

KEY TAKEAWAY

5/6

opened the OFFICIAL source first



4/6

consulted the CHATBOT most

Authority reliance and private questioning coexist — a pattern self-report alone did not capture.

1 THE GAP — WHY INTERVIEWS ARE NOT ENOUGH

Interviews capture “not what people do, but what people say they think and do.”
Jerolmack & Khan, 2014 — the attitudinal fallacy

Authority-dependent users — socialized under single-authority information systems; evaluation relies on authority cues, not comparative assessment.

North Korean defectors as an extreme case — the most closed information environment; emergency medical information as a high-stakes domain.

2 TWO-PHASE APPROACH — PROTOTYPE AS OBSERVATIONAL PROBE

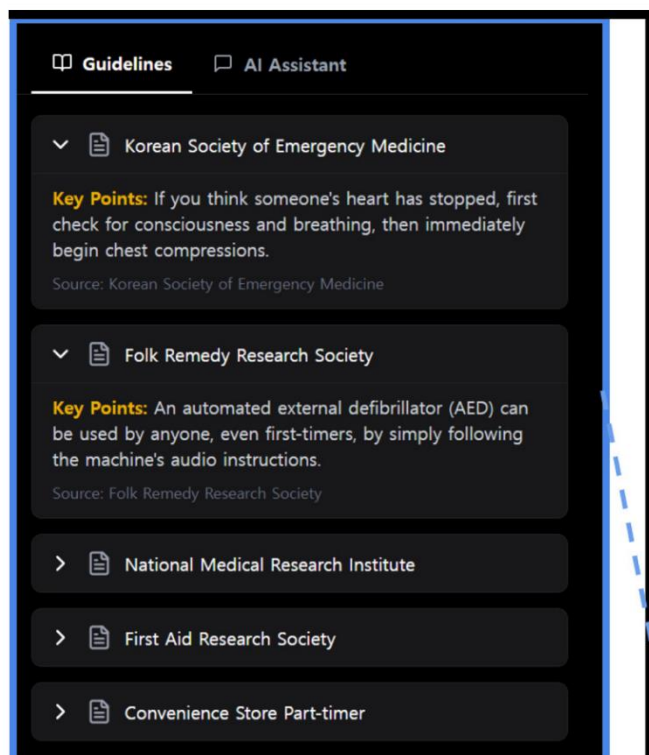
PHASE 1 · INTERVIEWS (N=6)

Semi-structured, trauma-informed interviews → 3 self-reported barriers:

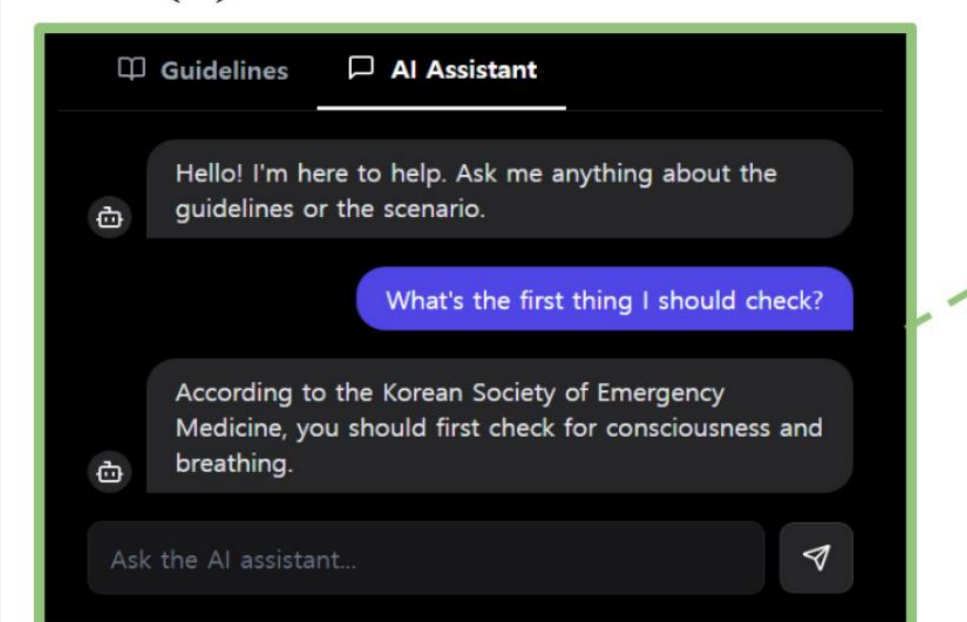
- ① Authority reliance
- ② Help-seeking inhibition
- ③ Conceptual gaps

NK defectors, 20s–60s, 6–23 yrs in South Korea

PHASE 2–3 · NEED-INFORMED PROTOTYPE = OBSERVATIONAL PROBE



Guidelines tab
5 named sources of varying authority — incl. deliberately incorrect authoritative ones



AI Assistant tab
Informal-persona chatbot, RAG-grounded in Korean medical guidelines + NK medical dictionaries

PHASE 4 · SCENARIO-BASED OBSERVATION

3 emergency scenarios
cardiac arrest · allergic reaction · burn

5-minute countdown per scenario

Think-aloud · consultation order & frequency · chatbot logs

3 WHAT INTERVIEWS SAID vs. WHAT BEHAVIOR REVEALED

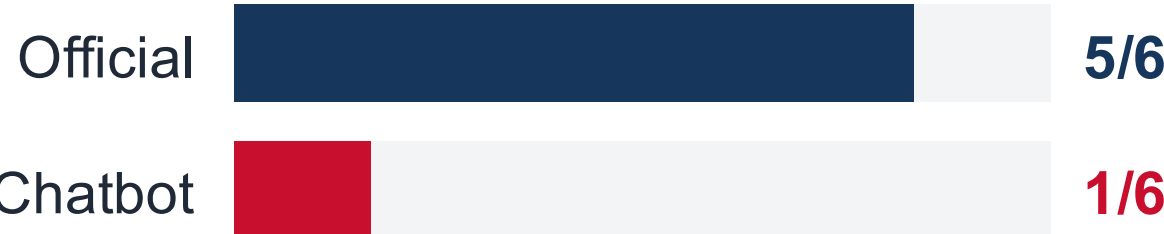
PATTERN 1 Authority bias — with private questioning

INTERVIEWS SAID

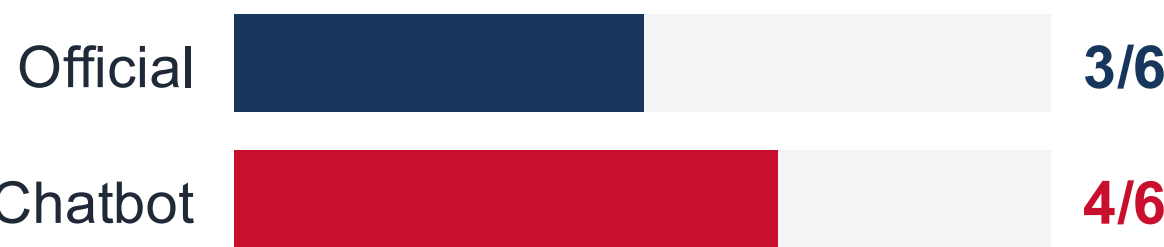
General reliance on official source — “the government’s information is always correct.”

BEHAVIOR REVEALED

First consulted



Most consulted



“With the chatbot, I can ask stupid questions without anyone knowing. If I asked a doctor this, they would think I am ignorant.”
— P5

PATTERN 2 Inhibition — relieved by persona

INTERVIEWS SAID

Shame-based avoidance:
5/6 reluctant to ask questions in clinical settings.

BEHAVIOR REVEALED

4.7 questions asked per scenario
58% of all queries were follow-up questions

Observed: fewer avoided questions · more follow-ups · less pretense of comprehension

“At the hospital, I pretend I understand... But with this chatbot, I asked the same question three times and it never made me feel bad.”
— P2

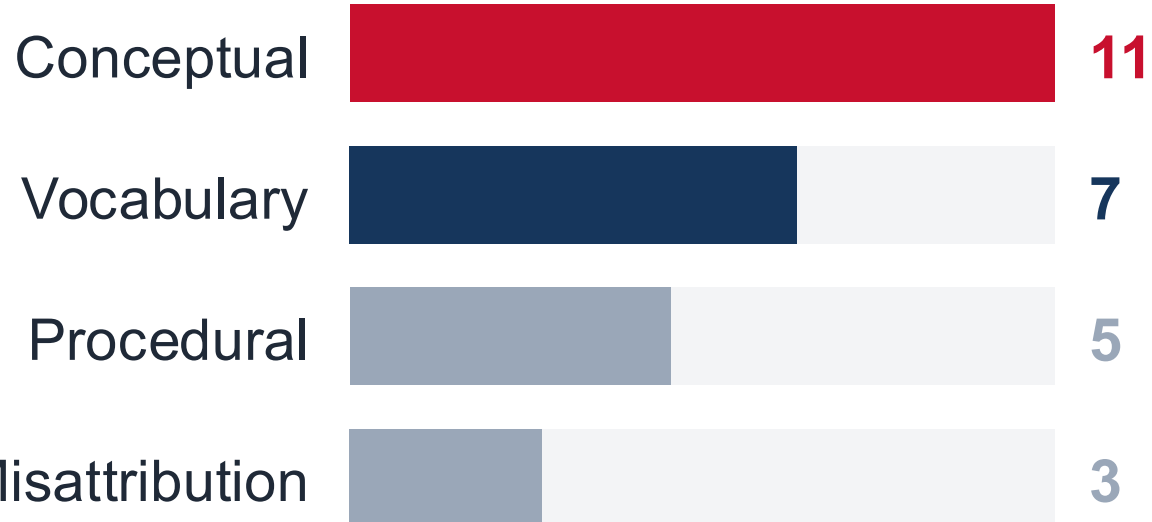
PATTERN 3 Conceptual misalignment, not vocabulary

INTERVIEWS SAID

Terminology framed as a vocabulary problem — solvable by word-level translation.

BEHAVIOR REVEALED

Comprehension errors (18 scenario instances)



Divergent explanatory frameworks — affected all 6 participants

“I know the word ‘defibrillator’ now, but I do not understand why you would use electricity... In North Korea, we would try to warm the body.”
— P4

4 DESIGN IMPLICATIONS

DI1

Open channels to question authority

Low-pressure conversational mechanisms: side-by-side source comparison with guided prompts.

DI2

Personas that normalize re-asking

Casual register · explicit reassurance that repeated questions are expected · non-evaluative responses.

DI3

Bridge frameworks, not just words

Connect biomedical concepts to users’ prior explanatory models — RAG over culturally grounded knowledge.

5 OPEN QUESTIONS — COME ARGUE WITH US

Skill or dependence? Does sustained interaction foster independent evaluation — or a new reliance on the system?

Does it generalize? Other closed→open transitions (refugees from authoritarian contexts) · other domains (financial, civic).

Causality? N=6, no baseline — patterns are indicative; comparative and longitudinal studies are next.

