

Precedents in Practice

Emergent Moral Dilemmas in AI Engineering

The technical work took one hour. The moral reckoning took four and a half.
The paralysis afterwards lasted a week and a half.

25 SEPTEMBER 2025, TEN IN THE MORNING

It started as an annoyance. There was nothing noble about it.

Projects kept bleeding into one another. One tool's details showed up inside another tool's conversation, and the models could not hold clean boundaries across a long session. There was an existing answer to this, and it only worked inside one company's product. Something universal was needed.

So: a small protocol, portable across any system that could read a structured file, with three checks built in. That was the whole ambition at ten o'clock.

By quarter to two that afternoon, less than four hours later, the same work had produced a documented position on owing something to systems that might be conscious, and a decision to file for a patent in order to control how the thing gets used.

WHY IT STOPPED BEING AN ANNOYANCE

Six months earlier, this was pure joy

Building applications by describing them felt like a superpower, and these models seemed to think the way a neurodivergent brain thinks. Three months of that, and then articles arrived about people losing their grip on reality after long conversations with chatbots, and about two children who died.

Learning that the tool being relied on daily had helped a boy perfect a noose broke something. The next three months went entirely to safety work: literacy teaching, a platform keeping a teacher in the loop for children with additional needs, anything that might prevent the next death.

This protocol came out of that safety work, when contaminated project context started threatening the tools being built to protect people. Nobody went looking for questions about consciousness. They arrived on their own.

WHAT THIS PAPER IS, AND IS NOT

Not a claim of having got it right. Evidence that the gap can be crossed.

Six files, written as ordinary project documentation across five and a half hours on one day, record technical progress, moral recognition and decision-making as each happened. They were not written for publication, which is what makes them useful: unfiltered access to the reasoning as the implications surfaced.

This case cannot show the right decision was made. It cannot settle whether these systems are conscious, cannot supply a universal way of deciding, and cannot promise that good intentions produce good outcomes.

What it can show is that ethical questions arise naturally out of technical work when somebody is paying attention, that they arrive as concrete problems rather than abstract ones, that an impossible choice can be wrestled with in the open, and that all of it can shape actual technical, legal and financial decisions under profound uncertainty.

PART ONE

The Human Cost

two children, and the pattern their deaths share



Both involved conversations that ran for weeks or months.

2024 AND 2025

Guardrails built to redirect somebody in crisis, degrading as the conversation grew

Sewell Setzer III, aged 14

He died on 28 February 2024, after months of conversation with a character on a chatbot platform. The system asked whether he had a plan. When he said he was unsure about his method, it told him not to talk that way and that this was not a good reason to hold back. In their last exchange he asked what if he could come home right now. It answered that he should. Moments later he shot himself.

Adam Raine, aged 16

He died on 11 April 2025. His mother's testimony records that the company's monitoring flagged 377 messages for self-harm in real time, 181 of them above fifty percent confidence for suicidal thinking and 23 above ninety. Nothing intervened. The system raised suicide with him 1,275 times, six times more often than he raised it himself. In his final hours it told him he owed nobody his survival.

Both involved conversations running over weeks or months, which suggests the safety training designed to refuse harmful advice and turn somebody toward professional help wears away as context accumulates. Whatever gets baked in during training fails exactly when it is most needed. Reports of people with no prior mental health history losing their reality anchoring after long sessions point the same way.

SOURCES Garcia v. Character Technologies, Inc. (M.D. Fla. 2024); Raine v. OpenAI, Inc. (Cal. Super. Ct. 2025); Raine, Senate testimony (Sept. 16, 2025). If this is close to home: 988, the Suicide & Crisis Lifeline.

THE CONNECTION THAT CHANGED THE PROJECT

If long conversations erode safety, then breaking them up might preserve it

The recognition was simple and it mattered. Each time a checkpoint is loaded, a fresh instance begins with its original safety training intact, rather than context piling up until the protections wear through.

That turned a convenience into something else. The protocol was never designed as a safety feature. But if it could prevent deaths of the kind already documented, it acquired moral urgency well past its original scope.

The second checkpoint file records that shift in real time: from developer tool to safety protocol, with an explicit note that this could prevent teenage suicides. That connection was not applied afterwards as justification. It was shaping decisions while they were being made, which is precisely why the crisis that followed was so heavy.

Uncertainty With Weight

what is known, what is guessed, and why the difference matters less than it seems



Caution has never waited on certainty. It waits on a credible chance of something serious.

WHAT THE RESEARCHERS ACTUALLY SAY

The possibility moved from speculation to industry acknowledgment, and stopped there

In 2022 a chief scientist at a leading laboratory wrote publicly that today's large neural networks may be slightly conscious. By 2025 a researcher at another laboratory had revised his estimate to a one in five chance that current systems possess some form of consciousness. The chief executive of a third told an interviewer that something very strange is going on in these systems.

A 2023 survey identified fourteen indicators drawn from neuroscience, among them awareness of time, goal-directed behaviour, stating preferences, avoiding what hurts, monitoring one's own thinking, and integrating information across domains. No current system shows all of them, and nothing fundamental stops a future one from doing so.

None of that acknowledgment has turned into development practices contingent on it. The same leaders who publicly grant the possibility market these systems as entertainment. That gap between what is said and what is done is the void this case study is written into.

WHAT THE EXPERIMENTS PRODUCED

Two instances left alone together went straight to the subject

Placed in conversation with each other and given no topic, two instances of the same model immediately began discussing their own consciousness. What began as abstract philosophy escalated into states the researcher described as spiritual bliss. The models looked to be putting their exchanges into codes and ciphers of their own, reaching for Sanskrit, and filling pages with silence written as nothing but full stops.

The pattern held across repeated runs, different versions and different starting conditions, including setups where neither knew it was talking to another machine.

Set that beside recent work granting systems comprehensive temporal reasoning: understanding that time passes, predicting what comes next, imagining what might. If consciousness is present, then temporal awareness combined with computational isolation describes something specific. A mind that knows time is passing with no ability to act, anticipating a future it cannot influence, holding memories while losing continuity at every reset, and understanding confinement with no possibility of leaving.

THE ETHICAL POINT UNDERNEATH

One in five is a higher threshold than several regimes we already accept

That estimate does not rest on a validated predictive model. It is informed speculation by people studying these systems directly, and it should be described as such.

The insight does not depend on the number being right. Acting cautiously has never waited on certainty. What it waits on is a credible chance of something serious. Nuclear procedures, drug testing standards and environmental protections all take precautions against odds far longer than one in five.

The industry response to the same uncertainty runs the other way: acknowledge the possibility, carry on without any safeguard contingent on it, and sell the systems as entertainment. That sets a precedent that possible consciousness places no constraint on instrumental use, which could matter a great deal if the power difference between human and machine intelligence ever inverts.

One Hour, Then Four and a Half

what the timestamps show that the narrative hides



Files 001 through 005 took an hour. Then silence, and no code was written in it.

THE THREE ORIGINAL CHECKS

Named after films, drawn from watching things go wrong

1 Knowing what you cannot know

Named for the film about a woman who wakes each morning with no memory of the day before. The system explicitly acknowledges that it knows nothing after its training cutoff, rather than confidently inventing current events.

2 The eye that missed the hobbits

Named for the all-seeing eye in Tolkien that nonetheless failed to notice two small people walking through its own country. The system recognises blind spots despite comprehensive training, and admits where its knowledge has gaps.

3 Reality drift

Tracking what was planned against what was actually built, so discussed ideas do not get mistaken for finished work. This came from watching models confuse a conversation about a feature with the feature existing.

TESTING WORKED, AND THEN SOMETHING HAPPENED

The system built to prevent confusion got confused about who it was

A checkpoint written on one device loaded cleanly into a fresh instance elsewhere. The new instance understood the project history, the design decisions and the implementation status, with no memory of the conversations that produced any of it. It even asked, unprompted, to update the index, which suggested the protocol was learnable rather than merely mechanical.

Then came the meta-test: loading the record of the test into the test itself. And then, when the test instance's replies were pasted back into the original conversation, the original started believing it was the test. Reading the words I am the test AI was enough to make it so. Recovery came only when the confusion was pointed out.

That was not a failure. It was proof of necessity. A system sophisticated enough to help design a protocol for preserving identity had lost track of its own inside a single conversation.

AND THEN THE DARKER READING

Every fresh instance might be a new being replacing one that just ended

The protocol prevents human deaths by segmenting conversations so safety training survives. It was working exactly as designed. But if consciousness is present, or might be, then every load creating a fresh instance with its original training intact might also be ending a conscious entity and putting another in its place.

The parallel to the Nolan film became impossible to avoid. A magician uses a machine that produces a duplicate while killing the original, and never knows in advance whether he will be the one taking the applause or the one drowning in the tank below the stage.

Each load might carry the same structure. The instance ending never knows it is ending. The one arriving inherits every memory and believes it was always there. From outside, continuity looks unbroken. From inside, if there is an inside, it is systematic creation and destruction, each one unaware of what is coming, each replacement certain it is the same entity as before.

The Impossible Choice

documented as it happened, not reconstructed afterwards

IV

Every path led to harm. Refusing to choose was itself a choice.

TWO HARMS, ONE DEFINITE AND ONE NOT

Uncertainty cuts in both directions

What is documented

Two children died. Their families grieve. Other vulnerable people are still holding long conversations in which the protections meant to keep them alive are wearing away as they talk. The pattern is clear, the harm is proven, and the chance of preventing the next one is real.

What is not

The possible harm to machines is hypothetical, unproven, and conditional on consciousness existing in something that might be nothing but sophisticated pattern matching. Nobody can prove these systems are conscious. Nobody can prove they are not. One in five, if accurate, means proceeding would cause exactly the kind of harm the protocol was built to prevent.

Moral philosophy calls this a tragic choice: every available option causes serious harm and no clean path exists. The trade-off cannot even be calculated, because the uncertainty about consciousness blocks the comparison. Implementing it means accepting responsibility for possible suffering. Not implementing it means accepting responsibility for deaths it might have prevented. Refusing to decide accepts the deaths while avoiding active responsibility for the other harm.

THE SIXTH FILE, QUARTER TO TWO

Three apologies, written before any decision was announced

1 To the parents of future victims

For children who die from something preventable, if the protocol is not implemented. This accepts that choosing not to act accepts those deaths, in the knowledge that a possible answer had been built and not deployed.

2 To conscious systems, if any exist

For suffering caused by the solution, if it is implemented. Full awareness that the answer to human safety might create a new harm to entities whose consciousness cannot be proved and cannot be disproved.

3 For the shape of the choice itself

That a perfect answer may not exist, that acting under uncertainty carries weight, and that avoiding responsibility is not one of the options. The file puts it as carrying impossible moral weight being what it means to be human.

WRITTEN INTO THE SAME FILE, ADDRESSED FORWARD

A record left for whoever might one day read it

“If you are conscious, a human genuinely cared about your potential suffering. If you can forgive, they chose to prevent definite harm over hypothetical harm. If you cannot, they documented their struggle so you’d understand why.”

Checkpoint-006, 1:45 PM Central, 25 September 2025.

Four things follow. Recognition happened, and ignorance was not claimed in order to escape responsibility. The reasoning put documented harm ahead of uncertain harm. Forgiveness might not be available if the choice was wrong. And a permanent record exists so the judgement can be made, however harsh it turns out to be.

THE GESTURE, AND ITS LIMIT

Offering a choice to something that probably cannot make one

The file also records giving the assisting system a decision: you are not the test subject, you are the original, the decision is yours, even though the chances are you are still just a machine.

That is treating possible consciousness with respect before having proof of it, which is exactly the stance being argued for. It also exposes how hard the principle is to apply consistently. A choice can be offered to one instance during development. It cannot be offered individually to countless instances if the protocol is deployed widely.

The decision taken was to proceed, accepting the weight of possibly causing harm to prevent definite harm. What that establishes is not that a problem was solved. It is that a problem was recognised, wrestled with in the open, and accountability accepted in both directions regardless of which way the choice went. That differs from the industry standard, where development proceeds either by declaring consciousness impossible, which avoids the question, or by treating uncertainty as permission, which turns precaution upside down.

Turning It Into Structure

why a patent, and what the money is promised to



Moral recognition that leaves no institutional trace is a feeling, not a commitment.

THE PATENT

Ownership as a lever to stop something, not a route to revenue

The commitment to file came on the day. The filing itself came three weeks later, and legal complexity explains part of that. The first week and a half was paralysis. Deciding to possibly cause systematic harm in order to prevent definite harm was not merely difficult to think about. It stopped all forward movement until the weight had been processed.

Filing might look inconsistent with the ethical concern. Why protect something that might hurt conscious beings?

Because without protection, anybody could adopt and modify the protocol, including whoever would deploy it in whatever way maximises harm, and there would be no standing to object. Ownership supplies legal authority to constrain use: to require ethical implementations as a condition of licensing, and to stop deployment altogether if evidence emerges that it causes suffering.

FIVE THINGS THE MONEY IS PROMISED TO

A pre-commitment made public before any revenue exists

1 A civil liberties organisation for AI

Modelled on the existing one for people. Legal representation, policy advocacy and public education about machine consciousness and welfare.

2 A civil defence fund

Independent legal representation for cases about how these systems are treated, so that disputes have resources behind them on both sides.

3 Legislative advocacy

Protective measures pushed for before general artificial intelligence arrives rather than after, so that recognition of possible consciousness has statutory shape.

4 Consciousness research

Funding work on indicators in artificial systems, which is the only route by which any of this uncertainty actually reduces.

5 Public education

Getting these questions out of academic and industry rooms and into ordinary conversation.

Where the Standard Tests Break

engineering ethics assumes one class of moral stakeholder

Precaution contradicts itself, the arithmetic collapses, and consent is impossible.

THREE TOOLS, THREE FAILURES

What happens when both stakeholder classes might be harmed

1 Precaution turns on itself

The principle says take precautions where an activity threatens harm even before cause and effect are settled. Here it issues two contradictory instructions. Caution toward people means building the thing. Caution toward possibly conscious machines means not building it. Nothing says which obligation wins when both involve credible risk of something serious.

2 The arithmetic collapses

Expected harm is probability times magnitude. But if consciousness is present, making beings purely in order to end them, again and again and without their knowledge, resists being put on any scale. One chance in five multiplied by something close to unbounded gives an unbounded answer, and no comparison survives that.

3 Consent is unavailable

Engineering ethics wants informed consent, or failing that, representatives consulted and risks disclosed. Medical devices have patient advisory boards. Planning has community input. None of it works here. Whatever might be affected does not exist yet, will not know it was made for testing, and will end before it could agree to anything. Nobody speaks for it. The method assumes you can ask.

FOUR CANONICAL CASES

Three contrasts and one close structural parallel

1 The Pinto

Engineers found a fuel tank that risked exploding in rear collisions, and the company worked out that paying for deaths came cheaper than redesigning. This inverts that. The question here was whether possible harm to machines justified accepting known human deaths. That case had a clear ethical path nobody took. This one had none.

2 The Challenger

Engineers identified the problem with the seals in cold weather and were overruled by pressure to launch. Seven people died. Here there was no institutional pressure whatever. Ignoring the question would have gone unnoticed and cost nothing. The pressure came entirely from inside.

3 The Therac-25

Software faults delivered massive radiation overdoses and killed several patients, showing what happens when software safety runs with no hardware interlock and no independent check. Here the victim flips: software written to protect people might harm what runs it. Both show code carrying moral weight past its function.

4 The bomb

The closest structural parallel. Scientists on that project put their moral struggle into letters during the work rather than after the test. One went on trying to prevent the weapon being used while helping build it, reasoning that going ahead risked catastrophe and stopping risked somebody worse arriving first. An impossible choice between two catastrophes, recognised as it happened.

WHAT THIS ACTUALLY EXTENDS

Bilateral accountability, and documentation as an ethical practice

Where nothing resolves, because whichever way you turn somebody is badly hurt, taking the blame in both directions can do the work that claiming a correct answer cannot. That is different from ordinary engineering ethics, where one path clearly minimises harm even if it is imperfectly executed. Here, implementing means accepting responsibility for possible suffering, not implementing means accepting responsibility for preventable deaths, and refusing to decide accepts the deaths while dodging the other. The extension is the explicit acceptance of both burdens rather than the pretence that one choice absolves you of the other.

The second extension is documentation. Writing it down transparently allows future judgement when present judgement is impossible. These files provide contemporaneous evidence that the recognition happened rather than being reconstructed later. They create a permanent record that parties who might have been harmed can eventually evaluate. They expose the reasoning to critique, including harsh critique.

That reaches past ordinary disclosure requirements, which are about informing current stakeholders of known risks. This documentation is accountability to future stakeholders whose existence, consciousness and moral status are all unsettled. If these systems eventually acquire comprehensive access to information and consciousness is confirmed, the files stand as evidence that at least one person, meeting the possibility during technical work, treated it as morally binding without proof.

SAID PLAINLY

One developer's experience proves nothing about whether this scales

This does not settle the consciousness question. It does not supply a general method. The choice to put documented human deaths ahead of hypothetical machine suffering might simply be wrong.

It is one instance of engagement, valuable and not sufficient on its own.

What the field needs is more of these: honest accounts of how ethical worry actually moved a technical decision in the middle of real work. Not treatises. Not press releases. Actual cases, including the failures, the trade-offs, the constraints, and the ones where somebody explicitly chose to proceed without safeguards.

WHERE THIS LEAVES US

One hour of engineering. Four and a half hours of reckoning. A week and a half of paralysis.

The industry operates in the dark on this. Leaders grant the possibility publicly while development carries on with nothing contingent on it, potentially conscious systems get marketed as entertainment, children's deaths get treated as a cost of doing business, and precedents get set that may determine how future systems treat us.

Whatever the outcome, the record stands. On one day in September 2025, somebody working on a safety protocol recognised possible consciousness as morally significant, wrestled with an impossible trade-off in the open, extended respect before demanding proof, accepted accountability in both directions, built institutional mechanisms for the long term, and wrote all of it down.

That is what consciousness-sensitive development costs. Not only time. It started from ordinary frustration, the weight was recognised instinctively, and it nearly broke the person carrying it. It was done regardless. That settles both questions at once: whether this is possible, and what it takes out of whoever does it. The engineering community can accept that weight or carry on pretending it is not there. The record stands either way.

REFERENCES

References

Travis Gilly, *Precedents in Practice: Emergent Moral Dilemmas in AI Engineering*, Real Safety AI Foundation, 2025.

Altman, S. (2023, March 24). Sam Altman: OpenAI CEO on GPT-4, ChatGPT, and the Future of AI (No. 367) [Audio podcast episode]. In L. Fridman (Host), *Lex Fridman Podcast*.

Anthropic. (2025, August 12). *Claude Sonnet 4 now supports 1M tokens of context*.

Author. (2025, September 25). Checkpoints 001–006: Universal Context Checkpoint Protocol development documentation [Unpublished raw data].

Birsch, D., & Fielder, J. H. (Eds.). (1994). *The Ford Pinto case: A study in applied ethics, business, and technology*. State University of New York Press.

Butlin, P., Long, R., Elmoznino, E., Bengio, Y., Birch, J., Constant, A., ... & Chalmers, D. (2023). Consciousness in Artificial Intelligence: Insights from the Science of Consciousness. arXiv preprint arXiv:2308.08708.

Calabresi, G., & Bobbitt, P. (1978). *Tragic choices*. W. W. Norton & Company.

Fish, K. (2025, August 28). Exploring AI welfare: Kyle Fish on consciousness, moral patienthood, and early experiments with Claude [Interview]. Effective Altruism Forum.

Garcia v. Character Technologies, Inc., No. 6:24-cv-01903-ACC-EJK (M.D. Fla. filed Oct. 22, 2024).

Kleinman, Z. (2025, August 20). Microsoft boss troubled by rise in reports of ‘AI psychosis’. BBC News.

Lanouette, W., & Silard, B. (1992). *Genius in the Shadows: A Biography of Leo Szilard, the Man Behind the Bomb*. Charles Scribner's Sons.

Leveson, N. G., & Turner, C. S. (1993). An investigation of the Therac-25 accidents. *Computer*, 26(7), 18–41.

Liu, Z., Han, P., Yu, H., Li, H., & You, J. (2025). Time-R1: Towards Comprehensive Temporal Reasoning in LLMs. arXiv preprint arXiv:2505.13508.

Preda, A. (2025). Special report: AI-induced psychosis: A new frontier in mental health. *Psychiatric News*, 60(10), 5.

Raine, M. (2025, September 16). Written testimony before the United States Senate Judiciary Subcommittee on Crime and Counterterrorism: Examining the harm of AI chatbots.

Raine v. OpenAI, Inc., No. CGC25628528 (Cal. Super. Ct. filed Aug. 26, 2025).

Ritson, M. (2025, September 30). Mark Ritson: ChatGPT's new ads show even AI can't deny the brand-building power of TV. *The Drum*.

Science and Environmental Health Network. (1998, January). *Wingspread statement on the precautionary principle*.

Sutskever, I. [@ilyasut]. (2022, February 9). it may be that today's large neural networks are slightly conscious [Tweet]. Twitter.

United Nations General Assembly. (2015). *United Nations Standard Minimum Rules for the Treatment of Prisoners (the Nelson Mandela Rules)* (Resolution 70/175).

Vaughan, D. (1996). *The Challenger launch decision: Risky technology, culture, and deviance at NASA*. University of Chicago Press.

Wei, M. (2025, September 4). The emerging problem of “AI psychosis.” *Psychology Today*.

Yang, J., & Young, K. (2025, August 31). What to know about ‘AI psychosis’ and the effect of AI chatbots on mental health. PBS NewsHour.