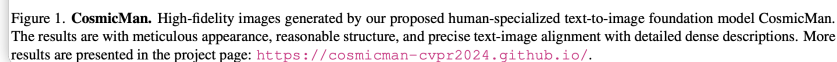


- **CosmicMan: A Text-to-Image Foundation Model for Humans**
- <https://arxiv.org/abs/2404.01294>
- CVPR 2024



The diagram illustrates the Decomposed-Attention-Refocusing framework. It shows a sequence of input images $h_1, h_2, \dots, h_i$ being processed by a module $\mathcal{L}_{HOLA}$ to produce a sequence of masked images $m(s_1, e_1), m(s_2, e_2), \dots, m(s_i, e_i)$. A text description "A full-body shot, a adult woman ... with brown hair ... short white blouse and ..." is used to generate a set of bounding boxes c_{body} , c_{outfit} , etc. These are used to crop the input images into z_t and z_{t-1} . These crops are then processed by a module with Q and K_i components to produce the final masked images. A noise loss $\mathcal{L}_{noise}$ is also shown.

caption mask caption
caption HOLA

$$\mathcal{L}_{\text{HOLA}} = \frac{1}{N} \sum_{i=1}^N \left(\sum_{j=s_i}^{e_i} \|m_j - h_i\|_2^2 + \left\| \frac{1}{e_i - s_i} \sum_{j=s_i}^{e_i} (m_j) - h_i \right\|_2^2 \right)$$

- CosmicMan-HQ 1.0
- FID CLIP User Study
- 32 A100/bs64
- <https://cosmicman-cvpr2024.github.io>

Older

PreciseControl

FlashFace

Made with [Montaigne](#) and [bigmission](#)