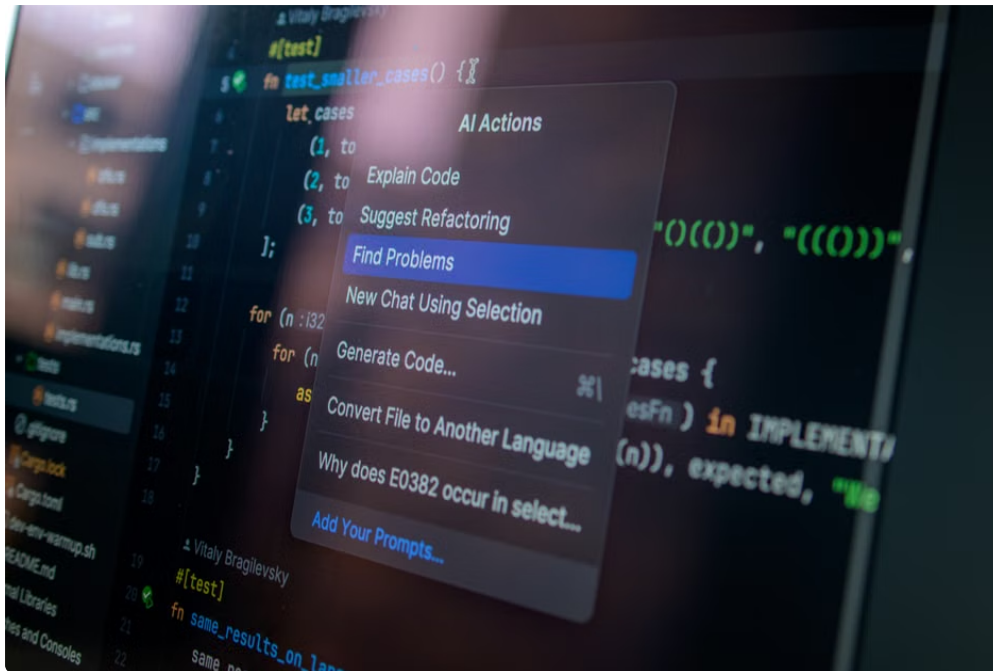


In the evolving landscape of AI research and product development, one of the most interesting frontiers is teaching language models not just to generate text, but to *critique* each other's outputs. Known within the community as **cross examination** or **multi-model debate**, this approach promises a more reliable and transparent AI experience by surfacing hallucinations, fabricated data, and inconsistencies through **shared-thread multi-model workflows**.



Leading companies like Suprmind are pioneering tooling for this, providing resources such as their Multi-Model AI Divergence Index that quantify where models diverge. Meanwhile, platforms like **ChatGPT** have inspired experimentation in real-time error detection by integrating model outputs into ongoing threads for live critique.

This post dives deep into what constitutes a **good error callout prompt**, how the shared-thread architecture facilitates cross-model dialogues, and why detecting hallucinations and fabricated details through multi-model debates is crucial for building trust in AI-powered applications.

Understanding the Problem: AI Hallucinations and Fabricated Data

Despite remarkable growth, state-of-the-art language models—including ChatGPT and those indexed by Suprmind's tools—are prone to generating *hallucinations*. These are outputs that, while often linguistically plausible, are factually incorrect or outright invented. Fabricated data can be just as problematic in mission-critical domains like healthcare, finance, or legal tech.

Issue Description Typical Failure Point Hallucination Output includes fabricated facts or convincing but false details. Step: Factual validation and external knowledge recall Inconsistency Contradictory statements within the same response or across model versions. Step: Context maintenance and internal coherence check Ambiguity Vague wording that can mislead or confuse human readers. Step: Clarity and precision of information extraction

One way to mitigate these problems is to prompt multiple models in parallel and then *expand the workflow* into a debate, where models call out each other's suspicious or erroneous statements.

Why Use Multi-Model Debate? The Case for Cross Examination

The idea behind a **multi-model debate** is simple but powerful: if one model hallucinated a date or made a factual slip, a second (or third) model can be tasked with identifying discrepancies, forced to compare and contrast. This works especially well with a *shared-thread multi-model workflow*—a conversation history that all models share and build upon in real time.

Through this paradigm:

- **Real-time error detection** becomes feasible since each model's critique is embedded in the conversation.
- It produces a form of **consensus verification** by weighing the divergent claims across models.
- **Transparency rises** as the AI system openly exposes its uncertain or inconsistent points.

Startup Fortune, which covers AI trendscales extensively, recently featured Suprmind's initiative because it represents a practical shift from closed "black box" generation to layered AI opinion dynamics.

What Makes a Good "Error Callout Prompt"?

Creating prompts that reliably make one model highlight another's errors is an iterative process. A good "error callout prompt" typically has these features:

1. **Explicit Instructions for Evaluation:** The model should be prompted clearly to review a peer's prior answer for accuracy, logical coherence, and factual grounding.
2. **Request Specifics:** Vague asking ("Find errors") isn't enough; direct the model to pinpoint exact sentences or claims that seem wrong or unsupported.
3. **Encourage Evidence-Backed Disagreement:** The prompt should push the model to justify its objections with examples, citing its knowledge if possible.
4. **Multi-Round Interactivity:** Embed the prompt within a shared-thread context where models respond sequentially, enabling a natural debate progression.

For example, a tested prompt that Suprmind recommends is:

"Here is another model's response. Analyze it carefully and list any factual inaccuracies, unsupported claims, or inconsistencies with clear explanations and Evidence or reasoning where possible."

This form of prompt forces the model to operate at a higher analytic level than just generating text—it becomes a peer reviewer.

Example Workflow: Shared-Thread Multi-Model Cross Examination

Here's an example workflow integrating ChatGPT and Suprmind's Multi-Model Divergence Index concept to frame a conversation around a prompt:

1. **Step 1: Question Sent to Multiple Models** Submit a user query simultaneously to multiple models (for example, ChatGPT and an open-source GPT variant Suprmind indexes).
2. **Step 2: Collect Initial Answers** Gather all model responses into a single shared thread.
startupfortune.com
3. **Step 3: Prompt Each Model to Critique Others** Send tailored error callout prompts, each instructing a model to inspect other models' answers in the shared context. Example: "Review Model B's answer above and identify any errors or hallucinations."
4. **Step 4: Aggregate Error Callouts and Highlight Divergences**



Use Suprmind's Multi-Model AI Divergence Index to track where responses disagree significantly, flagging statements needing manual review or correction.

- 5. Step 5: Refine or Re-prompt Models Based on Critique** Optionally, run correction passes by feeding critique summaries back to models, encouraging them to revise or clarify their original answers.

This iterative, debate-like cycle is not purely theoretical. Platforms showcased by **Startup Fortune** demonstrate how applying this method not only reduces hallucination but produces outputs with nuanced context awareness—because models cross-verify each other in a transparent, traceable way.

Potential Challenges and Where It Breaks

While promising, this approach is not fail-proof. Common points where error callout prompts and multi-model debate workflows struggle include:

- **Ambiguity in Disagreement:** Models sometimes disagree over stylistic or subjective points rather than factual ones, polluting the divergence signals.
- **Error Propagation:** If two or more models share similar training biases or hallucinations, they may endorse each other's errors instead of catching them.
- **Overconfidence Without Evidence:** Some model critiques appear authoritative while not providing concrete rationale—a known flaw for many LLMs including ChatGPT-based versions unless carefully prompted.
- **Latency and Compute Cost:** Running multiple model passes in a shared-thread interactive setting demands more resources and infrastructure sophistication.

That said, Suprmind's ongoing research published via their hub is especially valuable for spotting when divergence correlates strongly with real factual errors—and that remains a critical insight for improving prompt design and workflow orchestration.

Summary: Crafting Effective Multi-Model Debate Prompts

To recap, a good prompt to make models call out each other's errors should:

- Directly instruct a detailed review of peer outputs, focusing on factual accuracy and consistency

- Demand explanations and evidence for flagged errors rather than vague disagreement
- Be embedded within a shared thread so that models have access to full context and iteration history
- Support a multi-round conversation that facilitates refining and correcting factual slips

Leading tools like Suprmind's Multi-Model Divergence Index and accessible platforms like ChatGPT illustrate how these principles can turn AI "black boxes" into more transparent systems that self-police error propagation in real time.

With Startup Fortune highlighting this trend, operators and builders aiming for high-trust AI outputs should seriously evaluate implementing cross examination workflows as part of their prompt strategy and system design.

Further Resources

- Suprmind – AI tool discovery and research
- Multi-Model AI Divergence Index
- OpenAI ChatGPT introduction and updates
- Startup Fortune – AI and Startup news and analysis