

The Disclosure Deficit

What Survives the Fee Family

John Komkov, Independent Researcher

July 2026

LIFECYCLE technical-report | EVIDENCE formally-verified | RECORD /research/disclosure-deficit
BUILT FROM research-status.yaml @ 2026-07-22

Abstract. Six papers in this program advanced a scalar invariant of tool composition and the repair theory built on it. Five are withdrawn and one is archived. This report states what survived, at the strength the surviving evidence supports, and admits nothing else.

What survives is a rank identity and its matroid backbone. The disclosure deficit of a composition is the difference of two ranks of the same signed incidence structure, and equals the number of hidden fields minus the number of components those fields touch. It is a difference of ranks and a count of components; every stronger reading of it has been withdrawn. Because the incidence structure is totally unimodular, the rank results do not depend on the coefficient field.

The repair space is a product, not a spanning-tree count. The witness matroid is a direct sum of uniform matroids over carrier components, so the number of minimum disclosure sets is the product of component sizes. The published spanning-tree formula was wrong, and the reason is worth recording: it confused the graphic matroid of a graph with the column matroid of that graph's signed incidence matrix. The two have the same rank, different ground sets, and different basis counts, and they agree only when every component is a cycle.

One result is negative and is reported as such. The claimed operational sparsity of the repair space is not supported: the published measurement is consistent with a maximal-coverage null at every point, and the paper's own realizability theorem refutes the claim as quantified over its stated cost family. This is not evidence that sparsity is absent. It is the absence of evidence that it is present, and the instrument that would settle it is named.

No claim is made here about execution outcomes, safety, topology, or the necessity of cycles, and no connection is asserted to the accountability lane.

1. The scope lock

This report is bounded by a scope lock written before it, and the lock is part of the result. It admits exactly five classes of input, each only after line-by-line verification against the correction records:

1. the component-count and cokernel formula;
2. the column-matroid corank identity;
3. signed-incidence field independence;
4. the corrected repair-basis results;

5. null findings and falsifications.

It excludes, by name: topology branding, cycle necessity, execution prediction, universal safety, procurement, and market claims.

Two exclusions deserve their reasons stated rather than assumed. Execution prediction is excluded because a pre-registered, execution-labelled test of the predictor returned a likelihood ratio of approximately 1.07 over 759 cases — the quantity is a disclosure metric, not a detector, and the distinction is the difference between a measurement and a claim about the world. Topology branding is excluded because the surviving statement is a difference of two ranks; naming it after a cohomology group imports connotations the evidence does not carry, and this report accordingly writes the identity in its component and difference form throughout.

No connection to the compositional-accountability lane is asserted. The corank-reduction audit of 2026-07-14 records that reduction as *not proved* and forbids the reuse: admissibility there depends on adversary and survival schedules, while the corank here is computed over a fixed linear visibility structure. They are not presently the same object.

2. What the deficit is

Let a composition expose observable fields and hide others across a set of dimensions, and let δ be the signed incidence structure of the carrier graph. The disclosure deficit is the difference of two ranks:

$$\text{fee}(G) = \text{rank } \delta_{\text{full}} - \text{rank } \delta_{\text{obs}} = |H| - \#\{\text{components touched by hidden fields}\}.$$

Equivalently, under the disjointness and homogeneity hypotheses of the surviving construction, with carrier components $C_{d,j}$:

$$\text{fee}(G) = \sum_{d,j} (|C_{d,j}| - 1).$$

That is the whole statement. It is a difference of ranks, and a count of components. The rank identity itself is machine-checked, and it was verified on all 703 corpus compositions in the source study without exception.

Because the signed incidence matrix is totally unimodular, the rank results hold independently of the coefficient field. This is the one genuinely structural fact the family produced: the deficit is not an artifact of working over the rationals.

3. The repair space is a product

The second surviving result concerns how many ways a deficit can be repaired.

Under the stated hypotheses the witness matroid decomposes as a direct sum of uniform matroids over carrier components,

$$M/O \cong \bigoplus_{d,j} U(|C_{d,j}| - 1, |C_{d,j}|),$$

so the number of minimum disclosure sets is the product of component sizes, $\beta(G) = \prod_{d,j} |C_{d,j}|$, with sharp bounds $\text{fee} + 1 \leq \beta \leq 2^{\text{fee}}$, both attained. The deficit is a sum over components; the multiplicity is a product over them. Fixing the deficit does not fix the multiplicity.

3.1 The error that produced the wrong count

The published formula gave β as a product of spanning-tree counts. It is wrong, and the mechanism is a clean one to record because it is easy to repeat.

There are two different matroids naturally attached to the same graph. The **graphic matroid** has ground set the *edges*; its rank is $n - c$ and its bases are spanning forests, counted by the matrix-tree theorem. The **column matroid of the signed incidence matrix** has ground set the *vertices*; its rank is also $n - c$, but its bases are vertex subsets meeting every component, counted by a product of component sizes. Equal rank, different ground set, different bases.

The witness matroid is the second, because its ground set is hidden field columns. The smallest witness: for K_2 the signed incidence matrix is the 1×2 matrix $[-1, +1]$, whose column matroid has rank 1 and **two** bases, not one. Two such dimensions give four minimum disclosure sets where the published formula predicted one.

The two counts coincide exactly when every component is a cycle, since $\tau(C_n) = n$. That coincidence is why the error survived contact with several worked examples before a K_2 block exposed it.

The correction simplifies the theory rather than complicating it, and it leaves the rank, Kron-reduction, leverage, and effective-resistance results untouched — none of them counts bases.

4. A null result, reported as a null result

The withdrawn paper claimed that repair is *operationally sparse*: that under integer costs in $\{1, \dots, 10\}$, at $\beta = 378$ only 13.5% of bases are ever optimal. The claim does not survive, for two independent reasons.

The measurement cannot distinguish the hypothesis from its negation. The experiment drew 54 cost vectors, each retaining exactly one argmin basis, so it can never observe more

than 54 distinct optimal bases. At $\beta = 378$ the maximum measurable ratio was $54/378 = 14.3\%$; the reported figure was 13.5%. Against the maximal-coverage null $\mathbb{E}[D] = \beta(1 - (1 - 1/\beta)^{54})$, every one of twenty tested rows falls inside the central 95% interval, including the headline row: observed 51 against a null expectation of 50.38, a ratio of 1.01.

The paper’s own theorem refutes the claim as quantified. The realizing cost vector in its proof of full realizability sets $c(h) = 1$ on the basis and $c(h) = 2$ off it, using only values in $\{1, 2\} \subset \{1, \dots, 10\}$. Over the family of integer cost vectors in $\{1, \dots, 10\}$, then, every basis is the unique optimum of some member. The 13.5% was a property of 54 sampled vectors, not of the cost family.

The epistemic status is stated precisely, because the two available overstatements point in opposite directions. This is **no evidence for** operational sparsity. It is **not** evidence that sparsity is absent: 54 draws cannot separate “13.5% reachable” from “100% reachable” at $\beta = 378$. The instrument that would settle it is the one the withdrawn paper itself named — per-basis optimality-cone feasibility over the cost box, via the normal fan of the matroid basis polytope. Until that is computed, the question is open and is recorded as open.

The null is reproducible from a seeded stdlib script; exit code 0 means no evidence of sparsity.

5. What does not return

Five withdrawn papers and one archived paper stand behind this consolidation. Their retracted claims are not reinstated here in any form:

- the execution-failure interpretation of the deficit, falsified by a pre-registered execution-labelled test;
- the hierarchical-fee corollary and its repair prescription;
- the spanning-tree basis count and the “unique disclosure set” reading that followed from it;
- the operational-sparsity claim, per Section 4;
- the type-soundness theorem of the refinement calculus, including under its proposed static-fragment repair;
- the spectral-gap trichotomy and the linear communication bottleneck.

Each retains its own correction record, and each served PDF now opens on a publication-status notice. This report supersedes none of them; it collects what remained after each was corrected.

6. What this does not establish

The surviving results are conditional on the stated construction and its disjointness and homogeneity hypotheses. Within that scope:

- the deficit is a measurement of hidden convention obligation under a declared visibility structure — not a prediction, a score, or a safety property;
- the basis count is a count of minimum disclosure sets, not a claim that any particular one should be chosen;
- the sparsity question is open, with a named instrument;
- nothing here bears on occurrence, authority, remedy, or recourse.

The falsifier for this report is narrow by construction: any included result that turns out to depend on cycle necessity, topology branding, execution prediction, or a false repair-basis statement must be removed from it.