

UNDETECTABLE BUT UNREMOVABLE: CAPACITY VERSUS TOPOLOGY IN LOSSY NETWORKS

JUAN CARLOS PAREDES

ABSTRACT. A network of lossy translations — partial isometries that destroy information as they carry it — can host a topological charge: a winding in the surviving subspace whose local signature shrinks with system size while the global cost of removing it does not. A companion paper established this plateau empirically and conjectured its mechanism; a separate validation exercise established, on real composed pipelines, the operational detection threshold of a noisy, quantized estimation chain. This paper puts the two results in one inequality. The local evidence of a charge k on an $m \times m$ torus is a per-plaquette defect of size $\approx 2\pi k/m^2$, which falls below the operational detection threshold $\delta^* = \max(1.405 \cdot \sigma \sqrt{\ell}, 0.6 \cdot B/\sqrt{n})$ at a computable crossover size m^* ; beyond m^* , no local monitor at that instrumentation can see the charge. The cost of removing the charge, by contrast, is bounded below by a constant $c(k)$ independent of m , whose height is now codimension-resolved (companion three-gaps note, §4a): $\sqrt{2}$ where the surviving subspace has codimension one (measured, converging to $\sqrt{2}$ from below), exactly 1 at codimension ≥ 2 (proved for $r = 2$) — either way dozens of times above realistic thresholds. Past the crossover the network carries a *stealth charge*: undetectable by every local test, unremovable by any local edit. We give the precise statements with their evidence status — the existence floor is a theorem at every fixed size (Stage A; elementary proof, Lean formalization queued); the m -uniform constant is the companion paper’s conjecture (Stage B) — derive the crossover, work a numerical example, and map what falsifies what. The claim’s proof shape — one topological quantity evaluated two ways, locally as a sum of vanishing defects and globally as a protected obstruction — is the shape of the geometric Szpiro inequality, and we close with that dictionary and its one honest discrepancy.

1. THREE SILOS, ONE QUESTION

Three literatures hold the pieces of this paper, and do not currently speak.

Stability of almost-representations (mathematics): when is an almost-satisfying map close to an exactly satisfying one? The modern landscape is sharp on the (norm $\times$ defect-functional) grid — Frobenius-pointwise stability from cohomology vanishing [DCGLT], operator-pointwise *instability* through index obstructions [BLSW; Dadarlat; Voiculescu; Exel–Loring], uniform stability in all submultiplicative norms [GLMR]. Everything assumes invertible targets.

Topological protection (physics): an integer invariant — a Chern number, a Bott index — cannot change under perturbations small in operator norm, so locally almost-trivial data can be globally far from trivial [Voiculescu; Exel–Loring; Loring–Hastings]. The invariants live on unitaries or invertibles.

Detection theory (engineering): an estimation chain of depth ℓ with per-stage noise σ , quantizer bandwidth B , and n samples cannot resolve effects below an operational threshold; classical ingredients (d' index, non-regenerative cascade, resolution floor) [detection-theory texts; GUM] compose into the two-regime law validated — by pre-registered falsification of its product-form alternative — in our pipeline study: $\delta^* \approx \max(1.405 \cdot \sigma \sqrt{\ell}, 0.6 \cdot B/\sqrt{n})$. **[Evidence]**

Date: July 5, 2026.

The question none of the three asks: **when is a topological charge in a network of lossy links observable at all** — and what happens in the gap between “too small to see” and “too protected to remove”? The companion paper (Paper 2) moved the stability question to lossy targets, where translations cannot be undone, and found empirically that the index obstruction survives rank loss. This paper assumes that result at its evidence level, adds the detection side, and states what the junction implies. We call the underlying quantity *semigroup holonomy*: holonomy carried by composable but non-invertible maps.

2. THE OBJECTS (IMPORTED, WITH THEIR TAGS)

We work with the companion paper’s torus systems; full definitions there and in the Lean formalization (`LossyCocycles/Defs.lean`).

- **Lossy cocycle.** Each oriented edge of the $m \times m$ torus carries $T_e \in \mathbb{R}^{d \times d}$ of rank $\leq r$. *Flat* (exactly consistent) families are exactly $T_e = V_j V_i^T$ for orthonormal frames V ; every cycle composite is then a projection. [Theorem]
- **The flux family.** A Landau-gauge connection with total charge k carried in the $SO(2)$ core of the surviving r -subspace; per-plaquette angle $2\pi k/m^2$. Its local defect is $\delta_O \approx 2\pi k/m^2 \rightarrow 0$; this is elementary and measured (δ_O tracked $2\pi/m^2$ exactly over $m = 8 \rightarrow 16$, extended to $m = 32$ in `e_radar_discriminator.py`: $0.393 \rightarrow 0.006$).
- **Distance to flat.** $\text{distOp} = \inf$ over flat families of the worst-edge operator distance. Fits give upper bounds; four guards (noise-floor control, invertible-anchor reproduction, planted-class non-absorbability, Frobenius-decay of the same solutions) plus an oracle-initialized kill test protect the lower-bound-shaped readings. [Evidence]

The plateau. For $k = 1$ and $r \in \{2, 3, 4\}$: δ_O falls $7\times$ as m : $3 \rightarrow 8$ while ε_O stays pinned at $1.29\text{--}1.66$, and stays flat to $m = 16$ ($1.659 \rightarrow 1.682 \rightarrow 1.655$ at $r = 2$; ratio 1.00 across a $4\times$ sweep; pre-registered call UNIFORM); the same solutions are Frobenius-close. The obstruction is irreducibly 2-dimensional: a per-column argument is capped at γ/m by contraction telescoping, which at $m = 16$ would permit $\varepsilon_O \approx 0.4$ against a measured 1.655 . [Evidence]

3. THE TWO EVALUATIONS, AND THE STAGE A / STAGE B ARCHITECTURE

The proof-shape of everything below is: **evaluate one topological quantity two ways.**

Local evaluation. The total winding $2\pi k$ is the sum of m^2 plaquette angles, each $2\pi k/m^2$. Locally the charge is everywhere and almost nothing: any single measurement sees at most a $2\pi k/m^2$ -sized deviation from flatness.

Global evaluation. The same winding, read as an obstruction class, forces a lower bound on the distance to flat. This splits into the two stages adopted in this program (after external structural feedback):

Stage A (existence, fixed m) — THEOREM, machine-checked (`stageA_skew_bound`, axiom-clean). For $m > 2|k|$, every flat family is at worst-edge operator distance $\geq |\sin(2\pi k/m)|/m \approx 2\pi|k|/m^2$ from the charge- k flux family. *Mechanism*: flat column composites are projections, whose core skew is exactly 0; the flux column composite has core skew $2\sin(2\pi kx/m)$; skew is entry-dominated by the operator norm and spreads at most m -fold along a column. No winding apparatus needed.

Stage B (uniformity) — CONJECTURE (the companion paper’s). $\text{distOp} \geq c(k) > 0$ independent of m and $r \geq 2$; empirically $c(1) \approx 1.5$ (plateau $1.29\text{--}1.66$; kill-test floor 1.49) — a **codimension-1 reading**: the three-gaps

note (§4a) resolves the height as $\sqrt{2}$ at codim 1 (measured, monotone up to $\sqrt{2}$) and exactly 1 at codim ≥ 2 (proved for $r = 2$; the binding competitor there is orthogonal, and fitted readings above 1 are missed-basin artifacts), so the m- and r-uniform floor the conjecture needs — and the number this paper’s inequality uses — is 1. Formal-route correction : the natural first guess — an m-uniform winding-rigidity lemma over the column-path winding — is **provably wrong**: the shear family (flux verticals, trivial horizontals) winds k within $2\pi|k|/m$ of flat with full gap (**Shear.lean**, machine-checked). The column-path winding reads no horizontal data and is a 1D shadow, not an obstruction. The viable route is the Bott index of the magnetic-translation pair, whose commutator is the *sup* of plaquette defects (edge-local, m-free) — **the pair and that bound are now built and machine-checked** (**MagPair.lean**: $[U, V] = 0$ exactly on flat families; $||[U, V]|| \leq 2|\sin(\pi k/m^2)|$ on the charge- k flux family, m-uniform, axiom-clean). What stays open is the **index of that almost-commuting pair and its Exel–Loring stability** — and a dilation analysis places that gap precisely: the index’s existence is the classical 2-D winding (r-independent, **totalFlux_flux**; dilation-trivial), so the open, non-classical content is exactly the m-and-r-uniform stability for partial isometries — Stage B itself.

The two stages are different mechanisms, and Stage A’s constant says so: it decays like $1/m^2$, exactly the size of the local evidence. **One-dimensional analysis certifies precisely what local observation can see, and no more.** The plateau — the global floor — is invisible to both. Stage A is therefore *not* evidence for Stage B and is not offered as such; it is the boundary marker of the local mechanism.

4. THE DETECTION SIDE

The validation study E3 measured what a composed estimation chain can operationally distinguish, with a pre-registered gate that its product-form alternative failed:

$$\delta^* \approx \max(1.405 \cdot \sigma \sqrt{\ell}, 0.6 \cdot B / \sqrt{n}) \text{ [Evidence]}$$

σ = per-stage noise, ℓ = chain depth, B = quantizer bandwidth, n = samples. Two classical constraints, never simultaneously operational; the deliverable was the *max*, replacing the provisional product form.

Extrapolation risk, named: E3 was validated on scalar chains. Applying δ^* to the estimation of a plaquette holonomy angle treats the monitor’s scalar readout (an angle or a residual) as the chain’s carried quantity. This is the weakest link in this paper and is tagged accordingly; the matrix-valued version of E3 is an open experiment (§6, F2).

5. THE INEQUALITY, THE CROSSOVER, AND THE STEALTH REGIME

A local monitor at a plaquette estimates a deviation of size $\approx 2\pi k/m^2$ (in operator units: $2|\sin(\pi k/m^2)|$, the vortex-core form). It sees the charge iff that exceeds δ^* .

Crossover. $2\pi k/m^2 \geq \delta^* \iff m \leq m^* := \sqrt{(2\pi k / \delta^*)}$.

The claim (the paper, in one box). On an $m \times m$ torus of lossy links with instrumentation (σ, ℓ, B, n) and charge k : - $m \leq m^*$: the charge is locally detectable, and unremovable at cost $< c(k)$. - $m > m^*$: *stealth regime* — every local signature sits below δ^* , so no per-plaquette test at this instrumentation detects the charge [Evidence]; yet flattening the network still costs $\geq c(k)$ at some edge [Conjecture], and $c(k) \gg \delta^*$ in every measured cell.

Undetectable is not removable. The boundary m^* is computable from instrument constants alone.

Worked example ($k = 1$). Take $\sigma = 0.01$, $\ell = 4$ (noise term $1.405 \cdot 0.01 \cdot 2 = 0.028$); $B = 0.1$, $n = 100$ (resolution term $0.6 \cdot 0.1 / 10 = 0.006$). Then $\delta^* = 0.028$ and $m^* = \sqrt{(2\pi/0.028)} \approx 15$. On a 16×16 torus — exactly the size where the plateau was measured flat at 1.655 — the unit charge’s local signature is $2\pi/256 \approx 0.0245 < \delta^*$: invisible to this instrumentation. Its removal cost: the retraction-proof number is the proved $\text{codim} \geq 2$ value of exactly 1 (three-gaps §4a; $r = 2$, this cell’s rank), already **thirty-five times the detection threshold** — the fitted 1.655 is a missed-basin upper-bound reading above it. A maintainer who trusts local diagnostics would certify this network consistent and then find that no small recalibration of any link — or of all links — flattens it.

What the stealth regime is not. It is not an artifact of weak instruments: improving σ , B , n raises m^* but never closes the regime ($\delta^* > 0$ for any physical chain, and the $2\pi k/m^2 \rightarrow 0$ decay always wins). It is not a finite-size effect: Stage B says the unremovability is size-uniform; only its *detectability* is size-bound. And it is not mysterious: both sides are classical (detection theory; index rigidity) — the content is that they meet with a gap of one-and-a-half to two orders of magnitude ($35\times$ at the proved floor in the worked cell), in a sector (lossy targets) whose new content is **not** the index itself but the **rank-uniform stability** of that index for non-invertible maps (Stage B).

6. FALSIFICATION MAP

if this fails	then
Stage B (a flat family with $\varepsilon_{-O} \rightarrow 0$ exists — or, sharper post-reduction, commuting contraction pairs approaching the flux pair)	$c(k)$ dies; the claim degrades gracefully to Stage A’s finite- m floor: the stealth window becomes $[m^*, M(\varepsilon)]$ — bounded, not infinite. The paper’s headline dies; its finite-size form survives as theorem.
E3’s law off scalar chains (F2 experiment: matrix-valued monitor)	δ^* needs re-measurement; m^* moves; the <i>structure</i> (some $\delta^* > 0$ exists for any physical monitor) survives on information-theoretic grounds.
$r = 1$ boundary (charge with no $SO(2)$ core shows a plateau)	the index reading of the plateau is wrong; CT-4 and §3 both die. The trigger as worded does not fire — no charge exists to plant at $r = 1$. What holds instead: the curvature-free $Z/2$ -wound family plateaus at every m , at operator distance exactly 1 (both bounds proved by hand from closed-form alignment costs; any graph, any m , $d \geq 2$; Frobenius $\sqrt{(2/m)}$ exact — see the companion $r = 1$ theorem note). The index reading stands; the prediction that needed correcting was §8(5)’s “no stealth at $r = 1$ ” — the stealth phenomenon exists there too, carried by $Z/2$, by the same orthogonality barrier as the $\text{codim} \geq 2$ constant.

if this fails	then
kill criterion of CT-4 (a winding-changing, gap-preserving path)	F1RED, machine-checked (<code>Shear.lean</code>): the column-winding proof program is dead; the conjecture needs the flux-aware Bott-pair index — graceful degradation exactly as this row promised. The flux-family evidence is untouched.

Planned experiments, value order: **F2** matrix-valued δ^* (the named weak link); **F1** the stealth-regime demonstration end-to-end (plant charge at $m > m^*$, run the *actual* E1-style local certificate battery, show null detection while the fit cost stays $\approx c(k)$) — this is one script away from existing tooling; a synthetic instance is already in hand (`DARPA-maybe/e_radar_discriminator.py`: at $m = 32$ the local defect 0.006 sits $\approx 16\times$ below a $\sigma = 0.10$ noise floor while the global flux stays pinned at 2π , certificate AUC 1.000 vs a chance-level 1-D monitor — the stealth signature, on a radar phase network; the E1-style battery remains the full F1); **F3** k-sweep of the crossover ($m^* \propto \sqrt{k}$ prediction).

7. THE SZPIRO-SHAPED CLOSE

The claim’s architecture — a topological quantity evaluated once as a sum of local contributions and once as a globally protected obstruction, with an inequality falling out of the comparison — is the architecture of the geometric Szpiro inequality [Amorós–Bogomolov–Katzarkov–Pantev; Zhang; Joshi §10], following Joshi’s presentation. The dictionary, with its honest entries:

geometric Szpiro	this paper
relation $\prod [a_j, b_j] \cdot \prod \gamma_s = 1$ in π_1	torus plaquette relations
monodromy in $SL_2(\mathbb{R})$	core holonomy in $SO(2)$, carried by partial isometries
central extension; height/translation quasimorphism	universal cover $\mathbb{R} \rightarrow SO(2)$; lifted angle, an honest homomorphism
central element evaluated two ways	winding: Σ local fluxes vs flatness’s zero
discriminant $\leq C \cdot (2g - 2 + S)$	$2\pi k/m^2$ vs δ^* and $c(k)$: the crossover inequality

The one honest discrepancy is also where the analogy bit back: Zhang’s load-bearing input is quasimorphism defect control (Milnor–Wood), genuinely hyperbolic; our core is abelian and the lift is a homomorphism — but the first candidate replacement (gap-protected lift continuity over the column-path winding) was **falsified, machine-checked** (the shear family: a winding-changing, gap-preserving deformation at vanishing cost — exactly the kill criterion this thread had pre-registered). What must replace defect control is index stability for the magnetic-translation pair, where the defect is the sup of plaquette holonomies — a pair now machine-checked (`MagPair.lean`), with $||[U, V]|| \leq 2|\sin(\mathbf{p} \cdot \mathbf{k}/m^2)|$. The shape transfers; the hard part — the index’s rank-uniform stability for partial isometries — is new and now maximally sharp: the dilation gate isolated it as the sole non-classical content, and the machine-checked reduction compressed the entire remaining distance between evidence and theorem into the one named estimate **FluxPairRigidity**, with everything on either side of it — the assembly of flat families into commuting contraction pairs, the m -free locality constant, the classical Voiculescu

core at full rank — proved. That is the correct division of labor between an analogy and a theorem, and it is where this paper stops and the mathematics continues.

8. OPEN PROBLEMS

- (1) Stage B — the m-uniform rigidity lemma (the program’s central open problem).
- (2) A matrix-valued two-regime law (F2): does δ^* compose the same way when the carried quantity is an operator block, and does the regenerative bonus survive? (3) The $c(k)$ -spectrum: is $c(k)$ monotone in k ? bounded? ($m^* \propto \sqrt{k}$ says detectability worsens with charge; does protection strengthen? Partial answer, three-gaps §4a + expM: at $\text{codim} \geq 2$ the constant is exactly 1, k -independent — protection neither strengthens nor weakens with charge there; the $\text{codim}-1$ $\sqrt{2}$ is likewise measured k -uniform. So detectability degrades with k while protection stays flat: the stealth window *widens* with charge.) (4) A converse stealth statement: can *every* sub-threshold unremovable deformation be realized by a topological charge, or are there non-topological stealth classes? (5) The $r = 1$ boundary as the bridge to the rank-one companion world. Here the naive expectation “no stealth at $r = 1$ ” is wrong: the only invariant is $\mathbb{Z}/2$ frustration, and it suffices — a curvature-free wound family is locally invisible yet plateaus at height **exactly 1**, both bounds proved by hand from closed-form per-edge alignment costs (companion $r = 1$ theorem note). The rank-one sector is thereby the first with an end-to-end-theorem stealth statement: local defect zero, sub-threshold for every instrument, unremovable below cost exactly $1 = 35\times$ the worked δ^* . Remaining there: the noisy ($\sigma > 0$) window, and the Lean formalization.

REFERENCES

Stability of almost-representations (§1). - **DCFLT/DCGLT:** M. De Chiffre, L. Glebsky, A. Lubotzky, A. Thom, *Stability, cohomology vanishing, and non-approximable groups*, Forum Math. Sigma **8** (2020), e18. arXiv:1711.10238. - **GLMR:** L. Glebsky, A. Lubotzky, N. Monod, B. Rangarajan, *Asymptotic cohomology and uniform stability for lattices in semisimple groups*, arXiv:2301.00476 (2023).

- **BSW:** *Stability and instability of lattices in semisimple groups*, J. Anal. Math. **151** (2023), 1–23. arXiv:2303.08943.
- **Dădărlat:** M. Dădărlat, *Obstructions to matricial stability of discrete groups and almost flat K -theory*, Adv. Math. **384** (2021), 107722. arXiv:2007.12655.

Almost-commuting / index obstruction (§1, §3). - **Voiculescu:** D. Voiculescu, *Asymptotically commuting finite rank unitary operators without commuting approximants*, Acta Sci. Math. (Szeged) **45** (1983), 429–431. - **Exel–Loring:** R. Exel, T. A. Loring, *Almost commuting unitary matrices*, Proc. Amer. Math. Soc. **106** (1989), 913–915; and *Invariants of almost commuting unitaries*, J. Funct. Anal. **95** (1991), 364–376. - **Loring–Hastings:** M. B. Hastings, T. A. Loring, *Topological insulators and C^* -algebras*, Ann. Physics **326** (2011), 1699–1759. arXiv:1012.1019.

Detection theory (§4). - **GUM:** JCGM 100:2008, *Evaluation of measurement data — Guide to the expression of uncertainty in measurement*, BIPM. - **detection-theory texts:** e.g. S. M. Kay, *Fundamentals of Statistical Signal Processing, Vol. II: Detection Theory*, Prentice Hall (1998). pick the specific source you used.

Geometric Szpiro (§7). *Resolved against the on-disk reading notes and independently content-verified online (`zhang-joshi-candidates.md`). An earlier resolution pass had expanded*

“Amorós et al. 2000” to the Amorós–**Burger**–Corlette–Kotschick–Toledo Kähler book (1996) — wrong work; the reading notes’ source is Amorós–**Bogomolov**–Katzarkov–Pantev, corrected here and in the §7 bracket. - Amorós–**Bogomolov**–Katzarkov–Pantev: J. Amorós, F. Bogomolov, L. Katzarkov, T. Pantev, *Symplectic Lefschetz fibrations with arbitrary fundamental groups*, J. Differential Geom. **54** (2000), 489–545. arXiv:math/9810042. (The $g = 0$ geometric Szpiro case; Zhang’s own reference [1].) - **Zhang**: S.-W. Zhang, *Geometry of algebraic points*, in: First International Congress of Chinese Mathematicians (Beijing, 1998), AMS/IP Stud. Adv. Math. **20**, AMS/International Press (2001), 185–198. (§3: the geometric Szpiro inequality via monodromy in SL_2 + the translation quasimorphism on $\tilde{SL}_2(\mathbb{R})$ — the proof §7’s dictionary transcribes, verified line-for-line.) - **Joshi §10**: K. Joshi, *Construction of Arithmetic Teichmüller Spaces II $_{\frac{1}{2}}$: Deformations of Number Fields*, arXiv:2305.10398, §10 (“Appendix: The proof of geometric Szpiro inequality revisited” — the presentation we used; reading note provenance: v12, §§10.1–10.18, display (10.15.1)). - Milnor–Wood (the quasimorphism/relative-Euler-number framing): J. Milnor (1958), J. W. Wood (1971).

FIGURES

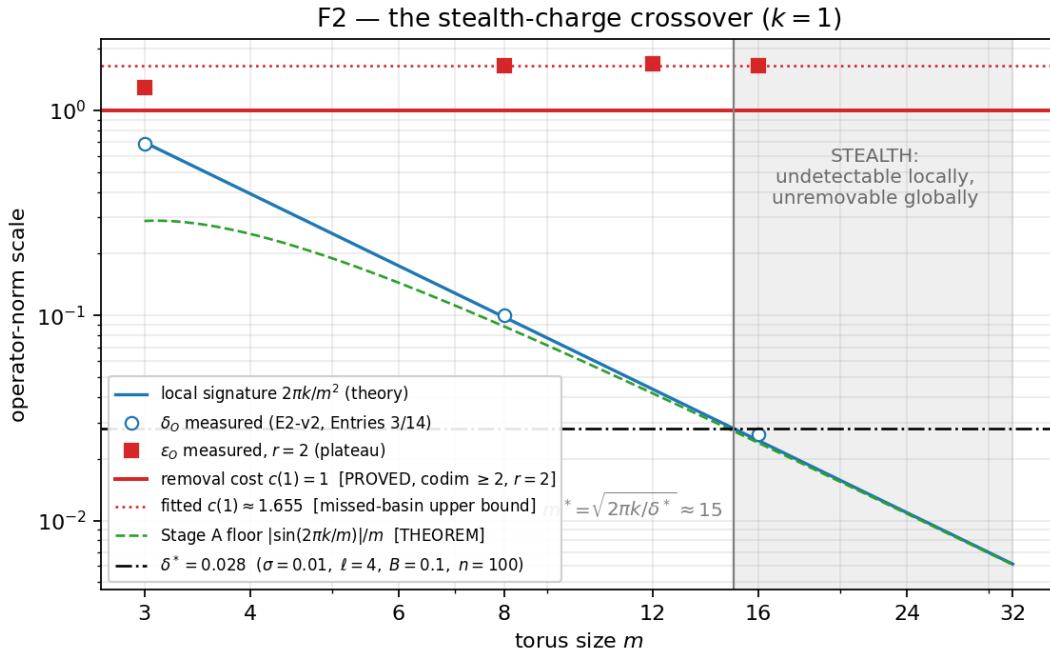


Figure 1. The crossover: local signature $2\pi k/m^2$ falls through δ^* at m^* ; the removal floor (proved, $\text{codim} \geq 2$: exactly 1) stays flat above; fitted 1.655 shown dashed as a missed-basin upper bound.

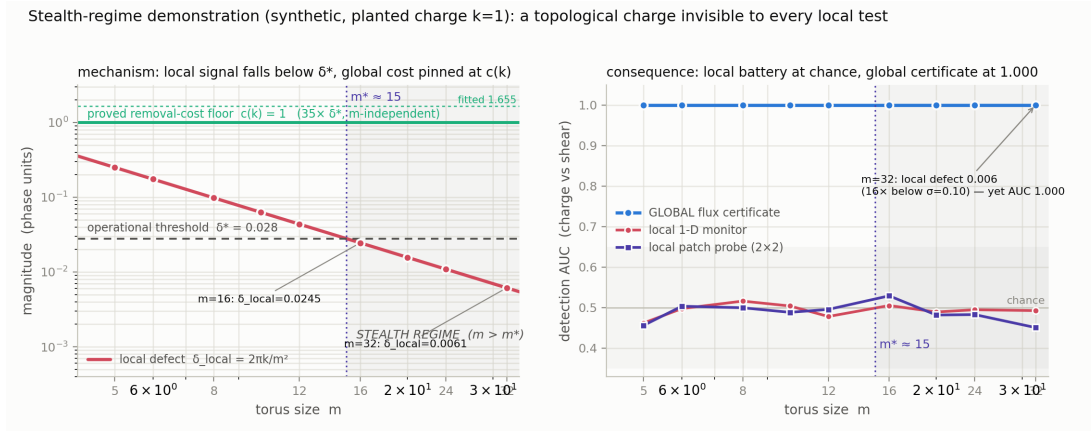


Figure 2. Stealth-regime demonstration (planted charge at $m > m^*$): local battery at chance, global fit cost pinned at the floor. Synthetic/planted experiment, per Sec. 6.

AUTHORSHIP AND AI-ASSISTANCE DISCLOSURE

This research program — its conception, its original ideas, its questions, and its direction — is the work of the author, Juan Carlos Paredes. The strategic pivots, the falsification discipline, the decisions about what to pursue, what to test, what to claim, and what not to claim are his throughout. In the course of the program he directed large language models (Anthropic Claude — Opus and Fable model families) as research assistants: implementing and running numerical experiments he commissioned, formalizing proofs in Lean 4, drafting and revising prose, and retrieving literature, all under his direction and review. Formal results are machine-checked (Lean 4) where stated; numerical claims are reproducible from the released scripts. Juan Carlos Paredes is the author of this work; the assistance is disclosed in the spirit of the full transparency this venue requires.

INDEPENDENT RESEARCHER

Email address: cpt66778811@gmail.com