



CREATOR RISK INTELLIGENCE

# The Signal

Monthly intelligence on Creator Risk, regulatory change, and brand safety incidents. Issue 29 · May 2026

The most dangerous brands in the creator economy are the ones that believe it won't happen to them. This month we examine the risk perception gap — why it exists, why it persists, and what the data from two and a half years of tracking this space consistently shows about the brands most likely to face a significant creator risk incident.

## IN THIS ISSUE

- 01 — THE MOST DANGEROUS BRANDS
- 02 — EU AI ACT LABELLING: ENFORCEMENT BEGINS
- 03 — REGULATORY SNAPSHOT
- 04 — WHAT WE'RE WATCHING
- 05 — OPINION

## 01 — THE MOST DANGEROUS BRANDS

*Risk perception · Industry data · May 2026*

Issue 17 (May 2025) introduced the risk perception gap — the finding that fewer than 5% of brands cite brand safety as a key influencer marketing challenge, while nearly two-thirds have experienced influencer fraud in prior campaigns. A year on, the gap has not closed. The enforcement environment has become significantly more active. The incident frequency has not decreased. And the proportion of brand-side teams treating creator risk as a priority remains low.

The explanation is structural. Brand safety incidents are episodic — they happen to brands, then they stop happening, then they happen again. In the gap between incidents, the organisational memory of the last incident fades, the risk framework built in response to it atrophies, and the process reverts to the pre-incident standard. The next incident then encounters the same gap that produced the previous one.

### The profile of the brand most likely to face a significant creator risk incident



|                                         |                                                                                                                                                                                                                                                                                                                                       |
|-----------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>High volume,<br/>low vetting</b>     | Brands running large creator programmes — fifty or more active creators — with no systematic vetting process are the statistical leaders. The probability of at least one creator in a large portfolio carrying a material risk is not low. It is near-certain. Volume without process is the highest-risk configuration.             |
| <b>Recent<br/>programme<br/>launch</b>  | Brands that have launched creator programmes in the past twelve to eighteen months have not yet experienced a significant incident. The organisational assumption that this means their process is working is the most common precursor to a first incident. The absence of a problem is not evidence that the process is adequate.   |
| <b>Values-led<br/>positioning</b>       | Brands whose commercial positioning is built on specific values — sustainability, inclusivity, authenticity, social good — face elevated risk when their creators' public conduct or private behaviour contradicts those values. The Djerf case. The Remi Bader case. The contrast between the claim and the reality is the incident. |
| <b>No monitoring<br/>infrastructure</b> | Brands that assess creators at activation and then cease monitoring until renewal are operating with no visibility of what develops in between. The Gonzalez signals were visible eight months before the arrest. The monitoring process would have caught them. The activation-only process did not.                                 |

#### THE SIGNAL

The risk perception gap is self-reinforcing. Brands that have not had a significant creator risk incident believe their process is working. The process may be adequate — or the incident may simply not have occurred yet. Those are different situations that look identical from the inside.

#### RECOMMENDED ACTION

Portfolio audit. Not creator-by-creator at activation. The whole portfolio, systematically, against the risk categories established across this publication. That is the assessment that surfaces what individual activation vetting misses.



## 02 — EU AI ACT LABELLING: ENFORCEMENT BEGINS

*EU AI Act · AI-generated content · May 2026*

The EU AI Act's content labelling requirements have been in force since late 2024. In May 2026, the Commission confirmed that enforcement actions for non-compliance are actively under way in two member states. The first formal enforcement decisions are expected before the end of 2026.

The cases under investigation involve AI-generated influencer content used in commercial contexts without the required labelling — a category that Signal Intelligence first identified as a risk vector in Issue 16 (April 2025) and that the Mia Zelu case at Wimbledon illustrated in Issue 19 (July 2025). The brands involved are not technology companies. They are consumer brands that commissioned AI-generated creator content and did not ensure the labelling obligation was met.

The Commission's position — confirmed in its August 2025 guidance and now reflected in the enforcement activity — is that platform-level AI detection labels do not satisfy the brand's disclosure obligation. The obligation runs to the brand. The platform label is a floor, not a ceiling, and not a substitute.

### THE SIGNAL

The EU AI Act enforcement activity confirms what the Commission stated in August 2025: platform labels are insufficient. Brand-level disclosure obligations exist independently of what the platform detects. Brands that have been relying on Meta or TikTok's AI labels as their compliance mechanism are exposed. The enforcement decisions expected before year end will establish the precedent.

### RECOMMENDED ACTION

Audit your current creator programme for AI-generated content — including AI-edited imagery, synthetic voices, and AI-generated testimonials. Confirm that EU AI Act labelling is being applied at the brand level, not deferred to platform detection. The enforcement decisions coming in 2026 will make this a priority regardless.



## 03 — REGULATORY SNAPSHOT

*Six markets · May 2026*

|                |                                                                                                                                                                                                                     |
|----------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>FRANCE</b>  | ARPP enforcement active. Audio and podcast content now in scope. H1 2026 brand-level focus continuing. Maximum fine €300,000. Speech-to-text operational.                                                           |
| <b>GERMANY</b> | UWG private enforcement at highest recorded level in Signal Intelligence's tracking period. Competitor-driven disclosure enforcement now the dominant risk mechanism for German market brands.                      |
| <b>ITALY</b>   | Ferragni Law and AGCOM Code fully operational. Post-acquittal landscape stable. Italy's framework is the most mature in Europe and the template for the Digital Fairness Act development.                           |
| <b>UK</b>      | First DMCCA enforcement actions under active investigation — decisions expected H2 2026. FCA finfluencer second round enforcement confirmed. ASA named non-compliers list at highest count since launch.            |
| <b>EU-WIDE</b> | EU AI Act enforcement active in two member states — first formal decisions expected H2 2026. Digital Fairness Act legislative proposal confirmed Q3-Q4 2026. DSA active.                                            |
| <b>USA</b>     | First major class action authenticity misrepresentation judgement expected imminently. FTC enforcement maintained. FINRA and FCA financial services enforcement active. Dual liability standard firmly established. |

## 04 — WHAT WE'RE WATCHING

*Forward signals · May 2026*

|                                         |                                                                                                                                                                                                                          |
|-----------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Digital Fairness Act</b>             | Legislative proposal now confirmed for Q3-Q4 2026. Signal Intelligence will publish a full analysis of the draft on release. This is the most significant regulatory development in EU creator marketing since the GDPR. |
| <b>First DMCCA enforcement decision</b> | The CMA's first published decision against a brand under the Digital Markets, Competition and Consumers Act will establish the enforcement standard for UK creator programmes. Expected H2 2026.                         |



**EU AI Act first  
enforcement  
decisions**

The first formal enforcement decisions in the EU AI Act cases currently under investigation will establish the penalties and disclosure standards applicable to AI-generated creator content across 27 member states.



## 05 — OPINION

*The gap between perceived and actual risk is where the next incident is being built*

Two and a half years of tracking creator risk incidents has produced one consistent finding: the gap between perceived and actual risk is where incidents originate. Not in the brands that know they have a problem. In the brands that don't know yet.

The ABC / Taylor Frankie Paul case cost \$55 million. The public record had been available since 2023. The Doritos / Hudson case ended a live campaign. The tweets had been live and searchable the day the brief was signed. The Ferragni case generated €3.4 million in civil liabilities and reshaped a continent's regulatory architecture. The commercial structure was visible to anyone who asked what the hospital actually received.

In none of these cases was the information unavailable. In all of them, the process did not look for it.

The next incident is being built right now in a creator programme where the risk perception is low, the last activation went well, and nobody has reviewed the full portfolio against what two and a half years of incident data has taught us about where the signals actually are.

### PORTFOLIO AUDIT

Signal Intelligence provides independent creator risk assessment and portfolio audits across all major markets. The cases in this issue — the risk perception gap, AI Act enforcement, the profile of brands most likely to face a significant incident — point toward the same question: when did you last look at your full creator portfolio systematically? Request a portfolio audit at [signalintelligence.consulting](https://signalintelligence.consulting).