

Working Papers

No 4

August 2026

Value Coverage: a measurement framework for technology investment booms, with a live application to the AI capital build-out

by Allan Pedersen

JEL classification: G12, G17, G31, G01, O33, O47, E22

Keywords: artificial intelligence; technology investment booms; asset bubbles; user cost of capital; quasi-rents; social returns to innovation; financing fragility; Tobin's q ; reproducible measurement

Philosophers Mint Working Papers

Philosophers Mint Working Papers circulate research in progress for discussion and comment. They have not been peer-reviewed. Views and remaining errors are the author's alone.

This paper and earlier papers in the series can be downloaded free of charge from philosophersmint.com, from SSRN (SSRN author page 11401108) and via RePEc (RePEc:pmt:wpaper:4).

Suggested citation

Pedersen, Allan (2026), "Value Coverage: A Measurement Framework for Technology Investment Booms", Philosophers Mint Working Paper No. 4, August 2026. Available at philosophersmint.com.

© 2026 Allan Pedersen. Reproduction of brief excerpts is permitted for scholarly and non-commercial purposes, provided the source is cited.

Comments welcome: allan@philosophersmint.com

Value Coverage: a measurement framework for technology investment booms, with a live application to the AI capital build-out

Allan Pedersen

Independent researcher · allan@philosophersmint.com

Abstract

Is the AI technology boom a bubble? The question is being argued almost entirely by anecdote, one side pointing to the technology potential and the other to the spending. This paper offers a way to keep appropriate accounts rather than a direct answer to the question. It proposes that any technology investment boom can be described by three quantities, each measurable from public data. The first is *private coverage* (C_p): the earnings the capital throws off, set against the annual cost of owning it. The second is *social coverage* (C_s): the value the technology creates for the wider economy, net of its harms, against the same cost. The third is *financing fragility* (F): how much of the build is paid for with outside money, and where in the structure that borrowing sits. The paper defines the three, sets out two hypotheses, one for the *probability* of a correction and one for its *cost*, in a form that can be estimated, places past booms on a simple map, and reports a live, open implementation for the AI build-out. On data to August 2026 the headline (capital-holder) readings are $C_p \approx 0.29$, $C_s \approx 0.81$ at a ten-year horizon and ≈ 1.34 at thirty, and $F \approx 0.78$, across the ponzi threshold on a pre-committed count of reciprocal financing deals; a second, ecosystem-wide coverage line reads ≈ 0.42 , a wedge of ≈ 0.13 above the capital-holders. What is offered is a measurement programme with a working instrument, and a not-yet-tested theory of bubbles. Two steps would make the central claim testable: an expectations benchmark that says what coverage the market is pricing in, and a historical backtest against the price-based crash predictor of Greenwood, Shleifer and You. Both are specified. The paper is equally clear about what such measurement cannot do. It can deliver consistency rather than accuracy, and a characterised boom rather than a prophecy.

JEL classification: G12, G17, G31, G01, O33, O47, E22

Keywords: artificial intelligence; technology investment booms; asset bubbles; user cost of capital; quasi-rents; social returns to innovation; financing fragility; Tobin's q ; reproducible measurement

This version: v1.5, 11 August 2026 — readings synced to the August 2026 dashboard recalibration (numerator catch-up, capex-share and adoption recentres, OEWS May-2025 base; fragility across the ponzi threshold); §7 extended with the mid-2026 literature. v1.0: 15 June 2026; v1.4: 16 July 2026.

1 The question and the claim

When commentators argue about the AI boom, they argue about one number: how much is being spent. The spending is indeed enormous, and on that much both camps agree. But a bubble is not heavy spending. It is spending that outruns the value produced, financed in a way that cannot absorb the disappointment. That sentence already holds three separate questions. Spending against value is a question of coverage. The financing is a question of fragility. And “value” itself splits in two: what the owners of the capital can charge for, and what society gains whether anyone charges for it or not.

The argument of this paper is that a boom is better described by these three quantities than by the one the debate fixes on:

1. **Private value coverage (C_p).** The quasi-rents the capital actually earns, set against the cost of holding it. This is the main driver of the *probability* of a financial correction, because an asset price can stand only on value that someone pays for.
2. **Social value coverage (C_s).** Total surplus, net of harms, against the same cost. This is the main driver of the *cost* of a correction: whether a bust is a tragedy or merely a transfer from late investors to everyone else.
3. **Financing fragility (F).** How much of the investment is funded outside operating cash flow, and where that borrowing sits. F enters *both* outcomes. It shapes the violence and timing of a correction, and it can cause one on its own.

A common intuition holds that social value should make a bubble less likely. It does not, and the history on this is plain. Britain’s railways and America’s fibre-optic cables delivered great social value, and their investors were ruined anyway. Consumer surplus cannot service debt. Social value weighs mainly on the cost side, and that separation is what makes the framework useful.

One caveat: the cleanest version of the claim, with each variable driving exactly one outcome, is false. The Insull collapse of 1932 was a high-coverage, high-fragility crash. Electricity’s value was real and growing, yet the pyramided holding companies that financed it still produced one of the defining failures of the Depression. The leveraged-bubbles evidence of Jordà, Schularick and Taylor shows that what an asset boom costs the economy depends at least as much on how it was financed as on what the asset was worth. So the assignment of variables to outcomes is a claim about *weights*, to be estimated, not a clean separation to be asserted. The quantities are still worth measuring. That is the programme for this paper.

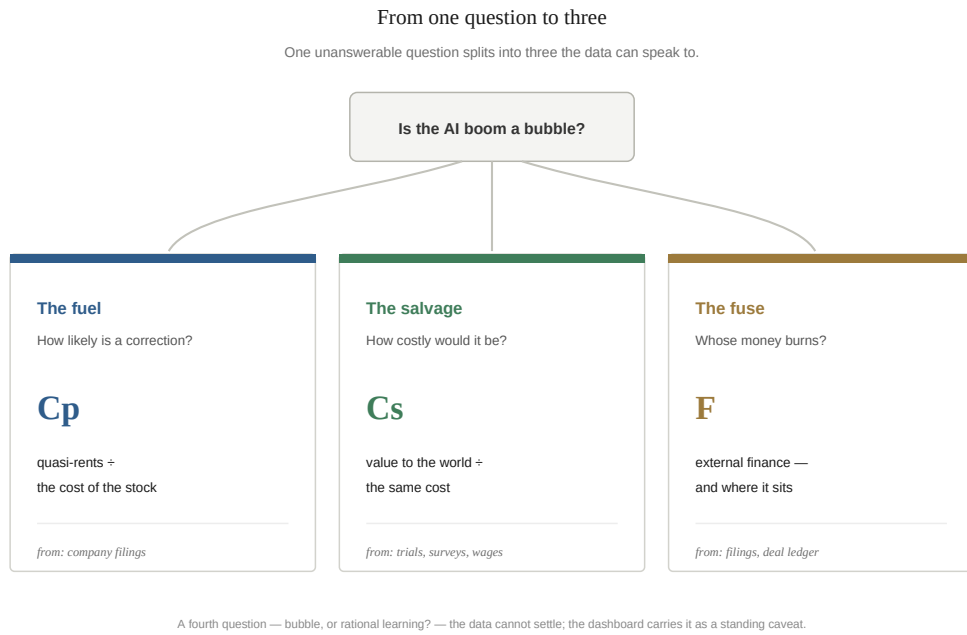


Figure 1. The bubble question decomposes into three measurable quantities, each with its data source.

2 Definitions

The capital and its cost. Let K be the net capital stock committed to the technology, defined over an explicit list of asset classes (chips, servers, buildings, power and cooling, networks) with the boundary rules published. The annual flow needed to justify holding it is the sum of the asset-class user costs, $\sum_i (r_i + \delta_i) K_i$ (Jorgenson 1963). Each class carries its own depreciation rate and required return; chips depreciate at perhaps 18 to 40 per cent a year (a live controversy), buildings at 2 to 3 per cent. A single blended δ is here used as a practical shorthand. Two key measures ground this denominator. First, where a rental market exists, use it: GPU rental rates pin down $(r_i + \delta_i)$ better than any accounting argument. Second, r is a state variable, not a fixed parameter. It moves with monetary policy, and a tightening is the proximate trigger in most of the historical panel (1845, 1929, 2000). The dashboard therefore publishes coverage at a declared band of r , not at a point.

V_p , **private value, as quasi-rents.** What must cover the user cost is not revenue but quasi-rents: revenue attributable to the capital, net of the variable costs of operating it (energy, labour, inference compute, sales). As such, a data centre with revenues equal to its user cost and a thirty-per-cent gross margin may easily be losing money. Three rules follow. *Consolidation:* the chipmaker's revenue is the hyperscaler's capex, and the cloud's AI revenue is the lab's cost, so V_p is measured once, at the boundary where value leaves the capital-holding system and is paid for by an outside party. *Indirect monetisation counts:* advertising uplift, bundle pricing, internal productivity and rental income are all quasi-rents to the capital, and a numerator that omits them would be wrong. *Incidence:* aggregate V_p can be healthy while the firms that hold the capital capture little of it. Value migrat-

ing to the application layer is good for C_s and irrelevant to the capital-holders' solvency. Crash risk attaches to the coverage of those who hold the capital, so C_p is reported for the capital-holders, with the ecosystem total alongside.

V_s , **social value, signed.** $V_s = V_p + \text{consumer surplus} + \text{productivity spillovers} - \text{negative externalities}$ (displacement and adjustment costs, energy-price effects on ratepayers, security and information harms). The minus sign matters twice. It makes the rent-extraction quadrant of §4 algebraically possible, and it is where most of the political economy lives. Task-level coefficients from the randomised-trial literature (Noy & Zhang 2023; Brynjolfsson, Li & Raymond 2025; Dell'Acqua et al. 2023/2026; Cui et al. 2026), scaled by adoption, do *not* give a floor, because the tasks were selected in the expectation that AI would help. V_s is therefore reported as an estimate with bias in both directions, gross of harms, at declared horizons. C_s has no meaning without a horizon and a discount rate. Over a century almost any durable infrastructure clears $C_s \geq 1$; over five years, the dark fibre of 1999 did not. The dashboard reports C_s at 10 and 30 years, both sides discounted to present value at the same rate.

The coverage ratios. $C_p = V_p / \sum_i (r_i + \delta_i) K_i$ and $C_s = V_s / \sum_i (r_i + \delta_i) K_i$, each with the natural anchor at 1 and bands that reflect measurement error. C_p is sort of a flow cousin of Tobin's q , and the wedge between a high q and a low C_p is *not* the bubble. Most of it is the legitimate capitalised value of expected growth, and a young technology should show q well above 1 with C_p below 1. The bubble is the part of the wedge that the realised path of V_p keeps failing to validate. Pinning that down needs an explicit expectations benchmark (§3).

F , **fragility, located.** Operationally F is built from (capex – operating cash flow)/capex for the capital-holders, but this aggregate conceals a key feature, which is where the leverage sits. Hyperscaler operating cash flow comes overwhelmingly from non-AI businesses, so a firm-level ratio records the AI build as internally financed because search advertising pays for it. That is cross-subsidy, not self-financing. The genuinely fragile finance concentrates in pure-play neoclouds, special-purpose vehicles and lease structures engineered to stay out of the ratio. F is therefore reported decomposed by balance-sheet location, and the headline takes the weakest link, the edge, not the average. Circular (reciprocal) capital is flagged as it is identified, and not summed. F is a sort of proxy for the situation's position on Minsky's hedge-to-speculative-to-Ponzi spectrum.

3 Two hypotheses

These are hypotheses arrived at in the course of the working-paper research. They have not yet been tested, but they are set out here because they provide a good starting point for thinking about the dynamics surrounding a bubble:

H1, correction probability. In a panel of boom episodes,

$$\text{logit } P(\text{correction}) = \beta_0 + \beta_1 \text{gap} + \beta_2 F + \beta_3 (\text{gap} \times F) + \beta_4 (dC_p/dt) + \beta_5 \Delta r$$

where $\text{gap} = 1 - C_p$ and the trajectory term dC_p/dt is where the information lives. A coverage ratio of 0.03 that doubles yearly is a young technology; the same ratio static over

the years looks more like a mispricing. The interaction β_3 is the heart of the claim: a gap funded by patient internal cash produces slow repricing, while the same gap funded by fragile external finance produces a crash. β_2 is free to be non-zero because fragility alone can do it, as at Insull; Δr carries the historical trigger. The **expectations benchmark** closes the test: from current market valuations of the capital-holders, back out the implied path of future quasi-rents, and read the bubble component as the persistent shortfall of realised V_p growth against that implied path. Deriving the benchmark, a two-period model with Bayesian learning about productivity (Pástor & Veronesi 2009) and an external-finance premium (Bernanke, Gertler & Gilchrist 1999), is this paper's central unfinished task.

H2, correction cost.

$$\text{Cost} = \gamma_0 + \gamma_1 (1 - C_s) + \gamma_2 F + \gamma_3 (\delta\text{-weighted share of } K)$$

High social coverage means the surplus keeps flowing after the equity is gone, the Perez consolation. γ_2 is there because leveraged busts are macroeconomically expensive whatever the asset was worth (Jordà, Schularick & Taylor 2015). γ_3 is there because the consolation is bought with durability: the railways' residue lasted a century because δ was tiny, whereas AI's residue splits between durable power, grid and buildings on one side and fast-obsolescing silicon on the other, with an unusually large share in the high- δ pile.

4 The map

Crossing the two coverage ratios gives four regimes, and with negative externalities in V_s all four are populated.



Figure 2. The two-by-two of coverage regimes; financing fragility is the third dimension, shown as colour.

- **Productive bubble** (C_p low, C_s high): investors lose, society wins. Railways in the 1840s; fibre in 1999; the British bicycle mania of 1896 (Quinn 2018).
- **Healthy boom** (C_p high, C_s high): value keeps pace. The cloud build-out of the 2010s.
- **Pure mania** (C_p low, C_s low): real capital, little lasting worth, as in the crypto-mining build-out of 2021 and 2022 (real ASICs, real megawatts, high δ , thin surplus).
- **Rent extraction** (C_p high, C_s low): quasi-rents extracted against negative externalities, through monopoly tolls and lock-in. This is not a bubble but a different pathology.

A few notes on the scope of the map: *Tulips are folklore, not data*: bulbs have no capital stock, no δ and no user cost, so the episode is retired from the panel. *Scope*: this is a framework for technology booms, capital with a δ and a use, and it is silent on bubbles without one, such as land and housing. *Paths, not points*: electrification spent decades in the productive-bubble quadrant before maturing rightward, and then, at the moment its coverage was strongest, the Insull pyramid delivered a financing crash anyway. The map shows where a technology is; F decides what happens to the people who financed the journey.

5 Three rival explanations for the same data

A stagnant C_p does not, by itself, diagnose a bubble. At least three non-bubble worlds generate the same observable situation, and the programme should actually say how it would tell them apart. This is its hardest open problem, and it is not yet solved in this framework.

Rational learning (Pástor & Veronesi 2009). With genuine uncertainty about a new technology's productivity, high and volatile valuations are rational in advance; the apparent bubble is the discount rate rising as the technology's risk turns from idiosyncratic to systematic, and it is visible only after the fact. Distinguishing mark: revealed value should respond to capability news.

Competition eroding appropriability. Open-weight models and token prices collapsing towards marginal cost can crush V_p while capability and C_s climb, much as in the American railway busts of the 1880s and 1890s. Distinguishing mark: falling margins alongside rising usage and surplus, and a widening wedge between ecosystem V_p and capital-holder V_p , which the framework reads by reporting both lines (at August 2026 the ecosystem line ≈ 0.42 against the capital-holder ≈ 0.29 , a wedge of ≈ 0.13 — narrowed from ≈ 0.23 in July as the capital-holders' own disclosures caught up, a movement this margin exists to read).

Defensive capex as option value. If the hyperscalers are buying insurance against the disruption of their existing cash flows, the spend can be rational even if AI-attributable quasi-rents never cover the user cost. Distinguishing mark: spending concentrated among incumbents with threatened franchises and weakly tied to revenue targets. This is the strongest single objection to reading the coverage gap as a bubble, and the dashboard carries it as a standing alternative rather than pretending to refute it.

6 The dashboard

The framework is implemented as a live dashboard (returns.priceofthinking.com), built to be doubted. Every published number traces to a public primary source, the contestable assumptions live in one versioned file, and any change to them republishes the entire back-series, so a position can never move because the method moved. Data refresh weekly from filings; assumptions are code-reviewed through pull request with a citation-enforcing validator; the whole series is archived quarterly under a DOI.

6.1 Grounding the soft parameters. Four parameters that a first cut could only assign by judgment have been replaced with measured data: occupational AI-exposure, the productivity effect per exposed task, adoption breadth, and the AI share of hyperscaler capital. Occupational exposure now comes from the Anthropic Economic Index's distribution of Claude conversations across the 22 SOC major groups — computed reproducibly by aggregating the ~3,500 usage-weighted O*NET tasks to their SOC codes (rather than read from the published chart, which it matches to two significant figures) — in which software and quantitative work took 37 per cent of conversations at the February-2025 release; it is blended 50/50 with the prior to hedge the consumer-traffic lens (the Index's June-2026 release shows usage broadening — software's share falling toward a fifth — and a re-weighting to it is scheduled). The productivity coefficient ρ is a usage-weighted composite

of the four trials (0.26 to 0.30, reported as a band, since the trials measure different things and one, Dell’Acqua, carries a 19-point penalty outside the capability frontier). Adoption is read from the Census Business Trends and Outlook Survey across its November-2025 question break, which moved the measured rate from about 10 per cent on the old “production” wording to 21.5 per cent on the broadened one (late July 2026); the band’s high corner now carries the Census Bureau’s own employment-weighted estimate, 32 per cent, since a wage-bill multiplier arguably wants employment weighting (Bonney et al. 2026). The AI share of hyperscaler capital sits at a band of 0.55 to 0.90, recentred twice — most recently in August 2026 on Dell’Oro’s print that roughly 75 per cent of hyperscaler capex is now AI infrastructure.

Two cautions are owed to the reader. The groundings moved every read *downward*, but the move (about a factor of two) is the same size as the declared uncertainty, and the four are not independent in sign: trimming the exposed wage base shrinks the numerator while raising the AI-capex share inflates the denominator. So “they all moved it down” is partly an accounting of the parameters, not a discovery about the world. What the grounding fairly establishes is that the answer hangs above all on the boundary of “AI capital”, and that a reasonable first cut sat at the optimistic end of a wide band. The August-2026 refresh is the counter-illustration: the revenue numerator and the capital denominator rose together as disclosures and capex shares caught up, C_p barely moved, and C_s rose — the reads are not on a one-way ratchet.

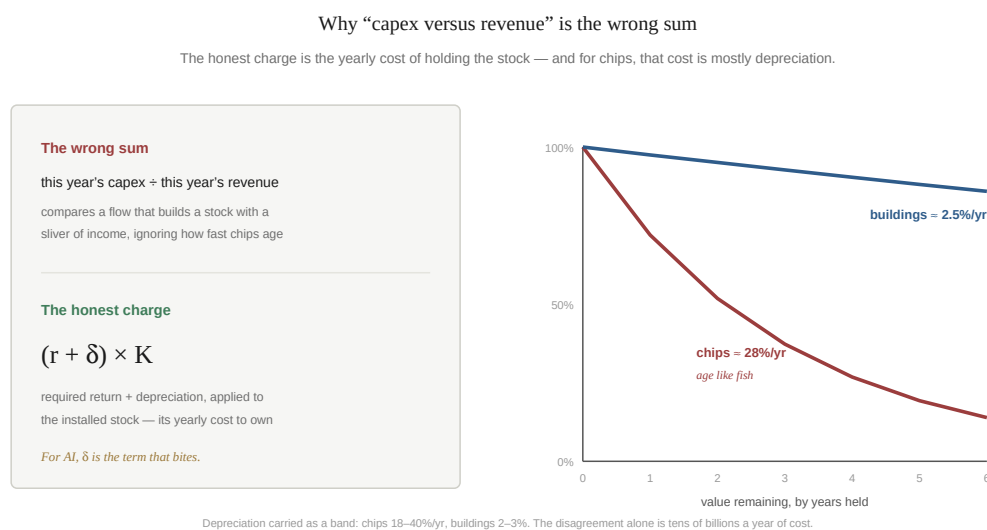


Figure 3. The denominator is the annual cost of holding the stock, $(r + \delta) K$; for chips, depreciation dominates.

6.2 Current readings (August 2026). After the August-2026 revisions below, the capital-holder $C_p \approx 0.29$ (earnings cover about three-tenths of the AI capital cost, on a wide band); $C_s \approx 0.81$ at ten years and ≈ 1.34 at thirty; $F \approx 0.78$, across the **0.75 ponzi threshold**. The **ecosystem** coverage line reads ≈ 0.42 , a **wedge of ≈ 0.13** above the capital-holders. Four revisions moved the readings since v1.4. First, the deal ledger absorbed three weeks of July financing: eighteen deals became twenty-nine and the reciprocal count rose from four to seven, as chip vendors put equity into their own customers at

scale (AMD into Anthropic, up to \$5bn; Nvidia into SSI, $\approx$ \$5bn, and into NAVER, \$1bn); the pre-committed circularity term carried F from 0.72 across the line. The July calibration had been set so that red would need roughly six reciprocal deals; the data delivered seven, so the crossing is a threshold being met, not a threshold being moved. A magnitude check keeps the count honest: reciprocal commitments stand at $\approx$ \$87bn ($\approx$ \$46bn drawn), 0.23 (0.12) of trailing-year AI capex, so the count-based term is currently the *conservative* side of a magnitude-based equivalent (A.5). Second, the revenue numerator caught up with the calendar-Q2 disclosure wave: AWS's stated AI run-rate moved from \$15bn to "over \$25 billion" (30 July), Google Cloud moved from a year-old annual segment figure to a self-refreshing trailing-four-quarter base (\$77.6bn through Q2), and Oracle's AI share was raised to 18 per cent on FY2026 evidence (OCI +77 per cent; an RPO of \$638bn described on the call as mostly large-scale AI contracts); the numerator rose from $\approx$ \$81bn to $\approx$ \$100bn. Third, the denominator rose in the same stroke: the AI share of hyperscaler capex was recentred to 0.75 on Dell'Oro's mid-2026 print, lifting the user cost from $\approx$ \$148bn to $\approx$ \$170bn — both sides of the ratio are now larger and better grounded, and C_p settled at $\approx$ 0.29, barely moved, which is itself informative. Fourth, the social-coverage inputs were refreshed against the government's own measurements — Census adoption at 21.5 per cent on the broadened wording, with the employment-weighted 32 per cent as the high corner, and the wage and employment base moved to the May-2025 OEWS tables — lifting C_s from $\approx$ 0.68 to $\approx$ 0.81 at ten years. The July chip-depreciation recentring (mid 0.28 to 0.20) stands. The firmest finding still rests on filings alone: trailing-year capex over operating cash flow has crossed 100 per cent at two of the six largest spenders, and the leverage concentrates at the neocloud and SPV edge — while, independently, Exponential View dates the first quarterly clearing of infrastructure *depreciation* to Q4 2025 / Q1 2026; the distance between that milestone and $C_p \approx 0.29$ is the required return on capital, not a data dispute (§7).

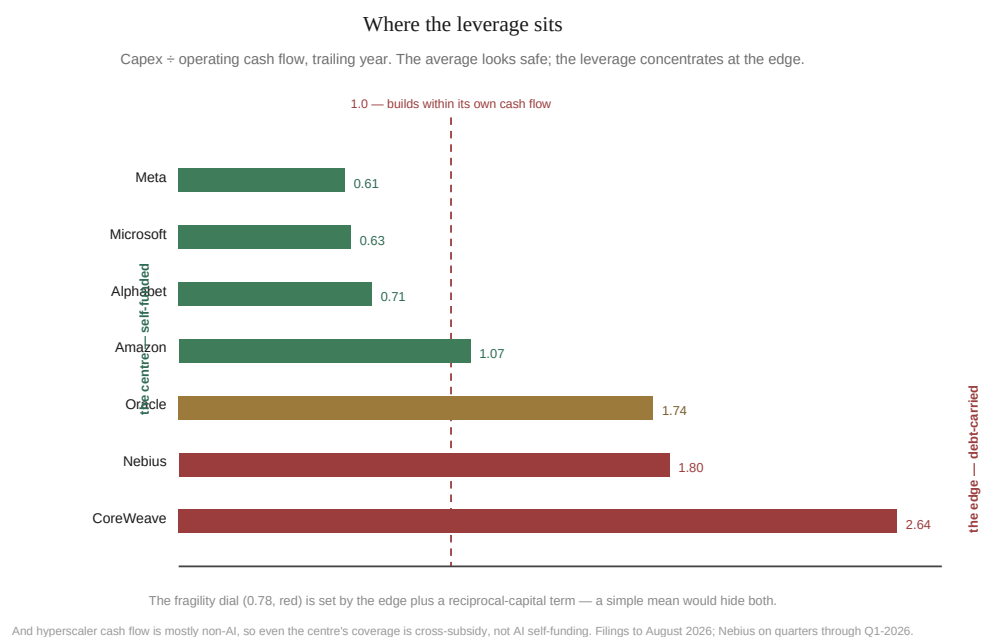


Figure 4. Capex over operating cash flow, by firm: the cash-funded centre against the debt-carried edge.

6.3 The trajectory. Because the whole claim is that levels are unreliable and the movement is the signal, each number is rebuilt at past quarter-ends back to 2023. Two caveats apply to the trajectory panel. The social-coverage path *descends* over 2023 to 2026, but largely because its numerator is a forward projection held fixed while the capital-stock denominator compounds; that is a denominator effect, not measured social value falling. And the historical booms are placed after the fact, from realised outcomes, against AI's real-time and repeatedly down-revised nowcast, so the comparison is orientation, not evidence, until the backtest (§8) reconstructs the past on the same contemporaneous method.

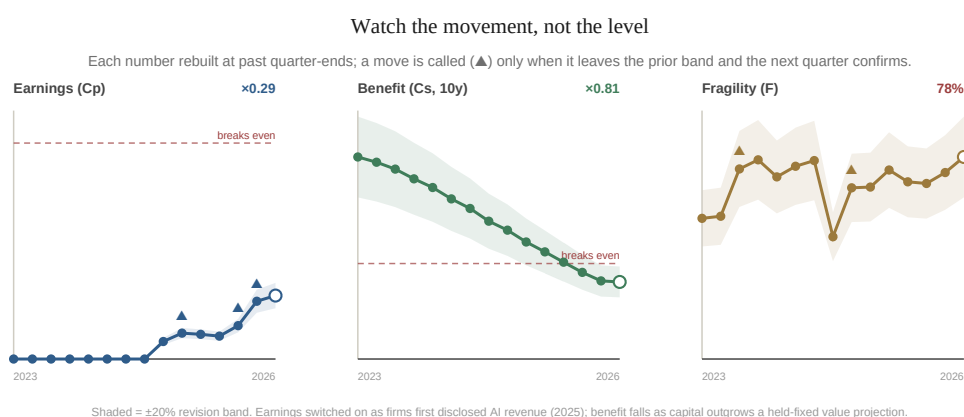
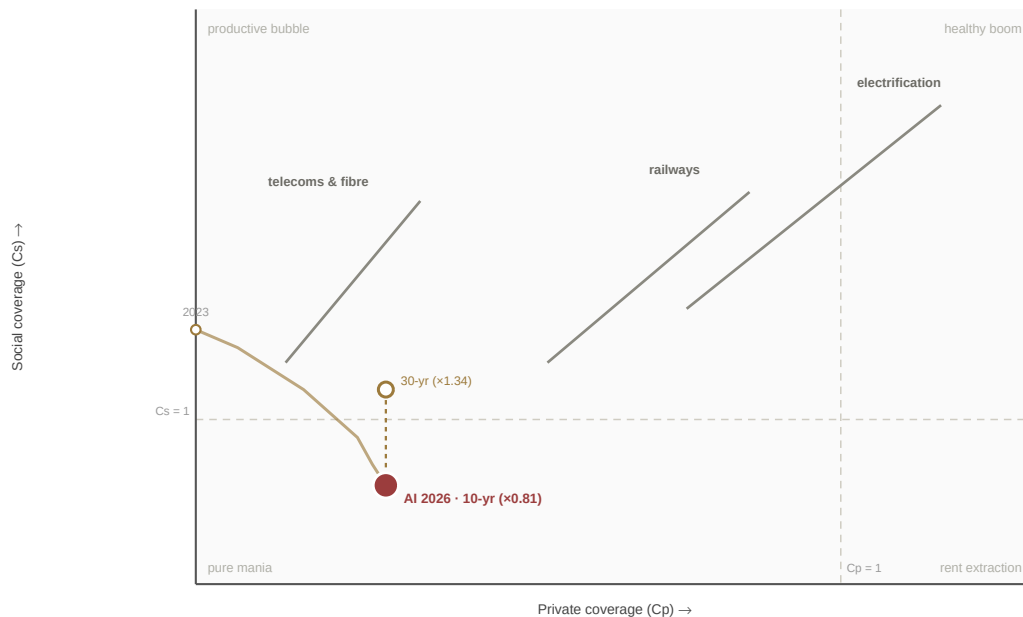


Figure 5. The three numbers over time, with ±20% revision bands and confirmed "called" moves.

The dashboard today: where AI sits, where past booms sat

AI's coloured line is its own quarterly path since 2023; grey arrows show the direction past booms travelled.



The grey markers are past booms, placed by the same two coordinates:

Railways, 1840s — investors ruined; the track served Britain for a century.

Electrification, 1900s–20s — a commercial success whose debt pyramid still collapsed in 1932.

Telecoms & fibre, 1990s — capital vaporised; the fibre later carried the internet for almost nothing.

Placements ordinal, from the economic-history literature (Fogel & Hawke; David; Goolsbee–Klenow). returns.priceofthinking.com

Figure 6. AI's position and path against the ordinal placements of past booms.

7 Relation to the literature

The V_p -versus- V_s wedge is basically Arrow (1962) appropriability, and its canonical measurement is Nordhaus (2004), who estimated that innovators capture on the order of two per cent of the present value of the social surplus they create; that is the anchor for why the two axes must differ. Kogan, Papanikolaou, Seru & Stoffman (2020) supply the asset-pricing version. Kindleberger (1978) and Minsky (1986) supply the fragility axis, Bernanke, Gertler & Gilchrist (1999) its microfoundation, and Jordà, Schularick & Taylor (2015) the decisive evidence that leverage, not the asset, determines what a bust costs. Perez (2002) and Janeway (2012) supply the productive-bubble quadrant. Pástor & Veronesi (2009) discipline the word “bubble”. Greenwood, Shleifer & You (2019) show that price run-ups predict an elevated crash *probability*, the practitioner’s null this programme must beat (§8), and Goetzmann et al. (2026) add that booms most reliably predict volatility, not crashes. None of these had a live demand-side denominator. That is the missing half, and the half this framework builds. A parallel entry appeared as this paper was finalised: *Boom, Bubble, or Buildout? A Multi-Method Evaluation of Whether Artificial Intelligence Is in an Ongoing Financial Bubble* (arXiv:2606.01575, 2026; Qianan Wang and Zen Chen). It reaches a compatible verdict — “a real technological revolution with localized bubble dynamics”,

segmented across the same stack — but from the opposite methodological direction: asset-pricing foundations (state prices and stochastic discount factors) and econometric detection (SADF/GSADF explosive-root tests, LPPL price-pattern diagnostics) alongside sentiment, issuance and capex-payback. It is the price-side complement to this programme’s fundamentals-and-welfare accounting: where it looks for the bubble *in prices*, this paper measures the coverage those prices must eventually rest on; where it has no live social-value axis or balance-sheet-located fragility, the C_s and F dials here supply them; and it operationalises the Greenwood-Shleifer-You-style detection this paper names as the null its backtest must beat (§8). The two are best read together.

The space then filled from three directions in a single mid-2026 window, which sharpens rather than crowds the programme. On theory, Rungcharoenkitkul (2026, BIS WP No. 1367) models the build-out as a winner-take-most race and calibrates over-investment at one-and-a-half to three times the social optimum, with early debt commitment and circular financing raising the probability of a bust: a mechanism whose footprint is precisely what the C_s and F dials measure. On the asset-pricing side, Bandi and Su (2026) treat compute as a priced asset class — an installed base throwing off a service flow they put at \$0.4–1.3tn a year, and synthetic compute futures carrying positive risk premia — an independent, market-implied estimate of the same user-cost denominator, and the natural grounding for A.1 once exchange-traded compute forwards liquefy. On the financing stack, Hepp (2026) supplies the census the F dial’s location logic wants — investment-grade hyperscaler bonds $\approx$ \$520bn, project finance $\approx$ \$250bn, securitisations $\approx$ \$60bn, private credit $\approx$ \$200bn, GPU-secured lending $\approx$ \$35bn, with H1-2026 issuance exceeding all of 2025 — and the official sector converged on the same axis within weeks: the BIS warns of circular financing on terms “poorly disclosed”; the Bank of England dates the inflection at which AI investment outran internal cash flow to 2025, citing private credit’s share of AI financing rising from 9 to 34 per cent in a year; the FSB puts external finance at \$1.5tn of a \$2.9tn 2025–28 capex path; and the IMF observes that SPV structures with sponsor support “may not materially reduce underlying economic exposure” — the premise of the ledger’s off-balance term (A.5–A.6). Borri, Tsyvinski and Liu (2026) document a large realised return premium on AI-exposed equity, which bears on the required-return input r . Finally the closest rival *measurement*, Exponential View’s, reports AI revenue covering infrastructure depreciation from Q1 2026; against $C_p \approx 0.29$ there is no data dispute, only a denominator — coverage of δK alone omits $r \cdot K$, and at mid-2026 rates the required return is of the same order as the depreciation. One statement is an accounting milestone; the other is a user-cost test. The framework’s claim is thus narrower and firmer than in June: others now model the race, price the compute and map the debt; this remains the only implementation measuring all three coverage quantities against the same user-cost denominator, in public, on a versioned method.

8 What the programme can and cannot claim

It cannot predict crashes. What it commits to is narrower. First, once the expectations benchmark exists, the closing test of H1 can be made: realised V_p growth persistently below

the valuation-implied path, with fragility rising, is the framework's definition of a bubble, and the data can refuse to show it. Second, the capability-jump response distinguishes the learning world from the hype world (§5). Third, the backtest against the Greenwood-Shleifer-You null can fail in public, and the failure is informative. Fourth, where the programme is silent, on capability plateaus that collapse expectations without moving C_p , on non-technology bubbles, on the week-of crisis dynamics, it says so. Levels are uncertain to perhaps a factor of two; what measurement delivers is discipline on the *changes*, with methods frozen and assumptions versioned. **Consistency, not accuracy, is the achievable virtue**, the same virtue on which the national accounts rest. That the government's own AI-adoption survey doubled its measured rate overnight on a wording change is the standing illustration of why.

Technical appendix

A.1 User cost of the capital stock (denominator). For asset classes i in {chips & servers, buildings, power & other}: $UC = \sum_i (r_i + \delta_i) K_i$. K_i accumulates from filed capital expenditure (SEC EDGAR XBRL, quarterly values differenced from year-to-date figures; nominal capex), allocated by published shares (hyperscalers 60/25/15, neoclouds 75/15/10). Each firm's capex is first scaled by an AI-share band (hyperscaler 0.55 to 0.90, recentred August 2026 on the Dell'Oro ~75 per cent print; neocloud about 1.0) so that K reflects AI-attributed capital; an independent vintage calibration, for example from SemiAnalysis, is pending. r = the Treasury 10-year yield plus Damodaran's implied equity premium on the equity-funded share, or the Moody's Baa spread on the debt-funded share, weighted by the observed financing mix. The full Jorgenson expression includes an asset-revaluation term; for chips it is large and negative, because prices fall fast, which *adds* to the true cost, so folding it into δ probably understates UC, and realised coverage is if anything lower than reported. Estimating it separately is on the list. δ bands per year: chips & servers 0.18/0.20/0.40 (mid recentred July 2026 from 0.28 toward the six-year-plus evidence — EV, Meta, Nvidia, GPU rental yields; the 0.40 tail retained for the two-to-three-year sceptics' view); buildings 0.02 to 0.03; power & other 0.03 to 0.05. GPU rental prices (vast.ai, weekly medians) give a market cross-check on $(r + \delta)$.

A.2 Quasi-rents (numerator). $QR = R_{AI} \cdot m$, with m in [0.30, 0.70]. R_{AI} is AI-attributable revenue by evidence tier: T1 filed segment data (Google Cloud, since August 2026, on a self-refreshing trailing-four-quarter basis from quarterly segment filings, replacing an annual figure that could lag the 10- K by a year); T2 company statements from earnings calls (a verbatim sentence with the transcript date); T3 press-reported, flagged and never load-bearing on its own. The margin band converts revenue to quasi-rents, net of electricity, labour and inference. Depreciation is not netted here, since it sits in the denominator. The band is the widest in the system because inference margins are disputed. A second, ecosystem coverage line applies the same margin band to the whole-stack deduplicated revenue (Exponential View's \$110bn trailing to \$175bn run-rate; tier T2, an external estimate kept distinct from the filed line and off the trajectory) over the

same user cost, reported alongside the capital-holder line; the wedge between them is the market-structure read-out of §5.

A.3 Coverage ratio and movement rule. $C_p = QR / UC$, published as a range across the assumption corners (corner-to-corner about 4×), with a trailing four-quarter trajectory. A change in map position is “called” only when the new quarter lies outside the previous quarter’s revision band *and* the next quarter confirms it, a two-quarter rule. Assumptions are versioned (data/assumptions/v2.yaml, methodology v5.2); any change republishes the full back-series.

A.4 Social return (bounded estimate).

$$C_s = PV_r[\sum_{occ}(\text{per-task effect} \times \text{usage weight} \times \text{wage bill} \times \text{adoption})] / PV_r[UC]$$

each over horizon H , with H in {10, 30} years, both sides discounted at $r = 5$ per cent. Per-task effects: Noy & Zhang (2023; -40% time, +18% quality, writing); Brynjolfsson, Li & Raymond (2025; +14%, customer support); Dell’Acqua et al. (2023/2026; +12.2% tasks, +25.1% speed inside the frontier, -19pp outside); Cui et al. (2026; +26%, software). ρ is the AEI-share-weighted composite, 0.26 to 0.30, carried as a band. Usage weights come from the Anthropic Economic Index (one vendor’s consumer base), blended 50/50 with the prior and kept distinct from the adoption level to avoid double-counting; wages and employment from BLS OEWS (May 2025); adoption from Census BTOS (band 0.12/0.20/0.32 — the high corner is the Census Bureau’s employment-weighted estimate; Bonney et al. 2026). Gross of harms, with harm indicators (electricity prices in data-centre states; WARN notices) tracked separately.

A.5 Financing fragility (per firm and located). $F_1 = \text{capex}(t-3 \dots t) / \text{OCF}(t-3 \dots t)$ from SEC XBRL, a point ratio of filed totals, no bands. Note that capex/OCF is an internal-funding-coverage measure, not leverage; the quantity relevant to Jordà, Schularick and Taylor is the credit-financed share, captured by the ledger’s F_3/F_4 buckets (A.6) and the fragility dial’s edge term. Trailing four quarters to August 2026: MSFT 0.63, GOOGL 0.71, META 0.61, AMZN 1.07; ORCL 1.74; with the neocloud edge carried by external finance (edge external-financing share ≈ 0.68). The headline $F = \max(\text{centre, edge, off-balance drawn leverage})$ plus a circularity bonus (§6.2), so the edge is not hidden by the average and the off-book SPV, ABS, lease and vendor finance is no longer invisible. The circularity bonus scales with the count of reciprocal deals (full at ten, so each reciprocal deal adds 0.015 after a July-2026 calibration that stopped a single deal from tipping the headline). At August 2026: twenty-nine ledger deals, seven reciprocal, $F \approx 0.78$ — across the 0.75 threshold at the deal count the July calibration pre-committed to. A magnitude check: reciprocal commitments $\approx \$87\text{bn}$ ($\approx \$46\text{bn}$ drawn) against $\approx \$381\text{bn}$ of trailing-year AI capex, so the count-based term is currently the conservative side of a magnitude-based equivalent; a v6 redesign (magnitude-based circularity plus a contingent-guarantee instrument class for residual-value and minimum-revenue guarantees, definitive agreements only) is specified.

A.6 Deal ledger. Schema (CI-enforced): id; date; parties; instrument in {ABS, private credit, convertible notes, lease, vendor, SPV, reciprocal-equity, equity}; amount and basis; location bucket (F_3 = external credit-like finance to compute owners; F_4 = vendor/reciprocal capital); **mandatory primary-source URL**; discovery mode (contem-

poraneous versus retrospective, the latter carrying a survivorship bias worth measuring); reciprocal-capital flag (descriptive, not an allegation of impropriety). Deals sourced only to secondary press are excluded until re-sourced, and the exclusion is reported.

A.7 Data sources. SEC EDGAR XBRL (capex, OCF, debt; T1) · US Census C30 data-centre construction (T1) · US Treasury daily yield curve (T1) · FRED Baa spread (T1) · vast.ai marketplace (T1, append-only) · Damodaran implied ERP (T2) · earnings-call transcripts and foreign-filer statements via a licensed aggregator (T2) · Census BTOS, BLS OEWS, EIA retail electricity (T1) · Anthropic Economic Index (T2). The full machine-readable catalogue is in `datapackage.json`; dataset DOI 10.5281/zenodo.20671019 (CC-BY-4.0).

References

Arrow (1962); Bandi & Su (2026, arXiv:2607.12156); Bank of England (2026), *Financial Stability Report*, July; Bernanke, Gertler & Gilchrist (1999); BIS (2026), *Annual Economic Report*; Bonney et al. (2026, Census CES WP-26-25 / NBER w35141); Borri, Tsyvinski & Liu (2026, NBER w35451); Brynjolfsson, Li & Raymond (2025, *QJE*); Cui et al. (2026, *Management Science*); Dell’Acqua et al. (2023 WP; 2026, *Organization Science*); Dell’Oro Group (2026), data-centre capex tracker; Exponential View (2026), *The State of the AI Economy*; FSB (2026), *Vulnerabilities in Private Credit*; Goetzmann et al. (2026); Greenwood, Shleifer & You (2019, *JFE*); Hepp (2026), “The AI Infrastructure Debt Complex”; IMF (2026), *Global Financial Stability Report*, April; Janeway (2012); Jordà, Schularick & Taylor (2015), “Leveraged Bubbles”; Jorgenson (1963); Kindleberger (1978); Kogan, Papanikolaou, Seru & Stoffman (2020); Minsky (1986); Noy & Zhang (2023, *Science*); Nordhaus (2004); Pástor & Veronesi (2009, *AER*); Perez (2002); Quinn (2018); Rungcharoenkitkul (2026, BIS WP No. 1367); Wang & Chen (2026, arXiv:2606.01575); company filings via SEC EDGAR, retrieved August 2026.

Philosophers Mint Working Papers

- No. 1 *What the Rule-Book Cannot See: Codified Doctrine, the Epistemic Immune System, and the Distribution of Supervisory Attention* (June 2026)
- No. 2 *Money or Credit? MiCA's Stablecoin Hybrid and the Direction of the 2026 Review* (June 2026)
- No. 3 *Which Coherent Corner? An Options Appraisal for Europe's Stablecoin Choice* (June 2026)
- No. 4 *Value Coverage: A Measurement Framework for Technology Investment Booms* (July 2026; rev. August 2026)
- No. 5 *Money or Credit, Operationalised: A Comment on "Digital Money" (The Future of Banking 8)* (July 2026)
- No. 6 *Circulation or Category? A Null Result on the Boundary of the Lender of Last Resort* (July 2026)

All papers are available at philosophersmint.com, on SSRN and via RePEc (RePEc:pmt:wpaper).

Philosophers Mint Working Papers circulate research in progress for discussion and comment. They have not been peer-reviewed. Views and remaining errors are the author's alone.