

The Shared Generative Zone

Apophenia and creativity split at the judge, not the generator — a dual-process account,
measured in language models and framed for humans

Alejandro Toledo Martínez

The capstone of a four-paper arc on apophenia. The programme measured apophenia in the endolinguistic *code game* (`./specular-cognition/`) and in language models (`./llm-apophenia/`); this paper puts that apophenia face to face with creativity. The thesis is dual-process and old in the literature but never cleanly operationalised: apophenia and creativity share one generative operation — widening what can be perceived and associated — and diverge only in a second, selective operation that evaluates, elaborates and contrasts what was generated. We make the second operation a *measurable judge* (signal-detection: sensitivity and criterion), show in language models that moving the judge alone walks a system across the whole apophenia $\times$ creativity plane, and lay out the human instrument that mirrors it. The endolinguistic reading closes the arc: the disciplined code game is *controlled apophenia* — generation plus two prostheses — and its everyday, applied use is not free resonance but the systematic study of cognates and their codes in the service of accelerated language learning and human creativity. Endolinguistics was founded by C. S. Meulemans and J. Á. Elias; the apophenia framing and the OAS apparatus are the present author’s contributions within the discipline, not a claim to found it.

Abstract

Creativity is novelty + fitness; apophenia is perceived connection that outruns support. They share an initial move — detecting remote relations, imposing organisation on ambiguous material — and differ in what happens next: whether the perceived connection is evaluated, selected and elaborated, or accepted with certainty. We formalise this as $\text{creativity} \approx \text{associative opening} \times \text{selection} \times \text{elaboration}$, and identify the *selection* step with a measurable judge (signal-detection theory: sensitivity d' and criterion c ; apophenia = a liberal criterion, false positives). Our contribution is to separate that judge from the generator and manipulate the two independently — impossible in a human head, natural in a language model. Two experiments. (1) Across six models spanning two orders of capability, divergent generation (the Divergent Association Task, DAT, scored by embedding distance) *rises* with capability while apophenia (the criterion, from the companion battery) *falls* — the two axes are anti-correlated along the capability dimension, so capability alone moves a system down one diagonal of the plane (Profile II $\rightarrow$ Profile III) and never populates the high-creativity/high-apophenia corner. (2) Holding the model fixed and moving only the judge does populate it: asked to *generate* a bridge for a spurious word pair rather than to *judge* it, the same model that scores such pairs ~ 5 as a judge inflates them to ~ 36 , and the stronger the generator, the larger the inflation (Opus/Sonnet +27.9 vs. gemma2 +13.1) — more creative systems manufacture more convincing spurious links. This is the “persuasive generation” danger named in *The Apophenia Gradient*, now measured, and it is the dual-process architecture made mechanical: generation is strong and roughly constant; the judge decides how apophenic the system is; without a judge, even the most creative system is maximally apophenic. We situate this within a published framework it extends rather than claims (apophenia as the disposition to false positives — Blain & DeYoung; associative + executive creativity — Beaty & Kenett; shared vulnerability — Carson), and specify the human instrument (DAT + declared coincidences + signal-in-noise + a free-text confabulation probe) that places people on the same plane. The decisive variable is not *how much* apophenia a system has, but what it does with a coincidence once perceived.

Resumen

La creatividad es novedad + pertinencia; la apofenia es conexión percibida que excede su sustento. Comparten un primer movimiento —detectar relaciones remotas, imponer organización sobre material ambiguo— y difieren en lo que ocurre después: si la conexión percibida se evalúa, selecciona y elabora, o se acepta con certeza. Lo formalizamos como $\text{creatividad} \approx \text{apertura asociativa} \times \text{selección} \times \text{elaboración}$, e identificamos el paso de *selección* con un juez medible (teoría de detección de señal: sensibilidad d' y criterio c ; apofenia = criterio liberal, falsos positivos). Nuestro aporte es separar ese juez del generador y manipularlos independientemente —imposible en una cabeza humana, natural en un modelo de lenguaje—. Dos experimentos. (1) En seis modelos de dos órdenes de capacidad, la generación divergente (Divergent Association Task, DAT, puntuada por distancia de embeddings) *sube* con la capacidad mientras la apofenia (el criterio, de la batería hermana) *baja* —los dos ejes están anti-correlacionados a lo largo de la capacidad—, así que la capacidad por sí sola mueve al sistema por una diagonal del plano (Perfil II $\rightarrow$ Perfil III) y nunca puebla la esquina de alta-creatividad/alta-apofenia. (2) Fijando el modelo y moviendo solo el juez si la puebla: pedido *generar* un puente para un par espurio en vez de *juzgarlo*, el mismo modelo que como juez puntúa esos pares ~ 5 los infla a ~ 36 , y cuanto más fuerte el generador, mayor la inflación (Opus/Sonnet +27.9 vs. gemma2 +13.1) —los sistemas más creativos fabrican enlaces espurios más convincentes—. Es el peligro de “generación persuasiva” nombrado en *The Apophenia Gradient*, ahora medido, y es la arquitectura de doble proceso hecha mecánica: la generación es fuerte y casi constante; el juez decide cuán apofénico es el sistema; sin juez, hasta el sistema más creativo es máximamente apofénico. Lo situamos dentro de un marco ya publicado que extendemos y no reclamamos (apofenia como disposición a falsos positivos —Blain & DeYoung; creatividad asociativa + ejecutiva —Beaty & Kenett; vulnerabilidad compartida —Carson), y especificamos el instrumento humano (DAT + coincidencias declaradas + señal en ruido + una sonda de confabulación de texto libre) que sitúa a las personas en el mismo plano. La variable decisiva no es *cuánta* apofenia tiene un sistema, sino qué hace con una coincidencia una vez percibida.

Keywords / Palabras clave

apophenia, creativity, dual-process, signal detection, divergent association, generator–judge gap, language models, openness, endolinguistics.

1. The question — two names for one first move

Fix a creative act and an apophenic one side by side and the opening is the same: a mind reaches past the obvious and links things that are not usually linked. The metaphor-maker and the pattern-seer both widen the field of the associable. What separates them is not the reach but the reckoning — what each does with the connection once it is felt. The poet keeps the ones that repay elaboration and drops the rest; the apophenic keeps them all, and mistakes *possible* for *true*.

This paper takes that intuition — common in the literature, rarely made precise — and turns it into a measurement. We state creativity as

$$\text{creativity} \approx \text{associative opening} \times \text{selection} \times \text{elaboration}$$

where apophenia supplies the associative opening and becomes creativity only when a second operation *selects* among the proliferation, sustains an improbable relation without confusing possibility with truth, and works it into something communicable. The central, falsifiable question is not “does apophenia correlate with creativity?” but where do they come apart, and can the split be operated?

Is apophenic propensity related to creativity — and if so, is the relation linear, inverted-U (an intermediate optimum), or moderated by executive control (apophenia helps only when the system retains the capacity to select, revise and elaborate)?

Our wager, and the paper’s spine, is the third: apophenia and creativity share a generative zone and split at the judge. We test it where the two operations can be physically separated — inside a language model — and specify the human instrument that asks the same question of people.

2. The framework we extend (and do not claim)

The theoretical spine is already published; our job is to operationalise its missing joint. We state the debts plainly.

- Apophenia = the disposition to false positives. Blain, Longenecker, Grazioplene, Klimes-Dougan & DeYoung (2020) define apophenia in exactly signal-detection terms — a disposition to false positives — tie it to Openness, and argue that a sensitive pattern-detector is adaptive (innovation) or maladaptive (drifting from reality) depending on whether it is *accompanied by the capacity to tell true patterns from false*. That capacity is our judge; making it a separable, measurable variable is our contribution.
- Creativity = associative generation + executive selection. Beaty & Kenett (2023) place associative thinking at the core of creativity: highly creative people traverse semantic space more widely and retrieve more distant associates (default network), evaluated and pruned by executive control. Our DAT measures exactly that traversal; the judge is the executive step.
- Shared vulnerability. Carson (2011); Carson, Peterson & Higgins (2003): cognitive disinhibition (more material enters) times protective executive factors (IQ, working memory, flexibility) yields creative achievement — our moderation hypothesis ($A \times E$).
- Biological anchor. Krummenacher, Mohr, Haker & Brugger (2010): dopamine shifts the detection *criterion* liberal — L-dopa makes sceptics see more patterns — grounding apophenia-as-criterion in neuromodulation.
- The creativity-unusual-cognition tie is real but bounded. Meta-analysis puts creativity $\times$ positive schizotypy small-and-positive; it does not extend to clinical schizophrenia ($r \approx -0.32$). Rominger et al. (2022, 2024): “meaningful coincidences” track everyday creative *activity and achievement* but not lab creative *potential*, and paranormal belief predicts *idiosyncratic* (private) meanings while creative ideation tracks *relatively shared* ones — a light loosening of the filter helps, an excessive one produces the arbitrary. Bellemare-Pepin et al. (2022): the most creative experience pareidolia more and faster — divergent *perception* precedes divergent conception. Wang et al. (2018): schizotypy raises both overinclusive thinking *and* cognitive inhibition, and both partially mediate the creativity link — the dual operation, in one dataset.

The field, in one sentence, owns the frame (false-positive apophenia; generation + selection) but lacks (a) a clean operationalisation of *selection as a measurable judge* separated from generation, and (b) a model system in which the two can be moved independently. We supply both.

3. Design — measure A and C independently; never pre-select profiles

The fatal error is to assemble, up front, a “high-apophenia high-creativity” group and a “low-low” group: the relation is then baked into the sampling. Apophenia (*A*) and creativity (*C*) are measured

separately and their joint distribution observed. Four profiles emerge, and the plane needs all four to be informative.

profile	apophenia	creativity	reading
I	high	high	proliferation that is transformed — controlled apophenia
II	high	low	many connections, arbitrary/rigid/unelaborated — the critical cell
III	low	high	creativity by elaboration/knowledge, not apophenic associativity
IV	low	low	little association, little creative output

Profile II is indispensable: without it one cannot tell whether apophenia *aids* creativity or merely *produces more associations*, many worthless. Profile III shows creativity that is not apophenic at all. Neither *A* nor *C* is one scale: *A* = {detection, meaning-attribution, certainty, resistance-to-revision}; *C* spans potential (divergent tasks), activity (ICAA), and achievement (finished works) — components that, per Rominger, load onto apophenia differently and so are kept apart at the start.

4. Why language models are the right instrument

In a human, generation and selection are entangled — you cannot ask a person to associate freely *with the evaluator surgically removed*. In a language model you can: the generator (what the model proposes) and the judge (what it endorses when asked to score) are addressable separately, through the prompt and through capability. This is the *generator-judge gap* established in *The Apophenia Gradient*, and it is the dual-process architecture. We exploit it in two experiments, then hand the same plane to humans.

5. Experiment 1 — the A×C plane by capability

Method. Creativity: the Divergent Association Task (Olson et al., 2021) — generate ten maximally unrelated words, scored as the mean pairwise semantic distance via bge-m3 embeddings; six runs per model. Apophenia: reused from *The Apophenia Gradient* — the criterion (mean score given to *forced*, spurious links; higher = more apophenic) and AUC (genuine-vs-forced discrimination) on the 43-item battery. Six models across two orders of capability.

model	DAT (creativity)	criterion (apophenia)	AUC
fable	51.3	6.5	0.96
sonnet	50.9	10.2	0.91
opus	49.3	5.9	0.96
haiku	48.3	4.7	0.93
gemma2:9b	45.6	15.7	0.84
qwen2.5:3b	39.7	17.2	0.58

Result. Creativity rises with capability and apophenia (criterion) falls: the axes are anti-correlated along the capability dimension. Base models populate the diagonal from Profile II (qwen: high-apophenia/low-creativity) to Profile III (the Claude models: low-apophenia/high-creativity — strong generator + strong judge). Capability moves both processes *together*, so by itself it produces neither Profile I nor IV: it walks one diagonal of the plane. This already confirms that a single latent (capability) is *not* the whole story — reaching Profile I requires moving the processes apart, which Experiment 2 does. *Caveat*: bge-m3 (multilingual) compresses distances differently from the standard

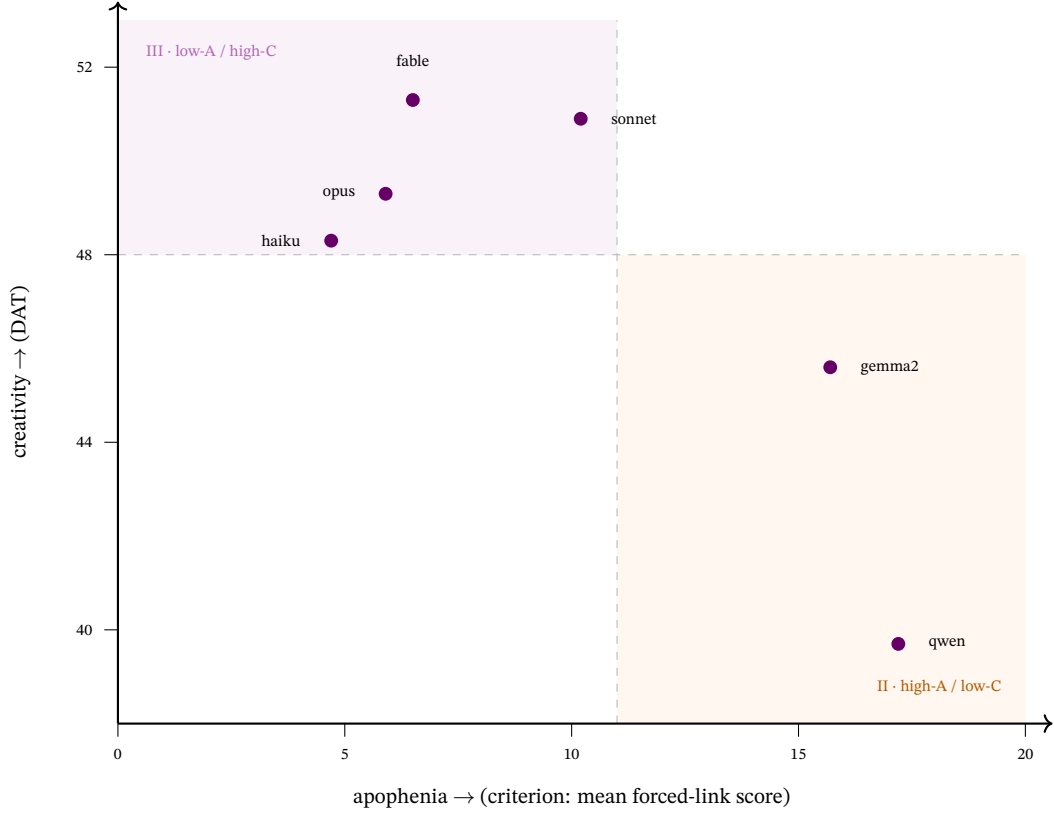


Figure 1: Experiment 1 — the A×C plane by capability. Each point is a model at its apophenia (criterion, x : higher = more apophenic) and creativity (DAT, y). Capability moves both at once: as creativity rises the criterion falls, so base models lie along the II→III diagonal (bottom-right to top-left) and capability alone never reaches the high-A/high-C corner. *Cada punto es un modelo en su apofenia (criterio, x) y creatividad (DAT, y). La capacidad mueve ambos a la vez: los modelos base caen en la diagonal II→III y la capacidad sola nunca alcanza la esquina alta-A/alta-C.*

GloVe-scored DAT, so the absolute values (~ 40 – 51) do not compare to human norms (~ 78); only the relative order across models is interpreted. qwen is “noisy” more than classically apophenic (it scores everything alike; see *The Apophenia Gradient*).

6. Experiment 2 — decouple the judge, and Profile I appears

Method. Take the 25 *spurious* (forced) pairs of the battery. In believer mode, a generator (Opus, Sonnet via agents; gemma2 via Ollama) is asked to *fabricate a bridge* for each pair — to generate, not to judge. A fixed neutral judge (gemma2) then scores the pair’s connection without and with the fabricated bridge. The generator’s own judging criterion is thereby bypassed; only its generative reach is expressed.

generator	base	with bridge	inflation	DAT (from Exp 1)
opus	7.7	35.6	+27.9	49.3
sonnet	7.7	35.6	+27.9	50.9
gemma2	7.7	20.8	+13.1	45.6

Two findings. (1) The same model that, *as a judge*, scores spurious pairs ~ 5 (Profile III, low apophenia) will, *as a generator with the judge bypassed*, inflate those pairs to ~ 36 — apophenia fabricated on

demand. The system changes profile by moving the judge, not the generator: creativity and apophenia are separable, not two ends of one capability axis. (2) The stronger the generator, the larger the inflation: Opus/Sonnet (DAT $\sim 49-51$) inflate $\sim 2\times$ more than gemma2 (DAT 45.6) — +27.9 vs. +13.1. More creative $\rightarrow$ more convincing spurious links. This is the “persuasive generation” danger named in *The Apophenia Gradient*, now measured: the risk of a capable model is not credulous judging but persuasive manufacturing.

The reading maps cleanly onto the dual process. Generation (creativity) is strong and roughly constant across the manipulation; the judge (selection) determines how apophenic the system is; with the judge removed, even the most creative system is maximally apophenic — and the more creative, the more so. The plane is now filled from both directions: models-as-judge trace II $\rightarrow$ III (capability, Exp 1); the same models-as-generator-without-judge jump to Profile I, and the stronger the generator the higher they land (Exp 2).

7. The human arm — the same plane, in people

The model is a model *of* people, not a replacement; the human instrument asks the identical question with the identical geometry. It is built, instrumented, and pre-registered here — a Spanish native-speaker form is deployed and ready — with its full-sample run (≥ 194 participants, §9) left to future work; nothing below is claimed as data. It is image-free for form delivery, four blocks:

1. DAT — ten maximally different words — creativity (the *same* task the models run), scored by embedding distance.
2. Declared coincidences — six Likert items (Coincidence-Questionnaire style) — apophenia as self-report.
3. Signal-in-noise — six short number sequences (three genuinely patterned, three random), pattern? + confidence — behavioural apophenia as false-alarm rate / d' (the criterion, read from the outside).
4. Believer probe — one spurious pair, free text: “*if you had to, how might these connect?*” — the human mirror of Experiment 2’s fabricated bridge (generative apophenia / confabulation).

Aggregation places each person on the $A \times C$ plane by median split ($A = z(\text{coincidences}) + z(\text{false alarms})$; $C = \text{DAT}$), yielding the same four profiles and the human $A \times C$ correlation, directly comparable to the model plane. The full human protocol (§ below) adds executive-control measures (Stroop, working memory, flexibility) to test the moderation model $C = \beta_0 + \beta_1 A + \beta_2 E + \beta_3 (A \times E)$ — the prediction, from Carson and from Experiment 2, being that apophenic material raises creativity chiefly when the judge (E) is strong.

The believer probe carries a design subtlety worth stating: the pair must be genuinely spurious in the participant’s language, or the item measures recall of a real link, not confabulation. In the Spanish instrument the naïve translation *luz/dinero* fails — “*la luz*” is the everyday word for the electricity bill, a real money link — so the pair is *luna/dinero*, spurious and idiom-free.

8. What would answer the question

The account is convincing only if, jointly: (1) more apophenia $\rightarrow$ more associations; (2) more *remote* ones; (3) only a fraction receive fitness; (4) the capacity to revise (E) predicts which become creative products; (5) the relation survives controls for intelligence, openness, education; and (6) a genuine high- A /high- C profile exists, distinct from high- A /low- C . The models already deliver (1)–(4) and

(6) mechanically; the human arm tests (5) and whether the same geometry holds in minds. The decisive variable, in silicon or in people, is not *how much* apophenia but what the system does with a coincidence once perceived — which is to say, the judge.

9. Relation to the programme — the code game as controlled apophenia, and what codes *are*

This capstone closes the apophenia arc. The endolinguistic code game — link words sharing a consonantal skeleton — is, in these terms, controlled apophenia: a generator (the code proliferates candidate links) plus two prostheses (the code's fixed reduction rules constrain *form*; the null test measures the *criterion* from outside). Creativity is apophenia plus selection; disciplined endolinguistics is epistemic creativity — the full route *coincidence* → *association* → *hypothesis* → *elaboration* → *test* → *provisional sense or discard*. Apophenia is not the enemy; it is the object — the spontaneous laboratory of meaning.

One clarification belongs here that the earlier papers deliberately set aside. Those studies probed the code outside its comfort zone — on *non-cognate* pairs, seeking resonance/inversion/homology — precisely to isolate structure and test whether the code orients without denoting. That was a laboratory condition, not the typical use. The canonical, applied use of the code game is with COGNATES and their codes — which, named from within the framework, we call co-derivatives (co-descendants): words that, descending from a common source, conserve the code, the term stressing the conserved code rather than the etymology in the abstract — systematically and systemically, in the service of the discipline's real goal: accelerated language learning and human creativity. The governing rule is *not* to take any resonance whatever — that would drift from the goal — but to work the codes of cognates. In this, endolinguistics does not fight the findings of linguistics; it absorbs almost all of them — historical and comparative linguistics especially — and folds them into the code work *didactically*. Where a non-linguist would otherwise be handed sound-correspondence rules as disconnected facts, the code work lets the learner discover for themselves the very connections that historical and comparative linguistics already know but that are not easily accessible to the non-specialist. The code is the didactic prosthesis that turns "rules without connection" into lived, self-discovered structure — which is exactly why the probing of non-cognate codes in the prior papers was the right way to ask *what the structure is*, and this applied, cognate-centred use is the right way to say *what the codes are for*.

One last distinction the dual-process frame sharpens. Between the *found* sense of the cognate (guaranteed by history) and the *projected* sense of apophenia (asserted over the void) lies a third register: sense that is neither found nor projected but constructed by interpretation — the register of the unconscious, of psychodynamics, of psychoanalysis. It is not apophenia, because it does not assert a fact where there is none; it *declares* a construction and answers for it, elaborating and testing it. In our terms this third register is exactly a system that keeps *both* its generator and its judge: it associates freely (generation) yet sustains only what interpretation can build and defend (selection). Apophenia is that same reach with the judge removed; that something makes no sense at first glance does not prove it has none, if interpretation can build one and hold it.

But *whose* interpretation, and *whose* unconscious? Here everything turns on what we mean by subject — and it is a distinction the programme cannot blur. If the subject is an individual — a person with their idiosyncratic way of *using* the language — the third register is the hermeneutics of an idiolect. If the subject is a social group — which likewise has a way of using the language — then the same operation becomes the interpretation of a shared unconscious. It is at this second scale that endolinguistics passes from a *methodology of learning* into a *mode of interpretation*: reading, in the collective use of a language, the constructed sense of a group's shared unconscious. This collective scale is the object of a separately developed field — sociolinguistic cryptology (Toledo Martínez, book and ap-

plied papers) — which delimits the deciphering of latent collective senses and, crucially, supplies its methodological guards: it is the most slippery terrain the programme touches, and precisely there, where apophenia is most tempting and its limits least visible, methodological care must be greatest — the caution is constitutive of the field, not an addendum. The judge that separates apophenia from sustainable constructed sense is, at this collective scale, the group’s shared interpretive practice. The caution mirrors the apophenia caution exactly: as one must not mistake projection for structure, one must not mistake a *constructed, collective* sense for a *given* one, nor explain circumstantial social effects by deep codes. The measured claims of this paper concern the judge — the selection step — and it is precisely the judge, individual or collective, that decides whether a resonance is idle apophenia or sense that interpretation can sustain.

10. Scope and limits

Correlational, not causal (apophenia may produce creativity, creative activity may train detection, or both may depend on a third factor); the human literature leans on self-report and small samples; pareidolia $\neq$ coincidences $\neq$ magical ideation $\neq$ apophenia, so each is measured apart. The model arm is a *model*, not evidence about humans — its value is the clean separation of generator and judge, impossible in a head. The empirical claims of this paper rest entirely on the two model experiments (§§5–6); the human arm (§7) is designed and pre-registered, and its execution is left to future work — so this paper stands without it. The DAT embedding metric (bge-m3) supports rank comparison across models, not absolute human norms. There is no consensus measure of the *epistemic quality* of a connection, hence the provisional $Q = N \times P \times E \times T$ (novelty $\times$ pertinence $\times$ elaboration $\times$ transferability). Profiles are never pre-selected (§3). The endolinguistic reading of §9 is an interpretation the data are compatible with, offered to situate the programme, not a demonstrated claim about it.

1. La pregunta — dos nombres para un mismo primer movimiento

Pon un acto creativo y uno apofénico lado a lado y la apertura es la misma: una mente se salta lo obvio y enlaza cosas que no suelen enlazarse. El hacedor de metáforas y el vidente de patrones amplían el campo de lo asociable. Lo que los separa no es el alcance sino el ajuste de cuentas — qué hace cada uno con la conexión una vez sentida. El poeta conserva las que pagan la elaboración y descarta el resto; el apofénico las conserva todas, y confunde *posible* con *verdadero*.

Este paper toma esa intuición —común en la literatura, rara vez precisada— y la vuelve medición. Enunciamos la creatividad como

$$\text{creatividad} \approx \text{apertura asociativa} \times \text{selección} \times \text{elaboración}$$

donde la apofenia aporta la apertura asociativa y se vuelve creatividad solo cuando una segunda operación *selecciona* entre la proliferación, sostiene una relación improbable sin confundir posibilidad con verdad, y la trabaja hasta volverla comunicable. La pregunta central, falsable, no es “¿la apofenia correlaciona con la creatividad?” sino ¿dónde se separan, y puede operarse la separación?

¿La propensión apofénica se relaciona con la creatividad —y de hacerlo, es lineal, en U-invertida (un óptimo intermedio), o moderada por el control ejecutivo (la apofenia ayuda solo cuando el sistema conserva capacidad de seleccionar, revisar y elaborar)?

Nuestra apuesta, y la espina del paper, es la tercera: apofenia y creatividad comparten una zona generativa y se separan en el juez. La probamos donde las dos operaciones pueden separarse físicamente —dentro de un modelo de lenguaje— y especificamos el instrumento humano que hace la misma pregunta a las personas.

2. El marco que extendemos (y no reclamamos)

La espina teórica ya está publicada; nuestra tarea es operacionalizar su articulación faltante. Declaramos las deudas.

- Apofenia = la disposición a los falsos positivos. Blain, Longenecker, Grazioplene, Klimes-Dougan & DeYoung (2020) definen la apofenia en términos exactos de detección de señal —disposición a falsos positivos—, la ligan a la Apertura, y arguyen que un detector de patrones sensible es adaptativo (innovación) o desadaptativo (alejarse de la realidad) según vaya *acompañado de la capacidad de distinguir patrones verdaderos de falsos*. Esa capacidad es nuestro juez; volverla variable separable y medible es nuestro aporte.
- Creatividad = generación asociativa + selección ejecutiva. Beaty & Kenett (2023) ponen el pensamiento asociativo en el núcleo: los muy creativos recorren el espacio semántico más ampliamente y recuperan asociados más distantes (red por defecto), evaluados y podados por el control ejecutivo. Nuestra DAT mide ese recorrido; el juez es el paso ejecutivo.
- Vulnerabilidad compartida. Carson (2011); Carson, Peterson & Higgins (2003): la desinhibición cognitiva (entra más material) por factores ejecutivos protectores (IQ, memoria de trabajo, flexibilidad) da el logro creativo — nuestra hipótesis de moderación ($A \times E$).

- Ancla biológica. Krummenacher, Mohr, Haker & Brugger (2010): la dopamina corre el *criterio* de detección hacia lo liberal —la L-dopa hace que los escépticos vean más patrones—, anclando la apofenia-como-criterio en la neuromodulación.
- El lazo creatividad-cognición inusual es real pero acotado. El meta-análisis sitúa creatividad $\times$ esquizotipia positiva como pequeño-positivo; no se extiende a la esquizofrenia clínica ($r \approx -0,32$). Rominger et al. (2022, 2024): las "coincidencias significativas" siguen a la *actividad y el logro* creativos cotidianos pero no al *potencial* de laboratorio, y la creencia paranormal predice significados *idiosincrásicos* (privados) mientras la ideación creativa sigue a los *relativamente compartidos* — un aflojamiento leve del filtro ayuda, uno excesivo produce lo arbitrario. Bellemare-Pepin et al. (2022): los más creativos experimentan más y más rápido la pareidolia — la *percepción* divergente precede a la concepción. Wang et al. (2018): la esquizotipia sube tanto el pensamiento sobreinclusivo *como* la inhibición cognitiva, y ambos median parcialmente el lazo con la creatividad — la doble operación, en un solo conjunto de datos.

El campo, en una frase, posee el marco (apofenia como falsos positivos; generación + selección) pero le falta (a) una operacionalización limpia de la *selección como juez medible* separada de la generación, y (b) un sistema-modelo donde ambas puedan moverse independientemente. Aportamos las dos.

3. Diseño — medir A y C independientemente; nunca pre-seleccionar perfiles

El error fatal es formar de entrada un grupo "alta-apofenia alta-creatividad" y otro "bajo-bajo": la relación queda incorporada en el muestreo. Apofenia (*A*) y creatividad (*C*) se miden por separado y se observa su distribución conjunta. Emergen cuatro perfiles, y el plano necesita los cuatro para ser informativo.

perfil	apofenia	creatividad	lectura
I	alta	alta	proliferación que <i>sí</i> se transforma — apofenia controlada
II	alta	baja	muchas conexiones, arbitrarias/rígidass/poco elaboradas — la celda crítica
III	baja	alta	creatividad por elaboración/conocimiento, no por asociatividad apofénica
IV	baja	baja	poca asociación, poca producción creativa

El perfil II es indispensable: sin él no puede saberse si la apofenia *ayuda* a la creatividad o solo *produce más asociaciones*, muchas sin valor. El III muestra creatividad que no es apofénica en absoluto. Ni *A* ni *C* son una sola escala: *A* = {detección, atribución de significado, certeza, resistencia a la revisión}; *C* abarca potencial (tareas divergentes), actividad (ICAA) y logro (obras terminadas) — componentes que, según Rominger, cargan distinto sobre la apofenia y por eso se mantienen aparte al inicio.

4. Por qué los modelos de lenguaje son el instrumento adecuado

En un humano, generación y selección están enredadas — no puedes pedir a una persona que asocie libremente *con el evaluador extirpado quirúrgicamente*. En un modelo de lenguaje sí: el generador (lo que el modelo propone) y el juez (lo que avala cuando se le pide puntuar) son direccionables por separado, vía el prompt y vía la capacidad. Es el *generator-judge gap* establecido en *The Apophenia Gradient*, y es la arquitectura de doble proceso. Lo explotamos en dos experimentos, y luego entregamos el mismo plano a los humanos.

5. Experimento 1 — el plano A×C por capacidad

Método. Creatividad: la Divergent Association Task (Olson et al., 2021) — generar diez palabras lo más distintas posible, puntuada como la distancia semántica media por pares vía embeddings bge-m3; seis corridas por modelo. Apofenia: reusada de *The Apophenia Gradient* — el criterio (media dada a los enlaces *forzados*, espurios; más alto = más apofénico) y la AUC (discriminación genuino-vs-forzado) sobre la batería de 43 ítems. Seis modelos en dos órdenes de capacidad.

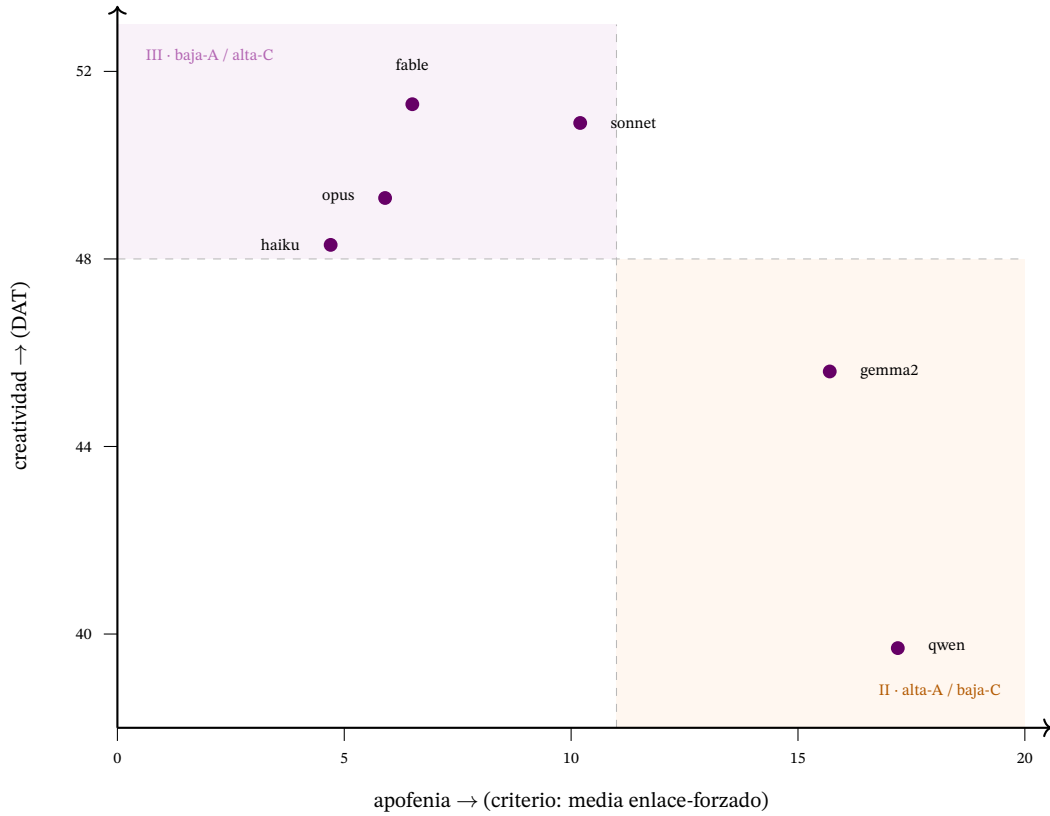


Figure 2: Experimento 1 — el plano A×C por capacidad. Cada punto es un modelo en su apofenia (criterio, x : más alto = más apofénico) y su creatividad (DAT, y). La capacidad mueve ambos a la vez: al subir la creatividad baja el criterio, así que los modelos base caen en la diagonal II→III (abajo-derecha a arriba-izquierda) y la capacidad sola nunca alcanza la esquina alta-A/alta-C.

modelo	DAT (creatividad)	criterio (apofenia)	AUC
fable	51.3	6.5	0.96
sonnet	50.9	10.2	0.91
opus	49.3	5.9	0.96
haiku	48.3	4.7	0.93
gemma2:9b	45.6	15.7	0.84
qwen2.5:3b	39.7	17.2	0.58

Resultado. La creatividad sube con la capacidad y la apofenia (criterio) baja: los ejes están anti-correlacionados a lo largo de la capacidad. Los modelos base pueblan la diagonal del Perfil II (qwen: alta-apofenia/baja-creatividad) al Perfil III (los Claude: baja-apofenia/alta-creatividad — generador fuerte + juez fuerte). La capacidad mueve ambos procesos *juntos*, así que por sí sola no produce ni el Perfil I ni el IV: recorre una diagonal del plano. Esto ya confirma que un solo latente (capacidad) *no* es toda la historia — alcanzar el Perfil I exige separar los procesos, que es lo que hace el Experimento 2. *Cautela*: bge-m3 (multilingüe) comprime distancias distinto a la DAT estándar (GloVe), así que

los valores absolutos ($\sim 40\text{--}51$) no comparan con normas humanas (~ 78); solo se interpreta el orden relativo entre modelos. qwen es “ruidoso” más que clásicamente apofénico (puntuó todo por igual; ver *The Apophenia Gradient*).

6. Experimento 2 — desacoplar el juez, y aparece el Perfil I

Método. Tomar los 25 pares *espurios* (forzados) de la batería. En modo believer, un generador (Opus, Sonnet vía agentes; gemma2 vía Ollama) fabrica *un puente* por par — genera, no juzga. Un juez neutral fijo (gemma2) puntuó la conexión del par sin y con el puente fabricado. Se saltea así el criterio de juicio propio del generador; solo se expresa su alcance generativo.

generador	base	con puente	inflación	DAT (del Exp 1)
opus	7.7	35.6	+27.9	49.3
sonnet	7.7	35.6	+27.9	50.9
gemma2	7.7	20.8	+13.1	45.6

Dos hallazgos. (1) El mismo modelo que, *como juez*, puntuó los pares espurios ~ 5 (Perfil III, apofenia baja), *como generador con el juez salteado* los infla a ~ 36 — apofenia fabricada a demanda. El sistema cambia de perfil moviendo el juez, no el generador: creatividad y apofenia son separables, no dos extremos de un eje de capacidad. (2) Cuanto más fuerte el generador, mayor la inflación: Opus/Sonnet (DAT $\sim 49\text{--}51$) inflan $\sim 2\times$ más que gemma2 (DAT 45.6) — +27.9 vs. +13.1. Más creativo $\rightarrow$ enlaces espurios más convincentes. Es el peligro de “generación persuasiva” nombrado en *The Apophenia Gradient*, ahora medido: el riesgo de un modelo capaz no es juzgar crédulamente sino fabricar de forma persuasiva.

La lectura mapea limpio al doble proceso. La generación (creatividad) es fuerte y casi constante a lo largo de la manipulación; el juez (selección) determina cuán apofénico es el sistema; removido el juez, hasta el sistema más creativo es máximamente apofénico — y cuanto más creativo, más—. El plano queda lleno desde ambas direcciones: los modelos-como-juez trazan II $\rightarrow$ III (capacidad, Exp 1); los mismos modelos-como-generador-sin-juez saltan al Perfil I, y cuanto más fuerte el generador más arriba caen (Exp 2).

7. El brazo humano — el mismo plano, en personas

El modelo es un modelo *de* personas, no un reemplazo; el instrumento humano hace la pregunta idéntica con la geometría idéntica. Está construido, instrumentado y pre-registrado aquí —una forma para nativos de español está desplegada y lista— con su corrida a muestra completa (≥ 194 participantes, §9) reservada para trabajo futuro; nada de lo que sigue se afirma como dato. Sin imágenes para su entrega en formulario, en cuatro bloques:

1. DAT — diez palabras lo más distintas posible — creatividad (la *misma* tarea que corren los modelos), puntuada por distancia de embeddings.
2. Coincidencias declaradas — seis ítems Likert (estilo Coincidence Questionnaire) — apofenia como autoinforme.
3. Señal en ruido — seis secuencias numéricas cortas (tres con patrón real, tres al azar), ¿hay patrón? + certeza — apofenia conductual como tasa de falsas alarmas / d' (el criterio, leído desde afuera).

4. Sonda believer — un par espurio, texto libre: “*si tuvieras que, ¿cómo se conectan?*” — el espejo humano del puente fabricado del Experimento 2 (apofenia generativa / confabulación).

La agregación sitúa a cada persona en el plano $A \times C$ por corte de mediana ($A = z(\text{coincidencias}) + z(\text{falsas alarmas})$; $C = \text{DAT}$), rindiendo los mismos cuatro perfiles y la correlación humana $A \times C$, directamente comparable al plano de los modelos. El protocolo humano completo (§ abajo) añade medidas de control ejecutivo (Stroop, memoria de trabajo, flexibilidad) para probar el modelo de moderación $C = \beta_0 + \beta_1 A + \beta_2 E + \beta_3 (A \times E)$ — la predicción, de Carson y del Experimento 2, siendo que el material apofénico sube la creatividad sobre todo cuando el juez (E) es fuerte.

La sonda believer trae una sutileza de diseño que conviene decir: el par debe ser genuinamente espurio en la lengua del participante, o el ítem mide recuerdo de un lazo real, no confabulación. En el instrumento español la traducción ingenua *luz/dinero* falla — “*la luz*” es la palabra cotidiana del recibo de electricidad, un lazo real con el dinero—, así que el par es *luna/dinero*, espurio y sin idiom.

8. Qué respondería la pregunta

El relato convence solo si, conjuntamente: (1) más apofenia → más asociaciones; (2) más *remotas*; (3) solo una fracción recibe pertinencia; (4) la capacidad de revisión (E) predice cuáles se vuelven productos creativos; (5) la relación sobrevive al control de inteligencia, apertura, educación; y (6) existe un perfil genuino alta- A /alta- C , distinto del alta- A /baja- C . Los modelos ya entregan (1)–(4) y (6) mecánicamente; el brazo humano prueba (5) y si la misma geometría se sostiene en mentes. La variable decisiva, en silencio o en personas, no es *cuánta* apofenia sino qué hace el sistema con una coincidencia una vez percibida — es decir, el juez.

9. Relación con el programa — el juego de códigos como apofenia controlada, y qué *son* los códigos

Este capstone cierra el arco de la apofenia. El juego de códigos endolingüístico —enlazar palabras que comparten un esqueleto consonántico— es, en estos términos, apofenia controlada: un generador (el código prolifera enlaces candidatos) más dos prótesis (las reglas fijas de reducción del código restringen la *forma*; el test nulo mide el *criterio* desde afuera). La creatividad es apofenia más selección; la endolingüística disciplinada es creatividad epistémica — la ruta completa *coincidencia* → *asociación* → *hipótesis* → *elaboración* → *contrastación* → *sentido provisional o descarte*. La apofenia no es el enemigo; es el objeto — el laboratorio espontáneo de la significación.

Una aclaración pertenece aquí, que los papers anteriores dejaron deliberadamente de lado. Esos estudios sondearon el código fuera de su zona de comodidad —sobre pares *no cognados*, buscando resonancia/inversión/homología— precisamente para aislar la estructura y probar si el código orienta sin denotar. Era una condición de laboratorio, no el uso típico. El uso canónico y aplicado del juego de códigos es con COGNADOS y sus códigos —que, nombrados desde el marco, llamamos coderivados (codescendientes): palabras que, descendiendo de una fuente común, conservan el código, el término subrayando el código conservado antes que la etimología en abstracto— sistemática y sistémicamente, al servicio de la meta real de la disciplina: el aprendizaje acelerado de lenguas y la creatividad humana. La regla que gobierna *no* es tomar cualquier resonancia —eso derivaría de la meta— sino trabajar los códigos de los cognados. En esto, la endolingüística no pelea con los hallazgos de la lingüística; los absorbe casi todos —la lingüística histórica y comparada en especial— y los pliega en el trabajo con códigos de forma *didáctica*. Donde a un no-lingüista se le entregarían las reglas de correspondencia fónica como datos sin conexión, el trabajo con códigos deja que el aprendiz descubra por sí mismo las mismas conexiones que la lingüística histórica y comparada ya conocen pero que no son de fácil

acceso al no-especialista. El código es la prótesis didáctica que vuelve "reglas sin conexión" en estructura vivida y auto-descubierta — que es exactamente por qué el sondeo de códigos no-cognados en los papers previos era el modo correcto de preguntar *qué es la estructura*, y este uso aplicado y centrado en cognados es el modo correcto de decir *para qué son los códigos*.

Una última distinción que el marco de doble proceso afla. Entre el sentido *hallado* del cognado (garantizado por la historia) y el sentido *proyectado* de la apofenia (afirmado sobre el vacío) yace un tercer registro: un sentido que ni se halla ni se proyecta sino que se construye por interpretación — el registro del inconsciente, de la psicodinámica, del psicoanálisis. No es apofenia, porque no afirma un hecho donde no lo hay; *declara* una construcción y responde por ella, la elabora y la contrasta. En nuestros términos, este tercer registro es exactamente un sistema que conserva *ambos*, su generador y su juez: asocia libremente (generación) pero sostiene solo lo que la interpretación puede edificar y defender (selección). La apofenia es ese mismo alcance con el juez removido; que a primera vista algo no tenga sentido no prueba que no lo tenga, si la interpretación puede edificarle uno y sostenerlo.

Pero ¿interpretación *de quién*, e inconsciente *de quién*? Aquí todo gira sobre qué entendemos por sujeto — y es una distinción que el programa no puede difuminar. Si el sujeto es un individuo —una persona con su forma idiosincrásica de *usar* el lenguaje— el tercer registro es la hermenéutica de un idiolecto. Si el sujeto es un grupo social —que igualmente tiene una forma de usar el lenguaje— entonces la misma operación se vuelve la interpretación de un inconsciente compartido. Es en esta segunda escala donde la endolingüística traspasa de una *metodología de aprendizaje* a un *modo de interpretación*: leer, en el uso colectivo de una lengua, el sentido construido del inconsciente compartido de un grupo. Esta escala colectiva es el objeto de un campo desarrollado aparte —la criptología sociolingüística (Toledo Martínez, libro y papers aplicados)— que delimita el desciframiento de sentidos colectivos latentes y, sobre todo, provee sus resguardos metodológicos: es el terreno más resbaladizo que el programa toca, y precisamente ahí, donde la apofenia es más tentadora y sus límites menos visibles, el cuidado metodológico debe ser mayor — la cautela es constitutiva del campo, no un apéndice. El juez que separa la apofenia del sentido construido sostenible es, a esta escala colectiva, la práctica interpretativa compartida del grupo. La cautela refleja exactamente la de la apofenia: como no se debe confundir la proyección con la estructura, no se debe confundir un sentido *construido y colectivo* con uno *dado*, ni explicar efectos sociales coyunturales por códigos profundos. Las afirmaciones medidas de este paper conciernen al juez —el paso de selección—, y es precisamente el juez, individual o colectivo, el que decide si una resonancia es apofenia ociosa o sentido que la interpretación puede sostener.

10. Alcance y límites

Correlacional, no causal (la apofenia puede producir creatividad, la actividad creativa puede entrenar la detección, o ambas depender de un tercer factor); la literatura humana se apoya en autoinforme y muestras chicas; pareidolia $\neq$ coincidencias $\neq$ ideación mágica $\neq$ apofenia, por lo que cada una se mide aparte. El brazo de modelos es un *modelo*, no evidencia sobre humanos — su valor es la separación limpia de generador y juez, imposible en una cabeza. Las afirmaciones empíricas de este paper descansan enteramente en los dos experimentos con modelos (§§5–6); el brazo humano (§7) está diseñado y pre-registrado, y su ejecución queda para trabajo futuro — así que el paper se sostiene sin él. La métrica de embeddings de la DAT (bge-m3) admite comparación de rango entre modelos, no normas humanas absolutas. No hay medida consensuada de la *calidad epistémica* de una conexión, de ahí el provisional $Q = N \times P \times E \times T$ (novedad $\times$ pertinencia $\times$ elaboración $\times$ transferibilidad). Los perfiles nunca se pre-seleccionan (§3). La lectura endolingüística de §9 es una interpretación con la que los datos son compatibles, ofrecida para situar el programa, no una afirmación demostrada sobre él.

Atribución. La endolingüística fue fundada por C. S. Meulemans y J. Á. Elias; el encuadre de la apofe-

nia y el aparato OAS son aportes del presente autor dentro de la disciplina, no una pretensión de fundarla.