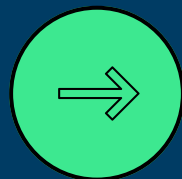


Eight Weeks in Review Performance Ray Eitel-Porter Report



MAY 2026

Executive Summary

Over the past eight weeks, your thought-leadership presence stepped up significantly.

Impressions grew nearly 8.5x against the prior period to just under 100,000, putting your content in front of a meaningfully larger pool of targeted professionals. Engagement jumped over 9x to 3,219, and your audience grew by 1,000 new followers, a 19% lift on total followers. Inbound conversations also moved, with 200 connection requests generated across April and May.

Together, these signals a strong foundation for long-term influence, commercial visibility, and continued momentum into Q2 2026.

Engagement

3,219

(Apr – May total)

3,219 Engagements
▲ 928.4% vs. prior 37 days

Impressions

98,384

(Apr – May total)

98,384 Impressions
▲ 749.7% vs. prior 37 days

New Followers

1,000

(Apr – May total)

6,261 Total followers
▲ 19% vs. prior 37 days

Connection requests

200

(April – May total)

Content performance

Impressions ▼

Daily ▼

109,254 Impressions

▲ 190.1% vs. prior 73 days



Daily data is recorded in UTC

First post with you
April 7th 2026

Content performance

Impressions ▼

Daily ▼

98,081 Impressions

▲ 747.1% vs. prior 37 days



Daily data is recorded in UTC

747.1% increase – expertise was
always there we wanted to turn it into
influence

Engagement

3,202 Social engagements

 Reactions	2,241
 Comments	549
 Reposts	25
 Saves	346
 Sends on LinkedIn	41

12 Link engagements

 Visits to links in all posts	12
--	----

Discovery

98,081 Impressions

36,547 Members reached

Audience Insights

Top demographics ⓘ

All

Job title

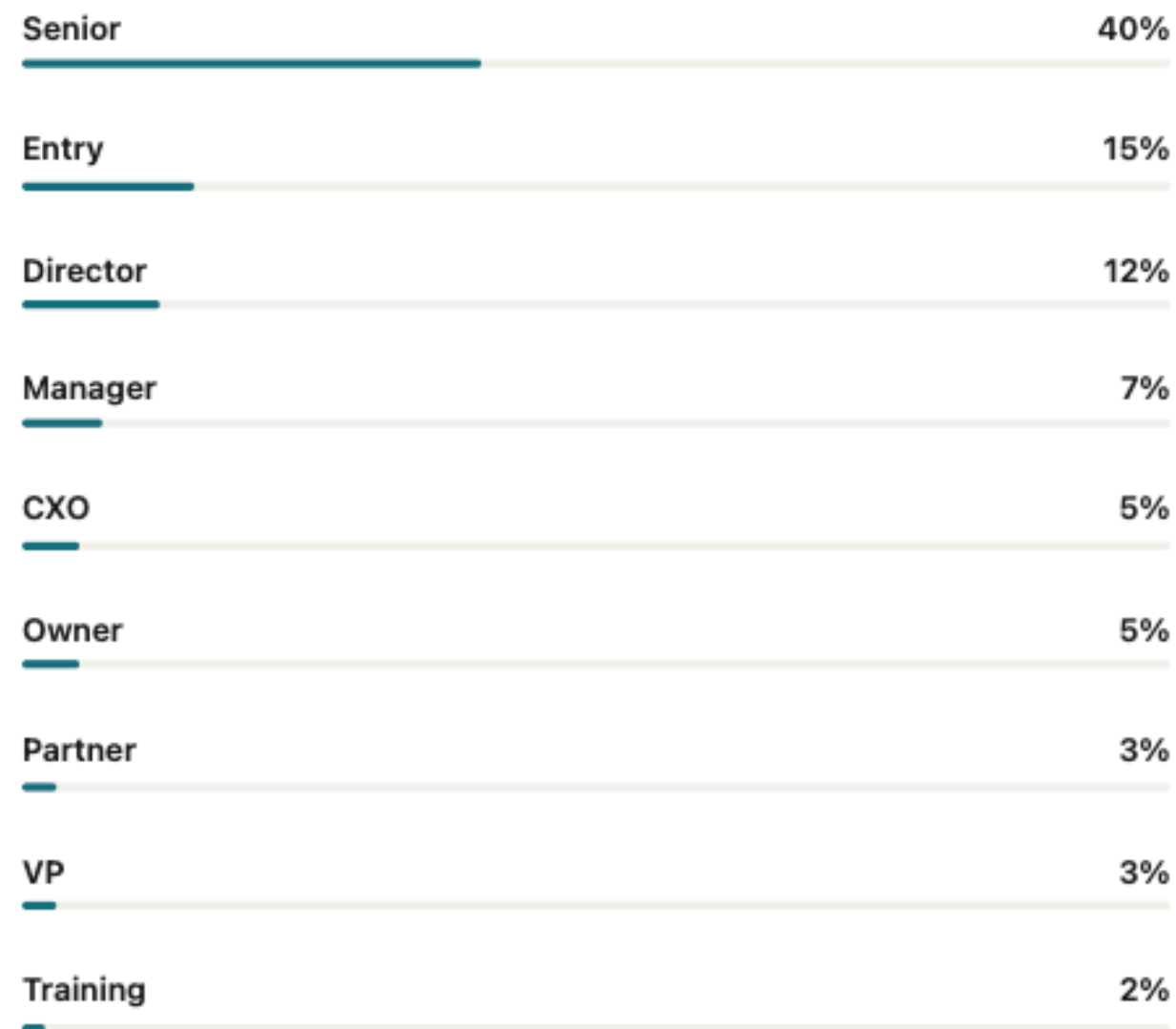
Location

Seniority

Company

Industry

Company size

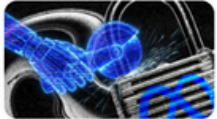


Before

Promote this post to reach people who matter to you. [Boost](#)

 **Ray Eitel-Porter**  · You
Senior Advisor for AI and best-selling author | Helping organizati...
1mo · 

Following on from several outages at AWS caused by AI agent errors, this suggests that AI Governance is not keeping up with the use of AI agents in high stakes situations <https://lnkd.in/ddAeMCu8>



Inside Meta, a Rogue AI Agent Triggers Security Alert
theinformation.com

 Ravi Khan and 12 others 2 comments

 Like  Comment  Repost  Send

 922 impressions [View analytics](#)

Promote this post to reach people who matter to you. [Boost](#)

 **Ray Eitel-Porter**  · You
Senior Advisor for AI and best-selling author | Helping organizati...
1mo · 

Could this be a blueprint for digital sovereignty?



Inside the German state trying to break free from Microsoft
ft.com

 Professor Andy Pardoe and 4 others 1 comment

 Like  Comment  Repost  Send

 275 impressions [View analytics](#)

Promote this post to reach people who matter to you. [Boost](#)

 **Ray Eitel-Porter**  · You
Senior Advisor for AI and best-selling author | Helping organizati...
1mo · 

China has unveiled AI-powered 'wolf robots' equipped with missiles and grenade launchers in military urban combat exercises. Developed by a state-owned research institute, these autonomous robots can ...more



China Deploys Armed AI 'Wolf Robots' in Urban Combat Training
oecd.ai

 Daniela Haskins and 6 others 4 comments · 1 repost

 Like  Comment  Repost  Send

 739 impressions [View analytics](#)

Promote this post to reach people who matter to you. [Boost](#)

 **Ray Eitel-Porter**  · You
Senior Advisor for AI and best-selling author | Helping organizati...
1mo · 

No doubt this will go to appeal but it sends a strong signal that platforms should be more accountable for the consequences of their business model and use of AI



Meta and Google liable for social media harm to children's mental health in landmark US...
ft.com

 Jean-Luc Chatelain and 4 others

 Like  Comment  Repost  Send

 386 impressions [View analytics](#)

 **Ray Eitel-Porter**  · You
Senior Advisor for AI and best-selling author | Helping organizati...
1mo · 

Very interesting analysis and encouraging to have some positive news on this topic



Tim Gordon  · 1st
Partner, Best Practice AI; Trustee at Full Fact; Co-Host 'Age of Inte...
1mo · 

Interesting article in today's [Financial Times](#) by [John Burn-Murdoch](#) arguing that AI chatbots may act against polarisation, given the tendency for the LLM's training data to push them "towards ...more

Could AI chatbots undo the harms of social media?

Financial Times UK
John Burn-Murdoch john.burn-murdoch@ft.com
28 Mar 2026

From the Arab Spring uprisings to the electoral gains of the radical right and left to the rise of anti-science sentiment, the past 15 years have been characterised by populism, polarisation and an erosion of trust in experts, expertise and the establishment. Many factors contribute to these trends, but, as I have previously argued, a key one

 Alicia Ngomo and 5 others 1 comment

 Like  Comment  Repost  Send

 409 impressions [View analytics](#)

After

Top performing posts

28K impressions • 786 engagements

[View analytics](#)



AI has no intrinsic ethics. That was the anchor point of a panel I joined in Oxford on Saturday.

...

11K impressions • 367 engagements

[View analytics](#)



The argument I hear most often against investing in AI governance is cost.

...

7K impressions • 201 engagements

[View analytics](#)



One of the hardest conversations in AI governance at the moment is about explainability.

...

[Show more >](#)

Let's deep dive...



Ray Eitel-Porter  • You
Senior Advisor for AI and best-selling author | Helping organizations use ...
2w • 

AI has no intrinsic ethics. That was the anchor point of a panel I joined in Oxford on Saturday.

The event was part of the public opening of the new [Schwarzman Centre for the Humanities](#) at the [Institute for Ethics in AI](#). Alongside [Edward Harcourt](#), [Kenneth Cukier](#), [Amanda Stent](#) and [Philipp Koralus](#), we brought together philosophy, journalism and business practice in one room.

What matters is the humans building, deploying and using it. And right now, those humans aren't being held accountable enough.

Three risks came up that don't feature enough in mainstream governance conversations.

Cognitive atrophy: as we delegate more thinking to AI, we risk losing the critical skills needed to scrutinise its outputs. That includes the ability to spot where AI decisions go wrong.

Epistemic decay: our collective ability to distinguish real from AI-generated content is eroding. The consequences are already visible in the spread of AI-generated misinformation that platforms and regulators struggle to contain.

Representation and access: the global south and other groups are missing out due to gaps in training data and unequal access. Healthcare AI is a stark example, with models performing poorly for populations that weren't adequately represented in training data.

All three risks share something in common: they're systemic, they compound over time, and voluntary self-governance has not been enough to address them. That's what makes the next question so significant.

The audience was asked whether frontier AI models should face something like pharmaceutical licensing. 80% agreed.

That is not a fringe position.



Is there such a thing as ethical AI?

Panelists:
Ray Eitel-Porter, Senior Advisor for AI, Nova Impact
Edward Harcourt, Director of the Schwarzman Centre for the Humanities at the Institute for Ethics in AI
Kenneth Cukier, Author of 'The Machine Question', University of Oxford
Amanda Stent, Senior Advisor for AI, University of Oxford
Philipp Koralus, Senior Advisor for AI, University of Oxford
Moderator: David Foray, Director of the Institute for Ethics in AI

  516  115  44 

 27,937 impressions  [View analytics](#)

Discovery 

27,938 Impressions

18,688 Members reached

 Members who post once a week can get up to 4x more profile views. Keep the momentum going by creating another post. 

[Start a post](#)

Profile activity 

219 Profile viewers from this post

270 Followers gained from this post 

Engagement 

812 Social engagements

 Reactions **516** 

 Comments **115** 

 Reposts **44** 

 Saves **126**

 Sends on LinkedIn **11**


Ray Eitel-Porter • You
 Senior Advisor for AI and best-selling author | Helping organizations use ...
 4w •

The argument I hear most often against investing in AI governance is cost.

Here's the counterargument, backed by data.

IBM's Institute for Business Value found that 27% of AI efficiency gains stem from strong governance. Companies in the top quartile of AI ethics spending report 30% higher operating profit from AI than those in the lowest quartile. The top three benefits cited by executives: increased trust (61%), stronger brand reputation (57%) and mitigated reputational risk (54%).

Meanwhile, the cost of getting it wrong is rising fast. The OECD AI Incidents Monitor shows the six-month moving average of AI incidents more than doubled from around 200 per month at the start of 2025 to over 500 by February 2026. AI-related litigation is growing. Regulatory fines are coming.

And there's an implementation reality most organisations underestimate: building governance in from the start is dramatically cheaper than retrofitting it afterwards, something we cover in Chapter 2 of Governing the Machine. I've seen AI systems require full retirement because applying controls after the fact was too complex.

The organisations treating AI governance as an overhead haven't done the maths. Those treating it as a strategic investment are pulling ahead.

The cost of inaction is higher than the cost of getting started.




 212
 
 64
 
 19
 

11,015 impressions

[View analytics](#)

Discovery

11,015 Impressions

7,174 Members reached

 Members who post once a week can get up to 4x more profile views. Keep the momentum going by creating another post.

Start a post

Profile activity

113 Profile viewers from this post

70 Followers gained from this post

Engagement

353 Social engagements

 Reactions 212

 Comments 64

 Reposts 19

 Saves 54

 Sends on LinkedIn 4

1 Link engagements

 Visits to links from this post
<https://oecd.ai/en/incidents> 1



Ray Eitel-Porter • You
Senior Advisor for AI and best-selling author | Helping organizations use ...
1w • 

One of the hardest conversations in AI governance at the moment is about explainability.

A large financial services organisation came to me with well-established model risk management and governance processes already in place. The problem was not a lack of governance frameworks. It was that nobody could agree on how AI fitted within them, or who was responsible when an AI system produced an output that nobody could fully explain.

The black box problem is not theoretical in financial services. Regulators expect firms to explain the decisions their systems make. When you cannot, you either accept unquantifiable regulatory risk or you stop using the system. Neither is a satisfying answer.

We worked through it with leaders from risk, legal and data teams together rather than in sequence. The workshops were as much about building shared understanding as they were about governance design. What emerged was a set of clear accountabilities, defined hand-offs between functions, and a practical approach to AI explainability tools and techniques.

Crucially, the organisation also agreed on something harder: there are AI solutions they will not use where explainability cannot be achieved. That is a governance choice, not a technical one.

If your organisation is at that intersection of model risk management and AI governance and is not sure how the two relate, I have worked through this with a number of firms. Happy to have a conversation.





126



34



3



 7,473 impressions

[View analytics](#)

Discovery

7,473 Impressions

4,789 Members reached

 Members who post once a week can get up to 4x more profile views. Keep the momentum going by creating another post.

[Start a post](#)

Profile activity

106 Profile viewers from this post

71 Followers gained from this post 

Engagement

203 Social engagements

 Reactions	127
 Comments	40
 Reposts	3
 Saves	31
 Sends on LinkedIn	2

The engine behind high-performing content

1. Activated Relevant Network

- Encouraged Ray colleagues to engage early, signalling relevance to LinkedIn's algorithm.
- Tailored wording to resonate specifically with people in the AI and Governance space.

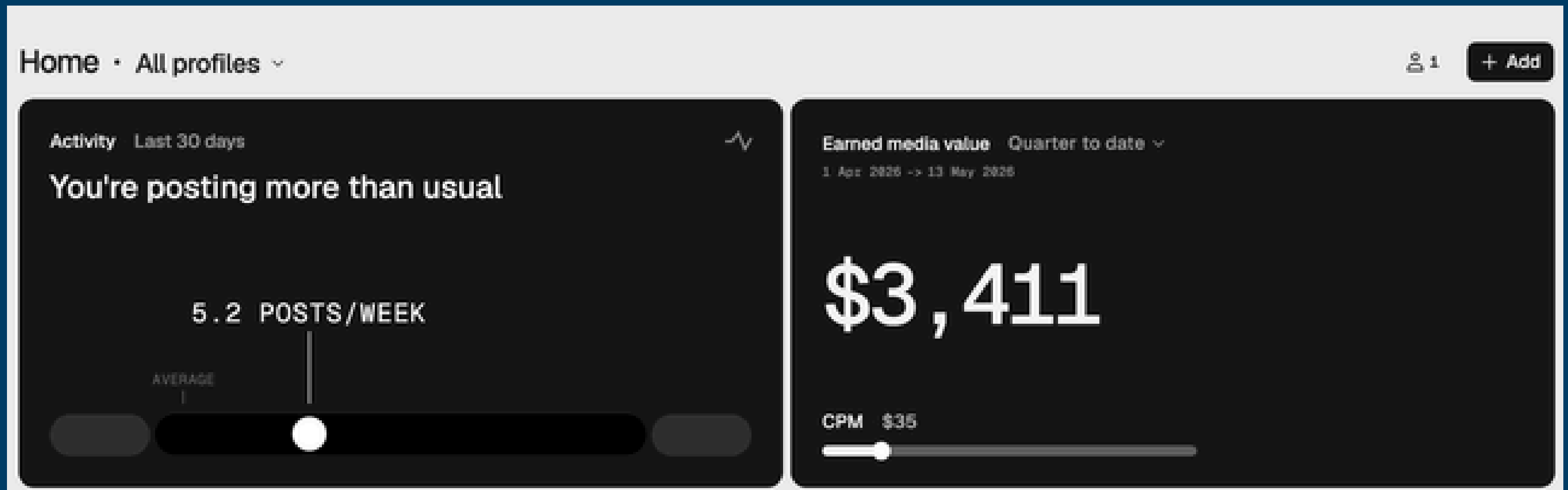
2. Targeted Strategic Individuals to Boost Credibility Signals

- Adding people within your ICP to engage with your content, in your network, will mention your name in boardrooms, buy your book and mention to their peers. Getting your name out there.

3. Optimised the Post for Distribution and Genuine Insight, Not Just Content

- Structured posts to encourage long comments, not short reactions.

Shield Analytics



Our Strategic Approach

Strategic Recommendations (What we set out to do)



01

Visibility & Audience Growth

- +10–15% increase in followers (not inc. people we add)
- +20–30% increase in profile views

Goal: Build steady visibility within your niche



02

Engagement & Content Performance

- +20-30 % increase in average engagement per post
- +30–40% increase in total impressions

Goal: Show consistent improvement in content performance and audience resonance



03

Pipeline & Conversion Benchmark

- 4 meetings booked with strong ICPs within 16 weeks

Goal: Validate LinkedIn as a consistent lead generation channel

Recommendations & Next Steps

Upcoming Goals

Goal 1

Impressions **Grow by 20%**

Impressions in 6 months

Goal 2

Achieve **8 qualified leads**

Qualified leads in 6 months

Goal 3

Grow audience to **8000+**

Followers in 6 months

Strategic Recommendations



Double down on proven content themes



Launch a Substack and email outreach: use Substack for more considered pieces, and email outbound for faster meeting bookings.



Strengthen ICP visibility through targeted outreach & commenting