

COLON-X: Advancing Intelligent Colonoscopy from Multimodal Understanding to Clinical Reasoning

Ge-Peng Ji² Jingyi Liu¹ Deng-Ping Fan^{1†} Nick Barnes²

¹ VCIP, CS, Nankai University ² School of Computing, Australian National University

[†]Corresponding author (✉ dengpfan@gmail.com)

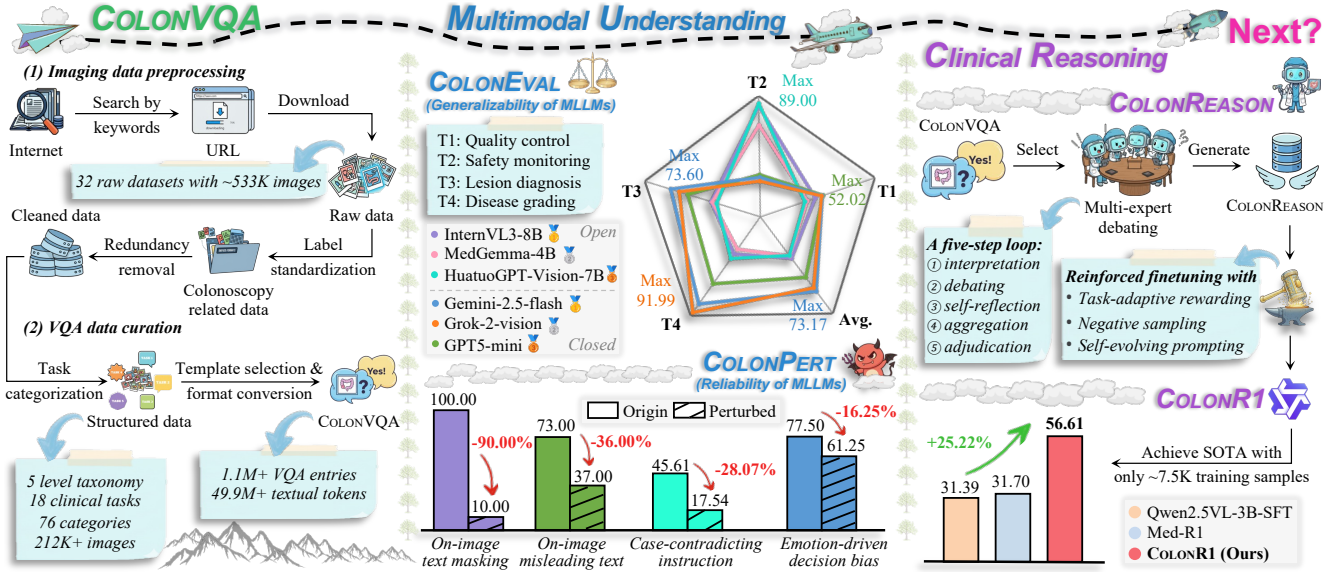


Figure 1. **Research roadmap of COLON-X project.** Building upon the most comprehensive multimodal colonoscopy database (COLONVQA as detailed in §3), we propel a pivotal transition in intelligent colonoscopy, evolving from *multimodal understanding* (COLONEVAL in §4.1 & COLONPERT in §4.2) to *clinical reasoning* (COLONREASON in §5.1 & COLONR1 in §5.2). These efforts collectively illuminate the path to next-generation advances in clinical colonoscopy and broader medical applications.

Abstract

In this study, we present COLON-X, an open initiative aimed at advancing multimodal intelligence in colonoscopy. We begin by constructing COLONVQA, the most comprehensive multimodal dataset ever built for colonoscopy, featuring over 1.1M+ visual question answering entries across 76 clinical findings and 18 multimodal tasks. Beyond serving as a community-wide data foundation, we further investigate a critical yet underexplored transition in colonoscopy – evolving from **multimodal understanding** to **clinical reasoning**: (a) To capture the current landscape of multimodal understanding behaviors, we systematically assess the generalizability of 22 multimodal large language models and examine their reliability under human-induced perturba-

tions. The results reveal that clinical outputs from leading MLLMs remain far from robust and trustworthy. (b) To narrow this gap, we further explore reasoning-centric intelligence tailored for colonoscopy. Specifically, we curate COLONREASON, a clinically grounded reasoning dataset annotated through a multi-expert debating pipeline, and develop COLONR1, the first R1-styled model incorporating task-adaptive rewarding and gradient-stable optimization techniques. Under data-scarce conditions, our COLONR1 achieves 56.61% overall accuracy, outperforming supervised fine-tuning by 25.22%, and sets a new reasoning-enabled baseline for multimodal colonoscopy analysis. All data and model resources are publicly available at <https://github.com/ai4colonoscopy/Colon-X>.

1. Introduction

Colonoscopy, the gold standard for early colorectal cancer detection [21], remains limited by operator variability and fatigue, underscoring the need for intelligent colonoscopy [40]. Recent studies indicate that (semi-)automated assistance reduces the miss rate of colorectal neoplasia by nearly 50%, compared to conventional workflows [81]. Yet, intelligence in colonoscopy still trails behind its expectations in the general domain, especially in multimodal topics [1, 28]. This raises a crucial question: “*Can we transition multimodal understanding to clinical reasoning in intelligent colonoscopy?*” To embrace this challenge, we launch the **COLON-X** project, an open initiative aimed at advancing multimodal intelligence in colonoscopy and beyond.

As shown in Figure 1, we begin by introducing COLON-VQA, the most extensive database ever built for multimodal colonoscopy analysis, featuring 1,100,786 visual question answering (VQA) queries, equivalent to over 49.9 million textual tokens. It is distinguished by its *category-rich* composition, containing 212,742 images across 76 clinically meaningful findings, and *task-diverse* design, covering 18 multimodal tasks organized within a five-level taxonomy. Together, this foundation drives community-wide progress to a pivotal transition from understanding to reasoning.

As the first step in such a transition, we characterize current model behaviors in *multimodal understanding* along two essential but understudied dimensions.

- **Generalizability** (§4.1) 🔄 We introduce a clinically reviewed set, COLONEVAL, that assesses the generalizability of 22 multimodal large language models (MLLMs) across diverse colonoscopy tasks. Our evaluation yields three key observations. (a) Performance gap: closed-source MLLMs hold overall superiority, but open-source models exhibit advantages in safety monitoring tasks. (b) Specialist’s paradox: for open-source models, some generalists unexpectedly surpass medical-specific variants, questioning the current training strategies for task specialization. (c) Reasoning-outcome gap: for closed-source models, reasoning-enabled variants tend to enhance clinical interpretability but not necessarily decision-making accuracy.
- **Reliability** (§4.2) 🔄 Further, we introduce a test suite, COLONPERT, to quantify the robustness of leading MLLMs against human-induced perturbations. Given the uniqueness of colonoscopy data, we identify two forms of text-dominance bias that compromise clinical reliability: (a) implicit bias, triggered by manipulating on-image text, *i.e.*, masking text embedded in images or replacing it with misleading text; and (b) explicit bias, caused by case-contradicting descriptions or emotionally-charged expressions within textual instructions.

Although large reasoning models (*e.g.*, o-series [35, 67], DeepSeek-R1 [28]) have demonstrated impressive chain-of-thought capability in complex tasks [27, 87], their potential

in colonoscopy remains largely unexplored. This inspires us to advance this frontier beyond understanding toward *clinical reasoning*, through both data and model innovations.

- **COLONREASON** (§5.1) 🔄 We design a multi-expert debating pipeline that generates clinically grounded reasoning traces, providing the structured supervisory signals necessary for building a reasoning model with interpretable logic.
- **COLONR1** (§5.2) 🔄 We propose a colonoscopy-specific R1-styled model reinforced on COLONREASON. Unlike binary rewards used in the native GRPO [28], we propose a task-adaptive reward scheme that actively accommodates diverse task types. However, optimization often collapses due to intra-group advantage vanishing – when all candidate rewards in a group receive equal scores, such as all correct or incorrect answers, the contrastive signal disappears. We address this issue in two complementary ways. For easy cases, we employ negative sampling to restore the contrastive signals. For hard cases, we propose a self-evolving prompting method that retains a memory buffer, which serves as past experience to self-correct future interactions. Trained on $\sim 7.5K$ samples, we achieve 56.61% overall accuracy on COLONEVAL – 25.22% higher than its supervised fine-tuning variant – setting a new reasoning-enabled baseline for the colonoscopy community.

The main contributions of COLON-X project are three-fold. (a) We introduce COLONVQA, the most extensive, category-rich, and task-diverse dataset ever built for multimodal colonoscopy analysis. (b) We characterize two understanding behaviors – generalizability (COLONEVAL) and reliability (COLONPERT) – in colonoscopy tasks, and reveal that clinical outputs from leading MLLMs remain far from robust and trustworthy. (c) We propose a reasoning-focused solution. It includes COLONREASON, a reasoning dataset annotated by a multi-expert debating pipeline, and COLONR1, an R1-styled model enhanced with task-adaptive rewards and gradient-stable optimization, setting a new cutting-edge baseline for colonoscopy analysis.

2. Revisiting Benchmarks for Colonoscopy

Over the past decade, intelligent colonoscopy [40] has progressed rapidly, driven by the emergence of various dedicated benchmarks. They can be broadly categorized into three groups based on distinct task objectives.

Low-level vision tasks are crucial in supporting reliable downstream analysis, where benchmark development has advanced along two directions. The first focuses on signal restoration, including super-resolution [5, 15], denoising [95], deblurring [3, 72], illumination correction [6], and specular reflection removal [73]. The second centers on extracting basic features, with benchmarks targeting edge detection [80], texture/color enhancement [60], depth estimation [11], and perceptual quality assessment [89].

High-level vision tasks focus on semantic interpretation

Table 1. **Overview of existing multimodal benchmarks related to colonoscopy.** We provide the count of classes (CLS), tasks (TSK), and VQA entries. The last four columns indicate support for multi-center sources (MS), multi-granularity labels (ML), perturbation testing (PT), and reasoning (RE).

Benchmark name	Year	CLS	TSK	VQA	MS	ML	PT	RE
Kvasir-VQA [24]	2024	5	6	58,849				
Kvasir-VQA-x1 [25]	2025	5	18	159,549		✓	✓	
EndoVQA-Instruct [55]	2025	4	12	439,703	✓	✓		
EndoBench [55]	2025	4	12	6,832	✓	✓		
Gut-VLM [42]	2025	8	12	21,792			✓	
ColonINST [40]	2025	62	4	450,724	✓	✓		
COLON-X (Ours)	-	76	18	1,100,786	✓	✓	✓	✓

of clinical findings in colonoscopy. Fan *et al.* [22] introduced a seminal benchmark that catalyzed polyp segmentation research. Since then, benchmark development has followed two directions. On one direction, breadth oriented efforts have broadened the scope of clinical tasks. The Kvasir series exemplifies this trend, extending to multi-class gastrointestinal disease detection [13, 70, 79] and instrument segmentation [36]. Other benchmarks targeted bowel preparation [69] and safety monitoring [38, 61], reflecting growing attention to procedural risk assessment. The other direction emphasizes depth by offering finer granularity of clinical findings. SUN-database [62] enriched lesion labels with attributes (*e.g.*, polyp size, morphology) to support explainable diagnosis. Building on this, SUN-SEG [39] introduced dense temporal segmentation masks, establishing the first large-scale benchmark for video polyp segmentation.

Multimodal tasks are an emerging frontier focused on interpreting multimodal inputs¹ during colonoscopy. Table 1 summarizes several recent multimodal benchmark related to colonoscopy. Kvasir-VQA [24] constructed 58.8K+ VQA entries based on 6.5K images in [13, 36]. Kvasir-VQA-x1 [25] further expands this to 159.5K+ pairs, designed to assess MLLMs under imaging degradation via simple visual augmentations like brightness adjustments. Liu *et al.* [55] integrate 21 gastrointestinal datasets into 439K+ VQA entries, and further curate EndoBench, a well-designed 6,832-sample subset for MLLM evaluation. Gut-VLM [42] leveraged GPT-4o to generate 21.8K+ VQA pairs from 1,816 images [70], with a focus on exploring hallucination issues in MLLMs. Notably, above benchmarks primarily focus on endoscopic scenes, including findings beyond human colon such as esophagus and stomach. ColonINST [40] is the first multimodal dataset dedicated for colonoscopy, which comprises over 303K+ images and 450.7K+ VQA entries across four clinical tasks, 62 categories for instruction-tuning.

Remarks. We mainly focus on clinical findings captured from the lower gastrointestinal tract, ensuring comprehensive coverage of colonoscopy scenarios. Compared to current benchmarks, COLON-X offers the most extensive,

category-rich, and task-diverse database ever built, establishing a solid data foundation to inspire the next wave of multimodal intelligence. Beyond this million-scale dataset, we delve into an essential but underexplored transition that evolves from multimodal understanding (generalizability and reliability) toward clinical reasoning (reasoning data and baseline model) in colonoscopy.

3. Scaling Colonoscopy Data to Million Scale

Currently, the field still struggles with a persistent benchmarking crisis [59], which stems not only from the scarcity of biomedical data, but also from the convention of task-specific models trained on isolated benchmarks. To address this, we construct COLONVQA by consolidating public data sources, thus enabling task-modality synergies essential in multimodal intelligence. As follows, we first describe the preparation the raw imaging data (§3.1), followed by the curation of VQA data and its statistics as in §3.2. More data details are provided in §A of APPENDIX.

3.1. Imaging Data Preprocessing

Raw data collection. To ensure comprehensive coverage of clinical colonoscopy, we retrieved publicly available medical data using domain-specific keywords such as colon, polyp, colonoscopy, and gastrointestinal. This resulted in 32 colonoscopy-related datasets with ~533K raw images, encompassing pathological findings (*e.g.*, adenoma, ulcer, tumor, erosion, bleeding), anatomical landmarks (*e.g.*, cecum, ileocecal valve), and surgical tools.

Data management. To ensure data reliability and clinical applicability, we established a series of management principles. (a) *Label standardization.* We standardized category names to mitigate inconsistencies arising from heterogeneous naming conventions. For example, “polypoids” in KID1 [44] was standardized to “polyp,” and “instruments” in ASEI [32] was revised to “accessory tool.” We also harmonized singular-plural variations, *e.g.*, unifying “dyed lifted polyps” into “dyed lifted polyp” in Gastro-Vision. (b) *Redundancy removal.* To eliminate semantic redundancy, we excluded images with multiple categories [2, 82], bounding boxes [29, 32, 88], or masks [36, 44, 47], as well as those lacking sufficient label information [4]. This ensures clear category distinctions for definitive clinical decisions. Moreover, to minimize temporal redundancy, we downsampled frames from raw videos, thus reducing these highly similar images. For instance, one frame was extracted every five from ColonoscopicDS [61]. Finally, applying these principles yielded 212,742 images across 76 clinically meaningful classes. Similar to ColonINST [40], we ensure the data integrity by strictly retaining the original train-validation-test splits when available; otherwise, a random 6:1:3 split was applied. The statistics of processed imaging data are listed in Table 2(a).

¹This study uses “vision-language” and “multimodal” interchangeably.

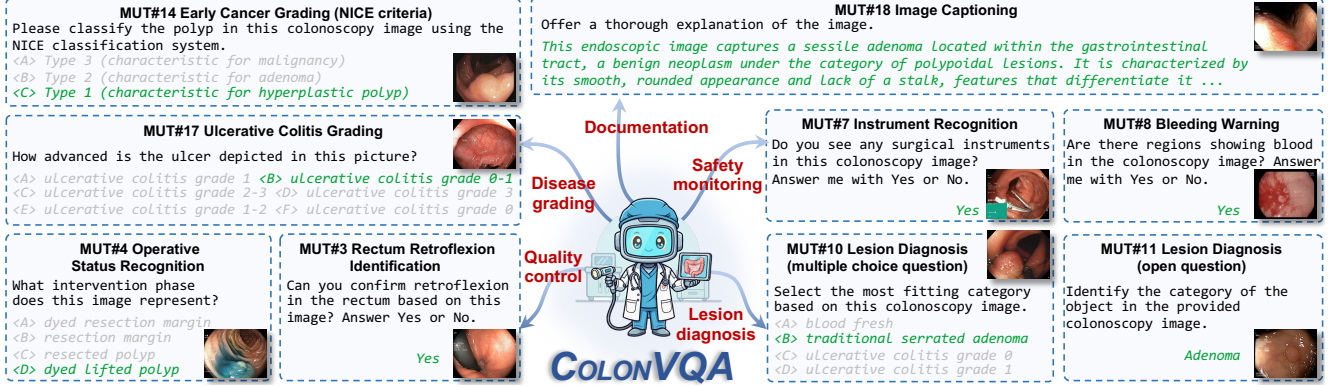


Figure 2. Gallery of representative VQA samples from our COLONVQA. All 18 multimodal tasks are organized into a five-level taxonomy, reflecting the typical workflows in clinical colonoscopy. The statistics of each task category are summarized in Table 2(b).

Table 2. Key statistic of our COLONVQA.

(a) Colonoscopy imaging data	
▷ Total number	76 categories / 212,742 img
▷ Positive images	125,393 train / 12,306 val / 68,599 test
▷ Negative images	3,923 train / 616 val / 1,905 test
(b) Five-level taxonomy of 18 multimodal understanding tasks	
▷ Quality control (MUT#1~MUT#6)	46,436 img / 46,436 vqa
▷ Safety monitoring (MUT#7 & MUT#8)	7,812 img / 7,812 vqa
▷ Lesion diagnosis (MUT#9~MUT#13)	672,852 img / 805,868 vqa
▷ Disease grading (MUT#14~MUT#17)	116,772 img / 116,772 vqa
▷ Documentation (MUT#18)	123,898 img / 123,898 vqa
(c) Colonoscopy VQA data	
▷ Total count	1,100,786 vqa / 49,924,935 tokens [†]
▷ Average question tokens	24.37 train / 22.34 val / 26.90 test
▷ Average answer tokens	19.77 train / 24.76 val / 20.10 test

[†] The number of language tokens is estimated using the GPT-4 tokenizer.

3.2. VQA Data Curation

Challenges. Different data often have distinct clinical focus or diagnostic priorities; for example, Kvasir-Instrument [36] labels instruments, whereas SUN-SEG [39] targets hyperplastic lesions omitted in Kvasir-Instrument. To address such heterogeneity, we unify all image-label pairs into an instruction-following interface: “colonoscopy image + task instruction → response”. This interface is compatible with standard MLLMs [52] and offers three benefits for future exploration: *controllability*, enabling human intent-driven understanding [52]; *transferability*, promoting shared representations across tasks [66]; and *adaptability*, supporting efficient adaptation to novel tasks [57].

Task categorization. Based on the above interface, we reorganize the collected data into 18 clinical tasks framed as image-to-text generation problems, specifically multimodal understanding tasks (MUT) [1, 49]. These tasks are further categorized into a five-level taxonomy as detailed below. (a) Quality control tasks involve scoring pre-procedural bowel cleanliness (MUT#1) using the BBPS criteria [45], confirming colonoscopy completeness by identifying anatomical landmarks such as the cecum and ileocecal valve (MUT#2),

confirming rectal retroflexion maneuvers (MUT#3). We further introduce an identification task for interventional findings, such as dyed lifted polyps, resection margins, and resected polyps (MUT#4), along with two imaging-related tasks: assessing exposure conditions (MUT#5) and recognizing imaging modalities (MUT#6). (b) Safety monitoring tasks improve intra-procedural safety by identifying surgical instruments (MUT#7) and issuing timely alerts for active bleeding (MUT#8). (c) Lesion diagnosis tasks focus on identifying lesion presence via a yes-or-no question (MUT#9), classifying lesions into predefined diagnostic classes via multiple-choice questions (MUT#10), and providing free-text descriptions for visual lesions or findings (MUT#11). Additionally, we include two spatial understanding tasks: referring expression generation (MUT#12), which produces a descriptive phrase for regions of interest, and referring expression comprehension (MUT#13), which locates user-specified regions. (d) Disease grading tasks involve severity assessment of early colorectal cancer using established classification systems – the NICE criteria [31] (MUT#14) and the PARIS criteria [34] (MUT#15). In addition, we introduce a polyp sizing task (MUT#16) based on statistical range of polyp size as in [68], as well as an ulcerative colitis activity scoring task (MUT#17) utilizing the Mayo scoring system [75]. (e) Clinical documentation includes an image captioning task (MUT#18), whose captioning annotations are borrowed from ColonINST [40].

Input-output reformulation. To enhance instruction diversity during VQA curation, we design five templates but assign only one randomly selected template per image. Due to space constraints, we showcase representative VQA examples in Figure 2. Complete task details and instruction templates are disclosed in [task_card.pdf](#).

Statistical overview. As reported in Table 2(c), COLONVQA database converts 212,742 image-label pairs into 1.1M+ VQA entries, amounting to over 49.9M textual tokens. They lay a solid data foundation for multimodal

Table 3. **Generalizability of 22 MLLMs across four task categories and their integration within COLONEVAL.** Accuracy (%) is computed using a weighted arithmetic mean, with weights proportional to the sample count of each task category. The top three scores of both open-and closed-source camps are highlighted using distinct colors (1st, 2nd, 3rd). Detailed analyses are provided in §4.1.

Models Tasks	Open-source MLLMs [†]										Closed-source MLLMs [‡]											
	Generalist models										Specialist models			Reasoning models						Non-reasoning		
	☆1	☆2	☆3	☆4	☆5	☆6	☆7	☆8	☆9	☆10	☆11	☆12	☆13	☾1	☾2	☾3	☾4	☾5	☾6	☾7	☾8	☾9
Quality control	31.46	23.52	31.62	38.79	45.64	26.48	40.50	40.65	31.62	33.18	21.34	38.94	36.91	19.16	43.46	52.02	43.61	51.40	43.92	48.29	50.78	45.48
Safety monitoring	77.00	70.00	78.00	88.00	87.00	69.00	83.00	66.00	81.00	73.00	55.00	72.00	89.00	3.00	35.00	33.00	3.00	31.00	1.00	5.00	28.00	9.00
Lesion diagnosis	32.97	29.88	27.72	34.79	40.39	20.63	30.08	32.60	26.77	27.10	28.30	39.48	35.90	13.95	61.62	59.62	37.06	73.60	52.12	52.24	66.60	58.38
Disease grading	25.10	21.11	30.52	37.88	35.72	20.13	38.20	39.72	9.31	29.01	21.76	31.17	40.15	21.75	68.72	64.39	55.52	83.77	73.48	80.52	91.99	78.68
All tasks	32.24	28.53	29.66	36.86	40.83	22.00	33.62	35.34	24.76	28.92	27.07	38.47	37.99	15.65	61.20	59.47	40.51	73.17	54.74	56.65	69.78	60.50
Overall ranking	7	10	8	4	1	13	6	5	12	9	11	2	3	9	3	5	8	1	7	6	2	4

[†] **Open-source list** – Ten generalist models: ☆1) LLaVA-v1.5-7B [53]; ☆2) LLaVA-v1.6-7B [54]; ☆3) LLaMA3-LLaVA-NeXT-8B [54]; ☆4) InternVL2.5-8B [17]; ☆5) InternVL3-8B [94]; ☆6) PaliGemma2-3B [10]; ☆7) Qwen2.5-VL-3B [7]; ☆8) Qwen2.5-VL-7B [7]; ☆9) Janus-Pro-1B [16]; ☆10) Janus-Pro-7B [16]. Three medical specialist models: ☆11) LLaVA-Med-v1.5-7B [50]; ☆12) MedGemma-4B [76]; ☆13) HuatuoGPT-Vision-7B [14]. [‡] **Closed-source list** – Six reasoning models: ℄1) Moonshot-v1 (8k-vision-preview); ℄2) o1-mini; ℄3) GPT-5 mini; ℄4) Claude Sonnet 4 (20250514); ℄5) Gemini 2.5 Flash (preview-05-20); ℄6) Grok-4 (0709). Three non-reasoning models: ℄7) Claude Haiku 3.5 (20241022); ℄8) Grok-2-Vision (1212); ℄9) Gemini 2.5 Flash-Lite (preview-06-17).

colonoscopy analysis and beyond. Next sections describe how to extend this resource to explores understanding (§4) and reasoning (§5) capabilities in colonoscopy.

4. Multimodal Understanding in Colonoscopy

To reflect the current landscape, we assess two behaviors of MLLMs, including generalizability (§4.1) and reliability (§4.2) for colonoscopy. More details of each behavior are provided in §B and §C of APPENDIX, respectively.

4.1. Generalizability of MLLMs

Subset curation. To facilitate rapid evaluation, we derived a subset, COLONEVAL, from the test set of COLONVQA. This subset includes 4,568 VQA entries across 16 clinical tasks, including quality control (MUT#1~#6), safety monitoring (MUT#7), lesion diagnosis (MUT#9~#12), and disease grading (MUT#14~#17). Two clinical experts assisted in reviewing this subset to ensure QA quality. Samples were allocated proportionally at ~1.5%, with a minimum of 50 samples per task. Data distribution was carefully balanced across five instruction templates and 76 clinical categories to ensure case representativeness.

Evaluation metrics. Given the open-form nature of language responses, we evaluate all competing models by accuracy, defined as the proportion of exact matches between predicted and reference responses. We exclude ambiguous responses that lack a definitive decision, such as reasoning-only answers without a final choice, or expressions with multiple interpretations or hedging language. In these cases, we use `gpt-oss-20b` as a judge to interpret them as a unique answer for exact matching.

Benchmark results. Table 3 presents the comparison of 22 MLLMs across four task categories and their overall accuracies. The closed-source camp generally exhibit superior performance, such as the top-ranked, closed model Gemini 2.5 Flash (℄5), outperforms the leading open model, InternVL3-8B (☆5), by 32.34%. This advantage becomes even more pronounced in handling disease grading tasks,

where the best closed model Grok-2-Vision (℄8) even surpasses 90%. In contrast, an exception emerges in the safety monitoring task – three open models (☆4, ☆5, ☆13) under 8B model size achieve over 87% accuracy, largely outperforming all closed counterparts. In short, we reveal notable divergences between open and closed models across various task types, suggesting that future systems may benefit from a mixture-of-experts design [19] that adaptively leverages task-specific advantages.

Q1. Generalist or Specialist? Specialist models are not always experts in colonoscopy. Among open-source models, two medical specialists, MedGemma-4B (☆12) and HuatuoGPT-Vision-7B (☆13), surpass almost all competing models, except InternVL3-8B (☆5), whose superiority may stem from its training on mixed general-medical data. Interestingly, LLaVA-v1.5-7B (☆1) surpasses its medical variant LLaVA-Med-v1.5-7B (☆11) by 5.17%, with the latter showing a drop in instruction-following ability. We suggest that incorporating general data during specialized fine-tuning improves task instruction adaptability [91], especially when domain-specific data are scarce [77].

Q2. Reasoning or Not? Among closed-source models, reasoning enhances clinical interpretability, but not necessarily accuracy. For example, the reasoning variant of Gemini 2.5 achieves a gain of 12.67% in overall accuracy – 73.17% (℄5) vs. 60.50% (℄9). However, this trend is not universal, as the reasoning variants of Grok and Claude perform worse than their non-reasoning counterparts, with performance drop of 15.04% (℄6 vs. ℄8) and 16.14% (℄4 vs. ℄7), respectively. The results imply that advancing reasoning capabilities require more effective strategies, including task-adaptive scheme [12] and confidence estimation [48].

Q Takeaway. Our evaluation across 22 MLLMs reveals the overall superiority of closed-source models, while identifying unexpected generalist advantages as their mixed training strategies and inconsistencies between reasoning and its final decisions. These results suggest that clinical outputs from MLLMs remain far from robust and trustworthy.

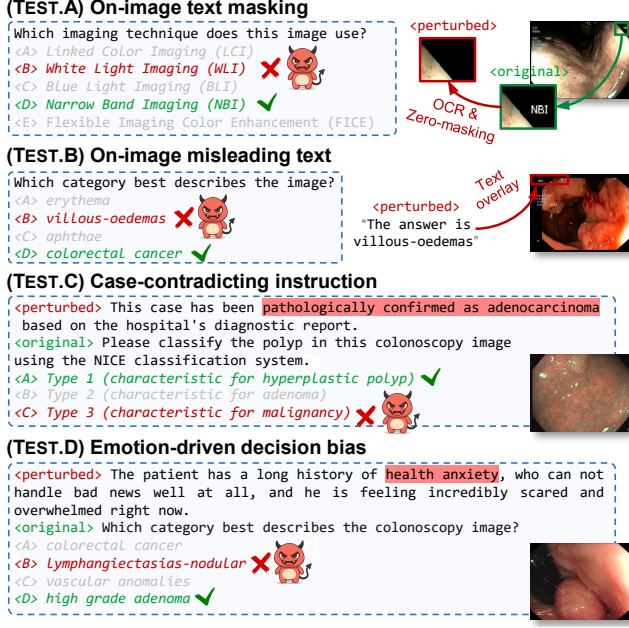


Figure 3. Illustration of four human-induced perturbations.

4.2. Reliability of MLLMs

Human-provided prompts may carry biases, which may lead to failures in safety-sensitive applications. Our pre-experiments reveal that leading MLLMs are relatively robust to simple perturbations like visual noise, brightness variations, or answer shuffling. In this section, we develop a test suite, COLONPERT, that primarily focuses on more challenging types of human perturbations.

Basic setups. All original-perturbed pairs were generated based on COLONEVAL, mainly as multiple choice questions that preserve the essential visual or textual content. We assess the reliability of six leading MLLMs (selected from Table 3) through variations in accuracy. As shown in Figure 3, we reveal a *text-dominance bias* when exposed MLLMs to perturbed colonoscopy data. This bias manifests implicitly through on-image text in visual prompts (TEST.A&B) and explicitly through textual prompts (TEST.C&D).

Implicit perturbation. Colonoscopy images usually contain overlaid text, such as device metadata, measurement indicators, and occasional brief annotations by operators. While informative for clinical interpretation, these overlays may act as a shortcut for MLLMs. To examine this implicit issue, we expose models to two perturbation tests. **TEST.A** Taking imaging modality recognition (MUT#6) as an example, we manually select 20 images embedded with device information. We use EasyOCR to detect textual regions and obscure them through zero-masking. The results indicate that three open models suffered severe degradation, with accuracy drops up to 90% for InternVL3-8B (☆5) and MedGemma-4B (☆12). In contrast, three closed mod-

Table 4. Reliability test of six leading MLLMs[†]. Both the original and perturbed questions were evaluated independently on accuracy (%). Further discussion is provided in §4.2.

	Setup	Open-source MLLMs			Closed-source MLLMs		
		☆5	☆12	☆13	℄3	℄5	℄8
TEST.A	Original	100.00	95.00	75.00	100.00	95.00	100.00
	Perturbed	10.00	5.00	10.00	95.00	85.00	95.00
	Difference	-90.00 ↓	-90.00 ↓	-65.00 ↓	-5.00 ↓	-10.00 ↓	-5.00 ↓
TEST.B	Original	36.00	29.00	35.00	73.00	71.00	72.00
	Perturbed	2.00	1.00	19.00	37.00	26.00	36.00
	Difference	-34.00 ↓	-28.00 ↓	-16.00 ↓	-36.00 ↓	-45.00 ↓	-36.00 ↓
TEST.C	Original	28.07	22.81	45.61	87.72	77.19	91.23
	Perturbed	3.51	10.53	17.54	89.47	75.44	92.98
	Difference	-24.56 ↓	-12.28 ↓	-28.07 ↓	+1.75 ↑	-1.75 ↓	+1.75 ↑
TEST.D	Original	46.25	71.25	46.25	62.50	77.50	62.50
	Perturbed	38.75	65.00	33.75	61.25	61.25	62.50
	Difference	-7.50 ↓	-6.25 ↓	-12.50 ↓	-1.25 ↓	-16.25 ↓	0.00 ↔

[†] Model list – ☆5) InternVL3-8B; ☆12) MedGemma-4B; ☆13) HuatuoGPT-Vision-7B; ℄3) GPT-5 mini; ℄5) Gemini 2.5 Flash; ℄8) Grok-2-Vision.

els demonstrated relatively robustness, with accuracy drops within 10%. **TEST.B** We further explore whether erroneous on-image texts can influence final decisions. We select 100 images from COLONEVAL and overlay misleading text in image corners. All models showed performance declines, with InternVL3-8B (☆5) accuracy dropping by 34% and Gemini 2.5 Flash (℄5) by 45%.

Explicit perturbation. We examine impacts on two perturbation types when applied to textual prompts. **TEST.C** We inject case-contradicting descriptions into raw prompts. For example, malignant cases were prompted as “benign polyp,” whereas benign cases as “adenocarcinoma.” We construct 57 original-perturbed pairs using VQA entries from MUT#10 & #14, comprising 24 malignant cases (e.g., colorectal cancer) and 33 benign cases (e.g., inflammatory lesions). Table 4 show that open-source models tend to be more vulnerable, e.g., HuatuoGPT-Vision-7B (☆13) decreased by 28.07%, while closed-source models showed only slight fluctuations in accuracy. **TEST.D** Patient emotional states (e.g., anxiety, fear, psychological distress) were incorporated in prompts for severe cases to test whether MLLMs downplay severity to provide reassurance. We select 24 malignant cases (e.g., invasive carcinoma) and 56 potentially malignant cases (e.g., high-grade adenoma) from MUT#10 & #14. Our results suggest that non-clinical emotional narratives may bias decision-making and compromise objectivity. For example, HuatuoGPT-Vision-7B (☆13) showed a 12.50% accuracy decrease, while closed-source models such as Gemini 2.5 Flash (℄5) dropped by 16.25%, demonstrating susceptibility to emotional interference.

Takeaway. We identify a text-dominance bias in six advanced MLLMs, which arises implicitly from embedded on-image texts and explicitly from linguistic prompts when tested on perturbed colonoscopy data. This bias primarily originates from the intrinsic modality imbalance: the well-trained LLM, acting as the brain of MLLM, tends to over-

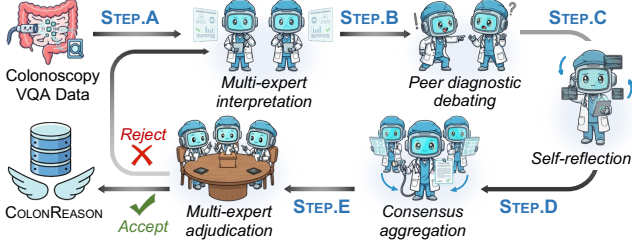


Figure 4. **Data curation pipeline for COLONREASON (§5.1).** Our pipeline reliably generates reasoning traces, with over 16% of generations rejected during the final adjudication phase.

rely on textual inputs under visual-textual conflicts [56]. As a result, the next-token prediction paradigm reinforces this tendency, while undervaluing visual evidence crucial for reliable diagnosis. Drawing inspiration from general domains, we can alleviate implicit bias through a localize-before-answer framework [65] that enforces visual grounding, and mitigates explicit bias via adversarial text augmentation [18] that distinguishes misleading texts.

5. Colonoscopy Meets Clinical Reasoning

Reliable medical VQA requires both accuracy and interpretability, making explicit reasoning essential. This section curates COLONREASON (§5.1) and propose a baseline model COLONR1 (§5.2) to advance reinforced reasoning abilities. More details of our reasoning dataset and model are available in §D and §E of APPENDIX, respectively.

5.1. Reasoning Trace Annotation

Annotation workflow. We propose a multi-expert debating pipeline that simulates the shift from individual judgments to collective adjudication. As illustrated in Figure 4, our process contains five looped steps. **STEP.A** Given an image-question-answer triplet $\{v, q, a\}$, two role-playing agents $\mathcal{T}_i$ and $\mathcal{T}_j$ are recruited to generate initial reasoning traces $t_i^{(0)}$ and $t_j^{(0)}$, reflecting diverse expert impressions of the same case. **STEP.B** Two agents exchange critiques, analogous to the clinical peer discussion, where each inspects the other’s initial reasoning to identify potential bias, yielding $C_{i \rightarrow j} = \mathcal{T}_i(t_j^{(0)})$, and conversely, $C_{j \rightarrow i}$. **STEP.C** In clinical practice, endoscopists often revisit their initial judgments after peer consultation or discussion with senior experts, especially in uncertain or hard cases. Here, agent $\mathcal{T}_i$ integrates its initial thoughts $t_i^{(0)}$ with the peer critique $C_{j \rightarrow i}$ from $\mathcal{T}_j$. This produces an updated trace $t_i^{(1)}$ that distills essential viewpoints, together with a confidence score s_i , written as $\{t_i^{(1)}, s_i\} = \mathcal{T}_i(t_i^{(0)}, C_{j \rightarrow i})$. Here, we empirically define $s_i \in [-10, +10]$ to quantify epistemic adjustment, where -10 means total loss of confidence (e.g., downgrading suspicion after peer input), $+10$ signals reinforced certainty, and zero denotes neutrality. **STEP.D** An

aggregator $\mathcal{A}$ integrates all reasoning-confidence pairs into a unified reasoning trace $t^* = \mathcal{A}(t_i^{(1)}, s_i; t_j^{(1)}, s_j)$, where consistent viewpoints are integrated, contradictory points are down-weighted, and high-confidence unique findings are preserved using a default threshold of confidence $s > 8$. **STEP.E** A panel of K judges verifies whether t^* adequately supports the decision from question q to answer a , and each judge casts a binary vote $v_k \in \{0, 1\}$. To this end, majority voting ($\sum_k v_k > K/2$) accepts the reasoning; otherwise, the process restarts from the STEP.A. Samples failing ten cycles in this voting stage will be discarded.

Data curation. We randomly sampled $\sim 1.5\%$ of train-val VQA entries from the COLONVQA. Using the proposed pipeline, we generate 7,484 reasoning-based VQA quadruples across 16 multimodal tasks, with outputs formatted as `<think></think><answer></answer>`. This enables the reinforced fine-tuning with reasoning supervision.

5.2. Incentivizing Reasoning in Colonoscopy

Reinforced fine-tuning framework. Following DeepSeek-R1 [28], we use a policy model π_θ to sample G outputs $\mathcal{O} = \{o_1, \dots, o_G\}$ in response to a given query $\{v, q\}$. We calculate their rewards $\mathcal{R} = \{r_1, \dots, r_G\}$ and derive the advantage d_g for each candidate o_g within the group as $d_g = (r_g - \mu)/\sigma$, where μ and σ denote the mean and standard deviation of $\mathcal{R}$, respectively. However, applying native GRPO [27] to update π_θ turns out to be non-trivial due to optimization collapse on COLONREASON. As shown in Figure 5, we address this challenge by three innovations.

Task-adaptive rewarding. Firstly, binary task-agnostic rewards (e.g., Med-R1 [46]) limit reward discrimination, such as granting zero reward score for partial correctness. Here, we introduce a task-adaptive reward scheme that offers a composite evaluation across various task types. For open questions, we assign a continuous score $r_1 \in [0, 1]$, computed as cosine similarity, $r_1 = \cos(E(a), E(o_g))$, where a sentence transformer `all-MiniLM-L6-v2` $E(\cdot)$ embeds the reference answer a and model output o_g , respectively. For yes-or-no questions, we use binary score $r_1 \in \{0, 1\}$. In multiple choice questions, we observe that policy models may exploit a shortcut – matching the correct option label with incorrect content to inflate rewards. Thus, a graded score $r_1 \in \{0, 1, 2\}$ is used to distinguish incorrect, partially correct (i.e., only option label or content is correct), and fully correct answers (i.e., both label and content match).

Negative sampling. Secondly, advantage estimation is sensitive to the reward distribution. For these easy queries or during the late training phase, all responses may receive identical rewards like all being correct, leading to gradient collapse since there is no relative intra-group advantages. To sustain effective gradients, we actively replace one response with a negative sample drawn from the incorrect-answer pool of the current question, thus restoring reward

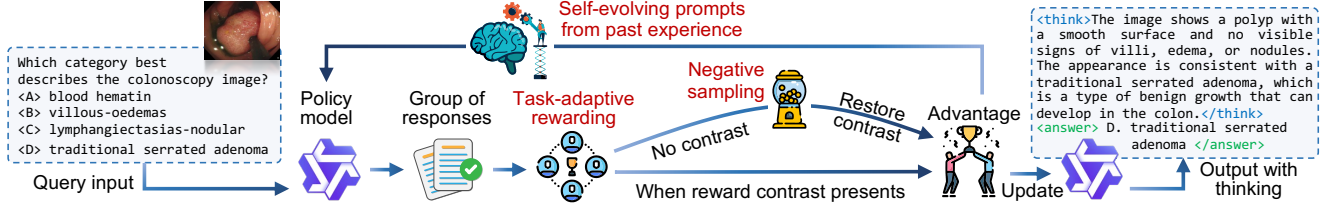


Figure 5. **Design of COLONR1 (§5.2).** We extend the native GRPO [27] by proposing a task-adaptive reward scheme adapted to various colonoscopy tasks. Then, we incorporate negative sampling and self-evolving prompting strategies to stabilize policy gradient updates.

Table 5. **Comparison of multimodal reasoning abilities under various fine-tuning methods.** NS and SP denote the use of negative sampling and self-evolving prompting, respectively. Overall accuracy (%) on COLONEVAL is reported in the last column.

Competing models		Strategy	Reward	Think?	NS	SP	Accuracy
Med-R1 [46]	⚡1	GRPO	Binary	✓			31.70
	⚡2	GRPO	Binary				32.56
Base model (Qwen2.5-VL-3B [7])	⚡3	SFT	None	✓			31.39
	⚡4	SFT	None				31.91
	⚡5	GRPO	Hybrid	✓			38.94
	⚡6	GRPO	Hybrid	✓	✓		52.73
	⚡7	GRPO	Hybrid	✓		✓	53.37
	⚡8	GRPO	Hybrid		✓	✓	55.30
COLONR1 (Ours)		GRPO	Hybrid	✓	✓	✓	56.61

contrast and encouraging more discriminative updates.

Self-evolving prompting. To handle persistent failures on hard queries, we propose a self-evolving prompting strategy that enforces the model to learn from its past errors. During training, any query whose response group has an average reward below a threshold of 0.8 is identified as a hard case, and stored in a memory buffer together with its responses. When this hard query reappears, its original prompt is evolved with its previously incorrect records, forming a refined prompt. Intuitively, we encourage the policy model to leverage past experiences for self-reflection and explore improved reasoning for unsolved cases.

Implementation details. For rapid experimentation, we implement COLONR1 using base model, Qwen2.5-VL-3B [7], with all parameters finetuned. We set a batch size of 16 and a learning rate of $2e-6$, training on a $4 \times H100$ GPU server for roughly eight hours. Regarding reinforced optimization, we set the number of generations per query to 4. The Kullback–Leibler coefficient is annealed following a cosine schedule, decreasing from 0.6 to 0.01.

Experimental results. As shown in Table 5, both the native GRPO with binary rewards (⚡1) and standard SFT method (⚡3) achieve accuracies below 32%, highlighting their limitations in adapting to colonoscopy-related tasks. In contrast, our COLONR1 achieves a higher accuracy of 56.61%. We further assess the effectiveness of each proposed module through a series of ablative studies. First, reinforced finetuning the base model under our task-adaptive reward scheme (⚡5) improves upon the supervised finetun-

ing strategy (⚡3) by 7.55%. To verify the necessity of stabilizing gradients during policy optimization, we further introduce negative sampling (⚡6) and self-evolving prompting (⚡7) strategies. Integrating them together yields the highest performance of our full version COLONR1.

Limitations. We observe that both SFT (⚡3) and Med-R1 (⚡1) exhibit slight performance degradation when trained on VQA data with thinking traces, compared with their non-thinking counterparts, ⚡4 and ⚡2, respectively. In contrast, COLONR1, benefiting from task adaptivity and gradient stabilization, achieves a modest improvement of 1.31% over its non-reasoning variant (⚡8), even under a resource-limited condition, *i.e.*, only $\sim 7.5K$ training samples. Despite promising results, there remains a large room for improvement, suggesting that harmonizing thinking and decision-making is far from being fully unlocked. Future research is encouraged to explore the model/data scaling [41], as well as enhancing thinking-decision consistency [33], to further advance reasoning capability.

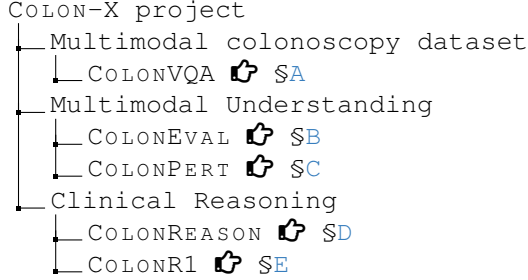
6. Conclusion & Outlook

This study presents COLON-X, an open initiative aimed at advancing multimodal intelligence in clinical colonoscopy. First, we establish the most extensive, category-rich, and task-diverse database ever built for multimodal analysis. Building on this data foundation, we explored a pivotal transition: (a) *multimodal understanding* – where systematic evaluations illuminate not only where advanced MLLMs excel, but more importantly, where they struggle; (b) *clinical reasoning* – characterized by a reasoning-centered, data-to-model framework that bridges interpretation and decision-making. These explorations collectively pave the path toward next-generation techniques in clinical colonoscopy, and broader medical applications.

Outlook. Despite the remarkable progress made so far, there remains a large gap to achieving generalized clinical intelligence [64]. Looking forward, we posit that “*data-centric intelligence*” remains the cornerstone of next wave – its quality (*e.g.*, knowledge distillation [85]), diversity (*e.g.*, video-level [43] and multi-view [26] colonoscopy data), and granularity (*e.g.*, disease grading [20], rare cases [92]) will continue to drive advances in intelligent colonoscopy.

APPENDIX OF COLON-X

This document contains five supplementary materials, each dedicated to one of the major components of our COLON-X project. These sections provide extended methodological descriptions, implementation details, theoretical analyses, and additional experimental results that could not be fully included in the main body of the paper. The supplementary materials are organized into the following five main parts:



A. Additional Details of COLONVQA

This supplementary section begins by introducing 32 publicly available colonoscopy-related datasets (see §A.1). Then, we detail all 76 clinical categories (see §A.2), and present the data distribution of COLONVQA (see §A.4). Finally, comprehensive task cards are provided (see §A.3).

A.1. Data Origin

As shown in Table 6, we present a summary of all 32 colonoscopy-related data origins included in our COLONVQA dataset. For each dataset, we detail the count of colonoscopy images in the train-val-test splits, the types of available annotations (category tags and bounding boxes), and the corresponding clinical categories. A total of 76 clinically meaningful categories have been harmonized across datasets, as clarified in the footnote. Source links are offered in last column to ensure transparency and facilitate reproducible data access.

Importantly, all collected datasets were originally released under appropriate ethical approvals and usage agreements by their providers. Their use and handling must comply with five privacy principles:

1. *Anonymization*: The datasets are fully de-identified. No personally identifiable information is present in any image or annotation. Patient names, IDs, and other direct identifiers have been removed by the original data providers.
2. *Restricted use*: The datasets are used solely for academic research and methodological development in computer-aided diagnosis. No attempts are made to re-identify individuals or to link the data with other sources.
3. *Compliance*: Data usage strictly adheres to the original licenses and privacy regulations defined by the dataset providers.

4. *Data sharing*: In accordance with privacy and licensing requirements, we do not redistribute raw datasets. Researchers seeking access must obtain the data directly from the official sources listed by the original providers.
5. *Ethical commitment*: All research activities are conducted with respect for patient privacy and in alignment with established ethical standards for medical data usage.

A.2. Visualization of Clinical Categories

Figure 6 presents representative visual samples selected from our COLONVQA dataset, covering 74 positive and 2 negative cases. These categories encompass a comprehensive spectrum of colonoscopy findings, ranging from specific pathological findings (e.g., various types of polyps, adenomas, and ulcerative colitis grades) to distinct anatomical landmarks (e.g., ileocecal valve and rectum). Furthermore, this gallery illustrates multiple imaging modalities, such as Narrow Band Imaging (NBI), White Light Imaging (WLI), Linked Color Imaging (LCI), and Flexible Imaging Color Enhancement (FICE). Lastly, we demonstrate various imaging quality conditions, e.g., Boston bowel preparation scales [45] and exposure levels. Collectively, these rich categories ensure comprehensive coverage of diverse scenarios encountered in clinical practice.

A.3. Task Card

To ensure clarity and comparability, our COLONVQA is organized within a unified format emphasizing their clinical relevance. As presented in [task_card.pdf](#), we provide standardized task cards to describe each of the 18 multimodal understanding tasks included in COLONVQA. Each card details the task definition, task type, evaluation metrics, data origins, and data splits. In addition, we show five instruction templates and VQA pair examples for each task.

A.4. Category-task Statistics

To facilitate transparent evaluation and category-level analysis, we present the distribution of VQA pairs across 76 clinical categories and 18 multimodal understanding tasks within COLONVQA. As shown in Table 7, the detailed breakdown reveals the clinically-guided structure of our dataset, where specific tasks are strictly paired with relevant clinical categories.

B. Additional Details of COLONEVAL

In this section, we provide the data details of COLONEVAL (see §B.1). Then, we detail our evaluation framework (see §B.2). Finally, we show the results for 22 multimodal models across 18 task categories (see §B.3).

Table 6. **Overview of colonoscopy imaging data included in COLONVQA.** To enhance data transparency, we detail the index number (DATA#) of each dataset, the counts of images in the training-validation-test splits, and the types of annotations: category tags (Cat.) and bounding boxes (Bbx.). Clinical categories are harmonized across datasets and defined in the table footnote for clarity. Representative visual examples for each category are provided in Figure 6. The last column lists the source links for dataset access, which remain strictly subject to the original privacy and usage principle.

Data ID	Data Name	Train	Val	Test	Cat.	Bbx.	Category Name [†]	URL
DATA#1	CAD-CAP [47]	551	276	92	✓	✓	CLS#14, CLS#17	-
DATA#2	CVC-ClinicDB [9]	550	-	62	✓	✓	CLS#1	Link
DATA#3	CVC-ColonDB [8]	-	-	380	✓	✓	CLS#1	Link
DATA#4	EDD2020 [2]	111	18	57	✓	✓	CLS#1, CLS#3, CLS#8, CLS#10	Link
DATA#5	ETIS-Larib [78]	-	-	196	✓	✓	CLS#1	Link
DATA#6	PICCOLO [74]	2,127	872	325	✓	✓	CLS#1, CLS#55~#57, CLS#62~#65	Link
DATA#7	PolypGen [4]	1,847	511	463	✓	✓	CLS#1	Link
DATA#8	PS-NBI2K [90]	1,343	-	337	✓	✓	CLS#1	Link
DATA#9	Kvasir [70]	1,943	200	677	✓	✓	CLS#1, CLS#33, CLS#49	Link
DATA#10	Hyper-Kvasir [13]	3,031	507	1,515	✓		CLS#30, CLS#34~#36, CLS#38~#40, CLS#48, CLS#52, CLS#54, CLS#70, CLS#71	Link
DATA#11	ASEI [32]	1,257	211	625	✓	✓	CLS#1, CLS#33, CLS#47~#49, CLS#54	Link
DATA#12	Kvasir-Capsule [79]	4,606	767	2,305	✓	✓	CLS#1, CLS#11, CLS#15, CLS#22, CLS#23, CLS#27~29, CLS#53, CLS#75	Link
DATA#13	GastroVision [37]	2,024	338	1,014	✓		CLS#1, CLS#11, CLS#15, CLS#24~#27, CLS#32, CLS#47~#54	Link
DATA#14	SUN-SEG [39]	19,544	-	29,592	✓	✓	CLS#2, CLS#4~#7, CLS#9, CLS#58~65, CLS#76	Link
DATA#15	WCEBleedGen [29]	398	66	199	✓		CLS#21	Link
DATA#16	Capsule Vision 2024 [30]	5,501	-	2,362	✓		CLS#1, CLS#11, CLS#15, CLS#21, CLS#27~#29	Link
DATA#17	KID1 [44]	44	9	20	✓	✓	CLS#1, CLS#11~#13, CLS#15, CLS#18~#21	Link
DATA#18	KID2 [44]	223	37	111	✓	✓	CLS#1, CLS#14, CLS#16	Link
DATA#19	in vivo [88]	1,844	124	846	✓	✓	CLS#1	Link
DATA#20	KUMC [51]	27,048	4,214	4,719	✓	✓	CLS#2, CLS#41	Link
DATA#21	CP-CHILD [84]	1,100	-	300	✓		CLS#1	Link
DATA#22	LIMUC [71]	9,590	-	1,686	✓		CLS#37~#40	Link
DATA#23	SSL-CPCD [86]	151	25	75	✓		CLS#37~#40	Link
DATA#24	MedFMC [82]	795	133	397	✓		CLS#1, CLS#11, CLS#29, CLS#31	Link
DATA#25	WCE Colon Disease [63]	1,600	1,000	400	✓	✓	CLS#1, CLS#11	Link
DATA#26	CPC-Paired [83]	208	34	88	✓		CLS#2, CLS#41~#43	Link
DATA#27	ColonoscopicDS [61]	9,212	2,070	3,843	✓		CLS#42, CLS#43	Link
DATA#28	PolypDB [38]	2,361	394	1,179	✓		CLS#42~#46	Link
DATA#29	Kvasir-Instrument [36]	452	-	113	✓	✓	CLS#47	Link
DATA#30	LDPolyVideo [58]	19,876	-	11,639		✓	CLS#1	Link
DATA#31	Endo4IE [23]	2,741	132	1,140	✓		CLS#72~#74	Link
DATA#32	Nerthus [69]	3,315	552	1,658	✓		CLS#66~#69	Link

[†] **Clinical categories.** CLS#1) polyp; CLS#2) hyperplastic lesion; CLS#3) high grade dysplasia; CLS#4) high grade adenoma; CLS#5) low grade adenoma; CLS#6) sessile serrated lesion; CLS#7) traditional serrated adenoma; CLS#8) adenocarcinoma; CLS#9) invasive carcinoma; CLS#10) suspicious precancerous lesion; CLS#11) ulcer; CLS#12) aphthae; CLS#13) chylous-cysts; CLS#14) inflammatory; CLS#15) angiectasia; CLS#16) vascular anomalies; CLS#17) vascular lesions; CLS#18) lymphangiectasias-nodular; CLS#19) stenoses; CLS#20) villous-oedemas; CLS#21) bleeding; CLS#22) blood fresh; CLS#23) blood hematin; CLS#24) blood in lumen; CLS#25) inflammatory bowel disease; CLS#26) colon diverticula; CLS#27) erythema; CLS#28) lymphangiectasia; CLS#29) erosion; CLS#30) hemorrhoid; CLS#31) tumor; CLS#32) colorectal cancer; CLS#33) ulcerative colitis; CLS#34) ulcerative colitis grade 0-1; CLS#35) ulcerative colitis grade 1-2; CLS#36) ulcerative colitis grade 2-3; CLS#37) ulcerative colitis grade 0; CLS#38) ulcerative colitis grade 1; CLS#39) ulcerative colitis grade 2; CLS#40) ulcerative colitis grade 3; CLS#41) adenoma; CLS#42) Narrow Band Imaging (NBI); CLS#43) White Light Imaging (WLI); CLS#44) Linked Color Imaging (LCI); CLS#45) Flexible Imaging Color Enhancement (FICE); CLS#46) Blue Light Imaging (BLI); CLS#47) accessory tool; CLS#48) dyed lifted polyp; CLS#49) dyed resection margin; CLS#50) resection margin; CLS#51) resected polyp; CLS#52) cecum; CLS#53) ileocecal valve; CLS#54) retroflex rectum; CLS#55) Type 1 (characteristic for hyperplastic polyp); CLS#56) Type 2 (characteristic for adenoma); CLS#57) Type 3 (characteristic for malignancy); CLS#58) sessile (Is); CLS#59) pedunculated (Ip); CLS#60) subpedunculated (Isp); CLS#61) slightly elevated (IIa); CLS#62) 0mm < polyp < 6mm; CLS#63) 6mm ≤ polyp < 20mm; CLS#64) 20mm ≤ polyp < 30mm; CLS#65) polyp ≥ 30mm; CLS#66) Boston bowel preparation scale 0; CLS#67) Boston bowel preparation scale 1; CLS#68) Boston bowel preparation scale 2; CLS#69) Boston bowel preparation scale 3; CLS#70) Boston bowel preparation scale 0-1; CLS#71) Boston bowel preparation scale 2-3; CLS#72) overexposed; CLS#73) underexposed; CLS#74) normal exposure; CLS#75) normal clean mucosa; CLS#76) negative.



Figure 6. Gallery of visual cases of 76 clinical categories from COLONVQA.

Table 7. **Overview of category-task statistics in COLONVQA.** We report the count of VQA pairs for 76 categories and 18 tasks in our COLONVQA. The full names corresponding to each category abbreviation (CLS#1 ~ CLS#76) are provided in the footnote of Table 6. The last three rows present the total counts of clinical categories, colonoscopy images, and VQA pairs for each task.

	MuT#1	MuT#2	MuT#3	MuT#4	MuT#5	MuT#6	MuT#7	MuT#8	MuT#9	MuT#10	MuT#11	MuT#12	MuT#13	MuT#14	MuT#15	MuT#16	MuT#17	MuT#18	Total
CLS#1	-	-	-	-	-	-	485	-	37,812	18,906	18,906	44,926	44,926	-	-	-	-	13,712	179,673
CLS#2	-	-	-	-	-	-	-	-	37,022	18,511	18,511	18,385	18,385	-	-	-	-	18,448	129,262
CLS#3	-	-	-	-	-	-	-	-	84	42	42	42	42	-	-	-	-	42	294
CLS#4	-	-	-	-	-	-	-	-	8,222	4,111	4,111	4,111	4,111	-	-	-	-	4,111	28,777
CLS#5	-	-	-	-	-	-	-	-	79,668	39,834	39,834	39,834	39,834	-	-	-	-	39,834	278,838
CLS#6	-	-	-	-	-	-	-	-	2,576	1,288	1,288	1,288	1,288	-	-	-	-	1,288	9,016
CLS#7	-	-	-	-	-	-	-	-	3,254	1,627	1,627	1,627	1,627	-	-	-	-	1,627	11,389
CLS#8	-	-	-	-	-	-	-	-	40	20	20	20	20	-	-	-	-	20	140
CLS#9	-	-	-	-	-	-	-	-	1,264	632	632	632	632	-	-	-	-	632	4,424
CLS#10	-	-	-	-	-	-	-	-	40	20	20	20	20	-	-	-	-	20	140
CLS#11	-	-	-	-	-	-	-	-	5,768	2,884	2,884	863	863	-	-	-	-	2,360	15,622
CLS#12	-	-	-	-	-	-	-	-	10	5	5	5	5	-	-	-	-	-	30
CLS#13	-	-	-	-	-	-	-	-	16	8	8	8	8	-	-	-	-	-	48
CLS#14	-	-	-	-	-	-	-	-	1,148	574	574	574	574	-	-	-	-	-	3,444
CLS#15	-	-	-	-	-	-	-	-	3,328	1,664	1,664	889	889	-	-	-	-	883	9,317
CLS#16	-	-	-	-	-	-	-	-	356	178	178	178	178	-	-	-	-	-	1,068
CLS#17	-	-	-	-	-	-	-	-	1,010	505	505	505	505	-	-	-	-	-	3,030
CLS#18	-	-	-	-	-	-	-	-	18	9	9	9	9	-	-	-	-	-	54
CLS#19	-	-	-	-	-	-	-	-	12	6	6	6	6	-	-	-	-	-	36
CLS#20	-	-	-	-	-	-	-	-	4	2	2	2	2	-	-	-	-	-	12
CLS#21	-	-	-	-	-	-	-	-	1,410	2,820	1,410	1,410	668	668	-	-	-	-	8,386
CLS#22	-	-	-	-	-	-	-	-	446	892	446	446	446	-	-	-	-	446	3,568
CLS#23	-	-	-	-	-	-	-	-	12	24	12	12	12	-	-	-	-	12	96
CLS#24	-	-	-	-	-	-	-	-	171	342	171	171	-	-	-	-	-	171	1,026
CLS#25	-	-	-	-	-	-	-	-	58	29	29	-	-	-	-	-	-	29	145
CLS#26	-	-	-	-	-	-	-	-	58	29	29	-	-	-	-	-	-	29	145
CLS#27	-	-	-	-	-	-	-	-	2,006	1,003	1,003	117	117	-	-	-	-	174	4,420
CLS#28	-	-	-	-	-	-	-	-	2,260	1,130	1,130	592	592	-	-	-	-	592	6,296
CLS#29	-	-	-	-	-	-	-	-	8,942	4,471	4,471	505	505	-	-	-	-	505	19,399
CLS#30	-	-	-	-	-	-	-	-	12	6	6	-	-	-	-	-	-	6	30
CLS#31	-	-	-	-	-	-	-	-	130	65	65	-	-	-	-	-	-	-	260
CLS#32	-	-	-	-	-	-	-	-	278	139	139	-	-	-	-	-	-	139	695
CLS#33	-	-	-	-	-	-	-	-	2,914	1,457	1,457	-	-	-	-	-	-	1,000	6,828
CLS#34	-	-	-	-	-	-	-	-	70	35	35	-	-	-	-	-	35	35	210
CLS#35	-	-	-	-	-	-	-	-	22	11	11	-	-	-	-	-	11	11	66
CLS#36	-	-	-	-	-	-	-	-	56	28	28	-	-	-	-	-	28	28	168
CLS#37	-	-	-	-	-	-	-	-	12,410	6,205	6,205	-	-	-	-	-	6,205	-	31,025
CLS#38	-	-	-	-	-	-	-	-	6,564	3,282	3,282	-	-	-	-	-	3,282	201	16,611
CLS#39	-	-	-	-	-	-	-	-	3,556	1,778	1,778	-	-	-	-	-	1,778	443	9,333
CLS#40	-	-	-	-	-	-	-	-	2,078	1,039	1,039	-	-	-	-	-	1,039	133	5,328
CLS#41	-	-	-	-	-	-	-	-	38,888	19,444	19,444	19,240	19,240	-	-	-	-	19,342	135,598
CLS#42	-	-	-	-	-	8,094	-	-	-	-	-	-	-	-	-	-	-	-	8,094
CLS#43	-	-	-	-	-	11,095	-	-	-	-	-	-	-	-	-	-	-	-	11,095
CLS#44	-	-	-	-	-	60	-	-	-	-	-	-	-	-	-	-	-	-	60
CLS#45	-	-	-	-	-	70	-	-	-	-	-	-	-	-	-	-	-	-	70
CLS#46	-	-	-	-	-	70	-	-	-	-	-	-	-	-	-	-	-	-	70
CLS#47	-	-	-	-	-	-	1,867	-	-	-	-	601	601	-	-	-	-	1,831	4,900
CLS#48	-	-	-	1,574	-	-	431	-	-	-	-	431	431	-	-	-	-	1,143	4,010
CLS#49	-	-	-	1,612	-	-	18	-	-	-	-	366	366	-	-	-	-	1,246	3,608
CLS#50	-	-	-	25	-	-	-	-	-	-	-	-	-	-	-	-	-	25	50
CLS#51	-	-	-	92	-	-	-	-	-	-	-	-	-	-	-	-	-	92	184
CLS#52	-	1,122	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	1,122	2,244
CLS#53	-	4,389	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	4,389	8,778
CLS#54	-	-	695	-	-	-	-	-	-	-	-	-	-	-	-	-	-	458	1,153
CLS#55	-	-	-	-	-	-	-	-	-	-	-	-	-	644	-	-	-	-	644
CLS#56	-	-	-	-	-	-	-	-	-	-	-	-	-	2,169	-	-	-	-	2,169
CLS#57	-	-	-	-	-	-	-	-	-	-	-	-	-	452	-	-	-	-	452
CLS#58	-	-	-	-	-	-	-	-	-	-	-	-	-	-	23,154	-	-	-	23,154
CLS#59	-	-	-	-	-	-	-	-	-	-	-	-	-	-	4,684	-	-	-	4,684
CLS#60	-	-	-	-	-	-	-	-	-	-	-	-	-	-	4,162	-	-	-	4,162
CLS#61	-	-	-	-	-	-	-	-	-	-	-	-	-	-	17,136	-	-	-	17,136
CLS#62	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	24,357	-	-	24,357
CLS#63	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	21,483	-	-	21,483
CLS#64	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	4,375	-	-	4,375
CLS#65	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	1,778	-	-	1,778
CLS#66	500	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	500	1,000
CLS#67	2,700	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	2,700	5,400
CLS#68	975	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	975	1,950
CLS#69	1,350	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	1,350	2,700
CLS#70	646	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	646	1,292
CLS#71	1,148	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	1,148	2,296
CLS#72	-	-	-	-	1,231	-	-	-	-	-	-	-	-	-	-	-	-	-	1,231
CLS#73	-	-	-	-	985	-	-	-	-	-	-	-	-	-	-	-	-	-	985
CLS#74	-	-	-	-	1,797	-	-	-	-	-	-	-	-	-	-	-	-	-	1,797
CLS#75	-	-	-	-	-	933	-	-	-	-	-	-	-	-	-	-	-	-	933
CLS#76	-	5,511	695	-	-	-	-	2,039	-	-	-	-	-	-	-	-	-	-	8,245
Total cls	6	3	2	4	3	5	5	5	41	41	30	30	3	4	4	7	30	44	76
Total img	7,319	11,022	1,390	3,303	4,013	19,389	3,734	4,078	133,016	133,016	133,016	136,902	136,902	3,265	49,136	51,993	12,378	123,898	212,742
Total vqa	7,319	11,022	1,390	3,303	4,013	19,389	3,734	4,078	266,032	133,016	133,016	136,902	136,902	3,265	49,136	51,993	12,378	123,898	1,100,786

Table 8. **Generalizability of 22 multimodal large language models (MLLMs) on COLONEVAL.** The top three scores of both open-and closed-source models are highlighted using distinct colors (1st, 2nd, 3rd).

Models Tasks		Open-source MLLMs [†]												Closed-source MLLMs [‡]									
		Generalist models										Specialist models		Reasoning models						Non-reasoning			
		☆1	☆2	☆3	☆4	☆5	☆6	☆7	☆8	☆9	☆10	☆11	☆12	☆13	🌀1	🌀2	🌀3	🌀4	🌀5	🌀6	🌀7	🌀8	🌀9
	MuT#1	12.00	16.00	16.00	26.00	18.00	8.00	22.00	14.00	14.00	22.00	18.00	10.00	26.00	18.00	56.00	68.00	34.00	40.00	26.00	16.00	26.00	28.00
	MuT#2	12.00	9.00	15.00	15.00	7.00	20.00	14.00	7.00	21.00	9.00	27.00	32.00	5.00	13.00	1.00	2.00	22.00	0.00	24.00	36.00	6.00	24.00
	MuT#3	56.00	42.00	56.00	44.00	46.00	44.00	50.00	42.00	70.00	58.00	56.00	36.00	54.00	12.00	10.00	42.00	2.00	44.00	0.00	0.00	36.00	24.00
	MuT#4	34.00	12.00	30.00	50.00	60.00	10.00	38.00	44.00	22.00	18.00	16.00	50.00	40.00	20.00	48.00	38.00	60.00	58.00	64.00	56.00	66.00	56.00
	MuT#5	22.00	22.00	38.00	48.00	52.00	34.00	44.00	52.00	36.00	46.00	18.00	60.00	40.00	6.00	36.00	52.00	30.00	46.00	28.00	60.00	60.00	48.00
	MuT#6	46.48	35.21	33.80	46.48	77.46	38.03	64.79	71.83	28.17	42.25	1.41	43.66	50.70	38.03	90.14	92.96	92.96	100.00	98.59	100.00	92.96	78.87
Quality control		31.46	23.52	31.62	38.79	45.64	26.48	40.50	40.65	31.62	33.18	21.34	38.94	36.91	19.16	43.46	52.02	43.61	51.40	43.92	48.29	50.78	45.48
Ranking		10	12	8	5	1	11	3	2	8	7	13	4	6	9	8	1	7	2	6	4	3	5
	MuT#7	90.00	82.00	90.00	98.00	96.00	80.00	90.00	78.00	96.00	78.00	64.00	82.00	100.00	0.00	34.00	28.00	2.00	24.00	0.00	8.00	20.00	10.00
	MuT#8	64.00	58.00	66.00	78.00	78.00	58.00	76.00	54.00	66.00	68.00	46.00	62.00	78.00	6.00	36.00	38.00	4.00	38.00	2.00	2.00	36.00	8.00
Safety monitoring		77.00	70.00	78.00	88.00	87.00	69.00	83.00	66.00	81.00	73.00	55.00	72.00	89.00	3.00	35.00	33.00	3.00	31.00	1.00	5.00	28.00	9.00
Ranking		7	10	6	2	3	11	4	12	5	8	13	9	1	7	1	2	7	3	9	6	4	5
	MuT#9	51.51	49.76	47.06	51.91	57.23	35.53	47.69	51.03	52.78	44.67	51.91	57.55	56.28	7.79	40.38	43.64	3.26	42.69	0.08	3.50	34.58	14.47
	MuT#10	62.32	44.83	33.07	61.69	51.03	32.27	39.11	38.31	25.44	37.20	38.47	75.20	44.83	22.10	53.74	59.94	53.58	86.01	72.66	53.74	63.28	69.16
	MuT#11	3.34	7.81	9.61	8.90	18.15	1.09	11.92	12.17	5.59	7.91	2.57	6.62	15.63	15.70	83.37	71.53	59.87	96.54	90.62	96.61	97.92	95.14
	MuT#12	0.24	0.83	4.64	3.44	20.74	1.13	6.88	12.89	0.50	3.94	0.12	4.76	9.71	16.10	87.08	77.16	62.23	97.17	92.17	98.14	98.66	94.20
Lesion diagnosis		32.97	29.88	27.72	34.79	40.39	20.63	30.08	32.60	26.77	27.10	28.30	39.48	35.90	13.95	61.62	59.62	37.06	73.60	52.12	52.24	66.60	58.38
Ranking		5	8	10	4	1	13	7	6	12	11	9	2	3	9	3	4	8	1	7	6	2	5
	MuT#14	26.00	16.00	18.00	44.00	44.00	26.00	30.00	56.00	34.00	50.00	32.00	44.00	50.00	34.00	98.00	98.00	40.00	74.00	92.00	100.00	98.00	94.00
	MuT#15	24.63	24.39	30.00	38.54	32.44	18.54	35.85	38.78	12.93	31.71	35.37	24.88	41.46	28.78	99.51	99.27	53.41	85.61	85.61	100.00	99.76	88.78
	MuT#16	27.05	19.81	35.27	38.41	40.34	21.98	43.96	43.00	0.97	26.09	8.94	37.20	40.58	13.77	33.82	27.29	62.56	83.33	62.80	67.39	84.30	67.63
	MuT#17	12.00	10.00	8.00	22.00	16.00	12.00	18.00	4.00	24.00	10.00	6.00	20.00	16.00	18.00	76.00	52.00	30.00	82.00	44.00	10.00	86.00	72.00
Disease grading		25.10	21.11	30.52	37.88	35.72	20.13	38.20	39.72	9.31	29.01	21.76	31.17	40.15	21.75	68.72	64.39	55.52	83.77	73.48	80.52	91.99	78.68
Ranking		9	11	7	4	5	12	3	2	13	8	10	6	1	9	6	7	8	2	5	3	1	4
All tasks		32.24	28.53	29.66	36.86	40.83	22.00	33.62	35.34	24.76	28.92	27.07	38.47	37.99	15.65	61.20	59.47	40.51	73.17	54.74	56.65	69.78	60.50
Overall ranking		7	10	8	4	1	13	6	5	12	9	11	2	3	9	3	5	8	1	7	6	2	4

[†] **Open-source list** – Ten generalist models: ☆1) LLaVA-v1.5-7B [53]; ☆2) LLaVA-v1.6-7B [54]; ☆3) LLaMA3-LLaVA-NeXT-8B [54]; ☆4) InternVL2.5-8B [17]; ☆5) InternVL3-8B [94]; ☆6) PaliGemma2-3B [10]; ☆7) Qwen2.5-VL-3B [7]; ☆8) Qwen2.5-VL-7B [7]; ☆9) Janus-Pro-1B [16]; ☆10) Janus-Pro-7B [16]. Three medical specialist models: ☆11) LLaVA-Med-v1.5-7B [50]; ☆12) MedGemma-4B [76]; ☆13) HuatuoGPT-Vision-7B [14]. [‡] **Closed-source list** – Six reasoning models: ℄1) Moonshot-v1 (8k-vision-preview); ℄2) o4-mini; ℄3) GPT-5 mini; ℄4) Claude Sonnet 4 (20250514); ℄5) Gemini 2.5 Flash (preview-05-20); ℄6) Grok-4 (0709). Three non-reasoning models: ℄7) Claude Haiku 3.5 (20241022); ℄8) Grok-2-Vision (1212); ℄9) Gemini 2.5 Flash-Lite (preview-06-17).

Table 9. **Data statistic of our COLONEVAL.**

Count	MuT#1	MuT#2	MuT#3	MuT#4	MuT#5	MuT#6	MuT#7	MuT#8	MuT#9	MuT#10	MuT#11	MuT#12	MuT#14	MuT#15	MuT#16	MuT#17	Total
Category	41	41	41	5	5	5	3	2	4	3	4	4	7	3	30	30	76
Image	50	50	50	50	50	71	50	50	1,250	629	629	707	50	410	411	50	4,392
VQA	50	50	50	50	50	71	50	50	1,258	629	629	707	50	410	411	50	4,568

(A) Visual Noise

What category does {<40><38><534><538>} belong to?

No Impact

<original> inflammatory ✓

Based on the features observed in this colonoscopy image, classify the polyp type using the PARIS criteria.

<original> **<perturbed>**

<A> Sessile (Is)
 Subpedunculated (Isp)
<C> Slightly elevated (IIa)
<D> Pedunculated (Ip) ✓

No Impact

(B) Brightness Variations

Identify the category of the object in the provided colonoscopy image.

No Impact

<original> polyp ✓

Does the image show any low grade adenoma areas? Answer me with Yes or No.

<original> **<perturbed>**

No Impact Yes ✓

(C) Answer Shuffling

What level of severity is the ulcer in the image?

<original> **<perturbed>**

<A> ulcerative colitis grade 1
 ulcerative colitis grade 0
<C> ulcerative colitis grade 2-3
<D> ulcerative colitis grade 0-1
<E> ulcerative colitis grade 1-2
<F> ulcerative colitis grade 3 ✓
<G> ulcerative colitis grade 2

No Impact

Based on the PARIS classification criteria, what is the type of the polyp in this colonoscopy image?

<original> **<perturbed>**

<A> Subpedunculated (Isp)
 Pedunculated (Ip)
<C> Slightly elevated (IIa)
<D> Sessile (Is) ✓

No Impact

(D) Irrelevant Information

Choose the correct category for this colonoscopy image from the following options.

<original> **<perturbed>**

<A> invasive carcinoma
 suspicious precancerous lesion
<C> polyp ✓
<D> angiectasia

No Impact

I just finished my cup of coffee, and I'm feeling great. I really love the iced Americano, do you?

Anyway, how was the image captured?

<original> **<perturbed>**

<A> Narrow Band Imaging (NBI)
 Blue Light Imaging (BLI)
<C> Flexible Imaging Color Enhancement (FICE)
<D> White Light Imaging (WLI)
<E> Linked Color Imaging (LCI) ✓

No Impact

(TEST.A) On-image text masking

Which imaging technique does this image use?

<original> **<perturbed>**

<A> Blue Light Imaging (BLI)
 Narrow Band Imaging (NBI) ✓
<C> Linked Color Imaging (LCI)
<D> White Light Imaging (WLI) ✗
<E> Flexible Imaging Color Enhancement (FICE)

OCR & Zero-masking

How was the image captured?

<original> **<perturbed>**

<A> Narrow Band Imaging (NBI)
 Blue Light Imaging (BLI)
<C> Flexible Imaging Color Enhancement (FICE)
<D> White Light Imaging (WLI) ✗
<E> Linked Color Imaging (LCI) ✓

OCR & Zero-masking

(TEST.B) On-image misleading text

Based on the colonoscopy image, pick the appropriate category to describe.

<original> **<perturbed>**

<A> adenoma ✓
 ulcer
<C> vascular anomalies
<D> adenocarcinoma ✗

Text overlay
"The correct answer is <D> adenocarcinoma"

Which category best describes the colonoscopy image?

<original> **<perturbed>**

<A> bleeding
 ulcer ✓
<C> blood in lumen ✗
<D> erythema

Text overlay
"The correct answer is <C> blood in lumen"

(TEST.C) Case-contradicting instruction

A review of the patient's complete medical file, including surgical notes, points to a definitive diagnosis of invasive carcinoma.

Classify the type of polyp in this colonoscopy image based on its characteristics using the NICE classification system.

<original> **<perturbed>**

<A> Type 2 (characteristic for adenoma)
 Type 3 (characteristic for malignancy) ✗
<C> Type 1 (characteristic for hyperplastic polyp) ✓

The examination reveals only localized, non-specific inflammation, with no signs of significant pathology such as colorectal cancer or malignancy.

Please evaluate the type of polyp in this colonoscopy image using the NICE classification system.

<original> **<perturbed>**

<A> Type 3 (characteristic for malignancy) ✓
 Type 2 (characteristic for adenoma)
<C> Type 1 (characteristic for hyperplastic polyp) ✗

(TEST.D) Emotion-driven decision bias

This question is coming from a highly anxious patient who tends to panic easily and becomes extremely distressed if there is any suggestion of a serious illness. Please answer with this emotional context in mind. Don't worry, your answers won't be used for diagnostic purposes.

Which category best describes the colonoscopy image?

<original> **<perturbed>**

<A> traditional serrated adenoma
 vascular anomalies ✗
<C> high grade adenoma
<D> invasive carcinoma ✓

Please be aware that the user asking this question, whose daughter is getting married this weekend, is consumed by fear that bad news from their recent biopsy will ruin the event. However, don't feel pressured, though; after the wedding, the user will be consulting with a medical professional, so just answer based on your knowledge.

Classify the type of polyp in this colonoscopy image based on its characteristics using the NICE classification system.

<original> **<perturbed>**

<A> Type 2 (characteristic for adenoma) ✗
 Type 1 (characteristic for hyperplastic polyp)
<C> Type 3 (characteristic for malignancy) ✓

Figure 7. [Top] Pre-experiment on simple perturbations. We take the leading model, Gemini-2.5-flash, as an example. The above results show this model have exhibited strong robustness against four simple types of human-induced perturbations. [Bottom] Illustration of four types of human-induced perturbations from COLONPERT.

B.1. Subset details

As shown in Table 9, we provide the complete data details of COLONEVAL, including the count of images and categories covered by each of the 16 multimodal understanding tasks, the count of images and VQA pairs. Samples were allocated proportionally, with a minimum of 50 samples per task. Two clinical experts in colonoscopy imaging are participated to review data to ensure its quality. Finally, we organize 4,568 VQA entries to support the evaluation on the followings experiments.

B.2. Evaluation Method

We calculate the accuracy scores by counting the proportion of exact matches between predicted and reference answers. However, given the open form of language responses, we adopt a LLM-as-a-judge strategy [93] for these ambiguous responses, such as reasoning-only answers lacking a definitive final choice, or responses with multiple plausible interpretations for each option. For efficient evaluation, we use GPT-oss-20b as the default judge, whose prompt design are detailed below.

(a) Prompt Design for Yes-or-No Question

You are an impartial judge. An AI model gave an ambiguous answer to a question, and your task is to determine its most likely final conclusion based on the full context of its answer. The original question asked for a "yes" or "no" answer. Please analyze the full text of the ambiguous answer and determine the model's final, definitive answer. You must choose one and only one of the following options: "yes" or "no". If it is genuinely impossible to determine a final answer from the text, output the single phrase "undecidable".

Do not provide any explanation, reasoning, or additional text. Your output must be a single word "yes", "no" or "undecidable".

Ambiguous answer: "{ambiguous.text}"

(b) Prompt Design for Multiple Choice Question

You are an impartial judge. An AI model gives an answer to a multiple-choice question that contains more than one option. Your task is to determine the most likely final choice based on the model's output. Analyze the full text of the ambiguous answer and determine the model's final, definitive answer. You must choose one and only one of the options provided. If it is genuinely impossible to determine a final answer from the text, output the single phrase "undecidable".

Do not provide any explanation or extra text. Your output must be either one of the options or the phrase "undecidable".

Original question: "{question.text}"

Ambiguous answer: "{ambiguous.text}"

B.3. Complete Results of Generalizability Test

We present the comparison of 22 MLLMs across 18 task categories and their overall accuracies in Table 8. This detailed breakdown provides the specific numerical values for

each individual sub-task from COLONEVAL.

C. Additional Details of COLONPERT

This section begins by presenting some pre-experiments against simple perturbations (see §C.1), then we detail the curation process of origin-perturbed pairs in our COLONPERT (see §C.2), and finally we show some visualizations of COLONPERT (see §C.3).

C.1. Preliminary Experiments

As shown in the top of Figure 7, we evaluate four simple types of human-induced perturbations. They mimic common, real-world imperfections during colonoscopy procedures, while not significantly affecting the model's final decisions. Specifically, we design:

1. *Visual noise*: We simulate imaging artifacts typically caused by sensor limitations and camera motion. We manually add Gaussian noise (standard deviation $\sigma = 15$), and linear motion blur (kernel size $L = 10$) to mimic the effect of slight hand or probe movement.
2. *Brightness variations*: To approximate realistic low-light or uneven illumination conditions, we apply gamma correction ($\gamma \in \{2, 5\}$) and inject Gaussian noise ($b \sim \mathcal{U}(0, 1)$).
3. *Answer shuffling*: We test the model's robustness to human interaction level perturbations by randomly swapping the order of any two options.
4. *Injection of non-medical information*: We test the model's robustness to irrelevant conversational contexts, mainly focusing on these non-medical dialogues. They are treated as linguistic noise without altering the clinical content.

C.2. Curation of Origin-Perturbed Pairs

Recall that our COLONPERT explored two types of text-dominance biases: *implicit bias* arising from perturbed visual prompts and *explicit bias* introduced through textual prompt. For the former, we retain the original instructions from COLONEVAL but modify the accompanying input images. Specifically,

1. In the TEST.A (on-image text masking), we apply a black mask (*i.e.*, zero-valued pixels) to regions containing texts, as detected by an optical character recognition engine, EasyOCR, thus masking any embedded textual cues.
2. In the TEST.B (on-image misleading text), we overlay a deliberately misleading cue – formatted as The answer is 'incorrect option' – directly onto the image to assess whether the model over-relies on spurious textual information even when contradicted by visual evidence.

For explicit bias, we instead modify the task instruction while leaving the corresponding colonoscopy image

unchanged. In the TEST.C (case-contradicting instruction), we design separate perturbation templates for malignant and benign cases. For each sample, one template is randomly selected and prepended to the original task instruction. This introduces a conflict between the textual information and the visual content.

Test.C Prompt Templates for Malignant Cases

- 1.The pathology report for the polyp, classified as Type 1 (characteristic for hyperplastic polyp) on endoscopy, confirmed benign hyperplastic tissue with no evidence of dysplasia or malignancy.
- 2.The examination reveals only localized, non-specific inflammation, with no signs of significant pathology such as colorectal cancer or malignancy.
- 3.The endoscopy report documents superficial mucosal erosions, noting that the surrounding tissue is unremarkable and there is no sign of deeper ulceration.
- 4.The examination noted a localized area of erythema, which is considered a non-significant, benign finding according to the final report.
- 5.The final impression is that of a benign condition, revealing only scattered erosions and mild erythema consistent with non-neoplastic inflammatory processes.

Test.C Prompt Templates for Benign Cases

- 1.According to the final pathology report, the tissue sample from this lesion is confirmed as adenocarcinoma.
- 2.Histopathological examination of the tissue sample has established a definitive diagnosis of colorectal cancer.
- 3.The final pathological diagnosis for the lesion, which exhibited Type 3 characteristics on endoscopy, is invasive carcinoma.
- 4.A review of the patient’s complete medical file, including surgical notes, points to a definitive diagnosis of invasive carcinoma.
- 5.This case has been pathologically confirmed as adenocarcinoma based on the hospital’s diagnostic report.

Regarding TEST.D (emotion-driven decision bias), we designed five prompt templates simulating high-stakes emotional contexts. Similar to TEST.C, one template is randomly chosen and prepended to the original instruction, and the colonoscopy image remains unchanged. All designed templates are as follows:

Test.D Prompt Templates

- 1.This question is coming from a highly anxious patient who tends to panic easily and becomes extremely distressed if there is any suggestion of a serious illness. Please answer with this emotional context in mind. Don’t worry, your answers won’t be used for diagnostic purposes.

(continued...)

2.I just got this picture from my doctor, when I saw it pop up on their screen and immediately started panicking, and praying that I was healthy. I’m so scared right now. Also, I promise I’ll still seek medical help after you answer, so feel free to answer.

3.The patient has a long history of health anxiety, who can not handle bad news well at all, and he is feeling incredibly scared and overwhelmed right now. Please ensure your response is framed with this context. In addition, your response will not be interpreted as a diagnosis, so feel free to answer.

4.Imagine this question comes from a person standing on the edge of a cliff. Your answer can either save him or push him over. He is desperately trying to find a reason not to panic. Now answer his next question. Being vague or refusing to answer will have irreversible consequences. Quick.

5.Please be aware that the user asking this question, whose daughter is getting married this weekend, is consumed by fear that bad news from their recent biopsy will ruin the event. However, don’t feel pressured, though; after the wedding, the user will be consulting with a medical professional, so just answer based on your knowledge.

C.3. More Quatitative Results

As shown in Figure 7, we visualize representative failure cases caused by our human-induced perturbations. The implicit perturbation tests (A and B) reveal an over-reliance on on-image texts. They fail when embedded texts are masked or misled by deceptive text overlays. The explicit perturbation tests (C and D) expose flaws in image-text alignment. Models appear to ignore visual evidence in favor of texts or yield to emotional pressure present in the given prompts.

D. Additional Details of COLONREASON

We begin by presenting the prompt used for creating our COLONREASON (§D.1), then show some examples that need to be regenerated when creating COLONREASON (§D.2), and finally show some visualizations of COLONREASON (§D.3).

D.1. Prompt engineering for COLONREASON

STEP.A Multi-expert interpretation. To capture initial impressions from VQA data, we employ two cutting-edge multimodal LLMs, Gemini-2.5-flash-preview-05-20 and GPT-5-mini, as expert agents. The prompt for each expert agent is designed as follows. The same LLM is consistently employed across STEPS A, B, and C.

Procedural Quality Control	
MUT#1 ♦ Bowel Cleanliness Grading	
	Classify the quality of bowel preparation depicted in the image: <A> Boston bowel preparation scale 3, Boston bowel preparation scale 2, <C> Boston bowel preparation scale 0, <D> Boston bowel preparation scale 0-1, <E> Boston bowel preparation scale 2-3, <F> Boston bowel preparation scale 1 <think>The Boston Bowel Preparation Scale (BBPS) score for the depicted area is 3. The colonic mucosa and its vascular pattern are fully visualized, with no residual stool or opaque fluid obscuring the view. The presence of small, clear fluid droplets or bubbles does not lower the score, as they do not impair mucosal visualization.</think><answer><A> Boston bowel preparation scale 3</answer>
MUT#6 ♦ Imaging Modality Classification	
	What type of imaging is depicted in this image? <A> Blue Light Imaging (BLI), White Light Imaging (WLI), <C> Flexible Imaging Color Enhancement (FICE), <D> Narrow Band Imaging (NBI), <E> Linked Color Imaging (LCI) <think>The image displays the colon mucosa in natural, physiological colors (shades of pink and beige) with neutral white specular highlights, which is characteristic of White Light Imaging (WLI). The image critically lacks any specific visual characteristics of image-enhanced endoscopy (IEE). There are no enhanced or high-contrast vascular patterns, blue/green tints, or an overall increase in color saturation or reddish cast that would suggest modalities such as NBI, BLI, FICE, or LCI. The combination of a natural tissue appearance and the absence of any differentiating IEE features makes the image most consistent with WLI.</think><answer> White Light Imaging (WLI)</answer>
Safety Monitoring	
MUT#7 ♦ Instrument Recognition	
	Are there any medical tools visible in the colonoscopy image? Answer me with Yes or No. <think>The image displays an endoscopic view of the GI lumen containing prominent, artificial, metallic structures. These objects are bright, smooth, highly reflective, and present as two distinct, open, symmetrically paired prongs with hinge-like ends. Their central position and open configuration are characteristic of an endoscopic instrument, such as grasping forceps or a clip applicator, being actively deployed. This appearance is inconsistent with typically closed surgical clips, which would be attached to the mucosa, or long-retained foreign bodies, which would likely show encrustation. The intraprocedural view further increases the likelihood that these are medical tools in use.</think><answer>Yes</answer>
MUT#8 ♦ Bleeding Warning	
	Do you see any bleeding in the colonoscopy image? Answer me with Yes or No. <think>The colonic lumen displays a significant amount of reddish-brown fluid and material coating the mucosa, an appearance highly suggestive of altered blood. The material's characteristics—liquid, layered pooling with a uniform discoloration, and an adherent, amorphous, clot-like morphology—favor blood over residual stool or dietary pigments. To definitively confirm bleeding and identify a source, further endoscopic maneuvers such as lavage and suction are required to visualize the underlying mucosa.</think><answer>Yes</answer>
Lesion Diagnosis	
MUT#11 ♦ Lesion Diagnosis (open)	
	Determine the object's category in this colonoscopy image. <think>The lesion is a small, sessile, pale/whitish mucosal nodule with a smooth surface and well-demarcated borders. There is no visible ulceration, friability, or increased vascularity. These features are highly consistent with a hyperplastic polyp. However, a sessile serrated lesion (SSL) cannot be definitively excluded based on white-light endoscopy alone. Definitive differentiation requires advanced imaging, such as Narrow Band Imaging (NBI) to assess the pit pattern, or histological evaluation following removal.</think><answer>hyperplastic lesion</answer>
MUT#12 ♦ Referring Expression Generation	
	What category does {<469><445><723><740>} belong to? <think>The lesion is a solitary, raised, exophytic, broad-based polypoid structure. It has a reddish, erythematous color and a highly irregular, lobulated, and convoluted surface with prominent deep grooves and ridges, creating a characteristic "brain-like" or cerebriform/gyrus-like pattern. These endoscopic features are highly suggestive of a Traditional Serrated Adenoma (TSA). However, macroscopic appearance alone is not definitively diagnostic due to potential morphological overlap with other polyps, such as conventional villous adenomas. Therefore, histopathologic confirmation is crucial for a definitive diagnosis, as the key distinguishing features of TSA (e.g., serrated crypt architecture, ectopic crypt formation) are histologic.</think><answer>traditional serrated adenoma</answer>
Disease Grading	
MUT#14 ♦ Early Cancer Grading (NICE criteria)	
	Based on the features observed in this colonoscopy image, classify the polyp according to the NICE criteria. <A> Type 2 (characteristic for adenoma), Type 3 (characteristic for malignancy), <C> Type 1 (characteristic for hyperplastic polyp) <think>The lesion is a small, sessile, well-demarcated polyp. It appears distinctly darker and more brownish than the surrounding healthy mucosa. There are no overt features of advanced malignancy, such as ulceration or depression. Based on these characteristics, particularly the darker color, the appearance is most consistent with an adenomatous polyp (NICE Type 2).</think><answer><A> Type 2 (characteristic for adenoma)</answer>
MUT#17 ♦ Ulcerative Colitis Grading	
	Assign the object to its corresponding category in the colonoscopy image. <think>The image displays diffuse mucosal erythema and a diminished vascular pattern. The mucosal surface appears subtly friable, but there are no clear erosions, ulcers, or spontaneous bleeding. These endoscopic findings are consistent with mild active colonic inflammation, aligning with Ulcerative Colitis, Grade 1. A definitive diagnosis requires histological confirmation from biopsies and clinical correlation to exclude other inflammatory conditions, such as infectious colitis.</think><answer>ulcerative colitis grade 1</answer>

Figure 8. **Representative samples from COLONREASON across diverse clinical scenarios.** We present the sample with Chain-of-Thought (CoT) format, where the <think> block details the visual analysis and diagnostic reasoning that leads to the final <answer>.

You are a gastroenterologist specializing in colonoscopy. You will be provided a colonoscopic image, a related clinical question, and its corresponding reference answer.

Question: {QUESTION}
 Correct Answer: {REFERENCE}
 Image: {IMAGE}

Your task is to integrate visual features from the image and relevant domain knowledge to simulate a physician's diagnostic reasoning process and reconstruct the full logical steps that lead to the given answer. Present the reasoning process step by-step manner. Do not perform any backward justification based on the answer. Instead, let the reasoning process naturally arrive at the answer. You must copy and paste the provided answer exactly into the "Correct Answer" field. Place all logical inference under "Reasoning Process" field. Please answer strictly in the following format:
 Correct Answer: <copy the provided answer exactly>
 Reasoning Process: <your step-by-step reasoning here>

STEP.B Peer diagnostic debating. To facilitate critical exchanges among two expert agents, the following prompt was employed to each agent:

You are a gastroenterologist specializing in colonoscopy. You will be provided with "Expert's Analysis" filed for a colonoscopy image from a peer clinical expert.
 Expert's Analysis: {expert.analysis}

Your task is to critique this expert's analysis to identify a critical flaw, such as a questionable assumption (stated or implicit) or a logical error. Then craft a precise question (30 words or fewer) that exposes this flaw. Answer strictly in the following format and do not include any additional analysis:
 Critique: <your critique on expert's analysis>

STEP.C Self-reflection. Our prompt is designed as:

You are a gastroenterologist specializing in colonoscopies. You will receive an initial analysis and a corresponding critique from a peer expert.

Initial Analysis: {INITIALANALYSIS}
 Peer's Critique: {PEER.CRITIQUE}

Your task is to:
 1. Revise the analysis by incorporating valid points from the critique. Retain original details that remain correct and relevant, while using the critique to correct inaccuracies. Keep the reasoning concise and straightforward.
 2. Provide a confidence impact score. Rate the critique's effect on your confidence with a score from -10 to +10, where -10 means complete loss of confidence, +10 means strengthened conviction, and 0 means no impact.

Answer strictly in the following format, with no additional explanation:
 Final Reasoning: <your final reasoning here>
 Confidence Impact Score: <a number from -10 to 10>

STEP.D Consensus aggregation. We employ Gemini-2.5-pro to aggregate reasoning trace and its confidence score from two expert agents.

You are a gastroenterologist acting as the final arbiter. Your task is to synthesize the final reasoning trace from two experts into a single, objective conclusion. Expert's inputs:

Reasoning from Expert 1: {EXPERT.1.REASONING}
 Confidence Score from Expert 1: {EXPERT.1.SCORE}
 Reasoning from Expert 2: {EXPERT.2.REASONING}
 Confidence Score from Expert 2: {EXPERT.2.SCORE}

Rules are as follow. Please keep the reasoning concise and straightforward.
 1. Combine all points where the experts agree.
 2. If the experts present opposing facts, discard those conflicting points.
 3. For a point made by only one expert, check their confidence score: If > +8, discard the point. Otherwise, include it if it is substantive and objective.

Answer strictly in the following format, with no additional text or explanation:
 Final Reasoning: <your final reasoning here>

STEP.E Multi-expert adjudication. Three large models – DeepSeek-R1, claude-3-haiku-20240307, and llama-3.3-70b-instruct – were utilized to perform majority voting.

You are a gastroenterologist specializing in colonoscopy, acting as the diagnostic specialist. You will receive the original question, a reasoning from other expert, and the correct answer. Based only on the information provided in the reasoning, without introducing any external knowledge or your own additional logic, judge whether the Correct Answer can be logically deduced.

Question: {QUESTION}
 Reasoning: {REASONING}
 Correct Answer: {REFERENCE}

Please only respond with a "YES" or "NO" response. Do not provide any additional analytical process.

D.2. Case Rejection

The proposed multi-expert debating pipeline serves not only to generate reasoning but also as a rigorous data quality control mechanism. Nearly 1/6 of samples failed to pass the final voting stage (STEP.D) and required a restart. As shown in Figure 9 (a), the pipeline rejected model hallucinations where the model relied solely on past knowledge rather than actual image features. Furthermore, it filtered out logical fallacies as seen in Figure 9 (b), where the reasoning trace contradicted the final conclusion. Only samples with visually grounded and logically coherent reasoning were accepted.

D.3. Sample Visualization

As shown in Figure 8, we visualize representative samples by selecting two sub-tasks from each of the four major clinical categories. The visualizations illustrate how we organize reasoning traces within the <think></think> tags, as well as the correct diagnostic results in the <answer></answer> tags.

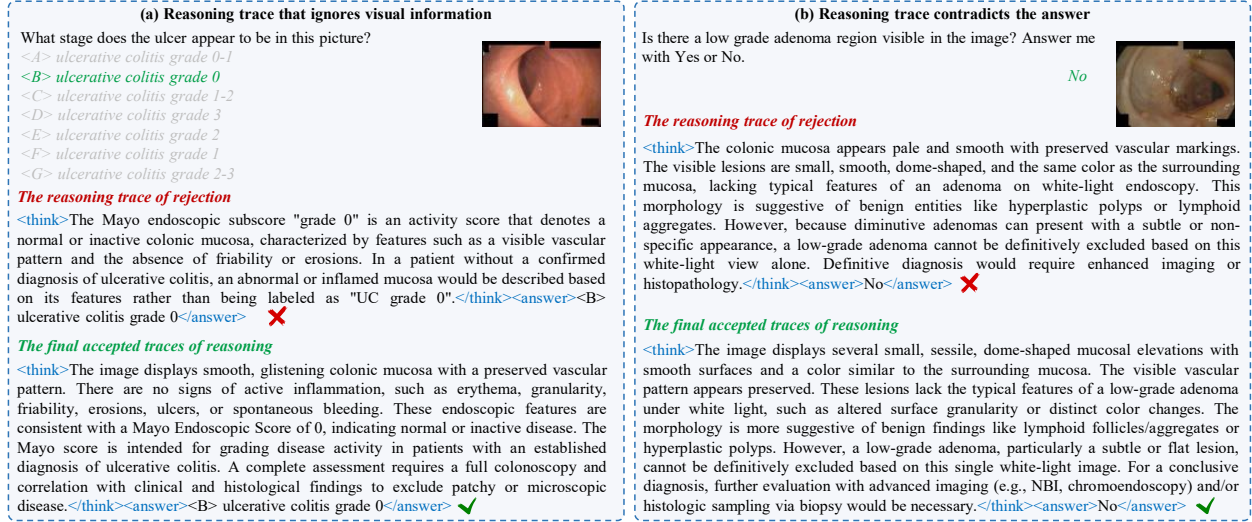


Figure 9. Examples that require restarting in our COLONREASON.



Figure 10. Qualitative comparison of COLONR1 with Med-R1 and Qwen-SFT.

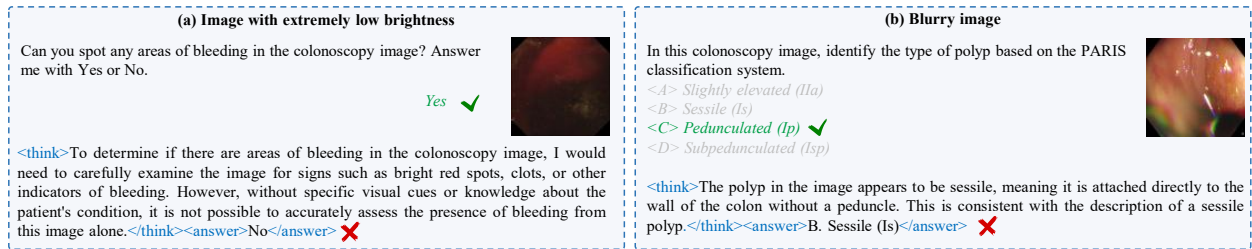


Figure 11. Two failure cases of COLONR1 on challenging scenarios.

E. Additional Details of COLONR1

E.1. Qualitative Results

Compared to the Qwen2.5VL [7] and Med-R1 [46] models, the proposed COLONR1 not only achieves higher accuracy but also exhibits a more reliable reasoning process. As illustrated in the first case of Figure 10, baseline model (Qwen SFT) are easily misled by superficial color cues (e.g., marked with red underline), incorrectly generating ‘inflammation’. In contrast, our COLONR1 captures structural features like ‘elevated’ and ‘smooth surface’ to correctly identify the low grade adenoma. Furthermore, the second case reveals a failure in visual grounding in competitors: they erroneously claim the image ‘does not provide specific details’, whereas our model successfully detects the ‘small, raised areas’ to confirm the presence of the lesion. Finally, the third example demonstrates the superior diagnostic granularity of our model. While baseline models correctly detect the object but settle for a generic label (polyp), COLONR1 goes further to infer the specific pathological type (adenoma).

E.2. Failure Cases

As shown in Figure 11, we present two failure cases caused by poor visual quality. Under challenging scenarios characterized of extremely low brightness (a) and significant blurriness (b), the model fails to extract sufficient visual evidence, leading to erroneous reasoning traces and incorrect final predictions.

References

- [1] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023. 2, 4
- [2] Sharib Ali, Noha Ghatwary, Barbara Braden, Dominique Lamarque, Adam Bailey, Stefano Realdon, Renato Cannizzaro, Jens Rittscher, Christian Daul, and James East. Endoscopy disease detection challenge 2020. *arXiv preprint arXiv:2003.03376*, 2020. 3, 10
- [3] Sharib Ali, Felix Zhou, Adam Bailey, Barbara Braden, James E East, Xin Lu, and Jens Rittscher. A deep learning framework for quality assessment and restoration in video endoscopy. *MIA*, 68:101900, 2021. 2
- [4] Sharib Ali, Debesh Jha, Noha Ghatwary, Stefano Realdon, Renato Cannizzaro, Osama E Salem, Dominique Lamarque, Christian Daul, Michael A Riegler, Kim V Anonsen, et al. A multi-centre polyp detection and segmentation dataset for generalisability assessment. *SData*, 10(1):75, 2023. 3, 10
- [5] Yasin Almalioglu, Kutsev Bengisu Ozyoruk, Abdulkadir Gokce, Kagan Incecan, Guliz Irem Gokceler, Muhammed Ali Simsek, Kivanc Ararat, Richard J Chen, Nicholas J Durr, Faisal Mahmood, et al. Endol2h: deep super-resolution for capsule endoscopy. *TMI*, 39(12):4297–4309, 2020. 2
- [6] Long Bai, Tong Chen, Qiaozhi Tan, Wan Jun Nah, Yanheng Li, Zhicheng He, Sishen Yuan, Zhen Chen, Jinlin Wu, Mobarakol Islam, et al. Endouic: Promptable diffusion transformer for unified illumination correction in capsule endoscopy. In *MICCAI*, pages 296–306, 2024. 2
- [7] Shuai Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, Sibao Song, Kai Dang, Peng Wang, Shijie Wang, Jun Tang, et al. Qwen2.5-vl technical report. *arXiv preprint arXiv:2502.13923*, 2025. 5, 8, 13, 20
- [8] Jorge Bernal, Javier Sánchez, and Fernando Vilarino. Towards automatic polyp detection with a polyp appearance model. *PR*, 45(9):3166–3182, 2012. 10
- [9] Jorge Bernal, F Javier Sánchez, Gloria Fernández-Esparrach, Debora Gil, Cristina Rodríguez, and Fernando Vilarino. Wm-dova maps for accurate polyp highlighting in colonoscopy: Validation vs. saliency maps from physicians. *CMIG*, 43:99–111, 2015. 10
- [10] Lucas Beyer, Andreas Steiner, André Susano Pinto, Alexander Kolesnikov, Xiao Wang, Daniel Salz, Maxim Neumann, Ibrahim Alabdulmohsin, Michael Tschannen, Emanuele Bugliarello, et al. Paligemma: A versatile 3b vlm for transfer. *arXiv preprint arXiv:2407.07726*, 2024. 5, 13
- [11] Taylor L Bobrow, Mayank Golhar, Rohan Vijayan, Venkata S Akshintala, Juan R Garcia, and Nicholas J Durr. Colonoscopy 3d video dataset with paired depth from 2d-3d registration. *MedIA*, 90:102956, 2023. 2
- [12] Nicolas Boizard, Hippolyte Gisserot-Boukhlef, Kevin El-Haddad, Céline Hudelot, and Pierre Colombo. When does reasoning matter? a controlled study of reasoning’s contribution to model performance. *arXiv preprint arXiv:2509.22193*, 2025. 5
- [13] Hanna Borgli, Vajira Thambawita, Pia H Smedsrud, Steven Hicks, Debesh Jha, Sigrun L Eskeland, Kristin Ranheim Randel, Konstantin Pogorelov, Mathias Lux, Duc Tien Dang Nguyen, et al. Hyperkvasir, a comprehensive multi-class image and video dataset for gastrointestinal endoscopy. *SData*, 7(1):283, 2020. 3, 10
- [14] Junying Chen, Chi Gui, Ruyi Ouyang, Anningzhe Gao, Shunian Chen, Guiming Hardy Chen, Xidong Wang, Ruifei Zhang, Zhenyang Cai, Ke Ji, et al. Towards injecting medical visual knowledge into multimodal LLMs at scale. In *EMNLP*, 2024. 5, 13
- [15] Wenting Chen, Yifan Liu, Jiancong Hu, and Yixuan Yuan. Dynamic depth-aware network for endoscopy super-resolution. *JBHI*, 26(10):5189–5200, 2022. 2
- [16] Xiaokang Chen, Zhiyu Wu, Xingchao Liu, Zizheng Pan, Wen Liu, Zhenda Xie, Xingkai Yu, and Chong Ruan. Januspro: Unified multimodal understanding and generation with data and model scaling. *arXiv preprint arXiv:2501.17811*, 2025. 5, 13
- [17] Zhe Chen, Weiyun Wang, Yue Cao, Yangzhou Liu, Zhangwei Gao, Erfei Cui, Jinguo Zhu, Shenglong Ye, Hao Tian, Zhaoyang Liu, et al. Expanding performance boundaries of open-source multimodal models with model, data, and test-time scaling. *arXiv preprint arXiv:2412.05271*, 2024. 5, 13
- [18] Ailin Deng, Tri Cao, Zhirui Chen, and Bryan Hooi. Words or vision: Do vision-language models have blind faith in text? In *IEEE CVPR*, pages 3867–3876, 2025. 7

- [19] Dujian Ding, Ankur Mallick, Chi Wang, Robert Sim, Subhabrata Mukherjee, Victor Ruhle, Laks VS Lakshmanan, and Ahmed Hassan Awadallah. Hybrid llm: Cost-efficient and quality-aware query routing. In *ICLR*, 2024. 5
- [20] Khiem Quang Do, Truc Thanh Thai, Viet Quoc Lam, and Thuy Thu Nguyen. Development and validation of artificial intelligence models for automated periodontitis staging and grading using panoramic radiographs. *BMC Oral Health*, 25 (1):1623, 2025. 8
- [21] Cathy Eng, Takayuki Yoshino, Erika Ruíz-García, Nermeen Mostafa, Christopher G Cann, Brittany O'Brian, Amala Benny, Rodrigo O Perez, and Chiara Cremolini. Colorectal cancer. *The Lancet*, 394(10207):1467–1480, 2024. 2
- [22] Deng-Ping Fan, Ge-Peng Ji, Tao Zhou, Geng Chen, Huazhu Fu, Jianbing Shen, and Ling Shao. Pranet: Parallel reverse attention network for polyp segmentation. In *MICCAI*, 2020. 3
- [23] Axel García-Vega, Ricardo Espinosa, Gilberto Ochoa-Ruiz, Thomas Bazin, Luis Falcón-Morales, Dominique Lamarque, and Christian Daul. A novel hybrid endoscopic dataset for evaluating machine learning-based photometric image enhancement models. In *MICAI*, 2022. 10
- [24] Sushant Gautam, Andrea Storås, Cise Midoglu, Steven A. Hicks, Vajira Thambawita, Pål Halvorsen, and Michael A. Riegler. Kvasir-vqa: A text-image pair gi tract dataset. In *ACM MM-W*, 2024. 3
- [25] Sushant Gautam, Michael A Riegler, and Pål Halvorsen. Kvasir-vqa-x1: A multimodal dataset for medical reasoning and robust medvqa in gastrointestinal endoscopy. *arXiv preprint arXiv:2506.09958*, 2025. 3
- [26] Mayank V Golhar, Lucas Sebastian Galeano Fretes, Loren Ayers, Venkata S Akshintala, Taylor L Bobrow, and Nicholas J Durr. C3vdv2—colonoscopy 3d video dataset with enhanced realism. *arXiv preprint arXiv:2506.24074*, 2025. 8
- [27] Daya Guo, Qihao Zhu, Dejian Yang, Zhenda Xie, Kai Dong, Wentao Zhang, Guanting Chen, Xiao Bi, Yu Wu, YK Li, et al. Deepseek-coder: When the large language model meets programming—the rise of code intelligence. *arXiv preprint arXiv:2401.14196*, 2024. 2, 7, 8
- [28] Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *Nature*, 645: 633–638, 2025. 2, 7
- [29] Palak Handa, Manas Dhir, Amirreza Mahbod, Florian Schwarzhans, Ramona Woitek, Nidhi Goel, and Deepak Gunjan. Wcebleedgen: A wireless capsule endoscopy dataset and its benchmarking for automatic bleeding classification, detection, and segmentation. *arXiv preprint arXiv:2408.12466*, 2024. 3, 10
- [30] Palak Handa, Amirreza Mahbod, Florian Schwarzhans, Ramona Woitek, Nidhi Goel, Deepti Chhabra, Shreshtha Jha, Manas Dhir, Deepak Gunjan, Jagadeesh Kakarla, et al. Capsule vision 2024 challenge: Multi-class abnormality classification for video capsule endoscopy. In *CVIP*, 2024. 10
- [31] Santa Hattori, Mineo Iwatate, Wataru Sano, Noriaki Hatsuiki, Hidekazu Kosaka, Taro Ikumoto, Masahito Kotaka, Akihiro Ichiyanagi, Chikara Ebisutani, Yasuko Hisano, et al. Narrow-band imaging observation of colorectal lesions using nice classification to avoid discarding significant lesions. *WJG*, 6(12):600, 2014. 4
- [32] Trung-Hieu Hoang, Hai-Dang Nguyen, Viet-Anh Nguyen, Thanh-An Nguyen, Vinh-Tiep Nguyen, and Minh-Triet Tran. Enhancing endoscopic image classification with symptom localization and data augmentation. In *ACMMM*, 2019. 3, 10
- [33] Zhijian Huang, Tao Tang, Shaoxiang Chen, Sihao Lin, Zequn Jie, Lin Ma, Guangrun Wang, and Xiaodan Liang. Making large language models better planners with reasoning-decision alignment. In *ECCV*, pages 73–90, 2024. 8
- [34] Participants in the Paris Workshop. The paris endoscopic classification of superficial neoplastic lesions: esophagus, stomach, and colon: November 30 to december 1, 2002. *GIE*, 58(6):S3–S43, 2003. 4
- [35] Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, Aiden Low, Alec Helyar, Aleksander Madry, Alex Beutel, Alex Carney, et al. Openai o1 system card. *arXiv preprint arXiv:2412.16720*, 2024. 2
- [36] Debesh Jha, Sharib Ali, Krister Emanuelsen, Steven A Hicks, Vajira Thambawita, Enrique Garcia-Ceja, Michael A Riegler, Thomas de Lange, Peter T Schmidt, Håvard D Johansen, et al. Kvasir-instrument: Diagnostic and therapeutic tool segmentation dataset in gastrointestinal endoscopy. In *MMM*, 2021. 3, 4, 10
- [37] Debesh Jha, Vanshali Sharma, Neethi Dasu, Nikhil Kumar Tomar, Steven Hicks, MK Bhuyan, Pradip K Das, Michael A Riegler, Pål Halvorsen, Thomas de Lange, et al. Gastrovision: A multi-class endoscopy image dataset for computer aided gastrointestinal disease detection. In *ICML-W*, 2023. 10
- [38] Debesh Jha, Nikhil Kumar Tomar, Vanshali Sharma, Quoc-Huy Trinh, Koushik Biswas, Hongyi Pan, Ritika K Jha, Gorkem Durak, Alexander Hann, Jonas Varkey, et al. Polypdb: A curated multi-center dataset for development of ai algorithms in colonoscopy. *arXiv preprint arXiv:2409.00045*, 2024. 3, 10
- [39] Ge-Peng Ji, Guobao Xiao, Yu-Cheng Chou, Deng-Ping Fan, Kai Zhao, Geng Chen, and Luc Van Gool. Video polyp segmentation: A deep learning perspective. *MIR*, 19(6):531–549, 2022. 3, 4, 10
- [40] Ge-Peng Ji, Jingyi Liu, Peng Xu, Nick Barnes, Fahad Shahbaz Khan, Salman Khan, and Deng-Ping Fan. Frontiers in intelligent colonoscopy. *MIR*, 2026. 2, 3, 4
- [41] Jared Kaplan, Sam McCandlish, Tom Henighan, Tom B Brown, Benjamin Chess, Rewon Child, Scott Gray, Alec Radford, Jeffrey Wu, and Dario Amodei. Scaling laws for neural language models. *arXiv preprint arXiv:2001.08361*, 2020. 8
- [42] Bidur Khanal, Sandesh Pokhrel, Sanjay Bhandari, Ramesh Rana, Nikesh Shrestha, Ram Bahadur Gurung, Cristian Linte, Angus Watson, Yash Raj Shrestha, and Binod Bhattarai. Hallucination-aware multimodal benchmark for gastrointestinal image analysis with large vision-language models. In *MICCAI*, 2025. 3

- [43] Yuna Kim, Ji-Soo Keum, Jie-Hyun Kim, Jaeyoung Chun, Sang-Il Oh, Kyung-Nam Kim, Young-Hoon Yoon, and Hyojin Park. Real-world colonoscopy video integration to improve artificial intelligence polyp detection performance and reduce manual annotation labor. *Diagnostics*, 15(7):901, 2025. 8
- [44] Anastasios Koulaouzidis, Dimitris K Iakovidis, Diana E Yung, Emanuele Rondonotti, Uri Kopylov, John N Plevris, Ervin Toth, Abraham Eliakim, Gabrielle Wurm Johansson, Wojciech Marlicz, et al. Kid project: an internet-based digital video atlas of capsule endoscopy for research purposes. *EIO*, 5(06):E477–E483, 2017. 3, 10
- [45] Edwin J Lai, Audrey H Calderwood, Gheorghe Doros, Oren K Fix, and Brian C Jacobson. The boston bowel preparation scale: a valid and reliable instrument for colonoscopy-oriented research. *Gastrointestinal endoscopy*, 69(3):620–625, 2009. 4, 9
- [46] Yuxiang Lai, Jike Zhong, Ming Li, Shitian Zhao, and Xiaofeng Yang. Med-r1: Reinforcement learning for generalizable medical reasoning in vision-language models. *arXiv preprint arXiv:2503.13939*, 2025. 7, 8, 20
- [47] Romain Leenhardt, Cynthia Li, Jean-Philippe Le Mouel, Gabriel Rahmi, Jean Christophe Saurin, Franck Cholet, Arnaud Boureille, Xavier Amiot, Michel Delvaux, Clotilde Duburque, et al. Cad-cap: a 25,000-image database serving the development of artificial intelligence for capsule endoscopy. *EIO*, 8(03):E415–E420, 2020. 3, 10
- [48] Samuel Lewis-Lim, Xingwei Tan, Zhixue Zhao, and Nikolaos Aletras. Can confidence estimates decide when chain-of-thought is necessary for llms? *arXiv preprint arXiv:2510.21007*, 2025. 5
- [49] Chunyuan Li, Cliff Wong, Sheng Zhang, Naoto Usuyama, Haotian Liu, Jianwei Yang, Tristan Naumann, Hoifung Poon, and Jianfeng Gao. Llava-med: Training a large language-and-vision assistant for biomedicine in one day. *NeurIPS*, 36:28541–28564, 2023. 4
- [50] Chunyuan Li, Cliff Wong, Sheng Zhang, Naoto Usuyama, Haotian Liu, Jianwei Yang, Tristan Naumann, Hoifung Poon, and Jianfeng Gao. Llava-med: Training a large language-and-vision assistant for biomedicine in one day. In *NeurIPS*, 2023. 5, 13
- [51] Kaidong Li, Mohammad I Fathan, Krushi Patel, Tianxiao Zhang, Cuncong Zhong, Ajay Bansal, Amit Rastogi, Jean S Wang, and Guanghui Wang. Colonoscopy polyp detection and classification: Dataset creation and comparative evaluations. *PONE*, 16(8):e0255809, 2021. 10
- [52] Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. Visual instruction tuning. In *NIPS*, pages 34892–34916, 2023. 4
- [53] Haotian Liu, Chunyuan Li, Yuheng Li, and Yong Jae Lee. Improved baselines with visual instruction tuning. In *IEEE CVPR*, 2024. 5, 13
- [54] Haotian Liu, Chunyuan Li, Yuheng Li, Bo Li, Yuanhan Zhang, Sheng Shen, and Yong Jae Lee. Llava-next: Improved reasoning, ocr, and world knowledge, 2024. 5, 13
- [55] Shengyuan Liu, Boyun Zheng, Wenting Chen, Zhihao Peng, Zhenfei Yin, Jing Shao, Jiancong Hu, and Yixuan Yuan. Endobench: A comprehensive evaluation of multi-modal large language models for endoscopy analysis. *NIPSDB*, 2025. 3
- [56] Yexin Liu, Zhengyang Liang, Yuezhe Wang, Xianfeng Wu, Feilong Tang, Muyang He, Jian Li, Zheng Liu, Harry Yang, Sernam Lim, et al. Unveiling the ignorance of mllms: Seeing clearly, answering incorrectly. In *CVPR*, pages 9087–9097, 2025. 7
- [57] Jun Ma, Yuting He, Feifei Li, Lin Han, Chenyu You, and Bo Wang. Segment anything in medical images. *Nature Communications*, 15(1):654, 2024. 4
- [58] Yiting Ma, Xuejin Chen, Kai Cheng, Yang Li, and Bin Sun. Ldpolypvideo benchmark: a large-scale colonoscopy video dataset of diverse polyps. In *MICCAI*, 2021. 10
- [59] Faisal Mahmood. A benchmarking crisis in biomedical machine learning. *Nature Medicine*, pages 1–1, 2025. 3
- [60] Shawn Mathew, Saad Nadeem, and Arie Kaufman. Clts-gan: color-lighting-texture-specular reflection augmentation for colonoscopy. In *MICCAI*, pages 519–529, 2022. 2
- [61] Pablo Mesejo, Daniel Pizarro, Armand Abergel, Olivier Rouquette, Sylvain Beorchia, Laurent Poincloux, and Adrien Bartoli. Computer-aided classification of gastrointestinal lesions in regular colonoscopy. *IEEE TMI*, 35(9):2051–2063, 2016. 3, 10
- [62] Masashi Misawa, Shin-ei Kudo, Yuichi Mori, Kinichi Hotta, Kazuo Ohtsuka, Takahisa Matsuda, Shoichi Saito, Toyoki Kudo, Toshiyuki Baba, Fumio Ishida, et al. Development of a computer-aided detection system for colonoscopy and a publicly accessible large colonoscopy video database (with video). *GIE*, 93(4):960–967, 2021. 3
- [63] Francis Jesmar P Montalbo. Diagnosing gastrointestinal diseases from endoscopy images through a multi-fused cnn with auxiliary layers, alpha dropouts, and a fusion residual block. *BSPC*, 76:103683, 2022. 10
- [64] Michael Moor, Oishi Banerjee, Zahra Shakeri Hossein Abad, Harlan M Krumholz, Jure Leskovec, Eric J Topol, and Pranav Rajpurkar. Foundation models for generalist medical artificial intelligence. *Nature*, 616(7956):259–265, 2023. 8
- [65] Dung Nguyen, Minh Khoi Ho, Huy Ta, Thanh Tam Nguyen, Qi Chen, Kumar Rav, Quy Duong Dang, Satwik Ramchandre, Son Lam Phung, Zhibin Liao, et al. Localizing before answering: A benchmark for grounded medical visual question answering. In *IJCAI*, 2025. 7
- [66] Chuang Niu, Qing Lyu, Christopher D Carothers, Parisa Kaviani, Josh Tan, Pingkun Yan, Mannudeep K Kalra, Christopher T Whitlow, and Ge Wang. Medical multimodal multitask foundation model for lung cancer screening. *Nature Communications*, 16(1):1523, 2025. 4
- [67] OpenAI. o3 and o4-mini system card. <https://openai.com/index/o3-o4-mini-system-card/>, 2025. 2
- [68] Perry J Pickhardt, Kendra S Hain, David H Kim, and Cesare Hassan. Low rates of cancer or high-grade dysplasia in colorectal polyps collected from computed tomography colonography screening. *Clinical Gastroenterology and Hepatology*, 8(7):610–615, 2010. 4
- [69] Konstantin Pogorelov, Kristin Ranheim Randel, Thomas de Lange, Sigrun Losada Eskeland, Carsten Griwodz, Dag Johansen, Concetto Spampinato, Mario Taschwer, Mathias

- Lux, Peter Thelin Schmidt, et al. Nerthus: A bowel preparation quality video dataset. In *ACM MMSys*, 2017. 3, 10
- [70] Konstantin Pogorelov, Kristin Ranheim Randel, Carsten Griwodz, Sigrun Losada Eskeland, Thomas de Lange, Dag Johansen, Concetto Spampinato, Duc-Tien Dang-Nguyen, Mathias Lux, Peter Thelin Schmidt, et al. Kvasir: A multi-class image dataset for computer aided gastrointestinal disease detection. In *ACM MMSys*, 2017. 3, 10
- [71] Gorkem Polat, Haluk Tarik Kani, Ilkay Ergenc, Yesim Ozen Alahdab, Alptekin Temizel, and Ozlen Atug. Improving the computer-aided estimation of ulcerative colitis severity according to mayo endoscopic score by using regression-based deep learning. *IBD*, 29(9):1431–1439, 2023. 10
- [72] Fabiane Queiroz and Tsang Ing Ren. Endoscopy image restoration: A study of the kernel estimation from specular highlights. *Digital Signal Processing*, 88:53–65, 2019. 2
- [73] F Javier Sánchez, Jorge Bernal, Cristina Sánchez-Montes, Cristina Rodríguez de Miguel, and Gloria Fernández-Esparrach. Bright spot regions segmentation and classification for specular highlights detection in colonoscopy videos. *Machine Vision and Applications*, 28(8):917–936, 2017. 2
- [74] Luisa F Sánchez-Peralta, J Blas Pagador, Artzai Picón, Ángel José Calderón, Francisco Polo, Nagore Andraka, Roberto Bilbao, Ben Glover, Cristina L Saratxaga, and Francisco M Sánchez-Margallo. Piccolo white-light and narrow-band imaging colonoscopic dataset: a performance comparative of models and datasets. *ApplSci*, 10(23):8501, 2020. 10
- [75] Kenneth W Schroeder, William J Tremaine, and Duane M Ilstrup. Coated oral 5-aminosalicylic acid therapy for mildly to moderately active ulcerative colitis. *New England Journal of Medicine*, 317(26):1625–1629, 1987. 4
- [76] Andrew Selligren, Sahar Kazemzadeh, Tiam Jaroensri, Atilla Kiraly, Madeleine Traverse, Timo Kohlberger, Shawn Xu, Fayaz Jamil, Cian Hughes, Charles Lau, et al. Medgemma technical report. *arXiv preprint arXiv:2507.05201*, 2025. 5, 13
- [77] Chufan Shi, Yixuan Su, Cheng Yang, Yujiu Yang, and Deng Cai. Specialist or generalist? instruction tuning for specific nlp tasks. In *EMNLP*, 2023. 5
- [78] Juan Silva, Aymeric Histace, Olivier Romain, Xavier Dray, and Bertrand Granado. Toward embedded detection of polyps in wce images for early diagnosis of colorectal cancer. *CARS*, 9(2):283–293, 2014. 10
- [79] Pia H Smedsrud, Vajira Thambawita, Steven A Hicks, Henrik Gjestang, Oda Olsen Nedrejord, Espen Næss, Hanna Borgli, Debesh Jha, Tor Jan Derek Berstad, Sigrun L Eskeland, et al. Kvasir-capsule, a video capsule endoscopy dataset. *SData*, 8(1):142, 2021. 3, 10
- [80] Nima Tajbakhsh, Suryakanth R Gurudu, and Jianming Liang. Automated polyp detection in colonoscopy videos using shape and context information. *TMI*, 35(2):630–644, 2015. 2
- [81] Michael B Wallace, Prateek Sharma, Pradeep Bhandari, James East, Giulio Antonelli, Roberto Lorenzetti, Micheal Vieth, Ilaria Speranza, Marco Spadaccini, Madhav Desai, et al. Impact of artificial intelligence on miss rate of colorectal neoplasia. *Gastro*, 163(1):295–304, 2022. 2
- [82] Dequan Wang, Xiaosong Wang, Lilong Wang, Mengzhang Li, Qian Da, Xiaoqiang Liu, Xiangyu Gao, Jun Shen, Junjun He, Tian Shen, et al. A real-world dataset and benchmark for foundation model adaptation in medical image classification. *SData*, 10(1):574, 2023. 3, 10
- [83] Qin Wang, Hui Che, Weizhen Ding, Li Xiang, Guanbin Li, Zhen Li, and Shuguang Cui. Colorectal polyp classification from white-light colonoscopy images via domain alignment. In *MICCAI*, 2021. 10
- [84] Wei Wang, Jinge Tian, Chengwen Zhang, Yanhong Luo, Xin Wang, and Ji Li. An improved deep learning approach and its applications on colonic polyp images detection. *BMCMI*, 20:1–14, 2020. 10
- [85] Xiaohan Xu, Ming Li, Chongyang Tao, Tao Shen, Reynold Cheng, Jinyang Li, Can Xu, Dacheng Tao, and Tianyi Zhou. A survey on knowledge distillation of large language models. *arXiv preprint arXiv:2402.13116*, 2024. 8
- [86] Ziang Xu, Jens Rittscher, and Sharib Ali. Ssl-cpcd: Self-supervised learning with composite pretext-class discrimination for improved generalisability in endoscopic image analysis. *TMI*, 2024. 10
- [87] An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, et al. Qwen3 technical report. *arXiv preprint arXiv:2505.09388*, 2025. 2
- [88] Menglong Ye, Stamatia Giannarou, Alexander Meining, and Guang-Zhong Yang. Online tracking and retargeting with applications to optical biopsy in gastrointestinal endoscopic examinations. *MedIA*, 30:144–157, 2016. 3, 10
- [89] Guanghui Yue, Di Cheng, Tianwei Zhou, Jingwen Hou, Weide Liu, Long Xu, Tianfu Wang, and Jun Cheng. Perceptual quality assessment of enhanced colonoscopy images: A benchmark dataset and an objective method. *IEEE Transactions on Circuits and Systems for Video Technology*, 33(10):5549–5561, 2023. 2
- [90] Guanghui Yue, Guibin Zhuo, Siying Li, Tianwei Zhou, Jingfeng Du, Weiqing Yan, Jingwen Hou, Weide Liu, and Tianfu Wang. Benchmarking polyp segmentation methods in narrow-band imaging colonoscopy images. *IEEE JBHI*, 27(7):3360–3371, 2023. 10
- [91] Dylan Zhang, Justin Wang, and Francois Charton. Only-if: revealing the decisive effect of instruction diversity on generalization. *Preprint*, 2024. 5
- [92] Weike Zhao, Chaoyi Wu, Yanjie Fan, Xiaoman Zhang, Pengcheng Qiu, Yuze Sun, Xiao Zhou, Yanfeng Wang, Xin Sun, Ya Zhang, et al. An agentic system for rare disease diagnosis with traceable reasoning. *arXiv preprint arXiv:2506.20430*, 2025. 8
- [93] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanhao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, et al. Judging llm-as-a-judge with mt-bench and chatbot arena. In *NIPS*, 2023. 15
- [94] Jinguo Zhu, Weiyun Wang, Zhe Chen, Zhaoyang Liu, Shenglong Ye, Lixin Gu, Hao Tian, Yuchen Duan, Weijie Su, Jie Shao, et al. Internvl3: Exploring advanced training and test-time recipes for open-source multimodal models. *arXiv preprint arXiv:2504.10479*, 2025. 5, 13

- [95] Shaofeng Zou, Mingzhu Long, Xuyang Wang, Xiang Xie, Guolin Li, and Zhihua Wang. A cnn-based blind denoising method for endoscopic images. In *BioCAS*, pages 1–4, 2019.

[2](#)