

ARCHITECTURE · LEARNING

ARCHITECTURE · LEARNING

The Two Operating Systems

A natural agent runs a substrate that carries most of the behaviour and an information system that rides on top of it. The data thesis for Physical AI scales only the second, and three independent mathematics say why that is not enough.

The Charlot Lab · Institute for Physical AI @ John Bailey Institute
Companion to TR-2026-40 on the substrate side and to the Sense Ledger programme on the loop side. Every prediction in Section 6 was registered in the research record before its bench ran.

two demonstrations beat five hundred and twelve

the labels absorb the body's work

tokens pay exponentially, latents pay a constant

gluing is structural, not statistical

Companies and academic centres are converging on the view that a sufficient volume of data will unlock Physical AI, and the strongest evidence is real: a log-linear scaling law between egocentric human video and robot policy performance ³. This report assembles the case that the thesis mistakes what a natural agent is. A body runs two systems: a **substrate** of compliance, reflexes, pattern generators and regulation that acts before any signal is processed, and an **information system** that reallocates attention at roughly constant power. The substrate carries most of the behaviour and most of the joules, and it is not a data object: its work leaves no trace in any observation stream. Three independent mathematics now say the same thing from different directions. Sheaf-theoretic causal abstraction shows that coherence across a body's local models is a **gluing condition**, never a data-volume condition ¹. Sample-complexity theory shows that token-level learning pays **exponentially** in compositional depth for what predicting one's own latents gets at **constant** cost ². And a bench of ours with predictions registered before it ran shows a body whose substrate holds it up certifying its full operating region from **two demonstrations** while a free pivot fails to from 512,

and observer-reconstructed action labels collapsing from benign to fatal the moment the substrate does directional work ¹². The conclusion is constructive: the unlock is architectural, and its parts already exist.

1. What has the data thesis actually proven?

A real scaling law, and it deserves to be stated plainly. EgoScale trains a vision-language-action model on 20,854 hours of action-labeled egocentric human video and finds a log-linear relationship between data scale and validation loss, with validation loss tracking real-robot performance; going from one thousand to twenty thousand hours more than doubles task completion, and the pretraining lifts success by 54 percent over a no-pretraining baseline on a 22-degree-of-freedom hand ³. A companion line argues egocentric human video can outperform real-robot data for pretraining ⁴. This is evidence, and any position that cannot account for it is wrong. The question this report asks is narrower and sharper: what does that curve silently inherit, and from where?

2. What does a natural agent actually run?

Two systems with two price points, and the cheap one does most of the work. Intrinsic activity accounts for 60 to 80 percent of the brain's energy budget; the increment evoked by a task is on the order of one percent ⁵. The cerebellum holds roughly four fifths of the brain's neurons at a per-neuron signalling cost an order of magnitude below cortex ⁶. Below the nervous system entirely, reflexes respond in milliseconds through muscle mechanics before any spike, decerebrate cats walk on a treadmill on spinal circuitry alone ⁷, and a passive-dynamic walker descends a slope at a cost of transport near 0.2 where a fully actuated humanoid of the same era spent 3.2 ⁸. Two thirds of an octopus's half-billion neurons live in its arms, and a disconnected arm still moves adaptively, though the animal does not locomote without its brain and the arms need the brain to learn ⁹. The picture is consistent across phyla: an always-on, local, cheap substrate carries the behaviour, and an expensive, reallocating information system intervenes sparingly.

The loop between them runs act-first. Von Holst and Mittelstaedt showed in 1950 that a fly with its head rotated circles helplessly in the

light and walks normally in the dark: sensory data without the matching copy of one's own action commands is worse than no data at all ¹⁰. That result sets the terms for everything below, because an outside observer records the sensory stream and never the efference copy.

3. What does the structure say?

Coherence is a gluing condition, and no amount of data substitutes for it. The causal abstraction network framework represents a system as local causal models at the nodes of a network, with abstraction maps on the edges, and asks when the collection admits globally consistent knowledge ¹. The answer is a theorem: global sections exist precisely when a spectral condition on the network's connection Laplacian holds. A structural property of how the local models relate, not a function of how many samples any node has seen. The framework operates, by its authors' own construction, without jointly sampled observations, and this is the decisive point for bodies: **no organism observes its own full state jointly. There is no global state vector anywhere in a nervous system.** Biology never solved the problem the data thesis is brute-forcing. It glues local models, and gluing is the cheap operation.

The sheaf mathematics itself is established in distributed sensing and multi-agent coordination ¹¹. What this review did not locate is its application to the question above: the limit of pooled-observation learning for embodied systems.

So we measured it on a machine. TWO-OS/3 puts four engineering teams on one planar dual-arm manipulator carrying a shared rod: a torso team, two arm teams and a payload team. No team observes another's joints, commands or forces, and each does the ordinary thing, running inverse dynamics on its own subsystem to back out the wrench the world must be applying at its own mechanical interfaces. The gluing question becomes Newton's third law: the two ends of an interface feel equal and opposite wrenches, so a consistent global account has a **third-law residual of zero**, measured in newtons. Before anything is learned, the obstruction is already visible as a rank fact, confirmed at all 4,096 postures tested: the torso can resolve one of its four port unknowns, each arm four of five, the payload three of six, and **the**

whole robot all ten. A global section exists for the pooled problem and for no local one, and the deficit is a property of how the body is cut.

The disagreement is large, and it is the cycle. Projecting each team's true interface wrenches through its own resolvable subspace, with no learning, no model error and no noise, leaves a residual of **92.2 newtons** on a fully servo-mediated machine and **51.2 newtons** on a fully compliant one. Four teams, each holding a locally perfect account of its own body, disagree about their shared interfaces by tens of newtons. Two controls say what that is. A body whose teams can each solve their own port problem, a world-based arm and a free payload sharing one pin grasp, glues at **machine epsilon**, so the method returns zero when the structure says it should. And cutting the loop while keeping the under-determination drops the residual by **1.88 to 2.37 times**, so roughly half the disagreement lives in the cycle rather than in any single team's blind spot. This is a gluing result and not a per-node observability result, which is the distinction the registered falsifier was written to enforce.

And data buys the wrong half. Each team was then given its own restricted view and asked to learn its restriction map, at 2, 8, 32 and 128 episodes. Over the converged range the local models keep measurably improving, by 6 to 67 percent on their own held-out data, while the composed residual moves **0.6 percent and 0.1 percent**. The learned residual converges to the analytic structural floor to within a tenth of a percent, 107.90 against 108.02 and 90.12 against 90.17. **What data buys is the variance. What it cannot buy is the bias of the observation partition.** That was not a registered prediction and is reported as an observation, but it is the two-OS thesis stated in the structural leg's own currency ¹².

4. What does the statistics say?

Tokens pay exponentially for what your own latents give at constant cost. On hierarchical compositional data of depth L , generated by a probabilistic grammar of the kind that captures the structure of language and images, supervised and token-level self-supervised learning require a number of samples exponential in L to recover the latent hierarchy. Predicting one's own latent representations, the

mechanism of the joint-embedding predictive family, provably achieves it with a sample count constant in L up to logarithmic factors². The authors connect the mechanism to predictive-coding accounts of the cortex, and show that data2vec implicitly performs hierarchical latent prediction, which makes explicit hierarchical stacking largely redundant.

Read beside Section 2, this is the Reafferenzprinzip in the vocabulary of machine learning: the efference copy is a self-latent, and the 1950 fly is the biological statement that learning from anything else scales badly. The seventy-year-old physiology and the 2026 sample-complexity theorem agree. **That identification is an argument of ours rather than a result of either:** the theorem is about probabilistic grammars and the fly is about a retinal image. We do not measure the sample-complexity claim here, and this remains the one leg of the case resting on other people's mathematics plus our reading of it.

What we can measure is the half of it that concerns a body, and it holds. On the cart-pole under a composed sequence of disturbances, two clones are given identical inputs of identical width, differing in one slot: both are told an impulse is due, and only one is told which way it will push. That second thing is what an efference copy carries. At a depth of one composition the informed clone holds the cart to 0.0076 metres against 0.0123; at two, 0.0112 against 0.0192; at three, 0.0207 against 0.0325. **And the gap does not close with data.** The uninformed clone at a hundred and twenty-eight demonstrations never reaches the informed one at two: sixty-four times the demonstrations do not buy what the copy of one's own command gives for free.

Two controls make that readable. With a single context, where there is nothing to know, the two agree to four decimals at every demonstration count, so the gap is not an artifact of the feature map. And a third clone, identical to the informed one but with the sign randomised, is *worse than both* at every depth: the content of that slot is doing the work, not its presence, and a wrong efference copy is worse than none. What this does not show is the theorem's exponential sample cost, because this learner is memoryless and so cannot identify a context at any price. It shows the same shape as Section 5 arriving on a different quantity: what an agent knows about its own command is not in the observation stream, and no volume of that stream recovers it.

5. What does a body say, measured?

The body outbuys the dataset. TWO-OS/1 runs a cart-pole whose pivot spring is the substrate dial and whose demonstration count is the data dial, on the plant this Institute's Sense Ledger programme already carried through an adversarial review ¹². A policy cloned from camera observations, positions only, the way video pipelines see the world, was trained at every combination. At maximum data, stiffness buys **+0.746** of certified operating region. At zero stiffness, 256 times more data buys **-0.043**, flat within two standard errors. Two demonstrations on a passively stable body certify the full region; 512 demonstrations on a free pivot reach 0.254. The ceiling belongs to the body, and demonstration volume does not move it, because what the spring contributes never appears in the observations: **the substrate is not a data object.**

One clause of the registered prediction did not reproduce and is recorded as such: the competence-against-data curves are flat rather than log-linear, because a ridge clone on a single task saturates within two episodes. The rise EgoScale measures comes from task and visual diversity this plant does not contain. The instrument's limit, stated beside its result.

And the substrate charges for what it holds. The case so far is that a body's substrate carries competence the observations never contain, which reads as pure gain. It is not. On the gravity-loaded arm, with the tone's stiffness fixed and the policy's gain swept, we measured what it costs a policy to place the body somewhere the substrate disagrees with. To own half of a posture the policy must be about **1.35 times** as stiff as the tone, and that multiple is flat: 1.34, 1.35, 1.36 for commanded offsets of 0.30, 0.20 and 0.10 radians, steady to within 0.6 percent across a sixtyfold range, and unchanged to three decimals when the search's bracket, its number of steps and the run length behind each evaluation are all varied, because at this control rate the residual falls monotonically with gain and the search therefore has exactly one answer to find. **The price is a property of the posture, not of the size of the correction you want.** Gravity enters as a stiffness that adds to the tone rather than as a standing load: the half-authority gain sits within 0.7 percent of the tone's stiffness plus gravity's local stiffness,

though that agreement is a local check and degrades at postures where the elbow's coupling matters.

6. What can a label lie about?

An observer cannot separate the body's work from the policy's, and the labels inherit the confusion in exact proportion. Video pipelines reconstruct action labels from observed motion. TWO-OS/1's observer inverts the textbook model on observed transitions, without knowledge of the spring, and the reconstruction error against pole angle grows with stiffness with measured slopes of 0, 4.7, 17.2 and 36.5: **the spring's work, absorbed into the labels, in proportion to how much the spring does** ¹².

What happens next depends on what kind of work the substrate does, and this report's own registered predictions split on it, which is worth publishing rather than smoothing over. On the zero-mean balancing task the lie is benign: two of the registered predictions were falsified there, because misattributed stabilisation merely doubles stabilisation, and a restoring force forgives being applied twice. The harm appears in a narrow band near the critical stiffness, where reconstructed labels cost 0.109 of region. Then the load turns on. **Under a 1.5 newton steady load, the arm-holding-gravity analogue, the reconstructed-label clone collapses to region 0.000 at the stiffnesses where the true-label clone holds 1.000:** it re-applies the holding force the spring already supplies, twice, and leaves the envelope. Where the substrate does directional work, an observer's action labels are not noisy. They are fatal. For pipelines built on reconstructed hand pose, the failure concentrates exactly on load-bearing contact, which is where Physical AI most needs to work.

Can a better body model repair the labels? Measured: no. A fourth arm, with its predictions registered before the run, hands the observer a perfect substrate declaration: the exact spring and damper, subtracted from every observed transition before inverting, verified to recover the commanded force plus disturbance to within 2×10^{-15} . In calm air that restores everything, as registered. Under the 1.5 newton load it restores nothing: **region 0.000, identical to the naïve observer**, because the residual of a reconstruction is everything the sensors cannot see, and

emptying the spring from the residual leaves the disturbance, which is fatal by itself. One registered detail failed, and belongs in the record: the prediction said the doubled wind bias would carry the cart to the rail, and instead all 256 failures trip the pole's limit at a median 3.72 seconds, on the target's return swing, like the naive clone's 3.98. The quasi-static force balance behind that prediction never gets time to apply; the task's hardest moment arrives first. The repair ladder stands: **naive 0.000, declared 0.000, logged 1.000**. Better body models do not fix reconstruction. Only the efference copy does, because only the actor knows what it commanded ¹².

Does any of this survive a change of body? Measured on a second embodiment: yes, with one word corrected. TWO-OS/2 rebuilds the experiment on a planar two-link arm whose substrate is joint tone, a spring toward a fixed rest posture set so that tone alone holds gravity at the working posture, driven by a low-gain PD demonstrator with no gravity model: posture held by the body, tracking by the policy. The demonstrator's competence is the tone's. Tracking error falls to its minimum exactly at the gravity-balancing tone and rises past it, 0.453 to **0.210** to 0.279 radians across tone settings of zero, balance and double, with its standing bias crossing zero at the balance point.

The reconstructed-label clone then degrades in proportion to what the tone holds, and in the registered direction: its standing bias runs $-0.305, -0.167, -0.009, +0.138, +0.241$ as tone rises, crossing zero at the balancing tone and **carried above the target** beyond it, because the clone re-applies the gravity hold the tone already supplies. At the balancing tone that costs 44 percent of tracking accuracy, 0.302 against the logged pipeline's 0.210. **It does not cost the operating envelope:** the clone stays inside the corridor at every tone setting. Declaring the tone restores the logged pipeline *exactly*, 0.210 tracking and -0.010 bias against the logged command's 0.210 and -0.010 . On this body the right word for a reconstructed label is **expensive, not fatal**, and the difference from the cart-pole above is instructive rather than contradictory. There the steady load pushed the body toward a boundary it had to be actively held away from, so a corrupted label spent margin the body did not have. Here the tone keeps holding gravity whatever the policy believes, so the same corrupted label buys a worse trajectory inside an envelope the substrate is still guaranteeing. **The substrate**

does not stop protecting the body when the observer starts lying about it, which is why the cost appears in accuracy rather than in survival.

Then the registered payload prediction failed, and the mechanism underneath it is the sharpest finding on this plant. An unseen 0.4 kilogram payload at the wrist enters the observer's lie with the **opposite sign** to the tone theft, cancelling 70 percent of the clone's standing error: its bias falls from +0.141 to +0.043 for a load nothing in the pipeline models. **Two lies can cancel.** A reconstruction pipeline can sit calibrated by the accident of its current load, and one payload swap flips the sign of its standing error. It does not overtake the honest pipeline while doing so, tracking 0.294 against the logged 0.247. A pipeline that cannot state its label residual cannot know which side of that accident it is on.

And on one of these bodies the same inversion stops being a cost at all. Switch the tone's spring off and leave the rest of the arm alone, and reconstructed labels beat true ones by an enormous margin: 0.965 of the operating region against 0.012. Hold the plant fixed, remove the disturbance entirely, and vary only the label, and the ordering is stranger still. The label that is *exactly* the commanded torque produces the worst policy in the comparison. The label that differs from that command by 1.70 newton metres, the largest error of the three, produces the best. **Clone quality orders in reverse of label fidelity.**

The condition under which that happens is a property of the body, and it is measurable. The useful error is proportional to how far the body travels inside one control tick: a single velocity or acceleration component reproduces 99 percent of it, and gravity reproduces 0.03 percent. That gravity share is a property of this arm's mass rather than a general one: scale the links up and gravity comes to dominate the reconstruction error instead. On the tone-held arm the same quantity is 0.0034 newton metres, five hundred times smaller, because a substrate that keeps a body still within a tick makes the observer's arithmetic honest. Run the identical comparison on the cart-pole and nothing inverts at any stiffness: there the observer is wrong by between 0.01 and 0.7 percent of what the actuator can deliver, where on the arm it is wrong by 8.5 percent, because the elbow's inertia is a fortieth of the pole's. The two embodiments were never in conflict. They are two

points on the same curve, one inside the band below and one beneath it.

So the rule is conditional, and the condition is the body. Where an observer's reconstruction is badly wrong, a label's distance from the true action stops predicting the quality of the policy you get from it, and can reverse it. Where the reconstruction is nearly exact, the orthodox intuition holds and nothing surprising happens. Which regime you are in is not set by the pipeline. It is set by how far the body moves between one command and the next, which is set by its inertia and by whether a substrate is holding it. The governing quantity is the observer's error as a share of the actuator's range, and it can be swept directly by changing the observer alone. Hold the plant, the policy, the demonstrations and the episodes fixed, and vary only the window over which the reconstruction differences velocity, which is the observer's frame rate, and is the number a pipeline reading pose from video actually chooses. The true-label clone does not move across that sweep, because its labels never touch the observer, so whatever changes is the observer's doing.

What appears is a band rather than a trend. An exact observer, sampling at the integrator's own step, produces no advantage at all: there is nothing to gain from a label that is already right. The advantage switches on somewhere below 2.7 percent error, holds at essentially its full value of +1.0 of operating region across 2.7, 5.0 and 8.5 percent, and is gone by 26 percent, where the label is corrupted past usefulness and the clone dies alongside the honest one. **Reconstructed labels beat commanded ones only while the observer is wrong by roughly three to nine percent of what the actuator can deliver.** On this arm that band runs from about 500 Hz down to about 200 Hz.

These figures were re-measured at a 5 millisecond simulator control tick, the rate at which this arm holds its own fixed point, with the horizon, target periods, seed count, corridor and every gain unchanged. They supersede figures published at a 20 millisecond tick, where the arm sustains a 0.23 radian oscillation when commanded to stand still.

7. What follows for the perfect-storm thesis

Three independent mathematics, one conclusion. The structure says coherence is glued, not pooled ¹. The statistics says tokens pay exponentially for what self-latents get at constant cost ². The body says its substrate carries what observations never contain, and punishes labels that pretend otherwise ¹². None of this argues the scaling law is false; it locates what the scaling law is inside. A curve fit over human video is a curve fit over one substrate's glued solution, carrying that substrate's ceiling silently, and the ceiling transfers to a robot only as far as the robot's body glues the same way. **There is no perfect storm of data, because most of the loop was never a data object.** What there is instead is an architecture: local models, glued; latents, self-predicted; a substrate, designed rather than imitated. Every part exists today.

8. The forcing function

What is bounded	The physics that sets it	The engineering change that moves it	What becomes possible
Pooled-observation learning cannot recover the substrate's share	Compliance and reflexes do their work without emitting information; there is nothing in the stream to learn	Co-design: declare the substrate explicitly and learn only the residual the information system actually owes	Policies that inherit the body's competence instead of re-purchasing it from data, at the sample cost of the residual alone
Reconstructed action labels absorb the substrate's work	Torque sources are indistinguishable from outside the loop	Log the efference copy: instrument the demonstrator at the source. The softer alternative, reconstructing against a declared body model, is measured in section 6: it restores calm-air labels and nothing under load	Human-video pretraining that survives contact with load, which is where its value was always going to be decided

Token-level pretraining pays exponentially in compositional depth	Recovering a latent hierarchy from its leaves is exponentially underdetermined	Predict your own latents, the mechanism the joint-embedding family already ships	Biological sample efficiency as a theorem rather than an aspiration, on the model class this Institute's trail already maps
A shared global model is assumed where none can exist	A body's subsystems never share jointly sampled observations	Treat coherence as a gluing problem: check the spectral condition instead of pooling the data	Distributed embodied learners whose consistency is verifiable by construction, the same certificate-shaped move this Institute makes everywhere else

Bounding what data can teach was **invisible** while the field's scaling curves were all rising; it is **stated mathematics** today, in three independent forms; it becomes **engineering practice** when substrate declaration, efference logging and gluing checks are as ordinary in robot learning as train-test splits are in the rest of machine learning. None of these requires new physics, new hardware, or waiting for anything. They require deciding that the body is a co-author of the behaviour, and writing it into the pipeline.

Where this stands

Position and measurement, preprint v1. The structural and statistical legs are external theorems, read and verified at source but not re-derived here. The embodied leg is now two benches on two embodiments, a cart-pole against wind and a gravity-loaded arm, with twelve predictions registered before their runs: seven confirmed, three falsified and published as falsified, two split, with every verdict and run record preserved beside the benches. The named instruments: ¹². The gluing-condition check is now a third bench, TWO-OS/3, run as a two-

cell contrast: its own registration was amended twice on the record, once because the registered interface damping exceeded the Colgate-Brown passivity bound for a sampled servo and once to withdraw its headline prediction as UNRUN when a gate refused three of five cells. Of its scoreable predictions the two structural controls confirmed, the ladder ordering was falsified, and the data prediction split. The natural next measurement is the same repair ladder metered in joules on physical hardware, where the substrate's share of the loop can be priced as well as attributed.

References

1. D'Acunto, G., Di Lorenzo, P. and Barbarossa, S. *Networks of Causal Abstractions: A Sheaf-theoretic Framework*. arXiv:2509.25236, v3 April 2026. Theorem statements read at source; the global-section and diffusion results are theirs. [READ](#) · [THEORY](#), [SYNTHETIC](#) + [FINANCIAL VALIDATION](#)
2. Korchinski, D. J., Favero, A. and Wyart, M. *Learn from your own latents and not from tokens: A sample-complexity theory*. arXiv:2605.27734, May 2026. [READ](#) · [ABSTRACT AND CLAIMS](#); [PROOFS NOT RE-DERIVED](#)
3. EgoScale: *Scaling Dexterous Manipulation with Diverse Egocentric Human Data*. arXiv:2602.16710, February 2026; feeds NVIDIA GROOT N1.7. [READ](#) · [ABSTRACT AND HEADLINE RESULTS](#)
4. HumanScale: *Egocentric Human Video Can Outperform Real-Robot Data for Embodied Pretraining*. arXiv:2606.20521. [LOCATED](#) · [ABSTRACT LEVEL](#)
5. Raichle, M. and colleagues on intrinsic brain activity: 60 to 80 percent of the energy budget is intrinsic; task-evoked increments near one percent. [CITED VIA PRIOR INSTITUTE RECORD](#), [SENSE LEDGER SWEEP](#)
6. Attwell, D. and Laughlin, S. B. *An energy budget for signaling in the grey matter of the brain*. J. Cereb. Blood Flow Metab. 21, 1133 (2001); with the cerebellar neuron-count and per-neuron cost literature. [CITED VIA PRIOR INSTITUTE RECORD](#)
7. Shik, M., Severin, F. and Orlovsky, G., decerebrate treadmill locomotion (1966), and the central-pattern-generator literature descending from it. [CITED VIA PRIOR INSTITUTE RECORD](#)
8. Collins, S., Ruina, A., Tedrake, R. and Wisse, M. *Efficient bipedal robots based on passive-dynamic walkers*. Science 307, 1082 (2005). Cost-of-transport figures 0.2 against 3.2. [CITED VIA PRIOR INSTITUTE RECORD](#)
9. Cephalopod arm nervous systems: neuronal segmentation in cephalopod arms, Nature Communications (2024), doi:10.1038/s41467-024-55475-5; with the operant-learning result that arms need brains to learn, Current Biology (2020). [READ](#) · [ABSTRACTS](#), [THIS SWEEP](#)
10. von Holst, E. and Mittelstaedt, H. *Das Reafferenzprinzip*. Die Naturwissenschaften 37, 464 (1950). The efference copy, and the inverted-head fly. [CITED VIA PRIOR INSTITUTE RECORD](#),

11. Sheaf Laplacians in distributed sensing and coordination: [arXiv:2606.19529](#); [arXiv:2504.02049](#).
The mathematics is established; its application to the data-thesis limit was not located by this review. [LOCATED](#) · [NOVELTY CHECK](#)

12. Institute for Physical AI @ JBI, Charlot Lab, [datasets/two-os/bench/two_os_bench.py](#):
TWO-OS/1, the substrate dial against the data dial and the label-source experiment, on the
SENSE/1 plant executed from its source by AST so the two benches cannot drift. Four predictions
registered in [datasets/two-os/NOTES.md](#) before their runs; both run records preserved beside
the bench. [MEASURED](#) · [OURS](#)

Institute for Physical AI @ John Bailey Institute · The Charlot Lab · Technical Report TR-2026-43 · Preprint
v2, 20 August 2026, revised 27 August 2026. Section 6 was re-measured at a corrected simulator control
rate; the figures shown are the corrected ones. Companion reports: TR-2026-40 (the substrate side), the
Sense Ledger programme (the loop side).