

WHAT AI CAN'T REPLACE AND WHY THAT MATTERS

Designing AI Around Human Preference at Work

Thought Leadership



Introduction

Most leaders ask what AI can do. That is only half the question. The other half—the one that often determines whether AI succeeds inside an organization—is what *should* AI do? Where should it lead, where should it assist, and where should it step back?

Despite the massive advancements in technology, including generative AI, there are still moments when people prefer to speak to another human. Not because the automated system failed, but because trust, familiarity, and accountability feel different when a human is involved.

This tension shows up everywhere: customer service, healthcare, financial advice, people decisions, and leadership communication. Organizations may automate to improve speed and scale only to find that trust, satisfaction, or loyalty hinge on human presence. This isn't resistance to change or technology, but how work is experienced and judged when a machine replaces a human. The question facing leaders is therefore shifting: where people are comfortable with AI doing the work, where they prefer human involvement, and why.

The Korn Ferry Institute proposes a taxonomy of work that helps explain where human preference will persist, where AI adoption may face friction, and where AI is actively preferred.

Why AI is Judged Differently Than Humans

Humans are wired for social connection. Because collaboration and trust were essential to survival, our brains evolved to build and maintain complex social networks, at home and at work. Evolution made social and collaborative behaviors rewarding: connection with other people actually feels good.

Those same systems still shape how we understand interactions today. When we engage with others, we run unconscious but powerful assessments of intention, effort, and care, reading tone, gesture, and expression to sense each other's emotional state. In this way, human interaction does more than exchange information; it shapes whether an experience feels trustworthy, legitimate, or meaningful.

The Korn Ferry Institute proposes a taxonomy of work that helps explain where human preference will persist, where AI adoption may face friction, and where AI is actively preferred.

In contrast to human connection, engaging with AI activates the same physical and cognitive systems far less. Machines rarely adjust to situational cues, which alerts the brain that we're not dealing with another social agent, and affects cognitive processing patterns. Even when AI sounds natural or empathetic, something feels off. So, we tend to experience AI interactions as lower in depth and connection, which is one reason people still prefer a human in many workplace settings.

But the story has a second side. The same wiring that makes connection feel good also underpins feelings of social risk. We feel shame, embarrassment, and the fear of being judged just as acutely, so we hold back from the things that invite judgment: asking a basic question, admitting we're lost, or practicing something we have not yet mastered. This is where the absence of a human can be a relief rather than a loss. For some people, and in many everyday tasks, talking to a machine removes the audience and with it, the fear of looking foolish.

Why the Same Output is Judged Differently When AI is Involved

Across domains, research shows that people evaluate identical outputs depending on whether they believe a human or a machine produced them. Even when accuracy and quality are the same, perceived value can shift. In some cases, people prefer AI-generated outputs—until they learn that AI produced them. Once the source is known, evaluations can shift, often becoming less favorable, with outputs perceived as colder, less effortful, or less meaningful.

This reflects social evaluation, not technical skepticism. People don't just ask, "Is this correct?" They also ask, "Who is behind this?" and "What does that say about the organization delivering the message?" Today, final deliverables are often the result of a hybrid approach, where humans and AI jointly produce different parts of one whole.

Perhaps a human has the original idea, uses multiple AI tools to explore angles and arguments, draft early iterations, code or prototype parts, or gather feedback.

Transparency plays an important role here. When the source of the work is revealed, it raises questions of provenance, which carries psychological meaning that shapes perception. Understanding where humans hold preferences and why can help leaders decide which parts to keep human in this hybrid model, and where AI can add impact without compromising trust.

Expectation Asymmetry and the Cost of Error

Expectations for humans and machines differ. When a human makes a mistake, people often consider context, constraints, or intent. When an AI system makes a mistake, it is more likely to be interpreted as a failure of design or judgment.

This dynamic contributes to what is often described as algorithm aversion. Machine errors, even if they are less frequent, feel less forgivable. These tendencies toward harsher judgments can undermine organizational trust and reputation. The social cost of "getting it wrong" is often higher for visible automation than for human judgment.

This penalty can extend to organizations and individuals who choose to use AI. Reliance on AI can be perceived as lower effort, care, or accountability—regardless of results. Even when algorithms perform as well, or better than human judgment, people may still prefer human involvement in the process, especially in consequential or personally meaningful decisions.

Social Meaning and AI Use

Many forms of work require more than the right answer. They require someone to explain trade-offs, own consequences, and be held accountable.

Where accountability matters, trust depends not only on accuracy, but on the presence of a decision-maker who is responsible for the results. When AI is positioned as an autonomous decision-maker, trust often erodes because accountability feels absent.

As AI produces work that resembles human talent, it is increasingly evaluated on social dimensions such as warmth, authenticity, and moral character. People implicitly ask whether an output reflects meaningful depth and connection and comes from a “genuine” source. When the source is perceived as artificial, perceived authenticity can decline, even if the surface-level language is empathic.

Taken together, these dynamics explain why it is critical to go beyond understanding AI's capabilities and consider human preference in the workplace.

The Taxonomy of Human Preference at Work

The point is not that humans are always preferred, or that AI should be avoided. Human preference is nuanced; it depends on the nature of the work.

To help leaders deploy AI more intentionally, we offer a taxonomy of work: where people strongly prefer human involvement, where AI adoption is likely to face friction, and where AI may actually be preferred.

1. Work Where the Source Is the Value

In some work, who creates something is part of what gives it value. Authorship, taste, and originality are an essential part of the outcome. This includes brand stories, executive communication, and culture-defining moments. In these cases, the question isn't “is this good?” but “who created it, and what does that signal?” AI can support drafting and research, but humans must remain visible owners.

Design implication: Use AI as an amplifier, not to replace the human voice.

2. Work Where Being Seen Is the Service

In relational work—coaching, therapy, leadership conversations, high-stakes sales—the value goes beyond information. It depends on perceived trust, empathy, connection, and the experience of being understood. Human-specific tone, timing, emotional attunement, and the ability to adjust in real time are essential to defining the experience.

Design implication: AI can assist operations and logistics, but replacing the human presence risks eroding trust.

3. Work Where Judgment Must Be Owned

Some decisions involve ambiguity, tradeoffs, and real consequences. In these cases, people want someone who can navigate nuances and will stand behind the outcome. This includes financial advice, medical decisions, and legal judgment.

Design implication: Let AI inform judgment, not replace the decision-maker.

4. Work Where Moral Authority Matters

Some work is not only technical but normative: it shapes opportunity, fairness, dignity, or consequence. Hiring, promotion, discipline, and resource allocation fall into this category. Here, consistency alone is not enough. People care not only whether decisions are applied evenly, but whether they feel fair, humane, and answerable to human values.

Design implication: Preserve human responsibility for value-laden decisions.

5. Work Where Removing the Human Helps

In some cases, the absence of a human is beneficial. Routine tasks, information retrieval, drafting support, translation, coding assistance, reviewing personal information, or working through practice problems can feel easier without an audience. Here, AI can lower the threshold for participation by making help feel immediate, private, and judgment-free. This can encourage more experimentation, learning, and participation.

Design implication: Deploy AI confidently where privacy and low stakes increase engagement.

What This Means for Leaders

The future of work will be shaped by how leaders design AI around human expectations. Human preference is a design constraint to respect, rich with information about what we value and why.

Organizations that succeed will:

- Keep humans visible where trust, legitimacy, and meaning matter
- Use AI to reduce friction where social presence adds little value
- Design human-AI systems with intention

The goal is not to choose between humans and machines. It is to use each where they create the most value (Figure 1: Human+ Solution Framework).

Human preference is one factor organizations need to consider when preparing to integrate AI into their workforce. What's more, they need to consider that this is a moment in time. In an AI-enabled world, organizations that possess the greatest competitive advantage will not ask only what can be automated. **They will ask what should remain human—and why.**

Korn Ferry's AI Impact Score is built on the Korn Ferry Success Profiles: a structured, calibrated library of profiles that describe what success in a particular role looks like when done well. It also looks at where humans are preferred and where legal obligations need to be pursued. This library spans industries and job functions. Rather than scoring profiles as a whole unit, the AI Impact score decomposes each profile into its component responsibilities and assesses each one individually for its exposure to automation, augmentation, and transformation by AI. Large language models are used to assess impact, using a consistent responsibility taxonomy and calibrated scoring criteria applied across profiles. Because scoring happens at the most granular levels, the findings make it possible to see AI's impact at whatever level of work a given decision requires.

AUTHORS

Amelia Haynes

Manager, Research & Partnership Development
Korn Ferry Institute

Hudson Pfister

Research Associate
Korn Ferry Institute

Petr Chocholka

Director, Org IP Development
Korn Ferry Institute

CONTRIBUTORS

Karin Visser

Vice President, Org IP
Korn Ferry Institute

Josh Royes, PhD

Associate Director, Org IP & Analytics
Korn Ferry Institute

Ondrej Vacha

Contractor
Korn Ferry

About Korn Ferry

Korn Ferry is a global consulting firm that powers performance. We unlock the potential in your people and unleash transformation across your business—synchronizing strategy, operations, and talent to accelerate performance, fuel growth, and inspire a legacy of change. That's why the world's most forward-thinking companies across every major industry turn to us—for a shared commitment to lasting impact and the bold ambition to *Be More Than*.

Appendix:
Figure 1: The Human+ Solution Framework

THE HUMAN+ SOLUTION FRAMEWORK

