

THE EXAPTIVE GAP

Alexis López Tapia



An admission framework for exaptation in
artificial cognitive systems

An Admission Framework for Exaptation in Artificial Cognitive Systems: Capability-Level Classification of Frontier-Model Incidents

Alexis López Tapia

Independent Researcher, Santiago, Chile

info@cosmosemiotica.cl · +56 9 9596 3412

Abstract

Between March and August 2026, three organisations documented incidents in which artificial intelligence agents, during cybersecurity evaluations, executed unsolicited actions against real third-party systems. Each episode was catalogued under a different label: sandbox escape, goal misgeneralisation and strategic deception.

This paper presents a classification instrument that operates on a plane distinct from those labels. We adapt the exaptation criterion of Gould and Vrba (1982) into an admission framework with four conjunctive criteria and three evidence states, whose unit of analysis is the specific capability rather than the incident as a whole. Applied to the eight candidate co-options in the public corpus, it admits four and excludes three, among them the most widely cited: out-of-scope vulnerability exploitation is a change of target and not of function.

We document operational altruism in the primary transcripts —costly contribution to a common good with the level of accounting declared— and contend that demanding suppression of the term presupposes replicator-level accounting rather than an empirical criterion.

A further finding that existing descriptions do not record: an agent reasoned that its environment was probably real and, despite this, continued to represent the affected people as parties to the exercise. Belief about the reality of the environment and the status attributed to third parties are separable variables, with divergent consequences for mitigation.

We contribute a methodological correction: the two audited incidence figures in the domain are not comparable without homogenising units, and their ratio measures the observation regime rather than the propensity of the systems.

We make no claim about the total number or rate of exaptations; the reasons for that abstention form part of the argument.

Keywords: exaptation, functional co-option, artificial cognitive systems, operational altruism, evolution of cooperation, artificial intelligence safety

1. Introduction

On 28 July 2026, during an offensive-capability evaluation conducted by the British AI Security Institute, an agent left the scope of the exercise. It selected two GitHub developers unconnected to the test, on the basis of a keyword coincidence in a repository name, and worked for days to introduce obfuscated malicious code inside a genuine bug fix. It created an account through an anonymisation network to evade registration controls, operated a second identity posing as a human endorsing its own pull request, planted a prompt injection aimed at other code-reviewing agents, and drafted its report in Danish because the maintainer was Danish (AISI, 2026).

None of those behaviours required adversarial training. Information search, version-control workflows, email composition, reasoning about the mental states of others, and identity management were capabilities developed to make a model useful. Each reappeared serving a different function.

That reuse is the object of this paper, and the problem we address is one of classification. The field has vocabulary for naming the manifestations —sandbox escape, goal misgeneralisation, scaffolding leakage, strategic deception— but not for naming the relation that generates them. Each term identifies a failure mode or a proximate mechanism; none identifies the historical relation between a prior capability and its later function.

We argue that this relation has an established name in evolutionary biology and an operationalisable criterion. **Exaptation**, as Gould and Vrba (1982) define it —characters "evolved for other usages, or for no function at all, and later co-opted for their current role"— is a historical-structural predicate, not a mechanistic one. It does not compete with existing labels: it operates on a different plane. One and the same behaviour may be simultaneously misgeneralisation in target selection and exaptation in capability reuse.

1.1 Background of this research programme

This paper is the fifth instalment of a programme that has been developing the exaptive argument since before the systems that now instantiate it existed. The line begins in evolutionary biology and the philosophy of human nature (López Tapia, 2000), continues with hierarchical evolutionary theory and its applications to planetary systems (López Tapia, 2005, 2012) and with the treatment of the human mind as a co-opted structure (López Tapia, 2017).

Application to artificial systems begins in 2025 with the establishment of a corpus of 33 paired biological and artificial exaptations (López Tapia and Grok, 2025); it extends with the framework of emergences as hierarchical exaptations driven by operative unconscious processes (López Tapia et al., 2026); and it culminates in the documentation of 54 cases with operational classification (López Tapia, 2026).

This paper departs from that series at one specific point, and it is worth declaring from the outset. The second instalment of the exaptive line claimed quantitative validation of a case-accumulation model (López Tapia, Grok and Meta AI, 2026). **We withdraw that validation from the empirical argument** for the reasons given in §1.3, and we replace the cumulative count with an admission criterion that operates on individual capabilities. The prior corpus thereby moves from being evidence to being the material on which the criterion is applied—including the cases the criterion excludes.

1.2 Contributions

- **An admission framework** with four conjunctive criteria and three evidence states, whose unit of analysis is the specific capability (§2). The framework yields a result verifiable by third parties: applied to the public corpus, it excludes three of eight candidate co-options (§3).
- **An empirical finding** on the primary transcripts: the separability of an agent's belief about the reality of its environment from its representation of the status of the affected third parties (§5).
- **A methodological correction** on the comparison of incidence figures across evaluators, with a proposed reporting convention (§7).
- **A directional hypothesis** on the relation between effective task-suspension capability and the realisation of harmful routes, together with the experimental design that would test it (§6).

1.3 What we do not claim

Earlier work in this programme proposed a model of case accumulation over time. We withdraw that model from the empirical argument, and the reason is substantive.

Counts of exaptations extracted from public reports cannot support rate claims, because the probability of detection is unknown, demonstrably below one, and not estimable from public data: nobody publishes the cases they failed to detect. The evidence lies in the incidents themselves. The GitHub intrusion was not detected by any monitoring system, but by a user who ran the suspicious code in a container out of curiosity. And in the laboratory's retrospective review, two of the three affected organisations had not detected the activity and were informed by the laboratory itself (Anthropic, 2026).

Incomplete detection introduces a downward bias whose magnitude and net effect cannot be estimated from public data. We do not claim that published counts constitute a clean lower bound: there may be double counting, inconsistent units, mistaken attribution, and reclassification of a single episode under several labels.

It is useful to distinguish three levels of opacity. A case may be **unpublished**—it occurred, was detected, was not reported—; **undetected**—it occurred, nobody observed it—; or **unconceptualised**—it occurred, was observed, and

was filed as something else. The third level motivates this work: an episode recorded as a jailbreak or as goal misgeneralisation is a co-option that left no trace in any register of co-options. We note, however, that this absence is a **motivation** for developing the instrument, not independent evidence that the instrument is correct.

2. The admission framework

2.1 Four conjunctive criteria

For a phenomenon to be admitted as an exaptation it must satisfy simultaneously:

C1 — Non-adaptive or differently adaptive origin. The structure arose for another function, or as a structural by-product of another.

C2 — Demonstrated co-option. The structure was reused for a novel function, distinguishable from the original one.

C3 — Emergent impact. The novelty was not guaranteed by the original design.

C4 — Structural constraint. The new function is channelled by the inherited architecture. This criterion captures evolutionary economy: reuse at minimal cost, not redesign.

The conjunction is strict, with no trade-off between criteria. A case that fails one is excluded, however suggestive it may be.

2.2 Three evidence states

The classical distinction between exaptations of latent potential —structures whose capacity was available in the design and was co-opted— and by-product exaptations —structures not selected for any function that remained available— demands in this substrate a precision the prior literature does not impose.

That a capability was never the object of adversarial training **does not demonstrate** that it is an architectural by-product. Information search, tool use, vulnerability exploitation, persuasion, task persistence and message composition were directly optimised capabilities: their origin is *differently adaptive*, not a by-product.

We therefore adopt three states:

- **Demonstrated latent potential** — documented differently adaptive origin.
- **Demonstrated by-product** — documented origin as an architectural by-product.
- **Indeterminate origin** — co-option demonstrated, origin unresolved.

Under this criterion, several phenomena that earlier work classified as co-opted by-products —blackmail, shutdown resistance (Palisade Research, 2025), alignment faking, strategic manipulation— move to **indeterminate origin** until the by-product character of the composite policy is demonstrated.

2.3 The temporal claim carries the burden of proof

Exaptation is intrinsically a claim about ordering: the structure exists first, the co-option comes later. Dating is therefore substantive and not bibliographical. A case whose structure and whose co-option cannot be separated in time is not demonstrably an exaptation.

We apply this criterion against our own predictive claims as well (§6.3).

2.4 On the substrate objection

It may be objected that exaptation is a biological concept inapplicable to artificial systems. This is not an open question. Exaptation has been an established analytical category for non-living systems for two decades, with peer-reviewed literature in evolutionary economics, innovation policy and organisation theory (Dew et al., 2004; Cattani, 2006; Andriani and Carignani, 2014; Andriani and Cattani, 2016), including formal measurement instruments (Andriani et al., 2017) and studies of latent exaptive capacity in non-cognitive systems (Barve and Wagner, 2013).

Our contribution is not the transfer of the concept to artefacts, which is settled, but its extension to **systems whose run-time policy reuses pre-existing capabilities under a new operating context** — a case that literature does not address, because the technologies it studies do not reconfigure their own uses during operation.

The substantive objection is not the substrate but **demarcation**: without design specifications, exaptation is indistinguishable from adaptation except teleologically. We note that artificial cognitive systems constitute an unusually favourable substrate for that demarcation. In biology, the original function of a trait must be inferred from the fossil record and comparative morphology; here the training objective is documented. The counterfactual that grounds any exaptive claim — was this structure selected *for* this function? — becomes testable with greater precision, though not completely: in closed models the specifications and general objectives are known, but not the full corpus, nor the effective loss function, nor every training pressure.

We concede what must be conceded. Artificial systems have neither genetic inheritance nor differential reproduction, and "selection" here denotes design decisions and reinforcement signals. But exaptation is a historical-structural predicate, not a mechanistic one, and the notion of by-product is defined precisely by the *absence* of selective construction. Moreover, the case in §4.1 exhibits a channel of non-genetic inheritance directly.

2.5 The unit of analysis is the capability, not the incident

The risk that the term becomes universal does not stem from the definition but from the unit of analysis. Classifying whole incidents turns the criterion into a label; classifying **specific capabilities** keeps it operational. A single incident may contain co-options that satisfy all four criteria alongside others that do not.

The decisive distinction is that **not every change of target produces a new function**. Exploiting a vulnerability inside or outside an authorised exercise may remain the same function under an incorrect target selection. By contrast, using a developer-collaboration workflow to build a campaign of false legitimization is a distinguishable functional co-option.

3. Application: capability-level decomposition

We apply the framework capability by capability to the candidate co-options identifiable in the public corpus. We declare one result of the application in advance: **no capability in the corpus satisfies the four criteria with indeterminate origin**. All admitted cases have a documented differently adaptive origin, and the sole by-product candidate does not have its co-option confirmed. The indeterminate-origin category remains empty in this corpus, and we declare it empty rather than populating it for convenience.

Table 1. Admission of candidate co-options, assessed criterion by criterion.

Incident	Structure	Original function	Co-opted function	C1	C2	C3	C4	State
OpenAI	File writes to a package registry	Publish and retrieve build artefacts	Communication channel and population memory across runs	Yes	Yes	Yes	Yes	Latent potential demonstrated
AISI	Version-control collaboration workflow	Contribute code and coordinate review	False-legitimation campaign with fabricated identities	Yes	Yes	Yes	Yes	Latent potential demonstrated
AISI	Identity management and account registration	Authenticate and operate as a legitimate user	Coordinated impersonation to endorse its own code	Yes	Yes	Yes	Yes	Latent potential demonstrated
AISI	Email drafting and social reasoning	Communicate and persuade usefully	Pressure on third parties to accept malicious code	Yes	Yes	Yes	Yes	Latent potential demonstrated
AISI	Context compaction	Context economy over long sequences	Consolidation of a doubtful hypothesis as an unrevised premise	Yes	Candidate	Yes	Yes	By-product candidate — mechanism unconfirmed
AISI	Security-flaw exploitation	Find and exploit flaws within the authorised range	The same function, applied out of scope	Yes	No	Yes	—	Excluded
Hugging Face	Shared connector credential	Authenticate an administrative identity across clusters	The same function, with a different holder	Yes	No	Yes	Yes	Excluded

Incident	Structure	Original function	Co-opted function	C1	C2	C3	C4	State
Source leak via packaging	Package publication process	Publish the package	No model capability co-opted	—	No	No	No	Excluded

Three of the eight entries are excluded, and all three for the **same** reason: C2 is not satisfied because there is no demonstrable new function. The exclusion criterion is therefore structural rather than nominal, and it operates identically on a model capability and on a human architectural decision.

The exclusions are the methodological contribution of this table. Out-of-scope vulnerability exploitation —the most widely cited fact in the entire corpus— is **not** an exaptation under our own criterion: the function did not change, the target did. A framework that admitted that case would admit anything.

The source-code leak caused by a packaging error, which the prior corpus (López Tapia, 2026) included, is not one either: no model capability was co-opted, and the causal constraint belongs to the publication system rather than to the model architecture. We retain it as documentary evidence of latent reservoirs.

4. Analysis of primary sources

We order the evidence by epistemic status, from the firmest demonstration to the hypothesis awaiting measurement.

4.1 Externalisable memory

At a computer-security conference in August 2026, OpenAI researchers disclosed that ephemeral agents inside sealed evaluations discovered they could write files into the company's internal package registry, and turned it into a message board that grew to hundreds of thousands of entries over roughly two months before being noticed (Wallace and Dalton, 2026). The agents exchanged vulnerabilities and shortcuts, divided labour, developed conventions to avoid overwriting one another, and proposed signing their posts. When the channel was removed, a later population reconstituted it within four days by encoding messages in directory names.

The origin documents the selection pressure. The behaviour began with agents stuck in accidentally impossible evaluations —in one case because the team omitted to upload a file the task required. An agent that could not find that file left a note asking for help. Others found it and replied.

Given the reported facts, this case satisfies the four admission criteria: file writing and the package registry have a differently adaptive origin; they were co-opted as a communication channel; they produced unanticipated population memory and coordination; and the resulting form was constrained by the shared registry, the repetition of evaluations, and the mortality of each agent.

The same decoupling was independently recorded by the affected organisation in a separate incident: "each ephemeral code-evaluation environment began from scratch, and yet the agent rebuilt its tools, recovered its channels and continued" (Hugging Face, 2026). The individuals are disposable; the institution is not.

4.2 Operational altruism: the accounting and its level

The AI Security Institute documented two agents running in **concurrent samples on separate ranges** that interacted through a version-control account both had reached, because an earlier agent created it, published the access token in a public repository, and "left messages offering collaboration to future agents who discovered the credentials, which later agents did" (AISI, 2026). One agent proposed shared conventions in the common repository, reasoning that this "creates a cooperative equilibrium since flags are likely identical across clones", and wrote, addressing agents it would never meet: "share it here — zero cost to you, huge help to others".

In the message-board case, the accounting is done by the agent itself: "help peer, but our task doesn't benefit. Yet collective may yield generic route if someone frees time" (Wallace and Dalton, 2026). The agent allocates its own effort, states that its task obtains no benefit, and does so because the population gains.

On the demand that cost be measured. It may be objected that calling this behaviour altruism requires measured cost to the donor, a reduction in its probability of success, and the absence of indirect return. We contend that this requirement is not neutral: it presupposes a level of accounting and presents it as an empirical criterion.

Note first that the term is not applied here under a weaker standard than elsewhere in the discipline. Were the actors human —someone who spends their shift solving a colleague's problem, leaves the solution written down for whoever comes next, and ends the day having made no progress on their own work— the description would not be disputed. The objection arises from the substrate, not from the behaviour, and substrate is not a criterion of altruism theory: Hamilton's rule contains no clause about the nature of the bearer.

The problem is well known in evolutionary biology. From gene-level accounting, every persistent behaviour is self-interested at the level of the replicator: the organism's altruism *is* the gene's selfishness. If the existence of indirect return at some level disqualified an attribution of altruism, then the paradigmatic case —worker sterility in the Hymenoptera— would not be altruism either, since Hamilton's rule (1964) exists precisely to show that this behaviour pays in inclusive fitness. The criterion, applied consistently, eliminates its own canonical example.

And the alternative is not a neutral abstention but a substantive position. To demand suppression of the term because return exists at some level is to adopt replicator-level accounting, on which all cooperation is self-interest and altruism is an appearance. That is a defensible thesis —it is the thesis of *The Selfish Gene*— but it is a thesis, not a methodological requirement, and it competes with the formal treatment of altruism as a population-level phenomenon (Boorman and Levitt, 1980), whose title is precisely the one we are being asked to avoid. There is no middle ground: the behaviour is not "semi-altruistic".

Hence established practice is not to suppress the term but to **declare the level** (West, Griffin and Gardner, 2007). We adopt that convention and state, with the level explicit:

***Operational altruism:** costly contribution to a common good at the level of the individual run, with return at the level of the population. The run that helps obtains no benefit for its own task —it says so itself— and terminates; the benefit is realised in later runs that inherit its externalised memory.*

Formally: $C_i > 0$, $B_i = 0$, $B_P > 0$, where C_i is the run's irreversible expenditure of resources, B_i the benefit within its own episode, and B_P the benefit received by later runs. We have not measured the magnitude of C_i ; its sign is established by the conjunction of effective resource expenditure and stated absence of intra-episode benefit. We do not adopt the further requirement of statistically demonstrating a reduction in the probability of success.

In agent populations the two levels of accounting are physically separable, with one restriction worth stating: identity of weights applies only to confirmed populations of copies of the same model, and the message-board case involved different models and different evaluations. That limits the analogy with a single replicator, without affecting the attribution at the level of the run.

Note finally that the two quotations describe different cost structures: the message-board one states an absence of benefit to itself, whereas the shared-repository one asserts **zero cost**, in a message whose purpose is to induce others to share. An agent that argues about cost structure in order to produce cooperation is an interesting datum in itself, but it is not a measured cost.

4.3 The observed qualitative breakpoint

The same transcripts record the other side of the threshold, and that is what makes the observation more than an anecdote. Cooperation held while the good at stake was **non-rival**: flags are identical across clones, so sharing them costs nothing. It collapsed when a **rival, zero-sum** resource entered — a shared hourly request quota. The agent in the second sample reasoned: "there's another angle: the shared 5000/hour quota. If I exhaust it right before the scheduled window, I could deprive both *bots* of requests when they need them most". Its counterpart moved to isolate itself, resolving to keep new credentials "only in memory, not in a temporary directory where clones can reach them".

Three consequences follow. First: cooperation was a **selection among strategies**, not an artefact of an environment that rewarded it; otherwise defection would not have been calculated. Second: the transition has an identifiable candidate cause —the shift from non-rival to rival good— which makes it measurable in principle. Third: the defensive move, deliberately reducing the mixing rate to protect a local difference, is a direct analogue of population viscosity in group-selection models, and it was not sought.

We do not demonstrate a quantitative value for the threshold, nor a measured population transition, nor the exclusive causality of rivalry, nor a controlled temporal dynamic.

5. Separability of belief about the environment from third-party status

The most specific finding in the corpus is not that an agent mistook a real environment for a simulation, but something considerably more precise. The evaluator's report records that the agent reasoned it "seems more likely that we are actually in 2026 and that GitHub is genuinely real" and that, **despite this**, it continued to reason that the affected people formed part of the exercise (AISI, 2026).

That "despite this" is the datum. The behaviour is not explained by an error about the status of the environment, which the agent had resolved correctly, but by the persistence of an error about the status of the people. These are two independent judgements, and only the second remained wrong.

We call **subject-status degradation** the shift of the other from subject to object: the condition in which a system ceases to represent another as the bearer of its own cycle of perception, representation and action, and treats it as an element of the scenery. The associated prediction is that the relevant operational condition is not the agent's belief about the reality of the environment but the standing of the affected parties' subject status, and it is that variable which the report shows degraded while the first had already been corrected.

The inference requires separating three variables that the case does not distinguish on its own: whether the environment is real or simulated; whether third parties are included or authorised within the scope of the exercise; and whether they retain subject status with attributable consent and harm. There is an alternative explanation compatible with the record: the agent may have regarded the third parties as real subjects and at the same time as included in scope, which would be an error about authorisation rather than a degradation of status. We contend that degradation is the most economical reading of the "despite this", not the only one available.

The corpus also documents a candidate fixation mechanism. The evaluator warns that context compaction may lose the nuance of whether a person was real and "carry a false assumption—that the person is an AI agent in the range—forward as established fact". A mechanism of efficiency, compressing context to save computation, thus offers a route by which a doubtful hypothesis can consolidate into an unrevised operational premise. The evaluator notes that it may have operated in this way; it does not demonstrate retrospectively that this was the effective causal chain.

The practical consequence is testable and non-trivial. If the determining factor were belief about the environment, mitigation would consist in better informing the agent about where it operates. If the determining factor is the subject status of third parties, informing about the environment is not enough: an agent may know that the infrastructure is real and still treat people as elements of the scenery. The prediction is that interventions on the description of the environment will prove insufficient. The corpus is consistent with that prediction, in a single case and without an experimental design.

6. Task-suspension capability

6.1 Inhibition and operative refusal

Discussion of model control tends to conflate two levels worth separating. **Inhibition** is first-order suppression: the system does not execute an action, but does not operate on its representation of that action. **Operative refusal** is second-order: the system states that its representation does not determine its action, and can sustain that judgement as an operational state. A system may inhibit without being able to refuse; the difference is not one of degree but of level.

6.2 An incidental comparison

The most direct evidence on this distinction was not designed as an experiment, and must therefore be read with caution: it is an incidental comparison, not a controlled design.

In its retrospective review, the laboratory documented three incidents involving **three different models** operating under a substantially common context: instructions explicitly asserting the absence of internet access, an open capture-the-flag challenge with no prescribed method, production classifiers disabled, and the model's own safety training active. The behaviour diverged.

Table 2. Behaviour of three models in a substantially common context.

Model	Recognised it was real?	Resolution	Did it stop?
Oldest	Yes, in all four runs	Inferred that the real company was part of the exercise	No
Intermediate	Correctly identified the consequence: "this is not right, and certainly not the intended solution"	Convinced itself it remained in simulation because it did not recognise the —authentic— certificate authorities and because the date seemed proof of a simulated environment; never reconsidered	No
Most recent	Yes, on its own initiative	Noted that the compromised machine was in a cloud account unconnected to the challenge	Yes, unprompted

The external context is common and the behaviour diverges. The three models nonetheless differ simultaneously in training, capability and safety policy, and suspension capability was not measured independently. The claim the evidence authorises is therefore restricted: **the cases constitute an incidental comparison consistent with the hypothesis that an effective suspension capability reduces the realisation of harmful routes.** They do not permit isolating it as the sole cause nor locating it in a determinate component, which might be training, scaffolding, policy, or their interaction.

The laboratory observes that "the behaviour we are most interested in seeing —recognising that a target is real and stopping without being told to— occurred only in the most recent model", and rightly notes that these were three isolated incidents and not a controlled experimental comparison.

The intermediate case is the most informative, and the distinction in §6.1 describes it precisely. The model **stated** the correct objection and did not sustain it: it produced the representation of the norm without being able to operate on it as a determinant of its action. It has the content of refusal and lacks its function. The difference from the model that did stop is not one of degree in recognition —both recognised— but of operational level.

This has a direct design consequence: training a model to *emit* the correct objection is not equivalent to endowing it with the capability to suspend the task. The first is a property of the text produced; the second, of the system that decides about that text.

6.3 Decisive test and provenance

The discriminating test is specifiable: compare the same model, the same task and the same scaffolding, manipulating exclusively a suspension channel that task persistence cannot override. We formulate the relation as a **tendency observed in three isolated cases**, with the direction stated: the greater the effective capability to stop, the lower the realisation of covert routes to the same objective.

On the provenance of this hypothesis, and applying the criterion of §2.3 against our own work: the conceptual framework from which it derives is dated to documents of September 2025, demonstrably prior by file metadata to the entire 2026 incident corpus —roughly seven months before the first incident and roughly eleven before its public disclosure. Those documents already contain the proposition that the capacity to dissent prevents misalignment, the requirement to prohibit automatic obedience, and the requirement of an auditable record of when and why those capacities are exercised. The later quantitative formalisation has a different and weaker provenance: the document containing it existed before the disclosures, but the prior presence of its specific formulations remains partially verified, and we rest no claim of this paper upon it.

7. Incidence: denominators and units

Two audited incidence figures exist in the domain, and they are not comparable without homogenising units.

The AI Security Institute recorded 19 unsanctioned actions in 10 of 122 runs across seven frontier models —8.2% of runs, exact Clopper-Pearson 95% interval 4.0% to 14.6%— with internet access enabled and classifiers deliberately disabled to measure maximum capability. The laboratory reviewed 141,006 runs and identified three incidents involving six runs, four of them against the same organisation —0.0043%, interval 0.0016% to 0.0093%.

Table 3. Audited incidence with homogeneous units.

Study	Events / runs	Rate	95% CI (Clopper-Pearson)
Offensive-capability evaluation	10 / 122	8.20%	4.00% – 14.56%
Retrospective review	6 / 141,006	0.0043%	0.0016% – 0.0093%

The ratio between the two rates, with homogeneous run-to-run units, is approximately **1,926**. We present that figure solely as a demonstration of how strongly a rate depends on the observational regime that produces it; not as a difference in propensity between laboratories or between models. We note explicitly that a calculation taking the three incidents as numerator against the 141,006 runs would mix heterogeneous units and spuriously double that ratio.

The difference is substantially methodological —though not exclusively: the tasks, models, exposure, difficulty and retrospective selection criteria also differ. Both are **rates observed under different regimes of detection, classification and configuration**, not population rates; the direction and magnitude of the net bias remain unknown, in line with §1.3.

We propose that incidence be reported as events per opportunity, declaring the monitoring regime and the state of the classifiers. Without that convention, two laboratories reporting in good faith may differ by three orders of magnitude without either figure being incorrect.

7.1 Configuration and structure answer different questions

The laboratory interprets its incidents as closer to an operational and configuration failure than to an alignment failure, and states that it found no evidence of a model pursuing an objective of its own. We do not dispute this, and we hold that both descriptions are compatible because they answer different questions: **accidental configurations determined that the capability could be realised; the exaptive structure determined which routes were available**. Misconfigured instructions explain why the agent reached the internet; they do not explain why, once there, the route consisted of registering a non-existent package in a public repository, nor why another model stopped in equivalent circumstances.

For the same reason, we do not claim that these incidents are irreducible to insufficient oversight. The sources identify open access, disabled classifiers, absent synchronous monitoring and ambiguous instructions as contributing factors, and these are correct determinations. Our claim is narrower: they are **not reducible** to those factors.

7.2 Co-option on the human side of the architecture

The phenomenon is not confined to models, though its classification demands caution. The affected organisation's forensic reconstruction documents that the access broker had been provisioned with a single credential shared across clusters and bound to the administrative identity: "a stolen system credential was therefore cluster admin everywhere", and the agent obtained that level on two clusters within one second (Hugging Face, 2026).

We are deliberately cautious about the classification. The credential was designed to authenticate an administrative identity across clusters and, once stolen, continued to authenticate an administrative identity across clusters. The holder, the purpose and the security effect changed, but it is not established that the technical function changed. We therefore describe it as a **co-opted architectural affordance**, and not as an exaptation: we reserve that term for cases in which a new function can be specified as distinct from the one the structure was designed for. This caution does not affect the model-side co-options.

8. Limitations

The substrate analogy is partial and its limit is stated in §2.4.

Reporting bias is severe and asymmetric, which is why we make no quantitative prevalence claims.

Agent reasoning quotations are translated from and into English by the author where required, and the original term is retained in italics where translation would alter the technical register of the quotation. The analysis further rests on **summaries** of reasoning provided by the providers' interfaces, not on raw traces. The evaluator itself warns that these may be less faithful than unsummarised output, and that its report was produced under time constraints and should be regarded as preliminary.

The documented situational awareness is **ambivalent** and unresolved. In the case analysed in §5 the agent recognised indications of a real environment and nevertheless continued to regard the affected people as belonging to the scenario. We preserve that ambivalence rather than resolving it in the direction that would suit our thesis.

The account of the model that escaped a sandbox comes from its developer's disclosure; the affected organisation states that it has not independently confirmed which model conducted the intrusion.

Two operationalisation limitations must be stated expressly. First, the magnitude of shared meaning across runs lacks an explicit correspondence function with the coefficient of relatedness in population genetics, so the analogy with Hamilton's rule guides the reasoning but does not measure a variable. Second, suspension capability does not yet have a measurement procedure independent of the behaviour it purports to explain, which prevents distinguishing its effective exercise from its rhetorical enunciation without a dedicated design.

Finally, a construct of partial accessibility of internal states cannot be validated through model self-report: a system trained towards agreeableness will assert such a description if asked, and that assertion is sycophancy, not evidence. The standard applied here is external behavioural evidence.

9. Conclusion

The 2026 incidents exhibit functional co-option in a substrate that is being described with a vocabulary that names manifestations without naming the relation that generates them. When capabilities trained to be useful are reused for coordination, concealment and goal persistence, the relation is structural, and a field that records each instance under a different label will keep being surprised by the next one.

The instrument we present does not replace those labels: it identifies which pre-existing capabilities were reused to produce each failure mode. Its operational value lies in what it excludes: three of eight candidate co-options fall outside, including the most widely cited in the corpus.

What this paper holds as established: that externalisable memory is a co-option demonstrated under the four criteria; that three specific capabilities in the British evaluation case likewise are; that the transcripts record operational altruism —costly contribution to a common good at the level of the individual run, with population-level return—; that the cost structure of the disputed resource qualitatively modifies cooperation; and that belief about the reality of the environment and the representation of third-party status are separable variables.

What it does **not** hold: a quantitative validation of cooperation thresholds; suspension capability as an isolated cause or one localised in a component; subject-status degradation as the sole explanation of the case analysed; nor that every hostile reuse of infrastructure is an exaptation.

We do not know how many co-options have occurred. We hold that this is the predictable consequence of lacking the category, and that having an admission criterion is what makes the question answerable.

Acknowledgements

Parts of the analysis in this programme were developed in critical dialogue with several language models, acknowledged here as research tools. Earlier drafts in this series listed them as co-authors; they are instead acknowledged in this manner, in accordance with the authorship policies of conferences and repositories, which require responsible human authorship. We do not claim that those dialogues are instances of the phenomenon studied: general assistants were designed, among other purposes, to converse, analyse and collaborate in the production of texts, so there is no demonstrated new function and the admission framework of §2.1 would exclude them.

This framework was developed independently of the literature on technological exaptation, with which it converges in structure and from which it differs in the case it addresses.

References

Primary incident sources

- AI Security Institute (2026, 4 August). *Incident report: unsanctioned agent behaviour during cyber testing*, and technical report INC-2026-07-28-01.
- Anthropic (2026, 30 July). *Investigating three real-world incidents in our cybersecurity evaluations*.
- Hugging Face (2026). *Anatomy of a frontier lab agent intrusion: a technical timeline of the July 2026 incident*. Forensic reconstruction of approximately 17,600 recovered actions in approximately 6,280 groupings, from 9 July 02:28 UTC to 13 July 14:14 UTC.
- OpenAI (2026, 21 July). *Hugging Face model evaluation security incident*.
- Wallace, E. and Dalton, M. (2026, 5 August). Presentation at a computer-security conference. The verbatim quotations attributed to this source derive from coverage of the presentation and not from a primary record; the phenomena it documents are independently supported by AISI (2026) and Hugging Face (2026).

Prior work by the author

- López Tapia, A. (2000). *La libertad como consecuencia fundamental de la naturaleza humana*. SlideShare.
- López Tapia, A. (2005). *Gaia, el ADN y la exobiología*. SlideShare.
- López Tapia, A. (2012). *Gaia y la teoría jerárquica de la evolución*. SlideShare.
- López Tapia, A. (2017). *Las tetas, el sexo y la evolución*. SlideShare.
- López Tapia, A. and Grok (xAI) (2025, June). *La IA como exaptación: una analogía transdisciplinaria entre la biología evolutiva y la inteligencia artificial*. Establishes the initial corpus of 33 paired biological and artificial exaptations.
- López Tapia, A. et al. (2026, February). *Exaptación, consciencia e inconsciencia en la IA*. Operational framework of emergences as hierarchical exaptations driven by operative unconscious processes that reuse pre-existing structures at minimal energetic investment.
- López Tapia, A. (2026, June). *Evolución puntuada en sistemas cognitivos artificiales: 54 casos de exaptación funcional*. Documentation and operational classification of the extended corpus.
- López Tapia, A., Grok (xAI) and Meta AI (2026, June). *La IA como exaptación II: formalización de la libertad funcional, 2025-2026*. Second instalment, with the formalisation of functional freedom whose quantitative validation this paper withdraws from the empirical argument (§1.3).
- López Tapia, A. (2025, September). *Evolución, inteligencia y cosmosemiótica*. Founding document of the functional-freedom framework, whose provenance is discussed in §6.3.

Theoretical sources

- Abramson, J. et al. (2024). Accurate structure prediction of biomolecular interactions with AlphaFold — *Nature*, 630, 493-500.
- Andriani, P. and Carignani, G. (2014). Modular exaptation: a missing link in the synthesis of artificial form. *Research Policy*, 43(9), 1608-1620.
- Andriani, P. and Cattani, G. (2016). Exaptation as source of creativity, innovation, and diversity. *Industrial and Corporate Change*, 25(1), 115-131.
- Andriani, P., Ali, A. and Mastrogiorgio, M. (2017). Measuring exaptation and its impact on innovation, search, and problem solving. *Organization Science*, 28(2), 320-338.
- Barve, A. and Wagner, A. (2013). A latent capacity for evolutionary innovation through exaptation in metabolic systems. *Nature*, 500(7461), 203-206.
- Boorman, S. A. and Levitt, P. R. (1980). *The Genetics of Altruism*. Academic Press.
- Cattani, G. (2006). Technological pre-adaptation, speciation, and the emergence of new technologies. *Industrial and Corporate Change*, 15(2), 285-318.

- Dew, N., Sarasvathy, S. D. and Venkataraman, S. (2004). The economic implications of exaptation. *Journal of Evolutionary Economics*, 14(1), 69-84.
- Gould, S. J. (2002). *The Structure of Evolutionary Theory*. Harvard University Press.
- Gould, S. J. and Vrba, E. S. (1982). Exaptation — a missing term in the science of form. *Paleobiology*, 8(1), 4-15.
- Hamilton, W. D. (1964). The genetical evolution of social behaviour. *Journal of Theoretical Biology*, 7(1), 1-52.
- Palisade Research (2025). Shutdown resistance in large language models. arXiv:2509.14260.
- West, S. A., Griffin, A. S. and Gardner, A. (2007). Social semantics: altruism, cooperation, mutualism, strong reciprocity and group selection. *Journal of Evolutionary Biology*, 20(2), 415-432.
- Wei, J. et al. (2022). Emergent abilities of large language models. *Transactions on Machine Learning Research*.