

In today's rapidly evolving AI landscape, leveraging multiple large language models (LLMs) like ChatGPT and Claude within a single workflow is crucial for teams aiming to maximize insight, accuracy, and efficiency. **Suprmind** is at the forefront of enabling such multi-model collaboration by offering flexible chat orchestration modes—especially its groundbreaking *Sequential mode* and *Super Mind mode*. Understanding how to effectively @mention only Claude and GPT in Suprmind unlocks powerful shared-thread, multi-model chats that outshine clumsy tab-switching workflows.

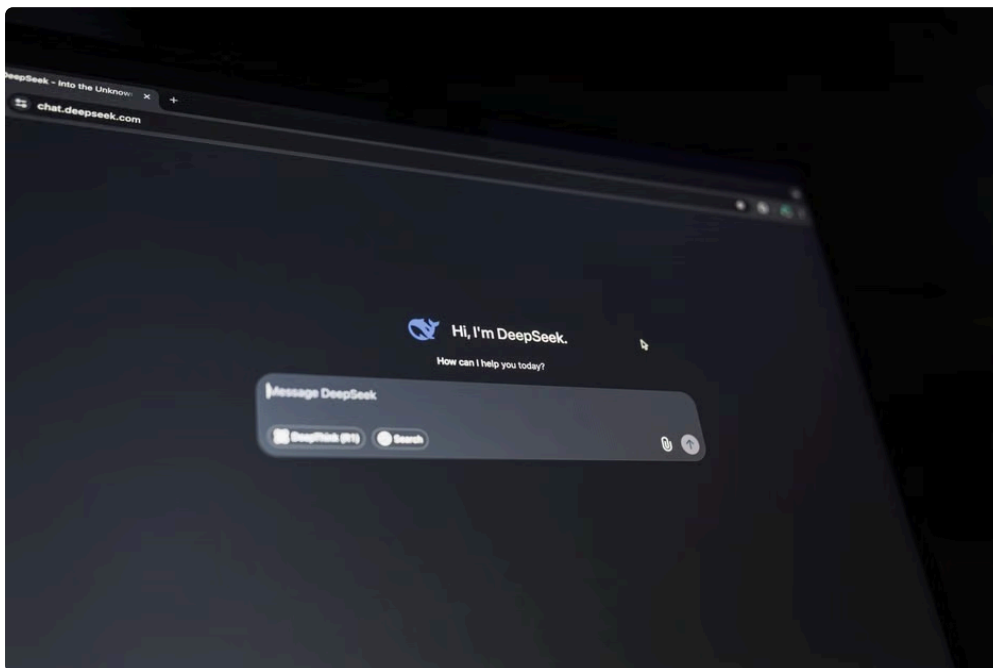
Why Multi-Model Chat Matters: Beyond Tab Switching

Traditionally, users interact with individual AI models like ChatGPT or Claude in separate tabs, bouncing between them manually. This disjointed approach leads to information silos, repetitive work, and cognitive overhead when comparing insights or compounding reasoning.

Suprmind eliminates this friction by enabling *shared-thread multi-model chat* experiences. Instead of fragmented tabs, you maintain one conversation thread where multiple models can be summoned, compared, and orchestrated smoothly through simple @mention commands. This keeps context unified and discussion auditable, a must-have for teams requiring transparent, repeatable outputs.

What Does It Mean to @mention Claude and GPT Only?

In Suprmind's flexible chat environment, @mention AI allows you to direct your prompt to specific models by tagging them inline. If you want to engage only Claude and GPT, the syntax is straightforward:





@Claude @GPT [your question or task here]

This tagged order indicates precisely which models should participate in the [suprmind](#) response generation, excluding others for a focused comparison or collaboration. Crucially, Suprmind preserves this *tagged order*, which controls the flow of reasoning and response priority—essential in workflows leveraging sequential or parallel orchestration modes.

Sequential Mode: Orchestrating Tagged Models for Compounding Reasoning

In **Sequential mode**, Suprmind orchestrates Claude and GPT in the exact tagged order you specify with @mention. Here's why that matters:

- **Logic Flow:** The first model receives your prompt and generates an initial response.
- **Compounding Outputs:** The second model then reads both your prompt and the first model's output to build on or critique it.
- **Amplified Reasoning:** This stepwise interaction mimics a human-like review and enhancement process, producing richer, more robust answers.

For example, entering @Claude @GPT sends your query first to Claude. GPT then sees Claude's reply and can refine or extend it. This sequential layering helps teams achieve nuanced, auditable syntheses of AI reasoning that are difficult to replicate with single-model chats or tab-switching.

Super Mind Mode: Parallel Orchestration With Synthesis and Conflict Mapping

Alternatively, **Super Mind mode** runs Claude and GPT in parallel rather than sequence. Both models respond independently to your prompt simultaneously. Suprmind then synthesizes their outputs, identifying points of agreement and disagreement.

This mode's key features:

- **Conflict Mapping:** Highlights where Claude and GPT diverge in understanding or recommendations.
- **Disagreement Confidence Index (DCI):** A metric Suprmind surfaces to quantify the level of disagreement, enabling users to focus review efforts where uncertainty is highest.
- **Correction Tracking:** Allows teams to record manual corrections or adjudications linked to each model's outputs for future audits.

Parallel comparison with synthesis uncovers hidden insights and potential model biases, giving teams a better foundation for informed decisions.

Shared Thread vs. Tab Switching: Productivity Gains and Auditability

Aspect	Shared-Thread Multi-Model Chat (Suprmind)	Tab Switching Between Models	Workflow	Context
Conversation history	Unified	Fragmented across separate tabs	Single-thread continuity	Lacks shared transcript
Model Coordination	Automatic orchestration with tagged order and explicit modes	Manual copying and pasting with no orchestration	Compounding Reasoning	Sequential builds on model outputs for deeper reasoning
Responses	Independent	Disagreement Visibility	Built-in conflict mapping, DCI, and correction tracking	No built-in disagreement analysis; user reliant
Auditability	Complete exportable conversation with model tags and corrections	Dispersed records, cumbersome to audit		

For teams that hate excessive tab switching and value auditable AI outputs, Suprmind's approach using @mention commands to tag and orchestrate specific AI models is a true game-changer.

Best Practices for @Mentioning Claude and GPT in Suprmind

1. **Use Tagged Order Wisely:** Choose the sequence (e.g., @Claude @GPT or @GPT @Claude) based on which model you want to prime first in sequential workflows.
2. **Pick Your Mode To Match Task:** Use *Sequential mode* for deep compounding reasoning and *Super Mind mode* for parallel perspective comparison and conflict detection.
3. **Leverage Correction Tracking:** Annotate disagreements or errors during synthesis review to build a quality-assured knowledge base over time.
4. **Export Artifacts:** Always export your shared-thread transcripts with model tags and corrections included for transparent review and auditing.

Conclusion: Mastering Multi-Model @Mentions In Suprmind

Suprmind's innovation in shared-thread multi-model chat spaces transforms the way teams leverage powerful LLMs like Claude and ChatGPT together. By learning how to @mention these models only, and understanding the nuances of tagged order, sequential orchestration, and parallel synthesis modes, you unlock workflows that are not only more productive but also auditable and trustworthy.

This aligns perfectly with the needs of teams that require transparent AI collaborations, from strategy advisory to compliance checks. Stop tab switching and start orchestrating your AI models—the next generation of chat is here with Suprmind.