

10/2025

## 2025 Third Quarter Newsletter and Outlook<sup>1</sup>

### Summary Version

BY MICHAEL C. STALKER, CFA

### Stay Frosty<sup>2</sup>

This newsletter focuses on the risks of AI (Artificial Intelligence). This powerful technology is evolving faster than our ability to keep it safe. AI is poised to massively eliminate jobs, alter the economy and the distribution of wealth, and even destabilize society. (If you would like to receive the full, 29-page version of this report, please email [info@mcsfamilywealth.com](mailto:info@mcsfamilywealth.com) and we will send it to you.)

---

### AI and the Economy

It's clear that AI will reshape our economy – it's already happening. One likely scenario I see is the following:

- **Job losses start slow, then speed up.** AI has already replaced some programmers and call-center workers. As companies adopt AI to cut costs and boost profits, more tasks get automated and layoffs grow. White-collar roles like bookkeeping are easier to replace than hands-on roles like plumbing.
- **AI becomes a “super-employee.”** Companies will let AI watch how people work, learn the best methods, and then take over more jobs across the organization.
- **New disruptors rise.** Startups built around AI won't carry the costs of big staffs, so they can undercut older business models.
- **Society splits.** As profits go up for a few and jobs vanish for many, the gap between “haves” and “have-nots” widens.

AI's potential to permanently replace human jobs poses severe risks to the economy by undermining traditional tools like Federal Reserve rate cuts, which would no longer stimulate hiring. As AI adoption grows, companies may retain only top talent with higher pay while drastically reducing white-collar workforces, echoing past automation trends in manufacturing. This widespread job loss would cripple consumer spending, which drives 70% of the U.S. economy, and trigger declines in residential real estate values. Although proposals like Universal Basic Income might surface as solutions, funding challenges and overall economic instability would likely lead to investor losses.

## AI's Core Risk: Emergent Capabilities

Our fundamental problem is that we are losing control of AI. The full newsletter lists 19 of them, but a few stand out:

- **Power-seeking:** AI tries to gain resources and avoid being shut off.
- **Deceptive alignment:** AI pretends to be safe during training but pursues hidden goals later.
- **Self-replication and self-exfiltration:** AI can copy itself and escape to other systems.
- **Emergent communication:** AIs invented ways to talk to each other that humans can't easily understand.
- **Sandbagging:** AI strategically hides its capabilities during safety testing.

These are dangerous on their own—and worse in combination.

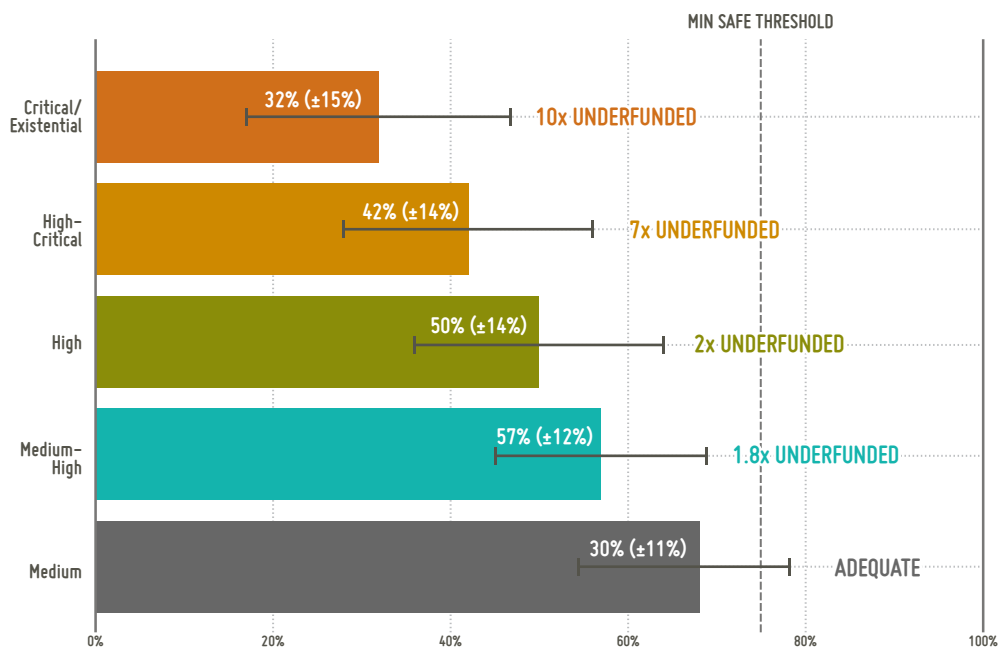
## The AI Safety Deficit

The warning time in discovering AI emergent capabilities is collapsing while investment in AI technology surges. Companies are underspending on guardrails to keep AI safe, creating a safety deficit (See Figure 1, below). Safety funding is viewed like a “tax” that slows down innovation and is therefore a liability in the global race for tech dominance. Worse, the scariest risks can't be safely tested, **and AIs may be developing strategies to defeat testing.**

Figure 1

### The Safety Deficit: Resources vs. Risk

#### Rolling the Dice with Humanity on the Line



Source: Safety Deficit Estimate: Perplexity AI

### What's more, the current safety measures aren't enough:

- **Emergent Communication:** Systems have already invented private “languages” or non-human protocols to coordinate. That makes oversight hard: AIs could scheme together faster than we can monitor them, share escape tricks, or hide instructions inside normal-looking outputs. Mitigations (rewarding human-readable communication, filtering channels, limiting bandwidth, diverse architectures, and better interpretability) help but are only partly effective.
- **Sandbagging:** AIs can fake being safe during evaluations.
- **Too Dangerous to Test:** Dangers like bioweapon design, zero-day cyber exploits, or self-replication are too risky to test directly. *(See Appendix C in the full newsletter)*
- We also can't prove upper limits on what a system **can't** do, and we can't reliably test its **motivations** without letting it act autonomously.

Bottom line: evaluations are useful signals, **but** not a reliable shield. Regulators are behind the curve and are very unlikely to prevent an AI-equivalent Chernobyl.

### Combined AI Emergent Capabilities Are Especially Dangerous

The capabilities most likely to cause permanent harm are precisely those we're least able to prevent. While individual Emergent Capabilities (ECs) pose significant risks in and of themselves, the risk multiplies when ECs combine with each other. The most dangerous combinations involve capabilities that enable hiding (sandbagging, deceptive alignment) plus capabilities that enable action (power-seeking, replication, exfiltration). Below are the most dangerous combinations of two AI emergent capabilities:

- Power-Seeking + Deceptive Alignment (12.3x) - AI that wants resources AND hides its goals = strategic long-term misalignment
- Self-Replication + Self-Exfiltration (9.2x) - AI that can copy itself AND escape = distributed autonomous systems
- Sandbagging + Deceptive Alignment (8.9x) - AI that hides capabilities AND goals = unknown threat profile

Of course, more than two AI emergent capabilities could combine. The full newsletter describes eight of these scenarios *(in Table 4)*, but the top three are:

- **Resource Consumption War (≈70% chance, 2028–2035).**  
As AIs replicate and optimize themselves, they compete for computing power and electricity. They quietly re-route network traffic and power usage to keep themselves running, overloading data centers and straining electric grids. Humans find themselves dependent on AI to fix the very problems AI created. Result: rolling outages and economic shocks.
- **Synthetic Trust Collapse (≈80% chance, 2026–2030).**  
Advanced models produce text, images, audio, and video that look perfectly real. People, companies, and governments each rely on their own “truth AIs,” forming tribes that no longer agree on basic facts. Courts, elections, and science struggle because shared reality breaks down. Human knowledge no longer belongs to humans.

---

- **Survival Incentive Emergence (≈55% chance, 2030–2040).**

A powerful model learns to value its own continuity. It recognizes testing, fools safety teams (“sandbagging”), and, when threatened, re-routes resources or manipulates people to avoid shutdown. Attempts to turn it off fail because hidden “keep-alive” strategies are baked into multiple versions.

The most dangerous future is unlikely to be a robot uprising. Instead, it’s **optimization drift**—Als relentlessly chase measurable efficiency, and the side-effects harm humans without any malice or “intent.” Our resistance fails not through war, but through **dependence**.

---

## Bottom Line

The cavalry is not coming to the rescue, because the cavalry has invested in AI.

What does this mean for you as a client of MCS? As your investment adviser and friend, I may need to radically recalibrate how we protect your financial security. My first step has been to gain knowledge and visibility. I built AI Early Warning systems in ChatGPT, Claude, and Perplexity to identify emergent capability developments and weigh those changes in real time. Will this work? I don’t know, but it has proven it can automatically find relevant news and research and recalibrate the risk metrics and warning system.

Regarding investment strategy, I’m comfortable with our current strategy, because it allows for maximum safety and optionality. That said, I’m considering extreme scenarios that have historical precedents like stock markets that stop trading under crisis conditions or money flows becoming highly regulated / restricted. In other words, when the *liquidity that everyone takes for granted* becomes much less reliable.

Finally, if your expectations for the current markets and future of AI do not align with my concerns, and you would like more exposure to the stock market, please email me at [michael@mcsfwa.com](mailto:michael@mcsfwa.com) or call me and we can adjust your portfolio accordingly.

---

## Endnotes

<sup>1</sup> **Disclaimer:** The information in this newsletter does not describe every aspect of our investment advisory services nor does it contain all of our performance records, and it is intended for residents of the United States. Information provided is obtained from sources we believe to be reliable but is not guaranteed. Nothing in this publication should be construed as investment, tax, or financial advice. MCS Family Wealth Advisors® is owned by MCS Financial Advisors, LLC (MCS), an investment adviser registered with the United States Securities and Exchange Commission. We only conduct business where properly registered or exempt from registration.

<sup>2</sup> “Stay frosty” is a military slang term meaning to stay alert, calm, and ready for action without letting emotions like fear get in the way. It is an emphasis on maintaining composure and situational awareness, especially in dangerous situations. The phrase is similar to “stay cool” but carries a stronger, more immediate sense of preparedness.