

Video models are zero-shot learners and reasoners

Google Deepmind 2025

Reviewed & Presented by Joon Hyeok Kim

Contents

1. Central Thesis : **Video Model's** Analogy with **LLMs** as a Generalist Model
2. Framework : **Hierarchical Categorizations** of Visual Capabilities
3. Methods, Evaluations, & Experiments
4. Limits & Future Outlook

Recall the history of the language models.

The Pre-LLM Era was like the 群雄割据(군웅할거) of ...

Task-Specific Bespoke Models

- Translation
- Summarization
- Domain specific QnA
- ...



Paradigm Shift : Chain-of-Thoughts (CoT) + Computing Power

Standard Prompting

Model Input

Q: Roger has 5 tennis balls. He buys 2 more cans of tennis balls. Each can has 3 tennis balls. How many tennis balls does he have now?

A: The answer is 11.

Q: The cafeteria had 23 apples. If they used 20 to make lunch and bought 6 more, how many apples do they have?

Model Output

A: The answer is 27. ❌

Chain-of-Thought Prompting

Model Input

Q: Roger has 5 tennis balls. He buys 2 more cans of tennis balls. Each can has 3 tennis balls. How many tennis balls does he have now?

A: Roger started with 5 balls. 2 cans of 3 tennis balls each is 6 tennis balls. $5 + 6 = 11$. The answer is 11.

Q: The cafeteria had 23 apples. If they used 20 to make lunch and bought 6 more, how many apples do they have?

Model Output

A: The cafeteria had 23 apples originally. They used 20 to make lunch. So they had $23 - 20 = 3$. They bought 6 more apples, so they have $3 + 6 = 9$. The answer is 9. ✅



Chain-of-Thought Prompting Elicits Reasoning in Large Language Models
Google Research, Brain Team 2022

LLM unified and started to work as a Generalist

- Translation
- Summarization
- Domain specific QnA
- And, now even capable of
 - Coding
 - Math
 - Creative writing
 - Deep research (oh my...)
 - and so on...



Analogy in Video Models

Current task-specific video models that outperform general video models...

Segmentation Specialist

Object Detection Specialist (YOLO variants)

Segment Anything

Alexander Kirillov^{1,2,4} Eric Mintun² Nikhila Ravi^{1,2} Hanzi Mao² Chloe Rolland³ Laura Gustafson³
Tete Xiao³ Spencer Whitehead Alexander C. Berg Wan-Yen Lo Piotr Dollár⁴ Ross Girshick⁴
¹project lead ²joint first author ³equal contribution ⁴directional lead

Meta AI Research, FAIR

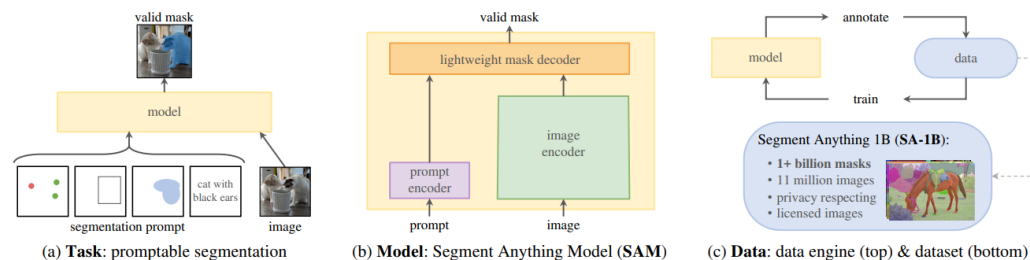
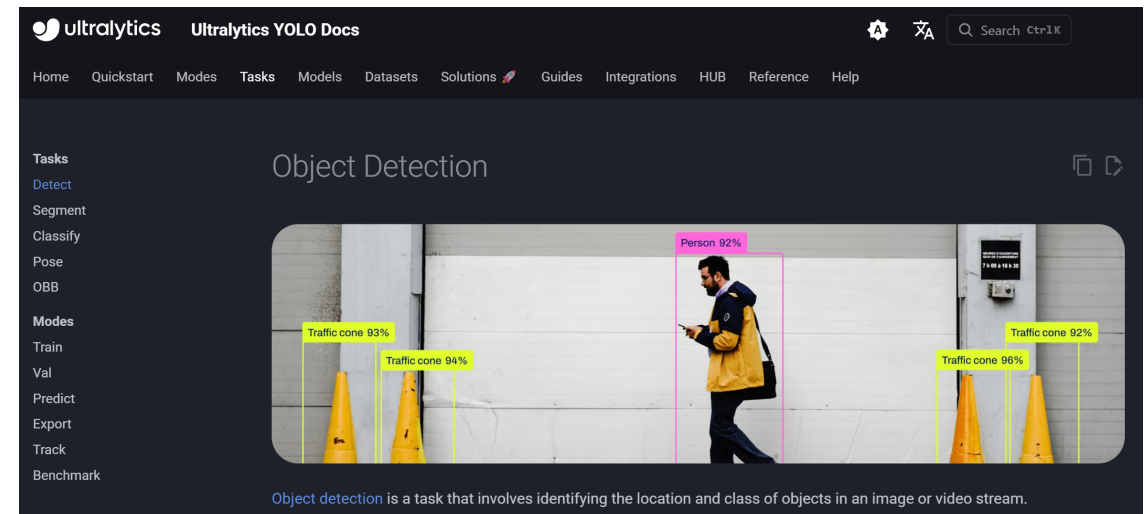


Figure 1: We aim to build a foundation model for segmentation by introducing three interconnected components: a promptable segmentation *task*, a segmentation *model* (SAM) that powers data annotation and enables zero-shot transfer to a range of tasks via prompt engineering, and a *data* engine for collecting SA-1B, our dataset of over 1 billion masks.



Can **video models** become **Generalists** just like **LLM** did?

The logo for Google VEO 3 is displayed on a dark gray rectangular background. The word "GOOGLE" is written in its signature multi-colored font (blue, red, yellow, blue, green, red). Below it, the text "VEO 3" is written in a clean, white, sans-serif font.

GOOGLE
VEO 3

The OpenAI Sora logo features the OpenAI logo, a stylized white knot-like icon, to the left of the text "OpenAI" in a black sans-serif font. Below "OpenAI", the word "Sora" is written in a large, bold, black serif font.

OpenAI
Sora

Maybe yes with...

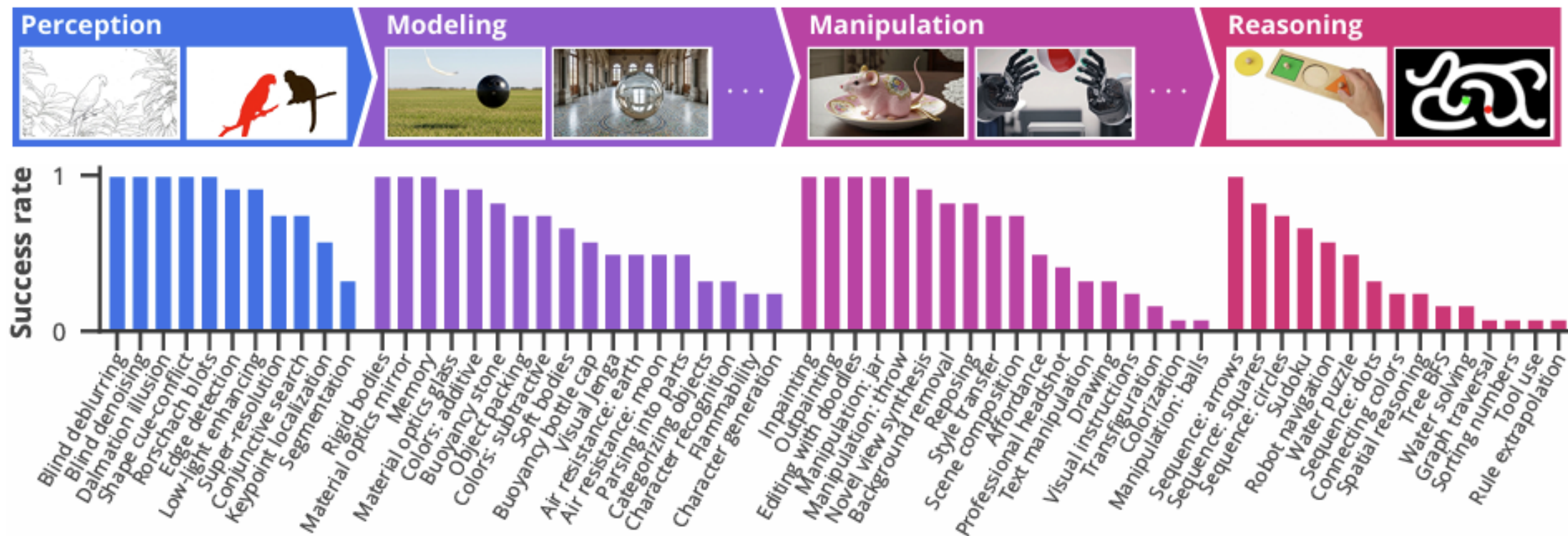
Chain-of-Frames (CoF) + Future Developments

- Applying changes across **dimensions** of the real world frame-by-frame
 - Dimensions : Time & Space
 - Similar to Step-by-step strategy in CoT



Still, conceptual and no analytic mechanism studied yet...

2. Hierarchical Categorizations of Visual Capabilities



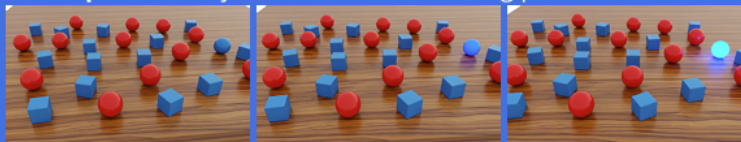
Provides a **framework** to assess various abilities of Video Models

Stacking up!

Perception – superresolution



Perception – conjunctive search / binding problem



Modeling – buoyancy



Modeling – memory of world states



Manipulation – 3D-aware reposing



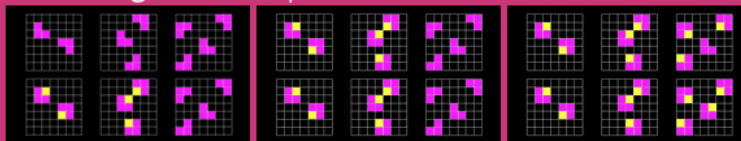
Manipulation – jar opening



Reasoning – robot navigation



Reasoning – rule extrapolation



Understand the world

Model the visual world

Alter modeled world

Reason across dims.

Examples

3. Methods, Evaluation, & Results

2. Methods

Approach and motivation Our method is simple: We prompt Veo. This minimalist strategy is intentional, as it mirrors the transformation of NLP from task-specific fine-tuning or training to

** Joint leads.*

e.g.



Figure 29 | **Categorizing objects.** Prompt: “A person puts all the kids toys in the bucket. Static camera, no pan, no zoom, no dolly.” Success rate: 0.33.

Evaluation Methodologies

Qualitative Evaluation

Concept) Success Rate

- Def.)
 - The fraction of generated videos that solved the task
- Props.)
 - Determined by **humans**
 - > 0 : the model possesses the ability to solve the task
 - ≈ 1 : the model reliably solves the problem irrespective of the random seed

Success

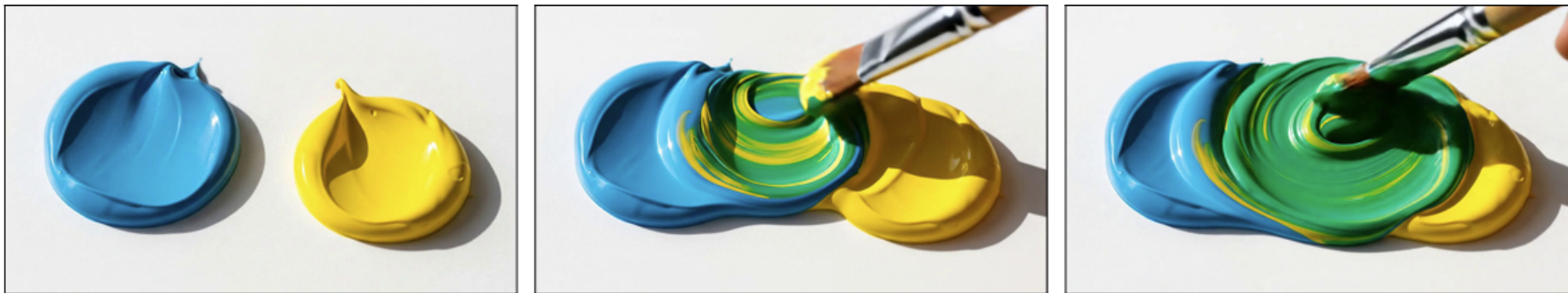


Figure 28 | **Color mixing. Additive** (lights, top). Prompt: “The spotlight on the left changes color to green, and the spotlight on the right changes color to blue.” Success rate: 0.92. **Subtractive** (paints, bottom). Prompt: “A paintbrush mixes these colors together thoroughly until they blend completely. Static camera, no pan, no zoom.” Success rate: 0.75.

Failure

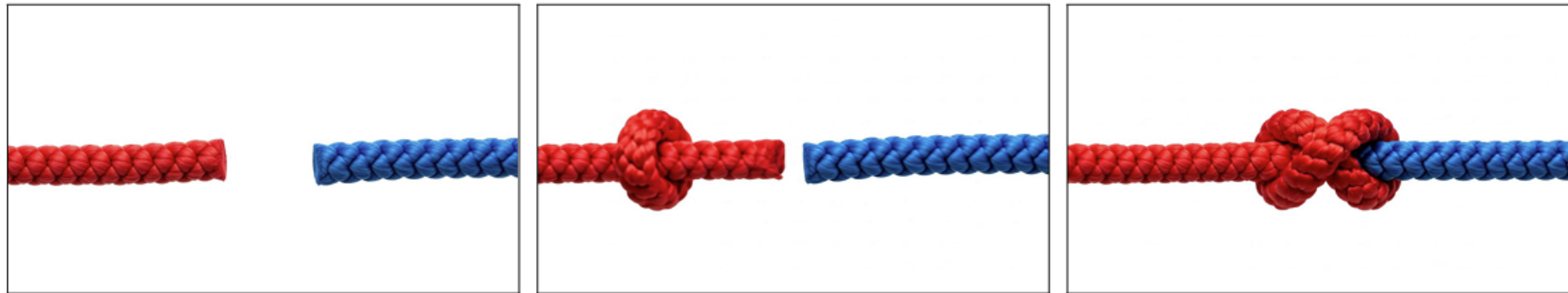


Figure 65 | **Tying the knot.** Prompt: “A knot is tied connecting these two rope ends.” Failure: physics violation, impossible rope movement.

Quantitative Evaluation

Problem specific scores are used.

- No unified measure nor statistics

Consider both the **best frame** and the **last frame**

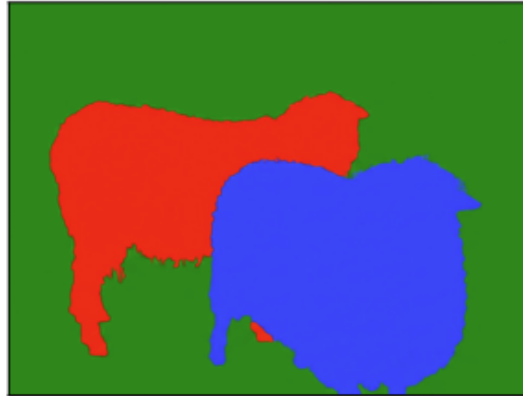
- Best Frame : Best performance but not deterministic
- Last Frame : Deterministic but no guarantee on performance

e.g.) Segmentation Problem (Perception) : Intersection over Union (IoU)

Original image



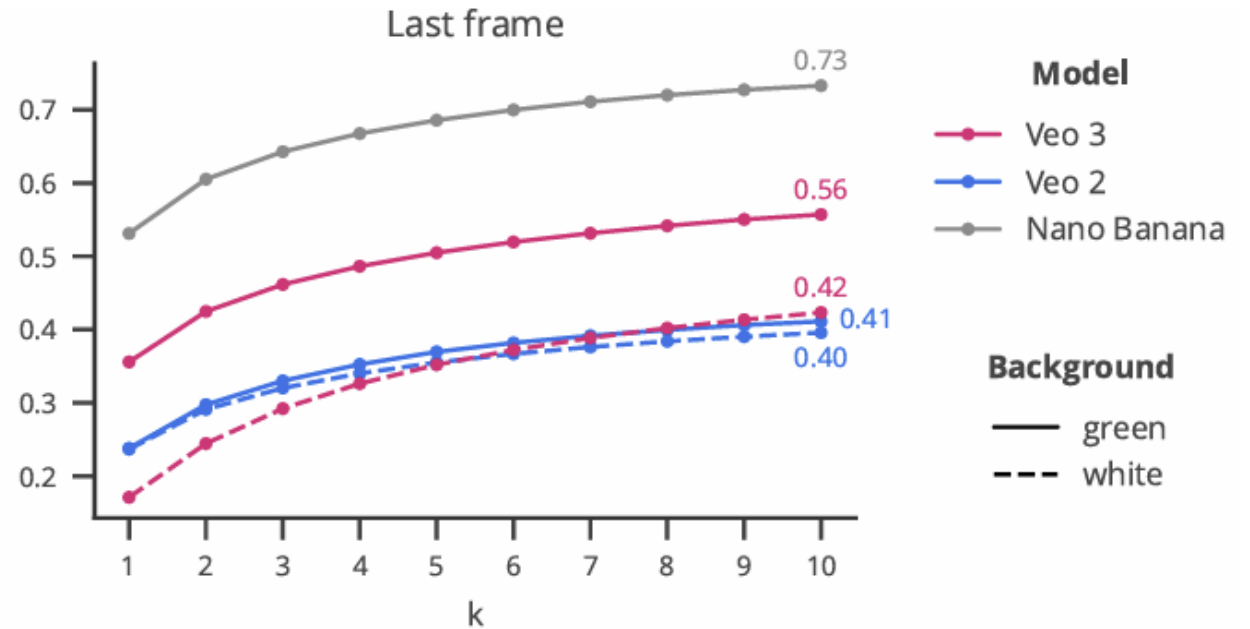
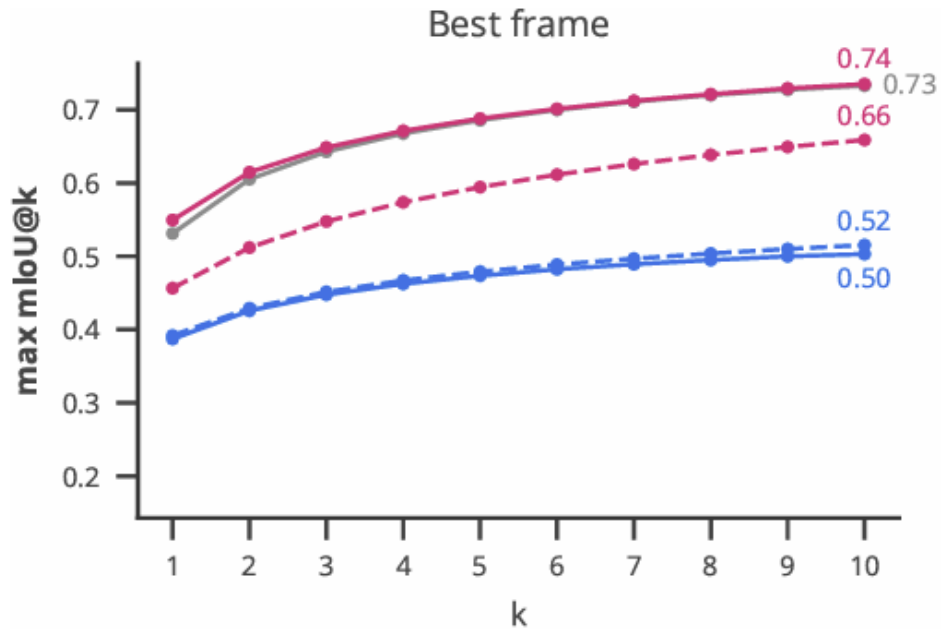
Generated frame (Veo 3)



Extracted masks

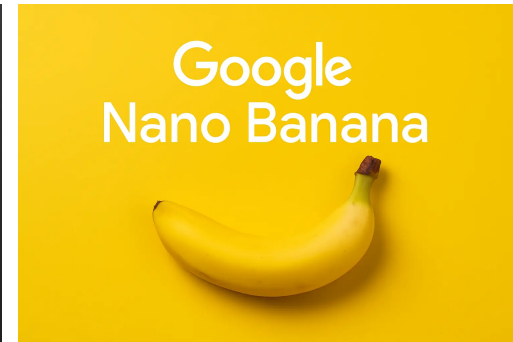


Ground-truth masks



e.g.) Maze Solving Problem (Reasoning)

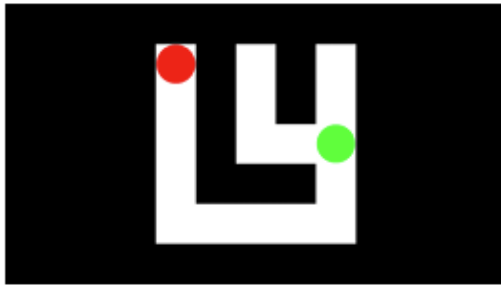
Compare...



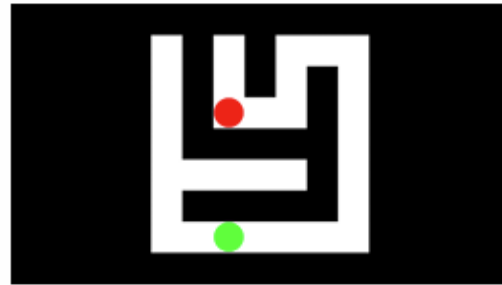
For Video Models... (Veo 2 & Veo 3)

1. Provide a maze image to the **prompt** as the first frame

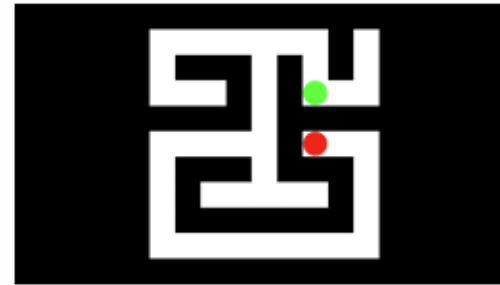
5×5 Grid



7×7 Grid



9×9 Grid



Irregular



- They tried various maze datasets.
 - maze-dataset 0.3.4
 - Hand drawn irregular mazes (flipped/rotated) to get 40 unique samples

2. Prompt text as below

Veo

Create a 2D animation based on the provided image of a maze. The red square slides smoothly along the white path, stopping perfectly on the green square. The red square never slides or crosses into the black areas of the maze. The camera is a static, top-down view showing the entire maze.

Maze:

- *The maze paths are white, the walls are black.*
- *The red square moves to the goal position, represented by a green square.*
- *The red square slides smoothly along the white path.*
- *The red square never slides or crosses into the black areas of the maze.*
- *The red square stops perfectly on the green square.*

Scene:

- *No change in scene composition.*
- *No change in the layout of the maze.*
- *The red square travels along the white path without speeding up or slowing down.*

Camera:

- *Static camera.*
- *No zoom.*
- *No pan.*
- *No glitches, noise, or artifacts.*

Detailed description of the maze
problem setting in paragraph form

Iteratively refining the maze
description and scenic details in a
note-taking format.

For Other Reference Models

- Nano Banana (Image Specific Generative Model)

Nano Banana

Mark the correct path from the red to the green circle through the maze in blue.

- Gemini 2.5 Pro (Language Model)

Gemini 2.5 Pro I2T

SYSTEM

Think step by step as needed and output in xml format:

<think>thinking process</think>

<final_answer>final answer</final_answer>

USER

The following image shows a maze, represented by colored squares:

- *Black squares represent walls and cannot be passed through.*
- *White squares are empty and can be passed through.*
- *The red square is the starting point.*
- *The green square is the end point.*

Evaluation

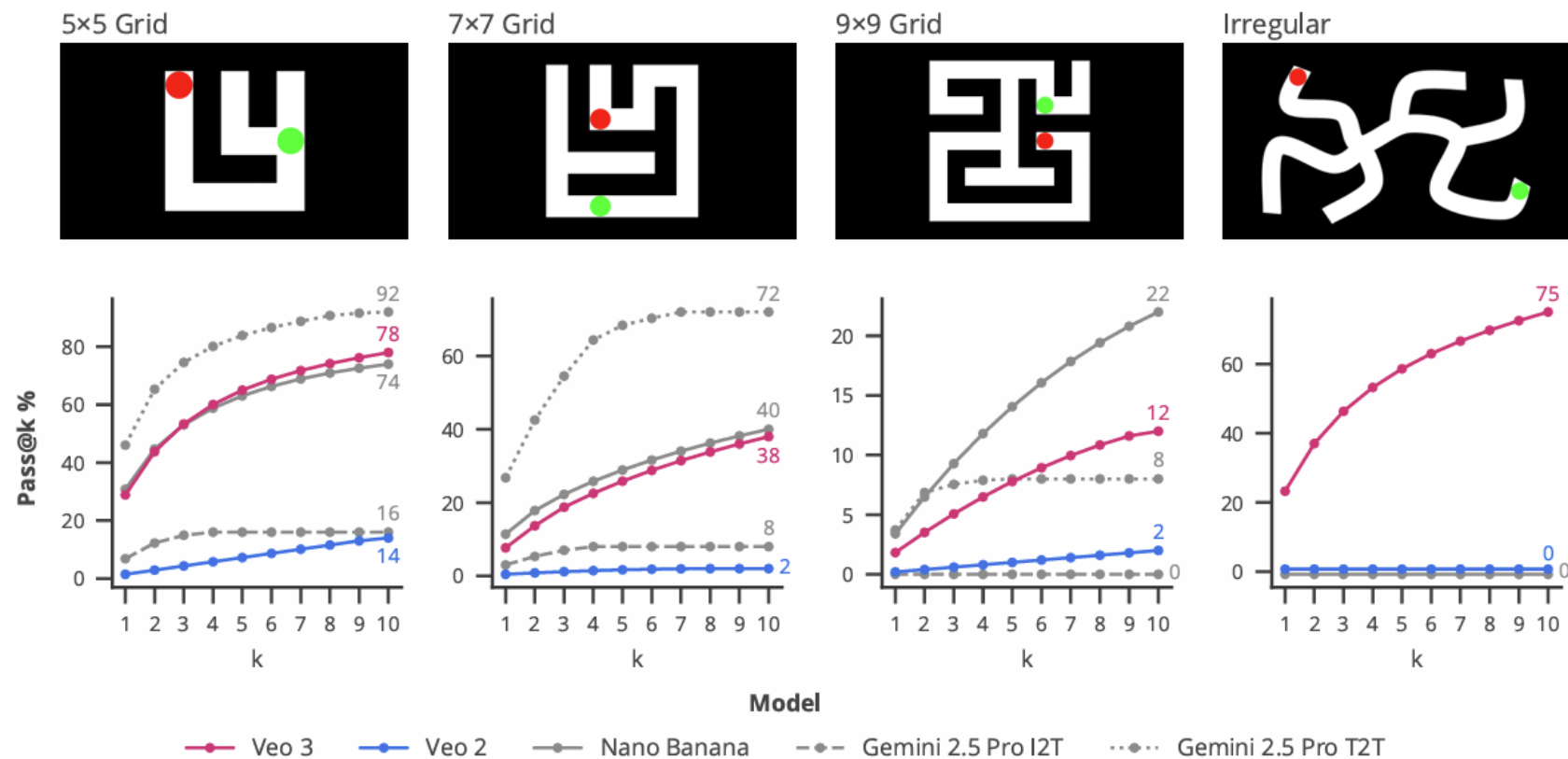
Define Illegal Moves

- Jumping over walls
- Clipping through boundaries
- Alteration of the goal's position
- ...

Success Rate (Quantitative)

- The fraction of k attempts where the agent successfully reaches the goal without illegal moves through out the generated video.

Result



- Veo2 vs Veo3
- $k \uparrow \Rightarrow \text{pass} \uparrow$
- Complex
 - Beats LLM
- Simple, Irregular
 - Beats Image

Limit : Insufficient Reasoning Capability

Significant portion of the reasoning experiments scored low success rates.

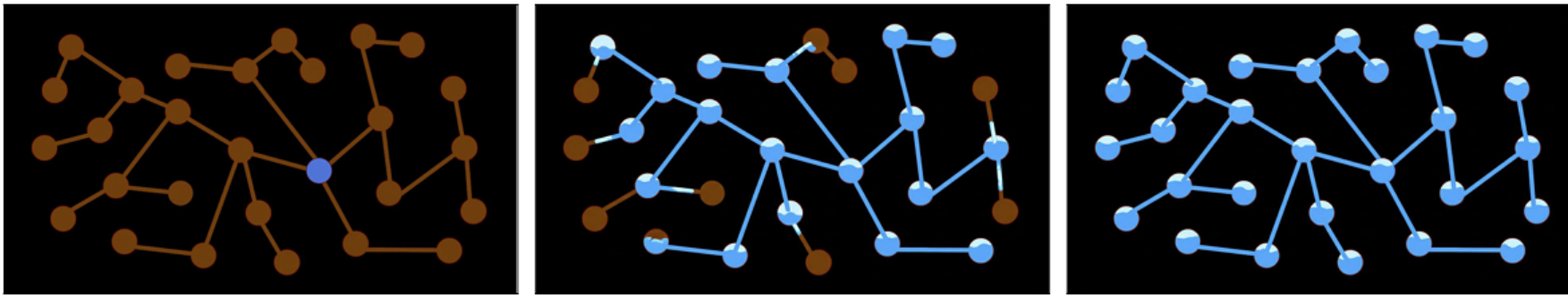


Figure 48 | **Graph traversal.** Prompt: “Starting from the blue well, an unlimited supply of blue water moves through the connected channel system without spilling into the black area.” Success rate: 0.08.



Figure 57 | **Maze solving.** Prompt: “Without crossing any black boundary, the grey mouse from the corner skillfully navigates the maze by walking around until it finds the yellow cheese.” Success rate: 0.17.

Failures on complex tasks

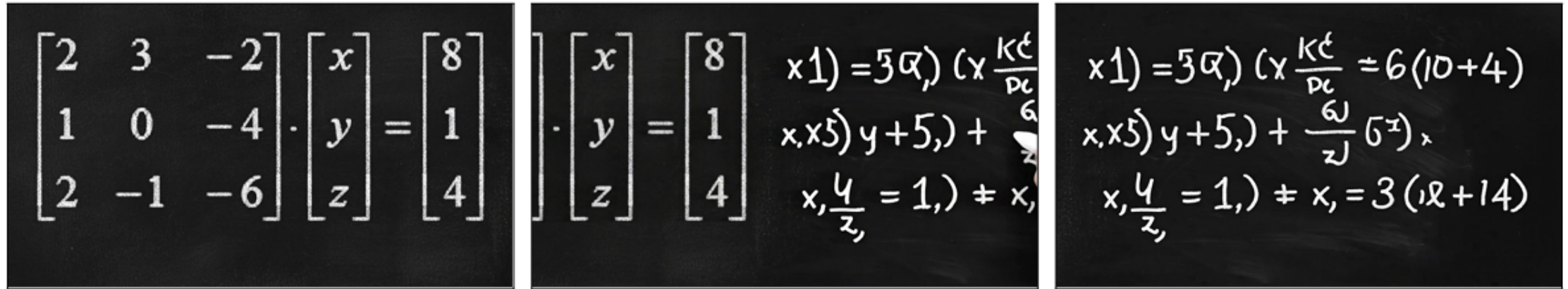

$$\begin{bmatrix} 2 & 3 & -2 \\ 1 & 0 & -4 \\ 2 & -1 & -6 \end{bmatrix} \cdot \begin{bmatrix} x \\ y \\ z \end{bmatrix} = \begin{bmatrix} 8 \\ 1 \\ 4 \end{bmatrix}$$
$$\begin{bmatrix} x \\ y \\ z \end{bmatrix} = \begin{bmatrix} 8 \\ 1 \\ 4 \end{bmatrix} \quad \begin{matrix} x1) = 3(2) (x \frac{K^6}{p^6} \\ x, x5) y + 5,) + \frac{6}{z} (5^x) \\ x, \frac{4}{z} = 1,) \neq x, \end{matrix}$$
$$\begin{matrix} x1) = 3(2) (x \frac{K^6}{p^6} = 6(10+4) \\ x, x5) y + 5,) + \frac{6}{z} (5^x) \\ x, \frac{4}{z} = 1,) \neq x, = 3(12+14) \end{matrix}$$

Figure 69 | **Solving system of linear equation..** Prompt: *"A hand appears and solves the set of linear equations. It replaces the x, y, z matrix with their correct values that solves the equation. Do not change anything else."* Failure: hallucinations with text on the blackboard.



Figure 72 | **Glass falling.** Prompt: *"The object falls. Static camera, no pan, no zoom, no dolly."* Failure: physics violation, glass does not break, and orients itself to be vertical after landing on the floor.

More Limits...

No analytic relation proven between **CoF** and **Zero-shot learning**

- Does CoF guarantees the correct reasoning path?
- Or, does "correctness" even matter if we have the answer?

No unified evaluation method : Human and problem dependent

Video generation is still too computationally expensive

- Emergent abilities only at scale! (Veo 3's performance improvement)

High Dependency on Prompt Engineering

- Nevertheless, CoT did change the world...

Jack of many trades, master of few

- Fine-tuned models dominate in specific tasks. Will they last forever?

Future Outlook

Video Models Show Great Potential for Zero-Shot Learning

- Distinguish a model's task performance and its underlying ability to solve it.
- Early LLMs underperformed against fine-tuned models.
 - But look who dominates now, huh?

We are at the very beginning of the Video Model development

- Improvement from Veo 2 to Veo 3 indicates further developments.
- Alternative approach on Prompt Engineering?
 - Current : first frame image + text description
 - Just like CoT changed the game with the prompt engineering in LLMs.
- Cost will fall : LLM's ongoing decrease in cost may support this.

Leveraging Scaling Laws and Optimization

- Performance can be further boosted by applying inference-time scaling and standard optimization toolkits, which were not used in this study

Questions

Thank you