

Token-level Response-visual Attention Guidance for Multimodal LLMs Knowledge Distillation

Jaehyun Jang¹, Eunseop Yoon¹, Hee Suk Yoon¹, SooHwan Eom¹, Mark A. Hasegawa-Johnson², and Chang D. Yoo¹
¹Korea Advanced Institute of Science and Technology (KAIST), ²University of Illinois Urbana-Champaign (UIUC)



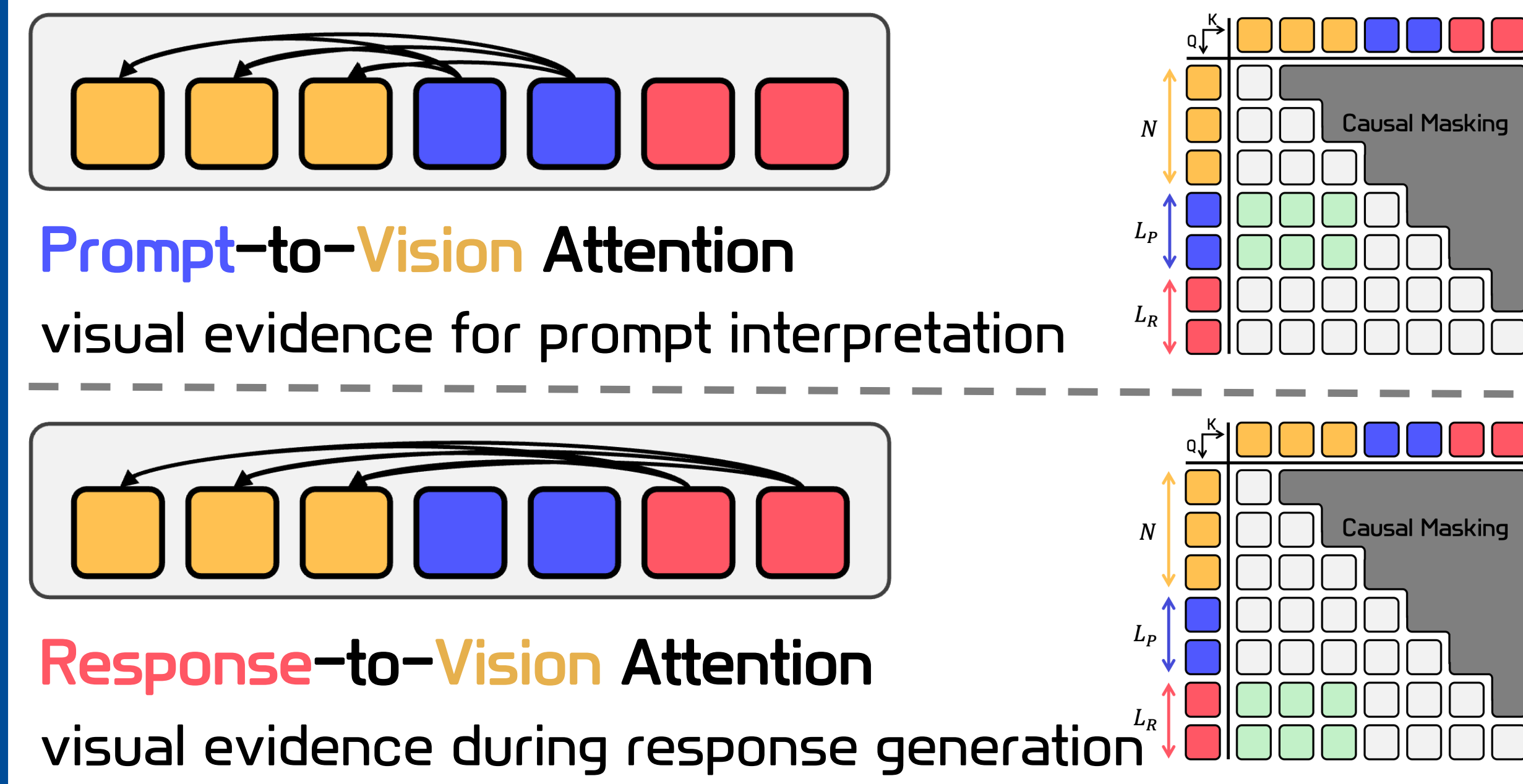
Paper



Video

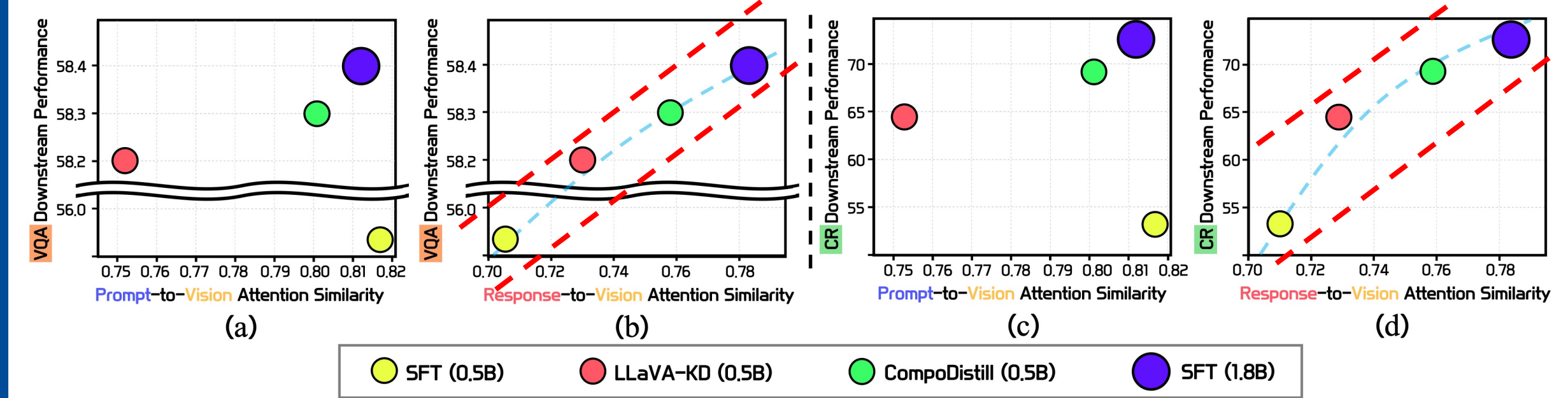
Takeaway) Distill **visual** grounding during **response generation**, **token by token**, with guidance adapted to how **diffuse** or **focused** the teacher's attention is.

Cross-modal Attention Extraction



Observations

Response-to-Vision is the more informative signal



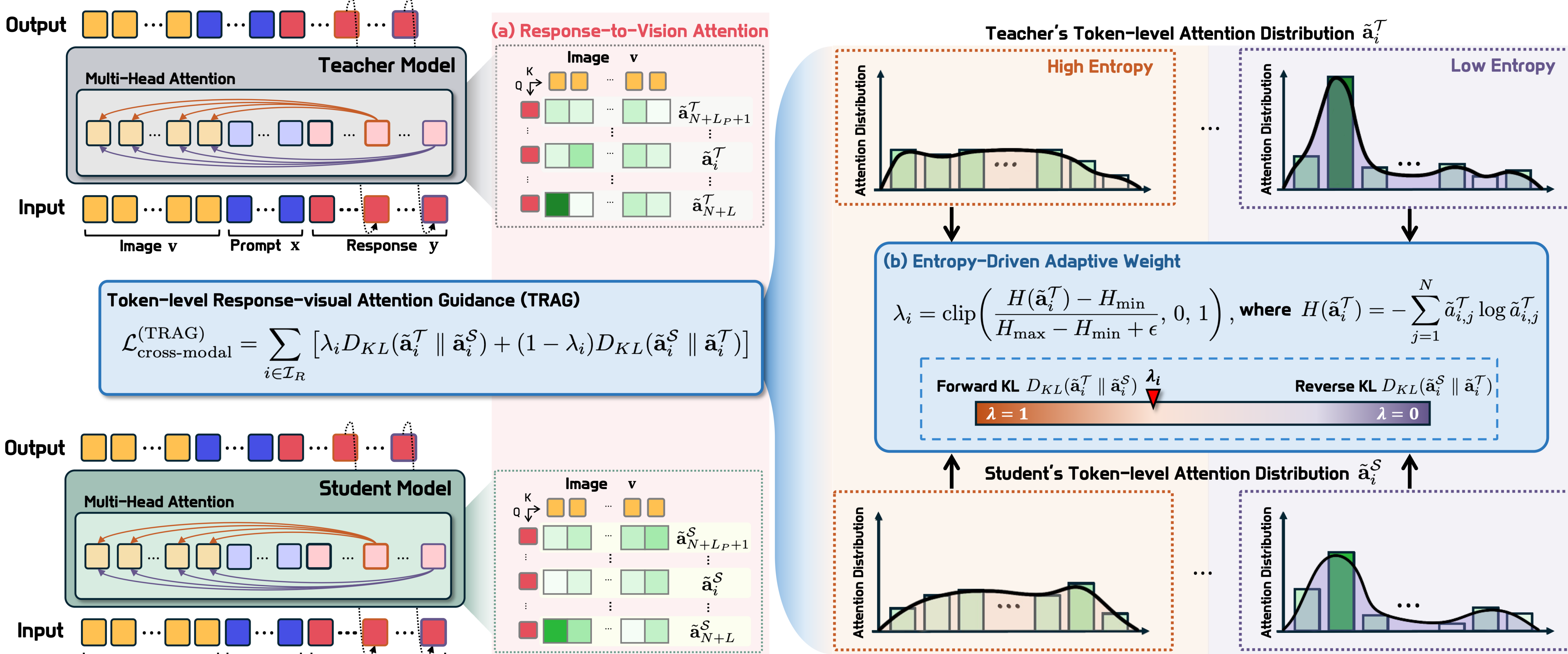
Prompt-to-Vision → **Weak** performance correlation
Response-to-Vision → **Strong** performance correlation
 → ① **Supervise the response phase**

Visual grounding varies across response tokens
Cross-modal interactions peak at intermediate depths



Across adjacent layers → **Redundant** attention
 Across response tokens → **Diverse** visual grounding
 → ② **Aggregate layers**, ③ **supervise each response token**

TRAG: Token-level Response-visual Attention Guidance



① Shifting Supervision to the Response Phase

$$\mathcal{L}_{\text{cross-modal}}^{(\text{CompoDistill})} = \sum_{l \in \mathcal{M}_S} \mathcal{D}_{\text{cos}}(A_{G_l}^T(I), A_l^S(I))$$

$I = I_P, I_P = \{N+1, \dots, N+L_P\}$

$\mathcal{L}_{\text{cross-modal}}^{(\text{prior})} = \sum_{l \in \mathcal{M}_S} \mathcal{D}_{\text{cos}}(A_{G_l}^T(I_P), A_l^S(I_P))$

Query range shift!
 $I_P \Rightarrow I_R$

$I = I_R, I_R = \{N+L_P+1, \dots, N+L\}$

$\mathcal{L}_{\text{cross-modal}}^{(\text{response})} = \sum_{l \in \mathcal{M}_S} \mathcal{D}_{\text{cos}}(A_{G_l}^T(I_R), A_l^S(I_R))$

② From Layer-wise Matching to Token-level Distribution

Aggregate intermediate-layer attention for each response token

$$\{a_{l,i}^S\}_{l \in \mathcal{M}_S} \rightarrow \tilde{a}_i^S = \frac{1}{|\mathcal{M}_S|} \sum_{l \in \mathcal{M}_S} a_{l,i}^S$$

$$\{a_{l,i}^T\}_{l \in \mathcal{M}_T} \rightarrow \tilde{a}_i^T = \frac{1}{|\mathcal{M}_T|} \sum_{l \in \mathcal{M}_T} a_{l,i}^T$$

Normalize each token's attention vector into a probability distribution

$$\tilde{a}_i^S := \frac{\tilde{a}_i^S}{\sum_{j=1}^N \tilde{a}_{i,j}^S + \epsilon}$$

$$\tilde{a}_i^T := \frac{\tilde{a}_i^T}{\sum_{j=1}^N \tilde{a}_{i,j}^T + \epsilon}$$

Define token-level guidance over normalized attention distributions

$$\mathcal{L}_{\text{cross-modal}}^{(\text{token-level})} = \sum_{i \in I_R} \mathcal{D}_{\text{guidance}}(\tilde{a}_i^T, \tilde{a}_i^S)$$

③ Entropy-Driven Adaptive Weighting

$$\mathcal{L}_{\text{cross-modal}}^{(\text{token-level})} = \sum_{i \in I_R} \mathcal{D}_{\text{guidance}}(\tilde{a}_i^T, \tilde{a}_i^S)$$

$$\lambda_i = \text{clip} \left(\frac{H(\tilde{a}_i^T) - H_{\min}}{H_{\max} - H_{\min} + \epsilon}, 0, 1 \right), \text{ where } H(\tilde{a}_i^T) = - \sum_{j=1}^N \tilde{a}_{i,j}^T \log \tilde{a}_{i,j}^T$$

$\lambda_i \uparrow$ **High Entropy**
 - Diffuse visual evidence
 - Forward KL for coverage

$\lambda_i \downarrow$ **Low Entropy**
 - Localized visual evidence
 - Reverse KL for precision

$D_{KL}(\tilde{a}_i^T || \tilde{a}_i^S)$ $D_{KL}(\tilde{a}_i^S || \tilde{a}_i^T)$

$$\mathcal{L}_{\text{cross-modal}}^{(\text{TRAG})} = \sum_{i \in I_R} [\lambda_i D_{KL}(\tilde{a}_i^T || \tilde{a}_i^S) + (1 - \lambda_i) D_{KL}(\tilde{a}_i^S || \tilde{a}_i^T)]$$

Main Results

General Visual Question Answering Benchmarks

Size	Method	LLM	# Samples	GQA	ScienceQA	TextVQA	MME	MMB	MMB ^{cn}	POPE	MMMU	Avg	Avg ₅
≤ 4B	TinyLLaVA ¹	Qwen1.5-4B	1.2 M	62.7	70.2	60.0	70.5	67.5	67.2	86.9	38.6	69.6	65.4
	TinyLLaVA ¹	Qwen1.5-1.8B	1.2 M	60.9	65.2	47.7	61.3	57.8	56.2	83.4	34.8	62.7	58.4
	KD ¹	Qwen1.5-1.8B	1.2 M	60.9	64.6	53.3	66.3	65.8	63.1	86.5	32.7	66.2	61.6
≤ 2B	CompoDistill ¹	Qwen1.5-1.8B	1.2 M	61.2	66.5	53.5	67.0	64.5	63.0	85.5	34.1	66.4	61.9
	LLaVA-KD	Qwen1.5-1.8B	1.2 M	62.3	64.7	53.4	69.1	64.0	63.7	86.3	33.6	66.6	62.1
	TRAG	Qwen1.5-1.8B	1.2 M	61.7	66.9	55.6	67.0	66.6	63.3	86.8	35.2	67.4	62.9
≤ 0.5B	TinyLLaVA ¹	Qwen1.5-0.5B	1.2 M	56.7	60.9	46.6	59.4	55.2	51.8	83.9	31.9	60.4	55.8
	KD ¹	Qwen1.5-0.5B	1.2 M	56.8	60.4	46.1	58.9	54.7	49.7	84.9	32.3	60.3	55.5
	CompoDistill ¹	Qwen1.5-0.5B	1.2 M	59.1	61.1	48.2	63.8	59.6	54.9	85.7	34.0	62.9	58.3
	LLaVA-KD	Qwen1.5-0.5B	1.2 M	59.6	60.6	49.9	64.5	60.1	55.5	85.9	30.2	63.4	58.2
	TRAG	Qwen1.5-0.5B	1.2 M	59.5	62.1	48.9	65.2	60.9	54.6	86.5	30.9	63.8	58.5

Across backbones & scales → **Consistent gains**

Compositional Reasoning Benchmarks

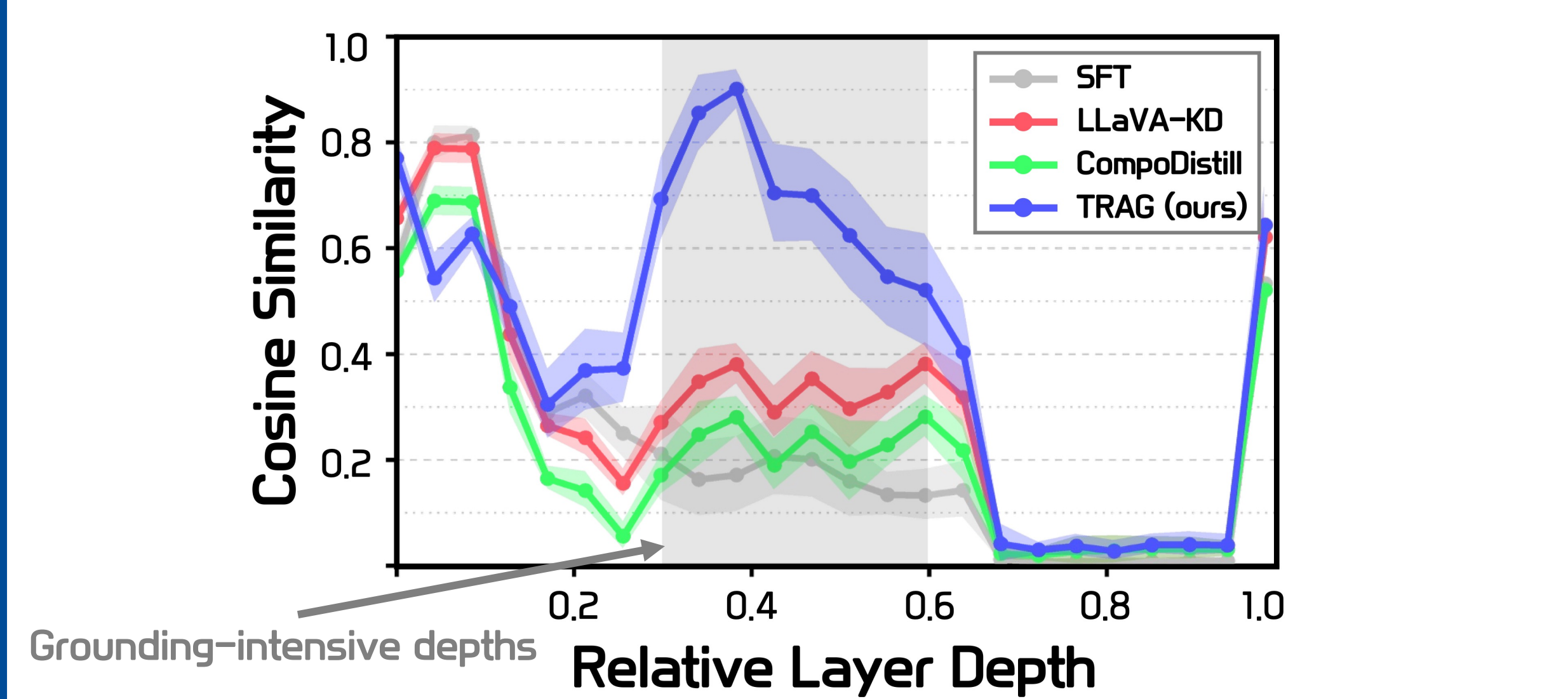
↳ Directly probe fine-grained visual grounding

LLM	Method	Size	Compositional Reasoning Benchmarks			Avg
			SugarCrep	BiVLC	Winground	
Qwen1.5	TinyLLaVA	4B	87.3	93.2	70.1	83.5
	TinyLLaVA		74.8	83.8	62.6	73.7
	KD		76.1	85.3	61.5	74.3
	LLaVA-KD	1.8B	76.6	85.5	60.9	74.3
	CompoDistill		81.7	86.5	62.2	76.8
	TRAG		83.1	89.1	67.2	79.8
	TinyLLaVA		52.8	56.5	50.2	53.2
	KD		51.0	50.1	52.1	51.0
	LLaVA-KD	0.5B	66.9	74.5	51.7	64.4
	CompoDistill		72.1	78.6	56.6	69.1
TRAG		73.4	80.2	57.3	70.3	

TRAG → **Up to +3.0 Avg points** over prior baselines

Response-to-Vision Attention Fidelity

Teacher-Student R→V Attention Cosine Similarity



30-60% Depth → TRAG ≈ **0.72**

SFT ≈ 0.18 | LLaVA-KD ≈ 0.32 | CompoDistill ≈ 0.41

→ **Stronger response-time grounding fidelity**

Check out the paper for more experiments and analyses!