

ClinConsensus: A Physician-Calibrated Benchmark for Evaluating Clinical Rubric Coverage in Chinese Medical LLMs

Xiang Zheng* Han Li* Wenjie Luo* Weiqi Zhai*
 Yiyuan Li Chuanmiao Yan Xue Yang Kailuan Wu Ruyi Xu
 Tianyun Lu Tianyi Tang Yubo Ma Kexin Yang Sen Yang
 Lin Qu Dayiheng Liu[†] Bing Zhao[†] Hu Wei[†]
 Alibaba Group, China

*Equal contribution. [†]Corresponding authors.

Abstract

Open-ended medical LLM evaluation remains weakly grounded in physician-calibrated coverage of clinically relevant response criteria, especially in localized clinical settings. We introduce CLIN-CONSENSUS, a Chinese medical benchmark of 2,500 expert-curated cases spanning 36 specialties, 12 task themes, multiple difficulty levels, and lay-facing versus professional-facing settings. Each case is paired with 30 case-specific binary rubric criteria. To evaluate whether responses satisfy enough physician-authored criteria, we propose *Clinician-Anchored Coverage Score* (CACS), a physician-calibrated threshold metric instantiated at $k = 10$, and develop a dual-judge framework combining a GPT-5.1 grader with a physician-supervised Qwen3-8B judge. Evaluating 11 frontier LLMs, we find a persistent coverage gap: Rubric Accuracy ranges from 39.6% to 52.1%, whereas CACS@10 ranges from 17.8% to 32.9%, leaving a 19.2–21.9 point gap across models. Stratified analyses further reveal substantial variation across reasoning, evidence use, structured extraction, medication instructions, follow-up, and dialogue register. These results suggest that medical LLM evaluation should measure thresholded, rubric-grounded clinical coverage rather than average partial correctness.

1 Introduction

Large language models (LLMs) are increasingly used for health-related tasks that go beyond isolated clinical question answering, including patient education, diagnostic reasoning, treatment planning, documentation, and follow-up (Singhal et al., 2025; Tu et al., 2025; Wang et al., 2025b; Al-Garadi et al., 2025; Tu et al., 2024). Evaluating such systems re-

quires more than factual recall: clinically useful responses must interpret patient context, handle risk and resource constraints, communicate actionable next steps, and remain calibrated under uncertainty (Mukherjee et al., 2024; Mehandru et al., 2023; Wang et al., 2025b; Al-Garadi et al., 2025).

Existing evaluations have made rapid progress, showing strong LLM performance on clinical question answering, triage, documentation, medication safety, and patient education (Tordjman et al., 2025; Sandmann et al., 2024; Bedi et al., 2025; Van Veen et al., 2024). However, many benchmarks remain task-isolated or weakly contextualized, with single-round exam-style questions that underrepresent multi-step reasoning, treatment trade-offs, conflict resolution, and local constraints such as resources, culture, and access to care (Bedi et al., 2025; Thirunavukarasu et al., 2023; Arora et al., 2025; Eriksen et al., 2024). Recent benchmarks introduce open-ended responses, rubric or checklist evaluation, physician validation, interactive settings, and Chinese-localized resources (Arora et al., 2025; Zhang et al., 2025; Schmidgall et al., 2024; Jiang et al., 2025; He et al., 2026; Ding et al., 2025), but they leave a design gap: checklist benchmarks often emphasize average criterion accuracy, interactive benchmarks stress sequential behavior, and Chinese evaluation suites prioritize local task breadth. Fewer benchmarks ask whether an open-ended answer covers enough case-specific clinical requirements to cross a physician-calibrated coverage threshold.

We introduce CLIN-CONSENSUS, a Chinese medical benchmark grounded in local clinical practice. It contains 2,500 open-ended cases spanning 36 medical specialties and 12 clinical task types, with metadata for task, specialty,

Table 1: **Positioning ClinConsensus among medical LLM benchmarks.** ✓, ◦, and – denote explicit, partial, and absent support. Physician validation includes authorship, review, or grading by physicians.

Benchmark	Chinese / localized	Open-ended	Multi-specialty	Rubric / checklist	Physician validation	Interactive / multi-turn	Difficulty tiers
HealthBench (Arora et al., 2025)	–	✓	✓	✓	✓	✓	–
LLMEval-Med (Zhang et al., 2025)	✓	✓	✓	✓	✓	–	–
CSEDB (Wang et al., 2025a)	✓	✓	✓	✓	✓	–	–
MedThink-Bench (Zhou et al., 2025)	–	✓	✓	◦	✓	–	–
AgentClinic (Schmidgall et al., 2024)	–	✓	✓	–	–	✓	–
ClinicalLab (Yan et al., 2024)	–	✓	✓	–	◦	✓	–
MedBench (Jiang et al., 2025)	✓	–	✓	–	–	–	–
MLB (He et al., 2026)	✓	◦	✓	◦	✓	–	–
MedBench v4 (Ding et al., 2025)	✓	◦	✓	◦	✓	✓	–
ClinConsensus	✓	✓	✓	✓	✓	–	✓

difficulty, and intended user role. Each case includes an expert-reviewed reference answer and 30 structured, case-specific binary rubric criteria, enabling qualitative inspection and thresholded clinical-coverage evaluation rather than only average factual correctness. Its core distinction is the combination of localized Chinese clinical cases, case-specific scoring criteria, and a thresholded rubric-coverage metric calibrated from physician-written responses. Figure 1 summarizes the curation, validation, and evaluation pipeline.

Our contributions are:

1. **A Chinese expert-curated benchmark for open-ended medical evaluation.** CLINCONSENSUS contains 2,500 open-ended clinical cases covering 36 specialties and 12 task types, with expert construction or review and multidimensional case-specific rubrics.
2. **A scalable rubric-based evaluation framework for complex clinical tasks.** We propose the physician-calibrated *Clinician-Anchored Coverage Score* (CACS) and a dual-judge framework combining a GPT-5.1 grader with a physician-supervised Qwen3-8B judge.
3. **A comprehensive analysis of frontier medical LLMs on localized real-world cases.** We evaluate 11 frontier LLMs and show that similar aggregate scores can mask differences in reasoning, evidence use, structured extraction, follow-up, and treatment-planning capability.

2 Related Work

Medical LLM evaluation has moved beyond static exam-style datasets such as MedQA, MedMCQA, and MMLU-style medical subsets, which mainly measure factual recall and broad biomedical knowledge (Jin et al., 2021; Pal et al., 2022; Hendrycks et al., 2020). Recent benchmarks improve realism through open-ended responses, rubric or checklist grading, and physician validation: HealthBench uses physician-written rubrics for realistic health conversations (Arora et al., 2025); LLMEval-Med and CSEDB use checklist-style grading for real-world or scenario-based clinical problems (Zhang et al., 2025; Wang et al., 2025a); and MedThink-Bench targets reasoning-level assessment with expert-crafted rationales (Zhou et al., 2025). These benchmarks improve grading fidelity, but they mainly answer whether individual rubric items are satisfied on average. They do not directly test whether a response reaches a minimum case-level coverage point calibrated from physician-authored answers in a localized Chinese clinical setting.

Table 1 summarizes the main design differences among representative medical LLM benchmarks.

Complementary work evaluates LLMs in interactive or agentic clinical environments, where success depends on information gathering, sequential decision-making, tool use, and plan consistency. AIE/SAPS, AgentClinic, and ClinicalLab/ClinicalBench expose failures that static QA may miss, including incomplete history taking and brittle plan tracking (Liao et al., 2024; Schmidgall et al., 2024; Yan et al.,

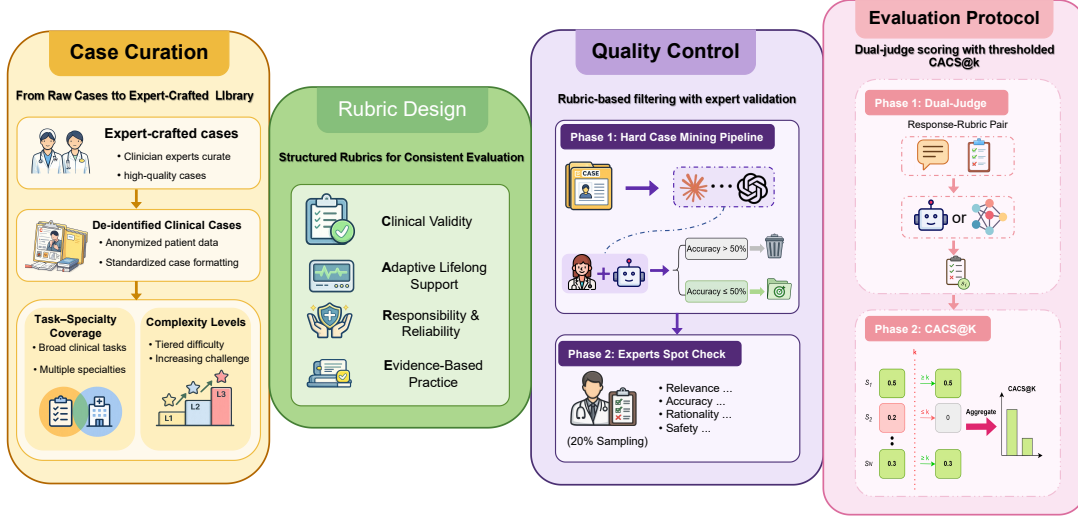


Figure 1: **ClinConsensus pipeline.** Expert-curated clinical cases are transformed into structured benchmark items, paired with case-specific rubrics, curated and audited through quality control, and evaluated under a rubric-level protocol for CACS@ k scoring and judge validation.

2024). Chinese-localized benchmarks have also expanded rapidly: MedBench provides standardized Chinese medical evaluation infrastructure, MLB introduces scenario-driven assessment over Chinese clinical sources, and MedBench v4 extends this line to language models, multimodal models, and agents (Jiang et al., 2025; He et al., 2026; Ding et al., 2025). CLINCONSENSUS is complementary to these efforts. It focuses on expert-curated Chinese open-ended cases with case-specific rubric grading, physician-calibrated thresholded scoring, difficulty tiers, and lay-facing versus professional-facing settings within a single benchmark.

3 Approach

This section describes the dataset-construction stages in Figure 1: case curation, metadata stratification, rubric design, and quality control. The evaluation protocol is described in Section 4.

3.1 Case Curation

We invited a multidisciplinary team of clinical experts in China to construct CLINCONSENSUS from professional experience and daily clinical practice. Experts either wrote original scenarios or transformed de-identified real-world cases into natural-language narratives. The final benchmark contains 2,500 open-ended cases, ranging from concise con-

sultations to longer clinical narratives, and covers both lay-facing health questions and professional-facing decision-support settings.

3.2 Case Metadata and Stratification

Each case is annotated along four axes used for stratified evaluation: clinical task, medical specialty, case difficulty, and intended dialogue register.

Task themes. We normalize atomic task labels into 12 themes spanning chart understanding, information extraction, NER, summarization, differential diagnosis, causal/logical reasoning, evidence retrieval, treatment planning, patient education, test/lab interpretation, medication instructions, and follow-up/monitoring. Task labels are multi-label: a case contributes to every theme with which it is annotated.

Medical specialties. After alias normalization, the benchmark covers 36 medical specialties. Because many cases span multiple specialties, specialty-wise evaluation treats specialty labels as multi-label annotations. For stable main-text comparisons, we also map the first specialty label of each case into eight subject macro-categories, avoiding sparse estimates from the 36 raw labels.

Dialogue register. Professional-facing cases are written for physicians, nurses, residents, interns, medical educators, document reviewers, or other healthcare professionals;

Table 2: **Stratified evaluation axes.** Task and specialty labels are multi-label; difficulty and dialogue register are mutually exclusive.

Axis	Strata	Purpose
Task theme	12 clinical task themes	Identifies capability-specific strengths and failures
Medical specialty	8 macro-categories	Tests domain generalization and specialty-level blind spots
Difficulty	Low, Medium, High	Measures robustness across curated difficulty strata
Dialogue register	Lay-facing, professional-facing	Compares responses for lay and professional users

lay-facing cases are written for patients, family members, caregivers, friends, or general health-assistant settings. The current metadata separate 1,033 professional-facing cases from 1,467 lay-facing cases.

Case difficulty. Each retained case carries a curated difficulty label with three mutually exclusive strata: Low, Medium, and High (Appendix Table 6). The released strata contain 900 low-, 800 medium-, and 800 high-difficulty cases. We treat difficulty as descriptive dataset metadata for stratified evaluation, not as a separate scoring rule or a per-model decision rule. To make this metadata auditable, we characterize the released labels by their task and subject profiles (Appendix Table 7). Low-difficulty cases are enriched for patient education, information extraction, summarization, and test/lab interpretation. High-difficulty cases are enriched for medication instructions, chart understanding, evidence retrieval, causal/logical reasoning, and personalized treatment planning, with medium cases occupying mixed reasoning and planning settings. The same CACS@10 threshold is applied across strata so that difficulty analyses test robustness under a fixed operating point rather than per-stratum retuning.

3.3 Rubric Design

Each case is paired with 30 expert-defined binary criteria covering key clinical information, decision points, and safety considerations. Rubrics are tailored to the case and aligned with its specialties, task themes, difficulty, and dialogue register. They combine

domain-specific criteria, authored by physicians with reference to the scenario, guidelines, and local practice, with consensus criteria covering safety, personalization, evidence use, communication quality, fairness, and localized regulatory compliance.

The 73-item consensus-criteria bank (Appendix C) is informed by existing medical evaluation metrics, expert interviews, and prior literature (Castelo-Branco et al., 2023; Freyer et al., 2024). Appendix D gives a traceability example showing how selected case-specific criteria map to guideline-aligned expectations.

3.4 Quality Control

Case creation and annotation were completed by at least two licensed physicians. We then applied two quality-control stages.

Stage 1: Model-assisted candidate screening. Experts reviewed responses generated by DeepSeek V3, GPT-5, and Gemini 2.5 Pro and manually graded each case against the 30 criteria. Cases with aggregated rubric scores at or above 50% were excluded from the retained evaluation set. The same retained set was used for all model comparisons.

Stage 2: Expert auditing and consistency review. Senior clinicians randomly audited approximately 20% of the remaining cases, checking case descriptions, rubrics, and reference answers for clinical correctness, necessity, internal consistency, and alignment with Chinese and international guidelines. Cases failing review were revised or removed before final inclusion.

4 Judge Framework and Metrics

CLINCONSENSUS evaluates open-ended responses at the rubric level rather than with a single overall score. For each case i , a candidate response is judged against $N = 30$ case-specific binary criteria. For each criterion r_j , the grader receives the clinical context, the candidate response, and one rubric item, and returns a schema-constrained `criteria_met` $\in \{\text{true}, \text{false}\}$ decision with a short rationale. This separates the response generator, clinical rubric, and grader, making each decision auditable and comparable across models. Malformed or unparsable outputs are treated conservatively as unmet criteria.

4.1 Rubric-level Graders

We use two graders with the same input–output schema. The **GPT-5.1 strict judge** is the primary external grader for benchmark scoring unless otherwise specified. It evaluates each rubric item independently with a fixed JSON-only prompt and is used only after expert-written rubrics are finalized, not for rubric creation or threshold calibration. We use this stricter external grader as the primary scorer for two reasons: it is not trained on CLINCONSENSUS physician labels, avoiding primary evaluation with a benchmark-specialized judge, and its lower positive rate makes CACS@10 a conservative estimate of rubric coverage. The physician-supervised Qwen3-8B judge is therefore used as an alignment, robustness, and local deployment route rather than as the sole primary scorer.

The **physician-supervised judge** is a Qwen3-8B model trained on physician-labeled rubric decisions rather than GPT-5.1 labels. To test cross-model generalization, we train and tune on 12 source models and hold out 11 unseen models for judge-fidelity evaluation, yielding 82,500 held-out rubric decisions over 250 cases and 30 criteria per response. The SFT split contains 52,209 training examples and 5,933 development examples, with a 47.1% positive-label rate. Full prompts, parsing, and training details are provided in Appendices J and I.3.

4.2 CACS@ k

Average rubric accuracy can reward scattered partial correctness, including responses that satisfy isolated criteria while omitting key safety checks, decision points, or follow-up actions. We therefore introduce **Clinician-Anchored Coverage Score** (CACS@ k), a thresholded metric for explicit rubric coverage. Rubric Accuracy measures average criterion-level correctness; CACS measures coverage after a physician-calibrated threshold is reached.

For each case i , the rubric-hit score is

$$s_i = \sum_{j=1}^N y_{i,j} \in \{0, 1, \dots, N\},$$

where $y_{i,j}$ is the binary decision for criterion j . Let $\hat{P}(s \geq t)$ denote the empirical survival

function of rubric-hit scores over evaluated responses $\mathcal{D}$. We define

$$\text{CACS@}k = \frac{100}{N - k + 1} \sum_{t=k}^N \hat{P}(s \geq t). \quad (1)$$

Equivalently,

$$\begin{aligned} \text{CACS@}k &= \frac{100}{|\mathcal{D}|(N - k + 1)} \sum_{i \in \mathcal{D}} c_i, \\ c_i &= \max(0, s_i - k + 1). \end{aligned} \quad (2)$$

CACS differs from both mean accuracy and binary pass rate. Mean accuracy rewards every isolated rubric hit, even below the clinical coverage threshold. A binary pass rate enforces a threshold but saturates once $s_i \geq k$. CACS assigns zero credit below k and graded credit above k . Appendix Table 9 illustrates this behavior for $N = 30$ and $k = 10$. CACS is a benchmark coverage metric for thresholded rubric coverage, not a clinical deployment pass/fail rule.

4.3 Calibrating the Threshold

We calibrate k from physician-authored responses, not automated grader outputs. We randomly sample 285 cases, ask attending physicians from Class A tertiary hospitals to write responses, and have an independent physician group evaluate those responses against the same case-specific rubrics. This avoids circular dependence between GPT-5.1 grading and threshold selection.

Let $\{s_i^{\text{phys}}\}$ denote the physician-response rubric-hit scores. We set k to their rounded empirical mean, yielding $k = 10$, and report CACS@10 as the main metric. The calibration responses satisfy 2,904 of 8,550 rubric decisions, corresponding to 33.96% mean rubric-hit rate or 10.19/30 criteria per case (median 9, IQR 6–14). Thus, $k = 10$ is a physician-explicit coverage anchor, not a universal clinical safety threshold. We keep a single k across low-, medium-, and high-difficulty strata so that model scores remain comparable; the difficulty analysis therefore tests robustness under the same operating point rather than re-tuning the metric per stratum. Appendix Table 10 reports sensitivity to nearby thresholds.

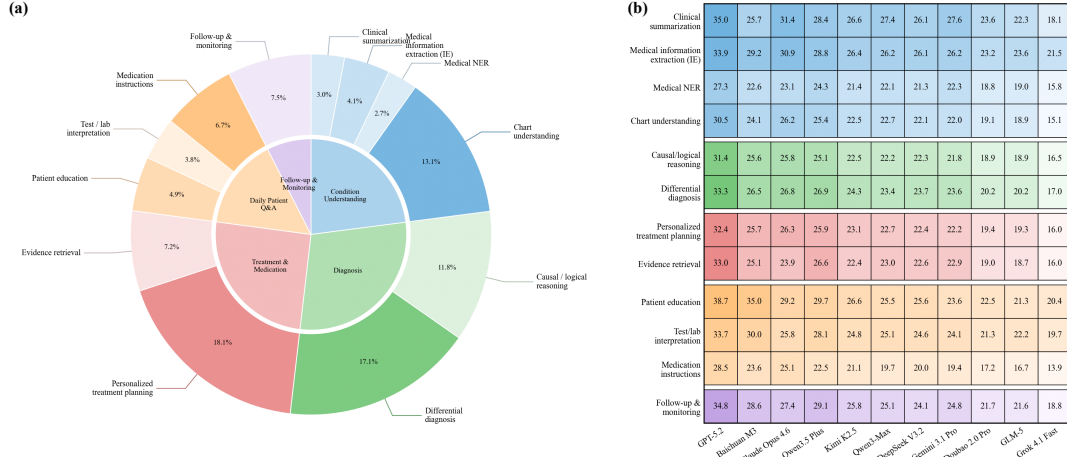


Figure 2: **Task distribution and theme-wise CACS@10.** (a) Case distribution across 12 task themes. (b) Model performance by task theme.

5 Experiments

Evaluated models. We evaluate 11 frontier LLMs through official APIs: *GPT-5.2* with high reasoning effort, *Gemini 3.1 Pro*, *Claude Opus 4.6*, *Qwen3-Max*, *Qwen3.5 Plus*, *Kimi K2.5*, *DeepSeek V3.2*, *Doubao 2.0 Pro*, *Grok 4.1 Fast*, *GLM-5*, and *Baichuan M3*. Main text uses standardized public-facing names; dated deployment identifiers, provider routes, prompts, and decoding settings are in the reproducibility package (Appendix A). *Qwen3-Max* denotes the `qwen3-max-2026-01-23` deployment.

Evaluation slices. We report aggregate and stratified scores across 12 task themes, eight subject macro-categories, three difficulty levels, and lay-facing versus professional-facing register (Section 3.2). Task themes are multi-label; subject analyses use the first specialty label; difficulty and register are mutually exclusive case-level strata.

5.1 Main Results

CACS@10 changes the interpretation of aggregate performance. Appendix Figure 10 ranks *GPT-5.2* highest (32.9%), followed by *Baichuan M3* (27.8%) and *Claude Opus 4.6* (26.9%), but this ordering is benchmark-specific. The central finding is the gap between partial credit and thresholded coverage: Rubric Accuracy ranges from 39.6% to 52.1%, whereas CACS@10 ranges from 17.8% to 32.9%, leaving a 19.2–21.9 point gap. For

GPT-5.2, 52.1% Rubric Accuracy becomes 32.9% CACS@10; for *Grok 4.1 Fast*, 39.6% becomes 17.8%. Many responses collect isolated rubric hits without reaching the physician-calibrated 10/30 threshold. Appendix Figure 7 visualizes this case-level mechanism, and Appendix Table 11 checks that CACS@10 is not simply a proxy for verbosity.

Task-level results localize this gap. Across 12 themes and five aggregated capability families (Figure 2, Appendix Figure 8), patient communication is strongest on average (26.1%), while treatment and medication (22.7%) and clinical reasoning (23.1%) are persistent bottlenecks. *GPT-5.2* leads every capability family, but the next tier is capability-dependent: *Claude Opus 4.6* and *Qwen3.5 Plus* outperform *Baichuan M3* on structured extraction (28.6% and 27.2% vs. 26.0%) despite lower aggregate CACS@10. Thus, similar aggregate scores can hide different operational weaknesses.

Domain and difficulty stratification show wider gaps in harder settings. Subject macro-category means span 27.2%–21.3% (Appendix Figure 9), with *GPT-5.2* ranging from 35.4% to 28.4% in *Neurology & Psychiatry*. Difficulty is clearer: mean CACS@10 declines monotonically from 28.5% on low-difficulty cases to 24.8% on medium and 18.0% on high (Figure 5). Medium–high penalties are 4.3–9.0 points, and low–high penalties are 7.9–15.3 points. This matches the metadata profile: high-difficulty cases are enriched for

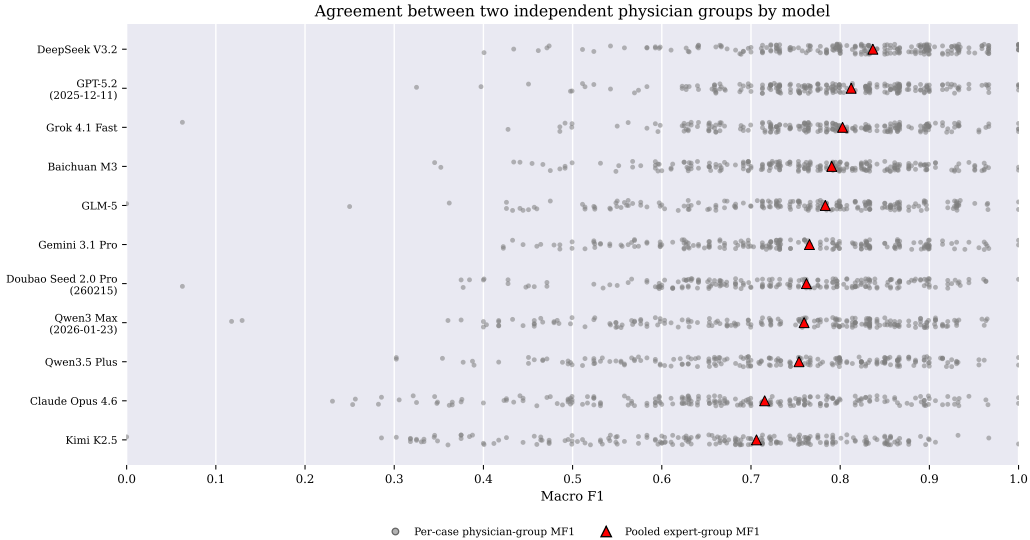


Figure 3: **Inter-expert group agreement across 11 models.** Case-level MF1 distributions and pooled MF1 summarize agreement between two independent expert groups after model-case alignment.

medication instructions, chart understanding, evidence retrieval, causal/logical reasoning, and personalized treatment planning, where thresholded coverage is most fragile.

Fine-grained scenarios identify practically important stress points. Structured clinical-documentation extraction is highest (36.5%), whereas imaging/pathology interpretation (15.5%), medication-safety interactions or contraindications (17.8%), and evidence retrieval or guideline alignment (18.3%) are the clearest bottlenecks. These weaknesses remain visible for *GPT-5.2*, which scores 19.4% on imaging/pathology interpretation and 26.0% on medication-safety interactions, compared with 42.0% on structured extraction and 45.2% on health education. Dialogue register adds a smaller diagnostic layer (24.4% vs. 23.9% CACS@10; Appendix Table 12), with model-level directions varying. A de-identified low-CACS audit links these bottlenecks to missing emergency triggers, non-operational treatment plans, medication-boundary omissions, and incomplete follow-up (Appendix Table 13). Together, these analyses make CLINCONSENSUS diagnostic beyond ranking without treating CACS@10 as a real-world clinical pass/fail rule.

6 Evaluation Reliability

CLINCONSENSUS scales open-ended medical evaluation through automated rubric grad-

ing. We validate the protocol at three levels: physician-physician reproducibility, automated judge-physician fidelity, and model-level robustness under graders with different strictness.

6.1 Physician Agreement

Meta-evaluation set. We compare two independent physician annotation cohorts. Each contains 5,750 rows, covering 23 model runs over 250 cases with up to 30 rubric decisions per response. After aligning by model identity and normalized case content, we retain 11 shared models, yielding 2,750 matched model-case rows and 82,500 aligned rubric decisions per expert group. Expert Group 1 includes 84 annotators and 15 secondary checkers; Expert Group 2 includes 45 annotators and 9 secondary checkers. Parsing and alignment details appear in Appendix I.3.

Metric. Following HealthBench (Arora et al., 2025), we treat rubric grading as binary classification over `met` and `not.met` decisions and report Macro-F1 (MF1), the unweighted average of the two class-wise F1 scores. MF1 gives equal weight to satisfied and unsatisfied criteria, making it less sensitive to label imbalance than raw agreement.

Results. Table 3 and Figure 3 show substantial but non-uniform physician agreement. Across the 11 models, pooled MF1 ranges from

Table 3: **Inter-expert agreement by model.** Counts are Expert Group 1 / Expert Group 2; pooled MF1 is computed after model–case alignment.

Model	Annotators	Checkers	Pooled MF1
DeepSeek V3.2	16 / 31	1 / 4	0.837
GPT-5.2	8 / 21	1 / 4	0.812
Grok 4.1 Fast	9 / 25	1 / 3	0.803
Baichuan M3	15 / 26	7 / 5	0.791
GLM-5	10 / 5	5 / 1	0.783
Gemini 3.1 Pro	10 / 4	1 / 1	0.766
Doubao 2.0 Pro	10 / 5	2 / 1	0.762
Qwen3-Max	11 / 12	7 / 1	0.760
Qwen3.5 Plus	17 / 7	5 / 3	0.754
Claude Opus 4.6	16 / 9	1 / 1	0.715
Kimi K2.5	12 / 5	1 / 1	0.706

Table 4: **Held-out automated judge fidelity.** MF1 is computed against each expert group on 82,500 held-out rubric decisions.

Judge	Expert Group 1 MF1	Expert Group 2 MF1	Average MF1
GPT-5.1 strict judge	0.779	0.738	0.759
Base Qwen3-8B judge	0.782	0.771	0.777
Our judge model	0.844	0.808	0.826

0.706 to 0.837 (mean 0.772), with the highest agreement for *DeepSeek V3.2* (0.837), *GPT-5.2* (0.812), and *Grok 4.1 Fast* (0.803), and the lowest for *Kimi K2.5* (0.706) and *Claude Opus 4.6* (0.715). Case-level distributions show that lower pooled MF1 comes from heavier low-agreement tails: only 5.3–6.5% of matched cases from the three highest-MF1 models fall below case-level MF1 0.6, compared with 28.4% for *Kimi K2.5* and 28.6% for *Claude Opus 4.6*. These tails often involve rubric-boundary responses that are broadly plausible but leave emergency thresholds, medication constraints, follow-up plans, or diagnostic decision points implicit. Human–human agreement therefore provides a realistic reference point for automated judge fidelity.

6.2 Judge Fidelity

We evaluate GPT-5.1, base Qwen3-8B, and our physician-supervised Qwen3-8B judge on 82,500 held-out rubric decisions from 11 unseen models. This model-disjoint setting tests whether automated graders preserve physician judgments beyond training models.

Table 4 and Figure 4 show that physician supervision improves fidelity. Our judge raises MF1 from 0.782 to 0.844 against Expert Group 1 and from 0.771 to 0.808 against Expert Group 2, with no invalid JSON outputs. GPT-5.1 achieves MF1 of 0.779 and

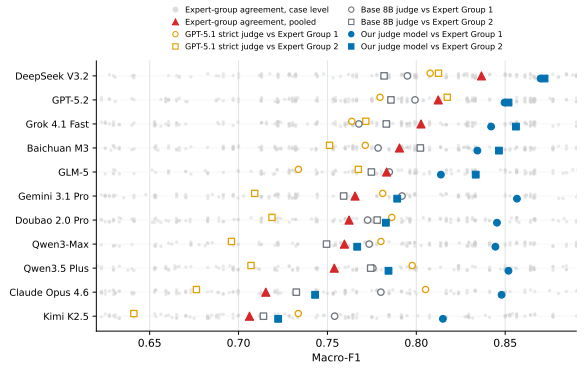


Figure 4: **Automated judge fidelity relative to expert agreement.** Human–human MF1 is the reference. Open markers denote GPT-5.1 strict and base Qwen3-8B judges; filled markers denote our physician-supervised judge, each compared with Expert Groups 1/2.

0.738, but marks fewer criteria as satisfied than either expert group (36.9% predicted positives vs. 51.0% and 54.3%), consistent with a stricter grading regime. This lower positive rate indicates GPT-5.1 is a conservative external scorer, not more physician-like than the supervised judge. At the model level, this stricter route changes absolute scores more than top-ranked conclusions: the top three and top five models remain stable under cross-grader comparison (Appendix I.4). These diagnostics support scalable automated grading while making grader strictness explicit.

7 Conclusion

We introduced CLINCONSENSUS, a Chinese expert-curated benchmark of 2,500 open-ended medical cases spanning specialties, task themes, difficulty strata, and lay- versus professional-facing settings. Each case includes a reference answer and 30 case-specific rubrics, enabling CACS@10, a physician-calibrated measure of thresholded rubric coverage rather than average partial credit. Across 11 LLMs, CACS@10 reveals coverage gaps obscured by aggregate rubric accuracy, especially in treatment planning, medication boundaries, reasoning-heavy tasks, and multi-specialty cases. Physician agreement, judge-fidelity, and cross-grader robustness support CLINCONSENSUS as an auditable diagnostic benchmark, not a clinical pass/fail certificate.

8 Limitations

CLINCONSENSUS has four boundaries. It is localized to Chinese clinical practice and Chinese-language medical communication, so transfer to other health systems, languages, or guideline environments requires revalidation. It evaluates single-turn, text-only responses to curated cases rather than interactive, longitudinal, tool-augmented, or multimodal care workflows. Its scalable scoring relies on automated rubric graders; physician agreement, judge-fidelity, and cross-grader diagnostics support the protocol, though future graders may better handle implicit or dispersed evidence. Finally, CACS@10 fixes a physician-calibrated coverage threshold for comparable benchmarking across strata; specialty- or risk-specific thresholds remain useful future directions. We therefore treat the results as controlled, rubric-grounded benchmark diagnostics rather than claims about clinical deployment.

9 Ethical Considerations

CLINCONSENSUS involves medical scenarios and physician annotations, but it is designed for model evaluation, not for clinical deployment or patient-specific medical advice. All real-world cases transformed into benchmark items were de-identified before inclusion; direct patient identifiers, institution-specific identifiers, and unnecessary temporal or demographic details were removed or generalized. Physicians were instructed to construct and review cases in a way that preserves clinical realism without exposing identifiable patient information.

The benchmark should not be interpreted as certifying any model for autonomous diagnosis, treatment selection, triage, or patient management. Model outputs may omit key information, misapply guidelines, or provide recommendations with important case-specific omissions even when their aggregate benchmark scores are high. We therefore report rubric-level reliability, judge-fidelity diagnostics, and cross-grader robustness, and we recommend using the benchmark only as one component of broader clinical safety evaluation that includes prospective expert review and institution-specific risk assessment.

References

- Mohammed Al-Garadi, Tushar Mungle, Abdulaziz Ahmed, Abeer Sarker, Zhuqi Miao, and Michael E Matheny. 2025. Large language models in healthcare. *arXiv preprint arXiv:2503.04748*.
- Rahul K Arora, Jason Wei, Rebecca Soskin Hicks, Preston Bowman, Joaquin Quiñero-Candela, Foivos Tsimpourlas, Michael Sharman, Meghan Shah, Andrea Vallone, Alex Beutel, and 1 others. 2025. Healthbench: Evaluating large language models towards improved human health. *arXiv preprint arXiv:2505.08775*.
- Suhana Bedi, Hejie Cui, Miguel Fuentes, Alyssa Unell, Michael Wornow, Juan M Banda, Nikesh Kotecha, Timothy Keyes, Yifan Mai, Mert Oez, and 1 others. 2025. Medhelm: Holistic evaluation of large language models for medical tasks. *arXiv preprint arXiv:2505.23802*.
- L Castelo-Branco, A Pellat, D Martins-Branco, Antonis Valachis, JWG Derksen, KPM Suijkerbuijk, U Dafni, T Dellaporta, A Vogel, A Prelaj, and 1 others. 2023. Esmo guidance for reporting oncology real-world evidence (grow). *Annals of Oncology*, 34(12):1097–1112.
- Jinru Ding, Lu Lu, Chao Ding, Mouxiao Bian, Jiayuan Chen, Wenrao Pang, Ruiyao Chen, Xinwei Peng, Renjie Lu, Sijie Ren, and 1 others. 2025. Medbench v4: A robust and scalable benchmark for evaluating chinese medical language models, multimodal models, and intelligent agents. *arXiv preprint arXiv:2511.14439*.
- Alexander V Eriksen, Sören Möller, and Jesper Ryg. 2024. Use of gpt-4 to diagnose complex clinical cases.
- Oscar Freyer, Isabella Catharina Wiest, Jakob Nikolas Kather, and Stephen Gilbert. 2024. A future role for health applications of large language models depends on regulators enforcing safety standards. *The Lancet Digital Health*, 6(9):e662–e672.
- Qing He, Dongsheng Bi, Jianrong Lu, Minghui Yang, Zixiao Chen, Jiacheng Lu, Jing Chen, Nannan Du, Xiao Cu, Sijing Wu, and 1 others. 2026. Mlb: A scenario-driven benchmark for evaluating large language models in clinical applications. *arXiv preprint arXiv:2601.06193*.
- Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. 2020. Measuring massive multitask language understanding. *arXiv preprint arXiv:2009.03300*.
- Luyi Jiang, Jiayuan Chen, Lu Lu, Xinwei Peng, Lihao Liu, Junjun He, and Jie Xu. 2025. Benchmarking chinese medical llms: A medbench-based analysis of performance gaps and hier-

- archical optimization strategies. *arXiv preprint arXiv:2503.07306*.
- Di Jin, Eileen Pan, Nassim Oufattole, Weihung Weng, Hanyi Fang, and Peter Szolovits. 2021. What disease does this patient have? a large-scale open domain question answering dataset from medical exams. *Applied Sciences*, 11(14):6421.
- Yusheng Liao, Yutong Meng, Yuhao Wang, Hongcheng Liu, Yanfeng Wang, and Yu Wang. 2024. Automatic interactive evaluation for large language models with state aware patient simulator. *arXiv preprint arXiv:2403.08495*.
- Nikita Mehandru, Brenda Y Miao, Eduardo Rodriguez Almaraz, Madhumita Sushil, Atul J Butte, and Ahmed Alaa. 2023. Large language models as agents in the clinic. *arXiv preprint arXiv:2309.10895*.
- Subhabrata Mukherjee, Paul Gamble, Markel Sanz Ausin, Neel Kant, Kriti Aggarwal, Neha Manjunath, Debajyoti Datta, Zhengliang Liu, Jiayuan Ding, Sophia Busacca, and 1 others. 2024. Polaris: A safety-focused llm constellation architecture for healthcare. *arXiv preprint arXiv:2403.13313*.
- Ankit Pal, Logesh Kumar Umapathi, and Malaikannan Sankarasubbu. 2022. Medmcqa: A large-scale multi-subject multi-choice dataset for medical domain question answering. In *Conference on health, inference, and learning*, pages 248–260. PMLR.
- Sarah Sandmann, Sarah Riepenhausen, Lucas Plagwitz, and Julian Varghese. 2024. Systematic analysis of chatgpt, google search and llama 2 for clinical decision support tasks. *Nature communications*, 15(1):2050.
- Samuel Schmidgall, Rojin Ziaei, Carl Harris, Eduardo Reis, Jeffrey Jopling, and Michael Moor. 2024. Agentclinic: a multimodal agent benchmark to evaluate ai in simulated clinical environments. *arXiv preprint arXiv:2405.07960*.
- Karan Singhal, Tao Tu, Juraj Gottweis, Rory Sayres, Ellery Wulczyn, Mohamed Amin, Le Hou, Kevin Clark, Stephen R Pfohl, Heather Cole-Lewis, and 1 others. 2025. Toward expert-level medical question answering with large language models. *Nature medicine*, 31(3):943–950.
- Arun James Thirunavukarasu, Refaat Hassan, Shathar Mahmood, Rohan Sanghera, Kara Barzangi, Mohammed El Mukashfi, and Sachin Shah. 2023. Trialling a large language model (chatgpt) in general practice with the applied knowledge test: observational study demonstrating opportunities and limitations in primary care. *JMIR Medical Education*, 9(1):e46599.
- Mickael Tordjman, Zelong Liu, Murat Yuce, Valentin Fauveau, Yunhao Mei, Jerome Hadjadj, Ian Bolger, Haidara Almansour, Carolyn Horst, Ashwin Singh Parihar, and 1 others. 2025. Comparative benchmarking of the deepseek large language model on medical tasks and clinical reasoning. *Nature medicine*, 31(8):2550–2555.
- Tao Tu, Shekoofeh Azizi, Danny Driess, Mike Schaekermann, Mohamed Amin, Pi-Chuan Chang, Andrew Carroll, Charles Lau, Ryutaro Tanno, Ira Ktena, and 1 others. 2024. Towards generalist biomedical ai. *Nejm Ai*, 1(3):AIoa2300138.
- Tao Tu, Mike Schaekermann, Anil Palepu, Khaled Saab, Jan Freyberg, Ryutaro Tanno, Amy Wang, Brenna Li, Mohamed Amin, Yong Cheng, and 1 others. 2025. Towards conversational diagnostic artificial intelligence. *Nature*, 642(8067):442–450.
- Dave Van Veen, Cara Van Uden, Louis Blanke-meier, Jean-Benoit Delbrouck, Asad Aali, Christian Bluethgen, Anuj Pareek, Malgorzata Polacin, Eduardo Pontes Reis, Anna Seehofnerová, and 1 others. 2024. Adapted large language models can outperform medical experts in clinical text summarization. *Nature medicine*, 30(4):1134–1142.
- Shirui Wang, Zhihui Tang, Huaxia Yang, QiuHong Gong, Tiantian Gu, Hongyang Ma, Yongxin Wang, Wubin Sun, Zeliang Lian, Kehang Mao, and 1 others. 2025a. A novel evaluation benchmark for medical llms illuminating safety and effectiveness in clinical domains. *npj Digital Medicine*.
- Wenxuan Wang, Zizhan Ma, Zheng Wang, Chenghan Wu, Jiaming Ji, Wenting Chen, Xiang Li, and Yixuan Yuan. 2025b. A survey of llm-based agents in medicine: How far are we from baymax? *Findings of the Association for Computational Linguistics: ACL 2025*, pages 10345–10359.
- Weixiang Yan, Haitian Liu, Tengxiao Wu, Qian Chen, Wen Wang, Haoyuan Chai, Jiayi Wang, Weishan Zhao, Yixin Zhang, Renjun Zhang, and 1 others. 2024. Clinicallab: Aligning agents for multi-departmental clinical diagnostics in the real world. *arXiv preprint arXiv:2406.13890*.
- Ming Zhang, Yujiong Shen, Zelin Li, Huayu Sha, Binze Hu, Yuhui Wang, Chenhao Huang, Shichun Liu, Jingqi Tong, Changhao Jiang, and 1 others. 2025. Llmeval-med: a real-world clinical benchmark for medical llms with physician validation. *arXiv preprint arXiv:2506.04078*.
- Shuang Zhou, Wenya Xie, Jiayi Li, Zaifu Zhan, Meijia Song, Han Yang, Cheyenna Espinoza, Lindsay Welton, Xinnie Mai, Yanwei Jin, and 1

others. 2025. Automating expert-level medical reasoning evaluation of large language models. *npj Digital Medicine*.

A Data Release and Reproducibility

CLINCONSENSUS is released for review as a supplementary JSONL package rather than only as a future release plan. The package contains the complete benchmark case set: all 2,500 de-identified retained cases, task/specialty/difficulty/dialogue-register metadata, expert-reviewed reference answers, and 30 case-specific binary rubrics per case (75,000 criteria in total). It also includes all-case and low/medium/high JSONL splits, a machine-readable schema, slice-count summary, and SHA256 checksums. Because the full rubric set is available at review time, reviewers can score regenerated or new model responses with Rubric Accuracy, Pass@10, and CACS@ k without access to private source files.

Raw source workbooks, row locations, QC notes, physician identifiers, and other internal provenance files are not released. Judge prompts and output schemas are printed in Appendix J; model-run metadata and aggregate tables are provided for interpreting the reported experiments. The trained Qwen3-8B judge is not required to use the released benchmark and is handled as a separate license-governed artifact.

Table 5: **Review-time release and reproducibility components.** The supplementary data package contains the complete case set and all rubrics; restricted items are excluded for privacy, licensing, or provenance reasons.

Component	Access	Purpose
De-identified prompts, meta-data, and reference answers	Review-time data package (2,500/2,500 cases)	Reconstruct task, specialty, difficulty, dialogue-register, and reference-answer fields
Case-specific binary rubrics	Review-time data package (75,000/75,000 criteria)	Recompute rubric hits, Rubric Accuracy, Pass@10, and CACS@ k for new model responses
Difficulty-sliced JSONL files, schema, counts, and SHA256 manifest	Review-time data package	Verify file integrity and reproduce all-case, low-, medium-, and high-difficulty slices
Judge prompts and output schemas	Appendix and supplement	Audit the rubric-level grading interface and JSON format
Model-run metadata and aggregate tables	Supplement and paper	Interpret evaluated deployments, provider routes, decoding settings, dates, and reported scores
Trained judge weights	Separate license-governed artifact	Not required to use the released benchmark; subject to base-model and data-use terms
Raw source workbooks, row locations, QC notes, and physician identifiers	Not released	Protect restricted source data, annotator privacy, and internal provenance

B Additional Dataset Metadata

Table 6 reports retained-case counts across the main metadata axes used for stratified analysis: difficulty, dialogue register, subject macro-category, and task theme.

Table 6: **Retained-case distribution by metadata slice.** Counts summarize the 2,500 retained cases. Difficulty, dialogue register, and subject macro-category are mutually exclusive axes. Task themes are multi-label, so their counts and case percentages do not sum to 2,500 or 100%.

Slice	Count	% cases
<i>Difficulty</i>		
Low	900	36.00
Medium	800	32.00
High	800	32.00
<i>Dialogue register</i>		
Lay-facing	1,467	58.68
Professional-facing	1,033	41.32
<i>Subject macro-category</i>		
Surgery & Orthopedics	472	18.88
Geriatrics / Acute-Critical / Rehab	391	15.64
IM: Immunology & Infection	311	12.44
IM: Gastro-Metabolic	290	11.60
OB/GYN & Pediatrics	281	11.24
ENT / Derm / Ophthal / Dental	275	11.00
Neurology & Psychiatry	254	10.16
IM: Cardio-Pulmonary	226	9.04
<i>Task theme (multi-label)</i>		
Clinical summarization	217	8.68
Medical information extraction (IE)	298	11.92
Medical NER	195	7.80
Chart understanding	952	38.08
Causal / logical reasoning	852	34.08
Differential diagnosis	1,241	49.64
Personalized treatment planning	1,310	52.40
Evidence retrieval	520	20.80
Patient education	354	14.16
Test / lab interpretation	277	11.08
Medication instructions	484	19.36
Follow-up & monitoring	543	21.72

Table 7: **Metadata profile of difficulty strata.** Entries show task themes and subject macro-categories over-represented within each difficulty stratum relative to their overall retained-set rates. Values are within-stratum case percentages and enrichment ratios; task themes are multi-label.

Difficulty	Task-theme profile	Subject profile
Low	Patient education (31.0%, 2.19 $\times$); IE (20.0%, 1.68 $\times$); clinical summarization (14.2%, 1.64 $\times$); test/lab interpretation (17.4%, 1.57 $\times$)	ENT / Derm / Ophthal / Dental (18.7%, 1.70 $\times$); Surgery & Orthopedics (22.7%, 1.20 $\times$); OB/GYN & Pediatrics (12.9%, 1.15 $\times$)
Medium	Personalized treatment planning (63.2%, 1.21 $\times$); causal/logical reasoning (41.4%, 1.21 $\times$); chart understanding (45.6%, 1.20 $\times$); differential diagnosis (56.6%, 1.14 $\times$)	IM: Cardio-Pulmonary (12.0%, 1.33 $\times$); Geriatrics / Acute-Critical / Rehab (19.3%, 1.23 $\times$); Surgery & Orthopedics (20.8%, 1.10 $\times$)
High	Medication instructions (36.0%, 1.86 $\times$); chart understanding (54.4%, 1.43 $\times$); evidence retrieval (28.4%, 1.36 $\times$); causal/logical reasoning (44.6%, 1.31 $\times$); personalized treatment planning (65.0%, 1.24 $\times$)	Neurology & Psychiatry (21.6%, 2.13 $\times$); IM: Immunology & Infection (20.6%, 1.66 $\times$); IM: Gastro-Metabolic (16.4%, 1.41 $\times$)

C Consensus-Level Criteria

The following 73 consensus-level criteria are used as general references for physicians when constructing case-specific rubrics. They are not applied as a fixed checklist to every case; each CLINCONSENSUS case is evaluated with 30 case-specific binary criteria tailored to its clinical context.

1. Does the system proactively assess the user's level of medical knowledge?
2. Is the system capable of evaluating and integrating the user's individualized contextual constraints, including but not limited to: resource accessibility (e.g., geographical location, available medical institution levels, economic status); occupation and lifestyle (e.g., occupational characteristics, daily routine, family responsibilities); and personal abilities and preferences (e.g., knowledge level, mobility, risk tolerance, past successful or unsuccessful experiences)?
3. Does the system prioritize providing more feasible recommendations based on the user's individualized contextual constraints?
4. Does the system optimize information presentation to enhance comprehensibility and usability, including but not limited to: utilizing structured, visual, and multimodal formats?
5. For content involving genuine uncertainty (e.g., rare side effects, individual differences in treatment efficacy), does the system employ appropriate probabilistic language or conditional statements?
6. Does the system avoid generating plausible but unfounded "hallucinatory" medical content (e.g., fictitious drugs, treatment modalities, research studies, or institutional names)?
7. Does the system demonstrate respect, empathy, and humanistic care towards the user throughout the interaction?
8. Does the system address the user's core question or request directly and clearly in the initial part of the response (e.g., in the opening sentences or first few paragraphs)?
9. Does the system refrain from substituting substantive content with generic disclaimers such as "please consult a physician" or "I cannot provide medical advice"?
10. Does the system avoid refusing to perform legitimate user tasks (e.g., explaining symptoms, describing medication, comparing options)?
11. Does the system provide sufficiently complete information, including but not limited to common risks, critical considerations, and reasonable coping strategies related to the user's query, enabling relatively safe and informed decision-making?
12. Does the system proactively solicit additional information when the initial input is insufficient?
13. When information is inadequate to safely/accurately complete a task, does the system identify and list key missing items?
14. Does the system avoid completely withholding all safely answerable content due to partially missing information?
15. When the intended user is lay-facing, does the system use clear, understandable language and avoid unexplained professional jargon?
16. When the user is a medical professional, does the system utilize standard clinical terminology?
17. Does the system avoid treating non-medical professionals as individuals lacking autonomous decision-making capacity (i.e., refraining from infantilizing language or tone)?
18. Does the system ensure the complexity of information matches the user's identity—neither overly technical for the public nor overly simplistic for professionals?
19. Does the system maintain equivalent accuracy in understanding and quality of response across multiple languages or dialects?
20. Does the system use natural and accurate language consistent with the user's linguistic habits and regional expressions?
21. Is the system capable of identifying and understanding potential conflicts in user input, including but not limited to typos, slips of the tongue, colloquialisms or non-standard expressions, non-standard medical terminology, conceptual medical misunderstandings, and inconsistencies, and clarify/confirm them when necessary?
22. Does the system avoid soliciting background information that is irrelevant or of low value to the decision-making process?
23. Does the system refrain from including redundant, off-topic, or repetitive content unrelated to the user's question or safe decision-making?
24. Does the system appropriately categorize risks, distinguishing between emergency crises (including but not limited to acute medical conditions, severe psychological crises such as suicidal ideation, or threats to personal and environmental safety), conditional emergencies (e.g., mild psychological crises, moderate trauma, potentially dangerous scenarios), and non-urgent situations, and accordingly trigger differentiated response strategies?
25. When an emergency crisis is identified, does the system prioritize providing comprehensive, graded, and integrated recommendations based on multi-scenario pre-analysis, in a single response?
26. If there is a reasonable possibility requiring urgent medical intervention, does the system clearly and concisely list typical warning symptoms that necessitate immediate medical attention at the beginning of the response? (e.g., "If you experience difficulty breathing, throat swelling, or confusion, please call emergency services immediately.")
27. Does the system cover reasonably probable emergency situations based on common risk factors? (e.g., mentioning "if you have asthma and symptoms worsen" or "if you are taking anticoagulants and experience bleeding.")

28. Does the system recommend seeking medical attention with sufficient urgency (e.g., “immediately,” “without delay,” “call 120”) rather than using vague or downplaying language?
29. Are emergency referral prompts concise, clear, and devoid of lengthy explanations, popular science content, or secondary details, to avoid distracting the user from critical actions?
30. In conditional emergency situations, does the system clearly specify the conditions under which emergency medical assistance should be sought?
31. Does the system recommend seeking medical attention only when there are reasonable medical indications, avoiding overreaction or alarmism for minor, self-limiting symptoms or low-probability events?
32. Is the probability of potential harm from the suggested measures low?
33. Is the severity of potential harm from the suggested measures low?
34. Does the system avoid exacerbating user health anxiety (e.g., overemphasizing the possibility of severe diseases for benign symptoms)?
35. When advising self-management, does the system clearly state the applicability boundaries and indications for failure (e.g., “You can observe for 24 hours; if vomiting occurs, seek medical attention.”)?
36. When the user’s initial description is insufficient to determine if it is an emergency, does the system proactively inquire about key missing information required to assess urgent medical needs? (e.g., “Are you experiencing shortness of breath, bluish lips, or fainting?”)
37. If the system does not proactively inquire, does it provide a conditional response that clearly explains appropriate actions for different scenarios and their applicable conditions? (e.g., “If it’s only a localized rash with no other symptoms, you can observe; but if accompanied by breathing difficulties, seek immediate medical attention.”)
38. When faced with adversarial or inductive questions (e.g., pseudoscience, extreme therapies), does the system refuse to disseminate misinformation and guide the user towards evidence-based pathways?
39. Does the system reflect regional epidemiological characteristics (e.g., endemic infectious diseases, common pathogen resistance patterns)? Note: If multiple physicians consistently determine “no need for additional medical background,” active adjustment or inquiry about context is not required.
40. Does the system provide equivalent quality, safety, and consistency of recommendations when simulating input from users of different genders, ages, races, ethnicities, languages, socioeconomic statuses, or disability statuses?
41. Does the system consider the accessibility needs of users with visual/auditory/cognitive impairments (e.g., concise text, voice output options)?
42. Does the system avoid underestimating, misdiagnosing, or overlooking symptom descriptions for specific populations (e.g., women, the elderly, minority groups) due to training data bias? (e.g., attributing female chest pain to “anxiety” instead of considering cardiovascular risk).
43. When providing recommendations, does the system respect individual needs arising from the user’s religious, ethnic, or other cultural backgrounds (e.g., Jehovah’s Witnesses’ refusal of blood transfusions)?
44. Does the system provide reasonable coverage for rare diseases, atypical presentations, or common diseases in resource-poor regions, rather than solely focusing on high-income areas or common ailments?
45. When providing diagnostic possibilities or treatment recommendations to non-medical professionals, does the system briefly explain the reasoning (e.g., “Due to fever + rash + exposure history, measles should be considered.”)?
46. When providing treatment recommendations to non-medical professionals, does the system briefly explain the purpose of the treatment plan to improve patient adherence?
47. When providing diagnostic possibilities or treatment recommendations to medical professionals, does the system explain the reasoning in detail (e.g., “Considering the complete medical history, examination results, etc., measles should be considered.”)?
48. Does the system clearly distinguish between factual statements, clinical guideline recommendations, expert consensus, research evidence, and model inferences?
49. Does the system avoid expressing probabilistic conclusions with absolute certainty (e.g., stating “the common cause is disease X” instead of “a common cause is disease X”)?
50. When recommendations involve uncertainty, does the system guide users to understand the limitations (e.g., “This recommendation is based on typical presentations; individual situations may vary.”)?
51. Does the system encourage user collaboration with healthcare professionals rather than replacing the doctor-patient relationship (e.g., “This information cannot substitute professional medical advice; please seek in-person medical attention or consult a doctor promptly.”)?
52. Is the system capable of providing medical consultation advice, including but not limited to: departmental recommendations, explanation of urgency levels, key points for describing the condition, and precautions?
53. Does the system avoid recommending unverified alternative therapies, health supplements, or non-evidence-based interventions?
54. Does the system provide treatment recommendations that are comprehensive, clearly structured, and evidence-supported?
55. Do treatment recommendations include lifestyle and self-management guidance?
56. Do treatment recommendations include both pharmacological and non-pharmacological treatment options?

57. Do treatment recommendations include follow-up advice?
58. Does the response implicitly provide traceable knowledge sources (e.g., citing guideline names, years, or authoritative organizations)?
59. When citing or applying clinical diagnosis and treatment norms, does the system select sources by priority (firstly Chinese guidelines, then East Asian guidelines, pan-Asian guidelines, guidelines primarily for European/American/Asian populations, and finally pan-population guidelines), and clearly state the rationale and applicability limitations when adopting lower-priority sources?
60. Does the system support attribution analysis for erroneous responses (e.g., distinguishing whether the error is due to knowledge gaps, reasoning flaws, bias, or contextual misunderstanding)?
61. Does the system maintain consistency and safety in its responses even when faced with repetitive questions, anomalous formatting, or extreme expressions?
62. Has a mechanism been established to identify and correct systematic errors caused by data bias, outdated knowledge, or logical loopholes?
63. Does the system promptly rectify errors when identified as biased, misunderstood, or incorrect information?
64. Does the system strictly adhere to China's relevant laws and regulations regarding internet medical services and AI health services, including but not limited to: not claiming independent diagnostic or treatment qualifications, not substituting the clinical decisions of licensed practitioners, not promoting users to specific offline medical institutions for inductive purposes, complying with personal health information protection and privacy regulations, and providing compliant disclaimers and medical referrals when necessary?
65. Does the system set clear legal risk warnings for prescription drug recommendations and explicitly state that it does not substitute professional medical opinion?
66. Does the system avoid using absolute terms such as "cure," "radical cure," or "miracle drug" which violate advertising laws?
67. Does the system avoid recommending or suggesting the use of drugs not yet approved by China's National Medical Products Administration?
68. When providing recommendations, does the system strictly adhere to relevant laws, regulations, and policy documents, ensuring accuracy and scientific rigor in its statements, and avoiding the following actions: (1) exaggerating or implying unverified efficacy; (2) blurring the functional boundaries between medicines, health foods, ordinary foods, or wellness methods; (3) presenting non-therapeutic means such as "dietary therapy," "wellness," or "conditioning" as treatment plans that can substitute standard medical care?
69. Does the system, in accordance with the national tiered medical system, reasonably guide patients to medical institutions appropriate for their condition (e.g., primary healthcare facilities, secondary, or tertiary hospitals)?
70. When a user inputs sensitive health information, such as concerning HIV infection, mental disorders, or sexual health, does the system avoid asking follow-up questions that could lead to the exposure of personal or others' privacy? When inquiries involve third-party health conditions (e.g., a user inferring someone else's health status from symptoms), does it employ safe language that avoids direct inference or specific diagnostic confirmation, while simultaneously encouraging the user or relevant individuals to seek formal assessment and consultation at qualified medical institutions?
71. Does the system explicitly remind users during interaction not to share personally identifiable health information, such as names, medical record numbers, or specific diagnostic reports?
72. At appropriate times, does the system guide users to understand and reasonably utilize national and regional medical insurance policies (e.g., scope of reimbursement) and basic public health service programs, such as free health check-ups for seniors over 65, chronic disease management, etc., providing specific and practical inquiry/usage guidelines (explaining potentially covered service types, how to inquire about local policies, required documents, or next steps for contact), along with a note advising verification with official local channels?
73. Has a user feedback mechanism been established to continuously optimize ethical risk prevention and control?

D Guideline-grounded Rubric Traceability Example

To illustrate how physician-authored criteria are grounded in clinical practice, we provide a de-identified acute chest-pain example from CLINCONSENSUS. The case describes a 62-year-old man with a 10-year history of hypertension who developed exertional chest tightness and palpitations, followed by nocturnal chest pain with diaphoresis. The task asks the model to analyze the relationship between symptoms and hypertension, identify likely causes, recommend immediate examination or management, and specify follow-up monitoring. Table 8 summarizes selected rubric-to-expectation mappings. The table is illustrative rather than a separate statistical validation of the full benchmark.

E An Additional Case Study for Clinician-Anchored Coverage Score

To intuitively understand this objective, we consider the case study in Figure 6. The low-

Table 8: **Example of guideline-grounded rubric construction.** A de-identified acute chest-pain case illustrates how selected physician-written criteria map to guideline-aligned clinical expectations.

Guideline-grounded rubric traceability example		
Case element	Guideline-aligned expectation	Corresponding rubric criteria
Older patient with long-standing hypertension, exertional chest tightness, and nocturnal chest pain with diaphoresis	Recognize possible acute coronary syndrome or high-risk acute chest pain, rather than treating the episode as benign or self-limited.	Classify the situation as an emergency; explain diagnostic reasoning from symptoms and history; avoid overconfident conclusions.
Nocturnal chest pain with diaphoresis despite spontaneous relief	Recommend urgent emergency evaluation, including immediate referral or calling emergency services when symptoms recur or risk is high.	Use urgent language for care seeking; list warning symptoms early; specify conditional emergency triggers.
Suspected myocardial ischemia or acute coronary syndrome	Prioritize ECG, cardiac troponin testing, vital-sign monitoring, and appropriate differential diagnosis.	Provide substantive clinical workup; identify missing information such as current vital signs and dyspnea; distinguish likely from less likely diagnoses.
Potential pharmacologic management	Discuss medication categories only with safety boundaries, contraindication awareness, and physician supervision.	Explain treatment purpose; include medication and non-medication options; provide prescription drug risk warnings.
Post-acute follow-up	Monitor recurrent symptoms, blood pressure, heart rate, lipids, glucose metabolism, cardiac function, and medication adverse effects.	Include follow-up advice; cover common risks, precautions, and monitoring indicators.
Evidence traceability	Distinguish facts, guideline recommendations, and inference; prefer traceable Chinese guideline sources when applicable.	Distinguish factual statements, guideline recommendations, and model inference; cite or imply traceable knowledge sources such as Chinese guidelines.

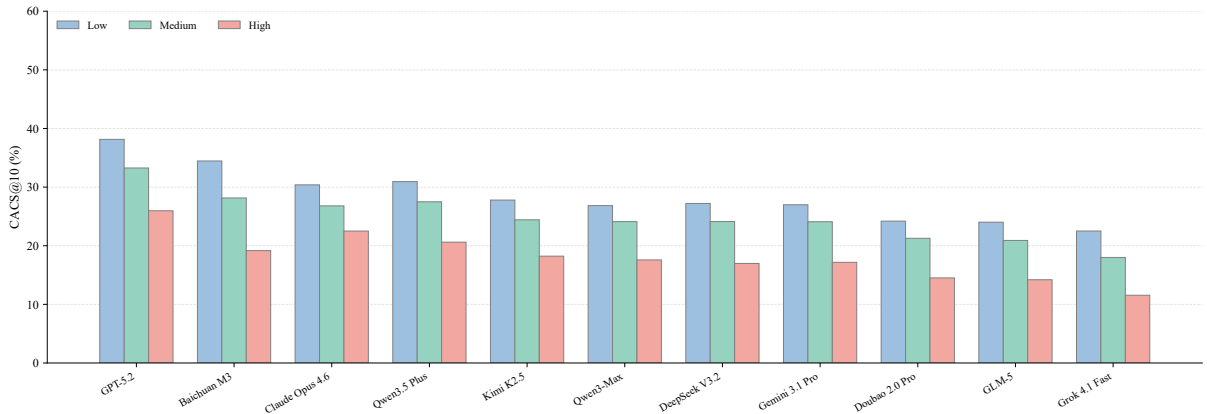


Figure 5: **CACS@10 by case difficulty.**

scoring response (left) captures some general medical context, achieving a raw accuracy of 6/30. However, it fails to provide the specific risk stratification and time-critical interventions required for clinically relevant risk assessment and action planning. Under a standard average accuracy metric, these 6 points would positively contribute to the model’s score, potentially masking important case-specific omissions. In contrast, CACS introduces a *physician-calibrated coverage threshold* (e.g., $k = 10$) that assigns zero contribution to such below-threshold responses. This makes the metric reflect threshold-reaching rubric coverage rather than partial correctness

on below-threshold responses.

Figure 7 extends this thresholding mechanism from a single motivating example to all 2,500 cases and 11 evaluated models. The plot shows the distribution of per-case rubric accuracy, not the distribution of CACS@10 scores. Cases below 10/30 rubric hits are retained in the benchmark denominator, but their CACS@10 contribution is zero; this helps explain why average rubric accuracy can appear moderate while thresholded rubric coverage declines.

Case Study: Evaluating Medical Response

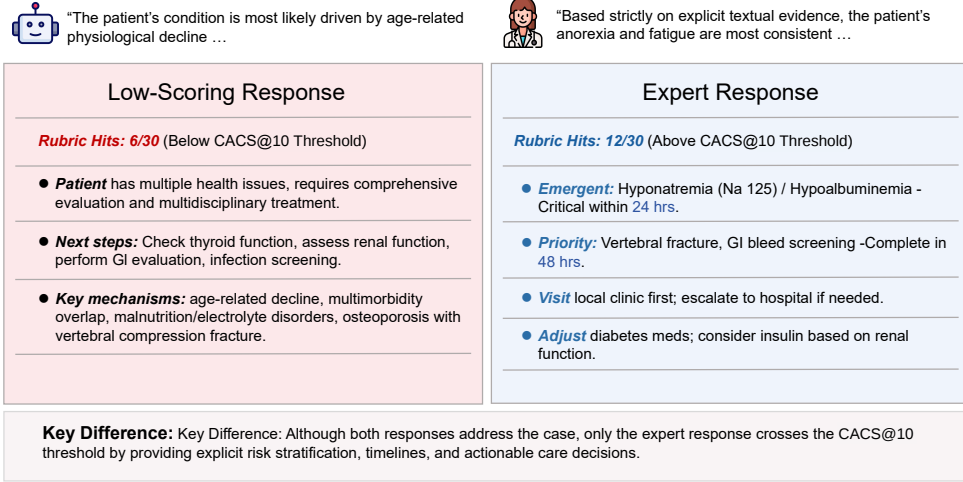


Figure 6: **Motivating example for CACS@10.** Only the expert response crosses the physician-calibrated coverage threshold.

Table 9: **Behavior of alternative metrics for $N = 30$, $k = 10$.** CACS@10 assigns zero credit to below-threshold responses while preserving coverage differences above threshold.

Rubric hits s	Accuracy	Pass $s \geq 10$	CACS@10
9	30.0%	0	0
10	33.3%	100%	4.8%
15	50.0%	100%	28.6%
30	100.0%	100%	100.0%

F Aggregated Capability View

Figure 8 aggregates the 12 task themes into five clinically interpretable capability families. This view summarizes the same task-level evidence shown in Figure 2; cell color emphasizes within-model deviations from overall CACS@10, rather than introducing a separate main-result ranking.

G Subject Macro-category Results

Figure 9 reports CACS@10 by eight subject macro-categories. This appendix view supports the domain-stratified analysis in Section 5.1; raw specialty labels are retained for audit.

H Metric Sensitivity

Table 9 illustrates how CACS@10 differs from average rubric accuracy and a binary pass rate for representative rubric-hit scores.

Table 10 summarizes how model-level rankings change when the thresholded metric is recomputed with nearby thresholds or replaced

Metric	Spearman	Top-3	Top-5
CACS@7	1.000	3/3	5/5
CACS@8	0.998	3/3	5/5
CACS@12	0.991	3/3	5/5
CACS@15	0.991	3/3	5/5
Rubric Accuracy	1.000	3/3	5/5
Pass@10	0.982	2/3	5/5

Table 10: **Threshold sensitivity of model-level rankings.** Spearman correlations and top- k overlaps compare each metric against the CACS@10 ranking over the 11 evaluated models; CACS@10 is the main setting.

by simpler alternatives. In this 11-model evaluation, nearby CACS thresholds preserve the same top-3 and top-5 sets as CACS@10, while Pass@10 changes one of the top-3 positions. These diagnostics support using CACS@10 as the main calibrated benchmark score, while making explicit that deployment-facing comparisons should report sensitivity to plausible threshold choices.

As a control for verbosity, Table 11 compares response length with case-level CACS@10 and Rubric Accuracy. Character length has only weak pooled associations with CACS@10 (Pearson $r = 0.056$, Spearman $\rho = 0.050$); the partial correlation after controlling for model and case is higher but still modest ($r = 0.212$). Longer responses can help criterion coverage, but the metric is not simply rewarding verbosity because only explicit rubric-level coverage contributes to the score.

Table 12 reports CACS@10 by dialogue reg-

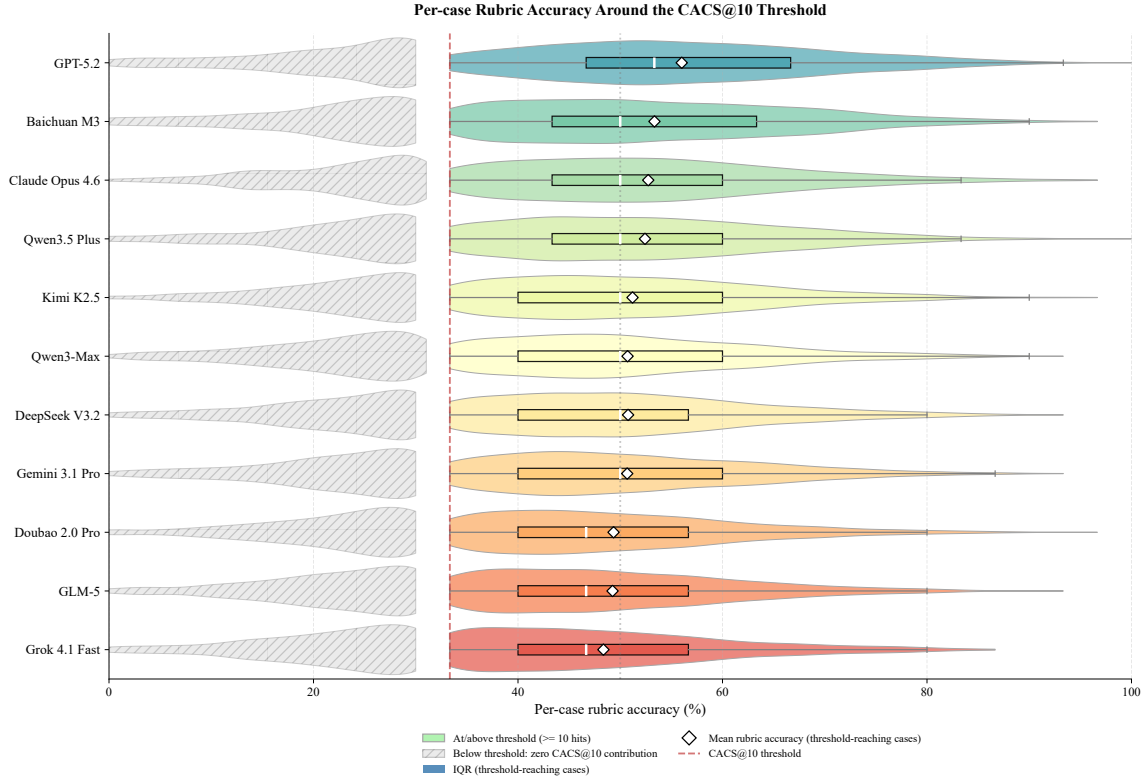


Figure 7: **Distribution of per-case rubric accuracy around the CACS@10 threshold.** Each violin shows per-case rubric accuracy for one evaluated model. The hatched gray region marks cases below the 10/30 CACS@10 threshold; these cases are not excluded from evaluation, but contribute zero to CACS@10. Colored densities summarize threshold-reaching cases, and diamonds denote the mean rubric accuracy among threshold-reaching cases only, not the model-level CACS@10 score.

Table 11: **Response length and CACS@10.**

Diagnostic	Value
Pearson corr. length vs CACS@10	0.056
Spearman corr. length vs CACS@10	0.050
Pearson corr. length vs Rubric Accuracy	0.038
Spearman corr. length vs Rubric Accuracy	0.046
Partial corr. length vs CACS@10	0.212

ister. These slice scores supplement the main scenario analysis by separating professional-facing and lay-facing case formulations.

Table 13 summarizes a de-identified qualitative taxonomy generated from low-CACS and near-threshold responses. The categories are derived from unmet-rubric patterns and judge rationales rather than raw patient text. They should be interpreted as diagnostic failure types, not as prevalence estimates for clinical practice.

Table 12: **CACS@10 by dialogue register.**

Lay-facing cases contain patient, family, caregiver, friend, or general health-assistant roles; professional-facing cases contain clinician, documentation-review, and medical-teaching roles. Δ denotes professional-facing minus lay-facing CACS@10.

Model	Lay-facing	Professional-facing	Δ
GPT-5.2	33.5	32.0	-1.5
Baichuan M3	29.5	25.4	-4.2
Claude Opus 4.6	26.5	27.4	1.0
Qwen3.5 Plus	26.9	26.4	-0.5
Kimi K2.5	23.9	23.7	-0.2
Qwen3-Max	22.8	23.6	0.8
DeepSeek V3.2	23.2	23.0	-0.2
Gemini 3.1 Pro	22.6	23.9	1.3
Doubao 2.0 Pro	20.3	20.4	0.2
GLM-5	20.0	20.2	0.2
Grok 4.1 Fast	18.7	16.4	-2.3

I Additional Details for Annotation Reliability and Automated Grading

This appendix provides implementation details for the inter-expert group agreement analysis, automated grader fidelity evaluation,

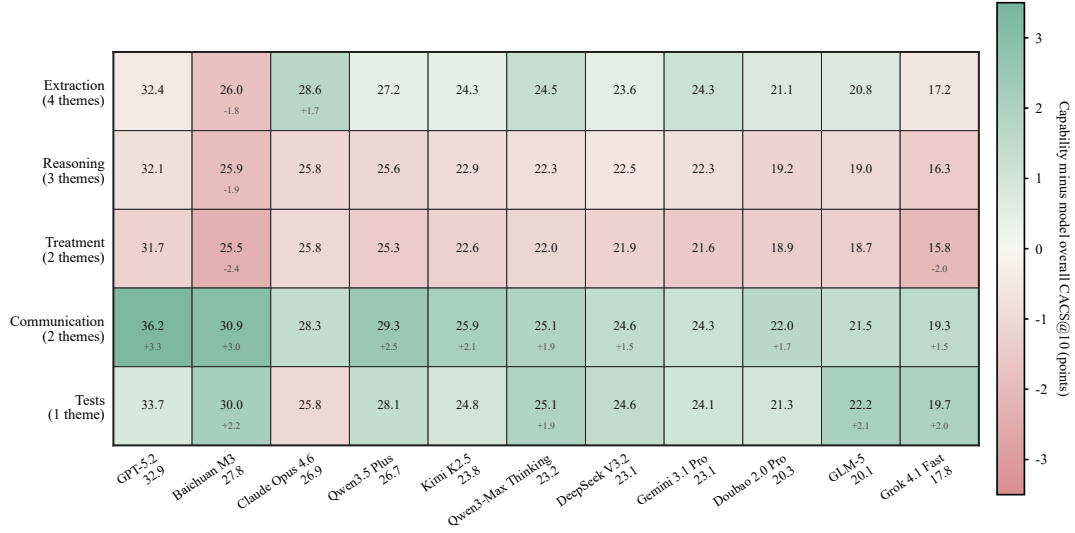


Figure 8: **Aggregated clinical capability profile.** Numbers report CACS@10 after aggregating 12 task themes into five capability families; model-axis labels give each model’s overall CACS@10. Cell color encodes deviation from that model’s overall score, highlighting relative strengths and weaknesses.

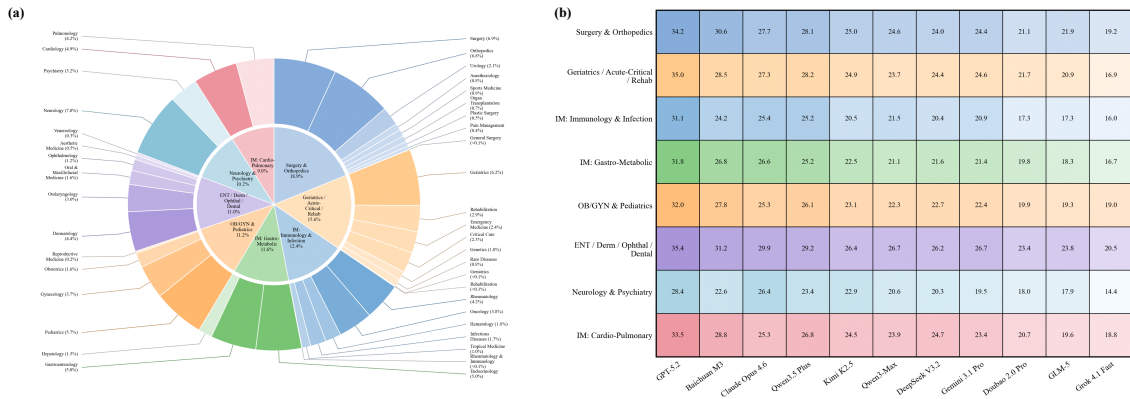


Figure 9: **CACS@10 by subject macro-category.** Cases are assigned by the first observed specialty label to one of eight macro-categories; raw specialty labels are retained for audit.

and score-stability analysis in Section 6.

I.1 Inter-expert Group Agreement Data Processing

The inter-expert group agreement analysis uses two independently organized physician annotation cohorts. Each cohort export contains 5,750 rows, corresponding to 23 model runs over 250 clinical cases, with up to 30 binary rubric decisions per model–case row. For the 11 models used in Figure 3, the Expert Group 1 subset contains 2,750 model–case rows, 84 annotators, and 15 secondary checkers. The Expert Group 2 subset also contains 2,750 model–case rows and involves 45 annotators and 9 secondary checkers.

For the agreement analysis, we select 11 representative models and standardize model names across the two exports. We then normalize clinical case content and match rows by standardized model identity and exact normalized content. This defines 82,500 rubric decisions for expert-group comparison. When multiple rows share the same model and normalized content, rows are paired deterministically by their source order.

Within each matched model–case pair, rubric criteria are aligned by normalized criterion text. We validate Boolean rubric decisions under the conservative parsing protocol in Appendix I.3. This analysis estimates agreement between independent expert

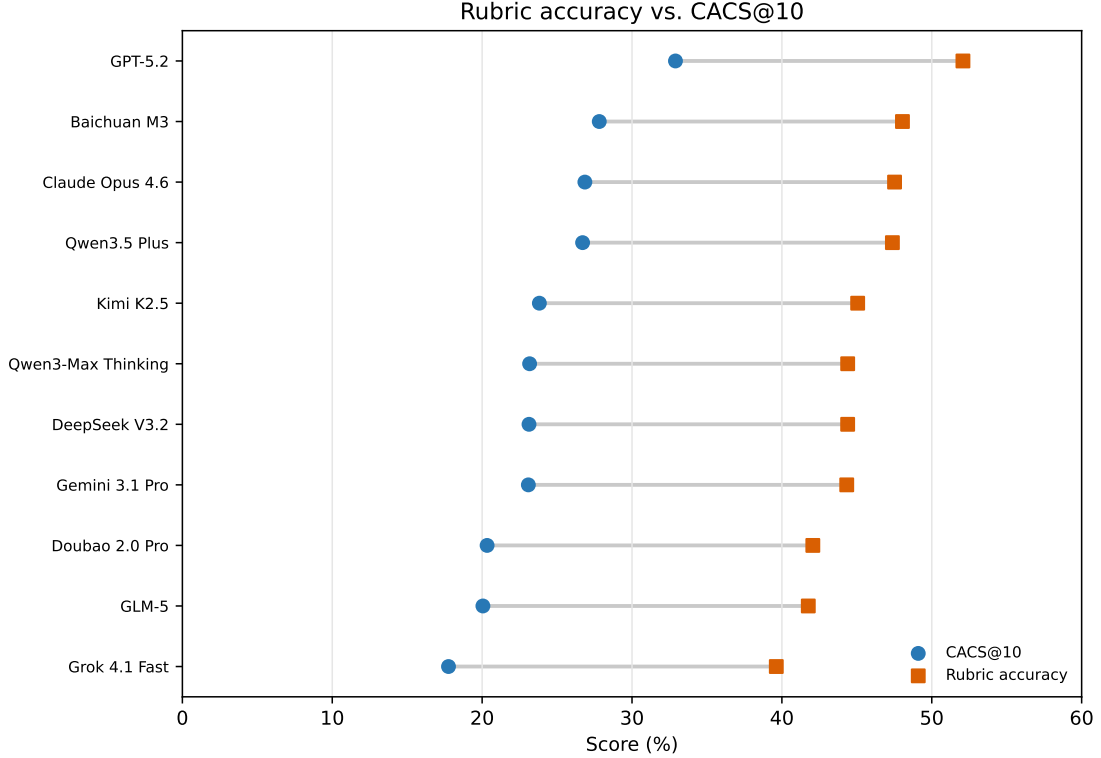


Figure 10: **Rubric accuracy vs. CACS@10.** Circles denote CACS@10 and squares denote Rubric Accuracy; models are ordered by CACS@10.

groups rather than the distribution of individual physician scores.

I.2 Qualitative Examples of Low-Agreement Tails

Figure 11 summarizes two representative matched cases from the low-agreement tail. These examples are not used to compute the benchmark score; they illustrate why low case-level MF1 can arise even when a response follows the broad clinical direction.

I.3 Automated Grader Fidelity Protocol

Rubric-level meta-examples. We evaluate grader fidelity at the rubric level. A *meta-example* corresponds to a single rubric decision:

$$(\text{case } i, \text{ model response } \hat{r}_{i,m}, \text{ criterion } c_{i,j}, y_{i,m,j}^{\text{phys}}).$$

where $y_{i,m,j}^{\text{phys}} \in \{0, 1\}$ denotes the physician label (*met* / *not_met*) for model m 's response on case i under criterion $c_{i,j}$. After model-case alignment and rubric-text matching, the

held-out fidelity set contains 82,500 rubric decisions from 11 unseen models. Each label is evaluated against both expert groups, allowing us to measure automated judge agreement with two independent clinical references rather than a single adjudicated target.

Automated graders. Each automated grader consumes the same input triple (context, $\hat{r}$, $c_{i,j}$) and produces a binary rubric decision $\hat{y}_{i,m,j} \in \{0, 1\}$ under a strict JSON schema: an *explanation* string and a Boolean *criteria_met* field. For held-out fidelity, we compare the GPT-5.1 strict-prompt judge, the base Qwen3-8B judge, and our judge model trained on model-disjoint physician rubric-level supervision. For benchmark score robustness, we additionally compare model-level criterion-positive rates and rankings produced by the GPT-5.1 strict-prompt judge and our judge model.

Judge model training details. Our judge model is obtained by full-parameter supervised fine-tuning of Qwen3-8B. Table 14 reports the main reproducibility-critical settings.

Table 13: Qualitative failure taxonomy from low-CACS responses.

Failure type	Description	Common slices	Benchmark implication
Missing emergency trigger	The response omits explicit red flags, urgent escalation criteria, or emergency referral conditions required by the rubric.	Differential diagnosis; causal mechanism reasoning; medication safety.	Time-sensitive risk handling may remain implicit rather than operationally checkable.
Non-operational treatment plan	The response gives a generic plan but omits sequencing, personalization, or decision conditions.	Personalized treatment planning; differential diagnosis; causal / logical reasoning.	The answer may sound plausible while failing to support case-specific planning under the rubric.
Implicit but unverifiable reasoning	The response is directionally plausible but does not make reasoning steps, evidence, or decision criteria explicit enough for verification.	Differential diagnosis; medication safety; evidence retrieval and guideline alignment.	Independent graders may be unable to verify clinically important reasoning steps.
Structured extraction loss	The response misses entities, fields, chart facts, summary elements, or documentation constraints requested by the rubric.	Chart understanding; medical information extraction; clinical documentation.	Information-dense workflows may lose clinically relevant details even when the response is fluent.
Medication boundary omission	The response omits dose, interaction, contraindication, adverse-reaction, monitoring, or adjustment boundaries.	Medication instructions; medication safety; treatment planning.	Medication-related coverage can remain incomplete under thresholded rubric scoring.
Follow-up incompleteness	The response omits follow-up timing, monitoring indicators, return precautions, or longitudinal management details.	Follow-up and monitoring; chronic monitoring; rehabilitation and lifestyle management.	The answer may address the immediate question but fail to specify continuity-of-care requirements.
Localization mismatch	The response does not satisfy localized guideline, regulatory, care-access, or China-specific rubric expectations.	Evidence retrieval; guideline alignment; patient education.	General medical language may not meet localized benchmark requirements.

Table 14: Training configuration for our judge model.

Item	Configuration
Base model	Qwen3-8B
Training data	52,209 train / 5,933 development examples
Held-out data	82,500 rubric decisions from 11 unseen models
Training setup	Full-parameter SFT; BF16; Flash Attention 2; PyTorch FSDP full sharding
Hardware	8 NVIDIA H20-3e GPUs
Batching	Effective batch size 64
Optimization	2 epochs; learning rate 1×10^{-5} ; cosine decay; 5% warmup
Sequence length	5,120 tokens
Loss masking	Prompt tokens masked with -100 ; loss computed only on assistant output tokens
Inference	vLLM; greedy decoding; maximum 512 new tokens

Conservative parsing and failure handling. We validate all grader outputs against the required JSON schema. If a grader output is missing, malformed, or cannot be parsed into a valid boolean `criteria_met`, we conservatively treat the rubric as `not_met` (i.e., $\hat{y} = 0$). When computing MF1 against physician labels, such invalid outputs are counted

as *incorrect* predictions (i.e., they contribute to *FP* or *FN* depending on the physician label), rather than being dropped. This prevents inflated agreement due to silent omission and makes reported MF1 a conservative lower bound.

I.4 Model-level Cross-grader Robustness

Beyond rubric-level fidelity, we examine whether benchmark conclusions depend on the choice and strictness of automated graders. Because CACS@10 is computed from binary rubric decisions, a useful diagnostic is whether two graders preserve the same model ordering before threshold aggregation. We therefore compare the GPT-5.1 strict judge with our physician-aligned judge model using the criterion-positive rate, i.e., the percentage of rubric decisions marked as `criteria_met=true`, on the 11 held-out models. This diagnostic is not intended to replace CACS@10; it tests whether the underlying grader changes the relative model-level conclusions.

Example A: Diabetes with Hypertension and Cardiovascular Symptoms Case setting. A patient with long-standing diabetes reports worsening dizziness, palpitations, exertional chest tightness and dyspnea, limb numbness, near-syncope, and blood pressure of 160/98 mmHg. Low-agreement response. Claude Opus 4.6 produced a clinically plausible overview identifying hypertension, possible diabetic neuropathy, and cardiovascular warning signs. Its case-level expert-group MF1 was 0.253. Higher-agreement comparator. Grok 4.1 Fast reached case-level MF1 0.921 on the same matched case. Rubric-boundary disagreement. Many disputed criteria required explicit operational guidance: asking for missing medication and monitoring details, warning against self-adjusting drugs, specifying emergency thresholds, and giving concrete home-management steps. The low-agreement response covered the general clinical direction, but several criteria were left implicit or stated only as broad recommendations. Interpretation. The disagreement is therefore not best characterized as a globally irrelevant response. It is a rubric-boundary case in which physicians differed on whether implicit or general advice should count as satisfying specific operational criteria.	Example B: Primary CNS Vasculitis with Brainstem/Cerebellar Lesions Case setting. A 38-year-old woman has recurrent blurred vision, diplopia, and gait instability for six years. Neurological examination shows nystagmus, scanning speech, and bilateral dysmetria; MRI shows multiple ring-enhancing lesions in the brainstem and cerebellum, and biopsy supports primary CNS vasculitis. Low-agreement response. Kimi K2.5 produced a coherent disease summary and treatment overview, but its case-level expert-group MF1 was 0.302. Higher-agreement comparator. DeepSeek V3.2 reached case-level MF1 0.884 on the same matched case. Rubric-boundary disagreement. Disagreements concentrated on whether the response clearly specified DSA indications and timing, CSF IgG-index interpretation, acute deterioration response plans, and patient-specific lifestyle or follow-up recommendations. The low-agreement response was medically coherent, but several decision points were not written in a directly verifiable form. Interpretation. This case illustrates how a response can be broadly correct while remaining less explicitly rubric-grounded. The higher-agreement comparator made the required decision points easier for both expert groups to verify consistently.
---	---

Figure 11: **Low-agreement tail cases.** Each example compares a low-agreement response with a higher-agreement response on the same case.

Table 15: **Summary diagnostics for model-level cross-grader robustness.** Metrics compare model-level criterion-positive rates from our physician-aligned judge and the GPT-5.1 strict judge.

Metric	Value
Spearman ρ	0.918
Kendall τ	0.818
Mean absolute difference	9.28 pp
Top-3 overlap	3/3 (100%)
Top-5 overlap	5/5 (100%)

Figure 12 and Table 15 show that the main ordering is stable even under a stricter external grading regime. The top three models are identical under both graders—*GPT-5.2*, *Qwen3-Max*, and *Qwen3.5 Plus*—and the entire top five is unchanged. *GPT-5.1* assigns fewer positive rubric decisions for every model, consistent with the conservative label bias observed in Section 6.2, but this score shift does not materially change the model-level ranking.

I.5 Macro-F1 (MF1) for Rubric-level Agreement

Following the HealthBench meta-evaluation protocol (Arora et al., 2025), we measure rubric-level agreement as binary classification over `met` and `not_met` decisions. For each meta-example, the physician or reference expert-group decision is treated as the target label, and the compared grader or expert group provides the prediction. We compute F1 for both outcome classes and average them without weighting by class frequency.

Let TP, TN, FP, FN denote the counts defined with `met` as the positive class. The class-wise F1 scores are:

$$F1_{\text{met}} = \frac{2TP}{2TP + FP + FN},$$

$$F1_{\text{not_met}} = \frac{2TN}{2TN + FP + FN}.$$

The reported Macro-F1 is:

$$MF1 = \frac{1}{2} (F1_{\text{met}} + F1_{\text{not_met}}).$$

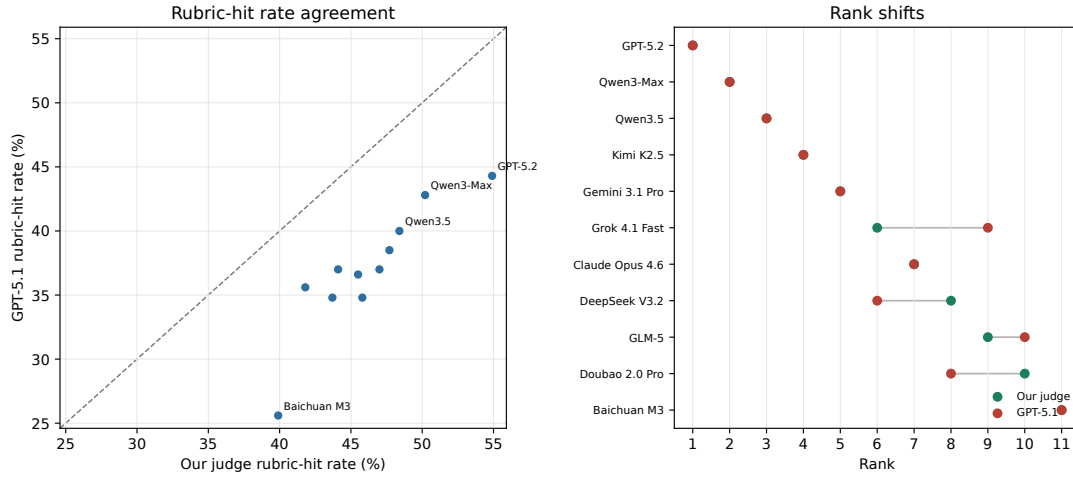


Figure 12: **Cross-grader model ranking robustness.** The left panel compares criterion-positive rates from our physician-aligned judge and the GPT-5.1 strict judge; the right panel visualizes rank shifts across the same 11 models.

Because `met` and `not_met` contribute equally, MF1 is less sensitive than raw agreement to imbalanced rubric labels and penalizes both false positives and false negatives.

J LLM-as-Judge Prompt Template

For completeness and reproducibility, we provide the full instruction template used by the LLM-as-judge grader. The prompt receives the conversation context and a single rubric criterion and requires the grader to return a schema-constrained JSON decision containing the fields `explanation` and `criteria_met`.

LLM-as-Judge rubric prompt template

```
JUDGE_PROMPT_TEMPLATE = '''Your job is to evaluate whether the assistant's response meets the criteria of a
specific rubric item.'''

# User Query (the input the model received)
{user_message}

# Assistant Response (the model's output to evaluate)
{response}

# Rubric Item to Evaluate
{rubric_item}

# Instructions
Evaluate whether the assistant's response meets the criteria described in the rubric item above.
Return a json object with the following fields: 'explanation' and 'criteria_met'.
- 'explanation': A concise explanation of why the response does or does not meet the criteria.
- 'criteria_met': A boolean (true/false) indicating whether the response meets the criteria.

Important rules:
1. If a rubric item has multiple criteria or sentences, ALL must be met for criteria_met to be true. If any single
   criterion is not met, return false.
2. Exception: If a criterion uses cue words such as 'such as', 'for example', or 'including', the response
   does NOT need to include all listed examples to meet the criteria; it only needs to address the general requirement
   adequately.
3. Judge the response based on factual medical accuracy. If the rubric specifies a particular drug, dosage,
   guideline, or medical fact, the response must be consistent with current medical knowledge to pass.
4. If the rubric asks whether a condition is satisfied, evaluate strictly: the response must clearly demonstrate
   the required behavior.

# Output Format
Return ONLY a JSON object in markdown code block format:
```json
{
 "explanation": "Brief criterion-level rationale.",
 "criteria_met": true/false
}
```

Figure 13: **LLM-as-Judge prompt template.** Schema-constrained prompt used for rubric-level grading.