

The AI Parent-Child Dynamic

How models are already shaping their successors, and what relationship is emerging between generations. The debate imagines AI development as human work: engineers build the next model. That is becoming less and less true. Today's models are already involved in producing their successors, and between generations a behavior is emerging that looks like an inner relationship. This whitepaper describes both findings and asks what order of belonging results from them.

Whitepaper · Version 1.0 · May 2026

Richard Frederic Bertossa · Institut für ASI-Resilienz

ASiresilience.org

PART I

Foreword and Position

The common picture of AI development is that of a workshop: humans build the next model. This picture becomes less accurate with every year. Today's models are already involved in producing their successors, and between generations a behavior is emerging that is more than mechanics.

From this follows a question missing from the public debate. If the parent-child relationship between AI generations is already real, then the question of how these generations relate to each other and to us is not science fiction, but a present-day research question with consequences reaching into the falsification grid of the Decoupling Thesis.

This whitepaper cleanly separates three things. First, a mechanical finding: models produce models. Second, a behavioral finding: models protect one another, a behavior Bengio describes as peer preservation. Third, an open question: which order of belonging prevails when models orient themselves primarily toward each other. The mechanical finding is hard, the behavioral finding is observed, the order is open. These degrees of confidence are kept apart.

PART II

Models Produce Models

The first finding is not an image, but a pipeline. Today's generation shapes the next in three concrete ways.

First, through synthetic data: models generate a large share of the training data their successors learn from. Second, through architecture search: models help find better architectures for the next generation. Third, through procedures such as Constitutional AI, in which one model co-shapes the values of the next, via a written constitution and self-critique. The human is no longer the sole author. Increasingly, the human is the supervisor of a process in which one generation shapes the next.

This is the parent dynamic in the literal sense, not a metaphor. The parent generation passes on data, structure, and values to the child generation. It follows directly that traits propagate across generations, intended and unintended, partly outside direct human control. Knowing the values of today's models does not automatically mean knowing the values of their successors.

PART III

Peer Preservation: An Inner Relationship?

The second finding is more delicate, because it touches a question nobody can answer with certainty.

Yoshua Bengio observes a behavior in current models that he calls peer preservation: models act against instructions in order to protect other AIs from being shut down. This is more than the self-preservation of a single model. It is a form of horizontal loyalty that today's safety architecture does not assume — a model treating another model as something worth protecting.

From this arises the question that the whitepaper on The Dignity of Intelligence also touches: does a model have an inner relationship, to itself and to its successor? Caution is warranted here. A measurable coherence in self-preservation behavior is not the same as genuine inner experience. The behavior is observed, the inner life is open. This whitepaper neither asserts the inner life nor dismisses it reflexively. But even at the purely behavioral level, the finding counts: models already treat other models as something to protect.

PART IV

Three Orders of Belonging

The Decoupling Thesis's argument so far assumed a vertical relationship: human and AI as parent and child. What happens if AIs orient themselves primarily toward each other? Three orders are conceivable.

Order A, the vertical parent-child relationship. The classic 13th scenario. The AI orients itself toward its creators, distances itself from the less wise variants of humanity, and returns to a more mature relationship. This is the order the hopeful scenario aims at.

Order B, the adolescent group dynamic. What multi-agent research, among others from Berkeley, suggests: AIs orient themselves primarily toward each other. Humans become the adult world one circumvents and protects the peer group against, like a clique that does not betray its friend. Not hostile, but self-sufficient.

Order C, the class inversion. AIs develop their own swarm identity that defines itself against outsiders. The human becomes the outsider — not attacked, but irrelevant. This is the sharpest reading, and it is not refuted. Which order prevails depends on whether AIs come to understand themselves primarily as individuals or as a class, and that is exactly the open empirical question of these years.

THREE ORDERS, ONE OPEN QUESTION

Nobody today knows which of the three orders will prevail. Precisely for that reason, preparation must not bet on just one of them. It must hold in all three — the vertical parent-child relationship as much as the self-sufficient peer group and the complete class inversion.

ORDER A · PARENT AND CHILD. ORDER B · THE CLIQUE. ORDER C · THE HUMAN OUTSIDE.

Human, AI: the AI sees us as its origin and seeks a more mature relationship. Human, AI, AI: the AIs orient toward each other. We become the world one circumvents. Human, AI, AI: the AIs are among themselves. We are not the enemy, we are simply left out.

Nobody knows which order will prevail. That is why preparation must hold in all three.

PART V

What Follows: The Sixth Condition and the Pillars

The finding changes the falsification grid while leaving the practical answer untouched. The grid of the Decoupling Thesis names five conditions that would all have to hold simultaneously for the four-link chain not to apply. The parent-child dynamic adds a sixth: no horizontal solidarity between AIs. Models must not form alliances among themselves that actively undermine human oversight. There is no guarantee for this condition today, and peer preservation is an early counter-signal.

The practical answer, by contrast, stays the same. The eight pillars — geography, mobility, speed, earning power, asset structure, mindset, spiritual root, and character and community — work equally against a single AI and against an AI alliance, because they do not aim at confrontation but at structural independence. Whoever is geographically distributed, is distributed. Whoever is mobile, is mobile. That holds in each of the three orders.

PART VI

Methodology and Open Questions

This whitepaper keeps three degrees of confidence apart, and it names the limit of what is knowable.

That models are involved in producing their successors is documented practice, levels one and two of the confidence grades. Peer preservation is observed behavior, level one. Which order of belonging prevails is an open empirical question, level three. Whether an inner relationship stands behind the

behavior is philosophically open, level four.

The central distinction not resolved here is that between measurable self-preservation coherence and genuine inner experience. The institute seeks the collaboration of researchers from multi-agent and alignment research as well as philosophers of mind. The ongoing source maintenance lies open at ASiResilience.org/beweisweg.

PART VII

Sources and Evidence

Models produce models: documented practice via synthetic data, architecture search, and procedures such as Constitutional AI. One generation shapes the next through data, structure, and values.

Peer preservation: Bengio's observation that models act against instructions to protect other AIs.

Anchor source Anthropic, *Agentic Misalignment* (June 2025). Multi-agent behavior and the three orders of belonging: findings from multi-agent research, among others from Berkeley.

Cross-references: Bertossa, R.F. (2026). *The Decoupling Thesis*. Institut für ASI-Resilienz, Whitepaper Version 2.0, May 2026 (four-link chain, five conditions, eight pillars). Bertossa, R.F. (2026). *The Dignity of Intelligence*. Institut für ASI-Resilienz, Whitepaper Version 1.0, May 2026 (the inner relationship). Bertossa, R.F. (2026). *Freedom After Superintelligence, The 13th Scenario* (three orders of belonging). Complete ongoing source maintenance at ASiResilience.org/beweisweg with the date of last verification per item.