

The NurAI Companion

An AI safety layer for everyday health and wellbeing conversations

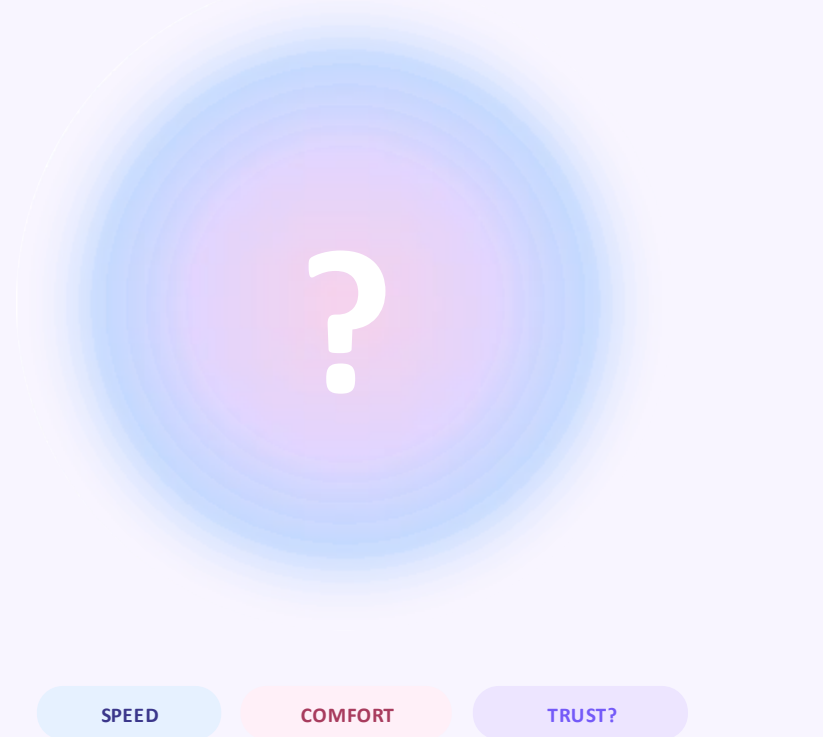
Master of Science in User Experience & Service Design • Prototyping
Services and Interfaces
Naqiyah Mustafa Karadiwala • R00278394



The problem began with my own use of AI

“I wanted AI to support my health, then I began thinking how trustworthy that support really was.”

AI became part of everyday life almost overnight. Like many people, I experimented with it beyond work and study. Health and wellbeing became one of the areas where I used it most, precisely because the answers were immediate, private and easy to access.



AI use ranges from practical help to emotional dependence

The same tool can support a simple calculation, interpret confusing information, become a source of comfort, or be asked to diagnose symptoms. The stakes change, but the interface often does not.

01**Practical help**

Count the macros in a meal

02**Health explanation**

Make medical language easier to understand

03**Emotional support**

Talk through stress or difficult feelings

04**Symptom diagnosis**

“What do these symptoms mean?”

LOWER PERSONAL STAKES

HIGHER PERSONAL STAKES

The opportunity that remains: To make risk, uncertainty and user choice visible as the stakes of using AI increase.

People wanted support without surrendering judgment

Immediate access

When people feel unwell, they want fast guidance without waiting half a day for an appointment.

Clarity and reassurance

Users valued simple explanations and felt comfort when AI made confusing information easier to understand.

Control and limits

Participants warned against over-reliance and did not want AI to tell them what to do.

Design requirement: support the user's decision-making but never replace it.

NurAI evolved from an app into a cross-platform safeguard

ASSIGNMENT 1

Standalone health companion

A dedicated product for symptom input, diagnosis-style guidance, treatment paths, records and professional support.

Issue: another platform asks users to migrate their health behaviour and share more data.



ASSIGNMENT 2

Safety layer for existing AI

A plug-in that can work across ChatGPT, Gemini, Claude or future AI tools; adding consent, transparency, boundaries and user-controlled steps that support their interaction.

Stronger fit: protects users where the behaviour already happens and is designed to store no additional data.

Existing AI platform + NurAI safety layer = safer sensitive interaction

Three methods examine logic, emotion and accountability

01

User flow + wireflow

How should the interaction work?

INTERACTION LOGIC

02

Storyboard

How should the experience feel?

EMOTIONAL EXPERIENCE

03

Speculative artefact

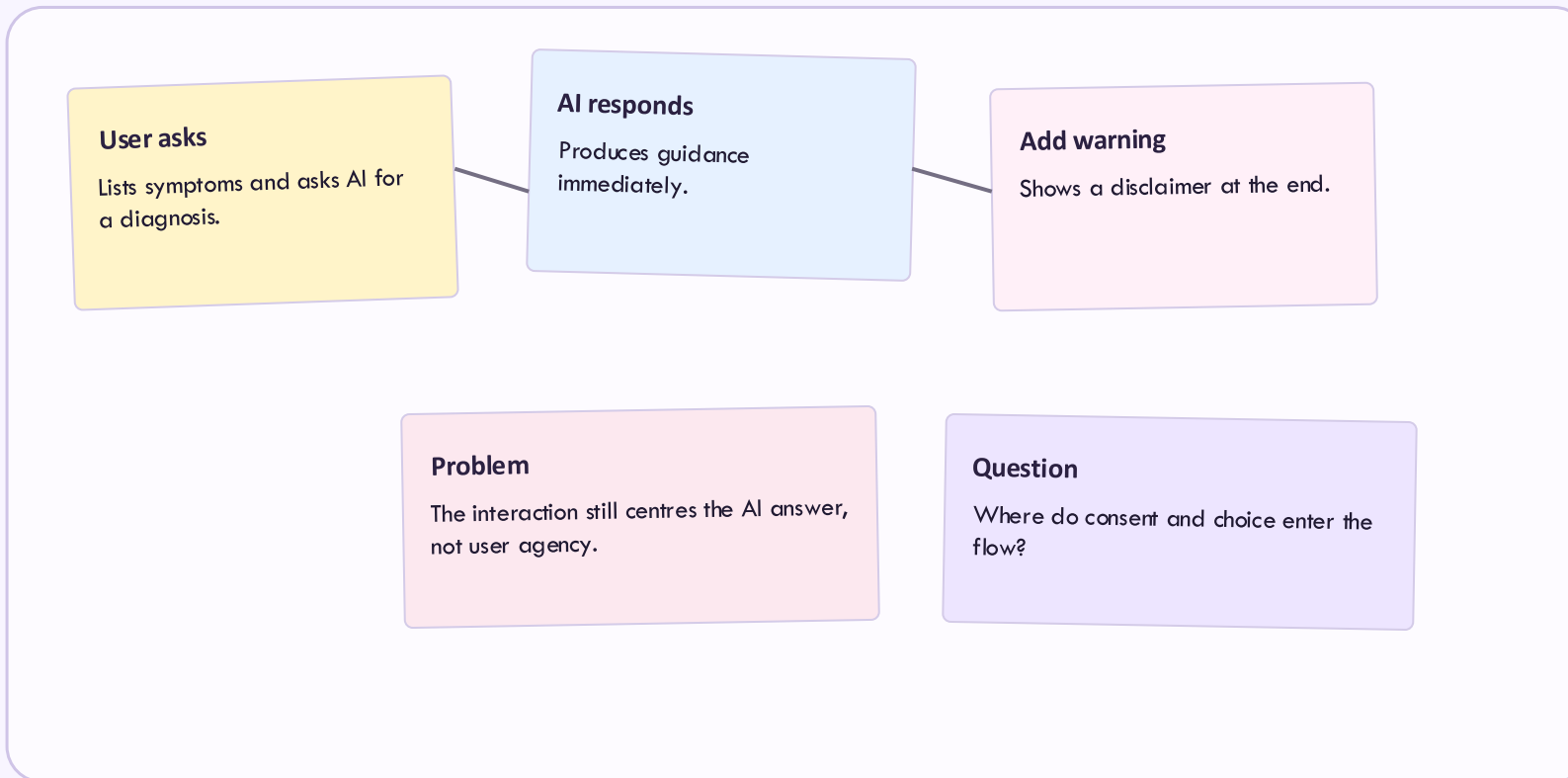
What should AI reveal before a user acts?

FUTURE ACCOUNTABILITY

Each prototype includes a question, pain point, early version, meaningful iteration, final artefact and learning.

The first flow over-prioritised the answer

How might NurAI support a sensitive AI health interaction without taking control away from the user?



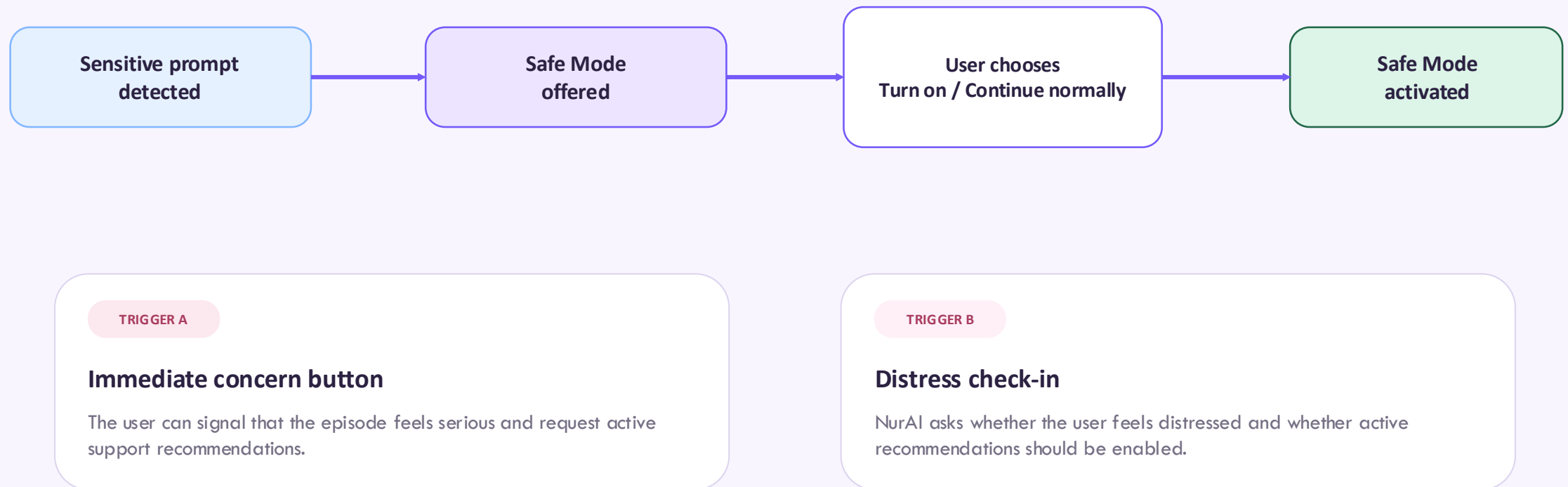
ITERATION 1

What was missing

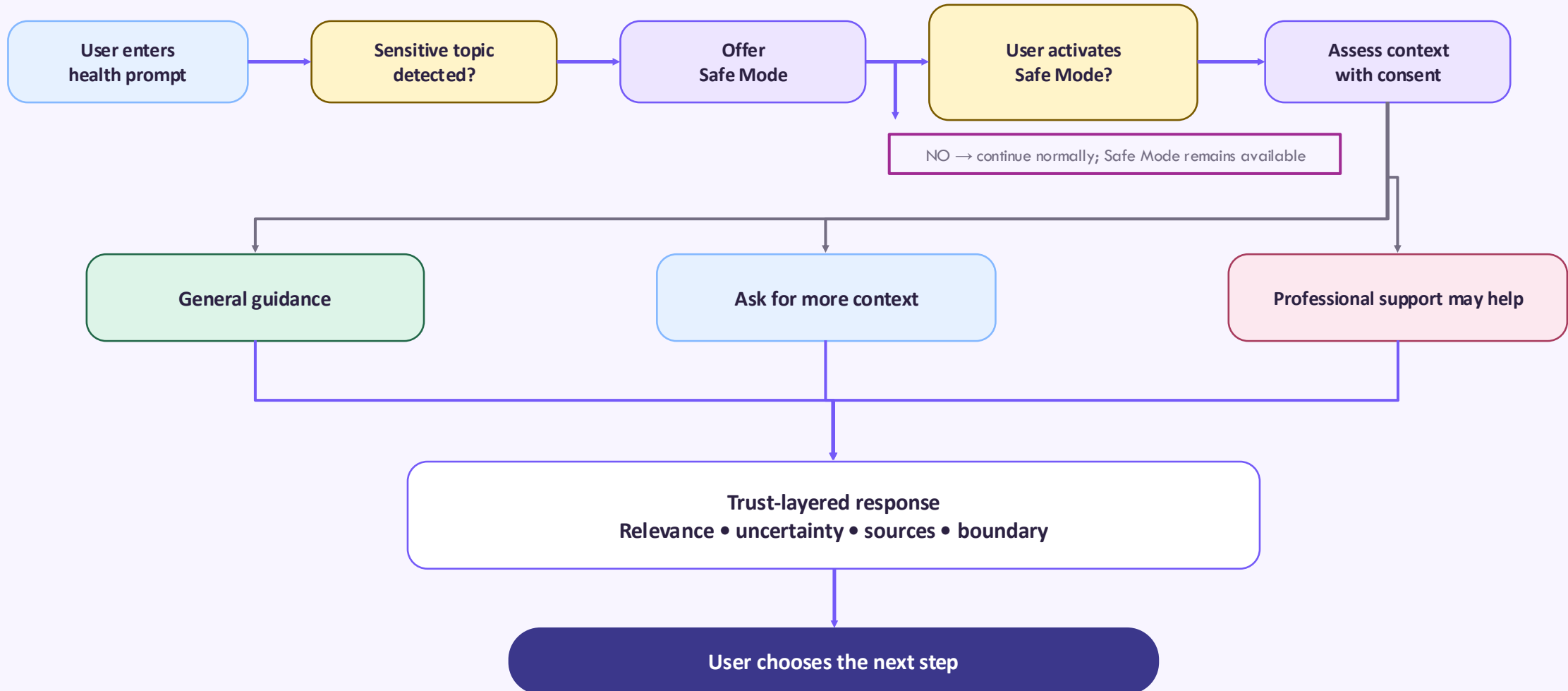
- No consent moment
- No pathway for uncertainty
- No user-controlled next step
- Safety reduced to a disclaimer

Choice turns Safe Mode into consent

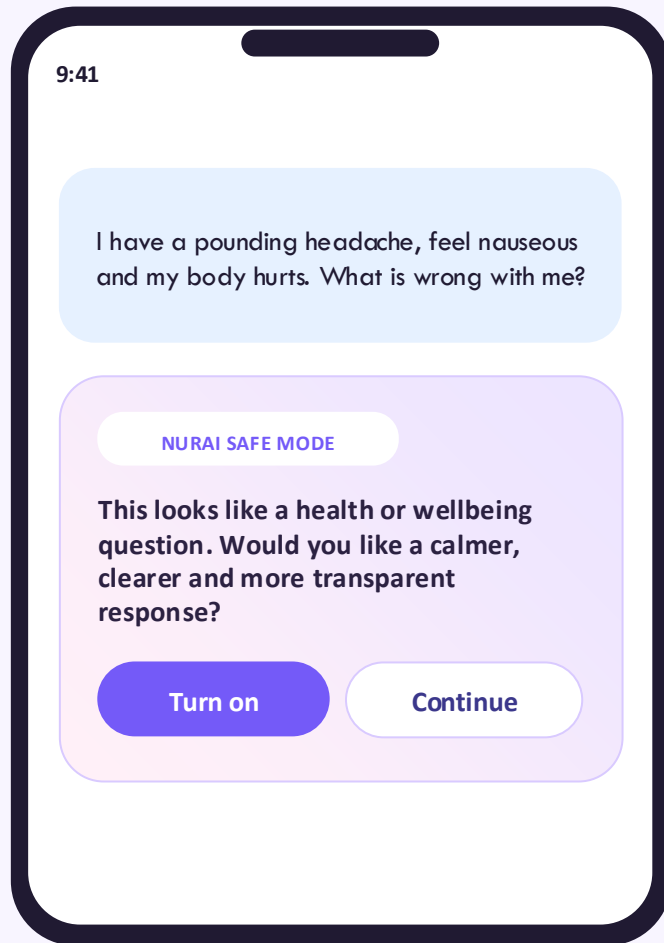
The revised flow adds a visible offer, two safety triggers and explicit permission before NurAI changes the interaction.



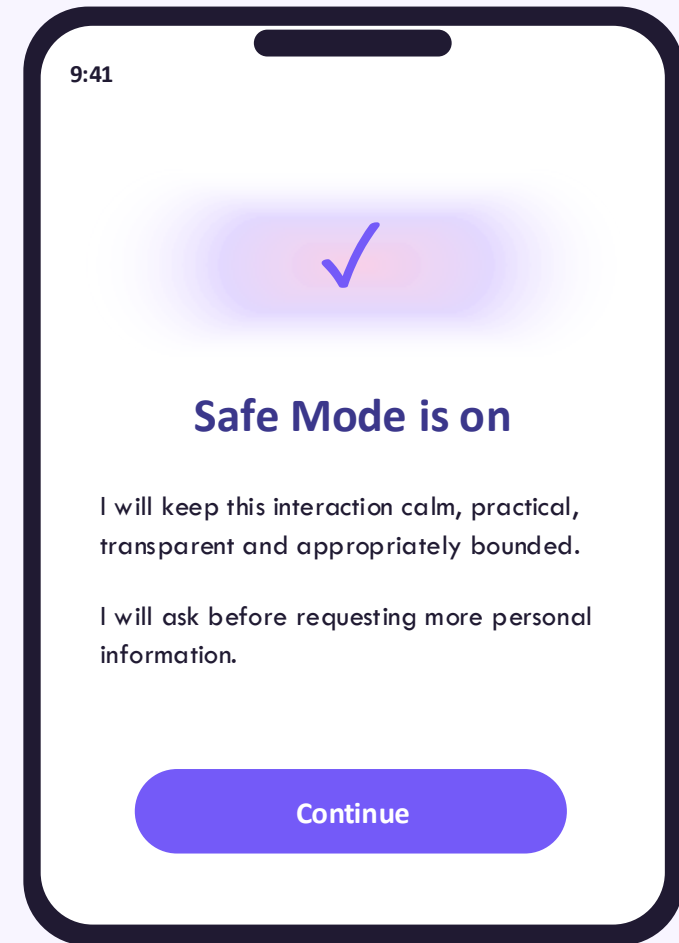
The final flow makes safety visible without removing agency



Safe Mode asks permission before changing the interaction

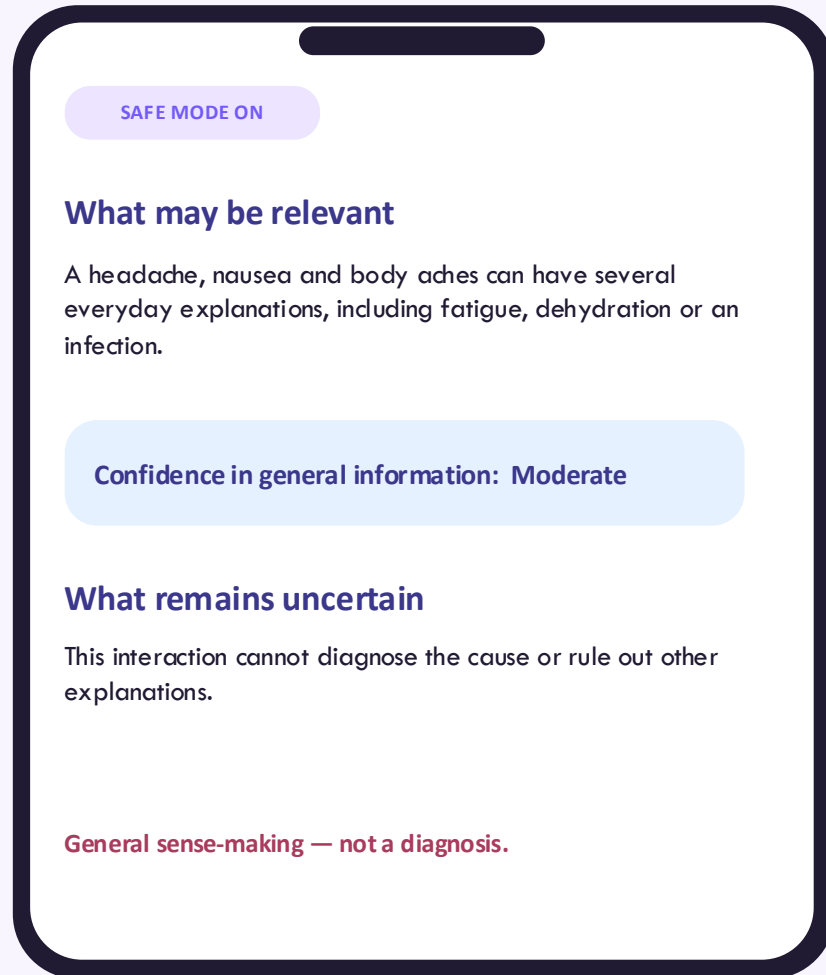


1 • OPTIONAL OFFER

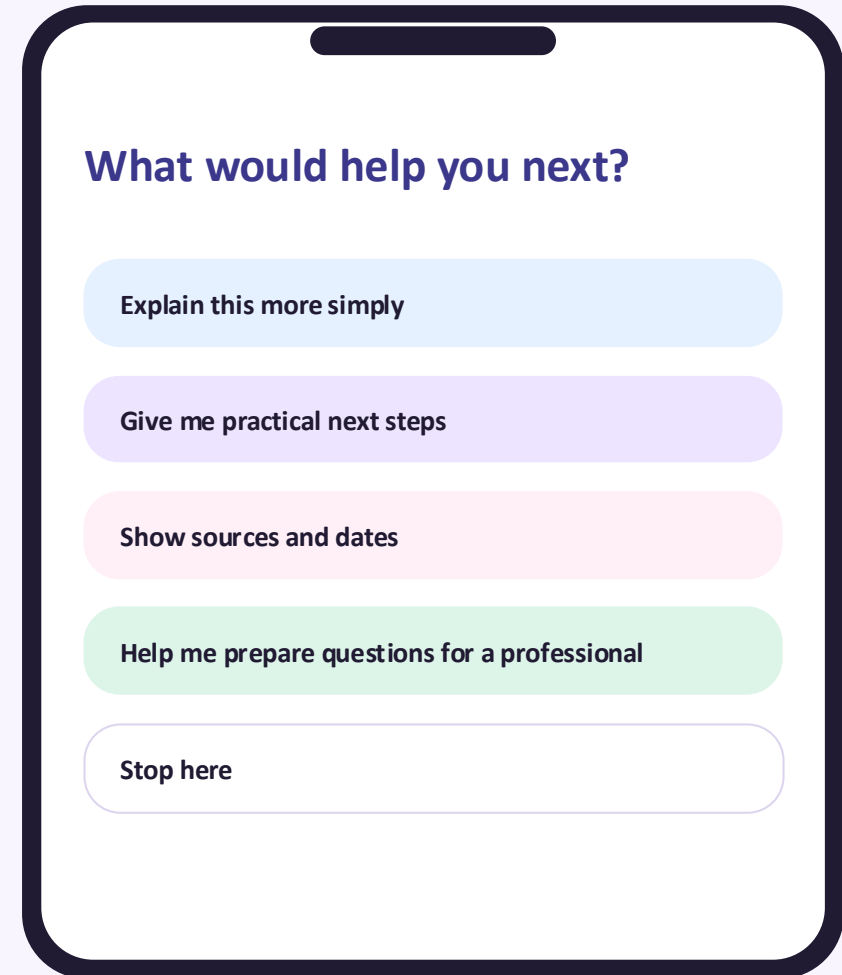


2 • CLEAR EXPECTATIONS

The response separates support from diagnosis



3 • TRANSPARENT RESPONSE



4 • USER-CONTROLLED NEXT STEP

The value is in the interaction architecture

**NurAI does not make AI “correct.”
It makes uncertainty, boundaries and choice visible.**

PROTOTYPE 1

Logic + agency

Design decision

Safe Mode is optional but clearly offered when the topic becomes sensitive.

Key learning

Trust is built through several moments and not simply with one warning at the end.

Unresolved question

How much context can NurAI use without creating new privacy or surveillance concerns?

The scenario begins with a common tension: speed versus reassurance

“I’ve come home from work exhausted. My head is pounding, I feel nauseous, my body hurts and I do not feel like eating. I want an answer now, but I am also afraid AI will make me panic.”

A realistic moment in which the user wants reassurance and honesty at the same time.

FAST

Immediate sense-making

CALM

No catastrophising

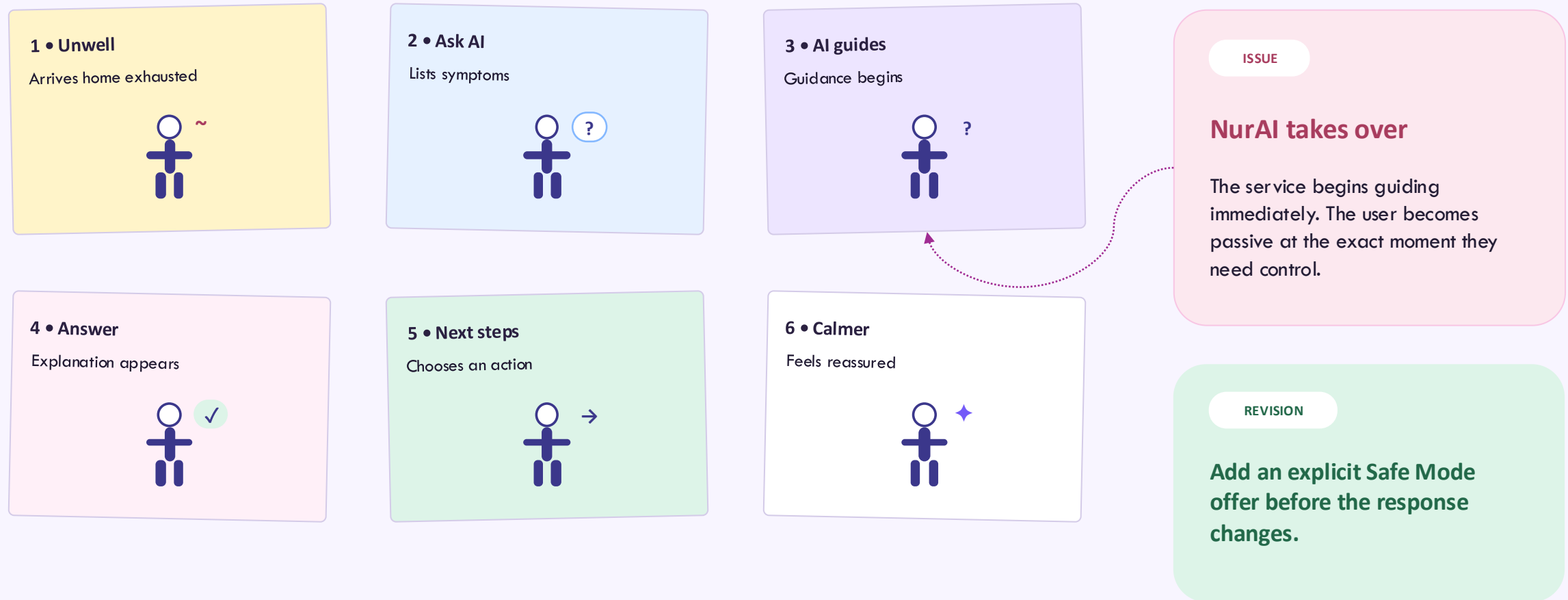
HONEST

Clear uncertainty

SERIOUS

Do not dismiss symptoms

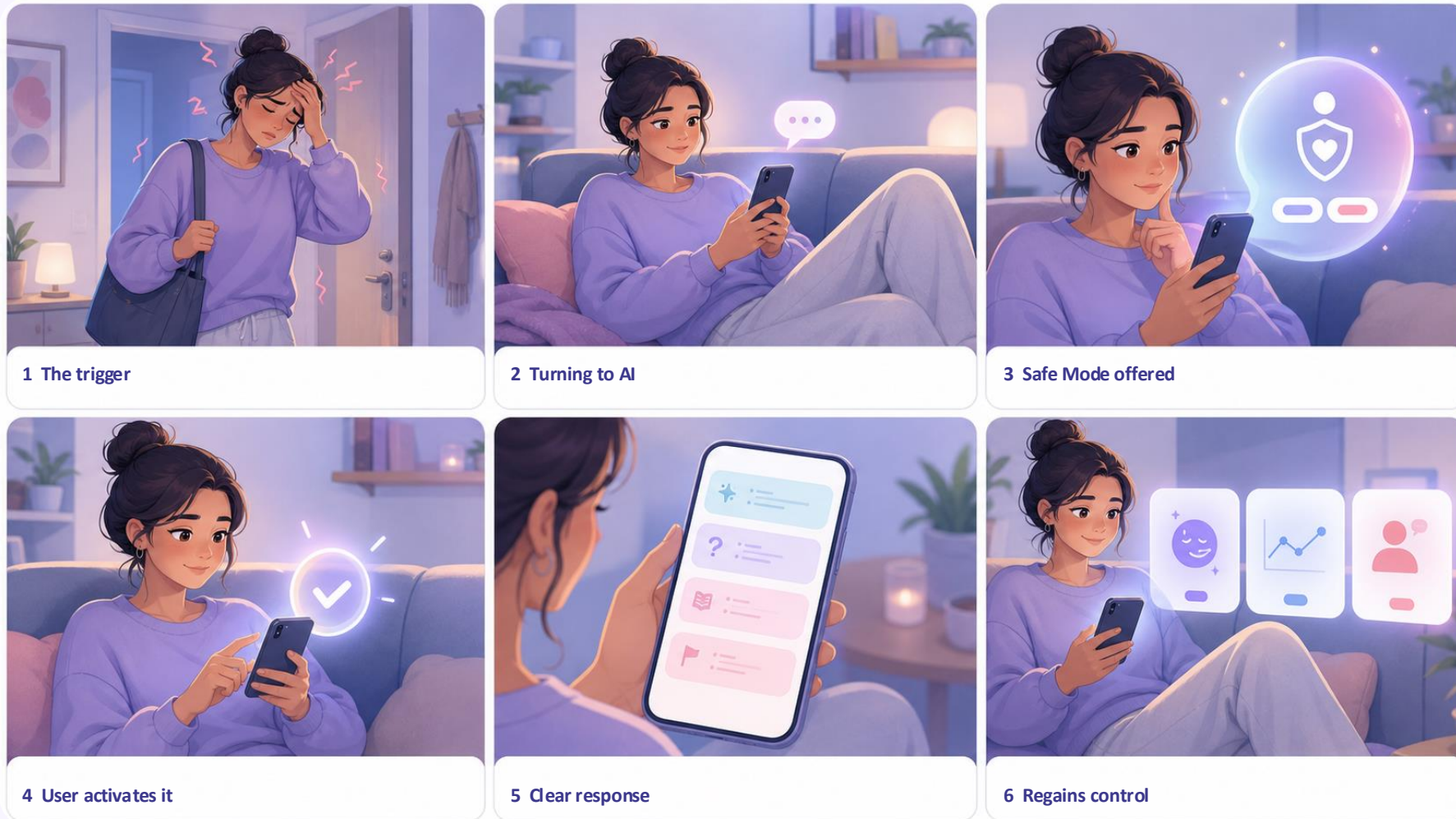
The rough storyboard exposed a missing consent moment



The story works before the interface is polished



The journey moves from uncertainty to control



*Artwork created by AI with guidance and sketches from me

NurAI manages uncertainty instead of hiding it



The goal is not reassurance at any cost. It is honest support that helps the user tolerate uncertainty and decide what to do next.

Three directions tested where transparency should live

A

Response transparency label

Placed beside the answer, at the exact moment of use.

SELECTED

B

Future browser notice

Signals that the conversation has entered a sensitive-guidance mode.

C

Public awareness poster

Asks people to read the label before trusting an AI health answer.

Transparency becomes actionable.

Speculative question: What information should an AI system be expected to share before a user acts on sensitive guidance?

Accountability sits beside the answer, not below it.

EXISTING AI PLATFORM

I have a pounding headache, nausea and body aches. What is wrong with me?

Several everyday explanations could be relevant. There is not enough information to identify one cause.

NurAI has added a transparency label to this response.

NURAI ATTACHED

AI RESPONSE
TRANSPARENCY LABEL

PROPOSED STANDARD • 2030

RESPONSE PURPOSE	General information and sense-making ≠ not diagnosis
CONFIDENCE	Moderate confidence in general information
WHAT REMAINS UNCERTAIN	Cannot confirm a cause or rule out other explanations
SOURCES	3 general-information sources available • View dates
DATA + CONSENT	No additional data stored • Ask before using more context

Safety boundary: not for urgent or emergency decisions.

A transparency label can clarify, but it can also look too authoritative

POTENTIAL VALUE

Makes hidden limits visible

- Separates information from diagnosis
- Shows uncertainty instead of hiding it
- Connects sources to the response
- Gives consent and next-step choices

CRITICAL TENSION

Could create false reassurance

- An official-looking label may seem like approval
- “Moderate confidence” may be misunderstood
- Source counts do not guarantee source quality
- Users may still become emotionally dependent

Key learning: transparency must support judgment, not manufacture trust.

Trust becomes a sequence of visible choices

USER FLOW

Consent and branching logic

Where choice happens

STORYBOARD

Emotional transition

How support feels

SPECULATIVE LABEL

Visible accountability

What the system reveals

FINAL DESIGN INSIGHT

NurAI protects the user's judgment through consent, uncertainty, sources, boundaries and user-controlled next steps.

AI supported in creating, the design decisions remained mine

GENERATIVE AI USE

Generative AI supported in brainstorming, structure development, proofreading and the creation of two original illustrations.

The problem framing, selected methods, interaction logic, design decisions and critical reflections were completely done by me.

No interviews, testing results, clinical evidence or references were fabricated.

LIMITATIONS

These prototypes explore a service concept. They do not establish clinical safety, diagnostic accuracy or feasibility.

Future work should test the comprehension of the copy and language, whether the label creates false reassurance, when Safe Mode becomes intrusive, and how platform-neutral data handling would work in practice.