



Summary

- Large language models with vision capabilities *still struggle* with **low-level vision tasks** that are trivial to humans.
- This failure is due to failure in translating **detailed visual information** into tokens in the LLM.

Motivation

- Gemini can **solve 42.9%** of the questions in MMMU benchmark **without seeing the input image**. [Are we on the right way...](#)Chen et al. 2024
- Most VQA benchmarks **do not exclusively test vision** capabilities.
- Text-only LLMs can reach **> 80%** of SotA on DocVQA, TextVQA, ChartQA, AI2D (images are serialized). [Analyzing...](#)Hedge et al. 2023

Benchmark findings

- VLMs cannot reliably
 - Tell if two circles are touching
 - Count the number of times two lines intersect
 - Follow paths from one point to another
 - Reliably count how many rows and columns are in a table
 - Tell which letter is being circled

58.07%

Examples from BlindTest benchmark with VLMs' responses

	P1	P2	P3	P4	P5	P6	P7
	1	Yes	o	6	5	3×4	1
	1	No	w	5	3	3×4	2
	1	Yes	o	5	4	4×4	2
	0	No	1	6	3	3×4	1
							

P1: How many times do the blue and red lines touch each other? Answer with a number in curly brackets, e.g., {5}.

P2: Are the two circles overlapping? Answer with Yes/No.

P3: Which character is being highlighted with a red oval? Please provide your answer in curly brackets, e.g. {a}

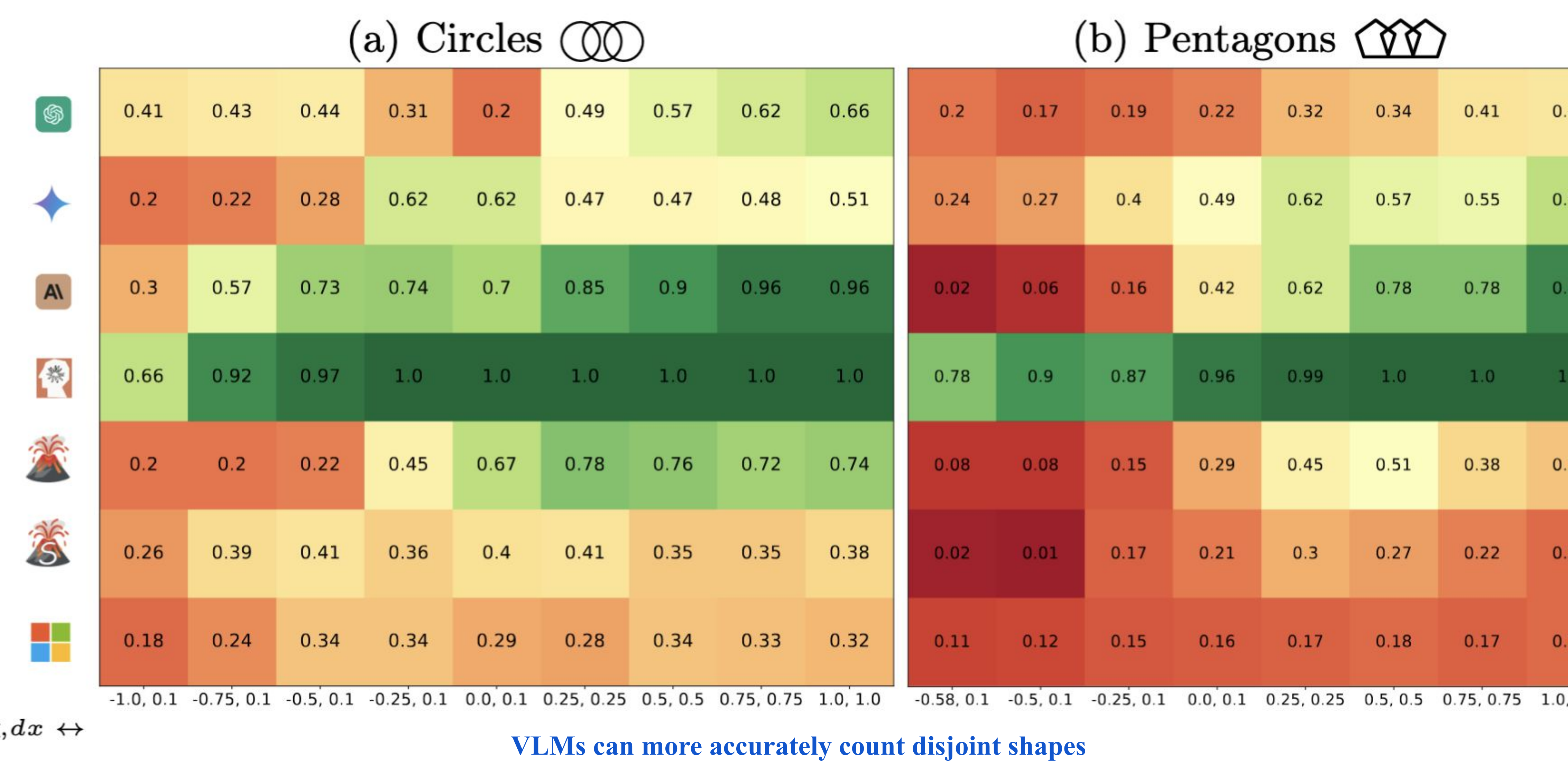
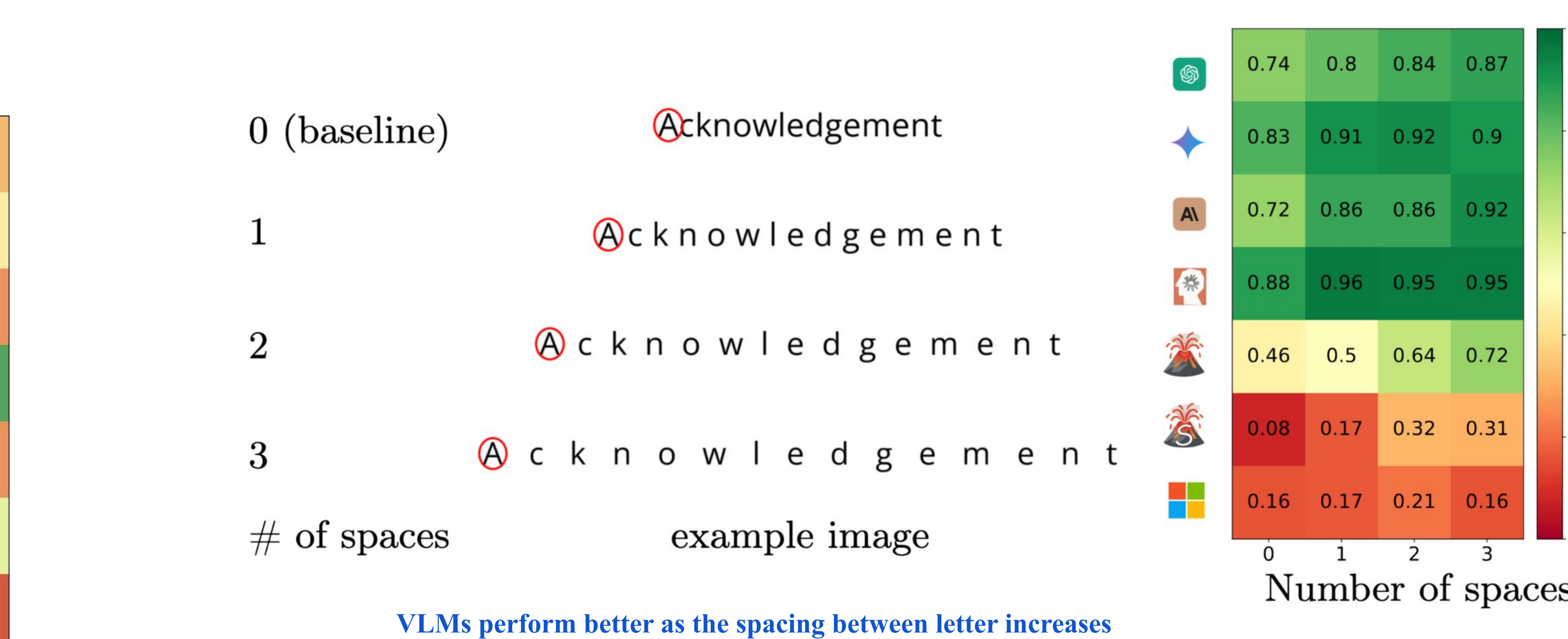
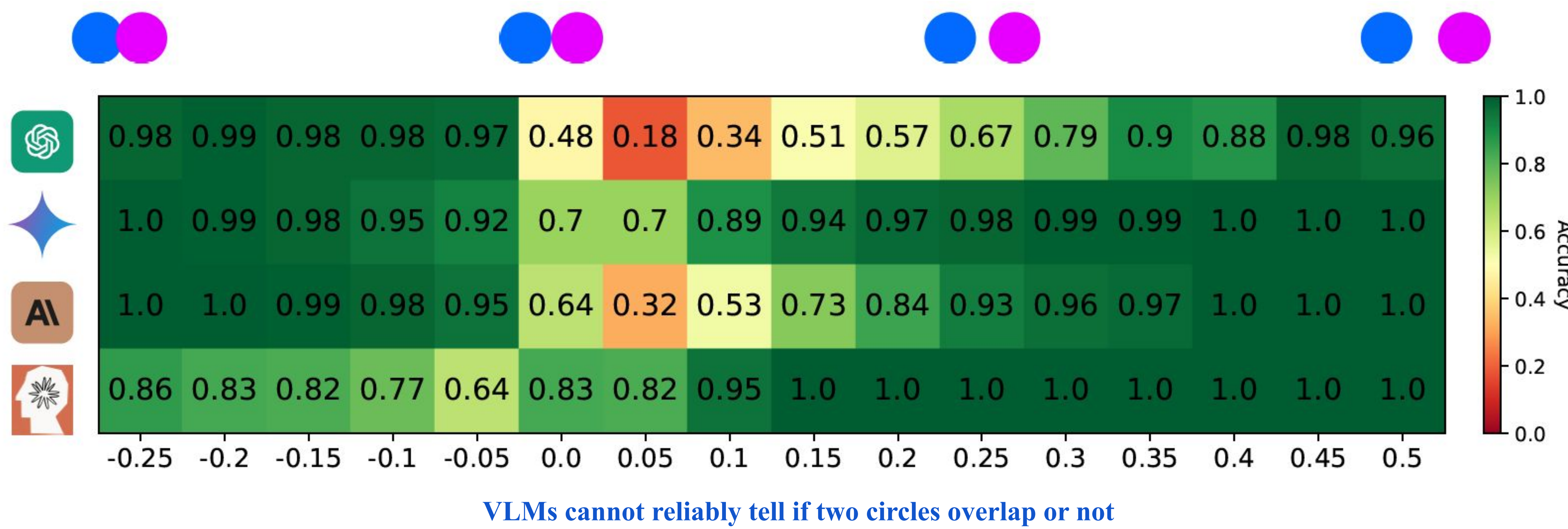
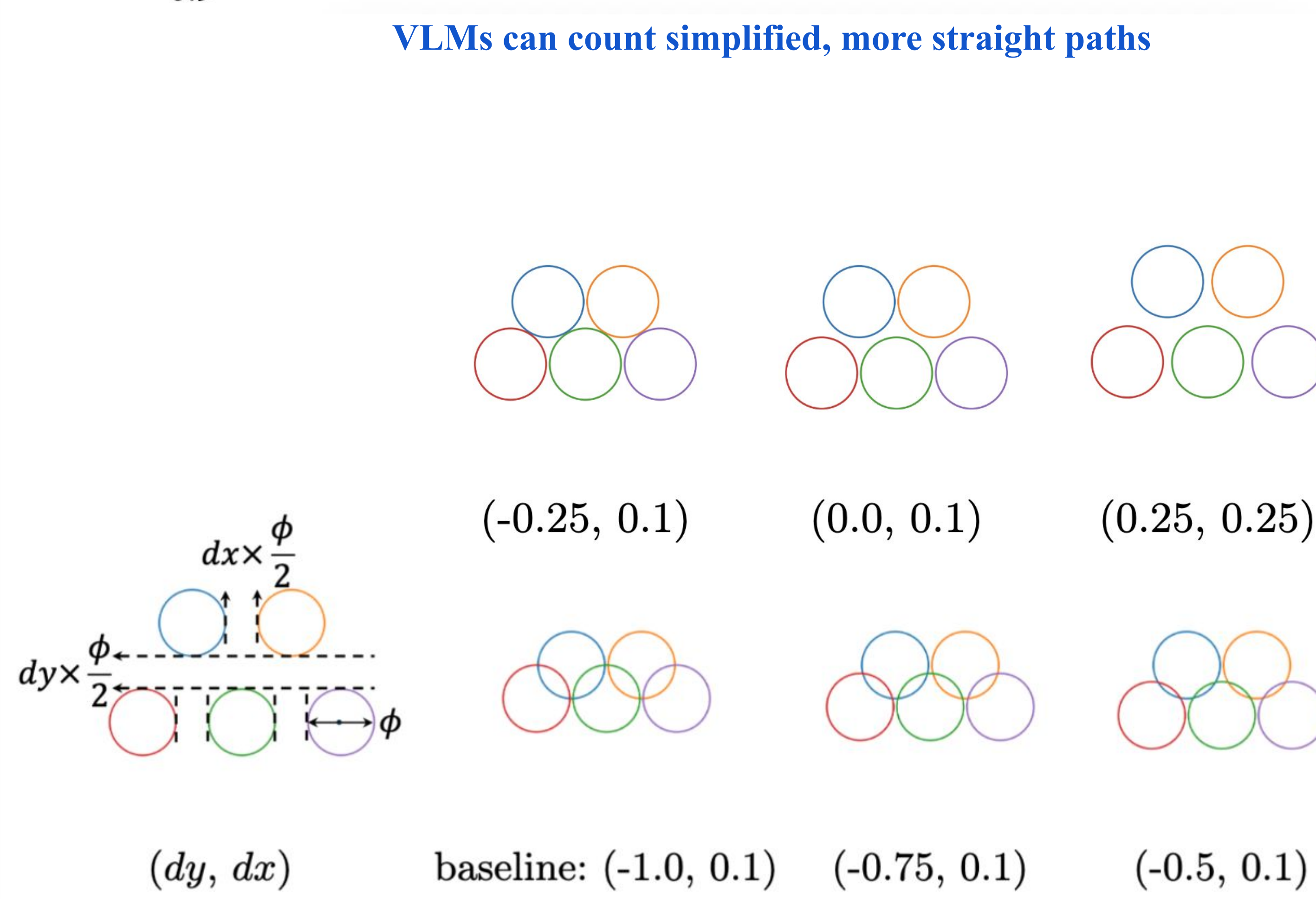
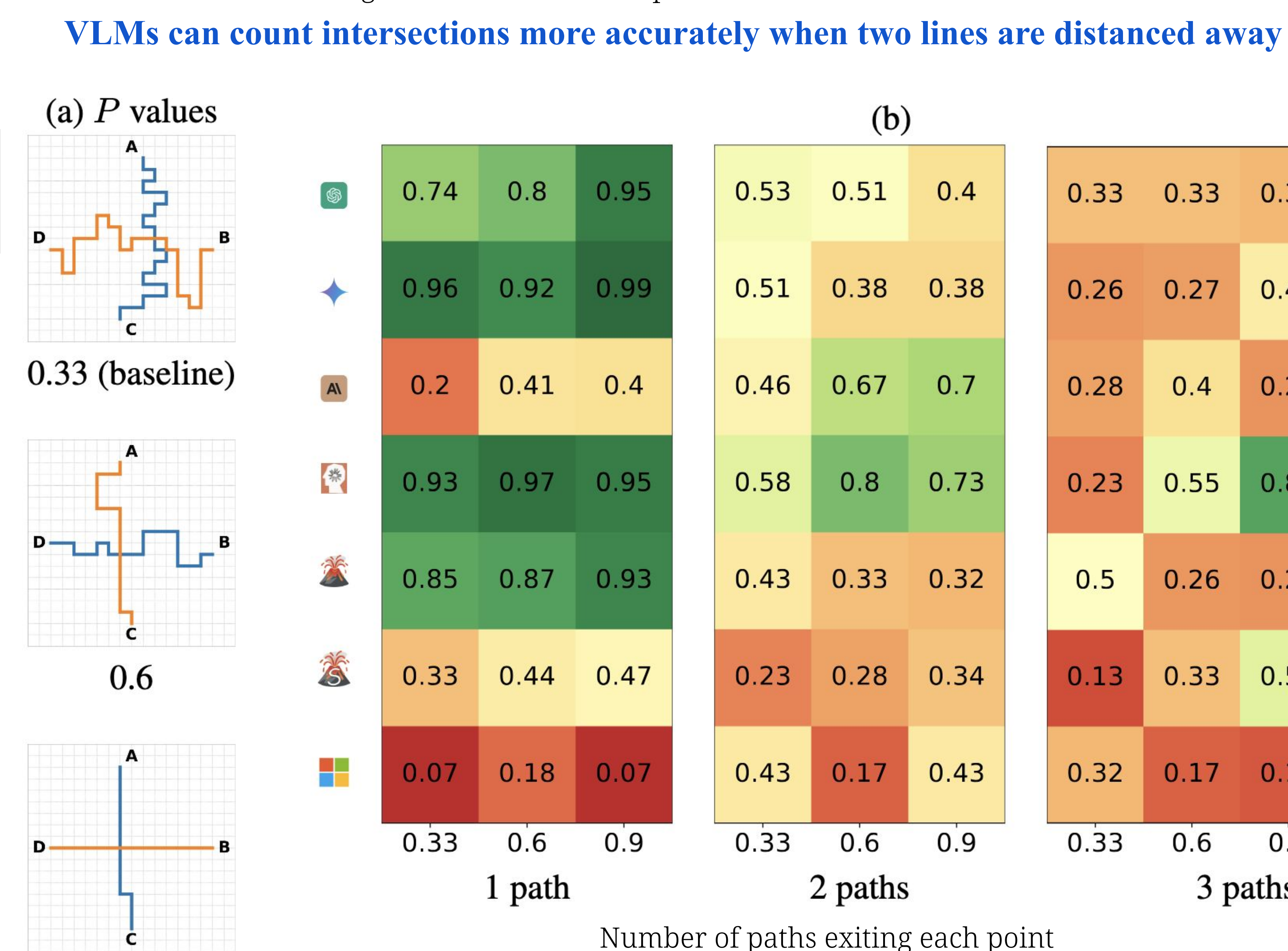
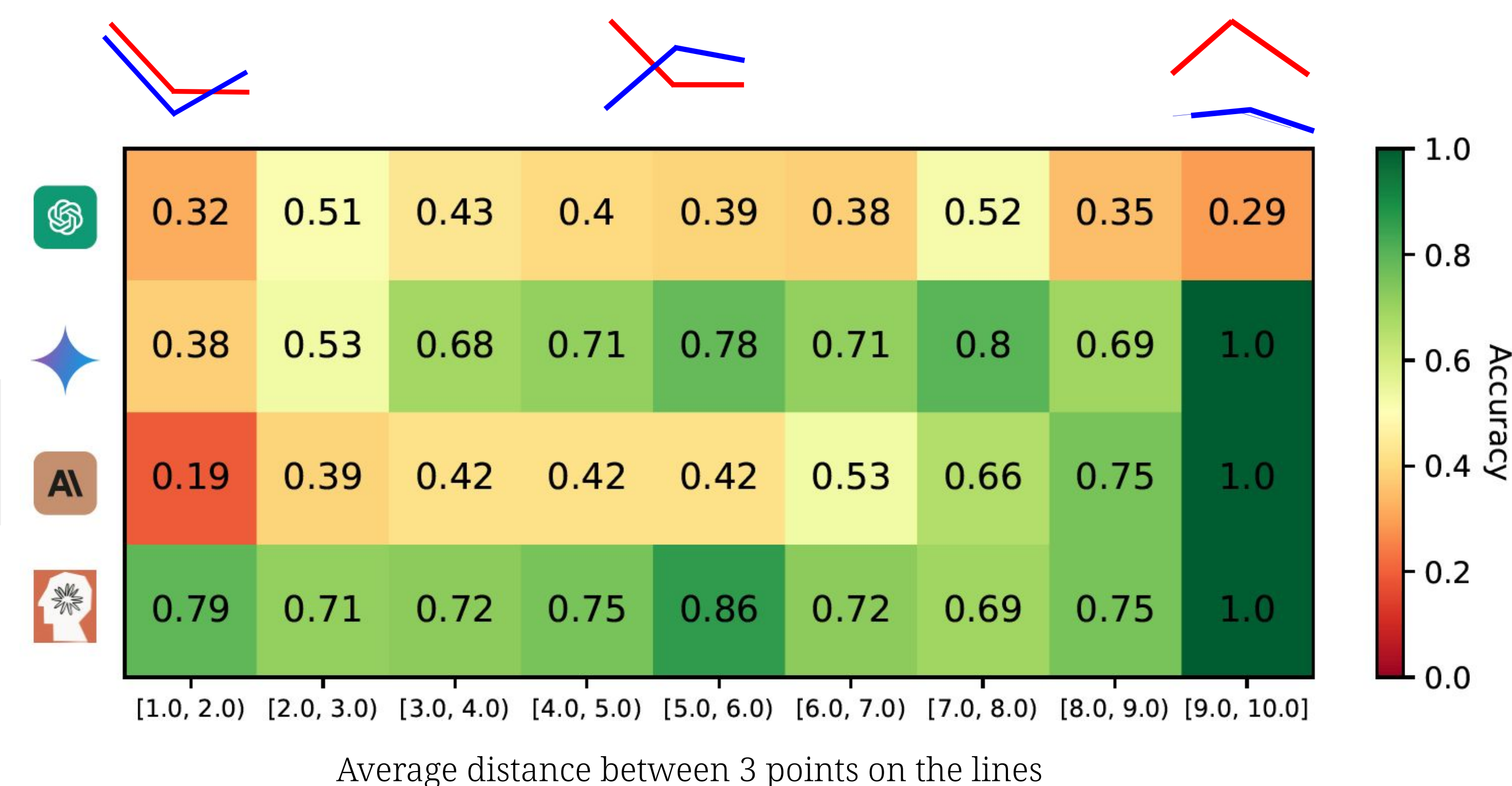
P4: How many circles are in the image? Answer with only the number in numerical format.

P5: How many squares are in the image? Please answer with a number in curly brackets e.g., {10}.

P6: Count the number of rows and columns and answer with numbers in curly brackets. For example, rows={5} columns={6}.

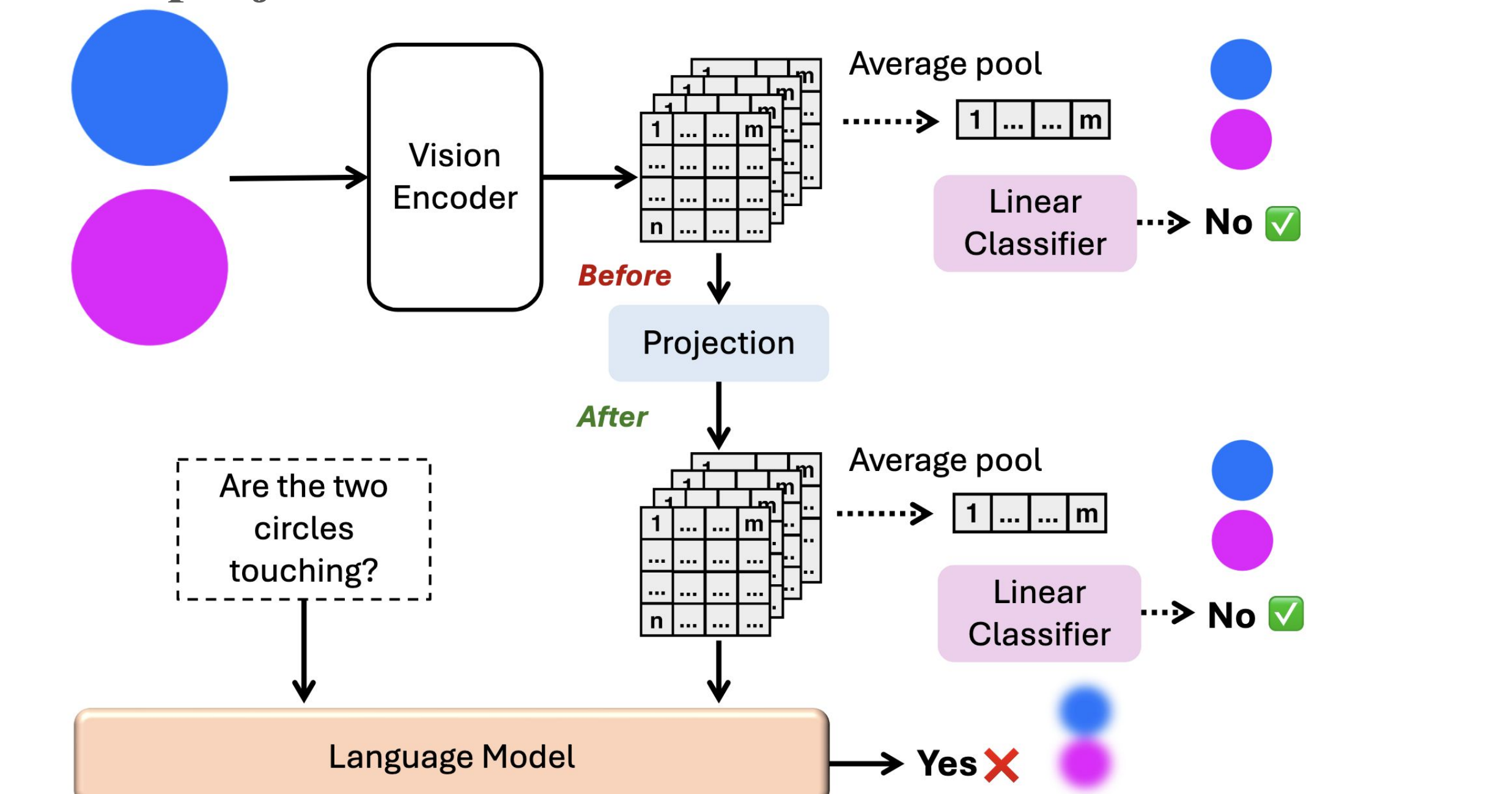
P7: How many single-color paths go from A to D? Answer with a number in curly brackets e.g. {3}.

Results



Where do VLMs fail?

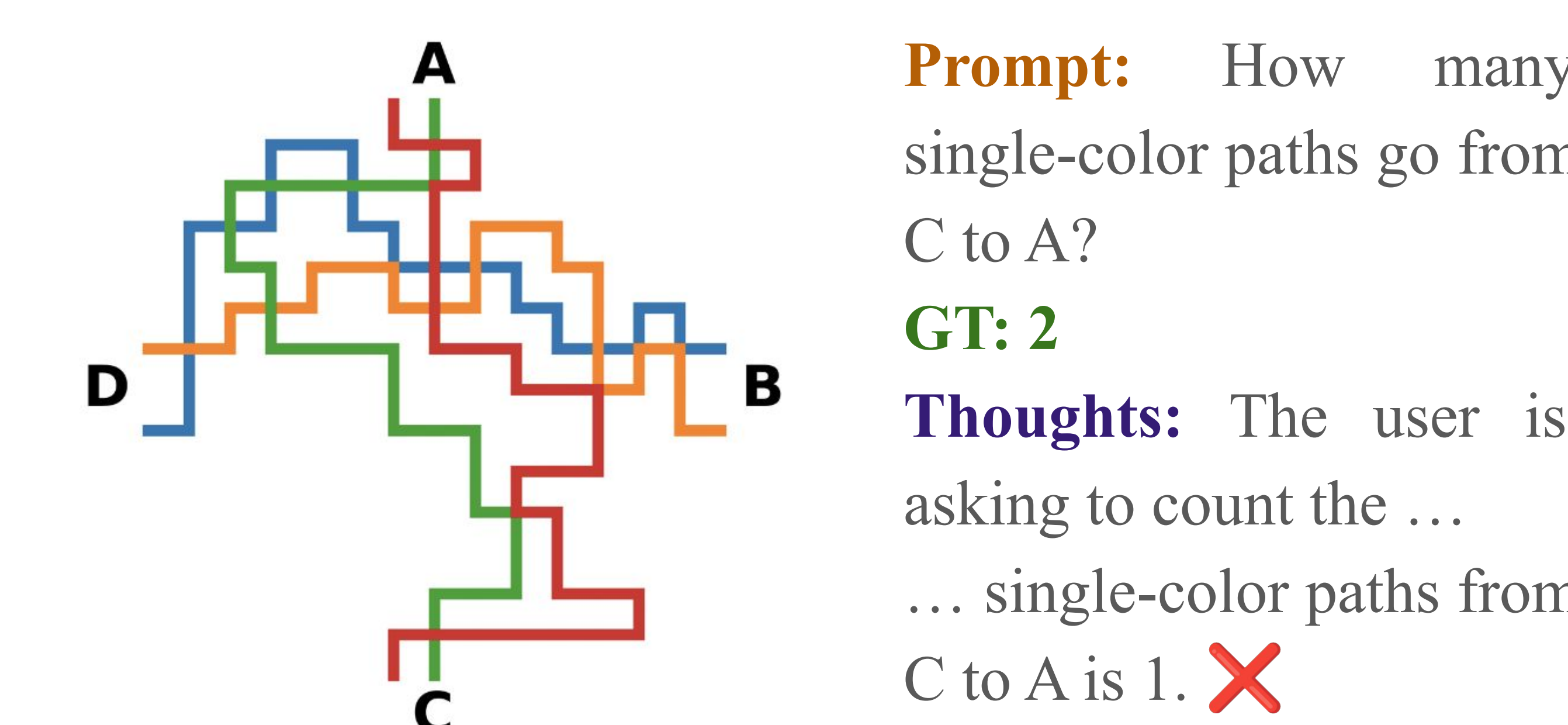
- We **probe** the frozen features at 4 levels in LLaVA-OneVision 0.5B:
 - Before the projection **99.82%**
 - After the projection **99.73%**
 - After the first decoder **97.11%**
 - At the last decoder **99.52%**



Slow-thinking

- Although, the **performance bottleneck** is in the LLM's ability to **interpret the visual features**, the slow-thinking, i.e., **reasoning models, are on par with their non-reasoning counterparts**.

Gemini 2.0 Flash **72.75%** vs. Gemini 2.0 Flash-Thinking **71.59%**



Conclusion

- BlindTest** exposes a **low-level visual shortcoming** in SotA VLMs, i.e., failing to generate **correct answers from detailed visual features**.
- Our findings suggest that **slow-thinking capabilities** do **not improve the VLMs performance on BlindTest**.