

VACE

- 介绍VACE: All-in-One Video Creation and Editing
- 介绍<https://arxiv.org/abs/2503.07598>
- ICCV2025

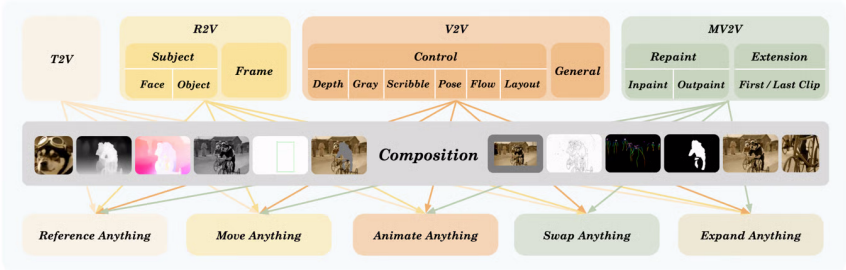


Figure 2. Task categories covered by VACE. Four basic tasks can be combined to create a vast number of possibilities.

video DiT T2V, R2V, V2V, MV2V Video Condition Unit [T, F, M] F M M

Table 1. The formal representation of frames (F s) and masks (M s) under the four basic tasks. Frames and masks are aligned spatially and temporally.

Tasks	Frames (F s) & Masks (M s)
T2V	$F = \{0_{h \times w}\} \times n$ $M = \{1_{h \times w}\} \times n$
R2V	$F = \{r_1, r_2, \dots, r_l\} + \{0_{h \times w}\} \times n$ $M = \{0_{h \times w}\} \times l + \{1_{h \times w}\} \times n$
V2V	$F = \{u_1, u_2, \dots, u_n\}$ $M = \{1_{h \times w}\} \times n$
MV2V	$F = \{u_1, u_2, \dots, u_n\}$ $M = \{m_1, m_2, \dots, m_n\}$

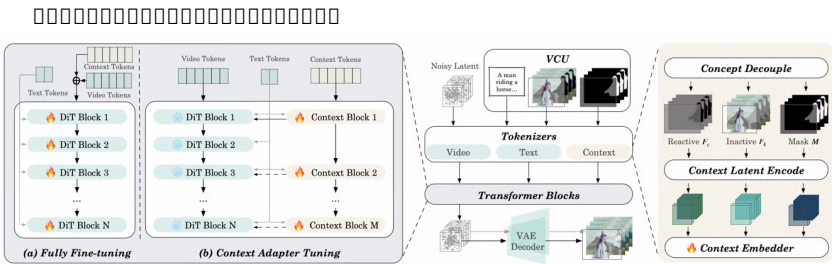


Figure 3. Overview of VACE Framework. Frames and masks are tokenized through Concept Decoupling, Context Latent Encode and Context Embedder. To achieve training with VCU as input, we employ two strategies, (a) Fully Fine-tuning and (b) Context Adapter Tuning. The latter converges faster and supports pluggable features.

VCU tokenizer embedder Context token (a) context token video token DiT (b) ControlNet block Context token block R2V Video Token ref image token

leicheng © 2022-2025

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

Made with [Montaigne](#) and by [anton](#) 