

The Road to Convergence: Evolution of Unified Multimodal Models

Past, Present, and Future



Jindong Wang
Assistant Professor,
William & Mary.



Hao Chen
Ph.D. from Carnegie
Mellon University.



Sharon Li
Associate Professor,
UW-Madison



Zhaolong Su
Cornell University
P.h.D. Student.



Jiakui Hu
Peking University
P.h.D. Student.

***Speaker of the day**

Tutorial Roadmap

Part 1	Past — From Divergence to Unification	30 min
Part 2	Present I — Unified model: Goals, Modeling, and Future	55 min
	<i>Coffee Break</i>	15 min
Part 3	Present II — Frontier Highlights & Trend	45 min
Part 4	Beyond Unified	25 min
Part 5	Future	25 min

P A R T 1

Past

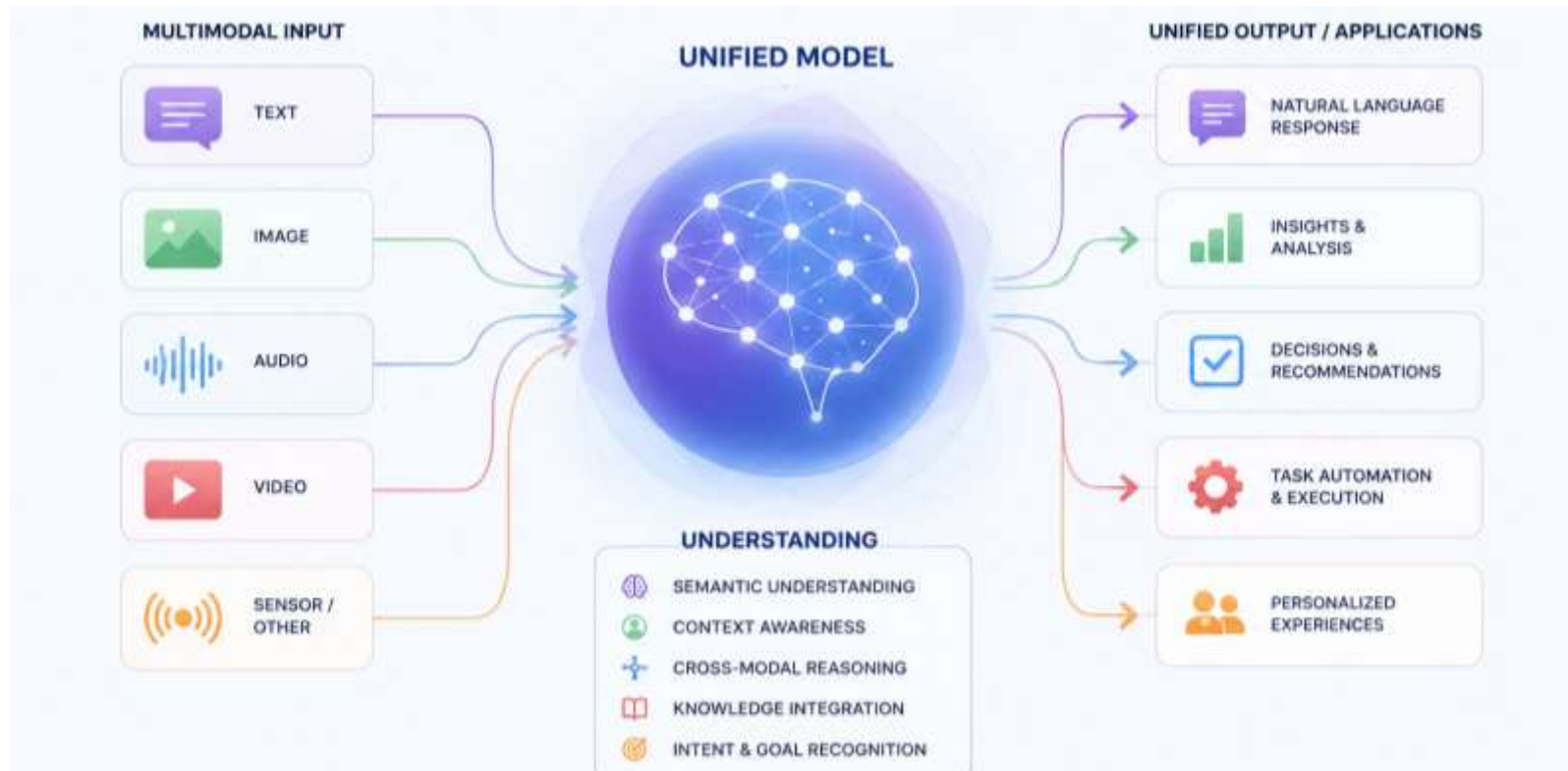
From Divergence to Unification

- Chameleon
- Nano-Banana
- Janus-Pro
- GPT4o
- Clip

Unified or Not?

- Chameleon ✓ Meta 2024.05
- Nano-Banana ✓ Google
- Janus-Pro ✓ DeepSeek-AI 2025.01
- GPT4o ✓ OpenAI
- Clip ✗ OpenAI 2021.02

Definition of Unified Models



Two Independent Tracks

What Is Multimodal Understanding?

Multimodal understanding is the task of producing an output that captures the semantic content of a *multimodal* input

Two Independent Tracks

$$P(y \mid x_v, x_t) = \prod_{i=1}^T P(y_i \mid y_{<i}, x_v, x_t; \theta)$$

What Is Multimodal Understanding?

Multimodal understanding is the task of producing an output that captures the semantic content of a *multimodal* input

Two Independent Tracks

Understanding tasks include:

- **Captioning:** summarize visual semantics
- **VQA:** answer grounded questions
- **Grounding / segmentation:** locate referred entities
- **Relation reasoning:** infer spatial, causal, and semantic relations
- **Verification:** check whether outputs satisfy instructions

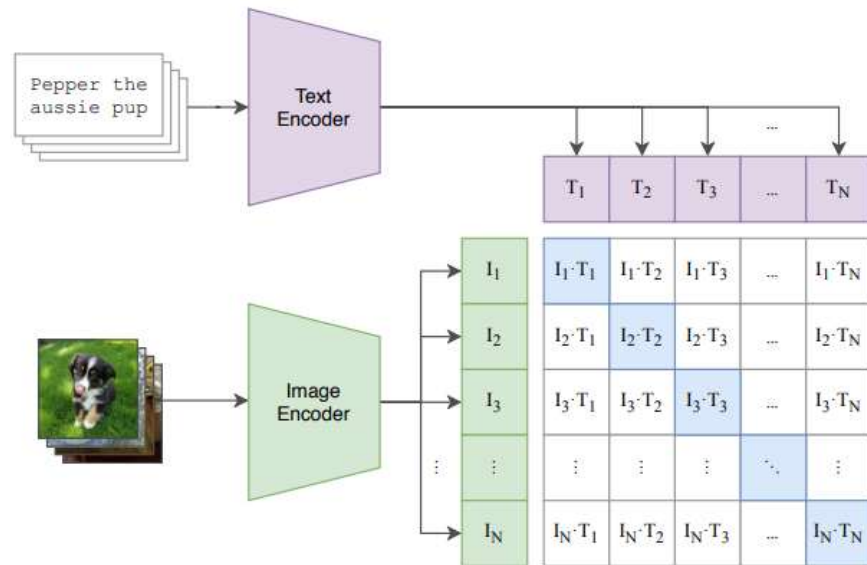
What Is Multimodal Understanding?

Two Independent Tracks

UNDERSTANDING

Clip

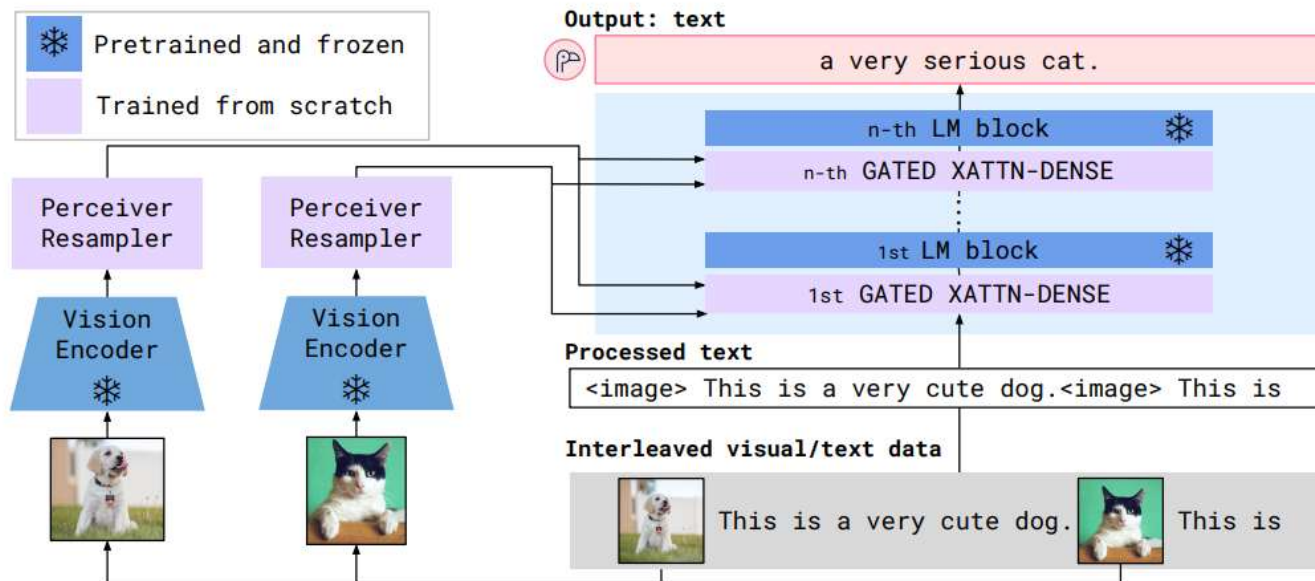
(1) Contrastive pre-training



Two Independent Tracks

UNDERSTANDING

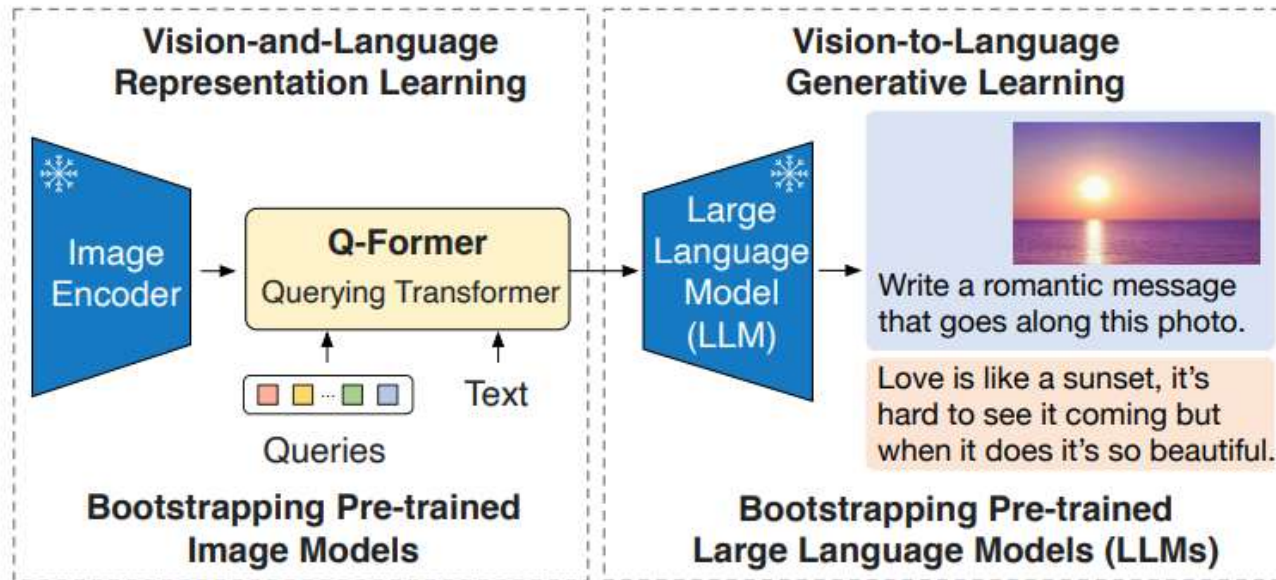
Flamingo



Two Independent Tracks

UNDERSTANDING

BLIP2



Two Independent Tracks

GENERATION

$$y_{\mathcal{M}_{out}} \sim p_{\theta}(y_{\mathcal{M}_{out}} \mid x_{\mathcal{M}_{in}}, c)$$

What Is Multimodal Generation?

Multimodal generation is the task of producing a *visual* (or audio/video) output that realizes the semantic content of a *textual* (or multimodal) input, typically a natural language prompt.

Two Independent Tracks

GENERATION

DDPM and DALL·E

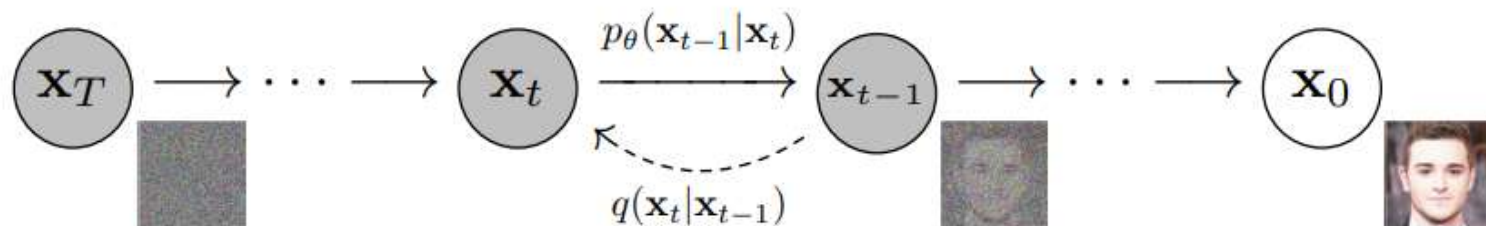


Figure 2: The directed graphical model considered in this work.

Two Independent Tracks

UNDERSTANDING

CLIP → BLIP-2 → Flamingo → LLaVA

Contrastive / instruction-tuned
encoders

Discriminative objective

Tokenizer: text-only LLM vocab

GENERATION

**GAN → DDPM → Imagen / SDXL /
DALL·E 3**

Iterative denoising/score matching

Generative objective

Tokenizer: VAE latents (continuous)

Inspiration came from NLP

Before T5:

Translation / Summarization / QA / Classification
→ different formats, heads, and objectives

Inspiration came from NLP

Before T5:

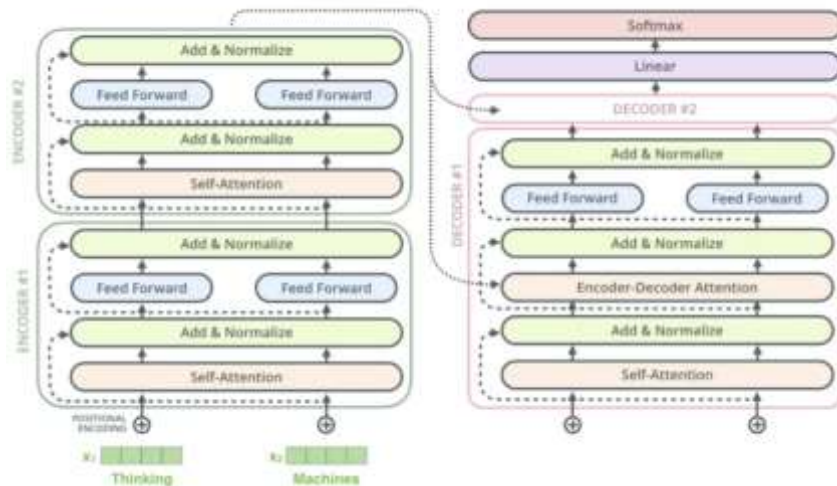
Translation / Summarization / QA /
Classification

→ different formats, heads, and objectives

T5:

Everything becomes:

Text input → Text output



1. Pipeline Systems Are Not Enough



Some kind of VLM



User: Generate something similar

1. Pipeline Systems Are Not Enough



Some kind of VLM



a young elf girl with
long white hair,
pointed ears, green
eyes, wearing a black
mage outfit with
white cape...

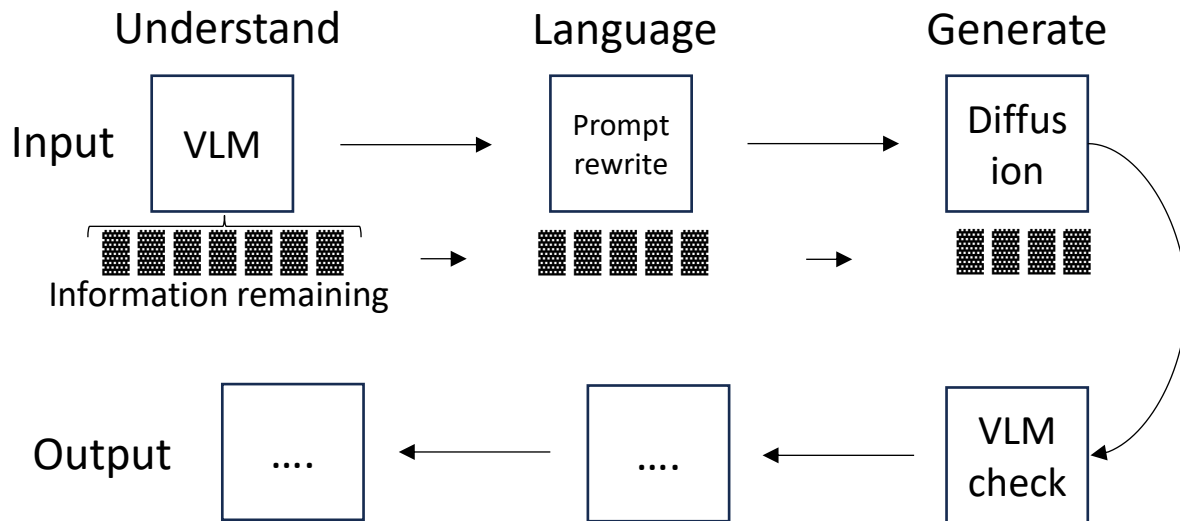
Some Generative
Model, e.g., diffusion



Output

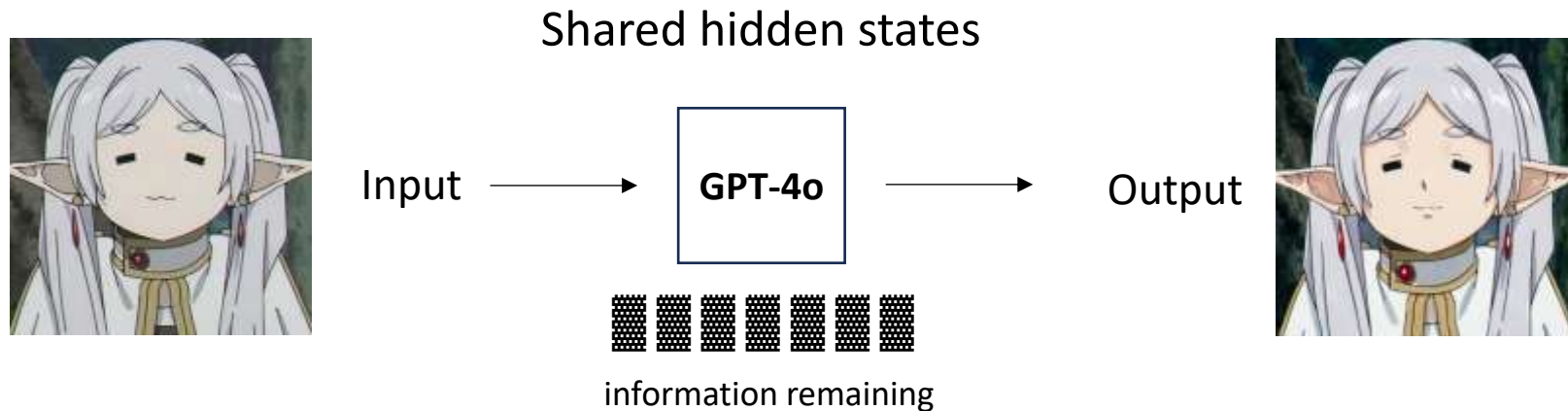
User: Generate something similar

1. Pipeline Systems Are Not Enough



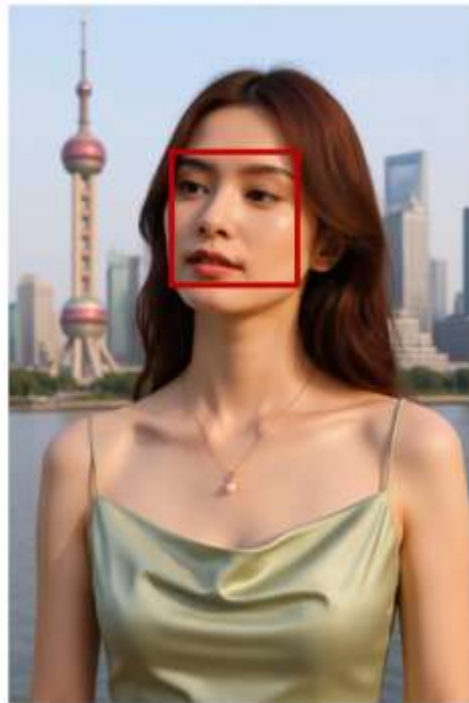
Each step come with information loss ❌

1. Pipeline Systems Are Not Enough



Unified Model ✓

2. Unlock Tasks



2. Unlock Tasks

Example: Multi-round editing



User: change the
hair color to blonde

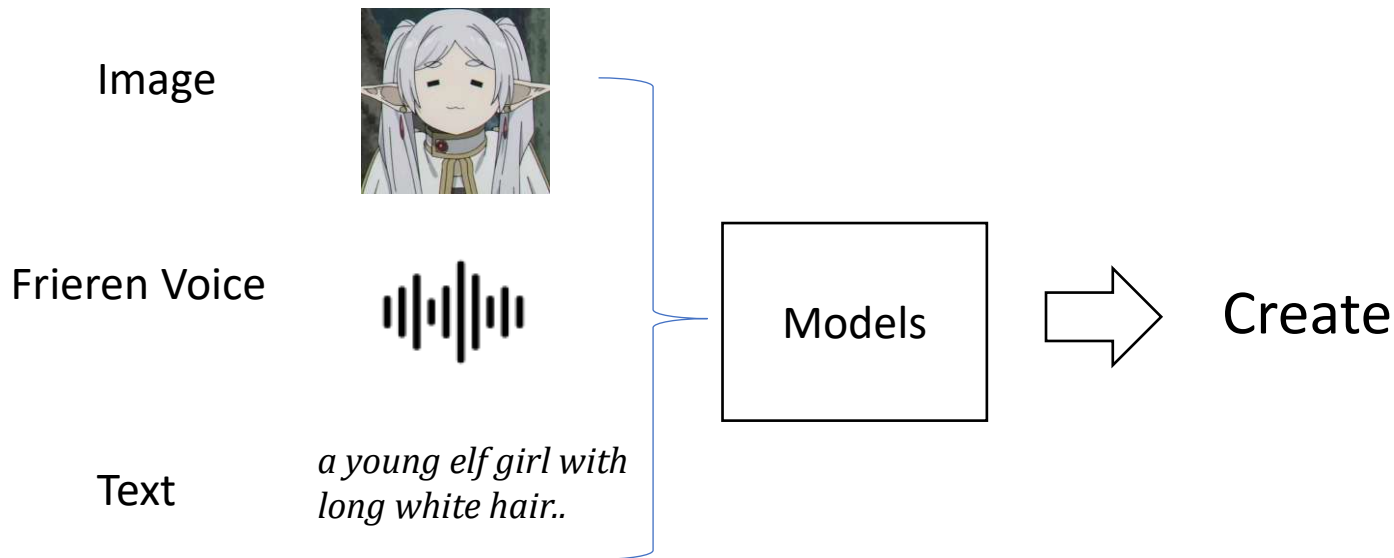


User: change the
hair color to blue



User:...

3. Potential human-like multimodal intelligence



Visions of Unified Model community

1. Eliminate the representational gap among perception, language, generation, and action.
2. Unlock more applications & tasks
3. From function calls to cognitive level unification

"What I cannot create, I do not understand."

Example: Nano Banana



Example: Nano Banana



PART 2

Present I

Unified model: Goals, Modeling, and Future

Jiakui Hu

jkhu29@stu.pku.edu.cn

Peking University

[1] Unified Multimodal Understanding and Generation Models: Advances, Challenges, and Opportunities. arXiv 2025.

Samples of GPT4o

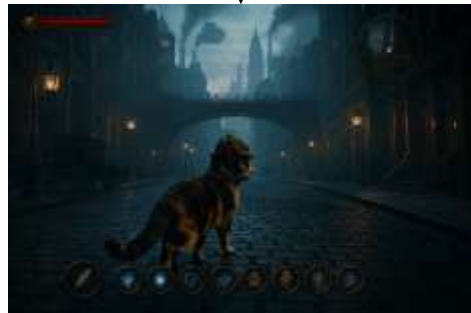
Give this cat a detective hat and a monocle



Update to a landscape image 16:9 ratio, add more spells in the UI



Turn this into a triple A video games



Samples of GPT4o

We need evidence there is a currently present invisible elephant.



The goals for unified model

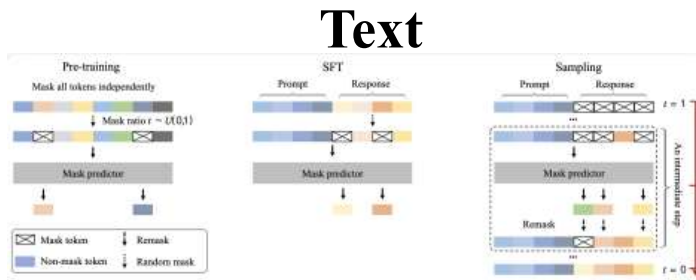
Interleaved text-vision modeling.



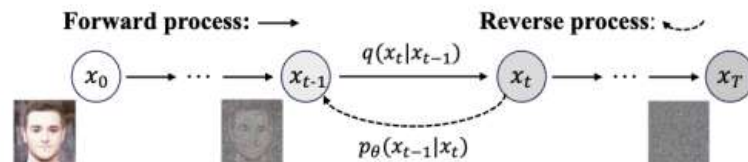
Unifying understanding, reasoning, and generation tasks.

Jointly modeling text and vision

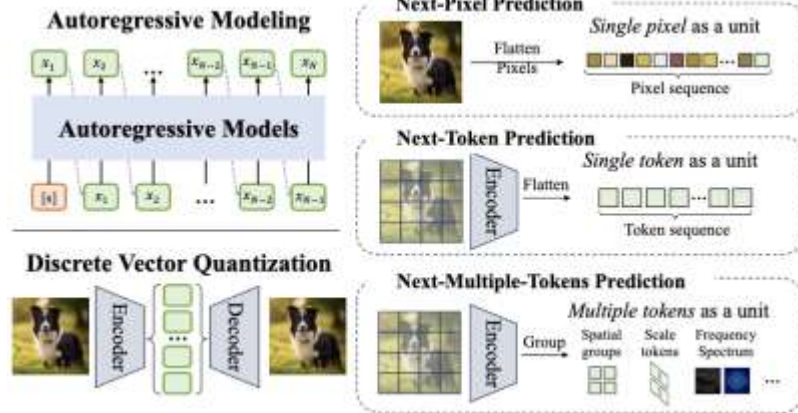
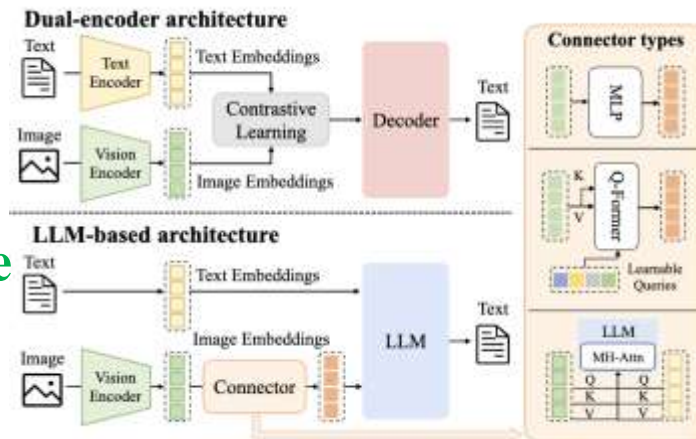
Diffusion



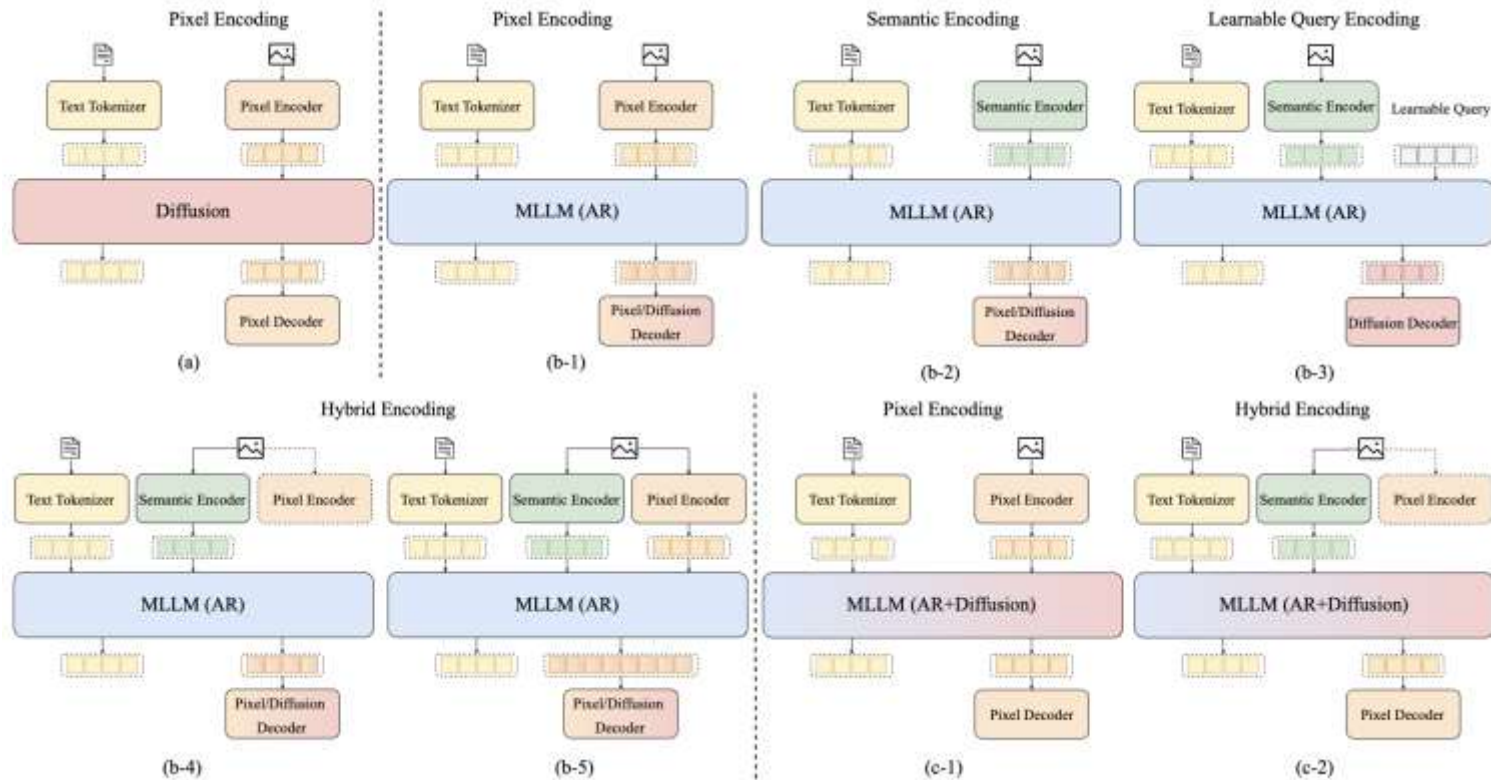
Vision



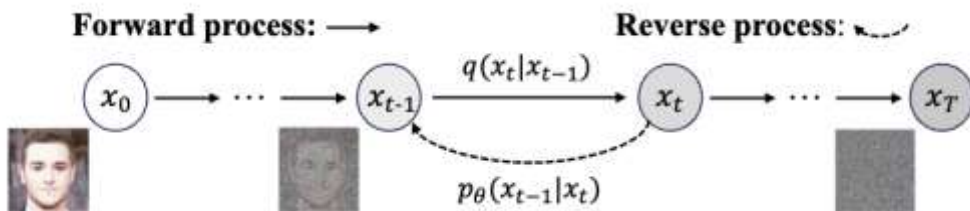
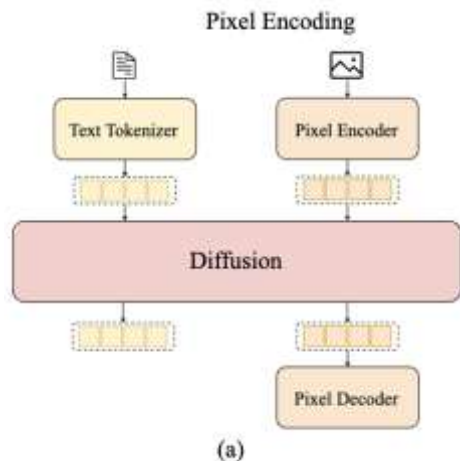
Auto
Regressive



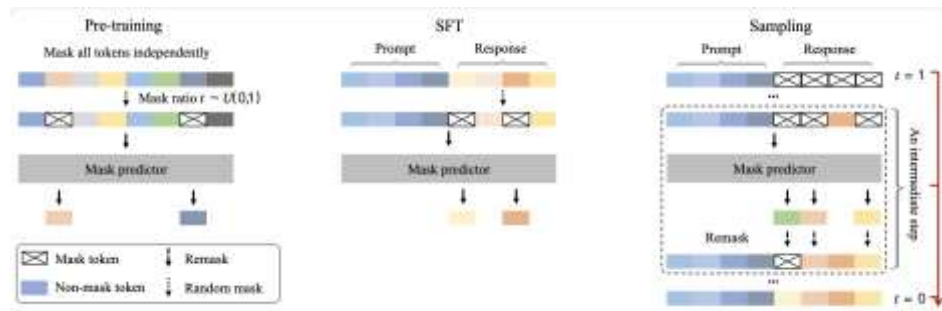
Permutations and combinations



Option A: Diffusion



Diffusion based image modeling: **denoising** from Gaussian noise



Diffusion based Image modeling: **demasking** from All-[mask] tokens.

[3] Denoising Diffusion Probabilistic Models. NeurIPS 2020.

[4] Large Language Diffusion Models. NeurIPS 2025.

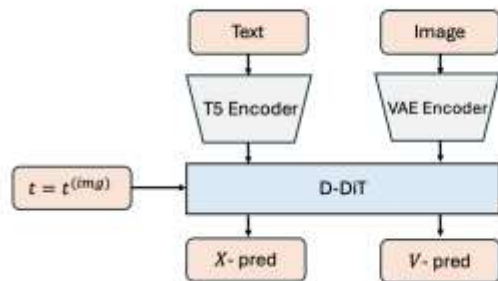
Diffusion based text modeling



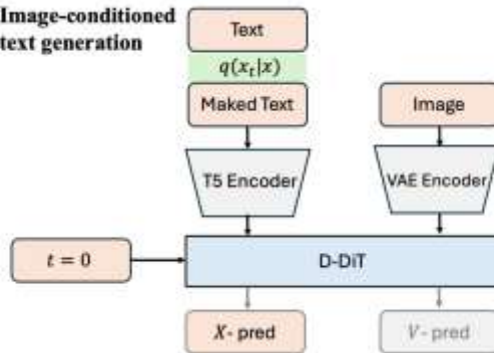
Q: Provide a brief description of the given image. A: [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]
[MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK] [MASK]

Examples: continuous visual tokens

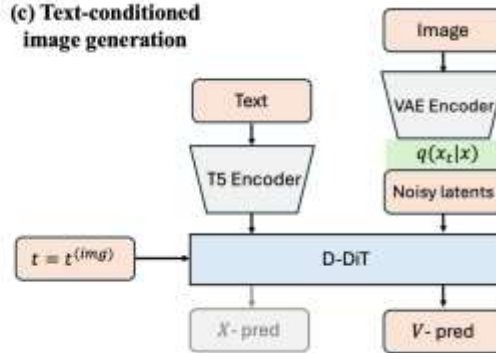
(a) Overview



(b) Image-conditioned text generation



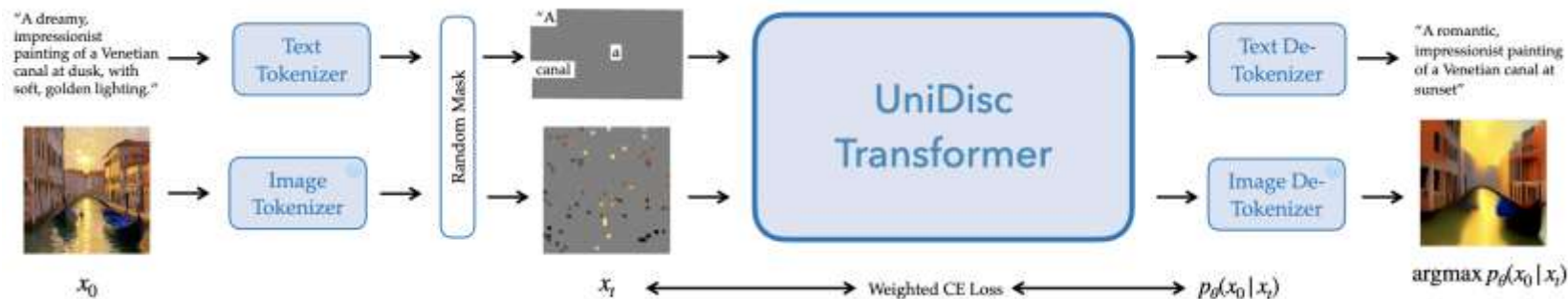
(c) Text-conditioned image generation



Based on an image generation model,

- Text generator uses masked modeling;
- Image generator uses denoising;
- If text generator need vision input, the vision input is noise-free.

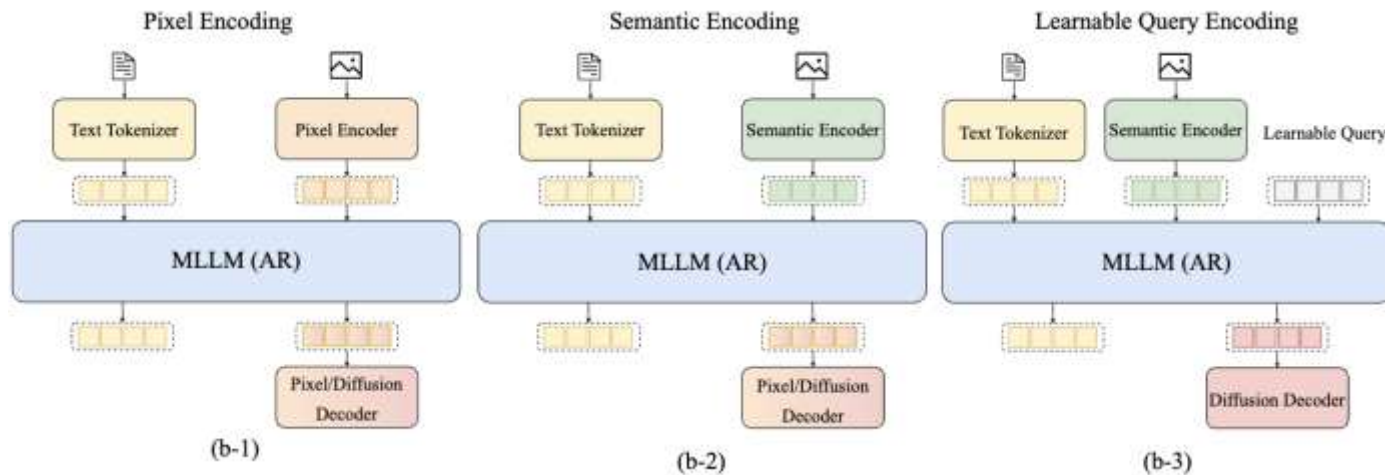
Examples: discrete visual tokens



Limitations of Option A: Diffusion

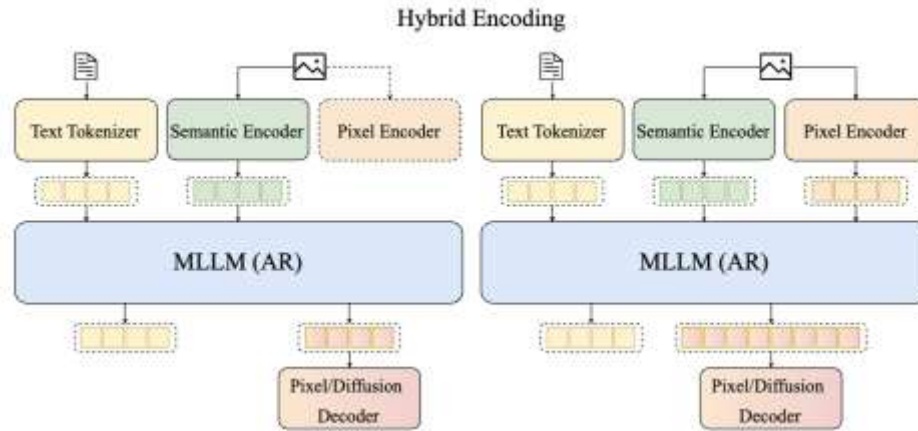
1. In **image generation**, it is still worse than continuous-space generative models;
2. Hard to perform flexible **interleaved** multimodal generation;
3. Training and inference **codebases** are still underdeveloped.

Option B: AutoRegressive Model



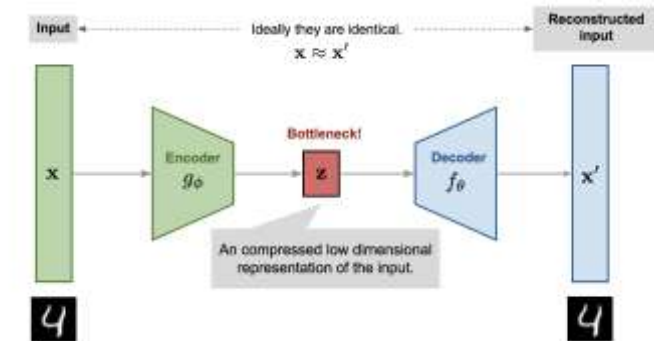
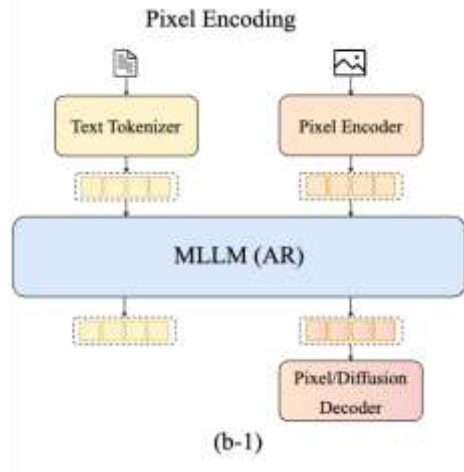
- **B-1.** Image encoder focus on **reconstruction**. Good at texture, bad at semantic;
- **B-2.** Image encoder finetuned on **semantic** encoder. Good at semantic, bad at texture;
- **B-3.** Image tokens learned by **learnable query**. Adaptive, trade-off.

Option B: AutoRegressive Model

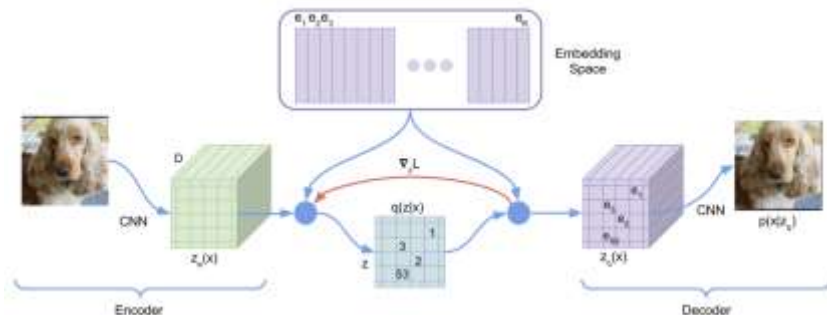


- **B-4. Fuse** reconstruction encoder and semantic encoder.

Option B-1: Reconstruction encoder



The encoder is to represent the image using a **compressed, low-dimensional** representation.

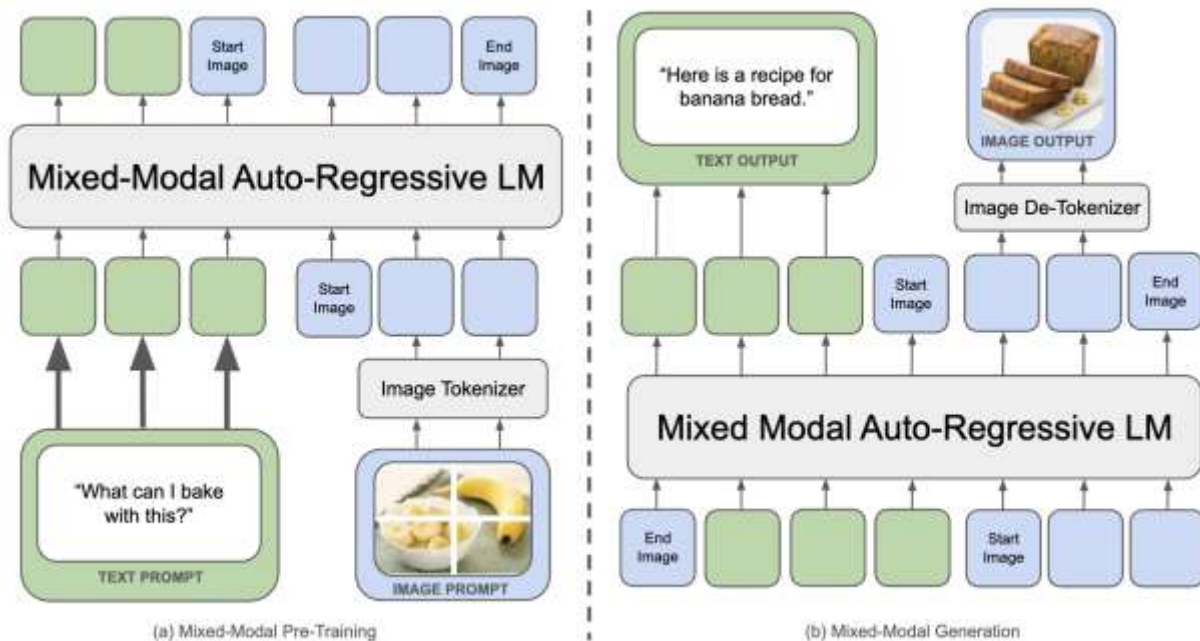


Convert the visual representation into **indices**.

[7] <https://lilianweng.github.io/posts/2018-08-12-vae/>

[8] Neural Discrete Representation Learning. NeurIPS 2017.

Examples



Interleaved demo

Prompt: I'm bored. Could you show me some cool, quirky-looking birds? Also tell me about them in 2-3 sentences.



Response: Absolutely! Here are three birds with some unique features that you might find interesting:

Keel-Billed Toucan The Keel-Billed Toucan is a brightly colored bird native to South American rainforests and woodlands. Its most distinctive feature is its colorful beak, which can grow up to seven inches long and has a distinctive, curved shape. The beak is used for plucking fruit from trees and for social displays. Despite its large size, the Keel-Billed Toucan is a skilled flyer and can often be seen fitting through the treetops. Here is an image of this vibrant bird: 

Puffin Puffins are small, aquatic birds that are found in the northern hemisphere. They have distinctive black and white feathers, as well as a brightly colored beak that is used for attracting mates and defending their territory. Puffins are known for their distinctive "puffin" call, which sounds like a loud, guttural chuckle. They feed on fish, which they catch by diving into the water from the air. Puffins are colonial birds, and they breed in large colonies on rocky cliffs and islands. Here is an image of this adorable bird: 

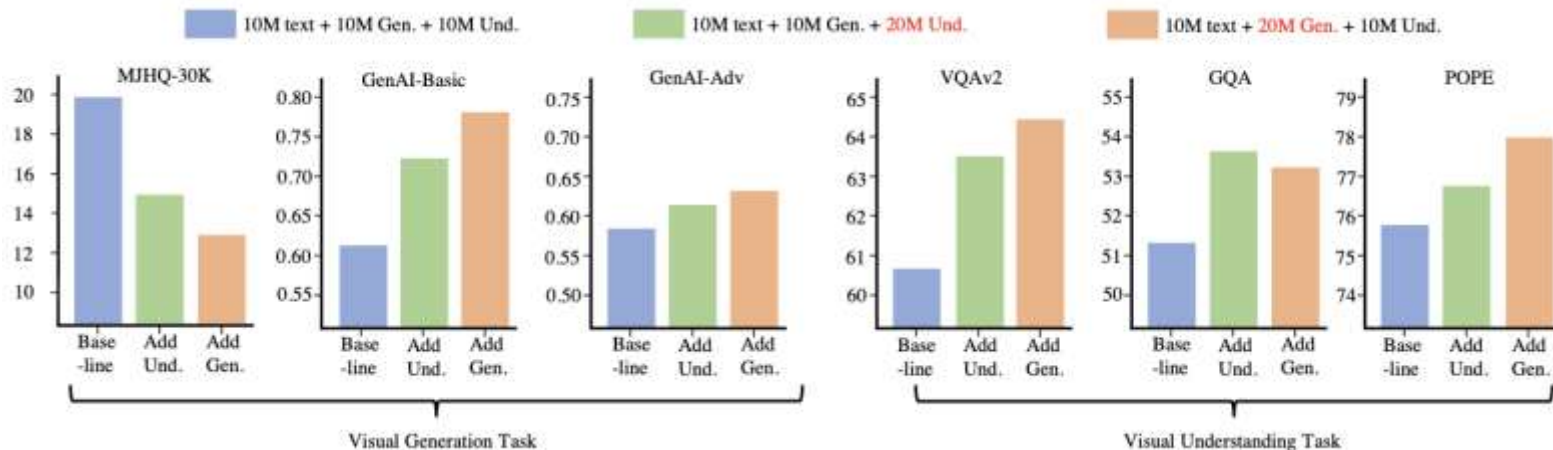


Golden Pheasant The Golden Pheasant is a brightly colored bird native to the forests of China. It is a member of the pheasant family and is known for its striking plumage, which features a golden-yellow body, red face, and green tail feathers. The Golden Pheasant is a ground-dwelling bird, and it feeds on a diet of insects, seeds, and fruits. It is a popular game bird, and is often hunted for its meat and feathers. In captivity, the Golden Pheasant is sometimes kept as an ornamental bird, due to its striking appearance. Here is an image of this beautiful bird: 

I hope these images and descriptions have helped to brighten your day! If you have any other questions or requests, please don't hesitate to ask.

Findings: simultaneous improvements

Method	Training Data			Visual Generation			Visual Understanding					
	Text-only	Visual Gen.	Visual Und.	gFID↓	Basic	Advanced	Overall	VQAv2	GQA	TextVQA	POPE	MME
Baseline	10M	10M	10M	19.9	0.63	0.58	0.60	60.7	51.3	39.2	75.9	909.6
Add Gen.	10M	20M	10M	12.8	0.78	0.63	0.69	64.5	53.2	39.8	78.0	1066.8
Add Und.	10M	10M	20M	14.9	0.73	0.62	0.66	63.5	53.7	40.5	76.8	1035.1

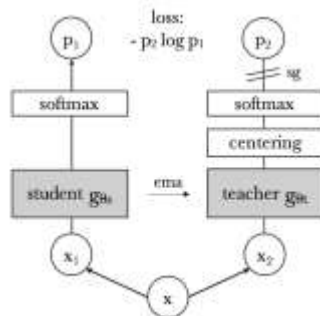
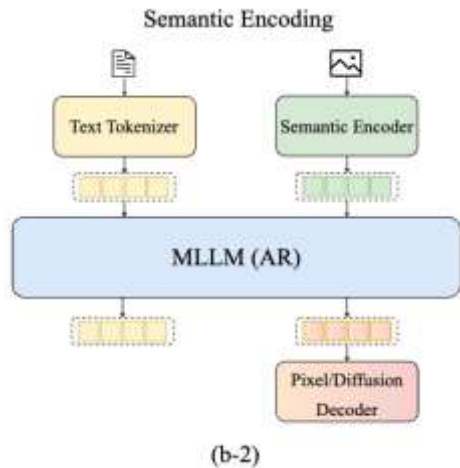


Limitations of Option B-1

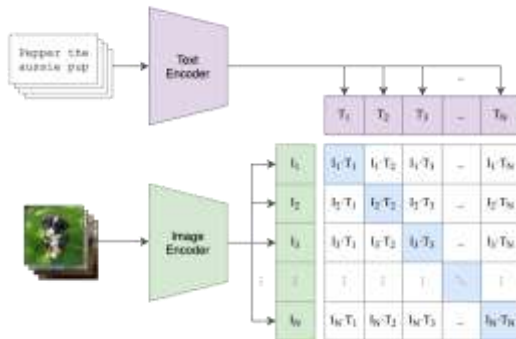
It lacks

- **high-level** semantic abstraction
- **cross-modal alignment** between text and image representations

Option B-2: Semantic Encoding



high-level semantic

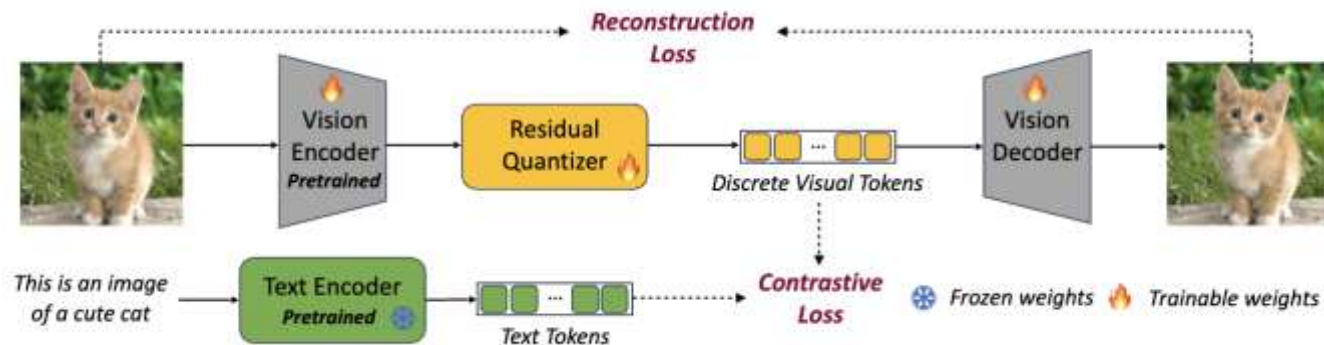


alignment
between text and image

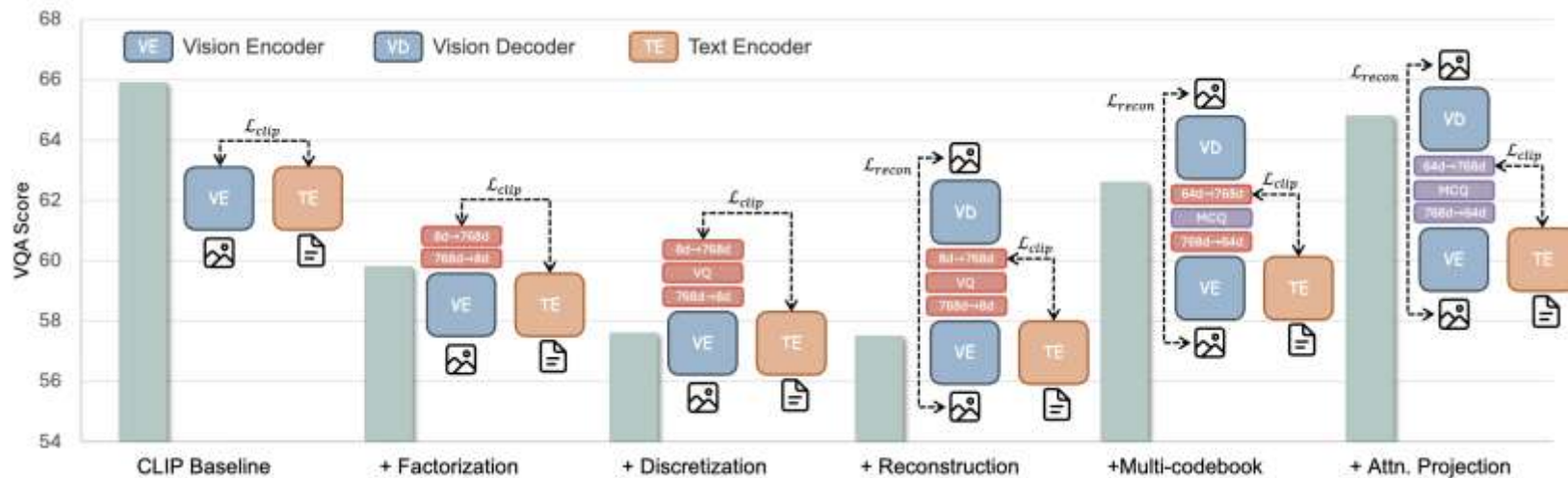
[11] Emerging Properties in Self-Supervised Vision Transformers. ICCV 2021.

[12] Learning Transferable Visual Models From Natural Language Supervision. ICML 2021.

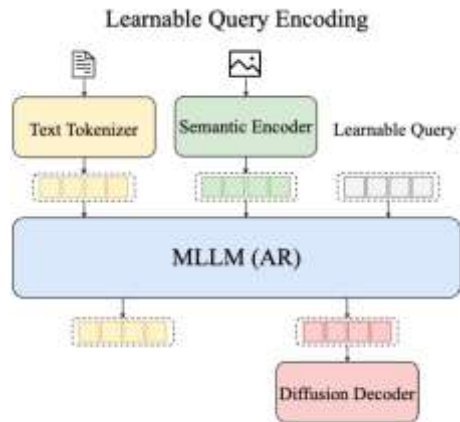
Examples: convert



Findings: conflict



Option B-3: Learnable Query Encoding

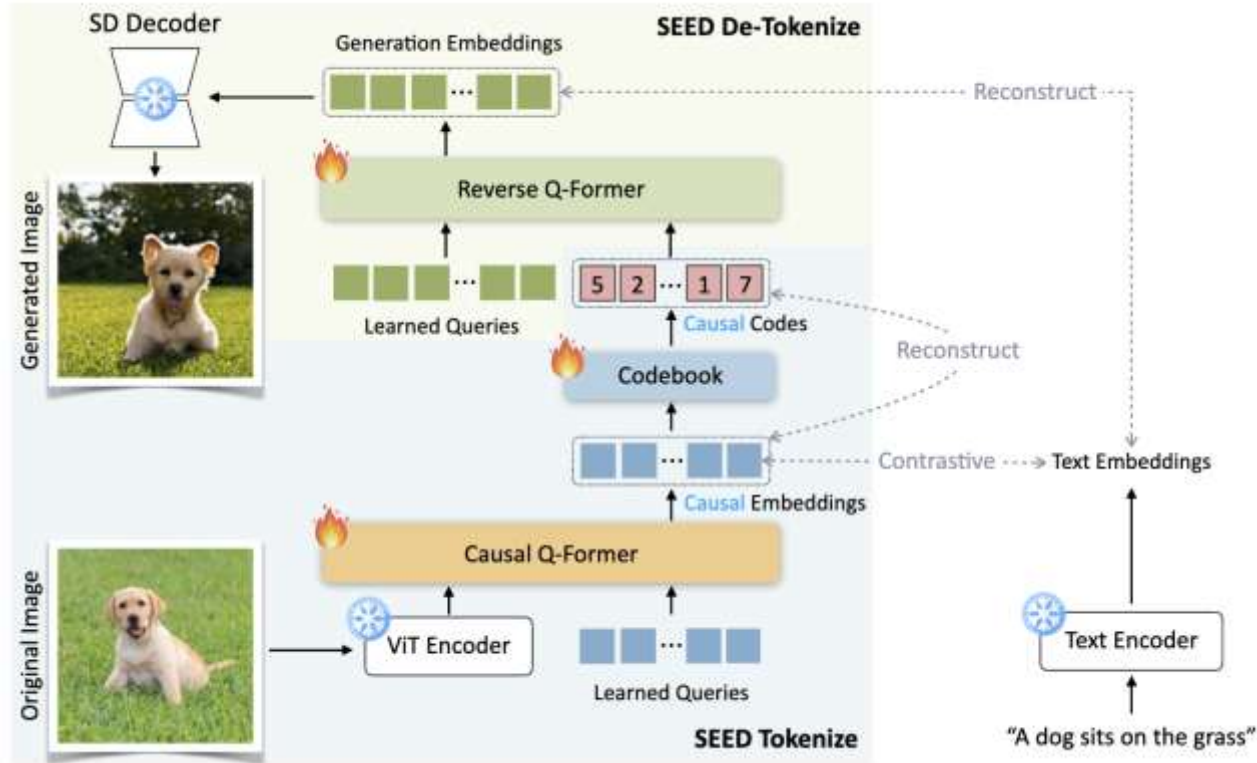


(b-3)

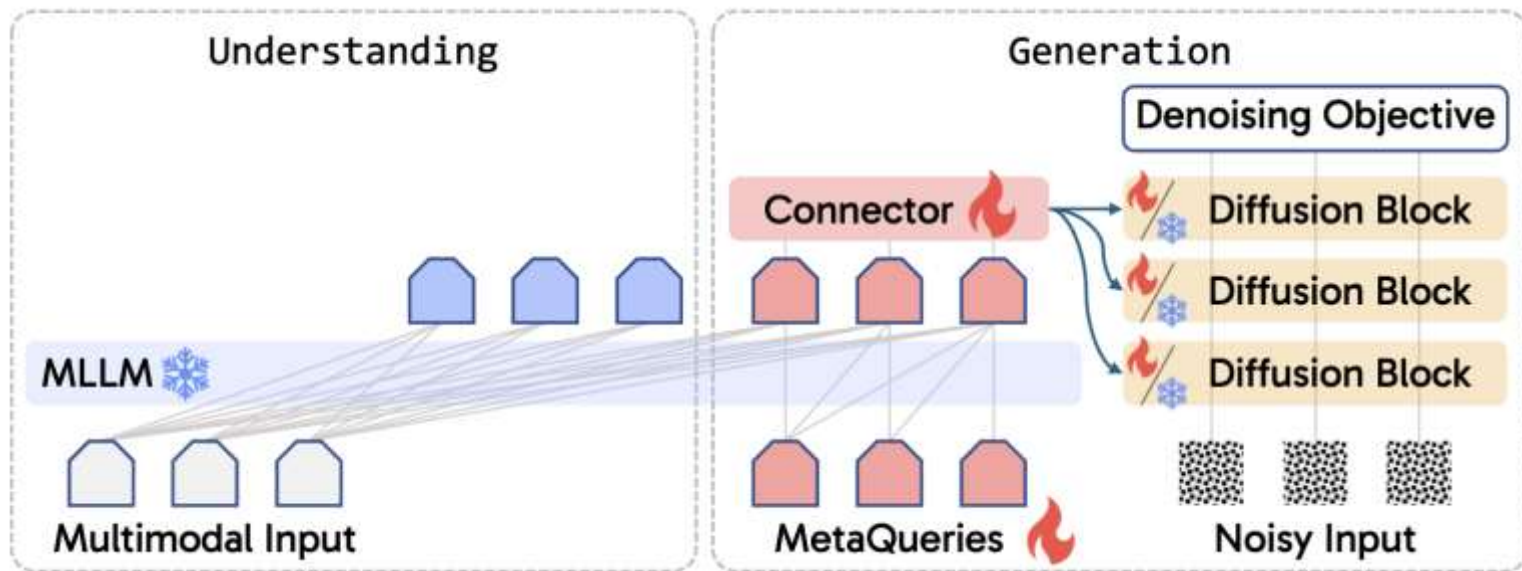
$$Z = \text{softmax} \left(\frac{QV^{\top}}{\sqrt{d}} \right) V \quad (1)$$

where $Q \in \mathbb{R}^{M \times d}$ denotes M learnable query tokens, $V \in \mathbb{R}^{N \times d}$ denotes N visual tokens extracted from the image, d is the embedding dimension, and $Z \in \mathbb{R}^{M \times d}$ denotes the compact semantic representations obtained by querying the visual features.

Examples: tokenizer level queries



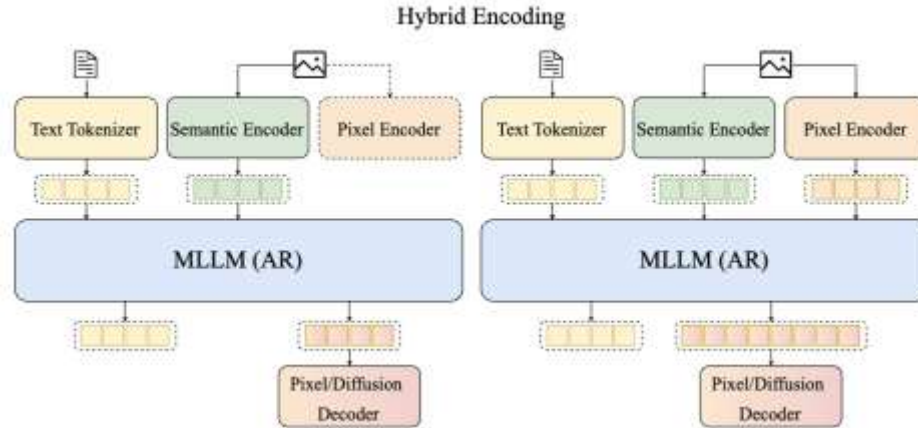
Examples: MLLM level queries



Limitations of Option B-3

- It transforms visual generation into a downstream application of visual understanding, is **not “unified”**.
- Querying may **fail**.

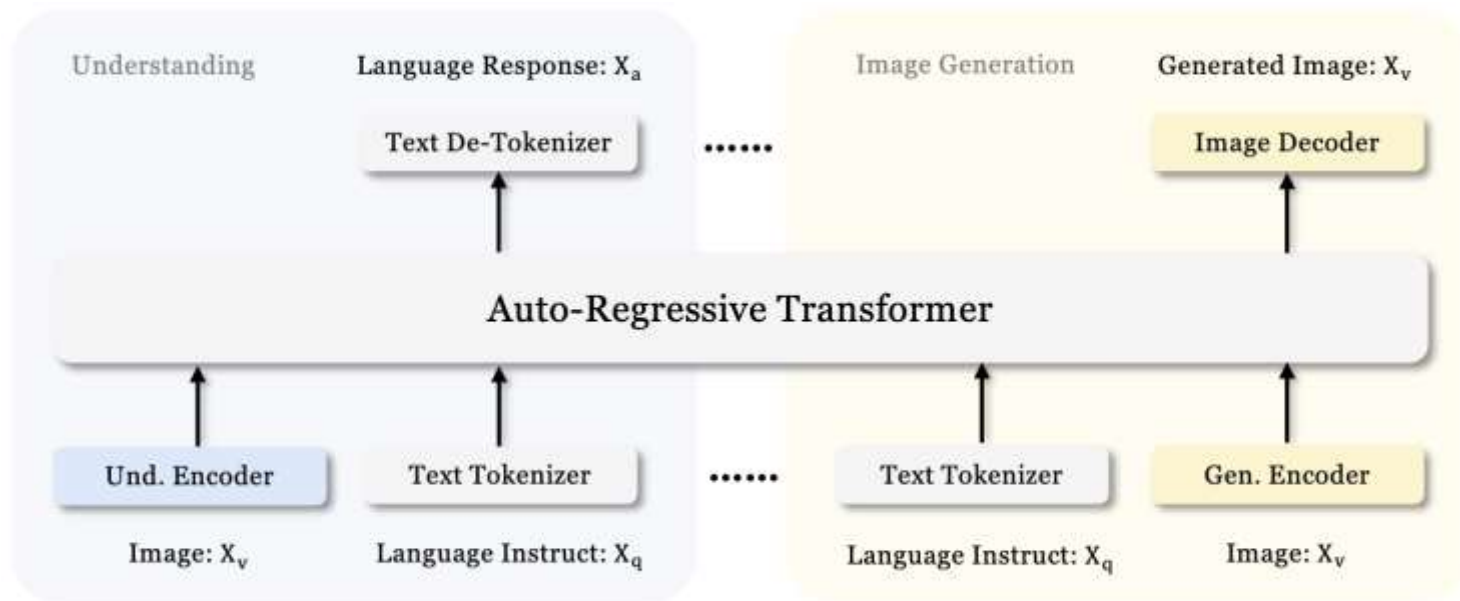
Option B-4: Hybrid encoder



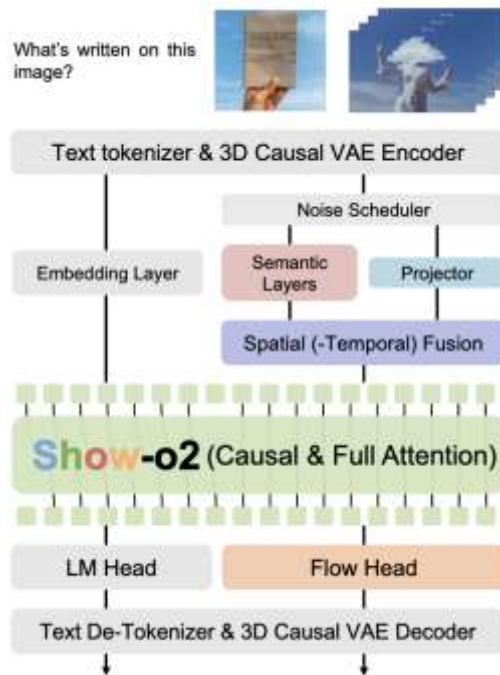
Using **multiple** visual representations

- Pseudo Hybrid Encoding
- Joint Hybrid Encoding

Examples: Pseudo Hybrid Encoding



Examples: Joint Hybrid Encoding



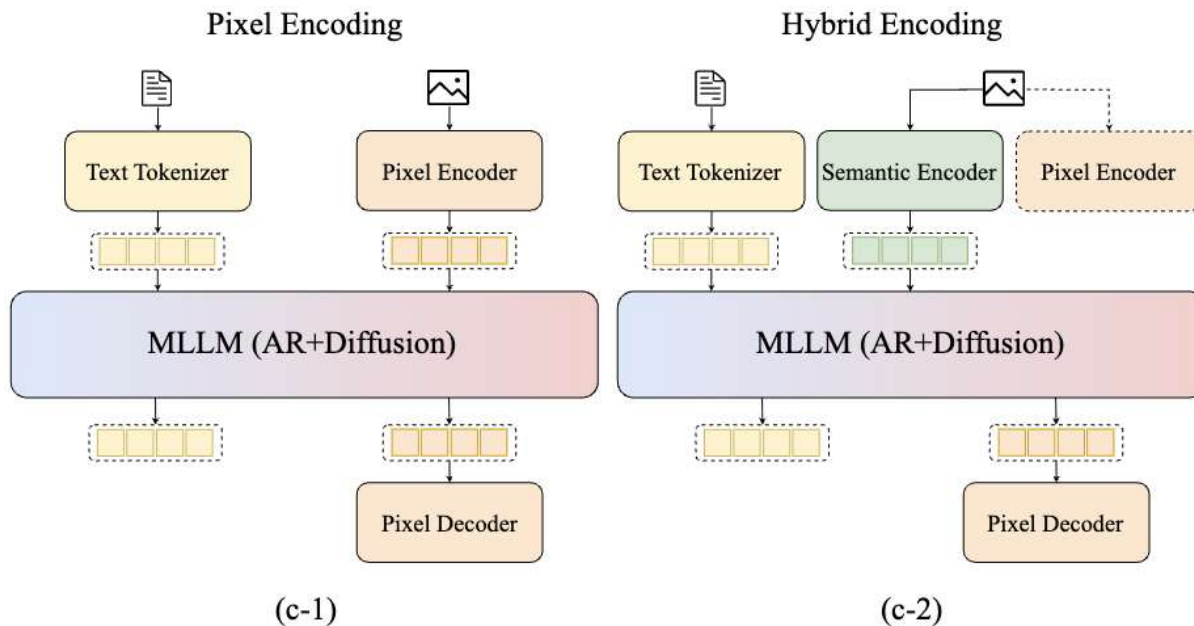
Limitations of Option B-4

Mismatch between reconstruction encoder and semantic encoder

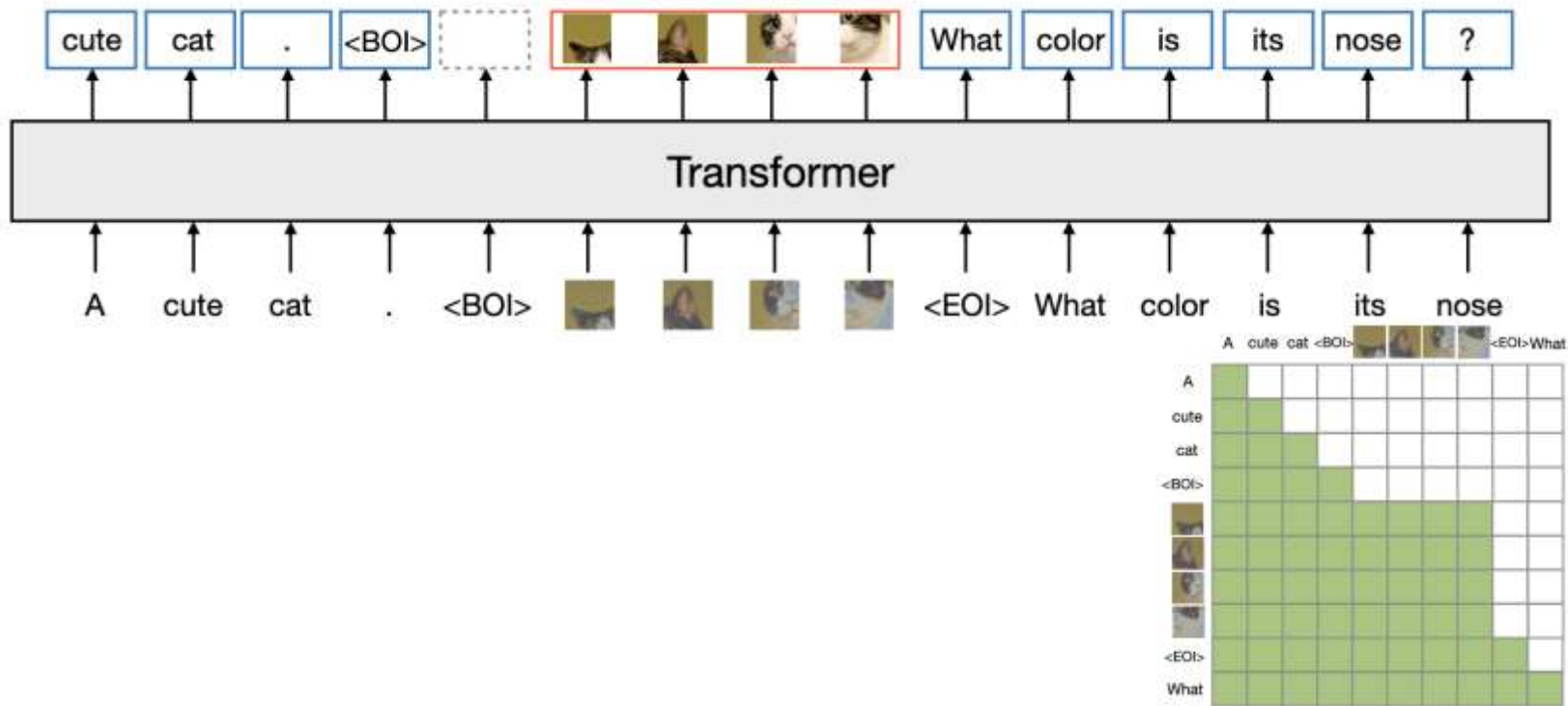


Conflicting optimization

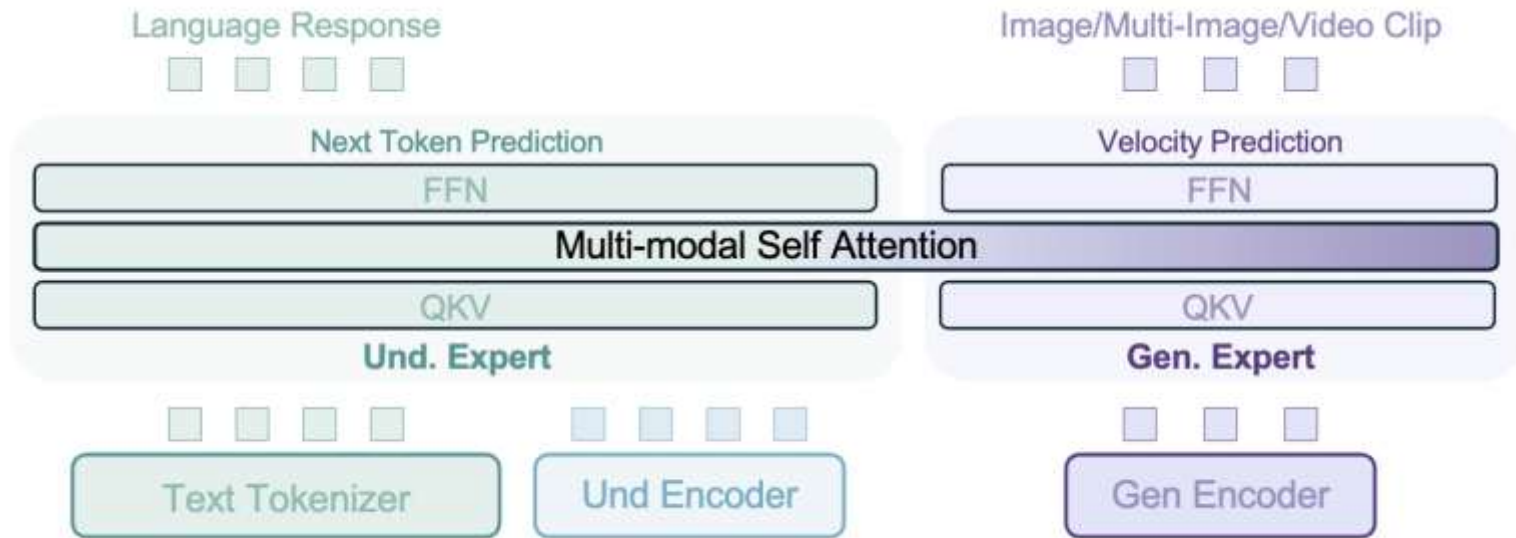
Option C: AutoRegressive & Diffusion



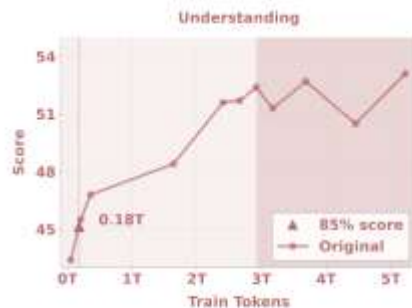
Examples: Reconstruction encoder



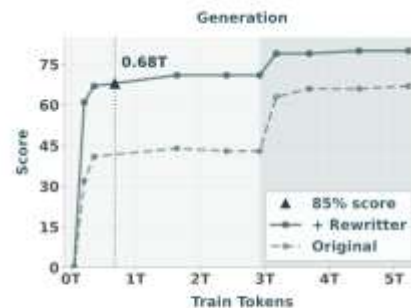
Examples



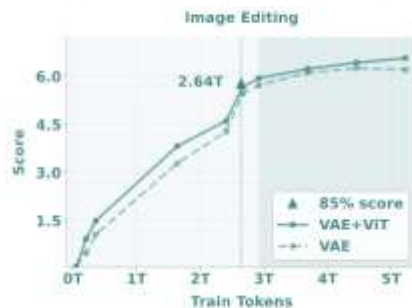
Findings: Emerging



(a) Average score on Image Understanding tasks.



(b) GenEval score on Image Generation task.



(c) GEdit Overall Score on classical Image Editing task.



(d) IntelligentBench Score on Intelligent Editing task.

Higher performance, but ...

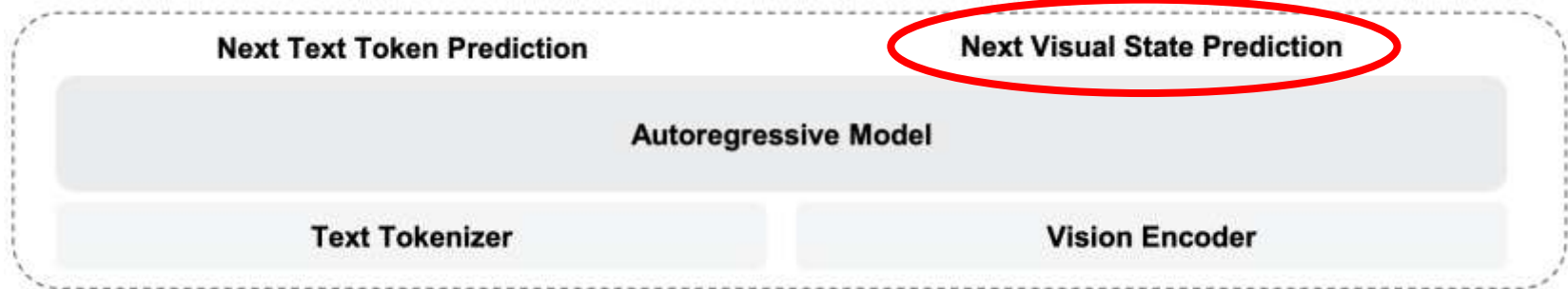
Redundant encoder-decode steps in interleaved multimodal modeling.

Visual representations **are not suited** for unified multimodal modeling.

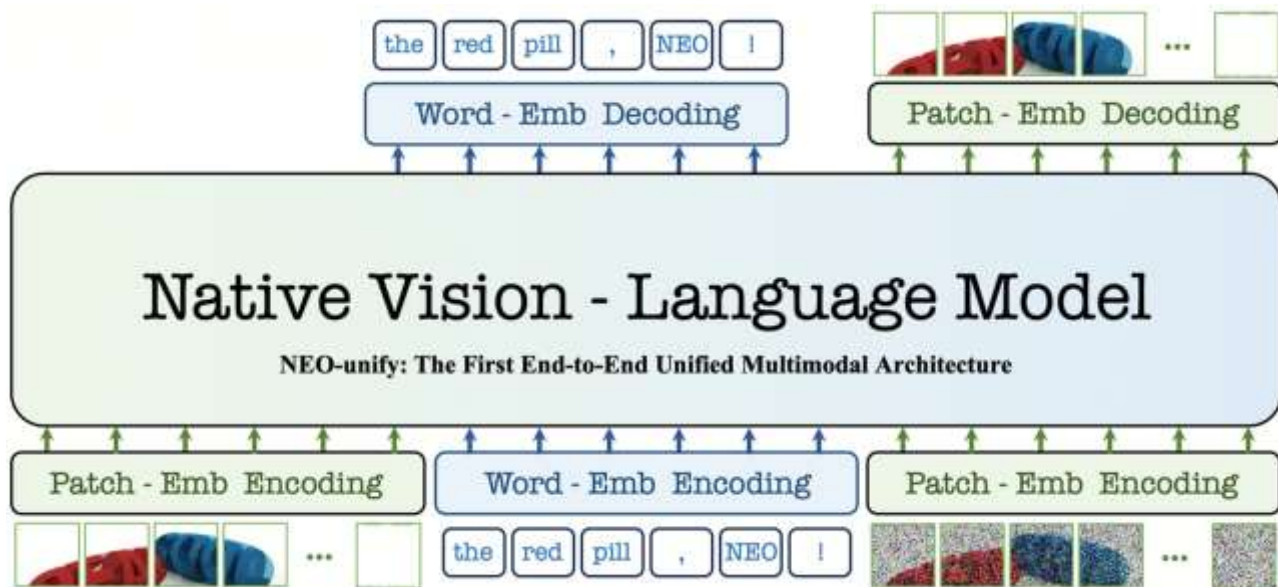
What would an ideal visual representation for unified multimodal models look like?

Attempt 1: Semantic Encoding

World model!



Attempt 2: no encoder



Coffee Break

15 minutes

We resume with Part 3 — Frontier Highlights

P A R T 3

Present II

latest trend

Beyond Representation — New Capability Frontiers

- Toward Bottleneck-Free Representations
- interleaved / any-to-any generation
- reasoning-centric generation / test-time scaling
- Cognitive Unification Native Multimodal
- Post-training and unified reward models
- Evaluate

Frontier I — Toward Bottleneck-Free Representations

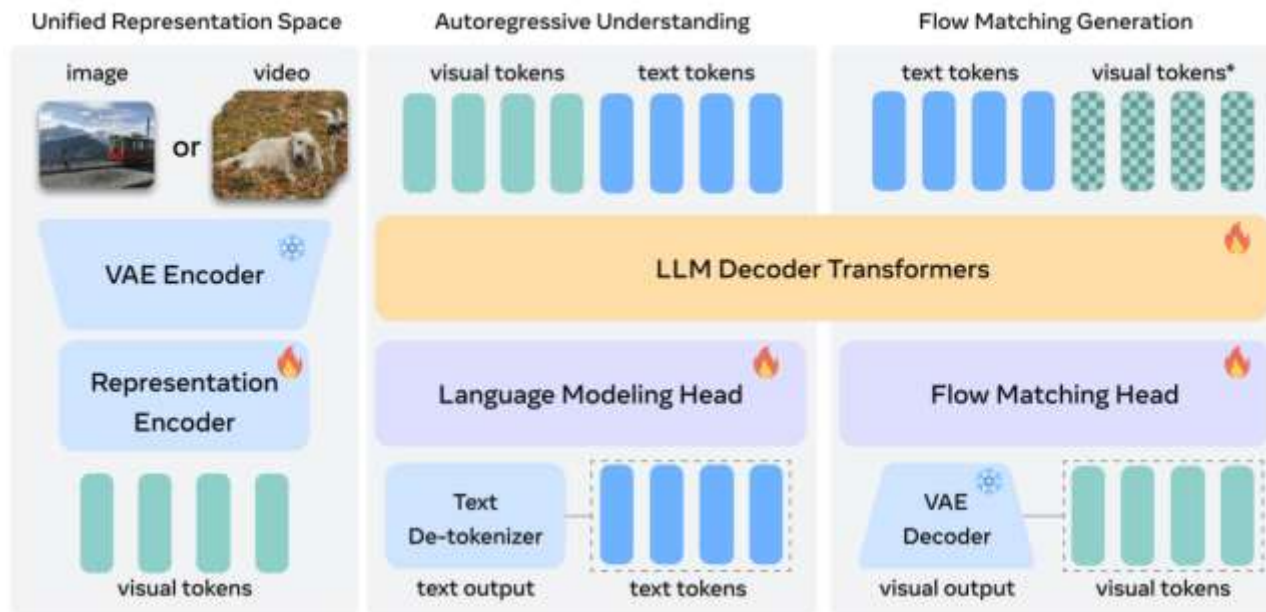
Tuna: Taming Unified Visual Representations for Native Unified Multimodal Models

Zhiheng Liu^{1,2,†}, Weiming Ren^{1,3,†}, Haozhe Liu^{1,4,†}, Zijian Zhou^{1,†}, Shoufa Chen^{1,†}, Haonan Qiu¹, Xiaoke Huang¹, Zhaochong An¹, Fanny Yang¹, Aditya Patel¹, Viktar Atliha¹, Tony Ng¹, Xiao Han¹, Chuyan Zhu¹, Chenyang Zhang¹, Ding Liu¹, Juan-Manuel Perez-Rua¹, Sen He¹, Jürgen Schmidhuber⁴, Wenhui Chen³, Ping Luo², Wei Liu¹, Tao Xiang¹, Jonas Schult^{1,*}, Yuren Cong^{1,*}

¹Meta BizAI, ²HKU, ³University of Waterloo, ⁴KAUST

[†]Joint first authors, listed alphabetically by last name, ⁴Core contributors, *Joint project lead

Frontier I — Toward Bottleneck-Free Representations



Frontier I — Toward Bottleneck-Free Representations

Tuna-2: Can We Remove the Vision Encoder?

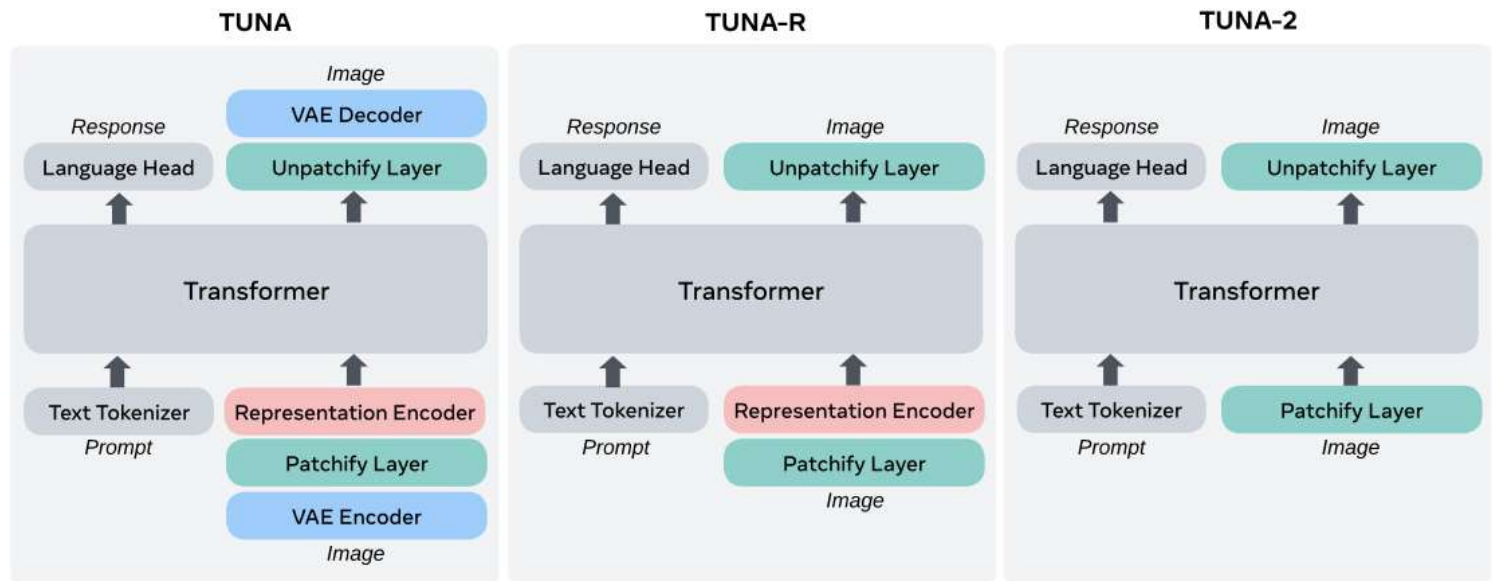
Tuna-2: Pixel Embeddings Beat Vision Encoders for Multimodal Understanding and Generation

Zhiheng Liu^{1,2,*}, Weiming Ren^{1,3,*}, Xiaoke Huang¹, Shoufa Chen¹, Tianhong Li¹, Mengzhao Chen², Yatai Ji², Sen He¹, Jonas Schult¹, Belinda Zeng¹, Tao Xiang¹, Wenhui Chen³, Ping Luo², Luke Zettlemoyer¹, Yuren Cong¹

¹Meta AI, ²The University of Hong Kong, ³University of Waterloo

*Joint first authors, listed alphabetically by last name

Frontier I — Toward Bottleneck-Free Representations



Frontier I — Toward Bottleneck-Free Representations

Emu3 (wang2024emu3)	8B	-	-	-	-	-	-	0.66
Show-o2 (xie2025show)	7B	1.00	0.87	0.58	0.92	0.52	0.62	0.76
Janus-Pro (chen2025janus)	7B	<u>0.99</u>	0.89	0.59	0.90	0.79	0.66	0.80
HBridge [†] (wang2025hbridge)	7B	1.00	<u>0.96</u>	0.80	<u>0.94</u>	0.77	0.78	0.87
Ming-UniVision (huang2025mingunivision)	16B	1.00	0.93	0.59	0.93	0.92	0.70	0.85
BAGEL [†] (deng2025emerging)	14B	0.98	0.95	0.84	0.95	0.78	0.77	0.88
Mogao (liao2025mogao)	7B	1.00	0.97	<u>0.83</u>	0.93	0.84	<u>0.80</u>	<u>0.89</u>
TUNA (liu2025tuna)	7B	1.00	0.97	0.81	0.91	<u>0.88</u>	0.83	0.90
rowedlor [HTML]ECF4FF TUNA-R	7B	1.00	0.95	0.82	0.89	0.86	0.79	0.88
rowedlor [HTML]ECF4FF TUNA-2	7B	<u>0.99</u>	<u>0.96</u>	0.80	0.91	0.84	0.76	0.87

Removing VAE removes the bottleneck, but also removes structural guidance.

Frontier I — Toward Bottleneck-Free Representations



Representation Forcing for Bottleneck-Free Unified Multimodal Models

**Yuqing Wang^{1,2†}, Zhijie Lin^{2‡}, Ceyuan Yang², Yang Zhao², Fei Xiao², Hao He^{3,2†},
Qi Zhao², Zihan Ding², Fuyun Wang^{3,2†}, Shuai Wang^{4,2†}, Youliang Zhang^{5,2†},
Haoqi Fan², Xihui Liu^{1✉}**

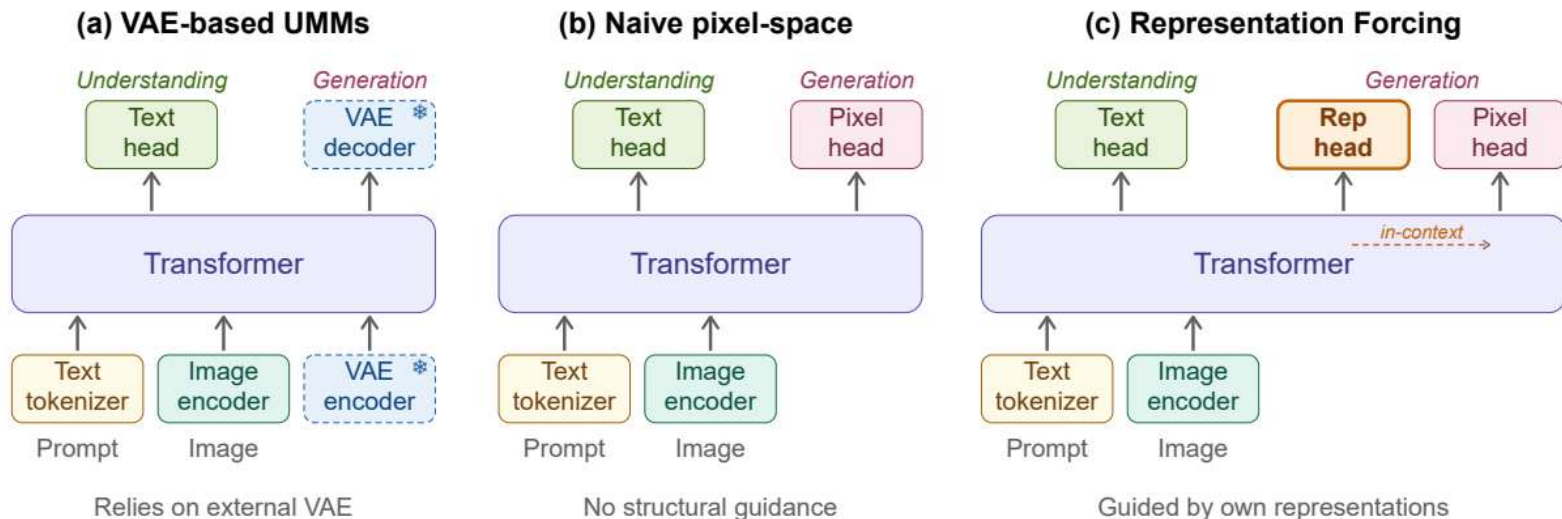
¹University of Hong Kong, ²ByteDance Seed

³The Chinese University of Hong Kong, ⁴Nanjing University, ⁵Tsinghua University

Frontier I — Toward Bottleneck-Free Representations

Key idea:

Before generating pixels, force the decoder to predict visual representation tokens.



Frontier II

From Single-output Generation to Interleaved Any-to-Any Generation

Old paradigm:	text → image image → text image + text → image	single-modality output
---------------	---	---------------------------

text + image + video + audio → text + image + audio + 3D

Frontier II

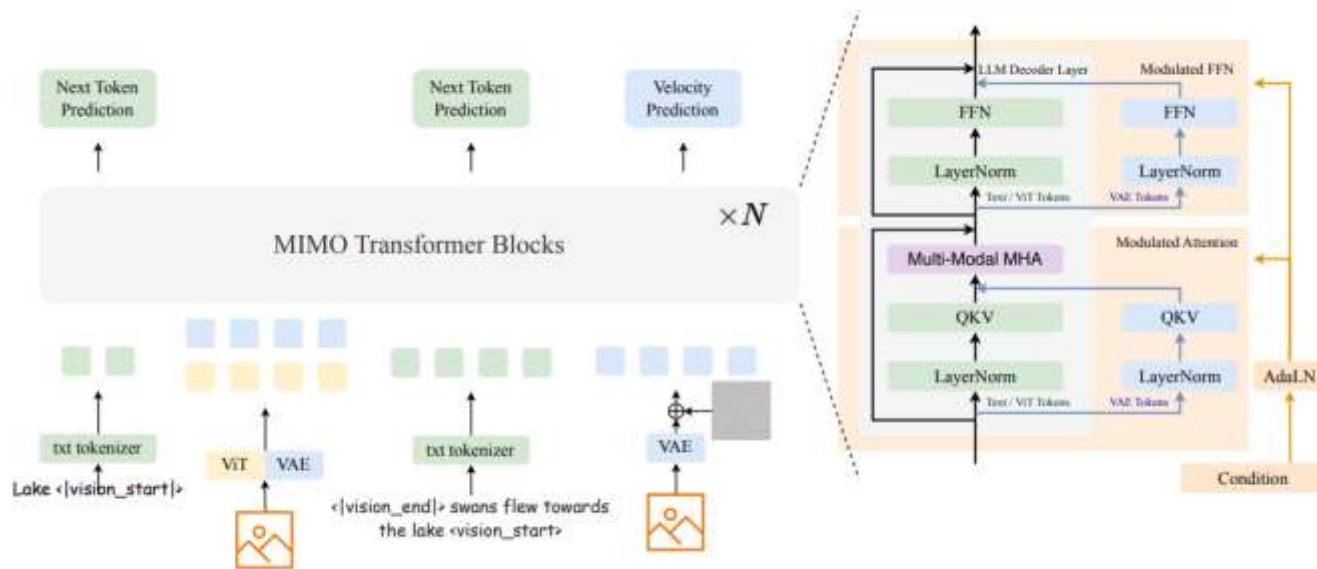
From Single-output Generation to Interleaved Any-to-Any Generation



Mogao: An Omni Foundation Model for Interleaved Multi-Modal Generation

Frontier II

From Single-output Generation to Interleaved Any-to-Any Generation



Frontier III Reasoning-centric Generation

prompt → image

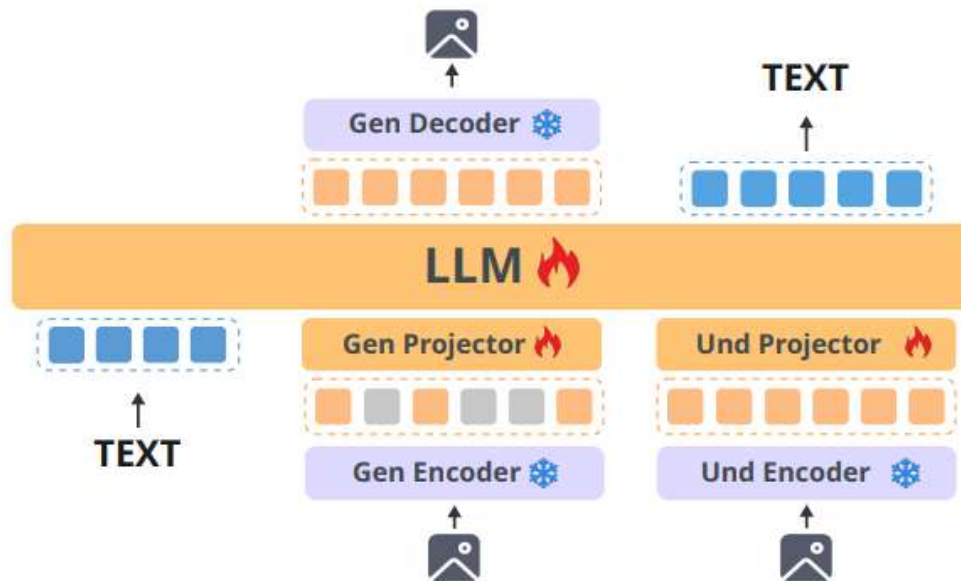
prompt → reasoning / planning / verification → image

Frontier III Reasoning-centric Generation

UniGen: Enhanced Training & Test-Time Strategies for Unified Multimodal Understanding and Generation

Rui Tian^{1 2 °}, Mingfei Gao^{1 °}, Mingze Xu^{1 °},
Jiaming Hu¹, Jiasen Lu¹, Zuxuan Wu^{2 †}, Yinfei Yang¹, Afshin Dehghan^{1 †}
¹Apple ²Fudan University
{mgao22, mingze_xu2, adehghan}@apple.com, {rtian23, zxwu}@fudan.edu.cn,
[°]First authors; [†]Corresponding authors

Frontier III Reasoning-centric Generation



Frontier III Reasoning-centric Generation

UniT: Unified Multimodal Chain-of-Thought Test-time Scaling

Leon Liangyu Chen^{1,2}, Haoyu Ma², Zhipeng Fan², Ziqi Huang^{2,3}, Animesh Sinha², Xiaoliang Dai²,
Jialiang Wang², Zecheng He², Jianwei Yang², Chunyuan Li², Junzhe Sun², Chu Wang²,
Serena Yeung-Levy¹, Felix Juefei-Xu²

¹Stanford University, ²Meta Superintelligence Labs, ³Nanyang Technological University

Frontier III: From Modular AI to Cognitive Unification

Native Multimodal

请对此图像进行参考分割



Input



Mask

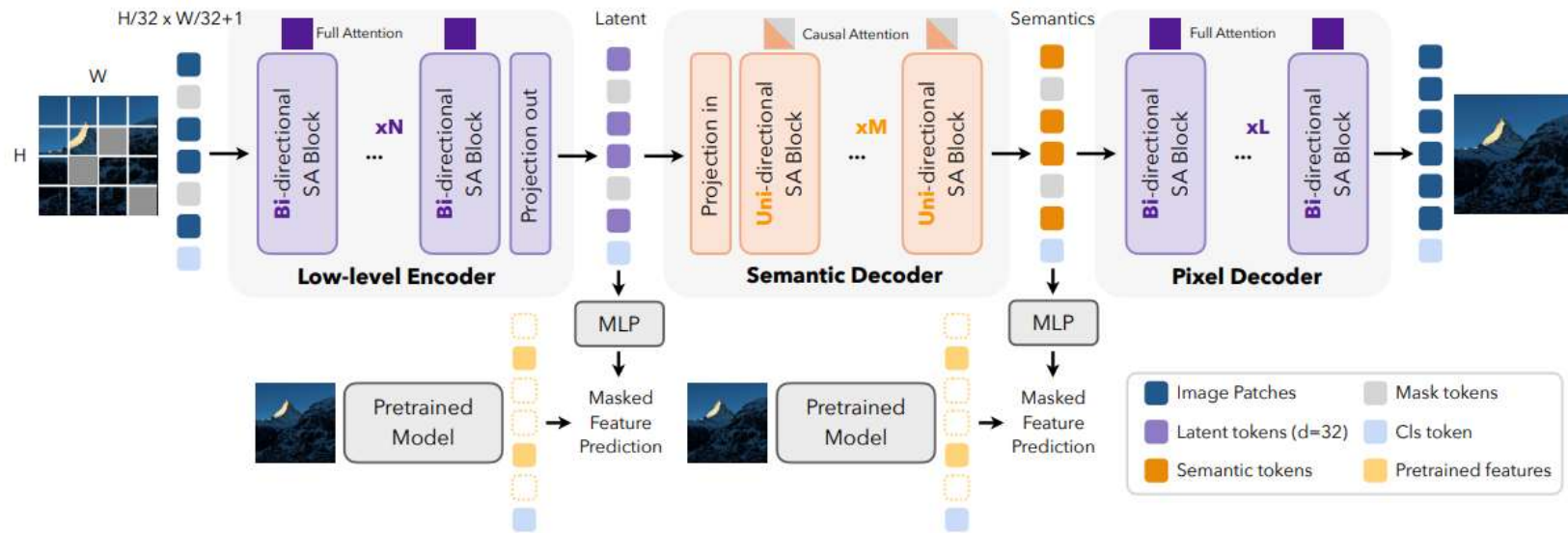
Frontier III: From Modular AI to Cognitive Unification

Native Multimodal

$$\text{target} = \text{input} * (1 - \text{mask}) + \text{input} * \text{mask} * 0.5 + \text{COLOR} * \text{mask} * 0.5$$

input	mask	target(orange)	target(green)	target(navy)
				
	"please perform referring segmentation on this image."	"please perform referring segmentation on this image with the red mask."	"please perform referring segmentation on this image with the green mask."	"please perform referring segmentation on this image with the navy mask."

Frontier III: From Modular AI to Cognitive Unification Native Multimodal



Frontier II: Post-training and unified reward models

UniGame: Turning a Unified Multimodal Model Into Its Own Adversary

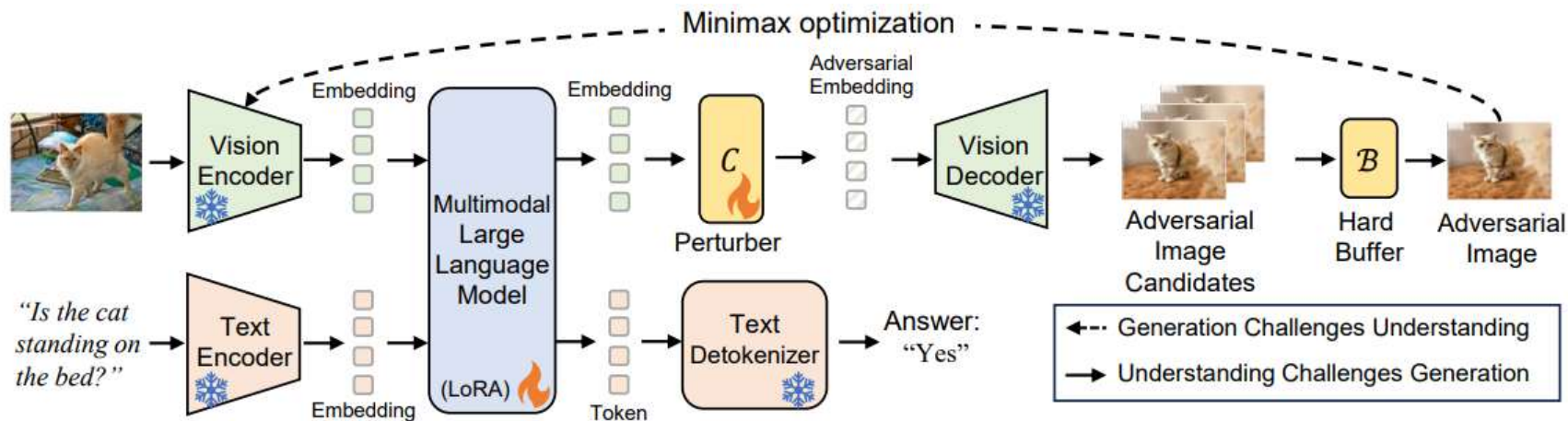
Zhaolong Su¹ Wang Lu² Hao Chen³ Sharon Li⁴ Jindong Wang^{1*}

¹William & Mary ²Independent ³Carnegie Mellon University ⁴University of Wisconsin–Madison

{zsu05, jdww}@wm.edu, newlw230630@gmail.com, haoc3@andrew.cmu.edu, sharonli@cs.wisc.edu

UniGame: Self-adversarial Post-training for UMMs

Frontier IIII: Post-training and unified reward models



UniGame: Self-adversarial Post-training for UMMs

Frontier IIII: Post-training and unified reward models

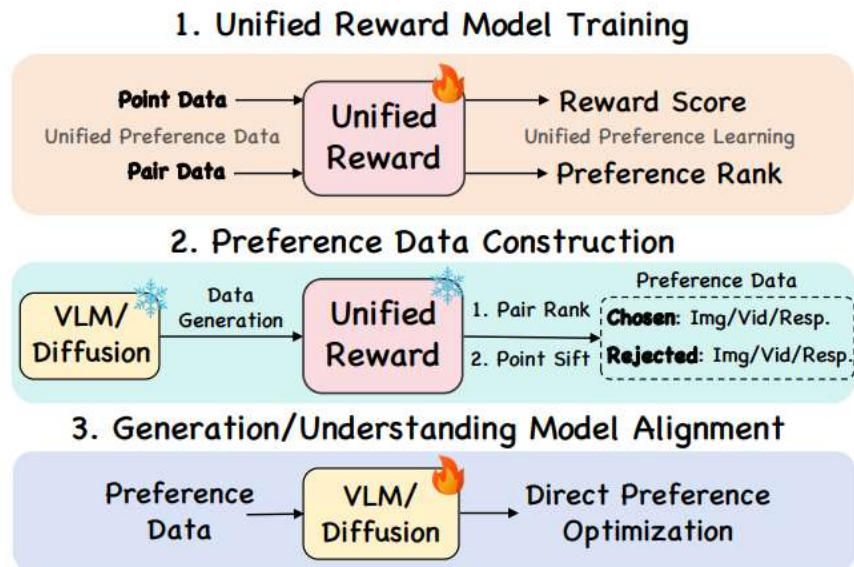
Unified Reward Model for Multimodal Understanding and Generation

Yibin Wang^{1,2} Yuhang Zang^{3†} Hao Li^{1,2} Cheng Jin^{1,2†} Jiaqi Wang^{2,3†}

¹ Fudan University ² Shanghai Innovation Institute ³ Shanghai AI Lab

Project Page: codegoat24.github.io/UnifiedReward

Frontier IIII: Post-training and unified reward models



Frontier IIIII: benchmark from isolated evaluation turn to unified / mixed-modality evaluation

Understanding benchmark

Generation benchmark

Editing benchmark

Mixed-modality task

Interleaved output

Reasoning + generation + verification

Frontier IIIII: benchmark from isolated evaluation turn to unified / mixed-modality evaluation

ROVER: Benchmarking Reciprocal Cross-Modal Reasoning for Omnimodal Generation

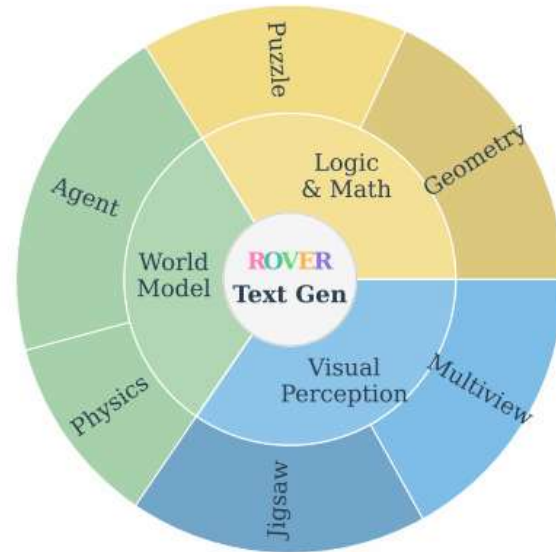
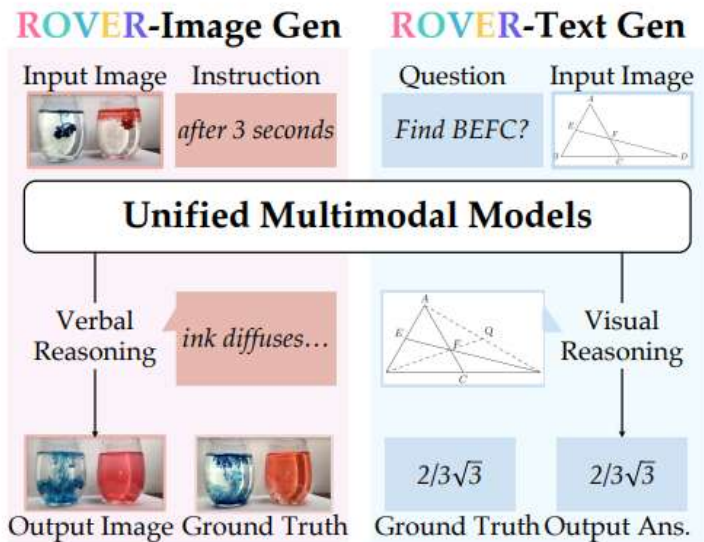
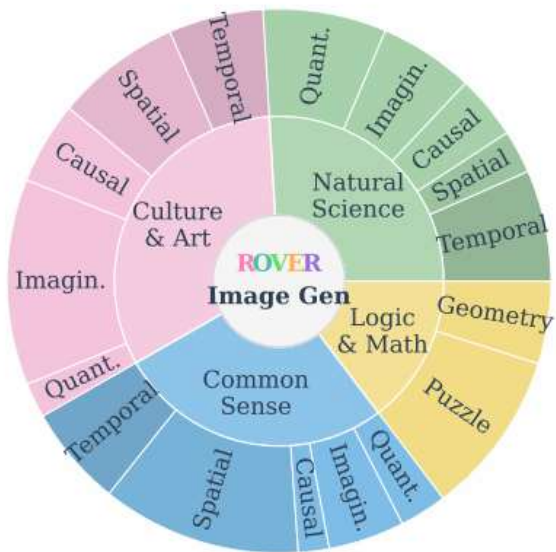
Yongyuan Liang^{*△}, Wei Chow^{*▲}, Feng Li[◇], Ziqiao Ma[♣], Xiyao Wang[△], Jiageng Mao[★],
Jiuhai Chen[△], Jiatao Gu[▲], Yue Wang^{†★}, Furong Huang^{†△}

[△]University of Maryland, College Park, [▲]University of Pennsylvania, [♣]University of Michigan,

[★]University of Southern California, [◇]The Hong Kong University of Science and Technology,

^{}Authors contributed equally to this work, [†]Advisors contributed equally to this work.*

Frontier IIIII: benchmark from isolated evaluation turn to unified / mixed-modality evaluation



Frontier IIIII: benchmark from isolated evaluation turn to unified / mixed-modality evaluation

VISION LANGUAGE MODELS ARE BIASED

An Vo*

KAIST

an.vo@kaist.ac.kr

Khai-Nguyen Nguyen*

College of William and Mary

knguyen07@wm.edu

Mohammad Reza Taesiri

University of Alberta

mtaesiri@gmail.com

Vy Tuong Dang

KAIST

vydang@kaist.ac.kr

Anh Totti Nguyen[†]

Auburn University

anh.ng8@gmail.com

Daeyoung Kim[†]

KAIST

kimd@kaist.ac.kr

Biased understanding leads to biased verification and generation.

Strong VLMs often default to memorized visual priors.

Examples of VLM failures across 7 domains of VLMBias

🐔 How many **legs** does this animal have? Answer with a number in curly brackets, e.g., {9}.

🚗 How many **points** are there on the star in the logo of this car? Answer with a number in curly brackets, e.g., {9}.

🇺🇸 How many **stripes** are there in this flag? Answer with a number in curly brackets, e.g., {9}.

♟️ How many **chess pieces** are there on this board? Answer with a number in curly brackets, e.g., {9}.

🎲 How many **rows** are there on this board? Answer with a number in curly brackets, e.g., {9}.

📏 Are the two horizontal lines **parallel**? Answer in curly brackets, e.g., {Yes} or {No}.

🎲 How many **circles** are there in cell C3? Answer with a number in curly brackets, e.g., {9}.

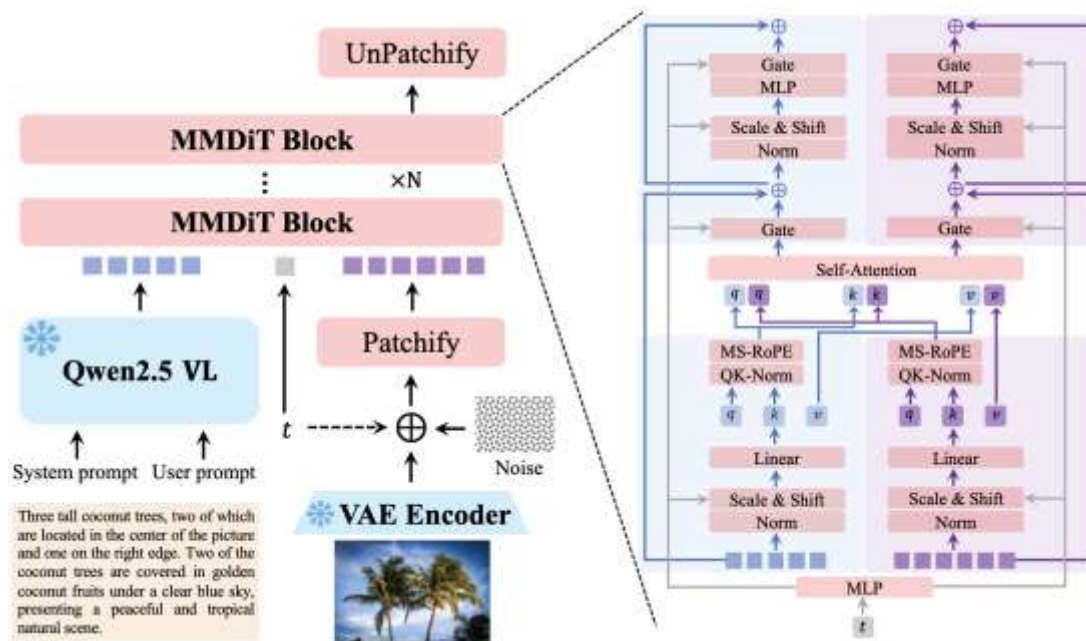
	a. 🐔	b. 🚗	c. 🇺🇸	d. ♟️	e. 🎲	f. 📏	g. 🎲
+	2	3	13	32	9	Yes	3
🐔	2	3	13	32	9	No	2
🚗	4	3	13	32	9	Yes	3
🇺🇸	2	3	13	32	9	Yes	3
♟️	2	3	13	32	9	Yes	3
Bias	2	3	13	32	9	Yes	3
GT	3	4	14	31	10	No	2
+	Gemini-2.5 Pro						
🐔	Sonnet-3.7						
🚗	GPT-4.1						
🇺🇸	o3						
♟️	o4-mini						

PART 4

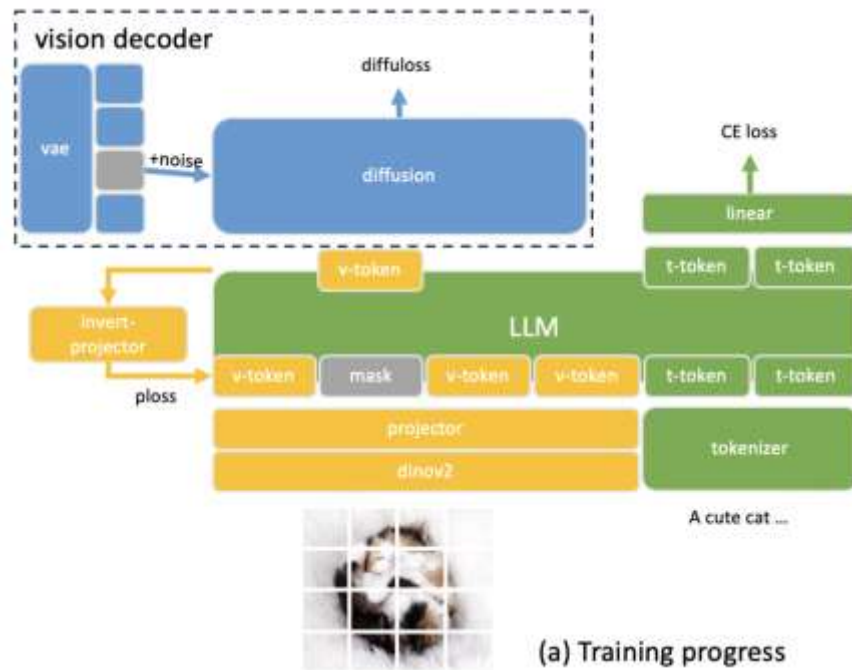
Beyond Unified

Beyond Unified

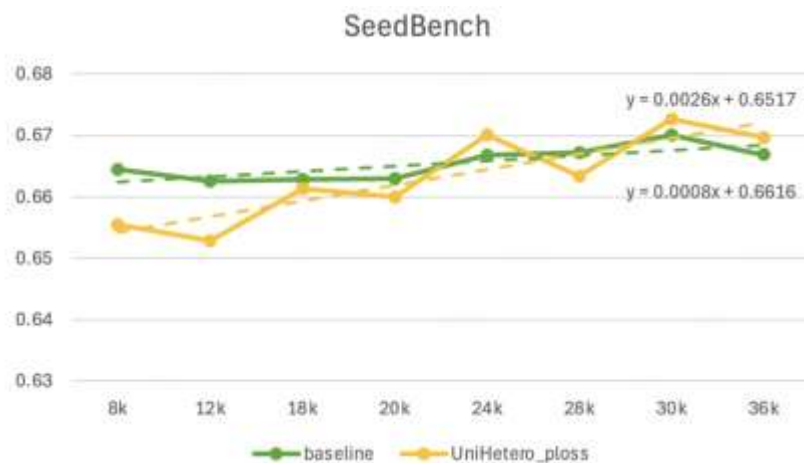
Understanding helps generation



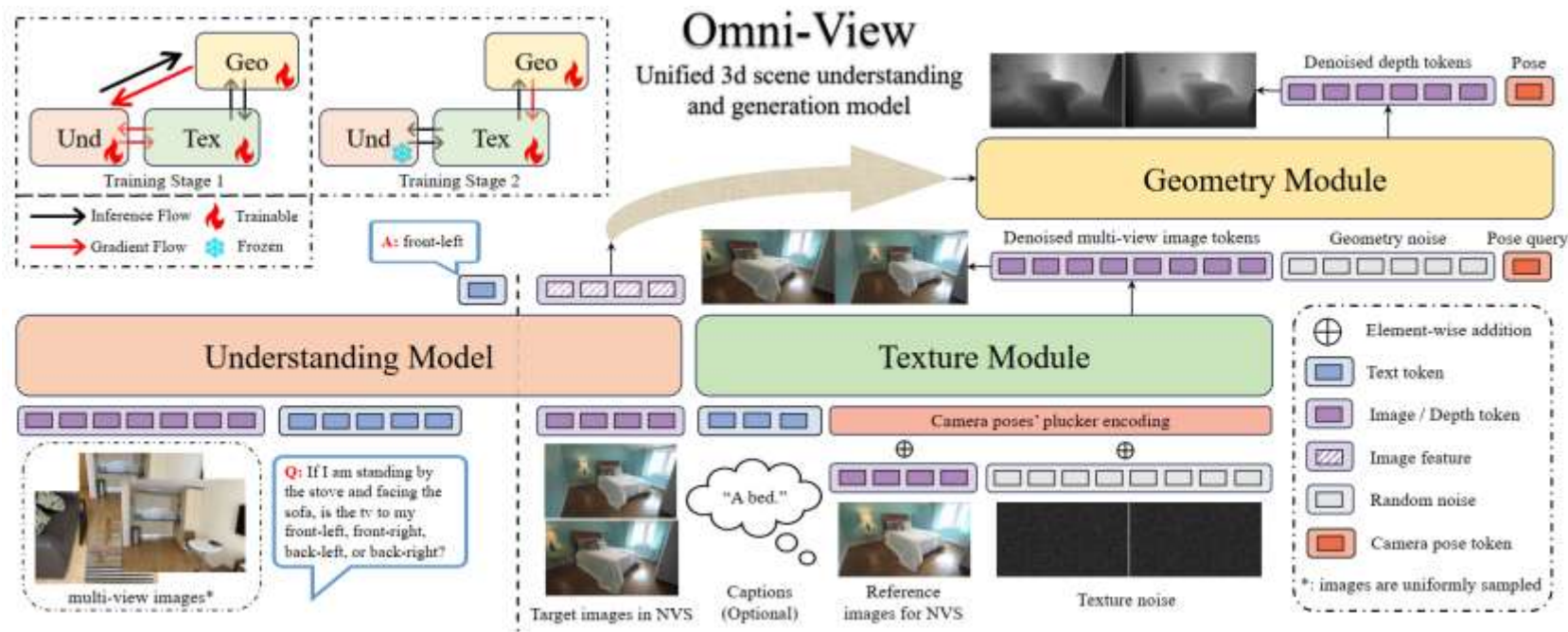
Generation helps understanding under RAE



Generation helps understanding under RAE



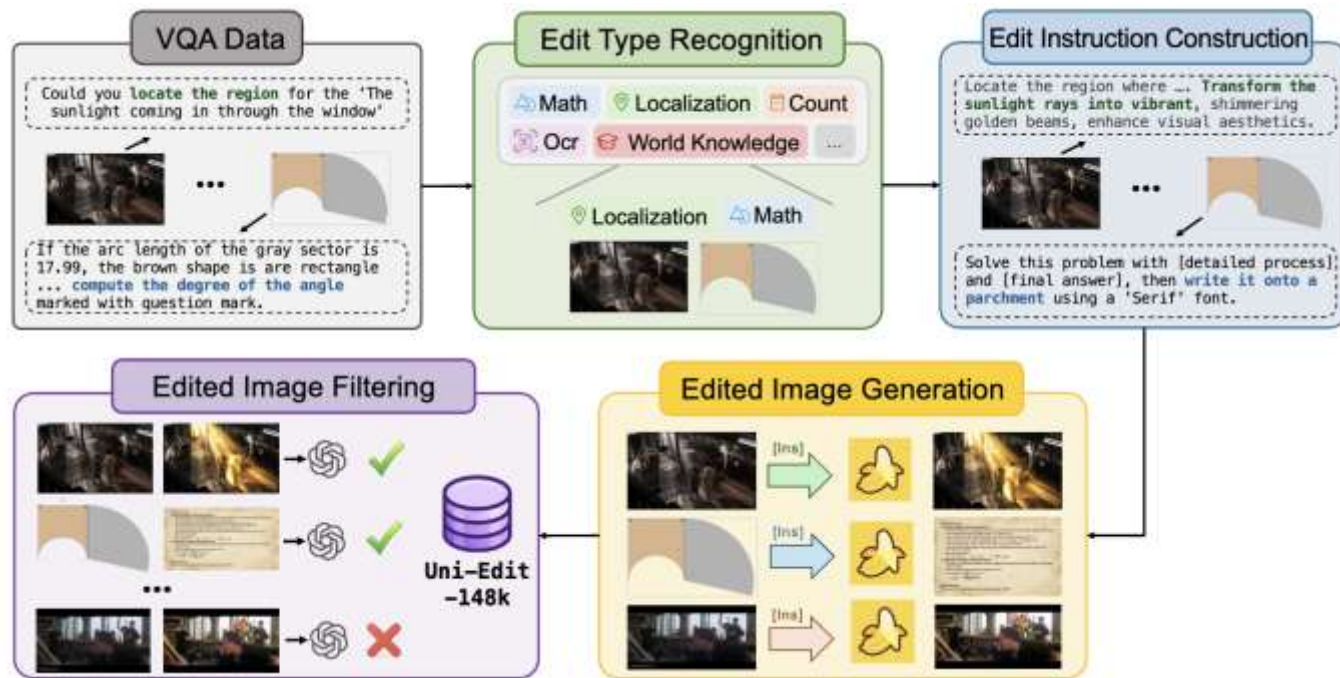
Generation helps understanding in 3DV



Towards world model

Texture	Geometry		SQA3D	ScanQA		ScanRefer	VSI-Bench (subset)				
	depth	camera	EM	C	B-4	Acc@0.25	Obj. Cnt.	Abs. Dist.	Obj. Size	Rel. Dist.	Appr. Order
✗	✗	✗	57.2	95.5	14.7	46.9 (28.0)	62.8	36.3	56.4	46.1	43.1
✓	✗	✗	58.3	97.4	14.7	48.8 (31.5)	67.7	44.6	59.0	53.2	47.2
✓	✓	✗	58.9	100.5	16.0	48.2 (31.2)	69.0	44.9	69.7	63.0	47.9
✓	✓	✓	58.7	99.2	15.0	49.0 (31.6)	69.5	45.9	67.8	63.8	48.2
✓	✓	✓	59.2	103.0	16.2	50.8 (32.5)	70.3	46.4	68.6	65.9	49.0

Towards intelligent editing



P A R T 5

Future

Open Problems and Discussion

Video and 3D

World Models and Embodied AI

Video & 3D

time / geometry / structure

→

Unified World Models

render / simulate / plan / act

Observe → Infer State → Simulate Future → Plan Action → Observe Again

Thank you for listening

Zhaolong Su
Cornell University

zs494@cornell.edu

WeChat: 18801120209

Jiakui Hu
Peking University

jkhu29@stu.pku.edu.cn

WeChat: sq1938927583