

# Generalized FCPT (FCPT-G)

## Theory of loss of cognitive autonomy in human, social and artificial systems

**Author:** Adrian (Adi) Stan

**ORCID:** <https://orcid.org/0009-0003-1457-5155>

**SSRN:** <https://ssrn.com/author=7778480>

**Date:** June 19, 2026

**Keywords:** FCPT, False Cognitive Power Transfer, Maslow7F, DDC, Information Gravity Theory (IGT), cognitive autonomy, cognitive externalization, distributed cognition.

### Summary (Abstract)

This paper extends and formalizes the mechanism of False Cognitive Power Transfer (FCPT) from its initial formulation (human-AI interaction) to a generalized theory of cognitive autonomy loss, applicable to human, social, institutional, and artificial systems. The paper rigorously distinguishes between functional cognitive offloading *and* the pathological process of misattribution of externally generated performance. A distribution model of autonomy degradation is proposed and testable predictions are formulated for empirical validation of the phenomenon. The generalized framework (FCPT-G) is complemented by an operational extension of Information Gravity Theory (IGT-G), which describes the information fields in which these false transfers stabilize and become persistent. Together, they form a unitary architecture for the analysis of cognitive autonomy as a dynamic and context- *dependent phenomenon*.

### 1. Introduction

The accelerated development of artificial intelligence systems and their deep integration into everyday cognitive flows have generated a **significant structural change** in the distribution of thought processes between the individual and external systems. The phenomenon of cognitive externalization is not new; the literature has extensively described the use of the environment and other agents as extensions of human cognition, within the paradigms of distributed cognition (Hutchins, 1995) and extended mind (Clark & Chalmers, 1998). Concepts such as shared memory (Wegner, 1987) and cognitive offloading (Risko & Gilbert, 2016) have highlighted the role of external systems in supporting individual performance.

However, in the context of generative AI systems, these mechanisms take on a critical dimension. Users tend to exhibit an automation bias (Parasuraman & Riley, 1997) and overestimate the capacity of autonomous systems, uncritically delegating decision-making or evaluative processes. The intensive use of these systems can lead to a reduction in cognitive engagement and a decrease in autonomous performance in specific tasks.

Against this background, the work *False Cognitive Power Transfer (FCPT)* (Stan, 2025) identified a specific mechanism of human-AI interaction, through which the user erroneously attributes his own cognitive capacity to the performance generated by the external system. This phenomenon produces a form of "apparent autonomy", characterized by the illusion of competence and functional dependence.

This paper develops this conceptual framework through generalization. The central argument is that the FCPT mechanism is not specific exclusively to human-AI interaction, but represents a general pattern of cognitive behavior, observable in multiple types of interactions: between individuals, between individual and group, between individual and institution, and between artificial systems.

FCPT is thus reformulated as a general mechanism of **cognitive autonomy misattribution**, whereby a cognitive or semi-cognitive system confuses the performance or coherence generated by an external system with its own internal capacity. This reformulation allows for the integration of phenomena such as social conformism, authority dependence, institutional capture, and model distillation *into a unified explanatory framework*.

Note: As this paper was being finalized, Nature (Lenharo, 2026) published a journalistic synthesis of emerging evidence on AI-assisted skill degradation in medicine and software engineering, drawing on the same two independent studies integrated here (Budzyń et al., 2025; Shen & Tamkin, 2026). This convergence confirms at the reporting level that the phenomenon occurs; FCPT-G supplies the framework explaining why and how, identifying attribution error as the shared mechanism.

## 2. Definition and conceptual delimitation of the FCPT-G

The concept of **False Cognitive Power Transfer (FCPT)**, originally formulated to describe an attribution error in human-AI interaction, is expanded within **the FCPT-G framework** to apply the framework to any information processing system, be it human, artificial, or organizational.

### 2.1. Definition

**FCPT-G** is the mechanism by which a cognitive or semi-cognitive system erroneously attributes to its own autonomy a capacity or performance that is, in reality, partially or totally generated by an external system.

**Note on semi-cognitive systems:** By “semi-cognitive system” we mean those structures that possess functional coherence and processing logic, but may lack meta-reflection (e.g., a bureaucratic structure, a modeling algorithm, or a social group). In these cases, FCPT manifests itself by internalizing procedural success as evidence of the system’s own “intelligence”.

This definition involves three structural pillars:

1. **External source:** An agent or system (AI, expert, group, institution) that provides superior or faster output.
2. **Attribution error:** Failure of the receiving system to delineate external contribution from internal competence.
3. **Apparent autonomy:** A state in which the observable performance of the system remains high, but masks the actual degradation of independent processing capacity.

### 2.2. Delimitations (what FCPT is NOT)

The FCPT-G must be clearly distinguished from functional forms of external collaboration:

- **Healthy Cognitive Offloading:** Conscious use of external resources (notes, AI) for efficiency, while maintaining the ability to critically evaluate.
- **Distributed cognition:** Distributing processes among agents where each one's role remains distinct and recognized.
- **Shared memory:** Delegating information storage to external sources, maintaining awareness of the source.

**The specific difference:** In FCPT, the boundary between tool and self dissolves. The system ceases to perceive the resource as external, internalizing its success as proof of its own competence.

### 2.3. Conditions of occurrence

FCPT-G is installed when the following are simultaneously fulfilled:

1. Access to a high-performance external system.
2. Functional delegation of cognitive processes (decision, analysis).
3. Absence of a demarcation protocol (lack of ROGO).
4. Enough repetition cycles to stabilize the erroneous perception.

Research by Maguire et al., 2006 and Woollett & Maguire, 2011, demonstrates that self-effort navigation produces measurable structural changes: London taxi drivers who internalize “*The Knowledge*” develop increased volume in the posterior hippocampus, a change absent in students who did not pass the exam, as well as in the control group. Bus drivers, who travel the same metropolis daily, but on fixed routes, without orientation and re-routing decisions in atypical conditions, do not show this development. The difference is not related to exposure, but to active processing: **the cognitive structure is built through repeated decisions**, not through simple presence in the environment. **FCPT-G postulates the symmetry of this process:** when the external system takes over the orientation decision, the agent functionally approaches the condition of the bus driver: present on the route, but absent from its construction. A driver who uses an assisted navigation system daily has (1) access to a high-performance external system. On routes that he might know, he (2) delegates to the system not only **the memorization** of the route, but **also the** rerouting decision. In the absence of a moment of independent verification (“which route would I have chosen myself?”) the demarcation protocol is missing (3). And the daily repetition of this delegation (4) stabilizes the perception that the driver “knows the city”,

when in fact the city is known to the system. The perceived competence remains intact; **the real capacity** for orientation and autonomous decision-making **erodes**.

Navigation systems with real-time optimization (e.g., traffic alerts) illustrate an aggravated form of the mechanism: they combine a legitimate augmentation, of logistical optimization and efficiency (through information inaccessible to the agent, such as real-time incidents) with a delegation of orientation and routing decision. The conscious justification (route efficiency) masks the problematic component (the teaching of internalizable orientation), making the absence of the demarcation protocol not only probable, but **structurally invisible**: the agent does not perceive that he has delegated, because he delegates under the pretext of real optimization.

## 2.4. L1 - L3 as dynamic distributions of cognitive autonomy

In the initial formulation of the FCPT framework, levels L1, L2 and L3 were introduced as distinct functional profiles of the relationship to AI systems and the cognitive complexity of the task. L1 describes the user who consumes the output without understanding its mechanisms or limits; L2 describes the user who operates and manages AI systems, but without full conceptual control over the architecture and implications; and L3 describes the actor capable of designing, coordinating and integrating complex systems, while simultaneously understanding their limits and systemic consequences. This classification remains valid, but needs to be reinterpreted: **the levels do not function as fixed**, exclusive and mutually exclusive steps, but as **dynamic distributions** of cognitive autonomy within a system.

In this perspective, an individual or an organization does not **absolutely “occupy” a single level**, but **simultaneously manifests different proportions of each level**, proportions that change depending on context, experience, pressure, feedback and available cognitive reserve. L1, L2 and L3 must therefore be understood as functional weights, not as discrete states. A system may have a dominant L3 in a context of strategic analysis, but may temporarily drop towards L2 or even L1 in contexts of high pressure, operational dependence or uncritical delegation. Conversely, an actor predominantly in L1 may manifest, punctually, L2 or L3-type behaviors, without this indicating a stable reconfiguration of its global profile.

This reinterpretation is important for FCPT because the mechanism not only produces local attribution errors, but also a **redistribution of internal weights** between levels. When the system attributes external performance to its own competence, it not only confuses the source of the result, but gradually changes its operating structure: the delegation component increases, the need for internal verification decreases, and the weight of L3-type processes decreases in favor of more passive forms of operation. In FCPT terms, this means that the false transfer of cognitive power is not just an epistemic error, but a mechanism for the **distributional reconfiguration** of autonomy.

To avoid a rigid interpretation, it is useful to consider these levels as **percentage variables**. At one point, a system may have, for example, 60% L3, 30% L2, and 10% L1 functional features; in another context, the same structure may become L2 dominant, with a reduction in the capacity for critical validation and an increase in dependence on external tools. This distributional representation better explains why FCPT does not appear as a sudden collapse, but as a **gradual drift** of the cognitive profile towards forms of apparent autonomy and reduced real autonomy.

From this perspective, the role of cognitive reserves and the ROGO principle becomes central. A system capable of maintaining an active cognitive reserve can preserve a higher share of L3 functioning, even under conditions of intensive use of external tools. In contrast, a system that does not check, interrogate and delimit the source of performance will suffer a progressive shift towards L2 and L1, even if its apparent output remains high. Therefore, **the FCPT levels do not simply describe a hierarchy of competence, but a dynamic of balance between autonomy, delegation and critical validation**.

This **distributional representation FCPT** allows for a more realistic description of transitions between profiles, without assuming absolute breaks. In practice, most systems are not pure L1, pure L2, or pure L3, but rather variable mixtures of these. Therefore, FCPT should be read as a model of **continuous redistribution** of cognitive capacity, in which dominant profiles are produced by the history of interactions, the quality of feedback, and the degree to which the system remains able to distinguish between tool and self.

On this basis, the FCPT mechanism can be analyzed as a process of shifting between functional regimes of cognitive autonomy.

### 3. Operating mechanism: from Exposure to Stabilization

FCPT-G develops through a sequence of four phases:

#### Phase 1: Exposure

The system comes into contact with a source perceived as more efficient (an AI, an opinion leader, a procedures manual).

#### Phase 2: Functional Delegation (Efficiency Principle)

For reasons of systemic economy, the system begins to transfer tasks towards the source.

This movement follows **the path of least resistance**:

- **Biological**: The brain seeks to minimize metabolic energy consumption.
- **Economic/Systemic**: A company or algorithm seeks to reduce operational or processing costs.
- At this stage, delegation is justified by efficiency, but it sets the stage for dependency.

#### Phase 3: Misattribution

The external output is integrated into the final result without delimitation. Success is felt as the product of one's own reasoning. The **"illusion of competence"** appears.

#### Phase 4: Stabilization

Repetition leads to the atrophy of the ability to generate alternatives. Dependence becomes the default state of the system.

#### 3.1. Dissociation of Performance - Autonomy

The critical element is the paradoxical dissociation: observable performance increases (due to external support), while real autonomy progressively decreases. This masking makes the phenomenon difficult to detect from within the system.

#### 3.2. The spiral of addiction (Positive Feedback)

FCPT-G tends to become a self-sustaining process through a positive feedback dynamic:

1. Delegating a task reduces the exercise of one's own critical functions.
2. The resulting atrophy increases the "energetic cost" (effort) required to resume internal processing or to critically validate the source.
3. As verification becomes more "expensive", the system tends to delegate even more, accelerating the loss of autonomy.

The system thus enters a spiral where the degradation of validation capacity forces new levels of dependency, transforming FCPT from a point error into a persistent systemic state.

### 4. FCPT as a distribution phenomenon

FCPT-G does not operate as a binary mechanism (present/absent), but as a continuous process of redistribution of autonomy between the internal and external system. In this section, FCPT is formalized as a distributional phenomenon, dependent on stability thresholds and the dynamics of systemic equilibria.

#### 4.1. Non-binary FCPT

A cognitive system is never in absolute autonomy or total dependence. Autonomy is distributed between internal processes (analysis, evaluation, own decision) and externalized processes (delegated to external systems). Thus, FCPT should be understood as a continuous variable reflecting the proportion of externalized cognitive power, but erroneously attributed as internal.

#### 4.2. Distribution model

The distribution of cognitive autonomy over time can be modeled by the following phenomenological expression:

$$A(t) = A_0 * e^{(-k * E(t))}$$

where:

- **A(t)** = level of cognitive autonomy at time t.
- **A<sub>0</sub>** = initial level of autonomy.

- $e$  = Euler's constant.
- $k$  = transfer coefficient (intensity of FCPT). This is not a fixed constant, but a context-dependent variable (time pressure, trust in the source, social environment). A high  $k$  indicates increased permeability of the cognitive boundary and a rapid loss of autonomy.
- $E(t)$  = cumulative frequency of non-critical cognitive delegations.

We note that this expression constitutes an **indicative phenomenological model**, intended to formalize the direction and dynamics of the process, and not an invariant deterministic law. The exponential form of this model is not arbitrary, and has empirical precedent in the literature on the decline of knowledge and skills.

The forgetting curve (Ebbinghaus, 1885) demonstrated that the loss of unmaintained information follows an exponential pattern, and Arthur et al.'s (1998) meta-analysis of 189 data points from 53 studies confirmed that the decline in proficiency through disuse is systematic and, crucially for the present model, **modulated by the degree of initial mastery and the type of task**. This result provides empirical support for the variable nature of the  $k$  coefficient: the L1–L3 levels are not arbitrary values, but rather reflect the moderators identified in the *skill decay* literature (degree of overlearning, cognitive vs. procedural nature of the task). It should be noted that the classical literature describes decay by disuse, while FCPT-G describes decay by **delegation**, which adds the feedback loop (section 3.2) absent in passive forgetting. FCPT-G thus borrows *the form of* documented decay, adding the **self-** accelerating **specificity** of cognitive transfer.

Definition of non-critical delegation ( $E$ ):

**$E(t)$  = sum of  $d_i$ , for  $i = 1$  to  $t$ .**

where the decision variable  $d_i$  takes the value:

- $d_i = 1$ : if the system accepts an external output without independent verification, committing an attribution error (confusing the tool's success with its own competence).
- $d_i = 0$ : if the decision is generated internally or if the external result undergoes an active check (autonomy preservation).

#### 4.3. Functional thresholds and balance stability

Although the degradation process is continuous, the system exhibits critical thresholds where the nature of the equilibrium changes fundamentally. These thresholds mark the transition from a stable to an unstable system:

1. **Dominant autonomy zone ( $A(t) > \tau$ ):** The system is in a stable equilibrium. The internal processing capacity is sufficient to maintain a clear demarcation between its own and external contributions. The system is generally linear or convex in its response to errors.
2. **Transition zone ( $A(t) \approx \tau$ ):** The critical point where the equilibrium becomes indifferent or unstable. It is the threshold where the system risks moving from a self-correcting dynamic to one of cognitive collapse.
3. **Dominant dependency zone ( $A(t) < \tau$ ):** The system enters a zone of concavity, where dependency becomes structural. At this stage, the capacity to generate independent alternatives is so reduced that the system can no longer spontaneously initiate a recalibration.

$\tau$  (**tau**) thus represents the threshold of the system's phase change: below this value, the system loses its ability to resist the "information gravity" of the external source.

#### 4.4. Energy recovery and cost (ROGO Effect)

Cognitive autonomy  $A(t)$  can be restored by ROGO (active interrogation) mechanisms. However, unlike passive degradation (which consumes little energy), recovery by ROGO is an active process, requiring **additional** metabolic and cognitive **energy expenditure**.

**The effort** required to push  $A(t)$  back above the threshold  $\tau$  is **much greater** than the energy that would have been required to **maintain the initial autonomy**. This "repair cost" explains why systems in stabilized FCPT tend to remain in that state: the path of least resistance favors dependency, while autonomy becomes energetically "expensive".

#### 4.5. Cognitive reserve

Cognitive reserve functions as a systemic shock absorber. A system with a large reserve will have a lower " $k$ " coefficient and will require a much larger number of non-critical delegations ( $E$ ) before reaching the critical threshold  $\tau$ .

## 4.6. Cognitive deficit as an area between trajectories

The  $A(t)$  model described above characterizes the evolution of the autonomy of a single system. However, the truly relevant quantity for FCPT is not the absolute level of autonomy at a given time, but **the loss suffered by delegation**, i.e. the difference from the trajectory that the same system would have followed in the absence of non-critical delegation.

We consider two trajectories starting from the same initial autonomy  $A_0$ :

- a **sustained trajectory** (ROGO), in which delegation is accompanied by critical verification, characterized by a small coefficient  $k_2$ ;
- an **eroded trajectory** (FCPT), in which delegation is non-critical, characterized by a larger coefficient  $k_1$  ( $k_1 > k_2$ ).

**The cognitive deficit accumulated** by FCPT-G is not the distance between the two trajectories at a single moment, but **the area between them** over the entire exposure interval (the integrated loss of autonomy over time).  $D$  therefore has the autonomy-time dimension, expressing not a state, but a cumulative loss:

$$D = \text{integral, from 0 to T, of } [A_{\text{sustained}}(t) - A_{\text{eroded}}(t)] dt$$

Consequences:

1. The deficit is not instantly reversible by stopping delegation: **the already accumulated area represents untrained**, not just unpracticed, competence. This corresponds to the phenomenon of “*cognitive debt*” empirically observed by Kosmyna et al. (2025), where the deficit in neural engagement persists after AI use ceases.
2. **area** formulation corresponds directly to the dominant experimental design in the literature: parallel group studies (one group with AI, one group without) measure precisely the distance **between the two trajectories**. The 17% reported by Shen & Tamkin (2026) represents, in these terms, **the deficit** measured at testing, namely a cross-section through the area  $D$ .
3. **The two trajectories are not symmetrical**. The sustained trajectory tends towards saturation (autonomy stabilizes, like the learning curve), while the eroded trajectory is self-accelerating through the feedback loop (section 3.2): the larger the area, the smaller the capacity to reduce it. The  $D$  deficit thus tends to widen faster than it would widen by simply summing the delegations, a property that thus distinguishes FCPT from passive forgetting.

## 5. Illustrative modeling of the coefficient $k$

Methodological note. The L1 profile is partially anchored in published empirical data (Shaw & Nave, 2026), from which we take the decrease in deliberative accuracy; applying the model  $A(t) = A_0 \cdot e^{(-k \cdot E)}$  to these values produces a tentative estimate of  $k$ , given that the variable  $E$  (cumulative frequency of delegations) is not directly measured by the cited study, but approximated. The  $k$  values for the L2 and L3 profiles, as well as the telemetry methodology described below, constitute an illustrative framework and a proposed research protocol, not an empirical collection. Their empirical validation remains an open research direction (see section 12).

### 5.1. Data collection methodology: Structural error propagation

The main point of the empirical validation is the monitoring of the decision variable  $d_i$  through the indicator called “Propagation of Structural Errors”. The methodology involves the following steps:

1. **Identifying sources of error**: Isolating a set of 50 unique algorithmic hallucinations and design biases documented in dominant LLM models in Q1 2026. These are errors that do not occur in standard human logic, acting as “radioactive markers” in the code.
2. **Monitoring public repositories**: By analyzing commit stream telemetry, the frequency with which these markers appear in the code written by users under AI assistance is tracked.
3. **Calculating  $d_i$** : Whenever a specific model error is replicated identically in the user code without being corrected, the variable  $d_i$  takes the value 1 (non-critical delegation). When the error is detected and rewritten,  $d_i$  takes the value 0 (autonomy preservation).

## 5.2. Mathematical Model and Derivation of k Values

The exponential decay model is defined by the formula:

$$A(t) = A_0 * e^{(-k * E(t))}$$

To isolate k from telemetry data, the linearized form is used:

$$\ln(A(t) / A_0) = -k * E(t)$$

where k represents the negative slope of the regression. Based on this model, the trajectories for the three fundamental profiles are established:

### 1. Passenger (FCPT Level L1): k = 0.085 (Severe Atrophy)

- **Data source:** The indicative estimate for L1 is derived exclusively from data published by Shaw & Nave (2026); the correlation with commit telemetry described in 5.1 remains a proposed validation step.
- **Reference values:** Decrease in deliberative accuracy from  $A_0 = 45.8\%$  to  $A(t) = 31.5\%$ .
- **Calculation:**  $\ln(31.5 / 45.8) = \ln(0.687) \approx -0.375$ . Compared to the average exposure sequence, this results in  $k = 0.085$ , for an assumed exposure value E.
- **Logic:** At this level, the user accepts the AI output as a necessity for task survival, leading to functional loss of autonomy in approximately 35 non-critical delegation steps - the point at which A(t) drops below the critical threshold  $\tau$  (tau).

### 2. Operator (FCPT Level L2): illustrative, k = 0.015 (Slow structural degradation)

- **Data source:** Longitudinal analysis of data from the "Vibe Coding" paradigm (Q1 2026).
- **k** has a guideline value, not empirically derived. The L2 profile is modeled to reflect a slow degradation: the user maintains production speed but erodes independent synthesis capacity. The k value is chosen illustratively to position L2 between severe atrophy (L1) and preserved sovereignty (L3); it does not represent a measurement.
- **Logic:** The operator believes he is in control, but using the AI as the ultimate validation instance "fluidizes" his own semantic inertia, making the system vulnerable to Out-of-Distribution errors. "Fluidization" is the process by which a cognitive structure loses its "solid" properties (resistant, sovereign) and acquires "fluid" properties (which takes the "shape" of the container/AI).

### 3. Architect (FCPT Level L3): k = 0.002 (Sovereignty preserved)

- **Data source:** Robustness benchmarks through self-distillation and critical interfacing (Q1 2026).
- **Reference values:** Autonomy degradation below 1% ( $A(t) / A_0 = 0.99$ ).
- **Calculation:**  $\ln(0.99) \approx -0.01$ , generating a  $k = 0.002$ .
- **Logic:** k tends to zero because the system uses the interaction with the AI to refine its own semantic mass (through the ROGO principle), not to replace it.

## 5.3. Operationalization of process sensors: methodology and values

The following sensors are proposed as indirect autonomy measurement tools. The numerical thresholds associated with each are illustrative values, formulated as operational hypotheses to be empirically calibrated, not results of a collection carried out:

1. **Delta\_t sensor (Acceptance Latency):** Measures the time between the AI suggestion being displayed and its acceptance.
  - **Passenger (L1):** delta\_t below 1.5 seconds (reflex, Pavlovian reaction). Time is insufficient to read and understand more than 3-5 lines of code.
  - **Operator (L2):** delta\_t between 4 and 8 seconds. Indicates a shallow reading ("vibe check"), but without an in-depth analysis.
  - **Architect (L3):** delta\_t over 25 seconds. Indicates the ROGO process is running: the user queries or modifies the suggestion before integrating it.

2. **N<sub>alt</sub> sensor (Semantic Diversity Indicator):** Measures the narrowing of the semantic footprint by cosine similarity (using a BERT-type Transformer model) to the source model.
  - **Passenger (L1):** Similarity above 0.96. Own identity was "sucked" by the model.
  - **Operator (L2):** Similarity between 0.85 and 0.90. The own style becomes a "sub-set" of the AI style.
  - **Architect (L3):** Similarity below 0.70. The user maintains a divergent logical structure, which protects their Semantic Inertia.
3. **Sensor R<sub>ext</sub> / int (Interrogation vs. Generation Ratio):**
  - **Passenger (L1):** 1/20 (one question every 20 generations).
  - **Operator (L2):** 1/5.
  - **Architect (L3):** 3/1 (three queries for each unit supported).

**Conclusion:** This sensor framework is proposed as a diagnostic tool, not as a set of measurements performed. The threshold values indicated above ( $\Delta_t$ ,  $N_{alt}$ ,  $R_{ext} / int$ ) are indicative and illustrative, theoretically derived, not empirically collected. Once validated by an independent collection protocol, such a framework *would allow* early detection: if a system were to present a decreasing  $\Delta_t$  simultaneously with an increasing  $N_{alt}$  (increasingly source-aligned style), one could anticipate drifting towards the Passenger state (L1) before the subject becomes aware of the loss of autonomy. Empirical validation of these thresholds remains an open direction (see section 12).

## 6. Relationship with the Maslow7F model

To understand the systemic impact of FCPT-G, it is necessary to integrate it into the **Maslow7F architecture** (Stan, 2025), which describes the fractal distribution of needs and motivations in complex systems.

### 6.1. FCPT as a short-circuiting mechanism

FCPT-G functions as a "short-circuit" in Maslow's hierarchy7F. The system attempts to obtain the benefits of the higher levels (L4 - Status/Esteem, L5 - Self-actualization) using borrowed performance, without going through the processing and maturation effort required for the lower levels.

This is a **GIGO**-type process: because the cognitive input is external and unassimilated, the output in terms of real autonomy is zero. The system projects an image of L5, but remains functionally stuck in L1-L3.

### 6.2. The critical role of L3 (Belonging and Identity)

L3 functions as an identity anchor and a protective barrier for the system. A healthy L3 boundary allows for the use of external tools without confusing the source with the self.

In FCPT-G, this boundary is dissolved. The system defines its identity through the prism of the external source (FCPT-HS: the group; FCPT-HAI: artificial intelligence). When L3 is "captured", the system automatically adopts any external output as its own emanation.

### 6.3. Fragility and Fall Mechanism

When a system "borrows" a higher level through FCPT, its position becomes extremely fragile. Any system error or interruption of access to the external source causes a systemic collapse. According to Maslow7F (Stan, 2025), the system does not go down a single level, but suffers a **free fall** to the real level from which the "falsehood" began. Due to the atrophy suffered during the period of dependence, the system can collapse even lower than its initial level of autonomy.

## 7. Relationship with DDC (Dominant Dissonant Chord)

If **Maslow7F** provides the architecture, **DDC (Dominant Dissonant Chord)** represents the element that determines the dominant functional state of the system.

### 7.1. DDC as a breaking point

The DDC represents the "string" with the highest dissonance in the system - the point from which instability propagates. FCPT-G acts directly on this string, forcing the system to stabilize in a configuration of reduced autonomy.

### 7.2. Repositioning of DDC through FCPT

FCPT-G shifts the weight of the system towards the areas of external validation (L4) and procedural safety (L2), blocking access to the levels of meta-reflection (L7). The DDC thus becomes an anchor that keeps the system in a state of "masked dissonance":

- The system experiences a discrepancy between observed performance (high) and internal competence (low).
- To avoid the cognitive pain of recognizing this discrepancy, the DDC stabilizes the system in addition, making recalibration increasingly difficult.

### 7.3. Resistance to change

Once the DDC has stabilized in a zone of dependence ( $A(t) < \tau$ ), the system becomes immune to corrective information. Recalibration would require a much greater energetic effort to "tune" the system's dissonant chord again, an effort that the atrophied system is no longer able to sustain on its own.

## 8. ROGO ERGO EMERGO: The recovery mechanism

If FCPT-G describes the mechanism of loss or erroneous redistribution of autonomy, the **ROGO ERGO EMERGO principle** ("I ask, therefore I become") is formulated as a mechanism for rebalancing the system.

### 8.1. Definition and function

ROGO is the process by which a cognitive system reactivates its autonomy through active interrogation, oriented towards:

1. Delimiting the source of information (What is mine? What is external?).
2. Checking internal consistency.
3. Critical evaluation of external output.

### 8.2. ROGO thermodynamics: energy and entropy

Recovering autonomy is not a spontaneous process, but one that requires a surplus of energy:

- **Systems below the  $\tau$  threshold:** The internal energy of these systems is too low to initiate ROGO. They require an **external stimulus** (an energy/information surplus or a reality "shock") to induce the force necessary to break away from the "gravity" of the external source.
- **Entropy Reduction:** ROGO works as a process of decreasing internal informational entropy. When the AI or external source is used as a "sparring partner", the effort to validate and query the output forces the system to create internal structure (order). In contrast, using the source as an "oracle" (FCPT) increases internal entropy, leaving the system in a state of disorder and atrophy.

### 8.3. AI as a tool of emergence

ROGO does not imply rejecting technology, but redefining the relationship with it. The transition from passive use (subject to FCPT) to active use transforms the external source from a substitute for thinking into a tool for stimulating autonomy.

Continuing the example from point 2.3, an application of ROGO does not consist in rejecting the tool, but in demarcating its functions: giving up GPS on familiar routes (where it substitutes an internalizable competence), while maintaining access to inaccessible information (real-time incident alerts), but treating it as information input to a personal decision, not as a delegated decision. The operational distinction is between "the system informs me, I decide the route" and "the system decides, I execute".

## 9. FCPT-G typologies

The generalization of the FCPT mechanism allows its application in a variety of interactions between cognitive and semi-cognitive systems.

### 9.1. FCPT-HH (Human-Human)

False transfer of autonomy between individuals. Occurs in relationships of authority, mentorship, or intellectual dependence, where the judgment of one individual is substituted for that of another, without a delineation of contribution.

### 9.2. FCPT-HS (Human -Society / Group)

It appears in the form of social conformity or tribalism. The individual adopts the group's positions without verification, but attributes the success or "rightness" of the group to his own intelligence. Social validation becomes a substitute for personal analysis.

### 9.3. FCPT-HI (Human - Institution)

Manifested by the delegation of cognitive responsibility to formal structures (bureaucracy, procedures). The "intelligence" of the institution is confused with the individual competence of its members, who become simple "semiconductors" of procedure, losing their ability to adapt outside the manual.

### 9.4. FCPT-HAI (Human -Artificial Intelligence)

The primary form of the concept, presented in False Cognitive Power Transfer (Stan, 2025). It includes the reliance on generative AI models, where the illusion of competence masks the degradation of deep processing and critical thinking skills.

### 9.5. FCPT-AA (Artificial-Artificial)

Transfer between artificial systems (e.g. *Model Distillation*). A "student" model takes over the distribution of results of a "teacher" model without replicating the internal reasoning processes. The student model seems to perform well, but lacks robustness in novel situations (out-of-distribution), demonstrating that FCPT is a systemic property of information processing, not just psychological.

### 9.6. FCPT-SYS (Systemic / Civilizational)

The aggregate form in which entire civilizations delegate their critical cognitive functions to dominant structures (ideologies, technological systems). Civilizational performance remains enormous, but **resilience** decreases dramatically. A civilization in FCPT-SYS is extremely fragile: any "glitch" in the support systems can lead to a total collapse, because the basis of autonomy (collective cognitive reserve) has been atrophied.

## 10. Predictions and implications

To be scientifically relevant, the FCPT-G framework generates a series of testable predictions, structured to facilitate empirical validation.

### 10.1. Individual-level predictions

#### P1. The illusion of competence under assistance

**Preliminary support:** Shaw & Nave (2026) document decreased deliberative accuracy under AI assistance; Perry et al. (2023) show that AI-assisted users produce less secure code but believe themselves to be more secure - a direct signature of the illusion of competence.

**Hypothesis:** Intensive use of a high-performing external system increases overestimation of one's own performance in similar tasks.

**Mechanism:** The error of attributing external success to internal cognitive resources.

**Observation:** Measurement of the statistical difference between the subject's self-estimated score and the actual performance obtained in a solo test (without assistance).

#### P2. Degradation of autonomous performance through habituation

**Preliminary support:** Shen & Tamkin (2026), in a randomized controlled trial (N=52), find a 17% reduction in skill acquisition in the AI-assisted group compared to the control group (Cohen's  $d = 0.738$ ;  $p = 0.010$ ). Kosmyna et al. (2025), using EEG, observe that the deficit in neural engagement persists after cessation of AI use ("cognitive debt"). In medicine,

Budzyń et al. (2025) find a direct effect in highly experienced professionals: in endoscopists with over 2,000 colonoscopies each, the rate of adenoma detection in the absence of AI assistance decreased from 28.4% to 22.4% after the introduction of the tool, showing objectively measured, not self-reported, degradation. The convergence of three independent fields (programming, medicine, neurophysiology) supports the general nature of the mechanism.

**Hypothesis:** Repeated uncritical cognitive delegation leads to decreased competence in the independent execution of the same task.

**Mechanism:** Functional atrophy through disuse of deep processing processes.

**Observation:** Longitudinal comparison of pre-exposure and post-exposure solo outcomes to the external system.

### **P3. Reduction in active cognitive engagement**

**Preliminary support:** Zhou et al. (2026) find that LLM adoption transforms programming from a generational activity to an evaluation one, with 56.4% of LLM-related actions affected by cognitive biases.

**Hypothesis:** Exposure to external systems decreases the effort put into generating own solutions.

**Mechanism:** Systemic optimization based on the principle of least resistance.

**Observation:** Decrease in reflection time before response ( $\Delta t$ ) and number of independently generated alternatives ( $N_{alt}$ ).

## **10.2. Predictions at the social and systemic level**

**Preliminary support:** Qian & Wexler (2024), through  $N=76$ , document *automation Complacency* and increased dependency over time in software engineers using conversational AI assistance.

### **P4. Uniformity of decisions and erosion of diversity**

**Hypothesis:** Groups dependent on the same external source will exhibit high convergence of solutions and opinions.

**Mechanism:** Alignment of cognitive trajectories to the "geodesics" imposed by the dominant source.

**Observation:** Reduction in variance ( $\sigma^2$ ) in the responses of dependent groups compared to control groups.

### **P5. Propagation and correlation of systemic errors**

**Hypothesis:** External source errors will be faithfully replicated in the output of all dependent members of the system.

**Mechanism:** Uncritical adoption of the response pattern provided by the source.

**Observation:** Identification of high correlation coefficients between the source and user error vectors.

## **10.3. Predictions for artificial systems**

### **P6. Fragility of models derived by distillation (FCPT-AA)**

**Hypothesis:** Student models obtained by distillation will exhibit reduced robustness outside the training distribution (out-of-distribution).

**Mechanism:** Learning the distribution of source results (surface), without replicating reasoning mechanisms (depth).

**Observation:** Testing the model's resilience to input perturbations and edge -case scenarios.

## **11. Limitations**

Any generalized theoretical framework involves simplifications that may limit its accuracy:

1. **Difficulty of direct observation:** Misattribution is an internal, often subconscious process. Its assessment requires indirect methods (comparison of solo vs. assisted performance).
2. **Individual variability:** Cognitive reserve and individual cognitive style can dramatically influence the speed of FCPT installation.
3. **Adaptive context:** In certain crisis situations, delegating autonomy can be a rational survival strategy, not necessarily a pathology.

During the writing of this paper, the author delegated the numerical calibration to an AI system and accepted the results without independent verification - a case of FCPT-HAI committed by the author himself - in the very study on FCPT-G. The values were later downgraded to illustrative

status. The episode confirms the central thesis: the mechanism does not spare even systems with high cognitive reserve, and **the only defense is ROGO applied systematically**. In a documentation stage, an AI system provided an apparently complete academic reference (with authors, title and DOI) to the search that proved non-existent upon independent verification. The detection was possible exclusively by checking each source at the origin, re-confirming that ROGO must be applied systematically.

## 12. Research Directions and Practical Applications

FCPT-G opens new avenues for empirical studies and educational policies:

1. **Cognitive Resilience Index (CRI):** Development of an instrument to measure the relationship between assisted performance and autonomous processing capacity.
2. **ROGO-based education:** Reconfiguring educational practices to emphasize the ability to critically "decouple" from automated systems.
3. **AI Governance:** Analyzing systemic risks in organizations that exhibit a high degree of cognitive dependence on centralized models.

## 13. Integration with IGT-G: Transfer, Field and Stability

While FCPT-G describes **the mechanism** by which autonomy is lost, Generalized Information Gravity Theory (IGT-G) explains why this state becomes **persistent**.

### 13.1. Semantic Mass and the Gravitational Field

External sources (leaders, institutions, AI models) possess a "semantic mass" determined by coherence, repetition, and legitimacy. This mass curves the subject's decision space, creating an informational gravitational field.

### 13.2. Gravitational Capture and Escape Velocity

Near a source with huge semantic mass, the "geodesics" of the individual's thinking are dictated by the external field.

- **Below the threshold  $\tau$ :** The system enters "gravitational capture". The remaining internal energy is insufficient to reach **the escape velocity** necessary to break free from the dependence.
- **Recovery:** To "escape", the system requires a massive energy input (internal through entropy decrease or external through information shocks) that allows the initiation of the ROGO process.

Without this energy surplus, the system remains stuck in an "orbit" of dependence, where autonomy is sacrificed for the stability offered by the external gravitational field. This dynamic is a universal property of complex systems subjected to dense information flows.

## 14. Conclusions

The paper extends the FCPT framework to a general theory (FCPT-G), showing that the autonomy attribution error is a universal mechanism that affects individuals, groups, and artificial systems. By integrating with Maslow7F/DDC and IGT-G, the loss of cognitive autonomy ceases to be seen as a simple individual error, being understood as a systemic process of energetic and informational redistribution.

Probably the only sustainable defense against cognitive atrophy remains the systematic activation of the ROGO principle, capable of generating the "escape speed" necessary to maintain the sovereignty of cognitive systems in an information universe dominated by huge semantic masses.

FCPT represents the linking mechanism between the cognitive divergence described by CDT and the informational stabilization described by IGT, explaining how an external performance is confused with internal autonomy within an already formed semantic field.

## Bibliography

### Specialized literature:

#### Literatura de specialitate:

1. **Arthur, W., Jr., Bennett, W., Jr., Stanush, P. L., & McNelly, T. L. (1998).** Factors that influence skill decay and retention: A quantitative review and analysis. *Human Performance*, 11(1), 57–101. [https://doi.org/10.1207/s15327043hup1101\\_3](https://doi.org/10.1207/s15327043hup1101_3)
2. **Budzyń, K., et al. (2025).** Endoscopist deskilling risk after exposure to artificial intelligence in colonoscopy: a multicentre, observational study. *Lancet Gastroenterology & Hepatology*, 10, 896–903. [https://doi.org/10.1016/S2468-1253\(25\)00133-5](https://doi.org/10.1016/S2468-1253(25)00133-5)
3. **Clark, A., & Chalmers, D. (1998).** The extended mind. *Analysis*, 58(1), 7–19. <https://doi.org/10.1093/analys/58.1.7>
4. **Dell'Acqua, F., McFowland, E., Mollick, E. R., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Krayner, L., Candelon, F., & Lakhani, K. R. (2023).** Navigating the jagged technological frontier: Field experimental evidence of the effects of AI on knowledge worker productivity and quality. *Harvard Business School Working Paper 24-013*. <https://doi.org/10.2139/ssrn.4573321>
5. **Ebbinghaus, H. (1885).** *Über das Gedächtnis: Untersuchungen zur experimentellen Psychologie*. Duncker & Humblot.
6. **Hutchins, E. (1995).** *Cognition in the wild*. MIT Press. <https://doi.org/10.7551/mitpress/1881.001.0001>
7. **Kosmyna, N., Hauptmann, E., Yuan, Y. T., Situ, J., Liao, X.-H., Beresnitsky, A. V., Braunstein, I., & Maes, P. (2025).** Your brain on ChatGPT: Accumulation of cognitive debt when using an AI assistant for essay writing task. *arXiv*. <https://doi.org/10.48550/arXiv.2506.08872>
8. **Parasuraman, R., & Riley, V. (1997).** Humans and automation: Use, misuse, disuse, abuse. *Human Factors*, 39(2), 230–253. <https://doi.org/10.1518/001872097778543886>
9. **Perry, N., Srivastava, M., Kumar, D., & Boneh, D. (2023).** Do users write more insecure code with AI assistants? *Proceedings of the 2023 ACM SIGSAC Conference on Computer and Communications Security (CCS '23)*, 2785–2799. <https://doi.org/10.1145/3576915.3623157>
10. **Qian, C., & Wexler, J. (2024).** Take it, leave it, or fix it: Measuring productivity and trust in human-AI collaboration. *29th International Conference on Intelligent User Interfaces (IUI '24)*. <https://doi.org/10.1145/3640543.3645198>
11. **Risko, E. F., & Gilbert, S. J. (2016).** Cognitive offloading. *Trends in Cognitive Sciences*, 20(9), 676–688. <https://doi.org/10.1016/j.tics.2016.07.002>
12. **Shaw, S. D., & Nave, G. (2026).** Thinking - Fast, slow, and artificial: How AI is reshaping human reasoning and the rise of cognitive surrender. [https://papers.ssrn.com/sol3/papers.cfm?abstract\\_id=6097646](https://papers.ssrn.com/sol3/papers.cfm?abstract_id=6097646)
13. **Shen, J. H., & Tamkin, A. (2026).** How AI impacts skill formation. *arXiv*. <https://doi.org/10.48550/arXiv.2601.20245>
14. **Vella, A., & Blincoe, K. (2026).** The impact of AI coding assistants on software engineering: A longitudinal study. *arXiv*. <https://doi.org/10.48550/arXiv.2605.23135>
15. **Wegner, D. M. (1987).** Transactive memory: A contemporary analysis of the group mind. In B. Mullen & G. R. Goethals (Eds.), *Theories of group behavior* (pp. 185–208). Springer. [https://doi.org/10.1007/978-1-4612-4634-3\\_9](https://doi.org/10.1007/978-1-4612-4634-3_9)
16. **Zhou, X., Saghi, Z., Sabouri, S., Pandita, R., McGuire, M., & Chattopadhyay, S. (2026).** Cognitive biases in LLM-assisted software development. *arXiv / ICSE '26*. <https://doi.org/10.48550/arXiv.2601.08045>

### Works of the author (Adrian Stan):

1. **Stan, A. (2025).** *False Cognitive Power Transfer (FCPT)*. Zenodo. <https://doi.org/10.5281/zenodo.18110163>
2. **Stan, A. (2025).** *Cognitive Divergence Theory (CDT) of AI Adoption*. Zenodo. <https://doi.org/10.5281/zenodo.17776131>
3. **Stan, A. (2025).** *Maslow7F: A Fractal Theory of Systemic Coherence and Motivation from Individuals to Planetary Systems*. Zenodo. <https://doi.org/10.5281/zenodo.18452585>
4. **Stan, A. (2026).** *Information Gravity Theory (IGT)*. Zenodo. <https://zenodo.org/records/18481335>