



Neural Rendering with Heterogeneous Scene Primitives

Chong Zeng¹ Yue Dong² Pieter Peers³ Lvmin Zhang¹ Maneesh Agrawala¹

¹Stanford University ²Microsoft Research ³College of William & Mary



primitives in one forward pass. RenderFormer-V1 stops at 4K

output pixels, without upsampling stage

PSNR increase over RenderFormer-V1 at 128K triangles, with 4.8× lower LPIPS

faster than RenderFormer-V1 and 4.7× faster than Cycles

- per-scene training, fine-tuning or hand-written transport code

A scene is a token sequence.

-
- The diagram illustrates the proposed framework, divided into three main stages:
- Scene Encoding:** Inputs include Triangle & Texture, Volume, Triangle Lighting, and Environment Map. These are processed through Triangle Embedding, Volume Embedding, Lighting Embedding, and Env Map Embedding respectively, resulting in a sequence of tokens.
 - View-Independent Stage:** The tokens are processed by Spatial Sorting (Hilbert Curve) and then a Global Attention Sink. The Global Attention Sink uses Register Tokens and Summarization Tokens, which are processed by Mean Pooling. A Local Sliding Window is also used for processing tokens. The output is a sequence of tokens.
 - View-Dependent Stage:** The tokens are processed by Sparse Self-Attention. The output is then used for Ray Embedding, Ray-Scene Cross-Attention, Swin Self-Attention, and finally the DPT Decoder to produce the output image.

One model, many light transport effects.

-
- Refraction
- Environment lighting
- Volume, participating media
- Textures, SIBRDFs
- Displacement mapping
- Complex interior scene

■ Triangle	3 vertices + 3 normals; a 9-D material code; a 32×32, 13-channel patch (9 material, 3 normal, 1 displacement) through a frozen VAE
■ Voxel medium	rotation + per-axis scale; a 4 ³ grid of RGB scattering, RGB absorption and phase anisotropy
■ Area light	3 vertices + 3 normals; RGB emittance, in its own token rather than fused into geometry
■ Environment map	LDR and log-HDR 512×256, each through a frozen VAE into an 8×4 token sequence; direction map; log-max intensity
■ Camera bundle	8×8 ray directions, in world space

A single positional encoding covers all of them: RoPE on the primitive centroid, with position-less tokens falling back to the scene centroid, so the whole embedding is translation invariant.

Data-driven reflectance model.

-
- Diagram illustrating the 9-D appearance code architecture:
- Top Path:** A sphere under a fixed light probe is processed by a **CNN encoder** to produce a **9-D appearance code**. The code is represented as $\mathbf{z}_{9d} \in [-1, 1]^9$, which is added straight into the primitive token.
 - Bottom Path:** Analytical BRDF parameters (Principled, GGX, transparent) are processed by an **MLP**. The MLP output is fitted afterwards, skipping the render.

Because the space is defined by rendered appearance, any material you can render is usable, including measured BRDFs never seen in training.

What is the attention sink?

-
- Diagram illustrating the sink's role in a distributed system:
- sink: always attended** (indicated by a bracket on the left)
 - ...and attends to everything** (indicated by a bracket on the right)
 - never computed** (indicated by a large gray area)
 - ±256 keys** (indicated by a diagonal band of blue area)
 - never computed** (indicated by a large gray area below the diagonal band)
 - geometry keys, in Hilbert order** (indicated by a horizontal axis at the bottom)
- Legend:
- 16 registers (gray)
 - ≤8 lights (yellow)
 - 1 summary per 64 primitives (orange)
- Complexity: $O(T^2) \rightarrow O(T(512 + T/32))$

Every row is one query, every column one key; sink and window widths are exaggerated for legibility.

Sorting the primitives along a Hilbert curve makes neighbours in space localized in the sequence, so a 1-D sliding window buys a spatial one.

16 registers

Free scratch space for whatever is globally true of the scene.

Each primitive sees each light whatever the distance, which is the attention analogue of importance sampling the lights.

Summaries

Mean-pooling every 64 geometry tokens keeps long-range occlusion in reach, coarsely: distant detail does not matter.

What does each part contribute?

- ### ABLATION OF THE SPARSE ATTENTION

Attention variant	PSNR $\uparrow$	SSIM $\uparrow$	LPIPS $\downarrow$	HDR-FLIP $\downarrow$
RenderFormer-V2	28.25	.8982	.0997	.4200
w/o attention sink	27.36	.8841	.1247	.4359
w/o sliding window	26.92	.8715	.1289	.4517
sink w/o lights	27.40	.8813	.1081	.4250
sink w/o summaries	27.09	.8754	.1210	.4485
sink w/o both	26.13	.8474	.1514	.4959
summaries 1 : 32	27.78	.8923	.1072	.4231
summaries 1 : 128	26.48	.8690	.1346	.4655
summaries 1 : 256	26.93	.8663	.1316	.4580
full attention	27.91	.8953	.1027	.4287

The graph plots the number of primitives (Y-axis, 0 to 0.4) against resolution (X-axis: 4K, 16K, 64K, 128K). RenderFormer-V1 (red line) shows a sharp increase in primitives as resolution increases, while RenderFormer-V2 (blue line) remains relatively stable.

Resolution	RenderFormer-V1 (Primitives)	RenderFormer-V2 (Primitives)
4K	~0.02	~0.01
16K	~0.05	~0.02
64K	~0.28	~0.04
128K	~0.48	~0.05

Primitives	Cycles (seconds/frame)	RenderFormer-V1 (seconds/frame)	RenderFormer-V2 (seconds/frame)
1.6K	~3.7	~0.1	~0.1
6.4K	~4.1	~0.1	~0.1
22K	~4.4	~0.2	~0.2
128K	~5.3	~4.0	~1.2