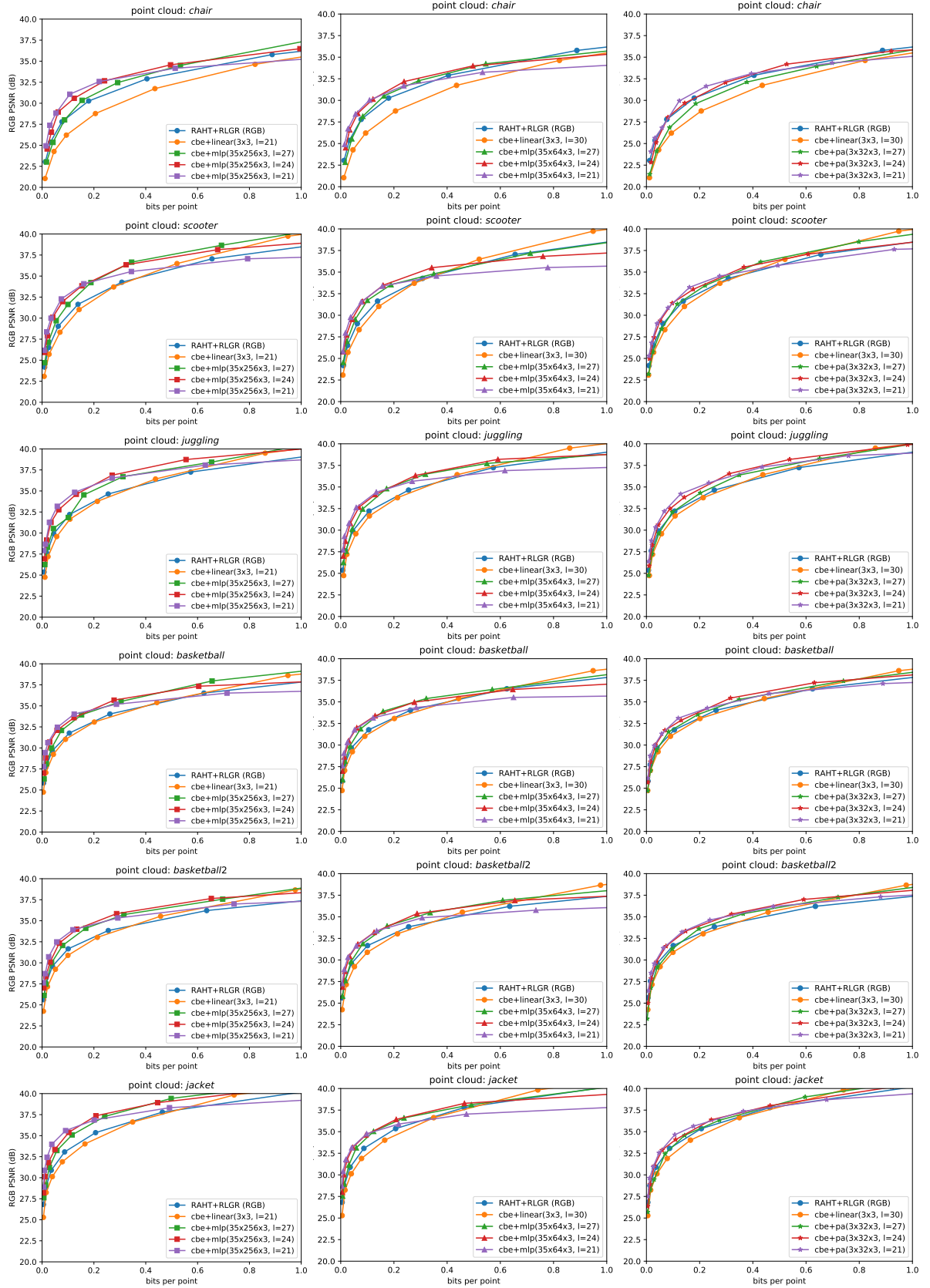


Figure S9: Coordinate Based Networks, by network. Each row is a different point cloud. Left, middle, right columns each show levels 27, 24, and 21, along with baselines, for CBNs $mlp(35 \times 256 \times 3)$, $mlp(35 \times 64 \times 3)$, and $pa(3 \times 32 \times 3)$. See Fig. S3 for point cloud *rock* and Fig. S10 for other point clouds.

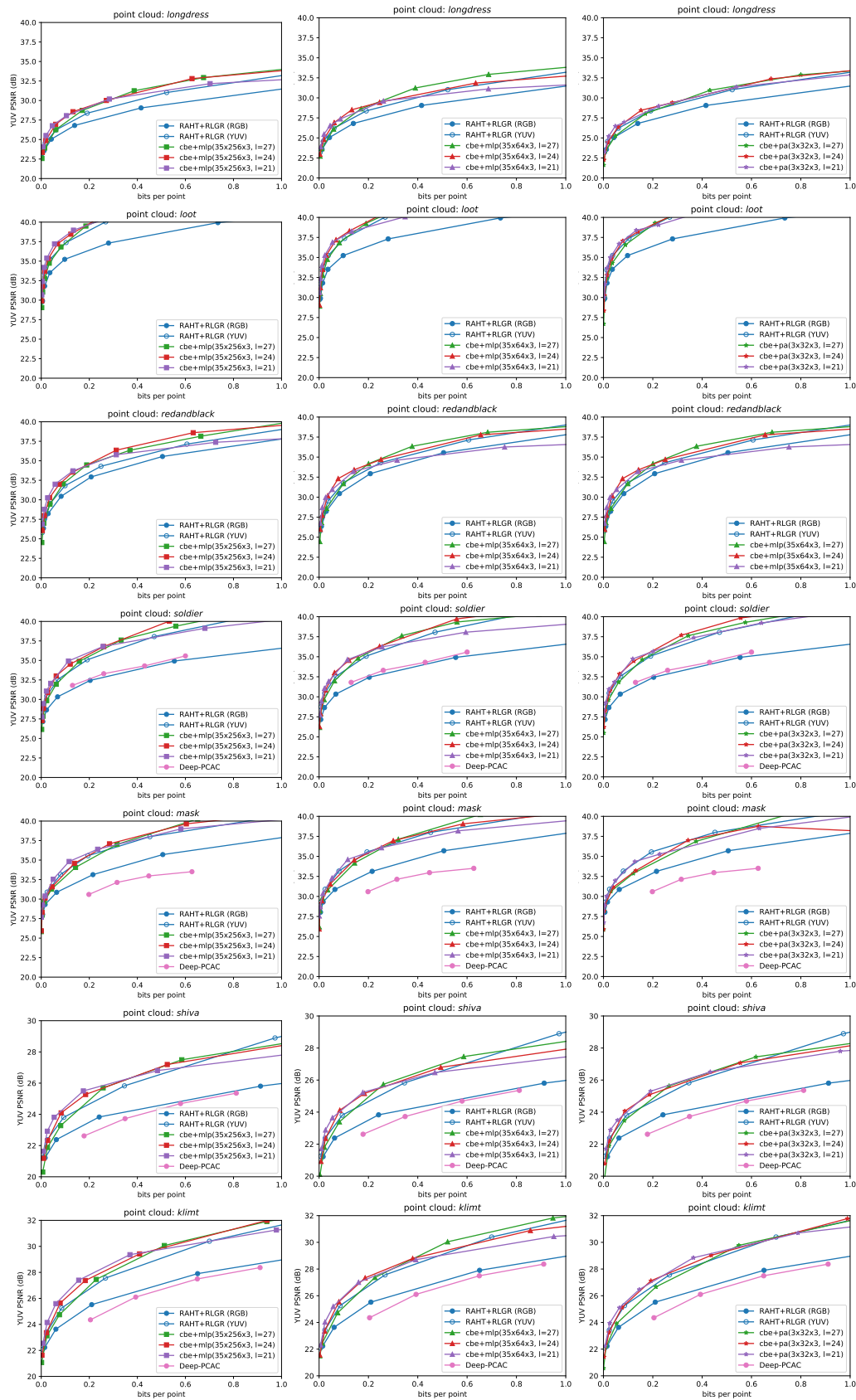


Figure S10: Coordinate Based Networks, by network, for MPEG point clouds. Each row is a different point cloud. Left, middle, right columns each show levels 27, 24, and 21, along with baselines, for CBNs *mlp(35x256x3)*, *mlp(35x64x3)*, and *pa(3x32x3)*. See Fig. S3 and Fig. S9 for other point clouds.

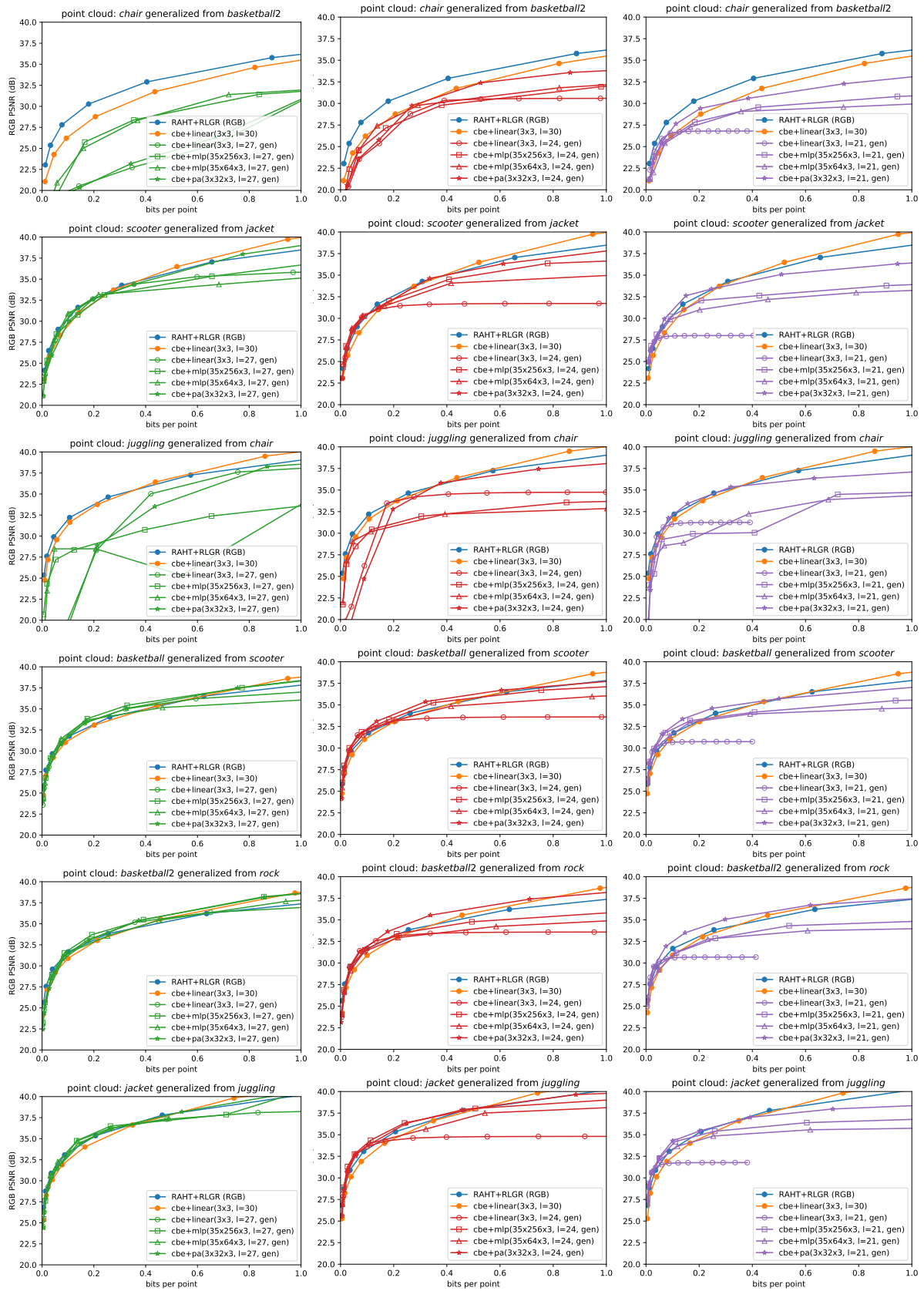


Figure S11: Coordinate Based Networks with generalization, by target level. Each row is a different point cloud. Left, middle, right columns each show $mlp(35 \times 256 \times 3)$, $mlp(35 \times 64 \times 3)$, and $pa(3 \times 32 \times 3)$ CBNs, along with baselines, at levels 27, 24, 21.

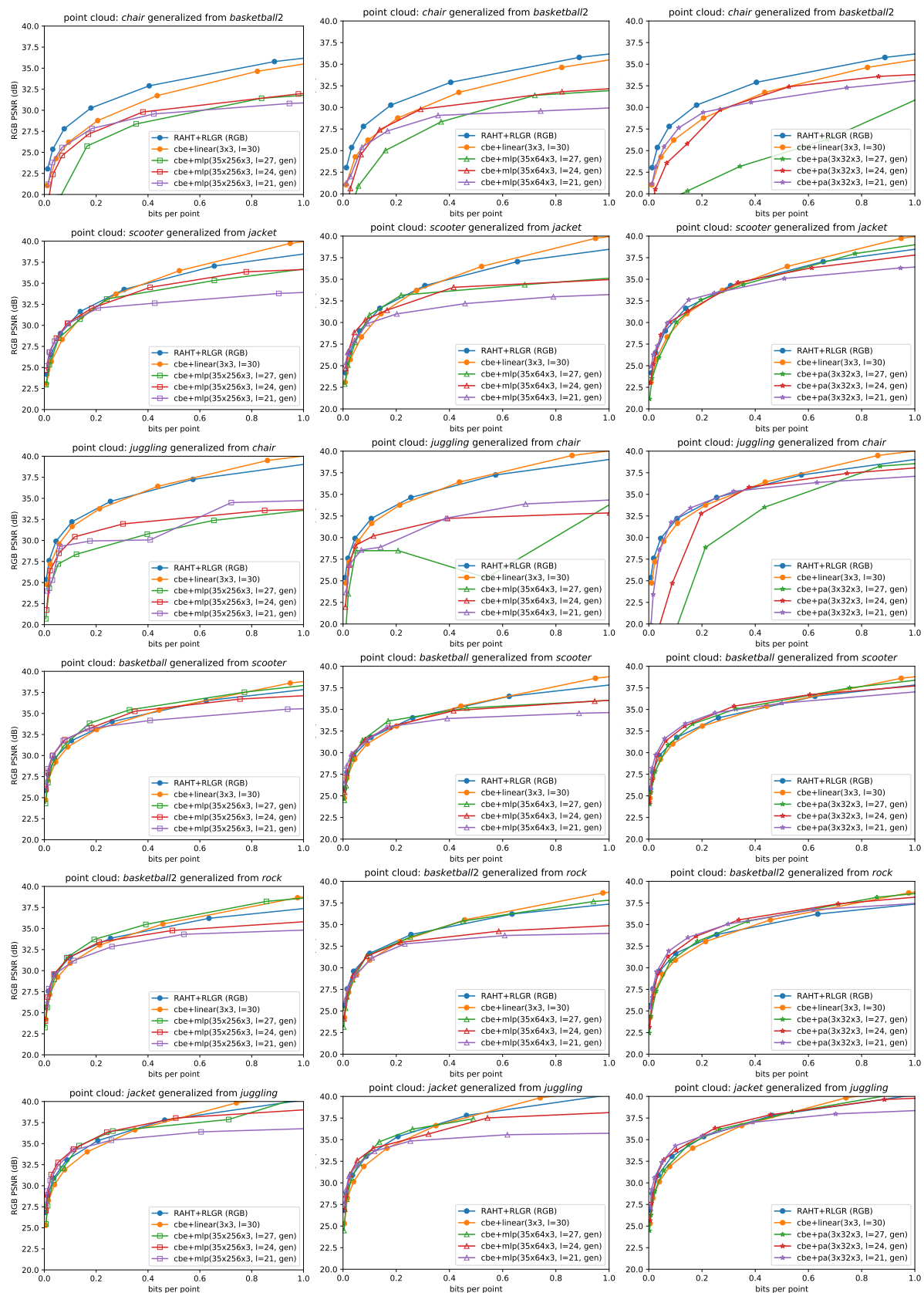


Figure S12: Coordinate Based Networks with generalization, by network. Each row is a different point cloud. Left, middle, right columns each show levels 27, 24, and 21, along with baselines, for CBNs $mlp(35 \times 256 \times 3)$, $mlp(35 \times 64 \times 3)$, and $pa(3 \times 32 \times 3)$.

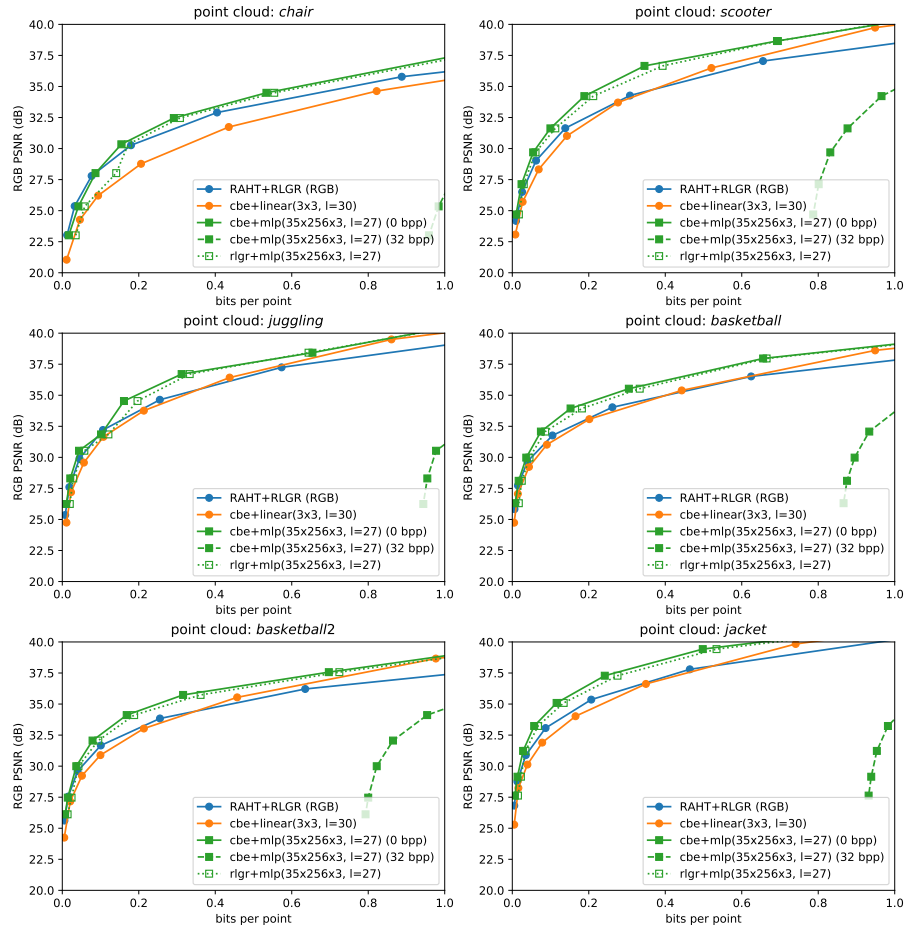


Figure S13: Side information for entropy model. Sending 32 bits per parameter for the *cbe* entropy model would reduce RD performance from solid to dashed green lines. But the backward-adaptive *RLGR* entropy coder (dotted, unfilled) obviates the need to send side information with almost no loss in performance.

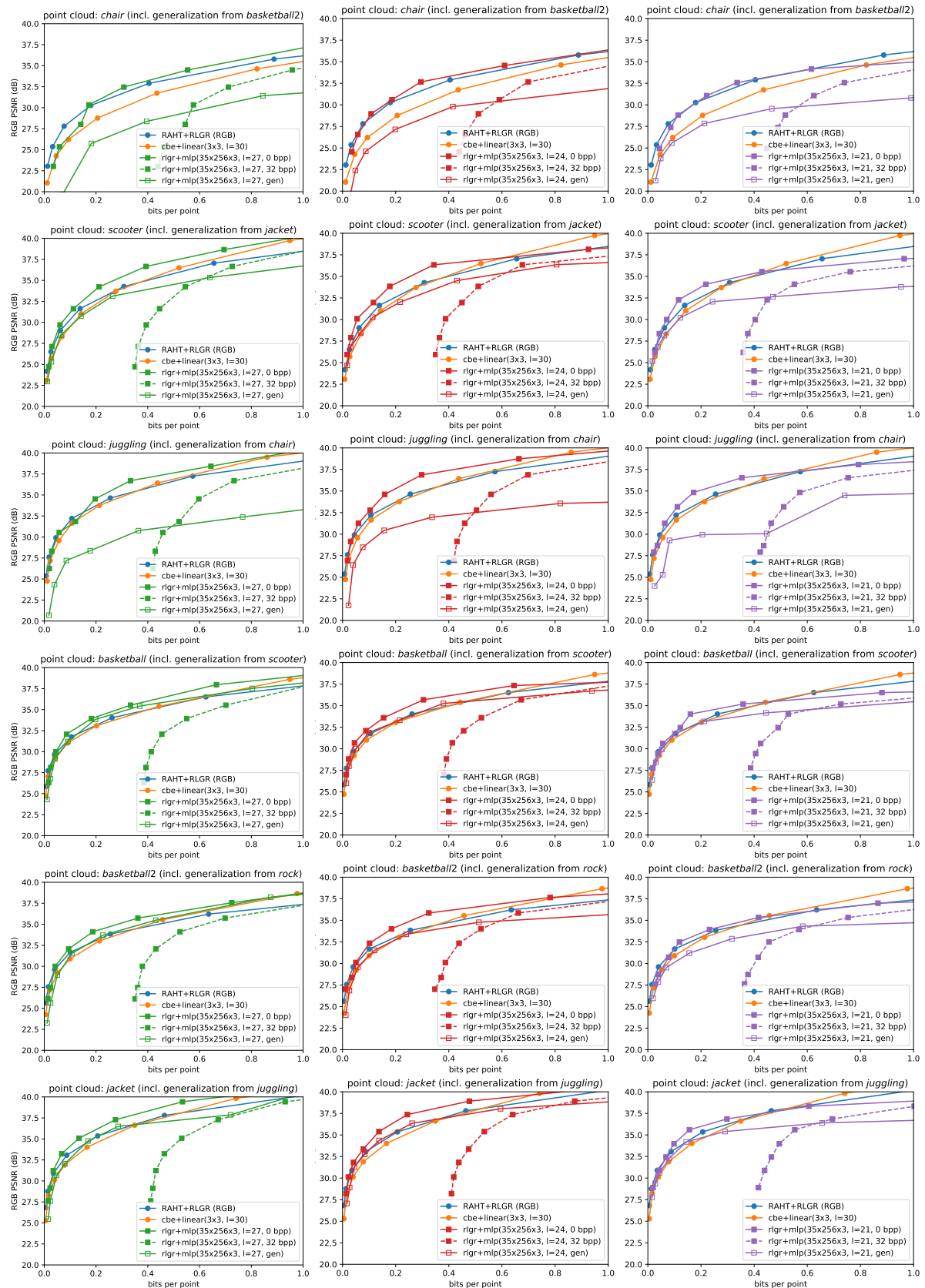


Figure S14: Effect of side information for coordinate based network $mlp(35 \times 256 \times 3)$ at levels 27 (left), 24 (middle), and 21 (right). Each row is a different point cloud.

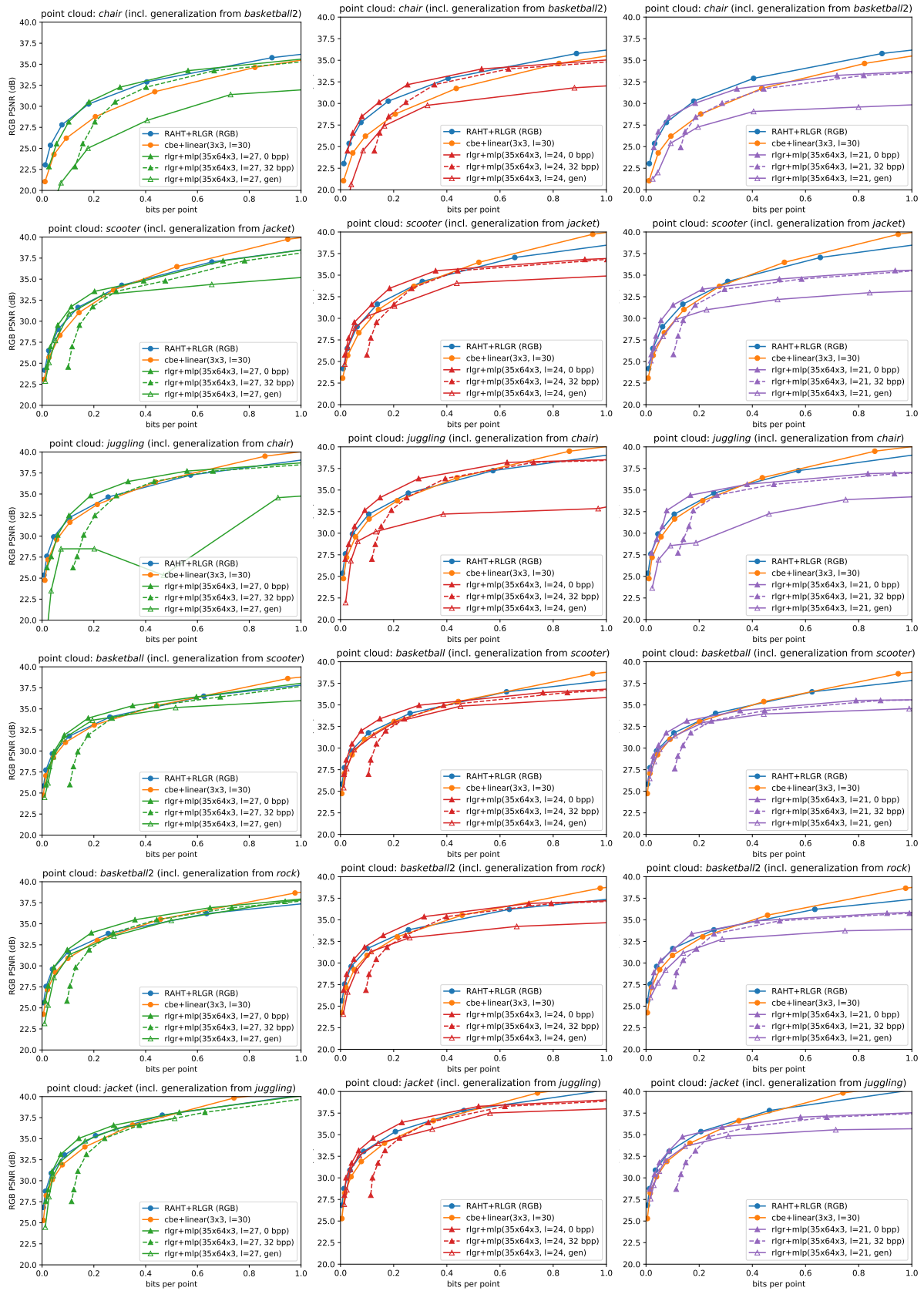


Figure S15: Effect of side information for coordinate based network $mlp(35 \times 64 \times 3)$ at levels 27 (left), 24 (middle), and 21 (right). Each row is a different point cloud.

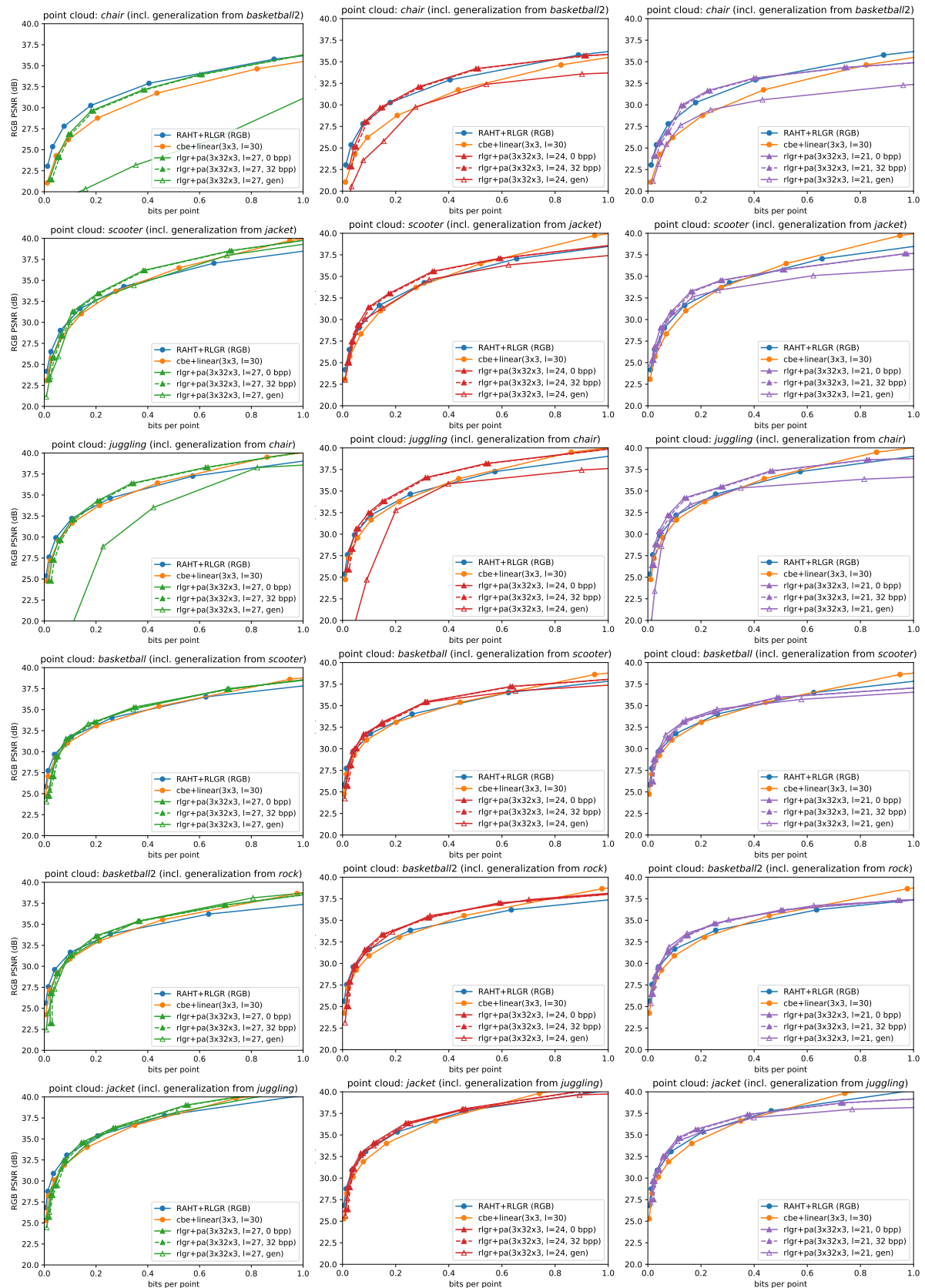


Figure S16: Effect of side information for coordinate based network $pa(3 \times 32 \times 3)$ at levels 27 (left), 24 (middle), and 21 (right). Each row is a different point cloud.

Table S3. BD-Rate reductions due to normalization.

<i>rock</i> CBN	level			
	30	27	24	21
linear(3x3)	-36.5%	-26.1%	-29.6%	-35.2%
mlp(35x256x3)	N/A	-35.8%	-39.4%	-28.7%
mlp(35x64x3)	N/A	-32.7%	-34.4%	-27.8%
pa(3x32x3)	N/A	-41.6%	-39.8%	-37.5%
<i>chair</i> CBN	level			
	30	27	24	21
linear(3x3)	-36.8%	-26.5%	-34.1%	-47.4%
mlp(35x256x3)	N/A	-34.8%	-39.5%	-28.0%
mlp(35x64x3)	N/A	-34.3%	-47.5%	-38.4%
pa(3x32x3)	N/A	-45.2%	-44.7%	-45.6%
<i>scooter</i> CBN	level			
	30	27	24	21
linear(3x3)	-32.6%	-19.5%	-34.6%	-41.8%
mlp(35x256x3)	N/A	-32.1%	-38.1%	-20.6%
mlp(35x64x3)	N/A	-27.1%	-30.2%	-39.1%
pa(3x32x3)	N/A	-41.9%	-42.2%	-40.6%
<i>juggling</i> CBN	level			
	30	27	24	21
linear(3x3)	-26.9%	-8.5%	-23.4%	-32.4%
mlp(35x256x3)	N/A	-22.7%	-35.0%	-27.7%
mlp(35x64x3)	N/A	-16.3%	-35.2%	-30.5%
pa(3x32x3)	N/A	-47.7%	-47.6%	-39.1%
<i>basketball</i> CBN	level			
	30	27	24	21
linear(3x3)	-30.0%	-19.7%	-27.7%	-28.9%
mlp(35x256x3)	N/A	-20.9%	-29.2%	-42.3%
mlp(35x64x3)	N/A	-5.8%	-28.4%	-32.8%
pa(3x32x3)	N/A	-31.1%	-27.7%	-24.5%
<i>basketball2</i> CBN	level			
	30	27	24	21
linear(3x3)	-29.2%	-14.6%	-20.8%	-36.5%
mlp(35x256x3)	N/A	-34.5%	-24.0%	-15.9%
mlp(35x64x3)	N/A	-28.7%	-27.8%	-24.1%
pa(3x32x3)	N/A	-41.2%	-42.5%	-41.7%
<i>jacket</i> CBN	level			
	30	27	24	21
linear(3x3)	-29.3%	-15.2%	-27.7%	-41.9%
mlp(35x256x3)	N/A	-28.8%	-35.1%	-28.4%
mlp(35x64x3)	N/A	-22.0%	-21.3%	-25.1%
pa(3x32x3)	N/A	-39.1%	-38.0%	-41.7%

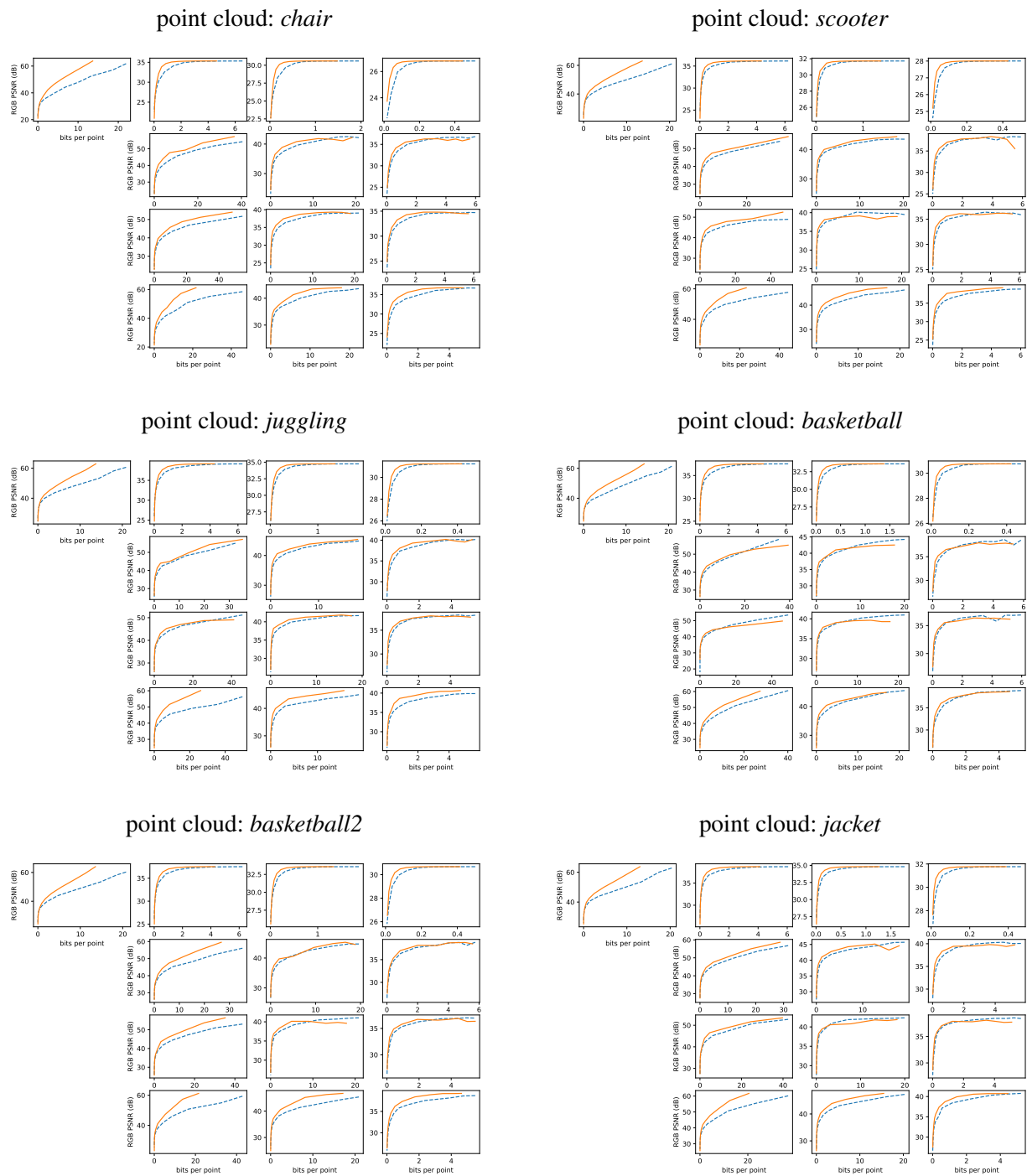


Figure S17: RD performance (RGB PSNR vs. bit rate) improvement due to normalization, corresponding to entries in Tab. S3.

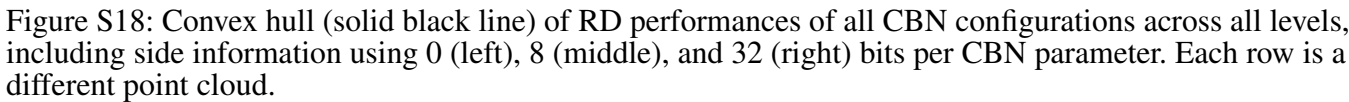


Table S4. Comparison to Sparse PCAC (Wang and Ma (2022)). Sparse PCAC (left) and LVAC (right) are each compared to references TMC13v14 and RAHT, in terms of BD-BD (average % bitrate change over YUV PSNR range) and BD-PSNR (average YUV PSNR dB change over bitrate range) relative to the reference. Negative/positive BD-BR means better/worse than the reference. Positive/negative BD-PSNR means better/worse than the reference. The references are the best-performing versions: TMC13v14 (YUV, predRAHT) and RAHT-RLGR (YUV). Both Sparse PCAC and LVAC are better than RAHT, not as good as TMC13v14, and about as good as each other. Sparse PCAC data taken from (Wang and Ma, 2022, Table 1). LVAC data computed from convex hulls of LVAC curves in Fig. S8 and Fig. S10. Bold values indicate which one of LVAC and Sparse PCAC performs better.

Point Cloud	Sparse PCAC (Wang and Ma (2022))				LVAC (ours)			
	vs. TMC13v14		vs. RAHT		vs. TMC13v14		vs. RAHT	
	BD-BR	BD-PSNR	BD-BR	BD-PSNR	BD-BR	BD-PSNR	BD-BR	BD-PSNR
longdress	+27 %	-0.81 dB	-42 %	+1.72 dB	+54 %	-1.25 dB	-33 %	+1.00 dB
loot	+142 %	-2.50 dB	+6 %	-0.10 dB	+61 %	-1.75 dB	-30 %	+0.31 dB
red&black	+55 %	-1.28 dB	-36 %	+1.30 dB	+41 %	-0.98 dB	-41 %	+1.40 dB
soldier	+66 %	-1.60 dB	-25 %	+0.75 dB	+59 %	-1.17 dB	-29 %	+1.27 dB
Average	+72 %	-1.55 dB	-24 %	+0.92 dB	+54 %	-1.29 dB	-33 %	+1.00 dB

REFERENCES

- 37 Alliez, P., Forge, F., De Luca, L., Pierrot-Deseilligny, M., and Preda, M. (2017). Culture 3D Cloud:
38 A Cloud Computing Platform for 3D Scanning, Documentation, Preservation and Dissemination of
39 Cultural Heritage. *ERCIM News* , 64
- 40 Cignoni, P., Callieri, M., Corsini, M., Dellepiane, M., Ganovelli, F., and Ranzuglia, G. (2008). MeshLab:
41 an Open-Source Mesh Processing Tool. In *Eurographics Italian Chapter Conference*, eds. V. Scarano,
42 R. D. Chiara, and U. Erra (The Eurographics Association). doi:10.2312/LocalChapterEvents/ItalChap/
43 ItalianChapConf2008/129-136
- 44 d'Eon, E., Harrison, B., Meyers, T., and Chou, P. A. (2017). *8i Voxelized Full Bodies — A Voxelized Point*
45 *Cloud Dataset*. input document M74006 & m42914, ISO/IEC JTC1/SC29 WG1 & WG11 (JPEG &
46 MPEG), Ljubljana, Slovenia
- 47 Guo, K., Lincoln, P., Davidson, P., Busch, J., Yu, X., Whalen, M., et al. (2019). The relightables:
48 Volumetric performance capture of humans with realistic relighting. *ACM Trans. Graph.* 38. doi:10.
49 1145/3355089.3356571
- 50 Malvar, H. (2006). Adaptive run-length/golomb-rice encoding of quantized generalized gaussian sources
51 with unknown statistics. In *Data Compression Conference (DCC'06)*. 23–32
- 52 Meka, A., Pandey, R., Haene, C., Orts-Escolano, S., Barnum, P., Davidson, P., et al. (2020). Deep
53 relightable textures - volumetric performance capture with neural rendering. vol. 39. doi:10.1145/
54 3414685.3417814
- 55 Sheng, X., Li, L., Liu, D., Xiong, Z., Li, Z., and Wu, F. (2021). Deep-pcac: An end-to-end deep lossy
56 compression framework for point cloud attributes. *IEEE Transactions on Multimedia* , 1–1doi:10.1109/
57 TMM.2021.3086711
- 58 Wang, J. and Ma, Z. (2022). Sparse tensor-based point cloud attribute compression. In *2022 IEEE Int'l*
59 *Conf. Multimedia Information Processing and Retrieval (MIPR)*