

Cognitive Offloading and the Atrophy of Generative Thinking in Frequent LLM Users

Akshat Kandpal, Gurpreet Kaur*, Janhvi Gupta, Satyam Kumar Roy, Mukund Keshav

HMRITM, Delhi

Abstract- The use of large language models became widespread to complete high-level generative work including brainstorming, drafting, and constructing arguments. This paper introduced a theoretical model called Generative Atrophy through Offloading (GATO) which posited that using these types of generative tools for repeated periods of time caused a gradual decrease in one's ability to generate independently. GATO drew from existing research within cognitive offloading theory, neuroimaging research into how cognition works while engaging in creative activities with technology assistance, and controlled research studies examining the effects of technology on creativity. Empirical evidence supporting GATO included four mechanisms: decreased ambiguity tolerance, atrophy of structural synthesis processes, contraction of originality, and decreased generative self-efficacy. Three large datasets were analyzed demonstrating that generative offloading occurred frequently and consistently across multiple platforms. A cognitive resilience framework was also developed based on empirical research.

Keywords: Cognitive offloading, generative thinking, large language models.

I. INTRODUCTION

The rapid integration of large language models (LLMs) into everyday cognitive tasks marks a qualitative shift in how people interact with artificial intelligence and collaborate with AI systems in knowledge work [9].

Unlike earlier tools that mainly supported memory retrieval -search engines, calculators, spell-checkers -tools like ChatGPT, Claude, and Gemini now actively participate in the generation of ideas, arguments, and written structures. Sparrow et al. [18] documented the Google Effect, wherein users externalise factual recall; yet LLMs facilitate something more profound: the delegation of the generative thinking process itself.

Cognitive offloading -using physical or digital actions to reduce internal cognitive demand [16] - has traditionally been seen as adaptive. External memory aids free working memory for higher-order reasoning, and tool use is central to human cognitive evolution [3]. What is less explored is whether consistently bypassing the generative phase of thinking may, over time, weaken the very capacity being bypassed. Bjork and Bjork [1] showed that removing cognitive difficulty undermines long-term

retention and capability: the effort of generation is itself a mechanism of consolidation.

The scale of this phenomenon is striking. As of mid-2024, ChatGPT alone reported over 100 million weekly active users, and three large-scale interaction datasets -WildChat-1M, LMSYS-Chat-1M, and ShareGPT-90K -collectively document over two million real-world conversations in which creative writing, brainstorming, and structural drafting were the most prominent categories. These are precisely the tasks that constitute what the present study terms generative cognition.

This paper introduced GATO (Generative Atrophy Through Offloading) as a structured theoretical account of how this atrophy may unfold, supported by four bodies of literature, three large-scale usage datasets, and nine primary empirical anchors. GATO is not a technophobic argument: it is a call for cognitive intentionality in an era when the default option is to outsource the blank-page moment.

II. LITERATURE REVIEW

Prior work across four interrelated streams converged on the GATO hypothesis without resolving it. Each stream provided a distinct but complementary lens.

Cognitive Offloading Theory

Risko and Gilbert [16] established cognitive offloading as a cost-benefit process in which agents delegate mental operations to the environment when external resources are available and the internal cost of computation is high. Clark and Chalmers [3] extended this with the Extended Mind Thesis: external tools and even social partners can constitute part of the cognitive system. Critically, Grinschgl et al. [8] demonstrated that while offloading boosts task performance, it measurably reduces memory retention for the offloaded content, establishing a carry-over cost precedent. Storm et al. [19] showed that using the internet to answer one question increases the tendency to use it again for subsequent questions, a self-reinforcing loop with direct relevance to habitual LLM use.

Technology-Induced Cognitive Change

Dahmani and Bohbot [4] conducted a landmark longitudinal neuroimaging study showing that habitual GPS use reduces hippocampal grey matter density and suppresses the use of spatial memory strategies, even in contexts where GPS is unavailable. This established that tool-induced cognitive atrophy is neurologically measurable over extended timescales. Paulin et al. [15] linked semantic memory integrity directly to divergent thinking capacity via voxel-based morphometry, providing the neural substrate for why any process that weakens semantic retrieval networks would also impair original ideation.

AI Effects on Generative Thinking

Kumar et al. [12] conducted the first pre-registered RCT (N = 1,100) demonstrating carry-over atrophy from a single LLM-assisted divergent thinking

session: participants who had their ideas generated by the LLM produced measurably less original output in a subsequent unassisted session. Doshi and Hauser [6] found that AI assistance raised individual output quality but homogenised collective creative production by 5.0–5.2%, constituting evidence of originality contraction at the population level. Dell'Acqua et al. [5] showed in an RCT of 758 BCG consultants that AI access degraded performance by 19 percentage points on out-of-frontier tasks -precisely those requiring novel synthesis rather than pattern retrieval. Choi et al. [2] documented anchoring: expert annotators exposed to LLM-generated framings reproduced those framings in their own independent outputs despite perceiving themselves as uninfluenced.

Human-AI Collaboration and Design

Gerlich [7] conducted a large cross-sectional survey of a diverse occupational sample and found a correlation of $r = +0.72$ between AI tool use and self-reported cognitive offloading, alongside $r = -0.68$ between AI use and critical thinking self-assessment. While correlational, the magnitude of these associations is striking. Kosmyrna et al. [11] conducted an EEG study (N = 54, 4 sessions) and found that LLM users showed 55% reduced brain connectivity in neural networks associated with deep cognition and 83% could not recall their own AI-assisted essays, providing the most direct neural evidence of cognitive debt to date. Park et al. [14] demonstrated that redesigning LLM interaction to ask questions rather than provide answers significantly enhanced reflective thinking, establishing empirical grounding for the cognitive resilience framework.

Table 1. Comparison of Reviewed Studies

Study	Method	Sample	Key Findings	GATO Relevance
Risko & Gilbert (2016) [16]	Theoretical review	Literature survey	Offloading is cost-benefit-driven; effects of technology underexplored	Foundational mechanism
Bjork & Bjork (2011) [1]	Empirical review	Laboratory experiments	Removing difficulty suppresses consolidation; effort drives storage strength	Ambiguity tolerance erosion
Grinschgl et al. (2021) [8]	Controlled experiment	Lab participants	Offloading boosts performance but reduces memory retention	Carry-over cost precedent

Dahmani & Bohbot (2020) [4]	Longitudinal neuroimaging	Habitual GPS users	Habitual GPS use reduces hippocampal grey matter and spatial strategy use	Primary atrophy analogue
Paulin et al. (2020) [15]	VBM neuroimaging	Semantic degeneration cohort	Semantic memory loss directly reduces divergent thinking	Neural basis for originality contraction
Kumar et al. (2024) [12]	Pre-registered RCT	N = 1,100	LLM-assisted divergent thinking degraded subsequent unassisted performance	First carry-over atrophy evidence
Doshi & Hauser (2024) [6]	Online RCT	N = 300 writers; 600 raters	AI assistance raised quality but homogenised output by 5.0–5.2%	Originality contraction evidence
Dell'Acqua et al. (2023) [5]	Field RCT	N = 758 BCG consultants	AI access degraded performance 19pp on out-of-frontier tasks	Structural synthesis atrophy
Choi et al. (2024) [2]	User study	Expert annotators + LLMs	LLM framings anchored expert outputs despite perceived independence	Anchoring mechanism evidence
Gerlich (2025) [7]	Cross-sectional survey	Diverse occupational sample	$r = +0.72$ AI/offloading; $r = -0.68$ AI/critical thinking	Population-level correlational support
Park et al. (2023) [14]	Design study	Ask-oriented LLM users	Ask-don't-answer LLM design significantly enhanced reflective thinking	Cognitive resilience basis
Kosmyrna et al. (2025) [11]	EEG study	N = 54; 3 groups; 4 sessions	LLM users: 55% reduced brain connectivity; 83% could not recall own essays	Neural self-efficacy signal

Table 2. LLM Interaction Datasets Used in This Study
Dataset Source & Scale Collection Generative

Dataset	Source & Scale	Collection	Generative Evidence	GATO Relevance
WildChat-1M [17]	Zhao et al. (2024) [20]; 1M convos; 204K users	Opt-in; GPT-3.5/4; Apr 2023–May 2024	Creative writing, role play, stories, poems named as primary category	Scale of generative offloading
LMSYS-Chat-1M [18]	Zheng et al. (2023) [21]; 1M convos; 210K users; 25 LLMs	Platform; Chatbot Arena; Apr–Aug 2023	Creative writing confirmed as major category via GPT-4 annotation	Cross-model replication
ShareGPT-90K [19]	RyokoAI (2023) [17]; ~90K user-selected convos	User-initiated API sharing	Creative interactions self-selected as most share-worthy	User valuation signal

III. METHODOLOGY

The research herein involved secondary data analysis and a comprehensive literature review. Primary data collection was absent; thus, no IRB approval was mandated. The study design followed a convergent synthesis approach: independent evidence streams

were evaluated for mutual consistency with the GATO framework rather than for direct causal proof.

Literature Search Protocol

A structured search across Google Scholar, ACL Anthology, PubMed, SSRN, and OSF Preprints/arXiv (2011–2025) was conducted following PRISMA 2020 guidelines [13]. Search strings combined terms from

two clusters: (a) generative cognition, creative cognition, divergent thinking, idea generation, and writing; and (b) LLM, ChatGPT, generative AI, AI assistance, and cognitive offloading. Boolean AND was applied between clusters; OR within. No language restriction was applied: all included studies are in English. Conference proceedings, peer-reviewed journals, and preprint servers were included; blog posts, opinion pieces, and non-peer-reviewed institutional reports were excluded.

Dataset Analysis

Three LLM interaction datasets -WildChat-1M [20], LMSYS-Chat-1M [21], and ShareGPT-90K [17] -were analysed documentarily using task-category findings published by the original authors. No re-analysis of

raw conversation logs was conducted. Together they documented generative task offloading across over 2 million real-world conversations on multiple platforms and models, providing ecological validity for the claim that generative offloading is a cross-platform, large-scale behaviour.

Framework Development

GATO was developed via theoretical synthesis: existing constructs from cognitive psychology, educational psychology, and neuroimaging were mapped to proposed mechanisms, and the empirical evidence base for each mapping was evaluated for consistency, sample size, effect size, and replication status.

Table 3. GATO Framework vs. Existing Frameworks

Framework	Focus	Mechanism	Longitudinal	Gen. Cognition	Design Response
Google Effect (2011)	Memory/retrieval offloading	Partial	No	No	No
GPS-Spatial Atrophy (2020)	Spatial navigation	Yes -hippocampal	Yes	No -spatial only	No
Cognitive Debt (2025)	Neural engagement in AI writing	Partial -EEG	Partial	Yes	No
Jagged Frontier (2023)	AI task performance	Partial	No	Partial	No
Thinking Assistants (2023)	LLM design for reflection	No	No	Partial	Yes -ask design
GATO (This Paper)	Generative cognition atrophy	Yes mechanisms	-4 Proposed	Yes -primary focus	Yes -3 principles

IV. The GATO Framework

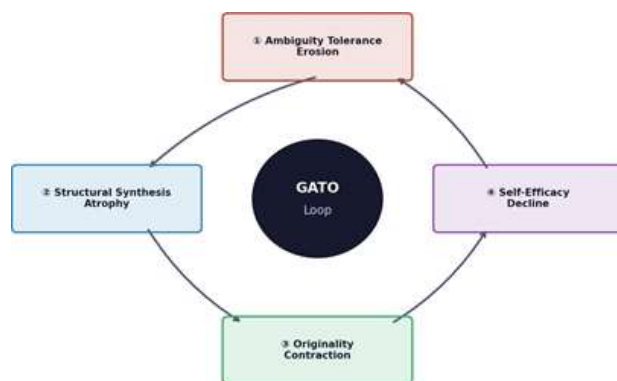


Fig. 1. GATO Self-Reinforcing Atrophy Loop

GATO proposed that habitual delegation of generative cognitive tasks to LLMs progressively eroded four interdependent capacities that formed a self-reinforcing atrophy loop. Each mechanism was

distinct, but deeply interconnected: erosion in one amplified erosion in the others. Figure 1 illustrated the loop structure.

Mechanism 1: Ambiguity Tolerance Erosion

Users who consistently invoked LLMs at the blank-page moment bypassed the consolidating effort of initiation [1]. The blank-page moment -the cognitively demanding, often uncomfortable state of confronting an unstructured problem space -was precisely where schema formation, analogical retrieval, and generative effort were exercised. Each bypass lowered the threshold at which discomfort triggered assistance-seeking, in a manner analogous to Storm et al.'s [19] self-reinforcing internet use loop. Over time, independent generation became an avoided state, not merely an effortful one. This mechanism was foundational because it governed the entry point to generative cognition: if the blank-

page moment was habitually bypassed, the remaining three mechanisms were starved of the triggering conditions that would keep them active.

Mechanism 2: Structural Synthesis Atrophy

Repeated exposure to LLM-generated framings anchored users' independent structural repertoire [2, 5]. LLMs, trained on large corpora, generated structurally modal outputs -the most statistically likely organisational patterns given a prompt. Users who regularly received and accepted these structures gradually narrowed their own repertoire toward the LLM norm, crowding out novel organisational patterns that would emerge from independent synthesis. Dell'Acqua et al.'s [5] finding -that AI-assisted BCG consultants performed 19 percentage points worse on out-of-frontier tasks - provided direct evidence of this structural impoverishment: precisely the tasks that required novel synthesis, rather than pattern retrieval, were degraded by AI exposure.

Mechanism 3: Originality Contraction

AI-assisted output converged toward LLM-generated norms at the population level [6]. While individual quality rose -the LLM raised the floor - collective diversity fell, by 5.0–5.2% in Doshi and Hauser's RCT [6]. At the neural level, Paulin et al. [15] established that semantic memory integrity directly supported divergent thinking capacity: the richer and more interconnected a person's semantic retrieval networks, the more original their ideation. If habitual LLM use reduced the active exercise of semantic retrieval, the neural substrate of originality would weaken over time, in a mechanism structurally parallel to Dahmani and Bohbot's [4] GPS-spatial atrophy finding.

Mechanism 4: Impaired Generative Self-Efficacy

The accumulation of cognitive debt eroded confidence in producing independent creative output, thereby reinforcing the externalisation cycle. Direct evidence supported this: 83% of LLM users in an EEG study could not recall their own AI-assisted essays, and brain connectivity in networks linked to deep cognition showed a 55% reduction [11]. A population-level survey further corroborated this, revealing a correlation of $r = -0.68$ between AI tool

usage and self-assessed critical thinking [7]. When individuals perceived their unassisted work as inferior -or were unable to recall it -their belief in independent generation waned, making LLM delegation increasingly appealing and completing the self-perpetuating feedback loop.

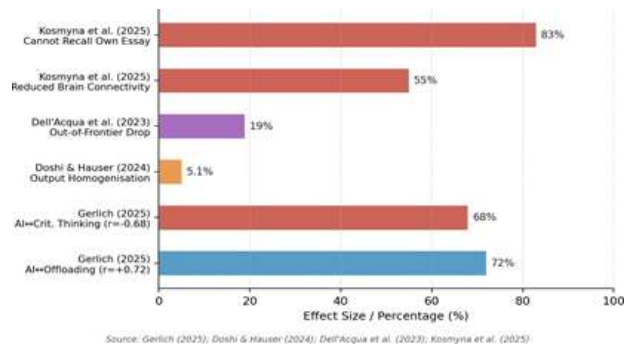


Fig. 2. Key Quantitative Findings

V. COGNITIVE RESILIENCE FRAMEWORK

GATO is not a case against AI use: it is a case for intentional AI use. The cognitive benefits of LLMs - rapid information synthesis, structural scaffolding for novice users, accessibility of expert-level language - are real and well-documented. The framework proposed here aimed not to eliminate LLM use but to preserve the generative cognitive capacities that LLMs are most likely to erode. Three empirically grounded practices can preserve generative capacity within an AI-assisted workflow.

Exercise 1: Scaffolded Generation

Before using an LLM on anything that requires generative effort, the user should take a small amount of time -even just 5 to 10 minutes -to attempt their own work first. It does not have to be polished. The goal is to experience the blank-page moment and thereby stimulate schemas that would otherwise degenerate. Bjork and Bjork [1] showed that desirable difficulty -the effort required to retrieve and generate under conditions of partial ambiguity -leads to stronger, longer-lasting memory storage than effortless completion. Even poor independent attempts, later revised with LLM support, preserve this important neurological exercise that habitual direct delegation circumvents.

Exercise 2: Interrogative AI Use

Instead of prompting an LLM to produce content, users should prompt it to ask questions, identify gaps, surface unstated assumptions, and challenge existing framings. Park et al. [14] demonstrated that this ask-don't-answer interaction design significantly enhanced reflective thinking compared to answer-provision prompting. It repositions the LLM as a Socratic interlocutor rather than a generative substitute, preserving the individual's role as the primary generator while leveraging the LLM's capacity for systematic questioning.

Exercise 3: Regular Unassisted Practice

Individuals should include unassisted generative sessions intentionally in their regular workflow, treating this practice as equivalent in importance to physical exercise for cognitive maintenance. Desirable difficulty [1] restricts the atrophy loop from consolidating into habitual avoidance. These sessions need not be long; their function is to maintain the activation threshold for independent generation below the level at which discomfort triggers assistance-seeking.



Fig. 3. Cognitive Resilience Framework

VI. DISCUSSION AND LIMITATIONS

Contributions

This paper made three significant contributions. GATO was the first framework to specify four mechanisms of generative atrophy, map them to supporting empirical data, and connect each mechanism to suitable neural substrates. Table 3 demonstrated that no prior research simultaneously satisfied all five criteria assessed. The second contribution was a thorough secondary analysis of three large-scale interaction datasets, with combined conversation data exceeding two million, which confirmed that generative offloading behaviour occurred across multiple platforms at

scale, in contrast to earlier research that treated this phenomenon as an artefact of laboratory settings. The third contribution was a cognitive resilience framework grounded directly in empirical design research rather than cautionary speculation.

Limitations

The primary constraint of this study was its inability to confirm causation. GATO assembled multiple lines of converging evidence for each component of the theoretical model; however, no long-term studies have tracked LLM users' generative capacity over time with matched control groups. Although the three identified atrophy mechanisms could be theorised based on analogous research, none had been directly studied within an LLM environment. Three of the studies cited were published online prior to peer review [11, 14, 10] and should therefore be considered preliminary findings. The correlational result reported by Gerlich [7], while striking, could not determine whether AI use caused cognitive offloading or whether individuals predisposed to offloading preferentially adopted AI tools. Finally, the framework primarily concerned individuals who habitually used AI tools for deep generative offloading; occasional users did not share the same risk profile.

Future Research Avenues

Pressing future investigations should encompass a pre-registered longitudinal study tracking divergent thinking performance over a minimum of six months in regular versus infrequent LLM users with matched control groups, neuroimaging analyses of default mode network and semantic network connectivity among habitual LLM users, task-type specificity evaluations to ascertain whether cognitive changes were broadly applicable or confined to particular generative tasks, demographic-specific vulnerability assessments exploring differential risks across student populations, creative professionals, and knowledge workers, and interaction design experiments testing the efficacy of interrogative AI interfaces in empirically mitigating the cognitive atrophy loop.

VII. CONCLUSION

The blank-page moment is not a problem to be solved: it is an exercise to be maintained. In the age of generative AI, maintaining it has become a deliberate choice rather than an unavoidable one. GATO provided a structured theoretical starting point for the empirical investigation of what LLMs augment versus what they replace in the human cognitive repertoire. The framework identified four interdependent mechanisms -ambiguity tolerance erosion, structural synthesis atrophy, originality contraction, and impaired generative self-efficacy - each grounded in converging empirical evidence and connected to a plausible neural substrate.

The convergent evidence reviewed established GATO as a theoretically coherent and empirically plausible hypothesis warranting rigorous longitudinal investigation. It must be noted, however, that the present study was limited by its theoretical nature and the absence of long-term causal data; future research should address these gaps through controlled longitudinal designs. As LLMs become more capable and more deeply embedded in cognitive workflows, the question of what they augment versus what they replace demands sustained empirical attention. The cost of not investigating it -continued expansion of habitual generative offloading without understanding its long-term cognitive consequences -is too high to accept.

VIII. DECLARATIONS

The authors declare no conflict of interest. This research received no specific funding from any public, commercial, or not-for-profit agency. The study involved secondary analysis of publicly available datasets (WildChat-1M, LMSYS-Chat-1M, and ShareGPT-90K), and no primary human subjects were involved; therefore, ethical approval was not required.

REFERENCES

1. E. L. Bjork and R. A. Bjork, "Making things hard on yourself, but in a good way," in *Psychology*

- and the Real World, M. A. Gernsbacher et al., Eds. New York, NY: Worth Publishers, 2011, pp. 56–64.
2. A. S. Choi, S. S. Akter, J. P. Singh, and A. Anastasopoulos, "The LLM effect: Are humans truly using LLMs, or are they being influenced by them?" in *Proc. EMNLP 2024*, pp. 22032–22054, 2024.
3. A. Clark and D. Chalmers, "The extended mind," *Analysis*, vol. 58, no. 1, pp. 7–19, 1998.
4. L. Dahmani and V. D. Bohbot, "Habitual use of GPS negatively impacts spatial memory during self-guided navigation," *Scientific Reports*, vol. 10, p. 6310, 2020.
5. F. Dell'Acqua et al., "Navigating the jagged technological frontier: Field experimental evidence of the effects of AI on knowledge worker productivity and quality," *HBS Working Paper 24-013*, 2023.
6. A. R. Doshi and O. P. Hauser, "Generative AI enhances individual creativity but reduces the collective diversity of novel content," *Science Advances*, vol. 10, no. 28, p. eadn5290, 2024.
7. M. Gerlich, "AI tools in society: Impacts on cognitive offloading and critical thinking," *Societies*, vol. 15, no. 1, p. 6, 2025.
8. S. Grinschgl, F. Papenmeier, and G. Jahn, "Consequences of cognitive offloading: Boosting performance but diminishing memory," *Quarterly Journal of Experimental Psychology*, vol. 74, no. 11, pp. 1927–1943, 2021.
9. J. Hutson, "Human-AI collaboration in writing: A multidimensional framework for creative and intellectual authorship," *International Journal of Changes in Education*, 2025.
10. A. Jadhav, "Distributed atrophy 2.0: Cognitive offloading and the reshaping of thought," *OSF Preprints*, doi: 10.31234/osf.io/8yr92_v2, 2025.
11. N. Kosmyna et al., "Your brain on ChatGPT: Accumulation of cognitive debt when using an AI assistant for essay writing," *arXiv preprint arXiv:2506.08872*, 2025.
12. H. Kumar et al., "Human creativity in the age of LLMs: Randomized experiments on divergent and convergent thinking," *arXiv preprint arXiv:2410.03703*, 2024.

13. M. J. Page et al., "The PRISMA 2020 statement: An updated guideline for reporting systematic reviews," *BMJ*, vol. 372, p. n71, 2021.
14. S. Park, H. Subramonyam, and C. Kulkarni, "Thinking assistants: LLM-based conversational assistants that help users think by asking rather than answering," *arXiv preprint arXiv:2312.06024*, 2023.
15. T. Paulin et al., "The effect of semantic memory degeneration on creative thinking: A voxel-based morphometry analysis," *NeuroImage*, vol. 220, p. 117073, 2020.
16. E. F. Risko and S. J. Gilbert, "Cognitive offloading," *Trends in Cognitive Sciences*, vol. 20, no. 9, pp. 676–688, 2016.
17. RyokoAI, "ShareGPT52K/ShareGPT90K," *HuggingFace*, 2023.
18. B. Sparrow, J. Liu, and D. M. Wegner, "Google effects on memory: Cognitive consequences of having information at our fingertips," *Science*, vol. 333, no. 6043, pp. 776–778, 2011.
19. B. C. Storm, S. M. Stone, and A. S. Benjamin, "Using the internet to access information inflates future use of the internet to access other information," *Memory*, vol. 25, no. 6, pp. 717–723, 2016.
20. W. Zhao et al., "WildChat: 1M ChatGPT interaction logs in the wild," in *Proc. ICLR 2024*, 2024.
21. L. Zheng et al., "LMSYS-Chat-1M: A large-scale real-world LLM conversation dataset," *arXiv preprint arXiv:2309.11998*, 2023.