

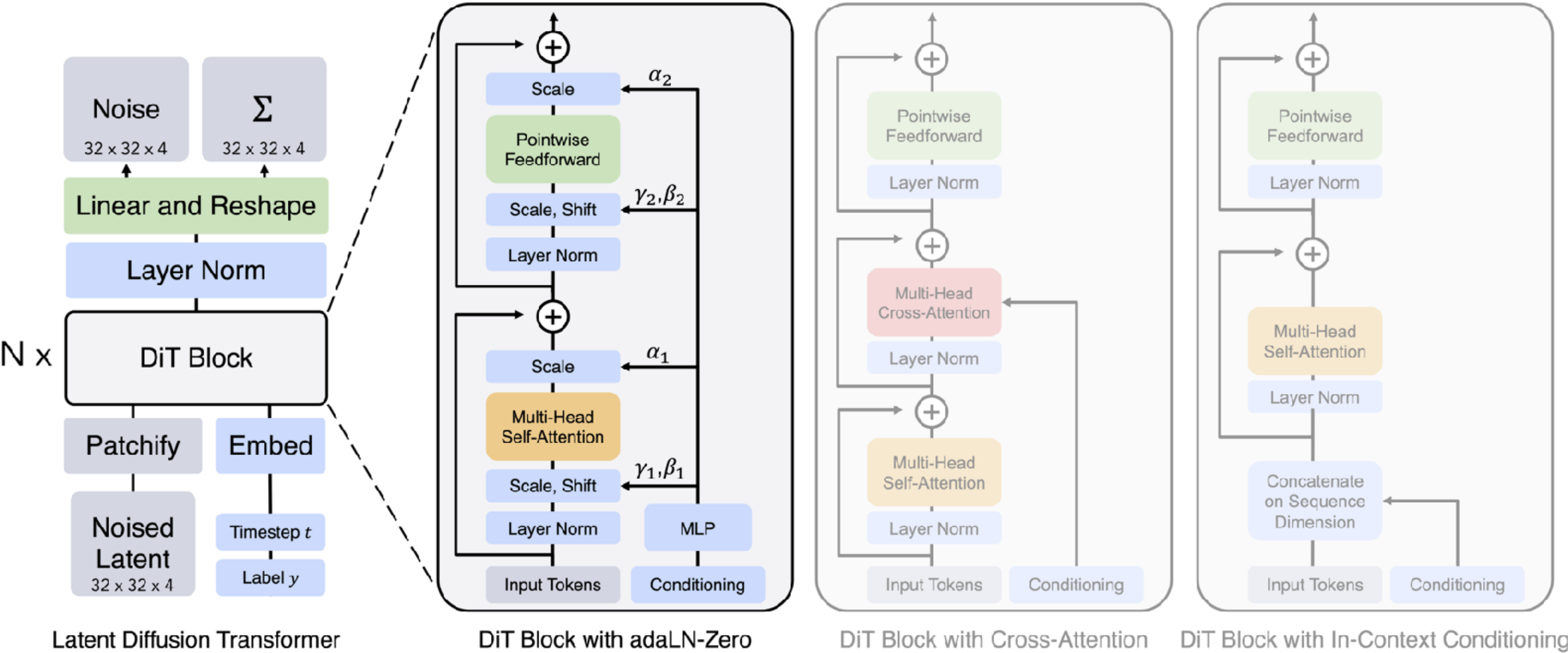
Streaming Video Generation as a Memory Problem

From Contextual to Parametric Memory

MA SHUAI

Diffusion based video generation

Limited generation duration

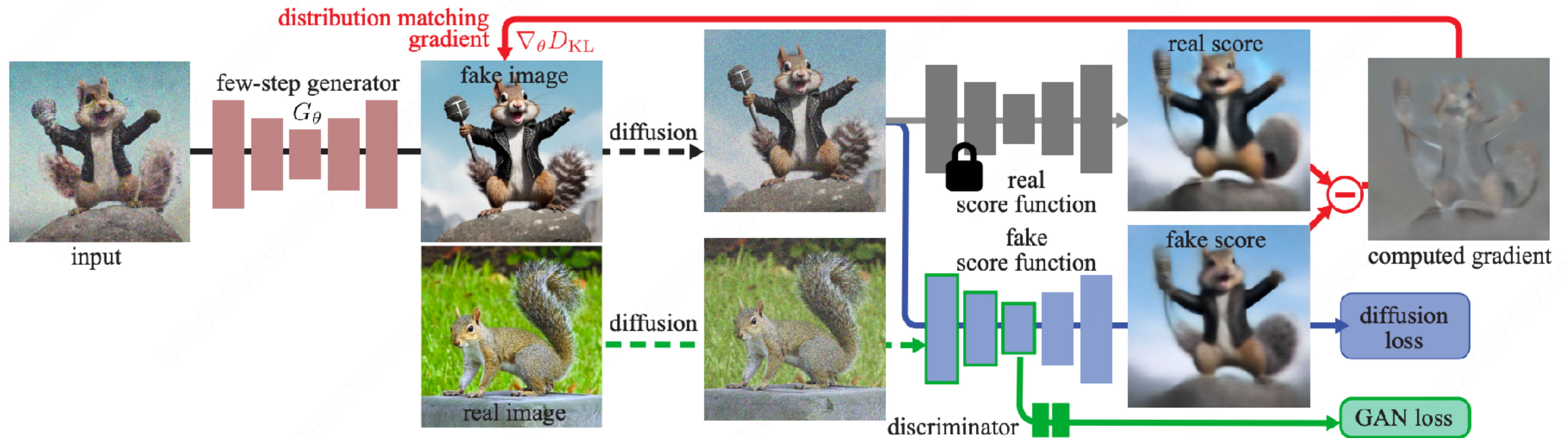


From Short Clips to Causal Rollout

AR based video generation

DMD

Distribution Matching Distillation



One-step Diffusion with Distribution Matching Distillation

Improved Distribution Matching Distillation for Fast Image Synthesis

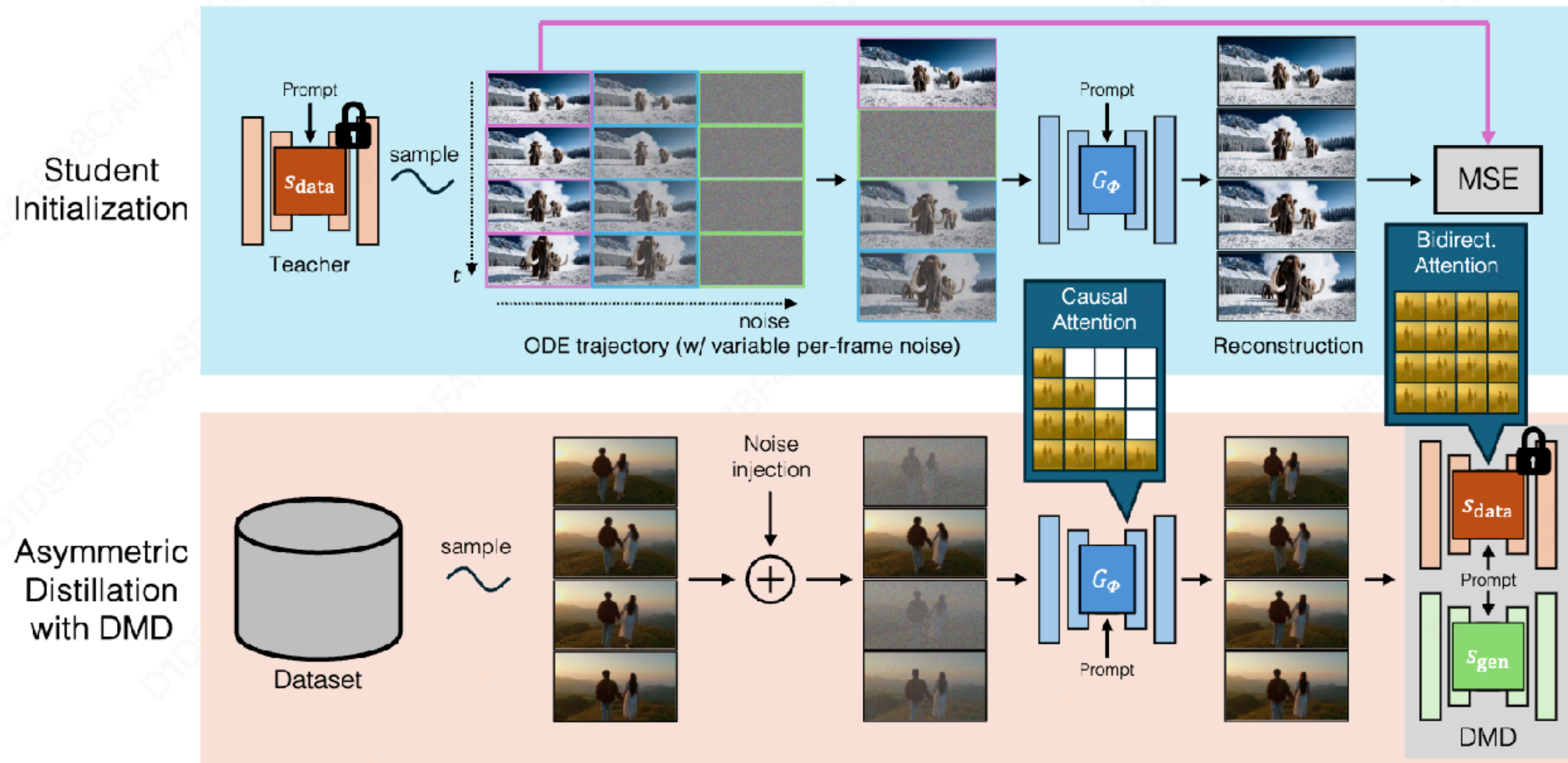
CVPR 2024

NeurIPS 2024 **Oral**

AR based video generation

CausVid

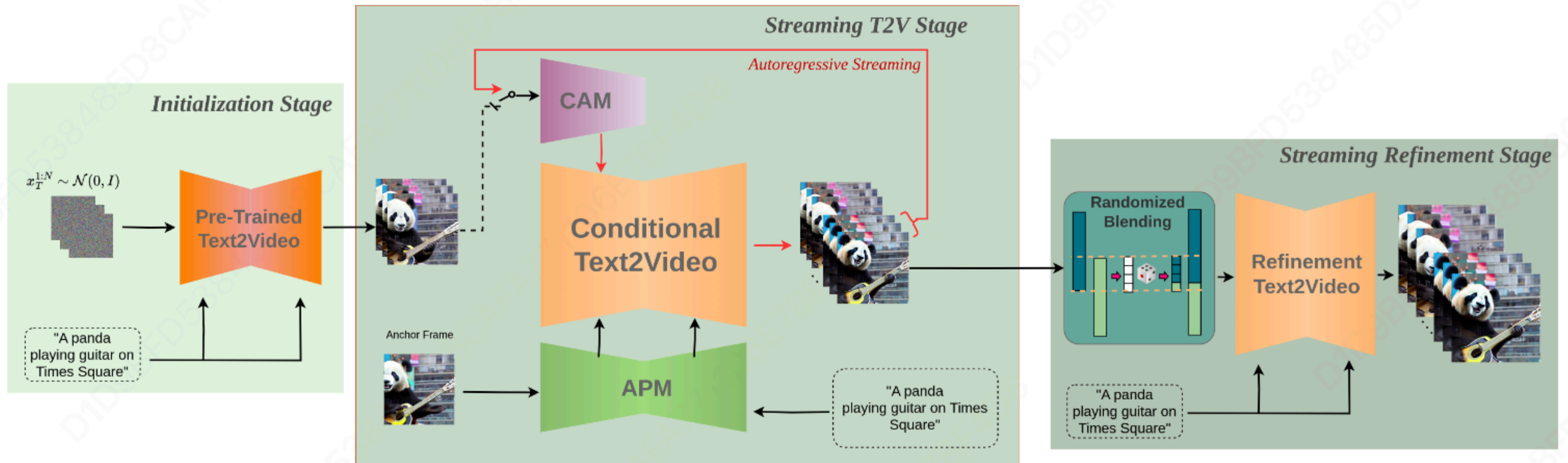
Distill from Bidirectional to Causal



History frames as context

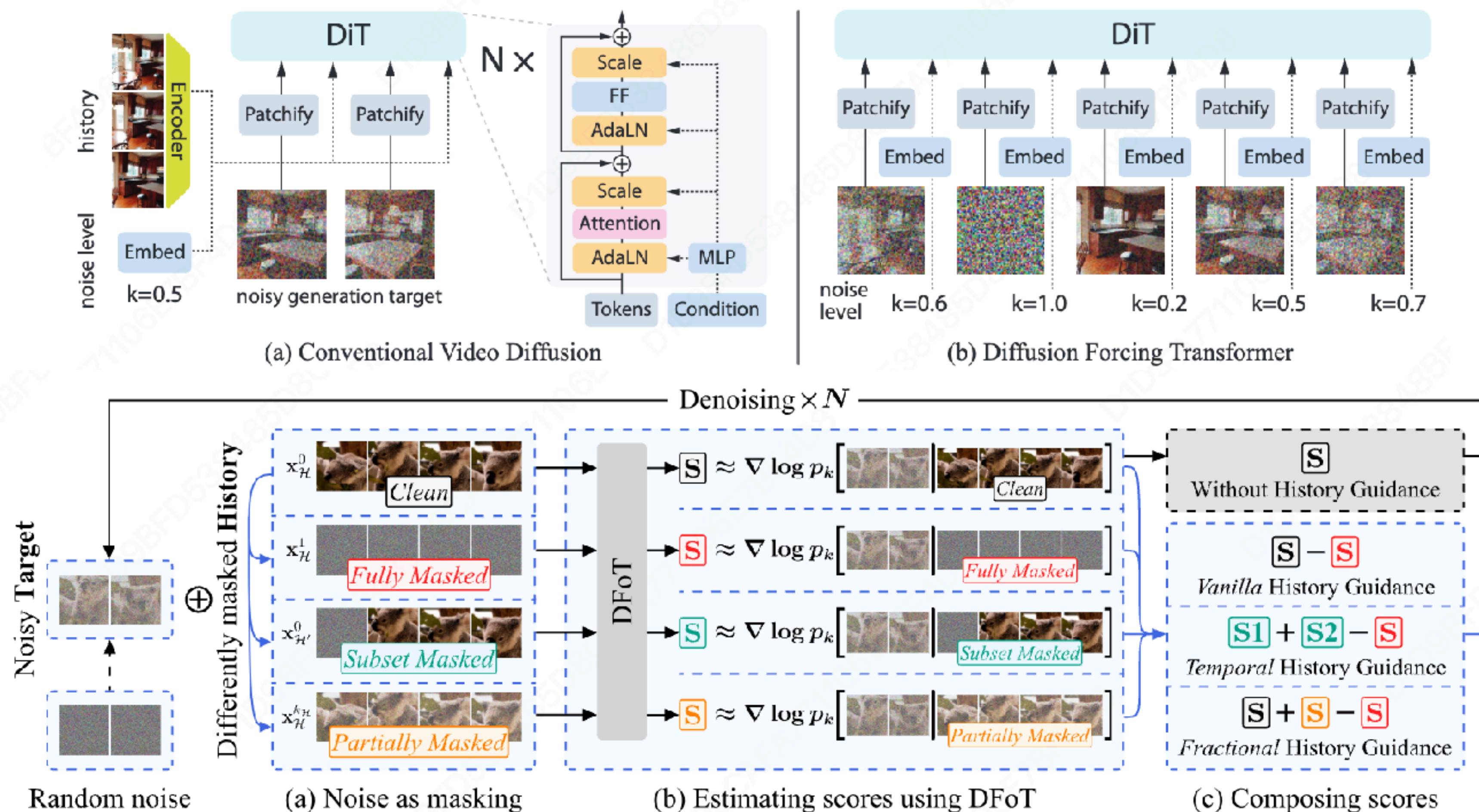
Streaming T2V

Dual-timescale Memory



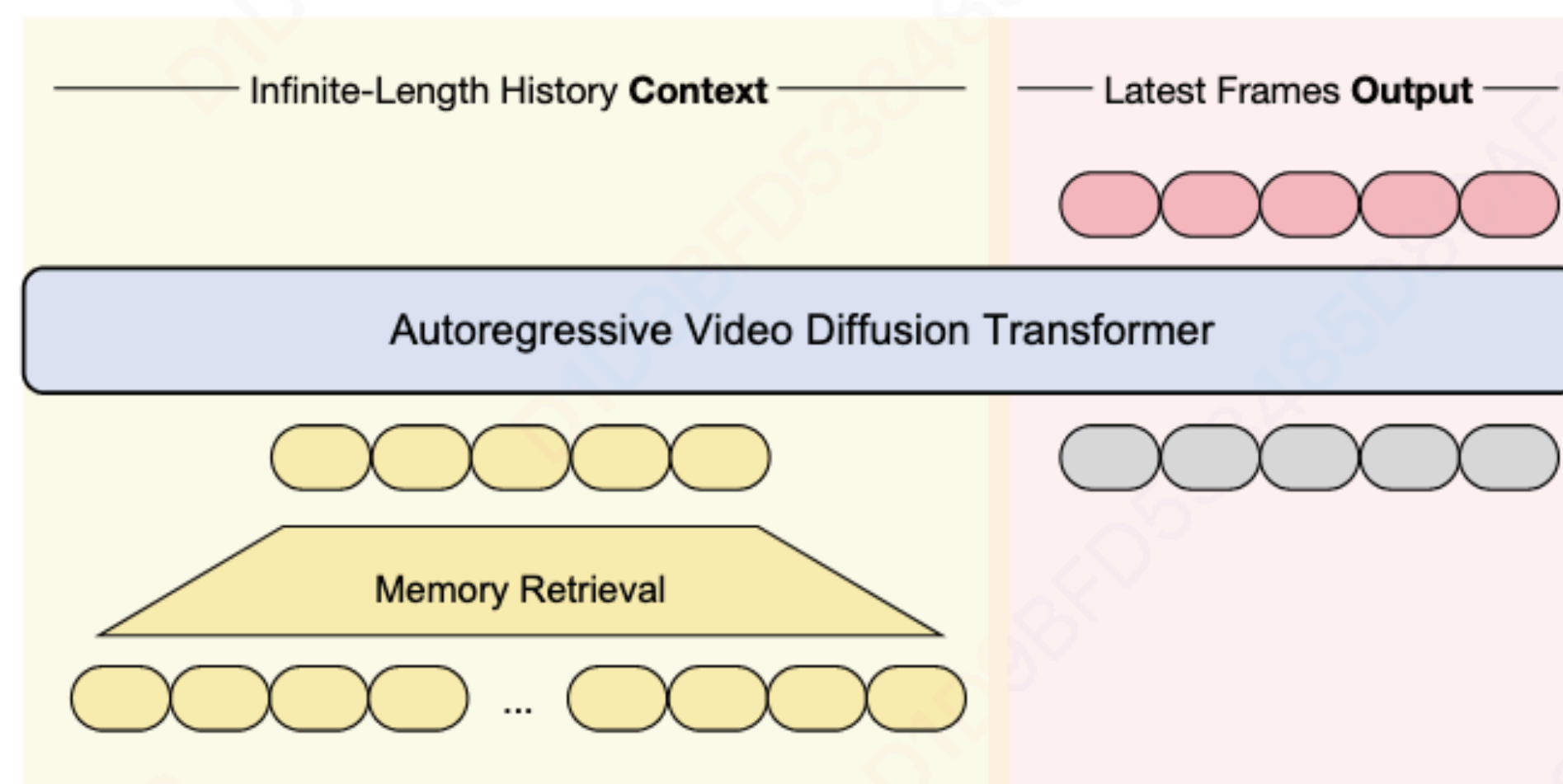
History frames as context

History guided video diffusion



History frames as context

Context as memory



(a) **Context learning:** long video generation with Memory Retrieval



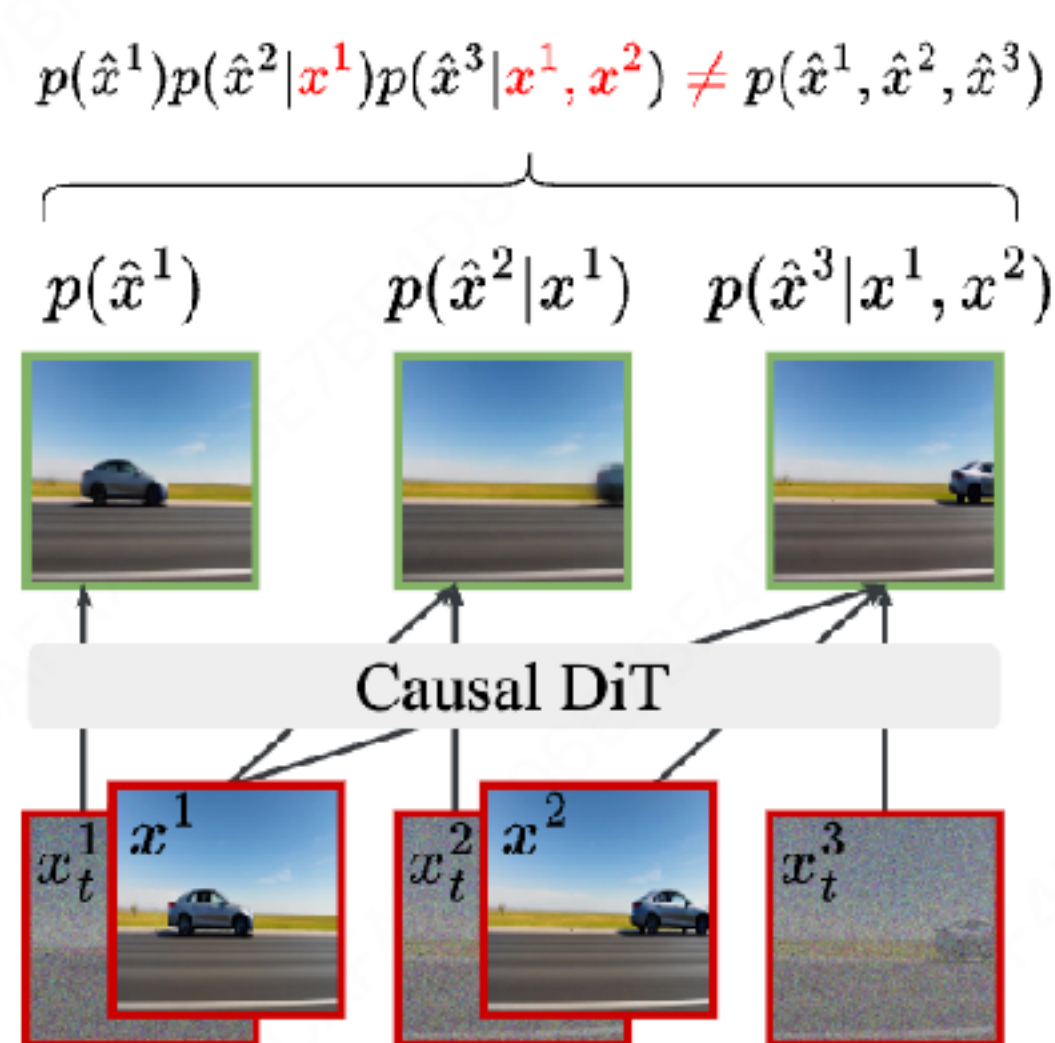
(b) **Memory Retrieval:** select **High-overlap Context** to condition the generation of **Predicted Frame**.

Towards Stable Rollout

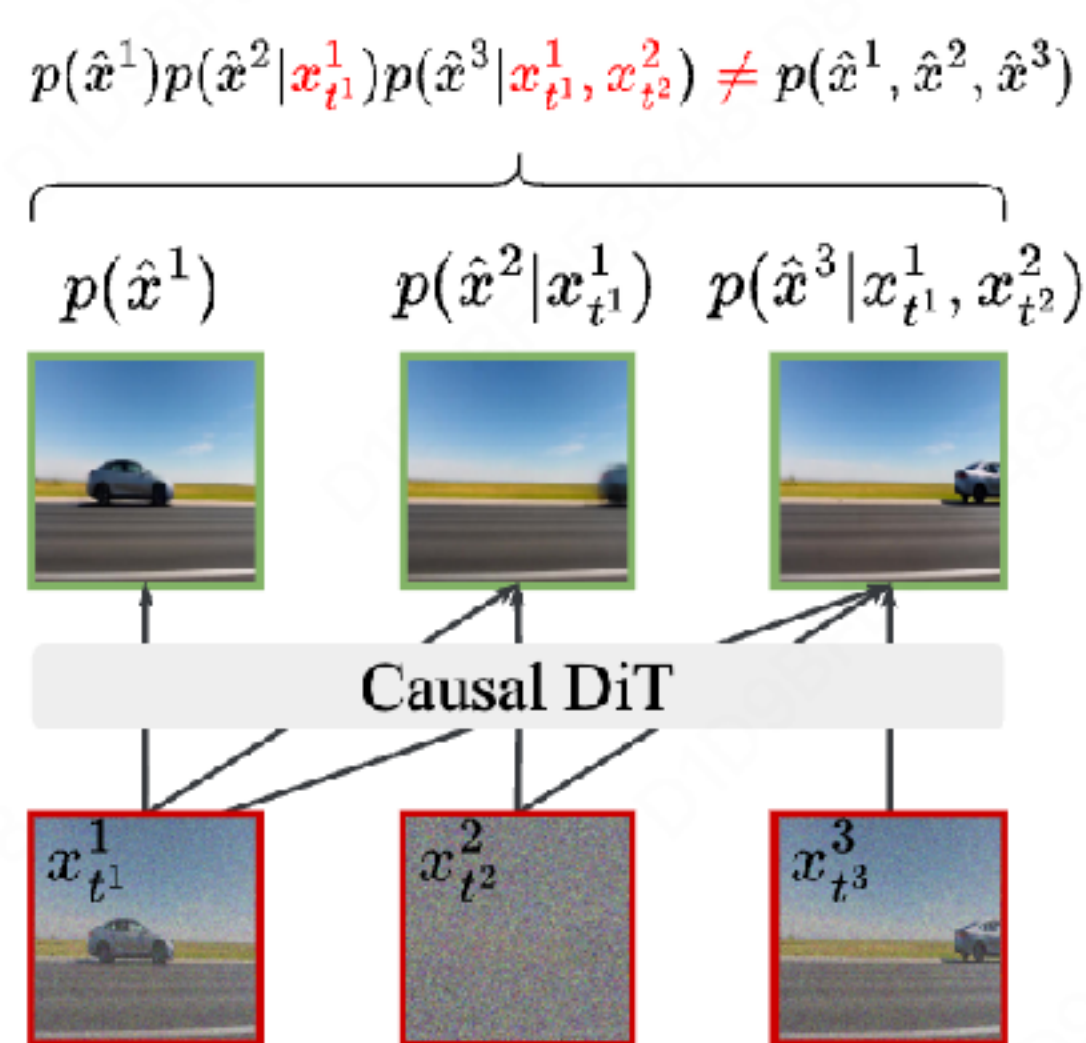
Stable Rollout

Self forcing

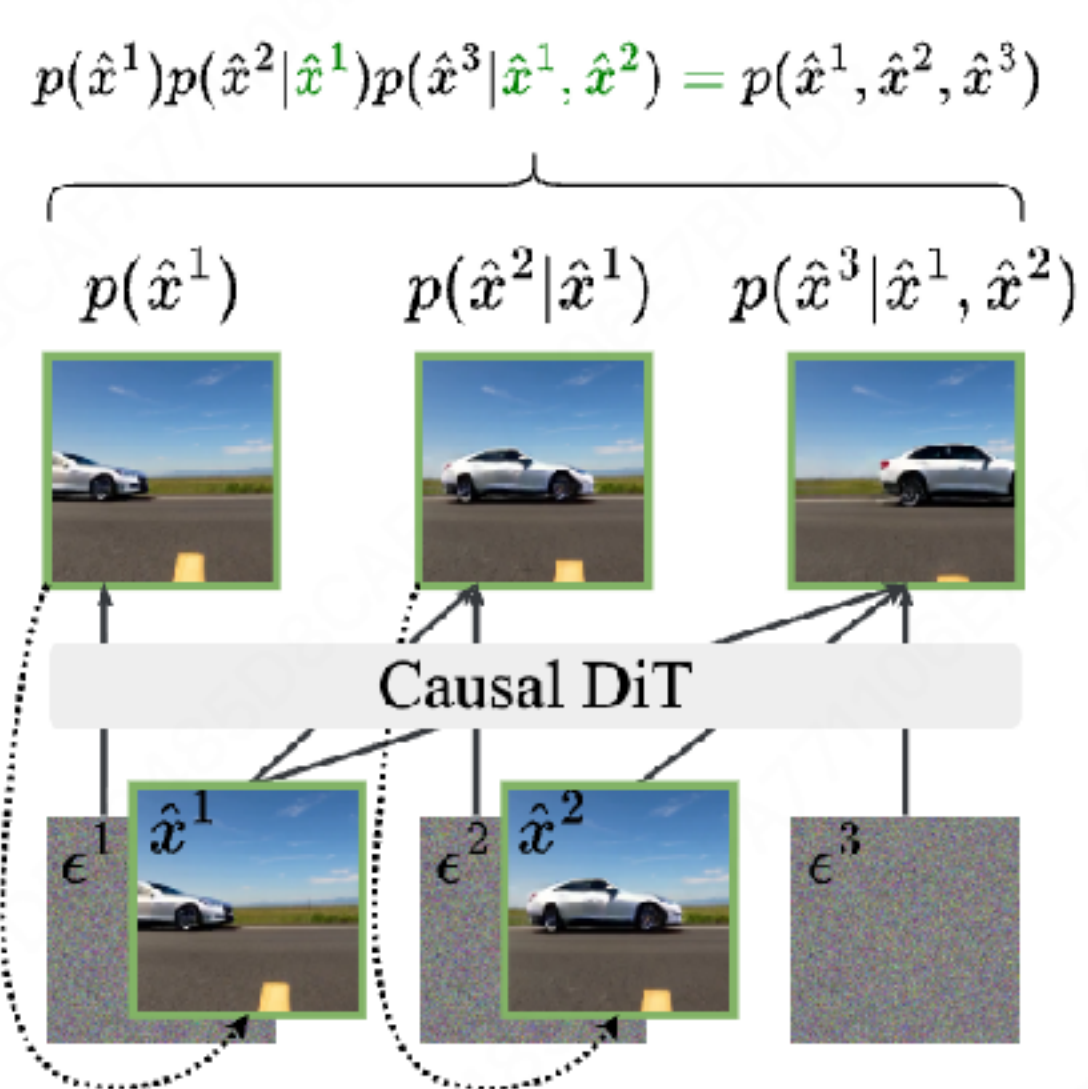
Test-Training gap



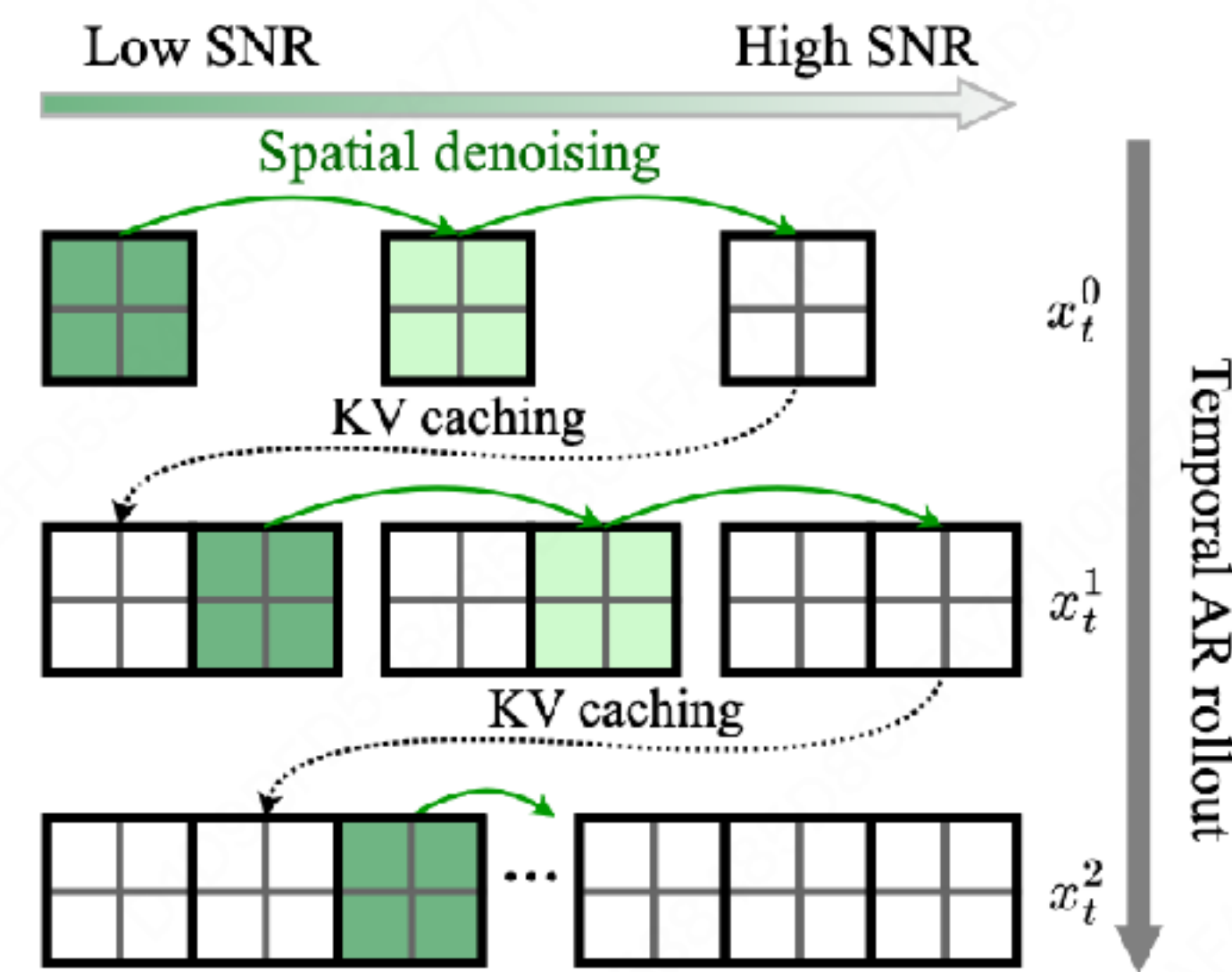
(a) Teacher Forcing Training



(b) Diffusion Forcing Training



(c) Self Forcing Training (ours)

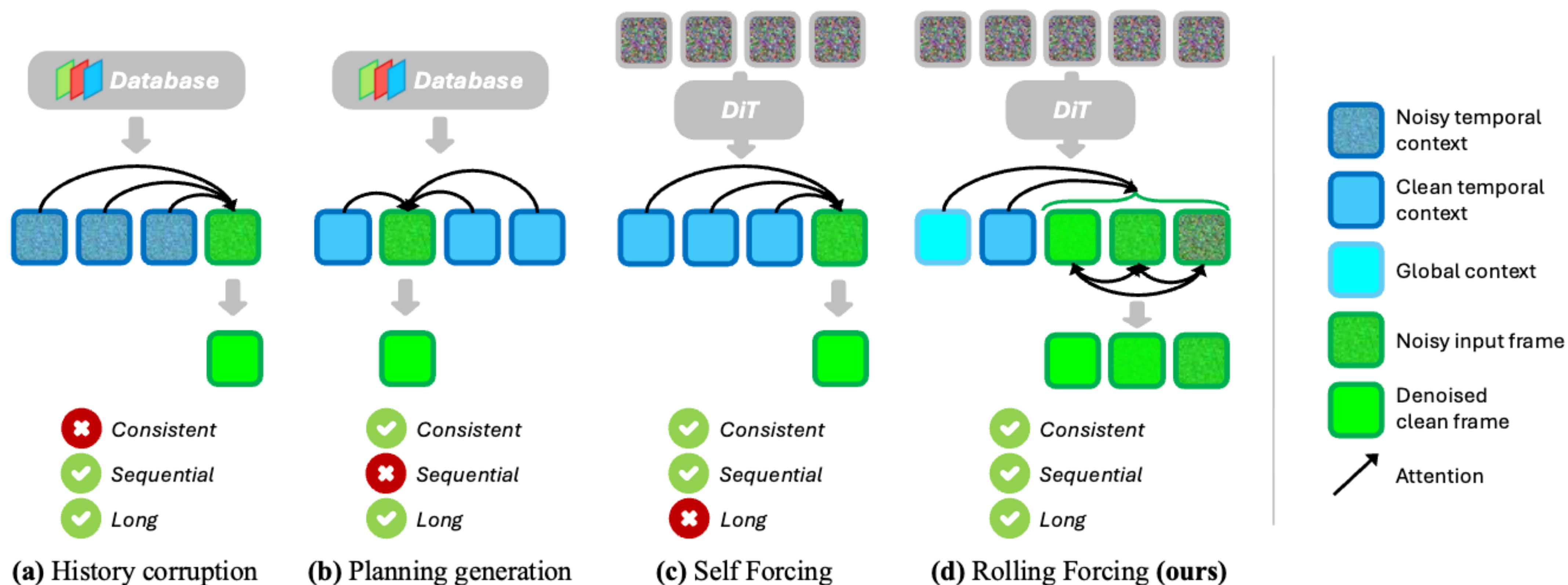


(c) Self Forcing/AR Inference

Stable Rollout

Rolling forcing

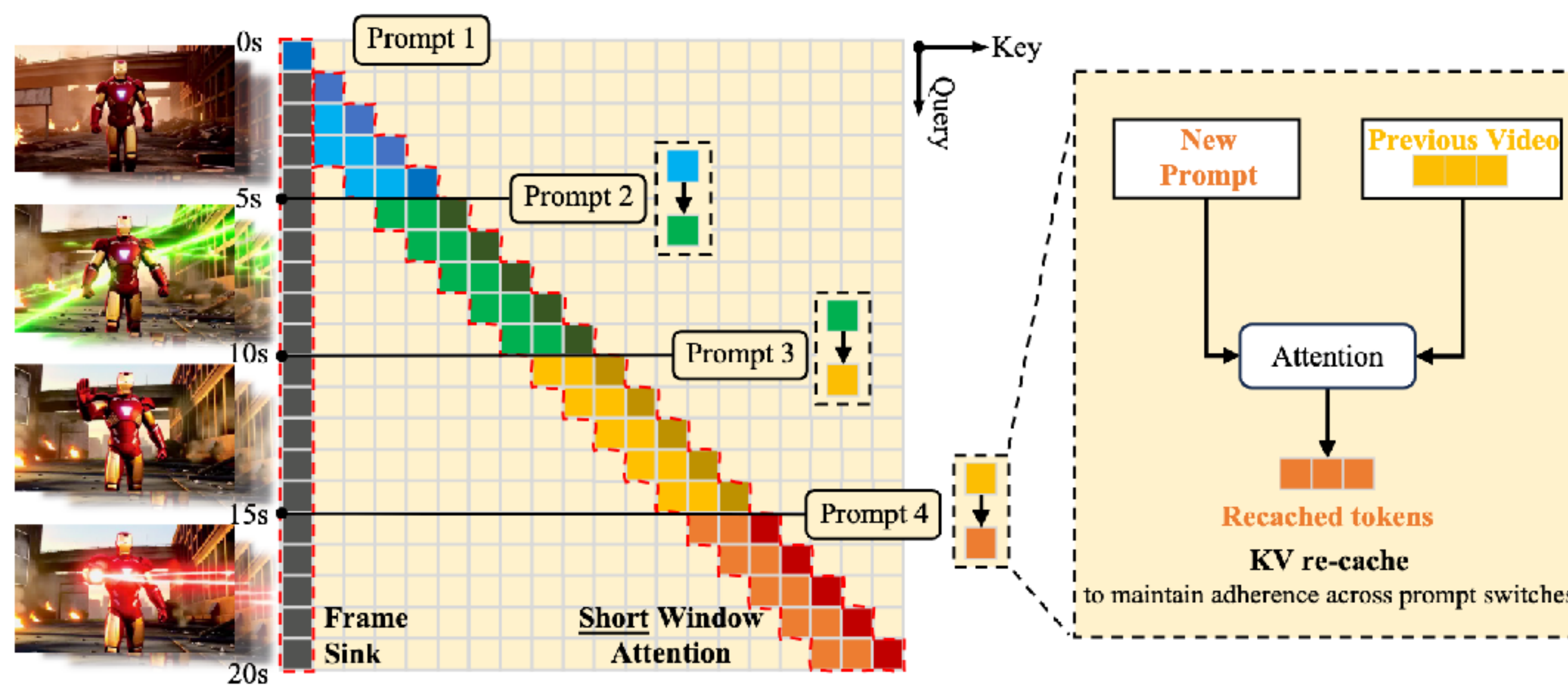
Rolling Window Denoising



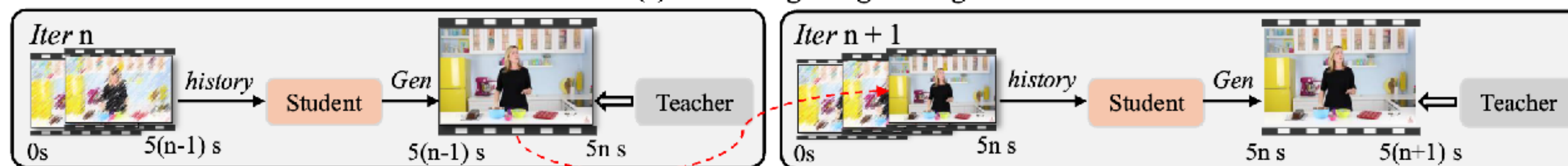
Stable Rollout

Longlive

Train-long-test-long



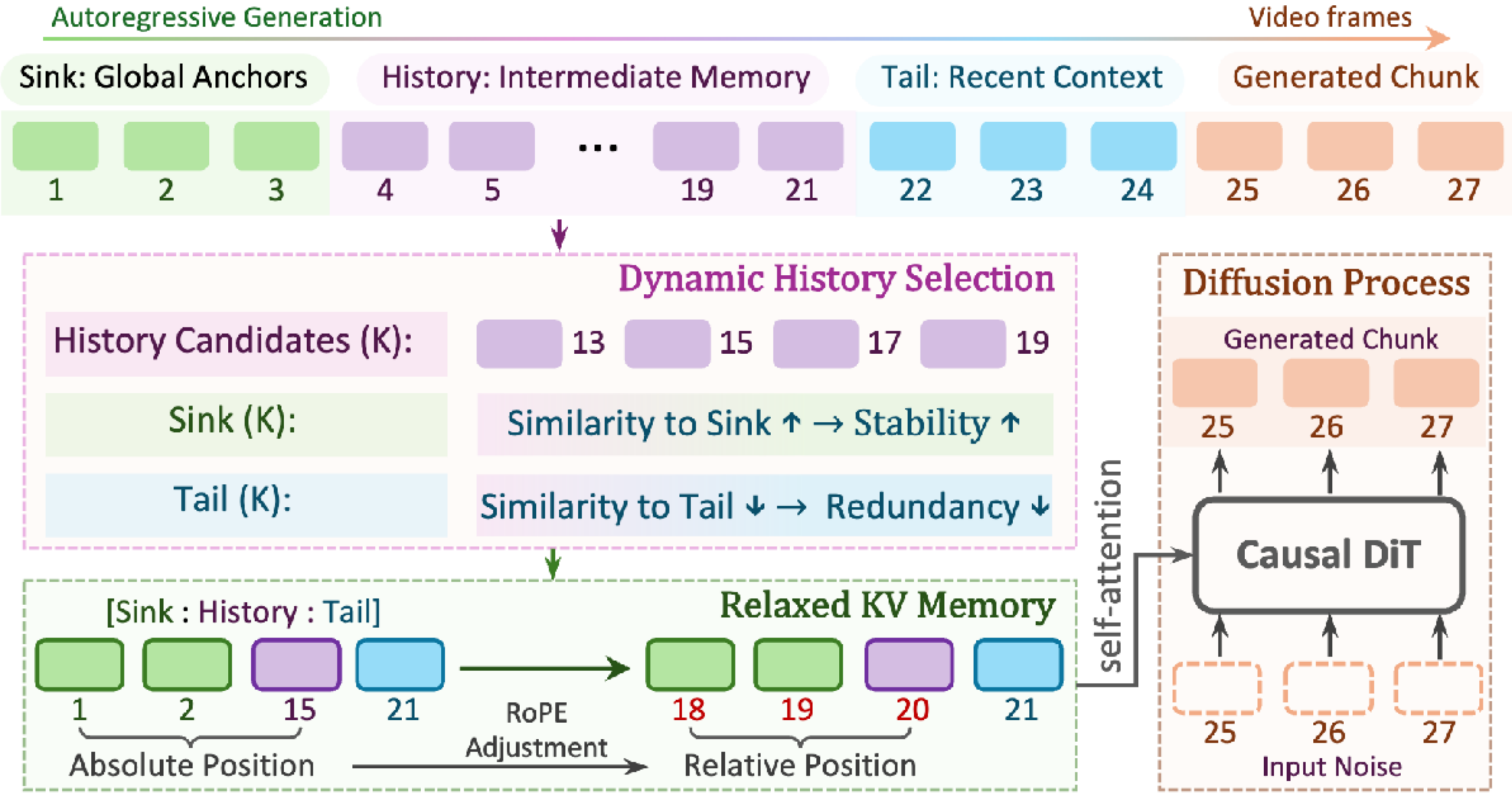
(c) Streaming Long Tuning



Stable Rollout

Relax forcing

Selecting Memory

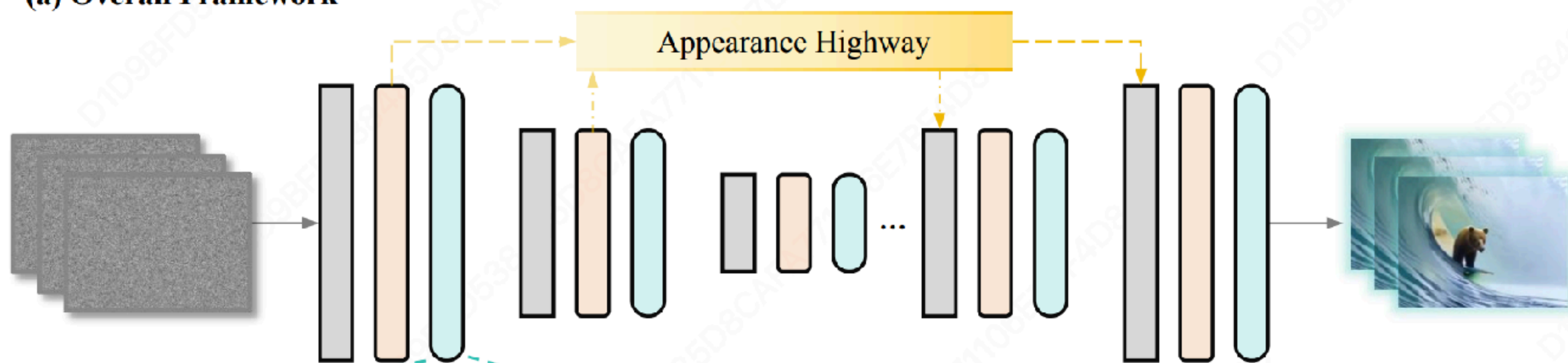


From Contextual to Parametric Memory

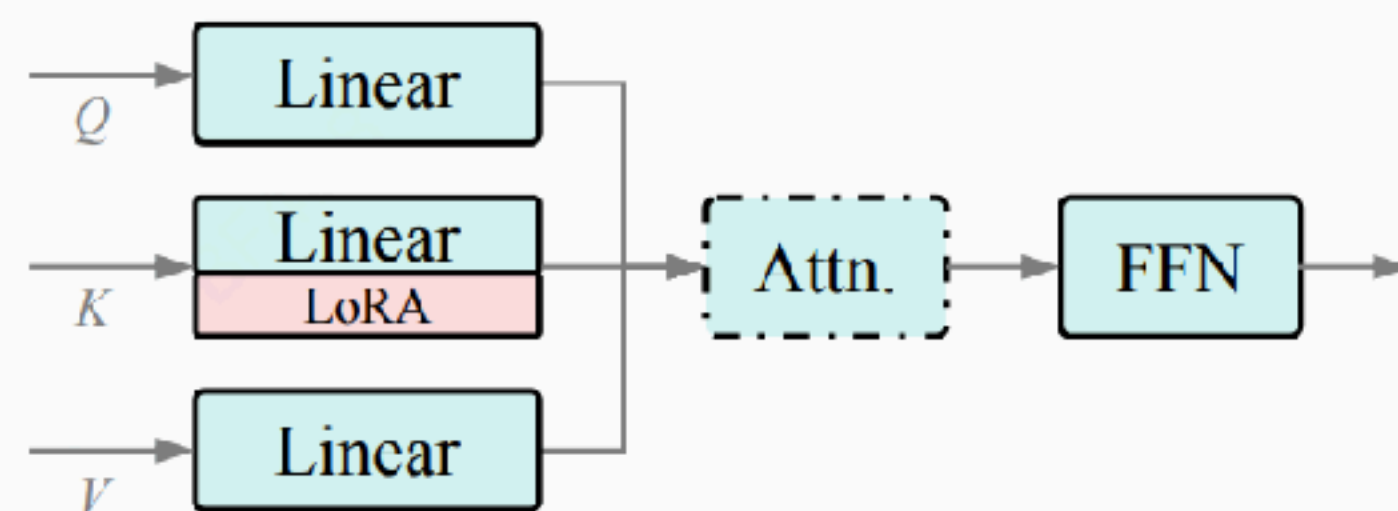
Parametric Memory

Separate Motion from Appearance

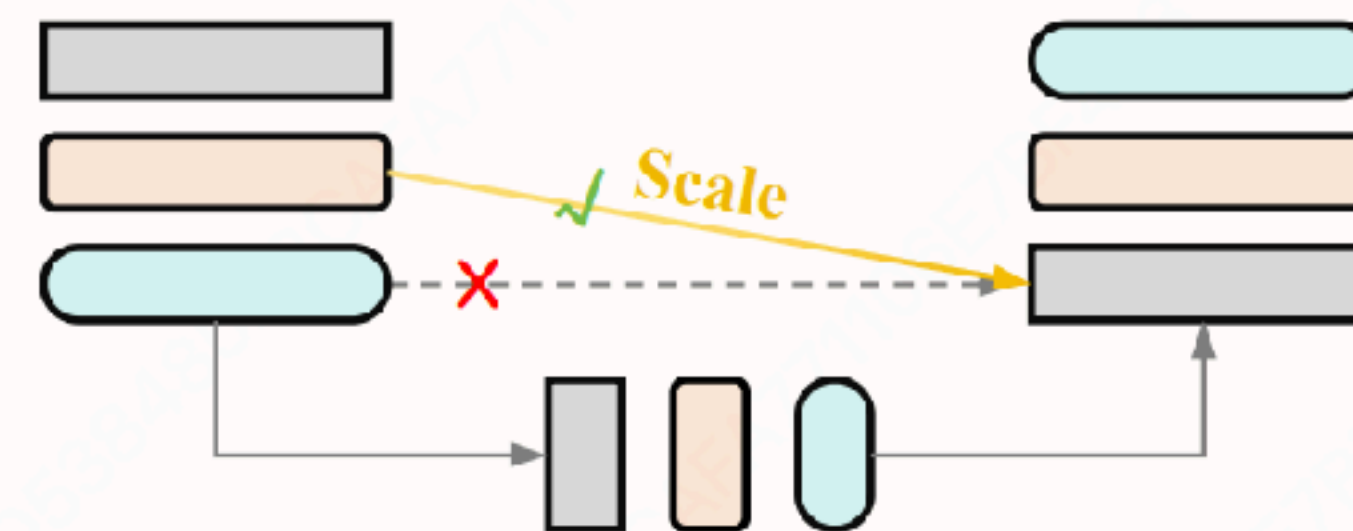
(a) Overall Framework



(b) Temporal Attention Purification



(c) Appearance Highway



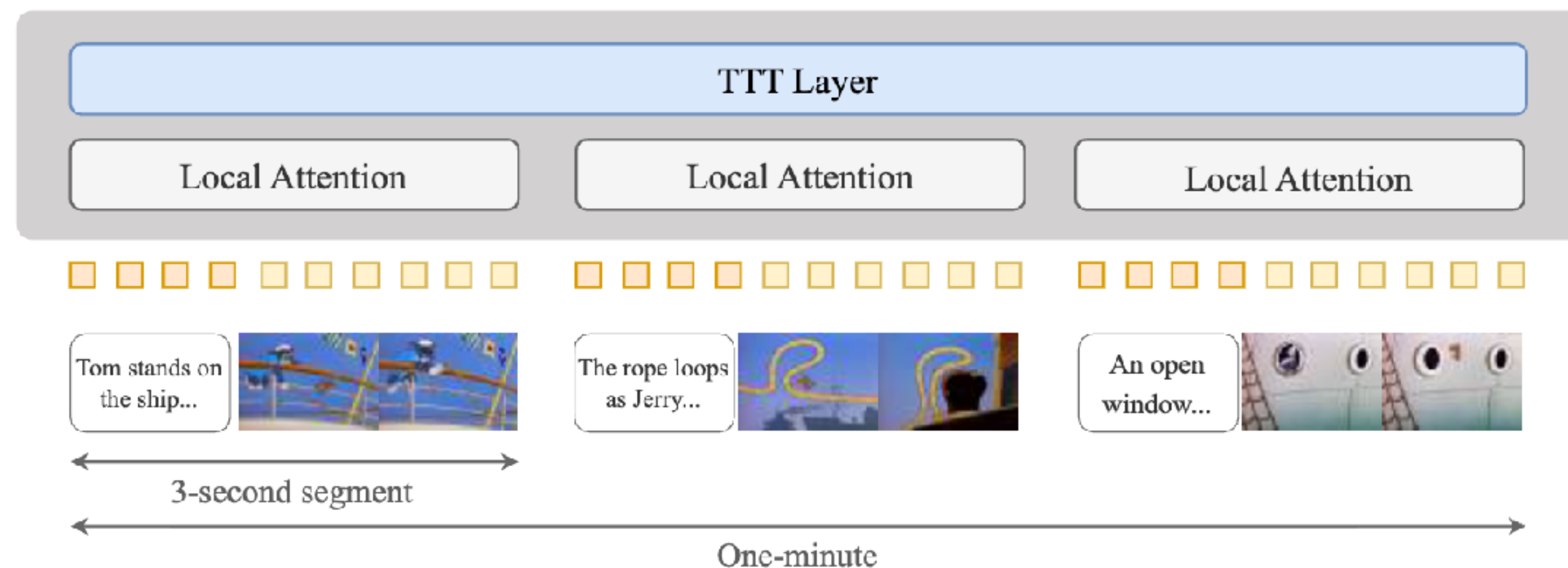
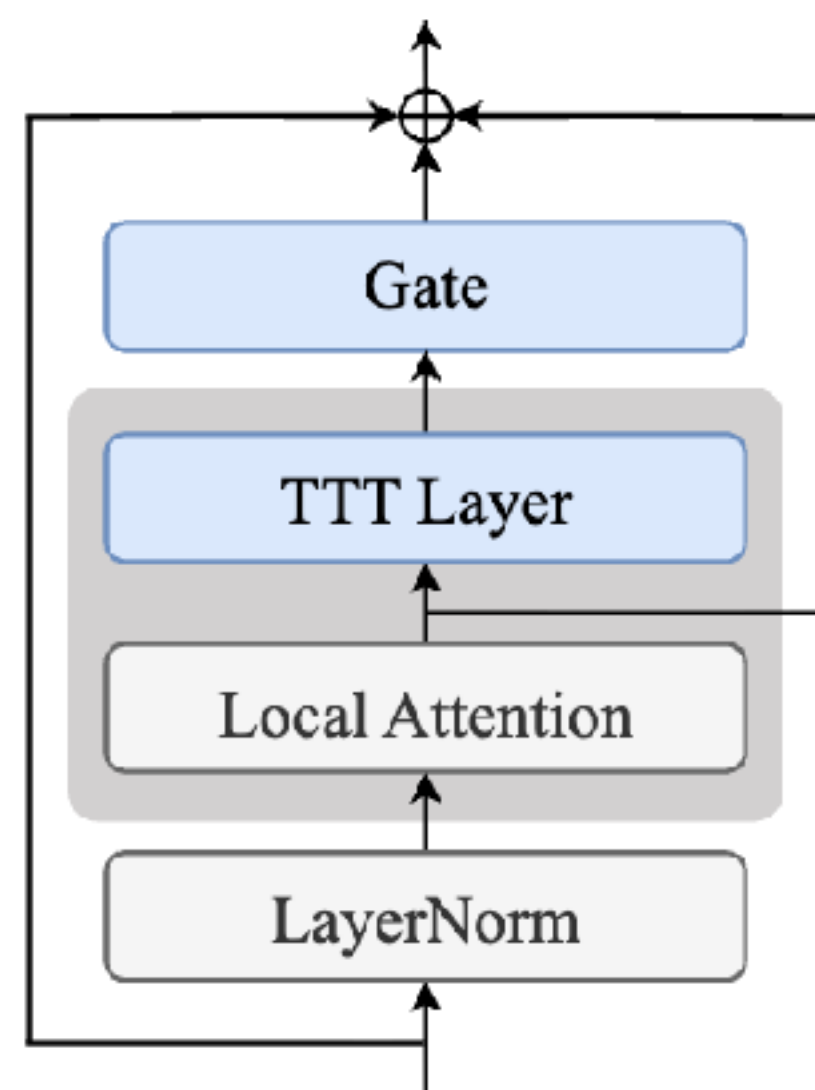
Conv Layers
 Spatial Transformer
 Temporal Attention Purification
 → Appearance Highway
 - - - → Skip Connection

Parametric Memory

One-Minute Video Generation with Test-Time Training

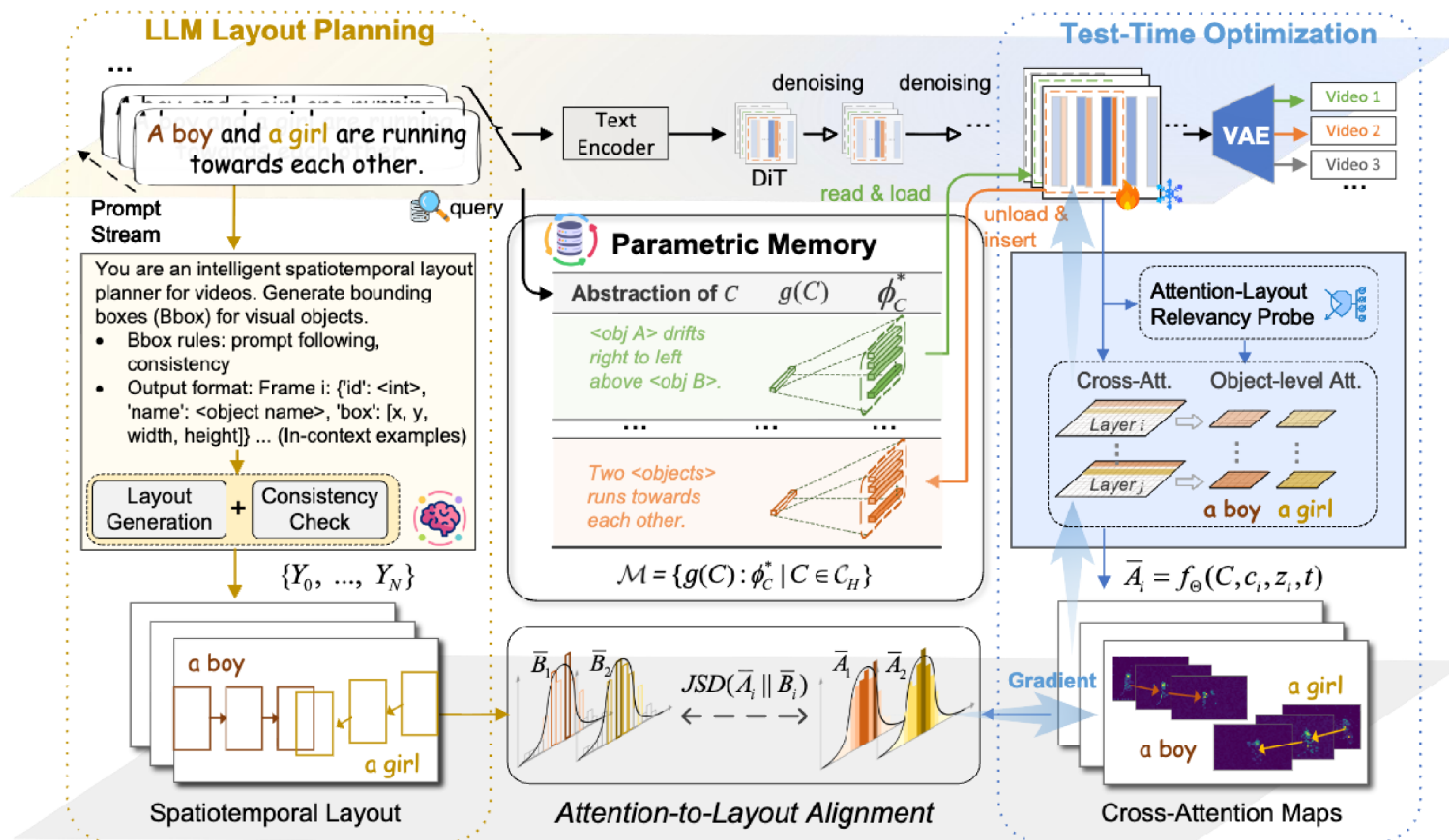
Diffusion based

$$\ell(W; x_t) = \|f(\theta_K x_t; W) - \theta_V x_t\|^2$$



Parametric Memory

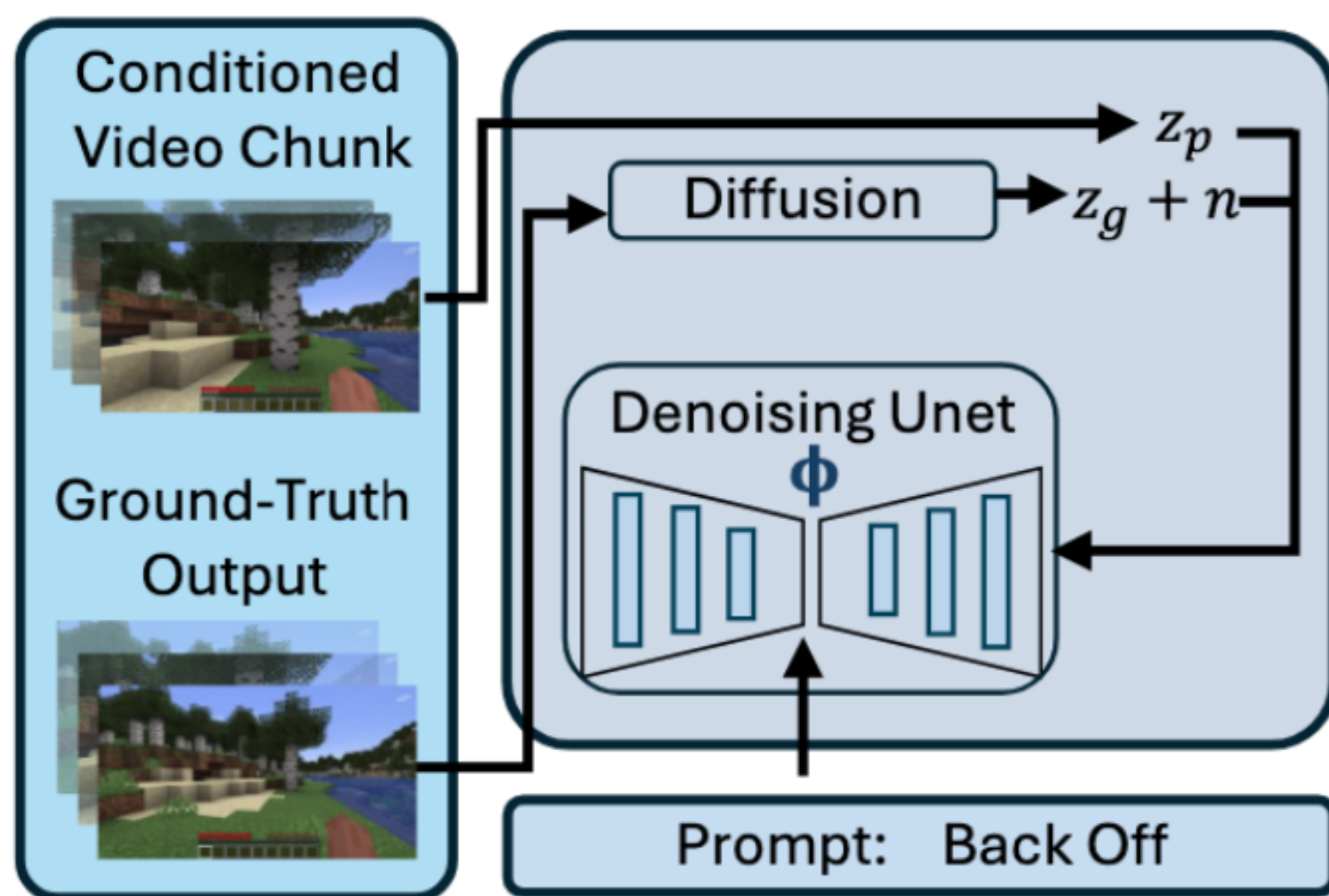
TTOM



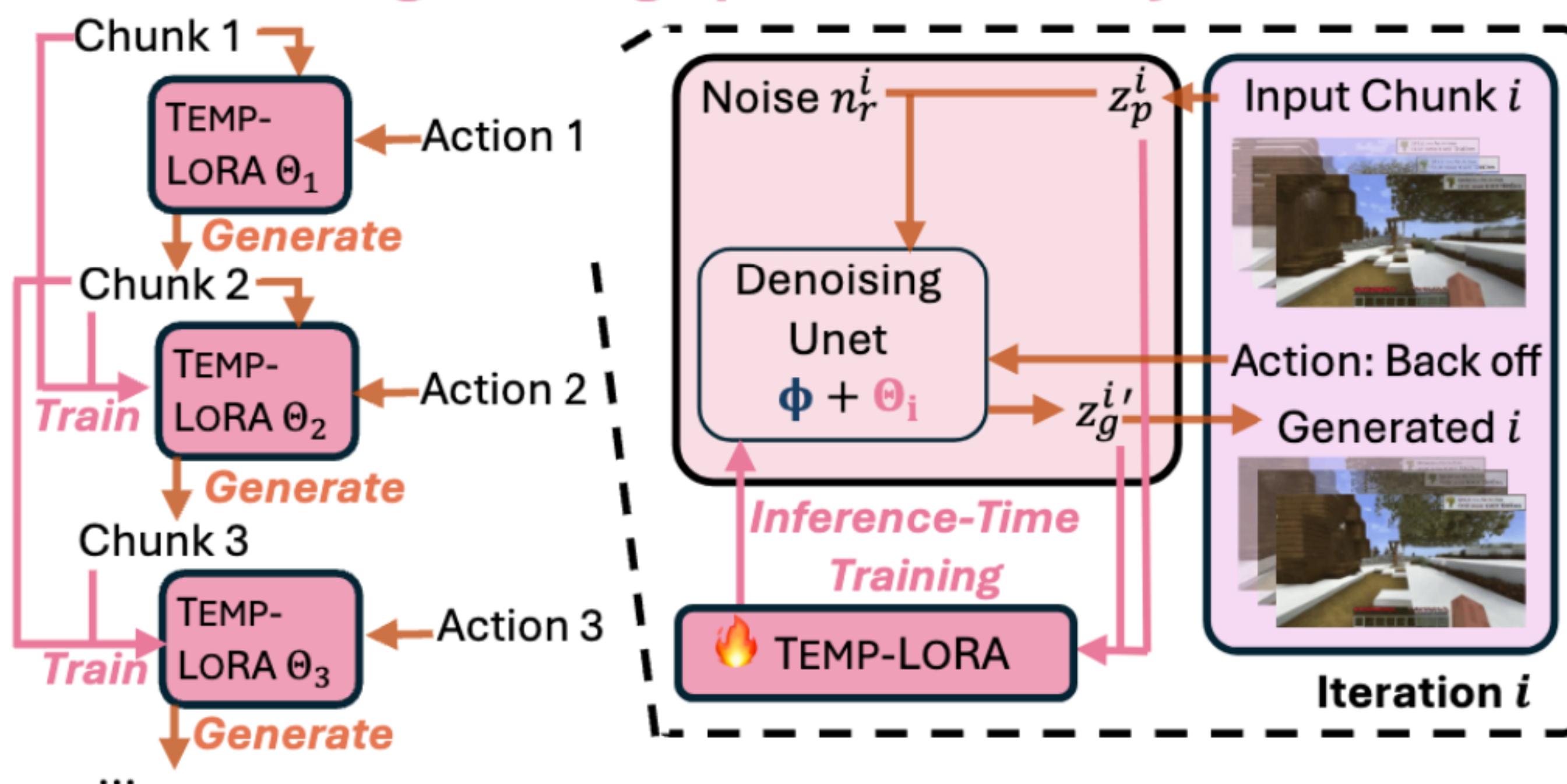
Parametric Memory

SlowFast-VGen

Slow Learning: Pretraining on All Data



Fast Learning: Storing Episodic Memory in TEMP-LORA

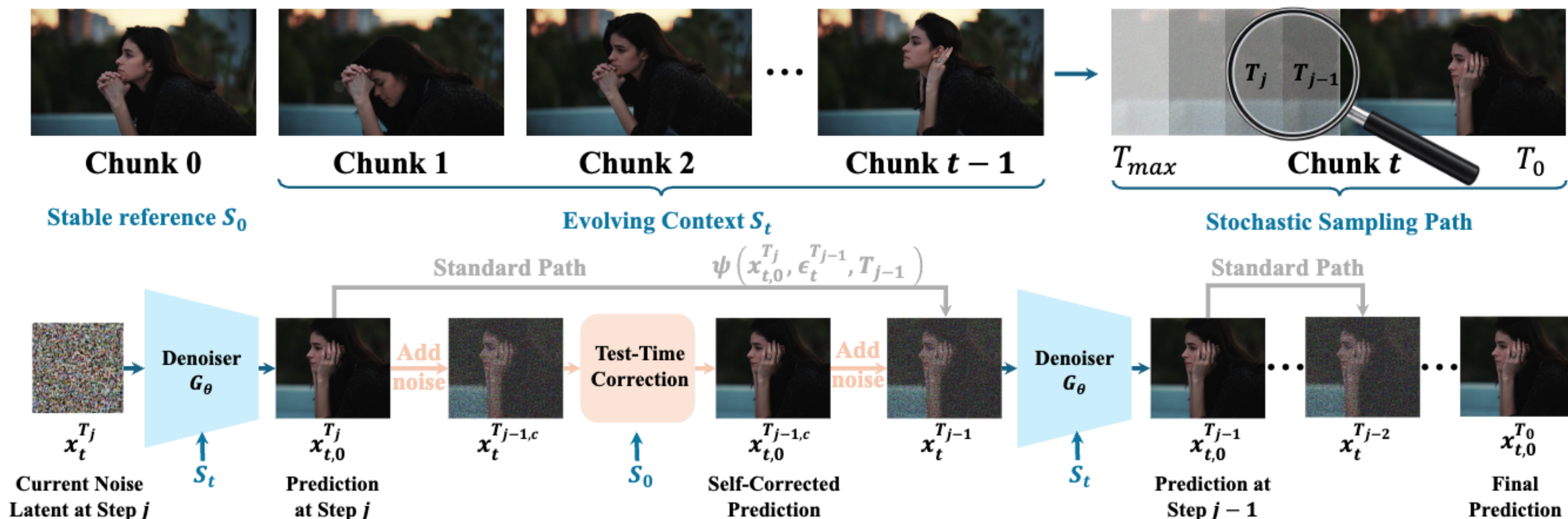


Inner Fast Learning Loop: Provide frozen θ_i for slow learning $\rightarrow \phi = \phi + \theta_i$.

Outer Slow Learning Loop: Provide frozen ϕ for Fast learning

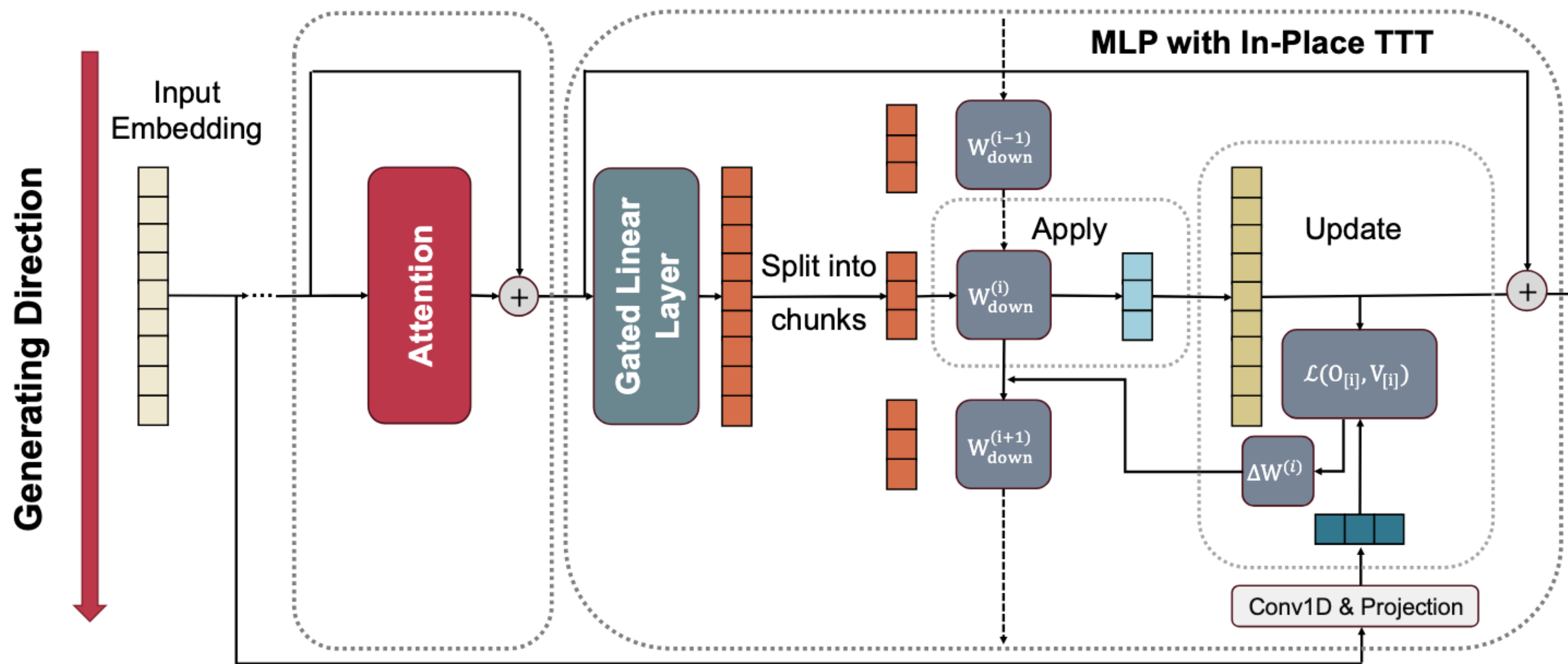
Parametric Memory

Pathwise Test-Time Correction for Autoregressive Long Video Generation



Parametric Memory

In-Place Test-Time Training



Parametric Memory

LoRA as Knowledge Memory

- *Hybrid Memory is better*
- *LoRA&ICL in video generation*
- *Test Time Training*

Embrace Hybrid Memory for Robustness. Our results indicate that LoRA-based parametric memory is best viewed not as a replacement for context-based methods, but as a complementary component. While LoRA can reduce inference costs for repeated access to a stable knowledge base, peak performance on long and multi-hop settings is achieved when LoRA is combined with external context (ICL/RAG), which mitigates fragmentation and restores global coherence. Together, these findings can position LoRA as a practical complementary axis of memory alongside RAG and ICL, and motivate hybrid systems as a promising direction for robust and efficient knowledge access.

核心问题

误差累积与长程锚点之间的trade off

- Anchor frame可抑制rollout shift，但是生成过程过分绑定在初始语义上、无法适应新语义
- attention sink过度关注首帧：sink collapse，后面记忆读不出来
- 如何管理context中的memory很重要，memory的重点是选择
- TTT：不在context中分配注意力预算，避免sink collapse，将部分长期信息写进参数中

LoRA

- 和历史帧的人物/背景一致性：与anchor frame的CLIP similarity
- diffusion loss重构：前期用anchor frame，后期转用lora：加噪训练重构前期用anchor frame生成的chunk，不过这次不用anchor frame；
- 生成两路：teacher（有anchor）+student（无anchor）

KV点积loss

需要预训练

$$W \leftarrow W - \eta \nabla_W \mathcal{L}(f_W(k), v)$$

$$\mathcal{L}(f_W(k_i), v_i) = -f_W(k_i)^\top v_i$$

$$f_W(x) = W_2 [\text{SiLU}(W_1 x) \circ (W_3 x)]$$

