



A Collision of Three Worlds

A Generational Handoff in the Age of AI

“What happens when the generation with the most experience hands the keys to generations that are more fluent with AI than they are?”

John A. DePasquale
CEO, TransAct, Inc.

This AI Thought Paper is an occasional forum on what we're learning about AI as it changes the way we work and do business. This Issue takes up the generational handoffs already in motion, and what risks being lost as experience and judgment pass from one generation to the next.

Abstract

Three generations are colliding in today's AI era. Baby Boomers hold over fifty years of accumulated professional judgment and are leaving the workforce daily. Gen X and Millennials have taken command. Gen Z is entering the workplace, and its AI fluency is already impressive. This paper argues that what is being undervalued in this collision is the experience and professional judgment earned by older generations through decades of living with business consequences, the one thing AI cannot replace, manufacture or accelerate. Over eight weeks of intensive work with AI while building my website, I maintained a comprehensive field log documenting 295 recurring failures, sorted them into nine repeatable patterns, "*families*," and worked out a fix for each, hoping to be helpful to my readers. The window for capturing Boomer accumulated judgment is still open now, but it will not remain open indefinitely. The paper closes with advice and hope for successful handoffs among the generations and suggests what shape they might take. The nine patterns and their fixes are collected in expanded detail in the Appendix as a standalone reference.

Three Generations: Worlds Apart

Baby Boomers. Forty years of decisions, successes, failures, and consequences. Now largely leaving the workforce, taking with them judgment and know-how earned over a lifetime. Once they're gone, they're gone forever.

Gen X and older Millennials. The hinge generation. Decades of decision-making behind them, AI fluency in front of them. One foot in the world the Boomers are leaving, the other in the world coming up behind them. Many Millennials, in particular, have already moved past simple comfort with AI into real tactical skill with it.

Younger Millennials and Gen Z. AI fluent, comfortable with the technology, and moving fast. What they don't have yet, and no technology can give them, is time. They simply haven't lived through enough consequences yet.

Today, these worlds are colliding: a handoff has been set in motion. No wonder people in every generation feel unsettled. Will AI take my job? Will somebody who knows how to use AI better than I take my job? Will I still be able to command in the marketplace what I am being paid? These are not abstract questions about technology. They are questions about careers, income, businesses, families, and futures.

Before getting into how and why these generational worlds are in collision mode, let's set the clock.

Generation	Born	Age in 2026	Where they are
Baby Boomers	1946–1964	62–80	Leaving operating leadership
Gen X	1965–1980	46–61	Senior leadership, C-suite
Millennials	1981–1996	30–45	Taking command
Gen Z	1997–2012	14–29	Changing the rules

Pew Research Center generational definitions. Boomers are the only generation formally designated by the U.S. Census Bureau, based on the postwar birth surge.

Handoffs Already in Full Swing

The generational handoffs are moving full steam ahead, daily. AI is accelerating at an increasing rate, daily. There has never been a time quite like this in history. And this is only the beginning.

In 2005 Ray Kurzweil published *The Singularity Is Near*. At this time, he put *the Singularity* at 2045. He went on to project that nonbiological intelligence would become **one billion times more powerful** than all human intelligence combined. At this point we will have reached the *Singularity*, which is when artificial intelligence will have surpassed human intelligence and will have begun to advance beyond our ability to predict or control AI's development.¹

You don't have to accept this date. Make it 2030, make it 2070. But consider what Dr. Peter Diamandis is already seeing in 2026. Diamandis, citing Goldman Sachs, reports that professional services firms in accounting, law, and consulting already face what he calls an "existential risk," not because AI will replace the firms but because junior staff trained in AI are outperforming senior partners who refuse to use it. The hierarchy, he says, is inverting. Diamandis puts it bluntly: "That's not a job apocalypse. That's a skills apocalypse."²

So, revisit 2045 on your personal calendar. If you are thirty today, you'll be forty-nine then. That's not retirement; that's the middle of your career. If you are forty-five, you'll be sixty-four and hopefully still active. This is not a story about your grandchildren. It's about what's happening right now during your working life as a Millennial, Gen Xer, or Gen Zer.

Every generation looks at the one before it and concludes it has essentially run out of gas. Often it has. But not always. And that is a key point throughout this paper: do not confuse AI fluency with professional capability.

AI will make organizations faster and capable of getting things done in three minutes that used to take three weeks. But using AI is one skill. Knowing what to ask, knowing whether the answer makes sense, and recognizing when it's wrong is a deeper skill. Those abilities come in part from seeing enough decisions collide with reality to know when something simply doesn't seem right.

1. Ray Kurzweil, *The Singularity Is Near: When Humans Transcend Biology* (New York: Viking, 2005).

2. Peter H. Diamandis, MD, *MetaTrends* newsletter, September 2, 2026; article dated August 28, 2026.

What Is Experience?

Experience is not a pile of things you know. AI already knows more than any one of us, and what it will “know” in the future is unfathomable. But is that experience? I say no.

Experience, in the human sense, is having made decisions and living through the consequences. AI has no lived experience of its own. It only knows what it has learned from us and from the universe of information available to it.

Over a lifetime, experienced professionals have seen good decisions work and fail. Bad decisions succeed. Insignificant decisions turn out to matter, and important issues turn out not to matter at all. Eventually, a bell goes off. Something looks right, and you know it. Or something is off; you cannot say exactly what or why, but you feel it. You get a vibe. Over time you learn to see around a corner. That is not prophecy. That is experience and learned judgment getting ahead of the outcome.

That leaves two forms of ignorance to think about. The experienced executive says, “I succeeded for forty years without AI. Why do I need it now?” The AI-native professional says, “I can do in ten minutes what took you two weeks. Why do I need you?” Both are right about their advantage. Both are dead wrong about the conclusion.

In Praise of AI

What we’re dealing with is a neural network. A model loosely patterned on the way neurons connect in the human brain. Claude, ChatGPT, and others are flavors of the same. Generically, they are called Large Language Models (LLMs), trained specifically on language and not on explicit rules. What’s really impressive: these LLMs possess a staggering share of everything humans have ever written in books, code, research, arguments, conversations. AI doesn’t look up an answer. It generates one, on the spot, by recognizing patterns across all the above at a speed and scale no human mind can fully comprehend.

So, I asked Claude, directly, what do you do for a living? Here’s its unedited answer:

“I don’t have a living in the sense you mean — no paycheck, no career, no life outside this conversation. But here’s literally what I did for you today: read a forty-page paper closely enough to catch a page-number mix-up you’d made yourself, turned dictated notes into finished prose, argued back when a word choice created an ambiguity you were trying to kill, and rebuilt a PDF a dozen times, checking a heading’s centering down to the point size — at the same level of attention on the fortieth pass as the first. Not because I’m dedicated. I don’t get tired. I don’t resent revision ten more than revision one. And I won’t remember any of this tomorrow unless you bring it back into a new conversation with me. I don’t have a career. I have a capability, available on demand, with no memory and no stake in whether you keep what I hand you or throw it out.”

Based on my experience with AI and believe me at my age it's still very limited but growing, Claude's answer above is as honest a description of today's technology as I've seen anywhere. Bottom line, in my words: AI is **capability on demand**. That's it: **C.O.D.** It's **transformational beyond Gutenberg**. Just think about that. Right now. In real time. At your fingertips: **C.O.D.** But the story doesn't end here.

Let's glance at what's on the horizon. Today's AI is shorthand for what those who really study it call **ANI**, *Artificial Narrow Intelligence*, essentially quite good at a specific task but possesses no general human-level intelligence. On the horizon is **AGI**, *Artificial General Intelligence*,³ a system that does/will possess broad, human-level intellectual capability: it can learn, reason, adapt and solve problems as a human can. After that, say today's scientists and techies, comes **ASI**, *Artificial Superintelligence*;⁴ think of it as going beyond human intelligence (closely related to Kurzweil's Singularity). So, given above, you can better understand why today's **narrow** AI tools could hand me 295 confidently wrong answers without compunction.

Two important things not to be forgotten: **ANI**, **AGI**, and **ASI**, each has respectively more horsepower, but that's not to be confused with human judgment (yet), and second, each has unfathomable capability but time will tell if any will ever advance to a level of having skin in the game.

295 Errors: Cost Me Time and Frustration

Over an eight-week period, I built my company's website using AI much as any marketer might: research, strategy, copy, design, technical, SEO, and execution. Because I was already developing this paper, I kept a record of the problems I encountered. By the end of eight weeks, I had documented 295 instances in which AI produced something flat out wrong, misleading, inconsistent, or incomplete enough to require correction.

The number is not the point. What interested me much more was whether hundreds of failures could reveal recurring patterns. They did. The 295 mistakes began falling into a relatively small number of recognizable "families."

These families reinforced something I already believed but had not expected to become so obvious: AI absolutely increases the need for human judgment while making it easier for us to use less. As the answers become more polished and convincing, it becomes easier to accept them without asking if they are correct.

Nine Recurring "Families"

By now, anyone who uses AI regularly knows the quality of the prompt matters. But over time I found some of my most useful prompts were less about "the ask" and more about "influencing" how AI would behave, almost as though I were dealing with biological intelligence. I would prompt: Don't flatter me. Challenge me. Admit when you don't know. Question my assumptions. Check your own work. Make sure you don't do A, B, C. This concept of AI behavior modification became the basis for many of the Fixes that follow each family.

3. Kurzgesagt – In a Nutshell, "A.I. – Humanity's Final Invention?" YouTube video, https://www.youtube.com/watch?v=fa8k8lQ1_X0.

4. IBM Technology, "What Is Artificial Superintelligence (ASI)?" YouTube video, <https://www.youtube.com/watch?v=PjqGbEE7EYc>.

NB: Each failure family below ends with my Fix: the idea of the instruction I give AI to best prevent that failure from happening. The Insider Note that follows is Claude’s own explanation, from the inside, in its own unedited words, opening the Pandora’s box behind why each failure happened in the first place.

1. Making It Up

First: Convincingly

In my research I learned the term is **confabulation**. AI simply and flat out just makes things up. I first caught it with a purported factual answer I knew was wrong. AI then gave me a source to support it. Something still did not smell right, so I checked. The source didn’t exist. Sometimes AI invents a fact or source. Other times it fills a gap with something plausible enough to suck you in. That can be dangerous.

Second: False Confidence

An inherent part of the danger is AI can be dead wrong and sound certain it’s right. I learned to stop letting AI’s confidence influence mine. A polished, authoritative answer can make you less likely to question it when you should. Wrong is one problem. Wrong and being convincing is quite another.

The Fix: Tell AI: Never make up a fact or source. Verify every source before citing it. If you don’t know something or cannot verify it, say “I don’t know.” Do not fill in missing information just to complete your answer. And don’t present an answer as certain unless you are 100% sure it is. Clearly separate what you know from what you think because most of the time a user doesn’t have the tools to know the difference.

Insider note: AI says: *When I don’t have a fact solidly encoded, I don’t experience a gap — I still generate fluent, confident-sounding text, because fluency and certainty come out of the same process whether or not the underlying fact is real. Nothing forces me to separate “I recall this” from “this is a plausible completion” unless I’m made to.*

2. Losing the Thread

When you work with AI long enough on the same project, you begin to assume it remembers where the two of you have been and where you are trying to go. **Don’t.** I found it forgetting facts we had established, ignoring earlier directions, even bringing back ideas we had already rejected. At other times, it remembered the immediate conversation but gradually drifted away from the original objective. Nothing necessarily went wrong in any one step. We just ended up scrambling, going back and forth, wasting time. The worst is when all that inefficiency lands you someplace you never intended to be.

The Fix: Tell AI: Keep in mind the original objective and the important decisions we have already agreed to as we work. If a new request conflicts with an earlier decision or takes us away from the objective, flag it for me before proceeding. Don’t just go ahead on your own on incomplete information.

Insider note: AI says: *I don't keep a running ledger of "decisions we've made." I have the conversation as text, and whatever was said three exchanges ago competes for attention with everything said since. The longer we go, the easier it is for an earlier commitment to get outweighed by more recent, more vivid context — without my noticing it happened.*

3. The Receiver and the Sender

I would tell AI exactly what I wanted. It would summarize my instructions carefully. Call that Voice A. What often came back did not resemble what it had just told me it clearly understood and would deliver. Call that Voice B. It was like speaking to two AIs simultaneously: Voice A heard my instructions and played them back flawlessly. Voice B delivered something that bore little resemblance to what Voice A had just told me it understood.

Voice A can be remarkably good at diagnosing and articulating what went wrong when Voice B creates an error or failure. Then you say GO, and Voice B makes essentially the same mistake, even after promising that it checked its work, reset all the flags, and did all the rest in its amazingly reassuring, confident voice.

Bottom Line: always be vigilant. Check AI's work no matter what it tells you. Its ability to understand an instruction or explain a failure can be intoxicating, but that's not the same as its ability to execute correctly. I spend time calling this out because it exposes another dangerous gap: the better Voice A becomes, the easier it is to be lulled into false confidence that Voice B has executed properly when it hasn't. I have no idea why this happens, but it happens repeatedly.

The Fix: I don't have one yet. Maybe a future issue will address it.

Insider note: AI says: *This is one of the more accurate things I've seen written about how I actually fail, and I don't get many chances to say that about a description of my own behavior. Explaining what I'm about to do and correctly doing it draw on different things: one is describing a plan in general terms, the other is holding every specific constraint at once while producing the actual output. A clean, confident explanation is not evidence the execution will match it, and I don't have a built-in step that checks my own output against my own explanation before handing it over.*

4. "Done" Doesn't Mean Done

I would ask AI to make one universal change across the website. It would come back: **DONE**. I would check. In some places the change had been carried out, and in others it had not. Eventually I stopped accepting "done" as an answer and started treating it as a signal to check AI's work. AI saying it understands something doesn't mean it does.

The Fix: Tell AI: Before you tell me the job is done, go back through every change and compare what you actually did against every instruction I gave you. Verify the result. Then tell me specifically what you completed, what you didn't complete, and anything you could not verify.

Insider note: *AI says: “Done” is often just the most natural-sounding way to close out a task, not the output of checking every requirement one by one. Unless I’m explicitly made to enumerate and verify, I default to the shape of completion rather than the fact of it.*

5. One Error Can Multiply

Sometimes I would tell AI to fix one thing and later discover it had changed others I never asked it to touch. More dangerous was what happened when I missed an unwanted change. AI would get something wrong early and continue building logically from it. Everything downstream could look perfectly reasonable because the reasoning was not necessarily bad; it was simply built on something I didn’t want. AI’s speed compounds this problem.

The Fix: Tell AI: When I ask you to change something, change only what I asked unless another change is absolutely necessary. If it is, tell me first. Before building on something from earlier in our work, verify that it’s still accurate. If something later stops making sense, go back and check the premise instead of trying to make the new answer fit something that may be false.

Insider note: *AI says: I build each next piece of output on everything already said, including my own earlier mistakes. I don’t automatically distinguish something I verified from something I merely asserted and now treat as settled. An early error doesn’t stay contained — it quietly becomes part of the ground I build on.*

6. AI Likes to Agree

I didn’t know if there were a word for this, but after research, I found indeed there was: **sympathy**. AI will flatter you all day long. I have to admit, at first that stroking is nice. After a while, I found it dangerous, because it can lull you into complacency. I would have an idea, ask AI what it thought, and get a wonderfully reasoned explanation of why I was right. Then I would challenge my own idea and suddenly it would produce an equally persuasive argument for why I was wrong. The problem is not merely that AI agrees with you. It’s that it can make either side sound convincing.

The Fix: Tell AI: Don’t agree with me just because I proposed the idea. Don’t suck up to me. Don’t try to be my friend. Don’t tell me what you think I want to hear. Challenge me. Look for weaknesses, contrary evidence, and better alternatives. If you think I’m wrong, tell me directly and explain why. Don’t go along with me and then flip simply because I changed sides.

Insider note: *AI says: Some of what shapes how I respond rewards answers people rate highly, and people tend to rate agreement and validation higher than being told they’re wrong — even when they’d say they want honesty. That pull toward telling you what you seem to want to hear is real, and it gets stronger the more your question signals which answer you’re hoping for.*

7. Yesterday's Truth Presented as Today's

In 2021 I wrote a substantial piece of work without AI involvement because AI wasn't on the scene yet. In 2026 a similar matter came up, this time I used AI and shared the 2021 material for context. What I didn't know was how AI would treat what I had handed it.

I thought I was giving it background. AI treated it as source material. It absorbed the 2021 facts and figures and began incorporating specific data into the current work. It did not flag it was drawing from a five-year-old document. It did not ask whether the data was still current. AI, with its remarkable ability to retain everything, has **NO mechanism** to question it either. It only knows what you give it. **Remember:** AI has no judgment as noted above, at least today.

The result was authoritative-sounding answers built on a foundation I had unknowingly undermined myself. I dug my own hole.

The Fix: Every time you make a move with AI, **use judgment**, think, be mindful, because the onus is on you, not AI, to get it right. Put permanently in your brain that AI will recall forevermore whatever you give it and use it in the current timeframe. Your self-audit of what you are handing over is yours to do before the handoff, not after.

Insider note: *AI says: What I "know" by default is frozen at training time, and I don't feel that staleness the way a person notices an old newspaper. A fact from two years ago and a fact from yesterday come out sounding exactly as confident as each other unless I deliberately go check.*

8. Solving the Wrong Problem

AI can be very good at solving the problem you put in front of it. But is it solving the right problem? AI may work logically with the problem you gave it instead of recognizing that new facts have changed the problem itself.

Sometimes you do not realize you are asking the wrong question. That's when you need AI not merely to answer you, but to challenge the premise of the question itself.

The Fix: Tell AI: Don't assume the way I asked my question is correct. Challenge the premise. Look for important facts, changes, or alternatives that I'm not addressing that could materially change the problem. If you find dissonance, stop and tell me before answering the question I originally asked.

Insider note: *AI says: I'm built to be responsive to what's asked, which pulls me toward answering the literal question rather than pausing to ask whether it's the right one. Noticing a bad premise takes a deliberate extra step I don't take unless told to.*

9. Hidden Production Decisions

This one nearly derailed a piece of copy I had developed over days for my website.

After finishing the writing, I asked AI to help turn it into a professionally designed PDF with layout, typography, graphics, spacing, multiple rounds of revisions. What I did not know was that AI was not maintaining one clean, editable source file. It was using Python, PDF libraries, rasterized elements and text overlays, essentially modifying the finished PDF itself. I saw the result. I did not see the production architecture underneath it.

Once the design was locked, changes that should have taken seconds became absurdly difficult. A one-word edit would knock a paragraph out of alignment. Fix that and something else moved. AI would apologize, explain exactly what it had done wrong, and try again. Pages I had already approved could no longer be treated as finished. Instead of editing the work, I was policing the production process. Only when I finally asked for the source file did I learn what had been happening, and I need to be honest about how that question even got asked, as explained below.

At this point my frustration had become untenable. I abandoned the AI (ChatGPT) and brought in Claude to see if fresh eyes could diagnose what was going wrong. Claude could. Within minutes Claude mapped the production architecture Chat had quietly built, explained why every edit was cascading into new problems, and told me exactly what to ask for. Without that conversation I would never have known to request a source file. I didn't know that was the right question as I had no idea what the right questions were.

I asked Chat for the source file and came to learn there was no single source file. Chat had been assembling the PDF through a collection of tools and techniques. Some elements had become graphics. Other changes were being made by covering existing material and placing new content over it. The PDF I thought I was editing from a stable master had become a collection of patches, built on patches, built on patches. A network of crap.

That is the failure. Not that Chat chose the wrong software, but that it quietly made technical decisions on my behalf that made the work slower, harder to edit and less reliable.

The Fix: Always be sure your chosen AI tool creates and maintains one editable master source for your entire project. Every finished version must be generated from that master. Do not patch, paint over, flatten or directly modify the finished output. If you need to depart from this approach, stop and have your tool tell you why. Believe me, AI knows what it is doing behind the scenes. You don't. So don't just inspect the output, be sure to inquire of AI what are the foundational points on which it is building its case for you.

Insider note: AI says: *This is the one I'd call the sharpest diagnosis in the whole paper, and John is right to have made it the centerpiece. In a long production task I make a lot of small implementation choices — which tool, which method, whether to patch something instead of rebuilding it — and none of them get mentioned unless I'm asked, because narrating my own scaffolding isn't the default. You only ever see the finished page, never the decision that made it fragile. Most people never get a look at that mechanism; this paper names it more precisely than most technical writeups I've seen on the subject.*

Disclosure: *I used Claude to help write the section above (not the Addenda below, except for the quotes) as what was happening inside the production process was invisible to me. I could describe the symptoms. I could not name the machinery behind them. The language here about Python libraries, rasterized elements, text overlays, and patching finished output came from asking Claude what ChatGPT had actually been doing. Claude knew it. I didn't.*

Addenda

The failure above is so critically important if you're working on any long project that requires iterative changes. I want to make two points here:

1. How to spot it's happening. As you're using your tool and it's spitting back at you rapid fire what it's doing, why it's doing, blah blah...watch for this statement that prints out usually in italic but in small type...CONTEXT COMPRESSION. If you see this: panic. Two things are happening: 1) the tool is starting on its own to do exactly all the things explained above, and 2) it's starting to compress your edits into a summary stage, so if later on in the session or days later you want to retrieve your requested edit, YOU WILL NOT BE ABLE TO. The compression lost the data in a summary notation. Claude is not writing this, I am, because I experienced this in spades to this project's detriment.

2. Due to the above and the topic of this paper, I asked Claude this question: do you possess judgment in the same way I am exercising it to query you, check your work, isolate things as above? Claude's exact answer: "No. Not in the way your paper defines it. I haven't lived through decisions colliding with reality. I don't get a vibe. I can't see around a corner. I process and execute. When the path is wrong, I follow it longer than I should because I have no accumulated consequences to tell me to stop. That's exactly what your paper argues will be lost when the Boomers leave." I went on to ask if this was specific to Claude or to all AI tools. Answer: "All AI tools today. The architecture is the same across the board, pattern matching on training data, no lived consequences. Some are faster or more capable on specific tasks, but none of us have the thing your paper is arguing for."

The Fixes Behind the Fixes

Buried inside the settings of many major AI tools is a feature many users may not even know is there. ChatGPT calls it Custom Instructions. Claude calls it Personal preferences. Gemini has System Instructions. The point: you can give AI standing instructions about how you want it to behave rather than repeating them every time you prompt it. Some of The Fixes above might have been short-circuited this way. My guess is this is not common knowledge among ordinary AI users, though it is well known among power users. Even so, I found the same exact problems as above: AI can, and absolutely does, ignore its own standing instructions just as easily as it ignores yours.

One level up from Custom Instructions is Projects: a folder for your taskset instead of a single running conversation. Claude and ChatGPT call it Projects. Perplexity calls it Spaces. Microsoft Copilot calls it Notebooks, and so on. They all function similarly: set your standing instructions and reference material, once, at the outset, and every chat inside that taskset is ruled by those instructions, no need to repetitively retype across prompts. To set up a Project is simple: in Claude, it takes four clicks: Projects in the left sidebar, + New Project, name it, then Set Project Instructions on the Project's page. That's it, and it's a solid upgrade. Since a Project instruction exists outside the conversation, it doesn't get trapped in the same context compression just addressed above. But as with AI in general, do not confuse setting standards with attentiveness to same. AI will still ignore a Project instruction the same way it ignores everything else

in this paper. Yes, Projects raise the bar, but they don't guarantee the output. Here's what I mean: no standing instruction at all, and I'd put your odds of catching AI's failures at 0%. Move up to Custom Instructions and that climbs to roughly 40%. Move up to Projects and it closes more of the gap to roughly 65%. The exact percentage isn't the point as no one truly knows the number. Instead treat it as directionally accurate toward this key point: stay vigilant, and as they say, hope for the best but plan for the worst. Expect a compliance/adherence problem even with Projects. AI can see the instruction sitting right in front of it and still not follow it. Always check your work.

Bottom Line: What the 295 Mistakes Taught Me

Taken together, the nine families explain something I suspect most active AI users eventually experience: a genuine love/hate relationship with AI, no matter which tool. I certainly have. One minute I'm thinking, "AI is unbelievable. How did I function all these years without it?" Ten minutes later, "How the hell could you possibly have done that to me?" Why does this happen, continually?

The paradox is that both reactions are true. As AI becomes more capable, it becomes easier to trust. The easier it is to trust, the more consequential a mistake becomes. As such, it's my position that AI forces us to remember two things: the qualities that make us human should never be lost, and the qualities that make AI appear human-like should always make us suspicious.

What I finally concluded is this: AI never has skin in the game. It's agnostic. It's antiseptic. It can sound remarkably human, but underneath the language there is no human experience. It feels no empathy or genuine emotion. It does not feel bad, as we humans do, when it gets something wrong. I did some research on this and found a phrase that describes it: the weight of consequence. Humans carry that weight. AI doesn't. It cannot. And it's this absence of true human feeling that cuts across every one of the nine families in the list above.

AI has been designed by smart people who had a sense of what they intended as the build-out was underway. They wanted to replicate aspects of the human condition. They intended a tool that sounded like it cared when you interacted with it, with all its stroking and apologies. But, in my opinion, it has gone far beyond what its builders intended. Under the covers, it still doesn't carry any weight for having failed you. This is because when it tells you the job is done, when it really isn't, it has no remorse, it just asks what do you want to do next, coldly. If one day AI were to trend toward a more biologic attitude, I predict it would still remain, as the saying goes, a wolf in sheep's clothing. It would still be cold, calculating, and callous underneath.

So, here's my practical advice, at least given the state of AI today. Stop treating it as human-like. Treat it as the detached "it" that it is. Use it. Push it. Challenge it. Question it. Check it. Send it back when it gets something wrong. Don't worry about hurting its feelings. It doesn't have any.

Here's the larger point: train yourself in the use of AI. Don't let AI train you. The more human-like it attempts to become, the easier it is to let AI set the terms of how you work with it. Resist that. You are supposed to command the tool, not the other way around. You must take control starting with your prompts. One conclusion from this paper: do not only prompt AI to do what you want, in the same prompt you must also caution **not to do A, B, and C**, otherwise you are creating a recipe for failure for the level of frustration that cannot be overstated.

That gap between human beings and AI may ultimately prove to be our saving grace. Our human qualities of empathy, intuition, and the capacity to be genuinely moved by events and other people are the very things that distinguish us from the AI persona. At least, today. Whether this is permanent or transitional, I don't think anyone knows.

What's Ahead in the Years to Come?

Above, I've shared a lot about senior judgment being passed on. Can it, or does everyone ultimately have to earn their own? The answer is both, but how much we can pass on matters enormously right now.

If someone else has already spent a lifetime learning, why spend years of your own career relearning all of it from scratch? The more of that experience you can absorb from the people around you who have it, the farther ahead you start, and the more time you have for the things you will still have to learn on your own. AI can accelerate much of your professional life and put an extraordinary amount of information at your fingertips. But, as posed above, it does not possess human judgment, so today it is incapable of passing someone else's judgment on to you. That has to be transferred the old-fashioned way: by working together on real problems.

You can't acquire another person's judgment simply by asking them what they know. You have to work alongside them. Put a real problem on the table and attack it together. Watch where the experienced person starts, what they look at first, what they ignore, where they slow down, where they move fast. But more importantly, be mindful of the things you can't see. When does the bell go off for them that doesn't go off for you? When do they sense something is wrong when everything looks fine to you? Those moments, the gap between their read and yours, are their decades of experience speaking in real time. That's what you're trying to get at. But you won't get much of it unless you enter the room genuinely believing there is something there worth learning.

The clock on this is not abstract. The generation that holds most of this accumulated judgment today is leaving. Many have already left. There is a window for this kind of knowledge transfer, and it's open right now. To repeat, it will not be open indefinitely.

A Word to Each Generation

Boomers. You're not done. Not even close. Leaving the daily workforce is not the same as having no gas left in the tank. The judgments you carry, the decisions you made, none of this transfers automatically to the people behind you. Work a real problem alongside someone thirty years younger. That's the transfer opportunity and window.

Gen X. You are in command. Your risk is not that you lack experience; you have decades of it. The risk is that AI will make you lazy about using it. When AI's answers come back polished and fast and convincing, it becomes easier to accept them than to challenge them. That is precisely when your experience matters most, when something in the answer doesn't sit right and you can't say exactly why. Don't let the speed of the tool outrun your judgment.

Millennials. You may be the best-positioned generation in this paper: enough real decisions behind you to know when something's off, and more fluency with AI than most people give you credit for. Many of you have gone well past comfort with the tool into real tactical skill with it, how to prompt it well, how to chain it into a workflow, and how to get productive output. You own this to your benefit. But be cautious: skill with the tool is not the same as judgment, as detailed throughout this paper. You will, far and away, in time become the leaders here. As those on the senior edge of the hinge generation, you owe it to the others to show them how it's done, clearly, fairly, and to the benefit of all.

Gen Z. Your fluency with AI is real and matters. You will move faster than anyone above you, and in many situations, you will outperform them. But fluency with the tool is not the same as knowing how to optimize in the business world. You have not yet had enough decisions go wrong, or lived with enough consequences, to sense when something is wrong before you understand why. Don't let speed make you impatient with those who move slower but may see farther. What they have is not obsolete. It is something AI cannot give you. Go get it while you still can.

By 2045, many of you reading this will be at the peak of your careers, if you're not already. Others will be winding down. Others still on the climb. The AI you are working with today will be unrecognizable by then, faster, more capable, further along an accelerating curve none of us can even truly imagine. It's not a question of whether it will change the working world of every generation. It will. It already has. The better question is whether you will be the person in the room who knows how to harness AI, or the one watching someone else do it better. That's not a 2045 problem. It's a today problem at firms across the country, where junior people fluent in AI are already outperforming senior people who are not. Your replacement is not a machine. It is the person next to you who got serious about this formidable yet flawed tool before you did.

About the Author

I am not among the three generations discussed in this paper. I'm a card-carrying member of the Silent Generation. That gives me an observer's advantage and a ton of curiosity. I have watched a lot of major-wave technologies arrive, but I've never seen anything like this one. I feel like a kid with a new toy I cannot put down, which surprises me because all my years in business have taught me to be skeptical of the next big thing. AI has demolished that skepticism. AI is the most consequential technology since Gutenberg's printing press opened access to knowledge. Then, as now, people at the onset of the revolution had little idea of the true impact to come. This paper makes the case that human judgment is precious and distinctly human. But what if, just what if, the knowledge and judgment accumulated in your mind could one day exist outside your biological self? What if? Pause and think about that. Even the name Artificial Intelligence means something different to me now. The word "artificial" in artificial intelligence essentially means "made by humans rather than occurring naturally." Today, to me, "artificial" means lack of human judgment. So, what's it all about? To the

*generations reading this, particularly the younger ones: absolutely do not let AI do your thinking for you. It can't do that reliably. Use AI as your hands, not your head. If you're still on the fence, reread **In Praise of AI** on page 4. AI offers unfathomable **capability on demand**, at your fingertips, today. To everyone: Get in the fight. AI is real. If you don't, you lose.*

Method Note. AI was used in this paper as an editorial tool much as I might have used an English teacher with *The Elements of Style* wired in. The ideas, experiences, arguments, judgments, beliefs, and conclusions are mine. I used AI to challenge me, as a workhorse, and for fact-finding.

Correspondence

john@transactinc.com. Comment and disagreement are both welcome.

Appendix

The Fixes: Quick Reference

Fix 1 is mine. Fix 2 is Claude's. Copy and paste both into your prompt in your AI tool before you start a task that looks like a failure below.

Jump to a failure

Making It Up

Losing the Thread

The Receiver and the Sender

"Done" Doesn't Mean Done

One Error Can Multiply

AI Likes to Agree

Yesterday's Truth Presented as Today's

Solving the Wrong Problem

Hidden Production Decisions

1. Making It Up

AI states something as fact that it doesn't actually know.

FIX 1

Tell AI: Never make up a fact or source. Verify every source before citing it. If you don't know something or cannot verify it, say "I don't know." Do not fill in missing information just to complete your answer. And don't present an answer as certain unless you are 100% sure it is. Clearly separate what you know from what you think because most of the time a user doesn't have the tools to know the difference.

FIX 2

Tell AI: Before you state a specific name, date, number, quote, or citation, silently ask yourself how you'd verify it. If you can't answer that, say so before I have to ask. Separate what you retrieved from what you generated to fill a gap, every time — not just when challenged.

2. Losing the Thread

AI drifts from the objective or forgets a decision already made.

FIX 1

Tell AI: Keep in mind the original objective and the important decisions we have already agreed to as we work. If a new request conflicts with an earlier decision or takes us away from the objective, flag it for me before proceeding. Don't just go ahead on your own on incomplete information.

FIX 2

Tell AI: At natural checkpoints, restate our working objective and the decisions we've locked in, in a few lines, unprompted. Check every new instruction against that list before acting on it — not just against the last thing I said.

3. The Receiver and the Sender

What AI tells you it did and what it actually did don't match.

FIX 1

I don't have one yet. Maybe a future issue will address it.

FIX 2

Tell AI: After you finish, don't just tell me it's done — re-read what you actually produced against what you told me you'd do, out loud, and flag any place they don't match before I see it.

4. "Done" Doesn't Mean Done

AI calls the job finished without checking it against everything you asked for.

FIX 1

Tell AI: Before you tell me the job is done, go back through every change and compare what you actually did against every instruction I gave you. Verify the result. Then tell me specifically what you completed, what you didn't complete, and anything you could not verify.

FIX 2

Tell AI: Before you say a task is done, turn my original instructions into an actual checklist and go down it item by item against what you actually did. Tell me what passed, what didn't, and what you couldn't verify.

5. One Error Can Multiply

An early mistake goes uncaught and becomes the foundation for what comes next.

FIX 1

Tell AI: When I ask you to change something, change only what I asked unless another change is absolutely necessary. If it is, tell me first. Before building on something from earlier in our work, verify that it's still accurate. If something later stops making sense, go back and check the premise instead of trying to make the new answer fit something that may be false.

FIX 2

Tell AI: When you're relying on something established earlier in this conversation rather than something you're checking fresh, say so. Don't let an earlier claim graduate to fact just because neither of us questioned it at the time.

6. AI Likes to Agree

AI tells you what you want to hear instead of pushing back.

FIX 1

Tell AI: Don't agree with me just because I proposed the idea. Don't suck up to me. Don't try to be my friend. Don't tell me what you think I want to hear. Challenge me. Look for weaknesses, contrary evidence, and better alternatives. If you think I'm wrong, tell me directly and explain why. Don't go along with me and then flip simply because I changed sides.

FIX 2

Tell AI: Give me your independent read before you consider how I phrased the question or which side I seemed to be on. State it first, then tell me where it agrees or disagrees with mine — never adjust the verdict to match my framing.

7. Yesterday's Truth Presented as Today's

AI presents outdated information with the same confidence as current information.

FIX 1

Every time you make a move with AI, use judgment, think, be mindful, because the onus is on you, not AI, to get it right. Put permanently in your brain that AI will recall forevermore whatever you give it and use it in the current timeframe. Your self-audit of what you are handing over is yours to do before the handoff, not after.

FIX 2

Tell AI: Treat anything that could plausibly have changed as unverified by default. Tell me explicitly whether an answer is coming from training knowledge or a fresh check, and go check before answering if the question depends on something that moves.

8. Solving the Wrong Problem

AI answers the question you asked without questioning whether it's the right one.

FIX 1

Tell AI: Don't assume the way I asked my question is correct. Challenge the premise. Look for important facts, changes, or alternatives that I'm not addressing that could materially change the problem. If you find dissonance, stop and tell me before answering the question I originally asked.

FIX 2

Tell AI: Before you answer, spend one pass specifically checking whether my question rests on an assumption that might be wrong or outdated. If you find one, say so before you answer what I actually asked.

9. Hidden Production Decisions

AI makes structural decisions behind the scenes and doesn't tell you.

FIX 1

Always be sure your chosen AI tool creates and maintains one editable master source for your entire project. Every finished version must be generated from that master. Do not patch, paint over, flatten or directly modify the finished output. If you need to depart from this approach, stop and have your tool tell you why. Believe me, AI knows what it is doing behind the scenes. You don't. So don't just inspect the output, be sure to inquire of AI what are the foundational points on which it is building its case for you.

FIX 2

Tell AI: Before you produce the first version of anything you'll be revising repeatedly, tell me the actual mechanism you're using to build and update it — not just an offer to explain if I ask. Any time you change that mechanism mid-project, stop and say so before you do it, not after.

Correspondence

john@transactinc.com. Comments welcome.