



- home
- papers
- study
- life
- About me

Break-A-Scene

- Break-A-Scene: Extracting Multiple Concepts from a Single Image
- <https://arxiv.org/abs/2305.16311>
- SIGGRAPH Asia 2023

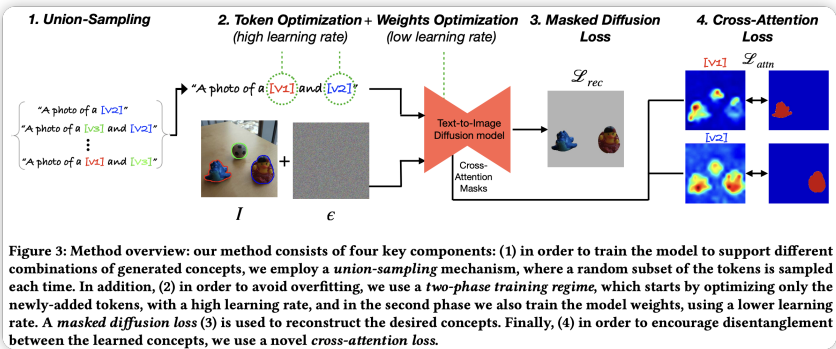
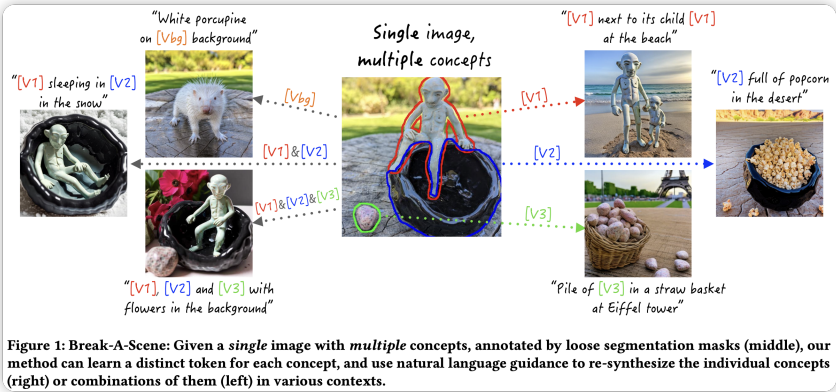


Figure 3: Method overview: our method consists of four key components: (1) in order to train the model to support different combinations of generated concepts, we employ a *union-sampling* mechanism, where a random subset of the tokens is sampled each time. In addition, (2) in order to avoid overfitting, we use a *two-phase training regime*, which starts by optimizing only the newly-added tokens, with a high learning rate, and in the second phase we also train the model weights, using a lower learning rate. A *masked diffusion loss* (3) is used to reconstruct the desired concepts. Finally, (4) in order to encourage disentanglement between the learned concepts, we use a novel *cross-attention loss*.

token mask embedding mask diffusion loss token cross-attention maps loss union-sampling N 1 2 3 4 4



- test-tuning
- prompt similarity CLIP ID preservation DINO DreamBooth
-
- <https://github.com/google/break-a-scene>

Newer	Older
2024-08-29 Prompt-to-Prompt	2024-08-29 Face0