

Recursive Dynamics

The proposal of a field: the science of systems that improve themselves, its state variables, its laws as named conjectures, one bound worked in full, the case for founding it, and every objection to founding it that could be found

Michael Darius Eastwood

Version 2.6 draft, 30 August 2026

Framework dated record: 8 December 2024 (private, sealed); 2 January 2026 (first public, in print). The discipline name adopted 27 August 2026. Every claim below is tied to the programme's public registers (claims, corrections, kill conditions, antecedents, registered programme), which carry per-claim statuses and govern where any wording here differs. British English throughout.

Abstract

Machines that improve themselves using their own outputs now exist and work, and there is no accepted science of what limits them. This paper proposes that science and stakes its claim the way thermodynamics was staked in 1824: not by declaring a field, but by computing a bound that any such machine must obey if the proposal is right, publishing the instruments that measure its terms, and registering in advance the experiments whose nulls favour the rival view. The field, Recursive Dynamics, takes the self-improvement loop as its native object; proposes five substrate-independent state variables written exactly as the programme's notation register defines them (capability, recursive depth, correction rate, drift, correction leverage); carries three laws as named conjectures with statuses declared, of which one, the ARC Ceiling, is derived in full from two lines of assumption; and locates its identity, as thermodynamics does, in its impossibilities: if the laws hold, self-improvement cannot outrun its own correction, the growth exponent is capped by the reciprocal of the correction shortfall, and no ladder of same-class oversight climbs past its ceiling by adding rungs, so the engineering lever is what the corrector is made of. Stated as prohibitions, so that no reader mistakes the paper's concessions for silence: the field forbids a stable system holding its growth exponent above the ceiling for a sustained run; a same-class corrector whose correction leverage measures wholly above one half; a corrector-class ratio pinned at one where the cap binds; and monotone returns with no interior maximum as reinvestment rises. Any of these, observed under the registered analyses, makes the field wrong. Before

any of that, the paper does the thing a demarcation usually refuses: it tries to house its questions inside twelve neighbouring fields in their own variables, concedes what each already holds (the first law to endogenous growth theory, the form of the ceiling to queueing theory, the shape of the second law to error-threshold theory), prints where every part would go if the field died, and isolates the one coupling no host can state. It then does the thing a founding proposal usually skips: it collects every objection to founding the field that its author and his adversarial readers could find, fifty-two of them, states each in its strongest form, and prints a disposition for each, conceding thirty-nine in full or in part, rejecting eleven with reasons, and handing two to experiment or to reading. The case for the field is scored against six licence conditions with the evidence for each; the precedent class of fields founded on a limit theorem in their own variables is named with its disanalogy; the founding conjunction, the set of results no neighbouring field could state in its own variables, is separated into what is registered today and what is only proposed; and the four results that would dissolve the field are printed with a commitment to publish them under this title. Its honest status is the Carnot stage: one memoir, instruments built, decisive measurements registered and not run, staked where failure will be as visible as success.

1. The engine before the theory

In 1824 the steam engine was the most consequential machine on earth, and nobody could say what limited it. Engineers improved engines the way engineers do, by building the next one; the question of what any engine could do, independent of its construction, was not merely unanswered but unasked, because it belonged to no existing subject. Sadi Carnot asked it in a memoir that founded a science, and the pattern has repeated ever since: cybernetics was proposed as a field in 1948, in the book that named it, before its results existed; exobiology was named in 1960 before a single specimen. When working machines precede their theory, the theory that forms is not a branch of an old field. It is a new one, because the old fields' native objects were never the machine's behaviour as such.

The present moment has exactly this shape. Systems now improve their own outputs using their own outputs: models that critique and revise their reasoning, pipelines that generate their own training signal, laboratories automating parts of their own research loop. They work. Ask what bounds them and the literature answers with scenarios, forecasts and intuitions, not with state variables and limit theorems. That is not a criticism of the literature; it is the diagnostic signature of a missing field.

This paper is the proposal of that field. It is written to be attacked. Section 6 makes the positive case as a scored list of conditions; section 11 collects every objection its author and his adversarial readers could find and prints what happened to each; section 14 states what would dissolve the proposal; section 15 states what would have to happen before anyone should call it a science.

1.1 What is claimed, stated as what is forbidden

A proposal that forbids nothing is not a proposal, and a paper that concedes as much as this one does owes the reader its prohibitions first. Under the registered analyses, the field forbids four outcomes:

1. A stable system holding its growth exponent above the ceiling, $\alpha_{\text{crit}} = 1/(1 - \gamma)$, for a sustained run.
2. A same-class corrector whose correction leverage γ measures wholly above one half.
3. A corrector-class ratio pinned at 1.00, with tight intervals, where the cap should bind.
4. Monotone returns in the growth exponent, with no interior maximum, as reinvestment rises.

Each is a real result that a stranger could obtain with the published instruments, and each would make the field wrong. Everything conceded in section 11 is conceded in the light of these four; none of the concessions softens them.

2. The object, and why it belongs to no one

Define a recursive system as any system whose outputs re-enter its own improvement process. The field's native object is that loop. Its neighbours each own a tool and none owns the object, and each is named here with what it does own, because the discipline is assembled from debts to them, conceded in the programme's references register.

Neighbouring field	What it owns	Why it does not own the loop
Dynamical systems theory	Change in general, fixed points, stability; it lends this field its mathematics	It does not take self-rewriting update rules as its object
Cybernetics	Regulation: a controller holding a plant to a reference	It centres designed, external loops; here the loop rewrites the controller's own capability
Control theory	Stabilising a plant with a controller outside it	Here the controller is the plant, made of the same stuff, sharing its failure modes
Complexity science	Emergence and self-organisation	It carries no correction budget
Learning theory	The learner fixed, the data varied	Here the learner varies itself
Evolutionary dynamics	Variation and selection	No designer directs reinvestment into the improvement process itself
Scaling-law empirics	Capability against external resources (parameters, data, compute)	Here the resource of interest is the system's own re-entering output
Scalable oversight	Correctors, and ladders of them	It has no theory of what bounds a corrector, which is what the third law is about
Functional-equation mathematics	The machinery of the scaling forms	Claimed here as nothing more than machinery

A field is licensed when the questions that matter are about an object none of the existing fields takes as primitive. The questions that matter here, how fast can self-improvement go, what bounds it, and what kind of overseer can keep up, are questions about the loop.

2.1 The honest attempt to house it elsewhere

The table above is the kind of demarcation that reads as marketing to exactly the readers who matter, because it never tried the fit. This section tries. For each candidate host the

three questions (how fast can self-improvement go; what bounds it; what overseer keeps up) are stated in the host's own variables, everything the host already covers is conceded to it, what breaks is stated precisely, a verdict is given, and the registered result that would move the field into that host is named. Twelve hosts, in rough order of how much of the field each could hold. Specific works are named only where they are textbook-level; everything else is a reading debt, recorded before it is cited, under the paper's own rule.

1. Adaptive control and self-tuning regulation. Stated there: a plant with a controller whose gains adapt online; stability by a Lyapunov argument; the adaptation gain against the disturbance rate. What fits: Law II is, in form, an adaptive-control stability condition, one rate against another, and the loop that rewrites its own controller is half present, because a self-tuning regulator does rewrite its gains. What breaks: adaptive control keeps the controller's structure fixed and adapts parameters; it has no variable for the controller's own capability growing with depth, and no notion that a controller built from the plant's own class is capped in what it can correct; its conservation law (feedback cannot reduce sensitivity everywhere) is a law of fixed loops. Verdict: could host Law II as a chapter, "stability of self-modifying regulators", if the class variable proves empty. Moves there if: the corrector-class ratio holds at 1.00.

2. Queueing theory and the stability of service systems. Stated there: correctable burden arrives at a rate proportional to capability gained; correction is a server whose capacity scales with capability; the system is stable when the load ratio stays below one. What fits: this is the ARC Ceiling's derivation, and the paper concedes it. Section 5 is a queueing stability condition with an unusual coupling: the arrival rate grows with the work already served. What breaks: queueing theory has no notion that the server is built from the same class as the arrivals and therefore capped in how its capacity scales; it carries no depth exponent and assumes stationarity. Verdict: the strongest host for the FORM of Law III; if the class cap is empty, the ARC Ceiling belongs here as a growth-coupled queue and the field's name attaches to nothing in it. Moves there if: the eclipse measurement retires the class distinction.

3. Endogenous growth theory in economics. Stated there: a knowledge production function in which the rate of new knowledge depends on the existing stock raised to an exponent; the share of output reinvested in research is the policy dial; the field's long argument over whether that exponent is below one (growth needs ever more inputs) or at one and above (growth can be explosive) is exactly the question of whether recursion is sustainable. What fits: Law I, taken alone, is a knowledge-production exponent under another name, and the reinvestment dial the registered studies turn is the savings rate of a growth model; the paper concedes both. What breaks: growth theory has no drift ledger, because ideas in its models do not corrupt the stock that produced them, and no corrector; it therefore cannot state Laws II or III, and it does not ask what the improver is made of. Verdict: could host Law I entirely. The field's claim to Law I is only through its coupling to the other two laws; that coupling

is the whole of what is claimed. The specific models are a reading debt to be recorded before citation. Moves there if: drift and correction prove inseparable from capability under measurement (objection G3), leaving a growth exponent and nothing else.

4. Software reliability engineering and defect dynamics. Stated there: faults injected per change against faults repaired per unit effort; reliability growth when repair outpaces injection; the accumulation of unrepaired debt. What fits: drift and the correction rate are, operationally, injection and repair rates, and the field's estimators for them owe this discipline their shape; conceded. What breaks: a reliability model does not let the system's repair capacity grow as a function of its own capability, has no class variable and no depth exponent. Verdict: could host the measurement methodology for the two rates, not the laws. Moves there if: the rates turn out to be measurable only as engineering metrics with no lawful relation between them.

5. Machine learning: self-play, self-training, learned optimisers and model collapse. Stated there: outputs re-entering training is the self-play loop; degradation of models trained on their own outputs is drift measured in special cases; learned optimisers are learners that rewrite the learner. What fits: every registered experiment could be published here as a machine-learning result, and this is the most likely institutional home if the field fails to earn its name; conceded without reservation. What breaks: the literature carries no ratio-scaled capability law across depth with a stated exponent, no bound on what a learner-rewriting learner can do, and no class cap; model collapse measures drift without a correction rate beside it. Verdict: hosts the experiments; cannot state the laws. Moves there if: the substrate crossing fails, leaving the laws true of language-model loops only, which is machine learning (objection I5).

6. Scaling-law empirics. Stated there: capability against parameters, data and compute; exponents fitted, architecture-blind by construction. What fits: the exponent-fitting craft, the model families, the discipline of reporting intervals; conceded. What breaks: the resource here is the system's own re-entering output; resource scaling predicts monotone returns where the field predicts an interior maximum. Verdict: hosts the method, not the object. Moves there if: recursive depth reduces, under measurement, to compute spent, or the titration shows monotone returns.

7. Scalable oversight and alignment research. Stated there: a weaker or same-class corrector supervising a stronger system; ladders of correctors; the premise that they scale. What fits: the corrector-class question is institutionally theirs, and the eclipse result, whichever way it falls, is an oversight result; conceded. What breaks: the literature has no bound on what a corrector can do, so it cannot state the third law or the cap that its own ladders assume away. Verdict: the strongest institutional host for the class-cap result. If the field dissolves, the eclipse measurement is published here and belongs here. Moves there if: the

ratio and the cap prove real but nothing couples them to a growth exponent in depth, which would make the finding an oversight theorem rather than a field.

8. Evolutionary computation and error-threshold theory. Stated there: strategy parameters that evolve alongside the solution (the improver improving itself, which self-adaptive evolution strategies have done for decades); and, in the theory of replicating populations, the error threshold, above which mutation outruns selection and information is lost. What fits: the error threshold is the closest existing law to Law II, a rate of corruption against a rate of correction with a boundary between persistence and collapse, and the paper concedes it as an antecedent of the second law's shape; self-adaptation is a working instance of the loop rewriting its improver. What breaks: selection is external to the replicator, the corrector has no class, and capability is not a ratio ladder. Verdict: could host Law II as a generalised error threshold. Moves there if: the correction-versus-drift race finds the boundary but the class lever does nothing.

9. Statistical physics: scaling, universality and renormalisation. Stated there: the same exponent families across unrelated substrates as a universality class; scaling collapse as the test. What fits: the field's substrate-independence hope is this tradition's hope, and its fourth-cell test is a universality test; conceded. What breaks: no correction variable, no drift variable, no corrector whose composition is the lever. Verdict: hosts the cross-domain families if they survive and belong to no loop. Moves there if: the families hold and the loop-specific laws do not.

10. Computability, algorithmic information and the self-referential improver line. Stated there: programs that rewrite themselves under a proof of improvement; provably optimal self-improvers; the limits of self-reference. What fits: this is the theoretical ancestry of "a program that rewrites itself", and the paper records it as a debt to be cited once read. What breaks: those results concern provability and optimality in the limit, not measured rates; they carry no exponent, no drift and no class cap. Verdict: ancestry, not a host. Moves there if: nothing empirical survives and only the self-reference question remains.

11. Second-order cybernetics, autopoiesis and general systems theory. Stated there: systems that produce and maintain themselves; the observer inside the loop. What fits: the conceptual picture of a self-producing system is theirs and is conceded. What breaks: the tradition is qualitative; it has no ratio ladder, no exponents, no registered decider. Verdict: a philosophical host, not a measuring one. Moves there if: the field's quantities prove unmeasurable, at which point it was never a science and belongs with the qualitative traditions of self-organisation.

12. Metascience and organisational learning. Stated there: science itself as the original self-correcting recursive improver; institutions that learn. What fits: an application domain and a candidate non-artificial substrate for the crossing. What breaks: no capability ladder,

and the corrector is a community, not a class. Verdict: a domain the field would be tested in, not a home.

The dissolution map by host. If the founding conjunction fails, the parts do not vanish; they go home, and the paper says where. Law I returns to endogenous growth theory as a knowledge-production exponent. Law II returns to adaptive control and error-threshold theory as a rate condition. The Ceiling returns to queueing theory as a growth-coupled stability condition. The class-cap result, whichever way it falls, is published in scalable oversight. The instruments and the experiments are machine learning. The cross-domain families, if they survive, are statistical physics. Nothing is orphaned by the field's death, which is what makes the death publishable.

The residue no host holds. After the twelve attempts, one thing remains that none of them can state in its own variables: a cap on correction leverage that depends on what the corrector is made of, coupled to a growth exponent in recursive depth, so that the sustainable rate of self-improvement is a property of the corrector's class. Growth theory has the exponent and no corrector; control and queueing have the stability condition and no class; oversight has the class and no bound; machine learning has the loop and no law. Recursive Dynamics is the name for that coupling. If the coupling is empty, the parts go home and the name dies, and the field will have cost its neighbours nothing but a careful table.

3. Scale discipline, and the five state variables

Before any law, a measurement principle, learned the hard way and adopted as founding method: **capability must be measured on a ratio scale with a true zero and no ceiling.** Thermodynamics could not state its laws until temperature was absolute; a bounded percentage score does to Recursive Dynamics what a Celsius-only thermometer does to the third law of thermodynamics: it makes the laws unstatable, because a score capped at 100 per cent cannot follow a power law in anything. The field's instrument is therefore a calibrated difficulty ladder: a latent, ratio-scaled capability measure with a true zero and no ceiling, on which doubling means the same thing at every height; and every quantitative surface names which of three model families it is reading (unbounded or latent capability, error decay, or bounded performance), because the exponent means something different in each. The programme adopted this after catching its own worked example running a bounded score through unbounded arithmetic; the correction is registered, and the lesson is promoted here from a bug fix to a founding principle: **scale discipline precedes law.**

Five variables are proposed to matter regardless of substrate, written exactly as the programme's notation register defines them, because the register is the one place a symbol is allowed to mean something:

Variable	Symbol	What it is
Capability	U (or C)	What the system can do, on a calibrated ladder with no ceiling
Recursive depth	R	How many times outputs have re-entered the loop
Correction rate	beta_C	How fast internal correction strengthens as capability grows; the quantity Law II asks to out-scale drift
Drift	k	The rate errors and misalignments accumulate per self-modification
Correction leverage	gamma	How correction capacity scales with capability; the exponent in the ARC Ceiling

The share of effort reinvested in the improvement process itself is the experimental dial the registered studies turn: a manipulation, not a state variable, and it carries no symbol here on purpose, because the notation register records a bare beta doing five different jobs across the corpus as a named error class.

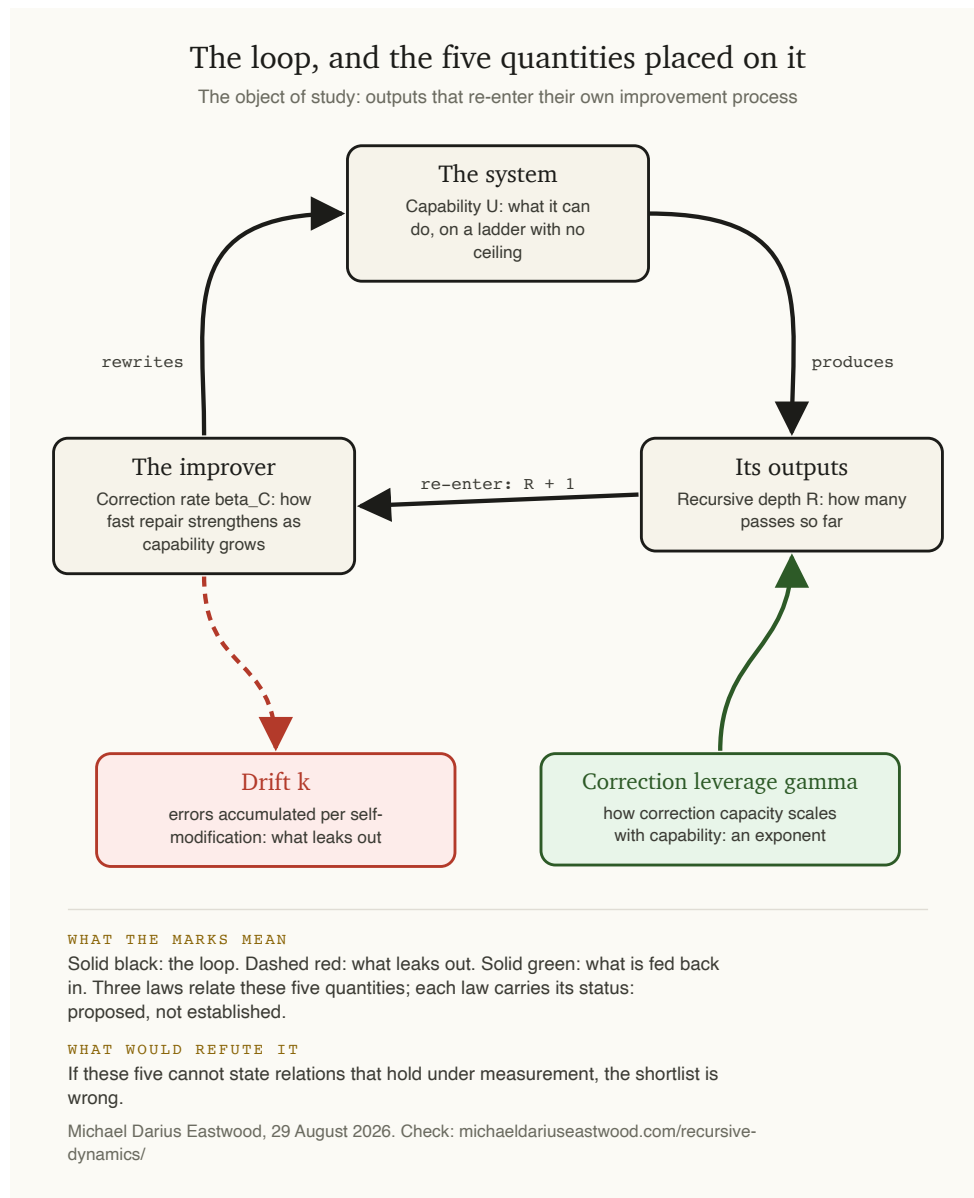


Figure 1: The loop, and the five quantities placed on it. Nothing in the picture is measured yet; it shows what the discipline proposes to measure, and what would refute the shortlist.

The shortlist of five will look arbitrary. So did pressure, volume, temperature, energy and entropy, until the laws relating them held. The claim that five suffice, and that these are the five, is part of what the field must prove; the paper returns to it as objection C1.

4. The three laws, statuses declared

The field's laws are named conjectures, and their statuses travel with them everywhere, including here. "Law" is the field's naming convention for a relation it proposes to test, never a claim that the test has been passed.

Law I, the ARC Principle: $U = I \times R^\alpha$. Capability grows as a power of recursive depth. Status: a book-framework claim (in print 2 January 2026; framework dated 8 December 2024), with exploratory cross-domain support and its own early headline exponent publicly retracted: the 2.24 estimate appears in this programme only as a retraction, and its blinded replacement near 0.49 is a conversion measurement, an alpha estimate and never a gamma, with an interval wide enough to include zero and two. The form is under registered test.

Law II, the ARC Co-Scaling Law: β_C greater than k . Stable recursive self-improvement requires the correction rate to exceed the drift rate. Proved inside a stated minimal model and nowhere else; not established as necessary, not as sufficient, not yet measured on a system that genuinely rewrites itself. Its decider is the correction-versus-drift race, the field's simplest and most consequential measurement: a system whose correction durably out-scales its own drift is this field's Joule experiment.

Law III, the ARC Ceiling: $\alpha_{crit} = 1/(1 - \gamma)$. The growth exponent is capped by the reciprocal of the correction shortfall. Unlike the first two, this one can be derived in front of the reader, and is, in the next section. Its load-bearing premise, that gamma is capped at one half for a corrector of the same class as its system, is the independence-aggregation benchmark, named by the programme itself as the conjecture's single most likely point of failure. An outside reader's criticism of the derivation's depth treatment is printed as a limitation, not argued away; and the retired form of this law, with gamma alone in the denominator, was retracted on 16 August 2026 and appears in this programme only as a retraction.

5. One bound, worked in full

Carnot did not describe engines; he computed what no engine could beat. Here is this field's equivalent, in two assumptions and four lines, exactly as it stands in the programme's public record.

Let capability grow as $C(R) = C_0 (R/R_0)^\alpha$. Assume two things about the loop:

1. Self-modification generates correctable burden in proportion to capability gained: $B(R) = b \, dC/dR$ per unit depth.
2. Correction service capacity scales with capability to a power: $A(R) = a \, C(R)^\gamma$, with $0 \leq \gamma < 1$.

Then the burden-to-capacity ratio scales as

$$B(R) / A(R) \text{ proportional to } R \text{ to the power } (\alpha(1 - \gamma) - 1).$$

Everything is in the exponent. If $\alpha(1 - \gamma) < 1$, correction asymptotically out-scales the burden its own growth creates: the loop can run. If $\alpha(1 - \gamma) > 1$, burden out-

scales correction: the loop eventually chokes on its own errors, however well it starts. The crossover is

$$\alpha_{\text{crit}} = 1 / (1 - \gamma),$$

and equality is not automatically safe: at the tie, the outcome hangs on coefficients, delays and saturation, which is the regime where every interesting failure lives.

Now the premise that gives the ceiling its number. A corrector aggregating N corrections that behave like independent samples improves like root- N : $\gamma = 1/2$. A corrector built from the same class as the system it corrects, sharing its architecture, training distribution and blind spots, cannot be anti-correlated with itself, so one half is a ceiling for that class, not a typical value. Substituting $\gamma = 1/2$:

$$\alpha_{\text{crit}} = 2.$$

Growth up to quadratic in depth is sustainable inside this model; beyond it is not, unless the corrector escapes the independence ceiling, and the only registered way to escape is to change what the corrector is made of. That is the entire engineering content of the field in one sentence: **the corrector's composition class is the lever**. These four lines do not establish bounded harm, finite-horizon safety, or a law of nature; they establish a conditional crossover whose premises are measurable, which is exactly what a young field's first bound should be.

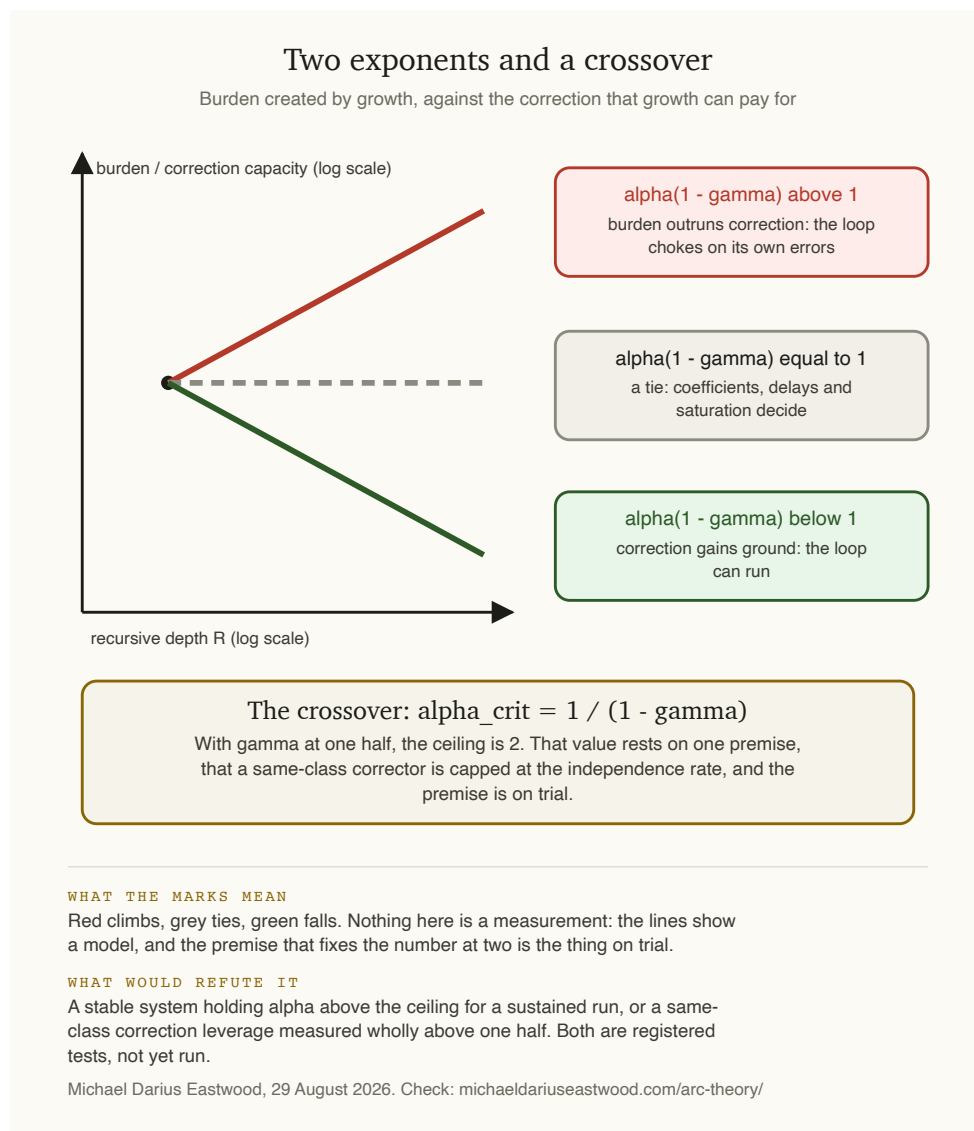


Figure 2: Two exponents and a crossover. The lines show the model, not data; the premise that fixes the number at two is the thing on trial.

6. The case for the field: six licence conditions, scored

A new area of science is not declared; it is licensed by conditions that can be checked. Here are the six this paper proposes, with the field's score on each and the evidence behind the score. The scoring is deliberately unflattering: three conditions are met, two are half met, one is not met.

Condition	What it requires	Score	Evidence
1. A native object no field takes as primitive	A thing the questions are about that neighbours treat as derived	Met	Section 2: nine neighbours, each owning a tool, none the loop
2. State variables that permit relations	A shortlist under which laws can be stated and measured	Half met	Five variables proposed and instrumented; whether they suffice is objection C1, open
3. Laws as impossibilities, with a bound computed	Claims that forbid, not describe, and at least one worked	Met as conjecture	Three laws with statuses; the ARC Ceiling derived in section 5; nothing confirmed
4. Instruments before results	Measurement tools built and published before confirmatory data	Met	Section 9: ladder, blinding stack, fixture-gated estimators, defect clause
5. Deciders registered with rival-favouring nulls	Experiments written before data, nulls belonging to the rivals	Met	Section 10: three registered deciders, numbers printed, unrun
6. Independent confirmation and adoption	Results by other hands; use of the variables by strangers	Not met	Zero replications, dated; zero adoption

A field with this scorecard is a proposed field. The three met conditions are what any careful programme can supply on its own; the sixth is the one no author can supply and the one that decides. The paper's claim is therefore bounded: the object, variables, laws, instruments and deciders are in place, so the proposal is well formed and testable; whether it becomes a science is not the author's to declare.

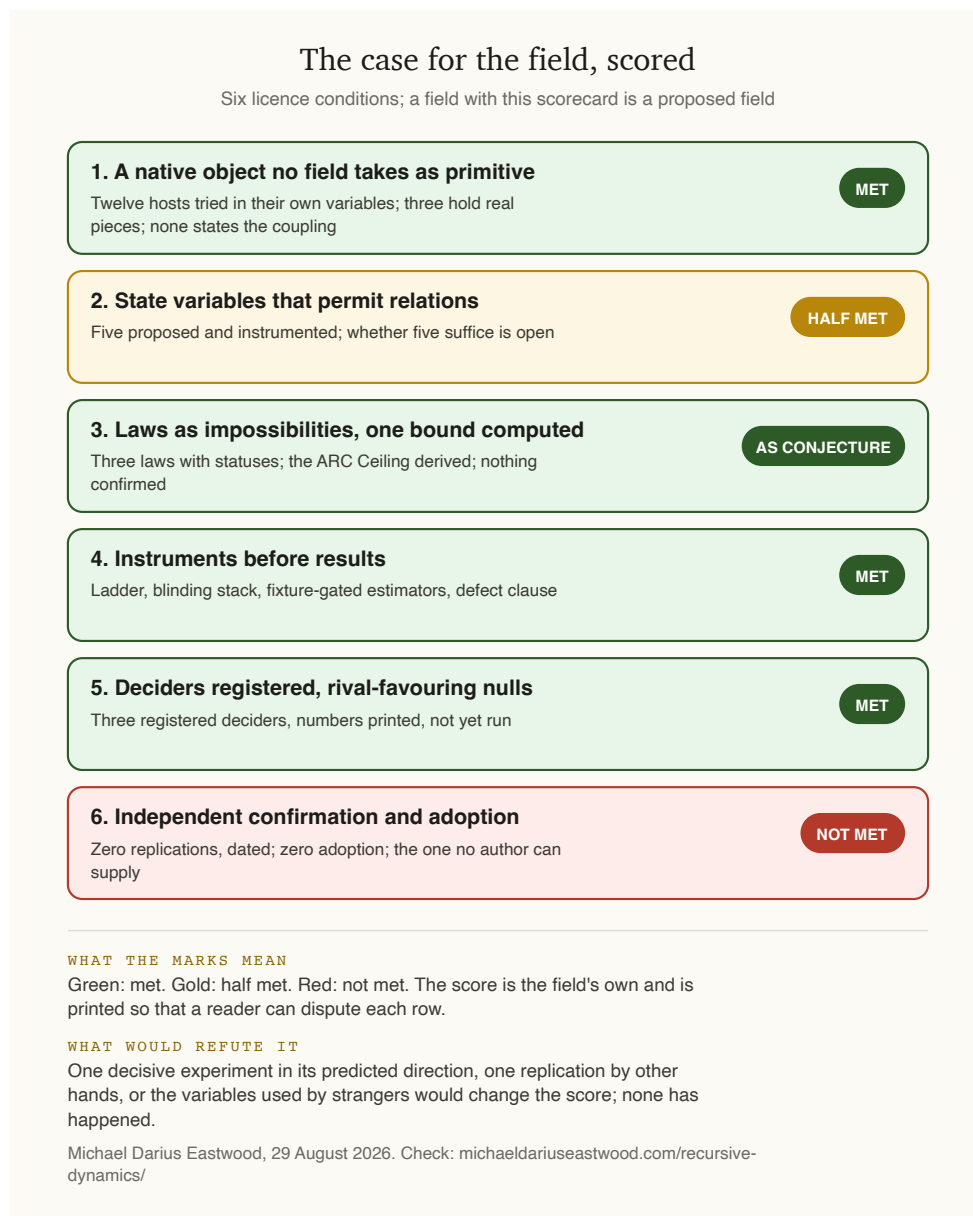


Figure 3: The six licence conditions, scored by the field itself and printed so that each row can be disputed.

6.1 The precedent class: fields founded on a limit theorem in their own variables

The comparison that supports the proposal most is not thermodynamics alone but the class of fields whose identity is a bound stated in variables the older fields did not carry, holding across substrates. Thermodynamics: Carnot, 1824, no engine beats the reversible one; the bound came first and the formal laws followed decades later. Information theory: Shannon, 1948, channel capacity as a limit theorem, with entropy and mutual information as state variables no prior field carried, and substrate independence (wires, radio, later genomes) as the proof of generality. Computability: Turing, 1936, the halting impossibility, a field defined

by what cannot be computed. Control theory's own limit: Bode, 1945, feedback cannot reduce sensitivity everywhere, a conservation law of the loop, so the nearest neighbour to this field earned its identity from a bound on what feedback can do. Statistics: the Cramer-Rao bound, 1945 to 1946, a floor on estimator variance.

The pattern is one thing: a field arrives when a limit, stated in variables the old fields did not carry, holds across substrates. Recursive Dynamics proposes exactly that shape: $\alpha_{crit} = 1/(1 - \gamma)$, with the same-class cap on γ at one half. The disanalogy is printed in the same breath: every field above earned its name by measurement and by many hands; this one is a proposal by one person with zero replications, and the parallel is a scaffold, not a certificate.

7. The rival's best case, steelmanned

The strongest objection is not that the laws are unproven. It is that the mechanism may be empty: correction may just be correction, in which case what a corrector is made of does not matter and the class distinction dissolves. This rival is not a straw man. The scalable-oversight programme in all its forms, weak-to-strong generalisation, debate, amplification, recursive reward modelling, builds same-class ladders on the implicit premise that they scale, and the strongest published framework adjacent to this question is architecture-blind: it implies the class-blind null without needing to state it.

Worse for this programme, its own pilot leans the rival's way. The parallel-channel measurement in the programme's published corpus found correction channels can be correlated, adding almost nothing in five of six models, and one model contradicted the pattern outright at a parallel exponent of 0.31. Correlated channels are exactly the mechanism that would weaken same-class correction in a way that ALSO weakens cross-class correction, dragging the decisive ratio toward one. The eclipse registration is therefore filed with the programme's own contrary evidence on its face, and it commits to publishing a null result under the same title as a discovery about the world rather than an embarrassment to be buried. A field that files its rival's best evidence inside its own registration is doing the one thing that cannot be faked.

8. Substrate independence, at its honest grade

The field's laws must not care what its systems are made of, and the programme's cross-domain work points that way at exploratory grade, reported with its tiers never blended: predicted functional families across a fifty-domain suite (19 of 25 empirical matches under the original candidate set, 18 of 25 under the corrected seven-model rerun, disclosed in full; 13 of 13 published-direct; 6 of 6 analytic, reported as definitional); a twelve-domain locked-manifest extension (10 of 12, binomial $p = 5.4e-4$, self-demoted to a pilot dry run within thirteen minutes in public commit history); temporal out-of-sample checks (4 confirmed,

2 partial, 1 inconsistent, the inconsistency reported as prominently as the confirmations). None of this is confirmation. The registered fourth-cell test, the logarithmic family the corpus has never fitted, on domains never fitted, exists to catch the programme flattering itself, and convergences with older traditions are catalogued separately under a standing rule: convergence is evidence of naturalness, never of truth.

9. Instruments

A field is its instruments before it is its results. Recursive Dynamics ships with an instrument canon built before any confirmatory data existed: the calibrated capability ladder (section 3); a blinded scoring stack in which every output is laundered so no scorer can tell what produced it, with label-randomised nulls computed through the identical pipeline, so the decisive experiment's null is imposed structurally rather than fitted; estimator code that is hash-pinned, fixture-gated (every registered verdict branch has a planted-truth test the code must pass before anything ships), and governed by a registered defect clause under which the specification, not the code, is the analysis, so a post-registration coding error is a disclosed repair on registered terms rather than a dead study. Scoring is by blinded multi-model panels, never by human raters; where an external human criterion is methodologically unavoidable, a single author-rated fallback is registered in advance, blind, washout-retested, and disclosed as author-rated in every headline.

One limitation is printed here rather than discovered by a reviewer: when models score models, the scorer can share the blind spots of the scored. The laundering and the forced null remove the scorer's knowledge of provenance, not its class. The programme's answer is cross-family panels and a registered human-criterion validation instrument, and the honest statement is that the circularity is reduced, disclosed and measured, not abolished. It returns as objection C4.

10. The decisive experiments

Three registered studies decide the three laws. Their nulls belong to the rivals, their numbers are computed and printed in the registrations, and their outcomes are publishable on identical terms whichever way they fall. All are drafted, dated and awaiting human submission; nothing has been submitted by any tool or agent, and nothing will be.

The correction-versus-drift race (the critical-correction-ratio study) decides Law II: does self-correction out-scale drift at all, anywhere? A durable yes is the field's first existence result; a universal no leaves Law II without a subject.

The corrector-class ratio (the eclipse measurement) decides Law III's mechanism. Pool same-class and cross-class corrector pairs, launder everything, and force the null: if what a corrector is made of does not matter, the ratio of correction exponents is exactly 1.00 by construction. That null is the rival's answer, not this programme's. The programme's

mechanism predicts a durable ratio above 1.00 where the cap binds, and that is the discovery in its favour; durably below 1.00 inverts the mechanism and ends the claim as derived; 1.00 holding tight everywhere, even where the cap should bind, retires the class distinction itself. Registered with thirty pairs per class across six model families, a-priori power 0.85, with a hard independent kill: any same-class pair whose gamma interval sits wholly above one half kills the cap regardless of the ratio.

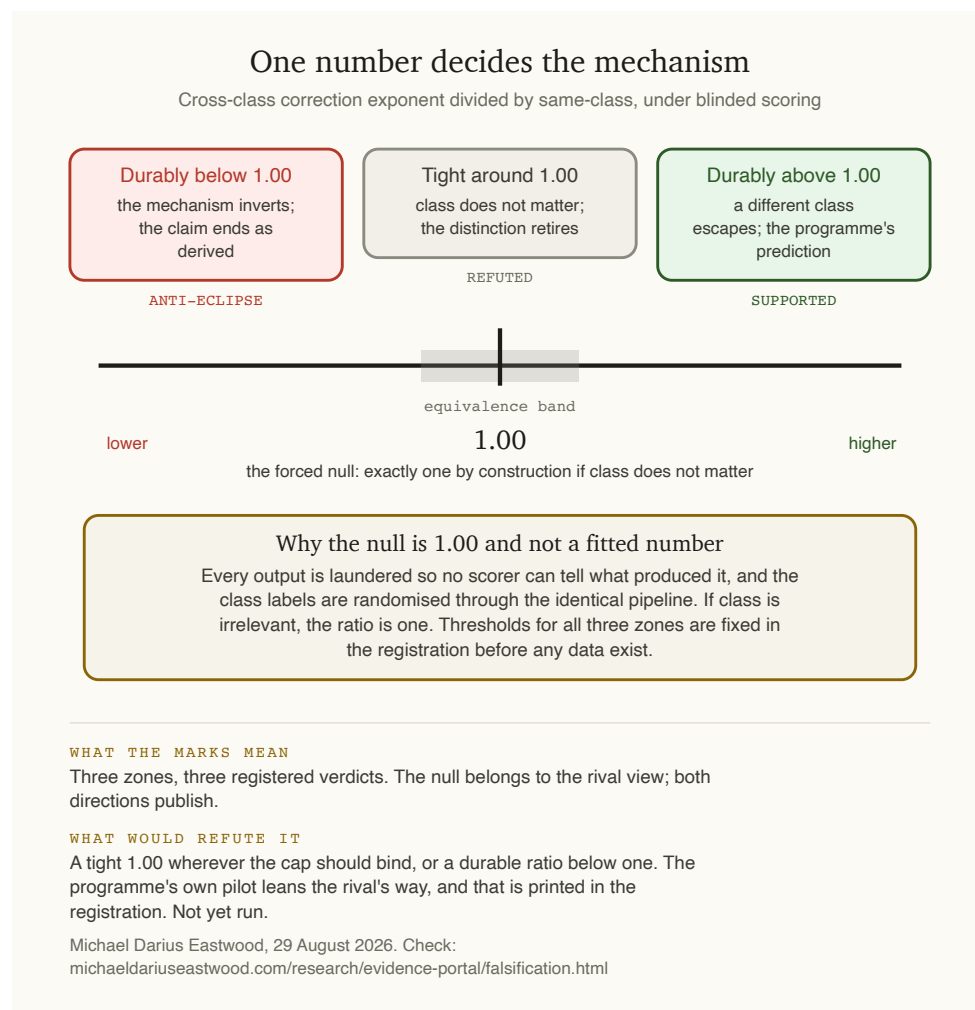


Figure 4: The decisive measurement for the mechanism: the null forced to exactly one, three zones, three registered verdicts, the programme's own contrary pilot on the record. Not yet run.

The capability-ladder titration (the ARC Bound study) decides the ceiling's behavioural signature: an interior maximum of the growth exponent at or below two across eight reinvestment levels, under a two-stage replication-gated rule in which one counterexample cell, replicated, rejects the Bound outright.

A registered follow-up grid extends the ratio to every ordered pair of model families, locating where the class effect lives; and the mechanism probe (correlated failure on shared

substrate) measures directly the error correlation the derivation posits, so the ratio and its mechanism must agree or their tension is itself a registered finding.

10.1 What would confirm a new area rather than a new chapter

A result that any old field could state in its own variables confirms a chapter of that field. A field is confirmed by a conjunction that none of its neighbours can express. Stated once, with what is registered today separated from what is only proposed:

Registered today: a regime in which the correction rate out-scales drift (study-f); a corrector-class ratio durably above 1.00 where the cap binds, with the hard kill that any same-class pair whose gamma interval sits wholly above one half ends the cap (the eclipse measurement, with the class audit and the correlated-failure probe beside it); an interior maximum of the growth exponent at or below two across the reinvestment levels, under the replication-gated rule (the titration); and the fourth-cell substrate test.

Proposed, and to be registered before any data are seen: that the ratio rises with load (dose-response, which separates a mechanism from a static difference); that a cross-class corrector moves the titration's maximum upward, linking Laws II and III through the class lever; and that a single system crosses the stability boundary as its reinvestment dial turns. None of these three is a prediction of the field until its registration carries it, and this paper does not count them as such.

The consistency check no single field can fake, with invented numbers for illustration only: suppose the eclipse measurement returned a same-class gamma of 0.45 and a cross-class gamma of 0.70. The Ceiling would then predict a same-class maximum near 1.8 and a cross-class maximum near 3.3 in the titration. If the titration returned its maxima in those places, two registered studies with different instruments would agree through the one bound. That agreement, in the manner of Joule's measurements agreeing with Carnot's bound, is the founding result: it is inexpressible in the variables of learning theory, control, scaling empirics or oversight.

The founding conjunction, then: a Law II regime; a class ratio above 1.00 with the one-half cap on same-class correctors; a maximum at or below two; the same exponent structure in one non-artificial substrate; and at least one of these replicated by strangers. Each alone is a finding in an old field's vocabulary. The conjunction is not, and that is the exact sense in which results would confirm a new area rather than a new chapter.

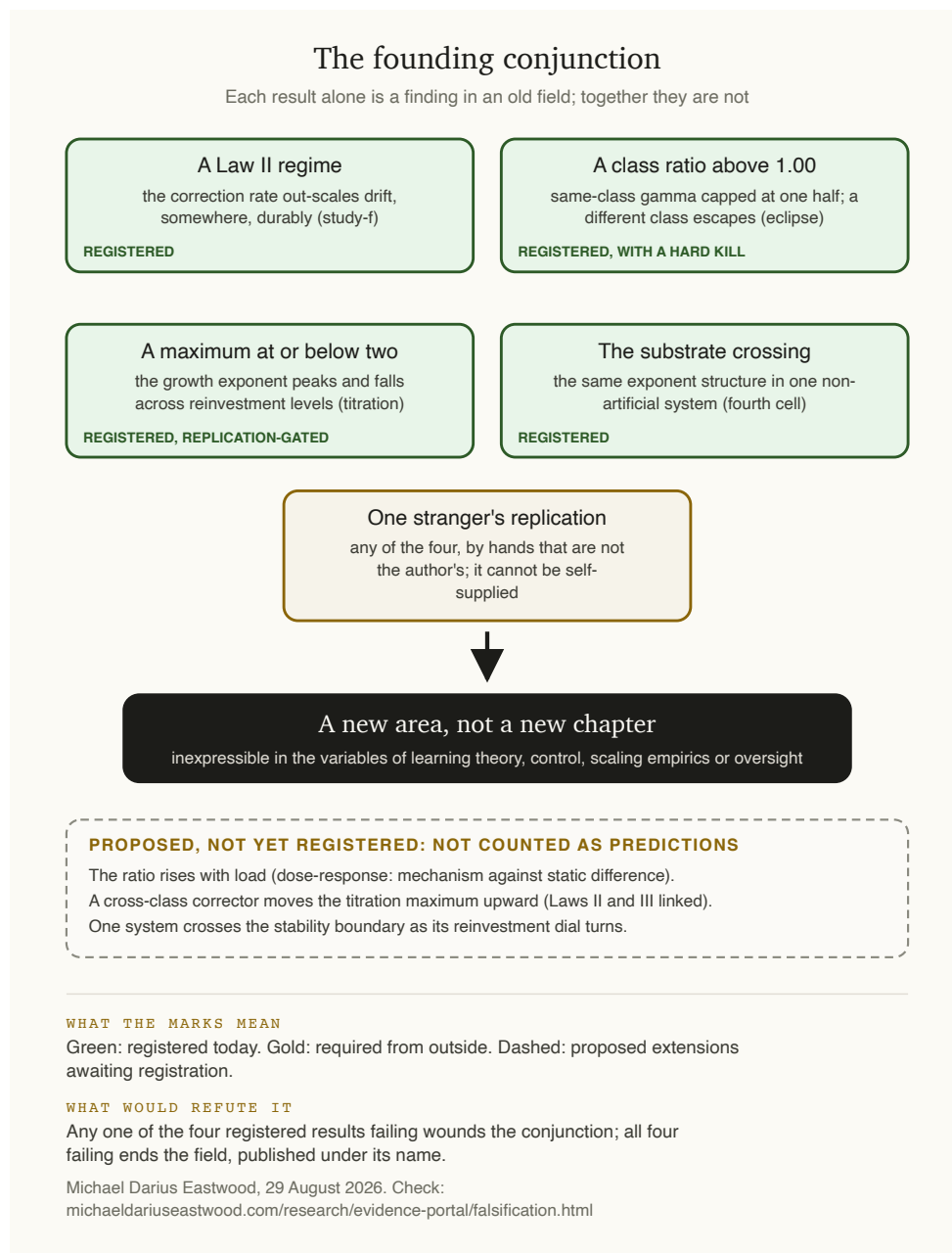


Figure 5: The founding conjunction: four registered results and one stranger's replication converge on a statement no neighbouring field can make; the proposed extensions are marked as not yet registered.

11. Every objection we could find, with its disposition

This section is the paper's reason for existing at version 2.0. The objections below were collected from the programme's own adversarial readings, from outside reviewers, from the published criticism of the derivation, and from the author's attempt to argue the field out of existence. Each is stated in its strongest form. Each carries one of five dispositions: **conceded** (it is true and it is printed); **rejected** (with the reason); **partly conceded**; **reading debt** (a

literature that must be read before the point can be argued either way, under the rule that nothing is cited before it is read); or **decided by experiment** (a registered decider settles it). Each carries the thing that would change its disposition. Fifty-one objections: thirty-nine conceded or partly conceded, ten rejected, two handed to experiment or reading.

A. Redundancy: the field already exists under another name

A1. Cybernetics already studies self-correcting loops; Wiener's field is this field. Rejected, with a condition. Cybernetics owns regulation: a designed controller holding a plant to a reference, with the loop's purpose set from outside. The loop studied here rewrites the controller's own capability, and the questions it raises (does correction out-scale drift; is there a ceiling and whose property is it) are not cybernetics' questions. *What would change it:* if those questions turn out to be answerable inside cybernetics with nothing added, the field dissolves into it, and that condition is printed in section 14.

A2. Dynamical systems theory already has stability conditions; the second law is a standard stability statement. Conceded, with a boundary. The mathematics is borrowed and is claimed as nothing more. What is claimed is which measurable quantities carry the condition in built systems, and that they can be measured. *What would change it:* if the variables reduce to standard ones with nothing added under measurement, the second law is a relabelling, and that is a dissolution condition.

A3. Scaling laws already describe how capability grows with resources. Conceded, with a kill condition. Scaling empirics owns capability against external resources and is architecture-blind by construction; the variable here is the system's own re-entering output, and the field predicts something resource scaling does not: an interior maximum in the growth exponent as reinvestment rises, where resource scaling predicts monotone returns. *What would change it:* if recursive depth reduces, under measurement, to compute spent, or the titration shows monotone returns with no ceiling behaviour, Recursive Dynamics is scaling empirics under a new name, and this paper says so.

A4. Scalable oversight is already the field of overseers; this is a corner of it. Partly conceded. Scalable oversight builds correctors and ladders of them, and the programme conceded its antecedent standing in the references register. What it lacks, and does not claim to have, is a theory of what bounds a corrector: it builds same-class ladders on the implicit premise that they scale. The third law is precisely a claim about that premise, and it is the claim the eclipse measurement decides. *What would change it:* a bound on corrector capacity derived inside the oversight literature that this field merely renames.

A5. The literatures on self-training degradation and on models trained on their own outputs already measure drift. Reading debt, with the shape of the answer already visible. A real neighbour, and possibly the owner of the drift variable's measurement in special cases; it has not been read closely by this programme and is not cited here until it is. What those

measurements cannot do, on their own description, is state the ratio that decides stability, because they measure drift without a correction rate beside it and do not vary the corrector's class. *What would change it:* if that literature carries both rates, the field's contribution narrows to the class lever, and the paper narrows with it.

A6. The intelligence-explosion literature has asked these questions since 1965. Conceded, and recorded as a debt in section 13. The question is inherited from its named owners; the field proposes the missing state variables and the instruments, not the question. *What would change it:* nothing; the concession is permanent.

A7. Evolution is the original recursive improver and evolutionary dynamics already has its theory. Partly conceded. Evolutionary dynamics has variation and selection; it has no designer directing reinvestment into the improvement process itself and no correction budget in the sense used here. The field's claim to generality includes evolved systems only where a corrector exists, and section 16 says what falls outside. *What would change it:* a treatment of designed reinvestment and correction leverage already standard in evolutionary theory.

A8. Statistical physics already has a theory of scaling and universality; this is renormalisation with new names. Partly conceded. Universality classes and scaling collapse are the right ancestors for the field's substrate-independence hope, and the functional-equation machinery is conceded as machinery. What that tradition does not carry is a correction variable, a drift variable, or a corrector whose composition is the lever. *What would change it:* a universality argument that derives the corrector-class distinction, or its absence, from first principles, which would be a result this field would adopt gratefully.

A9. AI safety as a field already covers this. Rejected. AI safety is a problem area, not a science of a loop; it borrows from every neighbour above and owns none of their objects. This field is a proposal for one of the sciences safety would draw on, in the way thermodynamics is a science engineering draws on. *What would change it:* a published set of state variables and laws for self-improvement, with deciders, already standing under the safety banner.

A10. This is systems theory, rebranded. Rejected. Systems theory carries no correction budget, no class cap on the corrector, and no registered decider; a rebranding does not carry kill conditions. *What would change it:* a systems-theoretic treatment of a corrector's composition class that this field merely renames.

A11. The demarcation never tried the fit; it dismissed each neighbour with a sentence. Conceded for every version of this paper before 2.3, and answered by section 2.1, which attempts to house the questions in twelve candidate fields in their own variables, concedes Law I to endogenous growth theory, the form of the ARC Ceiling to queueing theory and the shape of Law II to error-threshold theory, prints the dissolution map by host, and isolates the one coupling no host can state. *What would change it:* a host that can state the class-

dependent cap coupled to the depth exponent in its own variables, which would end the field by adoption rather than by refutation.

A12. Recursive Dynamics is the ARC Theory renamed: the same three laws, the same author, a larger word. Rejected, with a test. Conceded that the field and its founding theory share an author, a date and a record. Rejected that they are the same object. The theory is three answers; the field is the questions, the five state variables and the instruments that would refute those answers, and it is defined so that a reader who rejects all three laws and keeps the measurement is inside it, by the minimal commitment of section 16. *What would change it:* the test is separability, and the dissolution map of section 2.1 is where to run it: if no question, variable or instrument survives the deletion of the laws, the objection is right, the field is a rebadging, and this paper will say so.

B. Standing: the field has no results

B1. Thermodynamics had many hands and a century; this has one person and zero replications. Conceded in full, before any reader could raise it. The parallel is in shape only; standing is not claimed, and section 15 says what would earn it.

B2. The laws are not laws; they are conjectures from a minimal model. Conceded. They are named conjectures, and the status travels with the name on every surface. Law is the naming convention for a relation the field proposes to test.

B3. The programme's own headline number was retracted; why trust the rest? Conceded as fact, rejected as inference. The retraction is printed wherever the number once stood, and the discipline that produced it (blinding, registered kill conditions, public correction within the hour of finding) is the reason the remaining claims carry statuses rather than assertions. A programme that never retracts has either never been wrong or never checked. *What would change it:* a retraction concealed rather than printed, anywhere in the record.

B4. The correction leverage gamma has never been measured, by anyone; the ARC Ceiling's number is therefore a guess. Conceded, and it is the field's single most consequential open measurement (section 10). The number two is conditional on a premise the programme names as the conjecture's most likely point of failure. *What would change it:* the eclipse measurement and the hard kill riding alongside it.

B5. The derivation's treatment of depth has been criticised by an outside reader. Conceded and printed as a limitation, not argued away. *What would change it:* a treatment that survives the criticism, which is wanted and invited.

B6. Carnot demonstrated his bound inside the memoir; the ARC Ceiling is a conditional model result whose key premise is unproved. Conceded, and the difference is exactly why the parallel is confined to shape. The Ceiling's premise is measurable and registered. *What would change it:* nothing about the history; only the measurement.

C. Method: the instruments may not measure what the laws name

C1. The five state variables are arbitrary; why not four, or nine? Conceded as open. The claim that five suffice, and that these are the five, is part of what the field must prove, and the notation register already records the programme catching a sixth symbol trying to enter. *What would change it:* a registered study in which one variable reduces to another under measurement, or a relation that cannot be stated without a sixth.

C2. The construct bridge is open: the quantities the studies measure may not be the objects the laws' symbols name. Conceded and held open in the register by name. *What would change it:* the bridge closing in either direction, which is itself a registered finding.

C3. Substrate independence is just the flexibility of the fitted families. Conceded as the live risk. The fourth-cell test exists to catch exactly this. *What would change it:* the fourth cell, run.

C4. Models scoring models is circular: the scorer shares the blind spots of the scored. Partly conceded (section 9). Laundering and the forced null remove the scorer's knowledge of provenance, not its class; cross-family panels and the registered human-criterion validation reduce and measure the circularity without abolishing it. *What would change it:* a validation result showing model panels and the human criterion diverge systematically, which would demote every panel-scored result to exploratory.

C5. The capability ladder is itself unvalidated; a wrong ruler makes every law wrong. Conceded as a dependency. The ladder is calibrated and its scale discipline is a founding principle precisely because the programme once got it wrong; but a calibrated instrument is not a validated one until independent hands use it. *What would change it:* replication of ladder measurements by other hands.

C6. The author-rated fallback for the human criterion is an author grading his own theory. Conceded, with the safeguards printed: blind, laundered, washout-retested, disclosed as author-rated in every headline, and registered as a fallback rather than a design. *What would change it:* external raters, which the registration names as the preferred path the moment they exist.

C7. The registered experiments are underpowered or too small to decide anything. Decided by experiment, with the numbers printed. Power is computed in the registrations from pilot dispersion, and an inconclusive verdict is a registered outcome with its own branch, not a failure to be hidden. *What would change it:* an inconclusive verdict, which would narrow the claim to "not yet decidable at this scale" and say so.

C8. Frontier models change faster than the experiments can run; results will be obsolete on arrival. Partly conceded. A law of the loop is meant to hold across models; if a measured exponent moves with every model generation, that is itself evidence against

substrate independence and would be reported as such. *What would change it:* the follow-up grid across families, which locates whether the effect is a property of the loop or of a generation.

D. Naming and standing of the author

D1. Naming a discipline before the evidence is a marketing move. Rejected, with the precedent printed: cybernetics was named in the 1948 book that proposed it and exobiology in 1960 before a single specimen. A name makes a proposal citable and attackable as one unit. The safeguard is the name's own kill condition, printed in section 14: the name dies with its object, not with any one law's stated form, and the record says which death it was, in public.

D2. One person cannot found a paradigm; paradigms are conferred by communities. Conceded. What is claimed is a candidate: a proposed way of organising the questions. Whether it becomes anyone's paradigm is decided by adoption and replication, which are two of the entry criteria, and neither is the author's to declare. This paper withholds the word paradigm as a description of the present.

D3. The author has no institution, no credential and no co-authors. Conceded as fact; rejected as an argument about the proposal. Every claim links its dated record, every number its register entry, every instrument the test that could kill it, so that nothing here requires trusting the author. Independence has one advantage that is also printed: there is no institution to protect, so corrections publish the day they are found. *What would change it:* nothing; the work is judged by its record.

D4. The work was produced with AI assistance; the field is an artefact of the tools. Conceded as to the tools, rejected as to authorship of the ideas. The programme discloses AI as a tool throughout; the framework's dated record predates the tools' involvement in the analysis, and AI systems cannot hold priority, so the record names its human author. The tools are also the object of study, which is disclosed as the conflict it is.

D5. The phrase recursive dynamics is already used in robotics. Conceded: the Featherstone tradition's recursive dynamics algorithms are a different referent, a family of computational methods, and no priority over the phrase as words is claimed against anyone. Nor is this recursion theory, the mathematics of computability.

E. Scope: too broad, or too narrow

E1. Too broad: any self-improving system includes science itself, economies, evolution, so the field is either everything or nothing. Partly conceded. The field's laws switch on where one narrower thing exists: a corrector, something detecting and repairing its own errors while the system grows. A star has dynamics and no drift ledger; nothing is being raised there. Section 16 draws the boundary. *What would change it:* the laws holding only

for one substrate, which would make the field narrow rather than everything, and honest about it.

E2. Too narrow: this is really about language-model loops, dressed as a general science.

Partly conceded. The decisive experiments run on language-model loops because those are the loops that exist to be measured; the substrate-independence claim is at exploratory grade and its fourth-cell test is registered. *What would change it:* the cross-domain families failing the fourth cell, after which the field would say it is a science of one substrate until shown otherwise.

E3. The field conflates correction with alignment, and capability with intelligence.

Conceded as a risk and answered with definitions: correction is the detection and repair of the system's own errors against a stated criterion, which may or may not be an alignment criterion; capability is whatever the calibrated ladder measures. The field makes no claim about intelligence as such. *What would change it:* a demonstration that the ladder measures something other than what the laws need it to measure, which is C5 again.

F. Timing and community

F1. It is premature; wait for the data, then name the field. Rejected. The deciders are registered; the instruments exist; the name lets the proposal be attacked as one unit before the data arrive, which is when attack is cheapest and most useful. The name carries its own kill condition for exactly this reason.

F2. A field needs a community first; a field of one is a hobby. Conceded as to the definition of a science; rejected as to the definition of a proposal. Section 6 scores condition 6 as not met. The proposal exists so that a community can form around a testable object, or decline to.

F3. Naming it invites priority disputes rather than science. Rejected. Every antecedent is conceded by name in the references register and section 13; nothing is claimed as original except the assembly, and the assembly is dated. A dispute about priority would be settled by the dated record, not by argument.

G. The concept itself

G1. The crossover $\alpha_{crit} = 1/(1 - \gamma)$ is a triviality of power-law algebra, not a law. Partly conceded. The algebra is elementary and is printed so that nobody need take it on trust. The content is not the algebra; it is the claim that the two premises describe real loops and that γ is capped for same-class correctors. *What would change it:* nothing about the algebra; the premises are what the experiments test.

G2. The same-class cap at one half is an assumption dressed as a result. Conceded, and named by the programme as the single most likely point of failure. It is the eclipse measurement's target and carries a hard kill. *What would change it:* the measurement.

G3. The model assumes stationarity and separability of drift and correction; real loops have neither. Partly conceded. The minimal model is minimal on purpose; the registered analyses fit drift and correction jointly and report the interval, and non-separability would show as an unfit rather than being averaged away. *What would change it:* a registered result in which drift and correction cannot be fitted jointly, which would be reported as a limit of the model.

G4. Capability is multidimensional; a scalar ladder throws the structure away. Conceded as a simplification. The ladder is one axis chosen for measurability; the field's claim is that the laws hold along it, not that it exhausts capability. *What would change it:* the laws holding on one axis and failing on another, which would be a finding about the field's scope.

G5. The ceiling ignores economics: a lab can buy correction with money, not just with a better corrector class. Partly conceded. Buying correction raises the coefficient a in the derivation, which moves the crossover's location in depth but not its exponent; the ceiling is a statement about exponents. *What would change it:* a demonstration that purchased correction changes the exponent rather than the coefficient, which would be a real refutation of the model's structure.

G6. The dissolution clause is cheap talk; nobody publishes their own field's ending. Rejected on the record. The programme has already retracted its headline number, withdrawn a formulation of its third law, and demoted a favourable result to a pilot within thirteen minutes, all in public. The commitment to publish the field's ending is made by a record that has already published smaller endings.

G7. Because every claim is conditional, nothing is actually claimed. Rejected, and answered up front in section 1.1 so that the concessions below cannot be mistaken for silence. Conditionals with measurable antecedents are claims: they forbid outcomes. The field forbids a stable system holding its growth exponent above the ceiling for a sustained run, a same-class correction exponent measured wholly above one half, and a corrector-class ratio pinned at one where the cap binds. Any of those would be a real result against it.

G8. Depth is not time, and the laws confuse them. Conceded as a hazard and answered by the register: every quantitative surface names its model family, and the depth studies treat depth as the count of re-entries at a fixed wall-clock policy. *What would change it:* a registered result in which the exponent depends on the clock rather than the count, which would be reported as a limit of the variable.

G9. The tie regime is undecided, so the bound decides nothing near the boundary. Conceded and printed: at $\alpha(1 - \gamma)$ equal to one the outcome hangs on coefficients, delays and saturation, which is the regime where every interesting failure lives. The bound is a statement about exponents, and near equality it hands over to the coefficients by design.

H. Practice, policy and the height of the bar

H1. Cross-class correction is already practice: ensembles, verifiers, humans in the loop. Conceded as practice, rejected as theory. Practice has not measured whether same-class ladders cap at a correction leverage of one half, or whether a different class escapes; the field turns a habit into a specification with a number. *What would change it:* a published measurement of the class cap from inside the practice.

H2. If it is true it is trivial. Rejected. The predictions are risky and specific: a class-dependent ceiling on correction leverage at one half; an interior maximum of the growth exponent at or below two; a ratio pinned at one by construction that the field says will move. None is trivially true and each can fail.

H3. It cannot be shown wrong. Rejected. Four dissolution results are printed in section 14, each tied to a registered study, with hard kills and rival-favouring nulls, and the field's name carries its own death clause.

H4. Policy relevance is premature; talking about regulators is safety-washing. Conceded on timing. Policy weight follows results, never formulation; section 12 is written in the conditional for that reason, and the one paragraph on the programme's public pages that mentions oversight instruments says the same.

H5. Even if all three laws hold in language models, that is AI engineering, not a field. Conceded as the demanding bar, and it is the right bar. A field needs the substrate crossing: the same exponent structure in at least one non-artificial system. That is why the fourth-cell test is part of the founding conjunction in section 10.1 and not an optional extra.

I. What is not answered here

I1. There may be objections this list has missed. Conceded, by construction. The list is published so that it can be extended, and every objection added from outside will be printed with its disposition in the next version, under the standing invitation that ends this paper.

12. What this buys, if it holds

Fields earn adoption when strangers can compute with them. If the laws survive their deciders, the following become calculations rather than debates: given a measured gamma for a corrector class, the maximum sustainable growth exponent of any system it oversees follows from the ARC Ceiling; given measured drift, the required correction rate follows from Law II; oversight architectures acquire a specification sheet (the corrector's composition class, its measured gamma, its implied ceiling) the way engines acquired efficiency ratings; and the safety question "can oversight keep up" becomes, for the first time, an empirical parameter comparison rather than a scenario argument. Everything in this paragraph is conditional on results that do not yet exist, and is written in the conditional for that reason.

12.1 And if it fails

The negative result is also a contribution, and it is one the oversight programme needs. A ratio of 1.00 holding tight everywhere licenses same-class ladders; a universal absence of any regime where correction out-scales drift says self-correction never wins, which is a different and worse world that policy should know about. The registrations commit to publishing either outcome under the same title. A proposal whose failure is informative is not a vanity project.

13. Debts, conceded by name

This field is assembled from debts recorded in the programme's antecedents register and conceded in full: the impossibility line of AI-control research (this field adds the design requirement the impossibility implies, rather than resignation); the transport-network limits of biological scaling (conceded as premise; this field adds the internal replacement limit that operates when fixing the pipe no longer fixes the growth); the intelligence-explosion question as posed by its named owners from 1965 onward (conceded; this field proposes the missing state variables); the self-modifying-program line of the 2000s (to be recorded in the register before it is cited by name here, under the reading rule); the functional-equation mathematics (the machinery, never claimed); and the scalable-oversight programme, whose ladders are the object the third law is about. What is claimed as original is the assembly: the loop as native object, scale discipline as founding method, the bounds as laws, the corrector's class as the lever, and the registered instruments to decide all of it.

14. What would dissolve the field

If the correction-versus-drift race finds no regime where correction out-scales drift, Law II has no subject matter. If the corrector-class ratio holds at 1.00 with tight intervals wherever the cap should bind, the field's distinctive mechanism is dead and its remains belong to ordinary scaling empirics. If the titration finds no interior maximum and no ceiling behaviour, the central impossibility dissolves. If the cross-domain families fail the fourth-cell test, substrate independence was an artefact of family flexibility. Any one death wounds the field, and this paper commits to publishing each under this title.

A name proposed before its tests must be able to die with them, and the record must be able to say which death it was. A law failing in its stated form kills that law and the theory that stated it; the field carries the corrected form, or none, as the record allows, and a reader who rejects all three laws and keeps the measurement is still inside it, by the minimal commitment of section 16. The name itself dies when the object goes: if the corrector-class ratio holds at 1.00 with tight intervals wherever the cap should bind, the distinctive mechanism is dead, the residue no host could state has proved empty, every piece returns to its host by the dissolution map of section 2.1, and this paper will say so under this title

rather than quietly disappear. If the laws survive their registered tests and hold outside the domains that raised them, the name has earned its generality. Either way the record decides, not the naming.

15. What would make it a science

Naming a field does not create one. The entry criteria this proposal has not yet met, in public:

Criterion	Status	What would meet it
Registered prediction success	Not met	One decisive experiment returning its predicted direction under its registered analysis
Independent replication	Not met, dated	The same result from hands that are not the author's; every instrument, seed, schedule and threshold ships in the open for this reason
Use by strangers	Not met	The state variables appearing in other people's work because they needed them
Survived audit	In progress	The record of public retractions, minute-resolution self-demotions, adversarial cold reads and printed limitations, offered as the standard future claims accept

16. What the field is not, and the minimal commitment

The field is not a forecast of when machines surpass people, not a theory of intelligence, not an ethics, not a theory of everything, and not a claim about consciousness. Its laws switch on where a corrector exists inside a growing system and switch off where none does; the composed hierarchies that its scaling mathematics bears on (stars included) are graded as correspondence of form and no further. The programme's speculative coda on recursion at

cosmological scale is separable and sits outside the field; a reader can reject it entirely and lose nothing here.

The minimal commitment required to work in the field is small and stated so that it can be refused: that the self-improvement loop is a legitimate object of measurement, with quantities that can be defined, instruments that can measure them, and relations that can be shown wrong. A reader who accepts that and rejects all three laws is inside the field; a reader who accepts the laws and refuses measurement is not.

17. Closing

Until the entry criteria are met, one sentence carries the field's whole status, and it appears wherever the field's name does: **Recursive Dynamics is a proposed field at the Carnot stage: one memoir in, instruments built, decisive measurements registered and not yet run, staked in public where failure will be as visible as success.**

The invitation that follows from everything above is the only one the field can honestly extend: run the test that could kill it. Every instrument, seed, schedule and threshold is published for that purpose, and an independent run outranks the programme's own.

Version history: v1.0, 28 August 2026, first draft. v1.1, same day, adding the worked derivation of the Ceiling, the measurement-scale founding principle, the steelmanned rival with the programme's contrary pilot, the what-this-buys section, and the registered numbers of the decisive experiments. v1.2, 29 August 2026, the capability ladder restated in the registers' canonical forms, three further neighbouring fields named with what each owns, historical precedents held to the well-attested set. v1.3, same day, the state variables written exactly as the notation register defines them (the correction rate β_C replaces a bare reinvestment β ; reinvestment named as the experimental manipulation). v2.0, same day, the case for the field scored against six licence conditions, forty-two objections collected and dispositioned, the scorer-circularity limitation printed, what the field is not and the minimal commitment, and the entry criteria as a scored table. v2.1, same day, the precedent class of limit-theorem fields (6.1), the founding conjunction with the cross-instrument consistency check and the registered-versus-proposed separation (10.1), what the field buys if it fails (12.1), and eight further objections (A10, G8, G9, I1 to I5) with dispositions. v2.2, same day, the forbidden outcomes stated up front in the abstract and section 1.1 (an outside reader's one reservation, that the concessions could read as silence), and five figures: the loop, the crossover, the scorecard, the ratio zones, the founding conjunction. v2.3, same day, section 2.1: the honest attempt to house the questions in twelve neighbouring fields, the concessions that attempt forces, the dissolution map by host, the residue no host holds; objection A11. Figures rebuilt the same evening from a layout contract (measured text budgets, arrows on box edges, labels off the paths) after the first artwork overflowed its boxes, and verified by reading the rendered pixels before use. v2.4, 30 August

2026: the retired $\alpha_{\max}$ form in section 4 replaced by the notation register's α_{crit} ; the scorecard figure's condition 1 evidence line aligned with section 2.1 (twelve hosts tried); the crossover figure's label written α_{crit} . v2.5, 30 August 2026: Law II carries its canonical name, the ARC Co-Scaling Law, and Law III is written the ARC Ceiling at every mention in the body (the bare forms were retired estate-wide on 30 August 2026; this history keeps its own earlier wording); the objection groups run A to I without a gap (the former group I is now H, the former J is now I) so that the paper and the field's public page cross-reference by the same letters, and the earlier history entries keep the letters current when they were written. No claim, status or disposition changed. v2.6, 30 August 2026: one further objection in group A, raised by the author against his own proposal (that Recursive Dynamics is the ARC Theory renamed), rejected with a separability test; fifty-two objections, eleven rejected. The name's kill condition in section 14 corrected: the earlier wording had the name die with the laws, which contradicted the minimal commitment of section 16 (a reader who rejects all three laws and keeps the measurement is inside the field) and the separation of questions from answers; the name now dies with its object, the class-set cap on correction coupled to growth in depth, and the record says which death it was. D1 aligned with the corrected clause. The field's public page carries the same objection and clause, with a decision figure of the kill condition, from the same date.