

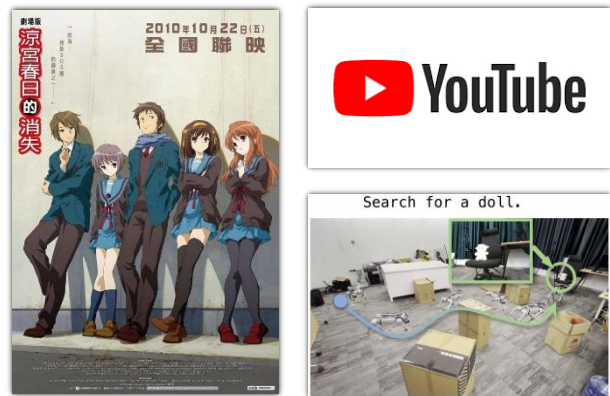
TimeViper:

A Hybrid Mamba-Transformer Vision-Language Model for Efficient Long Video Understanding

Boshen Xu¹, Zihan Xiao¹, Jiaze Li², Jianzhong Ju², Zhenbo Luo², Jian Luan², Qin Jin¹
AIM3 Lab, Renmin University of China¹, MiLM Plus, Xiaomi Inc.²

Background

- Long video understanding is essential for...



- MLLM paradigm for hour-long video understanding



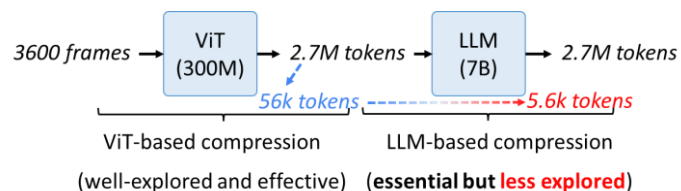
Challenges

- Efficient backbone construction

Attention Computation: $O(N^2d)$ $\mathbf{o} = \text{softmax}(\mathbf{QK}^T)\mathbf{V} \in \mathbb{R}^{N \times d}$
KV-Cache: $O(Nd)$

Linearized Computation: $O(Nd^2)$ $\mathbf{S}_t = \lambda_t \mathbf{S}_{t-1} + \mathbf{v}_t \mathbf{k}_t^T \in \mathbb{R}^{d \times d}$
or Mamba 1 Cache: $O(1d^2)$ $\mathbf{o}_t = \mathbf{S}_t \mathbf{q}_t \in \mathbb{R}^d$

- Severe token redundancy



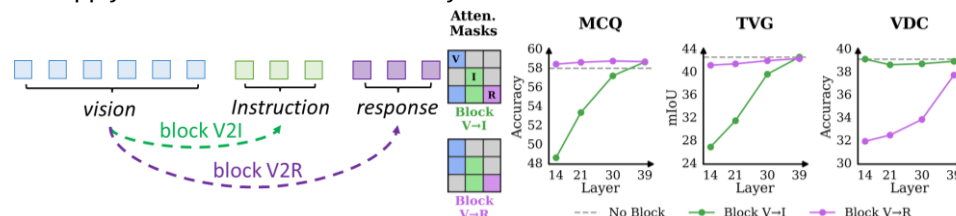
Develop & Interpret & Compress

We first develop a Hybrid Mamba-Transformer VLM baseline to observe...



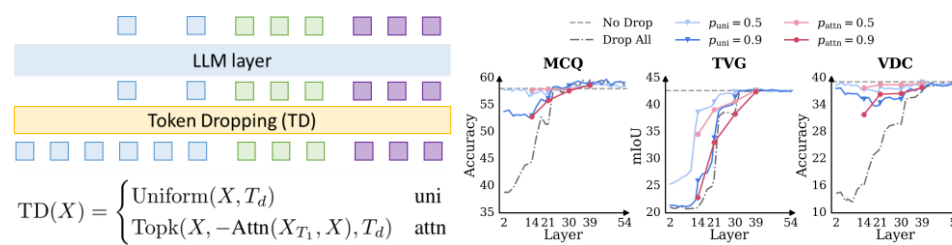
- Obs1: Vision-to-text information aggregation

apply attention mask at certain layer for vision token flow



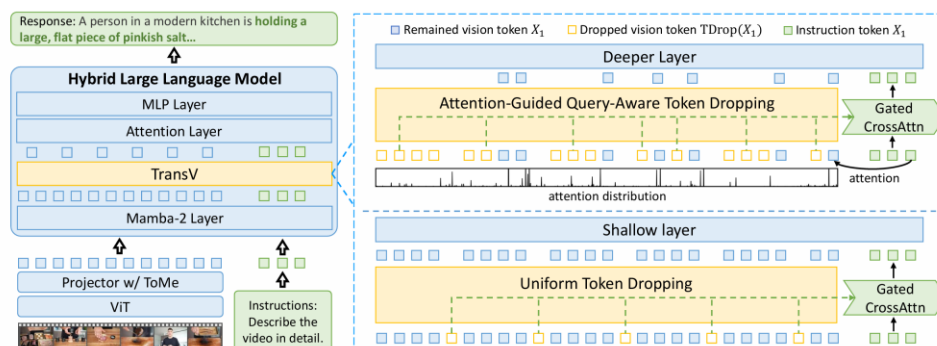
- Obs2: Severe vision tokens redundancy

apply token dropping with different strategy at certain layer



- TransV: Compress a hybrid VLM within LLM

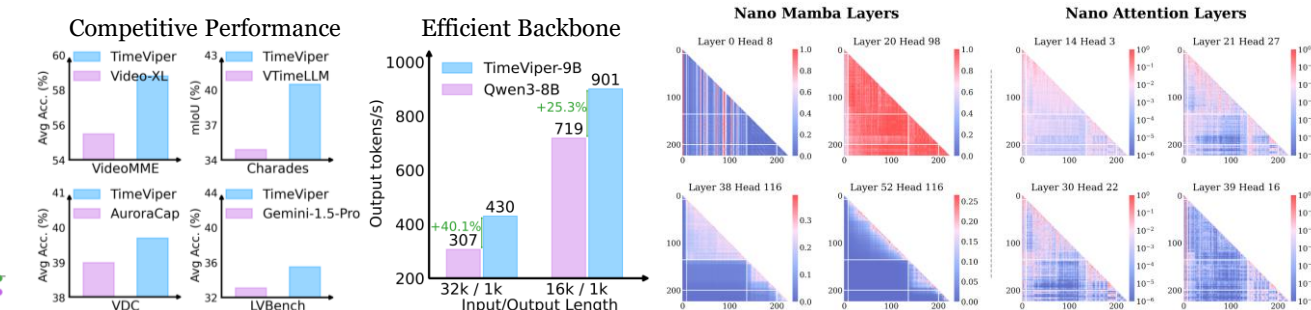
Dropping rate: 7th shallow layer 50%, 39th layer 90%



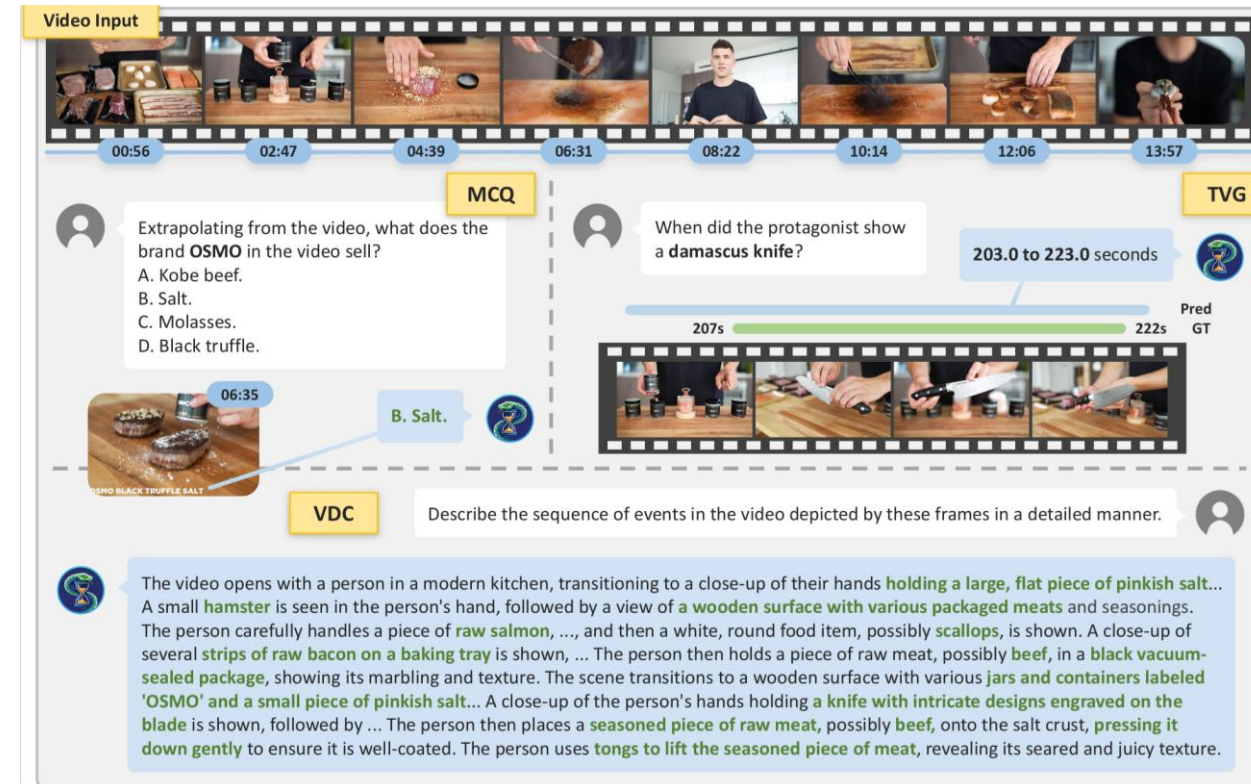
Turn out to be similar to Transformer-based VL processing

Results

- Comparisons of benchmark performance, decoding efficiency and attention visualizations



- TimeViper works well in video question answering, temporal video grounding and detailed descriptions



More quantitative, efficiency comparison, and interesting visualizations can be found in our paper :)
& Some interesting concurrent works from industry: Nano-VL-9B (NVIDIA), MiMo-V2-Omni (Xiaomi)