

ORIGINAL FRAMEWORK / AUGUST 2026

The AI Intent Gap Framework

A practical field guide for making outcomes, context, boundaries, and execution paths visible before AI acts.

CORE DEFINITION

The AI Intent Gap is the distance between what a person wants accomplished and the instruction an AI system can actually act on.

Published by BeforePrompt • beforeprompt.com/resources

1. Why this gap matters

AI can only act on the information, constraints, and execution path it receives. A capable model can still produce unusable work when the instruction leaves the intended outcome implicit.

The gap is not a demand for longer prompts. It is a diagnostic: identify the smallest missing piece that would materially change the result, then add that piece before execution.

Four places intent breaks

01

Outcome gap

The request names an activity but not the result, decision, or definition of done.

02

Context gap

The model lacks audience, background, evidence, or assumptions a colleague would already know.

03

Boundary gap

Constraints, exclusions, privacy limits, approval rules, or evidence requirements are missing.

04

Execution gap

The instruction does not say how the output will be checked, handed off, or acted upon.

Operational consequences

- Rework: correction loops repeat context that could have been present in the first instruction.
- Inconsistency: similar tasks produce different formats, assumptions, and quality levels.
- Risk: sensitive data, unsupported claims, or missing approvals reach the wrong path.
- Weak learning: teams see output quality but cannot identify which instruction habit caused the failure.

2. The five-stage maturity model

Maturity is measured by how reliably an organization converts intent into inspectable instructions—not by prompt length, model brand, or the number of templates it owns.

Level	Stage	Operating habit	Signal
1	Reactive	Correct weak output after it arrives.	Correction loops dominate.
2	Structured	Name outcomes, inputs, constraints, and output shape.	Common requests become clearer.
3	Reusable	Maintain good instructions as owned assets.	Teams stop rebuilding recurring work.
4	Governed	Add privacy boundaries, policy checks, and approvals.	Risk and accountability become explicit.
5	Adaptive	Learn from failure patterns without default raw-prompt surveillance.	Instruction quality improves systematically.

A practical progression

Move one recurring workflow at a time. First make the outcome and inputs explicit. Then preserve the instruction, name an owner, add approval and privacy boundaries, and only then measure recurring failure patterns.

GUARDRAIL

Do not use raw employee prompts as default telemetry. Prefer aggregate quality signals, explicit opt-in, de-identification, and purpose-limited review.

3. Check intent before send

Use these questions on a high-value instruction before it leaves the editor:

Check	Question
Outcome	What decision, artifact, or action should exist when this is complete?
Audience	Who will use the result, and what do they already know?
Inputs	What evidence, examples, data, or source material must the model use?
Boundaries	What must be excluded, protected, verified, or approved?
Output	What format, length, structure, or acceptance criteria make the result usable?
Next action	Will a person review it, will a system execute it, or is it only exploratory?

When to stop and ask

Ask a clarifying question when the missing information would change the decision, when a required source is absent, when sensitive data may be exposed, or when the output could trigger a consequential action.

4. Applying the framework

Before

RAW REQUEST

Create a launch plan for our new product.

Diagnose

- Outcome gap: launch plan could mean a calendar, campaign, operating checklist, or decision memo.
- Context gap: the product, market, audience, channel, budget, and launch date are unknown.
- Boundary gap: no claim rules, legal approvals, exclusions, or budget constraints are stated.
- Execution gap: ownership, review cadence, and go/no-go decision are absent.

After

ACTIONABLE INSTRUCTION

Create a six-week launch operating plan for [PRODUCT] aimed at [PRIMARY AUDIENCE]. Use the attached positioning, channel budget, and launch date. Return a week-by-week table with owner, deliverable, dependency, success signal, and approval gate. Flag missing inputs before drafting. Do not invent performance claims or legal approvals.

The stronger instruction is longer only because the work requires those decisions. For a simple request, the smallest useful fix may be one sentence.

5. Governance and evidence

BeforePrompt publishes this framework as an original operating model. It is a product-led perspective, not an independent academic standard. Any future benchmark using the framework must disclose the dataset, selection rules, scoring anchors, reviewer agreement, limitations, and BeforePrompt's commercial interest.

Evidence rules

- Separate observed facts from interpretation and product hypotheses.
- Do not claim benchmark results before the corpus and review are complete.
- Do not publish private customer or employee prompts without explicit permission.
- Report sample limitations and disagreement instead of hiding uncertainty.
- Version the rubric and preserve a holdout set before publishing comparative claims.

Current product boundary

Available now: a Chrome extension inside ChatGPT, Claude, Gemini, and Microsoft Copilot; local first-pass checks; deliberate cloud-assisted actions; and a supporting reusable-instruction library.

Currently testing: personal publishing and sharing workflows, team templates, policy guardrails, and aggregate coaching. Additional AI integrations are exploration. Broader routing, agent, desktop, IDE, API, and enterprise-control-plane concepts are long-term direction.

READ NEXT

Current capability status: beforeprompt.com/capabilities

Privacy architecture: beforeprompt.com/resources/privacy-architecture

Research methodology: beforeprompt.com/resources/prompt-quality-benchmark

BeforePrompt is an independent product built in India. It is not affiliated with, endorsed by, or sponsored by OpenAI, Anthropic, Google, or Microsoft.