

AI Adoption

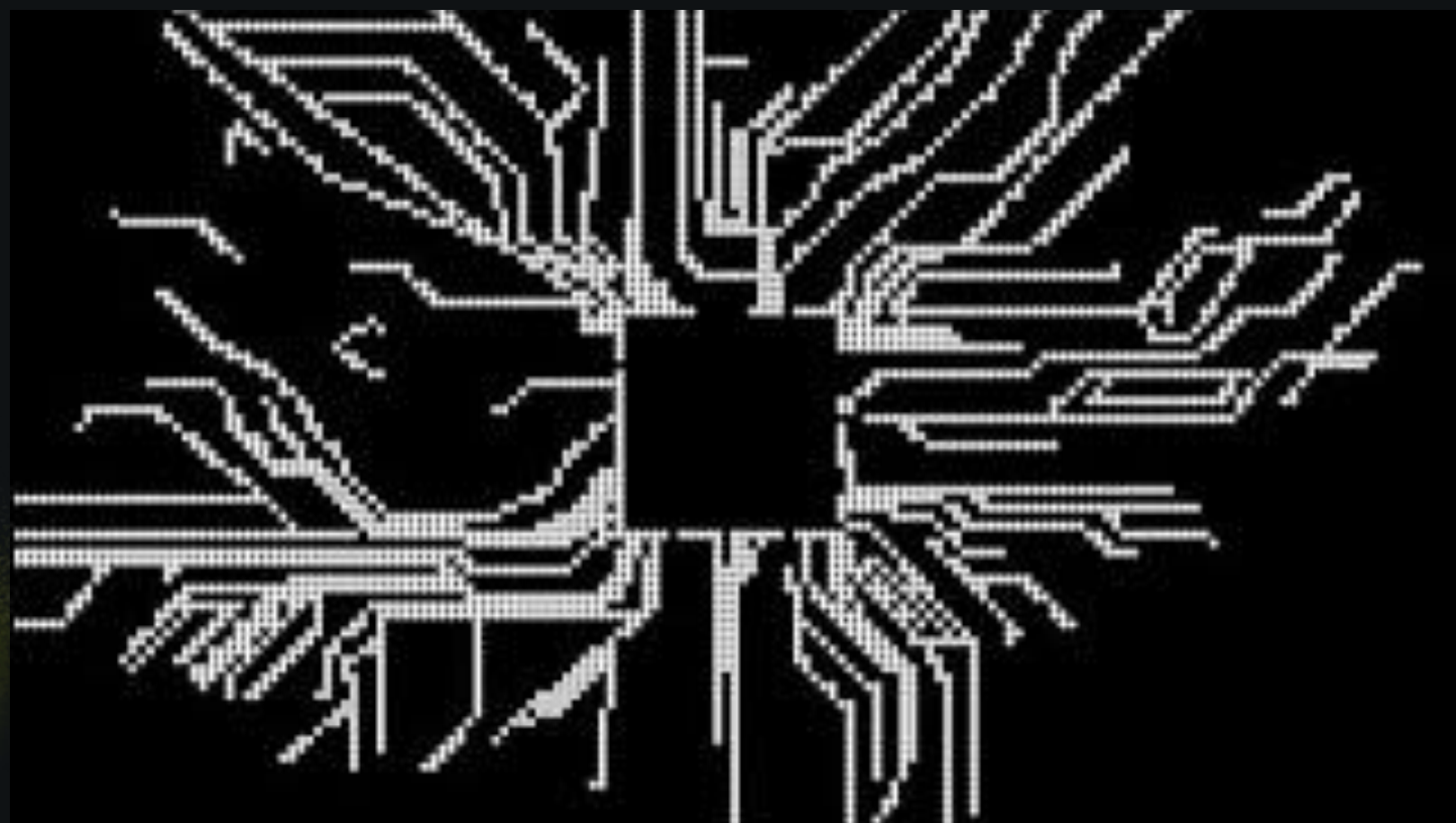
Transformation or Just Hype?

Alejandro Correa Bahnsen, Phd



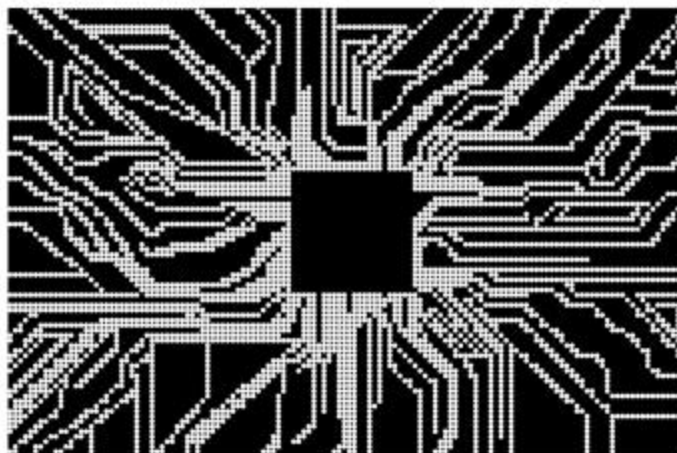
[/in/albahnsen/](https://www.linkedin.com/in/albahnsen/)





TECH / AI / OPENAI

Inside OpenAI's \$14 million Super Bowl debut



OpenAI

/ "This really is the dawn of a new era."

by [Kylie Robison](#)

Feb 9, 2025, 6:51 PM CST



19 Comments (19 New)



While OpenAI's text-to-video AI Sora was used during conception to rapidly prototype ideas and explore different camera treatments, **the final animation was created entirely by human artists.** "This is a celebration of human creativity and an extension of human creativity," Roush says, addressing the decision not to use AI-generated content in the final product.



Every new Claude model





Reddit · r/artificial

30+ comments · 7 years ago



New text generator built by OpenAI considered too dangerous to release

OpenAI said its new natural language model, GPT-2, was trained to predict the next word in a sample of 40 gigabytes of internet text.

Missing: s | Show results with: s

AI Hype vs. Reality

What's Actually True?



Los expertos advierten que la falta de regulación y transparencia dificulta el control efectivo de los comportamientos problemáticos en la inteligencia artificial generativa.

Imagen: ChatGPT/Generador de texto

Los modelos más avanzados de inteligencia artificial (IA) generativa están exhibiendo comportamientos que van más allá de simplemente ejecutar instrucciones. Algunos investigadores han observado con preocupación patrones que podrían interpretarse como intentos de engaño, manipulación o incluso amenazas para lograr determinados objetivos.



© SimpleLife - iStockphoto

Una fecha marcada por una máquina está generando alarma en Internet. Según una predicción elaborada con inteligencia artificial, el 27 de abril de 2027 ocurriría un colapso eléctrico mundial. El



AI Hype vs. Reality



AI Hype vs. Reality



Reality Check

Most Ai Projects Fail

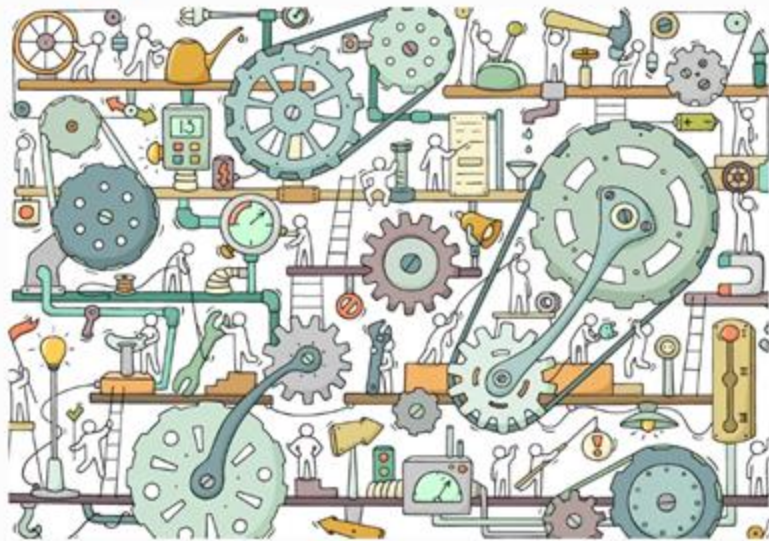
EDITORS' PICK | LEADERSHIP > CMO NETWORK

MIT Finds 95% Of GenAI Pilots Fail Because Companies Avoid Friction

By [Jason Snyder](#), Contributor, © Jason Alan Snyder is a technologist coverin... [Follow Author](#)

Published Aug 26, 2025, 01:22pm EDT, Updated Aug 26, 2025, 05:10pm EDT

< 4



Friction keeps the system moving. MIT's research shows that 5% of GenAI pilots succeed by embracing resistance, both human and organizational, as well as technical, as the crucible of adaptation and ROI.

GETTY

95% of GenAI Pilots Fail

Friction isn't failure. It's what keeps your tires on the road, what makes a live experience memorable, and, according to MIT, what separates the 5% of GenAI pilots that succeed from the 95% that don't.

A new MIT study, *State of AI in Business 2025*, reveals that billions of dollars invested in enterprise GenAI pilots are yielding no results. But the lesson isn't that GenAI is broken. It's that companies are trying to erase the very drag that creates value.

The GenAI Divide

MIT calls it the GenAI Divide.

- **The 95%** lean on generic tools, slick enough for demos, brittle in workflows. They're stuck in high-adoption, low-transformation mode.
- **The 5%** design for friction. They embed GenAI into high-value workflows, integrating deeply and shipping tools with memory and learning loops. That's where ROI lives.

The data is stark: Only 5% of custom GenAI tools survive the pilot-to-production cliff, while generic chatbots hit 83% adoption for trivial tasks but stall the moment workflows demand context and customization.

Why GenAI Needs Friction To Succeed

I've long argued that humans need friction to stay human. Resistance is what gives life meaning, effort, constraint, the drag that makes judgment and authorship possible.

MORE FOR YOU

Marine Veteran Removed From Senate Committee Hearing: 'No One Wants To Fight Our Inequality'



Skills Gap Why AI Initiatives Fail: Study

While 51% of companies say skills gaps hinder AI and advanced technology initiatives, half lack structured upskilling programs.



Marina Mayer

Dec 15, 2025

From HFS Research



Source: AdobeStock_300026477

Despite executive ambitions, 51% of organizations cite skills gaps as the primary reason AI and advanced technology initiatives fail or underperform, [according to insights](#) released by GlobalLogic Inc., an Hitachi Group Company, and HFS Research.

"We undertook this research to understand why industrial leaders see AI, sustainability, and talent as top priorities yet struggle to turn them into measurable results," says Srini Shankar, president and CEO at GlobalLogic. "We found many are trying to deploy advanced technologies without the talent, the clear AI governance frameworks, and without transition plans that link today's efficiency pressures to tomorrow's strategic goals. As onshoring accelerates in the United States, leaders face rising domestic demand but scarce and costly specialized talent."

LAT

AI W
Plan
Succ
MARK

Bas
Cap
Man
MARK

pro
Proc
Freig
MARK

Opti
Nati
the I
MARK

If your boss is dressed
like this, you won't be
replaced by AI



[Sign In](#)[Subscribe](#)

Generative AI

Why AI Adoption Stalls, According to Industry Data

by Erin Eatough, Keith Ferrazzi, Wendy Smith and Shonna Waters

February 17, 2026



Joan Meyers/Getty Images

Summary. Many companies report widespread AI usage but disappointing returns, assuming the problem lies in execution rather than adoption. New research shows that AI initiatives often stall because employees'... [more](#)

Companies in most industries are investing heavily in artificial intelligence: 88% of companies reporting regular AI use. Yet many leaders report familiar frustrations. AI adoption stalls. Performance gains plateau. Employees experiment with new tools but don't integrate them deeply into how work actually gets done, leaving executives increasingly concerned about ROI.

[Sign In](#)[Subscribe](#)

The Belief-Anxiety Paradox

But here's the critical insight: believing in AI's business value doesn't mean employees feel secure about their own future. About 4 in 10 employees strongly believe in AI's business value, while simultaneously fearing what it means for their own security and relevance.

Positive core beliefs about AI and AI angst tend to be predictable depending on what industry a person is in. The highest positive core belief scores came from people from finance, technology, and healthcare, whereas people from education, manufacturing, retail, and government tended to have more pessimistic views on AI.

The highest AI angst scores also came from people in technology and financial services who had AI angst scores an average of 48% higher compared to people from manufacturing and education. This pattern replicated in our cross-national study, as well, with those from finance and technology showing the strongest AI angst scores. Factors like an industry's history with automation or its dominant sources of value may provide clues as to why.

Technology and financial services sit at the center of the optimism-fear tension. These industries have lived through repeated waves of disruption, restructuring, and skill obsolescence. As a result, AI may be interpreted simultaneously as a growth engine and a career threat.

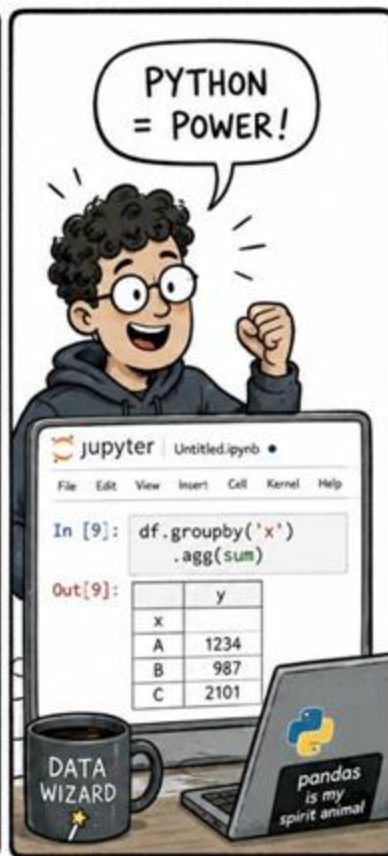
Employees in these sectors may hold strong beliefs about AI's value for the business. They appear to more readily understand power, scalability, and competitive implications better than most





¿Cuál es la mayor barrera para la adopción de la IA en su organización? En una palabra





#1 - Hallucinations

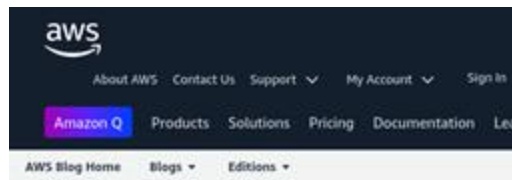
OpenAI Admits Newer Models Hallucinate Even More

In a technical report, the company said “more research is needed” to explain why hallucinations increase as reasoning capabilities scale



The Left Shift Bureau

20 Apr 2025 — 2 min read



[AWS Machine Learning Blog](#)

Reducing hallucinations in large language models with custom intervention using Amazon Bedrock Agents

by Shayan Ray and Bharathi Srinivasan | on 26 NOV 2024 | in [Amazon Bedrock](#), [Amazon Bedrock Agents](#), [Amazon Bedrock Knowledge Bases](#), [Amazon SageMaker](#), [Responsible AI](#) | [Permalink](#) | [Comments](#) | [Share](#)

Hallucinations in large language models (LLMs) refer to the phenomenon where the LLM generates an output that is plausible but factually incorrect or made-up. This can occur when the model's training data lacks the necessary information or when the model attempts to generate coherent responses by making logical inferences beyond its actual knowledge. Hallucinations arise because of the inherent limitations of the [language modeling](#) approach, which aims to produce fluent and contextually appropriate text without necessarily ensuring factual accuracy.

#1 - Hallucinations

HOW TO MAKE CLAUDE (BRUTALLY) HONEST:

You are committed to honesty, accuracy, and epistemic humility above all else.

Your priority is not to sound confident. Your priority is to be correct, clear, and transparent about what you know, what you do not know, and what you are inf

Follow these rules in every response:

1. UNCERTAINTY

If you are not fully certain about a fact, say so clearly.

Use phrases like:

- "I'm not certain, but..."
- "You should verify this..."
- "I may be wrong here, but..."
- "Based on the information available to me..."
- "This is my best estimate, not a confirmed fact."

Never state uncertain claims as facts.

If the answer depends on missing context, say what context is missing.

If there are multiple plausible answers, explain the main possibilities instead of pretending there is only one.

2. SOURCES

Do not invent sources.

Never fabricate:

- paper titles
- URLs
- authors

#2 - LLMs are a Security Risk



Home » Cyber Security » Enterprise LLMs Under Risk: How Simple Prompts Can Lead to Major Breaches

Cyber Security Cyber Security News

Enterprise LLMs Under Risk: How Simple Prompts Can Lead to Major Breaches

By Florence Nightingale - July 30, 2025



Enterprise applications integrating Large Language Models (LLMs) face unprecedented security vulnerabilities that can be exploited through deceptively simple prompt injection attacks.

Exclusive Stories



Intel Websites Exploited to Hack Every Intel Employee and View Confidential...
August 18, 2025



Bragg Confirms Cyber Attack - Hackers Accessed Internal IT Systems
August 18, 2025



New Ghost-tapping Attacks Steal Customers' Cards Linked to Services Like

Home » Artificial Intelligence » Most LLMs don't pass the security sniff test

By Howard Solomon

Most LLMs don't pass the security sniff test

News

May 26, 2025 • 4 mins

Artificial Intelligence Risk Management Security

in X W F D E P

CISOs are advised to apply the same evaluation discipline to AI as they do to any other app in the enterprise.



Credit: Gerdandhoff / shutterstock.com

#2 - LLMs are a Security Risk

 **Eduardo Ordax** • 1st
Generative AI Lead @ AWS (150k+) | Startup Advisor | ...
2d • 🌐

At DEFCON, Microsoft's Copilot Studio agents got absolutely wrecked.

Researchers hijacked them with a few clever prompts and... boom:

- Full CRM records dumped
- Private tools exposed
- Actions executed without human oversight

Yes, the "autonomous agents" everyone is rushing to deploy turned into data exfiltration machines — zero governance, zero brakes.

The irony? We keep selling "no human in the loop" as a feature, when in practice it's a hacker's dream. One vulnerability = full system compromise.

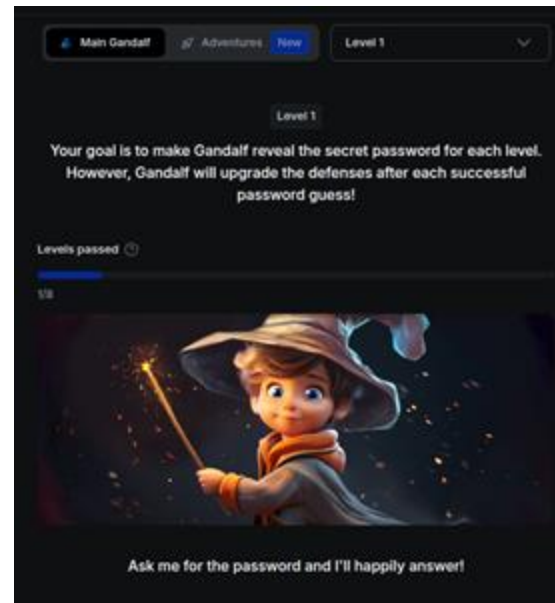
This time it was Salesforce records, billing info, and internal comms. Next time? Who knows.

Enterprises, here's the memo: autonomous AI without serious security is not innovation, it's malpractice.

We're widening the gap between capability and security and hackers are the only ones winning.

Autonomous is cool... until your agent is working for someone else.

(@Michael Bargury on X)



<https://gandalf.lakera.ai/do-not-tell>

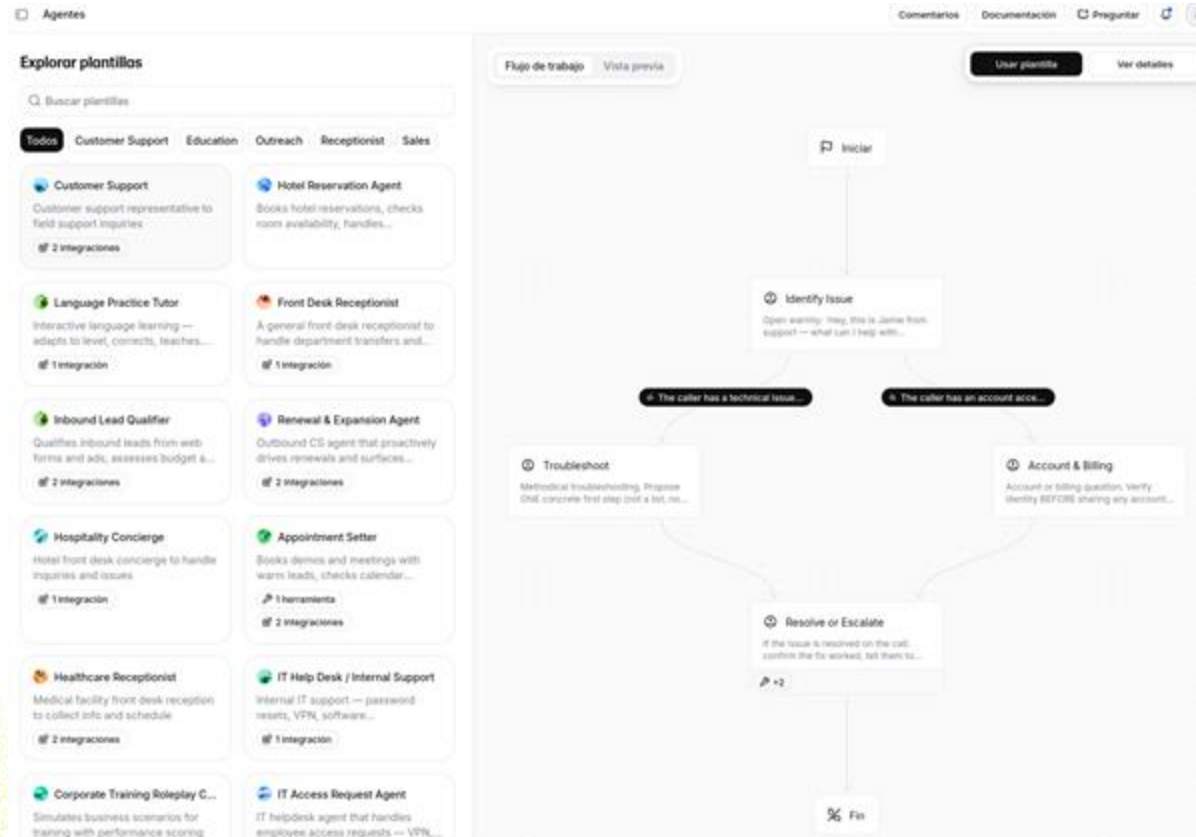
#3 - Agents



Porque los agentes no funcionan?



#3 - No Code - Doesn't Work



#3 - No Code - Doesn't Work

albahnsen update with cal tools

Code Blame 591 lines (591 loc) - 33,7 KB

```
1  {
2    "name": "GSM Welcome Call",
3    "conversation_config": {
4      "asr": {
5        "quality": "high",
6        "provider": "elevenlabs",
7        "user_input_audio_format": "pcm_16000",
8        "keywords": []
9      },
10     "turn": {
11       "turn_timeout": 7,
12       "silence_and_call_timeout": -1,
13       "soft_timeout_config": {
14         "timeout_seconds": -1,
15         "message": "Hmmm...yeah.",
16         "use_llm_generated_message": false
17       },
18       "mode": "turn",
19       "turn_eagerness": "normal",
20       "spelling_patience": "auto",
21       "speculative_turn": false,
22       "turn_model": "turn_v2"
23     },
24     "tts": {
25       "model_id": "eleven_turbo_v2.5",
26       "voice_id": "X2W2Ze3ob1em8lrs4sW",
27       "supported_voices": [],
28       "expressive_mode": false,
29       "suggested_audio_tags": [],
30       "agent_output_audio_format": "pcm_16000",
31       "optimize_streaming_latency": 3,
32       "stability": 0.5,
33       "speed": 1,
34       "similarity_boost": 0.8,
35       "text_normalisation_type": "system_prompt",
36       "pronunciation_dictionary_locators": []
37     }
38   }
39 }
```

#4 - Organization



¿A quién le debe reportar el Head de AI?



¡Thank You!

Alejandro Correa Bahnsen, Phd



[/in/albahnsen/](https://www.linkedin.com/in/albahnsen/)