

USER GUIDE

Delphi ecosystem — user guides

Plan · Facilitate · Analyse — one ecosystem, usable apart

English

Version 1.0 · 2026-09-24

Contact: Hannu Linturi · hannu.linturi@gmail.com

Contents

OVERVIEW

Ecosystem overview	6
Getting started	6
What do you want to do?	7
Three levels: operate → delegate → system	7
Descriptive vs. interpretive — a key boundary	8
Where to go next	8

INTEGRATION

Integration: Plan → Pronoia → DAE	9
Integration is the top of the capability arc	9
Stage 1 — Design (DPE)	10
Stage 2 — Collection (Pronoia)	10
Stage 3 — Analysis (DAE)	12
Closing the loop — learning across studies	13
The facilitator's end-to-end playbook	13

DELPHI VARIATIONS

Delphi variations	14
Templates for different purposes	14
Real-Time Delphi — how to run one	16
When to choose RTD	16
What the software does for you	16
Planning	17
While the study runs	17
Closing	18
What an RTD study does not have	18

PLANNING (DPE)

Orientation, variation and question architecture	19
Orientation — from which standpoint the future is studied	19
Variation — the cadence of rounds	20
Scoping the phenomenon — what is actually studied	20
Question architecture — the theses	21
When there is no real panel — simulation	21
Building the panel: epistemic space, perspective and gates	22
Orientation sets the panel type	22
The perspective layer: position + view	23
Perspective Health Score (PHS)	23
Epistemic coverage: $7 \times 3 = 21$ cells	24
Quality gates before collection	24

Building the panel in practice	25
Simulation: the Argument Simulator (Route B)	27
How it proceeds — six phases	27
Anomaly check — surfacing blind spots	28
Worked example — Food Paradigm 2050	28
 PANEL DESIGN	
Building the Delphi panel	30
The panel as a designed epistemic space	30
The two-axis matrix is primary	30
Orientation sets the panel type	31
Perspective: position + view (Hautamäki)	32
AI panels (Route B / C)	32
Reading the arguments: the epistemic lens	33
Panel diagnostics (coverage, PDI, lock-in, PHS)	33
Quality gates before collection begins	34
Building the panel in practice	34
 FACILITATION (PRONOA)	
The iterative process: rounds, phases and convergence	36
Why Delphi iterates	36
The Pronoa process at a glance	36
Round 0 — mapping and motivating (AI-assisted intake)	37
The phases inside a round	37
Controlled feedback: the visibility ladder	38
Policy presets — choosing an iteration stance	39
Revision and movement	40
The panelist Home — real-time data that fuels iteration	40
Tracking the panel — process metrics	41
Iterating across rounds — when to continue, when to close	41
The facilitator's iteration playbook	42
Anonymity: identity modes and group reporting	43
Principle	43
Three identity modes (per-panel setting)	43
Cross-cutting rules (always on, in every mode)	44
What is available in each mode	45
Group comparison — current behaviour	45
In practice: role play and the one-shot poll	46
Data model (implemented) — for the developer-partner	46
How Delphi Pronoa's two interfaces are built	47
1. One study, one truth	47
2. The panelist interface	48
3. The facilitator interface	50
4. The same round, from both sides	55

5. What the interfaces deliberately do not do	56
The panelist's process, round by round	57
1. What you are joining	57
2. Before anything opens: the invitation and the consent screen	58
3. Round 0: the orientation conversation	58
4. The frame that never changes	59
5. Your turn to answer	59
6. Answered, and the round is still open	60
7. Dialogue: the round where the panel talks to itself	61
8. Between rounds — the richest screen in the study	62
9. When the study ends	62
10. More services: the index	63
11. Two ways the same journey is named	64
12. If the study is real-time	64
13. Your answers, your data	65
14. For the facilitator: what is worth saying out loud	65
The facilitator's process, from planning to handover	67
0. The map you are working on	67
1. Planning: a completeness list, not a wizard	68
2. Round 0: orientation, run as a sequence	70
3. Collection: the round-management card	71
4. The response phase: your work is monitoring — and preparing	72
5. The dialogue phase: monitoring and moderation	73
6. The revision phase	74
7. Glancing at results without leaving your work	74
8. Between rounds: your richest screen too	75
9. Reporting and handover	76
10. Exceptions, and what they cost	77
11. What the panel sees while you work	78
12. Five things that cannot be undone	78
13. What is folded, and when	79
14. The line that says where on the arc you are	81
15. Retirement is the Barometer's way, not a general one	82
Role play: a role in a box of the panel matrix	83
What role play is for	83
Division of work: the facilitator gives the structure, the panelist chooses the box	83
The facilitator's steps	84
What the panelist sees	84
AI roles as panel supplements	85
Reading the results	85
Technical description	85
One-shot poll at a seminar: QR code and anonymous answers	87
What the one-shot poll is for	87
Before the session	87

During the session	88
Closing the poll	88
Results	89
Limits	89
Technical description	89
Taking part in the open Delphi study	91
No account needed	91
A Delphi study is not a survey	91
Anonymity is structural, not a promise	92
What else is on the page	92
If something doesn't work	92
Short answers	93
The facilitator trial in five stages	94
1. Open your own example study	94
2. See where the panel landed	94
3. Read the dialogue	94
4. Change something	95
5. Build a study of your own	95
The limits of the trial	95
Tell us what didn't work	96
 ANALYSIS (DAE)	
DAE analysis: the pipeline, phases and quality gates	97
Handoff in — what DAE reads	97
The pipeline: 21 checkpoints, eight phases	97
The quality engine — when the human enters the loop	98
Outputs	99
What if DAE is not in use?	99
 REFERENCE	
Glossary	100
Frequently asked questions (FAQ)	103
Role quick-guides	105
Researcher / designer	105
Facilitator	105
Developer-partner	105
Panelist	106
Decision-maker / stakeholder	106
Slides and deep dive	107
Presentation slides	107
Worked example	107

Ecosystem overview

The Delphi ecosystem is three products that form one pipeline but also work apart:

Planning (DPE)

Delphi Planning Ecosystem (DPE) — the research design: orientation, phenomenon, panel and question architecture.

Facilitation (Pronoia)

Pronoia, descriptive facilitation — running the panel: rounds, dialogue, the pause, and descriptive reports.

Analysis (DAE)

Delphi Analysis Ecosystem (DAE), interpretive analysis — in-depth analysis of the collected data under quality gates and human review.

The pipeline runs as handoffs: **Planning → Facilitation (Pronoia) → Interpretive analysis (DAE)**. Each part can also be used on its own.



Interactive ecosystem map: open the visual, clickable overview — tabs, phases, quality gates and foundations. Offline, projectable.

Getting started

Type `metodix` (in Claude Code `/metodix`) — it opens the Metodix start menu and routes you to the right part of the ecosystem (planning · pre-analysis · DAE analysis · single method · resume). Everything begins there.

What do you want to do?

Design a study (orientation, panel, questions)

→ Planning (DPE)

Run a panel: rounds, dialogue, reports

→ Facilitation (Pronoia)

Interpret collected data in depth

→ Analysis (DAE)

Understand how the parts connect end-to-end

→ Integration

Three levels: operate → delegate → system

You don't move up a level by gathering more knowledge — you move up by **ceasing to do something by hand**. That is the through-line of these guides: each part makes visible *the one next line whose crossing actually raises your level* — not "here is every feature".

LEVEL 1

Operate

The ecosystem is a tool you use. You stop doing nothing yet — you do it yourself, just faster: design theses, open rounds, read the results.

LEVEL 2

Delegate

You give the system tasks. Individual steps drop away: AI panelists, the R0 intake chat, PIRE round-gap suggestions, the AI helpdesk. You confirm.

LEVEL 3

System

Infrastructure your work runs on. The whole chain leaves your hands: the plan flows as handoffs into facilitation and analysis (Plan → Pronoia → DAE).

The methodological guardrail stays. Delegation does not mean losing control: AI suggestions are drafts (they require an audit), interpretation happens in DAE under quality gates, and a panelist's speech/dictation is edited before sending. The arc runs *operate → delegate → system*, not *human out of the loop*.

Descriptive vs. interpretive — a key boundary

The most methodologically important boundary runs between Pronoia and DAE:

	Facilitation (Pronoia)	Analysis (DAE)
Role	Descriptive facilitation	Interpretive analysis
Question	<i>What did the panel say and how did views move?</i>	<i>What does it mean?</i>
Output	Distributions, central tendencies, consensus/dissensus, arguments as-is	Theming, scenarios, causal layers, synthesis
Authority	Source material	Methodological interpretive authority

A Pronoia report is not an analysis — it is the source material for one. "What it means" belongs to DAE, under quality gates (κ / QVG) and human review.

Where to go next

- **As a designer** → Planning (DPE) (orientation, panel, questions)
- **As a facilitator** → Facilitation (Pronoia) (rounds, dialogue, reports)
- **As an analyst** → Analysis (DAE) (CP-0...CP-20, quality gates)
- **For the whole picture** → Integration guide (handoffs end-to-end)

Integration: Plan → Pronoia → DAE

A panelist experiences one study — its rounds, its feedback, its Home view. **The facilitator runs something larger.** Their process spans three stages — design, collection and analysis — and reaches across two connected ecosystems. Pronoia is the bridge between them.

Upstream sits the **Delphi Planning Ecosystem (DPE)**, where the study is designed. Downstream sits the **Delphi Analysis Ecosystem (DAE)**, where the collected material is interpreted. **Pronoia** is the data-collection layer in the middle: it receives a finished design, runs the study, validates the result, and hands a clean dataset on to analysis.

Why this matters for the facilitator. The quality of the analysis is decided long before analysis begins — in how the study is designed and how cleanly it is collected. Owning the whole arc, and the two handoffs that join it, is the facilitator's core responsibility.

Integration is the top of the capability arc

Assembling a single study by hand is **operating**. When the handoffs carry the design from stage to stage without rebuilding, the whole chain becomes a **system** — infrastructure the work runs on. The through-line of this guide is exactly that shift: you stop re-assembling the study in every tool, and let the handoffs flow.

HANDOFF 1

Design in

The Planning Handoff Package carries orientation, the panel matrix, the theses and the perspective context into Pronoia. You don't rebuild the study — you continue it.

COLLECTION

Pronoia runs

Round 0, the round phases, anonymity, pulse and metrics. The facilitator is active: collection is not a passive wait but spurring the panel into movement.

HANDOFF 2

Data out

The export contract hands the validated dataset to DAE. A malformed dataset cannot silently enter analysis — the gate stops it.

Stage 1 — Design (DPE)

Before Pronoia collects anything, the facilitator designs the study in the Planning ecosystem. Each step is a deliberate methodological choice that shapes everything downstream.

Planning step	What the facilitator decides	What it produces
Orientation	The study's logic (N / S / R / E / D / D+)	Expert vs observer panel; dialectical vs abductive theses
Pre-analysis	Futures Wheel (impacts), Relevance Tree (goals-means)	A scoped, structured view of the phenomenon
Phenomenon	What is actually being studied	The boundary of the inquiry
Panel design	The matrix, the perspective layer, the health score (PHS)	A designed epistemic space
Question architecture	The theses and their structure	The collection instrument
Simulation (optional)	Synthetic argumentation when no real panel exists	A test / seed argument set
Learning loop	Lessons carried from earlier studies	Better, evidence-based design choices

Planning ends at a quality gate: only a coherent design — orientation, phenomenon, panel and theses aligned — is packaged for collection.

Handoff in — the Planning Handoff Package

The design crosses into Pronoia as a structured **Planning Handoff Package**: the orientation, the panel matrix and its type, the theses, and the perspective/PHS context. This seeds Round 0 — so the facilitator does not rebuild the study in Pronoia, but continues it.

Stage 2 — Collection (Pronoia)

This is the facilitator's operational core. In brief, the facilitator opens **Round 0** as an AI-assisted intake (motivating participants and mapping the phenomenon, optionally with key informants), then runs each round through its phases — response → dialogue → revision → close — setting the iteration

stance with policy presets (classic / RTD / open, with Barometer arriving), protecting anonymity and identity, and reading the panel through the pulse, the group comparison and the process metrics.

What distinguishes the facilitator is an **active role**: collection is not a passive wait for answers but a continuous effort to spur the panel along — and the levers work at three levels.

Level	How the facilitator spurs movement	Effect
The whole panel	Pace the phase advance; open the result graph and pulse at the right moment; broadcast reminders and messages; choose the visibility stance; supplement the panel between rounds	Sustains momentum and collective movement
Groups	Surface the group comparison across the matrix axes; highlight that one group has moved while another has not; add perspective theses where a viewpoint is thin; recruit roles to cover blind spots	Drives differential movement and cross-position dialogue
Individuals	The live Home view with its "since your last visit" digest; direct nudges and personal messages; recognition of contributions; per-panelist reminders	Pulls each panelist back to revise and learn

Spurring is targeted, not generic. The most effective prompt is specific: show one group the gap between its position and another's, or show an individual how far the panel has moved since they last looked. Real-time data is what makes that precision possible.

The validation gate — earning the handoff

Collection does not flow into analysis automatically. Before any handoff, Pronoia validates the result, and a failing check blocks the export until it is fixed:

- **Panel-composition gates** — the orientation-specific minimums (matrix size, stakeholder groups or observation distances, domains or outside roles, decision-maker rules).
- **Epistemic coverage** — the argument space read as a coverage map; thin coverage is a prompt to supplement the panel before closing.

This gate is the facilitator's quality promise to the analysis stage: what crosses over is complete and well-formed, not raw.

Handoff out — the export contract

On a clean validation, Pronoia emits a structured payload for the DAE ecosystem. At the facilitator's level it carries:

- the **orientation** and **panel matrix type**;
- the **argument space** — theses × panel members and their arguments;
- the **epistemic coverage** — how much of the argument map is filled;
- the **perspective summary** — health score, dominant perspectives and named blind spots;
- the **validation status** — confirming the dataset passed its gates.

Downstream consumers verify this contract before they begin, so a malformed dataset cannot silently enter analysis.

Stage 3 — Analysis (DAE)

In the DAE ecosystem the collected material is interpreted through a gated pipeline. The facilitator **commissions and oversees** it rather than running every method by hand.

Foundation → Content → System → Futures → Synthesis → Report

DAE phase	What it does
Foundation	Validates the panel and builds the quantitative base
Content	Analyses arguments, dialogue and discourse, including the epistemic coding
System	Reads depth, dynamics and purpose (causal layers, transitions, soft systems)
Futures	Surfaces weak signals and maps structural tensions
Synthesis	Integrates every lens into a grand synthesis, with perspective depth
Report	Builds scenarios and audience-specific stakeholder reports

Quality gates and human judgement sit throughout: where automated confidence is not enough, the analyst is brought into the loop — and the facilitator decides how far interpretation should reach.

Closing the loop — learning across studies

The arc is not a straight line but a cycle. Both ecosystems keep a memory: the planning side records each design and its lessons, the analysis side keeps a case library of completed runs. The facilitator feeds these lessons back into the next study's design — so each Delphi run makes the next one better.

The facilitator's through-line. Design well, collect cleanly, analyse rigorously, and learn from the result. Pronoia is where that arc is held together — the point where a designed study becomes trustworthy data, ready for interpretation.

The facilitator's end-to-end playbook

1. **Design with the end in mind.** Choose orientation, scope the phenomenon (pre-analysis), design the panel and write the theses — knowing they define the whole argument space.
2. **Cross the first handoff intact.** Carry the Planning Handoff Package into Pronoia and continue the study rather than rebuilding it.
3. **Run collection deliberately.** Open the intake, set the iteration stance, advance the phases, protect anonymity, and watch the pulse and metrics.
4. **Earn the second handoff.** Clear the validation gate — panel composition and epistemic coverage — before closing; supplement a thin panel if needed.
5. **Commission analysis.** Export the validated dataset, oversee the DAE pipeline, and decide how far interpretation and reporting should go.
6. **Close the loop.** Capture the lessons and carry them into the next study's design.

Source: Pronoia — The Facilitator's Process (Metodix, June 2026). This block reframes the same arc along the capability arc.

Delphi variations

A study's **variation** decides how responses are collected and how movement is tracked over time. It sits alongside the orientation (N/S/R/E/D/D+) chosen in planning, and it sets the collection key and the lead indicator the analysis uses. The names are proper names and the same in every language:

Argument Delphi, Classic Delphi, Real-Time Delphi, Barometer Delphi. The value stored in the

contract is a different thing from the display name (the value for Classic Delphi is still `Classical`).

Variation	Collection	Lead indicator	What it is · when to use
Argument Delphi	round	CDI	Dialectical, multi-voice: required rationale + dialogue over rounds; aims at dissensus — maps tensions and reasoning; 7×3 epistemic coding in analysis. The ecosystem's signature when <i>reasoning</i> matters more than a number.
Classic Delphi	round	CCI	The same process and phases as Argument Delphi; aims at reasoned consensus and measures it with consensus indicators. The safest general foresight setup.
Real-Time Delphi	window	RCI	One round of responding and dialogue, open at the same time from the start: panelists see the group state and revise anytime; no comparison between rounds. Fast, asynchronous.
Barometer Delphi	wave	BCI	The same theses rated again from one wave to the next — trend and opinion-shift monitoring; a thesis may be retired between waves. For longitudinal studies.

The variation is an editable contract block; if absent, the default is Argument / round / CDI (backward compatible). Common to all four: Round 0 (orientation) belongs to every one, and dialogue exists in every one — the variations differ in where in the round the dialogue sits, not in whether it exists.

Templates for different purposes

Ready-to-run **templates** pair a variation with a panel and starter theses, so you can instantiate a complete study in one step — from the new-study wizard (which pre-fills every step) or by importing the package. Each is a contract-conformant starting point you then edit.

Template	Base variation	Purpose
Classic foresight Delphi	Classical	General foresight baseline (probability/desirability + rationale).
Argument Delphi	Argument	Dialectical, polyphonic process with 7×3 epistemic coding.
Barometer	Barometer	Longitudinal trend monitoring across waves.
Participatory civic Delphi	Classical (observer / E)	Participatory foresight (e.g. local development); Asianosaisuus × expertise.
Real-Time Delphi	Real-Time	Fast, asynchronous continuous window.
Corporate strategy Delphi	Argument (hybrid)	Argumentative rounds + a consensus final round; org groups × competence.
Full-AI perspective panel	Argument (Route B)	End-to-end AI-panelist run (Hautamäki perspectives) — for full-process testing and human-validated AI-analysis development.

Next: choose orientation and questions in Planning, and design the matrix in Panel design.

Real-Time Delphi — how to run one

Real-Time Delphi is the variation with **one round that is open from the start** — for answering and for dialogue at the same time. A panellist waits for nothing: they see the theses and the panel's discussion immediately, answer at their own pace, and can come back and change their position for as long as the study is open.

In the contract the variation is `Real-Time`, the collection key is `window` and the lead indicator is **RCI**. The three travel together: the collection key says the collection period is a window rather than a round, and the indicator says what the analysis reads.

When to choose RTD

RTD suits a study where **time is short and participants are scattered**. It is asynchronous: each person takes part when they can, and nobody waits for a round to close. It also suits a fast-moving phenomenon, where the gap between two rounds would already have dated the question.

RTD **does not fix the aim**. It can aim at consensus or at dissensus — the aim is resolved per thesis, not by the variation. This differs from Classic and Argument Delphi, where the aim is a property of the whole study. In practice: if you want the panel to seek common ground on *this* thesis but map disagreement on *that* one, RTD gives you that without splitting the study in two.

What RTD does not replace: if you want to **see how the panel moved** — that someone changed position after reading others' reasoning — you need two rounds. Movement is the difference between two measurements, and a single measurement has no difference to compute. That is the price of RTD, and it is worth knowing before the variation is chosen rather than after.

What the software does for you

The single-round nature is not a rule you have to remember to follow — it is built in:

- **The wizard does not ask for a number of rounds.** The round count is not a question but a property of the variation, and in RTD it is one. (Before 5 September 2026 the wizard asked every variation, defaulting to 2 — a choice that does not exist.)
- **The server refuses a second round.** `POST /rounds` answers `409 single_round_mode`. You cannot drift into a multi-round setting by accident.

- **All phases are open.** Response, dialogue and revision are all open at once in RTD. That is why the phase bar carries no round number: it says open round, not "Round 1".
- **The movement panel says outright that movement is not computed.** It is not an empty box but an explicit message — in a single-round study there is no between-round movement.

Planning

Planning follows the same route as in any other variation: orientation, phenomenon, panel, question architecture.

Round 0 belongs to RTD too. It is the orientation and key-informant stage, and it is the only numbered step before that one round. Round 0 material is intake — it is not panel data and does not travel into the analysis as responses.

Keep the number of theses moderate. Because everything is open at once, the panellist sees the whole workload at first glance. In a multi-round setting the work is spread across rounds; here it is not.

Require reasoning. The value of RTD comes from panellists reading each other's reasoning and responding to it within the same window. Without reasoning what remains is a survey with a comment box.

While the study runs

With no round comparison, you follow three things:

1. **The build-up of participation.** Who has answered, who has started, who has not opened the link. In a single window this is the only measure of progress.
2. **The density of dialogue.** Is discussion emerging, or are people only answering? In RTD the discussion is what separates it from a survey, so its absence is a finding rather than silence.
3. **The panel's pulse per thesis.** Where views sit close together and where they spread. This is read from one measurement, not from the difference between two.

The facilitator's role is lighter than in a multi-round Delphi: no rounds are opened or closed, no theses are rewritten between rounds. What remains is what matters most in RTD — **keeping the discussion alive**. If dialogue fades, you address it with a message to the panel, not with a new round.

Closing

The study ends when the facilitator ends it. The window does not close by itself and the round does not "fill up". Closing is therefore a real decision, and it is worth telling the panel in advance — otherwise a participant does not know whether there is still time to come back.

The discussion of a closed study remains readable.

What an RTD study does not have

These belong to multi-round variations and are not worth looking for in an RTD study:

- round comparison and movement analysis
- a pause between rounds
- a notification that a new round has opened
- a round number in the phase bar
- **retiring a thesis** — that belongs to Barometer Delphi only, and the server forbids it elsewhere

If you find yourself needing one of these, the variation is the wrong one. That is not a fault but a choice, and it can be made differently in the next study.

Orientation, variation and question architecture

Planning (DPE)

Planning determines what Pronoia collects and what DAE can interpret. The quality of the study is decided here — two locked starting choices (orientation and variation) plus the scoping of the phenomenon and theses shape everything downstream.

Orientation — from which standpoint the future is studied

Orientation is the study's **locked starting choice**. It affects how theses are framed, how the panel is composed, and how the facilitator acts between rounds (in Pronoia it also drives the PIRE round-gap assistant).

Orientation	Core	Relation to consensus
N — Normative	Desirability and the desired state. "What future is good?"	Leans toward consensus — seeking a shared view of the goal
S — Strategic	Feasibility and agency. "How is the goal reached, who acts?"	Emphasises means and alternatives
R — Radical	Ruptures and protecting deviance. "What could change fundamentally?"	Protects minority views — no forced convergence
E — Emergent / abductive	New, still-forming phenomena. An observing stance: what is emerging?	Does not seek consensus — it maps
D — Intentional combination	A mixed form where different questions may carry different logic	Default for multifaceted studies
D+ — Four-path	Intentional and emergent in parallel	Two panels, each with its own requirements

Orientation is not just a label — it **constrains methodological choices**. An R study does not use convergent "central-tendency only" feedback, because that would erase the very deviance it aims to protect. In an N study consensus is a natural target. Orientation also sets the **panel type**: N/S/R/D → expert matrix, E → observer matrix (see Building the panel).

Variation — the cadence of rounds

Alongside orientation, a **variation** is chosen that sets how rounds and feedback are paced. In Pronoia this is `delphi_mode` :

Variation	How it advances	When
Argument Delphi (default)	Staged rounds: response → dialogue → revision → pause; aims at dissensus — tensions and reasoning	When reasoning and the differences between views are the result
Classic Delphi	The same process and phases as Argument Delphi; aims at reasoned consensus and measures it (CCI)	When the study must arrive at a position
Real-Time Delphi	One round in which answering and dialogue are open at the same time; no comparison between rounds	Fast, continuous eDelphi
Barometer Delphi	The same theses from one wave to the next; a thesis may be retired between waves	Repeated monitoring — a trend over time, not converging one run

Variation and orientation are different axes: orientation gives the *standpoint*, variation the *cadence*. The same normative study can be run as any of the four. Round 0 (orientation) belongs to every variation.

Scoping the phenomenon — what is actually studied

Before the theses, the **phenomenon** is defined: the boundary of the inquiry. This is planning's most critical scoping step — it decides what is in and what is out. An optional pre-analysis sharpens it:

- **Futures Wheel** — an impact network: first-, second- and third-order consequences (STEPP+). Maps what follows from the phenomenon.
- **Relevance Tree** — a goals-means hierarchy: priorities and success criteria. Maps what the phenomenon aims at.

Pre-analysis produces a scoped, structured view of the phenomenon that feeds directly into panel and thesis design.

Question architecture — the theses

The theses are the collection instrument: every argument is one panel member's response to one thesis, so the set of arguments is **theses × panel members**. Thesis structure follows the orientation:

- **Dialectical theses (N / S / R / D)** — a claim taken a position on by probability and desirability; dialogue produces for/against arguments.
- **Abductive theses (E)** — more open and exploratory: what is emerging, what it looks like.

A single thesis's logic can be tuned with a **P/D tag**: probability (P) leans to consensus, desirability (D) allows dissensus — overriding the orientation default for that thesis. There are at least three theses.

Planning ends at a quality gate. Only a coherent design — orientation, phenomenon, panel and theses aligned — is packaged as the **Planning Handoff Package** and crosses into Pronoia for collection. The design is not rebuilt in Pronoia but continued (see Integration guide).

When there is no real panel — simulation

If no real panel exists yet, the **Argument Simulator** (Route B) generates synthetic, orientation-aware arguments for the theses — a test or seed dataset that can be carried into DAE. Simulation tests the theses before a real panel; it does not replace one.

Sources: help/02-orientaatiot, planning-core (M1a). Variations correspond to Pronoia's delphi_mode states.

Building the panel: epistemic space, perspective and gates

Planning (DPE)

A Pronoia panel is not a random set of experts. It is a **designed epistemic space** whose reach and gaps are deliberate and documented. Panel composition is a methodological decision, not a recruitment convenience — because the panel defines the entire argument space, its coverage and blind spots are known **before any data exists**.

Three working principles:

- A **homogeneous panel** (same education, sector, culture) is an epistemic risk and is recorded as such — even if every member is individually excellent.
- The target is **commensurability** — enough shared ground for genuine dialogue — not full consensus. Disagreement that is well covered is a feature, not a defect.
- Watch for **lock-in**: a single perspective claiming universality. It can masquerade as dialogue simply because it is loud.

Orientation sets the panel type

Orientation	Panel type	Selection basis	Practical requirement
N / S / R / D (intentional)	Expert matrix	Expertise, stakeholder representation, competence domain	At least 3×3; ≥3 stakeholder groups; ≥2 domains; power balance considered
E (emergent)	Observer matrix	Observation distance (inside / edge / outside), not expertise	At least 2×3; ≥3 observation distances; ≥3 "outside" roles; no decision-makers; each role names its observation position and "what it sees"
D+ (four-path)	Two parallel panels	Intentional and emergent in parallel	Both requirement sets apply independently

Emergent studies are about *where you stand to watch*, not credentials — which is why decision-makers are excluded and observation distance replaces expertise as the selection basis.

The perspective layer: position + view

The optional but recommended **perspective layer** rests on a simple idea: *perspective = position + view*. The same phenomenon genuinely looks different from different positions — not because anyone is wrong, but because the phenomenon is multidimensional. This is not relativism: the aim is to identify what each position *can* see and what it necessarily *cannot*.

When the layer is active, each role carries three attributes:

Attribute	What it captures	Values
Observation position	Where the member stands relative to the phenomenon	inside / edge / outside / below / above
Epistemic type	The kind of knowledge the member brings	formal / experiential / tacit
Value basis	The standpoint's underlying values	short free-text

Blind-spot analysis. The facilitator checks which positions identified in the phenomenon map are systematically absent from the panel matrix. The key distinction: a blind spot is a **systematic absence**, not a random gap. Naming blind spots before round 1 is part of designing the space, and they become candidates for supplementation later.

Perspective Health Score (PHS)

After the panel is designed, Pronoia summarises its diversity as a **Perspective Health Score** on a 0–12 scale, built from four dimensions (position diversity, epistemic diversity, value-basis variation, marginal voices — each 0–3):

PHS	Interpretation	Recommended action
0–5	Diversity too thin	Panel supplementation required before round 2
6–8	Adequate, with gaps	Hautamäki perspective theses recommended
9–12	Excellent diversity	No action needed

Epistemic coverage: $7 \times 3 = 21$ cells

Arguments are read on two dimensions — **where the claim draws from** (knowledge base) and **what it does to the thesis** (function):

Knowledge base (7)		Function (3)	
EMP Empirical	data, measurements, statistics	TUK Supporting	reinforces the thesis
KOK Experiential	personal/professional experience, tacit knowledge	HAS Challenging	questions, weakens
NOR Normative	values, ethical principles	AVA Opening	introduces a new viewpoint, widens the frame
TEO Theoretical	models, theories, frameworks		
INT Interest-based	actor/stakeholder viewpoint		
EPV Uncertainty	acknowledged ignorance		
IMA Imaginative	vision, alternative, radical reframing		

Seven knowledge bases $\times$ three functions = a **21-cell map** of the argument space. Empty cells are the panel's blind spots — viewpoints that never appeared. The **Epistemic Coverage Index (ECI)** is simply how many of the 21 cells are filled. A typical Delphi run fills roughly 9–14/21; very low coverage is a signal to add roles that bring the missing viewpoints.

Quality gates before collection

Before a panel is accepted, it must pass orientation-specific gates:

Check	Intentional (N/S/R/D)	Emergent (E)
Matrix size	$\geq 3 \times 3$	$\geq 2 \times 3$
Stakeholder groups / observation distances	≥ 3 groups	≥ 3 distances
Domains / outside roles	≥ 2 domains	≥ 3 outside roles
Decision-makers	Power balance noted	Exactly 0
Per-role attribute	Expertise named	Observation position + "what it sees" named
Theses	≥ 3	≥ 3

If a gate fails, the panel is not yet ready for collection — the facilitator adjusts composition or theses and re-checks. The gates are guardrails that keep coverage and balance honest before any arguments are gathered.

Building the panel in practice

1. **Start from orientation.** Confirm the orientation; it decides whether you build an expert or an observer matrix.
2. **Build the matrix.** Lay out the two axes (e.g. expertise $\times$ stakeholder, or observation position $\times$ context) and place members so the hard minimums are met.
3. **Activate the perspective layer** when standpoint matters: give each role its observation position, epistemic type and value basis.
4. **Map blind spots** against the phenomenon map — name the systematically missing positions rather than treating gaps as random.
5. **Read the PHS and act on it:** supplement a thin panel before round 2, or add perspective theses where recommended.
6. **Mind the coverage.** Watch which of the 21 epistemic cells stay empty; low coverage prompts adding roles.
7. **Guard against lock-in.** Keep checking that one loud perspective is not crowding out the rest; the aim is commensurable dialogue.
8. **Clear the gates.** Confirm the orientation-specific quality gates before collection — and re-confirm after any panel change.

Source: *Pronoia — Building the Panel* (Metodix, June 2026).

Simulation: the Argument Simulator (Route B)

Planning (DPE)

When no real panel exists yet, the **Argument Simulator (Route B)** generates realistic expert and observer arguments for the designed theses. It is **orientation-aware** (N/S/R/E/D/D+) and produces structured arguments that the DAE pipeline reads (CP-1...CP-20). The purpose is to **test the theses and panel logic before a real panel** — or to provide material to analyse when no panel is available.

This does not replace a real panel. The arguments are synthetic, GenAI-disclosed drafts that require an audit. The actual interpretation happens in DAE under quality gates — simulation provides source material, not conclusions.

How it proceeds — six phases

The simulator activates once the planning context has passed the **QG-SIM-TRANSFER** gate (orientation, question structure, panel logic, theses). Then:

1. **Individual arguments** — one argument per role per thesis. Dialectical (N/S/R/D): thesis + antithesis + synthesis view. Abductive (E/D+): observation → explanation → anticipation.
2. **Quick analysis** — distributions, the 7×3 epistemic profile (knowledge base × function), conditional asymmetry.
3. **Dialogues** — 5–8 thesis tensions or observation clusters: position → concessions → synthesis insight.
4. **Analysis dimensions** — tension architecture, CLA pre-profile (causal layers), the 7×3 epistemic cross-tab, MLP pre-profile (multi-level).
5. **Anomaly check & supplementation** — an automatic blind-spot check (see below). Where needed, the panel is supplemented with missing viewpoints (marked `source: supplement`, ≤25 % of the panel).
6. **Summary & handoff** — the overall result, a keystone candidate and readiness for the DAE pipeline (→ CP-1 panel validation).

Each transition produces a light **Learning Note** that carries lessons from planning through simulation into DAE.

Anomaly check — surfacing blind spots

After phase 4 the simulator automatically checks whether the material leans unhealthily:

Trigger	Threshold
Any 7×3 knowledge base = 0	empty cell
Opening (AVA) arguments < 3	frame not opened
CLA litany layer > 80 %	surface dominates
CLA deep layers (L3+L4) < 10 %	worldview/myth missing
MLP regime > 60 %	incumbent system dominates
D–P gap < 0.3 on all theses	too much consensus

A trigger leads to a **supplementation protocol**: roles are designed to bring the missing viewpoints and a supplementary simulation is run — the same logic as supplementing a real panel (see Building the panel).

Worked example — Food Paradigm 2050

The Documentation Sprint's worked example shows Route B in practice:

Field	Value
Orientation	D (N + S + R) — multi-actor adaptation
Route	B (Argument Simulator)
Time horizon	2050 (decisive window 2026–2035)
Perspective layer	Active (Hautamäki), PHS 12/12
Panel	19 simulated roles (post-audit)
Theses	8 R1 dialectical + 4 R2 perspective theses + invitation + blind-spot
Quality	94.5/100, 27/27 quality gates passed

The phenomenon: five competing food-system paradigms (techno-capitalist, distributed bio-economic, fragmentation, multipath cultural-pluralism, post-smallholder), entry tension *technological autonomy* ↔ *environmental sustainability*. The example shows how a D orientation propagates coherently from the goal hierarchy to the theses, and how the perspective layer integrates without friction.

Sources: *Argument Simulator v2.5 (metodix-planning-simulation); worked example CASE_202605_FOOD-2050-WE-001 (Documentation Sprint).*

Building the Delphi panel

The panel defines the **entire argument space** of a study: every argument is one panel member's response to one thesis, so the arguments a study can ever produce are *theses × panel members*. Design the panel well and its coverage — and its blind spots — are known **before any data exists**. This guide is the methodological centre of planning a Delphi study in the ecosystem.

Planning (DPE) chooses orientation and drafts the panel. **Facilitation (Pronoia)** is where the panel is finalised, validated and gated before collection. The panel is a **designed epistemic space**, not a recruitment convenience.

The panel as a designed epistemic space

A panel is not a random set of experts; its reach and gaps are deliberate and documented. The guiding principles the facilitator works with:

A **homogeneous panel** (same education, sector, culture) is an epistemic risk and is recorded as such — even if every member is individually excellent. The target is **commensurability** — enough shared ground for genuine dialogue — not full consensus. **Disagreement that is well covered is a feature, not a defect**. And watch for **lock-in**: a single perspective claiming universality, which can masquerade as dialogue simply because it is loud.

The two-axis matrix is primary

The panel's primary structure is a **two-axis matrix**. Its cells hold the **roles** that panelist behavior is grounded in. The **facilitator ultimately defines the axes and their classification** — the matrix is the design surface, and the templates below are only starting points.

The axes are **context-dependent**. The canonical choices:

Panel type	Axis 1 (primary)	Axis 2 (secondary)
Expert panel (N/ S/R/D)	Expertise (research / practice / administration)	Stakeholder affiliation (public / private / civil society)
Participatory panel (E)	Concerned-party position (resident / civil servant / decision-maker / NGO)	Expertise / theme area
Corporate strategy	Concerned-party group (leadership / staff / customers / owners)	Competence area (strategy / finance / operations / market)

Two properties matter in practice. A **cell may hold 0..n panelists** — several voices in one role, or none. **Empty cells are expected**, and a systematically empty cell is a **structural blind spot** (not a random gap). Naming blind spots before round 1 is part of designing the space; they become candidates for supplementation later.

Orientation sets the panel type

The study orientation chosen in planning (N, S, R, E, D or D+) decides whether the panel is built as an **expert matrix** or an **observer matrix**, and which axes, selection logic and minimums apply.

Orientation	Panel type	Selection basis	Practical requirement
N / S / R / D (intentional)	Expert matrix	Expertise, stakeholder representation, competence domain	At least 3×3; ≥3 stakeholder groups; ≥2 domains; power balance considered
E (emergent)	Observer matrix	Observation distance (inside / edge / outside), not expertise	At least 2×3; ≥3 observation distances; ≥3 "outside" roles; no decision-makers; each role names its position and "what it sees"
D+ (four-path)	Two parallel panels	Intentional and emergent in parallel	Both requirement sets apply independently

Emergent studies are about *where you stand to watch*, not credentials — which is why decision-makers are excluded and observation distance replaces expertise as the selection basis.

Perspective: position + view (Hautamäki)

The optional but recommended **perspective layer** rests on a simple idea: **perspective = position + view**. The same phenomenon genuinely looks different from different positions — not because anyone is wrong, but because the phenomenon is multidimensional. The aim is to identify what each position *can* see and what it *necessarily cannot*.

Following Hautamäki's viewpoint relativism, the deeper point is the **multiplicity of perspectives in general** — a panel should span the perspective space, not just tick boxes. The familiar four archetypes (constructivist optimist, empirical realist, critical-conservative, eco-social holist) are **one useful special case**, not the whole space; add or vary perspectives as the phenomenon demands.

When the layer is active, each role carries three attributes that make its standpoint explicit:

Attribute	What it captures	Values
Observation position	Where the member stands relative to the phenomenon	inside / edge / outside / below / above
Epistemic type	The kind of knowledge the member brings	formal / experiential / tacit
Value basis	The standpoint's underlying values	short free text

The perspective layer **complements** the matrix cell: the cell sets *what* a member speaks to (its role); the perspective sets *how* it reasons.

AI panels (Route B / C)

A study can run wholly or partly with **AI panelists** — a full-AI perspective panel makes an end-to-end run possible without recruiting humans, which is valuable for testing the whole process and for **human-validated AI-analysis development** (AI proposes → human validates → κ/calibration → improve).

An AI panelist sits in a **matrix cell** (its role) and carries an explicit **perspective profile** that complements it. Because the AI acts on that description, it must be **precise and sufficient** — this is enforced by the **AI Panelist Description Standard** (docs/ ai_panelist_description_standard.md): each AI panelist needs a filled cell (expertise × stakeholder), a distinct perspective label, the four Hautamäki dimensions (ontological / epistemological / axiological / temporal) as real descriptions, `perspective_axes` in [0,1], a bias

direction and the model used; across the panel the labels are unique and the perspectives diverse. Route B/C also requires a GenAI disclosure, and humans stay in the loop: every AI response is reviewed before a round closes.

Reading the arguments: the epistemic lens

As arguments arrive, Pronoia reads each one on two dimensions — **where the claim draws from** (its knowledge base) and **what it does to the thesis** (its function).

Knowledge base — seven categories: **EMP** empirical (data, measurement) · **KOK** experiential (personal/professional, tacit) · **NOR** normative (values, ethics) · **TEO** theoretical (models, frameworks) · **INT** interest-based (stakeholder/strategic position) · **EPV** uncertainty (acknowledged ignorance) · **IMA** imaginative (vision, radical reframing).

Function — three categories: **TUK** supporting (reinforces the thesis) · **HAS** challenging (questions, weakens, refutes) · **AVA** opening (introduces a new viewpoint, widens the frame).

Seven knowledge bases × three functions give a **21-cell map** of the argument space. Empty cells are the panel's blind spots — viewpoints that never appeared. The **Epistemic Coverage Index** is how many of the 21 cells are filled; a typical run fills roughly 9–14 of 21, and very low coverage signals a need to add roles that bring the missing viewpoints.

Panel diagnostics (coverage, PDI, lock-in, PHS)

Description quality is *per panelist*; **coverage** is a property of the whole panel. The ecosystem computes it deterministically — the computable core of **MaxPerspective (M5)**, aligned with Hautamäki's viewpoint relativism, available live in the new-study wizard and as a tool over any panel:

- **Matrix coverage** — filled vs empty cells (Absent cell = structural blind spot) and fill rate.
- **Dominance & diversity per axis** — aspect dominance **AD%** per category, the dominance class (DOMINANT >50% / STRONG / MODERATE / WEAK / TRACE / ABSENT), and the **Perspective Diversity Index** $PDI = 1 - \sum (AD_i/100)^2$ (Herfindahl). **Lock-in** is flagged when one category exceeds 50%.
- **Structural PHS (0–50)** — $PDI \times 25 + (1 - \max AD/100) \times 25$. The full Perspective Health Score also weighs *awareness* and *openness*, which are qualitative (the six-phase MaxPerspective analysis) and not inferable from a panel spec.
- **AI perspective spread** — over `perspective_axes`, the spread of the AI panel across the perspective space; low spread is perspective lock-in even when the matrix looks covered.

- **Supplement recommendations** — which categories to fill (Absent blind spots), where to reduce concentration, and how to broaden perspective spread.

As a fast human-facing summary the facilitator may also record a **PHS quick score (0–12)** from four diversity dimensions — position diversity, epistemic diversity, value-basis variation and marginal voices, each 0–3: **0–5** diversity too thin (supplementation required) • **6–8** adequate with gaps (Hautamäki perspective theses recommended) • **9–12** excellent diversity.

Quality gates before collection begins

Before a panel is accepted it must pass orientation-specific gates that turn the principles above into checkable minimums:

Check	Intentional (N/S/R/D)	Emergent (E)
Matrix size	≥ 3×3	≥ 2×3
Stakeholder groups / observation distances	≥ 3 groups	≥ 3 distances
Domains / outside roles	≥ 2 domains	≥ 3 outside roles
Decision-makers	Power balance noted	Exactly 0
Per-role attribute	Expertise named	Observation position + "what it sees" named
Theses	≥ 3	≥ 3

If a gate fails, the panel is not yet ready — the facilitator adjusts composition or theses and re-checks. The gates keep coverage and balance honest before any arguments are gathered.

Building the panel in practice

Ready-to-run **templates** give a validated starting point for each panel type — Classic foresight, Argument Delphi, Barometer, Participatory civic, Real-Time, Corporate strategy and a full-AI perspective panel — each a contract-conformant package you instantiate in one step (from the study wizard or by importing a planning package) and then edit.

A working sequence for the facilitator: pick a template (or start blank) → set the two axes for your phenomenon → place roles in the cells (leaving deliberate blanks) → for AI panels, write each perspective profile to the standard → read the live diagnostics (PDI, lock-in, blind spots, structural PHS, AI spread) and act on the recommendations → pass the quality gates → begin collection.

The panel is where a Delphi study's epistemic reach is decided. Everything downstream — the arguments, the analysis, the scenarios — is bounded by the space you design here.

The iterative process: rounds, phases and convergence

Why Delphi iterates

Delphi is not a one-shot survey. Its defining mechanism is **iteration with controlled feedback**: experts answer independently, then receive an anonymized picture of the whole panel's answers and reasoning, and reconsider in that light. Repeating this loop is what separates Delphi from a poll.

The iteration does specific work: it lets reasoning surface and be exchanged rather than just tallied; it gives people a structured chance to revise in response to arguments they had not considered; and across rounds it either moves the panel toward consensus or reveals a **stable, well-understood dissensus** — which is an equally valuable result.

The core sequence is: **independent judgement → feedback → reconsideration**. Anonymity and the staged release of feedback exist to protect that sequence from social pressure — bandwagon effects, dominance and anchoring — so that movement between rounds reflects better reasoning, not the loudest voice.

The Pronoia process at a glance

A study advances through a small state machine that the facilitator drives one step at a time:

Draft → Round 0 → Round 1 → Round 2... → Close → Analysis

Round 0 is the orientation, and it belongs to every variation; every numbered round runs the same internal sequence of phases; and the loop repeats until the facilitator closes the study and hands a clean dataset to analysis. The variation's own character begins when Round 1 opens: Real-Time Delphi runs one round in which responding and dialogue are open at the same time, and Barometer Delphi rates the same theses again from one wave to the next. Classic Delphi and Argument Delphi run the same process; they differ in the aim (consensus / dissensus) and in the indicators, not in the phases.

Round 0 — mapping and motivating (AI-assisted intake)

Round 0 is more than a setup screen. In Pronoia it is run as an **AI-assisted chat dialogue** that does two jobs at once:

- **Motivating.** The dialogue meets people where they are, explains what the study asks of them and why their view matters — so they arrive at Round 1 invested rather than cold.
- **Mapping.** It elicits how each participant frames the phenomenon — the issues, positions and tensions they see — which scopes the argument space and feeds directly into thesis design and the panel matrix.

Selective / key-informant variant. The dialogue does not have to be run with the whole panel. The facilitator can run it selectively with a handful of **key informants** — people who know the terrain — use their input to map the phenomenon and seed the theses, and then invite the wider panel straight into Round 1.

Whichever way it is run, Round 0 ends by fixing the two things that define the argument space: the **panel** (built as the orientation-appropriate matrix, cleared through its quality gates) and the **theses** (at least three, in the structure the orientation calls for). While in this setup state the facilitator can still edit and, where used, regenerate questions; once Round 1 opens, the instrument is **locked** so everyone answers the same questions.

The phases inside a round

Every numbered round runs the same four phases. Separating them is deliberate: it keeps "form your own view" apart from "engage with others", which is the heart of the Delphi method.

Phase	What the panelist does	Sees others?	Interaction	Why
Response (1a)	Answers each thesis independently	Not before submitting	Edit own answer	Capture genuine, un-anchored judgement
Dialogue (1b)	Reads the panel, comments and replies	Yes	Comment + reply, graph visible	Exchange reasoning, not just numbers
Revision	Updates positions in light of the dialogue	Yes	Edit own answer (tracked as a revision)	Let movement happen — and measure it
Closed	Round is locked	Yes (read-only)	None	Freeze a comparable snapshot before the next round

The response and dialogue phases are also called **1a** and **1b** in the facilitator controls — sub-rounds of the same round.

A phase adds a capability, it never removes one. Opening the dialogue phase does not close answering: a panelist who did not get to answer can still answer during the dialogue, and may take part with the full repertoire for the whole round. The idea of Delphi is that a panelist can change their answer when someone else's reasoning gives a good reason — and that requires answering to stay open.

Controlled feedback: the visibility ladder

What a panelist can see is set **per phase**, and it is the main lever Pronoia gives the facilitator over anchoring. Three settings climb from fully protected to fully transparent:

Visibility	Meaning	Effect on anchoring
Blind	Others' answers stay hidden for the whole phase	Maximum protection; use with deliberation
After-submit	No anchoring before you commit, then the panel and the aggregate graph open up	Protected where it matters, feedback where it helps (the eDelphi norm)
Open	The panel is visible throughout the phase	No protection by design; full transparency

Alongside visibility, each phase independently controls whether commenting, replies and the result graph are available — which is how the response phase stays quiet and focused while the dialogue phase opens up the argument.

Policy presets — choosing an iteration stance

Rather than set every phase by hand, the facilitator picks a preset that encodes a methodological stance; individual phases can still be overridden per round.

Preset	Response phase	Dialogue / revision	When to use
Classic	After-submit, no graph	Open	Staged Delphi with anchoring protection — the default
Real-time (RTD)	Open	Open	Continuous, always-open eDelphi; the round never closes for staged feedback
Open	Open, graph visible	Open	Full transparency from the first answer — workshops, teaching, low-stakes panels
Barometer	Open, recurring	Open	The same theses rated again from one wave to the next; tracks how positions move over time rather than converging in one study

Preset and variation are different things. The **Classic** preset is the shared stance of Classic Delphi and Argument Delphi — the same process, a different aim. Real-Time Delphi is a different cadence rather than a different method: instead of discrete rounds, the panel answers and sees feedback continuously for the length of one round, and the facilitator closes when the picture is stable. Barometer Delphi is a monitoring stance: the same instrument is measured again wave by wave, a

thesis may be retired between waves and replaced by a new one, and change over time is itself a finding.

Revision and movement

The point of feedback is to make **informed movement** possible. When a panelist changes an answer during the dialogue or revision phases, Pronoia records it as a revision rather than overwriting silently — so the facilitator can see *how much* and *where* the panel moved between rounds.

That movement is read against the **panel pulse** — the consensus / dissensus picture per thesis. Convergence shows where the panel is settling; persistent dissensus shows where a real, structured disagreement lives. Both are signals for the next decision.

Crucially, the learning movement is **not uniform** across the panel. Different groups — expertise areas, stakeholder positions or observation distances — typically behave differently: some converge quickly, others hold their ground, others shift only late. Reading movement **by group** is central to describing how the panel actually learns — and it is also where the most effective nudging happens: simply surfacing that one group has moved while another has not is itself a prompt for the lagging group to reconsider. Pronoia's group comparison along the panel-matrix axes makes these differential movements visible.

The panelist Home — real-time data that fuels iteration

Iteration only works if panelists come back and engage. The panelist **Home view** is designed to make that happen: it turns the panel's live data into something informative and absorbing, so that seeing the panel move makes the panelist want to move too.

The view adapts to where the panelist is. On the first visit, before they have answered, it stays deliberately minimal — a short "About this study" and a direct route into the questions — so nothing stands between them and their first independent answer. Once they have answered, the layout puts the whole panel first:

See the pulse → Notice movement & tension → Revise → See your effect & learn

Why real-time data fuels iteration. Seeing live consensus form, or a tension open up around one's own position, is what moves a panelist to reconsider and revise. Each revision then changes the picture others see, which prompts further movement — and through that loop the panel, and each panelist, learns.

Tracking the panel — process metrics

Throughout the whole process Pronoia follows the panel with a small set of activity metrics. Together they show, at a glance, whether the panel is alive and engaging:

Metric	What it counts	What it signals
Panelists	Members in the panel	The reach of the study
Respondents	Panelists who have answered	Response rate and coverage
Comments	Arguments and comments written	Depth of reasoning offered
Replies	Comments on comments	Genuine dialogue, not parallel monologue
Reactions	Lightweight responses to others	Low-effort engagement and attention
Revisions	Changes to answers and comments	The learning movement itself
Visits	Visits to the site	Sustained attention and return rate

Revisions and replies are the most telling: they are the traces of panelists actually reconsidering in light of one another — the iteration working as intended.

Iterating across rounds — when to continue, when to close

After a round closes, the facilitator makes one decision: run another round, or stop.

- **Run another round** when positions are still moving, when dialogue raised new considerations that deserve a fresh answer, or when a thin panel has just been supplemented.
- **Close** when the panel has converged, or when dissensus is **stable and well-understood** — further rounds would add cost without insight.

Between rounds the facilitator can also act on panel health: if diversity was too thin, supplement the panel (or add perspective theses) before the next response phase, so the new round genuinely widens coverage rather than repeating the same voices.

Closing hands off to analysis. On close, the collected and validated material becomes the input to the analysis ecosystem (the DAE pipeline and its methods). Iteration in Pronoia ends exactly where structured analysis begins.

The facilitator's iteration playbook

1. **Set the stance first.** Choose a preset (classic / RTD / open) that matches the stakes, and adjust individual phases only if you have a reason.
2. **Open with the intake dialogue.** Run the AI-assisted Round 0 chat — with the full panel or a few key informants — to motivate participants and map the phenomenon before the theses are locked.
3. **Open Round 0, lock the instrument.** Finalise panel and theses, then open Round 1 so everyone answers the same questions.
4. **Advance one phase at a time.** Move response → dialogue → revision deliberately; let the response phase do its un-anchored work before opening the panel.
5. **Keep the panel engaged.** Use reminders and messages, and point panelists to their live Home view — iteration only works if people return.
6. **Read the pulse before deciding.** Look at convergence and dissensus per thesis, and how much positions moved, before choosing to iterate or close.
7. **Tend panel health between rounds.** Supplement diversity where it was thin so the next round adds coverage.
8. **Close cleanly and hand off.** Stop when consensus or stable dissensus is reached, lock the study, and pass the validated dataset to analysis.

Source: Pronoia — *The Iterative Process* (Metodix, June 2026).

Anonymity: identity modes and group reporting

Principle

The core of the Delphi method is to reduce social-conformity pressure and dominance through **(semi-)anonymity**. In Pronoia, panelists are **strictly anonymous to one another yet identifiable through a pseudonym**, so that in dialogue a participant recognises the same conversation partner as the same person. Anonymity protects honest judgement; it must not break accidentally through small groups, AI participants, or comparisons.

Group-level reporting runs along the **panel matrix** — the expertise (primary) and stakeholder (secondary) axes — rather than a single legacy grouping. The same **k≥3 protection** applies to every axis and to the panel-pulse group dots.

Three identity modes (per-panel setting)

A. Pseudonymous (default)

A stable **word handle** per panelist (e.g. Aurora, Vega) — memorable and distinctive, giving no hint of identity or status, stable for the whole study. Real identity sits only behind the invite link (invite token). Full functionality: Home, dialogue, revisions, recognition, arena.

B. Role play (optional layer on top of A)

The panelist **writes a role** or perspective they answer from (e.g. *Municipal finance director in 2035*) and **places it in a box of the panel matrix**. Off by default. The division of work is clear: **the facilitator owns the structure** – the axes and categories of the panel matrix – and **the panelist chooses the box** their role fits best.

- **The role is a name, the box is a group.** Group comparison and analysis use the box the panelist chose, not the free role text. Freely written roles therefore do not break the panel into groups of one, and k≥3 protection stays in force.
- **Others see the role** next to the panelist's reasoning and comments, after the handle, marked *self-selected*. It is a *claimed* perspective, not verified expertise.

- **AI role.** The facilitator can commission a role for an AI panelist in the same matrix. It is shown marked *AI role*, and the AI answers from the role's perspective. A role a human chose and a role an AI plays are never confused.
- When role play is switched off, roles are shown to no one; the saved role and box are kept.

C. Fully anonymous one-shot poll

No identity, a single response – a seminar or mobile poll the facilitator shares as a **QR code** or link. No sign-in, own page, arena, revisions, recognition or persistent handle. Voters answer the theses and may write a short reason; the writer is never named. Each submission is stored as an anonymous participant who does not appear on the panel roster or in group comparison, and the answers go into the round's results and distribution like any other answers.

The duplicate-vote guard is per browser: the same browser answers once. It stops accidental repeats but not deliberate ones – another browser or a private window can answer again. The one-shot poll therefore suits taking the temperature of a room, not a vote whose outcome is binding.

Cross-cutting rules (always on, in every mode)

AI is always disclosed. Every AI panelist is marked with "**AI**" plus a short role description, 2–3 words (e.g. *AI · Economist*, *AI · Labour Researcher*). Shown everywhere an AI panelist appears: dialogue arguments and comments, panel lists, pulse group labels, reports. Never hidden, never presented as human.

k-anonymity (k = 3) across all group-level views. No sub-group with fewer than three people is reported at group level — comparisons, sentiment and panel-pulse group dots included — and the rule applies **independently on each matrix axis**. Complementary protection: when a group falls below the threshold it is folded (into an "Other" group) until every cell has at least three; if "Other" is still below three, the breakdown is not shown at all. Panelists with no value on an axis are excluded from that axis — they never form a residual bucket. The same threshold applies to cross-tabulations (group × thesis): a numeric cell is hidden when $n < 3$. On small panels the facilitator is steered to a 2–3 group split; anonymity wins. In the extreme, the pulse is shown at whole-panel level only.

Visibility & the after-submit reveal. Per-round phase policy controls when a panelist sees others. The default **classic** preset uses **after-submit** for the response phase: no anchoring before you commit your own answer, but once submitted the panelist sees others' anonymized responses and the aggregate graph. Strict **blind** remains available in the facilitator policy editor. All revealed material stays pseudonymous and $k \geq 3$ -protected; the reveal changes *timing*, not the identity rules.

What is available in each mode

Feature	A Pseudonymous	B Role-play	C Fully anonymous
Persistent handle	yes	yes	no
Recognisable in dialogue	yes	yes	no (nameless)
Home / arena / recognition	yes	yes	no
Revisions	yes	yes	no (one-shot)
Commenting	yes	yes	yes, anonymous (moderated)
Self-selected role in a box	no	yes (marked)	no
AI labelling (if any AI)	always	always	always
k≥3 group reporting (both axes)	always	always	always
Facilitator sees person	yes*	yes*	no (no identity)

*Unless blind-facilitator mode is enabled.

By default the facilitator sees real identities (for reminders and quality control); a per-panel **blind-facilitator** mode hides them, so the facilitator also sees only pseudonyms (the invite mapping still exists in the database).

Group comparison — current behaviour

- **Two axes, stacked.** Both the panelist Home and the facilitator Community view show the expertise (primary) comparison and, beneath it, the stakeholder (secondary) comparison — each with its own roster, numeric matrix and sentiment.
- **Numeric comparison where the question type allows.** Group means and their deviation from the panel mean (plus a divergence figure) are computed for scale questions only; other question types do not produce a numeric matrix.
- **Comment-based sentiment for every question type.** Each group's tone (positive / cautious / critical / neutral, with a net figure) is derived from arguments and comments across all theses, so open, ranking and multiple-choice questions are covered too.
- **Pulse group dots** follow the primary (expertise) axis, matching the group-comparison default.

- All of the above remain under **k≥3 with complementary suppression**; no "Out" / "Other" buckets appear for a fully classified matrix panel.

In practice: role play and the one-shot poll

Each has its own working guide:

- Role play: a role in a box of the panel matrix: when role play suits a study, the facilitator's steps, a template for the panel, AI roles, and reading results by box.
- One-shot poll at a seminar: QR code and anonymous answers: preparation, the QR code and presentation view, what to tell the room, disruptions, results and limits.

Data model (implemented) — for the developer-partner

- **Study:** `facilitator_blind`, `role_play_enabled`, `handle_theme`, one-shot fields (`oneshot_comment_window`, `oneshot_moderated`), `policy_preset` + per-round `phase_policy` (`visibility = blind | after_submit | open`), `panel_matrix` (JSON: primary/secondary axis names and categories), `matrix_type` (`expert | observer`), `language`.
 - **Panelist:** `display_handle` (themed, panel-unique, stable), `self_selected_role`, `participant_type` (`human | ai_agent`), `matrix_primary` / `matrix_secondary`, `description`, `is_decision_maker`, `language`.
 - **Axis grouping** (`_panelist_axis`): `primary` → `matrix_primary` (AI perspective as fallback; in role play the box the panelist chose – the role is a name, not a group); `secondary` → `matrix_secondary`; returns empty when the axis is empty, so the panelist is excluded rather than bucketed.
 - **Shared group reporting** (`grouped_display`): `k = 3` + complementary suppression ("Other") + None-exclusion + small-panel fallback to whole-panel level. Used by `/groups` (Home) and `/community` (facilitator) with both axes stacked, and by the pulse group dots.
 - **AI label** derived from `participant_type = ai_agent` plus the `ai_profile` role, rendered as "AI · {role}", role standardised to 2–3 words.
 - **Blind facilitator:** when `facilitator_blind = true` the UI shows only handles (names/emails hidden), though the invite mapping still exists in the database.
-

Source: Pronoia — Anonymity Design (Metodix, June 2026, status: implemented).

How Delphi Pronoia's two interfaces are built

Pronoia has two faces. A **panelist** sees one; a **facilitator** sees the other. They are not two applications and not two versions of the same screen. They are two questions asked of one running study, and the whole design follows from the fact that the two questions are different.

This guide explains the idea and the structure. The two guides beside it walk through the process — what actually happens, in order, with what on the screen.

1. One study, one truth

Both interfaces read the same two facts:

<code>current_round</code>	which round the study is on (0 = orientation, 1...N = Delphi rounds)
<code>current_phase</code>	<code>response</code> · <code>dialogue</code> · <code>revision</code> · <code>closed</code>

Everything visible is derived from those two. There is a third field, `Study.status`, and it is a **presentation** — a label derived from the round and the phase. Nothing decides anything from it.

This sounds like an implementation detail. It is the reason the interfaces are trustworthy, and it was learned the hard way. Before August 2026 the question "*which round are we on*" was answered by five surfaces reading three different sources, one of which recognised rounds by matching the strings `round1`, `round2`, `round3` — so **in a four-round study, round 4 never highlighted anywhere**. Two surfaces could disagree for fifteen seconds because one polled and the other did not.

The rule that replaced it: *round and phase are the truth; anything else is a label*. When you see the same fact in two places in Pronoia now — the round selector and a banner, say — they are two readings of one computation, not two computations.

2. The panelist interface

2.1 The idea

The panelist's question is "**what is open for me now?**" — not "what is the state of the study". The view answers that question directly and organises itself around the answer.

Five principles govern it:

P1 • The phase organises the view, not the menu. What you land on is what is open for you. The facilitator's switches decide exceptions, not the normal case.

P2 • The panel's *process* is always visible; the panel's *content* only after you have answered. "27 of 40 have answered", "two new comments since your last visit" — these say nothing about *what* the panel thinks, so they cannot anchor your judgement. The distribution, the average and other people's arguments do, so they wait until you have answered. The condition is **per thesis**, not per page: answer this thesis, see this thesis's distribution.

P3 • The frame never changes; the main column changes. Same heading, same menu button in the same place, the form always in the same spot. Not different pages — the same page in different states. A returning panelist needs recognisability more than a per-phase optimum.

P4 • An empty service is not in the menu. Every menu row is a promise that something is there. A "Group comparison" that opens empty is worse than no row. Emptiness is computed from the data, not switched on by hand.

P5 • The description and the notices are not extras. The study description sits above the form on first answering; the privacy and AI notices are always reachable. These are part of what the panelist consents to — not content that has to earn its place by being interesting.

2.2 The structure: six states

The panelist view has six states, defined by the panelist's own question:

	State	When	The panelist's question
S0	Not open yet	invited, nothing opened	"What am I coming into?"
S1	Your turn to answer	response phase, no answer from you	"What am I being asked?"
S2	Answered, round still open	response phase, you have answered	"Did that go through?"
S3	Dialogue turn	dialogue or revision phase	"What did the others think — do I change my mind?"
S4	Between rounds	round closed, next not open	"What came of it? What next?"
S5	Finished	study completed	"Where did this lead?"

S4 is the richest moment of the whole process, and it is easy to waste. It is where a panelist sees what the round produced and forms the view they will bring to the next one.

2.3 What is on the screen

The frame is constant: title, **Panel pulse** (how far the round has got), the menu, and the main column that changes with the state.

- **Panel pulse** shows process, never content: how many have answered, what has changed since your last visit, whether the round is live.
- **More services** is an index, not a drawer of extras — and by P4 it only lists what actually exists for this study.
- **AI panelists are marked wherever they have an effect.** If the panel includes AI participants, the panelist is told so, and AI contributions are marked in the distribution as well as in the text. This is not a footnote in the terms; it is a mark on every surface where an AI answer counts.

2.4 Two ways to name the journey

The same structure can be shown in two vocabularies, chosen per study:

	Story	Plain
Round 0	Threshold · Kastalia	Start · preparatory chat
Rounds	Lesche · panel rounds	Rounds · answering and dialogue
Analysis	Temple · Kassotis	Analysis · processing the results

This is naming, not structure: the states, the rules and the surfaces are identical. A study that wants the Delphic imagery gets it; a municipal panel that would find it strange gets plain words.

3. The facilitator interface

3.1 The idea

It is tempting to mirror the panelist view: the study is in the dialogue phase, so show the facilitator the dialogue phase. That is wrong, and the reason is the whole design.

The facilitator's phase is not the study's phase. The panelist asks "what is open for me"; the facilitator asks "**what is my work right now**", and the two do not coincide. While the response phase is open, the study's phase is **response** — but the facilitator's work is *monitoring*, and *preparing the next round*. A straight mirror would give the facilitator a screen that says "wait".

So the unit is not the phase. It is **(round, phase) → the facilitator's task**, and there are four kinds of task:

Study is...	The facilitator's task	Lands in
Draft	building	Planning
Round 0 open	monitoring + intervening	Collection
Round N, response	monitoring	Collection
Round N, dialogue	monitoring + moderating	Collection
Round N, revision	monitoring	Collection
Between rounds	preparing	Analysis
Finished	reporting	Reporting

Plus one thing that is not a place at all: **intervention** — advance, open, close, handle an exception. It is never its own screen; it is always present where it applies.

3.2 The structure: two levels and a bar

Both levels are at the top of the screen, in one navigation that stays in place while you scroll. A choice always brings the work directly beneath it, so there is nothing to look for. Only "What now?" sits above the navigation, and it scrolls away once the work begins (3.7).

Level 1 — the kind of work. Four buttons, always the same four:

Planning • Collection • Analysis • Reporting

Level 2 — where the work happens. Tabs that belong to the current kind of work:

Planning	Collection	Analysis	Reporting
Overview • Settings • Panelists • Theses • AI panelists • section <i>Rounds</i>	(<i>Round 0</i>) • Situation • Invitations • Facilitation • Community • Responses • (<i>Waves</i>)	Summary • Statistics & change	Export • sections <i>Reports</i> and <i>Maintenance</i>

Three newer blocks live inside these: **View as panelist** (the preview) is on the response-link card in Collection, the **Timeline — reasoning and dialogue** is in Collection's Responses, and the pick basket **Thesis material for the next round** is in Analysis above the thesis list — and opens beside the thesis editor in Design when a new thesis is written.

The bottom bar — "Always available". Ten things that belong to no round and no kind of work:

Settings • Responses • Panelists • Theses • Community • Summary • Team & access • Results • Export • Method library

A door carries its destination's name: *Responses* leads to the tab of that name in Collection, *Summary* to the one in Analysis. One place, one name — even when a place has two doors (decision K2).

Two of these deserve a note.

- **Export appears twice on purpose.** Reporting is a *phase* — what you do at the end. Export is a *function* — something you can do at any moment. Same view, two different questions, two different moments.
- **Results is the only item that is not a tab.** It opens a **Results so far** window *over* your current work. Glancing at an earlier round must not take you away from what you are doing, so it is a

window and not a destination. It contains no editing surface at all; there is exactly one way onward — **Open in Analysis** → — and all editing happens there.

Collection's level 2 is a process. The order is the order in which the work is done: *Situation* tells you where things stand, *Invitations* is the first act, *Facilitation* gathers the round's settings (deadline, dialogue phase, visibility and interaction), *Community* is communication — the same view as Community in the bottom bar — and *Responses* is the per-thesis result. *Round 0* comes first when the study runs it, and *Waves* last in the Barometer.

Collection shows the open round. There is no way back into a closed round: its material is read in Analysis, and the server refuses to change a closed round's dialogue, phase rules or schedule. Settings decided before a round begins — opening time and adding a round — are in Planning's *Rounds* section, and exceptions — changing the phase by hand and resetting an old round — in Reporting's *Maintenance*.

The bottom bar also shows your role and AI balance. They are information, not doors: they lead nowhere, but they are always visible.

3.3 Default, not compulsion

The phase decides where you **land** and what gets **weight** — it does not decide where you may go. The rail stays free, because in a Delphi study you return to planning in the middle of collection: round 2's theses are built while round 1 is still running.

When you are looking somewhere other than where the study is, the view says so. It does not stop you and it does not hide anything; it just refuses to let you believe you are looking at the current round when you are not.

3.4 "Applies to:" — the marking that prevents the worst mistake

Planning contains five things, and they are not all the same size:

What	Applies to
Description, phenomenon frame	the study
Settings (orientation, variation, language, AI mode)	the study
Rounds (dialogue, visibility policy, schedule)	each round separately
Theses	each round separately
Panel	the study — but participation is per round

So when you are preparing round 3, you must know whether you are editing round 3's theses or also round 1's. Every planning view carries a marking that says which: **Applies to: the study** / **Applies to: round 3**. Round-scoped editors take the round as an explicit parameter — never "whatever is current".

3.5 AI actions are in the background, behind a control

If the panel includes AI panelists, the Collection view shows a card that says how many there are. The **actions** — generating responses, comments, reactions, revisions — are behind **AI actions ▼**.

The reason is not tidiness. A block of AI operations rendered straight into the main column grows the very column the whole design is trying to shorten, and it puts machine operations in front of the work that is actually the facilitator's at that moment: watching a human panel answer.

If the actions are not available, the card **says why** — the first round has not been opened, round 0's AI sessions are generated from the Round 0 view, or the study has ended. An absent surface that gives no reason is indistinguishable from a broken one.

3.6 One panel, one list

The panel is a single list in which AI participants are marked — `participant_type` is a column on a panelist, not a different kind of object. But a default view listing every panelist would, in several existing studies, be data rather than a panel: measured on 25.8.2026, the database held 80 AI panelists and 33 people.

So the panel view has a filter — **SHOW: All • People • AI** — which is view state and changes nothing. The row always shows all three counts, so what is hidden is readable rather than inferred.

3.7 "What now?" — the card that also says why

The status card's top row names where you are in the arc — *Planning, Round 1 running, Pause after round 1 — prepare round 2, Ending*. Below it are the **next 2–3 steps**, and the order is itself information:

blocker before gap · gap before transition · transition before this phase's work

Each step carries three things. The **reason** says where the prompt comes from (*blocked · missing · next phase · this phase's work*); a blocking reason is amber. The **door** is a button that takes you there and scrolls the control into view. The **rationale** is a small ? that opens the method library at the right article — the rationale is a door, not a sentence written into the card.

The card blocks nothing. When nothing is required, it shows this phase's work: empty is an answer, not an empty card.

A blocker is visible before you hit it. The card asks the gates in advance, so it never prompts you past what would stop you — "Choose the instrument for round 2" appears before you press forward, not afterwards as an error.

The card is at the very top of the screen, above the navigation. You can collapse it when the space is needed for work — but even collapsed, the card shows the first step if that step is a blocker.

3.8 The phase narrows what you see

The Collection view has a row, **Informing and interaction**, below the AI runs: reminder, message to the panel, invitations. It is not a new view but **doors** into existing actions, and its contents follow the phase — in the response phase reminder · message · invitations, in dialogue only message, in the pause message · pause notice.

For the same reason some **sections fold** by phase: what does not belong to the work at hand is a chip rather than a block. Folding is a **default, not a rule** — your own choice wins, and a folded section leaves the chip that opens it.

3.9 The contents row — a view's sections on one line

Views with several sections (Responses, Community, Summary, Statistics & change, Waves) have a **CONTENTS** row at the top. Its chips are the view's sections, and a chip's name is a door: it takes you to the section and opens it if the phase has folded it (3.8). The small × after a chip hides the section. A dashed, faded chip means the section has no content in this study yet — it leads nowhere, because there is nothing there.

For example, Analysis's *Statistics & change* view holds seven analyses — thesis stats, convergence, key figures, dialogicity, shift directions, revisions and refresh — and *Summary* holds, among others, highlights and the epistemic 7×3. The contents row is their index.

3.10 Mail goes out as a letter

Every Pronoia mail goes out in two versions: as plain text and as a letter. The mail program shows the letter; a program that does not display HTML shows the same text as before. You still write only the text — the letter is built around it: the study's name at the top, `[link]` turns into a button where you wrote it, and at the end the contact person and a note that the recipient can reply directly. The letter has no images and no tracking: you will not learn when a panelist opened the message, and it was never promised to them.

4. The same round, from both sides

	The panelist sees	The facilitator sees
Round opens	S1 "your turn": theses, form, description	Collection · Responses: accumulation, reminders, advance
While answering	pulse (process only), own answer confirmed	who has answered, per thesis and per panelist
After own answer	S2, and the distribution for the theses they answered	the whole distribution
Dialogue	S3: others' arguments, anonymous, AI marked	comment stream, AI activity, moderation
Round closes	S4: what came of it, what next	preparing: what the round produced → next round's theses

Read that table downwards and the two designs are one design: **the panelist gets content when it can no longer anchor them; the facilitator gets it immediately, because their job is to see it.**

5. What the interfaces deliberately do not do

Five rules that explain absences you might otherwise read as gaps:

1. **No empty service.** A menu row that opens onto nothing is removed, not left as a promise.
2. **No second computation of the same fact.** There may be several buttons that advance a round; there is one computation of what "advance" means. Two buttons are fine, two answers are not.
3. **AI is marked on every surface where it has an effect** — not disclosed once and then forgotten.
4. **A closed round stays closed.** Its participants and its results are history. Changing them afterwards would silently change a response rate that has already been shown and possibly analysed, and comparison between rounds is the core of the method.
5. **Pronoia describes; it does not interpret.** Distributions, movement between rounds, who said what — those are description. Interpretation is the analysis ecosystem's job, and the boundary is deliberate.

Written 25.8.2026, against the interfaces as rebuilt on 23.–25.8.2026. Button and tab names are taken from the application's own English strings. The specifications behind this text are `docs/UI_Panelistin_nakyma_IA_v2.md` and `docs/UI_Vaihearkkitehtuuri_IA_v2.md` ; where they and this guide disagree, they are right and this is stale.

The panelist's process, round by round

This guide walks the panel from the invitation to the final results. It is written for the person who has been invited — *you* is the panelist — and the facilitator who prepares the panel can read it as a description of what their panelists will actually see.

The companion guide, *How Delphi Pronoia's two interfaces are built*, explains the design. This one is the walk: what is on the screen at each point, in order, and why some things are not there yet.

1. What you are joining

A Delphi panel is not a survey. A survey asks you once and adds you up. A Delphi asks you, shows you what the rest of the panel said and why, and asks again — and the interesting part is what happens in between.

Three things follow from that, and they are worth knowing before the first screen:

- **Your reasoning matters more than your number.** The rating locates you on a scale; the reasoning is what other panelists read and what the analysis finally works from. A rating without reasoning is half an answer.
- **You are anonymous to the other panelists.** They see you as a code, not a name: "*You respond anonymously — other panelists see you only as {code}*". The facilitator's visibility into who you are depends on how the study was set up, and the privacy notice says which arrangement applies.
- **Disagreement is not a failure of the panel.** It is often the finding. Nothing in the interface pushes you toward the middle, and the results screens show spread as a fact rather than a problem.

If the panel includes AI panelists, you are told so before you start. "*This panel includes AI panelists*" is stated up front, and every AI contribution is marked wherever it has an effect: beside the argument, and as a hatched share inside the distribution.

2. Before anything opens: the invitation and the consent screen

You arrive by a personal link. It is yours; it identifies you to the study without a password, which is why it should not be forwarded.

The first screen is **Before you begin**. It states what the study is, who is running it, how your data is handled, and — if the panel includes them — that some panelists are AI. You continue with **I have read this and I take part**, and **Read the privacy notice** stays reachable from the menu for the whole study.

This screen is not a formality that the design tries to get past. It is the thing you are agreeing to, so it comes before the content rather than under it.

3. Round 0: the orientation conversation

Most studies begin with **Round 0** — a conversation rather than a form. An AI assistant called **Kastalia** opens the phenomenon with you in your own words before any thesis is rated.

Kastalia says what it is, in its own words, at the start:

"I am an AI, not a person. I help you open up the questions in your own words before the round itself — I don't assess your answers and I won't tell you which view is right. I carry the data; the seeing is yours."

What is on the screen: the phenomenon under study, what the panel design already knows about your angle ("From the panel design, Pronoia knows about you" — correct it in the chat if it is wrong), and a progress line, "Answered {answered}/{total} orientation questions". You write freely; there is no scale.

You end it yourself with **I'm done →**. Then you land on your own home view, and the theses open later.

Round 0 is skipped in some studies. If it is, your first screen is the one in §5.

4. The frame that never changes

From here on, one thing is worth internalising because it makes everything else easy: **the page does not change, its middle does**. The heading is where it was, the menu button is where it was, the form is in the same place every round.

Three parts of the frame are always there:

Part	What it is
Panel pulse	how far the round has got — " <i>{n} have answered</i> ", what is new since your last visit
More services	the index of everything else (§10)
The main column	the only part that changes with the round and the phase

Panel pulse shows process, never content. It tells you that the panel is moving — how many have answered, how many new comments there are — and it deliberately does not tell you *what* the panel thinks until you have answered. There is no target number in it: it says "*27 have answered*", not "*27 of 40*". A target turns a sign of life into a scoreboard, and a scoreboard tells a late respondent that they are late.

5. Your turn to answer

This is the state you were invited for. The main column holds the theses of the current round, and for each one: the thesis, the scale, and a field for your reasoning — "*Explain your reasoning:*".

Answering, in practice

- **The rating and the reasoning travel together.** The placeholder says it plainly: "*Why do you rate this way?*" Some theses are open-ended and have no scale at all; some have two criteria (for example probability and desirability), and then you answer both — they are different questions about the same thesis.
- **Some theses are not scales.** Depending on the study you may be asked to rank options, choose one or several, place a point on a plane, distribute points, mark a time window ("*Earliest*" / "*Latest*", or "*Will never happen*"), or continue a time series. The pattern is the same in all of them: your position, then your reasoning.

- **One at a time, or all at once.** **One at a time** shows a single thesis with "*Thesis {i}/{n}*" and **Previous / Next**; **Show all** returns the full list. This is a reading preference, and it changes nothing else.
- **Your work is kept while you write.** A local draft is saved as you type and the form says "*Draft saved on this device*". If you leave and come back, "*Unsaved draft found*" offers **Restore draft** or **Discard**. The draft lives in your own browser, not in the study.
- **Saving.** Each thesis has **Save**; **Save all ({n}/{total})** submits the ones you have completed. When every thesis has an answer, the view says "*Thank you for your answers!*"

What you do not see yet, and why. During the response phase the panel's *content* — the distribution, the average, other people's arguments — is not shown before you have answered. Not because it is secret, but because seeing it first changes what you write. A first answer that has already been anchored to the panel's average is worth less to the study than a first answer that is yours. The condition is per thesis, not per page: answer this thesis, see this thesis's distribution.

Studies can be stricter or looser about this. The round's visibility setting is one of **Blind**, **After you answer**, or **Open**, and the usual setting for a response phase is *After you answer*.

6. Answered, and the round is still open

Your answer is in. Two things are now true.

You can still change it. "*You have answered all {total} theses. You can return to edit your answers during the round.*" — **Edit answers** reopens the form. Nothing is locked until the round closes — not even when the study moves into the dialogue phase: a thesis you did not get to answer can still be answered during the dialogue.

The theses you answered begin to show the panel. Under a thesis you have answered, the distribution appears: bars for the responses, a line for the mean, a dashed line for the median, a band for the interquartile range, and your own answer marked with ★ and an orange bar. Beside it, the numbers the panel uses — **Md, SD, IQR, n**.

If the panel includes AI panelists, the distribution says so: the AI share is drawn in a second colour **and hatched**, the legend names both groups — **People ({n})** and **AI panelists ({n})** — and n is given split, "*{h} human, {a} AI*". Colour alone is never the carrier of that fact, because colour alone fails for part of the panel.

7. Dialogue: the round where the panel talks to itself

The dialogue phase is what makes this a Delphi. *"You now see other panelists' anonymous answers and arguments. Explore perspectives and consider if you want to refine your own answer."*

Dialogue is open. Here the panel's content does not wait for your answer — during dialogue it is the thing being participated in, and a gate would shut out exactly the panelist who most needs to be let in. (A study can still choose a blind dialogue round; then the interface says so.)

Each thesis card is in one order, and the order is the argument:

1. **the thesis**
2. **your own answer** — so you read the others against what you actually said
3. **the distribution** — where the panel landed
4. **the discussion, behind a link with the numbers in its name** — *"Discussion · {a} arguments · {c} comments"*

The count is in the link on purpose: you can decide whether it is worth opening without opening it.

Inside the discussion

- **OTHERS' REASONING** is the arguments given with the ratings; comments are what the panel said about each other's arguments. They are different things, which is why the link counts them separately.
- Contributions from AI panelists carry the **AI** marker — the same fact as the hatched share in the chart, from the same field.
- If the study uses **role play**, the writer's role is shown next to the handle. A role a human wrote is marked *self-selected*, and a role the facilitator commissioned from an AI *AI role*. You write your own role and place it in a box on your own page, on the **Role play** card.
- On your own reasoning you do not reply but clarify: the button reads **Add a clarification**. Others see the clarification under your reasoning, marked *own clarification*, but it is not counted as dialogue – dialogue is what panelists say about each other's reasoning.
- You can **Reply**, and react with **Insightful**, **Agree** or **Challenge**. A comment that **another panelist** has replied to is locked from editing: *"Locked — the comment has replies"*. Editing it would silently rewrite what someone else answered. Refining your own comment yourself does not lock it. If your position changes after the lock, you say so in a new comment — then the change shows as dialogue, not as history edited after the fact.
- **Translate** renders a contribution in your language when translation is enabled for the panel; **Show original** puts it back.

Changing your mind is the point. If the round allows revision, you edit your answer and say what moved you — the field asks *"Why did you change your mind? (What argument or information influenced you?)"* — and you can classify the change: **More agreement, More disagreement, Nuance / clarification, Reframing.** Then **Save revision.**

Not changing your mind is equally a result. A held position with a reason that has survived the counter-arguments is stronger evidence than a position nobody challenged.

8. Between rounds — the richest screen in the study

When a round closes, the screen does not say "please wait". It says **Round {n} has closed**, and then it shows what the round produced:

"Below is where the panel landed this round: distributions, other people's reasoning, and what was left open. This is the round's interim result — not a waiting screen."

This is the moment the whole method is built around, and it is easy to waste. The round's result is in front of you with no obligation attached: nothing to fill in, nothing due. Whatever view you bring into the next round is formed here.

If the next round has been scheduled you are told when — *"Round {n} opens on {date}"* — and if it has not: *"The next round has not been scheduled yet."* The interim result does not depend on the schedule; time is an addition, not a condition.

In a **Barometer** study this is where a wave ends and the next begins; the interim result is the result, and there is no final screen.

9. When the study ends

"The study has ended" — all rounds are closed. The final distributions, the panel material, the discussion and your own path through it stay readable. If you revised your answers along the way, the interface says how often: *"You made {count} revisions to your answers based on the dialogue."*

What happens after this is outside the panel: interpretation — theming, scenarios, argument and discourse analysis — is done in the analysis ecosystem, and the results reach you as a message rather than as a new screen. Pronoia describes; it does not interpret.

10. More services: the index

More services is an index, not a drawer. It opens a list with one row per service: an icon, a name, one line about what it is, and a mark on the right.

Row	What it answers
About the study	<i>What it's about and where we are now</i>
Messages	<i>Facilitator messages and your replies</i>
My knowledge path	<i>Your contribution and activity</i>
Panel pulse	<i>What the panel thinks right now</i>
Group comparison	<i>How the groups differ</i>
Weak signals & wild cards	<i>Low-threshold foresight — flag an observation</i>
Activity and recognition	<i>Your activity compared with the panel</i>

Four things about this list are deliberate:

1. **A row only exists if the service exists.** Group comparison appears when the study has comparison axes; the signals channel appears when the study collects signals. A row that opens onto nothing is worse than no row.
2. **The mark is a fact, not a nudge.** An unread message is highlighted because it is waiting for you. Everything else is a count, and an empty service says *nothing yet* rather than pretending to be full.
3. **One service at a time, and one step back — Back to services.** Your work stays open behind it: *"Pick a service — your work stays open behind this."*
4. **The privacy and AI notice is always reachable.** It is not a service that has to earn its place by being interesting.

Two of these have a note worth reading.

Panel pulse draws one track per measure: a light band for mean $\pm$ standard deviation, a mark for the mean coloured by *where* it sits on the scale (the colour shows position, not whether the result is good), the group means as dots when the study has groups, and the ends of the scale underneath.

Consensus is not drawn as geometry anywhere: a consensus figure is a classification that needs a calibrated threshold, and that classification belongs to the analysis, not to a bar.

Group comparison is hidden when the groups are too small: *"there are not enough reportable sub-groups of at least 3 people. Results are shown at whole-panel level only."* Small groups identify people. The suppression is the promise working, not a fault.

11. Two ways the same journey is named

A study chooses one of two vocabularies, and the switch reads **View:** with **Story** and **Plain**:

	Story	Plain
Round 0	Threshold · <i>Kastalia</i> · <i>Round 0</i>	Start · <i>Round 0</i> · <i>preparatory chat</i>
The rounds	Lesche · <i>Panel rounds</i>	Rounds · <i>Answering and dialogue</i>
Analysis	Temple · <i>Kassotis</i> · <i>analysis</i>	Analysis · <i>Processing the results</i>

This is naming, not structure. The states, the rules and the surfaces are identical in both; a panel that would find the Delphic imagery strange gets plain words. In the Story view, tapping a station explains what that stage means — *"Tap a station — I'll tell you what it means"*.

Wherever you are, **You are here** marks the current station, and **Round {n} / {total}** says how far along the panel is.

12. If the study is real-time

In a **Real-Time Delphi** there is one round, and answering and dialogue are open at once in it from the start. The orientation conversation (§3) may be part of it too. *"Other panelists' answers appear continuously. You can change your answer at any time."*

The rule from §5 still does the work, thesis by thesis: a thesis shows you the panel as soon as you have answered that thesis. Updates arrive live where the connection allows, and the interface says which mode it is in rather than pretending — *"Updated just now"*, and a tooltip that distinguishes a live connection from fallback polling.

There is no "between rounds" screen in a real-time study, because there is no between. In a **Barometer Delphi**, by contrast, the same theses are returned to one measurement wave at a time: the gap between rounds is the gap between waves, and a thesis may go out of use or be replaced by a new one during it.

13. Your answers, your data

- **Anonymity toward the panel is structural**, not a setting you have to find. Other panelists see a code.
 - **Leaving is possible without erasing the discussion**. You ask the study's contact person, and they carry it out — the promise is kept by someone who is able to keep it. Withdrawal deletes your name, email address and personal link permanently; the answers you have already given stay in the material under the pseudonym, because other panelists have read and argued with them. That asymmetry is deliberate, and it is stated here rather than discovered afterwards.
 - **A closed round stays closed**. Answers to it are not edited afterwards — a response rate that has already been shown, and the comparison between rounds, would otherwise change quietly underneath the study.
 - **The privacy and AI notice** is the authority on how your data is handled in this particular study, and it is one click away from every screen.
-

14. For the facilitator: what is worth saying out loud

Three things panelists ask, that the interface answers but the invitation letter usually does not:

1. **"Do I have to finish in one sitting?"** No. Drafts survive, answers can be edited while the round is open, and the study waits.
 2. **"Why can't I see the results yet?"** Because you have not answered yet — and the moment you answer that thesis, its distribution opens. It is a rule about anchoring, not about access.
 3. **"What happens to what I wrote?"** It is read by the panel under a pseudonym during dialogue, and by the analysis afterwards. The reasoning is the material; the number is the index to it.
-

Written 26.8.2026, against the interfaces as rebuilt on 23.–25.8.2026. Button and tab names are taken from the application's own English strings, and a guard test (`scripts/test_guides_coverage.py`) fails if this guide names a control differently from the program. The

specification behind this text is `docs/UI_Panelistin_nakyma_IA_v2.md` ; where it and this guide disagree, it is right and this is stale.

The facilitator's process, from planning to handover

This guide runs a study from the empty draft to the handed-over result. It follows the order the work actually happens in, and it says at each point what is on the screen, what the panel sees at the same moment, and which actions cannot be taken back.

The companion guide, *How Delphi Pronoia's two interfaces are built*, explains why the facilitator view is organised the way it is. The one-sentence version, because everything below follows from it: **your phase is not the study's phase**. The panelist asks "what is open for me"; you ask "what is my work right now", and while the response phase is open your work is monitoring — and preparing the next round.

0. The map you are working on

Four kinds of work, always the same four buttons:

Planning · Collection · Analysis · Reporting

Ten things that belong to no round and no phase, in the bottom bar marked *Always available*:

Settings · Responses · Panelists · Theses · Community · Summary · Team & access · Results · Export · Share an observation · Method library

The names are the screen's own, not the guide's: a bottom-bar item is a door, and a door is named after the place it leads to.

The phase decides where you **land** and what gets **weight**. It does not decide where you may go: in a Delphi you return to planning in the middle of collection, because round 2's theses are built while round 1 is still running. The button marked *default* is where the study says you belong; when you are somewhere else, the view tells you so — "You are looking at {a}. The study is at {b}." — with a **Back** that returns you. It never hides anything; it only refuses to let you believe you are looking at the current round when you are not.

The navigation is at the top of the screen. The four buttons and the tabs of the chosen kind of work sit in one block that stays in place; a choice brings the work directly beneath it. Collection shows the open round — closed rounds are read in Analysis.

1. Planning: a completeness list, not a wizard

Planning is both a sequence (the first time through) and a place (round 3's theses, built mid-study). Pronoia resolves that by showing the first pass as a list of what is missing rather than as numbered steps — **Building the study**:

"Five things the first round is made of. The order is for reading; readiness is computed from what the study contains."

The five, each of which is also the way in:

Row	What it holds	Applies to
Description	what the study is about, the phenomenon frame	the study
Settings	orientation, variation, language, AI mode	the study
Rounds	dialogue, visibility policy, schedule	each round separately
Panel	who is in it, roles, AI panelists	the study — participation is per round
Theses	what is asked	each round separately

Role play and the anonymous one-shot poll are settings, and each has its own working guide: Role play: a role in a box of the panel matrix and One-shot poll at a seminar: QR code and anonymous answers. Role play needs the panel matrix to be built first.

The list is computed, not ticked off. Readiness comes from what the study contains, and the header states the situation without arithmetic on your part: *"Opening is waiting on {n} item(s)"*, or *"Everything that blocks opening is in place."* Items that are advice rather than obstacles are marked *recommended* — a thesis without a scale is open-ended, which is information, not a fault.

Two of the gates are enforced by the server, not by the screen. A study does not leave draft while the AI mode is unconfirmed, and invitations do not exist until the data-controller declaration is complete.

The gaps are named as they are: "AI mode not confirmed", "controller name", "controller contact", "legal basis for processing", "consent text still has placeholders". The list reads the server's answer instead of guessing it, so you never meet a gate you were not warned about.

And the gap is closed in the view the list leads to. In Settings, under the AI-mode choice, it says which of the two holds: "The AI mode is confirmed" or "The AI mode is not confirmed — a round cannot be opened until it is". **Confirm the AI mode** opens the same choice dialog that otherwise appears only after you hit the gate. The confirmation applies to the selected mode: change the mode and it has to be confirmed again, because a change does not carry the old confirmation with it.

Every planning view says what it applies to. "Applies to: the study" or "Applies to: round {n}", and the panel row says the hybrid out loud: "Applies to: the study · participation is per round". When you are preparing round 3 this is the difference between editing round 3's theses and quietly editing round 1's. Round- scoped editors take the round as an explicit parameter — never "whatever is current".

Rounds are not fixed at creation. "The number of rounds is not locked at creation: the length of a Delphi can be a result rather than a precondition." Dialogue is on by default for an added round; a pure feedback round is made by switching it off, deliberately:

"This round has NO dialogue phase — it collects positions but the panel never argues about them."

Building the panel. The panel is one list in which AI participants are marked — **SHOW: All · People · AI** filters the view and changes nothing. Under it, **PANEL COVERAGE** reports the **Coverage score**, the primary and secondary coverage, the cells that are empty, and — when they are — "Panel supplement recommended before the round".

Reserve enough rounds — an unused round costs nothing. The number of rounds is set when the study is created, and it is not a promise you have to keep: the program never requires a reserved round to be run. You may close the study after round two even if you reserved four; the rounds you did not open simply never existed for the panel. The reverse is harder — a round can be added mid-process, but that is the exception (see the round-management card), not the path. So when the number is in doubt, reserve the larger one.

2. Round 0: orientation, run as a sequence

Round 0 is the one place where the work genuinely is a sequence, and the view shows it as one, and the steps are named on the screen: **Configure, Open round, Send invitations, Monitor chats, Close & go to R1.**

- **Configure** sets the orientation questions, the opening text shown to the panelist, the AI assistant's deepening task and how many follow-ups it may ask. With no configuration yet: *"Create default config"*.
- **Who round 0 is for** is a choice of its own, made before opening. **Whole panel** is the default; **Selected key informants** narrows the round to the ones you tick. The scope covers AI panelists as well as humans, and the same choice drives the AI run, human access to the conversation and the cost estimate. The scope is an explicit choice, not derived from the ticks: one tick must not quietly shut the rest of the panel out.

From then on the scoped number appears everywhere round 0 is discussed — readiness, the R0 counters and the invitation card — beside the panel size. In the collect view the scope has a line of its own, and **Who round 0 is for** → leads back to the choice.

- **Open Round 0** publishes it. From this moment the panelists' links lead somewhere.
- **Invitations** are generated per panelist — **Generate invite links** — and sent from the communication centre or by hand. In a blind study the interface says what it cannot do rather than failing quietly: *"Blind mode hides the addresses: the links are created, but invitations cannot be mailed by hand."*
- **Monitor chats** gives the state of each session — **Started, In progress, Completed, Not opened** — and the average number of turns. You can read any session.
- **AI panelists' orientation chats** are generated here, not from the AI actions control elsewhere: this screen is a lifecycle you run in order, and the two buttons belong to that order. Regeneration deletes the existing AI chats and cannot be undone; the confirmation says so.
- **Close Round 0 → open Round 1** ends orientation and starts the study proper.

A study can skip Round 0 — *"Advance straight to Round 1. Sensible e.g. for a pure AI panel."*

Round 0 belongs to every variation. It comes before the panel opens, and the variation's own character begins only when Round 1 opens — Real-Time Delphi and Barometer Delphi get the orientation round too, and the only way to skip it is the same in all. Before 13.9.2026 the state machine jumped straight to Round 1 in those two, and the card was hidden.

3. Collection: the round-management card

Collection's tabs are in the order of the work: **Situation** • **Invitations** • **Facilitation** • **Community** • **Responses** (with *Round 0* first when it is run). The round-management card is in *Responses*. The settings mentioned below live in their own homes: **visibility & interaction** and the **deadline** in Collection's *Facilitation*, the **opening time** and **adding a round** in Planning's *Rounds* section, and the **phase exception** and **resetting an old round** in Reporting's *Maintenance*.

Everything about running a round happens in **Collection**, and the card marked **Round management** is where the round's state and its controls live: the round, the **Phase**, the accumulation of **Responses**, and a button whose label is the **act** — "*Move to the dialogue phase*", "*Open Round 0 (intake)*", "*Close round 2 (pause)*".

Closing a round is a pause, not an ending. Closing even the last round leaves the study open: rounds can be added. Ending the study has its own control, **End the study** (■), and it is a decision — it asks for a closing letter to the panel and marks the handoff final. Adding a round afterwards undoes it. Before 31 Aug 2026 finality was a side effect of a round, and the facilitator had no way to say *the study is complete*.

There may be several buttons that advance a round; there is exactly one computation of what advancing means, and it happens on the server. That is why the label reads as a destination rather than a verb.

A step says that it was taken. After a successful transition a notice names the destination — "*Round 0 opened · 4 key informants*". Without it, opening looked like it had failed: the only thing that changed was the button's own label, which became the next destination, and a second click in the same spot carried the study past a round.

And a step that leaves something unrun is asked first. If round 0 is unopened or unfinished, moving to round 1 asks for confirmation that names both the destination and what is being left out: "*Opening round 1. Round 0 (intake) will not be run.*" Round 0 is one-shot — orientation gathered afterwards is not orientation but post-hoc justification. Other transitions are not asked: a confirmation asked every time is learned as something to click through.

Before you open a round, two settings are worth a deliberate look, because both change what the panel experiences and neither can be repaired afterwards:

- **Visibility & interaction** — per phase. *Round 1a* — response is **blind by default**; *Round 1b* — *dialogue* is where the panel reads each other. The choices are **Blind**, **After you answer**, **Open**.

- **Deadline** — *not set* until you set it, then "*{n} days left*", **+7 days, Clear**, and *deadline passed* when it has. The deadline is your countdown; the panelist learns of it through a message, not through a clock on their screen.

View as panelist. The response-link card has a button that opens the survey in a new tab exactly as a panelist sees it at this stage — "*The preview identity is not a panel member and nothing is saved.*" The previewer is an identity of its own (an ESIK code), the consent card and the Round 0 conversation are skipped, and every attempt to write returns "*Not saved (preview)*". It does not appear in the response rate or in the reciprocity figures. **End preview** closes the link and removes the previewer; a new one can be opened at any time. This is where the question *what does this look like to a panelist* gets its answer before opening, not after.

4. The response phase: your work is monitoring — and preparing

The study's phase is **response**. Yours is **monitoring**, and the screen is built for watching rather than doing: accumulation per thesis and per panelist, and one intervention that matters — **Remind inactive ({n})**, with a confirmation that names the number, and a result that says what happened: "*Reminder sent to {n} panelists*".

Read the response rate carefully. If the panel is mixed, an AI panelist answers every thesis on every round; a combined rate is therefore partly a property of the panel's construction rather than of its engagement. Pronoia computes the human rate, the AI rate and the combined rate separately and treats the human panel as primary when there are humans — and where there is nobody to divide by, the figure is absent rather than zero, because *nobody was asked* and *nobody answered* are different facts.

The AI panel's actions live behind a control. **AI actions ▼** holds four: *Generate thesis responses*, *Generate dialogue comments*, *Generate reactions*, *Generate revisions*. They are behind a button because a block of machine operations rendered into the main column grows exactly the column this design is shortening, and puts machine work in front of the work that is actually yours: watching a human panel answer.

When the actions are not available the card says why rather than disappearing:

- "*Actions become available once the first round is open.*"
- "*Round 0 is open — the orientation conversations are started here.*"
- "*The study has ended — no new responses are generated.*"

The middle one is not only an explanation. While round 0 is open the card carries the command itself —  **AI panellists' chat**, and **Regenerate AI chats** when there are sessions to redo — next to **Open the Round 0 view**. It used to say the run was started from another view, and there was no way there from that sentence: a sentence pointing at another view is not a path to it.

Every generation passes a cost gate first. **AI run cost estimate** states the scope, the estimate in credits, the budget remaining, who pays and which limit binds — wallet balance, monthly cap, personal cap or this study's own cap — before **Continue and generate**. If the estimate is over: *"The estimate exceeds the remaining budget — some responses may be produced rule-based."*

Meanwhile, the other half of your work. This is when round N+1 gets built. The rail is free, *Applies to:* keeps the two rounds apart, and **Results** lets you glance at an earlier round without leaving what you are doing (§7).

5. The dialogue phase: monitoring and moderation

The panel now reads each other. Your view gains the comment stream, the dialogue threads — *"Dialogue threads (comment-on-comment)"* — and the AI panelists' activity in the same stream, marked.

Dialogue does not close answering. The phase adds a capability — commenting — and removes none: a panelist who did not get to answer can still answer during the dialogue, and changing an answer in the light of someone else's reasoning is the method's idea, not an exception. You need not open anything for it. A comment is locked against editing only once **another panelist** has replied to it; a panelist's own refinement of their own comment does not lock it.

Moderation here is mostly restraint. The dialogue phase is the study's learning phase, and a facilitator who answers arguments becomes a panelist with special powers. What the interface supports is the facilitator's own channel — the **COMMUNICATION CENTRE** — addressed to the **Whole panel**, a **Group**, or an **Individual**, with templates for *Reminder*, *Round open* and *Results* and placeholders ([link] , [handle] , [study]) resolved per recipient. **Group** reaches a panelist through both the home cell and any additional membership on either axis — a person who belongs to two groups receives the group message from both.

Two options there are worth understanding:

- **Also send by email** delivers the message to inboxes as well as the panel.

- **Notice only, no text** sends a mail that says a message is waiting and links to the panelist's own page, leaving the text inside the panel: *"Choose this when the message contains panel content that is not meant for inboxes."*

When a message goes to email it goes out as a **letter**: the study's name at the top, `[link]` as a button and the contact person at the end — the text version travels along for mail programs that do not display the letter. Two more choices:

- **Include »What we ask of you«** adds, after the button, a short three-point description of what the panelist is asked to do. The description is the same in every variation. It is on by default.
- **Preview the letter** shows the letter exactly as it will be sent — the same builder makes both. Placeholders are filled with sample values, and the preview sends nothing.

The delivery report is explicit about who did not get it — recipients without an address, withdrawn panelists skipped, and whether mail was sent at all.

6. The revision phase

Panelists revise their answers in the light of the dialogue, and say what moved them. Your view is the accumulation of revisions and the movement between rounds: **Revisions**, **Revised**, **Revision rate**, under the subtitle *"Panelist changes to R1 responses after dialogue"*.

Revisions is a count, not a share: there is no maximum number of revisions, and the raw figure tells you something about the panel's diligence that a percentage would flatten. **Revision rate**, by contrast, is one share, computed in one place: the *people* who changed their position divided by everyone who responded, humans and AI panelists together. Before 14.9.2026 it was computed in four places in four ways, and the community view showed 1600 %.

7. Glancing at results without leaving your work

Results in the bottom bar is the only item there that is not a tab. It opens **Results so far** over what you are doing: round pills for the rounds that have results, and a row per thesis with n, median and IQR.

The window contains no editing surface at all, and there is exactly one way onward — **Open in Analysis** → — which carries the round you were looking at with it: *"Viewing only — editing happens in Analysis."* When there is nothing to show it says which nothing: *"No results yet — the first round has not been opened."*, or *"Round 0 is the orientation; its results appear in the Round 0 view."*

Glancing is by definition something that does not interrupt. That is why it is a window and not a destination.

8. Between rounds: your richest screen too

The round closes. For the panel this is the interim result; for you it is **preparing**, and it is the only point in the process where the main column sits on the boundary between two kinds of work — analysis produces what planning consumes.

What Analysis gives you here

- **Summary** — distributions, central tendency and spread per thesis.
- **Statistics & change** (thesis stats, convergence, key figures, dialogicity, revisions) and **ROUND-OVER-ROUND MOVEMENT** — *Converged, Diverged, Stable, and "Attention — split theses": "Theses where views diverge the most — often the most interesting part of the analysis."*
- A readability verdict that is about your data rather than your patience: **Not readable yet** (*"Stability cannot be measured before the round is closed and there is something to compare against"*), **Needs action** (*"Coverage is thin. Remind the non-responders before reading the result"*), or **Readable**.

What builds the next round

- **PIRE** — the inter-round engine — proposes, one thesis at a time, how to continue: convergence, stable dissensus, noisy spread, an emergent theme, a bandwagon flag, and an operation to match. You accept, change the operation, flip the mode, or lock the thesis out of the next round. Its own help states the limit plainly: *"PIRE proposes — you confirm. It does not replace deep analysis."*
- **Build the round.** *"Build the next round's survey from the previous round's results — copy as a starting point and edit, or open with the same theses (re-rating)."* The first of these is the default in a round-based Delphi: the next survey is built from the previous results by editing, removing and adding. Plain re-rating is an exception worth justifying — in a barometer it is, conversely, the whole idea. Proposals that involve interpretation are flagged for you to reformulate before opening; you always edit the theses before the round opens. A thesis can also be **copied as the base of a new one** and **hidden** or shown mid-round behind the same gates; hiding does not delete answers.
- **Reading the discussion and picking.** The pause is the moment to go through the previous round's discussion question by question and branch by branch. **Open the whole discussion** shows a thread in full; the **Timeline** shows reasoning and dialogue in one stream, newest first.

Any reasoning or comment can be starred (☆) — *"Mark as thesis material for the next round"* — noting at the same time the **reason for the pick**. Picks collect in the basket **Thesis material for the next round**: *"A pick does not change the comment and is not visible to the panel. It is your memory between rounds."*

- **The basket is source material, not a thesis.** When you write a new thesis, the basket opens beside the editor under **Source material for the new thesis**: *"Write the thesis yourself. The contributions you select are recorded as its sources."* No button turns a pick into a thesis. The selected ones are recorded on the thesis card as **Source material**, which later tells where the question came from. Basket and source material are different things: the basket is your tool between rounds, the source material is bookkeeping. That is why they behave differently when a contribution is deleted — the basket skips it, the source material keeps the record and the reason but the content is emptied.
- **The panel's readiness for the next round** is a different question from the panel's composition: who answered last round, who withdrew, where the coverage gaps are. Composition is edited in Planning; readiness is read here.
- If a round was opened past the pause and you need its survey built after all, **Return Round {n} to preparation** closes it back to the pause and keeps the answers already given.

When you want real interpretation, it happens outside Pronoia: **Send round to DAE analysis** opens the export. Theming, scenarios, argument and discourse analysis and synthesis are the analysis ecosystem's work, under its own quality gates. Pronoia describes; it does not interpret, and this button is where that boundary is crossed on purpose.

9. Reporting and handover

Generate results report produces a descriptive report — summary, per-thesis result cards, the consensus–dissensus map, round comparison R1→R2, and highlights (greatest disagreement, most discussed, most revised). **Print / Save as PDF** and **Download Word (.docx)** take it out of the browser.

The report says what it is, in its own footer, and the sentence is worth quoting to whoever receives it:

"This report is not an analysis but its source material."

Export produces three files for the analysis pipeline: the CSV (quantitative data), the Handoff Package (responses, arguments, dialogue, revision history) and DAE_STATE.json (pipeline control data). Export

sits in two places on purpose: Reporting is a phase — what you do at the end; Export is a function — something you can do at any moment.

Anonymize study is the last step and it is irreversible: panelist names and email addresses are cleared and invite links revoked, leaving anonymous research data — pseudonyms without a key. Do it when the study is over and the material has been exported.

10. Exceptions, and what they cost

Three things sit outside the normal flow. Each is available, each is marked as an exception, and each is recorded.

Situation	What you use	What it costs
The study is stuck in the wrong phase	Change phase (exception)	<i>"The phase normally changes with the advance button. Changing it by hand is an exception: it is logged."</i> The panel sees the change immediately.
You need to look at a closed round	the round selector	<i>"You are viewing closed Round {n} — this view is read-only. Actions always target the current round."</i>
A panelist wants out	withdraw , not delete	Deleting would remove reasoning other panelists have already read and commented on. Withdrawal deletes name, email and link permanently; the answers stay in the material under the pseudonym. Cannot be undone.

11. What the panel sees while you work

You are...	The panel sees
Building the study	nothing — <i>"The study has not started yet. Please come back later."</i>
Running Round 0	the orientation conversation with Kastalia
Opening a round	their turn to answer: theses, form, study description
Monitoring the response phase	process only — <i>"{n} have answered"</i> — plus, per thesis they have answered, that thesis's distribution
Moderating dialogue	other panelists' anonymous reasoning, AI marked, and the discussion behind a link with its counts
Between rounds	the round's interim result in full, and when the next one opens if it is scheduled
Reporting	<i>"The study has ended"</i> — results, material and the discussion stay readable

Read that column downwards and the design is one design: the panelist gets content when it can no longer anchor them; you get it immediately, because your job is to see it.

12. Five things that cannot be undone

1. **Regenerating Round 0 AI chats** — the existing ones are deleted.
2. **Resetting a thesis, a round or a study's responses** — responses, revisions and comments for that scope are deleted. A round reset asks you to type the study name.
3. **Withdrawing a panelist** — name, email and link are permanently deleted.
4. **Anonymizing the study** — identities are removed and links revoked.
5. **Closing a round** — a closed round stays closed. Changing it afterwards would silently change a response rate that has already been shown and possibly analysed, and the comparison between rounds is the method itself.

Item 5 is not only guidance: the server refuses to change a closed round's dialogue, phase rules and schedule. That is why Collection shows only the open round, and a closed round's material is read in Analysis.

Everything else in this guide is reversible, and most of it is reversible by pressing the same button again.

13. What is folded, and when

Some sections are folded by default when they do not belong to the phase you are in. This is a promise, not a surprise: the table below is the whole rule, and it is generated from the same module the program reads (`lib/sectionFolding.mjs`). A guard test fails if the two disagree.

Three things are worth knowing before you read it.

1. **Your own choice always wins.** Open a section by hand and it stays open, whatever the phase says.
2. **A folded section is not a hidden section.** It leaves a chip and a line telling you how many sections the phase folded, so you can open them. In the **CONTENTS** row at the top of the view, a chip's name takes you to the section and opens it.
3. **A section not listed here is never folded.** The rule only removes; it never adds.

`task` is your work in the current phase (build · monitor · prepare · report), `phase` is the round's phase (response · dialogue · revision · closed), and `mode` is the Delphi variation. An empty column means that dimension does not restrict the section.

Section	id	Shown when task is	...and phase is	...and mode is
AI run	run_ai	monitor	—	—
R0 summary	run_r0	monitor	—	—
Respond as AI panelist	run_test	monitor	response revision	—
All responses	run_all	monitor prepare report	—	—
Round movement	ins_movement	prepare report	—	argumentative classical barometer
Epistemic 7×3	ins_epistemic	prepare report	—	—
Refresh	st_load	monitor	—	—
Convergence	st_conv	prepare report	—	—
Dialogicity	st_dialogue	—	dialogue revision closed	—
Shift directions	st_shifts	—	revision closed	—
Revisions	st_revlist	—	revision closed	—
Panel news	com_news	monitor	—	—
Participation	com_participation	monitor prepare	—	—
Signals	com_signals	prepare report	—	—

If a section you needed was folded, that is a fault in this table rather than in your work — and it is worth reporting, because the table is the rule.

14. The line that says where on the arc you are

Above the state card there is one short line. It is the only place that answers the question *where in the whole process am I* — not which screen you have open, but which stage of the study you are standing in.

It says one of these, and nothing else:

The line reads	You are
Design — build the panel and the instrument	before the first round; nothing is running yet
Round 0 in progress — key informant conversations	the orientation round, if the study uses one
Round n in progress — monitor and facilitate	a round is open and answers are arriving
Pause after round n — prepare round $n+1$ from what has accumulated	the round is closed and the next one is not open
Closing — results and next steps	no transition remains; what is left is the report and the letter

Real-Time Delphi and Barometer Delphi speak differently, because they do not walk the same arc (Classic Delphi and Argument Delphi walk it together):

The line reads	Variation
Open round in progress — response and dialogue side by side, no round comparison	Real-Time Delphi
Measurement wave n in progress — the same theses, monitor and facilitate	Barometer Delphi
Between waves (latest wave n) — thesis retirement and preparing the next wave	Barometer Delphi

Why it matters that this is a sentence and not a code. Internally each stage has a short code (A1, B1, A2 ...), and it is in the page's markup where a test can read it. It is deliberately *not* on the screen: a code would be a third vocabulary beside the sentence and the guide, and the codes only exist in round-based Delphi — in the Barometer and in Real-Time Delphi the code is empty, which would make the screen look as though the code were **missing** rather than as though it does not exist.

So the line is written to be read, not decoded. If it disagrees with what you believe the study is doing, believe the line: it is computed from the study's own state, not from the round number.

15. Retirement is the Barometer's way, not a general one

A thesis whose time has passed can be **retired** — but only in the Barometer, and only at the turn of a wave. Everywhere else the program refuses, and the refusal says why: *"Retirement is possible only in the Barometer, between waves. Let the thesis run the round to its end."*

This is a methodological rule, not a limitation of the software.

- **In a round-based Delphi a thesis runs the round to its end.** The whole point of the method is the comparison between rounds: the same thesis, asked again after the panel has seen each other's reasoning. Removing a thesis mid-round would delete the very measurement the round exists to produce.
- **In the Barometer the waves differ only by their date.** The same theses are rated from one time to the next, so a thesis can genuinely reach the end of its usefulness — the future it asked about has arrived, or the question no longer divides anyone. The facilitator decides, and a new thesis can take its place.

What retirement does, exactly:

1. The thesis **stays in the current wave to the end**. Panelists can still rate it, and the wave's results are archived into the wave series as usual.
2. **At the wave turn** it moves to the archive and is no longer asked from the next wave onwards.
3. **Its responses are kept** for analysis. The export marks it `lifecycle = retired`, so the analysis can see that the thesis is no longer being asked — a retired thesis is not a deleted thesis, and the difference matters when someone later wonders why a series stops.
4. **The decision can be cancelled** before the wave turn.

The reason is asked for and stored ("time is up", "the future is visible"). It is worth writing properly: at the turn of the next wave it is the only thing that explains why the series ends there.

Written 26.8.2026, against the interfaces as rebuilt on 23.–25.8.2026. Button and tab names are taken from the application's own English strings, and a guard test (`scripts/test_guides_coverage.py`) fails if this guide names a control differently from the program. The specification behind this text is `docs/UI_Vaihearkkitehtuuri_IA_v2.md` ; where it and this guide disagree, it is right and this is stale.

Role play: a role in a box of the panel matrix

Facilitation (Pronoia)

What role play is for

In an ordinary Delphi the panelist answers as themselves: their expertise is what they bring. In role play the panelist **answers from a role** they write themselves. The role can be their actual position, a future position (*Municipal finance director in 2035*) or a perspective they deliberately choose to represent (*Young rural entrepreneur*).

Role play suits a study when

- you want to **make assumptions visible**: once the role is written out, the reader sees from which position an argument is made
- you are looking at **future actors**: the panelist reasons as a decision-maker of 2035 rather than as today's official
- the panel is **missing perspectives** you want included: a human panelist can take a role, and the facilitator can commission a role from an AI panelist
- the study is used for **learning or workshops**, and changing perspective is the aim itself.

Role play does not suit a study whose value rests on **verified expertise**. A self-selected role is a *claimed* perspective, not a qualification, and it is always marked as such.

Division of work: the facilitator gives the structure, the panelist chooses the box

Role play is built on the **panel matrix**. The facilitator defines the matrix axes and categories, for example the main axis *Expertise* with categories *Technical, Economy, Administration, Experience* and a second axis *Stakeholder*. The categories are **boxes**. The panelist writes their role freely and places it in the box it fits best.

The role is a name, the box is a group. Group comparison, pulse group scores and reports group panelists by box, not by role text. This is the key design choice of role play: if every freely written role were its own group, every group would have one member and $k \geq 3$ protection would hide them all. The boxes are what make role-play results open to structured analysis.

The facilitator's steps

1. **Build the panel matrix** (Panelists → Panel matrix). Choose categories with role play in mind: broad enough that each gets at least three panelists, and clear enough that panelists can choose.
2. **Switch role play on** in the settings and save. The setting explains what the boxes are. If the matrix has no categories, the setting warns you: panelists can then write a role but cannot choose a box.
3. **Tell the panel what a role means in this study** (template below).
4. **Commission AI roles** if you want to supplement the panel. On the **AI panelists** tab, in the add form under **Role play: AI role**, write the role in the **AI role** field and choose a box from the same matrix. The AI answers from its role's perspective.
5. **Watch the boxes fill** before opening the round: a box that stays empty or gets only one or two panelists is not shown as a group.

Template for the panel

In this study you answer from a role. On your own page, on the **Role play** card, write the role you answer from, for example your own position in 2035 or a perspective you want to represent. Then choose the box your role fits best. Other panelists see your role next to your reasoning, but not your name.

What the panelist sees

The panelist's own page has a **Role play** card. The panelist writes their role in **Your role** (at most 120 characters), picks a box on the main axis under **Choose a box** (and optionally on the second axis), and presses **Save role**. The card says: *"Other panelists see your role next to your reasoning and comments."*

In the discussion the role appears next to the handle, marked:

- *self-selected*: the panelist wrote the role
- *AI role*: the facilitator commissioned the role from an AI panelist.

The markers never mix: a role a human chose is never shown as an AI role, or the other way round. When role play is switched off, roles are shown to no one, but the saved role and box are kept.

AI roles as panel supplements

An AI role is useful when the panel lacks a perspective you cannot get from human panelists: young people, future residents, actors from another field. The AI argues from the role's perspective, and its contributions are always marked.

Keep two things in mind:

- **An AI role is a proposal, not a voice.** It brings arguments into the discussion that people can respond to, but it does not tell you what people in that group actually think.
- **The AI counts as a member of the box.** The $k \geq 3$ threshold counts AI panelists in the group size (facilitator decision, 22 September 2026). A box with two humans and one AI is shown as a group. If you want the anonymity of the human group to be certain, keep at least three humans in each box.

Reading the results

- **Compare boxes, not roles.** Group comparison shows how the boxes differ. A single role can be read next to its reasoning, but it is not counted.
- **A self-selected role is a frame, not a sample.** A box's result tells you how the panelists who took that perspective argued, not what a whole profession thinks.
- **Separate the AI share.** In the distribution the AI share is hatched and the counts are split, so you can see how much of a box's result the AI produced.
- **Use roles in interpretation.** The combination of reasoning and role is often the richest material in the study: the same claim made from different roles shows where a view depends on position.

Technical description

- The panelist's role is stored in `self_selected_role` and the boxes in `matrix_primary` and `matrix_secondary`. Call: `POST /api/respond/{invite}/role` with the fields `role`, `matrix_primary`, `matrix_secondary`.
- An unknown box is rejected (`role_play_box_invalid`). When the matrix has categories, a role without a main-axis box is rejected (`role_play_box_required`).
- Adding an AI panelist (`POST /api/studies/{study}/ai-panelists`) takes the same fields; the role is stored as the profile's perspective.

- In the discussion response the role travels in `panelist_role` and `commenter_role` only while role play is on.

See also: Anonymity: identity modes and group reporting.

One-shot poll at a seminar: QR code and anonymous answers

What the one-shot poll is for

The one-shot poll is Pronoia's lightest way to collect answers: no invitations, no sign-in, no handle. The audience scans a QR code with their phones, answers the theses and submits anonymously. Everyone answers once.

The one-shot poll suits

- **taking the temperature of a seminar or workshop:** where the room stands on the theses before and after a discussion
- **the start or end of a Delphi:** a wider audience's view alongside the panel's results
- **teaching:** presenting the method so that everyone gets to try it.

It does not suit a vote whose outcome is binding. Anonymity means no one can verify who answered, and the duplicate-vote guard only works per browser (see **Limits**).

Before the session

1. **Write the theses.** Scale theses are answered with a slider, other theses in text. Short theses about one thing work best in a hall; three to six are enough.
2. **Choose Anonymous one-shot poll in the settings.** Tick **Poll open — accepting responses right now** and save. An unsaved poll will not open.
3. **Check the share card.** A card appears under the tick with the QR code, the **Public link** and the poll's state. The state should say that the code works now.
4. **Test with your own phone.** Scan the code, answer and submit. A second attempt from the same browser says you have already answered; that is correct. Your test answer stays in the results, so allow for it when you read them, or test on a separate copy of the study.
5. **Prepare the presentation view.** **Open presentation view** opens a white full-screen page: the study's name, a large code, the address and the prompt *"Scan the code with your phone and*

answer”. Open it in a browser tab you can project, or take a screenshot for a slide. **Copy link** puts the address on the clipboard if you want to send it to remote participants.

The address carries the study's language. Voters can switch language (Finnish, English, Swedish) at the top of the page.

During the session

What to tell the room:

Scan the code with your phone. Your answer is fully anonymous: no name, no sign-in, and no identifier is stored about you. Answer the theses and press **Submit anonymously**. Everyone answers once.

Keep the presentation view on screen throughout the poll: latecomers can still join, and the address is readable for anyone whose phone will not read the code.

Disruptions:

- **The service is busy.** The voter sees *“The service is busy right now. Wait a moment and submit again – your answers are still in the form.”* Ask people to wait a moment; nothing is lost. The server's limit is sized for a hall (see **Technical description**).
- **The network drops.** The message says the answer was not saved. The answers stay in the form, and submitting works once the connection is back.
- **Someone presses submit without answering.** The page asks them to answer at least one thesis.

Closing the poll

Remove the tick from **Poll open — accepting responses right now** and save. After that the code opens a page that reads *“The poll is closed.”* Anyone who submits at the very moment of closing sees a message that the poll was closed before their answer arrived.

You can reopen the poll the same way, for example for a second measurement after a discussion. Note that the same-browser guard applies to the whole poll: to measure the same audience twice, run the second measurement as a separate study.

Results

The answers go into the round's results and distribution like any other answers. Each submission is stored as an anonymous participant who does not appear on the panel roster or in group comparison. That is why one-shot poll results cannot be broken down by group: voters have no box in the panel matrix.

Read the result like this:

- **The distribution shows where the room stands**, not the expert panel's position. If you set it beside Delphi results, say in the report that it comes from a different group.
- **The reasons are anonymous and one-off**. They cannot be clarified or answered, so they are material, not dialogue.
- **The number of respondents is the number of submissions**. Without an identifier, it does not tell you how many different people answered.

Limits

- **The duplicate-vote guard is per browser**. The same browser answers once. The guard stops accidental repeats, but not deliberate ones: another browser or a private window can answer again.
- **No commenting**. Voters can write a reason, but cannot comment on others' answers.
- **No follow-up**. A one-shot poll creates no own page, no revision and no next round. To engage participants further, invite them to the panel in the usual way.

Technical description

- The poll page is `/poll/{study}?lang={language}` and the presentation view `/poll/{study}/qr`. The QR code is drawn in the browser (error correction level M, four-module quiet zone); the address never passes through an outside service.
- `GET /api/poll/{study}` returns the theses of an open poll and `POST /api/poll/{study}` stores the answers. Each submission creates an anonymous participant (`participant_type = anonymous`).
- The server rejects a closed poll (403) and overload (429); the page tells the voter about both.
- The per-address limit is 1,500 submissions and 3,000 page loads in ten minutes. Phones in a hall often reach the server from the same address, so the limit is sized for a hall, not for a single user.

- The duplicate-vote guard lives in the browser's storage; the server stores no identifier for the voter.

See also: Anonymity: identity modes and group reporting.

Taking part in the open Delphi study

Participation (Pronoia)

Metodix always has one open Delphi study running, and anyone interested in the topic can ask for an invitation. This guide tells you in two minutes what you are joining. Answering itself takes about fifteen minutes.

No account needed

A personal link arrived in your email, and that is all you need. There is no password and nothing to register. The same link brings you back to the panel round after round, so keep the message.

Please don't pass the link on. It is your seat in the panel, and if someone else answers with it, their views are recorded as yours.

A Delphi study is not a survey

A survey adds opinions up. A Delphi study reads them. A round goes like this:

1. **You answer first.** You rate the theses before you see what anyone else said. The order is deliberate: the first number you saw would stick to your own.
2. **You give your reasoning.** The reasoning is the most important part of your answer. The rating says where you landed; the reasoning says why, and that is what the analysis reads.
3. **You read the others' reasoning.** Under **OTHERS' REASONING** you see the rest of the panel anonymously. You can comment on a single piece of reasoning and ask its author for more.
4. **You may revise.** If some argument makes you think differently, change your answer and say what changed it. Changing your mind is not weakness but a result of the method, and it appears in the analysis as a finding of its own.

The aim is not unanimity. The aim is that those who disagree know why they disagree.

Anonymity is structural, not a promise

The other panelists see only a pseudonym. They see your rating and your reasoning; your name, your organisation and your email address stay out of sight. This is not politeness but a condition of the method: anonymous reasoning is judged by its content rather than by its author's standing.

In the distribution chart your own answer is marked **you**, and in the discussion your own turns appear to you as **You**. To everyone else the same row sits under a pseudonym.

Only the organisers at Metodix see your contact details, and they use them solely to send the invitation and the round messages.

What else is on the page

- **DISCUSSION** gathers the conversation around a thesis into a thread, and **YOUR REASONING** shows your own text, which you can refine.
- **More services** shows what has been stored about you, and it is also where you can withdraw from the study.
- **Weak signals & wild cards** is a place for observations that do not fit the theses but may still change the picture of what is coming.

If something doesn't work

Open **Share an observation** and write down what you tried and what happened. The message goes to the people who build Pronoia, not to the study's facilitator, and it has no effect on your answers. The most useful thing you can tell us is where you looked first — even if it was the wrong place.

Short answers

Question	Answer
Do I need an account?	No. The personal link is enough.
How long does this take?	About fifteen minutes per round. You can answer at your own pace.
Will anyone see my name?	Not the other panelists. The organiser sees your contact details for the invitation.
Can I change my answer?	Yes, for as long as the round is open. The change and its reason are part of the material.
Can I stop halfway?	Yes. Withdrawal is in the More services menu.

The facilitator trial in five stages

Facilitation (Pronoia)

You have a trial of Delphi Pronoia. In ten minutes this guide walks you through what running a Delphi study actually involves.

The stages are written as goals: each one says what you are trying to find out, not which button to press. Finding the route is part of the trial — and if you can't find it, that is the most valuable thing you can tell us.

1. Open your own example study

The invitation link lets you create an account and a password. After that your desk holds a study that already has a panel, theses and answers. It is yours: you can change anything in it without breaking anything.

Goal: find out what the study asks and how many people have answered.

Both answers are on the **Theses** and **Panelists** tabs. Note the basics of the method while you are there: a thesis is a claim to take a position on, not a question.

2. See where the panel landed

Goal: find one thesis the panel disagrees on, and find out on what grounds.

The distribution shows how the answers fell. Below it are the statistics: mean, median, standard deviation and the middle half. In a Delphi study the spread is often more interesting than the mean, because it tells you whether the panel is of one mind or split.

The reasoning, though, is why Delphi studies are run at all. Read it and see whether the disagreement lies in the ratings or only in the reasons behind them.

3. Read the dialogue

Goal: find a place where someone revised their position, and work out what changed it.

Panelists comment on each other's reasoning and may revise their own rating. The revision and its reason are stored separately. This is what separates a Delphi study from a survey: the material keeps a record of which argument moved whom.

4. Change something

Goal: add one thesis of your own and see how it looks to a panelist.

Use the panelist preview: you see the same view as your panel, and nobody is notified. It is the fastest way to check whether a thesis makes sense.

Keep one rule in mind: a thesis added while a round is open means that part of the panel answered a different set from the rest. Pronoia does not prevent it, but it does record it.

5. Build a study of your own

Goal: take one real question of yours as far as the theses.

The wizard on your desk asks for the variation (Argument, Classic, Real-Time or Barometer), the number of rounds and the default scale. On the **AI panelists** tab you can add synthetic perspectives if you want to see the dialogue move without a real panel.

When you invite real panelists, the **Invitations** tab creates a personal link for each of them. You can send them from Pronoia or copy them into your own email.

The limits of the trial

Limit	Amount	Why
Duration	14 days	A trial is for getting to know the tool, not for production use.
AI	5 euros' worth	AI panelists and analysis helpers are charged by use.
Real panelists	10	Across all of your trial's studies together. AI panelists you may add freely.

Hitting a limit is not an error but the end of the trial. Tell us what you were doing, and we will agree on what comes next.

Tell us what didn't work

At the bottom of the screen there is **Share an observation**. Use it freely. The valuable part is not that something was broken, but where you looked first and what you expected to happen. Pronoia's most common defect is a feature that exists but cannot be reached from the screen where it is needed — and only a new user can see that.

DAE analysis: the pipeline, phases and quality gates

Analysis (DAE)

DAE (Delphi Analysis Ecosystem) is the ecosystem's **interpretive** layer. Where Pronoia describes — *what the panel said and how views moved* — DAE answers *what it means*. DAE is the methodological interpretive authority, and it works under **quality gates (κ / QVG) and human review**.

The boundary is kept strict on purpose. If description and interpretation blurred — or if assistants generated "ready" conclusions — a false sense of precision would arise. A Pronoia report is not an analysis but its source material; interpretation is done in DAE.

Handoff in — what DAE reads

DAE receives Pronoia's **export contract**: an eDelphi CSV, a Handoff Package (JSON, read by DAE CP-0/CP-2) and DAE_STATE.json. The package carries the orientation, the panel matrix type, the argument space (theses × members), the epistemic coverage, the perspective summary and the validation status. Downstream verifies this contract before starting — a malformed dataset cannot silently enter analysis.

The pipeline: 21 checkpoints, eight phases

The DAE orchestrator runs the material through **21 checkpoints (CP-0...CP-20)**, one phase at a time, with operational discipline and User Choice Points. The phases:

Phase	Checkpoints	What it does
Foundation	CP-0...CP-2	Panel validation (MaxPanelist, PQS) + quantitative base (MaxQuant: distributions, clustering, ECI)
Content trilogy	CP-3...CP-5	Dialogue analysis (MaxDA), argument analysis (MaxAA: Toulmin + epistemic 7×3 coding), critical discourse analysis (MaxDiscourse)
System trilogy	CP-7...CP-10	Causal layers (MaxCLA / depth), multi-level perspective (MaxMLP / dynamics), soft systems (MaxCATWOE / purpose) + adaptive panel supplementation
Futures	CP-11... CP-13	Weak signals (MaxWS), signal validation, tension map (MaxTension: dissensus)
Synthesis	CP-14	Cross-trilogy integration (MaxGrandSynthesis): keystone, stability, polyphonic synthesis, falsification + perspective depth
Report	CP-15... CP-18	Scenarios (MaxScenarios) + audience-specific stakeholder reports (MaxReports)
Continuity	CP-19... CP-20	Case library: store the run and its lessons for the next study

The phases group into **trilogies**: content (what the panel argued), system (depth · dynamics · purpose) and futures (signals · tensions). Synthesis integrates every lens.

The quality engine — when the human enters the loop

DAE's quality engine ensures the qualitative checkpoints are valid:

- **Smoke and adversarial tests** plus a **gate validator** gate every phase.
- **κ (kappa)** measures coding agreement; low κ routes to **human coding** (QVG — Quality Validation Gate).
- **RED ZONE** phases (synthesis, scenarios, reports) require special human judgement.

When automated confidence is not enough, the analyst is brought into the loop — and the facilitator decides how far interpretation should reach. This is the arc's methodological guardrail: delegation is not a loss of control.

Outputs

The pipeline ultimately produces **scenarios** (several paths: causal layers, multi-level perspective, tensions, soft systems) and **audience-specific reports** with data-provenance markers. The synthesis is falsification-tested and keystone-validated — not an impression but a traceable, gated interpretation.

What if DAE is not in use?

Not everyone has the DAE ecosystem. For them, Pronoia offers **descriptive report formats** (summary, per-thesis result cards with arguments, round comparison) in HTML/PDF/Word. They give a proper summary but stay deliberately descriptive — they do not replace DAE's interpretation.

Sources: [help/06-dae](#); the DAE core ([dae-core](#)) and phase modules.

Glossary

Abductive thesis — an open, exploratory thesis (E orientation): what is emerging.

Argument Delphi — a variation: dialectical theses and dialogue over rounds; aims at dissensus — maps tensions and reasoning (lead indicator CDI). The same process and phases as Classic Delphi.

Argument Simulator (Route B) — generator of synthetic, orientation-aware arguments when no real panel exists.

Barometer Delphi — a variation: the same theses rated from one time to the next, wave by wave; change over time is the finding (BCI). The only variation in which a thesis can be retired.

Basket (thesis material for the next round) — the facilitator's picks between rounds; does not change the comment and is not visible to the panel. Opens beside the thesis editor as source material.

Classic Delphi — a variation: the same process and phases as Argument Delphi; aims at reasoned consensus and measures it with consensus indicators (CCI). The value stored in the contract is `Classical`.

Commensurability — enough shared ground for genuine dialogue; the goal of panel design (not full consensus).

Convergence — the narrowing of positions across rounds.

CP-0...CP-20 — the DAE pipeline's 21 checkpoints, grouped into eight phases.

Dialectical thesis — a claim taken a position on by probability and desirability (N/S/R/D).

Dissensus — structured, well-understood disagreement; in Delphi as valuable a result as consensus.

ECI (Epistemic Coverage Index) — how many of the 21 argument cells (7 knowledge bases × 3 functions) are filled.

Export contract — Pronoia's validated payload to DAE (eDelphi CSV, Handoff Package JSON, DAE_STATE).

Home cell and additional membership — a panelist's primary place in the panel matrix and any additional groups on the same axis; group size is counted in people, not memberships.

Handoff — a structured transfer between stages: Planning Handoff Package (in), export contract (out).

k-anonymity ($k = 3$) — no sub-group of fewer than three people is reported at group level.

κ (kappa) — a coding-agreement measure in DAE; low κ routes to human coding.

Lock-in — a single perspective claiming universality and crowding out the rest.

Closing — the study's final pause: the facilitator writes a letter about the results and next steps, and it appears on the panelist's home page. Not a state that is set but a task that is done.

Orientation (N/S/R/E/D/D+) — the study's locked standpoint; sets theses, panel type and between-round stance.

Panel matrix — a two-axis designed epistemic space (e.g. expertise $\times$ stakeholder).

Pause — the state between rounds (`round_closed`), in which the facilitator reads the discussion, picks and builds the next round; for the panelist an interim result, not a waiting room.

Perspective = position + view — the same phenomenon looks different from different positions; the aim is to identify what each can/cannot see.

PHS (Perspective Health Score) — panel diversity on 0–12; guides supplementation.

Pick — starring (☆) a reasoning or comment as thesis material for the next round, with a reason; a methodological move from dialogue to a new instrument, not a bookmark.

PIRE — Pronoia's orientation-aware round-gap assistant.

Planning Handoff Package — the planning payload into Pronoia: orientation, matrix, theses, perspective context.

PQS (Panel Quality Score) — DAE's panel-validation quality score (CP-1).

Preview (View as panelist) — the facilitator sees the survey through a panelist's eyes under a preview identity of its own; nothing is saved, and the previewer does not appear in the panel's figures.

Pseudonym (handle) — a stable word handle (e.g. Aurora); anonymous to others, recognisable in dialogue.

QVG (Quality Validation Gate) — DAE's gate that routes to human review when AI validity is not enough.

Real-Time Delphi — a variation: one round in which answering and dialogue are open at the same time from the start; no comparison between rounds (RCI). The aim is set thesis by thesis.

Retirement (retiring → retired) — a Barometer Delphi concept: a thesis that has lost its future relevance is archived between waves. The decision and its taking effect are two moments; the replacing thesis names the one it replaced.

Revision rate — the share of people who changed their position after the dialogue among everyone who responded, humans and AI panelists together; computed in one place. Not the same as the number of revisions.

Round 0 — the orientation: an AI-assisted intake that motivates participants and maps the phenomenon before the instrument is locked. Belongs to every variation; the only way to skip it is the same in all.

Source material (provenance) — the contributions recorded on a thesis card as the ones the thesis came from; a deleted contribution keeps its record, its content is emptied.

Trilogy — a DAE phase group: content, system (depth/dynamics/purpose), futures.

Variation — the cadence of rounds: Argument Delphi, Classic Delphi, Real-Time Delphi or Barometer Delphi. The name is a proper name and is not translated. `studies.variation` is the truth from which `delphi_mode` is derived.

Wave — one measurement in a Barometer Delphi: the same theses rated again, the waves differing only in their date. A wave is not called a round, so that the words do not blur.

Phases (1a/1b) — a round's response (1a) and dialogue (1b) phases. A phase adds a capability, it never removes one: answering stays open in the dialogue.

Visibility (blind / after-submit / open) — the per-phase control of when a panelist sees others.

Frequently asked questions

Can I use just one part of the ecosystem?

Yes. The three products work together as a pipeline but also apart: you can produce **only a design** (DPE), run **Pronoia without DAE** (descriptive reports), or **analyse an existing dataset with DAE**. The handoffs join the parts when you use them in sequence.

What is the difference between Pronoia and DAE?

Pronoia is **descriptive** (what the panel said and how views moved). DAE is **interpretive** (what it means), under quality gates and human review. A Pronoia report is DAE's source material, not an analysis.

Do I need DAE to get results?

No. Pronoia produces descriptive reports (summary, per-thesis result cards, round comparison) in HTML/PDF/Word. They are a proper summary but stay descriptive.

How do orientation and variation differ?

Orientation (N/S/R/E/D/D+) gives the **standpoint** — from where the future is studied. Variation (Argument Delphi, Classic Delphi, Real-Time Delphi, Barometer Delphi) gives the **cadence** — how rounds are paced. They are different axes. Classic and Argument Delphi run the same process and differ only in the aim (consensus / dissensus); Real-Time Delphi runs one round; Barometer Delphi rates the same theses again from one wave to the next.

Why is the response phase protected from others' answers?

To protect the core sequence *independent judgement* → *feedback* → *reconsideration* from social pressure (anchoring, bandwagon, dominance). The default preset (**Classic**, used by both Classic Delphi and Argument Delphi) uses **after-submit**: no anchoring before you commit, feedback right after. The protection concerns seeing, not answering: in the dialogue phase an unanswered thesis can still be answered.

Can the facilitator see panelists' real identities?

By default yes (for reminders and quality control). A per-panel **blind-facilitator** mode hides them, so the facilitator also sees only pseudonyms.

How are AI panelists marked?

Always and everywhere: "AI" plus a 2–3 word role description (e.g. *AI · Economist*). Never hidden, never presented as human.

How does group reporting protect anonymity?

With k-anonymity ($k = 3$): sub-groups of fewer than three are not reported, and the rule applies independently on each matrix axis. On small panels the pulse is shown at whole-panel level only.

When do I close the study?

When the panel has converged, or when dissensus is stable and well-understood — further rounds would add cost without insight. Read the pulse (convergence/dissensus per thesis) before deciding.

Are AI assistants' suggestions finished conclusions?

No. Generative assistants (e.g. PIRE's deepening suggestions) are always **drafts that require an audit** and that the facilitator owns. Interpretation is done in DAE under quality gates.

Role quick-guides

Short entry points: where to start, the most important guides, and the one next line whose crossing raises your level.

Researcher / designer

Planning (DPE)

Your job: the research design — orientation, phenomenon, panel, questions.

Start: Orientation, variation and question architecture → Building the panel.

Next line: from assembling a single study → the **Planning Handoff pipeline** that seeds Pronoia without rebuilding.

Facilitator

Facilitation (Pronoia)

Your job: running Pronoia — Round 0, rounds, dialogue, anonymity, reports.

Start: The iterative process → Anonymity. Whole picture: Integration guide.

Next line: from building each round by hand → let **PIRE suggest** the between-round operations (you confirm); from reading descriptive reports → hand a round to **DAE** for interpretation.

Developer-partner

Analysis (DAE)

Your job: integration, handoffs, MCP/plugins, installation.

Start: Integration guide (handoff schemas) → DAE pipeline → Anonymity: data model.

Next line: from handling individual transfers → the **handoff contracts** (Planning Handoff Package, export contract) that downstream verifies automatically — a malformed dataset cannot slip through.

Panelist

Facilitation (Pronoia)

Your role: you take part in the panel — answer the theses, read others' reasoning and revise your position.

What you need to know: you answer independently first (others' answers are hidden until you commit), then you see an anonymized picture of the panel and can comment and update. You are anonymous to others under a pseudonym (e.g. *Aurora*); the facilitator may see your identity unless blind mode is on. AI panelists are always labelled "AI". The Home view shows a "since your last visit" digest — come back to see how the panel moved.

Next line: from a one-shot answer → **coming back to revise:** Delphi's value comes from reconsidering in light of others' reasoning.

→ A light participation guide is also in the app's in-app help.

Decision-maker / stakeholder

Analysis (DAE)

Your role: you read the results and reports to support decisions — you don't run the study yourself.

What you need to know: distinguish **descriptive** from **interpretive**. A Pronoia report tells you what the panel said and how views moved (distributions, consensus/dissensus, arguments) — it is source material. A DAE report interprets: themes, scenarios, synthesis under quality gates. **Dissensus is a result**, not a failure: stable disagreement signals a real tension. Scenarios are alternative paths, not forecasts.

Next line: from reading a single figure → **ask what is description and what is interpretation**, and what a conclusion rests on (data provenance).

Start: Overview (descriptive vs. interpretive) → DAE guide (outputs).

Slides and deep dive

Presentation slides

The slides are **generated from the same markdown source** as these guides — in one build, as reveal.js decks. They are always in sync with the guides, not a separate slide library to maintain. Navigate with the arrow keys or space; the down arrow opens a section's inner slides.

Section	Slides
Whole guide	Slides
Overview	Slides
Integration	Slides
Delphi variations	Slides
Planning (DPE)	Slides
Panel design	Slides
Facilitation (Pronoia)	Slides
Analysis (DAE)	Slides
Reference	Slides

Worked example

One deeper learning artifact is **Food Paradigm 2050** — the whole planning pipeline run end to end (orientation D, Route B, 19 roles, PHS 12/12, 27/27 quality gates). It shows the method in practice; see also the Simulation block.

Module	Relates to	Formats
Worked example — Food Paradigm 2050	Planning (DPE)	Case card

Note — the worked example is from 2026-05. Its *method* (orientation, panel, theses, perspective layer, handoff) still holds, but product specifics have moved on since. The most up-to-date source is always this guide (DPE/Pronoia/DAE).