

The Unread Archive

Why the most valuable documents in your life and business are the ones nobody has read end to end — and what changes when the archive finally becomes answerable.

A position paper on document intelligence, verifiable AI, and the shift from storing documents to reasoning over them. Drawing on the public thinking of *Andrej Karpathy*, *Sam Altman*, and *Jensen Huang*, and on the 2024–2026 research literature on retrieval, long context, and citation reliability.

Read nothing. Know everything.

DocuStrata, LLC · docustrata.com · Every answer, checkable.

CONTENTS

-
- 01 Executive summary
 - 02 The document you keep forever and read never
 - 03 Why storage was never the answer
 - 04 A new mental model: the LLM as operating system
 - 05 The archive as working memory
 - 06 The retrieval debate
 - 07 The verifiability problem
 - 08 The enterprise stakes: a governed digital workforce
 - 09 What an answerable archive requires
 - 10 The DocuStrata position
 - 11 The strongest objections
 - References & attribution
-

SECTION 01

Executive summary

THE ARGUMENT IN BRIEF

Every person and every business keeps a class of document — leases, policies, contracts, filings, plan documents — with great care and reads it almost never. The information is preserved but inert. The cost surfaces only on the day a missed clause or forgotten deadline turns expensive.

For decades our answer was storage: digitize, index by keyword, move to the cloud. But keyword search only finds what you already know to look for, and a result is still a document you must then read. Storage made the unread archive larger, not readable.

A genuine shift is now underway, articulated most clearly by three figures. Andrej Karpathy reframes the model as an operating system whose scarce resource is the context window, and argues for compiling a knowledge base once and keeping it current rather than re-deriving it each query. Sam Altman describes a future in which a small reasoning model carries a lifetime of personal context. Jensen Huang frames the enterprise form: a digital workforce that must be governed like a human one.

Their synthesis is this paper's thesis: an archive that is reasoned-over, always-current, verifiable, and governed stops being a filing cabinet and becomes working memory. The load-bearing requirement is verifiability — every answer must cite its exact source, checkable in one click — because in 2026 even frontier models still fabricate in roughly one in ten single-document summaries. DocuStrata is built on that single promise: every answer is checkable.

The remainder of this paper develops each step of that argument, examines the technical debate beneath it, states the requirements an answerable archive must meet, and confronts the strongest objections to the position honestly.

SECTION 02

The document you keep forever and read never

There is a category of document in every life and every business that is kept with great care and read almost never.

The signed lease. The insurance policy. The plan document behind a retirement account. The operating agreement. The closing binder. The vendor contract with the auto-renewal clause. These are not trivial papers — they govern money, obligations, deadlines, and risk. And yet the median number of times any

one of them is read cover to cover, after the day it was filed, rounds to zero.

This is not negligence. It is economics. Reading a forty-page plan document to answer one question is a poor trade, so the question goes unasked and the document goes unread. The cost is invisible until the day it is not: a missed clause, an expired deadline, a term nobody remembered agreeing to, an asset nobody knew was covered.

The most expensive sentences in a business are the ones nobody read.

Multiply one governing document by every one a person or business accumulates over a lifetime and you arrive at a strange, universal condition: we are all sitting on a second library. Not the paperwork we are forced to keep, but the archive we built on purpose — the research, the references, the saved articles, the scanned filings — and almost none of it is ever read again. It is saved, and never used.

SECTION 03

Why storage was never the answer

For the entire history of computing, our answer to the unread archive has been storage. We digitized the filing cabinet. We made it searchable by keyword. We moved it to the cloud, then to faster clouds. Each step made the archive cheaper to keep and easier to accumulate. None made it easier to read.

Keyword search, the great advance of the last generation, has two structural limits that no amount of scale removes. First, it only finds documents you already know to look for: you must supply the right term, which means you must already know roughly what the document says. Second, a search result is not an answer — it is a document, which you must then open and read. Search narrowed the pile; it never removed the reading.

The three things storage never gave us

A system that actually addresses the unread archive must do three things storage cannot. It must answer questions rather than return documents. It must show its work, so an answer about money or law can be trusted. And it must stay current as the archive grows, without a manual re-indexing ritual that guarantees the archive falls out of date. Storage does none of these. It holds. Holding is not reading, and it never was.

What changed — recently, and suddenly — is that a different kind of system became possible. The change is best understood not through a product announcement but through the people who have articulated the underlying shift most clearly. Three voices map the territory.

SECTION 04

A new mental model: the LLM as operating system

Andrej Karpathy — a founding member of OpenAI and the former director of AI at Tesla — argues that large language models should not be understood as chatbots at all, but as a new kind of operating system. The model weights are the CPU: the fixed processing substrate. The context window is RAM: the finite, active working memory where all reasoning happens. External tools and data are peripherals the system reaches for. The precision of the analogy is the point — it changes what you build and what you stop building.

From it follows his most influential recent argument: the real discipline is no longer prompt engineering but context engineering. If the context window is RAM, then what you load into it is your program, and the craft that matters is deciding what to load.

“The delicate art and science of filling the context window with just the right information for the next step.”

✓ **Andrej Karpathy** — on context engineering, 2025

The word delicate carries weight. The research confirms it: a model's attention degrades as its context fills, with the sharpest losses for information stranded in the middle of a long context. More is not better. The right information, placed well, beats a larger pile every time.

The LLM Wiki: compile once, keep current

Karpathy extends this into a working pattern he calls the LLM Wiki. Rather than re-searching and re-deriving knowledge on every question, you compile source material once into a structured, interlinked form, and keep that structure current. The knowledge, in his framing, is compiled once and kept current, not re-derived on every query.

It is worth pausing on how close this is to a literal specification for a document intelligence system. An archive compiled once into an answerable, interlinked form and kept current as it grows is exactly what the unread archive has always needed.

The honest counterweight: the decade of agents

Karpathy is also the field's most credible skeptic of its own hype, and an honest paper must carry that alongside the optimism. In his October 2025 conversation with Dwarkesh Patel, he argued against the industry's 'year of agents' in favor of a 'decade of agents.' Today's systems, he said, still lack reliable memory, robust computer use, and full multimodality.

"It's the decade of agents. They just don't work. They don't have enough intelligence, they're not multimodal enough, they can't do computer use and all this stuff."

✓ **Andrej Karpathy** — on the Dwarkesh Patel podcast, October 2025

This is not a contradiction of the optimistic case; it is its discipline. His most counterintuitive point sharpens the whole thesis: agents fail when they go off the manifold of what exists on the internet — when a situation is genuinely specific to you. His prescription is a lean 'cognitive core' paired with external, structured memory that compounds.

The intelligence is not the bottleneck. Your specific context is.

Read carefully, this is the most important claim for anyone holding an archive. The frontier model is not what stands between you and an answer about your retirement plan — it has never seen your plan. What stands between you and the answer is whether your specific, private, off-manifold documents are organized into something a lean reasoning core can reliably reason over. That is a document problem.

SECTION 05

The archive as working memory

Sam Altman, the CEO of OpenAI, describes the destination in more personal terms. Asked at a Sequoia event how an AI system becomes genuinely personalized, he described an ideal in which a deliberately small reasoning model carries an enormous, continuously growing context drawn from a person's entire life.

“A very tiny reasoning model with a trillion tokens of context that you put your whole life into — every conversation you've ever had, every book you've ever read, every email — and your life just keeps appending to the context.”

✓ **Sam Altman** — at a Sequoia event, 2025

Notice what Altman does not emphasize. The reasoning model is ‘very tiny.’ The intelligence is almost an afterthought. The entire weight of the vision rests on the context — on getting a whole life's worth of material into a form the model can reason over, and keeping it current. The bottleneck is not cognition. It is organized, retrievable context.

This is the same conclusion Karpathy reaches from the opposite direction. He strips the model to a cognitive core and puts the weight on external structured memory; Altman imagines a tiny model and puts the weight on a trillion tokens of life. Both arrive at the same place: the prize is not a bigger brain, it is a better-organized, always-current, reasoned-over archive. The model is the reader. The archive is what gets read.

And it scales down as cleanly as up. You do not need a trillion tokens of your whole life to feel the effect. You need the forty-page plan document, the decade of scanned filings, the Evernote export you have carried between apps for years. The moment that archive becomes answerable, it stops being storage and starts being memory.

SECTION 06

The retrieval debate

Beneath these visions sits a genuine technical debate that determines whether an answerable archive is reliable or merely impressive. It is between two ways of getting your documents in front of a model: retrieval-augmented generation (RAG), which finds and injects the most relevant passages at question

time, and long-context, which loads large amounts of material directly into the model's expanding window.

The naive assumption is that long context wins as windows grow — just load everything. The research runs the other way. Long-context models face inference costs that scale poorly and, more importantly, their accuracy degrades as the window fills, especially in the neglected middle. Stuffing more into the window does not produce better answers; it often produces confident, coherent-sounding wrong ones.

The complementary truth

The mature 2026 consensus is not that one approach won, but that they are complementary — and the split maps almost perfectly onto a document archive's needs. Long context is strong when evidence is spread across many documents and you need holistic reasoning. Retrieval is strong when evidence is sparse and precise, and, critically, when you need to cite an exact source. Retrieval also has two properties that matter enormously for a private archive: it works over up-to-date and private information without retraining the model, and it can point to the exact passage an answer came from.

“What did we discuss last meeting? → long context. Cite the exact clause about termination? → retrieval, for precision and citations.”

✓ **The emerging 2026 architecture** — synthesized from the retrieval literature

For an archive of governing documents, the retrieval property is not a nicety. It is the difference between a system you can act on and one you cannot. When the answer concerns money or a legal obligation, ‘the model is fairly sure’ is worthless. ‘Here is the exact sentence, on this page, that says so’ is everything.

SECTION 07

The verifiability problem

A reasoning system that produces answers must be able to prove them. Without that, it is not an intelligence tool. It is a rumor engine with good grammar.

The problem is well documented and does not vanish with scale. Even in the simplest task — summarizing a single document placed directly in front of the model — frontier models still fabricate, inventing content in roughly one in ten single-document summaries as of 2026. And there is a subtler, more dangerous failure specific to systems that do cite: citation hallucination, where a model attaches a source to a claim the source does not support. The citation looks like verification. It is the opposite — a confident pointer to the wrong place.

Reasoning without citation is not intelligence; it is a rumor. A wrong citation is worse — it is a rumor wearing a badge.

The response cannot be to trust the model harder. It must be to make every answer checkable in one click — to place the exact source passage next to the answer, so the human verifies in seconds rather than trusting on faith. This catches the fabrications, because a claim with no real supporting passage has nowhere to hide; and it changes the relationship between human and machine: the human is no longer asked to believe, only to check. For documents that govern money and law, checkable-in-one-click is not a feature. It is the whole proposition.

This is also why the retrieval property from the previous section is load-bearing. A system that retrieves the exact passage can show you that passage. A system reasoning from a stuffed context window often cannot tell you which of a hundred thousand tokens produced its answer. Verifiability is not bolted on at the end; it is a consequence of how the archive is built.

SECTION 08

The enterprise stakes: a governed digital workforce

Jensen Huang, the CEO of NVIDIA, frames what happens when this capability meets the enterprise. At CES 2025 he declared the arrival of agentic AI in unambiguous terms.

“AI agents are the new digital workforce. The age of AI Agentics is here.”

✓ **Jensen Huang** — CES 2025 keynote

Huang's estimate of the opportunity runs to the trillions, but the part most relevant here is not the scale — it is the insistence on governance. A digital workforce that reads and acts on a company's documents cannot be turned loose without the same controls a human workforce operates under.

“You have human agents and AI agents, and they basically govern the same way — identity management, access controls, and governance policies.”

✓ **Jensen Huang** — on agent governance, 2026

For a document archive, this translates into hard requirements. If an AI system can read and answer over a company's contracts, filings, and financials, then access must be controlled, actions must be auditable, and — the non-negotiable one — the documents must never become training data for a model that other parties will use. The archive reasons over your documents. It does not surrender them.

Here the three thinkers converge into a single design brief. Karpathy: the machine is an engine for reasoning over loaded context, and the specific, off-manifold context is the bottleneck. Altman: the prize is putting a whole life or business into that context. Huang: at scale it works only if access is governed. An archive that is reasoned-over, always-current, verifiable, and governed is no longer a filing cabinet. It is institutional memory with a reader attached.

SECTION 09

What an answerable archive requires

Synthesizing the argument, an archive that is answerable rather than merely stored must satisfy five properties. Each follows from one of the observations above; each is a requirement, not a preference.

1 Every answer is checkable.

The single most important property. Every answer cites the exact source passage it came from, verifiable in one click — defeating both plain fabrication and citation hallucination, and converting the human's job from believing to checking.

2 **The archive stays current on its own.**

Compile-once-keep-current only holds if ‘keep current’ is real. New documents flow in and become answerable without a manual re-indexing ritual. A stale archive is a distrusted archive, and a distrusted archive goes unread.

3 **Retrieval is precise, not just capacious.**

Because attention degrades as context fills, and precise citation requires knowing exactly which passage produced an answer, the system retrieves the right material rather than stuffing in everything. Capacity is not accuracy.

4 **Access is governed.**

Identity, access control, and auditability, applied to the reasoning system the way they are applied to human staff. Who can ask, what they can reach, and a record of what was answered.

5 **Documents are never used to train models.**

The contents are sensitive by definition. They are encrypted at rest and in transit, and never become training data. The reasoning engine is general; the archive it reasons over stays private to its owner.

SECTION 10

The DocuStrata position

DocuStrata is built on a single promise that follows from the requirement outranking all others: every answer is checkable. It is a document intelligence system that ingests an archive — an Evernote export, a cloud drive, a decade of scanned filings — makes it fully searchable and answerable, and cites its sources on every response. It is designed so the expensive sentences finally get read: not by forcing a human to read forty pages, but by letting them ask the page a question and get a cited answer back.

The design choices line up one-to-one with the five requirements. Answers cite their exact source passages, verifiable in a click. Cloud connectors and watch folders keep the archive current as it grows. Retrieval is built for precision and citation. Access is authenticated and scoped to the owner. Documents are encrypted, private, and never used to train AI models.

In the language of the three thinkers: DocuStrata is the structured, always-current external memory Karpathy's cognitive core needs; the organized, appended context that makes Altman's tiny reasoning model useful on your life rather than the internet's; and the governed, private substrate Huang's digital workforce requires before it can be trusted with a company's documents.

This is what it means to give an archive a reader for the first time.

Not a better search box. A system that has, in effect, read everything you filed, and can answer for it — with receipts. The second library you built on purpose, finally readable. The forty-page plan document, finally answerable in a sentence. The unread archive, at last, read.

KEY TAKEAWAYS

- 1 The unread archive is a real, universal problem: valuable documents kept forever and read never, at quiet and recurring cost.
- 2 Storage never solved it. Answering, showing work, and staying current are what the problem actually required.
- 3 The bottleneck is not model intelligence — it is your specific, private context, organized into something a model can reason over.
- 4 Verifiability is load-bearing: every answer must cite its exact source, because models still fabricate.
- 5 DocuStrata's promise is exactly that: every answer, checkable — over an archive that stays current, governed, and private.

SECTION 11

The strongest objections

Three are worth stating at full strength: that the models still fabricate, that genuinely useful agents remain a decade away, and that none of this is new. Here is each, and what survives it.

OBJECTION 01

‘The models still hallucinate. Why trust any of this?’

Correct — which is precisely why verifiability is load-bearing rather than an afterthought. The claim is not that models stopped fabricating; the 2026 evidence says they have not. It is that a system built so every answer carries its exact source turns the model's fallibility from a hidden risk into a visible, checkable one. You are not asked to trust the model. You are asked to check its citation, which takes seconds. The honest claim is not ‘never wrong’ — it is ‘never wrong in a way you cannot immediately catch.’

OBJECTION 02

‘Karpathy says useful agents are a decade away.’

He does, and he is right about autonomous agents — systems you hand a goal and leave alone. But an answerable archive is not an autonomous agent. It does not act on the world; it answers questions about documents, with a human reading the cited result. The hard, unsolved problems Karpathy names — planning, durable autonomous memory, robust computer use — are the problems of acting without supervision. Answering a question about a document you already hold, and showing your source, is far more tractable, and available now.

OBJECTION 03

‘Isn't this just RAG with a nicer interface?’

Retrieval is a component, not the thesis. The thesis is that the unread archive is a real and universal problem, that verifiability is what makes a solution trustworthy, and that keeping the archive current and governed is what makes it durable. Retrieval is one honest technique in service of that, chosen because it supports exact citation and works over private data without retraining. The interface matters because a verification that takes one click gets used and one that takes ten does not — but the interface is downstream of the commitment: every answer, checkable.

APPENDIX

References & attribution

- Karpathy, A. — the LLM-as-operating-system framing (model = CPU, context window = RAM) and ‘context engineering’, Sequoia AI event and public writing, 2025; the LLM Wiki knowledge-base pattern, 2025–2026; the ‘decade of agents’, ‘cognitive core’, and compounding-memory arguments, interview with Dwarkesh Patel, October 2025.
 - Altman, S. — the trillion-token personal-context vision, Sequoia event, May 2025, as reported by TechCrunch and others.
 - Huang, J. — ‘AI agents are the new digital workforce’ and ‘the age of AI Agentics is here’, CES 2025 keynote; agent-governance remarks, 2026.
 - Research literature — retrieval-augmented generation versus long-context reasoning and their complementary strengths; context-length accuracy degradation; citation-hallucination detection; and public single-document hallucination measurements; drawn from the 2024–2026 arXiv literature.
-

Quotations are attributed to their speakers and drawn from publicly reported statements. Ideas are developed and extrapolated by DocuStrata from these public sources and the cited

research literature. This paper presents DocuStrata's interpretation and does not imply endorsement by any person quoted.

DocuStrata, LLC

docustrata.com · Every answer, checkable.