



This MICCAI paper is the Open Access version, provided by the MICCAI Society. It is identical to the accepted version, except for the format and this watermark; the final published version is available on SpringerLink.

# CerviThink: A Reinforced Visual Reasoning Framework for Cervical Cancer Cell Classification

Manman Fei<sup>1</sup>, Haotian Jiang<sup>2</sup>, Zhenyu Yi<sup>1</sup>, Qian Wang<sup>2</sup>, and Lichi Zhang<sup>1</sup>✉

<sup>1</sup> School of Biomedical Engineering, Shanghai Jiao Tong University, Shanghai, China  
lichizhang@sjtu.edu.cn

<sup>2</sup> School of Biomedical Engineering, ShanghaiTech University, Shanghai, China

**Abstract.** Cervical cancer remains a major global health challenge, with early diagnosis relying on accurate classification of cervical cell images. However, current computer-aided diagnosis systems often struggle with subtle morphological variations, overlapping cells, and imaging artifacts, particularly in ambiguous samples, due to their focus on single-cell views without contextual background analysis. Moreover, the lack of transparent reasoning in these models limits their interpretability. To address these issues, we propose CerviThink, an innovative framework that integrates reinforcement learning (RL) to support structured, context-aware diagnostic reasoning. CerviThink employs a large vision-language model, fine-tuned on a Retrieval-Augmented Generation constructed CerviCoT dataset. We construct a decoupled grounding-answering pipeline: the grounding phase uses focus, transform, and ignore operations to isolate key cellular features and mask suspected abnormal cells, while the answering phase optimizes classifications with the Dynamic Visual Hybrid Reward (DVHR). DVHR integrates the consistency reinforcement reward and background normality confirmation reward, enhancing interpretability through iterative refinement and uncertainty resolution. Extensive experiments demonstrate CerviThink’s superior performance over existing methods, underscoring its potential to advance computational pathology through context-aware, interpretable visual reasoning. The code is publicly available at <https://github.com/feimanman/CerviThink>.

**Keywords:** Cervical Cancer · Reinforcement Learning · Visual Reasoning · Vision-Language Models · Computational Pathology.

## 1 Introduction

Cervical cancer remains a leading cause of mortality among women globally, with early detection through microscopic cell analysis being pivotal for effective treatment [21]. Cytological examination, or the Pap smear test, is a widely used method for cervical malignancy diagnosis [13], where specialists collect cells from the cervix of a patient and examine them under a microscope on glass slides to detect any malignancies [9]. Pap smear screening plays a crucial role in helping women prevent cervical cancer. However, such manual screening is laborious and subjective for cytologists, considering that there are thousands of cells on each slide. As a result, it is necessary to apply automatic computer-aided diagnosis (CAD) systems for cervical cancer.

This task is complicated by subtle morphological variations, overlapping cellular structures, and imaging artifacts, posing significant challenges to automated systems [10]. Pathologists address these complexities through a systematic, step-by-step process, meticulously observing nuclear size, chromatin texture, and cellular arrangement to derive informed diagnoses. Replicating this expert reasoning in artificial intelligence models is essential for advancing medical image analysis.

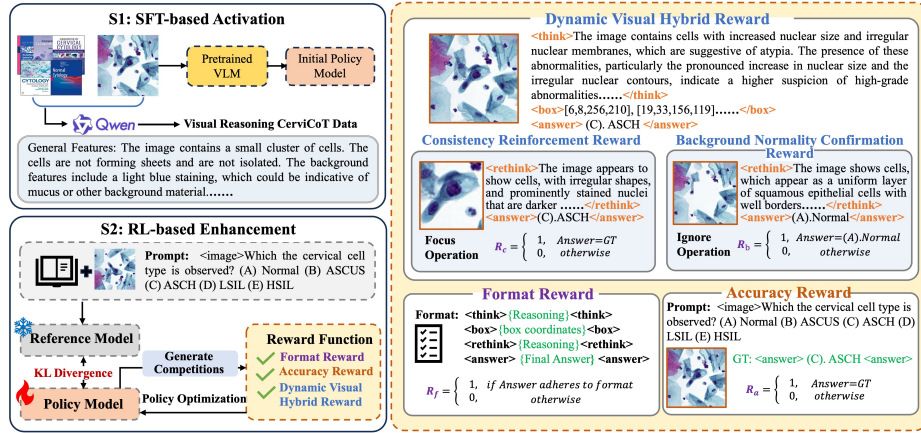
Traditional approaches, such as convolutional neural networks (CNNs) [5], have offered reasonable accuracy but lack interpretability and struggle with complex morphologies. In cervical cytology, prior methods have explored local consistency refinement, cytological knowledge, attribute descriptor matching, and weakly/semi-supervised lesion-cell detection to reduce annotation burden and improve cell-level recognition [8, 6, 7]. Recent advancements in large vision-language models (LVLMs), including Qwen2.5-VL [1] and GPT-4o [11], integrate visual and textual inputs for multimodal tasks. However, their application to medical imaging is hindered by unreliable reasoning, limited spatial perception, and reliance on large annotated datasets, which are costly and scarce in medical contexts.

To overcome these limitations, innovative visual reasoning frameworks have emerged. VILASR [22] enhances LVLMs’ spatial reasoning by interweaving textual reasoning with visual drawing operations. Liu et al. [17] introduce a closed-loop framework, enabling LVLMs to dynamically manipulate visual scenes through focus, transform, and ignore operations, addressing visual-to-textual conversion challenges in tasks like jigsaw solving. Ground-R1 [3] employs reinforcement learning (RL) for grounded visual reasoning without annotations, demonstrating uncertainty awareness. However, the aforementioned methods often struggle with cervical cell image classification due to challenges in accurately distinguishing nuclear sizes and chromatin textures in overlapping cells.

Motivated by these advances, we propose **CerviThink**, an RL-based framework that trains the Qwen2.5-VL model to emulate pathologists’ diagnostic reasoning for cervical cancer cell classification. CerviThink is fine-tuned on **CerviCoT**, a custom CoT dataset constructed via Retrieval-Augmented Generation (RAG) to embed medical expertise, comprising cervical cell images with detailed rationales. The framework employs a decoupled grounding-answering pipeline: the grounding phase applies focus, transform, and ignore operations to isolate key cellular features and mask suspected abnormal cells, while the answering phase optimizes classifications using the Dynamic Visual Hybrid Reward (DVHR) mechanism. This includes Consistency Reinforcement and Background Normality Confirmation Rewards, guiding the model to refine its reasoning through uncertainty resolution and iterative adjustments.

Our contributions are threefold:

- We introduce CerviThink, an RL-driven framework that mimics pathologists’ diagnostic process, avoiding ground-truth bounding-box supervision during RL training and inference.
- We design a decoupled grounding-answering pipeline with tailored ignore operations and a DVHR mechanism to enhance diagnostic specificity and accuracy.
- We validate CerviThink on the DST-derived CerviCoT benchmark and two external cervical cytology datasets, achieving significant accuracy gains and exhibiting



**Fig. 1.** Schematic of the CerviThink framework. The diagram is divided into two main stages: (S1) SFT-based activation from cervical cell image-question pairs to an initial policy model, and (S2) RL-based enhancement from policy-generated grounding and answer rollouts to the final cell-type prediction. The Reward Function includes Consistency Reinforcement Reward, Background Normality Confirmation Reward, Format Reward, and Accuracy Reward.

structured intermediate reasoning behaviors, advancing interpretable medical image analysis.

## 2 Methodology

CerviThink is a novel framework designed to emulate the diagnostic reasoning of pathologists for cervical cancer cell classification through reinforcement learning (RL). Inspired by Ground-R1’s RL-driven grounded reasoning [3], CerviThink leverages the Qwen2.5-VL model [1] as our backbone, fine-tuned on a custom Chain-of-Thought (CoT) dataset, **CerviCoT**, constructed via Retrieval-Augmented Generation (RAG). The framework employs a decoupled grounding-answering pipeline to guide the model step by step in observing cell morphology, inspired by pathologists’ systematic analysis of nuclear size, chromatin texture, and cellular arrangement. Dynamic visual operations (focus, transform, ignore) and the Dynamic Visual Hybrid Reward (DVHR) optimize the process, with the ignore operation tailored to mask suspected abnormal cells to verify if the remaining background is normal. Figure 1 illustrates the pipeline.

### 2.1 Dataset Construction with RAG

To embed pathologists’ diagnostic knowledge, we construct **CerviCoT**, a specialized CoT dataset for cervical cancer cell classification. According to the Bethesda system (TBS) [18], cervical squamous cells are classified into five classes: HSIL, ASC-H,

LSIL, ASC-US, and Normal. Using Retrieval-Augmented Generation (RAG), we retrieve medical literature from PubMed to generate detailed CoT rationales, outlining step-by-step reasoning processes such as identifying cell boundaries, assessing nuclear size and shape, evaluating chromatin patterns, and classifying abnormalities. Each sample in the dataset includes an image-question-answer triplet and a CoT rationale. This approach automates the integration of domain-specific knowledge, reducing manual annotation efforts while ensuring clinical relevance.

## 2.2 Grounding Phase

The grounding phase emulates a pathologist’s initial scanning process, identifying critical diagnostic evidence without requiring explicit bounding box supervision. Inspired by Ground-R1 [3], given an input image  $v \in \mathbb{R}^{H \times W \times 3}$  and a query  $q$ , the framework employs the reference policy  $\pi_{\theta_{old}}$  to generate a set of  $G_1 = 4$  grounding rollouts  $b = \{b_i\}_{i=1}^{G_1} \sim \pi_{\theta_{old}}(q, v)$ , where each  $b_i = [x_1, y_1, x_2, y_2]$  denotes the coordinates of an axis-aligned bounding box. Adhering to the GRPO paradigm [19], the model encapsulates its latent reasoning within `<think>` tags and formalizes spatial coordinates within `<box>` tags, ensuring high interpretability and structural consistency.

To replicate microscopic diagnostic techniques, we implement a suite of dynamic visual operations applied to the extracted evidence regions. Specifically, the **Focus Operation** acts as a precise localization mechanism by cropping the region  $b_i$  to isolate key morphological structures, such as hyperchromatic nuclei or irregular membranes, which effectively minimizes background interference for downstream reasoning.

This is complemented by the **Transform Operation**, which emulates microscopic magnification and enhancement through dynamic zooming with a scaling factor  $\sigma \in [1, 3]$  and contrast adjustments. Such augmentations are vital for resolving fine-grained pathological features like coarse chromatin patterns or nuclear pleomorphism, which are essential for accurate lesion grading. Furthermore, the **Ignore Operation** mimics the clinical “rule-out” strategy used in differential diagnosis; by leveraging OpenCV-based inpainting, the model selectively masks potentially abnormal cells and reconstructs the area using neighboring textures. This allows the framework to verify if the remaining substrate represents a truly normal background, thereby significantly enhancing diagnostic specificity and reducing false-positive interpretations.

## 2.3 Answering Phase

The answering phase simulates pathologists’ diagnostic decision-making by synthesizing classifications based on the original image  $v$ , question  $q$ , and manipulated evidence regions  $\{e_i\}$ . For each evidence region  $e_i$ , the policy  $\pi_{\theta_{old}}$  generates answer rollouts:

$$\alpha_i = \{\alpha_{i,j}\}_{j=1}^{G_2} \sim \pi_{\theta_{old}}(q, v, e_i), \quad \text{for } i = 1, \dots, G_1,$$

where  $G_2 = 2$  represents the number of answer rollouts, and responses are structured within `<answer>` and `</answer>` tags. This phase is driven by an overall training reward consisting of the Format Reward  $r_{i,j}^f$ , the Accuracy Reward  $r_{i,j}^a$ , and the Dynamic Visual Hybrid Reward (DVHR).

**DVHR Mechanism:** DVHR optimizes diagnostic reasoning through two visual diagnostic rewards: the Consistency Reinforcement Reward and the Background Normality Confirmation Reward. To compute these specific visual rewards, the framework performs auxiliary forward passes on the manipulated variations of  $e_i$ . This mechanism reflects pathologists’ confidence-building process:

- **Format Reward** ( $r_{i,j}^f$ ): Ensures outputs follow the required structure, calculated as:

$$r_{i,j}^f = \mathbb{I}(\text{output conforms to } \langle \text{think} \rangle, \langle \text{box} \rangle, \langle \text{rethink} \rangle, \text{ and } \langle \text{answer} \rangle \text{ tags}).$$

- **Consistency Reinforcement Reward** ( $r_{i,j}^c$ ): For each cropped region  $e_i$ , an auxiliary response  $\alpha_{i,j}^c$  is generated. The reward assesses consistency with the ground-truth label:

$$r_{i,j}^c = \mathbb{I}(\alpha_{i,j}^c = y_{\text{true}}),$$

rewarding consistent diagnostic observations akin to pathologists’ repeated feature checks.

- **Background Normality Confirmation Reward** ( $r_{i,j}^b$ ): Following the masking of suspected abnormal cells via inpainting, the response  $c_{i,j}^{\text{inp}}$  is generated via an independent forward pass (where  $c_{i,j}^{\text{inp}}$  denotes the classification of the inpainted evidence region  $e_i$ ). The reward is assigned based on normal classification:

$$r_{i,j}^b = \mathbb{I}(c_{i,j}^{\text{inp}} = \text{normal}),$$

incentivizing the model to verify whether the remaining context supports a normal-background interpretation. This reward is used as a soft regularizer rather than a strict assumption, since images with multiple abnormal regions may still contain diagnostically relevant background cues after inpainting.

- **Accuracy Reward** ( $r_{i,j}^a$ ): The final classification  $\alpha_{i,j}$  is rewarded for alignment with the ground-truth label:

$$r_{i,j}^a = \mathbb{I}(\alpha_{i,j} = y_{\text{true}}),$$

encouraging correct final diagnostic predictions.

The total training reward is computed as the sum:

$$r_{i,j}^t = r_{i,j}^f + r_{i,j}^c + r_{i,j}^b + r_{i,j}^a.$$

## 2.4 Reinforcement Learning Training

CerviThink is trained using the Group Relative Policy Optimization (GRPO) framework [19], which coordinates evidence grounding and answer generation through reinforcement learning to enhance reasoning for cervical cancer cell classification. The current policy is optimized by maximizing the following objective function:

$$\mathcal{J}_{GRPO}(\theta) = \mathbb{E}_{q,v,o \sim \pi_{\theta_{old}}} [\mathcal{L}(q, v, o)],$$

which is expanded as:

$$\mathcal{J}_{GRPO}(\theta) = \frac{1}{G_1 \cdot G_2} \sum_{i=1}^{G_1} \sum_{j=1}^{G_2} \min \left( \frac{\pi_{\theta}(o_{i,j} \mid q, v)}{\pi_{\theta_{old}}(o_{i,j} \mid q, v)} A_{i,j}, \right. \\ \left. \text{clip} \left( \frac{\pi_{\theta}(o_{i,j} \mid q, v)}{\pi_{\theta_{old}}(o_{i,j} \mid q, v)}, 1 - \epsilon, 1 + \epsilon \right) A_{i,j} \right),$$

where  $\pi_{\theta}$  denotes the current policy model,  $o_{i,j}$  represents the  $j$ -th rollout answer based on the  $i$ -th evidence region (derived from the grounding phase),  $\epsilon$  is a hyperparameter controlling the clipping range, and  $A_{i,j}$  are the advantages computed using the group reward  $\{r_{i,j}^t \mid 1 \leq i \leq G_1, 1 \leq j \leq G_2\}$  as follows:

$$A_{i,j} = \frac{r_{i,j}^t - \text{mean}(\{r_{i,j}^t\})}{\text{std}(\{r_{i,j}^t\}) + \delta}.$$

where  $\delta$  is a small constant for numerical stability.

Output regularization is provided by the format, consistency, background, and accuracy rewards. This process encourages structured reasoning behaviors, including uncertainty awareness by triggering transform operations on ambiguous regions, spatial perception through localization of nuclei, and iterative refinement by revisiting inaccurate regions.

### 3 Experiments

#### 3.1 Experimental Setup

**Dataset** To evaluate CerviThink’s performance in cervical cancer cell classification, we utilize three distinct datasets tailored to the task of microscopic image analysis, aligning with the framework’s focus on emulating pathologists’ diagnostic reasoning:

- **The DST dataset**, a private in-house collection sourced from collaborating hospitals, comprises 5,167 images of  $1024 \times 1024$  pixels, with 7,321 annotated cells. From this dataset, we construct **CerviCoT** by resizing the images to  $256 \times 256$  pixels using the annotated bounding boxes, integrating RAG to embed detailed Chain-of-Thought (CoT) rationales that mimic pathologists’ step-by-step analysis. A 4:1 train-test split is adopted at the patient level to prevent data leakage.
- **The ComparisonDetector dataset** [15] consists of 7,410 images cropped from whole slide images (WSIs), annotated with 48,587 bounding boxes representing diverse cervical lesion types. To standardize input for CerviThink, we resize these images to  $256 \times 256$  pixels based on pathologist-provided bounding boxes.
- **The publicly available HiCervix dataset** [2], a comprehensive cervical cytology dataset, encompasses 29 annotated categories. For our experiments, we focus on five clinically relevant categories, including ASC-US, ASC-H, LSIL, HSIL, and Normal, to align with CerviThink’s classification objectives.

**Table 1.** Performance Comparison of Different Models on the DST Dataset

| Method             | Precision    | Recall       | F1-score     |
|--------------------|--------------|--------------|--------------|
| LLaVA-1.5-7B [16]  | 41.89        | 42.35        | 42.23        |
| Qwen2.5-VL-3B [1]  | 43.21        | 44.56        | 43.88        |
| Gemini 2.5 [4]     | 47.98        | 52.10        | 49.98        |
| LLaVA-Med [14]     | 57.56        | 52.32        | 54.23        |
| Quilt-LLaVA [12]   | 48.43        | 51.42        | 50.09        |
| PathGen-LLaVA [20] | 54.11        | 59.78        | 55.63        |
| SFT                | 64.27        | 59.22        | 62.12        |
| RL                 | 67.38        | 64.13        | 65.78        |
| CerviThink         | <b>75.22</b> | <b>75.45</b> | <b>74.80</b> |

**Table 2.** Performance Comparison on ComparisonDetector and HiCervix Datasets

| Method            | ComparisonDetector Dataset |              |              | HiCervix Dataset |              |              |
|-------------------|----------------------------|--------------|--------------|------------------|--------------|--------------|
|                   | Precision                  | Recall       | F1-score     | Precision        | Recall       | F1-score     |
| <b>SFT</b>        | 48.92                      | 46.71        | 47.23        | 46.78            | 49.12        | 47.56        |
| <b>RL</b>         | 50.12                      | 47.45        | 48.98        | 51.23            | 49.45        | 50.91        |
| <b>CerviThink</b> | <b>61.77</b>               | <b>57.03</b> | <b>58.09</b> | <b>64.31</b>     | <b>57.14</b> | <b>60.67</b> |

The annotated bounding boxes are used only for preprocessing and standardized cell-crop construction. During RL training and inference, no ground-truth bounding boxes are provided as supervision; spatial coordinates are generated by the policy and optimized through downstream reasoning rewards. These datasets collectively support the development and validation of CerviThink’s decoupled grounding-answering pipeline across private and public cervical cytology images.

**Evaluation Metrics** To rigorously assess the performance of CerviThink in cervical cancer cell classification, we adopt three widely recognized evaluation metrics: Precision, Recall, and Weighted F1-score. These metrics are selected to comprehensively evaluate the model’s classification performance, aligning with the clinical need for accurate and balanced classification of cervical cells.

**Implementation Details** CerviThink is implemented using the Qwen2.5-VL-7B model [1], fine-tuned via supervised fine-tuning (SFT) with batch size 8 and learning rate  $1 \times 10^{-6}$  on CerviCoT, followed by RL with GRPO [19], inspired by Ground-R1 [3]. Trainable baselines are trained or fine-tuned with consistent hyperparameters, while closed-source LVLMS are evaluated under the same prompting and evaluation protocol. We use  $G_1 = 4$  grounding rollouts and  $G_2 = 2$  answering rollouts. The clipping range is set to  $\epsilon = 0.2$ , and all reward components are assigned equal weights of 1.0. For visual operations, the transform operation samples zoom factors from [1, 3] and

**Table 3.** Ablation Study on CerviCoT Construction

| Configuration    | Precision    | Recall       | F1-score     |
|------------------|--------------|--------------|--------------|
| Without CerviCoT | 67.90        | 70.47        | 69.14        |
| Partial CerviCoT | 70.56        | 72.32        | 70.98        |
| CerviThink       | <b>75.22</b> | <b>75.45</b> | <b>74.80</b> |

**Table 4.** Performance Comparison with Different Reward Functions

| Reward Functions |             |             |             | Performance Metrics |              |              |
|------------------|-------------|-------------|-------------|---------------------|--------------|--------------|
| $r_{i,j}^f$      | $r_{i,j}^a$ | $r_{i,j}^c$ | $r_{i,j}^b$ | Precision           | Recall       | F1-score     |
| ✓                | ✓           |             |             | 67.38               | 64.13        | 65.78        |
| ✓                | ✓           | ✓           |             | 70.48               | 68.21        | 69.35        |
| ✓                | ✓           |             | ✓           | 73.23               | 70.91        | 72.08        |
| ✓                | ✓           | ✓           | ✓           | <b>75.22</b>        | <b>75.45</b> | <b>74.80</b> |

contrast factors from  $[0.8, 1.5]$ , while the ignore operation uses OpenCV Navier-Stokes inpainting with radius 5.

### 3.2 Comparison with State-of-the-Arts

We evaluate **CerviThink** across the standard DST setting and few-shot settings on ComparisonDetector [15] and HiCervix [2]. For ComparisonDetector and HiCervix, the few-shot setting uses 10% of the training samples, while the test samples are kept disjoint. Baselines include general LVLMS (e.g., Qwen2.5-VL) and pathology-specific models (e.g., PathGen-LLaVA).

As shown in Tables 1 and 2, CerviThink consistently outperforms all baselines. In the DST benchmark, it achieves a 74.80% F1-score, surpassing PathGen-LLaVA by a large margin. This superiority stems from the **DVHR** mechanism and **ignore/transform** operations, which encourage more focused morphological feature analysis. In few-shot scenarios, CerviThink demonstrates robust adaptability, significantly exceeding the base RL performance. This validates that our complete RL-driven strategy effectively compensates for sparse annotations by refining feature resolution and validating background normality.

### 3.3 Ablation Studies

*Effectiveness of CerviCoT.* Table 3 validates the RAG-based expertise integration. Removing CoT rationales (Without CerviCoT) or using manual rationales (Partial CerviCoT) results in F1-score drops of 5.66% and 3.82%, respectively. Although manual rationales are accurate, they often lack the standardized, step-by-step deductive reasoning that RAG provides. This highlights RAG’s necessity in aligning model reasoning with clinical diagnostic logic.

*Effectiveness of Reward Function.* Table 4 decomposes the **DVHR** components. Integrating the Consistency reward ( $r_{i,j}^c$ ) improves the F1-score by 3.57% over the base configuration ( $r_{i,j}^f + r_{i,j}^a$ ), while the Background Normality reward ( $r_{i,j}^b$ ) independently provides a 6.30% boost. Combining all components in the full DVHR optimizes both diagnostic stability and differential reasoning, reaching the peak F1-score of 74.80%.

## 4 Conclusion

This study introduces CerviThink, a reinforcement learning-driven framework that is inspired by pathologists’ diagnostic reasoning for cervical cancer cell classification. Built on the Qwen2.5-VL model and fine-tuned on the CerviCoT dataset, which comprises cervical cell images paired with RAG-generated diagnostic rationales, CerviThink employs a decoupled grounding-answering pipeline with dynamic visual operations and the Dynamic Visual Hybrid Reward mechanism. The DVHR, integrating the consistency reinforcement and background normality confirmation rewards, enhances interpretability and accuracy by optimizing structured reasoning signals rather than relying on ground-truth box supervision during RL training. Future work will explore expanding the framework to other histopathological analyses.

**Acknowledgments.** This research is supported by the National Key R&D Program of China (Grant No. 2023YFB4706300) and the National Natural Science Foundation of China (Grant No. 62471288).

**Disclosure of Interests.** The authors have no competing interests to declare that are relevant to the content of this article.

## References

1. Bai, S., Chen, K., Liu, X., Wang, J., Ge, W., Song, S., Dang, K., Wang, P., Wang, S., Tang, J., et al.: Qwen2.5-vl technical report. arXiv preprint arXiv:2502.13923 (2025)
2. Cai, D., Chen, J., Zhao, J., Xue, Y., Yang, S., Yuan, W., Feng, M., Weng, H., Liu, S., Peng, Y., et al.: Hicervix: An extensive hierarchical dataset and benchmark for cervical cytology classification. IEEE transactions on medical imaging (2024)
3. Cao, M., Zhao, H., Zhang, C., Chang, X., Reid, I., Liang, X.: Ground-r1: Incentivizing grounded visual reasoning via reinforcement learning. arXiv preprint arXiv:2505.20272 (2025)
4. Comanici, G., Bieber, E., Schaekermann, M., Pasupat, I., Sachdeva, N., Dhillon, I., Blistein, M., Ram, O., Zhang, D., Rosen, E., et al.: Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities. arXiv preprint arXiv:2507.06261 (2025)
5. Fang, M., Fu, M., Liao, B., Lei, X., Wu, F.X.: Deep integrated fusion of local and global features for cervical cell classification. Computers in Biology and Medicine **171**, 108153 (2024)
6. Fei, M., Shen, Z., Liu, M., Song, Z., Sun, Y., Han, X., Liu, Z., Jiang, H., Bai, L., Wang, Q., et al.: Refining cervical cell classification with cytological knowledge and optimal attribute descriptor matching. In: International Conference on Medical Image Computing and Computer-Assisted Intervention. pp. 519–529. Springer (2025)

7. Fei, M., Song, Z., Shen, Z., Liu, M., Wang, Q., Zhang, L.: Weakly semi-supervised cervical lesion cell detection via twin-memory augmented multiple instance learning. In: International Conference on Medical Image Computing and Computer-Assisted Intervention. pp. 637–647. Springer (2025)
8. Fei, M., Zhang, X., Cao, M., Shen, Z., Zhao, X., Song, Z., Wang, Q., Zhang, L.: Robust cervical abnormal cell detection via distillation from local-scale consistency refinement. In: International Conference on Medical Image Computing and Computer Assisted Intervention. pp. 652–661. Springer Nature Switzerland, Cham (2023)
9. Gençtav, A., Aksoy, S., Önder, S.: Unsupervised segmentation and classification of cervical cell images. *Pattern recognition* **45**(12), 4151–4168 (2012)
10. Hemalatha, K., Vetriselvi, V., Aruna Gladys, A., et al.: Cervixfuzzyfusion for cervical cancer cell image classification. *Biomedical Signal Processing and Control* **85**, 104920 (2023)
11. Hurst, A., Lerer, A., Goucher, A.P., Perelman, A., Ramesh, A., Clark, A., Ostrow, A., Welihinda, A., Hayes, A., Radford, A., et al.: Gpt-4o system card. *arXiv preprint arXiv:2410.21276* (2024)
12. Ikezogwo, W., Seyfioglu, S., Ghezloo, F., Geva, D., Sheikh Mohammed, F., Anand, P.K., Krishna, R., Shapiro, L.: Quilt-1m: One million image-text pairs for histopathology. *Advances in neural information processing systems* **36**, 37995–38017 (2023)
13. Kessler, T.A.: Cervical cancer: prevention and early detection. In: *Seminars in oncology nursing*, vol. 33, pp. 172–183. Elsevier (2017)
14. Li, C., Wong, C., Zhang, S., Usuyama, N., Liu, H., Yang, J., Naumann, T., Poon, H., Gao, J.: Llava-med: Training a large language-and-vision assistant for biomedicine in one day. *Advances in Neural Information Processing Systems* **36**, 28541–28564 (2023)
15. Liang, Y., Tang, Z., Yan, M., Chen, J., Liu, Q., Xiang, Y.: Comparison detector for cervical cell/clumps detection in the limited data scenario. *Neurocomputing* **437**, 195–205 (2021)
16. Liu, H., Li, C., Li, Y., Lee, Y.J.: Improved baselines with visual instruction tuning. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. pp. 26296–26306 (2024)
17. Liu, J., Li, Y., Xiao, B., Jian, Y., Qin, Z., Shao, T., Ding, Y.X., Zhou, K.: Autonomous imagination: Closed-loop decomposition of visual-to-textual conversion in visual reasoning for multimodal large language models (2025), <https://arxiv.org/abs/2411.18142>
18. Nayar, R., Wilbur, D.C.: The bethesda system for reporting cervical cytology: a historical perspective. *Acta Cytologica* **61**(4-5), 359–372 (2017)
19. Shao, Z., Wang, P., Zhu, Q., Xu, R., Song, J., Bi, X., Zhang, H., Zhang, M., Li, Y., Wu, Y., et al.: Deepseekmath: Pushing the limits of mathematical reasoning in open language models. *arXiv preprint arXiv:2402.03300* (2024)
20. Sun, Y., Zhang, Y., Si, Y., Zhu, C., Shui, Z., Zhang, K., Li, J., Lyu, X., Lin, T., Yang, L.: Pathgen-1.6 m: 1.6 million pathology image-text pairs generation through multi-agent collaboration. *arXiv preprint arXiv:2407.00203* (2024)
21. Sung, H., Ferlay, J., Siegel, R.L., Laversanne, M., Soerjomataram, I., Jemal, A., Bray, F.: Global cancer statistics 2020: Globocan estimates of incidence and mortality worldwide for 36 cancers in 185 countries. *CA: A Cancer Journal for Clinicians* **71**(3), 209–249 (2021)
22. Wu, J., Guan, J., Feng, K., Liu, Q., Wu, S., Wang, L., Wu, W., Tan, T.: Reinforcing spatial reasoning in vision-language models with interwoven thinking and visual drawing. *arXiv preprint arXiv:2506.09965* (2025)