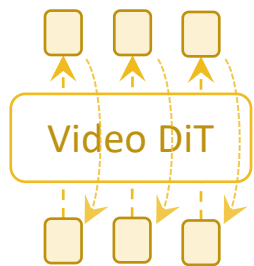


General Video Generation



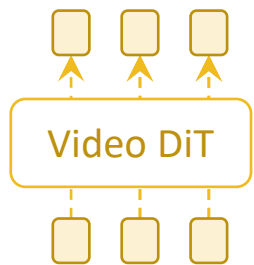
Denoising steps (τ)

1-step
($\tau = 1$)

2-step
($\tau = 2$)

4-step
($\tau = 4$)

DIDO (Ours)



1-step
($\tau = 1$)



Video Prediction

Future frame

$t + 1$ $t + 2$ $t + 3$ $t + 4$ $t + 5$



Dynamic content emerges progressively during denoising, with interaction-region motion sharpening step by step.



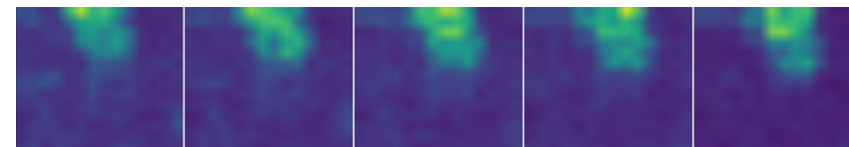
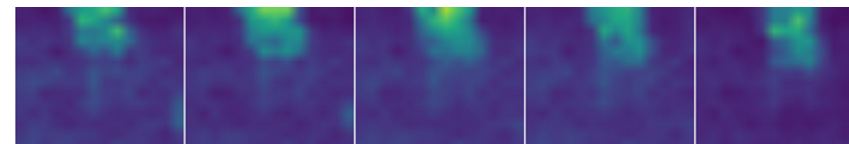
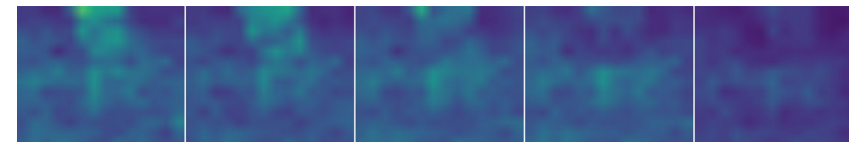
This motivates recovering interaction-centric dynamics with just one denoising step.



Latent Dynamics

Future frame

$t + 1$ $t + 2$ $t + 3$ $t + 4$ $t + 5$



Dynamics Change

Low

High