

AI and Competition Law: The Assumptions That No Longer Hold

Somewhere in the pricing systems of three competing firms, a price is rising. No executive proposed it, no email records it, and no one could be prosecuted for it, because the one thing cartel law is designed to find, the agreement, was never made. Elsewhere, an AI agent books a flight for a customer at a price quietly calibrated to what that person is willing to pay. And somewhere at the frontier of the technology itself, a lead may compound faster than any competitors, courts, or regulators can respond.

Three different scenarios with one property in common, and it is not that they involve artificial intelligence. It is that competition law, as currently written, cannot address them and almost nothing else in the crowded field called “AI antitrust” shares that property.

Two narratives currently dominate the field, and both are wrong in ways that cost money and cases. The first is alarm: AI is unprecedented, every development is a crisis, the system must be rebuilt. The second is complacency: we have seen it all before, a cartel is a cartel whatever runs it, the tools will cope as they always have. The alarmist buys intervention against harms the law already addresses. The complacent buys silence about the few it cannot. And both approaches make the same underlying mistake. They ask whether AI is anticompetitive, which is the wrong question. The right question is which of competition law’s assumptions AI dissolves.

Competition law rests on four assumptions so deep they are never written down. The law of cartels assumes that collusion involves an **agreement**: a meeting of minds an enforcer can find and prove. The law of monopolisation assumes an **author**: someone whose intent can be attributed and condemned. The consumer-welfare standard assumes a **consumer**: an individual who compares, switches, and walks away, and the readiness to walk away is what disciplines sellers. And the entire architecture of after-the-fact enforcement assumes **time**: when the case concludes, the harm will still be there to fix. None of these appear in any statute, because for a century none needed to. They were simply how the world worked.

AI rarely creates new anticompetitive conduct. Instead, what it does, in a small number of cases, is dissolve the assumptions beneath competition law.

Most of what dominates the headlines is the familiar toolkit doing familiar work at higher speed, and it is reachable precisely because the assumptions hold. RealPage is the defining example: pricing software recommending rents on pooled, confidential competitor data. The cases proceeded on a hub-and-spoke theory and resulted in a settlement addressing the features that made the conduct actionable; the pooled data, the granular recommendations and pressure to comply. There was a hub; there was an exchange; the categories fit. A faster cartel is still a cartel. Self-preferencing executed by a ranking model is still self-preferencing. This is old wine, and the honest word for the law’s position against it is *adequate*: strained in capacity, perhaps, but not in category.

The edge of that territory is now being mapped by the appellate courts. In *Gibson v. Cendyn Group*, the Ninth Circuit affirmed the dismissal of claims against Las Vegas hotels using common revenue-management software, holding that their individual licensing agreements did not restrain trade in the aggregate. The plaintiffs had abandoned their hub-and-spoke theory on appeal, so the court never reached it, and the court noted they had not alleged that the software fed one hotel’s confidential data into the recommendations generated for another. In July 2026 the Third Circuit went the other way in *Cornish-Adebiyi v. Caesars Entertainment*, reviving hub-and-spoke claims on the same software where the pooling of non-public competitor data was alleged. Read together, the two decisions mark where the doctrine’s grip depends on something more than shared software.

Past that edge sit three harms that are different in kind. (1) **Coordination without agreement**: independent pricing systems learn to sustain supra-competitive prices with no meeting of minds and no human author, dissolving the first two assumptions at once. (2) **The disloyal agent**: software that represents the buyer while monetising the buyer’s own information for the other side, dissolving the consumer whose walking-away was meant to keep markets honest. And (3) **frontier irreversibility**:

the possibility (so far modelled, not observed) a lead becomes permanent before any enforcement can act, dissolving the assumption of time.

This series will argue that your AI-competition worry is probably mis-sorted. Risk registers, litigation dockets, and policy workstreams are ordered by headline volume, and headline volume tracks the familiar harms the law can already address. The genuinely novel harms generate almost no headlines, because their defining feature is that they never surface. The autonomous cartel produces no leniency applicant, because there is no confederacy to defect from. The disloyal agent's extraction produces no complainant, because the victim cannot perceive the injury. The loudest problems are mostly the solved ones. The dangerous ones are quiet.

Over the coming weeks this series will walk through the argument step by step:

- Where power in the AI stack actually concentrates
- Whether the famous data flywheel exists at all
- The three novel harms in detail including which element fails and what it costs at the pleading stage
- A tool-by-tool verdict on the existing kit

One promise throughout: where a harm is modelled rather than observed, I will say so. I will argue that the frontier is not tipping, and that you should watch it anyway.

Food for thought. *For in-house counsel:* of the AI items on your risk register, which are familiar harms amplified and which, if any, would your regulator be structurally unable to see? *For practitioners:* when did you last read a theory of harm that quietly assumed an element that was not there? *For policymakers:* is your AI-competition workstream sorted by novelty, or by headline volume?

Join the conversation. If you think I'm wrong, please respond in the social media comments as your feedback will shape where this series goes.