

Can different parts of the same societal function need to be governed at different levels?

A blind test of the Governance Factorization Simulator v0.5

Björn Kenneth Holmström • September 2026
svensksubsidiaritet.se

Can different parts of the same societal function need to be governed at different levels?

A blind test of the Governance Factorization Simulator v0.5

Summary

The debate about centralisation and decentralisation often treats a societal function as if it had a **single natural level of governance**. But one and the same activity may simultaneously require local knowledge, regional specialist expertise, broad risk pooling, shared infrastructure, and sometimes binding decisions at the system level.

The Governance Factorization Simulator was developed to investigate this possibility. Instead of asking “*at what level should the activity be placed?*”, the model decomposes governance into six roles:

- **A — authority:** where final binding decision-making power needs to lie,
- **P — production:** at what scale the actual operation or delivery should be organised,
- **X — expertise:** at what scale specialist competence needs to be concentrated,
- **F — finance/risk pooling:** at what scale financial risk and capacity should be shared,
- **K_P — policy coordination:** at what scale policy needs to be harmonised,
- **K_I — interface coordination:** at what scale common interfaces and infrastructures need to be coordinated.

Version 0.5 was tested on eight entirely new Swedish societal functions. The model was frozen before the test. Cases were selected mechanically from a pre-declared candidate pool. All ten problem inputs — including two new variables for **binding integration needs**, (B_{14}) and (B_{420}) — were locked and hashed before any investigation was made of how the activities are actually organised. The model’s predictions were also locked before outcome research.

Compared with the earlier frozen v0.4 model, v0.5 improved on four central metrics:

Metric	v0.4	v0.5	Change
Average component hit rate	45.3%	50.5%	+5.2 percentage points
Probability of fully compatible architecture	10.4%	28.8%	+18.4 pp
Architecture distance	0.292	0.252	-0.040
Authority hit rate	60.4%	78.0%	+17.6 pp

At the same time, the model’s average bias towards broader scales increased slightly, and four locally organised functions still received zero probability of a fully compatible architecture.

The result should therefore not be described as the model being “validated”. The more defensible conclusion is:

v0.5 gave a clear improvement relative to the frozen predecessor in a genuinely new blind test, while important and recurring model errors remain.

The most interesting support concerns the new distinction between **coordination** and **binding authority**. The two cases that were blind-coded with the strongest need for shared binding decision-making power — operation of the transmission grid and air traffic control — also became the two cases where v0.5 most clearly corrected v0.4’s previous underestimation of authority scale.

1. The problem with the question “what level should govern?”

The principle of subsidiarity is often expressed roughly as follows:

Decisions should be made as close to those affected as possible, but at a higher level when necessary.

The principle is intuitively appealing but leaves a crucial question open:

What exactly is it that should be placed at a certain level?

Consider a hypothetical societal function. It may need to:

- be executed close to residents,
- use specialist expertise that only exists in a few places,
- be financed over a larger population base,
- follow common technical standards,
- and at the same time leave final decision-making authority locally.

If all these aspects are summarised as a single variable — *municipal, regional or national* — information is lost.

That leads to a different starting point:

[\boxed{\text{the level of governance may not be a property of the activity as a whole}}]

but rather:

[\boxed{\text{a property of different governance functions within the activity.}}]

The Governance Factorization Simulator is an attempt to make this hypothesis explicit and testable.

2. A six-part governance architecture

The model describes an architecture as:

[$G=(A,P,X,F,K_P,K_I).$]

Each role can, in the current model, take one of three stylised scales:

[{1,4,20}.]

They should be read as:

- **1 — local**

- 4 — **intermediate**
- 20 — **broad/system level**

The numbers are mathematical scale markers. They do not literally mean municipality, region and state. A regional Swedish organisation can, for example, be represented by the model's broad scale 20 if it functionally corresponds to the broadest relevant level in the analysed system.

2.1 Authority — (A)

Authority refers to **final binding decision-making power**.

It is not the same as coordinating, advising, financing or executing the activity. The question is whether several separate actors can continue to have the last word individually, or whether certain conflicts and trade-offs require a common decision point.

2.2 Production — (P)

Production refers to where operational delivery or production should be organised.

An activity may, for example, be financed and standardised broadly but still be executed locally.

2.3 Expertise — (X)

Expertise represents the scale required to sustain sufficient specialist competence.

Small units may lack the volume to support highly specialised professions, laboratories or technical resources.

2.4 Finance/risk pooling — (F)

Finance here refers primarily to **risk pooling and capacity sharing**, not formal budgetary power.

This is an important distinction. A municipality may have budgetary responsibility while certain risks need to be borne across a much larger collective.

2.5 Policy coordination — (K_P)

Policy coordination describes the need for common rules, priorities or harmonised policies.

It does not automatically imply that binding decision-making authority must be moved upwards.

2.6 Interface coordination — (K_I)

Interface coordination describes the need for common systems, standards, information flows or infrastructural interfaces.

There may, for example, be a national payment or information system while operational activity remains local.

3. What v0.5 changed

v0.5 introduced two main changes following earlier blind tests.

3.1 Binding integration as its own problem signal

Earlier models tended to keep authority too local in certain functions that in reality had broad binding governance.

But “increasing centralisation pressure” in general would have been a poor solution. Coordination and binding authority are not the same thing.

v0.5 therefore introduces:

[B_{14}]

and:

[B_{420},]

which represent two steps of **binding-integration pressure**.

(B_{14}) asks roughly:

How much functional value is created by moving from separate local final decisions to a common intermediate authority?

(B₄₂₀) asks:

How much of this need remains even after the intermediate level and requires a broad system level?

To prevent an artificial jump directly from local to system level, the following holds:

[0 ≤ B₄₂₀ ≤ B₁₄ ≤ 1.]

This means that a strong second step presupposes an at least equally strong first step.

3.2 Fixed institutional cost is counted per supra-local function, not per number of scales

v0.4 had a simple cost for the number of different scales used in an architecture.

The blind test showed that this could create an unwanted **co-location gravity**: if an intermediate level was already used, the model could favour moving several other functions there as well, even without a functional reason.

v0.5 replaces this with:

[M_{sup}.]

It instead counts how many functional domains actually require institutional organisation **above the binding authority level**.

The frozen cost is:

[$\mathrm{J}_{\mathrm{real}}$]

$\lambda_{\kappa_M M_{\mathrm{sup}}}$,]

with:

[κ_M]

0.028922276644969897.]

The value was derived from the model's own geometry before the new holdout and tested in a separate synthetic freeze protocol. It was not chosen by optimising fit on Swedish cases.

4. Before the blind test: analytical and synthetic control

Before v0.5 was allowed to meet new empirical cases, a full analytical/synthetic suite was run on the frozen implementation.

All **33 of 33** pre-registered hard gates passed.

Among other things, the following were verified:

- exactly 378 permitted architectures,
- the algebraic authority transitions,
- that (B_{420}) cannot create a second authority transition without support for the first,
- that policy and interface coordination can be separated,
- that local, intermediate and broad authority can all be optimal in different synthetic regimes,
- that the new institutional cost does not collapse the model into a single architecture,
- that the model reproduces v0.4 and v0.3 exactly in special compatibility modes.

In a balanced synthetic population of 100,000 cases, the realisation layer changed the selected representative architecture in roughly **23.4 percent** of cases and reduced the average number of supra-authority institutional domains from **1.77 to 1.15**.

At the same time, the model retained significant architectural diversity. This was important: an institutional cost addition that only makes everything local or everything broad would not be very informative.

The synthetic testing is, however, **not empirical validation**. It only shows that the model behaves according to its own structural intentions and does not collapse in obvious ways.

5. A genuinely new blind test

The empirical test was constructed to minimise the possibility of adjusting the model after seeing the answer key.

5.1 The model was frozen first

v0.5's:

- equations,
- constants,
- (κ_M),
- tie-breaking,
- architecture space,
- and implementation

were frozen before the holdout.

5.2 A new candidate pool was created

A pool of **16 previously unused societal functions** was created in four pre-declared stress categories.

No F-, H-, N-, G- or earlier G-pool cases were reused.

Two cases per category were selected mechanically by the lowest SHA-256 hash of the candidate's pre-declared identity.

The selection was therefore not made on the basis of which cases seemed favourable to the model.

The eight selected functions were:

1. school facility maintenance,
2. municipal graffiti removal,
3. school placement/admission,
4. parking enforcement,
5. blood culture diagnostics,
6. PET/CT diagnostics,
7. operation of the national transmission grid,
8. air traffic control in Swedish airspace.

5.3 Problem inputs were locked before outcome research

Each case was coded on ten dimensions:

[(L,S_{14},S_{420},C_{14},C_{420},E,R_{14},R_{420},B_{14},B_{420}).]

Crucially, even the two new authority inputs:

[B_{14},B_{420}]

were coded and hashed **before** we investigated the actual institutional structure.

5.4 Predictions were locked

For each function, 10,000 common-(x) realisations were generated and predictions were locked for:

- frozen v0.4,
- v0.5 functional optimum,
- v0.5 realisation optimum.

Only then was outcome research opened.

5.5 Outcome research was locked separately

The actual Swedish organisation was then investigated using primarily public sources.

Examples:

- Stockholm's school properties are managed by the municipally owned SISAB.[1]
- The City of Stockholm organises graffiti removal through, among others, the traffic administration, district administrations and the city's contractors.[2]
- The municipality is responsible for school placement in municipal compulsory schools.[3]
- Parking tickets are issued locally by the municipality or police, while the Swedish Transport Agency administers the payment chain.[4]
- Clinical microbiology at Sahlgrenska is part of Regional Laboratory Medicine in Region Västra Götaland.[5]
- PET/CT is available within Sahlgrenska University Hospital's nuclear medicine and radiology operations.[6]
- Svenska kraftnät is the system operator authority for the transmission system and is responsible for the transmission grid and real-time balancing.[7]
- Air traffic control is carried out through several control centres and several certified service providers, while the airspace requires system-wide coordination.[8][9]

The outcome scale was locked before any hit score was calculated.

6. Results

The primary comparison was frozen v0.4 against v0.5's realisation output.

Metric	frozen v0.4	v0.5 functional	v0.5 realisation	Change v0.5R – v0.4
Average component hit rate	45.3%	49.6%	50.5%	+5.2 pp
Joint-compatible probability	10.4%	25.1%	28.8%	+18.4 pp
Architecture distance	0.292	0.256	0.252	–0.040
Signed scale deviation	0.068	0.109	0.093	+0.025
Authority hit rate	60.4%	75.1%	78.0%	+17.6 pp

Metric	frozen v0.4	v0.5 functional	v0.5 realisation	Change v0.5R – v0.4
Modal full architecture compatible	1/8	2/8	2/8	+1

6.1 What the metrics mean

Component hit rate asks how often the six individual roles land on a scale judged compatible with the observed organisation.

Joint-compatible probability is stricter: how often is the entire six-role vector simultaneously compatible?

Architecture distance measures how far the predicted six-role vector is from allowed outcome scales.

Signed scale deviation shows the direction of errors. A positive value means the model on average predicts somewhat broader scales than the outcome.

The result is therefore not unambiguously positive.

v0.5 comes **closer** to the observed architecture in absolute terms and greatly increases the probability of full compatibility, but at the same time acquires a slightly larger average bias towards broader scales.

7. What the eight cases show

J01 — School facility maintenance

Blind-coded binding pressure: (B=(0,0))

The outcome was mainly local/municipal.

v0.5 nevertheless continued to most often choose:

[(1,4,4,4,4,4).]

The component hit rate was only about **20.5 percent**, and joint compatibility was zero.

This is fresh evidence of a known model problem:

moderate scale advantages in production, expertise and finance can still overcome locality for activities that are in practice sharply local.

J02 — Graffiti removal

Blind-coded binding pressure: $(B=(0,0))$

This case is also mainly municipal.

v0.5 improved the component hit rate from about **20.5 to 37.6 percent**, suggesting that the new realisation cost actually dampens unnecessary supra-local organisation.

But the modal architecture remained:

[(1,4,4,4,4,4)]

and joint compatibility was still zero.

The new institutional geometry thus helps, but does not solve the deeper locality problem.

J03 — School placement

Blind-coded binding pressure: $(B=(0.25,0))$

The municipality has clear responsibility for school placement in municipal schools.

The model nevertheless predicted strongly intermediate P/X/F/K roles. The component hit rate was about **15.8 percent**.

This shows that the sharp-local residual does not only apply to physical operations. It also occurs in local administration.

J04 — Parking enforcement

Blind-coded binding pressure: $(B=(0.25,0))$

This turned out to be a particularly interesting factorised outcome.

The parking ticket itself is issued locally by the municipality or police, while the Swedish Transport Agency administers the national payment chain.

This roughly corresponds to:

local production/authority but broad interface coordination.

v0.5 missed this structure and was again drawn towards:

[(1,4,4,4,4,4).]

This is a clear remaining problem for the model:

[$\boxed{\text{broad interface coordination need not imply broad production.}}$]

J05 — Blood culture diagnostics

Blind-coded binding pressure: $(B=(0.25,0))$

Clinical microbiology is regionally organised within Region Västra Götaland's laboratory medicine.

v0.5's modal became:

[(4,4,4,4,4,4).]

The model thus recognises the need for pooling and larger scale, but often stops at the intermediate level.

Joint compatibility was about **30.2 percent**, compared with **19.0 percent** for v0.4.

J06 — PET/CT

Blind-coded binding pressure: (B=(0.25,0))

PET/CT is a highly specialised activity within the regional hospital structure.

v0.5 almost always placed:

- production,
- expertise,
- finance,
- policy coordination,
- interface coordination

on a broad scale.

But authority remained local:

[(1,20,20,20,20,20).]

Five of six roles were therefore essentially correct, while the authority hit rate was **0 percent**.

This is important because we **must not correct the model by retrospectively raising (B)**. (B=(0.25,0)) was already locked before outcome research.

The case therefore suggests another possible mechanism:

an activity can acquire its binding authority from a broader institutional container — for example a regional healthcare system — even if the task itself does not have strong functional binding pressure.

This is not explicitly present in v0.5.

J07 — Operation of the transmission grid

Blind-coded binding pressure:

[$B=(1,1)$.]

Svenska kraftnät has system-wide responsibility for the transmission grid and the balancing and operational security of the power system.

This became v0.5's clearest success.

v0.5 gave:

[(20,20,20,20,20,20)]

with **100 percent component compatibility and 100 percent joint compatibility**.

v0.4 had already understood that many other roles needed to be broad, but still had significant probability of too local authority.

The new (B)-system corrected precisely this.

J08 — Air traffic control

Blind-coded binding pressure:

[$B=(1,0.75)$.]

Air traffic control requires system-wide airspace coordination and binding operational decisions, even though production occurs through several control centres and several certified providers.

v0.4's authority hit rate was only about **0.9 percent**.

v0.5 achieved:

- **100 percent authority compatibility**
- **100 percent joint compatibility**

in the outcome set.

This is particularly important methodologically because the high binding signal was coded before outcome research.

8. The most important new evidence: coordination is not authority

The most interesting result from v0.5 is not the total five-percentage-point increase in component hit rate.

It is that a previous conflation appears to have been better identified.

A system may need:

- shared information,
- shared standards,
- shared financing,
- or shared policy,

without separate actors having to give up their final decision-making power.

Conversely, there are situations where coordination is not enough.

Transmission grid and air traffic control involve hard, system-wide trade-offs where someone must be able to make a final binding decision.

The two cases received the highest blind-coded (B)-values — and they are also the two cases where v0.5 most dramatically corrected v0.4's authority errors.

With only eight cases, this is far from a final statistical result.

But it is precisely the kind of observation a blind test is designed for:

[\boxed{ \text{a new mechanism made the right kind of difference on new cases coded without access to the answer key.} }]

9. What the model still misses

The result is at least as interesting where the model fails.

9.1 Sharp-local false factorization

J01–J04 all had:

[$\text{joint compatibility}=0$.]

The model still overvalues moderate scale advantages for certain activities that in practice remain local.

This may mean that the model lacks some form of:

- locality benefit,
- diseconomies of scale,
- spatial execution constraint,
- or task decomposition.

It would, however, be methodologically wrong to choose a mechanism by fitting it against just J01–J04.

9.2 Authority can follow the institution rather than the task

J06 shows that broad authority does not always need to be explained by strong binding pressure in the analysed function itself.

PET/CT lies within a broader regional healthcare system.

This opens a new distinction:

[$\boxed{\text{task-required authority} \neq \text{institutionally inherited authority}}$.]

A future model may need to separate these.

9.3 Broad interface with local production

J04 shows another pattern:

[$K_I \gg P$.]

National registers, payment systems or information interfaces can exist on top of highly local activities.

The current interface layer can represent this in the architecture space, but the model does not yet produce this structure sufficiently naturally.

9.4 Broad authority with lower production is structurally difficult

An earlier algebraic review of v0.5 showed that the current functional core in practice makes:

[$P < A, \quad X < A, \quad F < A$]

dominated as strict optima.

This means the model has difficulty describing:

broad binding authority with genuinely decentralised operational production.

This is a **known validity limit**, not something to be concealed in reporting.

10. What this says about subsidiarity

The result does not support a simple thesis that Sweden should centralise or decentralise.

Rather, it points towards the question itself often being misposed.

A governance architecture may, for example, be:

[$A=1, \quad P=1, \quad X=4, \quad F=20, \quad K_P=4, \quad K_I=20.$]

This could correspond to:

- local decision-making authority,
- local delivery,

- shared specialist expertise,
- broad risk pooling,
- some policy harmonisation,
- and a national technical interface.

Calling the whole construction “municipal” or “national” then says relatively little.

A more precise subsidiarity principle could therefore be formulated as:

Place each governance function at the lowest scale that can bear that function’s requirements, and link the levels together without automatically moving other functions with it.

This is a stronger principle than “as local as possible”.

It also says:

[\boxed{ \text{do not centralise an entire activity just because one part of it needs centralisation.} }]

And conversely:

[\boxed{ \text{do not keep binding authority local if the system function itself requires shared final decisions.} }]

It is this kind of **functional subsidiarity** that the simulator attempts to operationalise.

11. What the report does not show

There are several reasons for caution.

Eight cases are few

The holdout is small. Individual cases can therefore strongly affect the aggregate metrics.

Outcome coding involves judgement

Real institutions do not fit perfectly into three synthetic scales.

Outcomes are therefore sometimes specified as a set of compatible scales rather than a single exact number.

The model does not predict historical causality

The fact that a Swedish institution is at a certain level does not mean that the level is functionally optimal.

Institutions are also shaped by:

- history,
- law,
- politics,
- professional boundaries,
- budget systems,
- organisational inertia,
- and previous reforms.

The simulator attempts to estimate functional architecture, not reconstruct this entire historical process.

No success threshold was pre-registered

Before the test, no arbitrary threshold was defined at which, for example, 50 or 60 percent would be called “validation”.

Therefore such a label should not be invented afterwards.

The correct summary is:

mixed with strong paired improvement — not validation.

12. Research status after v0.5

v0.5 should now be treated as a **frozen research checkpoint**.

There are three reasons.

First, the model has undergone:

1. design review,
2. algebraic review,
3. separate derivation of (κ_M),
4. pre-registered freeze audit,
5. implementation audit,
6. full analytical/synthetic suite,
7. development face check,
8. and then an entirely new blind test.

Second, the blind test produced both successes and clear residuals.

There is therefore no methodological reason to “fix” the model immediately after seeing the results.

Third, the model’s failures can now be as valuable as its hits. They define questions for future research:

- How does one represent genuinely strong locality?
- How does one distinguish task-required authority from institutional-container authority?
- How does one model broad interfaces on top of local production?
- How does one allow broad authority with decentralised production without building it in as a desired outcome?

The next model version should therefore begin with these questions as **open problems**, not with J01–J08 as training data.

13. Conclusion

The Governance Factorization Simulator began with a simple question:

Can different parts of the same societal function need to be governed at different levels?

After several model versions, the answer is still not “proven”.

But v0.5 gives stronger reasons to take the question seriously.

A frozen model faced eight entirely new functions. The new authority inputs were locked before outcome research. The model improved on its predecessor in component hit rate, full architecture compatibility, authority hit rate and absolute architecture distance.

Particularly clear was that a high blind-coded need for binding system integration coincided with the two new cases where broad authority was actually decisive.

At the same time, the model consistently failed on several local functions and on certain more complex combinations of local production and broad institutional infrastructure.

Perhaps the most important result is therefore not a specific hit percentage.

It is a change in how the subsidiarity question can be formulated:

[\boxed{ \text{not “what level should govern?”} }]

but:

[\boxed{ \text{“which governance function needs to be at which level — and why?”} }]

That is a more demanding question.

But probably also a more useful one.

Technical appendix

A. Frozen v0.5 architecture

$$[G_5=(A,P,X,F,K_P,K_I)]$$

with 378 permitted architectures over:

$$[\{1,4,20\}.]$$

B. Problem vector

$$[\theta_5=(L,S_{14},S_{420},C_{14},C_{420},E,R_{14},R_{420},B_{14},B_{420})]$$

with:

$$[0\leq B_{420}\leq B_{14}\leq 1.]$$

C. Binding fragmentation

$$[J_B]$$

$$w_B[B_{14}d(1,4)\mathbf{1}_{\{A<4\}}+B_{420}d(4,20)\mathbf{1}_{\{A<20\}}].$$

D. Supra-authority realisation

$$[M_{\mathrm{sup}}]$$

$$\mathbf{1}_{[P>A]}+\mathbf{1}_{[X>A]}+\mathbf{1}_{[F>A]}+\mathbf{1}_{[(K_P>A)\vee(\Gamma_I>0)]}.$$

$$[J_{R5}]$$

$$J_{F5}+\lambda_I\kappa_MM_{\mathrm{sup}}.$$

Frozen:

[\kappa_M

0.028922276644969897.]

E. Primary blind-test metrics

- component P(allowed),
- joint-compatible probability,
- six-role architecture distance,
- signed scale deviation,
- authority P(allowed),
- modal full-tuple compatibility.

No single metric was defined in advance as a validation threshold.

Sources for outcome research

1. **SISAB — Our history.** SISAB describes how the company was formed to streamline the management of Stockholm's school buildings.
<https://sisab.se/sv/om-sisab/var-historia/>
2. **City of Stockholm — Action plan for reduced graffiti.** The traffic administration is responsible for cleaning the city's and district administrations' objects within its area of responsibility and procures external contractors.
<https://start.stockholm/globalassets/handlingsplan-for-minskat-klotter-2021.pdf>
3. **Swedish National Agency for Education — Choosing preschool class and compulsory school or adapted compulsory school.** The municipality is responsible for school placement at municipal schools within the framework of the Education Act.
<https://www.skolverket.se/styrning-och-ansvar/regler-och-ansvar/ansvar-i-skolfragor/valja-forskoleklass-och-grundskola-eller-anpassad-grundskola>

4. **Swedish Transport Agency — Parking ticket.** Parking tickets on street land are issued by the municipality or police while the Swedish Transport Agency administers the parking tickets and payment information.
<https://www.transportstyrelsen.se/sv/vagtrafik/fordon/skatter-och-avgifter/parkeringsanmarkning/>
 5. **Sahlgrenska University Hospital — Department of Clinical Microbiology.** Clinical microbiology is part of Regional Laboratory Medicine with laboratory operations at several hospitals in Region Västra Götaland.
<https://www.sahlgrenska.se/omraden/omrade-4/verksamhet-klinisk-mikrobiologi/>
 6. **Sahlgrenska University Hospital — Imaging and Intervention Centre / Nuclear medicine.** PET/CT is conducted within Sahlgrenska's nuclear medicine and radiology operations.
<https://www.sahlgrenska.se/forskning-utbildning-innovation/samverkan/verksamhetebild--och-interventionscentrum-boic/>
 7. **Svenska kraftnät — Svenska kraftnät's responsibility in the power system.** Svenska kraftnät is the system operator authority for the transmission system, responsible for the transmission grid and balances the system.
<https://www.svk.se/om-kraftsystemet/oversikt-av-kraftsystemet/svenska-kraftnats-ansvar-i-kraftsystemet/>
 8. **LFV — Here is LFV / Air traffic control.** LFV conducts air traffic control through several control centres and local air traffic services at several airports.
<https://www.lfv.se/om-oss/dethararlfv/har-finns-lfv>
 9. **Swedish Transport Agency — Organisations providing ATM/ANS services in Sweden.** Several certified organisations provide various air traffic services in Sweden.
<https://www.transportstyrelsen.se/sv/luftfart/flygplatser-flygtrafiktjanst-och-luftrum/Flygtrafiktjanst/organisationer-som-utovar-atmans-tjanst-i-sverige/>
-

Reproducibility

The complete research chain includes, among other things:

- frozen v0.5 engine,

- design and structural review documents,
- (κ_M) freeze protocol,
- implementation audit,
- analytical/synthetic suite,
- blind input lock,
- prediction lock,
- outcome lock,
- integrity audit,
- and full J01–J08 synthesis.

The public report summarises these results. The technical artefacts should be published or linked separately for readers who wish to review methodology and reproducibility.