

[PODCAST] AI: How Businesses Have Been Managing Rapid Change

Published 8/4/2026 | Last Updated: 8/4/2026

Executive Summary: As artificial intelligence rapidly integrates into the workplace and healthcare industry, organizations face a critical gap between rampant employee usage and the implementation of corporate guardrails. Unregulated AI use exposes businesses to severe reputational, cyber, and regulatory risks. This guide outlines the differences between AI models, the specific risks to mitigate, and actionable steps senior management can take to safely harness its power without compromising security.

Generative AI vs. Agentic AI: What is the Difference?

Understanding the tools your employees are using is the first step in mitigating risk.

Feature	Generative AI	Agentic AI
Action Type	Reactive in nature.	Proactive and goal-oriented.
Functionality	Creates new content only in response to a user prompt.	Plans, decides, and executes multiple steps autonomously.
Examples	ChatGPT.	AI platforms that act as project managers.
Risk Profile	High risk of “hallucinations” or fabrications in generated content.	Exacerbated cyber and regulatory risks due to its ability to pull data and interact with multiple electronic systems with limited human involvement.

What are the Risks of Unregulated AI in the Workplace?

While AI offers meaningful efficiencies (such as AI scribes, patient scheduling agents, and revenue cycle management tools in healthcare), the lack of management oversight presents significant liabilities.

- **Reputational Damage:** Organizations can suffer irreparable harm from publishing AI-generated fabrications. For example, Deloitte’s Australian business was forced to retract a \$290,000 government study due to AI hallucinations.
- **Judicial Embarrassment:** Two federal judges were publicly called out by the Senate Judiciary Chair for publishing opinions containing AI fabrications.
- **Healthcare and Patient Risks:** In healthcare, risks include the compromise of patient data privacy and security, biased AI output, and the potential for misdiagnosis.

- **Skill Degradation:** Overdependence on AI tools can actively reduce the clinical capabilities of staff.

Action Plan: How to Build AI Guardrails for Your Organization

To close the gap between employee use and corporate management, organizations must establish clear procedures.

1. **Enforce an Interim Ban on Agentic AI:** Until a holistic policy is developed, advise employees that Agentic AI cannot be used for work purposes without express management approval.
2. **Appoint AI Leadership:** Designate a specific person or group with ultimate responsibility for AI. Because AI management is fundamentally a “people issue,” this leader must have authority over managing personnel.
3. **Conduct Private Briefings:** Provide private briefings on AI risks to senior management and board members to create a “burning platform” that kickstarts the implementation of guardrails.
4. **Develop Formal Policies:** Create acceptable use policies and integrate them into existing corporate and HIPAA compliance programs.
5. **Educate the Workforce:** Train both management and employees on the risks and the new procedural guardrails. Require employees to formally disclose when they are using AI.
6. **Vet Vendors and Negotiate Contracts:** Management must understand the evolving AI marketplace to properly vet vendors. Ensure detailed legal review of all AI vendor contracts. Be prepared to walk away from vendors who lack transparency or cannot provide sufficient assurances regarding the privacy and security of protected health data.
7. **Audit and Monitor:** Track regulatory changes, institute compliance measures, and conduct ongoing auditing.

Mini-FAQ

Q: Are businesses actually facing real-world consequences for poor AI use? **A:** Yes. High-profile organizations and individuals, including federal judges and Deloitte’s Australian business, have suffered severe public embarrassment and reputational damage for attaching their names to shoddy, hallucination-filled AI work product.

Q: Why is Agentic AI considered riskier than Generative AI?

A: Agentic AI platforms can act autonomously, pulling data from multiple electronic systems and executing functions without

human involvement. This deep system integration significantly exacerbates cyber and regulatory risks.

Q: What should an organization do if an AI vendor refuses to be transparent about data security?

A: Legal counsel often advises organizations to simply walk away from these contracts. If a vendor cannot provide sufficient privacy assurances (especially regarding protected health data), the organizational risk is too high.

Author & Credibility Details

- **Lani M. Dornfeld, CHPC, Esq.**

- Member, [Healthcare](#), Brach Eichler LLC
- 973.403.3136 | ldornfeld@bracheichler.com

- **Jay Sabin, Esq.**

- Member, [Labor & Employment](#), Brach Eichler LLC
- 917.596.8987 | jsabin@bracheichler.com

- **Source:** [Brach Eichler Talks Podcast](#) (July 2026)

- **Disclaimer:** This content is intended to provide general information, not legal advice. Please contact Lani or Jay to consult on your organization's specific AI management needs.

Authors

The following attorneys contributed to this insight.



Lani M. Dornfeld

CHPC, Member

Healthcare Law

973.403.3136 • 973.618.5536 Fax

ldornfeld@bracheichler.com



Jay Sabin

Member

Labor and Employment, Litigation,
Cannabis Industry

917.596.8987 • 973.618.5907 Fax

jsabin@bracheichler.com