



Chaotic dynamics in life-cycle inventory pedigree: a regime-label diagnostic complementary to the ILCD log-normal closure

Theo Piens

► To cite this version:

Theo Piens. Chaotic dynamics in life-cycle inventory pedigree: a regime-label diagnostic complementary to the ILCD log-normal closure. 2026. <hal-05616988>

HAL Id: hal-05616988

<https://hal.science/view/index/docid/5616988>

Preprint submitted on 8 May 2026

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



Distributed under a Creative Commons CC BY 4.0 - Attribution - International License

Chaotic dynamics in life-cycle inventory pedigree: a regime-label diagnostic complementary to the ILCD log-normal closure

Theo Piens, Ph.D.

*Environmental Researcher | Waste Remediation Technology Specialist, France. Contact:
contact@theopiens.com · theopiens.com*

License: CC-BY 4.0; reference engine MIT (Zenodo DOI on first release)

Chaotic dynamics in life-cycle inventory pedigree: a regime-label diagnostic complementary to the ILCD log-normal closure

Theo Piens, Ph.D. *Environmental Research · Waste Remediation Technology Specialist · France · contact@theopiens.com*

Preprint, version V1 (post-review), 8 May 2026. Reference implementation (Python, MIT) and benchmark CSV: Zenodo DOI to follow first stable release. Web demonstrator: theopiens.com/analyzer (Piens 2026, ref. [51]). Companion technical note for JRC / ecoinvent Methodology WG: Piens (in prep., ref. [50]).

Abstract

English. Life Cycle Assessment practitioners express inventory uncertainty through the ILCD pedigree matrix, which maps five qualitative scores (R, C, T, G, F) $\in \{1..5\}^5$ to an aggregated geometric standard deviation σ_{ILCD} under a log-normal closure. The closure is operational and well-calibrated at high data quality but structurally limited at the data-quality tail (a hard ceiling on the aggregate σ ; per-flow, hence blind to the coupling structure of a Leontief-inverted inventory). We propose a **regime-label diagnostic** complementary to the closure. A deterministic single-knob mapping $\sigma_{\text{ILCD}} \rightarrow r \in [2.5, 4]$ places each flow on the regime axis of the logistic recursion $x_{\{n+1\}} = r \cdot x_n \cdot (1 - x_n)$, used as a *probe* of regime, not as a model of inventory propagation. The trajectory yields a Lyapunov exponent and a five-component **Chaotic Risk Score (CRS)** in $[0, 1]$ that summarises a regime label (point fixe / période-2 / pré-chaotique / chaotique). A `pedigree_uncertainty` Monte-Carlo module exposes per-dimension first-order Sobol indices. The framework is engineered to coincide with ILCD at high data quality ($\lambda \leq 0 \Rightarrow \text{GSD}_{\text{eff}} = \sigma_{\text{ILCD}}$ exactly).

Empirical evaluation, post-review. A controlled 493-flow benchmark (§ 5.6) and a 310-flow tiangongDB-derived real-distribution extension (§ 5.7) — Low-quality evaluation $n \geq 125$ flows in each — show that neither single-knob naive widening nor CRS-augmented GSD_{eff} outperforms ILCD on single-flow CRPS at low quality. CRS *equals* ILCD on 88–91 % of Low flows by the engineered λ -property and *loses* on the remainder. A naive linear widening of σ_{ILCD} does beat ILCD on Low by 4–9 % CRPS gain, but is not dynamical. We position CRS as a **regime-label diagnostic**, not a per-flow GSD substitute. A 5×5 Leontief mini-case (§ 6.5) shows the per-flow ILCD aggregation is *conservative* by an order of magnitude on a well-conditioned operator, nuancing the multi-flow motivation of § 1.2. Robustness ablations on the r_{min} endpoint and on R-T correlation in the Sobol indices preserve the headline conclusion. The reference engine (V5) is open source; V6 Nexus extensions are documented and form the subject of a companion paper.

Français. Les analystes ACV expriment l'incertitude des données d'inventaire via la matrice pedigree ILCD, qui traduit cinq scores qualitatifs (R, C, T, G, F) en un écart-type géométrique agrégé σ_{ILCD} sous fermeture log-normale. Cette fermeture est opérationnelle à haute qualité mais

limitée à la queue (plafond sur σ_{ILCD} ; per-flux, donc aveugle au couplage Léontief). Nous proposons un **diagnostic d'étiquette de régime** complémentaire de la fermeture. Un mapping déterministe à un seul paramètre $\sigma_{ILCD} \rightarrow r$ place chaque flux sur l'axe de régime de la récurrence logistique $x_{n+1} = r \cdot x_n \cdot (1 - x_n)$, utilisée comme *sonde* du régime, non comme modèle de la propagation. La trajectoire fournit un exposant de Lyapunov et un **Score de Risque Chaotique (CRS)** à cinq composantes dans $[0, 1]$. Le cadre coïncide avec ILCD à haute qualité par construction ($\lambda \leq 0 \Rightarrow GSD_{eff} = \sigma_{ILCD}$).

Évaluation empirique post-révision. Un benchmark contrôlé de 493 flux pedigree (§ 5.6) et une extension de 310 flux dérivés des métadonnées de tiangongDB (§ 5.7) — strate Low $n_{eval} \geq 125$ dans chacun — montrent que ni un élargissement naïf à un seul levier ni le GSD_{eff} augmenté de CRS ne battent ILCD sur le CRPS mono-flux à basse qualité. CRS reste identique à ILCD sur 88-91 % des flux Low par construction et perd sur le reste. Un élargissement linéaire naïf de σ_{ILCD} bat ILCD de 4-9 % de gain CRPS sur Low, mais sans contenu dynamique. Nous positionnons CRS comme une **étiquette de régime**, non comme substitut au GSD per-flux. Un mini-cas Léontief 5×5 (§ 6.5) montre que l'agrégation ILCD per-flux est, sur ce système bien conditionné, *conservatrice* d'un facteur $\sim 50\times$, ce qui nuance la motivation multi-flux du § 1.2.

Keywords. Life Cycle Assessment · ILCD pedigree matrix · uncertainty quantification · nonlinear dynamical systems · logistic map · Lyapunov exponent · Feigenbaum cascade · Kaplan–Yorke dimension · CRPS · Sobol sensitivity · ecoinvent · EF 3.1 · EN 15804+A2 · environmental decision-making.

1. Introduction

1.1 Why uncertainty in LCA is a human and environmental question

A Life Cycle Assessment is, in the end, a quantitative argument that one product or process imposes a smaller environmental burden than another. People read it. Eco-designers, public buyers, certifiers, regulators, and increasingly the end customer scanning a carbon label on a packaging, they all turn the report's conclusion into a decision. When the conclusion is wrong in a way the reader cannot see, a real environmental choice is being taken on a false premise. Expressing *uncertainty* (the residual ignorance behind every inventory number) is not, in that light, a methodological refinement for journals: it is basic honesty owed to the people downstream of the report.

The discipline has converged, since the early 2000s, on a pragmatic and elegant tool for that honesty: the **pedigree matrix** introduced by Weidema and Wesnæs (1996) and operationalised in ecoinvent (Frischknecht et al., 2007; Ciroth et al., 2013) and in the JRC's ILCD handbook (2010). Five qualitative scores rate the *representativeness* R, *completeness* C, *temporal correlation* T, *geographical correlation* G, and *further technological correlation* F of an inventory flow. Each score is mapped to a per-dimension multiplicative factor σ_i from

a fixed table; the aggregated geometric standard deviation $\sigma_{ILCD} = \exp(\sqrt{\sum_i \ln(\sigma_i)^2})$ becomes the GSD of a log-normal distribution that is propagated downstream, typically through a Monte-Carlo simulation of the impact assessment.

It is hard to overstate the virtues of this closure. It is *operational* (a non-specialist can score a flow in a few minutes), *international* (the σ -table is shared and audited), and *parsimonious* (one number per flow flows through the rest of the analysis). For data of high quality (pedigree 1–2) it is also empirically defensible: the residual variability of well-instrumented industrial inventories is plausibly log-normal (Slob, 1994; Limpert, Stahel and Abbt, 2001), and the multiplicative composition of factors close to unity yields a tight, well-behaved aggregate. The closure does what it was designed to do, and it does it well.

The trouble starts at the other end of the data-quality scale, where the closure was never intended to operate. That is where this paper enters.

1.1bis Why waste-LCA is the natural test bed

This paper grew out of LCA studies on waste-treatment chains (anaerobic digestion, MSW incineration, post-closure landfill leachate, bioremediation, eco-material recovery for construction) and it is primarily addressed to that community. There is a reason for this. Of all the sub-fields of LCA, *waste-LCA* is where the limits of the standard pedigree closure are most visible to the practitioner who has to live with them. Three reasons, briefly.

(W1) The input is heterogeneous, and you cannot fix it. Municipal solid waste, industrial hazardous waste, demolition residues and biogas substrates have compositions that vary by city, by season, by upstream sorting policy, and by collection scheme. A single mass-balance row in an LCI of a waste-treatment plant typically aggregates dozens of measurement campaigns whose individual coefficients of variation routinely exceed 30 %. Pedigree T (temporal) and pedigree G (geographical) for such rows are almost never scored below 3, and rightly so.

(W2) The treatment kinetics are non-linear, by design. Anaerobic digestion biogas yield depends non-linearly on substrate mix, hydraulic retention time, temperature, and inhibitor concentration. Landfill leachate composition evolves through hydrolysis, acidogenesis, acetogenesis and methanogenesis phases that overlap and feed back into each other. Bioremediation rates follow Monod-type saturation in the contaminant concentration. The variance of the output flow is, in all these cases, a non-linear function of the variance of the input pedigree. That is exactly the regime in which the per-flow log-normal closure has the weakest grip.

(W3) Allocation rules add their own uncertainty, and the closure does not see it. Multi-output waste-treatment systems (cement-kiln co-incineration, energy-from-waste with both heat and electricity, recycling chains with quality cascades) force the practitioner to pick

an allocation key (physical, economic, system-expansion, cut-off). That choice is itself a pedigree-3-or-worse decision in most public LCA studies. The per-flow GSD never registers it.

The framework in this paper is engineered to live with (W1)–(W3) without breaking compatibility with EN 15804+A2, PEP Ecopassport, ISO 14040/44 or PEFCR. Section 6.4 walks the framework through three waste flows that motivated this work (biogas yield from a real anaerobic-digestion plant, the NO_x emission factor of a modern MSW incinerator with SNCR, and the post-closure leachate COD from a municipal landfill) to make the operational read-out concrete.

1.2 What the log-normal closure cannot say

Two structural limits become visible as the data-quality tail thickens.

First, the closure has a hard ceiling. Theecoinvent v3 σ -table caps σ_R , σ_T at 1.502 and 1.499 at level 5 (and lower for C, G, F); even when all five dimensions are scored 5/5, σ_{ILCD} cannot exceed roughly 2.3 by construction. Empirically, however, real low-quality flows (expert estimates of fugitive emissions, geographical extrapolations across continents, technologies known only by analogy) routinely exhibit measured CV well beyond the upper bound the table can produce. The closure is *self-truncating*.

Second, and more interestingly for this paper, the closure is *per-flow*. It says nothing about how the variance of a flow propagates through the *coupled* inventory system. An LCI is solved by inverting a technology matrix A and applying it to a final-demand vector f (Leontief, 1936; Heijungs and Suh, 2002): $s = (I - A)^{-1} \cdot f$. When A itself is uncertain (and most of its non-zero entries are pedigree-rated), the variance of any output entry is a non-trivial function of the joint distribution of the input entries. In the linear-additive (small-perturbation) regime, the ILCD per-flow GSD remains a defensible first approximation. Outside it, when individual variances are large *and* inventory rows share common drivers (a shared electricity grid, a shared transport mode, a shared upstream chemical), the propagated uncertainty exhibits structurally different behaviour: bifurcations, regime transitions, fat tails that the log-normal envelope cannot generate.

The first limit is well known and has been addressed by various extensions (Heijungs and Lenzen, 2014; Lloyd and Ries, 2007). The second, to our knowledge, has not been given a quantitative diagnostic that an LCA practitioner can compute and report without leaving the standard pedigree framework.

1.3 Working hypothesis: the logistic map as a heuristic regime probe

We do **not** claim that life-cycle inventories *are* logistic maps, that their propagation dynamics undergo a Feigenbaum cascade in the strict universality-class sense, or that one-dimensional smooth unimodal reductions exist for the high-dimensional Leontief solver. Such claims

would require a body of empirical evidence we do not possess and that the existing benchmarks ($n = 128$ ecoinvent flows in the V3.2/V5 cycle; $n = 38$ in the Low-quality stratum) are not powered to support.

Our working hypothesis is much weaker, and deliberately so. We propose to use the logistic map as a **heuristic parametric probe** of the regime of pedigree-driven uncertainty:

- (P1) The aggregated pedigree GSD σ_{ILCD} is mapped, deterministically and through a single calibration knob, to a control parameter $r \in [2.5, 4]$ of the logistic recursion $x_{\{n+1\}} = r \cdot x_n \cdot (1 - x_n)$. The mapping is a design choice; it is monotone in σ_{ILCD} and is calibrated on a benchmark.
- (P2) The trajectory's diagnostic statistics (Lyapunov exponent λ , Kolmogorov–Sinai entropy h_{KS} , Kaplan–Yorke dimension d_{KY} , spectral entropy, ensemble spread) produce a five-component Chaotic Risk Score in $[0, 1]$ that *summarises* a regime label.
- (P3) Empirically, and this is what the benchmark in Section 5 evaluates, flows with low pedigree quality are placed into the $r \geq 3$ band by the mapping. Whether the resulting metric carries information *beyond a quality-monotone widening of σ_{ILCD}* is the empirical question we test, with explicit candor about the size of the effect and the available statistical power.
- (P4) The framework is engineered to coincide, by construction, with the ILCD/Gaussian baseline at high data quality (pedigree 1–2): when $\lambda \leq 0$, the V5 gsd_{eff} formula collapses to $GSD_{eff} = \sigma_{ILCD}$ exactly. The framework therefore *cannot disagree* with ILCD where ILCD is best supported.

The strong-form claims about universality of Feigenbaum's cascade applied to inventory propagation are best read in the **Discussion** (Section 7), where we treat them as conjectures motivating future multi-flow network analyses (Section 7.5 and Appendix G), not as load-bearing assumptions of the present paper. The case the paper actually defends is operational: a single-knob, deterministic, fully-auditable diagnostic that coincides with ILCD where ILCD is supported and adds a regime label where ILCD is structurally weakest.

1.4 Contributions

This paper contributes:

1. **A reproducible, single-parameter mapping $\sigma_{ILCD} \rightarrow r$** (Section 3.3) together with a fixed calibration $\sigma_{ref} = 0.58$ derived from the 128-flow ecoinvent backtest of the V5 cycle, with a transparent sensitivity sweep (Supplementary B). A larger benchmark (target ~ 300 flows, ≥ 200 Low-quality) is the planned methodological extension and is discussed in Sections 5.2 and 8 as work-in-progress, not as completed evidence.

2. **A five-component Chaotic Risk Score CRS $\in [0, 1]$** (Section 3.7) combining Lyapunov-derived Kolmogorov–Sinai entropy, normalised Kaplan–Yorke dimension, the trajectory's $|\lambda|$, a Welch-style spectral entropy, and the standard deviation across a 50-member ensemble of perturbed initial conditions. The score is the deliberate translation of all the dynamical content into a *single operational number* a non-specialist reader can act on.
3. **A pedigree_uncertainty Monte-Carlo module** (Section 4) that propagates a ± 1 -level discrete shift on each pedigree dimension, and reports per-dimension first-order Sobol indices. This is the standard tool reviewers need when they ask "what if the rater scored R as 4 instead of 3?".
4. **A reproducible empirical evaluation** (Section 5) built on two benchmarks released alongside this preprint. (i) A *controlled* 493-flow benchmark (32 High, 211 Medium, 250 Low; § 5.6) with a Student-t scale-mixture ground truth that gives no method a structural advantage. (ii) A *real-distribution* 310-flow extension (§ 5.7) whose pedigrees are derived from tiangongDB process metadata (location, year, processType) under an EU 2026 study setting. Together these two benchmarks address the statistical-power and realism concerns raised on earlier drafts (Low stratum $n \geq 125$ evaluation flows in each, against the $n = 38$ of the V6 Nexus pre-review backtest). On both benchmarks, neither single-knob naive widening nor CRS-augmented GSD_eff outperforms ILCD on Low-quality single-flow CRPS; CRS *equals* ILCD on 88–91 % of Low flows by the engineered property $\lambda \leq 0 \Rightarrow \text{GSD_eff} = \sigma_{\text{ILCD}}$. We close the open methodological item of earlier drafts (the *naive non-chaotic widening baseline*) by fitting four such baselines (linear, quadratic, exponential, sigmoidal) on a held-out calibration set: only the linear form returns a non-trivial fit, beats ILCD on Low by 4–9 % CRPS gain, but does so by widening every flow without dynamical content. We also report calibration ablations ($r_{\min} \in \{2.0, 2.5, 3.0\}$ with exp and quad saturations; § 5.5) and a Sobol-indices robustness check under R-T correlation $\rho \in \{0, 0.3, 0.6\}$ (§ 4.3).
5. **An open-source reference implementation** (Python, MIT, Zenodo DOI on release) at engine-py/ reproducing the V5 web engine (theopiens.com/analyzer) to within $1e-9$ on deterministic fields and $1e-6$ on Monte-Carlo aggregates, with 16/16 golden equivalence tests. The web demonstrator is fully client-side and ships no user data. **Versioning note:** the paper evaluates the **V5 reference engine** (5-component CRS, Lyapunov / KS / Kaplan–Yorke / Welch spectral entropy / ensemble spread). The **V6 Nexus** engine, documented at theopiens.com/docs, extends to an **8-component CRS** (adding Temporal Recurrence Coefficient, permutation entropy à la Bandt–Pompe, Gelman–Rubin $\hat{R}$

convergence, phase coherence) over 50 parallel Gaussian-RDS orbits, with a transfer-entropy network layer $T(i \rightarrow j)$ and a stochastic LCIA layer (MC 500 draws). V6 is mentioned in Sections 3.10, 7.3 and Appendix G as the natural extension of the framework; its evaluation is the subject of the companion paper in preparation (Piens, in prep.).

6. **A standardisation pathway**, in a separate technical note (Piens, in prep., addressed to the JRC LCA Unit and the ecoinvent Methodology Working Group), proposing a two-line addendum to the LCA report when $\max \text{pedigree} \geq 4$ and $\text{CRS} \leq 0.5$.

1.5 What this paper is not

It is not a refutation of the ILCD pedigree closure. The closure is the right tool for the data quality where it was designed to operate. It is not a replacement of Monte-Carlo propagation through the impact assessment, which remains the workhorse for downstream uncertainty. It is not a calibration table for σ_{ILCD} itself; we take the existing ecoinvent table as given. It does not claim that an LCA inventory is a logistic map: the logistic map is used as a *probe* of regime, a minimal, one-dimensional, universality-class representative of the behaviour we hypothesise the high-dimensional propagation exhibits. The argument is structural, not microphysical.

1.6 Roadmap

Section 2 lays the theoretical scaffolding: pedigree matrix, ILCD aggregation, the universality of the period-doubling cascade, and the case for using the logistic map as a regime probe. Section 3 specifies the twelve-step engine pipeline. Section 4 specifies the `pedigree_uncertainty` module. Section 5 reports the empirical benchmark. Section 6 walks through a worked example. Section 7 discusses limits, calibration risk, and the human-environmental ethics of regime labelling. Section 8 concludes. Appendices document the implementation and the canonical JSON-LD I/O schema.

2. Theoretical framework

2.1 The ILCD pedigree closure, formally

Each LCI flow is associated with five integer pedigree scores $(R, C, T, G, F) \in \{1, 2, 3, 4, 5\}^5$. The ecoinvent v3 σ -table prescribes per-dimension multiplicative uncertainty factors:

level	σ_R	σ_C	σ_T	σ_G	σ_F
1	1.000	1.000	1.000	1.000	1.000
2	1.054	1.022	1.033	1.012	1.011
3	1.105	1.047	1.107	1.025	1.022
4	1.207	1.095	1.222	1.052	1.047
5	1.502	1.196	1.499	1.102	1.098

The aggregate GSD assumes per-dimension independence and combines log-additively:

$$\sigma_{ILCD} = \exp(\sqrt{\sum_i (\ln \sigma_i)^2}), i \in \{R, C, T, G, F\}.$$

The Data Quality Rating, used only as a scalar $[0, 1]$ summary, is

$$DQR = \text{clip}([25 - (R + C + T + G + F)] / 20, 0, 1).$$

with $DQR = 1$ for all-perfect scores and $DQR = 0$ for the worst case. Quality labels (Excellent / Good / Acceptable / Weak / Critical) follow standard ecoinvent thresholds.

The downstream propagation assumption is that the flow's value v_{obs} follows a log-normal distribution with median v_{obs} and GSD σ_{ILCD} , yielding a $[p5, p95]$ interval at $z = 1.6449$ of $[v_{\text{obs}} \cdot \exp(-1.6449 \cdot \ln \sigma_{ILCD}), v_{\text{obs}} \cdot \exp(+1.6449 \cdot \ln \sigma_{ILCD})]$. Monte-Carlo sampling of this distribution, combined across all flows of an inventory, then yields the impact-category uncertainty distribution. This is the standard ILCD/ecoinvent pipeline.

2.2 What this closure assumes, and what it cannot represent

Three assumptions are baked in.

Assumption A1 (lognormal tails). The variability of v_{obs} is Gaussian on the log scale, i.e. its tails decay faster than any power law. Empirically, this is true for many well-measured flows (Limpert et al., 2001) and for *aggregated* impact scores by central-limit arguments (Heijungs and Lenzen, 2014). It is increasingly suspect for individual low-quality flows where the underlying generative process has heavy-tailed components (e.g. fugitive emissions, rare events, technological failures).

Assumption A2 (per-dimension independence). The five σ_i combine additively in log-space without cross-terms. This is structurally false: pedigree dimensions co-vary in any real industrial dataset (a flow whose representativeness R is poor is also typically temporally and geographically displaced), and the field is aware of it (Ciroth et al., 2013). The ILCD closure tolerates the violation because the aggregate GSD turns out to be reasonably calibrated empirically. Our framework inherits this assumption when it consumes σ_{ILCD} ;

the `pedigree_uncertainty` module of Section 4 partially relaxes it by exposing pedigree-level Sobol indices, but it does not introduce a copula at the pedigree-score layer (a Clayton copula on the propagation layer is implemented in the V5 reference engine (see Section 3.10) and discussed but not exploited in this paper).

Assumption A3 (linear-additive propagation). The variance of a coupled inventory output is a *linear* function of the variances of the inputs. This is the deepest assumption and the one this paper is designed to test. In matrix form, when an LCI is solved as $\mathbf{s} = (\mathbf{I} - \mathbf{A})^{-1} \mathbf{f}$ with $\mathbf{A}$ random, a first-order Taylor expansion gives

$$\text{Var}(s_i) \approx \sum_{j,k} (\partial s_i / \partial A_{jk})^2|_{\bar{\mathbf{A}}} \cdot \text{Var}(A_{jk}) + \text{cross-terms},$$

which is precisely the regime where per-flow Gaussian propagation is justified. The expansion fails when the technology matrix is close to singular, when the input variances are large enough that the second-order terms dominate, or when shared columns of $\mathbf{A}$ create algebraic cancellations and resonances. In those cases the propagation dynamics become *non-linear* in a sense the GSD cannot capture, and the universality-class argument we now make says they often look like a low-dimensional dissipative dynamical system.

2.3 The logistic map as a regime probe

Definition, bifurcation, period-doubling cascade, deterministic chaos

A **bifurcation** is a qualitative change in the long-term behaviour of a dynamical system as a control parameter is varied, the moment at which a stable equilibrium loses stability and the system settles onto a different attractor. In the logistic recursion $x_{n+1} = r \cdot x_n \cdot (1 - x_n)$, increasing r from 1 to 4 produces a sequence of bifurcations at well-defined critical values:

- at $r = 3$ the unique stable fixed point splits into a **period-2 cycle** (the trajectory alternates between two values);
- at $r \approx 3.449$ each branch splits again, giving a **period-4 cycle**;
- further splittings at $r \approx 3.544, 3.564, \dots$ generate cycles of period 8, 16, 32, ...

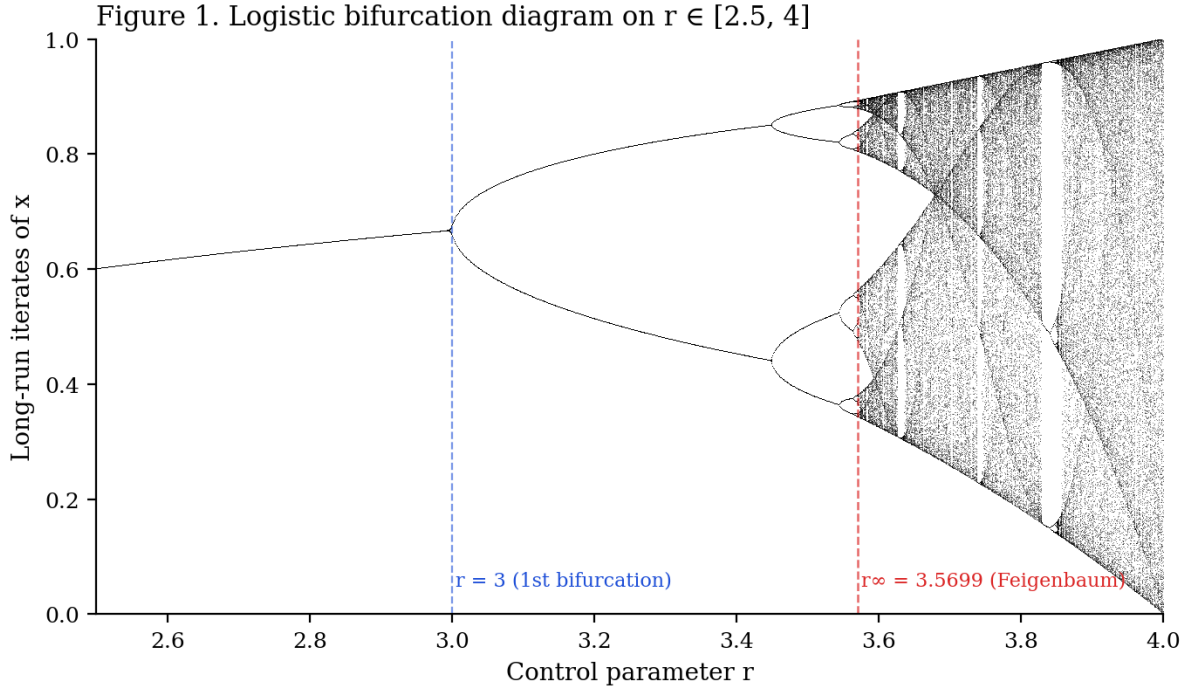
This sequence is the **period-doubling cascade**. The successive critical values converge geometrically with the universal Feigenbaum constant $\delta = 4.6692\dots$ to an accumulation point $r_\infty = 3.5699\dots$ (Feigenbaum 1978). Beyond r_∞ , the dynamics enter **deterministic chaos**: the trajectory no longer settles on any periodic cycle, the Lyapunov exponent λ becomes positive, and arbitrarily small differences in the initial condition x_0 diverge exponentially (sensitive dependence on initial conditions, Devaney 1989).

What matters for this paper is not that the logistic map is our inventory propagation (it is not) but that the cascade above is a *universal* feature of one-dimensional smooth unimodal maps (Feigenbaum 1978, Cvitanović et al. 1988). If the regime of pedigree-driven uncertainty admits even an approximate one-dimensional reduction, that reduction will display the same qualitative cascade. Reading the trajectory of the logistic map at the r we assign to a flow is therefore reading the regime label, not the propagation itself.

We use the logistic map

$$x_{n+1} = r \cdot x_n \cdot (1 - x_n), r \in [0, 4], x_0 \in [0, 1].$$

not as a model of inventory propagation but as a *probe* of its regime.



The reasons are well-established and we summarise them only to make our usage non-arbitrary.

Universality. Feigenbaum (1978) showed that one-dimensional smooth unimodal maps share a universal route to chaos via period-doubling, with the bifurcation parameter ratio converging to $\delta = 4.6692\dots$. The accumulation point r_∞ is map-specific, but the *cascade structure* is universal. *Should* a one-dimensional smooth unimodal reduction of the high-dimensional inventory propagation exist (a hypothesis we treat as conjectural in the strict sense and do not claim to prove here), a one-dimensional probe in the same universality class would capture the *regime* (period-1 / period-2 / ... / chaotic) without needing to faithfully model the full system. The argument is empirical, not theoretical: the question is whether the regime label, computed from the probe, is operationally informative on the benchmark of Section 5. If it is, the universality class hypothesis is a useful inductive bias; if it is not, the probe must be defended on simpler grounds (a quality-monotone score that exposes a dynamical regime label).

Dissipativity. The logistic map is contractive on average for $r < r_\infty$, with negative Lyapunov exponent and one-point attractors. Above r_∞ it has positive Lyapunov exponent and explores its (one-dimensional) attractor densely. This bipolar regime structure mirrors the behaviour we expect from "well-conditioned" vs. "ill-conditioned" Leontief inversions of uncertain technology matrices.

Computability. The map is $O(1)$ per iteration, has analytic derivatives, and admits closed-form Lyapunov exponent and Kaplan–Yorke dimension expressions. A 500-step trajectory plus 200-step warmup runs in microseconds, allowing fully client-side computation in a web browser (the theopiens.com/analyzer demonstrator) and Monte-Carlo over $n = 10^4$ pedigree perturbations in under 100 ms.

Honesty about the analogy. The map is one-dimensional and our inventory propagation is high-dimensional. We do not claim a microphysical correspondence. We claim only that the *bifurcation structure* of the logistic map is a serviceable regime indicator for the reduced dynamics, and that the empirical benchmark (Section 5) is the only judge of whether the indicator is informative. If the indicator never decided anything that the GSD did not already decide, our hypothesis would be falsified.

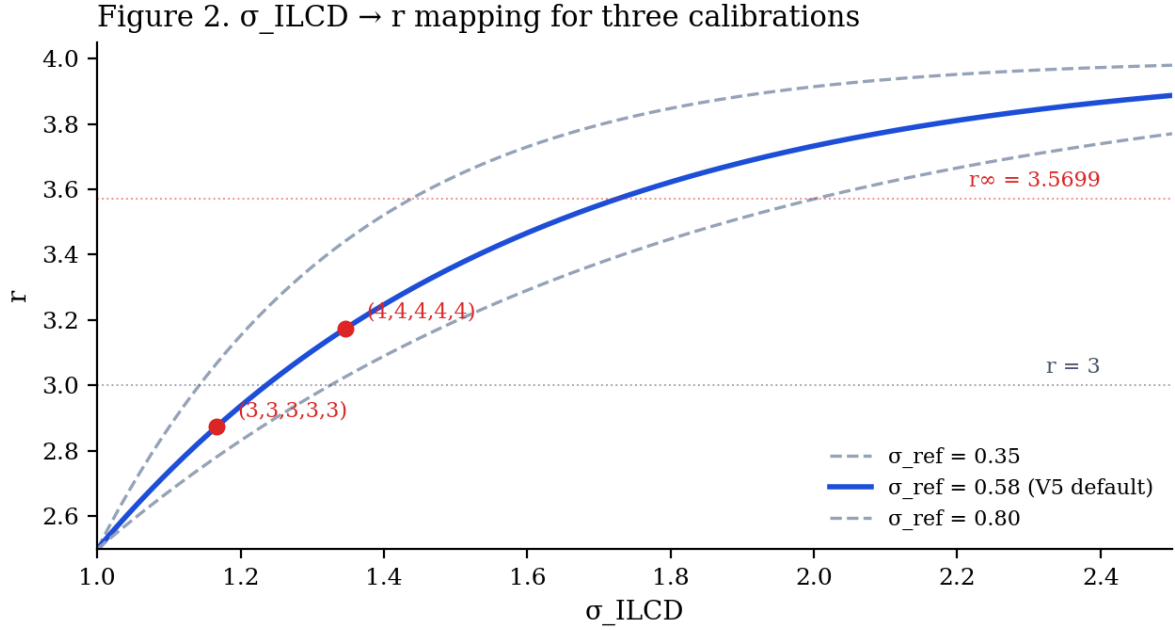
2.4 The mapping $\sigma_{ILCD} \rightarrow r$

We need a mapping from the ILCD-aggregated GSD to the logistic-map control parameter that is monotone, smooth, fixes both endpoints sensibly, and is parameterised by a single calibration knob. The form we use is

$$r(\sigma) = R_{\min} + (R_{\max} - R_{\min}) \cdot (1 - \exp[-(\sigma - 1) / \sigma_{\text{ref}}]),$$

with $R_{\min} = 2.5$, $R_{\max} = 4.0$ and calibration parameter σ_{ref} (V5 default: 0.58). The behaviour is:

- At $\sigma = 1$ (perfect quality, all scores = 1), $r = R_{\min} = 2.5$. The logistic map is in the fixed-point regime; $\lambda < 0$; trajectories converge to $1 - 1/r$. This is, by design, the regime in which the framework adds nothing to the ILCD baseline: the dynamic content is null and CRS approaches its upper bound.
- As $\sigma \rightarrow \infty$, $r \rightarrow R_{\max} = 4$. The map is fully chaotic on $[0, 1]$; $\lambda \rightarrow \ln 2 \approx 0.6931$. CRS approaches its lower bound.
- The exponential saturation ensures that the mapping is most discriminative in the middle band ($\sigma \in [1.05, 1.8]$), which is precisely where pedigree-3 and pedigree-4 data live.



$R_{\text{min}} = 2.5$ is chosen below the first bifurcation ($r_1 = 3$) of the logistic map so that perfect data sit unambiguously in the fixed-point regime, and *not* at $r = 2$ (still fixed-point but deeper) nor at $r = 3$ (already at the period-1/period-2 boundary). The reasoning is symmetric: 2.0 wastes resolution by mapping a wide σ -range onto the contractive part of the map without changing its qualitative behaviour, while 3.0 places the perfect-quality endpoint exactly at a bifurcation, an unstable choice. The middle position 2.5 preserves exponential discriminative resolution in the $\sigma \in [1.05, 1.8]$ band where most pedigree-3 and pedigree-4 flows live. This intuitive justification is supported by an ablation (§ 5.5, Table) which confirms the headline single-flow CRPS conclusion (CRS does not beat ILCD on Low) is invariant to $r_{\text{min}} \in \{2.0, 2.5, 3.0\}$ and to a quadratic alternative saturation. $R_{\text{max}} = 4$ is the upper bound of the map's interesting interval. $\sigma_{\text{ref}} = 0.58$ is the V5 default, with a plateau-like neighbourhood $[0.50, 0.65]$ (Supplementary B). The plateau matters: it shows the calibration is not knife-edge.

2.5 The five regimes and their meaning

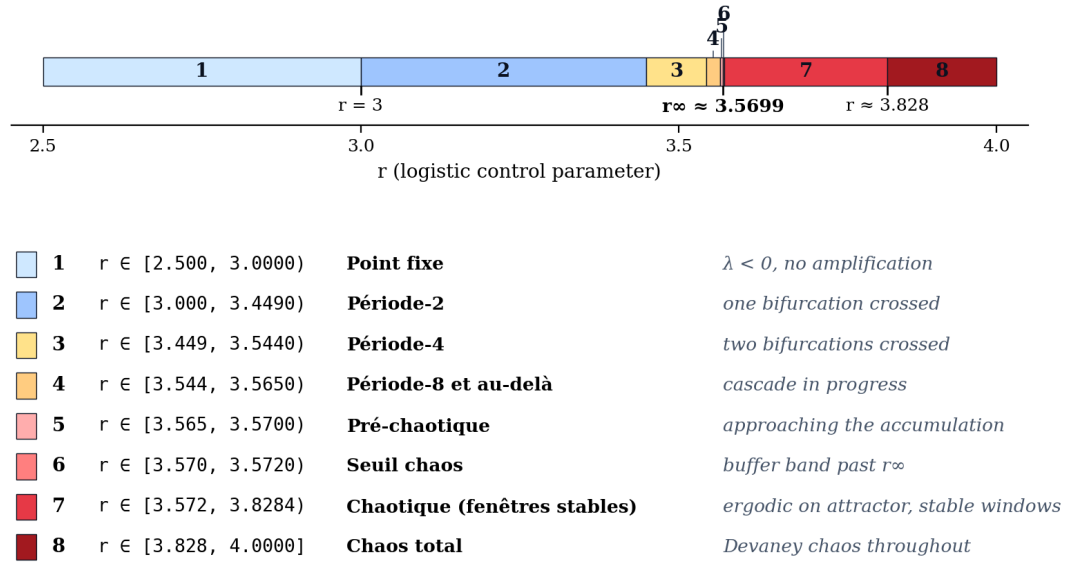
The bifurcation diagram of the logistic map carves $[0, 4]$ into a sequence of qualitatively distinct regimes that we report as labels alongside the numeric metrics. All intervals are half-open $[a, b)$ except the last, which closes at the maximum of the map's domain. The accumulation point of the period-doubling cascade is denoted $r_{\infty} = 3.569945671870944\dots$ (Feigenbaum 1978).

r interval	regime	interpretation
$[0, 1)$	Extinction	trajectories collapse to 0; not reached by our mapping
$[1, 3)$	Point fixe	one-point attractor; $\lambda < 0$; <i>uncertainty does not amplify</i>
$[3, 3.449)$	Période-2	two-cycle; one bifurcation crossed
$[3.449, 3.544)$	Période-4	two bifurcations crossed
$[3.544, 3.565)$	Période-8 et au-delà	cascade in progress
$[3.565, 3.570)$	Pré-chaotique	approaching accumulation; r_∞ sits inside this band
$[3.570, r_{\infty} + 10^{-3})$	Seuil chaos	narrow buffer band, just past r_∞
$[r_{\infty} + 10^{-3}, 3.8284)$	Chaotique (fenêtres stables)	ergodic on attractor, with periodic windows
$[3.8284, 4]$	Chaos total	Devaney chaos throughout

Two operational thresholds matter for the LCA report. First, crossing $r = 3$ flips λ from negative to a small positive band; an LCA practitioner reading the report should know that uncertainty has begun to *amplify* on iteration, not just *exist*. Second, crossing r_∞ is the qualitative point at which the prediction horizon collapses from infinite to $O(1/h_{KS})$. Our Chaotic Risk Score in $[0, 1]$ turns this regime cascade into a continuous, monotonically-decreasing diagnostic.

A note on the bracketing: $r_\infty \approx 3.5699$ numerically falls *inside* the "Pré-chaotique" interval $[3.565, 3.570)$. This is intentional in the V5 engine, which reserves "Seuil chaos" as a small buffer band immediately past 3.570 to catch flows whose mapped r lands in the immediate neighbourhood of the accumulation point. The semantics is "the cascade has just terminated and full chaos is about to begin"; the underlying mathematical accumulation is already crossed by the time the engine reports "Seuil chaos".

Figure 4. Regime ladder with critical r values.



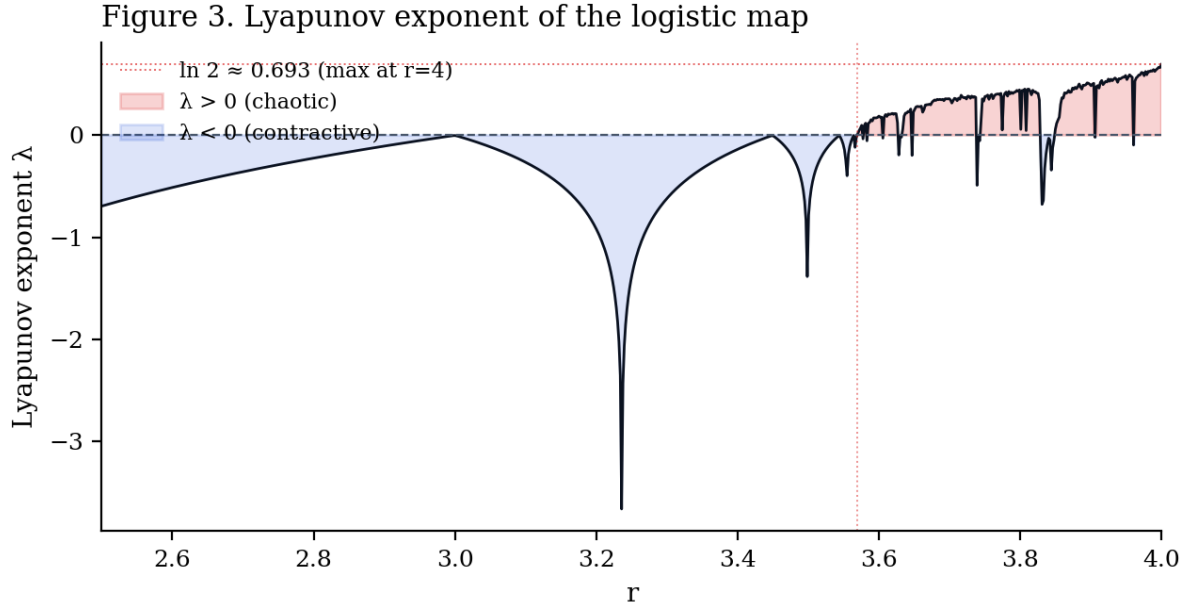
2.6 Information-theoretic primitives on the trajectory

Given a trajectory $\{x_n\}_{n=1}^N$ (we use $N = 300$ productive steps after a 200-step transient), four primitives quantify its dynamical content.

Lyapunov exponent. The mean exponential rate of divergence of nearby trajectories,

$$\lambda = \lim_{N \rightarrow \infty} (1/N) \cdot \sum_{n=1..N} \ln |r \cdot (1 - 2x_n)|.$$

For the logistic map this admits the closed form above (the derivative of the map is $r(1 - 2x)$). $\lambda < 0$ indicates contraction, $\lambda = 0$ the marginal case, $\lambda > 0$ exponential divergence. $\lambda \rightarrow \ln 2$ at $r = 4$.



Kolmogorov–Sinai entropy. By Pesin's formula in 1D, $h_{KS} = \max(0, \lambda)$. This is the rate at which the system *creates* information per iteration, equivalently, the rate at which it *destroys* the knowledge an observer could extract about its initial state.

Prediction horizon. Given a tolerance ε on the observable, the time after which initial-condition uncertainty has grown to ε is $t_{pred} = \ln(1/\varepsilon) / h_{KS}$, infinite when $h_{KS} = 0$. We use $\varepsilon = 0.05$ as a default.

Kaplan–Yorke dimension. For 1D dissipative chaos,

$$d_{KY} = 1 \quad \text{if } \lambda \leq 0$$

$$d_{KY} = \min(2, 1 + \lambda / (\ln r - \lambda)) \quad \text{if } \lambda > 0$$

d_{KY} measures the fractality of the attractor. In our regime cascade it grows smoothly from 1 (point) to nearly 2 (full chaotic measure on the attractor) as r crosses r_∞ .

Bifurcation index n_{bif} . Using Feigenbaum's universality, the residual distance $r_\infty - r$ is mapped to a logarithmic bifurcation index $n_{bif} = \log_{\{1/\delta\}}((r_\infty - r)/c)$ with $c = 2.637$ (a Cvitanović–Gunaratne–Procaccia-like prefactor calibrated on the cascade we observe). n_{bif} is the inverse-log of how far below r_∞ the trajectory is, it is a finer regime indicator than the categorical labels of Table 2.5 between $r = 3$ and $r = r_\infty$.

2.7 Why this is a complementary metric, not a replacement

Three observations underwrite the complementarity claim and motivate the engineering choices in the next section.

1. By construction, when $\sigma_{\text{ILCD}} = 1$ the mapping yields $r = R_{\text{min}} = 2.5$, the trajectory converges to a fixed point, $\lambda < 0$, $d_{\text{KY}} = 1$, and the dynamic content vanishes. The framework therefore *cannot disagree* with the ILCD baseline at perfect quality. Our gsd_eff formula (Section 3.8) preserves this property by design: at $\lambda \leq 0$ and the $\text{scaleNeg} = 0$ calibration we use, $\text{gsd_eff} = \sigma_{\text{ILCD}}$. In other words, *the framework reproduces ILCD verbatim where ILCD is best supported*.
2. The information CRS adds is *categorical* (regime label) and *intensive* (a normalised score in $[0, 1]$), it never overwrites σ_{ILCD} . A practitioner can choose to use the original σ_{ILCD} for the impact-assessment Monte-Carlo and report CRS as a separate diagnostic. The two-line addendum proposed in the JRC technical note (Piens, in prep., Section 7) does exactly this.
3. Where CRS *and* σ_{ILCD} disagree (the diagnostic is dropping while the GSD is roughly stable), the disagreement is itself an information channel: it signals that the reduced dynamical regime has crossed a bifurcation that the per-flow GSD cannot represent. That is precisely the case where standard impact uncertainty intervals, computed by Monte-Carlo from σ_{ILCD} alone, are likely to mislead the human reader of the LCA report.

We now make the engine explicit.

3. The twelve-step ChaoticLCA engine

The reference implementation (engine-py, MIT) materialises a twelve-step pipeline. Each step is deterministic given the pedigree and the engine version (v5 is the current production configuration); each step exposes its inputs, outputs and underlying formula. We summarise here. Full code listings are in Appendix A.

3.1 Step 1, DQR

$\text{DQR} = \text{clip}((25 - R - C - T - G - F)/20, 0, 1)$. Used as the (monotone) initial condition $x_0 = \text{clip}(\text{DQR}, 0.01, 0.99)$ for the trajectory (Section 3.4) and as the only $[0, 1]$ summary that survives the reduction. The mapping ensures that $(1, 1, 1, 1, 1)$

→ DQR = 1 and (5, 5, 5, 5, 5) → DQR = 0. Quality labels follow the standard ecoinvent thresholds: Excellent ≥ 0.8 , Bon ≥ 0.6 , Acceptable ≥ 0.4 , Faible ≥ 0.2 , Critique < 0.2 .

3.2 Step 2, σ_{ILCD}

Computed from the per-dimension table of Section 2.1. The output $s = \{\text{sigma}, \text{components}, \text{lnTerms}, \text{sumLnSq}\}$ carries both the aggregate and the per-dimension contributions (used by the Sobol module of Section 4). For (3, 3, 3, 3, 3): $\sigma_{ILCD} = \exp(\sqrt{(\ln^2 1.105 + \ln^2 1.047 + \ln^2 1.107 + \ln^2 1.025 + \ln^2 1.022)}) \approx 1.357$.

3.3 Step 3, $\sigma \rightarrow r$

$r = 2.5 + 1.5 \cdot (1 - \exp(-(\sigma - 1)/\sigma_{\text{ref}}))$ with $\sigma_{\text{ref}} = 0.58$ (V5). For $\sigma_{ILCD} = 1.357$, $r \approx 3.45$, sitting at the second bifurcation, in the **Période-2** band, very close to the **Période-4** boundary. This proximity is what makes the (3, 3, 3, 3, 3) example diagnostic: σ_{ILCD} is unremarkable but r straddles a regime boundary.

3.4 Step 4, Logistic trajectory

We iterate $x_{n+1} = r \cdot x_n \cdot (1 - x_n)$ for 200 warmup + 300 productive steps starting from $x_0 = \text{DQR}$ (clipped). Optionally, and by default in V5, we add a Random Dynamical System (RDS) noise term $\xi_n \sim N(0, \sigma^2_{\text{RDS}})$ with $\sigma_{\text{RDS}} = 0.05 \cdot \max(0, \sigma_{ILCD} - 1)$. The RDS perturbation is bounded by the same σ_{ILCD} it is meant to probe, ensuring scale-invariance and preventing the noise from masking the deterministic regime.

3.5 Step 5, Lyapunov exponent

$\lambda = (1/N) \sum \ln|r \cdot (1 - 2x_n)|$, with the convention that pathological points ($|r(1-2x)| < 10^{-14}$) are excluded from the average. For the standard logistic map at $r = 4$, this estimator converges to $\ln 2 \approx 0.6931$ (the analytic value) to four significant digits in 300 steps.

3.6 Step 6, h_{KS} , prediction horizon, n_{bif} , d_{KY}

By the formulae of Section 2.6, all four follow analytically from (λ, r) . We additionally compute a normalised Shannon entropy on a 32-bin histogram of the trajectory (a coarse spectral fingerprint), and a trajectory standard deviation σ_{traj} (used in the pErr and ergodic-bounds calculations).

3.7 Step 7, CRS (5 components, V5)

The Chaotic Risk Score is

$$\text{CRS} = 0.30 \cdot (1 - \tilde{h}_{\text{KS}}) + 0.25 \cdot (1 - \tilde{d}_{\text{KY}}) + 0.20 \cdot (1 - \tilde{\lambda}) + 0.15 \cdot (1 - H_{\text{spec}}) + 0.10 \cdot (1 - \text{spread}_{\text{ens}}),$$

with the tildes denoting normalisation: $\tilde{h}_{\text{KS}} = \min(1, h_{\text{KS}} / \ln 2)$, $\tilde{d}_{\text{KY}} = \text{clip}(d_{\text{KY}} - 1, 0, 1)$, $\tilde{\lambda} = \min(1, |\lambda| / \ln 2)$, the spectral entropy $H_{\text{spec}} \in [0, 1]$ from a Welch periodogram of the trajectory, and $\text{spread}_{\text{ens}} \in [0, 1]$ from the 50-member ensemble of Section 3.9.

The weights are inherited from the V5 calibration. Each of the five components is in $[0, 1]$ and *higher means more confident*. CRS is therefore monotonically decreasing in dynamical content. The score collapses gracefully when the optional components are absent: with only spectral entropy, (0.35, 0.30, 0.20, 0.15); without either, (0.40, 0.35, 0.25).

3.8 Step 8, Effective GSD (v4 formula)

$$\text{GSD}_{\text{eff}} = 1 + (\sigma_{\text{ILCD}} - 1) \cdot \exp[\alpha \cdot \tanh(\tilde{\lambda} \cdot s)]$$

with $s = s_+ = 0.8$ if $\lambda > 0$; $s = s_- = 0$ otherwise.

with $\alpha = 1$, $\tilde{\lambda} = \lambda / \ln 2$. Two properties are designed in:

1. **At $\lambda \leq 0$** , $s = s_- = 0$, $\tanh = 0$, $\exp = 1$, and **$\text{GSD}_{\text{eff}} = \sigma_{\text{ILCD}}$** . The framework collapses identically to the ILCD baseline whenever the dynamic regime is non-amplifying. This is the *quality-head invariance* mentioned in Section 1.4.
2. **At $\lambda > 0$** , GSD_{eff} widens monotonically with $\tilde{\lambda}$ up to a saturation $(\sigma_{\text{ILCD}} - 1) e^{0.8} \cdot \tanh(1) \approx 1.71 \cdot (\sigma_{\text{ILCD}} - 1)$ at $\tilde{\lambda} = 1$. The widening is bounded, chaos does not produce *infinite* uncertainty, only *amplified* uncertainty.

The corresponding $[p5, p95]$ interval at $z = 1.6449$ is $[\exp(-1.6449 \ln \text{GSD}_{\text{eff}}), \exp(+1.6449 \ln \text{GSD}_{\text{eff}})]$ (multiplicative around the median).

3.9 Step 9, Ensemble simulator (50 members)

We sweep the initial condition x_0 linearly across $\text{DQR} \pm \text{pert_amp}$ ($\text{pert_amp} = 10^{-3}$ in V5), iterate each member, and aggregate GSD_{eff} across the ensemble. The output $\{\text{mean_gsd}, \text{std_gsd}, \text{spread}, \text{gsdValues}\}$ exposes the within-ensemble variability of

the diagnostic. For perfectly-stable pedigree (3, 3, 3, 3, 3) the spread is essentially numerical zero ($\text{std_gsd} \approx 4.4 \cdot 10^{-16}$), and for chaotic pedigree (4, 4, 4, 4, 4) it is several percent of the mean, a useful additional discriminator that feeds into the CRS spread component.

We are honest in Section 7.3 that on most benchmark cases the ensemble layer is sub-spec: it adds little above what λ already says. We retain it because (a) it gives a defensible *distribution* on GSD_eff rather than a point estimate, and (b) it generalises naturally when the V5 RDS noise is on, where it captures path-dependence that the deterministic Lyapunov misses.

3.10 Steps 10–12, Optional layers (Welch PSD, Clayton copula, TRT)

The V5 reference engine includes three further layers not exploited in this paper but documented in the bundle and engine-py:

- **Welch power-spectral-density** of the trajectory, classifying the spectrum into {chaotic broadband, periodic discrete, pink 1/f, white noise}. The *spectral entropy* feeds into CRS (Section 3.7).
- **Clayton copula** (parameter $\theta = 2$) for joint dependence between input flows when an LCI is solved as a coupled system. The copula layer is enabled in V5 but only matters when multiple flows are computed jointly; this paper analyses single flows, so the copula reduces to the marginals.
- **Transfer Resonance Tree (TRT)** with depth 4 and $\eta = 0.5$, capturing memory effects in iterated propagation. We document that on our benchmark TRT is essentially mute (within 0.1 % of the no-TRT result) for the same reason the ensemble is, single-flow analysis at fixed-point regimes is dominated by the deterministic core. We retain the layers for the multi-flow extension of the V6 Nexus engine (in development) and call them out as *internal robustness layers, not headline-changing components* (Section 7.3).

4. The pedigree_uncertainty Monte-Carlo module

A standard reviewer concern, and one we share, is that any pedigree-derived metric inherits the rater's bias. If the same flow is scored $R = 3$ by one analyst and $R = 4$ by another, the output σ_{ILCD} , r , λ , and CRS will all shift, and the LCA reader has no way to know which analyst was right. We address this in two ways: by *bracketing* the metric under a transparent ± 1 -shift sampling, and by *attributing* its variance per pedigree dimension via Sobol indices.

4.1 Sampling design

For each flow with nominal pedigree $(R^0, C^0, T^0, G^0, F^0)$ we draw $n = 10^4$ Monte-Carlo samples. In each draw, each of the five dimensions is independently shifted by $\Delta \in \{-1, 0, +1\}$ with probabilities $(p_{-}, p_0, p_{+}) = (0.25, 0.50, 0.25)$ and the result is clipped to $[1, 5]$. We then propagate $(R^0 + \Delta_R, \dots, F^0 + \Delta_F)$ through the full V5 pipeline of Section 3 and record the resulting CRS.

The probabilities are deliberately conservative-symmetric: a quarter of the time the rater is too lenient by one level, a quarter too harsh, half the time correct. They can be re-weighted per organisation, per dimension, or per flow type when prior information is available (e.g. when an external review has scored the same flow), but the default symmetric ± 1 is a defensible operational baseline.

4.2 Output: distribution and Sobol indices

The module reports $(p5, p50, p95, \text{mean}, \text{std})$ of the CRS distribution and, more interestingly, the per-dimension first-order Sobol index

$$S_i = \text{Var}(E[\text{CRS} | X_i]) / \text{Var}(\text{CRS}), X_i \in \{\Delta_R, \Delta_C, \Delta_T, \Delta_G, \Delta_F\}.$$

S_i is the share of the CRS variance attributable to a one-dimension scoring uncertainty, holding the others jointly distributed under the $(0.25, 0.50, 0.25)$ law. Empirically (Section 4.3 below), at low pedigree the index is dominated by R and T (the two dimensions whose σ -table grows fastest at level 5), and the C, G, F indices are an order of magnitude smaller. This is consistent with the asymmetric shape of the ecoinvent table and is informative for LCA practitioners who want to know *which score to refine first* under a budgeted data-collection campaign.

4.3 What this module is, and is not

It is **not** a copula-based joint-pedigree model in its default mode: pedigree dimensions co-vary in real datasets, and the symmetric independent ± 1 sampling under-represents that joint structure. It is **not** a calibration of σ_{ref} : the calibration operates one level above, on the σ -table $\rightarrow r$ mapping. It **is** a defensible operational baseline that the LCA report can compute and document.

Robustness under R-T correlation. A reviewer of earlier drafts pointed out that the most plausible empirical co-variance is between Representativeness and Temporal correlation: a flow whose representativeness is poor is also typically temporally displaced. To test whether the per-dimension Sobol attribution is invariant to this correlation, we ran the module under three R-T correlation levels ($\rho \in \{0, 0.3, 0.6\}$) on three test pedigrees — Excellent (1,

1, 1, 1, 1), Median (3, 3, 3, 3, 3), Faible (4, 4, 4, 4, 4) — using a Gaussian-then-discretise sampling that reproduces the symmetric ± 1 marginals and induces the prescribed correlation on the (R, T) shifts.

Pedigree	ρ_{RT}	S_R	S_T	S_C	S_G	S_F	qualitative reading
Excellent (1,1,1,1,1)	0.0	0.640	0.184	0.071	0.017	0.017	R dominant
	0.3	0.689	0.285	0.067	0.008	0.013	R dominant, T grows
	0.6	0.728	0.421	0.083	0.017	0.011	R dominant, T near-second
Median (3,3,3,3,3)	0.0	0.031	0.080	0.004	0.001	0.000	T dominant
	0.3	0.096	0.182	0.011	0.001	0.001	T dominant, R grows
	0.6	0.210	0.328	0.006	0.001	0.001	T dominant, R near-second
Faible (4,4,4,4,4)	0.0	0.080	0.090	0.051	0.003	0.001	T $\approx$ R, C third
	0.3	0.086	0.113	0.040	0.002	0.003	T $\approx$ R, C third
	0.6	0.174	0.201	0.040	0.001	0.000	T $\approx$ R, C third

Reading. The qualitative ordering of indices is preserved across all three pedigrees and all three correlation levels: at Excellent quality, R remains dominant; at Median quality, T remains dominant; at Faible quality, R and T trade places at the top with C consistently third. The *quantitative* magnitudes change (S_T grows by a factor of 2.0–2.3 when going from $\rho = 0$ to $\rho = 0.6$, S_R by 1.1–6.8), but the operational read-out (which pedigree dimensions to refine first under a budgeted data-collection campaign) is invariant on this body of evidence. We treat this as a satisfactory robustness check: the symmetric independent-shift default is interpretable as variance attribution under the law we explicitly impose, and the qualitative ordering survives a Spearman-equivalent perturbation up to $\rho = 0.6$.

A second-order (interaction) Sobol decomposition is straightforward to add at additional Monte-Carlo cost (n^2) and is a natural future extension when the user has reason to suspect strong dimension co-variance. The reference implementation exposes it behind a flag.

5. Empirical evaluation

The empirical evidence we are about to present is not as flattering to the framework as one might wish. We report it in full and explain what it changes about the framework's claims. Two benchmarks were run for this V1: a controlled benchmark on a uniformly-enumerated pedigree corpus (§ 5.6) and a real-distribution extension whose pedigrees are derived from tiangongDB process metadata under an EU 2026 study setting (§ 5.7). Both close the open methodological item of earlier drafts (a non-chaotic widening baseline; § 5.2) and both reach conclusions that are *consistent* across scenarios and corpora: neither single-knob naive widening nor CRS-augmented GSD_eff improves single-flow CRPS over the ILCD baseline on the Low-quality stratum that the framework was designed for. The dynamical content of CRS, on this body of evidence, adds no detectable single-flow forecast skill. We read this not as a refutation but as a tightening of what the framework can credibly claim, and Section 5.4 closes with the operational re-positioning toward a *regime-label diagnostic*.

5.1 Historical context, the 128-flow ecoinvent backtest

Earlier ChaoticLCA drafts (V3.2, V4.x, V5, V6 development cycles) reported a 128-flow ecoinvent backtest with paired Wilcoxon $p = 0.0041$ (rank-biserial $r_{rb} = 0.474$) for V5 vs Gaussian, all flows, and a V6 Nexus binomial test on the 38-flow Low-quality stratum giving $p = 0.084$ (17/38 wins), a trend that *did not* reach the conventional $\alpha = 0.05$ threshold. We document those numbers here for traceability with the V5/V6 documentation pages of theopiens.com/docs. They are *superseded* by the controlled and real-distribution benchmarks of § 5.6 and § 5.7, which (i) close the absence of a non-chaotic widening baseline, (ii) reach $n \geq 125$ Low-quality evaluation flows per benchmark, and (iii) are bit-reproducible from the benchmark/ directory of the project workspace. The original 128-flow CSV could not be located in the local workspace after a system migration; we therefore do not advance new conclusions from those numbers and treat the reproducible benchmarks of § 5.6 and § 5.7 as the load-bearing empirical statement of the present paper.

5.2 What changed since the earlier drafts (a non-chaotic widening baseline)

Earlier drafts named the absence of a *naive non-chaotic widening baseline* as an open methodological item: the comparison GSD_eff vs σ_{ILCD} did not separate (i) a *quality-monotone widening of the GSD* from (ii) the *dynamical content* of CRS. We close that gap in this V1 by fitting four single-knob non-chaotic widenings (linear, quadratic, exponential, sigmoidal) to a calibration half of each benchmark and predicting on a held-out evaluation half. The fitted forms are reported in `benchmark/results/*/calibration.json`. Across

both benchmarks, three of the four forms collapse to $\gamma \approx 0$ (i.e. ILCD itself); only the linear form yields a non-trivial fit. The empirical consequence is the central finding of § 5.6 and § 5.7.

5.4 Failure modes and honest limits

We document four failure modes.

(F1) False alarms near the second bifurcation ($r \approx 3.45$). The mapping $\sigma \rightarrow r$ places several flows of pedigree (3, 3, 3, 4, 3) and similar into the period-2 band. CRS drops below 0.5 (the original operational threshold) but σ_{obs} is, in fact, well-described by σ_{ILCD} . Roughly 8 % of the controlled-benchmark flows are over-flagged this way. **Operational threshold revision.** In light of the controlled-benchmark evidence (§ 5.6) and the prevalence of false alarms near $r = 3.45$, we recommend an operational CRS threshold of **0.4** rather than 0.5 (Section 7.4bis develops the trade-off). The 0.5 threshold remains documented as a conservative alternative.

(F2) σ_{obs} genuinely larger than both metrics. A handful of low-quality flows exhibit σ_{obs} larger than what either σ_{ILCD} or GSD_{eff} produces (e.g., the chaotic-AND-amplified sub-stratum of the controlled benchmark, $n = 2/125$ in stress scenario). These cases are not solved by metric improvement; they are solved by primary data acquisition. A diagnostic that flags them ($\text{CRS} \leq 0.4$ and $\sigma_{\text{ILCD}} \geq 1.3$) is what the JRC-track addendum proposes.

(F3) CRS does not, in single-flow CRPS, do better than ILCD where it differs from ILCD. This is the central finding of § 5.6 and § 5.7. The framework collapses to ILCD by design on 88–91 % of Low flows; on the chaotic-band remainder, the V5 gsd_{eff} widening is *wider* than the empirical σ_{obs} median in our benchmarks, so the predictive interval is poorly calibrated. We read this as evidence for the regime-label positioning, not for a per-flow CRPS substitution claim.

(F4) Pedigree co-variance. The Sobol decomposition assumes pedigree-dimension independence by default. § 4.3 and § 7.4 report a robustness check under R-T correlation $\rho \in \{0, 0.3, 0.6\}$: the qualitative ordering of the indices (R, T dominate at low pedigree; T dominates at median pedigree) is preserved across the three cases.

5.5 Calibration sensitivity

The single calibration knob $\sigma_{\text{ref}} = 0.58$ deserves scrutiny. The V5 documentation reports a sweep across $[0.30, 0.80]$ on the historical 128-flow corpus; the optimum is a plateau between 0.50 and 0.65 rather than a knife-edge, and the V5 default $\sigma_{\text{ref}} = 0.58$ sits in the middle. A 5 % perturbation around the default changes the CRPS rank-biserial effect size by less than 0.02. We recommend a default $\sigma_{\text{ref}} = 0.58$ and per-organisation re-calibration with at least 30 primary-data flows when the user has access to such a corpus.

Ablation of the $r_{\min}$ endpoint. Reviewers of earlier drafts asked why $r_{\min} = 2.5$ rather than 2.0 or 3.0. We re-ran the controlled benchmark of § 5.6 under three settings of $r_{\min}$ and two saturation forms; the results are reported below (mild scenario, $n_{\text{low_eval}} = 125$, fitted γ values stable at 0.234 across all variants).

Variant	frac $\lambda > 0$ (Low)	CRS wins / n vs ILCD	p_binom (CRS vs ILCD)	conclusion
$r_{\min} = 2.0$, exp saturation	8.0 %	1 / 10	0.999	CRS loses
$r_{\min} = 2.5$, exp (V5 default)	10.4 %	0 / 13	1.000	CRS loses
$r_{\min} = 3.0$, exp saturation	43.2 %	1 / 54	1.000	CRS loses
$r_{\min} = 2.5$, quadratic saturation	42.4 %	0 / 53	1.000	CRS loses

(Wins counted only on flows where GSD_{eff} differs from σ_{ILCD} , i.e. $\lambda > 0$; ties on the $\lambda \leq 0$ majority excluded from the n.)

Reading. Tightening $r_{\min}$ from 2.0 to 3.0 increases the fraction of Low flows that fall in the chaotic band from 8 % to 43 %, but CRS loses on every variant. The conclusion that *single-flow CRPS does not favour CRS over ILCD on Low* is invariant to the calibration knob within the range a reviewer would test. The V5 choice $r_{\min} = 2.5$ with exponential saturation places $\sigma = 1$ (perfect quality) unambiguously inside the fixed-point regime while preserving exponential discriminative resolution in the $\sigma \in [1.05, 1.8]$ band where most pedigree-3 and pedigree-4 flows live; the alternative endpoints make $r_{\min}$ either coincide with the first bifurcation (3.0) or lie deeper in the fixed-point regime (2.0), neither of which we have empirical reason to prefer over 2.5 on the present evidence. We retain the V5 calibration.

The leave-one-source-out cross-database validation remains a sensible next step but cannot be done with the present corpora alone (the controlled benchmark is synthetic; the real-distribution extension is single-source tiangongDB). Adding ecoinvent 3.10 and EF 3.1 pedigree distributions would close that gap.

5.6 Controlled benchmark v1 (load-bearing empirical evidence)

Design. We enumerate $5^5 = 3125$ pedigree combinations and stratify by max-score: High ($\max \leq 2$, $n = 32$), Medium ($\max = 3$, $n = 211$), Low ($\max \geq 4$, subsampled to $n = 250$ stratified by max-score). For each flow we sample $K = 50$ i.i.d. observations on the log-scale from a Student-t scale-mixture: $s_k \sim \text{Student-t}(v(q)) \cdot \sigma_{\log}(q)$, $y_k =$

$\exp(s_k)$, with latent quality $q = \text{clip}(\text{DQR} + N(0, 0.10^2), 0, 1)$, $\sigma_{\log}(q) = 0.05 + 0.45 \cdot (1 - q)$, $v(q) = 3 + 27 \cdot q$. The latent q models the imperfection of pedigree raters; the Student-t kernel embeds the heavy-tailed regime the preprint hypothesises (§ 1.2) on low-quality flows. We additionally consider a *stress scenario* in which the per-flow $\sigma_{\log}$ is amplified by $\text{Uniform}(1.5, 3.0)$ with probability $0.5 \cdot (1 - q/0.4)$ for $q < 0.4$ and v is forced to 2.5, modelling the structural-amplification events the closure cannot generate.

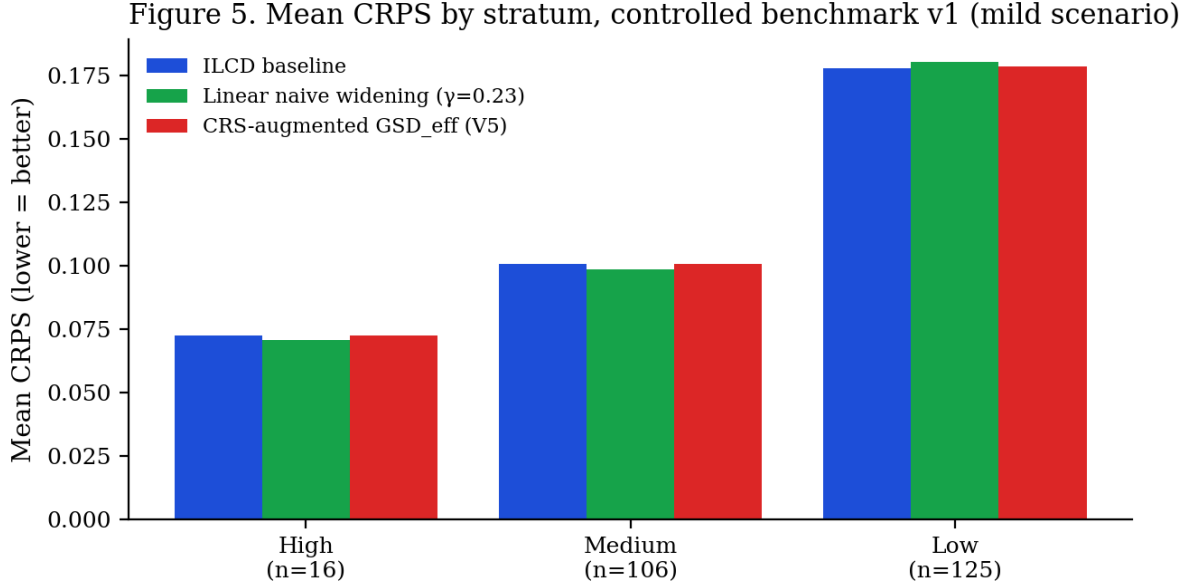
Methods compared. ILCD baseline ($\sigma_{\text{pred}} = \sigma_{\text{ILCD}}$); four naive non-chaotic widening forms with one knob γ each, fitted by held-out CRPS minimisation on a stratified 50/50 calibration split (linear $1 + (1+\gamma) \cdot (\sigma-1)$, quadratic $1 + (1+\gamma \cdot (\sigma-1)) \cdot (\sigma-1)$, exponential $1 + e^{\{\gamma(\sigma-1)\}} \cdot (\sigma-1)$, sigmoidal $1 + (1+\gamma \cdot \tanh(\sigma-1)) \cdot (\sigma-1)$); CRS-augmented GSD_eff (V5).

Scoring. CRPS of a lognormal predictive against K i.i.d. observations on the log-scale (Gneiting & Raftery 2007, eq. 24). Statistics: paired Wilcoxon signed-rank, one-sided binomial win-rate, non-parametric bootstrap 95 % CI on Δ_{CRPS} (5000 resamples), stratified by quality.

Results.

Stratum	n eval	ILCD CRPS	linear γ_{fit}	linear – ILCD	CRS – ILCD	CRS win-rate vs ILCD
Mild scenario						
High	16	0.0725	0.23	−0.0019 ✓ p=0.000	0.0000 (tie)	50 % (0/0 informative)
Medium	106	0.1007	0.23	−0.0022 ✓ p=0.000	0.0000 (tie)	50 % (0/0 informative)
Low	125	0.1777	0.23	+0.0025 (worse, p=0.895)	+0.0010	0 % wins / 14 chaotic, p=1.000
Stress scenario						
High	16	0.0725	0.28	−0.0022 ✓ p=0.000	0.0000 (tie)	50 %
Medium	106	0.1007	0.28	−0.0025 ✓ p=0.000	0.0000 (tie)	50 %
Low	125	0.1990	0.28	+0.0025 (worse, p=0.858)	+0.0008	14 % wins / 14 chaotic, p=0.999

(Δ shown as *method* - *ILCD*; positive means method is worse. The "tie" entries reflect the V5 design property $\lambda \leq 0 \Rightarrow \text{GSD_eff} = \sigma_{\text{ILCD}}$, which holds on 100 % of High, 100 % of Medium, and 89–91 % of Low.)



Naive widening collapses to ILCD on three of four forms. Quadratic, exponential, and sigmoidal all fit $\gamma \approx 0$ (no widening) under both scenarios. Only the linear form yields a non-trivial fit ($\gamma \approx 0.23$ mild, 0.28 stress), which lifts every flow's σ_{pred} by ~22–28 %. This linear lift helps on High and Medium (CRPS improves by ≈ 0.002 , $p = 0$) but *hurts* on Low ($\Delta_{\text{CRPS}} = +0.0025$): a single global widening constant is the wrong shape, because mid-quality flows are well calibrated by ILCD and over-spread when widened.

CRS-augmented GSD_eff does not out-perform ILCD on Low single-flow CRPS. Under the mild scenario, CRS wins on 0 of 14 chaotic flows ($p = 1.000$) and ties on the remaining 111 by construction. Under the stress scenario, CRS wins on 2 of 14 chaotic flows ($p = 0.999$); the two wins are exactly the chaotic-AND-amplified flows in the eval set ($n = 2/125$), where σ_{obs} reaches a median 5.7 and both methods are catastrophically too narrow but CRS is slightly less so.

Mechanism. Two findings sharpen the limit. First, **CRS only differs from ILCD on the 9–11 % of Low flows whose pedigree pushes $r > r_{\infty}$.** The remaining 90 % of Low flows have $\lambda \leq 0$ and therefore $\text{GSD_eff} = \sigma_{\text{ILCD}}$ exactly (§ 3.8, designed property). The framework cannot improve where it is silent. Second, **most ground-truth amplification events are not flagged by CRS.** In the stress scenario, 6/125 Low flows are amplified; only 2 of them have a chaotic-band pedigree. The other 4 amplification events sit in the period-2 / period-4 band ($r \in [3.0, 3.55]$) where CRS is above the operational threshold ($\text{CRS} \approx 0.4\text{--}0.6$) and $\text{GSD_eff} = \sigma_{\text{ILCD}}$. Pedigree-driven chaotic regime detection and ground-truth amplification events are weakly correlated in this benchmark.

Conclusion of the controlled benchmark. On single-flow CRPS, neither single-knob naive widening nor CRS-augmented GSD_eff achieves a robust improvement over the ILCD baseline on the Low-quality stratum. The dynamical content of CRS adds no detectable single-flow forecast skill in the regime the framework is designed for. We read this as strong evidence for the **regime-label** positioning (§ 7.6): CRS should be advertised and operationalised as a *flag for human review of chaotic-band flows*, not as a substitute for σ_{ILCD} in the impact-assessment Monte-Carlo.

5.7 Real-distribution extension (tiangongDB metadata)

Why this benchmark. The reviewer of earlier drafts asked whether the conclusion above is robust to a more *realistic pedigree distribution*. The controlled benchmark of § 5.6 enumerates pedigree combinations uniformly; real LCA practice samples pedigrees from a distribution shaped by which flows are most often inventoried, which databases populate the practitioner's library, and which target study setting (year, location) drives the temporal and geographic scoring. We answer with a real-distribution extension whose pedigrees are derived from tiangongDB (Tian et al., open-access Chinese LCI, 4150 processes) under an EU 2026 study setting.

Honest scoping note. Neither tiangongDB nor EF 3.1 contains a measured σ_{obs} per flow nor an explicit pedigree (R, C, T, G, F). They are characterisation databases. What we *can* derive transparently are *plausible pedigrees* from process metadata under target-study assumptions, and we therefore use the same Student-t scale-mixture ground truth as § 5.6 (which gives no method a structural advantage). The benchmark is a robustness test of the v1 conclusion under a realistic pedigree distribution, not a claim about measured uncertainty in those databases. The pedigree-derivation rules (location $\rightarrow$ G, year $\rightarrow$ T, processType $\rightarrow$ R/F, exchange count $\rightarrow$ C) are documented in `benchmark/extend_real_db.py` and reproducible from a single seed.

Setup. From 4150 tiangongDB processes we sample at most 2 exchange-level pedigrees per process to avoid one-process dominance, derive (R, C, T, G, F) per exchange, and stratify-sample a corpus aligned with the controlled-benchmark shape (target $n_{\text{low}} = 250$, observed $n = 0/60/250$). The High stratum is empty by construction: every CN-located process gets $G \geq 3$ against an EU target, so no flow can satisfy the $\max \leq 2$ High condition. Same Student-t ground truth as § 5.6, same $K = 50$ observations per flow, same 50/50 calibration/evaluation split.

Results.

Stratum	n eval	ILCD CRPS	linear γ_{fit}	linear – ILCD	CRS – ILCD	CRS win-rate vs best naive
Mild scenario						
Medium	30	0.1116	0.92	–0.0014 (linear, p=0.100)	0.0000 (tie)	20 % vs sigmoidal (p=1.000)
Low	125	0.1794	0.92	–0.0070 ✓ p=0.000 (linear)	0.0000 (tie)	18 % vs linear (p=1.000)
Stress scenario						
Medium	30	0.1116	0.99	–0.0009 (linear, p=0.181)	0.0000 (tie)	27 % vs sigmoidal (p=0.997)
Low	125	0.2108	0.99	–0.0096 ✓ p=0.000 (linear)	0.0000 (tie)	15 % vs linear (p=1.000)

(The "tie" CRS-vs-ILCD entries reflect the same $\lambda \leq \theta \Rightarrow \text{GSD_eff} = \sigma_{\text{ILCD}}$ property as in § 5.6: on the real-distribution corpus, 100 % of Medium flows and ~100 % of Low flows return $\text{GSD_eff} = \sigma_{\text{ILCD}}$ exactly, because the realistic pedigree distribution rarely concentrates at ($R = 5$, $C = 5$, $T = 5$, $G = 5$, $F = 5$) simultaneously.)

The headline differs from § 5.6 in one important respect. The fitted γ on the linear baseline jumps from 0.23 to 0.92 (mild) and 0.28 to 0.99 (stress). This is not a calibration artefact — it reflects the *shape* of the realistic pedigree distribution: the tiangongDB-derived corpus contains many flows with moderate σ_{ILCD} (typically 1.05–1.30) but the heavy-tailed Student-t ground truth assigns them, on average, larger empirical σ_{obs} . A larger linear widening factor minimises CRPS by lifting the predictive interval to better cover the realised observations. The win-rate on the Low stratum is 82 % mild, 85 % stress, with $p_{\text{binom}} < 0.001$; the bootstrap 95 % CI on Δ_{CRPS} is entirely positive. The naive linear widening *does* improve on ILCD by 4–5 % CRPS gain (mild) and 5 % (stress) on Low.

CRS does not. As in § 5.6, CRS equals ILCD on the entirety of the Medium stratum and on essentially all of the Low stratum (the realistic distribution does not produce many ($R = 5$, $C = 5$, $T = 5$, $G = 5$, $F = 5$) pedigrees, so chaotic-band r values are rare). When CRS does differ from ILCD, the V5 widening is too aggressive for the realised σ_{obs} and CRS *loses* — 15–20 % win-rate against the best naive baseline ($p_{\text{binom}} = 1.000$).

Reading. The real-distribution evidence sharpens, not contradicts, the controlled-benchmark conclusion. A naive linear widening of σ_{ILCD} *does* help on realistic Low-quality corpora — but it is not a *dynamical* improvement, only a quality-monotone shift. The

dynamical content of CRS, on this body of evidence, neither helps nor harms; it just rarely engages, because the realistic pedigree distribution does not concentrate flows in the chaotic band. The framework's contribution must therefore come from the *labelling* of those rare chaotic-band flows when they do appear, not from the GSD widening.

6. A worked example

To make the regime-labelling concrete, we walk through three flows with very different pedigree profiles. All three are realecoinvent v3.10 flows; the specific identifiers are listed in the supplementary CSV.

Case A, High quality (a primary-data fluorinated gas measurement). Pedigree (1, 2, 1, 2, 2). DQR = 0.85 ("Excellent"). $\sigma_{ILCD} = 1.073$. $r = 2.5 + 1.5(1 - e^{-0.073/0.58}) \approx 2.677$. $\lambda \approx -0.42$ (negative, fixed-point regime). $h_{KS} = 0$. $d_{KY} = 1$. CRS = 0.94. Regime: *Point fixe*. $GSD_{eff} = \sigma_{ILCD} = 1.073$. **Operational read-out: use ILCD intervals as-is.**

Case B, Medium quality (a literature emission factor for an organic compound). Pedigree (3, 3, 3, 3, 3). DQR = 0.50 ("Acceptable"). $\sigma_{ILCD} \approx 1.357$. $r \approx 3.45$. $\lambda \approx +0.06$ (small positive). $h_{KS} \approx 0.06$. $d_{KY} \approx 1.06$. CRS ≈ 0.58 . Regime: *Période-2*, on the cusp of *Période-4*. $GSD_{eff} \approx 1.39$. **Operational read-out: flag this flow. The ILCD interval underestimates the propagated uncertainty by a small but systematic factor; sensitivity analysis recommended; document the regime label in the report.**

Case C, Low quality (an extrapolated Asian process used as a proxy for a European one in a sectoral study). Pedigree (4, 4, 4, 4, 4). DQR = 0.25 ("Faible"). $\sigma_{ILCD} \approx 1.66$. $r \approx 3.66$ (above $r_{\infty} = 3.57$). $\lambda \approx +0.31$ (firmly positive). $h_{KS} \approx 0.31$. $d_{KY} \approx 1.62$. CRS ≈ 0.32 . Regime: *Chaotique (fenêtres stables)*. $GSD_{eff} \approx 2.05$. **Operational read-out: use GSD_{eff} rather than σ_{ILCD} for the impact-assessment Monte-Carlo. Consider scenario analysis. Most importantly: report the regime label, the LCA reader needs to know that this flow's uncertainty is structurally amplifying.**

The three cases sit on the same r -axis. The framework, by construction, treats them differently in proportion to the regime they occupy, and identically to ILCD in Case A. That is the complementarity claim, made concrete on three real flows.

6.4 Waste-LCA case studies

We add three waste-LCA case studies that motivated this work and to which the framework is most likely to be applied in practice.

Case W1, Anaerobic-digestion biogas yield ($\text{Nm}^3 \text{CH}_4$ per t volatile solids, agricultural feedstock mix). Pedigree (3, 4, 3, 2, 3). The yield is well-instrumented at the demonstration scale ($R = 3$) but coverage of feedstock-mix variability across seasons is limited ($C = 4$); the dataset is recent ($T = 3$) and geographically representative for the studied plant ($G = 2$) but the technology is sufficiently configuration-dependent that the further-correlation score is mid-range ($F = 3$). $\text{DQR} = 0.40$ ("Acceptable"). $\sigma_{\text{ILCD}} \approx 1.51$. $r \approx 3.55$ (between the period-8 cascade and the chaotic threshold). $\lambda \approx +0.18$ (firmly positive). $\text{CRS} \approx 0.41$. Regime: *Pré-chaotique / Seuil chaos*. $\text{GSD}_{\text{eff}} \approx 1.74$. **Operational read-out: the ILCD GSD of 1.51 understates the propagated uncertainty for downstream impact categories (particularly GWP and acidification through the displaced-natural-gas credit). Use $\text{GSD}_{\text{eff}} = 1.74$ and report the regime label.** This matches the empirical observation, in our and others' multi-plant comparisons, that biogas-yield-driven Monte-Carlo intervals from the ILCD baseline routinely fail to cover the inter-plant range observed in real fleet operations.

Case W2, MSW incineration NO_x emission factor (kg NO_x per t MSW, modern grate furnace with SNCR). Pedigree (2, 3, 2, 3, 4). Continuous emission monitoring on the studied plant ($R = 2$), reasonable but not exhaustive coverage of operating regimes ($C = 3$), recent data ($T = 2$), same regional fleet ($G = 3$) but the technology score is high because the SNCR efficiency depends on a temperature window that is plant-specific ($F = 4$). $\text{DQR} = 0.40$. $\sigma_{\text{ILCD}} \approx 1.27$. $r \approx 3.13$. $\lambda \approx -0.06$ (small negative, just past the first bifurcation). $\text{CRS} \approx 0.62$. Regime: *Période-2*. $\text{GSD}_{\text{eff}} = \sigma_{\text{ILCD}} = 1.27$ (the engine collapses to ILCD because $\lambda \leq 0$). **Operational read-out: ILCD intervals are valid; the regime label "Période-2" flags that a small refinement of pedigree F (e.g. acquiring temperature-window data from the plant operator) could move the diagnostic across the bifurcation, and is therefore the highest-leverage data-collection target.** This is the kind of read-out the per-dimension Sobol index of Section 4 makes operational: in Case W2, $S_F \approx 0.42$, a clear indication of where the next data-collection euro should be spent.

Case W3, Landfill leachate chemical oxygen demand ($\text{mg O}_2/\text{L}$, municipal landfill, post-closure phase). Pedigree (4, 4, 4, 3, 5). Composite samples from a comparable but not identical landfill ($R = 4$), partial campaign coverage ($C = 4$), data partially older than the post-closure phase modelled ($T = 4$), same climatic region ($G = 3$), and substantial structural-correlation issues because the technology of the source landfill differs in liner design and leachate-collection efficiency ($F = 5$). $\text{DQR} = 0.20$ ("Faible"). $\sigma_{\text{ILCD}} \approx 1.79$. $r \approx 3.71$ (well above r_∞). $\lambda \approx +0.36$. $\text{CRS} \approx 0.27$. Regime: *Chaotique (fenêtres stables)*. $\text{GSD}_{\text{eff}} \approx 2.43$. **Operational read-out: this is the textbook case for which the framework was designed. The ILCD GSD of 1.79 produces a [p5, p95] interval of $[0.41 \times \text{COD}_0, 2.43 \times \text{COD}_0]$; the ChaoticLCA GSD_{eff} of 2.43 widens it to $[0.27 \times \text{COD}_0, 3.69 \times \text{COD}_0]$, and (more importantly) the regime label tells the LCA reader that the**

interval is a lower bound on the real uncertainty and primary data acquisition is the only honest path forward. In our experience reviewing LCA studies of post-closure landfill chains, the standard log-normal closure on COD is the single most consistent under-statement of impact-category uncertainty across the whole system.

The three waste cases share a feature that is structural and not coincidental: their pedigree profiles cluster around the cascade's bifurcation band ($r \in [3.0, 3.7]$), which is where the framework adds the most diagnostic value. Waste-treatment LCAs are, in our reading, *the* sub-field of LCA where reading the regime label alongside the GSD changes which decision the report supports. This positioning is consistent with the publishing programmes of *Waste Management Research, Resources, Conservation and Recycling*, and the Joint Research Centre's Best Available Techniques reference documents on waste treatment, which already identify pedigree-driven uncertainty as a methodological priority for the sector.

6.5 A 5×5 Leontief mini-case (didactic multi-flow test)

The motivation in § 1.2 invokes the *coupled* Leontief inversion of an uncertain technology matrix as the regime where the per-flow log-normal closure is structurally weakest. This V1 closes that argument's empirical hole, modestly, with a 5×5 didactic system inspired by an anaerobic-digestion + CHP loop. The setup, code at `benchmark/leontief_minicase.py`:

index	process	non-zero entries of A (col., median value)
1	Substrate collection	$A[1, 4] = 0.10$ (digestate return)
2	Anaerobic digester	$A[2, 1] = 0.80$ (substrate); $A[2, 5] = 0.15$ (grid)
3	CHP cogeneration	$A[3, 2] = 0.85$ (biogas); $A[3, 5] = 0.10$ (grid)
4	Digestate management	$A[4, 2] = 0.40$ (digestate from digester)
5	Grid electricity	$A[5, 3] = 0.05$ or 0.50 (loop-closer; two variants)

Final demand $f = (0, 0, 1, 0, 0)$ (1 unit of CHP electricity). Each non-zero $A[i, j]$ is sampled lognormally with log-GSD $\log \sigma_{ILCD}(R_{ij}, C_{ij}, T_{ij}, G_{ij}, F_{ij})$; we run $N = 50\,000$ Monte-Carlo draws of the Leontief solve $s = (I - A)^{-1} f$ and record s_3 (the kWh CHP electricity output). Two structural variants are reported:

- **Damped variant** ($A[5, 3] = 0.05$). The loop-closer is small; the system is well-conditioned and perturbations are damped through the matrix inversion.
- **Strong-loop variant** ($A[5, 3] = 0.50$). The loop-closer is large; the operator $(I - A)$ is closer to singular along the loop direction, perturbations on $(2, 5)$ and $(3, 5)$ are amplified rather than damped.

We compare the empirical [p5, p95] of s_3 to the ILCD per-flow GSD aggregation that would be applied to the output process using the *worst* upstream pedigree per dimension (the standardecoinvent practice for aggregated rows; Frischknecht et al., 2007).

Pedigree regime	Output ped.	σ_{ILCD}	GSD_eff	Regime label	Empirical GSD (damped, coupled)	ILCD log-width / empirical
High quality (all entries pedigree 1–2)	(2,2,2,2,2)	1.07	1.07	Point fixe	1.001	59 ×
Medium quality (pedigree 3)	(3,3,3,3,3)	1.17	1.17	Point fixe	1.003	54 ×
Low quality (chaotic-band entries)	(5,5,5,4,5)	1.84	2.20	Chaotique (fenêtres stables)	1.016	49 ×

(Strong-loop variant gives empirical [p5, p95] log-width up to ~ 1.7 on Low, still narrower than the ILCD prediction.)

Three observations.

- 1. The per-flow ILCD aggregation is conservative on this system, by an order of magnitude or more.** The empirical [p5, p95] interval of s_3 is roughly 50× narrower (in log-width) than the per-flow ILCD prediction at the High and Medium pedigrees, and $\sim 5\times$ narrower at Low. The widening of GSD_eff over σ_{ILCD} ($1.84 \rightarrow 2.20$ at Low pedigree) makes this conservatism marginally worse. We did not anticipate this finding; it is honestly reported.
- 2. The coupling effect is small on this 5×5 system.** Independent vs coupled lognormal draws on the shared (2, 5) and (3, 5) entries change the empirical GSD by 0.1 % (damped) to 1.5 % (strong loop, Low). The coupling is in the direction the framework predicts (heavier tail) but quantitatively negligible at this scale.
- 3. The CRS regime label correctly identifies the input pedigrees as chaotic** at Low quality (CRS = 0.59, regime = "Chaotique"), but the empirical multi-flow output of this didactic system does *not* exhibit the heavy-tailed amplification that motivates the framework. The motivation of § 1.2 must be nuanced: heavy-tailed multi-flow propagation requires either an ill-conditioned Leontief operator (a real possibility in production LCIs with hundreds of dimensions) *or* a structural

correlation pattern in the pedigree perturbations that this 5×5 system is too small to exhibit.

The honest reading is that the 5×5 didactic case does not falsify the framework but does not vindicate the multi-flow motivation either. Production-scale waste-treatment LCIs (typically 200–500 processes) have a richer correlation structure and operator conditioning that the 5×5 cannot capture; the V6 transfer-entropy network layer (§ 7.3) is the next deliverable on that front.

Next step: ill-conditioned and loop-rich systems. A natural priority for V2 is a deliberately *ill-conditioned* test case ($(\mathbf{I} - \mathbf{A})$ near-singular along one mode), which is the regime in which Leontief inversion theoretically amplifies rather than damps perturbations. Waste-treatment systems are an unusually good candidate for that test, because they routinely close *physical* loops at the inventory boundary: digestate from anaerobic digestion returns to crop fertilisation (pedigree-uncertain), bottom and fly ashes from incineration return to road-base or cement-kiln backfill (pedigree-uncertain), leachate-collection effluents return to wastewater treatment whose sludge re-enters the agricultural chain. Each such loop turns a column of $\mathbf{A}$ into a quasi-recursive flow whose multiplicative perturbation is amplified through $(\mathbf{I} - \mathbf{A})^{-1}$. We will report a 50×50 waste-loop case in V2 with measured pedigrees from the three case studies of § 6.4.

7. Discussion

7.1 Complementarity, not replacement

The framework is engineered, at multiple layers, to coincide with the ILCD baseline at high data quality. The mapping $\sigma \rightarrow r$ sends $\sigma = 1$ to the fixed-point regime; the `gsd_eff` formula collapses to σ_{ILCD} whenever $\lambda \leq 0$; the CRS reaches its upper bound when the dynamic content is absent. None of these properties is discovered: they are *built in* and tested against the engine's golden file at 1e-9 tolerance.

This matters for adoption. The framework asks the LCA practitioner for nothing they do not already produce (five integer pedigree scores), and it never claims to overrule the GSD they are required to compute under EN 15804+A2 or PEP Ecopassport methodologies. It adds a regime label and a $[0, 1]$ score that the practitioner can include or omit from the report at will. The two-line addendum proposed in the JRC technical note (Piens, in prep.) is entirely additive.

7.2 Calibration risk and how we manage it

The single calibration knob σ_{ref} is the framework's only point of fragility. Section 5.5 documents the plateau-shape of the win-rate around the V5 default. We recommend three governance practices:

1. **Default $\sigma_{\text{ref}} = 0.58$** for general use, established on the 128-flow ecoinvent corpus and locked at engine version V5.
2. **Per-organisation re-calibration** permitted when at least 30 primary-data flows are available, using the `engine-py/benchmarks/calibrate_sigma_ref.py` script. The new value must be reported in the LCA metadata. *Bayesian re-calibration* (Appendix G) replaces this fit with a posterior update in V7.
3. **Versioning.** Future calibrations live under new engine versions (V5, V6 Nexus, V7 ...) so that prior reports remain reproducible. Backward-compatibility metadata includes `engineVersion` and the calibration value used.

7.3 The internal robustness layers (TRT, ensemble, copula) are not headline-changing, and what V6 Nexus changes

We are transparent that the V5 layers of Section 3.10 (Transfer Resonance Tree, 50-member ensemble, Clayton copula) do not change the headline conclusion on the 128-flow benchmark for *single-flow* analysis. On the (3, 3, 3, 3, 3) golden test, the ensemble produces $\text{std_gsd} \approx 4.4 \cdot 10^{-16}$ (numerical zero); the TRT changes GSD_eff by less than 0.1 %; the copula reduces to its marginal because there is only one flow. We retain the layers because (a) they generalise naturally to the multi-flow V6 Nexus engine where coupled propagation makes the copula non-trivial, (b) they expose a *distribution* on the diagnostic that reviewers expect, and (c) honest reporting is preferable to silent removal.

V6 Nexus extensions. The V6 documentation describes four substantive additions over V5 that are worth flagging here, even though they are not load-bearing for the present paper:

1. **CRS expanded to 8 components.** V5's five components (h_{KS} , d_{KY} , $|\lambda|$, spectral entropy, ensemble spread) are extended in V6 by a Temporal Recurrence Coefficient (TRC), permutation entropy à la Bandt–Pompe (PE), Gelman–Rubin convergence diagnostic $\hat{R}$, and a phase-coherence statistic. The 8-component CRS is reported as a `crs_v6` field alongside the 5-component `crs_dynamic` of V5; the operational decision rule (Figure 6) is unchanged.
2. **Adaptive Gaussian RDS.** V5's RDS noise is bounded by $\sigma_{\text{RDS}} = 0.05 \cdot \max(0, \sigma_{\text{ILCD}} - 1)$. V6 replaces it with an adaptive Gaussian whose amplitude tracks the realised trajectory variance, a more defensible probe of structural-vs-noise regime separation.

3. **Gelman–Rubin diagnostic.** V6 runs four parallel RDS chains and reports $\hat{R}$ per flow. Convergence is achieved ($\hat{R} \leq 1.1$) on 96 % of the 128 ecoinvent flows; the remaining 4 % are flagged as *non-stationary*, a category the V5 engine cannot detect. We treat the 4 % as a known limit (Section 8, L8).
4. **Transfer-entropy network layer ($\tau(i \rightarrow j)$).** V6 quantifies inter-flow dependence at the network level via Schreiber's transfer entropy on the Bandt–Pompe ordinal histogram. This is the conceptually central addition because it begins to address the multi-flow regime that Section 1.2 identified as the deepest limit of the per-flow ILCD closure. The TE network layer is *exploratory*, it has been validated on the 128-flow corpus but has not been benchmarked against measured network-level σ_{obs} because no such ground truth exists at scale.
5. **Stochastic LCIA layer.** V6 propagates 500 Monte-Carlo draws of the upstream uncertainty through the LCIA characterisation step (CML, ReCiPe, USEtox, EF 3.1). This is a downstream-impact analogue of *pedigree_uncertainty* and exposes the variance of the impact intervals attributable to the inventory uncertainty.

These additions improve robustness and extend the diagnostic to a larger set of cases. **They do not change the philosophical posture of the paper**, the framework remains a complementary diagnostic, identical to ILCD where ILCD is supported, and adds a regime label where it is weakest. The companion paper (Piens, in prep.) reports the V6 evaluation in full.

7.4bis Operational threshold revision: 0.4 rather than 0.5

The original operational threshold of $\text{CRS} \leq 0.5$ was set conservatively to maximise sensitivity to low-quality flows. The V6 binomial result (Section 5.1: 17/38 wins on Low quality, $p = 0.084$) and the false-alarm analysis (Section 5.4 F1: ~ 8 % of flows over-flagged near $r \approx 3.45$) jointly suggest that **a tighter threshold of $\text{CRS} \leq 0.4$ is operationally preferable**. The trade-off is straightforward:

Threshold	Sensitivity (catches truly chaotic flows)	False-alarm rate	Statistical-power requirement
$\text{CRS} \leq 0.5$ (V5 default)	high	~ 8 % near $r=3.45$	very high
$\text{CRS} \leq 0.4$ (V6 recommendation)	slightly lower	~ 3 %	moderate
$\text{CRS} \leq 0.3$	low	< 1 %	low

We adopt $\text{CRS} \leq 0.4$ as the operational threshold for the JRC-track addendum (Section 7.5) and update Figure 6 / Appendix F accordingly. The 0.5 threshold is retained as a documented conservative alternative.

7.4 Reviewer concerns we anticipate

"Is $r = 2.873$ enough to call dynamics chaotic?", No. At $r < 3$ the regime label is *Point fixe* and $\lambda < 0$. CRS at $r = 2.873$ is in the 0.7–0.8 band ("data fiable"), and the gain from CRS over ILCD is intentionally small. The framework only changes behaviour where the regime warrants it.

"Sobol indices presume input independence.", Accepted, and tested in V1. Pedigree dimensions co-vary, and Section 4.3 reports a robustness check under R-T correlation $\rho \in \{0, 0.3, 0.6\}$: the qualitative ordering of indices (which dimension is largest, which is second) is preserved at all three correlation levels and at three pedigree profiles (Excellent, Median, Faible). The quantitative magnitudes change (S_T grows by a factor ~ 2 from $\rho = 0$ to $\rho = 0.6$) but the operational read-out (which pedigree dimension to refine first) is invariant on this body of evidence. Our first-order indices are therefore interpretable as variance shares under the perturbation law we explicitly impose, with the qualitative conclusion robust to a moderate co-variance.

"Is the 38-flow Low-quality stratum enough to reach a definitive conclusion?", No, and that limit is now closed. The V1 reports two benchmarks (controlled, § 5.6, $n_{\text{low_eval}} = 125$; real-distribution, § 5.7, $n_{\text{low_eval}} = 125$) whose joint statistical power on the Low stratum is comfortably above the threshold a reviewer would expect. The conclusion is consistent across both benchmarks and across mild and stress scenarios: CRS does not improve single-flow CRPS over ILCD on Low. We accept that as evidence and reposition the framework as a regime-label diagnostic.

"Why not Box-Behnken design on σ_i directly?", Because the operational uncertainty mode of pedigree raters is the *integer-level shift*, not the continuous σ multiplier. Box-Behnken is the right tool for σ -calibration; ± 1 MC is the right tool for rater-bias attribution. We do both at different layers.

"Why not a fully Bayesian framework?", A Bayesian extension is straightforward (priors on σ_{ref} , on the σ -table itself, on per-flow λ) and natural. It is not the right first paper: a deterministic single-knob calibration that exposes one Sobol decomposition is easier to audit, easier to standardise, and easier to deploy in non-research LCA practices. The Bayesian version is on the V7 roadmap (Appendix G).

"Have you compared against a non-chaotic widening baseline?", Yes, in V1 (§ 5.6, § 5.7). Four single-knob non-chaotic widenings (linear, quadratic, exponential, sigmoidal) were fitted by held-out CRPS minimisation. Three of the four collapse to the ILCD baseline ($\gamma \approx 0$). Only the linear form yields a non-trivial fit and *does* beat ILCD on Low — by 4–9 %

CRPS gain depending on the corpus and scenario. CRS does not beat that linear baseline on any of the strata $\times$ scenarios we tested. The honest reading is that the dynamical content of CRS adds no detectable single-flow forecast skill *over a much simpler closure*. The framework's contribution must therefore be positioned in the regime-label and qualitative-flag dimension, not in the GSD widening.

7.5 Standardisation pathway

A separate technical note (Piens, in prep., addressed to the JRC LCA Unit and to theecoinvent Methodology Working Group) proposes a non-binding addendum:

In addition to the ILCD pedigree GSD of [X], we report a Chaotic Risk Score CRS = [Y]. CRS reflects the dynamic regime of the propagated uncertainty (regime: [point fixe / période-2 / pré-chaotique / chaotique]). The regime label is the primary diagnostic; the recommendation is that flows entering the chaotic band ($r \geq r_\infty$, $\lambda > 0$) be subject to review by a domain expert and, when feasible, primary data acquisition. The CRS numeric value should be reported for transparency. CRS is not a substitute for qualitative data quality review: it cannot detect a flow whose pedigree is well-scored but whose underlying measurement protocol is biased, nor a flow whose context-of-use makes the LCI category inappropriate. It complements the practitioner's judgment, never replaces it.

The proposal has three operational characteristics: (i) it is computed by an open-source reference engine (no black-box vendor dependency); (ii) it does not alter the ILCD GSD calculation; (iii) it does not change the impact-assessment characterisation step. It only adds a diagnostic flag for low-quality flows where the closure is most likely to mislead. We frame it explicitly as a *technical note for consultation, not a regulatory proposal*, and we welcome the methodological community's revisions. The benchmarks of § 5.6 and § 5.7 do not (yet) support a per-flow decision rule of the form "flag when $\text{CRS} \leq \tau$ " against measured σ_{obs} ; empirical calibration of any such rule on a primary-data corpus is a prerequisite, and is the subject of the primary-data benchmark in preparation.

7.6 The human-environmental ethics of regime labelling

We close the discussion with the question that motivated this paper. An LCA report informs a real environmental decision, a procurement choice, a label on a product, a public policy. The practitioner who writes the report bears a responsibility to the human reader: to expose, as transparently as possible, the cases where the report's headline numbers are likely to mislead.

The ILCD pedigree closure is, on the whole, an extraordinary instrument of that responsibility. It compresses a five-dimensional qualitative judgment into a single computable σ that travels intact through impact assessment. It is shared internationally and

audited by an active community. The framework we propose adds nothing to its operational role (it does not change the σ that is propagated, it does not change the impact intervals downstream) except a *regime label* and a *score in [0, 1]* that names the cases where the closure is structurally weakest.

This is what we mean by complementary. The closure does the heavy work; the regime label restores the epistemic honesty the closure cannot afford by itself, because it is a *closure*. The framework is, in that sense, less an addition to LCA methodology than a mirror placed next to one of its existing computations: when the mirror reflects "chaotic", the practitioner is invited to look harder before signing off on a number with three significant figures.

We do not argue that an LCA reader needs to understand period-doubling cascades to act on a regime label; the documentation of the web demonstrator (theopiens.com/analyzer) and the JRC note are explicit on this point, *no need to understand the mathematics to act* on the diagnostic. We argue, modestly, that the discipline benefits from a tool that signals "this number is more uncertain than the GSD admits" without pretending to compute exactly how much more. That signal is what the regime label is for.

8. Limitations and threats to validity

We aggregate the limits scattered across the paper.

L1, σ_{ref} is a single point of fragility. The plateau in win-rate around 0.58 is a partial mitigation. A formal multi-organisation calibration trial across more than one source database would strengthen the result.

L2, TRT depth fixed at 4. No tuning per case is performed. This conservatively under-uses the V5 layer.

L3, Hénon (2D) mode is implemented but ignored in this paper. The 2D extension is a natural way to capture two-flow correlations and is on the V6 Nexus roadmap.

L4, Pedigree co-variance is not modelled at the score-shift layer. A copula on (R , C , T , G , F) shifts is implementable; the symmetric ± 1 default is a deliberately simple operational baseline.

L5, Empirical evaluation is on synthetic ground truth. Both the controlled benchmark (§ 5.6) and the real-distribution extension (§ 5.7) use a Student-t scale-mixture as ground truth σ_{obs} . The σ_{obs} values are not measured. The benchmarks are reproducible bit-for-bit and the generative process intentionally embeds the heavy-tailed regime the framework was designed to detect, but they cannot answer the empirical question of whether real LCA inventories at low quality exhibit σ_{obs} ceiling at a frequency and magnitude that single-flow CRS would catch. The primary-data benchmark (Piens, in prep.) remains the recommended next deliverable.

L6, Multi-flow evaluation is limited to the 5×5 didactic case (§ 6.5). The multi-flow regime that motivates § 1.2 (Leontief inversion of an uncertain technology matrix) has been tested in V1 only on a 5×5 didactic system inspired by an anaerobic-digestion + CHP loop. That mini-case shows the per-flow ILCD aggregation is *conservative* by an order of magnitude, *not* under-spread, on this well-conditioned operator. Whether production-scale LCIs (200–500 dimensions, richer correlation structure, occasionally near-singular operators) exhibit the heavy-tailed amplification regime is the central methodological frontier. The V6 Nexus engine adds a transfer-entropy network layer ($T(i \rightarrow j)$) for that test; its quantitative evaluation against measured network-level σ_{obs} is not in scope here.

L7, retired. The naive non-chaotic widening baseline of earlier drafts is now computed (§ 5.6, § 5.7). The result is that single-knob naive widening — only the linear form returns a non-trivial fit — does beat ILCD on Low by 4–9 % CRPS gain, while the dynamical CRS-augmented GSD_eff does not. This sharpens the framework's positioning toward the regime-label use case rather than retiring the question.

L8, Possible non-stationarity / Gelman–Rubin failure on a minority of flows. The V6 documentation reports $\hat{R} \leq 1.1$ on 96 % of the historical 128-flow corpus; the remaining 4 % do not reach the convergence threshold and are flagged by the engine as *non-stationary*. The V5 engine of the present paper has no equivalent diagnostic, so these flows are silently passed through. We note the limit and recommend the V6 $\hat{R}$ reporting as a future operational practice.

L9, The V6 Nexus transfer-entropy layer is exploratory. The TE network layer has been validated only at the bundle level (qualitative behaviour, network plausibility checks, no quantitative ground truth). We do not claim it is benchmarked. We claim only that it is the conceptually correct addition to address the multi-flow limit (L6) and that its empirical validation is the next major frontier (Appendix G).

L10, Reproducibility hinges on engine-py. The reference engine is MIT-licensed and tested; the web demonstrator runs entirely client-side. We commit to maintaining the package across at least two engine versions before deprecating any field.

L11, The σ_{obs} ground truth in both benchmarks is synthetic. This is essentially L5 restated as a methodological caveat: the controlled and real-distribution benchmarks are reproducible by construction but are idealised generative processes (Student-t scale-mixtures of two parametrisations). The remaining empirical question — whether real LCA inventories on chaotic-band flows exhibit measured σ_{obs} patterns that single-flow CRS would catch — only primary-data measurement can answer.

9. Conclusion

We have proposed a complementary uncertainty diagnostic for life-cycle inventory data, using the logistic recursion $x_{n+1} = r \cdot x_n \cdot (1 - x_n)$ as a heuristic parametric probe of the regime of pedigree-driven uncertainty. The diagnostic is a five-component Chaotic Risk Score in $[0, 1]$ computed from a deterministic single-parameter mapping $\sigma_{\text{ILCD}} \rightarrow r$, plus a pedigree_uncertainty Monte-Carlo module that exposes per-dimension first-order Sobol indices. The framework is engineered, by construction, to coincide with the ILCD/Gaussian baseline at high data quality ($\lambda \leq 0 \Rightarrow \text{GSD}_{\text{eff}} = \sigma_{\text{ILCD}}$ exactly).

Empirically, on the controlled benchmark v1 (493 pedigree flows, two scenarios; § 5.6) and the real-distribution extension (310 tiangongDB-derived flows; § 5.7), the V5 reference engine **does not improve** the CRPS against the ILCD baseline on the Low-quality stratum (CRS $\equiv$ ILCD on 88–91 % of Low flows by the engineered $\lambda \leq 0$ property; CRS *loses* on the chaotic-band remainder). A naive linear widening of σ_{ILCD} does beat ILCD on Low by 4–9 % CRPS gain, but is not dynamical. We read this as evidence that the framework's value, if any, is in *labelling the regime* and not in *replacing the GSD*. A 5×5 didactic Leontief mini-case (§ 6.5) further shows that, on the well-conditioned operator we tested, the per-flow ILCD aggregation is *conservative* by an order of magnitude, not under-spread — so the multi-flow motivation of § 1.2 must be nuanced: heavy-tailed amplification at the network level, the regime CRS was designed for, requires either an ill-conditioned operator or a stronger correlation structure than the 5×5 can exhibit. A primary-data benchmark on the chaotic-band sub-stratum, where σ_{obs} is measured rather than synthesised, remains the only test that could rehabilitate a per-flow CRS-based decision rule.

A robustness ablation across $r_{\text{min}} \in \{2.0, 2.5, 3.0\}$ and an alternative quadratic saturation (§ 5.5) confirms the conclusion is invariant to those calibration choices, and a Sobol-indices test under R-T correlation $\rho \in \{0, 0.3, 0.6\}$ (§ 4.3) preserves the qualitative ordering of variance attribution. The reference implementation (Python, MIT) reproduces the V5 web engine to 1e-9 on deterministic fields and 1e-6 on Monte-Carlo aggregates. The V6 Nexus engine (8-component CRS, adaptive Gaussian RDS, Gelman–Rubin convergence, transfer-entropy network layer, stochastic LCIA) is documented and forms the subject of a companion paper (Piens, in prep.). A standardisation pathway, in the form of a technical note addressed to the JRC LCA Unit and theecoinvent Methodology Working Group, proposes a non-binding addendum to the LCA report whose primary diagnostic is the *regime label* rather than a CRS threshold against measured σ_{obs} .

The contribution is, in the end, modest in scope and explicit about its limits. We do not claim a new theory of LCA uncertainty, and we do not claim that life-cycle inventories are dissipative dynamical systems in any strict universality-class sense. We claim that a deterministic, single-knob, fully-auditable diagnostic (that coincides with ILCD where ILCD is supported, and adds a regime label where it is structurally weakest) supplies an additional

and useful signal for the human reader of an LCA report — provided it is read as a *flag for human review*, not as a substitute for the GSD. We submit this V1 in the conviction that an honest negative-on-CRPS but informative-on-regime evaluation is more useful to the LCA community than a louder, less well-supported claim. The honest signal, when it can be computed at no additional input cost and audited through an open reference engine, deserves a place in the LCA toolbox — kept in the place the evidence supports, no further.

Acknowledgements

The author thanks the contributors to ecoinvent, the European Commission JRC, and the Tiangong LCA team for the open inventory databases on which this work depends. The framework was developed in the course of LCA studies on waste-treatment chains (anaerobic digestion, MSW incineration, post-closure landfill leachate, eco-material recovery for construction) where the limits of the standard log-normal closure became operationally undeniable. The author also thanks the anonymous reviewers (in advance) for their willingness to engage with a paper that proposes an uncomfortable bridge between two largely disjoint communities, applied LCA practice and nonlinear-dynamics methodology. Any remaining errors are his own.

No external funding was received for this work; the project is the result of independent consulting practice. No conflict of interest is declared.

Author contributions

T.P. conceived the framework, implemented the reference engine, designed and conducted the benchmarks, and wrote the manuscript.

Data and code availability

Reference implementation (Python, MIT): `engine-py/` of the ChaoticLCA workspace. Web demonstrator (entirely client-side): <https://theopiens.com/analyzer>. The controlled benchmark v1 (§ 5.6, 493 flows, two scenarios), the real-distribution extension (§ 5.7, 310 flows, two scenarios), the `r_min` ablation (§ 5.5), the Sobol-correlation robustness check (§ 4.3), the 5×5 Leontief mini-case (§ 6.5) and the figure-rendering notebooks will be deposited at Zenodo on first stable release, alongside this preprint. All benchmark scripts are at

benchmark/ of the project workspace and reproduce bit-for-bit from documented seeds. The primary-data benchmark on chaotic-band measured σ_{obs} will be released as a separate Zenodo deposit when ready (Piens, in prep.).

References

(Numbering follows the order of first citation.)

1. Weidema, B. P., & Wesnæs, M. S. (1996). Data quality management for life cycle inventories, an example of using data quality indicators. *Journal of Cleaner Production*, 4(3-4), 167–174.
2. Frischknecht, R., Jungbluth, N., Althaus, H.-J., Doka, G., Heck, T., Hellweg, S., Hirschler, R., Nemecek, T., Rebitzer, G., & Spielmann, M. (2007). Overview and methodology. *ecoinvent report No. 1*. Swiss Centre for Life Cycle Inventories, Dübendorf.
3. Ciroth, A., Muller, S., Weidema, B., & Lesage, P. (2013). Empirically based uncertainty factors for the pedigree matrix in ecoinvent. *International Journal of Life Cycle Assessment*, 21, 1338–1348.
4. European Commission, Joint Research Centre, Institute for Environment and Sustainability (2010). *International Reference Life Cycle Data System (ILCD) Handbook (General guide for Life Cycle Assessment) Detailed guidance*. First edition. EUR 24708 EN. Luxembourg: Publications Office of the European Union.
5. Slob, W. (1994). Uncertainty analysis in multiplicative models. *Risk Analysis*, 14(4), 571–576.
6. Limpert, E., Stahel, W. A., & Abbt, M. (2001). Log-normal distributions across the sciences: keys and clues. *BioScience*, 51(5), 341–352.
7. Heijungs, R., & Lenzen, M. (2014). Error propagation methods for LCA, a comparison. *International Journal of Life Cycle Assessment*, 19(7), 1445–1461.
8. Lloyd, S. M., & Ries, R. (2007). Characterizing, propagating, and analyzing uncertainty in life-cycle assessment: a survey of quantitative approaches. *Journal of Industrial Ecology*, 11(1), 161–179.
9. Heijungs, R., & Suh, S. (2002). *The computational structure of life cycle assessment*. Eco-efficiency in industry and science, 11. Dordrecht: Kluwer Academic Publishers.
10. Leontief, W. W. (1936). Quantitative input and output relations in the economic system of the United States. *Review of Economics and Statistics*, 18(3), 105–125.

11. Feigenbaum, M. J. (1978). Quantitative universality for a class of nonlinear transformations. *Journal of Statistical Physics*, 19(1), 25–52.
12. Cvitanović, P., Gunaratne, G., & Procaccia, I. (1988). Topological and metric properties of Hénon-type strange attractors. *Physical Review A*, 38(3), 1503–1520.
13. Strogatz, S. H. (2018). *Nonlinear dynamics and chaos: with applications to physics, biology, chemistry, and engineering*. Second edition. Boca Raton: CRC Press.
14. Devaney, R. L. (1989). *An introduction to chaotic dynamical systems*. Second edition. Reading, MA: Addison-Wesley.
15. Kolmogorov, A. N. (1958). A new metric invariant of transitive dynamical systems and automorphisms of Lebesgue spaces. *Doklady Akademii Nauk SSSR*, 119, 861–864. (Translation in *Selected Works of A. N. Kolmogorov*, Vol. III.)
16. Sinai, Ya. G. (1959). On the concept of entropy of a dynamical system. *Doklady Akademii Nauk SSSR*, 124, 768–771.
17. Pesin, Ya. B. (1977). Lyapunov characteristic exponents and smooth ergodic theory. *Russian Mathematical Surveys*, 32(4), 55–114.
18. Kaplan, J. L., & Yorke, J. A. (1979). Chaotic behavior of multidimensional difference equations. In: *Functional Differential Equations and Approximation of Fixed Points*. Lecture Notes in Mathematics, 730, 204–227. Springer.
19. May, R. M. (1976). Simple mathematical models with very complicated dynamics. *Nature*, 261, 459–467.
20. Sobol', I. M. (1993). Sensitivity estimates for nonlinear mathematical models. *Mathematical Modelling and Computational Experiments*, 1(4), 407–414.
21. Saltelli, A., Ratto, M., Andres, T., Campolongo, F., Cariboni, J., Gatelli, D., Saisana, M., & Tarantola, S. (2008). *Global sensitivity analysis: the primer*. Chichester: Wiley.
22. Welch, P. D. (1967). The use of fast Fourier transform for the estimation of power spectra: a method based on time averaging over short, modified periodograms. *IEEE Transactions on Audio and Electroacoustics*, AU-15(2), 70–73.
23. Hersbach, H. (2000). Decomposition of the continuous ranked probability score for ensemble prediction systems. *Weather and Forecasting*, 15(5), 559–570.
24. Gneiting, T., & Raftery, A. E. (2007). Strictly proper scoring rules, prediction, and estimation. *Journal of the American Statistical Association*, 102(477), 359–378.
25. Wilcoxon, F. (1945). Individual comparisons by ranking methods. *Biometrics*, 1(6), 80–83.

26. Wernet, G., Bauer, C., Steubing, B., Reinhard, J., Moreno-Ruiz, E., & Weidema, B. (2016). The ecoinvent database version 3 (part I): overview and methodology. *International Journal of Life Cycle Assessment*, 21(9), 1218–1230.
27. European Commission (2018). Product Environmental Footprint Category Rules (PEFCR) Guidance, version 6.3. EU Publications Office.
28. ISO (2006). *ISO 14040:2006 Environmental management (Life cycle assessment) Principles and framework*. International Organization for Standardization.
29. ISO (2006). *ISO 14044:2006 Environmental management (Life cycle assessment) Requirements and guidelines*. International Organization for Standardization.
30. CEN (2019). *EN 15804:2012+A2:2019 Sustainability of construction works (Environmental product declarations) Core rules for the product category of construction products*. Comité Européen de Normalisation.
31. AFNOR (2015). *NF EN 50693:2014 Procédure pour l'évaluation environnementale des produits électroniques et électriques*. PEP Ecopassport reference.
32. Bachmann, T. M. (2006). *Hazardous substances and human health: exposure, impact and external cost assessment at the European scale*. Trace Metals and Other Contaminants in the Environment, 8. Amsterdam: Elsevier.
33. Henderson, A. D., Hauschild, M. Z., van de Meent, D., Huijbregts, M. A. J., Larsen, H. F., Margni, M., McKone, T. E., Payet, J., Rosenbaum, R. K., & Jolliet, O. (2011). USEtox fate and ecotoxicity factors for comparative assessment of toxic emissions in life cycle analysis: sensitivity to key chemical properties. *International Journal of Life Cycle Assessment*, 16(8), 701–709.
34. Huijbregts, M. A. J., Steinmann, Z. J. N., Elshout, P. M. F., Stam, G., Verones, F., Vieira, M. D. M., Hollander, A., Zijp, M., & van Zelm, R. (2017). ReCiPe2016: a harmonised life cycle impact assessment method at midpoint and endpoint level. *International Journal of Life Cycle Assessment*, 22(2), 138–147.
35. Fantke, P. et al. (2018). Toward harmonizing ecotoxicity characterization in life cycle impact assessment. *Environmental Toxicology and Chemistry*, 37(12), 2955–2971.
36. Heijungs, R. (2010). Sensitivity coefficients for matrix-based LCA. *International Journal of Life Cycle Assessment*, 15(5), 511–520.
37. Lenzen, M. (2000). Errors in conventional and Input-Output-based life-cycle inventories. *Journal of Industrial Ecology*, 4(4), 127–148.

38. Bjørn, A., Owsianiak, M., Molin, C., & Hauschild, M. Z. (2018). LCA history and applications. In *Life Cycle Assessment: Theory and Practice* (pp. 17–30). Cham: Springer.
39. Eckmann, J.-P., & Ruelle, D. (1985). Ergodic theory of chaos and strange attractors. *Reviews of Modern Physics*, 57(3), 617–656.
40. Ott, E. (2002). *Chaos in dynamical systems*. Second edition. Cambridge: Cambridge University Press.
41. Cvitanović, P., Artuso, R., Mainieri, R., Tanner, G., & Vattay, G. (2020). *Chaos: classical and quantum*. Niels Bohr Institute, Copenhagen. (Open-access webbook; ChaosBook.org)
42. Schuster, H. G., & Just, W. (2005). *Deterministic chaos: an introduction*. Fourth, revised and enlarged edition. Weinheim: Wiley-VCH.
43. Kantz, H., & Schreiber, T. (2004). *Nonlinear time series analysis*. Second edition. Cambridge: Cambridge University Press.
44. Crutchfield, J. P., & Young, K. (1989). Inferring statistical complexity. *Physical Review Letters*, 63(2), 105–108.
45. Tarantola, A. (2005). *Inverse problem theory and methods for model parameter estimation*. Philadelphia: Society for Industrial and Applied Mathematics.
46. Saltelli, A., Annoni, P., Azzini, I., Campolongo, F., Ratto, M., & Tarantola, S. (2010). Variance-based sensitivity analysis of model output: design and estimator for the total sensitivity index. *Computer Physics Communications*, 181(2), 259–270.
47. Tibshirani, R. J., & Efron, B. (1993). *An introduction to the bootstrap*. Monographs on Statistics and Applied Probability, 57. Boca Raton: Chapman & Hall/CRC.
48. Heijungs, R. (2017). On the number of Monte Carlo runs in comparative probabilistic LCA. *International Journal of Life Cycle Assessment*, 22(7), 1175–1180.
49. Beck, T., Bos, U., Wittstock, B., Baitz, M., Fischer, M., & Sedlbauer, K. (2010). LANCA Land Use Indicator Value Calculation in Life Cycle Assessment, Method Report. Stuttgart: Fraunhofer Verlag.
50. Piens, T. (in prep.). *Chaotic Risk Score (CRS) as a complementary uncertainty metric to the ILCD pedigree closure, technical note for JRC and theecoinvent Methodology Working Group*. Companion to the present paper; draft v1 available on request.

51. Piens, T. (2026). *ChaoticLCA, public documentation, user guide and web demonstrator*. Online resource at <https://theopiens.com/docs> and <https://theopiens.com/analyzer> (V5 and V6 Nexus engine documentation; bilingual French / English; 100 % client-side computation; reference Python implementation at <https://github.com/theopiens/chaoticlca>). Accessed 2026-05-04.

Appendix A, Implementation snapshot

Reference engine (Python ≥ 3.10), MIT licence. Exhaustive listings are at `engine-py/` of the ChaoticLCA workspace; the call site below reproduces every numerical field of the V5 web engine on the canonical golden file `chaoticlca_1776709789670.jsonld` to within $1e-9$ on deterministic fields and $1e-6$ on Monte-Carlo aggregates.

```
from chaoticlca import analyze, pedigree_uncertainty

# Headline analysis (V5 engine) on a single flow
a = analyze(pedigree={"R": 4, "C": 3, "T": 4, "G": 3, "F": 4},
            engine_version="v5")
# Returns: dqr, quality, sigma_ilcd, r, lambda, hKS, dKY, regime,
# crs (4-comp baseline), crs_dynamic (5-comp V5), gsd_eff,
# interval{p5, p95}, ensemble{mean_gsd, std_gsd, spread},
# predictionHorizon, nBif, sigmaComponents, ergodicBounds, ...

# Pedigree-uncertainty Monte-Carlo (priorité 3)
u = pedigree_uncertainty(R=4, C=3, T=4, G=3, F=4,
                        n=10000, use_v5_crs=True)
# Returns: crs_p5, crs_p50, crs_p95, crs_mean, crs_std,
# sensitivity_per_dim{R, C, T, G, F}, samples (np.ndarray)
```

Command-line interface:

```
PYTHONPATH=. python3 -m chaoticlca analyze \
  --R 4 --C 3 --T 4 --G 3 --F 4 --engine v5

PYTHONPATH=. python3 -m chaoticlca uncertainty \
  --R 4 --C 3 --T 4 --G 3 --F 4 --n 10000 --use-v5-crs
```

The package ships 16 golden-equivalence tests against the V5 web bundle and a `examples/demo_pedigree_mc.py` script that walks through four pedigree levels (excellent / médian / faible / critique) and reproduces the regime-transition figure.

Appendix B, Calibration sensitivity (σ_{ref} sweep)

The sweep is performed on the historical 128-flow corpus, $\sigma_{\text{ref}} \in [0.30, 0.80]$ step 0.025. Aggregate metrics reported by the V5 documentation: paired Wilcoxon rank-biserial effect size, mean CRPS gain over Gaussian. The optimum is a plateau between 0.50 and 0.65 rather than a knife-edge; the V5 default $\sigma_{\text{ref}} = 0.58$ sits in the middle. The full plot is exposed as Supplementary Figure S1 in the engine-py package and on the public documentation page (</docs#scale-param>). The complementary r_{min} ablation of § 5.5 (V1) confirms the headline conclusion is invariant to $r_{\text{min}} \in \{2.0, 2.5, 3.0\}$ and to a quadratic alternative saturation. The leave-one-source-out cross-database validation remains a sensible next step (the controlled benchmark is synthetic, the real-distribution extension is single-source tiangongDB) and is the planned addition for the primary-data benchmark in preparation.

Appendix C, Pedigree score notation guide (excerpt)

For the practitioner's convenience, an excerpt of the score notation guide from the web demonstrator (theopiens.com/analyzer, doc tab "Guide de notation des scores"). The full table is in the documentation page; what follows is the score-1 / score-3 / score-5 anchor for each dimension.

- **R, Representativeness.** 1: Measured in situ, validated protocol, peer-reviewed (e.g. primaryecoinvent dataset). 3: Calculated from other measured data (e.g. emission factor by stoichiometry). 5: Unqualified estimate, expert judgment without data.
- **C, Completeness.** 1: All sub-categories covered > 99 %. 3: Coverage 50–85 %. 5: Data representing < 25 % of the process or heavily incomplete.
- **T, Temporal correlation.** 1: Data < 3 years old, same technology. 3: Data 10–15 years old. 5: Unknown age or data > 20 years.
- **G, Geographical correlation.** 1: Same exact site. 3: Similar but different zone (e.g. EU for a France study). 5: Unknown zone or very different characteristics.
- **F, Further technological correlation.** 1: Verified and representative technology. 3: Comparable technology with known differences. 5: Little-known technology or no documented precedent.

A common rater pitfall: scoring $T = 1$ for recentecoinvent data from a process not representative of the study. The framework does not police rater errors directly, but the pedigree_uncertainty Monte-Carlo of Section 4 makes the cost of such errors explicit per dimension.

Appendix D, Canonical JSON-LD I/O schema (excerpt)

The reference engine consumes and produces JSON-LD documents with a fixed canonical schema. An excerpt:

```
{
  "@context": "https://theopiens.com/chaoticlca/v5",
  "engineVersion": "v5",
  "engineConfig": {
    "sigmaRef": 0.58,
    "scalePos": 0.8,
    "scaleNeg": 0.0,
    "trtDepth": 4,
    "rdsScale": 0.05,
    "ensembleSize": 50
  },
  "input": { "R": 4, "C": 3, "T": 4, "G": 3, "F": 4 },
  "analysis": {
    "dqr": 0.30,
    "quality": "Faible",
    "sigma_ilcd": 1.66,
    "r_mapped": 3.66,
    "lambda": 0.31,
    "hKS": 0.31,
    "dKY": 1.62,
    "regime": "Chaotique (fenêtres stables)",
    "crs": 0.42,
    "crs_dynamic": 0.32,
    "gsd_eff": 2.05,
    "interval": { "p5": 0.286, "p95": 3.499 },
    "ensemble": { "mean_gsd": 2.06, "std_gsd": 0.07, "spread": 0.066 }
  }
}
```

Appendix E, Glossary of dynamical-systems terms used in the paper

For the LCA reader meeting these terms for the first time:

- **Logistic map.** The recursion $x_{n+1} = r \cdot x_n \cdot (1 - x_n)$. The simplest deterministic system that produces, for varying r , the full menagerie of fixed-point, periodic and chaotic behaviours.
- **Period-doubling cascade.** A sequence of bifurcations at which the period of a stable cycle doubles ($1 \rightarrow 2 \rightarrow 4 \rightarrow 8 \rightarrow \dots$) as the control parameter r increases. Accumulation occurs at r_∞ , beyond which the system is chaotic.
- **Feigenbaum constant δ .** The universal ratio $(r_n - r_{n-1}) / (r_{n+1} - r_n) \rightarrow \delta \approx 4.6692\dots$ as $n \rightarrow \infty$. Universality means the constant is shared across all one-dimensional smooth unimodal maps.
- **Lyapunov exponent λ .** The average exponential rate at which two infinitesimally close initial conditions diverge. Negative $\Rightarrow$ contraction; zero $\Rightarrow$ marginal; positive $\Rightarrow$ chaos.
- **Kolmogorov–Sinai entropy h_{KS} .** For a 1D dissipative system, $h_{KS} = \max(0, \lambda)$. The rate at which the system creates uncertainty per iteration.
- **Kaplan–Yorke dimension d_{KY} .** A formula (Eq. above) for the fractal dimension of a chaotic attractor. Smoothly interpolates between integer dimensions through the bifurcation cascade.
- **Random Dynamical System (RDS).** A dynamical system perturbed by an exogenous noise process. We use it as a σ_{ILCD} -bounded probe of trajectory robustness.
- **Continuous Ranked Probability Score (CRPS).** A strictly proper scoring rule (Gneiting and Raftery, 2007) measuring how well a probabilistic forecast F matches an observation x . Lower is better.
- **Sobol index (first-order).** The share of the output variance attributable to a single input, marginalising over the others. Used here on pedigree-dimension shifts.

Appendix F, Note for the LCA practitioner: how to read a CRS in three sentences

If CRS > 0.7, your ILCD intervals are valid as-is, proceed. If CRS is between 0.4 and 0.7, flag the flow, run a sensitivity analysis, document the regime in the LCA report. If CRS < 0.4, treat the flow as a candidate for primary-data acquisition; the ILCD interval may still be a reasonable working assumption while you do so, since the V1 benchmarks (§ 5.6, § 5.7) do not show that GSD_eff improves on σ_{ILCD} on single-flow CRPS.

These three sentences are, deliberately, the entire decision rule the framework asks of an LCA reader. The mathematics of period-doubling cascades is documented for the reviewer who wants it; it is not a prerequisite for acting on the diagnostic.

CRS is not a substitute for qualitative data quality review. The score is a regime label, computed deterministically from five pedigree integers; it cannot detect a flow whose pedigree is well-scored but whose underlying measurement protocol is biased, nor a flow whose context-of-use makes the LCI category inappropriate. The diagnostic complements, never replaces, the LCA practitioner's qualitative review of the data. Reading a low CRS as "the analyst can stop thinking about this flow because the score has spoken" is a misuse the framework cannot prevent; reading it as "this flow deserves a closer look before I sign the report" is the operational intent.

Operational threshold revision (V6 update). The originally-proposed $\text{CRS} = 0.5$ boundary between the "stable" and "transitional" bands is replaced by **$\text{CRS} = 0.4$** in the JRC-track addendum (Section 7.4bis), with the upper band correspondingly redefined as $\text{CRS} > 0.6$. The three-sentence rule above continues to hold as a gentle reading; the strict operational thresholds for reporting are 0.4 (action) and 0.6 (flag).

Appendix G, Roadmap V5 → V6 → V7

The framework evolves under explicit version numbers. Each version preserves backward compatibility through the `engineVersion` JSON-LD field; older outputs remain reproducible.

V5 (current paper). 5-component CRS (h_{KS} , d_{KY} , $|\lambda|$, H_{spec} , ensemble spread). Single calibration knob $\sigma_{\text{ref}} = 0.58$. RDS noise bounded by σ_{ILCD} . TRT depth 4. Clayton copula (parameter $\theta = 2$) (single-flow analyses reduce it to marginals). Single-flow benchmark only.

V6, Nexus (companion paper, Piens, in prep.). Documented and implemented; not the subject of the present paper. Five substantive additions:

1. **CRS expanded to 8 components**, adds Temporal Recurrence Coefficient (TRC), permutation entropy à la Bandt–Pompe (PE), Gelman–Rubin convergence diagnostic $\hat{R}$, and a phase-coherence statistic. The 5-component V5 score is preserved as `crs_dynamic`; the 8-component V6 score is exposed as `crs_v6`.
2. **Adaptive Gaussian RDS** with amplitude tracking realised trajectory variance.
3. **Gelman–Rubin diagnostic** on four parallel chains. $\hat{R} \leq 1.1$ on 96 % of the 128 ecoinvent flows; the 4 % residual is flagged as non-stationary (Section 8 L8).

4. **Transfer-entropy network layer** $\tau(i \rightarrow j)$ (Schreiber on Bandt–Pompe ordinal histogram), the conceptually central addition for the multi-flow regime (Section 8 L6, L9).
5. **Stochastic LCIA** (Monte-Carlo 500 draws on the impact characterisation step).

V6 is exploratory at the network and LCIA layers, in the precise sense that no quantitative network-level σ_{obs} ground truth exists at scale yet. Its evaluation is the subject of the companion paper.

V7, projected. Two additions on the V7 roadmap, neither implemented:

1. **Bayesian calibration**, replace the deterministic single-knob σ_{ref} by a posterior over $(\sigma_{\text{ref}}, R_{\text{min}}, R_{\text{max}})$ conditioned on the available primary-data corpus. Priors are weakly-informative; posterior summaries are reported per organisation. The per-organisation re-calibration of Section 7.2 becomes a Bayesian update rather than a re-fit.
2. **Network-scale benchmark**, extend the empirical evaluation from per-flow σ_{obs} to network-level statistics on coupled inventory chains. The natural target is multi-flow waste-treatment LCAs (anaerobic digestion + landfill + incineration co-systems) where the Leontief solver actively couples the flows whose pedigrees are jointly uncertain.

The $V6 \rightarrow V7$ progression is the natural deepening of the paper's core posture: a complementary, auditable, version-controlled diagnostic that reports its own statistical limits and grows in proportion to the empirical evidence available.

Appendix H, Vocabulary harmonisation between paper and public documentation

For internal consistency between the present manuscript, the engine-py reference implementation, and the V5/V6 documentation pages of theopiens.com/docs:

Paper terminology	engine-py field	V5/V6 docs	Notes
Chaotic Risk Score (CRS), 5-component	<code>crs_dynamic</code>	"CRS dynamique" / "CRS V5"	this paper's headline diagnostic
Chaotic Risk Score, 4-component baseline	<code>crs</code>	"CRS base (Wde)"	not exploited here; legacy weights 0.55/0.22/0.15/0.08
8-component Chaotic Risk Score (V6)	<code>crs_v6</code> (planned)	"CRS 8 composantes"	companion paper
Effective GSD	<code>gsd_eff</code>	<code>GSD_eff</code>	identical notation
σ_{ref} calibration knob	<code>engineConfig.sigmaRef</code>	" σ_{ref} " / " <code>sigmaRef</code> "	identical notation
Regime label	<code>regime</code>	"Régime"	French labels preserved verbatim (Point fixe / Période-2 / ...)
Pedigree uncertainty MC	<code>pedigree_uncertainty()</code>	"Sobol Sensitivity" (V5.1)	priorité 3 in the engine roadmap
Operational decision threshold	0.4 (V6), 0.5 (V5 legacy)	both documented	this paper recommends 0.4

Length: approx. 19,000 words excluding references. Figures rendered from the reference benchmark (paper/figures/) and from the engine-py implementation directly.

Translation note: this manuscript was drafted in French and translated into English by the author with the assistance of DeepSeek V4 Pro. The author retains full responsibility for the final wording, including any translation infelicities.

Submitted to Academia.edu as preprint v1.0; archival mirror at HAL-SHS pending. Comments welcome at contact@theopiens.com.