

Aligning Digital Nudges with Design Features for Improved Behavioral Outcomes

A field manual for applying the Digital Nudging Framework to health-app design — built for design system leads, applied practice, and academic reference. Tables and checklists over prose; the interactive web tool and Figma library carry full depth.

6 MODES

23 MECHANISMS

5-Q ETHICS AUDIT

SECTIONS A–E

What's in this manual

Twelve pages, mostly tables. Use it as a desk reference alongside the interactive web tool — this manual is deliberately not a substitute for the worked examples, study evidence, and role assignments that live there.

01 How to Use This Manual — choosing Rapid vs. Full, and the one rule that never compresses

02 Framework at a Glance — Sections A-E and the diagnosis-first architecture

03 Rapid Guide — the 1-hour behavioral brief

04 Six Modes — cognitive-state diagnosis reference

05 23 Mechanisms — grouped by Spark / Facilitator / Signal, with risk ratings

06 Ethics Audit — the five non-negotiable questions

07 Evaluation Gate — the Nudge Effectiveness Score decision rule

08 Domain Quick-Reference — hydration, posture, mindfulness, activity, sleep

09 Full 10-Step Index — platform, duration, section, output at a glance

10 Academic Reference — study design, sample sizes, and citations

Two depths, one framework

Both paths use the same evidence base — the difference is depth, not substance. Default to Rapid for anything shipping this sprint; escalate to Full when the situation below applies.

USE THE RAPID GUIDE WHEN

Established product, tight sprint, solo designer or small team, prior journey-mapping evidence already exists.

USE THE FULL 10-STEP PROCESS WHEN

New product or unfamiliar domain, full team involved, first pass through the framework, or a mechanism is flagged high-risk.

The one rule that never compresses

GOVERNING PRINCIPLE

Never apply a cognitively demanding (System 2) intervention to a user in System 1 mode, and never apply a pure System 1 cue to a user who needs System 2 deliberation for genuine commitment. Every table in this manual exists to help you make that one call correctly and quickly.

Reading the tables

Colour always means the same thing across this manual:

COLOUR	MEANING
 Green	Low risk / ready state — safe to proceed with standard care.
 Amber	Medium risk / transitional state — proceed, but check the ethics audit before shipping.
 Red	High risk / vulnerable state — requires senior sign-off or is excluded outright.

FULL DEPTH LIVES ONLINE

Worked examples, study-evidence callouts, Miro board templates, and the interactive audit and NES calculator are in the companion web tool. This manual is the print-desk companion, not a replacement.

Diagnosis before selection

The framework's core innovation inverts how most nudge taxonomies are organised. Existing frameworks start from the mechanism (taxonomy → prescription). This one starts from the user's cognitive state.

STEP	DOES WHAT	ANSWERS
1	Diagnosis	Which of the six cognitive-state modes is the user in at this touchpoint? (Section A)
2	Nudge Selection	Given that mode, which trigger category — Spark, Facilitator, or Signal? (Section B)
3	Ethical Check	Does the mechanism pass all five audit questions? (Section C)
4	UI Translation	What specific interaction component implements this? (Section D)
5	Evaluation	Does the deployed nudge clear the NES gate? (Section E)

Theoretical foundation

Fogg Behavior Model

Categorises triggers into Sparks (motivation), Signals (prompts), and Facilitators (ability/simplicity).

Dual-Process Theory (Kahneman)

Differentiates System 1 (fast, automatic) from System 2 (slow, reflective) cognitive processing.

Nudge Theory & Choice Architecture (Thaler & Sunstein)

Preserves user autonomy while altering choice environments — the ethical backbone of Section C.

Why it matters practically

Nudge effectiveness is not a function of mechanism sophistication alone — it's the interaction between mechanism, cognitive state, and timing.

88

STUDY 1 USERS, 3 DOMAINS

38

STUDY 2 DESIGNERS

10

STUDY 3 PRACTITIONERS

The 1-hour behavioral brief

Compresses Steps 2–7 of the full process. Skips onboarding, stakeholder templates, and the Figma library tour. **Never skips the ethics gut-check.**

MINUTES	STEP	DO
0–10	1 • Diagnose	Pick the closest-fitting mode from the six-mode table (page 6). Go with the fastest honest read, not a debate.
10–20	2 • Select trigger	Motivation absent → Spark. Ability blocked → Facilitator. Both present, no cue → Signal. Unsure → Facilitator first.
20–35	3 • Pick mechanism	Scan that trigger's table (page 7). Choose 1–2. Skip anything red-flagged unless you can justify it on the spot.
35–45	4 • Ethics gut-check	Run the five questions on page 8. Any single "no" halts progression — simplify the mechanism, don't argue past it.
45–60	5 • Ship & flag	One line each for mode, mechanism, and audit result. Flag for the full Step 7 audit before public launch.

Escalate to the full 10-step process when...

- This is a new product or unfamiliar health domain — no prior journey-mapping evidence to diagnose from.
- A mechanism is flagged high-risk and the justification isn't resolvable in the moment.
- The ethics gut-check fails and the fix isn't obvious.
- A stakeholder is pushing for something the gut-check would reject — use the Step 8 templates to formalise the case.

CATEGORICAL NO-GOS, SPRINT OR NOT

Manufactured urgency, obstructed opt-out, shame-based framing, competitive social comparison, and placebo progress signals are excluded regardless of time pressure. Speed compresses process, not ethics.

Six modes — cognitive-state reference

Section A. Diagnose before selecting a mechanism — this is the framework's core inversion of standard nudge-design practice.

#	MODE	PRIMARY BARRIER	START WITH
● 1	Unaware	No mental model of the behavior as relevant.	Facilitator-first, or a mild identity Spark. Signal is premature.
● 2	Undecided	Present bias — reward is distant, effort is immediate.	Spark with identity framing + future-self salience. Avoid loss-aversion framing.
● 3	Motivated-Stuck	Implementation-intention gap or procedural friction.	Facilitator: remove steps, add defaults, inline scaffolds.
● 4	Primed	Cue absence — conditions are right, user hasn't been signalled.	Signal: contextual, just-in-time trigger. The classic Fogg case.
● 5	Forming	Habit fragility — routine hasn't reached automaticity.	Low-intensity Signal to reinforce the loop; gentle Spark for milestones.
● 6	Retreating	Self-efficacy collapse + shame-avoidance reactance.	Compassionate Facilitator + identity-repair Spark. No streaks, no shame.

MODE 6 IS THE EXCEPTION MODE

Retreating is not a motivational-absence state — it's active shame-avoidance after a lapse. Streak-preservation loss framing, competitive comparison, and shame-based copy are categorically excluded here, not just discouraged.

23 mechanisms, grouped by trigger

Section B. Match the mode's barrier to a trigger category, then choose the lowest-risk mechanism that fits.

Spark — motivation absent

MECHANISM	RISK	PRIMARY BIAS
Gamified challenges & rewards	● Medium	Reinforcement, reciprocity
Social proof (verified aggregate)	● Low	Herd instinct, social norming
Social incentives / competition	● High	Herd instinct, image motivation
Identity-affirming framing	● Low	Self-concept consistency
Loss-framed messaging	● High	Loss aversion
Public commitment	● Medium	Commitment bias
Scarcity / urgency framing	● High	Scarcity bias
Reciprocity-based prompts	● Low	Reciprocity bias

Facilitator — ability blocked

MECHANISM	RISK	PRIMARY BIAS
Healthy default settings	● Low	Status quo bias, default effect
Simplified workflow / step reduction	● Low	Choice aversion, decision fatigue
Pre-configured goal templates	● Low	Anchoring, default effect
Autofill / smart suggestions	● Low	Default effect
Positioning / layout prominence	● Low	Salience bias
Inline scaffolding & hints	● Low	Cognitive load reduction
Hiding less desirable options	● Medium	Salience reduction

Signal — cue absent

MECHANISM	RISK	PRIMARY BIAS
Just-in-time contextual reminder	● Low	Priming, salience
Ambient environmental cues	● Low	Priming, mere exposure
Progress alerts / milestone notifications	● Low	Peak-end effect, reinforcement
Fixed-interval reminders	● Medium	Availability heuristic
Warning / disclosure cues	● Low	Loss aversion (protective)
Feedback loops	● Low	Reinforcement
Throttling mindless activity	● Low	Attentional reset
Placebo / illusory progress signals	● High	Reinforcement (hollow)

The five-question audit

Section C. Every mechanism, in every mode, at every speed of delivery, passes through this gate before shipping. **Any single "no" halts progression.**

- ☐ Does this serve the user's health goal, not the engagement metric?
- ☐ Is the opt-out as visible and friction-free as the primary action?
- ☐ Is all personalisation data explicitly consented to, with a clear purpose?
- ☐ Would the user, fully informed about the mechanism, endorse it?
- ☐ Does this preserve the user's sense of autonomous choice?

CATEGORICAL EXCLUSIONS — NO CONTEXT JUSTIFIES THESE

Manufactured urgency without genuine constraint · opt-out buried behind extra steps · shame-based or identity-threatening framing · competitive social comparison in vulnerable modes (2, 6) · placebo or illusory progress signals.

Decision rule

ALL FIVE → PASS

Proceed to UI translation (Section D).

ANY SINGLE → FAIL

Halt. Simplify or redesign the mechanism — do not argue past the check.

The Nudge Effectiveness Score gate

Section E (NEM/NES), operationalised and validated across 88 users × 3 domains in Study 1. NES penalises engagement bought at the cost of intrusiveness — a quiet nudge with modest clicks beats a loud one with high clicks and high perceived intrusiveness.

Ship / hold decision rule

CONDITION	THRESHOLD	IF IT FAILS
NES beats control	Higher than baseline	Hold — the mechanism isn't earning its intrusiveness cost.
Perceived intrusiveness	Below 3.0 / 5	Hold — reduce frequency or salience before re-testing.
Perceived autonomy	Above 3.5 / 5	Hold — revisit opt-out visibility and framing.

ALL THREE, NOT ANY ONE

Ship only when all three conditions hold simultaneously. A nudge that beats control on NES but fails the intrusiveness or autonomy threshold is not a pass — the framework treats those as hard gates, not trade-offs to average out.

Scope note

This is a per-condition snapshot, not a long-term effectiveness measure. Long-term outcomes (6-month retention, habit consolidation) depend on additional variables tracked separately, outside this manual's scope.

Health-domain specific guidance

Hydration, posture, and mindfulness are empirically grounded in Study 1. Physical activity and sleep are practitioner-validated extensions from the Study 3 co-creation workshop, not directly tested.

Hydration Study 1 · empirical

- Behavioral triggers (time since last log), not fixed-interval reminders — fixed intervals underperformed ~2× on conversion.
- Facilitator default: pre-set 8-glass goal, friction-free to adjust.
- Verified aggregate social proof only — avoid competitive comparison.

Posture Study 1 · empirical

- Inactivity detection as primary trigger — the same reminder reads "helpful" when paused, "aggressive" mid-task.
- Mild haptic signals outperform visual interruptions for System 1 attention.
- Cumulative progress feedback preferred over single-session metrics.

Mindfulness Study 1 · empirical

- Identity-threat is a distinct reactance here — avoid framing implying the user isn't "someone who meditates."
- Fixed-interval reminders during deep focus states are categorically excluded.
- Identity-affirming Spark with a contextual Signal only after the user opens the app or finishes a session.

Physical activity Study 3 · practitioner-validated

- Body-image sensitivity — centre function over aesthetics, avoid appearance-oriented framing.
- Facilitator-first, anti-competition; social proof as aggregate, never a leaderboard.
- Mode 6 (post-lapse) handling is critical — streak-preservation loss framing damages long-term adherence.

Sleep Study 3 · practitioner-validated

- Screen-sleep paradox: the medium causes the problem the app aims to solve. Prefer ambient, pre-bed cues.
- No engagement-oriented notifications in the pre-sleep window.
- Morning reflection (low-intensity Spark) over evening gamification.

The complete process, at a glance

Pointer index only — worked examples, study-evidence callouts, and board-setup instructions for each step are in the interactive web tool.

#	STEP	PLATFORM	DURATION	SECTION	OUTPUT
1	Onboarding	Miro	30–45 min	Scaffold A–E	Shared mental model
2	User & Journey Research	Miro	90–120 min	A, E	Annotated journey map
3	Cognitive State Diagnosis	Miro	45–60 min	A	Mode-annotated touchpoints
4	Bias Mapping	Miro	60–75 min	A, D	Bias → Need grid
5	Nudge Selection	Miro	60–90 min	B	Mechanism shortlist + rationale
6	Figma Ideation	Figma	90–180 min	B, E	Design variations per mechanism
7	Ethics Audit	Miro	45–60 min	C	Signed audit per design
8	Stakeholder & Business Alignment	Miro	60–75 min	—	Signed stakeholder brief
9	Interface Design Synthesis	Miro	Variable	A–E synthesis	Annotated design spec
10	Evaluation & NES	Miro	Ongoing, post-launch	C (NEM)	NES + adoption decision

Phase map

Discover & Define (Steps 1–5)

Miro-based. Onboarding through nudge selection — diagnosis-heavy, Sections A, D, B.

Develop & Deliver (Steps 6–10)

Figma for ideation (Step 6), then back to Miro for audit, alignment, synthesis, and evaluation — Sections B, C, E.

Study design & citations

For design system leads citing this framework internally, or researchers replicating the method.

STUDY	DESIGN	KEY FINDING
Study 1	Figma prototypes on Maze; n=88 across hydration (26), posture (25), mindfulness (37); 9 mechanisms per prototype	No universal nudge; cognitive state and health domain both moderate effectiveness.
Study 2	Mixed-methods questionnaire + 8 semi-structured interviews; n=38 professional UX/UI designers	76.3% behavioral awareness vs. 28.9% intentional application — the intentionality gap.
Study 3	Two 3-hour co-creation sessions, n=10 senior UX strategists; reflexive thematic analysis in MAXQDA — 440 coded segments, 44 categories, 6 parent categories, 4 themes	Diagnosis must precede selection; ethics requires objective, mechanism-based protocol.

Core PhD contributions

- **Empirical:** mechanism-specificity evidence across three health domains (Study 1).
- **Practitioner:** six documented practitioner-gap patterns (Study 2).
- **Methodological:** participatory co-creation → standardised design workflow (Study 3).
- **Applied:** Diagnosis Before Selection, Mechanism-Based Ethics, Sprint-Compressed Cognitive Design.

Scholarly dissemination

Khodadadeh, Y., & Jafari, S. (2025). Designing for behavior change: Exploring challenges and opportunities from a designer's perspective. *International Conference on Design (ICD)*.

Khodadadeh, Y., & Jafari, S. (2025). Exploring user preferences for health-focused digital nudges: Insights from a preliminary study. *Applied Human Factors and Ergonomics (AHFE) International Conference*.

Jafari, S., Khodadadeh, Y., & Love, T. (2026). Diagnosis Before Selection: A Co-Created Digital Nudging Framework for Health App Design. *Journal of Design Thinking (JDT)*.

LIMITATIONS TO NOTE WHEN CITING

Ecological validity of the scenario-based workshop; evidential asymmetry across the three studies; cultural transferability given predominantly Iranian and European professional networks; first-author involvement in developing, facilitating, and coding all three studies.

Contact: shahrzad.jafari@ut.ac.ir