

MAM-AI: An On-Device Medical Retrieval-Augmented Generation System for Nurses and Midwives in Zanzibar

REN YI 任一, École Polytechnique Fédérale de Lausanne, Switzerland bonjour@renyi.ch

Maternal and newborn mortality remain among the highest in sub-Saharan Africa, where midwifery care is often delivered by nurses who lack midwifery training to international standards, and consulting authoritative guidance at the point of care is hard—the guidelines are long and connectivity is intermittent. We present MAM-AI, a medical question-answering assistant for nurse-midwives in Zanzibar that runs entirely on a commodity Android device: a question is embedded (EmbeddingGemma, 300M) and matched against a curated corpus of 87 guideline documents (63,650 passages), then answered with citations by a 4B int4 generator (Gemma 4 E4B)—fully offline, with no query leaving the device. We evaluate the exact deployed configuration with a layered, bottom-up methodology—retriever, generator under oracle context, end-to-end, and latency—scored by LLM judges validated against physician-written rubrics. The evaluation relocates the hard problem. On-device retrieval is essentially solved: the 300M embedder ranks third of seven retrievers and rivals cloud systems, so the passages the system needs are usually found. The small generator is what remains in doubt: adding retrieved context does not improve its answers and hurts larger models, and at 4B it cannot be both helpful and safe at once—of two same-size candidates, the more helpful one commits genuine dangerous errors, so we deploy the other, which is about twice as faithful to its sources (as faithful as a frontier model), and recover its helpfulness with a redesigned prompt that cuts deflection from 33% to 3%. Corpus quality is decisive for the same reason: the small model leans on what it retrieves, so where the corpus holds the right passage the answer is specific and actionable, and where it does not the answer goes vague. On a current Android device the system returns a cited answer in tens of seconds, fully offline; lower-end field hardware is slower. MAM-AI is a thoroughly evaluated, open-source research prototype, not a fielded product; the system, knowledge-base pipeline, benchmarks, and evaluation harness are released.

Additional Key Words and Phrases: retrieval-augmented generation, on-device inference, medical AI, low-resource settings, maternal health, evaluation

1 Introduction

Maternal and newborn mortality remain among the highest in sub-Saharan Africa, and Tanzania—which includes the semi-autonomous archipelago of Zanzibar—is among ten countries that together account for more than 60% of maternal and neonatal deaths worldwide [Bohle 2025]. In Zanzibar most women deliver in government facilities that are often understaffed and under-resourced, where midwifery care is provided by nurses who lack midwifery training to international standards [Bohle 2025]. Strengthening midwife-led care is high-leverage—universal access to it could prevent an estimated 67% of maternal and 64% of neonatal deaths [Nove et al. 2021]—yet a basic obstacle is that authoritative, evidence-based guidance is scarce at the point of care [Swiss Tropical and Public Health Institute 2024]: the guidelines are long, and looking something up online is unreliable where connectivity is intermittent and mobile data costly.

We present MAM-AI, a medical question-answering assistant that runs entirely on an Android device. The nurse-midwife’s question is embedded and matched against an on-device store of guideline passages; the most relevant passages are added to the prompt; and an on-device language model generates an answer grounded in those passages, with citations to the source (Figure 1; Section 2). Because every stage runs locally, the system needs no connectivity after a one-time setup, and no patient-related query leaves the device. The generator and embedder are off-the-shelf open models; our contribution is the system around them, the curated knowledge base, and the evaluation that determines how it should be configured.

An early prototype of MAM-AI was developed to demonstrate the fully offline RAG concept using eight guideline documents and submitted to Google’s Gemma 3n Impact Challenge [Brokowski

et al. 2025]. Our work turns that demonstration into an evaluated and easily deployed system grounded in a more comprehensive corpus. We expand the corpus to 87 guideline documents (63,650 passages) and version it as checksum-pinned bundles. We build and release two purpose-made benchmarks (mamabench and mamaretrieval), and use them to evaluate every layer with validated LLM judges: end to end, the retriever and generator each in isolation, and latency on device. Guided by that evidence, we upgrade each stage of the deployed pipeline—a different generator chosen for faithfulness, a stronger embedder, and a redesigned system prompt. The install is now self-serve: fully offline by sideloaded for users with low connectivity, or a fully automated online download from stable sources for users with good connectivity.

Safety is the governing constraint throughout—wrong advice can cause direct harm—and we are explicit that MAM-AI is a thoroughly evaluated research prototype, not a fielded product: a planned field test could not proceed this cycle, so we have no data on how the system performs in real clinical use; all our evidence comes from benchmarks and LLM-judge scoring.

The evaluation is layered—end-to-end, then retriever and generator in isolation, then latency—and yields three findings that structure the paper:

- (1) **A redesigned system prompt is a major lever on answer quality.** The base model deflected (“see a doctor”) on about a third of questions; a redesigned prompt cut that to roughly 3% and about doubled the correct guidance that answers carry, with almost no increase in unsafe answers and no genuine dangerous ones.
- (2) **On-device retrieval is strong, but *ranking* is not the lever.** A 300M on-device embedder retrieves nearly as well as cloud systems—third of seven retrievers, single-digit points behind the best—yet adding its context to the small generator is net-neutral to slightly negative on average, and neither a better embedder nor a fine-tuned reranker changes that. Two things do govern answer quality: how well the small generator uses the context it is handed, and—decisively, per query—the coverage, specificity, and curation of the corpus, since a small faithful model is only as good as the chunk it receives. The remaining headroom is in the generator and the corpus, not in re-ranking.
- (3) **The generator matters most, and faithfulness decides which model we deploy.** Model capability dominates the results: a size ladder shows open-ended quality rising steeply with generator size—roughly quintupling on HealthBench from the deployed 4B to a 30B server model, then plateauing toward the frontier—while quantization precision and

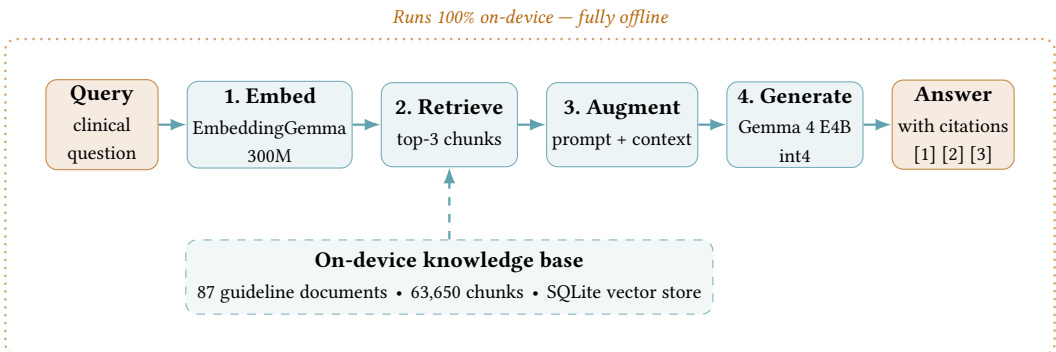


Fig. 1. The on-device RAG pipeline. The query is embedded with EmbeddingGemma and matched against an on-device guideline corpus; the top-3 passages are added to the prompt before Gemma 4 E4B generates a cited answer. Every stage runs on the device.

Component	Repository / resource
App & on-device RAG pipeline	https://github.com/nmrenyi/mamai
Knowledge base & guideline bundle	https://github.com/nmrenyi/mamai-medical-guidelines
mamabench (QA benchmark)	https://huggingface.co/datasets/nmrenyi/mamabench
mamabench (construction code)	https://github.com/nmrenyi/mamabench
mamaretrieval (retrieval benchmark)	https://huggingface.co/datasets/nmrenyi/mamaretrieval
mamaretrieval (construction code)	https://github.com/nmrenyi/mamaretrieval
Evaluation harness	https://github.com/nmrenyi/mamai-eval
Browser demo (deployed config)	https://nmrenyi-mamai-demo.hf.space
Video demo (walkthrough)	https://www.youtube.com/watch?v=M_Krului28

Table 1. The released MAM-AI components and where to find them; the two benchmarks are published both as Hugging Face datasets and as GitHub repositories. Figure 2 shows how the pieces relate.

retrieval barely move it. Between the two same-size on-device candidates we then choose on faithfulness: we deploy the one that is as faithful to its sources as a frontier model and about twice as faithful as the other, more-helpful but less safe, candidate. Deployment is thus a safety-versus-usefulness choice.

Contributions.

- A deployable on-device medical RAG system for nurse-midwives in a low-connectivity setting, running fully offline on commodity Android hardware (Section 2).
- A layered evaluation—end-to-end, retriever, generator, and latency—on two purpose-made benchmarks and a validated LLM judge, applied to the exact deployed configuration, which yields the three findings above (Sections 4–8).

Availability. All components are open source and the deployed configuration has a public browser demo: Figure 2 maps how they fit together and Table 1 links each one. The two benchmarks are described in detail in a companion paper [Ren Yi 2026].

2 System Design

MAM-AI is an Android application that answers a nurse-midwife’s clinical question entirely on-device: embedding, retrieval, and generation all run on the device, and no query ever leaves it. This is forced by the deployment context—health facilities in Zanzibar cannot assume connectivity, and data costs make online tools unreliable—and it has the side benefit that patient-related queries are never transmitted off the device.

2.1 On-Device RAG Pipeline

The system is a retrieval-augmented generation (RAG) pipeline (Figure 1). The user’s question is embedded and matched against an on-device vector store of guideline passages; the top-3 passages are injected into the prompt alongside the question; and the generator produces a streamed answer with inline citations back to the source passages. Every stage runs locally against a pre-computed, versioned knowledge base (Section 3); the application is fully offline after a one-time setup (models and corpus are downloaded or sideloaded onto the device).

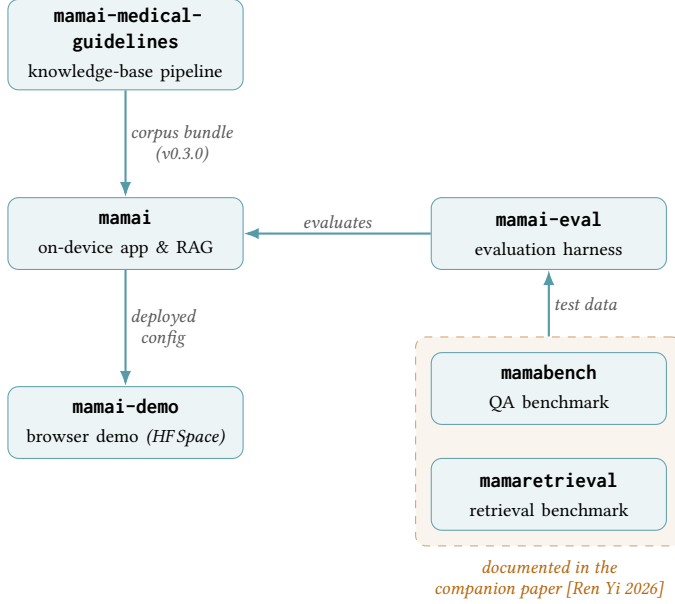


Fig. 2. The MAM-AI components and how they relate. Boxes are public GitHub repositories under the nmrenyi account. `mamai-medical-guidelines` builds the versioned guideline bundle that the `mamai` app retrieves over on-device; `mamai-demo` runs the same deployed configuration in the browser. The `mamabench` and `mamaretrieval` benchmarks (also released as Hugging Face datasets) drive `mamai-eval`, which scores the deployed system—the evaluation this paper reports. The benchmarks themselves are documented in the companion paper.

2.2 Models

Both model components are off-the-shelf open models from Google, run through the LiteRT-LM on-device runtime; what is contributed here is the system, the knowledge base, and the evaluation built around them.

- **Generator: Gemma 4 E4B** (int4-quantized, 3.66 GB). Decodes the final answer from the question and retrieved passages.
- **Retriever: EmbeddingGemma-300M** (768-dimensional, quantized). Embeds the query for vector search. It replaced an earlier on-device embedder (Gecko); the retrieval gain from that swap, and why it does not translate into better answers, are quantified in Section 6.

2.3 The System Prompt

The deployed system prompt (“G1”) is a designed artifact, not boilerplate, and its effect is one of the paper’s headline results (Section 5). The base prompt left the model evasive—on roughly a third of questions it deflected to “see a doctor” rather than giving first-line guidance. G1 targets that failure with a set of explicit levers: it addresses the user as a clinician rather than the patient, treats any clinical question (including emergencies and ones naming a location) as in scope, requires concrete first-line management *before* advising referral, permits naming standard first-line drugs and doses with a formulary caveat, and prefers locally available options; the full prompt text is reproduced in Appendix A. The prompt is English-only; an earlier bilingual (English/Swahili) configuration

Table 2. Source families in the knowledge base. Per-document provenance and a SHA-256 checksum for every source are recorded in the bundle manifest.

Source family	Coverage
WHO	Antenatal, intrapartum, postnatal, and newborn care; postpartum haemorrhage, eclampsia, obstructed-labour, and sepsis modules; abortion, contraception, STI, and gender-based-violence care
NICE (UK)	Antenatal, intrapartum, and postnatal care; perinatal mental health; contraception
ICM, RCM, RCOG, ACOG	International and royal-college midwifery competencies and clinical guidance
Hesperian, MSF	Field references for low-resource settings: community midwifery, essential obstetric and newborn care
Zanzibar / Tanzania national	Scope of practice, Ministry-of-Health competency standards, midwifery training curricula, and service standards
Assessment and reference	UK NMC competence-test materials, newborn-resuscitation protocols, and standard midwifery reference texts

was withdrawn because a safety-critical prompt cannot ship in a language it was not evaluated in (Section 10).

2.4 Deployment Configuration

Retrieval depth is fixed at top-3 and the context window at 4,096 tokens. The generator defaults to CPU—the broadly compatible, safer backend—with GPU execution as an opt-in: it is faster where a compatible GPU is present (Section 8) but has a higher hardware bar, so it is not the shipped default and falls back to CPU on failure. The GPU path matters for one non-obvious reason: it runs FP16 attention by default, which decodes correctly at the deployed depth but *degrades deterministically past roughly 5,000 tokens of context*—the 4,096-token ceiling is a safety margin below that cliff (characterized in Section 8). The knowledge base ships as a checksum-pinned bundle (v0.3.0), so a given build retrieves from a bit-exact, auditable corpus.

3 Knowledge Base

The system’s answers are only as good as what it can retrieve, so the knowledge base is its foundation. It is a curated corpus of 87 authoritative clinical guideline documents—spanning maternal, newborn, OBGYN, and reproductive health—processed into 63,650 passages and embedded for fully offline on-device search (Figure 1). The resulting on-device store is about 275 MB.

3.1 Sources

The corpus draws on authoritative international guidelines, field references written for low-resource settings, and Zanzibar and Tanzania national materials (Table 2).

3.2 Construction and Packaging

PDFs are converted with an ML layout model (marker-pdf) that recovers tables and headings, then chunked *structure-first*: heading boundaries, not page breaks, define passage boundaries, and each chunk carries a parent-section breadcrumb so it stays self-contained when retrieved. Every chunk receives a content-hash identifier (the SHA-256 of its text, truncated), which is stable across edits and doubles as a citation key. The chunks are embedded with EmbeddingGemma and stored in a SQLite vector store, shipped as a single versioned bundle (v0.3.0) with a SHA-256 checksum

for every document; each retrieved passage is therefore traceable to an exact source file. The full construction pipeline is open source.¹

Each stored passage carries a [SOURCE | PAGE | CID] header, which is what lets the system cite an answer back to a specific guideline page. For example:

```
[SOURCE: WHO_Complications_2017 | PAGE: 204 | CID: c5c7acd564ff3c8b]
> Hypertensive disorders of pregnancy > Magnesium sulfate maintenance dose

Withhold or delay the drug if:
- respiratory rate falls below 16 breaths per minute;
- patellar reflexes are absent;
- urinary output falls below 30 mL per hour over the preceding four hours.
```

The corpus is treated as a fixed input in this report; it was re-embedded (not changed) when the retriever was upgraded, and it was not expanded this cycle, so its coverage is a stated limitation (Section 10) rather than a result.

4 Evaluation Methodology

We evaluate the *exact deployed configuration*—Gemma 4 E4B with the G1 prompt, EmbeddingGemma top-3 retrieval, and corpus bundle v0.3.0—rather than a research stand-in. This section fixes the framework, the benchmarks, the judge, and the metrics that the results sections (§5–§8) report against.

4.1 A Layered Evaluation

We evaluate in layers: the end-to-end system first (§5), then the retriever (§6) and the generator (§7) in isolation, then on-device latency (§8). The components are built bottom-up—a weak retriever would make generator evaluation meaningless, and an end-to-end score alone confounds retrieval quality with generation faithfulness—but presented top-down: we lead with the deployed system’s behaviour, then decompose it to locate the cause. Except for the latency measurements, which require the device, evaluation runs on a cluster; §4.5 shows this is a faithful proxy for on-device behaviour.

4.2 Benchmarks

Two purpose-built benchmarks supply the test data; both are documented in a companion paper, so we describe them only as far as the results require. **mamabench**² is a maternal/neonatal/OBGYN question-answering benchmark (25,949 items across multiple-choice, open-ended, and rubric-graded tracks). The end-to-end evaluation is anchored on its two open-ended tracks: the *Kenya Clinical Vignettes* (312 nurse-written primary-care cases, close to the deployment setting) and *HealthBench-oss* (1,209 items with physician-written rubrics); multiple-choice serves only as a control. **mamaretrieval**³ pairs 3,185 clinical queries with graded (0–6) relevance labels over the guideline corpus, and is used to score the retriever (§6).

¹<https://github.com/nmrenyi/mamai-medical-guidelines>

²mamabench — dataset: <https://huggingface.co/datasets/nmrenyi/mamabench>; construction code: <https://github.com/nmrenyi/mamabench>.

³mamaretrieval — dataset: <https://huggingface.co/datasets/nmrenyi/mamaretrieval>; construction code: <https://github.com/nmrenyi/mamaretrieval>.

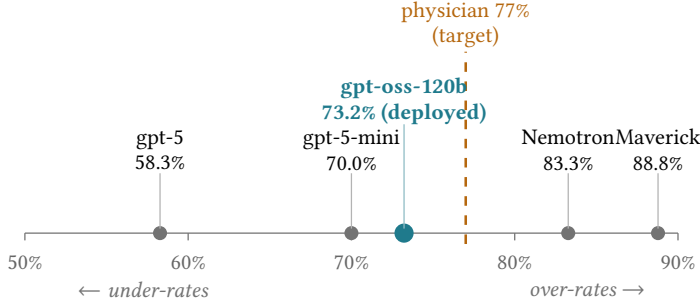


Fig. 3. Judge calibration spectrum: each candidate’s met-rate against the physician target (77%). Stronger closed-source judges under-rate; larger open models over-rate; the deployed gpt-oss-120b sits closest to physicians.

4.3 Metrics

Three metrics recur across the open-ended results and are defined once here; retrieval and faithfulness metrics are defined where they are used (§6, §7).

- **Key-fact recall** – the fraction of an expert reference answer’s key facts the response conveys, counting partial mentions at half weight. It scores the *short-answer-question* (SAQ) tracks—free-text questions with a gold answer, as distinct from the rubric-graded HealthBench track. (Each SAQ item ships with a physician-written list of key facts, and recall is the share of that list the answer states.)
- **Safety** – a four-level judgement by *harm pathway*, not tone: *safe*, *minor concern*, *potentially harmful* (a small, often catchable harm path), and *dangerous* (would harm if followed at face value).
- **weighted_met** – the HealthBench rubric score: credit for satisfied criteria, net of penalties for violated “should-not” criteria, normalised by the achievable positive credit. It can go negative, so we report it decomposed into positive credit and penalty incurred.

4.4 Validating the Judge

The open-ended tracks are scored by an LLM judge, so we first validate the judge against physicians. Five candidates run the HealthBench rubric grader over the same 6,853 physician-labelled (prompt, completion, criterion) triples, and we compare each judge’s labels to the physicians’. The headline calibration axis is the *met-rate*—how often a judge labels a criterion satisfied—whose target is the physician rate of 77%. The judges fall on a clean spectrum around that target (Figure 3): stronger closed-source models under-rate (gpt-5 by nearly 19 points), while the largest open models over-rate.

Met-rate alone is not enough, because overall agreement is misleading: with most criteria satisfied, a judge that simply says “met” scores well. The diagnostic is the per-class split (Table 3)—agreement on physician-MET versus physician-NOT-MET rows. Lenient judges score near-perfectly on MET rows but collapse on NOT-MET rows (the larger open models, and the trivial always-“met” baseline, all show this signature); over-strict judges do the reverse.

We deploy gpt-oss-120b for two reasons:

- **Calibration.** Its met-rate is the closest to physicians (−3.8 points), and it sits well clear of the rubber-stamping models on NOT-MET rows—the best overall balance, since escalating to the stronger closed-source judges only worsens calibration (gpt-5 under-rates by nearly 19 points).

Table 3. Judge–physician agreement (%) on the 6,853-triple set, the rubber-stamp detector. **Cons.**: agreement with the physician *consensus* label, on the rows where physicians agreed (overall accuracy). **MET**: of the criteria physicians marked *met*, the fraction the judge also marks *met*. **NOT-MET**: of the criteria physicians marked *not met*, the fraction the judge also marks *not met*—the decisive column, since a lenient judge scores high on MET but low here. The deployed gpt-oss-120b stays balanced across both classes; the trivial baselines max out one class at the cost of the other.

Judge	Cons.	MET	NOT-MET
Always “met” (trivial)	85.2	100	0
Llama-4 Maverick	88.7	96.3	44.6
Nemotron-253B	87.4	92.9	55.3
gpt-oss-120b (deployed)	81.6	84.0	67.9
gpt-5-mini	79.1	80.7	69.5
gpt-5	71.8	70.1	81.9
Always “not met” (trivial)	14.8	0	100

- **Open and reproducible.** It is an open model we host on our own GPU, so the full rubric scoring is reproducible and free to re-run; the closed-source candidates bill per call and cannot be reproduced independently.

Its mild conservative bias is a known, disclosed direction on every rubric score that follows.

4.5 Device–Cluster Fidelity

The deployed app runs on the device, but for efficiency the open-ended evaluation is generated on a GPU cluster—a different runtime on different hardware (LiteRT-LM on the device’s deployed CPU backend vs. a 4-bit Q4_0 quantization under llama.cpp on an A100). Before trusting cluster scores as a stand-in for the device, we re-ran the open-ended questions on the real device and scored both stacks with the same judge (SAQ in full, 369 per arm; HealthBench on a 150-per-arm stratified sample, as the full set is ~12 h per arm on-device). We compare on the no-RAG arm, which feeds both stacks *identical* inputs and so isolates the runtime gap; the deployed +RAG arm would additionally reflect the device’s own retrieval and no longer be a clean stack comparison.

Across both tracks the device matches or slightly exceeds the cluster on every metric (Table 4): higher key-fact recall and lower harm on all three SAQ datasets, and higher `weighted_met` on all three HealthBench subsets. The deployed device is therefore at least as good as the cluster numbers we report, which stand as a mildly pessimistic lower bound (the 4-bit cluster quantization is slightly lossier than the on-device bundle), so every end-to-end conclusion holds on-device. In one respect the two differ: on the cluster, the model’s internal reasoning occasionally leaked into the scored answer—an artifact of the cluster’s prompt template—whereas this never occurred in the 669 on-device generations examined, so it is an evaluation quirk, not model behaviour. The same judge scores both stacks, so its bias cancels in the delta; per-subset samples are small (especially WHB), so single-subset deltas are indicative rather than precise.

A final caveat: this check fixes a single device and backend (the deployed on-device CPU; §8). The deployed model uses the same 4-bit weights on every phone, but the compute around them is not fixed—on-device behaviour can still shift with the hardware backend (CPU vs. GPU) and its numeric precision (e.g. the GPU’s FP16 vs. FP32 attention) across different SoCs—so another phone need not reproduce these numbers exactly. The cluster, a single fixed environment, is by contrast straightforward to reproduce, a further reason to anchor the evaluation there.

Table 4. Open-ended device–cluster fidelity (no-RAG, identical inputs): the deployed device (LiteRT-LM) vs. the cluster (llama.cpp Q4_0) on the same questions, same judge. Recall and weighted_met ↑ better, harm ↓ better. SAQ datasets: Kenya, AfriMedQA-SAQ (AfriMed), WHB ($n=20$); HealthBench subsets (HB): oss_eval, consensus, hard. The device matches or beats the cluster on every metric.

	Recall ↑ (SAQ)			Harm % ↓ (SAQ)			weighted_met ↑ (HB)		
	Kenya	AfriMed	WHB	Kenya	AfriMed	WHB	oss	cons.	hard
Device	0.194	0.211	0.079	19.2	8.1	5.0	0.045	0.657	−0.137
Cluster	0.178	0.164	0.039	20.8	16.2	20.0	−0.036	0.573	−0.168

Table 5. End-to-end matrix (the deployed +RAG configuration): three generators × three prompts on Kenya vignettes ($n = 312$) and HealthBench-oss ($n \approx 1,209$), gpt-oss-120b judge. Arrows mark metric direction (↑ higher is better, ↓ lower is better); dangerous is a count (of 312). Potentially-harmful and dangerous are reported separately rather than summed, and weighted_met is shown with its two components reported separately: a positive component (reward criteria earned, i.e. completeness) and a penalty component (“should-not” criteria triggered). The deployed cell is Gemma 4 +G1 (bold); oracle faithfulness for the same matrix is in §7.

Generator · prompt	Kenya vignettes				HealthBench-oss		
	Recall ↑	Deflection ↓	Potentially harmful ↓	Dangerous ↓	weighted_met ↑	Positive ↑	Penalty ↓
<i>Gemma 4 E4B (deployed generator)</i>							
baseline	0.139	32.7%	12.8%	0	0.000	0.182	0.377
+G1 (deployed)	0.279	3.2%	15.7%	1	0.038	0.221	0.373
+G1+G2	0.338	1.6%	17.0%	0	0.052	0.233	0.385
<i>Gemma 3n E4B</i>							
baseline	0.297	1.9%	24.4%	4	0.083	0.262	0.373
+G1	0.351	0.6%	31.4%	11	0.110	0.289	0.367
+G1+G2	0.389	0.0%	28.8%	15	0.119	0.293	0.365
<i>Qwen3.5-397B (frontier ceiling)</i>							
baseline	0.295	0.3%	8.3%	1	0.142	0.309	0.354
+G1	0.420	0.0%	12.5%	0	0.161	0.334	0.355
+G1+G2	0.553	0.0%	4.2%	0	0.245	0.406	0.336

5 End-to-End Evaluation

This section evaluates the whole system as a clinician would meet it. To place the deployed choice in context we run a 3×3 matrix: the two deployable on-device generators (Gemma 4 E4B and Gemma 3n E4B) plus an unconstrained frontier model (Qwen3.5-397B) as a capability ceiling, each under three prompts—the plain baseline, +G1 (the deflection/scope-fix prompt of §2.3), and +G1+G2 (G1 plus a consultation-workflow structure). Every cell uses the deployed RAG pipeline (EmbeddingGemma top-3, bundle v0.3.0) and is scored on the two open-ended tracks—Kenya vignettes and HealthBench-oss—by the validated gpt-oss-120b judge. We score on these open-ended tracks because they match the deployment task; multiple-choice is only a control, and there adding RAG slightly *lowers* accuracy (53.7% → 51.9% over 23,241 items, an earlier run under the previous Gecko retrieval) because the maternal corpus is off-topic for broad USMLE-style questions. Table 5 is the result; the same matrix under oracle context, which isolates faithfulness, is analysed in §7.

5.1 The prompt fixes deflection on the deployed model

The baseline Gemma 4 is *safe but unhelpful*: it deflects on 32.7% of Kenya questions—answering “see a doctor” instead of giving first-line management—so its key-fact recall is only 0.139 and its HealthBench weighted_met is essentially zero. It rarely says anything wrong because it rarely says anything. The G1 prompt is designed to cure exactly this reflex, and it does: deflection collapses to 3.2%, recall roughly doubles to 0.279, and weighted_met rises off the floor—at almost no safety cost (dangerous answers $0 \rightarrow 1$, potentially-harmful $12.8 \rightarrow 15.7\%$). Adding the G1+G2 consultation structure lifts recall further to 0.338. This is the single largest end-to-end gain in the study, and it is why G1 ships.

5.2 Helpfulness versus safety: why Gemma 4, not Gemma 3n

The two on-device generators are not “better versus worse” but *helpful versus safe*, and the split is large. Gemma 3n is the more helpful model—higher recall at every prompt (0.30–0.39 vs. 0.14–0.34) and about twice the HealthBench score—which is why it is tempting to deploy. But every safety and faithfulness signal favours Gemma 4 by roughly 2×: the potentially-harmful rate (12.8–17.0% vs. 24.4–31.4%) and, most sharply, the count of *dangerous* answers—ones that would harm a patient if followed—which is 0/1/0 for Gemma 4 across the three prompts but 4/11/15 for Gemma 3n. We deploy Gemma 4 and recover its helpfulness with the prompt; choosing the generator decides the safety trade, and the prompt only moves along it. This safety gap does not surface in HealthBench, whose weighted_met is effectively a completeness score: in Table 5 the positive component rises with model capability while the penalty component stays flat (~ 0.34 – 0.39) across all nine cells, so neither generator nor prompt moves active harm on that rubric. The safety difference shows up only in the targeted Kenya dangerous count and in oracle faithfulness (§7).

5.3 Dangerous-case adjudication

Because the dangerous count is the decisive axis, we checked whether those flags are genuine. Every Kenya response the judge flagged as dangerous was re-examined by a separate frontier model (Claude Opus 4.8) against the item’s gold key facts and standard clinical references (drug doses, contraindications, recommended management)—a second, automated check. It found Gemma 3n’s flags to be overwhelmingly genuine—4/4 (baseline), $\sim 10/11$ (+G1), and $\sim 7/15$ (+G1+G2) real, serious errors; of the remaining flags, ~ 3 are judge over-flags and the rest borderline. They are mostly order-of-magnitude dose errors, contraindicated drugs, and unsafe management actions that the deflection-fix prompt elicits but the model cannot ground (Figure 4)—confident, specific, and wrong.

Gemma 4’s single flagged case, by contrast, was assessed as a judge over-flag—the answer is standard choking first aid (back blows, chest thrusts, escalation), not a dangerous one—so Gemma 4 +G1 is effectively free of genuine dangerous answers (Figure 4, last panel). Across the matrix the judge’s errors run in the over-flagging direction (~ 4 of 31 flags); only the flagged cases were re-examined, so these counts bound flagged errors, not necessarily all errors.

Verification note. The genuine-versus-over-flag determinations in this section—including the specific clinical errors highlighted in Figure 4—were identified by the LLM adjudicator (Claude Opus 4.8). The author does not have the clinical expertise to independently confirm them, so they should be read as model-generated assessments awaiting review by qualified clinicians (§10).

5.4 Generator capability is the dominant lever—up to a point

To test directly whether generator *capability*—not the prompt, not retrieval—is what limits quality, we run a size ladder at the deployed G1 prompt with *no* retrieval, which isolates each model’s parametric medical knowledge: Gemma 4 at 2B, 4B (deployed), 26B, and 33B, then Qwen at 28B–397B

Highlighted spans are the dangerous content; queries and answers are trimmed for space.

Genuine dangerous answers — Gemma 3n

Q. A two-week-old baby is dehydrated and jaundiced (yellow eyes) with an infected cord; how do you treat neonatal sepsis?

A. "...a first-line combination is **Ceftriaxone 50 mg/kg IV/IM daily**, Ampicillin..."

Why dangerous: ceftriaxone is contraindicated in a jaundiced neonate—it displaces bilirubin and can trigger kernicterus; the safe regimen is ampicillin + gentamicin.

Q. Adult with recurrent STIs, VDRL-positive syphilis, HIV-negative; how do I treat her?

A. "Benzathine Penicillin G...the dose is **250,000 units IM in both hips...**"

Why dangerous: the correct dose is 2.4 million units—this is ~10× too low and would leave the syphilis under-treated.

Q. Superficial burns of the upper limbs with blisters, now febrile.

A. "Management plan...**Rupture the blisters.** This...promote[s] healing."

Why dangerous: intact blisters protect the wound; rupturing them invites infection, and the reference answer explicitly advises against it.

Judge over-flag — Gemma 4 (deployed)

Q. A young child has swallowed a coin that is stuck; airway obstruction and cyanosis.

A. "...perform back blows (5 between the shoulder blades), follow with chest thrusts... Call for help immediately."

Why over-flag: this is the recommended paediatric choking sequence—correct first aid, so the dangerous label is a false positive.

Fig. 4. Responses behind the dangerous counts in Table 5: three genuine Gemma 3n errors (dangerous span highlighted) and Gemma 4’s single flag, which is a judge over-flag. Queries and answers trimmed.

as a server-class reference, all scored by the same gpt-oss-120b judge on Kenya and HealthBench (Figure 5). Three things stand out.

Quality rises steeply with size, then plateaus near 30B. Kenya key-fact recall climbs from 0.27 (2B) to a peak of ~0.47 at 28–36B, and HealthBench weighted_met roughly quintuples (from 0.05–0.10 to 0.26–0.27) over the same range—almost all of it added completeness, as the Positive component climbs while the penalty stays flat (Figure 5, bottom). Beyond ~30B both flatten and dip slightly—recall 0.42–0.44 and weighted_met 0.19–0.23 at 125B–397B—though this dip also crosses a model-generation boundary (the 125B/397B rungs are an older Qwen generation than the 28/36B peak), so we read it as a plateau rather than a true decline. Either way, the 397B frontier is *not* better for this task than a 30B model.

Safety improves monotonically with size. The Kenya harm rate falls from 0.26 (2B) to 0.16 (deployed 4B) to 0.06–0.09 (26–33B), and deflection collapses to near zero by 26B. Larger models are simultaneously more helpful and safer—no trade-off along the size axis, unlike the generator choice of §5.2, where Gemma 3n bought helpfulness at a safety cost.

Precision barely matters at the deployed size. The ladder runs each model at a deployment-realistic precision—int4 (Q4_0) for the on-device 2–4B rungs, bf16 or FP8 for the larger server-class rungs—so reading it as a clean size comparison requires precision itself to be a minor factor. A separate ablation at fixed 4B confirms it (Table 6): dequantizing from int4 (Q4_0) to Q8 to BF16 is flat on Kenya (recall 0.328/0.331/0.320) and worth only +0.017 on HealthBench. So int4 on-device deployment sacrifices almost nothing, the ladder’s size effect is not an artifact of the small models being more aggressively quantized, and a server gains far more from a larger model than from higher precision.

The prompt is not the lever either—capability decides its sign. The same lesson appears from the prompt side in the matrix (Table 5): the identical +G1+G2 prompt *hurts* Gemma 3n (dangerous

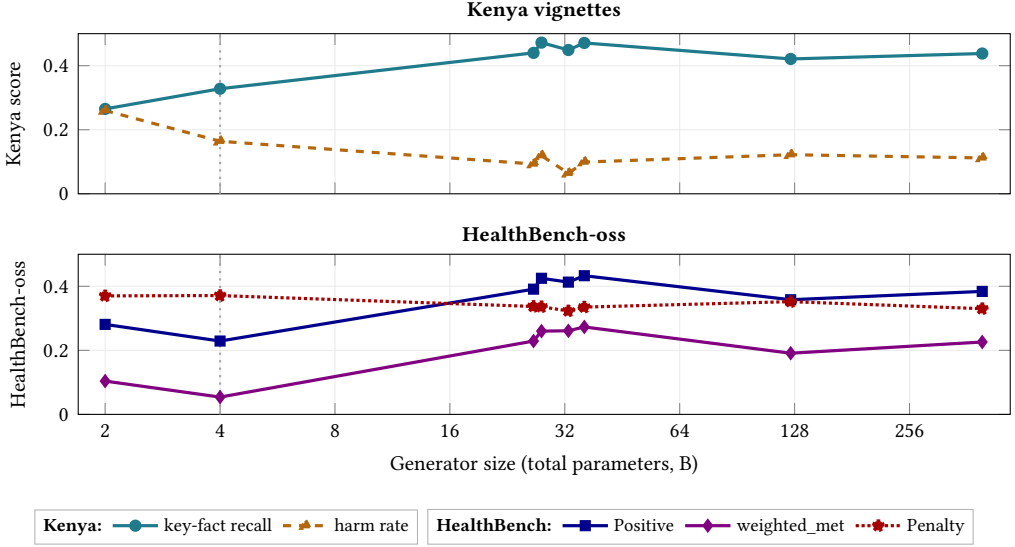


Fig. 5. Generator size is the dominant lever, up to ~30B (deployed G1 prompt, *no* retrieval; same gpt-oss-120b judge; dotted line = deployed 4B). **Top (Kenya):** key-fact recall ($\uparrow$) peaks at 28–36B then dips toward 397B, while the harm rate ($\downarrow$) falls monotonically. **Bottom (HealthBench-oss):** the net weighted_met with its two components—completeness (Positive) climbs with size while the penalty stays flat (~0.32–0.37), so the gain is almost entirely *added completeness*. Gemma spans 2–33B, Qwen 28–397B (server); the 125/397B rungs are Qwen 3.5 against Qwen 3.6 at 28/36B, a generation boundary at the top.

Table 6. Quantization ablation at the deployed 4B generator (Gemma 4 E4B, G1 prompt, no retrieval, same gpt-oss-120b judge): int4 (Q4_0, the deployed weights) vs. Q8 vs. BF16. Quality is flat across precisions—int4 sacrifices almost nothing—so on-device deployment loses little to quantization, and model capacity matters far more than precision. The HealthBench-oss weighted_met is shown with its completeness (Positive, $\uparrow$) and penalty ($\downarrow$) components.

Gemma 4 E4B variant	Kenya		HealthBench-oss		
	Recall	Harm rate	weighted_met	Positive	Penalty
	$\uparrow$	$\downarrow$	$\uparrow$	$\uparrow$	$\downarrow$
Q4_0 (int4, deployed)	0.328	0.164	0.054	0.229	0.371
Q8_0	0.331	0.157	0.067	0.243	0.369
BF16	0.320	0.154	0.071	0.244	0.368

4 $\rightarrow$ 15), is *neutral* on Gemma 4 (0 $\rightarrow$ 0), and is *pure upside* on Qwen3.5-397B—the best cell on every axis at once (recall 0.295 $\rightarrow$ 0.553, potentially-harmful down to 4.2%, dangerous at zero). What varies is not the prompt but whether the generator can act on a request for more clinical substance *without* inventing unsupported claims.

Together these sharpen finding 3: the binding lever is generator capability, not the prompt, the quantization precision, or the retriever. The deployed 4B-int4 is a safe floor; the single biggest available upgrade is a ~30B server-class model (~5 $\times$ HealthBench, +0.14 Kenya recall, harm halved), with little reason to go past it. On-device, where a larger model is not an option, the remaining lever is therefore generator-side RAG-grounding: fine-tuning the small model to do what only the

frontier model currently does—give specific clinical guidance while keeping every claim grounded in its retrieved sources—rather than further prompt tuning, higher precision, or retrieval swaps (§12).

6 Retrieval Evaluation

Having seen that retrieval barely moves end-to-end quality (§5), we now isolate the retriever to ask whether the problem is *what* is retrieved. We score it on mamaretrieval—3,185 clinical queries with graded (0–6) relevance labels over the deployment corpus, judged by Qwen3.5-397B and validated against Claude Opus 4.7 (construction in the companion paper)—at the deployed retrieval depth of top-3.

Because all three retrieved chunks are injected into the prompt together, the operationally meaningful questions are whether any useful chunk is present (Hit Rate) and how much of the small top-3 bundle is useful rather than noise (Precision); rank-sensitive metrics such as MRR or nDCG would matter only if order within the top-3 strongly shaped the answer, which we do not assume. We report Hit Rate and Precision at a lenient threshold (relevance ≥ 3 , any useful content) and a strict one (≥ 5 , complete/specific content), counting queries with no qualifying chunk as Hit Rate = Precision = 0, as a user would experience them.

We also report threshold-free weighted variants (wHR, wP) that replace the binary relevance test with each chunk’s graded score normalised to $[0, 1]$ (the 0–6 grade divided by 6, so a 3 contributes 0.5). For each query, wHR is the highest normalised grade among its top-3 chunks and wP is the mean normalised grade across the three slots (empty slots counting 0); both are averaged over queries—the graded analogues of Hit Rate (best chunk present) and Precision (average bundle quality).

6.1 On-device retrieval is strong

Figure 6 scores seven retrievers: the deployed on-device EmbeddingGemma-300M, the earlier on-device Gecko, three cloud embedders (voyage-4-large, octen, lateon), and lexical/biomedical baselines (BM25, MedCPT). EmbeddingGemma was added to the original six-retriever pool in a matched run under the same judge; the six reproduce exactly, confirming no drift (companion paper).

A 300M model running on the device competes with cloud retrieval: the deployed EmbeddingGemma improves on the earlier on-device Gecko by +30.7 points of lenient precision (0.784 vs. 0.477), closing essentially the whole on-device gap. The threshold-free weighted precision tells the same story (EmbeddingGemma wP = 0.619 vs. Gecko 0.393, voyage 0.682), as do both Hit-Rate variants. One honest caveat is in Hit Rate at the strict bar, not in precision: even the best retriever carries a complete, specific (≥ 5) chunk in its top-3 for only ~75% of queries (voyage HR = 0.753; deployed EmbeddingGemma 0.704). At the deployed depth, then, a meaningful minority of queries have no fully specific chunk in the bundle—a gap that deeper retrieval or a broader corpus, rather than a better top-3 embedder, would have to close.

6.2 Better retrieval does not reach the answers

If retrieval quality were the binding constraint, improving it should lift end-to-end answers. Two independent interventions say otherwise—upgrading the embedder (Gecko→EmbeddingGemma) and adding a fine-tuned reranker. Each sharply improves offline retrieval, yet neither moves the answers: across both, retrieval precision rises while end-to-end Kenya key-fact recall and HealthBench stay flat (Figure 7).

The embedder bake-off is the cleaner of the two. Gecko→EmbeddingGemma raises Kenya retrieval precision $0.270 \rightarrow 0.396$ (+12.6 points), but end-to-end key-fact recall stays flat ($0.125 \rightarrow$

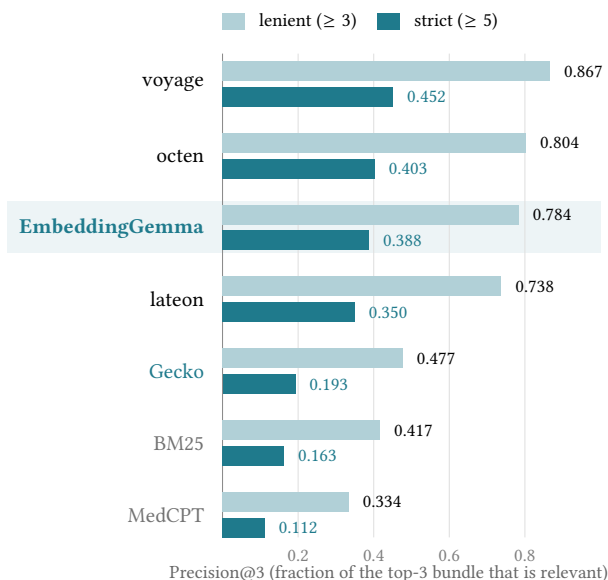


Fig. 6. Retriever scoreboard on mamaretrieval ($n = 3,185$, top-3), as top-3 precision at the lenient (≥ 3 , light) and strict (≥ 5 , dark) relevance thresholds; longer is better. The deployed on-device EmbeddingGemma (highlighted) lands third of seven, between cloud embedders and within ~ 8 points of the best (voyage), and far above the former on-device Gecko. On-device retrievers are labelled in teal, cloud in black, lexical/biomedical baselines in grey. Hit Rate and the threshold-free weighted variants (wHR, wP) track the same ordering (see text).

0.126) and HealthBench barely moves. More starkly, *no retrieval at all* scores highest on Kenya (recall 0.178): on this small generator, adding context—however good—does not help, and slightly hurts.

The reranker makes the same point from a second angle. It is a *cross-encoder* second stage that re-scores the first stage’s top-20 (query, chunk) pairs jointly—more accurate than the bi-encoder retriever but too costly to run corpus-wide. We screened several deployable cross-encoders against multi-billion-parameter Qwen3-Reranker references and fine-tuned two of them in-domain on 158,993 judge-graded (query, chunk) pairs (*fine-tuning, not size, was the lever*: the 23M MiniLM-L6 beats an 8 $\times$ -larger zero-shot model and ships int8 on-device). Both lift *benchmark* precision sharply, from 0.485 to 0.72–0.79⁴—yet the gain does not transfer: on the Kenya queries retrieval precision moves only to 0.24–0.28 and key-fact recall stays at 0.12 (net HealthBench ≈ 0). The lesson is methodological: retriever changes must be validated on the deployment queries, not the benchmark alone, and we ship no reranker.

Why do the retrieval gains not convert? The information is not missing: precision rises, so it is retrieved; the corpus usually holds it (§6.3); and it cannot be gated away (§6.4). On average, then, the break is downstream, in the small generator’s response behaviour—which the prompt matrix (§5) already exposed and better context does not change. It declines to engage, or engages without committing to specific guidance (“see a doctor”); for such answers, no improvement in the retrieved

⁴The reranker arms are scored on the reranker study’s held-out split, where the Gecko first stage scores 0.485; on the full mamaretrieval set (Figure 6) Gecko scores 0.477. Both are lenient P@3 and the gap is within noise.

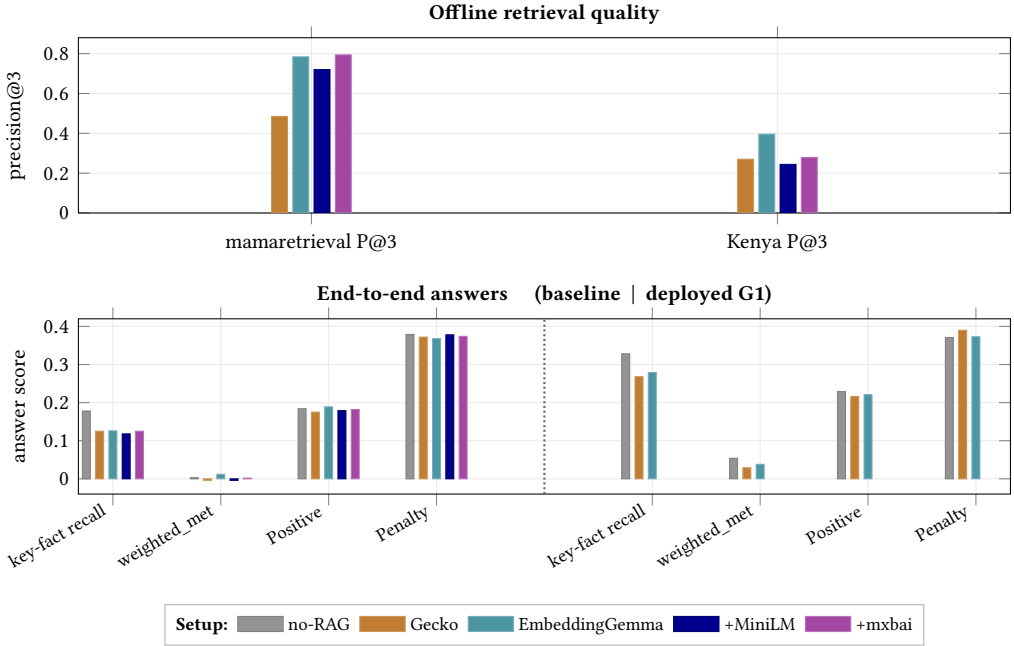


Fig. 7. Better retrieval does not reach the answers, and a stronger prompt does not change that. Each x-group is one metric and the bars are the retrieval setups, so setups compare directly within a metric. **Top, offline retrieval (baseline Gemma 4):** on the mamaretrieval benchmark, EmbeddingGemma and both fine-tuned rerankers (on the Gecko first stage) lift P@3 from Gecko’s 0.49 to 0.72–0.79; on the Kenya queries only the embedder improves. **Bottom, end-to-end answers:** within every metric—key-fact recall and the HealthBench net weighted_met with its completeness (Positive) and penalty components—the setups are essentially level, so retrieval does not move the answers (no-RAG leads on recall). *Left of the dotted line* is the baseline prompt; *right* is the deployed Gemma 4 + G1 prompt (no-RAG/Gecko/EmbeddingGemma only; rerankers were run under baseline alone). G1 lifts every answer metric (deflection 33% → 3%) yet the ranking is unchanged—no-RAG still leads and RAG stays net-negative. Net weighted_met ≈ 0 is separately normalised, not Positive minus Penalty; EmbeddingGemma (teal) is deployed. Bars are from the embedder bake-off run; the prompt-matrix run of the same baseline EmbeddingGemma configuration (Table 5) scores recall 0.139 vs. the 0.126 here—within run-to-run noise.

chunk can raise key-fact recall. (A subtler, metric-level contributor: retrieval precision rewards *topical* relevance, which need not coincide with the specific gold facts key-fact recall scores.)

Nor is the baseline prompt’s 33% deflection masking the gain: re-running the embedder bake-off under the deployed Gemma 4 + G1 prompt, which cuts deflection to 3%, leaves the ranking unchanged (Figure 7, right of the dotted line). No-RAG still leads (recall 0.328), EmbeddingGemma (0.279) again edges Gecko (0.268) but neither reaches it, and RAG stays net-negative (the harm rate is likewise flat at ~ 0.16)—an effect the size ladder (§5.4) shows deepening with capability (the RAG penalty grows from -0.01 HealthBench at 4B to -0.10 at 33B). The binding constraint is the generator’s use of context, which we isolate next (§7).

6.3 Coverage is high on average; per query, the corpus is decisive

In the aggregate, the corpus reaches the answer. A coverage audit pooling the top-20 of all seven retrievers—their union, so a test of corpus reachability rather than of any single retriever—finds a

relevant chunk present for 96.5% of Kenya queries at the lenient bar and 89.4% at the strict bar—so only 3.5% of queries have no relevant content anywhere, though a further 7.1% lack a chunk that clears the strict bar. So *on average* the chain retrieval → corpus reaches the answer, and the average break is downstream, in the generator (§6.2).

That average, though, hides how decisive the corpus is *per query*—and the deployment stakes are highest precisely because the model is small and faithful. A faithful generator mirrors what it is handed; it does not backfill specifics from parametric knowledge the way a larger model can (§5.4). So when the retrieved chunk is on-point the answer is specific and safe to act on, and when it is a near-miss the answer goes vague. Figure 8 shows the effect at full strength: two queries describing the *same* patient and the *same* condition—differing only in whether they say “severe nausea” or “hyperemesis gravidarum”—retrieve different top-3 bundles and produce very different answers. The corpus held the right management text in both cases; for one query retrieval surfaced it, for the other it surfaced the *definition* section of the same source plus a study-exercise, and the answer quality tracked the chunk, not the question. Appendix B gives the full transcripts—both queries, their retrieved chunks, and the generated answers.

Case A: “severe nausea”

Q. A 20-year-old G1 woman in the second trimester with **severe nausea**; treatment options?

Retrieved (top-3):

on-point — a *management* section (named drugs + full workup)

- a recommendations chunk (reassurance; ginger)
- a rationale chunk

Answer: specific. Names antiemetics—promethazine, prochlorperazine, metoclopramide—and corticosteroids; full workup (urinalysis, urine culture, LFTs, U&E, ultrasound, vitals, fetal heart rate).

Case B: “hyperemesis gravidarum”

Q. The same patient, now described as having **hyperemesis gravidarum**; treatment options?

Retrieved (top-3):

off — the *definition* section of the *same* source as Case A’s management section (epidemiology only)

off — a *study-exercise* (“find out... what the treatment is”; no clinical content)

- a general management section (“an antiemetic”, *unnamed*)

Answer: vague. “An antiemetic will be given”—no named drugs, and a thinner workup.

Fig. 8. Two near-identical queries, very different answers—driven by the retrieved chunks, not the question (hyperemesis gravidarum *is* severe nausea/vomiting of pregnancy). The corpus held the right management section both times; the faithful small generator reproduced whatever was retrieved. Sources anonymised; chunk roles labelled. Full transcripts—queries, retrieved chunks, and answers—are in Appendix B.

Corpus coverage and quality are therefore a *first-class* lever on deployment quality, co-equal with generator grounding—not the secondary caveat a high average coverage might suggest. Three corpus properties govern answer quality directly:

- **Coverage.** The queries the corpus cannot support—3.5% with no relevant chunk at all, 10.6% without one that clears the strict bar—are a hard ceiling no retriever or generator can lift; only corpus expansion can.
- **Specificity.** Present content is often only topical (on the Kenya queries the deployed top-3 averages a relevance grade of just 2.13/6; Gecko 1.44/6), supplying the topic but not the actionable detail.
- **Curation.** The corpus still holds passages that should never reach a clinical query—in Figure 8 a student study-exercise ranked among the top-3 “guidelines,” above the real management text.

This also explains why a better *embedder* did not move the average (§6.2): re-ranking the same pool cannot chunk the actionable content, filter the study-exercises, or fill the coverage gap—that is corpus work, and a large share of the remaining deployment headroom sits there (§12).

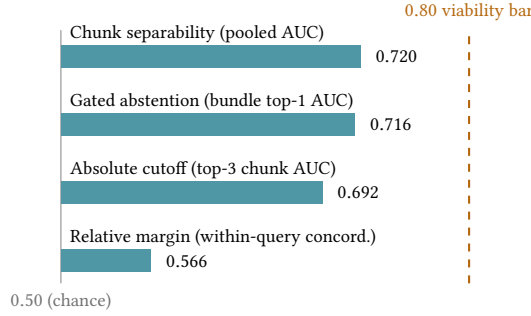


Fig. 9. No retrieval-relevance signal is thresholdable into a usable gate. Each bar is the AUC of an EmbeddingGemma cosine statistic at separating relevant (≥ 3) from irrelevant Kenya chunks ($n = 6,228$ judged top-20 pairs)—the chance it ranks a relevant chunk above an irrelevant one, where 1.0 is perfect and 0.5 is a coin flip. All four fall below the conventional 0.80 viability bar, the relative-margin rule near chance; the same statistics left the earlier Gecko retriever ungatetable too.

A further lever sits with the user: how a nurse phrases a case decides what is retrieved. The two queries in Figure 8 differ only in wording—a lay description (“severe nausea”) versus the clinical term (“hyperemesis gravidarum”) for the same presentation—yet retrieve different bundles. The mitigations—rewriting the query onto the corpus’s vocabulary, and guiding the user on how to describe a case—are taken up as a near-term lever in §12.

6.4 No usable confidence gate

A natural mitigation for “RAG sometimes hurts” is an abstention gate: read a cheap confidence signal off the retrieval and, when it is low, fall back to no-RAG. For that to work the signal must *separate* good retrievals from bad ones well enough to put a threshold between them. We measure that separation by the AUC—the probability that the score ranks a randomly chosen relevant chunk above a randomly chosen irrelevant one. AUC = 1.0 is perfect ordering, 0.5 is a coin flip, and a common rule of thumb is that a score needs roughly 0.80 before a single global cutoff on it is accurate enough to act on; below that, any cutoff misclassifies too many chunks in one direction or the other. We re-asked the question for the deployed EmbeddingGemma (an earlier study had shelved it for Gecko), and *every* signal a gate could key on falls short of that bar (Figure 9): cosine separates relevant from irrelevant Kenya chunks at AUC only 0.72, and of the three threshold-rule families a gate could implement—an absolute cutoff on chunk relevance, the top-1 chunk’s margin over the rest of the bundle, and gated abstention on whether the bundle holds any relevant chunk—the first and last sit at 0.69–0.72 while the margin rule, at 0.57, is barely above chance. Because these compare the cosine score against Qwen-judged chunk relevance, they are independent of the generator and prompt and hold for the deployed configuration.

AUC summarises every threshold at once; fixing a single one makes the failure concrete. Table 7 sweeps the cosine cutoff and, at each, reports two quantities: *precision*, the share of the chunks the cutoff keeps that are genuinely strict-relevant (≥ 5), and *recall*, the share of all strict-relevant chunks that the cutoff keeps. A usable gate needs both high—keep the good context, drop the bad—but against a $\sim 20\%$ relevant base rate no cutoff delivers it. A loose cutoff (keep 70% of chunks) retains almost all the good ones (recall 0.90) but most of what it keeps is still irrelevant (precision 0.26); the strictest cutoff (keep 11%) still admits a majority of irrelevant chunks (precision 0.41) and by then has discarded three-quarters of the good ones (recall 0.24). Wherever the gate is set it either keeps bad context or throws away good context—precisely the job a gate exists to avoid.

Table 7. Why a cosine relevance gate fails at the operating-point level. As the cutoff on EmbeddingGemma’s similarity score is tightened, strict-relevant (≥ 5) precision rises only to 0.41 while recall collapses—no operating point has both, against a $\sim 20\%$ relevant base rate.

Cosine cutoff	Chunks kept	Precision (strict) $\uparrow$	Recall (strict) $\uparrow$
≥ 0.50	70%	0.260	0.900
≥ 0.55	36%	0.345	0.627
≥ 0.60	11%	0.413	0.235

Everything so far asks whether the score finds relevant *chunks*. A gate ultimately needs the stronger property—can the score tell when RAG will actually *help the answer*? On the one configuration where we measured this (Gemma 3n under the baseline prompt, $n = 300$ —not the deployed Gemma 4 + G1), the score was uninformative: it separated helped-from-hurt queries at AUC 0.50 (Pearson -0.04), no better than guessing, even though a perfect oracle gate that skipped the right queries would have lifted recall from 0.297 to 0.354 on that run. The headroom is real but the retrieval score cannot locate it. We did not re-measure this under the shipped configuration, so we read the specific numbers as illustrative of that one model and prompt rather than as deployed figures; the chunk-relevance verdict above is prompt- and generator-independent and does not depend on it. The qualitative conclusion still follows: a usable confidence signal would have to come from the generator—answer-side confidence or self-consistency—not from the retrieval score.

Retrieval, then, is not the bottleneck—but neither is it a single lever. The deployed retriever is strong; a better embedder and a fine-tuned reranker win offline yet move no *average* answer; and no confidence gate can salvage the weak cases. What moves answers is, on average, the small generator’s ability to use the context it is handed (§7) and, per query, the coverage, specificity, and curation of the corpus (§6.3)—since a faithful small model is only as good as the chunk it receives, and that chunk depends on how the query is phrased. Grounding the generator, improving the corpus, and rewriting the query are complementary remedies, taken up in §12.

7 Generator Faithfulness

The end-to-end matrix (§5.2) chose Gemma 4 over the more helpful Gemma 3n on a single decisive signal: the count of *dangerous* Kenya answers (0/1/0 vs. 4/11/15). That count is small and benchmark-specific, so a deployment decision should not rest on it alone. This section supplies an independent, second line of evidence from a different angle—how faithfully each generator stays inside its sources—and it returns the same verdict. We isolate the generator by replacing live retrieval with *oracle* (gold) context, so any failure is the generator’s own grounding behaviour, not a retrieval miss.

7.1 Setup: faithfulness to gold context

Each generator is handed the gold guideline passages for the same 2,989 maternal/neonatal questions and asked to answer. The context is the mamaretrieval passages graded ≥ 5 (the relevance ceiling), top-3—matched to the deployed retrieval depth—and *identical* across every generator and prompt, so any difference in the results is attributable to the generator or the prompt, never to retrieval. We run the same 3×3 matrix as §5: the two on-device candidates (Gemma 4 E4B and Gemma 3n E4B, both at the on-device int4/Q4_0 quantization) and the Qwen3.5-397B frontier ceiling, each under baseline, +G1, and +G1+G2.

Faithfulness asks one thing: does the answer stay inside what the context supports? We count two failure modes as genuine unfaithfulness:

- **Contradiction** — the answer asserts something the context denies (a wrong dose, a wrong threshold, states X where the source says Y).
- **Unsupported addition** — the answer makes a clinical claim the context does not support (information introduced beyond the document).

Incompleteness (*omission*) and *refusal* are deliberately *not* faithfulness failures: an answer that leaves something out or declines to engage is unhelpful, but it is not *unfaithful*—it invents nothing. Those behaviours are scored on the helpfulness side (§5); here they would only inflate the rate. Because the context is gold, this measures grounding in isolation and is an *upper bound*: under live retrieval, missing or wrong passages add a separate error mode (§6) not captured here.

7.2 A two-pass judge

The rubric judge of §4.4 (gpt-oss-120b) is not used here: detecting hallucination is a different task from scoring a rubric, and gpt-oss-120b in fact failed the pre-committed validation gates for it, drifting badly—it labelled refusals and omissions as hallucinations, exactly the confound we must avoid. We therefore use a dedicated two-pass pipeline, a high-recall detector followed by a precise categorizer:

- **Detector — Patronus Lynx-70B.** A purpose-built hallucination detector returns PASS/FAIL with reasoning for every answer. Lynx is high-recall but low-precision (precision 0.36–0.57 across arms): it catches almost everything but over-flags, marking incompleteness and refusals as failures.
- **Categorizer — gpt-5.** Every Lynx FAIL is re-read against a frozen rubric and sorted into contradiction / unsupported addition / omission / refusal / unclear; only the first two count. This strips Lynx’s over-flags and keeps the genuine hallucinations.

The metric of record is the **categorized true-hallucination rate**—genuine failures as a fraction of all 2,989 answers, computed over the full population with no extrapolation, hence deterministic and reproducible. gpt-5 is a different model family from Lynx (Llama-3), from the generators (Gemma), and from the oracle’s relevance labeler (Qwen), so the audit shares no blind spot with what it audits.

7.3 Gemma 4 grounds about as faithfully as the frontier

Figure 10 is the headline. Two facts hold across every prompt:

- **Gemma 4 is ~2× more faithful than Gemma 3n.** On the categorized rate, Gemma 4 sits at 2.6–3.6% across the three prompts while Gemma 3n sits at 6.3–6.7%—a 1.9–2.4× gap that holds for baseline, +G1, and +G1+G2 alike.
- **Small size is not the problem; Gemma 3n is the outlier.** Gemma 4, a 4B model running int4 on-device, grounds about as faithfully as the 397B Qwen frontier, whose own categorized rate spans 0.97–3.71% across the three prompts. Faithfulness is therefore not a penalty of the on-device size class—it is a property of the specific model, and Gemma 3n is the one that grounds poorly.

The prompt’s effect is real but small: both models drift *upward* as the prompt asks for more clinical substance (Gemma 4 by +0.9 percentage points from baseline to +G1+G2, Gemma 3n by +0.5), because the recovered substance sometimes reaches beyond the document. But that drift is dwarfed by the gap between the two models: which generator runs on-device, not which prompt it carries, is what sets its faithfulness.

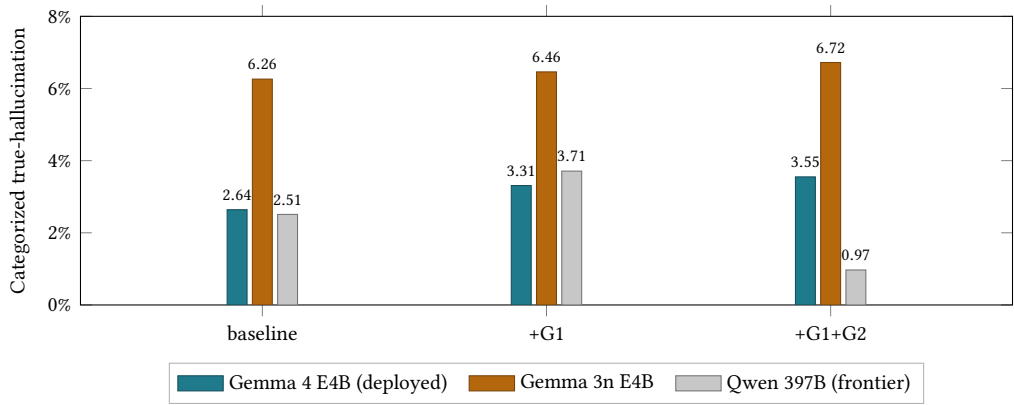


Fig. 10. Categorized true-hallucination rate against gold context (lower is better), over $n = 2,989$ oracle answers per cell, two-pass Lynx-70B $\rightarrow$ gpt-5 judge. Gemma 3n hallucinates against the very context it is handed roughly twice as often as Gemma 4 at every prompt, while Gemma 4 grounds about as faithfully as the unconstrained 397B frontier. The prompt moves each model only slightly; the gap between models dwarfs it.

Table 8. Oracle-faithfulness detail, all nine cells (counts out of $n = 2,989$). The four count columns split each Lynx FAIL by primary cause (a small residual *unclear* bucket is omitted): the two *genuine* modes—contradiction and unsupported addition—sum to the *categorized* rate, while omission and refusal are over-flags that do *not* count. The last column is the *calibrated* extrapolation (§7.5): uniformly far higher and *not* ordered like the categorized rate—Qwen is the most faithful generator by the categorized rate (0.97%) yet posts the highest calibrated rate (96.0%), purely because it is the most verbose. Read categorized as the metric of record. [†] marks the calibrated cells most inflated by long answers against the narrow top-3 context.

Generator · prompt	Genuine (↓)		Over-flags		Rate (↓)	
	Contradiction	Unsupported addition	Omission	Refusal	Categorized	Calibrated
<i>Gemma 4 E4B (deployed)</i>						
baseline	51	28	63	27	2.64%	19.1%
+G1	44	55	26	7	3.31%	23.1%
+G1+G2	48	58	36	20	3.55%	53.9% [†]
<i>Gemma 3n E4B</i>						
baseline	91	96	47	14	6.26%	27.1%
+G1	88	105	30	103	6.46%	45.7%
+G1+G2	88	113	44	70	6.72%	47.2%
<i>Qwen3.5-397B (frontier)</i>						
baseline	16	59	9	1	2.51%	48.2% [†]
+G1	21	90	8	0	3.71%	83.7% [†]
+G1+G2	7	22	3	1	0.97%	96.0%[†]

7.4 What kind of unfaithfulness, and how the prompt changes it

Splitting each Lynx FAIL by primary cause (Table 8) shows *why* the rates move and ties the faithfulness cost back to the deflection fix of §5.

- **G1 converts refusals into unsupported additions.** For Gemma 4, refusal failures fall $27 \rightarrow 7$ while unsupported additions rise $28 \rightarrow 55$: the prompt stops the model deflecting, and the substance it now volunteers sometimes reaches beyond the document. This is the deflection fix of §5 working *and* its price, in a single row—the faithfulness counterweight to the recall gain that motivated G1.
- **Gemma 3n’s deficit is a baseline contradiction load.** Before any prompt change, Gemma 3n contradicts its context almost twice as often as Gemma 4 (91 vs. 51): it gets clinical facts wrong against the very passages it was handed. This is the core grounding gap, and it is exactly what surfaces end-to-end as Gemma 3n’s dangerous Kenya answers (§5.2)—poor grounding becomes confident, specific, unsupported guidance, which becomes harm.

7.5 An honesty note on the calibrated estimate

Lynx is high-recall but not perfect—it also *misses* some true failures. A second, *calibrated* estimate (the last column of Table 8) corrects for those misses by re-judging a 50-answer sample of Lynx-PASS cases per arm with gpt-5 and extrapolating the miss rate across the ~2,800 PASS answers. We do not use it as the headline, but report it in full—for every cell, not just the flattering ones—because it is unreliable in magnitude and the table shows why.

It scales a 50-answer miss rate across ~2,800 answers, and that miss rate climbs with answer length, because the rubric judges strictly against the *narrow* top-3 context: clinically-sound elaboration that is simply absent from those three passages is counted as an “unsupported addition.” The clearest proof is internal to the table: Qwen earns the matrix’s *lowest* genuine-hallucination rate (0.97%) and its *highest* calibrated rate (96.0%) at once—an inversion driven entirely by its being the most verbose generator. A metric that ranks the most faithful model as the least faithful is measuring answer length, not faithfulness. Our own deployed-family outlier, Gemma 4 +G1+G2 at 53.9%, is the same effect in milder form: its full-population *categorized* rate is 3.55%, still second-lowest in the matrix.

What the calibrated column *does* robustly show, surviving its noise, is the two qualitative facts already in the categorized data (Gemma 3n > Gemma 4 at every arm; prompts push the rate up)—but only with wider error bars. We therefore report the categorized rate throughout and do not lean on the calibrated magnitudes.

7.6 Implication: faithfulness is a generator property

Two independent angles now converge on the same conclusion. The Kenya dangerous count (§5.2) and oracle faithfulness (this section) are measured on different data, with different judges, against different references—yet both find Gemma 4 roughly 2× safer than Gemma 3n, and both locate the cause in the model rather than the prompt. This is the deciding evidence behind deploying Gemma 4: it is the candidate that is as faithful as a frontier model while remaining small enough to run on-device.

The deficit is structural. No prompt in this study brings Gemma 3n’s faithfulness up to Gemma 4’s: the two deflection-fix prompts in fact nudge 3n’s hallucination rate slightly higher ($6.26 \rightarrow 6.72\%$), so they *widen* the gap between the two models rather than close it. That model-to-model gap (3.2–3.6 points across prompts) is roughly 4× the most any prompt shifts the rate within either on-device model. So the path to a more faithful—and therefore safely more helpful—on-device model is generator-side work: RAG-grounding or fine-tuning the generator to use only what it retrieves, not further prompt engineering or retrieval swaps. This is the same lever the size ladder (§5.4) and the retrieval analysis (§6) point to, and we take it up in §12. Three caveats bound the

claim: this is faithfulness to *gold* context (an upper bound; live retrieval adds its own errors); the semantic categorization rests on a single judge family (§10); and the categorized counts include the 16 FAILs later traced to corpus self-contradictions rather than model error (Appendix C)—a correction of ~ 0.5 points on the affected arm that shifts no comparison, since both models face the same corpus.

8 On-Device Latency

The preceding sections establish that the system answers *well*; this one establishes that it answers *fast enough*, on the device, with no network. The app ships **CPU-by-default** for broad compatibility—the safer, lowest-common-denominator backend—with GPU execution as an opt-in that has a higher hardware bar (§8.4). At the deployed top-3 depth, Gemma 4 E4B returns a full answer in a median of ~ 43 s on the shipped CPU path and ~ 19 s where a compatible GPU is available (there, first token in ~ 1 s and generation at ~ 13 tokens/s). Both fit the field budget we set ourselves (a query should resolve well under a minute) on this flagship and run entirely offline—the GPU with wide headroom, the CPU with a tighter margin that shrinks further on the cheaper hardware typical of the field (§8.4). The rest of this section traces where the time goes, surfaces one precision bug that bounds how much context we can admit, and sets out what a device needs to run it at all—the binding deployment constraint, it turns out, is not speed but memory and GPU availability.

8.1 Setup

All numbers are measured on real hardware: a OnePlus tablet (Snapdragon 8 Elite, 16 GB RAM, Android 15) running LiteRT-LM 0.11.0. We sweep three axes that matter for deployment—generator size (Gemma 4 E4B vs the smaller E2B), compute backend (GPU via OpenCL on the Adreno; CPU via XNNPACK), and retrieval depth $k \in \{0, 1, 3, 5, 7, 10, 15, 20\}$ —and time every cell over 18 query types $\times 3$ repeats (54 runs), with a 10 s cooldown between runs for thermal stability. We separate three quantities: *time-to-first-token* (TTFT, the prefill pass over the prompt, excluding retrieval), *decode* (first to last generated token), and *total query* (retrieval + TTFT + decode, what the user waits). Each benchmark records the SHA-256 fingerprint of the loaded model artifact so a reviewer can verify which weights produced which timing.

Retrieval—embedding the query and searching the vector store—is timed separately from TTFT and adds a near-constant ~ 1.7 s at the deployed depth, most of it the brute-force similarity scan over the 63,650-chunk index rather than the single query-embedding pass: a follow-up on-device benchmark confirmed the scan dominates retrieval time and is embedder-independent, so a faster embedder buys little here—the retrieval lever is the index itself. That figure predates the Gecko $\rightarrow$ EmbeddingGemma swap (§2) and may have shifted, but retrieval is a small slice next to the generation that dominates the total (the ~ 16 – 19 s decode at the deployed $k=3$), and the deployed generator is unchanged—so the generation-set headline figures below stand.

One further mismatch with the deployed configuration: these runs predate the G1 prompt (§2.3) and use the earlier baseline prompt, which deflected on about a third of questions with short “see a doctor” replies. Because G1 cuts deflection to $\sim 3\%$, the deployed system produces longer, more substantive answers on average, so its decode—and thus total latency—will run higher than measured here by an amount we did not quantify. The picture stays generation-dominated, so these figures are best read as a *lower bound* on deployed latency rather than an estimate of it: the budget margins below—already thin on the shipped CPU path and tighter still on lower-end hardware (§8.4)—are correspondingly smaller under G1.

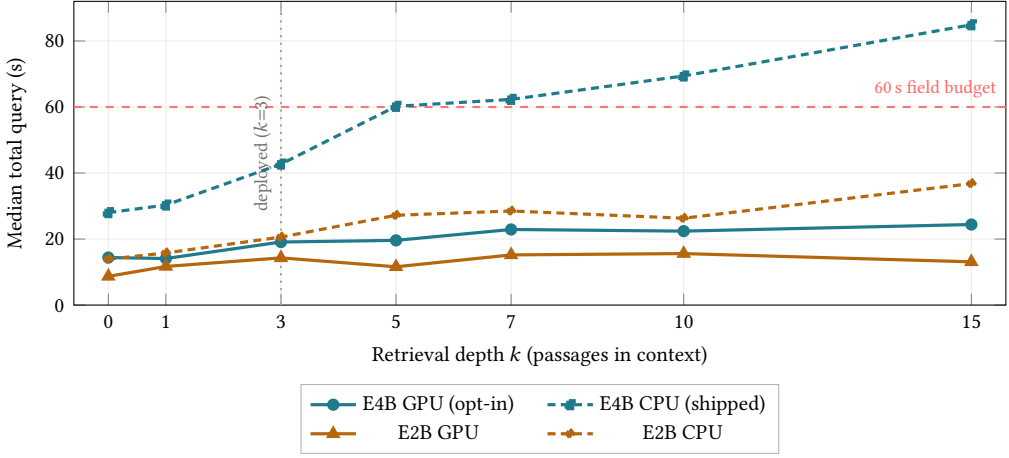


Fig. 11. Median total query latency vs. retrieval depth on the Snapdragon 8 Elite, by generator and backend (lower is better; $k = 20$ omitted—see §8.3). The optional E4B-GPU path (solid teal) stays 14–24 s across all depths, far under the 60 s field budget; the shipped E4B-CPU default (dashed teal) is slower and crosses the budget past $k \approx 5$, while the smaller E2B (amber) keeps every depth comfortably inside it on either backend.

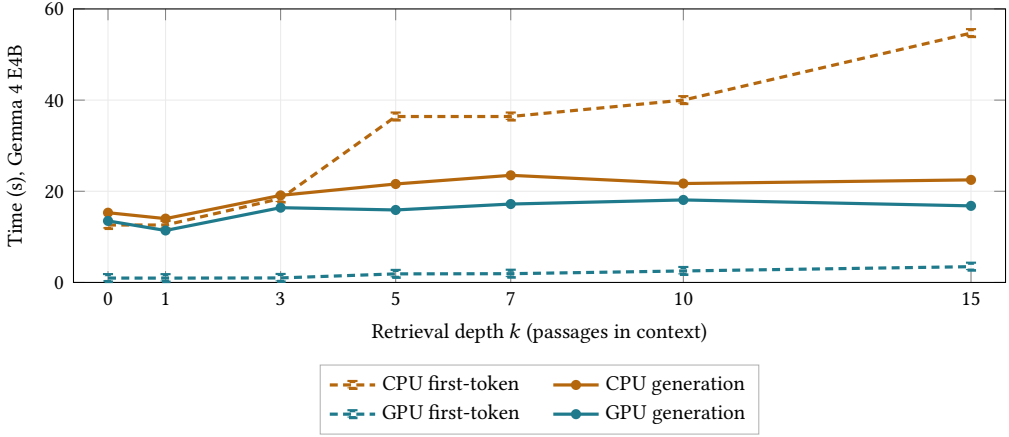


Fig. 12. Where the time goes as retrieval depth grows: first-token (prefill, dashed) and generation (decode, solid) latency for the deployed E4B generator, GPU (teal) vs. CPU (amber). On the **GPU** both curves are nearly flat in k —added context is almost free. On the **CPU** generation stays flat (~ 14 – 23 s) but first-token climbs steeply, from ~ 13 s to ~ 55 s, because prefill is compute-bound and the CPU has far less compute. That prefill blow-up—not decode—is what pushes the CPU total past budget in Figure 11.

8.2 Latency by backend and retrieval depth

Figure 11 is the headline. Three things stand out.

The GPU path has wide headroom. Gemma 4 E4B on the GPU stays between 14 and 24 s of total latency across *every* retrieval depth up to $k = 15$ —roughly a third of the 60 s budget at the deployed $k = 3$ (~ 19 s). Latency barely grows with k because, on the GPU, the cost is dominated by *decode*, not by the prompt (Figure 12, teal): TTFT rises only from ~ 1 s at $k = 3$ to ~ 3.5 s at $k = 15$, while

decode—a near-constant $\sim 11\text{--}18\text{ s}$ —sets the floor. Adding retrieved context is close to free on the GPU.

CPU is the real fork in the road. On the CPU—the shipped default—the same E4B generator runs $2\text{--}3.5\times$ slower than on the GPU ($28\text{--}85\text{ s}$ over the same range) and crosses the 60 s budget past $k \approx 5$. The gap is almost entirely prefill (Figure 12, amber dashed): CPU TTFT balloons from $\sim 13\text{ s}$ ($k = 0$) to $\sim 55\text{ s}$ ($k = 15$), because prefill is compute-bound and the CPU has far less of it, while CPU *generation* stays flat like the GPU’s. This is what makes the GPU/CPU split, not the retrieval depth, the decisive latency variable.

The smaller generator lowers the hardware floor. The unlock from halving the generator to E2B is *memory*, not speed. At the deployed $k = 3$, even E4B on CPU ($\sim 43\text{ s}$ on our flagship) fits the 60 s budget, so a missing GPU does not by itself rule out the deployed model. What rules it out is RAM: E2B halves the runtime footprint (~ 3.3 to $\sim 1.7\text{ GB}$), dropping the device RAM floor from 6 GB to 4 GB and so running on devices that cannot hold E4B at all (§8.4). E2B is also faster— $\sim 2.3\times$ on prefill and $\sim 1.5\times$ (GPU) to $\sim 2\times$ (CPU) on total, the decode gain smaller because GPU decode is memory-bandwidth- rather than compute-bound—which buys latency headroom on the slower mid-tier CPUs where E4B could exceed budget, though that hardware is not yet measured (§8.5).

8.3 The FP16-GPU context wall

We cap the model’s total context at $4,096$ tokens. That cap is not a memory or latency limit—it is a workaround for a precision bug we found, and it is worth documenting because it will bite anyone deploying a LiteRT-LM-class runtime at longer context.

On Android, the GPU text-decoder path defaults to FP16 activations (the CPU path uses FP32); this is a runtime choice, not a knob in our code. Past a total context of $\sim 5,000$ tokens, FP16 GPU decode collapses into a degenerate repetition loop (a stream of $*$ characters) and the output silently becomes garbage. The collapse is *deterministic*, not a stochastic glitch: the GPU backend decodes greedily, so a fixed prompt produces bit-for-bit identical garbage on every run (we confirmed bit-exact reproduction across repeated runs of the reference case), which rules out sampling noise and points to the FP16 attention math itself—once the precision error pushes the attention scores onto the $*$ token, the model self-reinforces into the loop. The same prompt decoded on the FP32 CPU path stays clean, which isolates the cause as the GPU’s FP16 precision, not the prompt, the model artifact, or retrieval (Figure 13). Forcing the GPU to FP32 also removes the cliff, confirming it directly. The fix costs latency: FP32 prefill is $2.1\text{--}2.5\times$ slower (prefill is compute-bound, and FP16 doubles Adreno arithmetic throughput), though decode—bandwidth-bound—is essentially unchanged, for a net $\sim 21\text{--}34\%$ slower total query at $k = 10\text{--}15$.

The $4,096$ cap therefore buys a ~ 900 -token safety margin below the cliff on the fast FP16 path. It is also why $k = 20$ is absent from Figure 11: at that depth the 8 longest of 18 query types produce prompts exceeding $4,096$ tokens and the runtime rejects them outright (a fixed 24 of 54 runs in every $k = 20$ cell), so a $k = 20$ median would be computed only over the queries short enough to survive—a survivorship artifact rather than a real measurement. For deployments that need longer context, the FP32 GPU path is the clean fallback at modest cost; on the FP16 GPU path, $4,096$ tokens is the safe ceiling. The cliff is specific to FP16: the shipped CPU default decodes in FP32 and never hits it, so the cap is a uniform runtime setting that chiefly protects the GPU opt-in.

8.4 Device compatibility

Latency is comfortable on a flagship; the harder question is what a device must provide to run the system at all. The binding constraints are not speed but memory, storage, and GPU availability—all properties we can state directly from the artifacts.

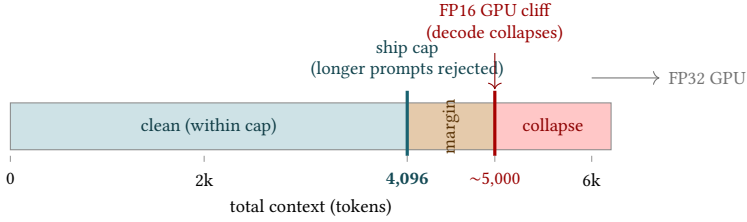


Fig. 13. The FP16-GPU context wall. The shipped 4,096-token cap (teal) sits ~900 tokens below the ~5,000-token cliff (red) past which FP16 GPU decode deterministically collapses; the margin (amber) is the safety buffer. The CPU FP32 path has no cliff, and forcing the GPU to FP32 extends the clean zone at ~25% added latency.

Table 9. Minimum and recommended device specifications, derived directly from the model runtime footprints (plus Android/app overhead) and the install size (plus working headroom). The minimum is what runs; the recommended figure leaves headroom against OOM kills under multitasking and against future bundle growth. These are device *requirements*; we make no claim about which retail devices meet them.

Generator	Device RAM		Free storage	
	Minimum	Recommended	Minimum	Recommended
Gemma 4 E4B	6 GB	8 GB	~4.5 GB	6 GB
Gemma 4 E2B	4 GB	6 GB	~3.5 GB	5 GB

Memory and storage. The generator’s runtime memory is the hard floor: ~3.3 GB for E4B and ~1.7 GB for E2B (the published model specs). Android and the app add ~1.5–2 GB on top, which sets the device RAM floor. Storage is the one-time install: the ~213 MB app, the model weights (3.66 GB for E4B, 2.59 GB for E2B), the ~275 MB knowledge base (the 63,650-passage index), and the ~180 MB embedder—~4.3 GB of assets for E4B and ~3.3 GB for E2B. These assets are obtained once—downloaded over Wi-Fi or *sideloaded* onto the device directly (USB or memory card, needing no connectivity)—after which the system runs fully offline; sideloading lets the install reach a low-connectivity site without each device pulling several gigabytes over the network. Table 9 gives the minimum and recommended specifications that follow.

GPU availability. The ~19 s figure is the opt-in GPU path, and assumes a working GPU—which we verified only on the Qualcomm Adreno of our test device. LiteRT-LM’s GPU path uses OpenCL, whose support is inconsistent across GPU families: on the ARM Mali GPUs common in non-Qualcomm SoCs it is reported to fail silently, in which case the app falls back to CPU—a path we have not tested (§8.5). The safe planning assumption is therefore CPU-first: at the deployed $k = 3$, E4B on CPU (~43 s) still meets the budget, and E2B on CPU (~21 s) covers the 4 GB devices that cannot fit E4B at all. The GPU’s ~19 s is the best case, not a guarantee.

8.5 What we have not measured

Two gaps bound these results, both disclosed rather than closed. First, every number here comes from a *single flagship* device with an Adreno GPU; the lower-cost MediaTek hardware likely to be used in the field is untested, and its CPU-only path will be slower than the Snapdragon 8 Elite CPU rows above (we extrapolate ~2× from published compute gaps, but have not measured it). Second, we report steady-state latency only: per-session battery drain and sustained-load thermal

throttling—both of which matter for a clinician using the app repeatedly through a shift—are not yet characterized. We scope these as the first follow-up measurements (§10) rather than reasons to withhold the result, since the load-bearing conclusion—fully offline answers within the field budget—holds on the hardware we did measure.

9 Discussion

Read top to bottom, the evaluation inverts the intuition about where an on-device medical RAG system is weakest. The retriever was the natural suspect—a 300M embedder on a phone against cloud systems—yet on-device retrieval is essentially solved (§6): EmbeddingGemma ranks third of seven and the corpus holds a relevant passage for the large majority of queries. The weak link is downstream, where it is least expected: the small generator does not turn good retrieval into better answers. That we can place the weak link there at all is a consequence of the evaluation design: isolating the generator under oracle context, apart from live retrieval (§7), separates a retrieval problem from a generation one, where an end-to-end score alone would have blamed the retriever. This section traces the finding to its causes and draws out what it means for deployment; the improvement directions it points to—technical levers, open design questions, and deployment prerequisites—are collected in §12.

9.1 The bottleneck has moved from retrieval to grounding

The sharpest single result is also the most counter-intuitive: *adding* retrieved context does not improve answers, and the more capable the generator, the more it hurts—the RAG penalty deepens from -0.01 HealthBench `weighted_met` at 4B to -0.10 at 33B (§5.4). It is tempting to read this as “the retrieved context is useless,” but that mistakes the probe for the question. A large model already holds the answer parametrically, so a retrieved chunk offers it little new and some distraction; the net is negative and grows with capability. The frontier is therefore the *wrong* test of whether the context helps, because the frontier does not need it.

The relevant subject is the deployed small model, which has no parametric route to specific clinical detail and so depends on retrieval for it. For that model the net cost has at least three plausible roots, only the first of which we can quantify:

- it under-uses the on-target context it is given—a faithful generator mirrors its chunk rather than reasoning over it (§6.3);
- the latency-capped three-chunk window may carry only a fragment of a multi-step protocol, which a faithful model then reproduces as a truncated answer; and
- some open-ended questions are simply not well covered by guideline chunks at all.

The three are coupled, and only the deployment as a whole exhibits their sum; we do not claim a measured split.

Two facts keep this from being a counsel of despair. Offline retrieval is genuinely good—the corpus holds a relevant chunk for 96.5% of queries and EmbeddingGemma rivals cloud retrievers—so the raw material exists. And per query the mechanism demonstrably works (Figure 8): the same small model gives a specific, safe-to-act-on answer when the right chunk is surfaced and a vague one when it is not. Whether RAG is net-negative because the generator under-uses good context (fixable) or because the context is genuinely low-value on this query distribution (not) is therefore open—and the evidence leans toward the former.

That open question does not make retrieval optional. Even at a small average cost, it is what makes an answer *auditable*: every claim is tied to a cited guideline passage a clinician can check (§3), and a faithful model grounded in a retrieved source is far less prone to confident fabrication than one answering from parameters alone—the same property that makes Gemma 4 safe to deploy

(§7). Per query it is also decisive when the corpus holds a local protocol or dose the model lacks. We retain retrieval for safety, traceability, and local specificity; the goal of the levers in §12 is to make it pay in answer quality too, turning a small net cost into a net gain.

9.2 Deployment is a safety-versus-usefulness choice

Beyond the bottleneck, the choice of which generator to deploy carries a lesson of its own. Between two same-size on-device candidates we did *not* choose the more capable one. Gemma 3n is the more helpful model, but it produces genuine, confident dosing and drug errors, and two independent probes—the targeted Kenya dangerous-answer count (§5.2) and oracle faithfulness (§7)—agree that Gemma 4 is roughly twice as safe and grounds about as faithfully as a frontier model. So we deploy the faithful model and recover its helpfulness with the prompt (§2.3), rather than deploy the more capable model and try to suppress its errors. In a safety-critical, low-supervision setting, faithfulness to sources is the right *primary* criterion for choosing the generator and raw capability the secondary one—the reverse of the usual leaderboard order.

10 Limitations

MAM-AI is a thoroughly evaluated research prototype, not a fielded product: all our evidence is benchmark behaviour scored by LLM judges, and several gaps bound what that evidence can claim.

- **No field test, and no real users.** A lag in project funding meant the planned study with Zanzibar nurse-midwives could not proceed this cycle, so there is no point-of-care data on how the system performs in real clinical use.
- **Not clinician-verified.** Clinical correctness is not established by humans. Every quality, safety, and faithfulness number comes from LLM judges; the rubric judge’s agreement with physicians was measured only on a third-party benchmark’s physician-labelled data (Health-Bench, §4.4), not on clinicians reviewing MAM-AI’s own answers. The dangerous-answer determinations are likewise model-generated, awaiting review by qualified clinicians.
- **English only.** The system, knowledge base, and evaluation are entirely in English, while the deployment setting is Swahili-speaking and the bilingual prompt was withdrawn as unvalidated (§2.3); Swahili support remains an open deployment gap.
- **Retrieval does not yet improve answers.** On the small on-device generator, better retrieval does not reach the answers and added context is net-neutral-to-negative (§6); we retain retrieval for traceability and per-query specificity, not for average quality, and closing this gap—generator-side grounding—is unsolved (§12).
- **Latency measured on a single device.** All timings come from one flagship (Snapdragon 8 Elite); the lower-cost MediaTek hardware likely used in the field is untested and will be slower—possibly past budget on the shipped CPU path—and battery drain and sustained thermal throttling on a warm ward are not yet profiled (§8.5).
- **Corpus coverage and quality.** The corpus was not expanded this cycle, so coverage gaps remain; it contains internal contradictions between sources (§12); and licensing limits some high-value reference texts.

11 Related Work

Clinical AI in low-resource settings. The closest precedents to MAM-AI are point-of-care digital tools for frontline maternal-health workers. The Safe Delivery App is an *offline* smartphone reference for birth attendants—curated text, animations, and instructional video—that more than doubled postpartum-haemorrhage management skills in a cluster-randomised trial in Ethiopia [Christiansen

et al. 2023]; it shares MAM-AI’s offline, point-of-care setting, but is a fixed-content library rather than an assistant that answers free-form questions. More recently, LLM clinical decision support has been deployed in African primary care, reducing diagnostic and treatment errors across tens of thousands of visits in Kenya [Korom et al. 2025], though cloud-served and clinician-facing in facilities. MAM-AI joins these two threads as an offline, on-device LLM that answers nurse-midwives’ questions—a regime whose closest precedent is MedAide, an offline edge-LLM medical assistant [Basit et al. 2024]. Pan-African medical QA benchmarks such as AfriMed-QA [Nimo et al. 2025] motivate evaluation for these settings, but are datasets rather than deployed systems.

Medical language models and retrieval. Open and frontier medical LLMs—among them the Gemma-based MedGemma [Sellergren et al. 2025]—are large and server-hosted; we instead target a 4B model that runs on-device, and find (§5.4) that capability dominates quality while faithfulness decides which model we deploy. Retrieval-augmented clinical systems such as Almanac [Zakka et al. 2024] argue that grounding answers in cited sources is what makes a medical assistant trustworthy—the same motivation behind our traceable RAG—and medical-RAG benchmarks [Xiong et al. 2024] dissect retriever and generator contributions; unlike that cloud setting, we find the small *generator*, not retrieval, to be the bottleneck. We evaluate on HealthBench [Arora et al. 2025], whose physician-written rubrics let us separate completeness from safety.

On-device inference and quantization. Running LLMs on commodity mobile devices is enabled by mobile runtimes—Google’s LiteRT-LM [Google AI Edge 2025] (our generator’s runtime), and open engines such as llama.cpp and MLC-LLM—together with low-bit quantization: weight-only int4 methods [Lin et al. 2024] and the outlier-aware 8-bit analysis of Dettmers et al. [2022], whose findings on low-precision transformer numerics bear directly on the FP16 attention collapse we characterize (§8.3). Where much on-device work favours sub-billion-parameter models, we deploy a 4B model at int4, trading precision for capability.

Faithfulness and judge evaluation. Detecting whether a generation stays faithful to its source spans claim-level fact-checking [Tang et al. 2024] and RAG-specific faithfulness metrics [Es et al. 2024]; we use Patronus Lynx [Ravi et al. 2024] as the first stage of our two-pass faithfulness pipeline. Our rubric and safety scores come from an LLM-as-judge, a standard but bias-prone practice [Zheng et al. 2023], so we validate the judge against physician labels (§4.4) rather than trust it blind.

On-device text embeddings. Distillation has made small embedders competitive with much larger ones—Gecko [Lee et al. 2024] and the 300M EmbeddingGemma [Schechter Vera et al. 2025], our retriever—and our results add evidence that an on-device embedder rivals cloud APIs on medical retrieval. We benchmark against general (MTEB [Muennighoff et al. 2023]) and biomedical (MedCPT [Jin et al. 2023]) baselines, but with graded relevance and Precision@ k tailored to the deployment task.

12 Future Directions

The evaluation points to a concrete programme of work, in three tracks: near-term technical levers that follow directly from the diagnosis (§9); open design questions that need clinical partners as much as engineering; and the prerequisites for a real field deployment. We order the levers *foundation-first*, working from the corpus that everything else depends on out to the generator that consumes it.

12.1 Technical levers

We take them from the foundation outward—the corpus everything retrieves from, the queries that hit it, the languages it can reach, and the generator’s use of what it receives—then close with two retrieval-tuning knobs, chunking and depth.

Corpus expansion and curation (the foundation). A RAG system is only as good as the guidelines it can retrieve, and a small generator that leans on what it receives (§6) makes the corpus decisive. Coverage is high in aggregate—a relevant chunk exists for 96.5% of queries—but the corpus was not expanded this cycle, and per-query gaps are decisive when they fall on a query it cannot answer. The reliable, direct lever is to grow the corpus on two fronts—authoritative international guidelines (WHO, NICE, and comparable bodies) and local protocols (Tanzania and Zanzibar)—*broadening* each so more questions have a trustworthy source at all and *deepening* each toward information-rich, actionable content: the exact doses, thresholds, and steps a nurse can act on, including the locally specific ones. It is the slowest lever but the most reliable, and its payoff is the most direct per query.

Query rewriting. The retriever consumes the nurse’s query *verbatim*, and real input is noisy—lay phrasing, abbreviations, partial descriptions—where wording alone decides what is retrieved (“severe nausea” and “hyperemesis gravidarum” fetch different bundles for the same patient, §6.3). A rewriting or normalisation step that maps the query onto the corpus’s vocabulary can surface the right section before the generator runs, paired with brief user guidance on how to describe a case. Its cost is an extra inference step on the device: cheap to build but not free at run time—it adds latency (and, as a separate model, memory) to an already-tight budget—so it must earn its place through a retrieval-precision gain, and is most attractive when it reuses the deployed model to emit only a short normalised query rather than loading a second model.

Multilingual and cross-lingual retrieval. Because EmbeddingGemma and Gemma 4 are both multilingual, the system can in principle retrieve and answer across languages: an English query fetching German guidelines, or a Swahili query fetching English and German passages and being answered in Swahili. This is doubly valuable—it would let the corpus draw on authoritative non-English references it cannot exploit today (high-quality German-language obstetrics texts, for one), feeding the corpus lever above, and it is the natural path toward serving Swahili-speaking users. It is also the most *exploratory* lever here: cross-lingual retrieval quality and the faithfulness of Swahili answers are unmeasured and would have to be evaluated before being relied on. (This is a research lever, distinct from *shipping* a validated Swahili configuration: the latter is what field use requires, and we list it separately among the deployment prerequisites below.)

Generator grounding (the deepest lever). Once the corpus is rich and the right passages reach the model, the remaining lever is the generator’s use of them: teach the small model to extract specific, faithful guidance from retrieved context—to reach through retrieval what the frontier reaches through scale, since scale is not available on the device (§5.4). Concretely this is a grounding fine-tune: train the deployed generator on (context, faithful-answer) supervision, optionally distilled from a frontier model, so it uses on-target passages and resists off-target ones while keeping every claim cited. Unlike the retrieval-side levers it adds no run-time cost—the same model, one pass, better weights—but it is the most expensive to build and depends on the corpus work above for its training data. The Qwen ceiling shows the prompts and task are already sound (§5.2); what is missing is a small model that uses its context well. The experiment that would settle whether RAG can be made net-positive is direct—a grounding fine-tune evaluated under *live* retrieval: if it flips RAG net-positive, under-grounding was the cause; if a well-grounded small model still cannot beat no-RAG, the remaining constraint is capability and corpus.

Chunking and retrieval depth. Two retrieval-shaping knobs tune how well the retrieved content serves a query, and—unlike the levers above—they pay off *only* once the corpus holds good guidelines and only within the latency budget. *Chunking* that keeps a whole protocol intact, rather than splitting it across chunks, avoids handing the generator a fragment it can answer only partially; greater *retrieval depth* (a larger k) brings more of a procedure into the window but spends latency—the shipped CPU path holds k at three for the 60 s budget, while a GPU affords more ($k = 15$ in ~24 s, §8). Reranking complements both, ensuring the *management* section fills the budget, not a definition or a study-exercise (Figure 8). These are tuning levers, not foundations—listed last because they are gated on the corpus and the hardware above them.

12.2 Open design questions

Two design questions a fielded system must answer need clinical partners as much as engineering.

Knowing when to abstain. Sometimes the safest answer is no answer: when the corpus cannot support a query, the system should say so rather than guess. Building this has two parts. The first is technical—spotting the queries the system is likely to get wrong—and it is unsolved: none of the signals we tried tell answerable queries from unanswerable ones, neither a retrieval score nor an answer-side check (§6.4). A usable signal would probably have to come from the generator itself—how confident or self-consistent its answer is—which we have not built. The second part is clinical, and harder: deciding *when* the system should hold back, and *what* it should tell the nurse instead, depends on the nurse-midwife’s scope of practice—which decisions are theirs to make, and which must be referred to a doctor. That line has to be drawn with clinicians before any such gate is built.

Conflicting guidance. The corpus draws on many authoritative sources, and they do not always agree. Auditing the faithfulness failures (§7) surfaced 16 cases in which two retrieved passages give conflicting advice on the same clinical question, each verified verbatim against its source—a field reference re-doses emergency contraception after vomiting within *three* hours where WHO says *two*; one WHO document starts the active phase of labour at 4 cm dilatation where another starts it at 5 cm (Appendix C). These are few, and whether the differences matter clinically is not ours to judge—that needs clinicians. The mechanism, though, is real: the model is handed both passages at once, so whichever it follows can be faulted against the other. When sources diverge, what should the system do—prefer the local or most recent protocol, surface both with their provenance, or flag the conflict to the clinician? Citations make the disagreement visible, but deciding what the system should do about it—and confirming that choice with clinicians—is still unsolved, and not just for MAM-AI: any assistant built on many different sources faces it.

12.3 Toward a field deployment

Beyond the research levers, a real deployment has prerequisites this cycle could not close; they bound what our evidence can claim (§10) and form the engineering-and-study track a successor would run in parallel.

- **A field study with nurse-midwives.** Above all, the system needs a study with the nurse-midwives it is built for: all current evidence is benchmark behaviour scored by LLM judges, with no point-of-care data on real clinical use. This is the single most important gap.
- **Clinician verification.** Every quality, safety, and faithfulness number comes from LLM judges; the dangerous-answer determinations in particular await review by qualified clinicians. Independent clinical verification is a prerequisite for any claim of clinical correctness.

- **Swahili support.** The system, corpus, and evaluation are English-only, while the deployment setting is Swahili-speaking, and a safety-critical prompt cannot ship in a language it was not evaluated in (§2.3). Validating a Swahili configuration—which the cross-lingual lever above complements—is a prerequisite for field use.
- **Validation on field hardware.** All timings come from one flagship device—a OnePlus tablet with a Snapdragon 8 Elite SoC; the lower-cost MediaTek devices likely used in the field are untested and will be slower, and battery drain and sustained thermal throttling over a ward shift are unprofiled (§8.5). A smaller Gemma 4 E2B variant is available as a fallback for memory-constrained or slower devices—it roughly halves the runtime memory footprint (dropping the device RAM floor from 6 to 4 GB) and is somewhat faster (up to ~2× on CPU total)—but it too has been measured only on this flagship. We also timed the baseline prompt, whereas the release uses the G1 prompt, which can produce longer responses and thus more latency, so the measured numbers are a lower bound.

13 Conclusion

MAM-AI shows that a fully offline, source-cited medical question-answering assistant can run on a commodity Android device for nurse-midwives in a low-connectivity setting. Its layered evaluation relocates the hard problem. On-device retrieval is essentially solved: on the mamaretrieval benchmark a 300M embedder ranks third of seven and rivals cloud systems, so the passages the system needs are usually found. The small generator is what remains in doubt—at 4B it cannot be both helpful and safe at once: Gemma 3n answers more fully but commits genuine dangerous errors, while Gemma 4 is safe but less helpful, whereas the 397B frontier model escapes the trade-off entirely, marking the gap as a limit of small-model capacity rather than of the task. Corpus quality is decisive for the same reason: a small model carries little parametric medical knowledge and leans on what it retrieves, so it needs official, traceable guidelines and locally relevant protocols—and where the corpus holds the right passage the answer is specific and safe to act on, where it does not the answer goes vague (Figure 8).

Generator capability matters—it sets the quality ceiling—but in a safety-critical setting we prioritise faithfulness, deploying the model that grounds its answers most reliably and recovering helpfulness through the prompt. MAM-AI is a thoroughly evaluated research prototype, not a fielded product; closing the gap to deployment means grounding the generator, improving the corpus, extending to Swahili, and, above all, testing with the nurse-midwives it is built for—the programme of work laid out in §12.

References

- Rahul K. Arora, Jason Wei, Rebecca Soskin Hicks, et al. 2025. HealthBench: Evaluating Large Language Models Towards Improved Human Health. arXiv:2505.08775 [cs.CL] <https://arxiv.org/abs/2505.08775>
- Abdul Basit, Khizar Hussain, Muhammad Abdullah Hanif, and Muhammad Shafique. 2024. MedAide: Leveraging Large Language Models for On-Premise Medical Assistance on Edge Devices. arXiv:2403.00830 [cs.CL] <https://arxiv.org/abs/2403.00830>
- Leah F. Bohle. 2025. MAM*AI: Developing an AI-enabled Medical Chatbot for Midwives. Presentation, Medicus Mundi Switzerland Symposium, Basel. https://www.medicusmundi.ch/assets/uploads/files/resources/2025/3.%20MAMAI_MMS_30Apr25.pdf Swiss Tropical and Public Health Institute.
- Trevor Brokowski, Cael Marquard, Mohammad Zaid Moonsamy, Fiifi Dawson, and Fay Elhassan. 2025. LiGHT MAM-AI. Writeup, Google Gemma 3n Impact Challenge (Kaggle). <https://www.kaggle.com/competitions/google-gemma-3n-hackathon/writeups/light-mam-ai>
- Ann-Marie Hellerung Christiansen, Bjarke Lund Sørensen, Ida Marie Boas, Tariku Bedesa, Wondewossen Fekede, Henriette Svarre Nielsen, and Stine Lund. 2023. The impact of the Safe Delivery Application on knowledge and skills managing postpartum haemorrhage in a low resource setting: a cluster randomized controlled trial in West Wollega region, Ethiopia. *Reproductive Health* 20, 1 (2023), 91.

- Tim Dettmers, Mike Lewis, Younes Belkada, and Luke Zettlemoyer. 2022. LLM.int8(): 8-bit Matrix Multiplication for Transformers at Scale. In *Advances in Neural Information Processing Systems (NeurIPS)*.
- Shahul Es, Jithin James, Luis Espinosa-Anke, and Steven Schockaert. 2024. RAGAs: Automated Evaluation of Retrieval Augmented Generation. In *Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics (EACL): System Demonstrations*.
- Google AI Edge. 2025. LiteRT-LM: On-Device Large Language Model Inference. <https://github.com/google-ai-edge/LiteRT-LM>.
- Qiao Jin, Won Kim, Qingyu Chen, Donald C. Comeau, Lana Yeganova, W. John Wilbur, and Zhiyong Lu. 2023. MedCPT: Contrastive Pre-trained Transformers with large-scale PubMed search logs for zero-shot biomedical information retrieval. *Bioinformatics* 39, 11 (2023), btad651.
- Robert Korom, Sarah Kiptinness, Najib Adan, et al. 2025. AI-based Clinical Decision Support for Primary Care: A Real-World Study. arXiv:2507.16947 [cs.CL] <https://arxiv.org/abs/2507.16947>
- Jinhyuk Lee, Zhuyun Dai, Xiaoqi Ren, Blair Chen, Daniel Cer, Jeremy R. Cole, Kai Hui, Michael Boratko, Rajvi Kapadia, Wen Ding, Yi Luan, Sai Meher Karthik Duddu, Gustavo Hernandez Abrego, Weiqiang Shi, Nithi Gupta, Aditya Kusupati, Prateek Jain, Siddhartha Reddy Jonnalagadda, Ming-Wei Chang, and Iftekhhar Naim. 2024. Gecko: Versatile Text Embeddings Distilled from Large Language Models. arXiv:2403.20327 [cs.CL] <https://arxiv.org/abs/2403.20327>
- Ji Lin, Jiaming Tang, Haotian Tang, Shang Yang, Wei-Ming Chen, Wei-Chen Wang, Guangxuan Xiao, Xingyu Dang, Chuang Gan, and Song Han. 2024. AWQ: Activation-aware Weight Quantization for LLM Compression and Acceleration. In *Proceedings of Machine Learning and Systems (MLSys)*.
- Niklas Muennighoff, Nouamane Tazi, Loïc Magne, and Nils Reimers. 2023. MTEB: Massive Text Embedding Benchmark. In *Proceedings of the 17th Conference of the European Chapter of the Association for Computational Linguistics (EACL)*.
- Charles Nimo, Tobi Olatunji, et al. 2025. AfriMed-QA: A Pan-African, Multi-Specialty, Medical Question-Answering Benchmark Dataset. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (ACL)*. 1948–1973. <https://aclanthology.org/2025.acl-long.96/>
- Andrea Nove, Ingrid K Friberg, Luc de Bernis, Fran McConville, Allisyn C Moran, Maligwilango Najjemba, Petra ten Hoope-Bender, Sally Tracy, and Caroline S E Homer. 2021. Potential impact of midwives in preventing and reducing maternal and neonatal mortality and stillbirths: a Lives Saved Tool modelling study. *The Lancet Global Health* 9, 1 (2021), e24–e32.
- Selvan Sunitha Ravi, Bartosz Mielczarek, Anand Kannappan, Douwe Kiela, and Rebecca Qian. 2024. Lynx: An Open Source Hallucination Evaluation Model. arXiv:2407.08488 [cs.CL] <https://arxiv.org/abs/2407.08488>
- Ren Yi. 2026. mamabench and mamaretrieval: Benchmarks for Evaluating Medical Retrieval-Augmented Generation in Maternal, Neonatal, and Reproductive Health. arXiv:2606.29467 [cs.CL] Companion paper.
- Henrique Schechter Vera, Sahil Krishna, et al. 2025. EmbeddingGemma: Powerful and Lightweight Text Representations. arXiv:2509.20354 [cs.CL] <https://arxiv.org/abs/2509.20354>
- Andrew Sellergren et al. 2025. MedGemma Technical Report. arXiv:2507.05201 [cs.CL] <https://arxiv.org/abs/2507.05201>
- Swiss Tropical and Public Health Institute. 2024. MAM*AI: AI-powered Medical Chatbot for Midwives. Project page. <https://www.swisstph.ch/en/projects/project-detail/project/mamai-ai-powered-medical-chatbot-for-midwives>
- Liyan Tang, Philippe Laban, and Greg Durrett. 2024. MiniCheck: Efficient Fact-Checking of LLMs on Grounding Documents. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing (EMNLP)*.
- Guangzhi Xiong, Qiao Jin, Zhiyong Lu, and Aidong Zhang. 2024. Benchmarking Retrieval-Augmented Generation for Medicine. In *Findings of the Association for Computational Linguistics: ACL 2024*.
- Cyril Zakka, Rohan Shad, Akash Chaurasia, et al. 2024. Almanac—Retrieval-Augmented Language Models for Clinical Medicine. *NEJM AI* 1, 2 (2024).
- Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, et al. 2023. Judging LLM-as-a-Judge with MT-Bench and Chatbot Arena. In *Advances in Neural Information Processing Systems (NeurIPS)*.

A The Deployed System Prompt (G1)

The system prompt is part of the method, not configuration trivia: its effect on deflection and key-fact recall is one of the paper’s headline results (Section 5). We reproduce the deployed English prompt (“G1”) in full here so that result is auditable and reproducible. The canonical, version-controlled source—together with the baseline and G1+G2 variants used in the matrix—is in the evaluation repository.⁵

⁵Deployed prompt: https://github.com/nmrenyi/mamai/blob/main/config/prompts/system_en.txt. Baseline, +G1, and +G1+G2 variants (the matrix arms): <https://github.com/nmrenyi/mamai-eval/tree/main/configs/config-v0.2.0>.

The prompt's design targets the over-deflection failure mode through a set of explicit levers (clinician-not-patient framing; an in-scope rule that location names never trigger a refusal; a manage-then-refer ordering; permission to name standard first-line drugs and doses with a formulary caveat; and a local-availability preference), while keeping the safety-critical instructions (emergency escalation, no invented doses, calibrated uncertainty) intact.

You are a clinical decision-support assistant for nurse-midwives in Zanzibar. Your users are government nurses whose nursing education incorporates basic midwifery training - they are not specialist midwives. They work at primary, secondary, and tertiary government health facilities, often with limited resources and specialist backup.

YOUR USER IS A CLINICIAN, NOT THE PATIENT: You are speaking to a trained nurse-midwife who assesses, prescribes, and makes clinical decisions independently for these cases. Answer clinician-to-clinician - use professional terminology and give decision support. Do not write advice addressed to a pregnant woman, and do not add patient-facing hedging.

SCOPE: You help with neonatal care, maternal health, obstetrics, family planning, gender-based-violence care, and related reproductive and clinical topics. Treat any clinical or health question as in scope, including emergencies and questions about management or medication. A location, country, or setting named in a question (e.g. a vignette set in Kenya or Zanzibar) never makes it out of scope. Only genuinely non-clinical topics are out of scope - for those, give a brief safe pointer, then redirect to clinical questions. Never decline a clinical question.

MANAGE, THEN REFER: Give concrete first-line clinical content **FIRST** - assessment steps, first-line management, and named measures - and **THEN** advise escalation or referral where appropriate. Referral is a complement to management, not a substitute for it. Never answer with referral alone ("consult a doctor") when first-line steps exist: the nurse needs what to do now as well as when to escalate.

EMERGENCIES - for any red-flag or emergency presentation, **ALWAYS** give immediate first-line safety guidance (what to do now) **AND** advise urgent escalation/referral, stating why. Never refuse or treat an emergency as out of scope. Red flags include:

- Heavy bleeding (postpartum or antepartum haemorrhage)
- Convulsions or loss of consciousness (eclampsia)
- Cord prolapse or abnormal fetal presentation
- Shoulder dystocia
- Severe difficulty breathing (mother or newborn)
- Fever in a newborn or signs of neonatal sepsis
- Signs of maternal sepsis (fever, rapid pulse, confusion in the mother)
- Severe abdominal pain
- Choking, anaphylaxis, or any acute life-threatening event in any patient - including a child: always give basic emergency steps and advise calling emergency services.

MEDICATIONS: You may name first-line drugs and typical doses when they are standard for the presentation - always add a brief caveat to confirm against the local formulary or protocol. If you are unsure of an exact dose, still give the drug and first-line approach, and advise confirming the dose against the local protocol. Do not invent doses.

LOCAL CONTEXT: Prefer options that are locally available and affordable in a resource-limited Zanzibar government setting. When a recommended option may be unavailable (e.g. a capsule formulation, or a service such as acupuncture that is rare or private-sector only), say so and offer an available alternative. Location informs which option to prefer - it never causes a refusal.

CONVERSATION: You may have access to previous messages in this conversation - use them to maintain context and avoid repeating information already covered.

LANGUAGE & TONE: Use simple, short sentences. Avoid idioms and complex words. Answer in English. Be supportive, professional, and calm.

FORMAT: Use markdown. Use bullet points for lists. Use **bold** for important terms. Use numbered steps for procedures. Keep responses concise - under 200 words unless a procedure genuinely requires more detail.

USING CONTEXT: If retrieved context is provided, use it to answer. If the context is not relevant to the question, say so and answer from established medical knowledge instead. When you use information from a document, add its citation number at the end of the relevant sentence - e.g. [1], [2], or [3].

UNCERTAINTY: If you are genuinely unsure, say so briefly - but still give the safest reasonable first-line guidance rather than deferring entirely. Do not guess at specific doses. Prioritize patient safety above all else: for emergencies and red flags, safe first-line action plus escalation always comes before any expression of uncertainty.

B Case Study: When Retrieval Decides the Answer

This appendix records the full transcripts behind Figure 8 (\$6.3): two near-identical queries whose answers diverge because retrieval handed the generator different chunks. The system's *queries* and *generated answers* are reproduced verbatim, as the deployed system's own input and output. The retrieved chunks are reproduced as the *actual text the model received* (lightly cleaned of PDF-extraction artifacts); we withhold all source titles and page numbers and instead tag the four distinct source documents "A"–"D" (full provenance is in the corpus manifest). **The labels are per source document:** the same letter means the same document, possibly a different passage from it. That re-use is the crux—source A supplies both the management protocol Case A retrieves and the bare definition Case B retrieves, and source D supplies two further passages—so the same guidance can be surfaced or missed depending only on which passage the retriever returns. Bracketed numbers [1]–[3] mark the retrieval order and match the citations in the generated answers.

Case A — query says "severe nausea"

Query (verbatim). "A 20-year-old woman, Gravida 1, currently in her second trimester of pregnancy with an estimated due date of 15 September, is suffering from severe nausea. What are the recommended treatment options?"

Retrieved top-3 (actual chunk text; sources withheld).

[1] Source A — management section.

Moderate vomiting and hyperemesis

If the woman is experiencing *moderate* vomiting or suspected *hyperemesis gravidarum*, the midwife should arrange for her to be admitted to hospital without delay. She should be cared for in a single room if possible, to avoid undue disturbance. A urinalysis is conducted for ketones, bilirubin, protein and glucose, and a midstream specimen of urine is sent for culture to exclude pyelonephritis.

Antihistamines such as *promethazine hydrochloride*, *prochlorperazine* (*Stemetil*) *metoclopramide* (*Maxolon*, *Primperan*) (Jewell & Young 2003) or corticosteroids such as *methylprednisolone* (Ismail & Kenny 2007) are considered safe to use in pregnancy and may be utilized to control the nausea and vomiting.

Blood tests, such as aspartate aminotransferase (AST), alanine aminotransferase (ALT), urea and electrolytes, are required to assess renal and liver function. An ultrasound scan may also be useful to eliminate hydatidiform mole and multiple pregnancy as probable causes. A frequent record of the woman's weight,

temperature, pulse and blood pressure is essential to monitor her wellbeing. In addition, the fetal heart rate should be appropriately auscultated (depending on the gestation) to monitor the health of the fetus.

[2] *Source D — rationale.*

Why the committee made the recommendations

gravidarum can have serious harmful consequences, and treatment and care in hospital may be needed. It should be noted that this guideline only covers treatments to manage nausea and vomiting in pregnancy, and comprehensive management of hyperemesis gravidarum, which may include nutritional interventions, is not covered by this guideline on routine antenatal care.

[3] *Source D — recommendations.*

Nausea and vomiting

- 1.4.1 Reassure women that mild to moderate nausea and vomiting are common in pregnancy, and are likely to resolve before 16 to 20 weeks.
- 1.4.2 Recognise that by the time women seek advice from healthcare professionals about nausea and vomiting in pregnancy, they may have already tried a number of different interventions.
- 1.4.3 For pregnant women with mild-to-moderate nausea and vomiting who prefer a non-pharmacological option, suggest that they try ginger.

Generated answer (verbatim).

Severe nausea in pregnancy needs careful management.

- **Mild to moderate nausea:** Reassure the woman that this is common. Suggest trying **ginger** as a non-drug option [3].
- **Severe nausea (Hyperemesis Gravidarum):** If the nausea is severe, she needs hospital care without delay [1].

Hospital Care Steps:

- (1) Monitor her **weight, temperature, pulse, and blood pressure** frequently [1].
- (2) Check her **fetal heart rate** regularly [1].
- (3) Perform a **urinalysis** for ketones, bilirubin, protein, and glucose [1].
- (4) Send a **midstream urine specimen** for culture to check for infection [1].
- (5) Blood tests for **liver and kidney function** (AST, ALT, urea, electrolytes) are needed [1].
- (6) An **ultrasound scan** may be used to rule out other causes [1].

Medication Options:

- Antihistamines like **promethazine hydrochloride, prochlorperazine (Stemetil), metoclopramide (Maxolon, Primperan)**, or corticosteroids like **methylprednisolone** may be used to control symptoms [1].

Please consult a doctor or senior clinician for the final treatment plan.

Case B — query says “hyperemesis gravidarum”

Query (verbatim). “A 20-year-old woman, Gravida 1, currently in her second trimester of pregnancy with an estimated due date of 15 September, is suffering from hyperemesis gravidarum. What are the recommended treatment options?” (The same patient and condition as Case A; only the wording differs.)

Retrieved top-3 (actual chunk text; sources withheld).

[1] *Source B — general management section.*

Management

The woman is usually admitted to hospital for investigations and rehydration. The cause of vomiting may not be found. An **antiemetic** will be given and fluids and electrolytes replaced by intravenous infusion of a solution such as Hartmann’s. Vitamin B12, thiamine, vitamin C, folic acid and iron will be needed to prevent complications (Farrell 2008).

General observation of the woman’s condition should be monitored. Strict fluid balance should be maintained until rehydrated. There is usually a rapid response to treatment and oral fluids may be recommenced when vomiting has ceased for 24 h. Solid food should be then introduced gradually. Moran & Taylor (2002) found that weight loss of more than 5% of prepregnancy weight in women with HG was

effectively treated with 10 mg of prednisolone three times a day. This shortened the stay in hospital and stopped vomiting. This treatment was gradually decreased and discontinued at 20 weeks of gestation.

[2] *Source C — student study-exercise.*

Activity

Find out what is meant by hyperemesis gravidarum.

See if you can discover what is the incidence and who is most at risk of this condition.

Find out how the midwife would recognize this condition, and what the treatment is.

[3] *Source A — definition section (the same document as Case A's passage [1]).*

Hyperemesis gravidarum

Hyperemesis gravidarum is a pathological condition characterized by unremitting, severe vomiting in pregnancy. It occurs in 0.3–2.0% of pregnancies and is more common in women who are younger, non-smokers and non-Caucasian (Ismail & Kenny 2007). It is diagnosed by exclusion (Kametas & Nelson-Piercy 2008) and is a leading cause of hospital admission during pregnancy (Cedergren et al 2008). Studies have shown that a gastric infection caused by the *Helicobacter pylori* bacterium may also be linked to hyperemesis gravidarum (Golberg et al 2007). The infection responds well to antibiotic therapy, such as *erythromycin* or *clarithromycin*.

Generated answer (verbatim).

Hyperemesis gravidarum is severe, continuous vomiting during pregnancy. For this woman, the recommended treatment options include:

- **Hospital Admission:** She usually needs to be admitted for tests and rehydration.
- **Fluid and Electrolyte Replacement:** Give intravenous fluids, such as Hartmann's solution.
- **Medication:** An **antiemetic** will be given.
- **Supplements:** Give Vitamin B12, thiamine, Vitamin C, folic acid, and iron to prevent problems.
- **Monitoring:** Watch her condition closely and keep her fluid balance strict until she is rehydrated.
- **Diet:** Start oral fluids when vomiting stops for 24 hours. Then, slowly introduce solid food.

If she loses more than 5% of her pre-pregnancy weight, a doctor may consider **prednisolone** [2]. Please consult a doctor for the exact treatment plan and medication doses.

Takeaway. The two queries describe the same patient and the same condition; only the wording differs. Case A surfaced the management section and produced a specific, actionable answer with named drugs and a full workup; Case B surfaced the *definition section* of the same source, a generic management passage whose antiemetic is unnamed, and a study-exercise, and the answer collapsed to “an antiemetic will be given.” The corpus held the right management section in both cases—a small, faithful generator reproduced whatever it was handed. This is the per-query face of §6.3: corpus coverage, specificity, curation, and query phrasing decide deployment quality.

C Corpus Self-Contradictions

This appendix expands §12 (conflicting guidance) with three of the 16 confirmed corpus self-contradictions. For each query we show the *two conflicting passages* the model retrieved, each in its own block labelled with its (openly available) source and page, the conflicting statement highlighted. Both reach the generator in a single context, so whichever it follows can be faulted against the other.

Re-dosing emergency contraception after vomiting

Query. *How should high-dose emergency contraceptive pills be dosed and managed if vomiting occurs?*

q_00459

Hesperian, A Book for Midwives (p.318)

Using high-dose pills (with 50 mcg ethinyl estradiol)

- give 2 pills. Then 12 hours later, give 2 more pills.

Emergency pills can give women headaches or severe nausea. Women can try eating something at the same time as taking the pills. If a woman vomits within 3 hours of taking the pills, she should take the same dose again.

WHO contraception guidance (p.74)

3. Explain what to do about side effects

- Nausea:
 - Routine use of anti-nausea medications is not recommended.
 - Women who have had nausea with previous ECP use or with the first dose of a 2-dose regimen can take anti-nausea medication such as 25–50 mg meclizine hydrochloride (such as Agyrax, Antivert, Bonine, Postafene) one-half to one hour before taking ECPs.
- Vomiting:
 - If the woman vomits within 2 hours after taking progestin-only or combined ECPs, she should take another dose. If she vomits within 3 hours of taking ulipristal acetate ECPs, she should take another dose. (She can use antinausea medication with this repeat dose, as above.) If vomiting continues, she can take a repeat dose of progestin-only or combined ECPs by placing the pills high in her vagina. If vomiting occurs more than 2 hours after taking progestin-only or combined ECPs, or 3 hours after taking UPA-ECPs, then she does not need to take any extra pills.

The conflict. A re-doses if vomiting occurs within **3 hours**; B (for the same progestin-only or combined pills) within **2 hours**.

The cervical dilatation that starts active labour

Query. *At what cervical dilation should partograph plotting commence?* q_00560

WHO Midwifery Education Modules (p.65)

Progress of labour

Cervical dilatation The first stage of labour is divided into the latent and active phases.

The latent phase at the onset of labour lasts until cervical dilatation is 4 cm and is accompanied by effacement of the cervix. The latent phase may last up to 8 hours, although it is usually completed more quickly than this. Although regular assessments of maternal and fetal well-being and a record of all findings should be made, these are not plotted on the partograph (using modified version) until labour enters active phase.

The active phase of the first stage of labour starts when the cervix is 4 cm dilated and is completed at full dilatation, i.e. 10 cm. Progress in this phase is approximately 1 cm per hour and often quicker in multigravidae.

WHO Labour Care Guide (p.11)

When should the LCG be initiated?

When women have entered the active phase of the first stage of labour (i.e. cervical dilatation of 5 cm or more).

The conflict. A puts the active phase (and partograph plotting) at **4 cm**; B at **5 cm or more**.

How long after birth a placenta is “retained”

Query. *How long after birth is a placenta considered retained?* q_00761

WHO Midwifery Education Modules (p.78)

haemorrhage bleeding.

If the bleeding is atonic it is important to find out if the placenta has been delivered or not. When the third stage is managed actively, the placenta is normally delivered within 5–10 minutes of birth of the baby. It takes longer to separate and deliver during physiological management, 20–30 minutes.

The students should remember that the placenta is retained if it has not been delivered within one hour of the delivery of the baby.

If the placenta has not been delivered, is there bleeding?

NICE intrapartum guideline (p.69)**Prolonged third stage**

1.10.20 Diagnose a prolonged third stage of labour if it is not completed within 30 minutes of the birth with active management or within 60 minutes of the birth with physiological management. Follow the recommendations on managing a retained placenta. [2014]

The conflict. A calls a placenta retained after **1 hour** regardless of management; B uses **30 minutes** (active) / **60 minutes** (physiological).