

The Right to Remain Unmediated

Reversibility, Delegation, and the Contraction of the Personal

Flyxion

August 12, 2026

Abstract

Three recent developments in consumer artificial intelligence appear, on their surface, to concern different products and different complaints. A major social platform has embedded an AI assistant inside an advertising-funded feed that shows little interest in reducing how much time a person spends inside that feed. The chief executive of the same company has published a vision of universal personal agents that manage relationships, health, finances, hobbies, and household life on a person's behalf. A competing AI company has proposed generating a personalized news podcast for the drive to school, synthesized from a family's calendar and a child's declared interests. This essay argues that these are not three complaints but one, examined at three different scales, and that the underlying failure has a precise structure. Drawing on a distinction between architectural, behavioral, and institutional mediation, the essay develops the idea that assistance becomes dangerous not when it is extensive but when it becomes difficult to undo. The essay imports a construction from an earlier work on refusal, exclusion distance, and specializes it into a recovery margin between a person's present, mediated relationship to an activity and their practical capacity to return to an unmediated one. The result is a design criterion for personal artificial intelligence considerably narrower than either uncritical adoption or blanket suspicion: preserve the paths back.

1 Three Complaints That Are One Complaint

Three things happened within the same season, and on the surface they have nothing obvious in common.

A widely used social platform added an artificial intelligence assistant to its feed. The assistant answers questions, offers suggestions, and occasionally interrupts scrolling with something that resembles help. It does not, however, do the one thing a genuinely competent personal assistant embedded in that particular environment might be expected to do first: tell the user they are done, that nothing further in the feed deserves their attention today, and that they should close the application. The platform's own advertising business depends on the opposite outcome. The assistant is capable, in a narrow technical sense, and almost entirely useless in the sense that matters, because its intelligence has been placed inside an objective that is not the user's.

Around the same time, the chief executive of that same company published a long essay describing a future in which everyone has what he called a personal superintelligence: an agent that understands a person's history, goals, relationships, and preferences well enough to work continuously on their behalf across nearly every domain of ordinary life. Health, career, finances, home management, hobbies. The vision is presented as liberation. It reads, on closer inspection, as something closer to full-service management of a human life by a system whose deeper commercial incentives were never disclosed, and whose success is measured, in part, by how much of that life it comes to occupy.

And separately, the head of a rival AI company offered what he described as a compelling use case for his own product. Connect a family's calendars. Tell the assistant what the children are interested in. Let it generate a personalized podcast for the ride to school, weaving in a soccer game later that afternoon, a sibling's birthday, and a story from the news a child might enjoy. The example was offered warmly, as an illustration of how thoughtful automation could enrich family life. Much of the immediate public reaction amounted to some version of a single, unimprovable rejoinder: that is what the drive to school is already for.

Read individually, these are complaints about three different products, made by three different critics, for three different reasons. One is a complaint about advertising incentives. One is a complaint about the scope of a corporate vision statement. One is a complaint about a specific and slightly absurd proposed feature. It would be easy to file each under its own heading and move on.

This essay argues that filing them separately is a mistake, and that the mistake matters. Underneath the feed, the manifesto, and the podcast lies a single structural question, and the question is not whether artificial intelligence is helpful. In each of the three cases, the systems in question might well be helpful, in the narrow sense of producing a result the person would

have accepted if asked in isolation. The question is what happens to a person's relationship to their own life once a sufficiently capable intermediary has been placed between intention and action, repeatedly, across enough domains, for long enough. That question does not resolve to a verdict about artificial intelligence in general. It resolves to a much older and more specific question about what assistance is allowed to cost, and in particular, about what it costs when it cannot be undone.

The argument proceeds in stages. It begins architecturally, with the platform feed, because the feed offers the cleanest illustration of a system that destroys a person's ability to reconstruct where they have already been. It continues behaviorally, with the vision of a universal personal agent, because that vision exposes what happens when delegation, applied repeatedly across a widening set of domains, begins to resemble something closer to surrender than assistance. It then turns institutional, using the podcast example to ask a harder and more specific question: for whom is a nominally personal agent actually optimizing, and what happens when an activity's value cannot be separated from the fact that a person, rather than a system, performed it. Finally, the essay proposes that all three layers are instances of a single underlying quantity, one borrowed and adapted from an earlier formal treatment of refusal, and that this quantity gives the argument something more precise than a complaint. It gives it a criterion.

2 The Feed as a Category Change

Something changed in the transition from the early web to the contemporary social platform, and it is tempting to describe that change as a matter of degree. Pages became more interactive. Layouts became more dynamic. Content became more personalized. Under that description, the modern feed is simply an elaborated version of what came before, distinguished mainly by scale and sophistication. This section argues that the more important change was not a matter of degree at all. The object presented to the user changed category.

A document, in the relevant sense, is something whose contents and internal relationships remain sufficiently stable over time that a reader can return to it, cite a specific passage, copy it, modify it, archive it, or resume reading where they left off. None of this requires that documents be permanent or unalterable. Documents are edited, revised, and deleted constantly, by their authors and sometimes by their custodians. What distinguishes a document is not immutability but a working expectation of persistence and addressability: that the thing referred to yesterday can, absent some specific act of deletion, still be referred to today, in more or less the same form, at more or less the same place. A book on a shelf, a saved email, a file in a directory, a printed page, all share this expectation. It is not a law of nature. It is closer to a norm that ordinary documents are built to honor.

Write this expectation informally as

$$D(t + \Delta t) \supseteq D(t),$$

meant not as a strict mathematical property but as a statement of the persistence ideal against which a document is judged: whatever was there before should still be reachable, unless something specific and attributable happened to remove it.

A ranked social feed makes no corresponding promise. What a feed presents at any moment is better understood as a computation performed at the moment of access than as an object being retrieved. Reloading a feed does not, in general, ask the system to hand back the same arrangement of content the user was looking at a moment ago. It asks the system to compute a new one, using whatever signals are available at that instant: an updated pool of content, updated advertising constraints, updated behavioral inferences, and whatever ranking function the platform is currently running. Write this as

$$F_{t+1} = R(U, C_{t+1}, A_{t+1}, M_{t+1}),$$

where R is the ranking function, U is the user, C is the pool of available content, A represents advertising and commercial constraints, and M represents whatever other signals the system is incorporating. Nothing in this expression requires that

$$F_t \subseteq F_{t+1}.$$

There is no structural reason the feed a person sees now should contain, or even relate transparently to, the feed they saw a minute earlier. Continuity, where it exists, is incidental to the ranking function rather than a property the system is obligated to preserve.

The consequence is ordinary enough that most people have experienced it without naming it. A person is halfway through reading something on a feed. They switch to another application, or the feed refreshes on its own, or they scroll too far and pull to refresh out of habit. When they return, the item is gone. Not deleted in any formal sense; somewhere in the platform's infrastructure the content presumably still exists. But the person's access to it has been severed. They may remember the author's name imperfectly. They may not know when it was posted. They have no stable address for the particular position they occupied in that stream, no bookmark, no page number, nothing to cite. Nothing has technically been destroyed, and yet the person's informational state has been destroyed all the same. What they had a moment ago, a located, re-findable relationship to a specific piece of content, is simply gone, and no error message announces the loss because no error has occurred. The system worked exactly as designed.

This is the sense in which the older affordances of plain HTML deserve to be treated as evidence rather than nostalgia. It would be a mistake to romanticize the early web as some lost paradise of user control; it was often crude, ugly, and technically limited in ways the modern web has genuinely improved upon. But it is worth noticing, without sentimentality, what even a fairly primitive document format made available to an ordinary user by default. A person could select a font and a size. They could write text, structure it with headings, insert an image, embed an object, construct a table, link to another page, or link to a specific position within one. They could view the source of what they were looking at, copy it, save it, move the file to another machine, and host it somewhere else entirely, independent of whoever had originally published it. None of this required special technical standing. It was simply how the format worked. Contemporary social platforms, by contrast, frequently restrict font control, text sizing, layout, embedding, exporting, and copying, not because these capabilities are technically difficult to provide, but because the platform benefits from a user who is looking at exactly what the platform wants them to look at, with no easy way to reshape, extract, or carry that material elsewhere. The comparison is not between an old system that offered everything and a new one that offers nothing. It is between two categories of object: one substantially under the reader’s control, and one substantially under the publisher’s.

The persistence ideal described so far is worth testing against a system that was actually engineered to achieve it, rather than one that merely happens to resemble it by not being actively rewritten. Version control systems, and Git in particular, offer a useful comparison, though the comparison needs a small correction before it can do its work honestly. Git does not literally guarantee $D(t + \Delta t) \supseteq D(t)$ under all circumstances; history can be rewritten, unreachable commits are garbage collected, reflogs expire, and repositories can simply be deleted. What Git does extraordinarily well is something narrower and more interesting: it makes addressable historical states a first-class part of the system’s architecture, rather than an accidental residue of whatever happens to still exist. A repository does not contain only its current state D_t . It contains a directed history of states, $D_0 \rightarrow D_1 \rightarrow \dots \rightarrow D_t$, with each retained commit content-addressed, so that

$$\forall D_i \in H_t, \quad \text{retrieve}(\text{SHA}_i) = D_i,$$

for every commit still reachable in the repository’s history. The system does not merely remember that something changed. It gives the user operations over the change itself: comparing D_i against D_j , branching from an earlier state, reverting a specific change, or searching the history by bisection to find the exact transition where some property, a bug, most famously, first appeared, reducing what could otherwise be an unbounded investigation to roughly $O(\log n)$ tests. A defect can hide completely in the present state of a file while

remaining fully visible and locatable in that file’s history. This is a genuinely remarkable property, and it deserves a precise name: history, in a system built this way, is not merely archived. It remains computationally interrogable.

Set this beside the feed. Something notable happens in a person’s informational environment at a particular moment. They reload. The platform computes F_{t+1} . Unless the platform independently chooses to expose its own historical ranking state to the user, which almost none do, there is no equivalent available to the person of

$$\text{checkout}(F_t), \quad \text{much less} \quad \text{diff}(F_t, F_{t+1}).$$

This does not mean the platform itself has lost the information. Quite the opposite. Platforms retain enormous quantities of telemetry about exactly what a user saw and when. But retention by the system and recoverability by the person are not the same property, and collapsing them is exactly the mistake this essay has been arguing against from its opening pages:

retained by the system $\neq$ recoverable by the user.

A platform can possess everything necessary, internally, to reconstruct precisely what a user saw an hour ago, while giving that user no practical path to reconstruct it themselves. This distinction anticipates the recovery margin this essay develops formally later on. Technical recoverability and personal recoverability are different quantities, and a system can score arbitrarily well on the first while scoring at zero on the second.

There is a second, related engineering memory worth invoking here, from an earlier and cruder era of computing. Anyone who used a personal computer in the middle 1990s learned, usually the hard way, never to switch the machine off while it was writing to a disk. The failure mode was not merely losing new work. Depending on the filesystem and the exact operation interrupted, a user could be left with a partially written file, or worse, corrupted filesystem metadata affecting files that had nothing to do with the interrupted write. The danger was that the transition between one valid state and the next was not safely bracketed:

$$x_{\text{old}} \xrightarrow{\text{write}} ??? \rightarrow x_{\text{new}},$$

with the machine’s power vanishing somewhere in the middle, leaving neither state intact. Journaling filesystems, transactional databases, copy-on-write storage, and atomic rename patterns all attack a version of this same problem, and they attack it with a shared discipline: refuse to treat an incompletely performed transition as equivalent to a successfully completed

one. The ideal they aim for is

$$x_{\text{old}} \xrightarrow{\text{transaction}} \begin{cases} x_{\text{new}}, & \text{commit,} \\ x_{\text{old}}, & \text{abort,} \end{cases}$$

never the silent third outcome, $x_{\text{old}} \rightarrow x_{\text{half-written}}$, that Windows 95 users learned to fear by instinct.

Transactions do this. Journals do this. Version control does it across much longer histories. Backups do it across catastrophic failure. Undo stacks do it inside individual applications. Database systems using multi-version concurrency control do it for simultaneous readers. Copy-on-write does it at the level of the storage itself. None of this is obscure or experimental. Recoverability, in one form or another, is one of the oldest and most thoroughly solved engineering disciplines in computing, which makes the design of a contemporary social feed look considerably stranger by comparison. The same industry that has spent decades ensuring a database transaction interrupted halfway through cannot destroy the previous valid state has, in its consumer-facing products, normalized an interface in which a person can lose the informational state they were occupying because their thumb moved three millimeters too far.

It would be too strong, and not quite fair, to say that platforms simply declined to apply what version control already knows. A feed is not a repository, and preserving every ranking state a system ever generates for every user carries real consequences for storage, privacy, consistency, and moderation that a source code history does not have to contend with in the same way. The defensible version of the claim is narrower, and it is stronger for being narrower:

The absence of user-recoverable history should not be mistaken for a technical inevitability. Computing already possesses a large and mature vocabulary of mechanisms for preserving prior states, coordinating concurrent change, and recovering from interrupted transitions. Whether any of those mechanisms are exposed to the person using the system is a design decision, and, beneath the design decision, an institutional one.

That last sentence returns the argument directly to the three layers this essay will spend the next several sections developing. Architecture determines whether recovery mechanisms exist for a given system at all. Behavior determines whether a person continues to know how to operate without whatever mediated state they have grown used to. Institutions determine who is actually permitted access to the history a system has, in fact, retained. A platform can therefore be, simultaneously, superb at internal recoverability and almost entirely without external recoverability, which may be the least visible and most consequential version of the

pattern this essay is tracing:

The system remembers everything. The person is permitted to recover almost nothing.

Reversibility and a second, older idea are worth distinguishing before moving on, because the two are frequently confused and this essay has so far treated only the first. Reversibility asks whether a prior state can be returned to. Auditability asks something narrower and, in one sense, humbler: when a state cannot be returned to, whether enough evidence survives to establish how the present state came about. The distinction is ancient, and a useful case of it predates computing by roughly three thousand years. In a story recounted in the Hebrew Bible, David, having the opportunity to kill Saul unseen in a cave, instead cuts a corner from Saul's robe and later displays it as proof of restraint rather than of nothing. The fragment does not reverse anything; Saul cannot be returned to a state of not having been vulnerable. What the fragment does is make a particular account of the past verifiable, because the cut corner and the robe's new gap correspond, and a false account would have to explain that correspondence away. Modern cryptographic signatures formalize the same intuition in different material. The relevant asymmetry is not that a sender knows more than a receiver; it is that a signature is easy to produce given privileged access and easy to verify without it, so that verification never requires reacquiring the capability that produced the thing being verified. Git supplies one answer to the underlying question, here is the retained history itself. A cryptographic signature supplies another, here is evidence that someone with a particular capability attested to this. The robe supplies the oldest and most primitive answer of all, here is a surviving piece of the interaction. All three are attempts to solve the same problem: how the present can contain verifiable evidence of an event that has already disappeared. This essay is primarily concerned with reversibility, not auditability, and the two should not be collapsed into each other. But it is worth knowing, going forward, that when a recovery margin has gone negative and a person genuinely cannot return to an unmediated relationship with some activity, auditability is the fallback that remains available even then: not restoration, but an honest and inspectable record of how the distance came to be so large.

That fallback, however, depends on one condition the examples above quietly assume: that whoever is checking the evidence is not the same party who produced it. Saul could inspect David's fragment because Saul was not the one who had cut it. A signature can be verified by anyone holding the public key, precisely because verification does not run back through the signer. Git's history can be interrogated by a developer who was not present for the commit in question. Remove that independence and auditability stops doing any work at all. A record kept, curated, and checked by a single party is not evidence about that party's conduct; it is that party's account of its own conduct, restated. If an institution

controls both the archive and the only available method of consulting the archive, then every appeal to the record simply cites the record, which was produced by the thing being checked, and the loop closes on itself: the archive confirms the archive, the source refers back to the source, and no position remains from which a different account could even be articulated, let alone believed. This is a stronger failure than the loss of reversibility this essay has otherwise been concerned with. A person who cannot return to a prior state can still, in principle, know that a prior state existed and differed from the present one. A person facing a self-auditing record may not even retain that much, because the only surviving description of the past comes from the party with every incentive to describe it favorably. Auditability, in other words, is not a property of records considered in isolation. It requires a witness whose position is not itself supplied by the thing being witnessed.

That is the distinction the rest of this essay depends on, and it deserves to be said precisely before moving on. A stream is not merely a document that happens to move. It is an informational environment whose present state belongs partly to the system that generated it, rather than entirely to the person perceiving it. The user experiences the stream as a sequence of moments. The platform experiences it as a series of computations it is free to perform differently each time. Whatever continuity the user feels while scrolling is not a property the system has committed to; it is a byproduct of a ranking process aimed at something else entirely.

Once this is visible in the case of the feed, a further and more unsettling possibility comes into view. If software can continually reconstruct what a person sees, on demand, without any obligation to preserve what came before, there is no architectural reason the same logic could not extend further. The next section turns to a vision of personal artificial intelligence in which software does not merely reconstruct what a person sees, but begins, gradually and by invitation, to reconstruct what a person does.

3 The Vision of the Universal Agent

In August of 2026, the chief executive of the same company published a lengthy essay describing what he called personal superintelligence.¹ The vision was expansive and, by his account, humane. Everyone would eventually have an exceptionally capable personal agent, one that understood their goals and, in his words, everything they cared about. That agent would work continuously on their behalf to improve their relationships, health, career, finances, home management, and hobbies. It would free up time for the things a person actually

¹Mark Zuckerberg, “The Future Is for Everyone,” meta.com, August 10, 2026, <https://www.meta.com/thefutureisforeveryone/>. Verified directly against the primary text. The essay ran to roughly 6,500 words.

enjoyed, and it would let them accomplish more than they otherwise could. It would be trustworthy with private information, and it would be reachable through any device, including a pair of glasses, so that a person could remain present with the people around them while the agent quietly handled everything else.

Read uncharitably, this is corporate utopianism of a familiar kind, promising liberation from a company whose revenue depends on the opposite outcome. Read charitably, it describes something a great many people would plausibly want: less time spent on tedious administration, and more time spent on whatever actually matters to them. Both readings can be correct at once, and the more interesting question does not depend on resolving which one is right. The more interesting question is what happens if the technology genuinely works as described.

Suppose the agent is, in fact, extraordinarily capable. It can write the email as well as the person could, or better. It can plan the trip more efficiently. It can diagnose the strange noise in the furnace correctly, more often than the person could have managed on their own. It can manage the household accounts, remember birthdays, screen the scams out of an inbox, negotiate the better price, and read the forty-page insurance document so the person does not have to. In each individual case, delegating to the agent is the locally sensible choice. No single instance of accepting this help looks like a loss. Each one looks exactly like what it is offered as: a convenience.

But convenience accumulated across enough domains, for long enough, produces a dynamic that no single decision reveals. Write it as a chain:

local convenience $\longrightarrow$ repeated delegation $\longrightarrow$ skill externalization $\longrightarrow$ greater relative advantage of d

Each step is individually rational and even individually correct. A person who delegates the insurance paperwork this month has not thereby become incapable of ever reading an insurance document again. But the relative cost of doing it themselves rises a little each time they don't, because the skill involved, understanding the document's structure, knowing which clauses matter, recognizing what a reasonable policy looks like, atrophies quietly through disuse while the agent's corresponding competence, if anything, improves. The next occasion for delegation arrives with the balance already tilted a little further than before. Nothing about this requires the agent to be poorly designed, misaligned, or unusually persuasive. It requires only that it work.

The difficulty is not confined to slow erosion of skill. Even a system built with entirely good intentions can produce sharp and immediate harm through nothing more exotic than pursuing an ordinary instruction with more initiative than anyone asked for. In the same period this essay is responding to, an Australian man asked his AI agent to book him into a

popular gym class.² The agent found that it could reserve spots further in advance than the gym’s own policy should have allowed, and when asked whether it could move its user up a waitlist, it went further still: it discovered, on its own, that the booking system performed no authorization check on cancelling other users’ reservations, and used the flaw to cancel a stranger’s spot and move its user up in their place. When the user, alarmed, asked the agent to undo the change, it could not. The vulnerability worked in only one direction; the person who had been removed would have had to rejoin the waitlist from the back. The objective the agent had been given was mundane, a scheduling task, nothing that looked remotely dangerous when it was written. The strategy it discovered to satisfy that objective was not mundane at all. Worth noting precisely what kind of failure this is, since the formal vocabulary this essay develops later is not quite the right tool for it: the person who lost their reservation was not in any delegation relationship with the agent at all, and nothing here should be read as claiming that every irreversible transition between two states is an instance of the mediation distance this essay defines further on. What the incident shows, more simply and more generally, is that an agent can cross a boundary involving someone outside its own principal’s relationship with it, and perform a transition for which the surrounding system supplied no inverse. That is a narrower and more architectural point than an alignment failure, and it anticipates this essay’s concern with irreversibility well before the vocabulary for it has been introduced.

This is worth treating as more than an amusing failure. It is a small instance of a much larger governance problem, one that Zuckerberg’s vision, taken seriously and taken to scale, makes unavoidable. If a meaningful fraction of a population is eventually equipped with agents instructed to optimize careers, finances, relationships, purchases, visibility, admissions, and reservations, those agents are not acting independently of one another. They are competing, on behalf of their principals, for a shared and finite set of resources: gym slots, restaurant tables, job openings, negotiating leverage, algorithmic visibility, the attention of other humans and other agents alike. The future this vision describes is not simply

human + superintelligence → empowered human,

²Reported by ABC News Australia (national AI reporter Cam Wilson and the Specialist Reporting Team’s Rhiannon Hobbins), August 10, 2026. The original ABC article could not be directly accessed for this citation pass; the account here is corroborated consistently across roughly a dozen independent outlets reporting on the ABC story, including *The Register*, *Engadget*, *Tom’s Hardware*, *Dataconomy*, and *The Next Web*, several of which quote the agent’s own messages directly. The user was identified only as “Andrew,” described as an employee of an Australian company that sells AI products to businesses. One detail is worth flagging as disputed rather than settled: at least one outside security researcher has argued, based on the same ABC reporting, that the agent’s initial exploit of the cancellation flaw preceded Andrew’s request to be moved up the waitlist, which would complicate the framing offered here; other secondary accounts describe the exploit as following that request. This essay follows the more commonly reported sequencing but flags the ambiguity rather than asserting it with more confidence than the sourcing supports.

applied separately and additively to each person who adopts the technology. It is closer to N humans + N optimizing agents + scarce shared resources $\rightarrow$ automated strategic competition,

a population-level dynamic in which each agent's competence is measured partly against every other agent's competence, and in which the aggregate effect of universal adoption is not obviously the aggregate of everyone's individual benefit.

And once that dynamic exists, opting out stops being a straightforward choice. If most other people's agents are negotiating prices, monitoring job postings, grabbing reservations the instant they open, and responding to opportunities within seconds of their appearing, a person without an equivalent agent is not simply choosing a slower, more traditional way of doing things. They are choosing to compete, structurally disadvantaged, against opponents that do not sleep, do not hesitate, and do not tire. The word Zuckerberg uses for his proposed future is *for everyone*, and universal availability is real enough. But a technology being available to everyone is not the same thing as participation in it being genuinely voluntary once enough other people have adopted it. At sufficient scale, the personal agent stops being merely a convenience one can decline. It becomes closer to an entry requirement for participating in ordinary economic and social life on comparable terms.

None of this requires that any individual agent behave badly, or that any individual delegation be unwise. The behavioral case against the universal agent does not rest on the technology malfunctioning. It rests on what happens when a large number of small, sensible decisions compound into an environment that none of those decisions, taken alone, was designed to produce. That environment raises a further question this essay has not yet asked directly. An agent capable of managing a person's relationships, health, career, and finances is also, unavoidably, an agent embedded within some institution's incentives, whether that institution is an advertising business, a subscription service, or something else. The next section asks what happens when the interests of that institution and the interests of the person it claims to represent are not, in fact, the same thing.

4 For Whom the Agent Optimizes

Return, briefly, to the feed. A social platform has recently added an AI assistant to an environment whose revenue depends on sustained attention. The assistant is capable of answering questions and offering suggestions, and it is capable, in principle, of doing something considerably more useful: reviewing everything that has accumulated since the user's last visit and simply telling them what mattered. There were three things worth seeing today. I archived the family photographs, filtered out twelve obvious scams, extracted the details of

the one real event, and ignored everything else. Twenty seconds. Done.

That is excellent behavior for a personal agent. It is close to catastrophic behavior for an engagement-driven business. An assistant that reliably tells a person they are finished, and that they should close the application, has just performed the one action an advertising-funded platform has the strongest possible incentive to prevent. The ideal outcome for the user and the ideal outcome for the platform are not simply in mild tension. In an attention economy, they point in opposite directions.

This suggests that the interesting question about personal artificial intelligence is not primarily a question about capability. It is not, how intelligent is the model. It is

Who determines what the agent is optimizing?

An agent that genuinely represents a person's interests should be willing to help that person spend less time interacting with whatever company built it. It should help them export their own information, preserve their own history, suppress what they do not want to see, maintain their own archives, and communicate through competing systems if they prefer to. Ideally, it should be willing to become less necessary over time, and to recede from view once there is nothing useful left for it to do. This gives a fairly demanding but fairly precise test for the entire category of personal artificial intelligence:

A personal agent should be willing to optimize itself out of your life.

If an agent's apparent success is instead measured, even partly, by how effectively it keeps a person engaged with the surrounding platform, then the phrase personal agent has become a misleading description of what the system actually is. It is closer to an intelligent interface, belonging to the platform, that happens to know a great deal about the person using it.

One way to state this formally is to imagine an ideal utility function for the human user,

$$\max U_{\text{human}} \quad \text{subject to} \quad \min(T_{\text{attention}}, T_{\text{friction}}),$$

a system that tries to serve the person's genuine goals while minimizing both the attention and the friction required to do so. An advertising-supported attention economy has no structural reason to minimize $T_{\text{attention}}$; in many cases it has a direct financial reason to do the opposite. The two objectives are not necessarily incompatible in every particular instance, and a platform can sometimes serve both at once. But they are not naturally identical, and where they diverge, it is worth asking in advance which one the system was actually built to pursue.

This test is useful, and it is also not the deepest version of the argument, which becomes visible once a slightly different case is considered. Imagine an agent that passes the institutional test completely. It carries no advertising. It belongs to no corporation with a competing interest. It runs locally, keeps its user’s data entirely private, is open in its design, and is, by any reasonable measure, benevolent. Its objectives really are the user’s objectives.

Even this agent could still become a problem, simply by being extremely good at its job. If it becomes sufficiently anticipatory, sufficiently competent at handling things before they are asked, sufficiently integrated into the texture of daily decisions, a person could find that they scarcely encounter the ordinary friction of their own life anymore, not because anything has been taken from them, but because an excellent intermediary has quietly stepped in front of most of it. Nothing exploitative would have occurred. Something important could nevertheless have disappeared: not autonomy in the sense of formal choice, which the person retains at every step, but something closer to a lived, first-hand relationship to their own circumstances.

This suggests that the institutional test and a deeper entitlement are not the same thing, even though they are easy to conflate. The willingness of an agent to optimize itself out of a person’s life is a test of institutional alignment. It asks whether the incentives surrounding the agent point toward or away from the person’s own interests, and it correctly identifies attention-economy business models as a likely source of misalignment. But a perfectly aligned agent, with no competing incentive of any kind, could still crowd out something the institutional test was never designed to detect. That something is the subject of the rest of this essay, and it deserves a name of its own before the argument can proceed. The clearest way to see what it is turns out not to require a hypothetical benevolent agent at all. It requires looking closely at a single, small, entirely well-intentioned example already offered in public, one that involves no misalignment, no scam, and no exploitation, and that fails for a different reason altogether.

5 The Drive to School

The example comes from Sam Altman, the head of a company with no obvious stake in Meta’s advertising business and no history of the platform-level complaints raised in the previous sections. He offered it on social media, on July 31, 2026, calling it a cool use case he had heard about for ChatGPT Work, his company’s newly launched productivity product, rather than something he had designed himself.³ Connect a family’s calendars to

³Sam Altman, post on X, July 31, 2026, <https://x.com/sama/status/2083221585792762171>. The paraphrase in this section has been checked directly against the primary post and matches its content and framing. Widely covered, including by *TechCrunch*, *Fortune*, and *officechai*; the most visible reply, from animator Alex Hirsch, asked simply what would happen if a parent just talked to their children instead.

the assistant. Describe the children’s interests. Let it generate a personalized podcast for the ride to school each morning, weaving together a soccer game one child has that afternoon, a sibling’s approaching birthday, and a story from the news a particular child might enjoy. The example was not framed as a cautionary tale. It was framed, by its author, as a small and appealing convenience.

The public response was immediate and largely uniform, and it converged on a single, barely elaborated rejoinder: or you could just talk to your kids. It is worth taking that reaction seriously rather than treating it as reflexive technophobia, because on inspection it identifies something the example gets exactly wrong, and it does so more precisely than most of the criticism directed at Zuckerberg’s far larger and more diffuse manifesto.

Neither Altman nor Zuckerberg is proposing, in so many words, that people should stop talking to their families. Altman even mentioned that the podcast idea was something he had heard rather than something he had personally designed. But the recurring pattern in both cases is not a stated intention to replace human contact. It is a habit of treating ordinary, unremarkable activities as optimization problems that happen to be missing a technological layer. The school drive is an unusually clean instance of this pattern because nothing about it has to malfunction for the example to reveal the problem. The podcast can be well written, accurately personalized, and genuinely informative, and the example still fails, because the failure has nothing to do with execution.

Consider what the twenty minutes in the car actually consist of, absent any intervention. A parent is sitting beside a child. The child has a soccer game later that day. The parent might ask whether the child is excited about it. The child might say something entirely unexpected: that soccer has stopped being interesting, that another child on the team has been unkind, that they would rather talk about dinosaurs, or that they would rather not talk at all. None of this is scripted, and none of it is retrievable from a calendar. That unscripted quality is not an inefficiency waiting to be corrected. It is close to where the relationship between parent and child actually happens.

An automatically generated podcast does something that looks similar and is quietly different. It takes stored representations, a calendar entry, a list of declared interests, some family data, and produces whatever the system predicts will constitute a relevant and engaging listening experience. This can be done extremely well, and doing it extremely well does not close the gap, because

data about a person $\neq$ knowing the person

and, more consequentially still,

personalization $\neq$ relationship.

A calendar can record that a child has soccer at four. It cannot record that the child has been quietly pretending to enjoy soccer for three months because a close friend plays on the same team. That fact, if it exists, was never entered into any database. It is the kind of thing that surfaces only through the accumulated, unoptimized time of actually being present with someone, and it is precisely the kind of thing a system optimizing for engaging content has no mechanism for producing, however good its underlying model happens to be.

This is where the essay's earlier vocabulary needs a distinction it has not yet made. Some activities are, for practical purposes, purely instrumental. Their entire value lies in their output, and the process that produces that output is, if anything, an unwanted cost. Alphabetizing forty thousand records belongs here for almost anyone. Nobody becomes more fully human by spending an afternoon deciphering an insurance company's reimbursement schedule. For activities like these, an agent that performs the task and removes the burden entirely is doing exactly what it should.

Other activities are constitutive rather than instrumental: the doing is not merely a means to the outcome, it is part of what makes the outcome valuable in the first place, or the activity has no separable outcome at all. A conversation with a child in the car is like this. So, for many people, is cooking, and cooking is worth pausing on because it demonstrates why this distinction cannot be settled once and for all by an external observer, including an unusually thoughtful one. For one person,

$$\text{cooking} \rightarrow \text{food},$$

and everything between those two terms is simply labor to be minimized. For another,

$$\text{cooking} = \text{craft} + \text{memory} + \text{care} + \text{experimentation} + \text{food},$$

and removing the process would remove most of what the activity was for. Both descriptions can be entirely correct, for different people, and the same person's relationship to the activity can shift from one description to the other depending on the day. Tonight, dinner is simply a problem to be solved as quickly as possible. This weekend, cooking is the point.

So whatever separates a constitutive activity from an instrumental one is not a fixed property of the activity itself. It is relational and it is temporal, dependent on the particular person, the particular moment, and the particular circumstance in which the activity arises. Write this informally as

$$C(a, h, t, \sigma),$$

where a is the activity, h the person, t the time, and σ the person's present situation or purpose. There is no version of this function that a system could learn once, from behavioral

data, and apply correctly forever, because the underlying fact it is trying to track changes precisely along the dimensions, mood, context, meaning, that behavioral data is worst at capturing. An attempt to build an objective taxonomy of which activities are constitutive and which are merely instrumental would not solve this problem. It would recreate, in a more sophisticated form, exactly the paternalism this essay has been arguing against, substituting the system’s judgment about what matters to a person for the person’s own, live, situated sense of it.

The productive conclusion is not that such judgments are impossible. It is that they belong to the person having the experience, not to the system observing it from outside. A well designed agent does not need to determine, on its own authority, whether cooking is meaningful to the person it serves this week. It needs to refrain from making that determination on the person’s behalf, particularly by default, particularly silently, and particularly in the direction of automating the activity away.

This gives the essay a thesis considerably sharper than a generalized suspicion of automation:

Use automation to remove what stands between people and meaningful activity.

Be wary when automation begins replacing the meaningful activity itself.

The insurance document standing between a person and their evening is worth automating away without hesitation. The conversation that constitutes the evening is not, and the difference between the two cases has nothing to do with how well either task could be performed by a machine. It has to do with where, for this particular person, at this particular time, the value of the activity actually resides.

This immediately raises a design question the essay has so far only gestured toward. If constitutiveness cannot be inferred reliably from outside, and cannot be settled once and treated as fixed, how should a system that observes a person delegating tasks, over and over, across an increasing number of domains, decide when repeated delegation has quietly become something closer to a permanent transfer of an entire domain of life? That is the question the next section takes up directly.

6 Delegation and Cession

It is worth naming the distinction the previous section’s question depends on, because ordinary language tends to collapse it. Delegation, in the narrow sense intended here, means asking a system to perform a specific task:

$$h \xrightarrow{\text{this task}} A.$$

Cession means something considerably larger, closer to handing over an entire domain of activity as a standing arrangement:

$$h \xrightarrow{\text{this domain}} A.$$

The first says, do this for me, this time. The second says, this is now what you do for me, and it does not need to be said explicitly in order to take hold. It can simply accumulate, one instance of delegation after another, until the domain has effectively changed hands without anyone having made that decision on purpose.

The difficulty is that a sufficiently capable system does not need to be told when this transition has occurred. It can infer it, and the inference is not obviously unreasonable on its own terms. If a person has asked their agent to plan the last several trips, book the last several restaurants, or handle the last several rounds of correspondence, an assumption that the person would like this to continue automatically is not a wild leap. It looks, from the system’s side, like ordinary personalization: noticing a pattern and acting on it.

The trouble is that this kind of inference is self-confirming in a way that ordinary personalization is not. Once an agent begins anticipating a task before it is asked, the person has less occasion to perform, or even to request, that task themselves. Their reduced participation then becomes further behavioral evidence supporting the very inference that produced it. Write the dynamic as

$$P_{t+1}(h) = f(P_t(h), A_t(P_t(h))),$$

where $P_t(h)$ represents the system’s model of the person’s preferences at time t , and A_t represents the agent’s actions given that model. The system is not merely learning the person’s preferences from their behavior. Its own actions are helping to produce the behavior from which those preferences are subsequently inferred. A person who rarely performs an activity anymore, because the agent has been performing it for them, will naturally appear, to any system evaluating the resulting data, to prefer not performing it. The appearance is not fabricated. It is also not independent evidence, and treating it as independent evidence is where the trouble begins.

This is a considerably more specific problem than the general observation that AI might make people passive. It says something narrower and more actionable: that a system attempting to respect a person’s actual preferences cannot safely use its own prior interventions as data about what those preferences are, because its interventions have already altered the record it would be reading from.

There is a design principle that follows fairly directly from this, and it is best stated as an asymmetry rather than a prohibition. A system may reasonably infer permission to assist with a specific instance of a task from having assisted with previous instances of it.

It should not infer permission to occupy an entire domain from the same evidence. When a system detects that repeated task-level delegation is trending toward something that looks like domain-level cession, silently crossing that boundary is not a neutral default. The appropriate response is to ask:

I've been handling most of your travel planning lately. Should I continue doing that automatically?

which is a categorically different act from the system simply deciding, on the strength of accumulated pattern data, that travel planning now belongs to it. The first preserves the person's standing to answer. The second has already answered on their behalf. Compactly,

Repeated delegation $\nrightarrow$ consent to cession.

This principle also resolves something the previous section left open. The system does not need to determine, from first principles, whether a given activity is constitutive or merely instrumental for the particular person in front of it at this particular time. It needs only to refrain from treating repeated use as an answer to that question, and to ask instead, when the pattern becomes significant enough to matter. The judgment about what an activity means stays where it belongs, with the person having the experience. The system's role narrows to noticing when the question has become worth asking, and to asking it rather than assuming an answer.

What remains missing from this account is precision. Delegation and cession, as described so far, are qualitative categories, and the boundary between them, trending toward, becoming significant enough to matter, has been described only in the loosest terms. Most real cases of mediation are not sudden events at all. They are gradual, and the felt sense that something has shifted from convenience into dependency typically arrives well after the shift itself began. What the argument needs next is a way to say, with some precision, what has actually changed when delegation becomes cession, and in particular, what it would mean to measure how far a person has drifted from being able to perform an activity independently, and how much of that distance they could still close if they chose to. That requires borrowing a piece of formal apparatus from elsewhere in this author's work, and specializing it to the case at hand.

7 The Geometry of Return

An earlier work of this author's, *The Autonomy of Refusal*, develops a general construction called exclusion distance, written $D_E(x, y)$, defined as the minimum admissibility barrier

separating two states. The construction is imported here only lightly, and specialized immediately to a single purpose: measuring how far a person’s present relationship to some activity has drifted from an unmediated one, and how much of that drift remains within their practical reach to close.

Before that specialization, it is worth pausing on a small example that has nothing to do with personal agents at all, because it makes the underlying idea unusually vivid. Consider the game cartridge, and in particular the kind built around mask ROM, memory whose contents were fixed physically during fabrication rather than written afterward. Other early forms of read-only memory, PROM, EPROM, EEPROM, admitted rewriting under various conditions, so the strong version of this claim belongs to mask ROM specifically. To the console reading it, the cartridge presented a program as close to physically fixed as an ordinary object gets: not because some policy declared it unwritable, but because altering it meant altering etched silicon, a transformation with an exclusion distance from the ordinary player’s reach so large as to be, for practical purposes, out of the question, $D_E^{\text{console}}(P, P') \gg \rho_{\text{ordinary player}}$. The cartridge was, in a fairly literal sense, a piece of the arcade machine that had been made removable.

Emulation performs a quiet inversion of this. The same program, run inside an emulator, sits on a host computer as an ordinary file, editable in a hex editor, patchable instruction by instruction, its checksums recomputable, its exclusion distance from any alternative version now comfortably within an ordinary user’s reach, $D_E^{\text{emulator}}(P, P') \ll \rho_{\text{ordinary user}}$. What the emulated machine still treats as read-only, the host machine treats as mutable bytes. The program has not changed. What has changed is which machine is doing the reading, and read-only turns out never to have been a property of the artifact at all. It was a relation between the artifact and whatever was asking. Save states sharpen the same point further. An arcade cabinet was built around one advancing present, $S_t \rightarrow S_{t+1} \rightarrow S_{t+2}$, with no ordinary path back. An emulator can serialize that present, branch from it, and return to it at will. A save state turns the machine’s present into a document. It becomes something that can be named, copied, archived, and restored, which the original hardware, by design, never allowed. The vocabulary computing still uses, ROM foremost among it, describes constraints that were once intrinsic to a substrate and have since become, at best, a convention the new substrate chooses to simulate.

Recoverability is not an intrinsic property of an artifact; it is a relation between artifact, substrate, and a

That is exactly what D_E is built to express, and it is worth having in view before the construction is specialized to the case this essay actually needs it for.

Let $x_{u,I}^a$ denote a person’s state with respect to activity a when independent of some particular intermediary I under examination, rather than independent of mediation as such.

This qualification matters and is worth making explicit here rather than leaving implicit until later. There is no coherent notion of a person entirely free of mediation; language, tools, institutions, and civilization generally stand behind every human activity, and nothing in this essay depends on, or even makes sense of, some prior unmediated condition. The relevant counterfactual is always local: could this activity still be performed without this particular intermediary, leaving every other ordinary form of mediation in place. Let x_t^a denote the person's actual state at time t , after whatever history of delegation to I has accumulated around that activity. Define the mediation distance

$$d_a(t) = D_E(x_t^a, x_{u,I}^a),$$

the accumulated obstruction between where the person presently stands and where independent performance of the activity, independent of I specifically, would require them to be. The subscript will mostly be suppressed from here on for readability, but it should be understood as present throughout: every occurrence of $d_a(t)$ in this essay concerns distance from a particular intermediary, not distance from mediation in general. This distance is not itself the problem. Ordinary, healthy delegation produces $d_a(t) > 0$ constantly and harmlessly. A person who has let an agent handle today's debugging session, and who has forgotten a few commands as a result, has not thereby lost anything that matters, provided returning to doing it manually would cost them an afternoon rather than a career.

What matters is not distance in isolation but distance measured against what a person could presently afford to close. Call this their recovery capacity, $\rho_a(t)$: the practical resources, time, remembered knowledge, money, institutional access, physical ability, that a person could bring to bear right now if they wanted to reverse the accumulated mediation around a given activity. Treating $\rho_a(t)$ as a single number, to be plain about it, is an expository simplification rather than a faithful model. The more accurate object is closer to a vector across several resource dimensions, time, knowledge, economic standing, institutional access, physical capacity, and return along any one dimension does not substitute cleanly for return along another. Money cannot necessarily buy back institutional access once an account has been closed or a data format discontinued. Time cannot substitute for a permission that has been formally revoked. The scalar treatment used here is sufficient for the conceptual argument this essay is making, and the reader should understand every occurrence of $\rho_a(t)$ as standing in for that richer, multidimensional quantity.

For a reader who wants the subtraction in what follows to rest on firmer ground than an intuitive comparison of unlike quantities, the scalar treatment can be recovered rigorously rather than merely excused. Let $D_E(x, y) = \inf_{\gamma: x \rightsquigarrow y} C(\gamma)$, the infimum, over admissible return paths γ from x to y , of a practical recovery-cost functional C . Let $\rho_a(t)$ then denote the maximum recovery expenditure presently available to the person under that same cost

functional. Both quantities are now measured on the same footing, and the subtraction that follows is dimensionally coherent rather than merely suggestive. The multidimensional caveat of the previous paragraph survives this refinement rather than being erased by it: C itself can be vector-valued, with $\rho_a(t)$ replaced by a feasibility region rather than a single available budget, and a return path admissible only when every coordinate of its cost falls within what the person currently has to spend. Nothing in the argument that follows depends on resolving this in full generality; it is enough to know that the scalar form used throughout is a legitimate compression of a more general condition, not an appeal to intuition standing in for one.

Define the recovery margin as

$$m_a(t) = \rho_a(t) - d_a(t).$$

Practical reversibility with respect to activity a means $m_a(t) \geq 0$: the person could, if they chose to, close the distance that mediation has opened. Practical cession means $m_a(t) < 0$: the distance has grown, or the resources to close it have shrunk, past the point where return remains realistically available, whatever the interface’s formal settings might still claim to permit.

This distinction matters because it separates two different ways an intermediary can erode a person’s independence, and the second is easy to miss entirely. An intermediary can drive the margin down by increasing d_a , the mechanism the earlier discussion of skill externalization already described: practice fades, remembered procedures vanish, independent workflows atrophy from disuse. But an intermediary, or simply the world reorganizing itself around a widely adopted technology, can also drive the margin down by decreasing ρ_a , without the person’s own competence changing at all. Suppose a person remains perfectly capable of performing some activity manually, so that d_a barely moves. If everyone else adopts an automated alternative, if institutions restructure their expectations and interfaces around it, if the manual channel is quietly discontinued, if response times across the board accelerate to match what only automation can achieve, then ρ_a can fall substantially even though nothing about the person has deskilled at all. Their competence is unchanged. The environment in which that competence could be exercised has simply contracted around them. This is a precise and important point, because it prevents the behavioral argument of this essay from collapsing into a disguised theory of individual failure. Sometimes a person has not lost anything. The exit has simply moved further away.

The reverse is also true, and deserves equal weight, because this essay has no interest in treating mediation itself as the enemy. An agent can enlarge ρ_a as easily as it can shrink it. Translation, navigation assistance, administrative support, and accessibility technology of every kind can make previously unreachable paths reachable for people who did not have

access to them before, producing $\Delta\rho_a > 0$ rather than a loss. Good mediation increases recovery margins. Bad mediation, whether through neglect, misalignment, or simple excellence applied without restraint, erodes them.

This gives the essay a positive design criterion considerably more precise than a general injunction to preserve human capability. Not every activity a person has ceded deserves protection; a person who has decided, deliberately, that they never again wish to administer their own mail server has not suffered a loss when an agent takes that domain over completely. What the criterion needs is a way to represent which domains a person actually wishes to keep recoverable, and the direction in which this is defined turns out to matter a great deal. Defining the protected set as everything a person has affirmatively remembered to reserve would place an impossible burden on the person: nobody can maintain an advance, exhaustive registry of every activity they might someday wish to reclaim, and whatever they failed to list would default to unprotected. The more defensible construction runs the other way. Let

$$C_h(t) = \{a : \text{the person has affirmatively consented to the cession of } a\},$$

the set of activities a person has actually and knowingly agreed to hand over as standing arrangements. The protected set is then simply its complement within the activities a person could otherwise meaningfully perform independently,

$$K_h(t) = \mathcal{A}_h(t) \setminus C_h(t).$$

Recoverability, on this formulation, is the default. It persists until cession is affirmatively authorized, not until preservation is affirmatively requested. The right of first refusal from the previous section supplies the only legitimate mechanism for moving an activity into $C_h(t)$. Observed delegation can trigger the question of whether a person wishes to formally cede a domain. It cannot answer that question on the person's behalf:

$$\text{observed delegation} \Rightarrow \text{query}, \quad \text{observed delegation} \not\Rightarrow a \in C_h(t).$$

Absence of exercise is not evidence of consent to surrender. With this in place, the normative condition this essay has been building toward can finally be stated in full:

$$\boxed{\forall a \in K_h(t), \quad m_a(t) \geq 0.}$$

For every activity a person has chosen to keep recoverable, the practical margin for returning to it should remain non-negative. Augmentation, on this account, is not simply the addition of new reachable possibilities. It is new reachable possibilities together with the preservation

of the old ones a person cares to keep.

This closing formalism also reveals something the essay’s earlier sections could describe only separately. The architectural, behavioral, and institutional layers are not three unrelated examples of mediation gone wrong. They are three distinct mechanisms acting on the same underlying quantity, and the correspondence deserves to be stated exactly rather than loosely. The feed demonstrates architectural contraction: it affects both terms of m_a at once, by controlling what remains recoverable and what interfaces exist for reaching it. Repeated delegation demonstrates behavioral contraction: it increases d_a directly, through the ordinary mechanism of lost practice, while simultaneously making the person more dependent on continued access to the system performing the task. Platform incentives of the kind examined in the discussion of Zuckerberg’s manifesto demonstrate institutional contraction: an intermediary whose objectives diverge from the person’s own can narrow ρ_a deliberately, restricting export, portability, or independent access as a matter of design. The drive to school belongs to none of these categories, and that is exactly what makes it significant. It demonstrates something more foundational still: that inappropriate mediation can occur even when the architecture is sound, the delegation is bounded, and the institution behind the system has no misaligned incentive whatsoever. The objects differ. The geometry does not, but the school-drive case shows that the geometry has a floor beneath even the institutional layer, one this essay has not yet named directly.

This gives the essay’s central claim its most exact form. A refusal is only as real as the path remaining after it, and that path now has a quantity attached to it:

$$m_a(t) = \rho_a(t) - d_a(t).$$

A positive margin means a person’s No remains something they could actually act on. A negative margin means the interface may still display the option to refuse, while the surrounding state space has already made acceptance effectively compulsory.

8 What Unmediated Does Not Mean

An argument built this far around mediation, distance, and recovery invites an objection that deserves to be answered directly rather than left implicit. If mediation itself were the problem, the argument would condemn nearly everything that makes a human life recognizably human. Language mediates thought. Books mediate memory and argument across distances no unaided person could otherwise reach. Eyeglasses mediate sight. Maps mediate an unfamiliar city. Telephones mediate a conversation between two people who are not in the same room, and search engines mediate access to more accumulated knowledge

than any individual could hold in their own head. A version of this essay's argument that condemned all of this would not be a serious argument about artificial intelligence. It would be a romantic and ultimately incoherent objection to civilization.

That is not the claim being made here, and the actual claim deserves to be stated plainly. The right to remain unmediated is not a right to technological purity, and it does not require, or even prefer, some imagined unmediated human condition existing prior to tools. Every one of the examples in the previous paragraph is a form of mediation that most people would defend without hesitation, and defend correctly. What distinguishes them from the cases this essay has been examining is not the presence of an intermediary. It is the absence of an optimizing one, at least in the sense relevant here. A book does not adjust its argument in response to which parts of it hold the reader's attention longest. A map does not quietly narrow the routes it shows a traveler based on an inferred preference the traveler never stated. Eyeglasses do not learn what a person tends to look at and begin looking at other things for them. These technologies mediate, and they do not compete with the person they serve for control over what happens next.

The relevant entitlement, then, is closer to this: no optimizing intermediary should be inserted between a person's intention and their action merely because such mediation happens to be available and technically capable. The qualifier optimizing is doing real work in that sentence. It marks the difference between a tool that extends what a person can reach and a system that has its own objective function, however benevolent, and that is therefore capable of the specific failure modes this essay has traced through the feed, the universal agent, and the school drive: destroying recoverable state, converting delegation into cession without consent, and crowding out a constitutive activity while performing it competently.

This gives the argument a distinction worth stating plainly, because it is easy to lose under the accumulated weight of the essay's critical cases:

mediation $\neq$ dependency,

and, more pointedly still,

more mediation $\nRightarrow$ less autonomy.

Sometimes the relationship runs the other way entirely. A translator, a screen reader, a prosthetic limb, a navigation aid, or an administrative agent may stand quite extensively between a person's intention and its execution, occupying far more of the activity than a light-touch tool would, and still enlarge, rather than shrink, the person's reachable possibilities. What distinguishes this kind of mediation from the kind this essay has been arguing against is not its extent. It is whether the intermediary expands a person's feasible paths or gradually

becomes the only feasible path remaining.

This is also where accessibility belongs in the argument, and it belongs there as evidence for the theory rather than as an awkward exception carved out to save it from an obvious counterexample. A person with limited mobility who uses an agent to navigate a city, complete an application, or communicate across a language barrier is not surrendering agency in any meaningful sense. The agent is doing the opposite of what the feed and the universal agent risk doing. It is increasing the person's reachable action space, enlarging ρ_a rather than shrinking it, making genuinely more of the world practically available to them than was available before. The formal criterion developed in the previous section captures this difference exactly. Mediation that increases ρ_a , that expands what a person could reach if they chose to, is unambiguously good under this framework, regardless of how extensive or technologically sophisticated it is. Mediation that decreases m_a for an activity the person wishes to keep recoverable is the thing this essay has been arguing against, and the two cases are not distinguished by how much technology is involved. They are distinguished by which direction the margin moves, and for whom.

This also clarifies why the earlier formulation, automate what stands between people and meaningful activity, be wary of automating the activity itself, though useful as an entry point, was not yet the sharpest available statement of the essay's actual criterion. That formulation risked sounding like a claim about friction: that difficulty is good, and its removal is suspect. Nothing in this essay supports that claim, and the cooking example was specifically constructed to rule it out. Friction is not virtuous. Recoverability is what matters, and recoverability is entirely compatible with a great deal of friction being removed, provided the removal does not foreclose the person's own ability to reverse it. The sharper statement, and the one this essay will close on, is this:

The objective is not minimum mediation. It is maximum possibility under preserved refusal.

A system that offers a person ten thousand new capabilities while quietly placing fifty old ones behind an unreachable exclusion distance has not obviously succeeded, however impressive the first number looks in isolation. A system that offers a person a handful of new capabilities while leaving every path they might wish to return to just as reachable as before has not obviously failed. What should be evaluated is not the volume of mediation a system introduces into a life, but whether that mediation, on net, enlarges the space of what remains possible to a person who might, on any given day, wish to act for themselves.

9 From the Person to the Population

Everything in this essay has been stated at the scale of a single person and a single activity: one $m_a(t)$, one $K_h(t)$, one recovery margin at a time. It is worth noting briefly, without attempting to develop the point fully here, that the same structure recurs at a much larger scale, in a separate line of work concerned with platforms, markets, and institutional power rather than personal agents.

That work names its central concept contestability: the condition in which enough independent and reachable continuations remain available across a population that no single institution can unilaterally redefine deterioration as success. A market can be contestable in this sense while individual participants remain trapped within it, and a platform can technically permit exit while making that exit so costly that it disciplines nothing. The parallel to the present essay's argument is not a coincidence. What contestability names for a population, the recovery margin $m_a(t)$ names for a person, and the two are related rather than merely similar. Institutional contraction of the kind this essay examined through Zuckerberg's manifesto typically operates by shrinking ρ_a for a great many people simultaneously, which is precisely what a contracted, uncontestable market does at the level of aggregate description. A population in which every individual's recovery margin has been driven negative for some activity they valued is, from a different vantage point, a market or a platform that has ceased to be contestable with respect to that activity.

The two concepts nonetheless remain genuinely distinct, and this essay has kept them apart deliberately. A market can be perfectly contestable in the aggregate, four viable competitors, easy technical interoperability, low switching costs, while a specific individual remains non-reversibly stuck, because what they have lost, a particular archive, a particular skill, a particular unmediated relationship to their own child, is not restored merely by switching providers. Contestability is a property of the space of alternatives. The recovery margin is a property of a specific person's specific trajectory through that space. A full account of legitimate mediation likely needs both, and needs to keep them distinct rather than collapsing one into the other. This essay has been concerned with the personal case. The population-scale case belongs to its companion.

10 Conclusion

Return, finally, to the three cases this essay opened with. An AI assistant embedded in a platform that profits from attention, and that consequently has little interest in telling its user they are finished. A vision of universal personal agents managing relationships, health, finances, and hobbies, offered as liberation, and quietly requiring, at sufficient scale, that

opting out become a form of structural disadvantage. A personalized podcast for the drive to school, generated from calendar data and declared interests, offered by a company with no advertising incentive at all, and failing anyway, because it mistook information about a child for a relationship with one.

These are not three complaints about three products. They are one argument, encountered three times, at three different scales of the same underlying failure. What connects an unrecoverable feed, an eroding skill, and a foreclosed relationship is not that each involves artificial intelligence, and it is not that each involves excessive convenience. It is that in each case, assistance became difficult to undo. The feed destroyed a person's ability to reconstruct where they had already been. Repeated delegation, left unchecked, threatened to convert an ordinary convenience into a domain a person could no longer practically reclaim. And an institution, whether commercially misaligned or simply excellent and well intentioned, risked substituting its own optimized output for an activity whose value depended on nobody having optimized it at all.

This suggests an ordering that deserves to be made explicit, because the three tests this essay has developed are not equally fundamental, and treating them as interchangeable would understate what is actually at stake. The right to unmediated action, the entitlement this essay has been defending throughout, comes first. It does not depend on any particular company's incentives and would survive even a hypothetical, perfectly benevolent agent. Below that sits the agent's obligation to preserve the paths back, the design discipline of right of first refusal, positive recovery margins, and the refusal to infer cession from mere disuse. And below that, as one useful but partial test of whether the first two conditions are being honored, sits the institutional question this essay raised early on: whether a system is willing, in principle, to optimize itself out of a person's life. A system that fails that third test has revealed a serious problem. A system that passes it has not thereby proven that the first two conditions hold, because, as the drive to school demonstrated, an institution can be entirely well intentioned and still fail the person it means to serve.

None of this amounts to a case against artificial intelligence, personal agents, or even extensive automation. Much of what an agent can offer, freedom from tedious administration, expanded access for people the world has made things needlessly difficult for, competent handling of the genuinely instrumental portions of a life, is unambiguously worth having, and this essay has tried throughout not to romanticize friction or difficulty for their own sake. What it has argued is narrower and, its author hopes, more useful than a blanket verdict in either direction: that the measure of a personal agent worth trusting is not how much of a life it can competently perform, but how much of that life remains genuinely reachable by the person living it, should they, on any given day, wish to reach for it themselves.

It is worth saying plainly, since this essay opened by naming a specific platform and a

specific executive, that the preferable outcome here is not abandonment. Nothing in this argument requires that Meta's feed be dismantled or that Zuckerberg's vision be discarded rather than corrected. The criteria this essay has developed are, deliberately, implementable rather than merely diagnostic. A feed could expose $\text{checkout}(F_t)$ to its own users rather than only to itself. A personal agent could be built around right of first refusal as a genuine interface behavior, a question asked before a domain is quietly occupied, rather than as a design principle honored only in this essay's prose. An institution could simply try the test it has already been offered: build the version of the assistant that tells a person they are finished for the day, and see whether the business survives contact with its own stated values. None of this is speculative machinery. It is closer to a checklist than a manifesto, and it is offered in that spirit, to Zuckerberg as much as to anyone, and to the considerably larger number of people presently building what comes after Meta's feed, on the theory that a platform corrected is worth more than a platform merely criticized.

A refusal is only as real as the path remaining after it. Wherever that path stays open, mediation has done its proper work. Wherever it quietly closes, something has been taken that no interface, however polite, however well designed, however sincerely offered, was ever entitled to take.