



WHAT A GENERAL MODEL CAN'T TELL YOU

Why Lou answers differently than ChatGPT, Grok, or DeepSeek, and where it stops on purpose.

By Joe Worden

IDEA IN BRIEF

THE QUESTION

Why is Lou different from ChatGPT, Grok, or DeepSeek? A general model answers marketing questions with competent averages: clear, reasonable, identical for every firm, with no source to point to when pressed.

THE ANSWER

Lou is a general model inside a harness we built: our doctrine, our corpus, our retrieval, our rules for what gets said and in what order. Nothing was retrained. Every answer traces to a published entry a stranger can open.

WHERE LOU STOPS

Lou diagnoses what's wrong with how a firm earns trust, and says plainly what it can't know without research: your specific market, buyers, and positioning. The boundary is the point.

The most honest question anyone asks Lou comes in some form of the same sentence: why is this different from ChatGPT, Grok, or DeepSeek?

Fair question, and one that deserves a straight answer. The general models are really good, and starting this page anywhere else would be the wrong way to open a page about trust.

Query a general model on how a wealth management firm earns trust with skeptical prospects and you'll get a decent response. Clear. Well formed. Reasonable. It will also fit any firm, any market, and any year. The model has read the internet, which means its answer is an average of everyone who ever wrote about the subject, without a point of view and without a source to point to if you press it on where a claim came from.

That's the first difference: Lou can point.

"B2B buyers are progressing through critical buying tasks in more autonomous ways, and sellers can't rely on static collateral to carry influence in those moments."

Alyssa Cruz, Senior Principal Analyst, Gartner Sales Practice, March 2026

01 What Lou works from

Lou is a marketing associate built on a specific published body of work: The Decision Sequence, the research library, and the decision science beneath them all. The thinking covers one kind of business, the expert-led firm that sells earned trust: accounting, law, wealth management, capital, insurance, architecture, consulting, nonprofit development. If Lou gives you a read, the logic traces to published entries you can open and evaluate. Ask for something to read on the subject and Lou hands you the entry behind the answer. Stay on a subject long enough and Lou will offer one unprompted, once, with the link appearing only after you say yes to it.

The research is proprietary only in the sense that nobody else wrote it. Access is open, and I published it that way on purpose. A firm that sells trust should be easy to check, and so should the associate who speaks for it.

The architecture matters. A general model is a mile wide and an inch deep. Lou was built an inch wide, expert-led B2B marketing and engagement, and a mile deep, running on a published corpus, the decision science beneath it, and thirty years of being hired instead of the incumbents.

02 Nothing was retrained

The engineering deserves the same straight talk the question got, because the vocabulary in this category is loose, and I dislike sloppy vocabulary even more than sloppy marketing.

Nothing was retrained. Lou is a general model held in a harness we built: the doctrine, the corpus, the retrieval, the routing, and the rules governing what gets said and in what sequence.

The intelligence is general. The judgment, the corpus, and the sequence are ours.

The harness constrains Lou to one belief: earn trust by aligning every interaction with how people make complex decisions.

That distinction matters, and a buyer can feel it as consequences. A model trained on proprietary material absorbs it into weights it can't cite, and the knowledge goes rotten the day training ends. A harness reads the library live. Tomorrow's new entry becomes tomorrow's reasoning, and every sentence stays traceable to a page a stranger can open.

The harness does one more thing: it enforces the philosophy, which training never could. Lou believes value before the ask, and the wiring reflects that: no link leaves the conversation until you've said yes to one. The sequence runs as a constraint ahead of whatever Lou remembers about it. Trust-sequenced by construction.

03 The conversation is different too

A general model answers the question you asked. Lou spends the first act of every conversation determining whether that's the right question.

People come in saying they need more content. Sometimes they do, and Lou will tell you that. More often the content is just fine, but the order is wrong: the firm is asking for commitment before it has earned belief, and no publishing calendar can fix an ordering problem. A general model will joyously produce the content calendar. Lou will tell you when the calendar is the wrong purchase.

04 Where Lou stops

Lou can tell you what's wrong with how your marketing earns trust, and the pattern behind it. Lou can't tell you what's true about your specific market, buyers, or positioning, because that research hasn't been done yet. When a question crosses that line, Lou tells you instead of guessing past it.

A general model has no such boundary, because it never runs out of plausible sentences. The boundary is the point. An advisor who tells you what it doesn't know is the only kind worth asking about what it does.

05 What Lou keeps

During the beta, conversations are retained so we can improve how Lou serves visitors. Nothing you share trains any model, and nothing you share enters the library. The policy sits on this page because the firms Lou serves carry real compliance obligations, and a plain answer beats a reassuring one.

06 How to test this

Don't take the page's word for it. I ran this comparison myself, more than once, while building Lou, and it's why this page is here. Ask Lou the question your firm is wrestling with right now, then ask the same question of any general model, and evaluate two things: whether the answer could be anyone else's, and whether anything in it can be checked.

Here's what the test looks like in practice. Ask both about a referral slowdown. The general model returns a numbered list of tactics, reasonable and anonymous. Lou asks whether referrals slowed because people aren't thinking of you, or because they're no longer sure how to explain who you help. One is a handout. The other is a diagnosis starting.

One answer will be agreeable. The other will have a spine.

PRIMARY SOURCES

Alyssa Cruz, Senior Principal Analyst, Gartner Sales Practice, March 2026. Quoted from Gartner research on autonomous B2B buying behavior; also cited in Proximity Is No Longer Measured in Miles.

FOUNDATIONS

The behavioral research underneath this argument is documented in the Decision Architecture Research Library, the corpus Lou reasons from.