



JOHNS HOPKINS

WHITING SCHOOL  
of ENGINEERING

# Does Rethinking Text as Vision Allow us to Scale LLMs Better?

**DeepSeek-OCR:** Contexts Optical Compression

**Glyph:** Scaling Context Windows via Visual Text Compression

## Authors:

- Wei et al. @ DeepSeek
- Cheng et al. @ Tsinghua University

# Outline

---

1. Background: OCR and VLMs
2. **Glyph**: Introduction & Motivation
3. **Glyph**: Method & Novelty
4. **Glyph**: Results
5. **DeepSeek OCR**: Introduction & Motivation
6. **DeepSeek OCR**: Method & Novelty
7. **DeepSeek OCR**: Results
8. **Discussion and Open Questions**



JOHNS HOPKINS

WHITING SCHOOL  
of ENGINEERING

# Some Background

# Long Context Problem

## The Problem: Long context problem.

- OOD for the model. (performance drop).
- Inference cost increase.
- GPU CUDA memory limit.

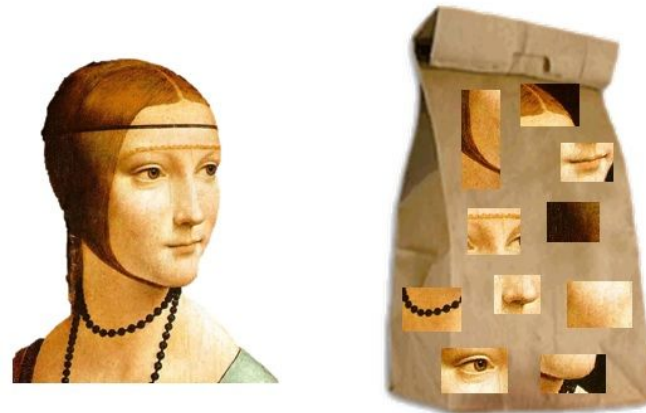
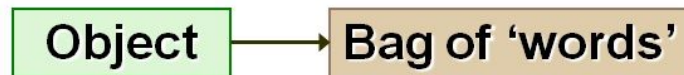
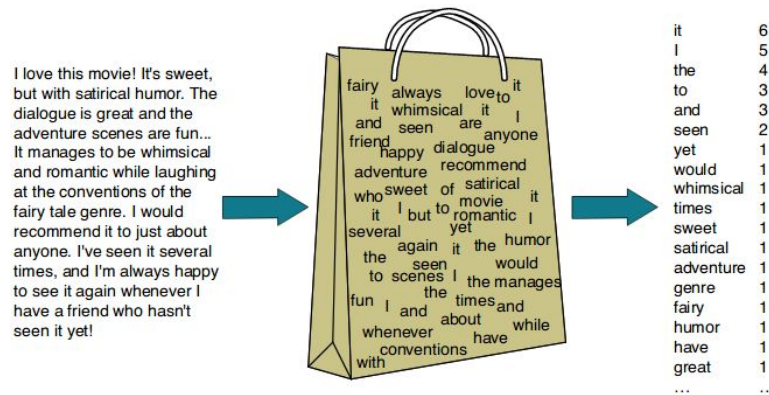
## Solutions:

- Extend position encoding. e.g. YaRN
- KV cache eviction, sparse attention.
- OCR: compress text into visual tokens. ← Today's papers

## Visual Tokens vs Text Tokens:

A picture have many patches, each patch has its own “local” semantic (i.e. a cat, a table edge), and have dependencies / connections in between. This is similar to tokens in natural language.

=> We can treat image as “bag of word”



# Background: VLM

We have good text LLM already. Can we leverage natural language data to understand image?  
Key: Representation can be shared.

Solution: Align visual token (of an image) to word token (of the caption) via contrastive pretraining.  
(=>CLIP)

Then we can fuse the visual encoder and the text encoder and just use a single LLM . (VLM)

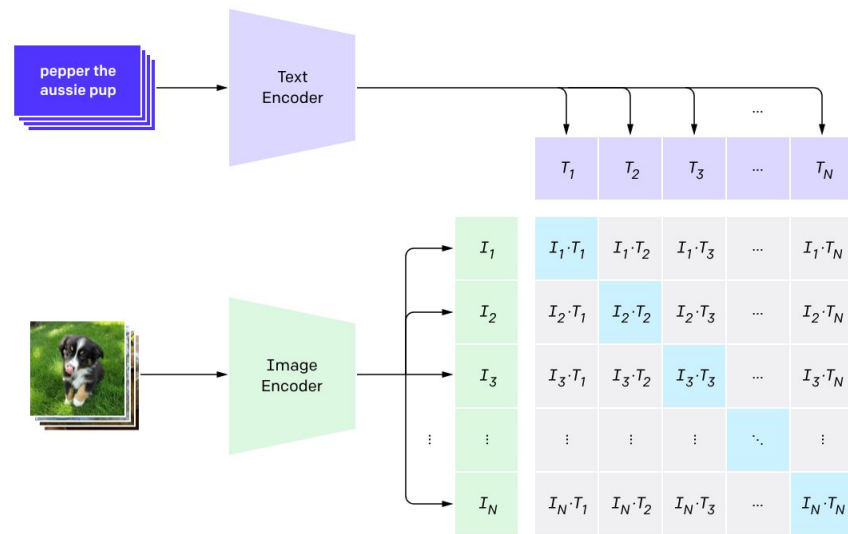
Now with a VLM, what can we get?

=> Generic image understanding.

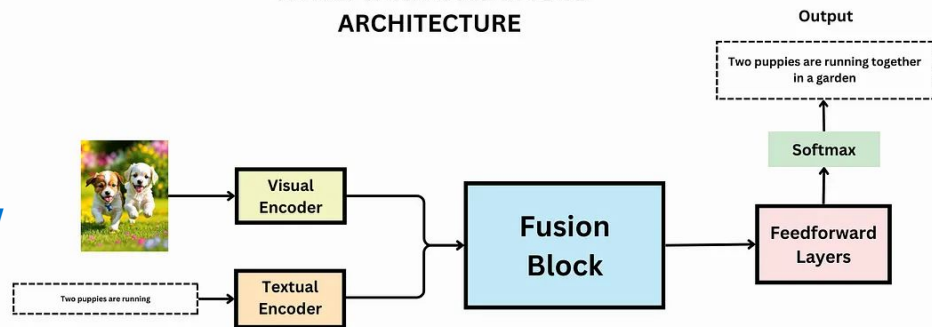
-> OCR ability!

-> Can we represent text as image? Just like how human read books.

## 1. Contrastive pre-training



## VISION LANGUAGE MODEL ARCHITECTURE





JOHNS HOPKINS

WHITING SCHOOL  
of ENGINEERING

# ***Glyph:** Scaling Context Windows via Visual Text Compression*

# Glyph: Introduction

## Motivating the Problem

Reading isn't a character-by-character computation; the **visual stream** aggregates shapes into **word forms** that the **linguistic stream** maps to meaning.

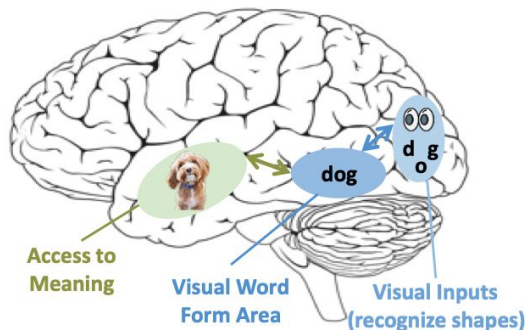
Pure text tokenizers ignore this pipeline and over-segment (especially across scripts and morphologies), inflating sequence length with units that are visually redundant.

If we let the model ingest the *visual* bundles directly, a few visual tokens can carry the same semantic payload as many subwords—compressing context without discarding meaning.

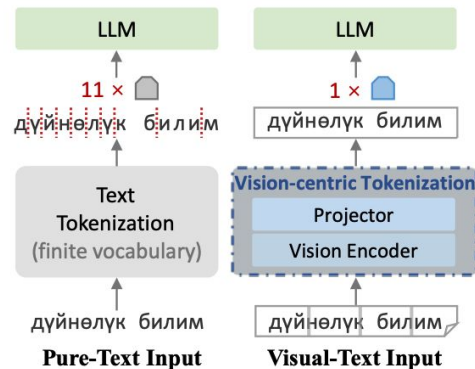
## Why Current Solutions Don't Work!

- **Extractive prompt compression (LLMLingua family):** Removes “less important” tokens to hit budgets; helpful but **task-dependent**.
- **Context-extension via RoPE scaling (YaRN / NTK-aware):** Increases the **window size** but **doesn't reduce tokens**—compute and memory still grow with the raw sequence length.

**Analogy** from prior work: **Visual processing** can *aggregate shape/word* cues into **fewer** units: used here only to motivate Glyph's visual-text compression.



Reading Activity in the Brain



# Glyph: Introduction & Motivation

## How Glyph Presents the Problem

Modern LLMs need to reason over **long inputs**, but compute and memory grow with **sequence length**, not semantic content density!

Scaling to **million-token** contexts is impractical due to prohibitive compute and memory; existing approaches either extend windows or prune tokens with trade-offs.

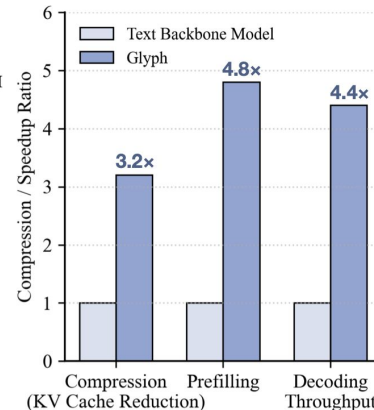
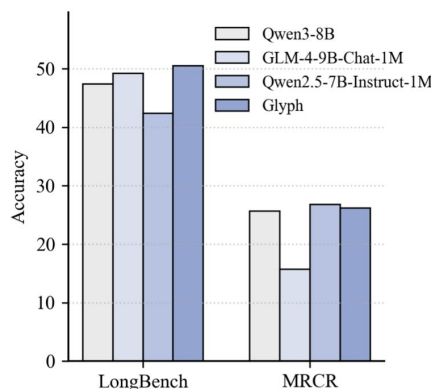
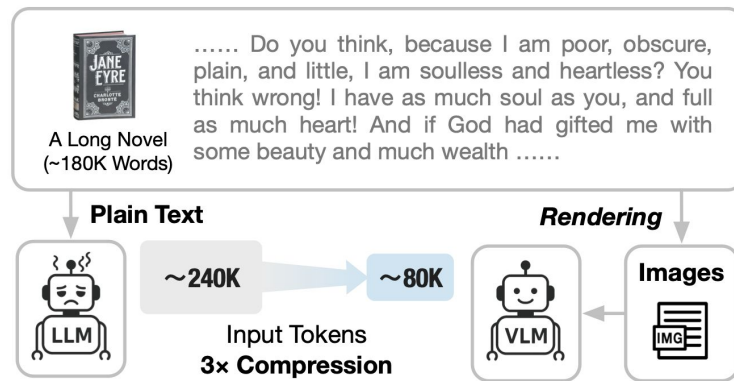
**Render** → **VLM**: Convert long text to images so each **visual token aggregates many words**, attacking the bottleneck by changing the **representation**, not just the window.

## What it Proposes to Solve the Problem

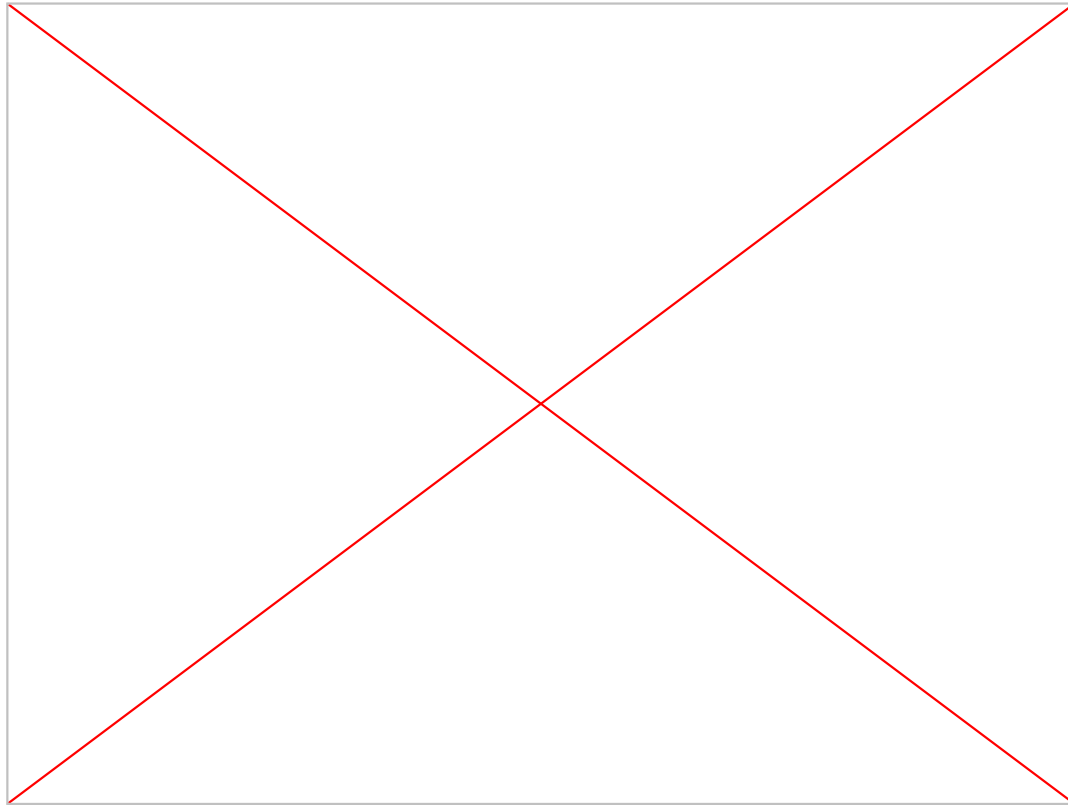
**Framework**: A practical pipeline that **renders long text to images** and processes them with a VLM, enabling large effective context with fewer tokens and stable accuracy.

**LLM-driven rendering search**: A **genetic search** over fonts/layout/DPI to optimize the **readability** ↔ **compression** trade-off automatically.

**Training recipe + scaling result**: **Continual pre-training** on rendered corpora → **post-training (SFT/RL)** with an **auxiliary OCR-style loss**.



# Glyph: Demo to Motivate



# Some Basics

Any task is defined by the triplet:  $(\mathcal{I}, \mathcal{C}, \mathcal{R})$  - **I** is a user instruction **that specifies** the core goal **C**, where  $\mathcal{C} = \{c_1, \dots, c_T\}$  is an **ultra-long** textual context.

**R** is target response.

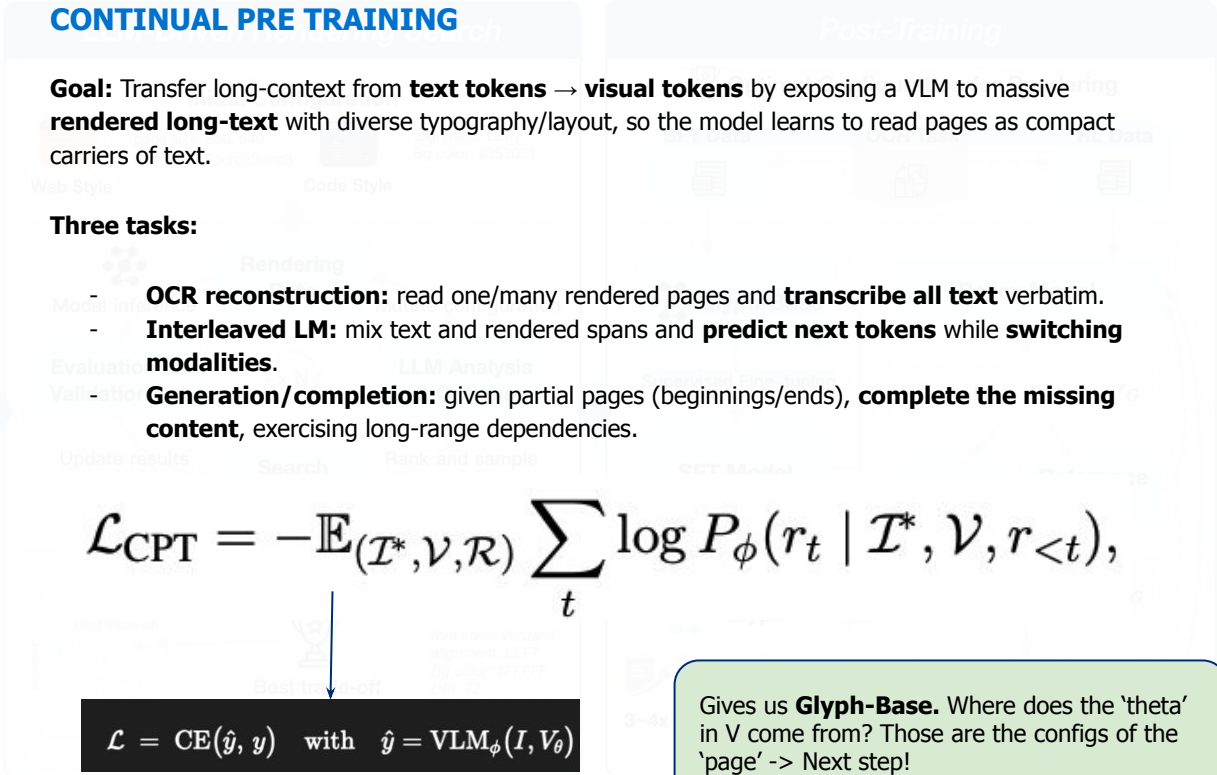
Traditionally, we would do:  $P(R \mid I, C)$ .

**GLYPH: Render**  $\mathcal{C}$  as visual pages  $\mathcal{V} = \{v_1, \dots, v_n\}$  each containing *glyph's* of multiple text segments. Hence, Glyph does:

**P (R | I, V)** and a task is defined by  $(I, V, R)$



© 2006 The Authors  
Journal compilation © 2006 Blackwell Publishing Ltd



Gives us **Glyph-Base**. Where does the 'theta' in V come from? Those are the configs of the 'page' -> Next step!

# Glyph: Method

## Rendering Search

**Why?** Remember the photos are not 'real' but instead generated from text. Need configs for how the photo 'looks'.

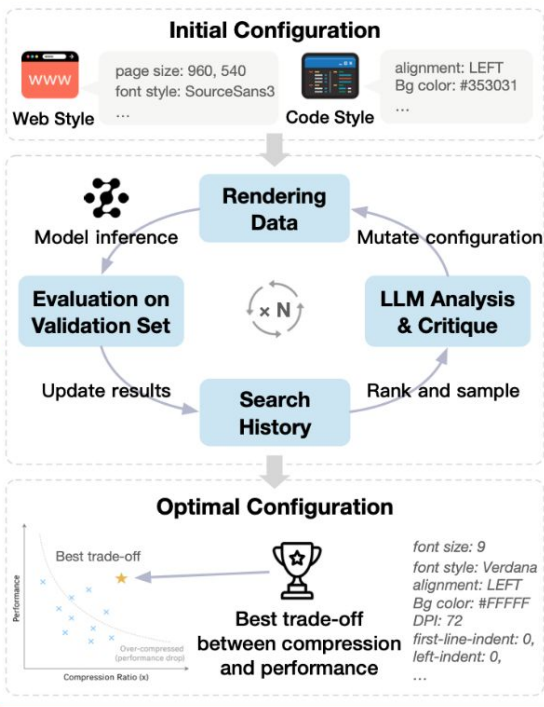
Examples:

- DPI & Page Size
- Font Family and Size
- Borders and Spacing
- Colors

Parametrized by  $\theta$ . Controlling this **controls compression ratio from text to vision**.

$$\rho(\theta) = \frac{|\mathcal{C}|}{\sum_{i=1}^n \tau(v_i)},$$

## LLM-Driven Rendering Search



## Genetic Algorithm Approach

**What is being evolved.** A rendering genome  $\theta$ . Each  $\theta$  renders the same long text into images  $V\theta$ , changing **compression** (visual tokens/page) vs **readability**.

**Mutation-only evolutionary loop.** Start from seed styles (web/code).

For each generation: **Render**, external LLM **critiques** outputs (where/why errors happen, readability issues), **Rank & sample** promising candidates, **Mutate** parameters.

**Objective/selection signal.** Keep candidates that **maximize task utility under high compression**, effectively searching the **Pareto front** between accuracy and compression ratio; the chosen  $\theta^*$  is the **best trade-off**.

# Glyph: Method

## Post Training

Two tasks to improve “**Glyph-Base**”: Supervised FT and RL & Auxiliary OCR Alignment Task

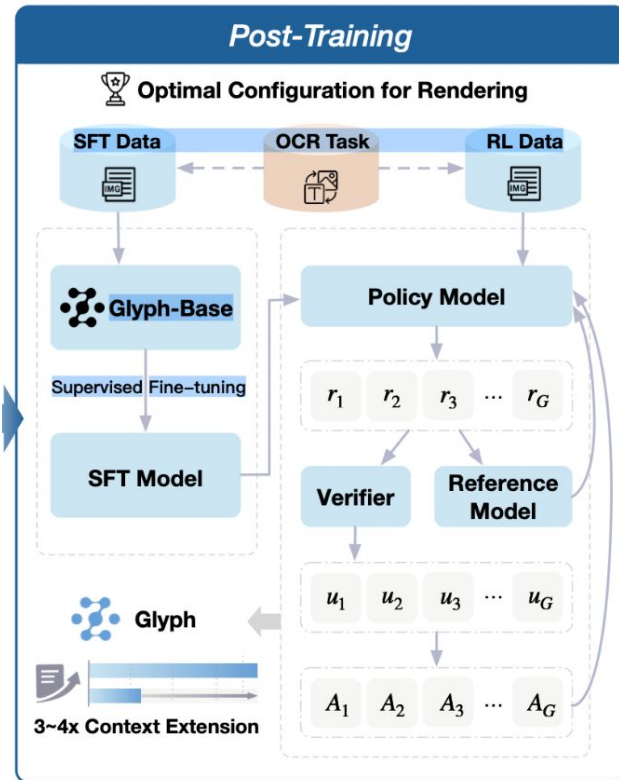
**SFT**: Similar to what you would do with a Language Model **except C is now V**. Loss function is written as:

$$\mathcal{L}_{\text{SFT}} = -\mathbb{E}_{(\mathcal{I}, \mathcal{V}, \mathcal{R})} \sum_t \log P_{\phi}(r_t | \mathcal{I}, \mathcal{V}, r_{<t}),$$

**RL**: Again, like basic LLM training, they **use GRPO** (C is now V). Importance sampling weight defined as:

$$w_i = \frac{\pi_{\phi}(r_i | \mathcal{I}, \mathcal{V})}{\pi_{\phi_{\text{old}}}(r_i | \mathcal{I}, \mathcal{V})} \cdot \begin{matrix} \text{Reward } r \text{ is } \{0,1\} \\ \text{integrating Verifiable} \\ \text{Rewards and Format} \\ \text{Rewards.} \end{matrix} \rightarrow A_i = \frac{u(r_i) - \text{mean}(\{u(r_j)\}_{j=1}^G)}{\text{std}(\{u(r_j)\}_{j=1}^G)},$$

$$\mathcal{J}_{\text{GRPO}}(\phi) = \mathbb{E}_{(\mathcal{I}, \mathcal{V}) \sim P, \{r_i\}_{i=1}^G \sim \pi_{\phi_{\text{old}}}} \left[ \frac{1}{G} \sum_{i=1}^G \left( \min(w_i A_i, \text{clip}(w_i, 1 - \epsilon_l, 1 + \epsilon_h) A_i) - \beta D_{\text{KL}}(\pi_{\phi} \| \pi_{\text{SFT}}) \right) \right],$$



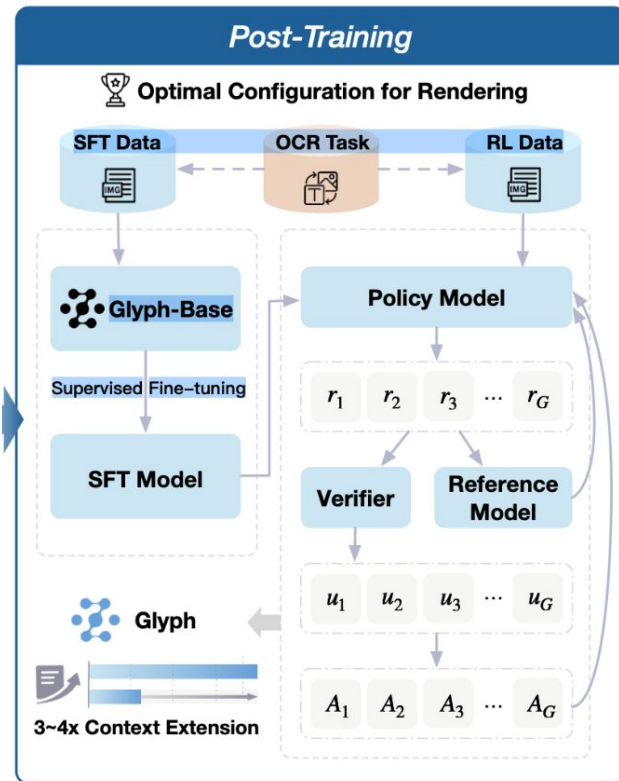
# Glyph: Method

## Post Training OCR Alignment

Reconstruct all text from a generated 'V' and fine tune to improve on that task - similar to **Continual Pre-Training Step**.

They do this in both the SFT and RL Step:

1. For SFT - **Its same** as Pre-Training. Just minimize CE.
2. For RL - **Its** basic GRPO but reward is Levenshtein Distance.



# Glyph: Controllable Factors and Sampling

| Factor      | Specification / Sampling Strategy                                                                                                                                                                                              |                |                                                                                                                                  |
|-------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------|----------------------------------------------------------------------------------------------------------------------------------|
| dpi         | Mixture of sets: <i>lowest</i> (45–59), <i>low</i> (60–71), <i>medium</i> (72–119), <i>normal</i> ({72,80,96,100,110,120,144,150,300}), <i>high</i> (over 300); favor normal/medium with small probability spikes to extremes. | alignment      | LEFT/JUSTIFY (dominant) with small-prob. RIGHT/CENTER.                                                                           |
|             |                                                                                                                                                                                                                                | margins        | Three patterns: all-equal; vertical-larger; horizontal-larger; values in 10–40pt ranges.                                         |
|             |                                                                                                                                                                                                                                | indent         | Modes: none; first-line indent ( $\approx 1$ –2.5 em); block/hanging with left/right indents.                                    |
| page_size   | (i) Fixed paper sizes (A4, Letter, Legal, A5, B5, A3, B4, Tabloid) with priors; (ii) common aspect ratios (e.g., 1.414, 1.333, 1.5, 1.778); (iii) fully random aspect via piecewise distribution (narrow $\rightarrow$ tall).  | spacing        | space-before/space-after use a multi-mode prior (none, small, large).                                                            |
|             |                                                                                                                                                                                                                                | h_scale        | Horizontal glyph scaling (0.75–1.0) with decaying probabilities.                                                                 |
|             |                                                                                                                                                                                                                                | colors         | Page/background/font palettes for light/dark themes; document/web/code styles inherit coherent triplets (page, paragraph, font). |
| font_family | Pooled and deduplicated families across serif/sans/mono/pixel; italics sampled by filename heuristics (suffixes, italic/oblique).                                                                                              | borders        | Optional paragraph borders with width/padding; disabled by default.                                                              |
|             |                                                                                                                                                                                                                                | newline_markup | With small probability, explicit markers (e.g., <code>\n</code> , tags, or tokens) inserted to preserve structure.               |
| font_size   | {7, 7.5, 8, 9, 9.5, 10, 11, 12, 14}; line_height tied as font_size + {0, ..., 3}.                                                                                                                                              | auto_crop      | Optional white-margin cropping and last-page trimming.                                                                           |



JOHNS HOPKINS

WHITING SCHOOL  
of ENGINEERING

# Experiments and Results

# Glyph: Experimental Setup

## Benchmarks:

- **LongBench**: a bilingual , multitask long context understanding benchmark
- **MRCR**: A multi-document / “needle in haystack” style long-context retrieval QA benchmark.
- **Ruler**: long-context benchmark with various sub-tasks including UUID recognition etc (where they ignore the UUID recognition).
- **MMLongBench-Doc**: A long document understanding benchmark (multimodal PDF documents with images/layout) used for cross-modal generalisation.

Backbone VLM: GLM-4.1V-9B-Base

# Glyph: Results

Main performance  
on LongBench =>  
Matches GPT-4.1

(And in Appendix  
the rest results)

More degradation in  
Single-Doc QA  
(Glyph has a more  
blurry  
understanding?)  
and summarization.

| Model                  | Single-Doc QA |              | Multi-Doc QA |              | Summarization |              | Few-shot     |              | Synthetic     |               | Code         |              | Avg          |
|------------------------|---------------|--------------|--------------|--------------|---------------|--------------|--------------|--------------|---------------|---------------|--------------|--------------|--------------|
|                        | QP            | NQA          | HQA          | 2QA          | QSUM          | GovRep       | TREC         | TriQA        | PR Zh         | PR En         | RB           | LCC          |              |
| GPT-4.1                | 51.60         | 35.73        | 69.10        | 74.15        | 23.50         | 33.36        | 77.00        | 93.36        | 100.00        | 100.00        | 67.94        | 68.43        | 56.03        |
| LLaMA-3.1-8B-Instruct  | 44.56         | 26.34        | 56.88        | 46.67        | <b>23.28</b>  | <b>32.36</b> | 19.25        | <u>89.12</u> | 62.20         | <u>99.50</u>  | 42.81        | 46.35        | 41.34        |
| Qwen2.5-7B-Instruct-1M | <b>45.29</b>  | 25.61        | 60.70        | 40.51        | <u>22.95</u>  | <u>29.97</u> | 59.37        | 86.93        | <u>98.5</u>   | <b>100.00</b> | 29.80        | 21.72        | 42.42        |
| Qwen3-8B               | <u>44.67</u>  | 26.13        | <u>65.83</u> | <b>73.92</b> | 19.60         | 26.85        | <u>70.50</u> | 87.98        | <b>100.00</b> | 97.26         | 40.89        | 44.87        | 47.46        |
| GLM-4-9B-Chat-1M       | 43.75         | <u>26.72</u> | 58.98        | 50.89        | 22.84         | 27.60        | 61.50        | <b>90.07</b> | <b>100.00</b> | <u>99.50</u>  | <u>55.64</u> | <b>59.54</b> | <u>49.27</u> |
| Glyph                  | 40.64         | <b>28.45</b> | <b>66.42</b> | <u>72.98</u> | 19.78         | 25.53        | <b>82.62</b> | 88.54        | 89.03         | <u>99.50</u>  | <b>60.80</b> | <u>48.85</u> | <b>50.56</b> |

Table 1: Performance comparison of Glyph with leading LLMs on LongBench (%). Our model achieves competitive results in the overall average score. Best results are **bolded**, and second-best are underlined. Refer to Table 10 for the rest of the results.

| Model                  | Single-Doc QA |              | Multi-Doc QA |              | Summarization |              | Few-shot     |              | Synthetic    |
|------------------------|---------------|--------------|--------------|--------------|---------------|--------------|--------------|--------------|--------------|
|                        | QA Zh         | QA En        | Mus          | Dur          | News          | VcSum        | Sam          | Lsht         |              |
| GPT-4.1                | 63.90         | 51.27        | 55.63        | 24.58        | 23.70         | 14.66        | 41.25        | 50.00        | 26.5         |
| LLaMA-3.1-8B-Instruct  | 62.20         | <b>54.98</b> | 31.61        | <b>33.75</b> | <b>24.21</b>  | <b>16.23</b> | 7.61         | 0.00         | 7.13         |
| Qwen3-8B               | 60.98         | 49.78        | <u>45.54</u> | 16.69        | 18.55         | 12.08        | <u>36.47</u> | 42.00        | <u>12.81</u> |
| GLM-4-9B-Chat-1M       | <b>63.17</b>  | 52.88        | 39.14        | <u>28.27</u> | 23.90         | 16.21        | 36.15        | <b>47.38</b> | 2.39         |
| Qwen2.5-7B-Instruct-1M | <u>62.98</u>  | <u>53.62</u> | 34.72        | 21.85        | 21.02         | 12.20        | <b>39.17</b> | 28.68        | 3.50         |
| Glyph                  | 37.23         | 45.89        | <b>56.18</b> | 26.87        | 21.52         | 12.43        | 32.49        | <u>44.43</u> | <b>30.50</b> |

Table 10: The rest of the results on LongBench benchmark (%), which encompasses Single-Document QA, Multi-Document QA, Summarization, Few-shot Learning, and Synthetic task.

# Glyph: Results

Main performance  
on MRCR => small  
gap below GPT-4.1  
match open-source.

2 needle is less  
performant (so they  
defer to appendix  
and didn't explain  
why)

| Model                  | 4 Needle     |              |              |              |              |              | 8 Needle     |              |              |              |              |              |
|------------------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|
|                        | 0k-8k        | 8k-16k       | 16k-32k      | 32k-64k      | 64k-128k     | Avg          | 0k-8k        | 8k-16k       | 16k-32k      | 32k-64k      | 64k-128k     | Avg          |
| GPT-4.1                | 50           | 38           | 29           | 42           | 38           | 39.4         | 33           | 26           | 17           | 22           | 19           | 23.4         |
| LLaMA-3.1-8B-Instruct  | <u>33.42</u> | <u>25.97</u> | <u>22.73</u> | <b>26.97</b> | 12.68        | <u>24.35</u> | <u>23.80</u> | 17.69        | <b>19.85</b> | <u>17.72</u> | 11.79        | <b>18.17</b> |
| Qwen2.5-7B-Instruct-1M | 25.96        | 20.13        | 19.93        | 24.25        | <u>17.29</u> | 21.51        | 17.64        | 19.48        | 12.41        | 14.80        | <u>14.24</u> | 15.71        |
| Qwen3-8B               | 29.34        | 22.67        | 20.34        | 23.63        | <b>19.11</b> | 23.02        | 18.75        | <u>19.69</u> | <u>16.81</u> | <b>17.86</b> | <b>15.00</b> | 17.62        |
| GLM-4-9B-Chat-1M       | 15.17        | 13.78        | 9.18         | 20.27        | 15.05        | 14.69        | 14.55        | 9.65         | 9.34         | 9.47         | 8.97         | 10.40        |
| Glyph                  | <b>35.44</b> | <b>26.82</b> | <b>24.15</b> | <u>25.69</u> | 16.37        | <b>25.81</b> | <b>25.12</b> | <b>21.22</b> | 16.43        | 13.91        | 13.51        | <u>18.14</u> |

Table 2: Performance comparison of our model against leading LLMs on the 4-needle and 8-needle sub-tasks of the MRCR benchmark (%). Our method consistently ranks first or second across most settings while preserving about  $3\times$  compression ratio. Performance on the 2-needle task is deferred to the Appendix.

| Model                  | 2 Needle     |              |              |              |              |              |
|------------------------|--------------|--------------|--------------|--------------|--------------|--------------|
|                        | 0k-8k        | 8k-16k       | 16k-32k      | 32k-64k      | 64k-128k     | Avg          |
| GPT-4.1                | 83           | 72           | 67           | 62           | 59           | 68.6         |
| LLaMA-3.1-8B-Instruct  | <u>54.27</u> | <b>53.21</b> | <b>51.05</b> | <u>29.81</u> | <u>24.98</u> | <u>42.66</u> |
| Qwen3-8B               | <b>58.95</b> | 41.18        | 36.18        | 24.99        | 20.89        | 36.44        |
| GLM-4-9B-Chat-1M       | 39.77        | 15.87        | 18.42        | 18.63        | 18.42        | 22.22        |
| Qwen2.5-7B-Instruct-1M | 45.92        | <u>51.07</u> | <u>46.97</u> | <b>34.67</b> | <b>37.57</b> | <b>43.24</b> |
| Glyph                  | 41.51        | 40.78        | 39.58        | 29.67        | 22.41        | 34.85        |

Table 9: Performance of various models on the MRCR task (%) with the 2 Needle setting across different context length intervals (0k–8k, 8k–16k, 16k–32k, 32k–64k, 64k–128k) and the average score.

# Compression Ratio

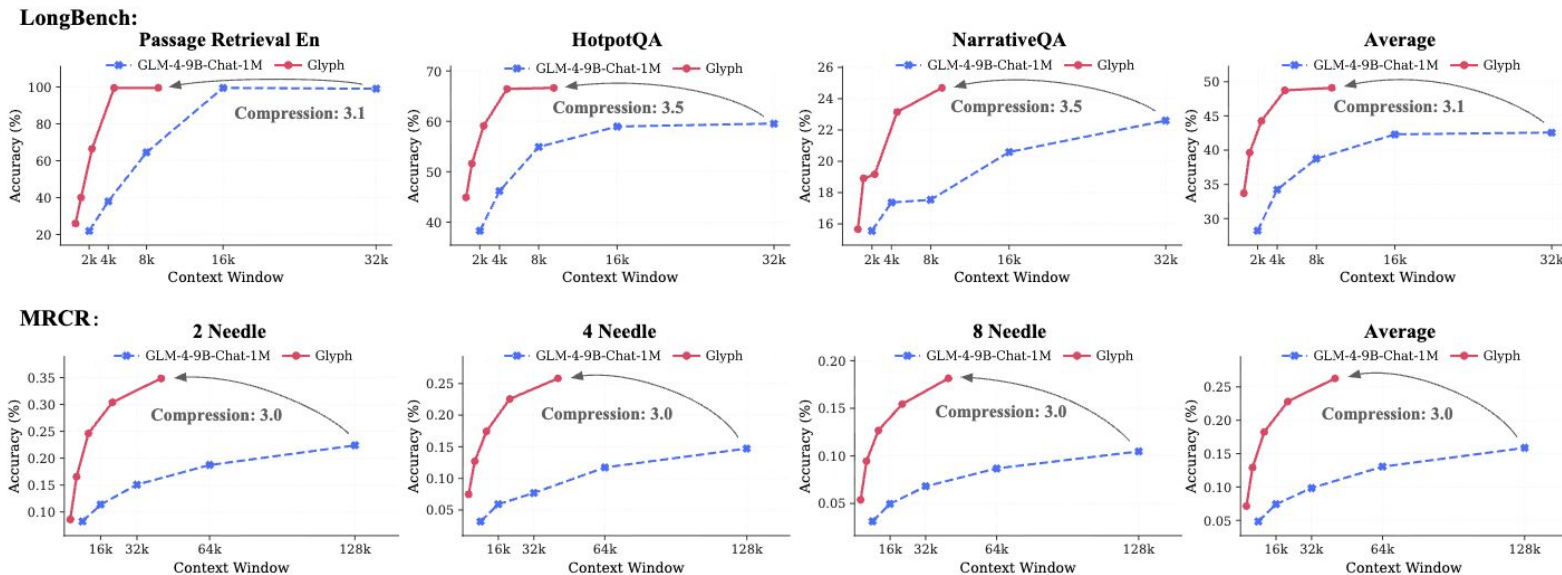


Figure 3: Performance comparison of Glyph and the baseline across different context windows, demonstrating that Glyph achieves performance equivalent to longer contexts with substantially shorter context windows.

# Varying Compression Ratio.

Higher DPI → larger rendered doc → More 14\*14 pixel patch → More visual token.

| Model                                               | Niah-S1       | Niah-S2       | Niah-M1       | Niah-M2      | Niah-V       | Niah-Q        | VT           | CWE          | FWE          | QA-1         | QA-2         | Avg          |
|-----------------------------------------------------|---------------|---------------|---------------|--------------|--------------|---------------|--------------|--------------|--------------|--------------|--------------|--------------|
| GPT-4.1                                             | 100.0         | 98.85         | 100.0         | 100.0        | 99.67        | 100.0         | 100.0        | 97.87        | 98.66        | 86.82        | 77.47        | 96.30        |
| LLaMA-3.1-8B-Instruct                               | 99.33         | 99.33         | 99.33         | <b>99.00</b> | 98.17        | <u>99.67</u>  | 87.07        | 57.30        | 81.85        | <b>84.00</b> | 58.00        | 87.55        |
| Qwen2.5-7B-Instruct-1M                              | <b>100.00</b> | <u>99.67</u>  | <u>99.67</u>  | <b>99.00</b> | 93.83        | 98.75         | 85.40        | 72.10        | 85.67        | 80.00        | 60.67        | 88.61        |
| Qwen3-8B                                            | <b>100.00</b> | <b>100.00</b> | 95.33         | 84.67        | 97.42        | 99.33         | <u>98.47</u> | 74.67        | 86.67        | 70.33        | 53.33        | 87.29        |
| GLM-4-9B-Chat-1M                                    | <b>100.00</b> | <b>100.00</b> | 92.67         | <b>99.00</b> | 95.00        | <b>100.00</b> | 98.20        | 49.50        | 83.22        | 72.67        | 56.67        | 86.08        |
| DPI: 72 / Compression rate: average 4.0, up to 7.7  |               |               |               |              |              |               |              |              |              |              |              |              |
| Glyph                                               | 73.33         | 64.67         | 67.33         | 56.00        | 73.42        | 71.42         | 77.93        | 94.40        | 92.67        | 59.33        | 63.33        | 72.17        |
| DPI: 96 / Compression rate: average 2.2, up to 4.4  |               |               |               |              |              |               |              |              |              |              |              |              |
| Glyph                                               | 98.00         | 95.33         | 95.67         | 85.00        | 96.33        | 95.83         | 94.93        | <u>94.80</u> | <u>98.00</u> | 79.00        | <u>70.67</u> | <u>91.23</u> |
| DPI: 120 / Compression rate: average 1.2, up to 2.8 |               |               |               |              |              |               |              |              |              |              |              |              |
| Glyph                                               | <u>99.67</u>  | 99.00         | <b>100.00</b> | <u>93.67</u> | <b>99.00</b> | 99.58         | <b>99.33</b> | <b>98.97</b> | <b>99.11</b> | 79.00        | <b>74.00</b> | <b>94.67</b> |

Table 3: Performance on the Ruler benchmark (%). We demonstrate the impact of different DPI settings on our model’s performance and the resulting compression ratios. For each configuration, the table includes both the average compression ratio across all sub-tasks and the maximum compression achieved for specific sub-task types.

# Efficiency

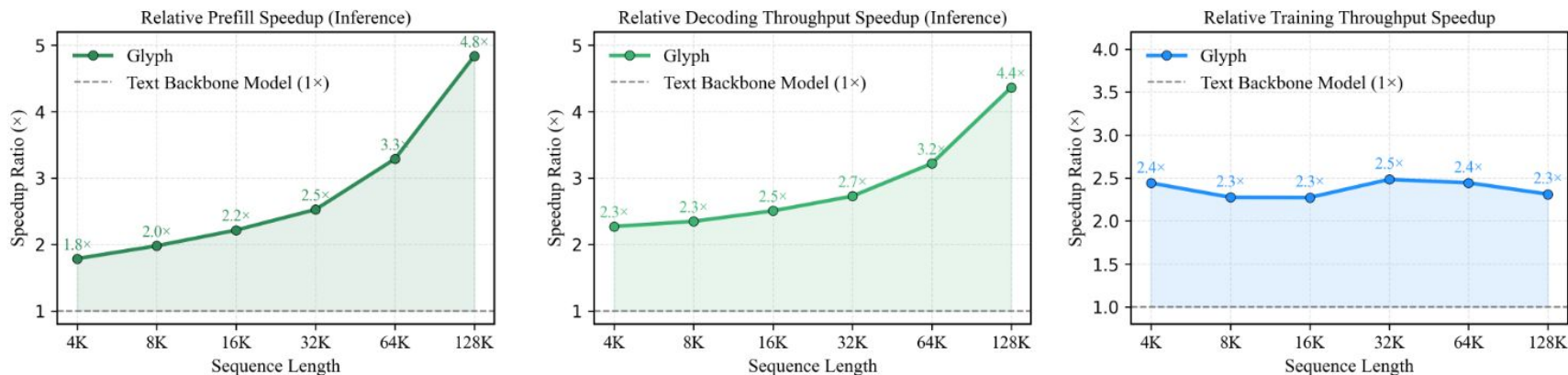


Figure 4: Speedup ratios of Glyph over the text backbone model for prefill, decoding, and training across different sequence lengths.

# As context grow...

Glyph has slower degradation.

(what about making axis the number of token? Is the degradation tied deeply to the attention mechanism - some retrieval efficiency upper bound)

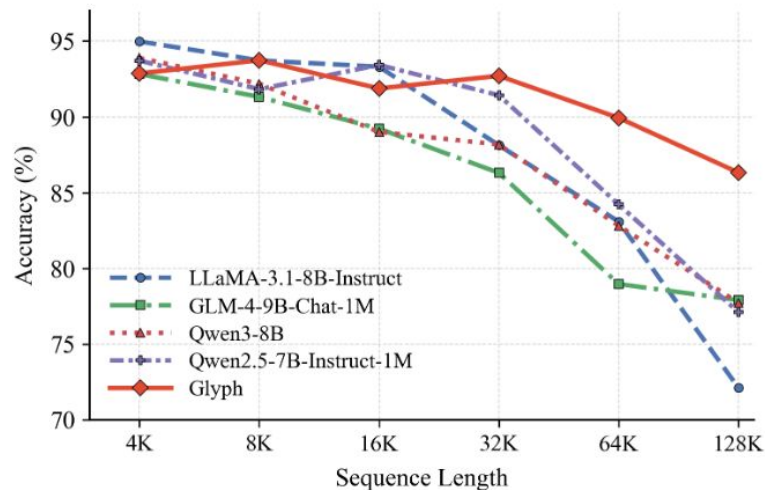


Figure 5: Model performance degradation across different sequence lengths on the Ruler benchmark.

# Better than base VLM

On MMLongBench-Doc  
(multimodal document  
understanding)

The full Glyph version  
outperform Base VLM.

→ Cross modality  
post-training/RL actually  
improves model's ability.

| Model            | SP           | CP           | UA           | Acc          | F1           |
|------------------|--------------|--------------|--------------|--------------|--------------|
| GLM-4.1V-9B-Base | 36.76        | <u>23.41</u> | <u>21.52</u> | 29.18        | 28.78        |
| Glyph-Base       | <u>47.91</u> | 22.24        | 14.80        | <u>32.48</u> | <u>34.44</u> |
| Glyph            | <b>57.73</b> | <b>39.75</b> | <b>27.80</b> | <b>45.57</b> | <b>46.32</b> |

Table 4: Results on MMLongBench-Doc (%). SP, CP, UA, and Acc denote Single-page, Cross-page, Unanswerable, and Overall Accuracy, respectively.

# Ablation: Is Genetic Search necessary

| Configuration       | LongBench    | MRCR         | Ruler        | Avg.         |
|---------------------|--------------|--------------|--------------|--------------|
| Random Config       | <u>41.78</u> | 15.82        | 65.13        | 40.91        |
| Manual Config       | <b>43.45</b> | <u>19.33</u> | <u>68.09</u> | <u>43.62</u> |
| Search-based Config | <b>43.45</b> | <b>22.10</b> | <b>71.24</b> | <b>45.60</b> |

Table 5: Ablation study comparing randomly combined, manually designed, and search-based configurations on three benchmarks under SFT setting. The search-based configuration achieves the best overall performance.

# Ablation: Is Aux. Objective necessary

| Model              | LongBench | MRCR  | Ruler |
|--------------------|-----------|-------|-------|
| Glyph              | 50.56     | 26.27 | 72.17 |
| – w/o OCR (in RL)  | -1.40     | -2.00 | -0.35 |
| – w/o RL           | -7.11     | -4.17 | -0.93 |
| – w/o OCR (in SFT) | -8.12     | -8.42 | -1.23 |

Table 6: Ablation study showing the performance drop (%) relative to the final Glyph model when components are progressively removed.

# Limitations

Sensitivity to rendering parameters.

- Model's performance noticeably affected by rendering configurations such as resolution, font, and spacing.

OCR-related challenges.

- UUID task especially challenging. (quote - even the strongest models [e.g., Gemini-2.5-Pro] often fail to reproduce them correctly.)

Task diversity.

- More tasks besides long-context understanding.
- e.g. agentic, reasoning.
- "We also observe that, compared with textual models, the visual-text model tends to generalize less effectively across tasks"

# Questions?

---

Is patch based visual token necessary or sufficient? (What are we actually gaining through using VLM, is it just a dense tokenization that encompass multiple token. )



JOHNS HOPKINS

WHITING SCHOOL  
of ENGINEERING

# *DeepSeek-OCR: Contexts Optical Compression*

# DeepSeek-OCR: Introduction & Motivation

- Significant compute challenges with long textual content.
  - One solution: use single image containing document text → rich information while using substantially fewer tokens.
- OCR tasks establish natural compression-decompression mapping between visual and textual representations
  - Motivated **DeepSeek-OCR**

## Traditional LLMs

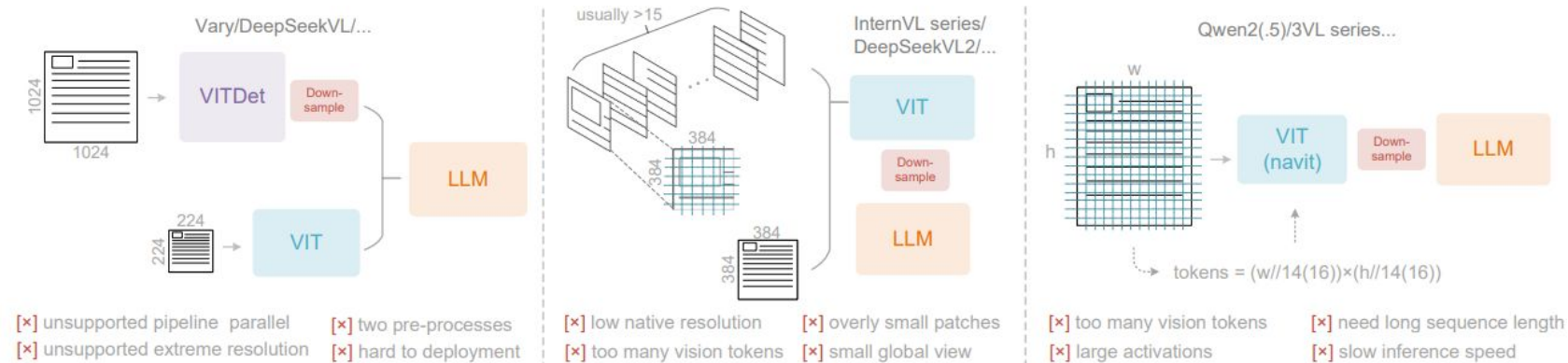


## DeepSeek OCR



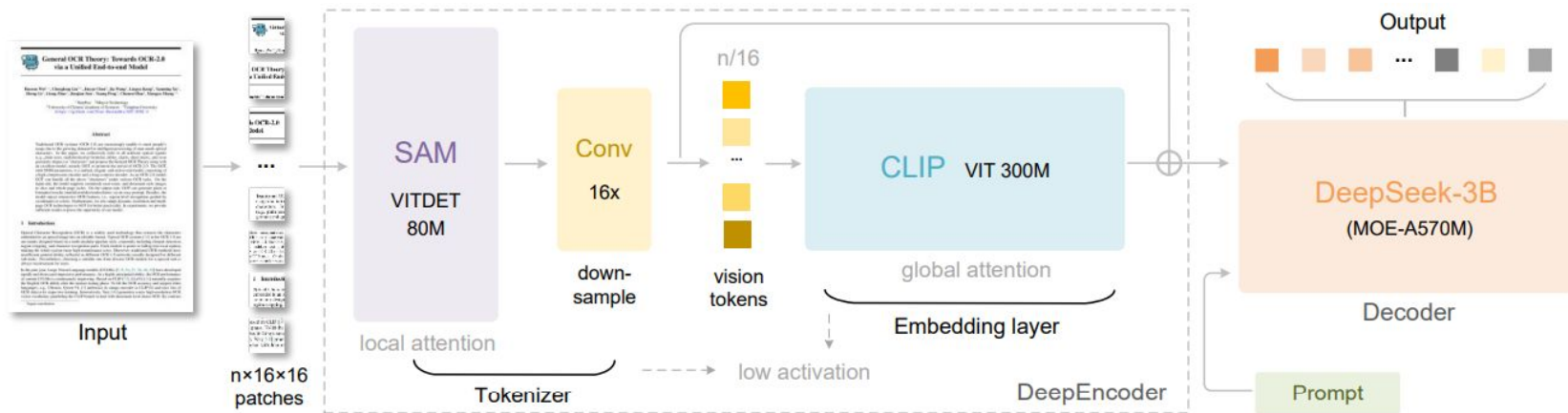
# DeepSeek-OCR: Introduction & Motivation

- Types of Vision Encoders in current VLMs:
  - Dual-tower architecture: two parallel encoders
  - Tile-based method: split image into small tiles
  - Adaptive-resolution encoder: process full image using patch segmentation
- One question not addressed: What is the least amount of vision tokens we need to decode a document of  $x$  words?



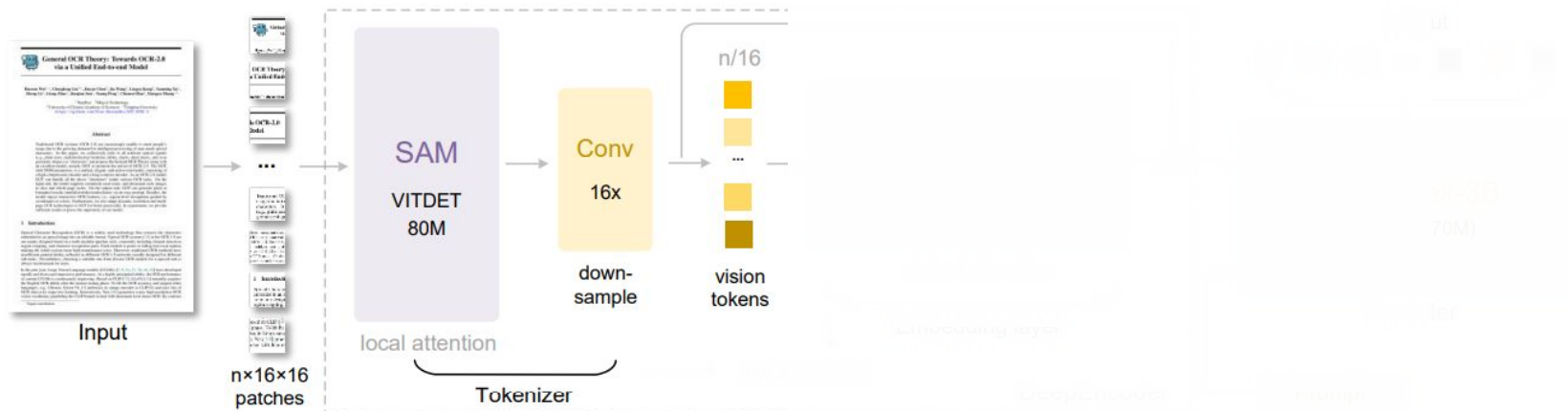
# DeepSeek-OCR: Methodology

- Encoder three part architecture:
  - SAM base
  - 2 layer convolutional model (downsampler)
  - CLIP
- DeepSeek-3B MoE Decoder



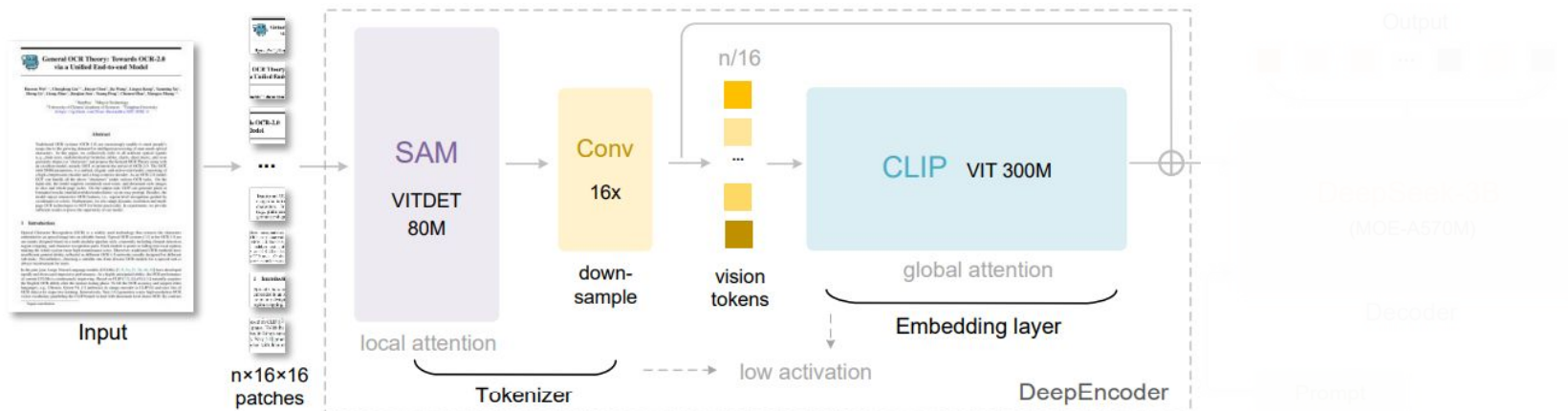
# DeepSeek-OCR: Methodology

- Segment Anything Model (SAM) base: vision model designed to separate/segment regions or objects in an image and extract visual perception features.
- Breaks image into  $16 \times 16$  patches and creates local feature tokens to then be sent into downsampling module → Downsampler reduces number of vision tokens



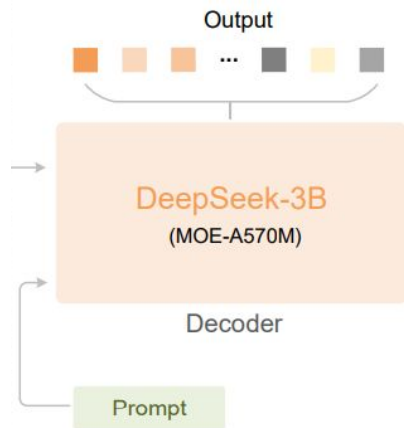
# DeepSeek-OCR: Methodology

- CLIP: Embedding layer with global attention
  - Vision Encoder that learns to connect images and text.
  - The downsampled tokens are fed into CLIP which applies dense global attention, how pieces fit together as “text”



# DeepSeek-OCR: Methodology

- During inference, model activates 8 out of 64 experts (specialized sub-networks)
- $Z$  are the compressed latent (vision) tokens from DeepEncoder
- $\hat{X}$  is the reconstructed text representation.
- $f_{\text{dec}}$  is a non linear mapping that can be learned.

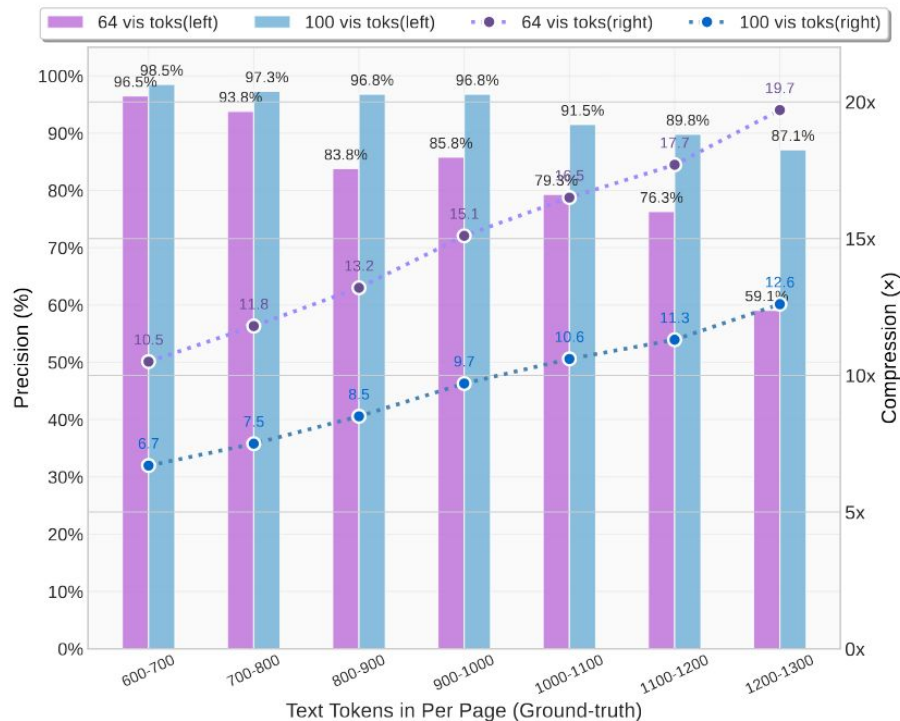


$$f_{\text{dec}} : \mathbb{R}^{n \times d_{\text{latent}}} \rightarrow \mathbb{R}^{N \times d_{\text{text}}}; \quad \hat{X} = f_{\text{dec}}(Z) \quad \text{where } n \leq N$$

# DeepSeek-OCR: Data

- OCR 1.0 data: ~30M PDF pages + natural scene text (English, Chinese, ~100 languages)
- OCR 2.0 data: synthetic charts, geometry, and chemical formulas for structured parsing
- General vision data: small portion of image-caption, grounding, detection
- Text-only data: 10% of training to maintain language modeling ability

# DeepSeek-OCR: Results

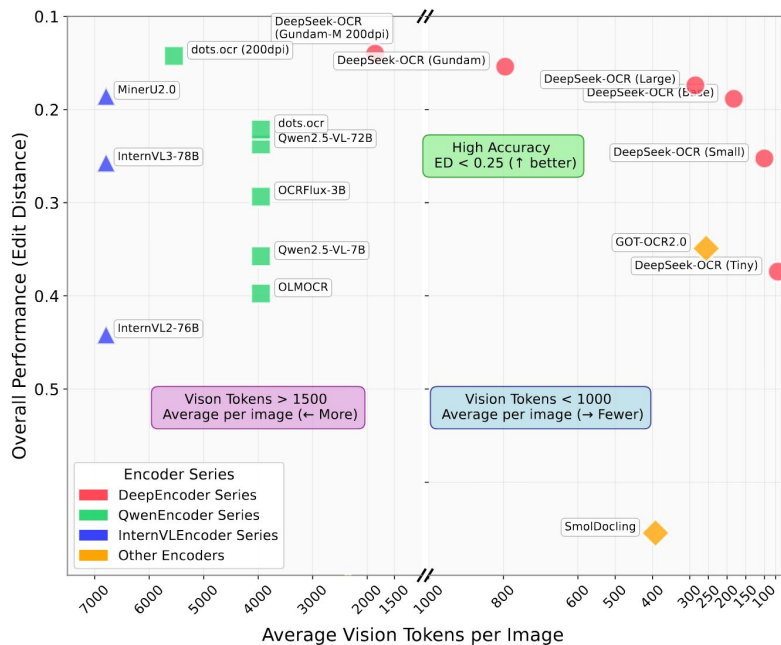


(a) Compression on Fox benchmark

- 97% OCR precision with text compressed by up to 10 times.
- 60% OCR precision even at a 20x text compression ratio
- Promising for handling long context

| Text Tokens | Vision Tokens =64 |             | Vision Tokens=100 |             | Pages |
|-------------|-------------------|-------------|-------------------|-------------|-------|
|             | Precision         | Compression | Precision         | Compression |       |
| 600-700     | 96.5%             | 10.5×       | 98.5%             | 6.7×        | 7     |
| 700-800     | 93.8%             | 11.8×       | 97.3%             | 7.5×        | 28    |
| 800-900     | 83.8%             | 13.2×       | 96.8%             | 8.5×        | 28    |
| 900-1000    | 85.9%             | 15.1×       | 96.8%             | 9.7×        | 14    |
| 1000-1100   | 79.3%             | 16.5×       | 91.5%             | 10.6×       | 11    |
| 1100-1200   | 76.4%             | 17.7×       | 89.8%             | 11.3×       | 8     |
| 1200-1300   | 59.1%             | 19.7×       | 87.1%             | 12.6×       | 4     |

# DeepSeek-OCR: Results



(b) Performance on Omnidocbench

- Achieves **better** performance with **less** token needed compared to other model.
- Shows **adaptive** performance for diverse content.

| Mode      | Type | Book  | Slides | Financial Report | Textbook | Exam Paper | Magazine | Academic Papers | Notes | Newspaper | Overall |
|-----------|------|-------|--------|------------------|----------|------------|----------|-----------------|-------|-----------|---------|
|           |      |       |        |                  |          |            |          |                 |       |           |         |
| Tiny      |      | 0.147 | 0.116  | 0.207            | 0.173    | 0.294      | 0.201    | 0.395           | 0.297 | 0.94      | 0.32    |
| Small     |      | 0.085 | 0.111  | 0.079            | 0.147    | 0.171      | 0.107    | 0.131           | 0.187 | 0.744     | 0.205   |
| Base      |      | 0.037 | 0.08   | 0.027            | 0.1      | 0.13       | 0.073    | 0.052           | 0.176 | 0.645     | 0.156   |
| Large     |      | 0.038 | 0.108  | 0.022            | 0.084    | 0.109      | 0.06     | 0.053           | 0.155 | 0.353     | 0.117   |
| Gundam    |      | 0.035 | 0.085  | 0.289            | 0.095    | 0.094      | 0.059    | 0.039           | 0.153 | 0.122     | 0.083   |
| Guandam-M |      | 0.052 | 0.09   | 0.034            | 0.091    | 0.079      | 0.079    | 0.048           | 0.1   | 0.099     | 0.077   |

Edit distances for different categories of documents in OmniDocBench

# DeepSeek-OCR: Results

- **Deep Parsing:** Extracts structured data from non-textual visual elements, such as converting text into Markdown, chemical formulas into SMILES strings, and geometric figures into structured data.
- **Multilingual Recognition:** Supports nearly 100 languages in PDF documents.
- **General Vision Understanding:** Describing image, object detection, grounding, etc.

<image>\n<[grounding]>Convert the document to markdown.

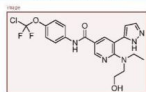
WO 2013/171642

PCT/IB2013/053771

[00369] The title compound was prepared in an analogous fashion to that described in Stage 22.1 using 5-bromo-6-chloro-N-(4-(chlorodifluoromethoxy)phenyl)nicotinamide (Stage 22.2) and 2-methylamino-ethanol to afford a white crystalline solid. HPLC (Condition 4)  $t_R = 5.72$  min, UPLC-MS (Condition 3)  $t_R = 1.14$  min,  $m/z = 452.2$   $[M+H]^+$ .

Example 24

N-(4-(Chlorodifluoromethoxy)phenyl)-6-(ethyl(2-hydroxyethyl)amino)-5-(1H-pyrazol-5-yl)nicotinamide

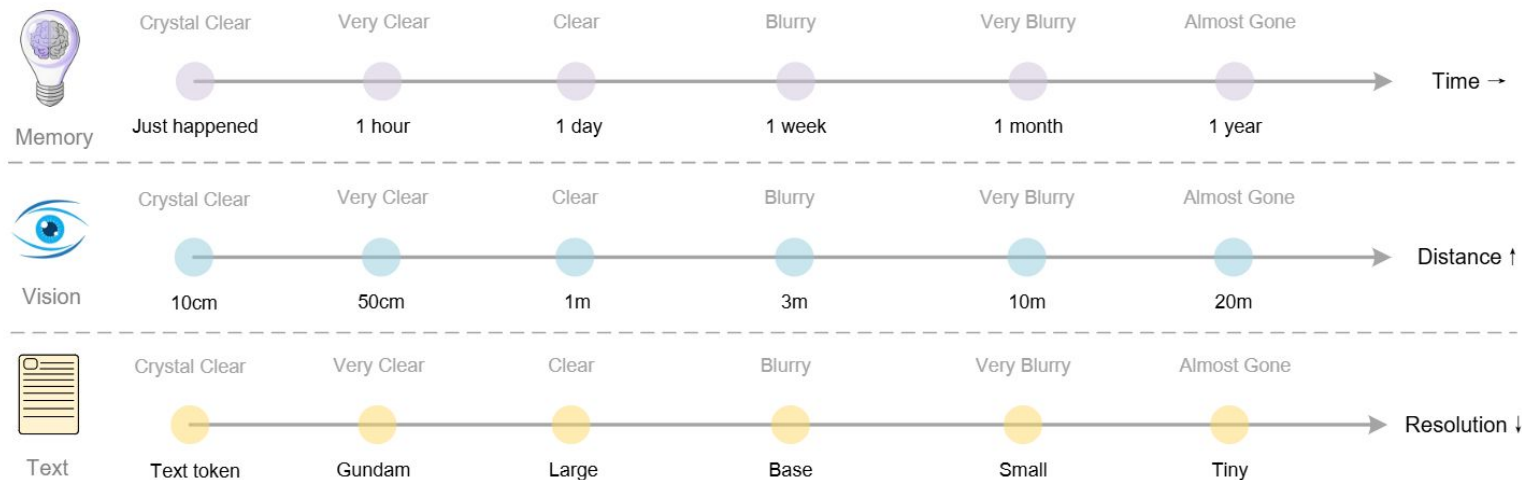


<image>\nIdentify all objects in the image and output them in bounding boxes.



# DeepSeek-OCR: Discussion

- Simulates human memory forgetting by blurring and reducing tokens for older visual contexts, allowing distant memories to naturally fade.
- Enables scalable, theoretically unlimited context architectures.



# Questions?



JOHNS HOPKINS

WHITING SCHOOL  
of ENGINEERING