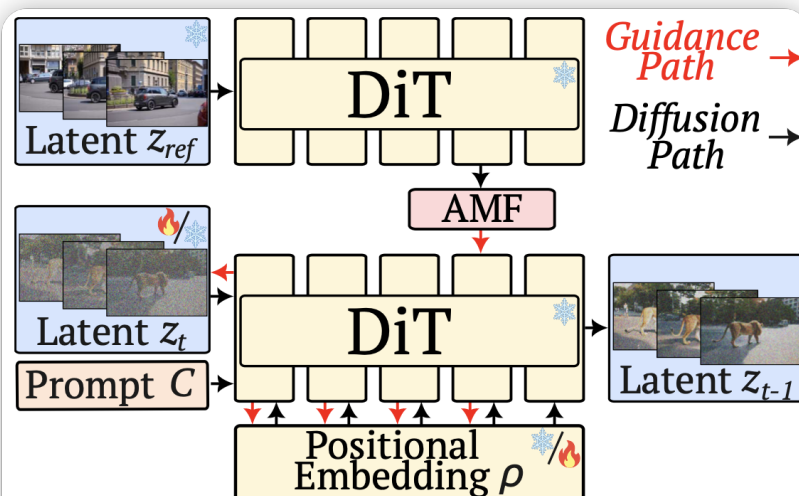


- Video Motion Transfer with Diffusion Transformers
- <https://arxiv.org/abs/2412.07776>
- CVPR 2025

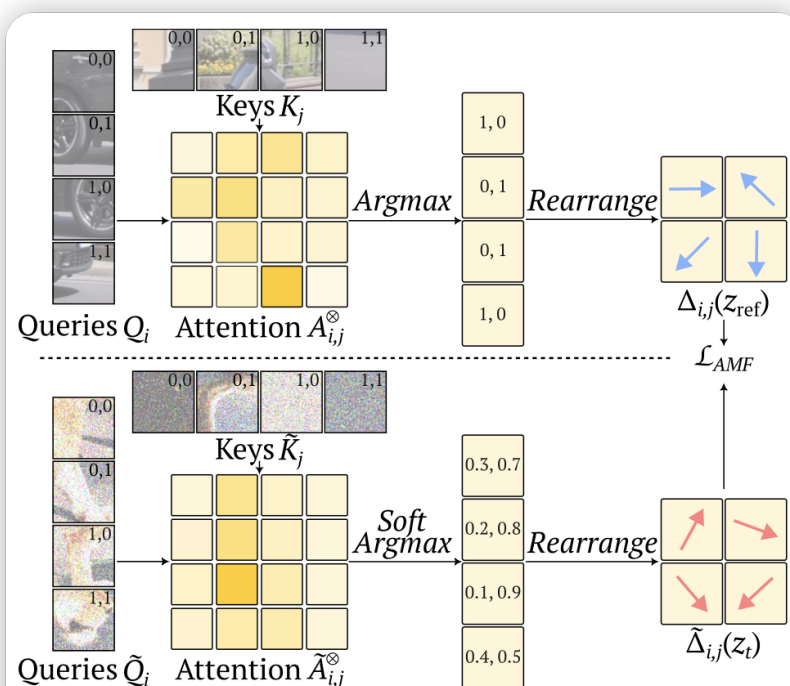


**Figure 2. Core idea of DiTFlow.** We extract the AMF from a reference video and we use that to guide the latent representation  $z_t$  towards the motion of the reference video. In our experiments, we also tested optimizing positional embeddings for improved zero-shot performance.

```

0000000000DiT00000000DiTFlow000000000000000000000000t=0000DiT0000
000000000000000000000000DDIM0000000000000000DiT_block00000000
attention0000000000patch000000000000000000000000AMF000000000000
00attention0000argmax00000000patch0000000000AMF00gt00000000000000
00000000000000AMF00attention0000000000000000000000zt00000000000000

```



**Figure 3. Guidance.** We compute the reference displacement by processing cross-frame attentions with an argmax operation and rearranging them into displacement maps, identifying patch-aware cross-frame relationships. For video synthesis, we do the same operation with a soft argmax to preserve gradients, and impose reconstruction with the reference displacement.

reconstruction with the reference displacement.

[illegible]

- CogVideoX
- 
- MFQCLIP
- 1 A40
- <https://ditflow.github.io>

Newer

Older

2025-06-05

AnalysisAttentionVDiT

2025-06-04

FreeTraj

leicheng © 2022-2025

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

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