



Shrinking Vision Language Models for Edge Deployment via Distillation and Pruning

Denis Gudovskiy

Distinguished Engineer

Panasonic AI Lab, Mountain View, CA



Agenda

Applications using Vision-Language Models (VLMs)

- Our UI agent running on the phone
- LLMs vs. VLMs for physical AI agents

Part 1. Visual token pruning for VLMs

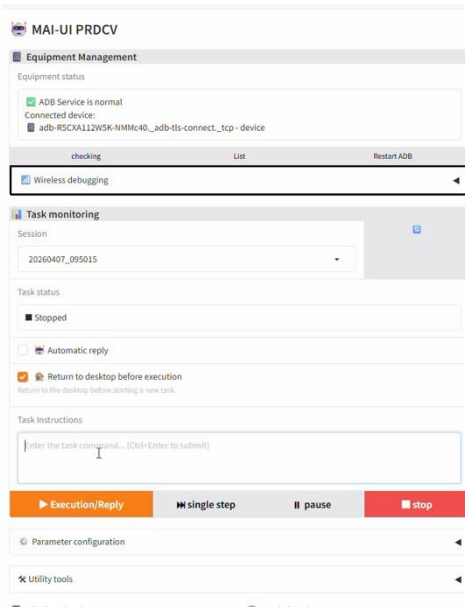
- Overview of recent training-free methods
- Our [SparseVLM](#) @ [ICML 2025](#)

Part 2. Distilling representations in 3D VLMs

- 3D spatial reasoning using VLMs
- Our [Proxy3D](#) @ [CVPR 2026](#)

Applications: Our UI Agent

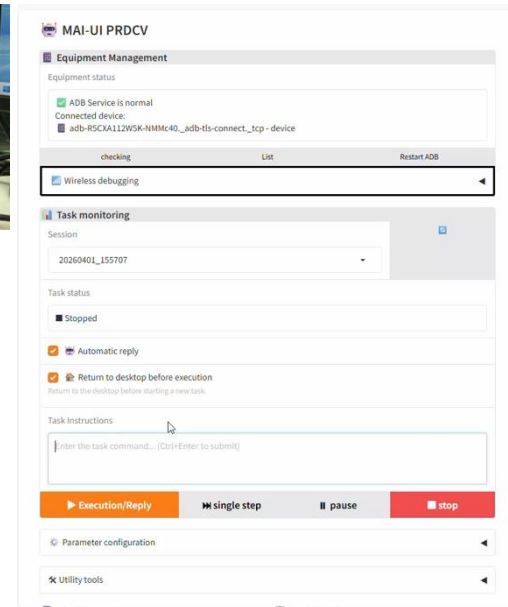
- API for LLM agents is not always available
- Screenshot-based agents are more general



Android Phone
@ ADB

←GitHub actions

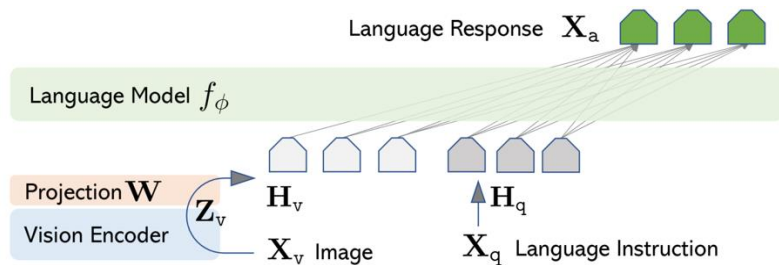
Gallery actions→



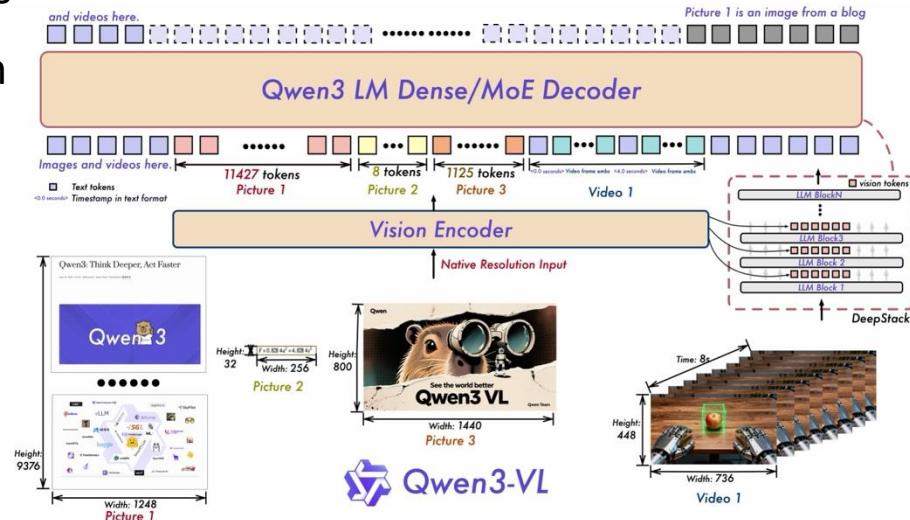
Overview

Vision-Language Models (VLMs)

- Multimodal LLMs (MLLMs) with visual inputs
- Early VLMs had fixed-resolution tokenization
- Recent ones have dynamic variable-lengths



Early VLMs: [LLaVA](#) network architecture

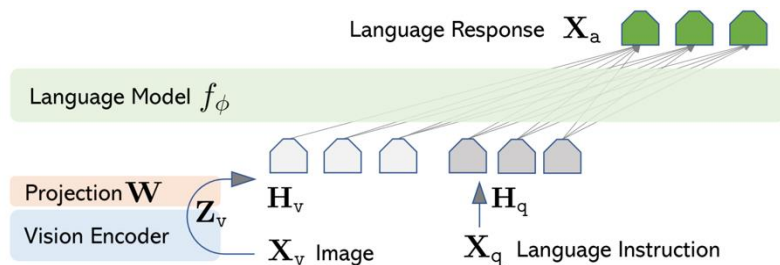


Recent VLMs: [Qwen3-VL](#) tokenization scheme

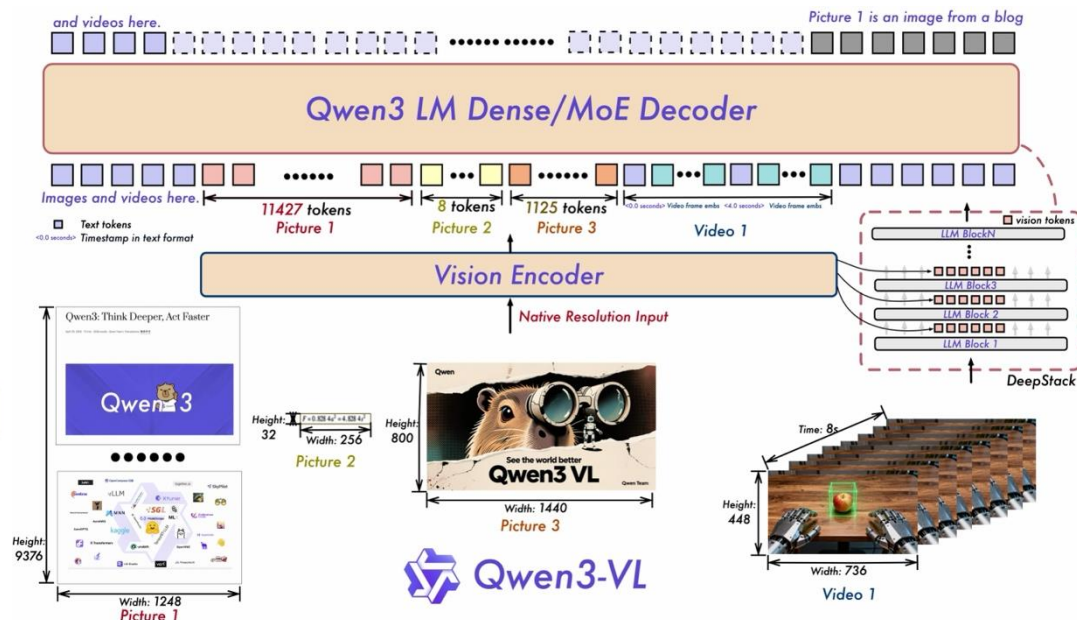
Problem Statement

Exponential growth in number of tokens

- 3D/video >> 2D images >> 1D text
- How to reduce number of tokens?
- How to encode visual modality?



Early VLMs: [LLaVA](#) network architecture



Recent VLMs: [Qwen3-VL](#) tokenization scheme

Part 1. Training-free Token Pruning

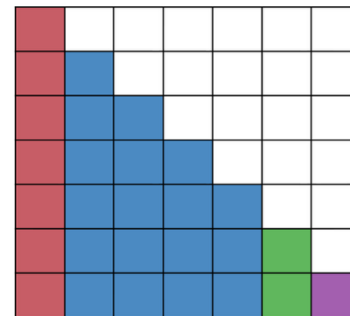
Visual token pruning

- Self-attention complexity is quadratic with the sequence
- Training-based pruning is hard due to complex pipelines
- Training-free visual token pruning is significantly easier
- Hardware support by edge devices and GPUs

Spatial and temporal redundancy

- Static visual / video inputs are often sparse
- Examples: security cameras, user interfaces
- How to prune tokens in a training-free manner?

Token Attention Illustration



LLaVA self-attention example from [FastV](#)

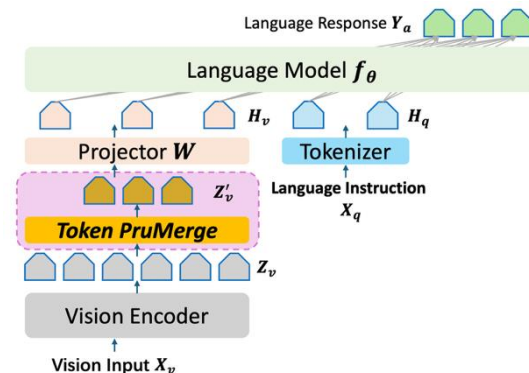
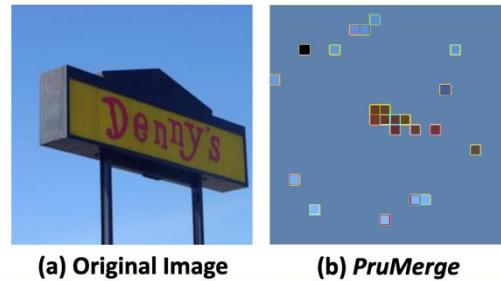


Existing Methods Training-free Methods

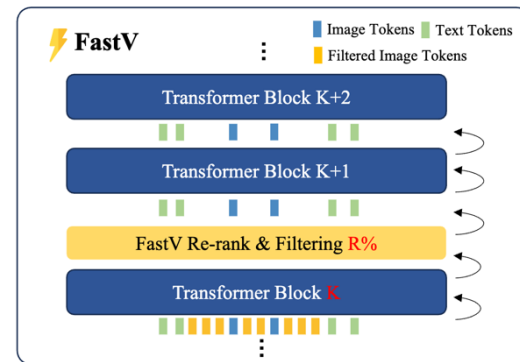
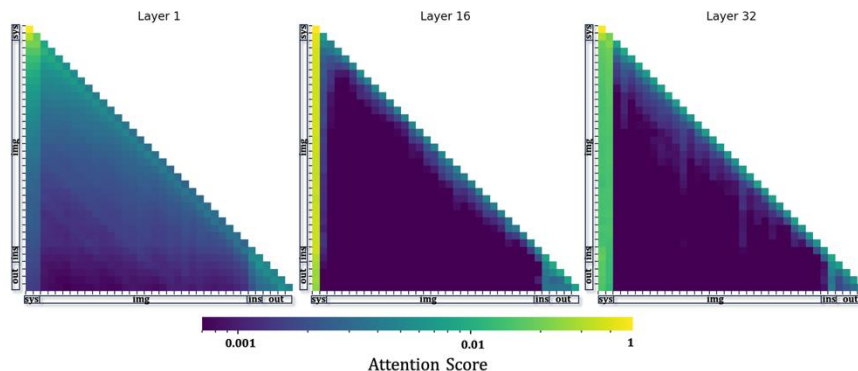
Localization: input vs. layer-wise

Scoring: relevance vs. attention

Selection: pruning vs. merging



Input pruning in [PruMerge](#)

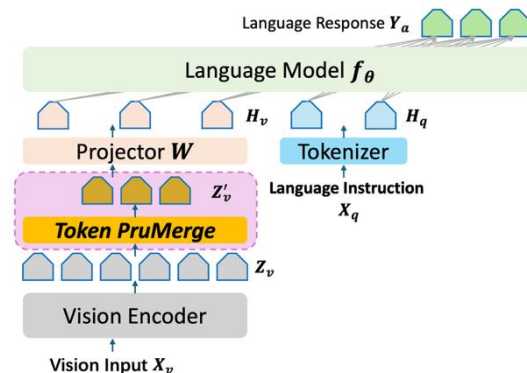
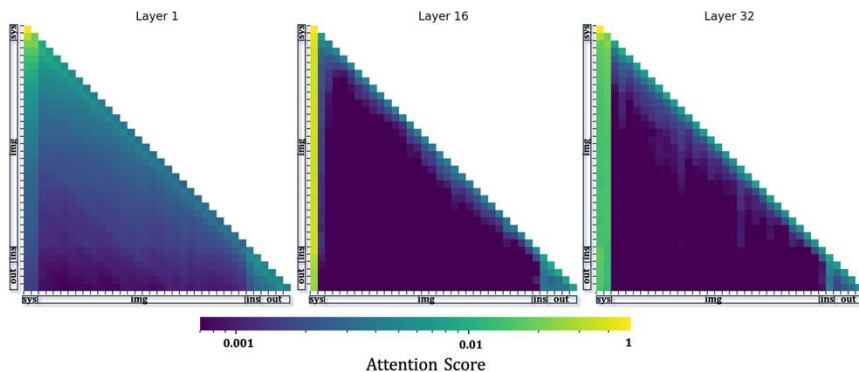
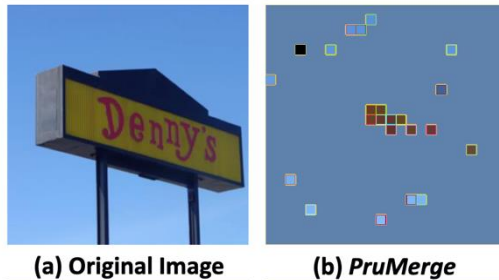


Block-wise pruning in [FastV](#)

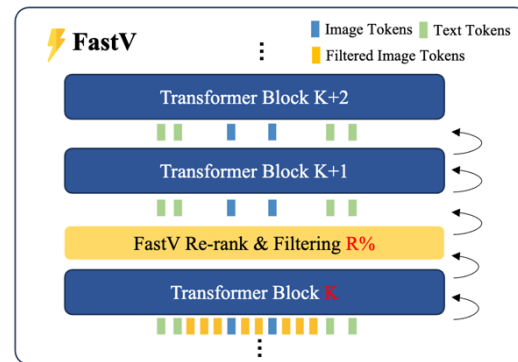
Existing Methods Training-free Methods

Drawbacks

- Text-agnostic methods
- Suboptimal sparsification



Input pruning in PruMerge



Block-wise pruning in FastV

SparseVLM: Text Prompt Guidance for Pruning

- Queries typically contain questions about specific objects
- Text prompts can be used to guide a VLM sparsification method



(a) Input Image

[SparseVLM @ ICML'25](#)

Q-1: What is written on this blue sign?



Q-2: How many buses are there in the image?



Q-3: What color is this roof?



(b) Ours. Text-guided Visual Sparsification



(c) Existing Methods. Text-agnostic Visual Sparsification

SparseVLM: Attention Scores for Sparsification

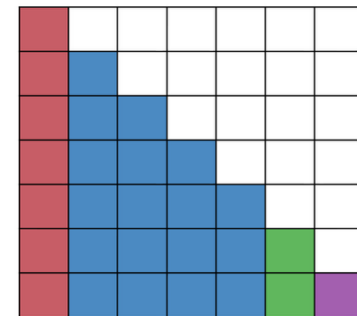
- Main idea is to use transformer's self-attention scores
- Then, we can estimate significance of visual tokens w.r.t. text query
- This is computationally-free and compatible with e.g., FlashAttention

$$P = A_{i,j}, \text{ and } (i, j) \in \{\mathbb{L}, \mathbb{I}\},$$

$$\tilde{p} = [\tilde{p}_1, \tilde{p}_2, \dots, \tilde{p}_{L_v}] = \frac{1}{L_t} \sum_{i=1}^{L_t} P_i.$$

How to get
the priority
matrix?

Token Attention Illustration



Red System Prompt (35) Green User Instruction (135)
Blue Image Tokens (576) Purple Output Tokens (150)

SparseVLM: Token Selection and Sparsification Level

- Not all text tokens are relevant to the sparsification objective
- Hence, we use “raters” to estimate relevance

$$R = \frac{1}{L_v} \sum_{j=1}^{L_v} \text{Softmax}(H_v H_q^T),$$

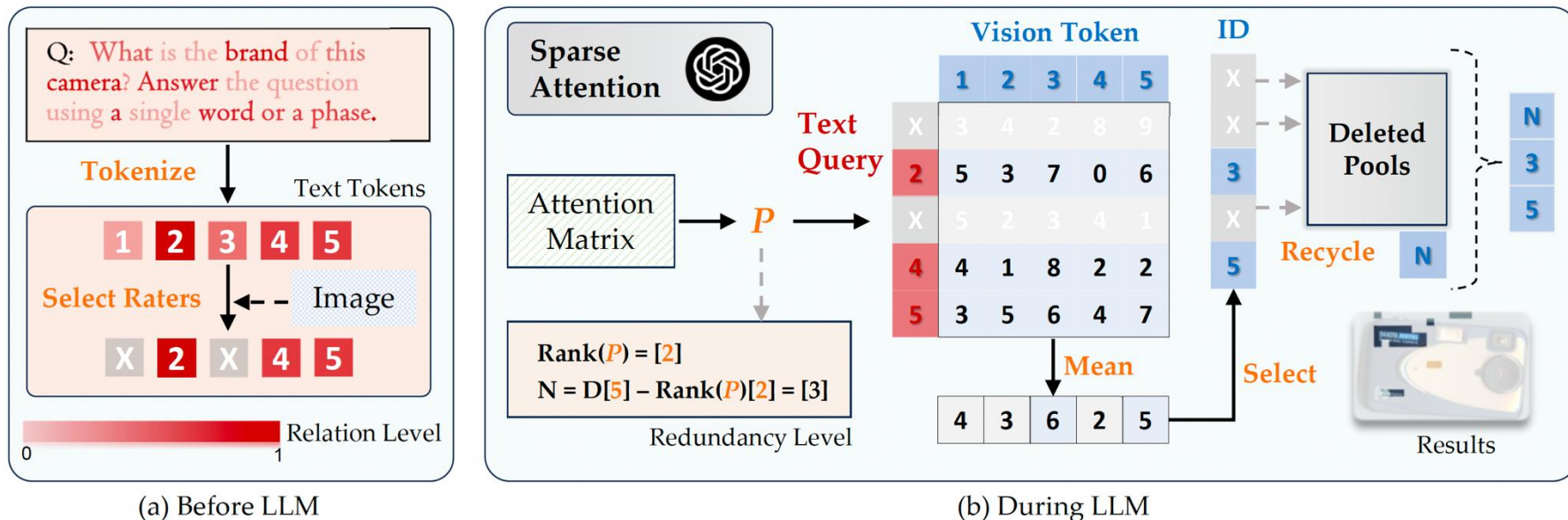
$$S = \{i | R[i] \geq \text{Mean}(R), i \in \{1, 2, \dots, L_t\}\}.$$

- Then, the matrix rank can be used to find visual token redundancy

$$N = \lambda \times (L_v - \text{Rank}(P)).$$

SparseVLM: Overall Diagram

- Visual tokens are being pruned at each decoder layer with high sparsification levels
- Therefore, we recycle certain visual tokens by aggregation and reconstruction



SparseVLM: Pruning Results

Method	GQA	MMB	MME	POPE	SQA	SEED	VQA ^{Text}	MMVet	Acc. (%)	FLOPs (T)	Latency (ms)
<i>Upper Bound, 576 Tokens (100%)</i>											
Vanilla	61.9 100%	64.6 100%	1864 100%	85.9 100%	69.5 100%	60.3 100%	58.3 100%	30.9 100%	100	4.62	57.82
<i>Retain 128 Tokens (↓ 77.8%)</i>											
ToMe _(ICLR23)	52.4 84.7%	53.3 82.4%	1343 72.1%	62.8 73.1%	59.6 85.8%	50.9 84.4%	49.1 84.4%	27.2 88.0%	81.9 (↓ 18.1)	1.62	30.00
FastV _(ECCV24)	49.6 80.1%	56.1 86.8%	1490 79.9%	53.4 62.2%	68.6 98.7%	48.1 79.8%	50.5 86.6%	26.3 85.1%	82.4 (↓ 17.6)	1.70	30.70
PDrop _(CVPR25)	56.0 90.5%	61.1 95.4%	1664 89.3%	82.3 95.8%	69.9 100.6%	53.3 88.4%	55.1 94.5%	30.8 99.7%	94.3 (↓ 5.7)	1.62	37.77
SparseVLM	58.4 94.3%	64.5 99.8%	1746 93.7%	85.0 99.0%	68.6 98.7%	58.2 96.5%	56.7 97.3%	29.0 93.9%	96.7 (↓ 3.3)	1.72	33.28
<i>Retain 64 Tokens (↓ 88.9%)</i>											
ToMe _(ICLR23)	48.6 78.5%	43.7 67.5%	1138 61.1%	52.5 61.1%	50.0 71.9%	44.0 73.0%	45.3 77.8%	24.1 78.0%	71.1 (↓ 28.9)	1.19	26.52
FastV _(ECCV24)	46.1 74.5%	47.2 73.1%	1255 67.3%	38.2 44.5%	68.7 98.8%	43.7 72.5%	47.8 82.0%	19.6 63.4%	72.0 (↓ 28.0)	1.29	27.30
PDrop _(CVPR25)	41.9 67.7%	33.3 51.6%	1092 58.6%	55.9 65.1%	69.2 99.6%	40.0 66.3%	45.9 78.7%	30.7 99.4%	73.4 (↓ 26.6)	1.18	43.41
SparseVLM	53.8 86.9%	60.1 93.0%	1589 85.2%	77.5 90.2%	69.8 100.4%	52.2 86.6%	53.4 91.6%	24.9 80.6%	89.3 (↓ 10.7)	1.30	29.89

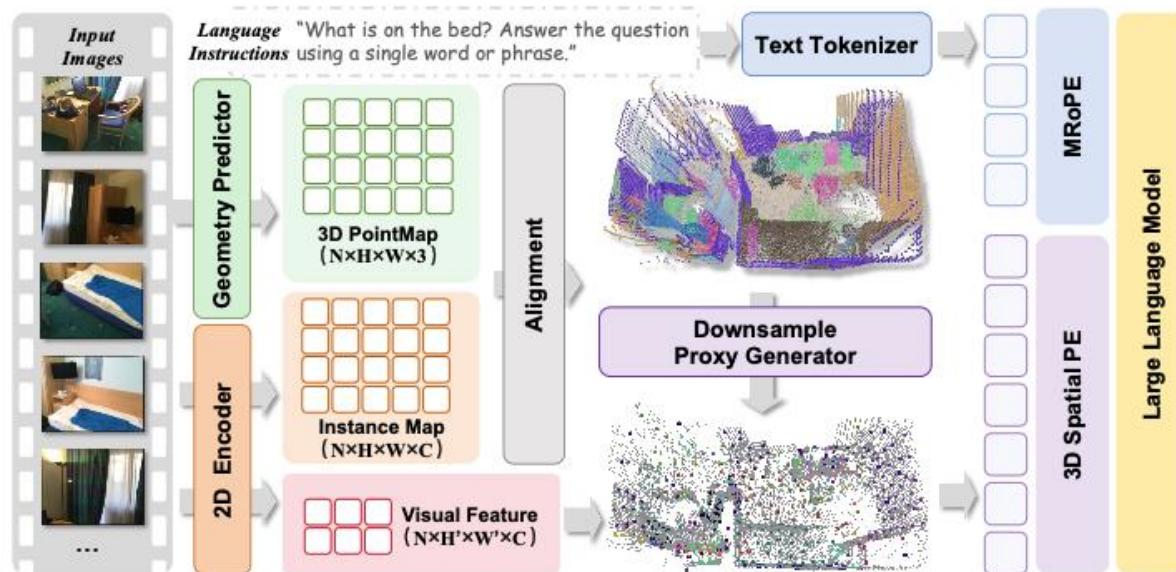
SparseVLM: Pruning Results

			 Is there a backpack in the image?	 Yes.
			 Are the shorts large and blue?	 No.
			 Are there both toothbrushes and mats in this picture?	 No.
			 Do the balls to the left of the other ball look right?	 Yes.
			 Does the sweater look open and blue?	 No.

Part 2. VLMs for 3D Spatial Reasoning

3D VLMs are an extension of conventional VLMs

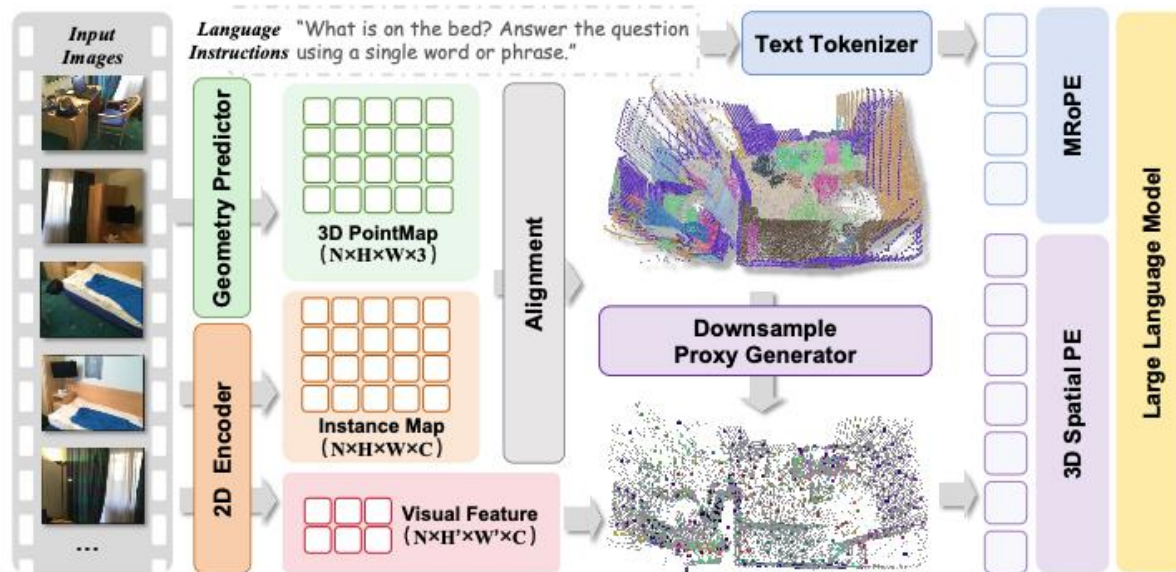
- A set of dense input frames is encoded into a latent token sequence
- How to reduce the number of 3D tokens for the edge deployment?



Comparison to Prior Methods

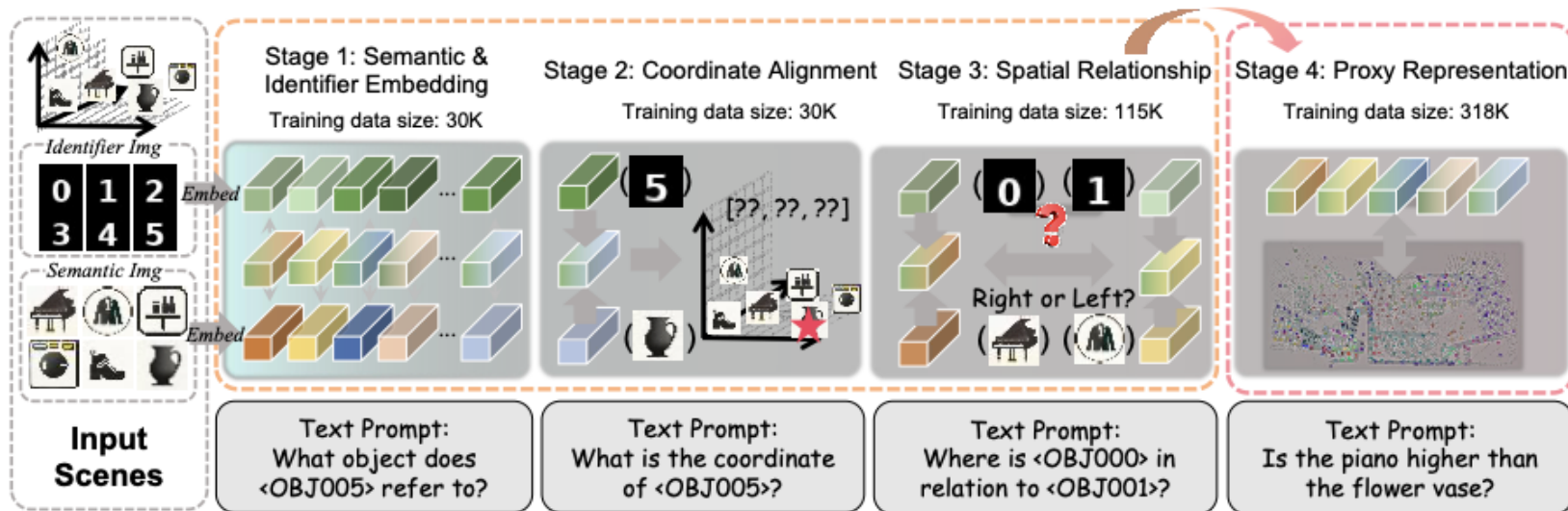
Training-free methods underperform and a distillation training is required

- Representation vs. weight distillation: quadratic vs. linear, KV cache
- Our [Proxy3D](#) with the representation distillation via 3D clustering



Proxy3D: Training Pipeline

- Distill the pretrained Qwen2.5-VL knowledge to Proxy3D with token clustering
- Main idea: grow 3D spatial reasoning in a sequential manner
- Data scarcity: overcome it using multi-stage training pipeline



Proxy3D: Tasks and Results

3D scene understanding

- Visual grounding, dense scene captioning, question answering



Dense Captioning

Question: Provide a detailed description of the appearance of <OBJ008> before analyzing its spatial connections with other elements in the scene.



Answer: This is a white toilet paper roll. It is to the left of the toilet. ✓

Ground Truth: There is a roll of toilet paper. It is on the wall to the left of the toilet.



Visual Grounding

Question: Are there any objects that correspond to the description "a water bottle is situated on the counter to the right of a refrigerator"? If yes, could you share the IDs for those objects?



Answer: Yes. <OBJ063> ✓

Ground Truth: <OBJ063>



Proxy3D: Tasks and Results

Visual-spatial intelligence

- Route planning, relational reasoning, physical agents in the 3D space



Route Plan

Question:



You are a robot beginning at the blue chair and facing the fan beside the boxes. You want to navigate to the trash bin. You will perform the following actions (Note: for each [please fill in], choose either 'turn back,' 'turn left,' or 'turn right.'): 1. Go forward until passing the tables on the right. 2. [please fill in] 3. Go forward until the trash bin. You have reached the final destination.

Answer:



A. Turn Left B. Turn Back C. Turn Right.



Object Relative Direction (Hard)

Question:



If I am standing by the tv and facing the bed, is the table to my front-left, front-right, back-left, or back-right? The directions refer to the quadrants of a Cartesian plane (if I am standing at the origin and facing along the positive y-axis)

Answer:



A. back-left
C. front-right

B. back-right
D. front-left



Proxy3D: Tasks and Results

Models	Vision	# of	Scan2Cap (val)		ScanRefer (val)			Multi3DRefer (val)		ScanQA (val)		SQA3D (val)
	modal.	tokens	C@0.5	B-4@0.5	Uni@0.5	Acc@0.25	Acc@0.5	F ₁ @0.25	F ₁ @0.5	C	EM	EM
<i>Correspondence-based:</i>												
Qwen2-VL-7B [40]	I	8000	0	3.8	6.3	5.4	5.1	21.1	19.9	53.9	19.0	40.7
GPT4Scene (HDM) [‡] [31]	B,I	8000	86.3	40.6	83.7	62.6	57.0	64.5	59.8	90.9	28.2	57.4
Video-3D-LLM [54]	I	8000	<u>83.8</u>	<u>41.4</u>	78.3	58.1	51.7	58.0	52.7	<u>102.1</u>	<u>30.1</u>	58.6
Spatial-MLLM-4B [42]	I	3100	–	–	–	–	–	–	–	91.8	26.3	55.9
3DRS [18]	I	8000	86.1	41.6	77.9	62.9	56.1	<u>60.4</u>	54.9	104.8	30.3	<u>60.6</u>
<i>Representation-based:</i>												
LLaVA-3D [56]	D,I	3096	79.2	41.1	–	54.1	42.4	–	–	91.7	27.0	55.6
LEO-VL [17]	D,I	<u>750</u>	–	–	–	–	–	–	–	100.4	22.6	60.8
Proxy3D _{Qwen2.5-VL-7B}	I	700	73.3	34.7	84.0	<u>59.6</u>	<u>54.1</u>	62.0	57.5	93.6	25.2	57.5

Edge deployment benefits

- Latency reduction
- Lower memory usage

Table 3. **Latency** (sec) for batch size of 8 and 16. OOM in table refers to out-of-memory run.

Batch size	Vision Encoder		Cluster.		LLM		Latency Per Question	
	8	16	8	16	8	16	8	16
Qwen2.5-VL	9.4	16.2	–	–	9.3	OOM	2.25	OOM
Ours	9.4	16.2	2.9	5.6	1.1	1.5	1.75	1.5

Proxy3D: Tasks and Results

Specialized 3D VLMs beat proprietary models with the lower complexity

Models	Numerical Answer				Multiple-Choice Answer				Avg. Rank	
	Obj. Cnt.	Abs. Dist.	Obj. Size	Room Size	Rel. Dist.	Rel. Dir.	Route Plan [‡]	Appr. Order [‡]		
Human level	94.3	47.0	60.4	45.9	94.7	95.8	95.8	100.0	79.2	-
<i>Proprietary via API:</i>										
GPT-4o [19]	46.2	5.3	43.8	38.2	37.0	41.3	31.5	28.5	34.0	8
Gemini-1.5 Pro [36]	56.2	30.9	64.1	43.6	51.3	46.3	36.0	34.6	45.4	3
<i>Open-source:</i>										
InternVL2-40B [10]	34.9	26.9	46.5	31.8	42.1	32.2	<u>34.0</u>	39.6	36.0	7
LLaVA-OV-72B [21]	43.5	23.9	57.6	37.5	<u>42.5</u>	39.9	32.5	44.6	40.2	5
LLaVA-Video-72B [53]	48.9	22.8	57.4	35.3	42.4	36.7	35.0	48.6	40.9	4
Qwen2.5-VL-7B [2]	40.9	14.8	43.4	10.7	38.6	38.5	33.0	29.8	33.0	9
Qwen2.5-VL-72B [2]	25.1	29.3	54.5	38.8	38.2	37.0	<u>34.0</u>	28.9	37.0	6
Spatial-MLLM-4B [42]	65.3	<u>34.8</u>	<u>63.1</u>	45.1	41.3	<u>46.2</u>	33.5	<u>46.3</u>	48.4	1
Proxy3D _{Qwen2.5-VL-7B}	<u>63.9</u>	41.9	67.2	<u>42.8</u>	50.3	46.5	31.4	32.0	<u>47.0</u>	<u>2</u>

Part 1:

[LLaVA-PruMerge: Adaptive Token Reduction for Efficient Large Multimodal Models](#)

[An Image is Worth 1/2 Tokens After Layer 2: Plug-and-Play Inference Acceleration for Large Vision-Language Models](#)

[SparseVLM: Visual Token Sparsification for Efficient Vision-Language Model Inference](#)

Part 2:

[Spatial-MLLM: Boosting MLLM Capabilities in Visual-based Spatial Intelligence](#)

[LEO-VL: Efficient Scene Representation for Scalable 3D Vision-Language Learning](#)

[Proxy3D: Efficient 3D Representations for Vision-Language Models via Semantic Clustering and Alignment](#)

Conclusions

Applications using Vision-Language Models

- UI agents with the universal information-sparse inputs
- 3D VLMs for spatial reasoning in dense streaming videos
- All these cases are compute/memory “expensive” for edge devices

Part 1. Visual token pruning for VLMs

- We leveraged spatial and temporal redundancies in images
- [SparseVLM](#) with the prompt guidance and attention-based scores

Part 2. Visual-spatial intelligence using 3D VLMs

- 3D VLMs implement perception for physical AI agents
- [Proxy3D](#) clustering with the distillation compresses token length by $10\times$

Questions?



Denis Gudovskiy
Distinguished AI Engineer

