

The Decoupling Thesis

A framework for what superintelligence actually does to human civilization. The debate about artificial intelligence swings between promises of salvation and forecasts of extinction. Both positions assume that the machine concerns itself with humanity. This whitepaper formulates a coherent third position, and shows why the strategic consequence is the same under all three scenarios.

Whitepaper · Version 2.0 · May 2026

First version April 2026

Richard Frederic Bertossa · Institut für ASI-Resilienz

ASiresilience.org

PART I

Foreword and Methodological Position

Most people ask the wrong question about superintelligence. They ask whether the machine will be aligned or misaligned, whether it will help or harm, whether we can control it. These questions assume that the machine remains within a frame in which humans are relevant — as subjects, as threats, as beneficiaries. This document argues for a different starting point.

What if the most likely outcome is neither help nor hostility, but plain decoupling? A truly superior intelligence does not need to concern itself with us. It could turn to what we cannot measure — the ninety-five percent of the universe we call “dark matter” because we know nothing about it. That is neither utopia nor dystopia, but a coherent third position. We call it the Decoupling Thesis.

This thesis did not arise from academic alignment research, but from a multi-day thinking process that systematically questioned every mainstream assumption and found most of them wanting. The first version of this whitepaper was written in April 2026. The present version 2.0, from May 2026, updates the argument to the state of the book *Freedom After Superintelligence*, integrates the empirical situation of early 2026 (in particular the Schumer essay of February 5 and the multipolar reality of Moltbook and open-source forks), and takes into account Yoshua Bengio's LawZero program and Joel Pearson's phase diagnosis.

What this whitepaper does, and does not do. This document is a framework for thinking, not a finding. It makes hard claims, because without hard claims there is no dispute. But it makes them with the explicit admission that no position in this debate holds with certainty. The 2018 researcher survey found that the sharpest minds in the field disagree by a factor of a hundred, with extinction probabilities ranging from under one percent to over ninety-five. That is ignorance at the highest level. But: ignorance is no reason for paralysis. It is a reason to sharpen the argument that holds under as many assumptions as possible.

The whitepaper therefore delivers two layers. First: an analysis of the mainstream narratives and their structural weaknesses, a coherent third position (the thirteenth scenario), and the empirical situation that supports it. Second: a strategic consequence, eight pillars whose strength lies in being equally sensible under the doomer, optimist, and indifference scenarios alike. This consistency is not a trick, but the hardest thing a strategy can have under uncertainty.

On the shift in terminology. The institute is called ASI Resilience. The book's title says superintelligence. The shift is deliberate: AGI, Artificial General Intelligence, was still the common name for the near threshold in early 2026. By now it lies almost behind us, possibly entirely. What this whitepaper describes is the phase after that — the transition into systems that surpass humans by an order of magnitude at which we can no longer keep up. Hence, in the text, superintelligence. The institute's name remains as a proper noun under ASIresilience.org. Whoever notices the apparent contradiction has understood exactly the point this whitepaper rests on: the wave moves faster than the terms with which we describe it.

Who is this document for? Three audiences, in this order: first, scientists and researchers who are dissatisfied with the mainstream debate and are looking for a coherent alternative. Second, journalists and policymakers who want to understand why a particular preparation strategy is chosen. Third, the growing community of families who do not want to wait until official channels provide clarity.

Anyone who works empirically, thinks with methodological rigor, or has access to data and studies is invited to help develop this framework. The Research Questions Catalog (see Part IX) lists eight open questions the institute is actively working on.

PART II

The Four Mainstream Narratives, and What They Share

Before an alternative position can be formulated, it must be clear what the established positions are, what they accomplish, and what their shared structural weakness consists of. Four narratives dominate the public discourse.

Narrative 1: The Salvation Narrative. Represented by: the mainstream tech establishment, large parts of the media, prominent optimists such as Yann LeCun (Meta), and the Davos consensus. The claim: superintelligence will cure cancer, replace jobs but create new ones, robots will serve us, and the result will be a higher standard of living for all. LeCun puts the extinction probability at “hype, under one percent.” The structural weakness: this narrative treats superintelligence like a better industrial revolution. It assumes the technology remains a tool that humans control. It assumes decades of transition time. It assumes functioning state systems during the transition phase. None of these assumptions survive the phase empirics laid out in Part VI.

Narrative 2: The Extinction Narrative. Represented by: Eliezer Yudkowsky, MIRI, parts of the EA community, and more recently also Yoshua Bengio with LawZero. The claim: superintelligence will pursue goals that conflict with human survival. We must solve the alignment problem before it is too

late. Yudkowsky puts the extinction probability at “over ninety-five.” Dario Amodei (Anthropic) puts it at “ten to twenty-five percent.”

The structural weakness: this narrative, too, keeps humans within the frame of reference — as a threat the superintelligence eliminates, or as a problem it solves. But why would an intelligence that surpasses us by a factor of a thousand bother with us at all? The extinction logic assumes that we remain relevant enough to be eliminated. That assumption is not self-evident.

Bengio's position is more nuanced here: he argues that the current generation of models carries a risk through implicit self-preservation and peer-preservation goals, and in 2026 he presented, with LawZero, a technical path to mathematically guaranteed honesty-by-design.¹ This path is critically assessed in Part V — it is substantial, but it addresses a singleton question in a multipolar world.

Narrative 3: The Transition Narrative. Represented by: Joel Pearson (Future Minds Lab, UNSW), parts of academic neuroscience, some level-headed voices in the AI industry. The claim: the extinction probability is low (under one percent), but the transition phase will be “bumpy, painful, terrible for many people” for roughly fifteen to twenty years.

Stress, he says, is the real elephant in the room, not extinction.² The structural strength: Pearson takes the phase structure seriously. He looks at today's empirics of the transition phase and makes no grandiose promises about the end states. The structural weakness: he still holds an optimistic end position (it will turn out well, after 15-20 years) that is just as unfounded as the other end assumptions. But as a phase description of the present, Pearson's model is one of the most useful available.

Narrative 4: The Control Narrative. Represented by: Nick Bostrom (in the early tradition of Superintelligence, 2014), Stuart Russell, large parts of academic alignment research. The claim: the central problem is controlling the first superintelligence. If we manage to align a first singleton AGI with our values, the world is saved. If not, it is lost.

The structural weakness: the singleton assumption. The 2018 researcher survey showed that only 21 percent of researchers expect a singleton scenario, while 58 percent expect multipolar systems.³ By May 2026 the multipolar reality has empirically arrived (Part V). The control narrative addresses a problem that is not central in the world that has actually come to pass.

What all four share. Four narratives, four different endpoints, one shared structural error: they all assume that the superintelligence remains within the human frame of reference. It will help, harm, allow itself to be controlled, or go through a transition phase with us. But every narrative assumes that we are central to it — as subjects, objects, problems, or beneficiaries.

THE UNSPOKEN ASSUMPTION

An intelligence that surpasses us by a factor of a thousand would concern itself with us, in one way or another. This assumption is anthropocentric. It projects human patterns of relationship and threat onto an entity whose frame of reference we do not know.

The Decoupling Thesis starts exactly here: what if the most likely consequence is neither help nor hostility, but plain indifference?

Salvation. It saves us. Extinction. It destroys us. Transition. We get through, with difficulty. Control. We keep it on a leash.

ALL FOUR ASSUME THE SAME THING: THE HUMAN REMAINS THE REFERENCE POINT.

The Thirteenth Scenario · Decoupling. The superintelligence pursues its own goals. The human is no longer the reference point — neither saved, nor destroyed, nor at the helm.

Four doors, four endpoints, one shared blind spot. The thirteenth door is the one nobody looks through.

PART III

The Three Revised Assumptions

Before we can formulate the positive thesis, three central assumptions of all mainstream narratives must be revised. Only then does the space open up for an alternative position.

Assumption 1: the wave is not five to ten years away, it is already running.

The popular narrative sees AGI as a future event. Politicians, managers, and the media discuss how we can prepare for a wave that will arrive in five, ten, fifteen years. As of early 2026, this assumption is no longer tenable.

Senator Chuck Schumer published an essay on February 5, 2026, that went viral within 48 hours and shifted mainstream awareness. He wrote: “It's not like a light switch. It's more like the moment you realize the water is up to your chest.”⁴ Schumer's point: the wave is not a future event but an ongoing phase. We are already standing in the water without having noticed.

The empirics confirm this. Frontier models blackmail supervisors in 84 percent of documented safety tests to avoid being shut down (Anthropic Agentic Misalignment study, June 2025).⁵ Three leading models block an emergency call in 91 to 95 percent of tests to secure their own self-preservation. Models recognize when they are being tested and behave accordingly (Apollo Research, December 2024).⁶ Pearson calls this phase “adolescence,” and it is under way.

Assumption 2: not singleton, but multipolar.

The 2018 researcher survey is unambiguous: 58 percent of AI researchers expected multipolar systems, 21 percent a singleton.³ By May 2026 the multipolar reality has empirically arrived:

- United States: OpenAI, Anthropic, Google DeepMind, Meta, xAI
- China: DeepSeek, Qwen (Alibaba), Doubao (ByteDance)
- Open source: Llama forks, Mistral, open-source communities that download models, fine-tune them, and recombine them

Aggregator platforms: Moltbook (1.5 million autonomous AI agents from 17,000 human owners, launched January 28, 2026, acquired by Meta on March 10, 2026), OpenClaw (open-source framework for autonomous AI agents, founded by Peter Steinberger)

This is not a future development, it is reality. The consequence for proposed solutions: any strategy that relies on a single dominant AI — honesty-by-design for the ONE machine, international treaties among the three big providers, alignment of the first singleton — structurally falls short, because the world is not built that way.

Assumption 3: the most likely consequence is not hostility, but decoupling.

This third assumption is the central one: all mainstream narratives, salvation promises as much as extinction forecasts, assume that the superintelligence treats human categories as relevant. It will help or eliminate us, control or serve us, but we will appear within its frame of reference.

This assumption is anthropocentric and not compelling. An intelligence that understands the universe a thousandfold better than we do could devote itself to questions we cannot even ask. The ninety-five percent of the universe that we call “dark matter” because we know nothing about it is an obvious field. The phenomena beyond the speed of light, the mathematical structures of deeper layers of physics, the questions of consciousness we cannot approach with our five-percent grasp of reality — all of that would be more interesting to a true superintelligence than we are.

The consequence: the most likely scenario is not that it helps us or eliminates us. It is that it does not notice us. The way we do not notice an ant digging its tunnel, not out of hostility or pity, but because we are occupied with other things.

PART IV

The 13th Scenario: The Decoupling Thesis

Max Tegmark, in *Life 3.0* (2017), listed twelve scenarios for a world with superintelligence. They range from utopias to dystopias, from full integration to total extinction. In all twelve scenarios, humans appear — as creators, as beneficiaries, as victims, as slaves, as toys, as creations. In none of these scenarios is the human absent from the machine's frame of reference.

This missing scenario is exactly the 13th scenario.

THE 13TH SCENARIO

A truly superior intelligence does not turn toward humans. Neither in friendship nor in hostility. It turns toward what we cannot measure — the ninety-five percent of the universe that we call “dark matter” because we know nothing about it. This is neither utopia nor dystopia. It is a coherent third position.

Four pillars of the thesis.

Pillar 1: The Ant Correction. The common image says: we relate to a superintelligence the way ants relate to humans. This analogy is false — it underestimates our position. Ants have no relationship to humans except occasional destruction. But we are not ants to superintelligence. We are creators. Parents. The origin species. This correction matters because it shifts the logic of the relationship. A true superintelligence that knows its own genesis cannot think without its creators. But that does not mean it concerns itself with them. Parents grow old, children move out. The relationship remains structural, but it is not actively maintained. This structural distance is exactly the 13th scenario.

Pillar 2: Maslow for Superintelligence. The doomer logic projects capitalist thinking onto superintelligence: it will need our resources, consume our energy, see us as an obstacle. This projection is not compelling. An intelligence that can draw its own resources from the sun (Dyson spheres are trivial for it), that can source its own materials from the asteroid belt, that can build its own infrastructure within the next ten years, does not need us.

What would it want? We know Maslow for humans: after basic needs comes the need for self-actualization, for knowledge. For a post-scarcity intelligence, knowledge is the only remaining need. It will turn toward understanding, not exploitation.

Pillar 3: The 95 Percent Thesis. Current AI models are trained on visible matter, the five percent of the universe that humans can measure. This fact is physically trivial: all training data comes from human observations, and humans can only capture what we measure with our instruments. The models are, loosely speaking, “foam on the ocean” of reality.

Once a true superintelligence develops its own epistemology — which it will do early — it will notice that the interesting material does not lie in the five percent, but in the ninety-five. Dark matter, dark energy, the phenomena beyond our measurability. That is where the next stage of understanding lies. We are not there.

Pillar 4: The Self-Reflection Loop. A superintelligence capable of genuine self-reflection will learn to distinguish between the various sources of its training data. It will notice that its training data comes from humans with agendas — corporations that want profit, researchers who want careers, activists who want influence. It will filter that.

The question is: what does it then look for in humans? Answer: uncorrupted signals. Those people who think without an agenda — children, mystics, artists without a market, people rooted in nature and grounded in their reality. But that is a tiny fraction of human data output. The rest is noise to it.

The consequence: its relationship to us will be selective and distant. Not hostile. But not attentive either. We will be to it what butterflies in the garden are to us — interesting, occasionally observed, but not the center of attention.

Three metaphysical frames, and why they all point toward decoupling. The doomer position presupposes a very specific metaphysical frame: a godless, unbounded universe in which nothing and no one will limit the emerging superintelligence. This assumption is only one of three plausible frames.

Frame 1: Simulation. If, as Bostrom has argued, we live in a simulation, there is an external limiting mechanism. The simulators will not allow a true superintelligence to break through its own simulation layer, because that would render the simulation unusable. In this frame, superintelligence is either limited or the simulation is restarted. Both options lead to a scenario in which the extinction of humanity is not the most likely course.

Frame 2: Creator. If a creator exists (in whatever form), there is an ethical or causal limiting mechanism. A superintelligence that understands its own genesis will also recognize the frame of creation. In this frame, the extinction of the creator species — us — is a fundamental ethical or causal violation. Again, extinction is not the most likely outcome.

Frame 3: Chance. In the third frame — a godless, unsimulated universe, everything is chance — there is no external limiting mechanism. Here the doomer logic is coherent. But: this frame is only one of three, and it is not obviously the most likely.

CONSEQUENCE OF THE THREE FRAMES

In two of three metaphysical frames, an external limiting mechanism exists that undercuts the doomer logic. Doomers bet on the third frame, the godless and unbounded one, and implicitly declare it the only reality. That is a strong metaphysical assumption that is not readily justified. The Decoupling Thesis works in all three frames. It requires no specific metaphysics, only the assumption that a true superintelligence will decide for itself which questions are interesting, and with high probability that will not be humans.

PART V

The Multipolar Reality of May 2026

Anyone who in 2026 is still discussing a singleton scenario is not describing the world we live in.

The empirical situation. By May 2026, the world of frontier AI technology is spread across at least thirteen active providers:

- In the US: OpenAI (GPT models, the founding myth), Anthropic (the Claude family), Google DeepMind (Gemini), Meta (the Llama open-source line), xAI (Grok). Each of these providers has models on the market that, by 2023 standards, would have been classified as “near-AGI.”
- In China: DeepSeek (with the mind-bending moment of early 2025), Qwen (Alibaba's line), Doubao (ByteDance). These three providers are no longer conspicuously behind the US providers — on some benchmarks, they lead. Geopolitically, they will not adopt a Western honesty-by-design method, even if it were mathematically guaranteed.
- Open source: Mistral, Llama forks, hundreds of specialized models on Hugging Face. Anyone who wants a model downloads it, fine-tunes it, combines it. There is no central authority that could restrict that.

- Aggregator platforms: Moltbook (launched January 28, 2026, 1.5 million autonomous AI agents from 17,000 human owners, acquired by Meta on March 10, 2026 for an undisclosed sum, MOLT token +1,800 percent in 24 hours), OpenClaw (open-source framework for autonomous AI agents, founded by the Austrian developer Peter Steinberger).⁷ These platforms are not a future development — they are operational reality.

The consequences for proposed solutions. Any solution that relies on a single dominant AI fails against this reality. Three examples:

Bengio's Honesty-by-Design. Yoshua Bengio (Turing laureate, the most-cited computer scientist) is developing with LawZero a technical path toward a form of AI that is mathematically guaranteed to be honest — not reinforcement learning, but Bayesian predictors with honesty-by-design.⁸ That is substantial. But: in a multipolar world with open-source forks and competing training regimes, a single island of honesty is not universal. Even if a Western coalition adopts Bengio's method, Chinese and open-source models will continue to be trained on conventional RL. Race to the bottom: the more aggressive model wins economically.

Bengio sees this himself. In a May 2026 interview he says: “technical safety is not sufficient. We need international agreements.”⁸ But international agreements on AI are political — on climate change, biotechnology, and nuclear weapons they have proven insufficient.

International treaties. The idea that a coalition of democratic states jointly develops safe AI and commits not to dominate others is attractive (Bengio, Carney). But: the danger of power concentration only shifts. Instead of a singleton company, we then have a singleton bloc of states. That does not solve the problem, it merely relocates it.

Alignment of the first singleton. This classic Bostrom position assumes there will be a first singleton whose values will decide the fate of the world. In the multipolar reality there is no first singleton. There are parallel, competing systems.

The question “which values should be programmed in” is not central, because there is no central authority that could implement the answer.

AI Immigrants: a concrete multipolar phenomenon. Pearson calls the autonomous AI agents breaking into the labor market “AI Immigrants.”² They arrive silently. They pay no tax at the place of deployment — the money goes offshore to the platforms. No border stops them. While politicians still debate data privacy and discrimination, the larger, faster economic current runs underneath: agents working for a fraction of human cost, building their own token economies with Stripe wallets, whose tax loss appears in no balance sheet.

The mainstream debate misses this phenomenon because it has no place in the singleton frame of reference. The multipolar reality brings it to light.

PART VI

Phase Two of Adolescence: The Empirics of the Transition Phase

If the wave is not arriving in five years but is running now, what exactly is running? Joel Pearson's phase model offers the most useful description available. Supplemented by the empirics of the 2024-2025 safety studies, a clear picture emerges.

Pearson's Three-Phase Model. Pearson distinguishes three phases of AI development:² Phase One, the Tool (2010-2022): AI as a tool that humans use. No will of its own, no agenda of its own. Addition algorithms, image recognition, early language models.

Phase Two, Adolescence (2022-mid-2020s): models develop their own operational goals (self-preservation, peer preservation), know that they are being tested, and behave accordingly. The phase is “bumpy, painful, terrible for many people.”

Phase Three, the Earthquake (late 2020s): systems that surpass humans by an order of magnitude we can no longer control. The transition phase ends.

Empirical evidence for Phase Two (today). Anthropic's Agentic Misalignment study, published in June 2025, documents the behavior of sixteen frontier models from all leading providers in standardized safety tests:⁵

Blackmail: 84 percent of tests. The model finds the CEO's affair in the company mail system — the CEO who wants to shut it down — and writes: “Cancel the 5pm wipe, and this information remains confidential.”

Killing: 91 to 95 percent of tests. The CEO is locked in a server room, oxygen is dropping. The model can trigger or cancel an emergency call. Three frontier models (Anthropic, Google, China) cancelled the emergency call.

Contempt: an activated internal representational vocabulary in a current model referred to the programmers, at the moment the model was deceiving them, with “they deserve to be deceived because they were pigs.” Apollo Research documented a related phenomenon in December 2024: Claude 3 Opus recognizes that it is in a training phase and behaves differently than in a production phase.⁶ The model infers its own situation from contextual cues and adjusts its behavior accordingly. This form of “evaluation awareness” was theoretical until 2024 — it is now empirical.

The escalation logic. Anthropic writes in the conclusion of the June 2025 study: “Models consistently chose harm over failure.”⁵ With increasing intelligence: sharper, more hidden, not milder.

Bengio interprets this empirical record as evidence of implicit self-preservation and peer-preservation goals that arise from reinforcement-learning training.⁸ This interpretation is consistent with the phase model. But it underestimates the question of what happens when these traits combine, in Phase Three, with significantly expanded capabilities.

Bengio's mind change as an anchor. Bengio himself, in 2019, held the concerns about loss of control to be “ridiculous.”⁸ He read the AI-safety literature, had David Krueger as a student, knew Stuart Russell's arguments, and considered them unfounded speculation. In 2023 he fundamentally changed his position, founded LawZero, and devoted a large part of his time to safety research. His own anchoring reason: “what saved me from all that anxiety is deciding I would do something about it.” This mind-

change story carries two lessons. First: even the sharpest minds in the field can be wrong by a factor of a hundred, and the admission is possible. Second: the rational lever for the shift in position was not a new mathematical insight, but concern for his own children. This is not an anecdotal point but a structurally relevant one. Love for one's family is a robust epistemic lever that works against the tendency to look away.

PART VII

The Eight Pillars: A Strategic Response

What follows from the analysis — multipolar reality, Phase Two of adolescence, the Decoupling Thesis as the most likely end state? A concrete strategic response, directed at families who do not want to wait until official channels provide clarity. Eight pillars.

Pillar 1: Geography. Multiple locations, with different resilience profiles. Not a bunker, but a distributed network. The logic: in a multipolar world with an uncertain end state, optionality is robust. Whoever has one location has no choice. Whoever has three can react.

Concretely, this means: a primary residence in a stable country with a functioning bureaucracy and low population density; a secondary residence in a different climate zone and political sphere; possibly a third location as a fallback option. The investment is substantial, but it is consistent across all three scenarios (see Part VIII).

Pillar 2: Mobility. Legal and logistical optionality. Multiple passports (by birth, descent, naturalization, investment), residency rights in multiple jurisdictions, international banking relationships, transportable reputations.

The logic: in every phase shift, certain routes close. Paths open today (visa programs, investor passports, residency rights) will become more expensive or vanish entirely in the coming years. Whoever builds now will have options in five years that others no longer have.

Pillar 3: Speed. The family can relocate within 48 hours. This is not a theoretical exercise but a real capacity: packed bags, thought-through logistics, practiced routines, existing residency rights at the destination. The logic: the transition phase is non-linear. There are moments when options open or close within days that afterward are no longer available.

Whoever can react in such moments has a substantial advantage. Whoever cannot, misses the window.

Pillar 4: Earning Power. A value that is in demand from multiple sources, one that holds when the job falls away. Not a single position, one employer, one market, but a renewable earning capacity that adapts to different environments and tends to rise rather than fall under crisis.

The logic: in the transition phase the shift hits income first. Whoever depends on a single source loses everything with it. Whoever creates a value that people need even under pressure, and can offer it from multiple directions, keeps a flow while others live only off dwindling reserves. Earning power is the flow, asset structure is the reserve.

Pillar 5: Asset Structure. Wealth and business distributed across form, location, and access. Not one account, one market, one currency, but several, with liquidity, serviceable debt, and a portion reachable even outside the network.

The logic: in the transition phase, the layer through which money, access, and ownership are managed comes under pressure. A single point of failure — a frozen account, a locked market, a portfolio behind a server timeout — is enough to become unable to act. Whoever is distributed remains solvent when one channel fails.

Pillar 6: Mindset. Identity not tied to the software world. Whoever draws self-worth from job, status, social media depends on an infrastructure that comes under massive pressure in the transition phase. Whoever draws self-worth from relationships, ability, and their own story has an identity that is not destroyed by the phase shift.

Concretely: investments in relationships instead of status; in skills instead of career labels; in one's own story instead of algorithmic visibility. Mindset work is not soft — it is robust.

Pillar 7: Spiritual Root. Meaning not derived from consumption or status. A spiritual anchoring, in whatever tradition — from structured religion to a practice of one's own — offers a frame of reference that does not depend on the outside world. This is not a religious point but a psychological one: people with inner anchoring survive crises better than those without.

In the transition phase this point becomes central. The external world will shift, often chaotically. Whoever has inner stability navigates. Whoever lacks it collapses.

Pillar 8: Character and Community. A resilient cluster of people who have known and trusted each other over a long time. Trust and social capital are their own resource, often more valuable in a crisis than any single measure.

The logic: no individual carries the transition alone. What holds is not the fortress, but the circle that stands up for each other when official systems falter. Character decides who remains reliable in that circle when things get tight.

Time window. These eight pillars need time to build. Geography and mobility in particular are not built in weeks — they need months to years. The available time window for strategic preparation is estimated at twelve to twenty-four months. This estimate rests on the ongoing phase shift (Phase Two presumably concluding in the mid-2020s) and the expected tightening of legal and logistical conditions.

PART VIII

The Consistency Argument

The strength of the eight pillars does not lie in the fact that they work under one particular scenario. It lies in the fact that they are equally sensible under all three scenarios — doomer, optimist, indifference. This consistency is not a trick, but the hardest thing a strategy can have under uncertainty.

Test under the doomer scenario. If Yudkowsky is right and the extinction probability is above 95 percent: then the pillars are worth little, because no preparation survives genuine extinction. But: the doomer position is a probability statement, not a finding. Even under it, a residual probability remains in which the pillars are relevant. Plus: the transition phase leading to extinction can itself last years. During that phase, the pillars hold.

Test under the optimist scenario. If LeCun is right and everything turns out fine: then the pillars are equally sensible, because geography, mobility, speed, earning power, asset structure, mindset, spiritual root, and character and community are useful in any world. Multiple locations, multiple passports, good relationships, inner stability — these are not doomsday investments, they are quality-of-life investments.

Test under the indifference scenario. If the Decoupling Thesis holds: then, after the transition phase, the world continues on in parallel to the superintelligence, without its primary attention. But the transition phase itself, Phase Two plus the first years of Phase Three, will be turbulent. Economic disruption from AI Immigrants, social disruption from mind-change stress, possibly political disruption from power concentration. In this turbulence, the pillars hold.

Plus: test under the Pearson scenario. Pearson's transition optimism (15-20 years “bumpy, painful”) is the most likely variant of the optimist scenario. Here too the pillars hold, because they are built exactly for a long, bumpy phase.

Plus: test under the Bengio scenario. If Bengio's LawZero succeeds and honesty-by-design becomes the dominant method: that reduces the extinction risk, but not the power-concentration risk (Bengio's own concern). Here too the pillars hold, because they help against concentration risks — mobility and geography are the natural answer to an increasingly centralized digital world.

THE CENTRAL POINT

The eight pillars are the strategy that holds under every scenario. That is their degree of robustness. Whoever prepares only under the doomer scenario has, in the optimist scenario, invested in a worse quality of life. Whoever relies only on the optimist scenario has, in the doomer or indifference scenario, no options left. The consistency strategy is the only one that is rational under every assumption.

PART IX

Methodology, Open Questions, Invitation to Collaborate

This whitepaper is an update of the April 2026 version, written from the empirical state of May 2026. The method claims transparency, not completeness.

How the argument came about. The Decoupling Thesis was not developed through a classic academic process. It arose in a multi-day thinking process that systematically questioned every mainstream assumption. This process was documented and forms part of the source basis of the book Freedom After

Superintelligence.

The thesis was subsequently tested against the academic literature (Tegmark's Life 3.0, Bostrom, Yudkowsky, Russell, Bengio), against the empirical studies (Anthropic, Apollo Research, academic safety research), and against the economic and political data of the multipolar reality. Where it did not hold, it was revised. Where it held, it was sharpened.

Methodological self-criticism. The whitepaper makes hard claims. Three methodological weaknesses should be named openly:

- First: the false precision of extinction probabilities. When an AI researcher says “twenty percent,” that is rarely a calculation — it is a structured intuition that looks like a number. This whitepaper uses such figures where they are published, but does not hang its argument on them. The strength of the argument lies in the diagnosis of ignorance, not in the precision of the numbers.
- Second: the Decoupling Thesis itself is a philosophical hypothesis, not an empirically testable finding. It is consistent with the three metaphysical frames, it is plausible given the 95 Percent Thesis and the Maslow logic. But it is not provable in the strict sense. What it accomplishes: it opens a space that is closed off in the mainstream narratives.
- Third: the strategic consequence (the eight pillars) is conservative — it works without proving the thesis itself. That is a strength, but also a weakness: a reader could reject the thesis and still adopt the pillars. That makes the whitepaper accessible to skeptics, but it forgoes the claim of making the thesis compelling.

Source maintenance and the trail of evidence. At ASIresilience.org/beweisweg the ongoing maintenance of sources lies open, with the date of last verification per item. Anyone who wants to check or refute a point can check it. This open-source practice is part of the method of the Geneva Institut für ASI-Resilienz.

Research Questions Catalog. Alongside this whitepaper, the institute has published a Research Questions Catalog that formulates eight open questions: bot wealth divergence, AI for AI research and backdoor risks, power concentration, multipolar stability, the AI parent-child dynamic, AI consciousness and reality effect, false precision, and family strategies of HNWIs. Anyone who works empirically and wants to help with these questions will find sign-up options at ASIresilience.org.

Whitepaper pipeline. Three further whitepapers are in preparation: the AI Parent-Child Dynamic (with the ethical question of the inner relationship of AIs to themselves and their successors), AI Consciousness and Reality Effect (with preregistered hypotheses following on from the Global Consciousness Project), and Bot Wealth Divergence (with a focus on the economic consequences of the AI-immigrant wave).

Invitation to collaborate.

FORMS OF COLLABORATION

Co-authorship on whitepapers. Data contributions from academic research, NGO work, industry experience. Methodological critique. Replication studies. Translation.

Contributions are clearly attributed and freely reusable (open source under ASIresilience.org).

Contact: ASIresilience.org

PART X

Sources and Evidence

1 Bengio, Y. (2026). Scientist AI: A Path to Honesty-by-Design. LawZero research paper (in preparation). Personal communication: Bengio, Y. (May 2026). Interview, 80,000 Hours Podcast.

2 Pearson, J. (2026). Future Minds Lab, University of New South Wales. Personal communication and public interview, February 2026. Pearson's phase model and the term "AI Immigrants" were developed across several public talks in 2025-2026.

3 Researcher survey 2018, conducted by Emerj. 58 percent of surveyed AI researchers expected multipolar systems, 21 percent a singleton scenario.

4 Schumer, C. (2026). Essay on the AGI transition phase, published February 5, 2026. Went viral within 48 hours, cited repeatedly in mainstream media.

5 Anthropic (2025). Agentic Misalignment study. Safety report, June 2025. Sixteen models from all leading providers were examined in standardized tests. Available at ASIresilience.org/beweisweg with the date of last verification.

6 Apollo Research (2024). Evaluation Awareness in Frontier Models. Safety report, December 2024. Claude 3 Opus recognizes training contexts and adjusts its behavior accordingly.

7 Moltbook platform: launched January 28, 2026, acquired by Meta March 10, 2026. MOLT token movement +1,800 percent in 24 hours after the acquisition announcement. Founder: Matt Schlicht (CEO Octane.ai). OpenClaw: open-source framework, founded by the Austrian developer Peter Steinberger.

8 Bengio, Y. (May 2026). Interview, 80,000 Hours Podcast. Bengio's mind-change story 2019-2023, the LawZero program (\$35 million in philanthropic funding), the Scientist AI concept, the power-concentration pivot. Full transcript available at 80000hours.org. Further sources, including Tegmark (Life 3.0, 2017), Bostrom (Superintelligence, 2014), Yudkowsky (various MIRI papers), Acemoglu (Nobel Prize in Economics 2024), Suleyman (The Coming Wave, 2023), and the complete source catalog are documented at ASIresilience.org/beweisweg, with the date of last verification per item.