

With the rapid evolution of AI language models, running multiple distinct models like **GPT**, **Claude**, **Gemini**, **Grok**, and **Perplexity** side by side is an exciting frontier. Instead of bouncing between tabs or separate interfaces, many teams want a *shared conversation thread* that stitches responses seamlessly, enabling powerful multi-model synergy.

In this post, we'll break down best practices for **multi model AI chat** workflows based on my experience consulting on AI rollout in research and strategy teams, highlighting Suprmind's recent innovations. We'll cover the benefits of a *shared conversation thread* vs tab switching, explore orchestration modes like *Sequential mode* and *Super Mind mode*, and dive into advanced techniques around surfacing disagreements using *DCI* (Disagreement, Correction, and Integration) frameworks.

Why Run Multiple AI Models in One Shared Conversation Thread?

It's tempting to open multiple tabs—one for ChatGPT, another for Claude, a third for Gemini—and then manually synthesize results. But this tab switching has major drawbacks:

- **Cognitive load:** Juggling different contexts and writing styles from various UIs is draining.
- **Fragmented history:** No centralized log ties together the whole dialogue, making auditing or backtracking frustrating.
- **Slower iteration:** Each AI interaction is siloed, adding latency to workflows that require cross-model validation or compounding reasoning.

Moving to a single *shared conversation thread* for multiple models offers compelling wins:

- **One unified transcript:** Every prompt, response, disagreement, and correction lives in one place—easy to export or audit when needed.
- **Context continuity:** Long-running reasoning chains stay intact across AI model boundaries, enabling complex workflows.
- **Rapid multi-model comparisons:** You send one prompt to multiple AIs simultaneously and see their responses side-by-side for immediate synthesis.
- **Seamless orchestration:** Automate workflows that leverage the unique strengths of each model in sequence or parallel.

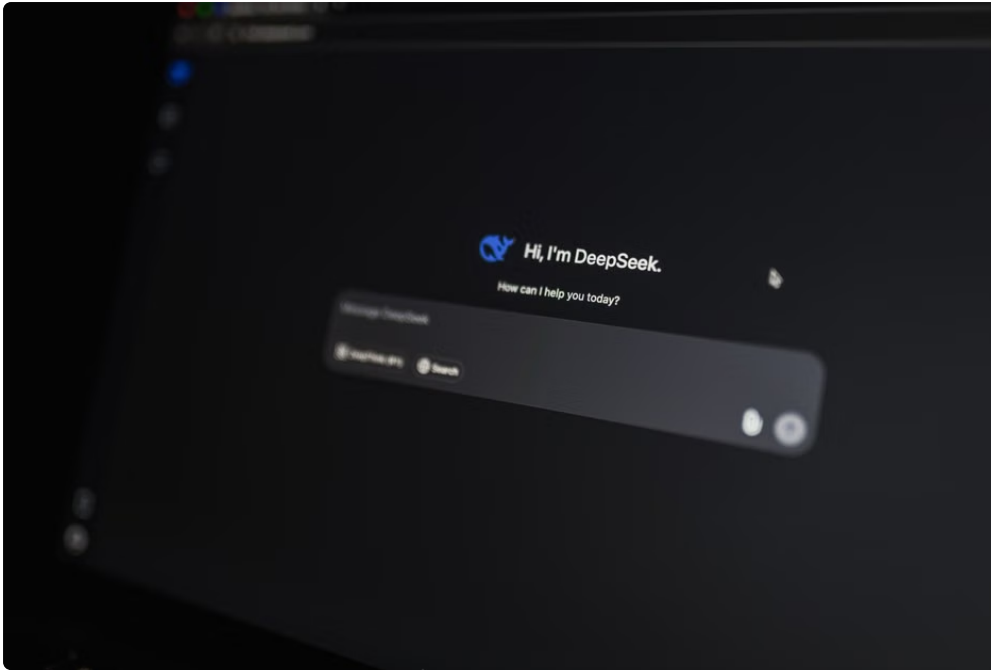
Introducing Suprmind's Orchestration Modes: Sequential & Super Mind

Suprmind—a company pioneering multi-model AI experiences—has developed two cutting-edge interaction modes to maximize the value of multi-model AI chats within one thread:

Mode Description **Ideal Use Case** **Sequential mode** Models respond one after another, each building on the previous AI's output, allowing *compounding reasoning*. Complex workflows where you want GPT, Claude, and Gemini to refine and extend an answer iteratively. **Super Mind mode** Parallel calls to models with aggregated synthesis and conflict-mapping to present consensus or highlight disagreement. Quickly compare Grok, Perplexity, and GPT to surface diverse perspectives or spot factual conflicts.

How to Send One Prompt to Multiple AIs in a Shared Conversation Thread

From a user workflow standpoint, here's how Suprmind's platform enables this with minimal tab switching:



[AI risk register](#)

1. Type your prompt once in a unified chat interface.
2. Select the AI models you want involved (e.g. GPT-4, Claude 2, Gemini-3).
3. Choose orchestration mode: *Sequential* for stepwise reasoning or *Super Mind* for parallel perspective synthesis.
4. Hit send; Suprmind distributes the prompt accordingly and aggregates responses into the shared thread.
5. Review model outputs side-by-side or in sequence right below the prompt.
6. Optionally highlight where AI responses conflict, triggering DCI workflows to track corrections.

This approach avoids tab switching entirely and creates an auditable conversation history—a huge win for compliance and strategy teams.

Sequential Orchestration and Compounding Reasoning

Sequential mode is powerful when addressing nuanced problems requiring progressive refinement. Imagine you want Gemini to draft a market strategy, then have Claude critique it, followed by GPT-4 synthesizing a final recommendation:

- **Step 1:** You prompt Gemini for an initial market entry approach.
- **Step 2:** Claude reviews Gemini's draft, points out weaknesses or gaps.
- **Step 3:** GPT-4 takes the previous two responses and composes a consolidated strategy incorporating feedback.

This controlled pipeline transforms multiple model outputs into a compounding reasoning chain versus fragmented chat bubbles across tabs.

Benefits of Sequential Orchestration

- Clear provenance showing how each model contributed over time
- Enables layering expertise—some models better at creativity, others at critique
- Improved answer quality via iterative refinement
- Simplifies export of the entire dialogue for compliance audits or stakeholder sharing

Parallel Orchestration with Synthesis and Conflict Mapping

Parallel orchestration (aka Super Mind mode) excels when you want breadth rather than depth, gathering multiple viewpoints simultaneously:

- Send the same prompt to GPT, Claude, Grok, and Perplexity.
- The system aggregates responses and visually maps points of agreement and disagreement.
- A summary synthesizes consensus and calls out where models conflict—ideal for fact-checking or diverse ideation.

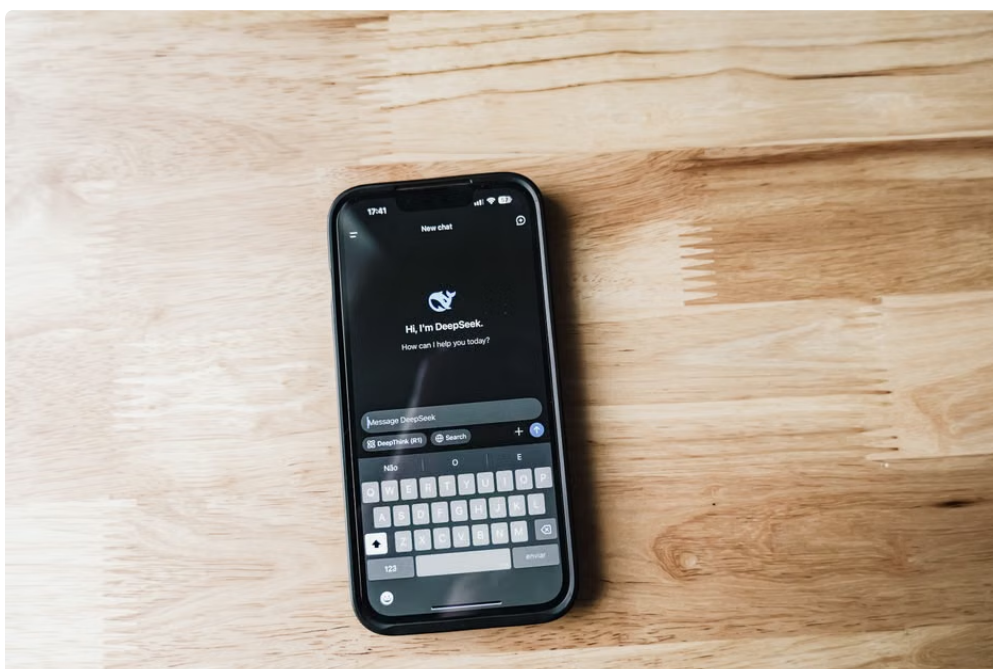
To illustrate, say you ask “What are the major risks for launching X product in market Y?” Grok might emphasize regulatory risks, while Perplexity highlights supply chain issues, and GPT provides a general overview. The Super Mind mode will flag these differing angles clearly for your review.

Why Conflict Mapping Matters

AI models trained on varying datasets and architectures often disagree. It’s crucial to surface these divergences explicitly rather than glossing over them. Suprmind’s Disagreement, Correction, and Integration (DCI) framework helps teams:

- **Identify disagreements:** Which facts or interpretations conflict?
- **Track corrections:** Who or what resolved the disagreement?
- **Integrate learning:** How does the team’s final artifact incorporate multiple views?

By codifying disagreements and corrections directly in the shared thread, compliance-heavy environments (like finance or legal) get a transparent audit trail of how the final AI-driven decision was formed.



Practical Tips for Implementing Multi-Model AI Chat with Suprmind

- **Define your AI roles:** Determine which model is your creativity engine (e.g. GPT), your critic (Claude), your data oracle (Perplexity), etc. This guides orchestration.
- **Use Sequential mode for projects that benefit from reasoning chains** and stepwise validation.
- **Opt for Super Mind mode when you need rapid, multi-angle feedback quickly.**
- **Leverage DCI workflows:** Annotate disagreements to build a culture of continuous improvement and auditability.
- **Export artifacts regularly:** Always ask “What is the artifact I can export and send?” to keep stakeholders aligned.

Conclusion

The future of AI workflows is multi-model and multi-context—no single model has all the answers. Using tools built on Suprmind’s foundational work with **multi model AI chat**, you can *send one prompt to multiple AIs* and harness the best of GPT, Claude, Gemini, Grok, and Perplexity in a single shared conversation thread.

This approach eliminates inefficient tab switching, supports complex compositional reasoning with Sequential mode, and surfaces divergent perspectives for richer synthesis in Super Mind mode. Incorporating DCI practices ensures auditability and trust in your AI outputs.

As you experiment with these workflows, keep a running list of moments when “AI said this confidently and it was wrong” — those milestones sharpen your correction tracking and improve the final artifact’s quality.

Ultimately, the goal is auditable, collaborative, and frictionless AI-assisted workflows—not just flashy feature lists or marketing fluff. Suprmind’s innovations make that future tangible today.