

- Attention Calibration for Disentangled Text-to-Image Personalization
- <https://arxiv.org/abs/2403.18551>
- CVPR 2024

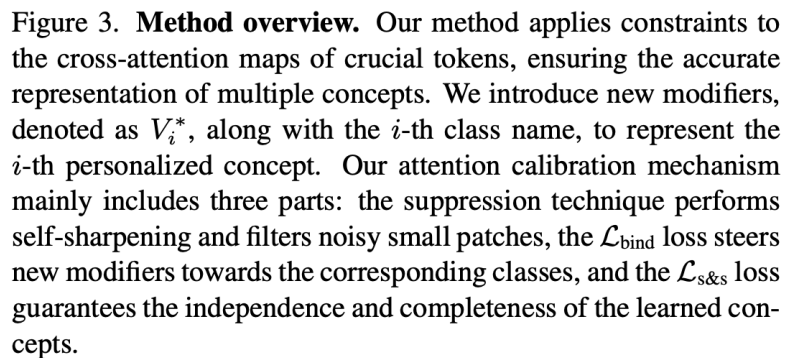


Figure 3. **Method overview.** Our method applies constraints to the cross-attention maps of crucial tokens, ensuring the accurate representation of multiple concepts. We introduce new modifiers, denoted as  $V_i^*$ , along with the  $i$ -th class name, to represent the  $i$ -th personalized concept. Our attention calibration mechanism mainly includes three parts: the suppression technique performs self-sharpening and filters noisy small patches, the  $\mathcal{L}_{\text{bind}}$  loss steers new modifiers towards the corresponding classes, and the  $\mathcal{L}_{\text{s\&s}}$  loss guarantees the independence and completeness of the learned concepts.

Single Input

Inference Combined Concepts  
 $V_1$  man and  $V_2$  woman

Inference Independent Concepts  
 $V_1$  man and  $V_2$  woman

- □□□□□□□□□□□□10□□□□□□

- [OpenCLIP](#)
- [OpenCLIP](#)
- <https://github.com/Monalissaa/DisenDiff>

[Newer](#)

[Older](#)

2024-12-03

PersonalizedResidual

2024-11-14

CrossInitialization

leicheng © 2022-2025

[Archive](#) [RSS feed](#) [GitHub](#) [Email](#) [QR Code](#)

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