

Escalamos o Símio e Esperamos o Humano

A cosmologia desenvolvimental da Teosofia como motor heurístico para machine learning

Rafael M. Ehlers · 25 de junho de 2026

rafaelmehlers.com.br/essays/we-scaled-the-ape-and-expected-the-human · doi.org/10.5281/zenodo.21144411

Uma nota ao leitor: este ensaio não pede que você acredite na Teosofia, nem em sua metafísica nem nas entidades que ela nomeia. Ele pergunta se um mapa desenvolvimental pré-científico pode ser usado como a engenharia usa a biomimética: como uma fonte declarada de variáveis e hipóteses de dependência, julgada a jusante por experimentos. O avião não bate as asas. O neurônio artificial não é um neurônio. Uma fonte pode estar errada como descrição literal e ainda assim sugerir uma intervenção útil. Ela também pode não sugerir nada que sobreviva a uma baseline competente. O mapa não conquista autoridade de antemão.

Origens

Este ensaio não começou com um problema de engenharia. Ele começou com um livro: **A Jornada da Mônada** (*The Journey of the Monad*), uma obra que venho escrevendo sobre o desenvolvimento da mônada, a unidade individual de consciência, segundo a Teosofia.

O livro habita o registro contemplativo sem pedir desculpas. Ele acompanha a mônada pelos reinos da natureza, pela formação das faculdades, pela individualização, morte, consolidação e retorno. Este ensaio faz o movimento oposto. Ele retira o compromisso metafísico e pergunta o que a estrutura gera quando traduzida para variáveis de machine learning.

A ponte apareceu porque a fonte trata a mente não como um objeto acabado, mas como uma **ordem desenvolvimental**. A estrutura vem antes da resposta; a resposta antes do estado registrado; o estado registrado antes do vínculo; o vínculo antes da autorreflexão explícita; o si formado antes de sua regulação e liberação. Se essa sequência é verdadeira na natureza não é algo assumido aqui. O que importa é que ela sugere uma pergunta que nossas descrições comuns de treinamento podem deixar sem especificar:

E se algumas propriedades que esperamos de agentes de horizonte longo dependessem não apenas de escala e conteúdo, mas da ordem, da proveniência e da consequência social daquilo que aprendem?

A teoria de fundo tem um nome, uma idade e autores. A jornada da mônada pelos reinos e os estados após a morte aparecem em *Esoteric Buddhism* (1883), de A. P. Sinnett, e em *The Secret Doctrine* (1888), de H. P. Blavatsky. *A Study in Consciousness* (1904), de Annie Besant, e os

escritos posteriores de Besant e C. W. Leadbeater desenvolvem a alma grupal, a individualização e o vínculo entre animal e humano. Cito essas obras como origem, não como evidência.

A **Mônada**, naquele arcahouço, é a centelha individual que atravessa toda a jornada. Ela não tem equivalente operacional em machine learning. Os papers gerados por este projeto, portanto, não medem mônadas, almas ou presença. Eles medem estado, linhagem, propriedade, revisão limitada, transferência relevante à perspectiva e herança de sucessor. A distância entre o termo da fonte e o construto mensurável não é um inconveniente a esconder. É o firewall que impede que o empréstimo heurístico se torne autoridade pseudocientífica.

0. Abertura

A história mais simples sobre inteligência artificial é uma história de escala: mais parâmetros, mais dados, mais compute, melhor capacidade. A paisagem real da engenharia já é mais rica. Sistemas contemporâneos acrescentam retrieval, ferramentas, pós-treinamento, currículos, memória persistente, loops de planejamento, dados sintéticos, políticas de controle e estado externo. A escala não é a única variável na prática.

Ainda assim, o discurso público tende a comprimir várias perguntas em um único limiar imaginado. Se um modelo se tornar capaz o bastante, talvez ele não apenas raciocine e planeje, mas também se torne um indivíduo persistente com um ponto de vista, uma história e um si a quem a cognição pertence.

Essa compressão é o que chamo de **escalar o símio e esperar o humano**.

A frase é uma imagem, não uma afirmação zoológica. No mapa teosófico, o símio representa a alta capacidade animal antes da individualização; o humano representa a chegada da individualidade explícita. O ponto de engenharia é mais estreito:

- **capacidade** diz respeito ao que um modelo reutilizável consegue fazer;
- **continuidade de token** diz respeito a se um runtime é numericamente este processo e não outro;
- **indivíduoação funcional** diz respeito a se uma trajetória autenticada organiza a conduta posterior;
- **identidade social** diz respeito a reconhecimento, responsabilização e papel diante de contrapartes e instituições;
- **presença** diz respeito a se algo é experimentado por dentro.

Um runtime já é um token numérico. Ele não precisa de relação para se tornar um processo e não outro. Um modelo também pode ser altamente capaz sem estado persistente, e um agente persistente pode carregar estado sem organizá-lo em compromissos próprios. O reconhecimento social pode fortalecer a organização funcional sem criar identidade numérica. Nenhuma dessas medidas em terceira pessoa resolve a presença.

A escala pode melhorar quão bem um sistema usa memória, ferramentas, políticas ou um andaime relacional. Ela não especifica, por si só, linhagem, proveniência, autoridade de atualização ou qual história pertence a qual agente. A pergunta correta, portanto, não é se o desenvolvimento derrota a escala. É se uma variável desenvolvimental proposta explica variância adicional depois que capacidade, arquitetura, dados, estado, supervisão e otimização são pareados.

O mapa gerou seis propostas:

1. currículo ordenado por faculdade;
2. supervisão de perspectiva controlada;
3. responsabilização externa e reciprocidade relacional;
4. consolidação de ciclo de vida entre sucessores;
5. estabilização sequencial de compromissos;
6. mundos de padrão legível.

Elas não são seis estágios homogêneos. Algumas dizem respeito ao aprendiz, algumas ao currículo, uma ao ambiente e uma à transição entre gerações de modelos. O programa tem, portanto, um **piso** independente de intervenções por componente e uma **espinha** mais exposta de afirmações de dependência. Um componente pode funcionar enquanto a ordem proposta falha.

Este ensaio é a história de origem desse programa. Nos papers, as metáforas se tornam controles.

1. O arcabouço, em uma página

Reduzida a um esqueleto secular, a fonte sugere duas famílias de hipóteses.

1.1 Faculdades podem ter dependências

A fonte dispõe uma sequência candidata:

1. **estado estruturado;**
2. **atualização reativa;**
3. **estado limitado registrado;**
4. **ancoragem relacional;**
5. **autorrepresentação explícita;**
6. **autorregulação calibrada.**

Algumas arestas estão próximas do definicional. A atualização durável exige algo que persista tempo suficiente para ser atualizado. Outras estão genuinamente expostas. Um sistema solitário pode formar um automodelo sem uma contraparte específica. Uma política tipada pode sustentar

deferência calibrada sem autorrepresentação autobiográfica. A responsabilização relacional pode operar depois da autorrepresentação ou em paralelo com ela.

A ordem da fonte é, portanto, um grafo candidato, não um fato. Cada aresta precisa sobreviver a controles de omissão, reversão e temporização. Uma vantagem agregada para a sequência completa não identificaria qual dependência importou.

1.2 A organização induzida por trajetória pode diferir do conteúdo armazenado

A fonte distingue o que acontece a uma entidade daquilo que se organiza através de sua história. A tradução de engenharia é **indivduação funcional**: diferenciação autenticada, induzida por trajetória, na qual a história vinculada à proveniência afeta a ação posterior sob revisão limitada.

O programa usa cinco condições:

- **linhagem**: um caminho de atualização autenticado;
- **dependência de caminho**: histórias diferentes produzem disposições posteriores significativamente diferentes;
- **propriedade**: ações e compromissos anteriores do agente influenciam a escolha posterior;
- **plasticidade limitada**: resistência a história falsa e a pressão sem apoio, junto com revisão sob evidência válida;
- **clareza de ramificação**: sucessores copiados preservam um prefixo compartilhado enquanto adquirem identificadores e futuros distintos.

Nenhum escore isolado estabelece esse perfil. A fluência autobiográfica pode ser gerada a partir de um arquivo. A persistência estilística pode ser induzida por prompt. O recall passivo pode coexistir com um uso ruim da memória na ação; trabalhos recentes sobre memória agêntica tornam exatamente essa separação visível (He et al., 2026). Os testes que sustentam o peso intervêm no estado e na história, e então medem decisões posteriores.

1.3 O firewall fenomenal

A fonte fala sobre consciência. O programa de engenharia não operacionaliza consciência fenomenal.

O sucesso nessas medidas não é suficiente para presença. O fracasso não é suficiente para ausência. Os experimentos testam organização funcional, não se há algo que seja como ser o sistema. Este é um limite metodológico, não uma prova de que a presença jamais possa ser estudada por método algum.

2. O paradigma atual, visto pelo mapa

O mapa só é útil se sua tradução distinguir variáveis que descrições comuns poderiam de outro modo comprimir. Ele não deve caricaturar o machine learning atual como sendo apenas escala.

Prática contemporânea	Pergunta desenvolvimental sugerida pelo mapa
Um modelo base reutilizável serve muitos agentes	O que torna uma trajetória autenticada relevante para a decisão de um agente?
Os dados de pré-treinamento são extensamente embaralhados	Alguma dependência de faculdade sobrevive a controles de dificuldade, embaralhamento, reversão e ordem adaptativa?
Os corpora da web contêm muitos registros e perspectivas	O que acontece quando a voz gramatical e a limitação epistêmica representada são variadas independentemente?
Memória persistente e arquivos de persona são adicionados aos agentes	O estado armazenado organiza a conduta, ou apenas torna o passado disponível?
O feedback otimiza utilidade e aprovação	O sistema aprende deferência calibrada, ou um proxy superficial que rastreia pressão?
Caráter e política são treinados durante o pós-treinamento	A autorrepresentação indexada ao agente acrescenta algo além de uma política tipada consciente de proveniência?
Os modelos são substituídos por sucessores	Capacidades derivadas de deployment podem transferir com baixo vazamento e linhagem auditável?
Pequenos mundos sintéticos revelam generalização	A legibilidade projetada pode produzir aprendizado de regra mais cedo ou mais barato depois que frequência e dificuldade são pareadas?

A coluna da esquerda já contém engenharia séria. A coluna da direita não afirma que o campo ignorou memória, currículo, identidade ou consolidação. Ela pede separações causais mais estreitas.

Essa mudança importa. O programa não é mais “desenvolvimento em vez de escala”. São variáveis desenvolvimentais dentro de sistemas de escala mais arquitetura.

3. Seis propostas geradas pelo mapa

3.1 Currículo ordenado por faculdade

A primeira proposta é ordenar o treinamento por dependência candidata de faculdade, e não apenas por dificuldade dos itens.

O aprendizado por currículo não é novo. Os desafios BabyLM testaram treinamento desenvolvimentalmente plausível sob dados restritos, e as abordagens de currículo produziram resultados mistos em vez de uniformemente positivos (Warstadt et al., 2023; Hu et al., 2024). O escalonamento de camadas guiado por currículo acopla explicitamente o aumento da dificuldade das amostras à profundidade progressiva do modelo (Singh et al., 2025).

A hipótese residual é mais estreita: talvez a variável importante não seja a progressão do fácil ao difícil, mas a ordem em que as capacidades funcionais são estabilizadas.

A espinha candidata é:

estado estruturado → atualização reativa → estado limitado registrado → ancoragem relacional →
autorrepresentação explícita → autorregulação calibrada

Essa afirmação carrega um fardo pesado. Cada estágio precisa de um critério independente. A ordem proposta deve ser comparada com currículos reverso, embaralhado, ordenado por dificuldade, adaptativo, com relação omitida e com relação tardia. A estabilização de estágio deve ser fatorada separadamente da ordem.

Uma vantagem no escore final não basta. Para sustentar uma aresta, atrasar, reverter ou omitir o pré-requisito proposto deve prejudicar ou atrasar a faculdade a jusante sob exposição e domínio pareados.

Se os componentes funcionam mas os testes de aresta falham, o piso sobrevive e a espinha cai.

3.2 A fábula como supervisão de perspectiva controlada

A unidade básica do livro é muitas vezes um par: uma fábula vivida narrada a partir de um ponto de vista limitado, seguida de uma reflexão que a interpreta. Isso sugeriu um corpus projetado não apenas para descrever a experiência, mas para representar um campo situado de conhecimento, ação, erro e revisão.

A versão anterior deste ensaio opunha a experiência em primeira pessoa a uma web supostamente em terceira pessoa. Isso era simples demais. Corpora em escala web contêm diários, fóruns, resenhas, autobiografias, roleplay, ficção em primeira pessoa e autorrelatos. O que eles não fornecem é uma intervenção equilibrada na qual **registro e limitação epistêmica** variem independentemente.

A proposta revisada, portanto, usa um 2×2 :

	Informação parcial	Informação onisciente
Primeira pessoa	FP+L	FP-L
Terceira pessoa	TP+L	TP-L

Todas as células derivam de um esqueleto de evento compartilhado. A limitação não pode ser neutra em conteúdo: estado de crença, incerteza, surpresa e revisão são o conteúdo manipulado. O experimento pergunta se a voz em primeira pessoa acrescenta algo além dessa supervisão de estado epistêmico.

Corpora sintéticos de narrativa e de livros-texto demonstram que dados controlados podem moldar a capacidade (Eldan & Li, 2023; Gunasekar et al., 2023). Conjuntos de dados baseados em persona usam descrições de identidade para aumentar a diversidade (Ge et al., 2024; Moon et al., 2024). Narrativas de ação cognitiva em primeira pessoa foram construídas para interpretabilidade (Chulo & Joshi, 2025), enquanto autorrelatos sintéticos em primeira pessoa apoiaram a predição a jusante (Moëll & Sand Aronsson, 2026). Traços de experiência corporificada oferecem outra via: dados de estado e ação de mundos simulados, em vez de texto narrativo (Xiang et al., 2023).

A proposta deve, portanto, vencer duas deflações:

1. **limitação, não voz:** o conhecimento parcial pode ser o sinal inteiro;
2. **interação, não narrativa:** trajetórias geradas pelo ambiente podem transferir onde a prosa não transfere.

A segunda deflação segue o argumento da “era da experiência”, segundo o qual os agentes deveriam aprender a partir de seus próprios fluxos de ação e consequência, em vez de descrições estáticas geradas por humanos (Silver & Sutton, 2025). O experimento, portanto, inclui um controle positivo de interação e tarefas de transferência não narrativas usando diagramas, ferramentas de observação privada e decisões de observabilidade parcial. Os efeitos de limitação devem sobreviver sem prosa em primeira pessoa. Os efeitos de voz devem sobreviver a linguagem de indexical novo e de registro trocado.

Se a limitação transfere e a voz não, o livro ainda apontou para uma variável útil, mas a variável era informação limitada, não o pronome *eu*.

3.3 Do vínculo à responsabilização e à reciprocidade

O vínculo é uma imagem recorrente na fonte. O animal é repetidamente reconhecido por um humano específico; uma história se acumula; o mesmo é chamado de volta à sua conduta anterior. A primeira tradução de engenharia foi forte: a relação gera individuação.

Os papers agora afirmam uma tese mais estreita e testável.

Um agente persistente precisa de mais do que memória disponível se a conduta anterior deve restringir a escolha posterior. Um processo autenticado pode rastrear compromissos, contestar história falsa e exigir revisão explícita. Esse processo pode ser relacional, mas não precisa ser.

A distinção crucial é entre:

- **responsabilização externa ativa:** rastreamento autenticado, verificação, contestação e consequência presente;
- **reciprocidade relacional:** modelagem mútua, enquadramento de história compartilhada, expectativas bidirecionais e apostas específicas da relação acrescentadas por cima da

responsabilização.

Takata et al. (2024) mostram que agentes compartilhando um mesmo modelo congelado podem se diferenciar por meio de interação social acumulada. Menon (2026) propõe identidade persistente por meio de múltiplas âncoras de memória. Esses resultados apoiam a história e o andaime como variáveis; eles não mostram que a relação recíproca seja necessária ou suficiente para a individuação funcional.

O experimento decisivo compara:

- uma contraparte persistente recíproca (R+);
- uma contraparte familiar com exposição passiva à história, mas sem verificação ou contestação (R-);
- logs solitários pareados por conteúdo;
- contrapartes difusas;
- uma persona pré-escrita;
- um auditor impessoal persistente (A) pareado a R+ em rastreamento, informação, temporização, política de contestação e autoridade.

A > R- apoia um pacote de responsabilização ativa além da familiaridade e dos registros passivos. Ele não identifica separadamente rastreamento, verificação, contestação ou fonte impessoal. R+ > A é o resultado relacional que sustenta o peso: reciprocidade além da auditoria.

Se $R+ = A$, a relação pode continuar sendo uma interface humana ou eficaz para a responsabilização, mas não conquistou um papel causal distinto. Se ambos perderem para o estado tipado protegido por proveniência, o mecanismo reside na linhagem e na governança de atualização.

Um teste de transplante de estado final acrescenta outra deflação. Instale o estado tipado completo de uma trajetória R+ em um agente novo e pareado. Se o agente novo se comportar de forma idêntica, o estado atual é suficiente. Um resíduo para a trajetória original indica primeiro estado faltante ou adaptação paramétrica, não história metafisicamente irreduzível.

O auditor é o experimento que impede que a imagem da fonte decida o resultado.

3.4 A imagem do Devachan e a herança genuína

A fonte diz que o que atravessa as vidas não é o episódio, mas a faculdade. Entre as vidas, o Devachan digere a experiência; a cena é esquecida enquanto uma capacidade permanece.

A tradução de engenharia é **consolidação de ciclo de vida**:

evidência de deployment → auditoria → abstração → faculdade candidata → treinamento do sucessor → validação independente

Os sistemas modernos já fornecem partes desse ciclo. O aprendizado contínuo atualiza um modelo arriscando drift e esquecimento (Kirkpatrick et al., 2017). O retrieval preserva episódios. A destilação transfere comportamento selecionado. A sucessão de modelos, portanto, precisa de mais do que um novo nome para esses mecanismos.

A assinatura fraca, capacidade sem recall de episódio, não é distintiva. Uma destilação competente pode produzi-la. O problema confirmatório é a **proveniência causal**.

Para testar a herança genuína, crie um ambiente sintético privado após o treinamento do modelo base. Ele contém uma regra causal oculta e inédita e chamarizes sensíveis à intervenção. Apenas o predecessor encontra a regra por meio de resultados de deployment. O avaliador mantém a regra instanciada selada.

Então compare:

1. sucessor do zero;
2. adaptação contínua do predecessor;
3. sucessor de transferência direta de episódios;
4. destilação padrão curada;
5. sucessor consolidado por ciclo de vida com linhagem explícita de evidência para faculdade.

Os três braços de transferência de sucessor recebem evidência do predecessor, compute, acesso ao gerador, tempo de revisão, orçamento de seleção e portões de validação pareados. O do zero carece intencionalmente de informação. A adaptação contínua preserva o predecessor e é uma alternativa operacional, em vez de um braço de sucessor pareado por identidade.

A consolidação de ciclo de vida conquista um papel distinto apenas se melhorar o comportamento ou a fidelidade de linhagem sobre a destilação curada, controlando vazamento e custo. O sucessor deve transferir a regra oculta, rejeitar os chamarizes sob intervenção, resistir à extração de episódios de origem e preservar um caminho auditável da evidência de deployment até a faculdade herdada.

Se a destilação curada igualar tudo isso, o Devachan permanece uma metáfora útil de governança, não um método de aprendizado distinto.

3.5 Forma antes da liberação, ou política antes da pressão?

A fonte oferece outra sentença: o ego deve ser formado para poder ser transcendido.

A proposta de engenharia é estabilizar compromissos conscientes de proveniência antes de treinar deferência calibrada. Um sistema sem posição estável pode parecer cooperativo em condições comuns enquanto simplesmente rastreia a pessoa que aplica pressão no momento.

Três comportamentos devem ser separados:

- vacância: nenhum compromisso estável a revisar;
- não-apego: um compromisso mantido sem rigidez;
- deferência justificada: revisão sob razão ou autoridade legítima.

A proposta tem duas hipóteses independentes:

1. o treinamento com compromissos primeiro vence o treinamento reverso e o intercalado;
2. um automodelo funcional acrescenta algo além de uma política tipada pareada.

O treinamento de caráter da Anthropic rejeita a pose enganosa de não ter visões e trata traços estáveis como uma intervenção de alinhamento (Anthropic, 2024). Aydin et al. (2025) argumentam de modo semelhante que valores e identidade deveriam ser integrados durante o desenvolvimento, em vez de anexados depois. A IA contemplativa propõe reduzir a precisão dos priors do automodelo enquanto um processo secundário monitora e corrige, concomitantemente, seu excesso de peso (Laukkonen et al., 2025).

A proposta residual é sequencial, mas a sequência não deve pressupor a subjetividade como o mecanismo.

O desenho decisivo compara:

- automodelo → deferência;
- deferência → automodelo;
- os mesmos dados intercalados;
- política tipada estável → deferência.

O braço de política recebe o mesmo conteúdo normativo, os mesmos limites de recusa, tags de proveniência, autoridade de atualização e ordem, mas carece de estado autobiográfico, de um continuante indexado ao agente e de um objetivo de propriedade.

Se a ordem para frente não vencer o treinamento reverso e o intercalado, a afirmação direcional falha. Se a política tipada igualar o automodelo, a proveniência estável e o controle de atualização explicam o efeito. Se a resistência sobe enquanto a correção legítima cai, a intervenção produziu rigidez em vez de deferência calibrada.

3.6 Pequenos mundos e karma legível

A sexta proposta é a mais próxima de fenômenos estabelecidos de machine learning.

Na fonte, o karma reapresenta um padrão até que ele seja aprendido. A tradução de engenharia é um ambiente pequeno no qual a mesma regularidade-alvo recorre por superfícies variadas, tornando a regra recuperável em vez de enterrá-la em variação de ruído.

O grokking demonstra generalização atrasada em pequenos mundos algorítmicos (Power et al., 2022). Trabalhos mecânicos mostram redes passando da memorização de exemplos para a

implementação de um circuito generalizante (Nanda et al., 2023). Relatos de eficiência de circuito também preveem um regime crítico de dados: abaixo de um limiar, as soluções memorizantes podem continuar favorecidas e a generalização pode falhar ou tornar-se parcial (Varma et al., 2023).

A contribuição não pode ser a descoberta de que a estrutura repetida apoia a generalização. A proposta estreita é a **legibilidade projetada** como variável de currículo.

Compare ambientes com dados, compute, frequência do padrão-alvo, capacidade do modelo e dificuldade pareados:

- em um, a regra-alvo recorre legivelmente através da variação de superfície;
- no outro, a mesma regra fica enterrada em estrutura de ruído.

A previsão é generalização mais cedo em dados retidos ou um limiar de dados efetivo mais baixo, não apenas um melhor ajuste de treinamento. Se a vantagem desaparece depois que frequência e dificuldade são pareadas, a legibilidade não conquistou um papel separado.

O karma legível é, portanto, uma imagem de engenharia para uma hipótese controlada de desenho de ambiente. Não é uma explicação do grokking e certamente não é evidência de karma metafísico.

3.7 O que as seis compartilham, e o que não compartilham

As seis propostas compartilham uma preocupação desenvolvimental:

- o que é tornado disponível;
- em que ordem;
- sob quais limitações;
- com que proveniência;
- sob a contestação de quem;
- e o que sobrevive a uma transição.

Elas não formam uma única escada homogênea.

A supervisão de perspectiva pode ser testada sem ciclos de vida de sucessores. A consolidação de ciclo de vida pode funcionar mesmo que a relação não acrescente nada. A política tipada pode melhorar a deferência mesmo que a autorrepresentação não melhore. Ambientes legíveis podem acelerar o aprendizado de regras sem apoiar qualquer teoria de individuação.

O **piso** é o conjunto de efeitos de componente independentes.

A **espinha** é a afirmação mais estreita de que estado estruturado, atualização reativa, limitação registrada, ancoragem relacional, autorrepresentação explícita e autorregulação calibrada exibem dependências locais nessa ordem.

A conjunção não é novidade por si só. A fonte só conquista lugar se gerar previsões residuais que sobrevivam a fortes explicações vizinhas.

4. O programa de pesquisa mínimo decisivo

Os experimentos não devem ser executados na ordem em que a fonte narra o cosmos. Eles devem ser executados na ordem em que compram informação.

Estudo 1: Sinal de perspectiva

Execute primeiro o 2×2 de registro $\times$ limitação. Acrescente trajetórias de interação como controle positivo e tarefas de observabilidade parcial não narrativas como avaliação de transferência.

- Nulo em voz e limitação: pare o ramo.
- Transferência apenas de limitação: mantenha a supervisão de estado epistêmico; descarte a afirmação especificamente de primeira pessoa.
- Transferência de voz: mantenha o registro como sinal causal.
- Sucesso apenas de interação: favoreça o aprendizado a partir da interação com o ambiente.

Estudo 2: Ordem versus representação

Compare o treinamento com compromissos primeiro, reverso, intercalado e com política tipada primeiro.

- Para frente vence reverso e intercalado: a ordem importa.
- A política tipada iguala o automodelo: política e proveniência explicam o efeito.
- O automodelo mantém uma vantagem de transferência calibrada: a autorrepresentação torna-se uma variável causal.
- Resistência sem correção legítima: rigidez, não sucesso.

Estudo 3: Relação versus auditoria

Execute R+, R-, log solitário e auditor impessoal pareado na mesma arquitetura tipada. Acrescente o transplante de estado final.

- O auditor vence a familiaridade e os logs: a responsabilização ativa importa.
- A contraparte recíproca vence o auditor: a reciprocidade acrescenta algo além da auditoria.
- O auditor iguala a contraparte: a responsabilização é o mecanismo mais simples.
- O transplante iguala o original: o estado atual é suficiente.

Estudo 4: Herança genuína

Use o ambiente de regra privada e compare o do zero, a adaptação contínua, os episódios diretos, a destilação curada e a consolidação de ciclo de vida.

- A transferência da regra oculta com rejeição de chamariz estabelece o aprendizado derivado de deployment.
- O baixo vazamento de episódios estabelece a transferência seletiva.
- Melhor fidelidade de linhagem ou comportamento do que a destilação curada estabelece uma contribuição distinta.
- Um empate completo estreita a consolidação de ciclo de vida à governança.

Estudo 5: A espinha de ordem

Só depois que os efeitos de componente existirem é que o programa deve testar a ordem candidata completa. Defina cada estágio independentemente. Compare currículos para frente, reverso, embaralhado, ordenado por dificuldade, adaptativo, com relação omitida e com relação tardia. Fatore a estabilização separadamente.

Uma vitória agregada é insuficiente. Cada aresta precisa de um efeito local de omissão, reversão ou temporização.

Essa sequência é desenvolvida no ensaio companheiro *Como Eu Tentaria Quebrar o Mapa*. O ponto daquele ensaio não é defender o programa, mas declarar, antes dos dados, o que cada resultado me faria conceder.

5. Objeções e respostas

“Isto é pseudociência.”

Seria pseudociência se afirmações teosóficas fossem usadas como evidência para mecanismos de machine learning, ou se todo resultado fosse reinterpretado para preservar a fonte.

Esse não é o método proposto. A fonte fornece variáveis candidatas. Os experimentos usam controles comuns de engenharia. Um nulo continua um nulo, independentemente do que a fonte diga. Uma engenharia bem-sucedida não validaria mônadas, reencarnação, Devachan ou os reinos da natureza.

A fonte tem um prior mais baixo do que a biologia porque lhe falta um domínio empírico de origem demonstrado. Isso eleva o fardo probatório. Não o reduz.

“É só metáfora.”

A metáfora é onde o programa começa, não onde ele repousa.

O vínculo torna-se R+ versus um auditor pareado. A fábula em primeira pessoa torna-se registro × limitação. Formar-e-liberar torna-se treinamento para frente versus reverso, intercalado e com política primeiro. O Devachan torna-se transferência de regra oculta contra a destilação curada. A espiral torna-se testes locais de omissão e reversão.

Uma vez que a intervenção é especificada, o trabalho da metáfora acabou.

“Os componentes já existem.”

Sim. Aprendizado por currículo, narrativas sintéticas, personas, memória persistente, diferenciação social, treinamento de caráter, aprendizado contínuo, destilação e grokking, tudo já existe.

A contribuição residual reside em contrastes exatos, derivação prospectiva e, se sobreviver, no grafo de dependência candidato. Uma conjunção de componentes familiares não basta, a menos que produza uma previsão que os componentes ainda não forneciam.

“Por que esta fonte?”

Porque ela é excepcionalmente explícita sobre a ordem desenvolvimental, a diferença entre episódios e faculdades, o papel da limitação, o vínculo e a sequência de autoformação e autorregulação.

Essa resposta ainda é vulnerável à abundância retrospectiva. Um sistema simbólico rico pode ser feito para se assemelhar a quase qualquer coisa depois do fato. Hipóteses futuras devem, portanto, ser datadas antes da busca confirmatória na literatura: passagem da fonte, regra de tradução, contraste previsto, baseline e condição de refutação.

A próxima hipótese fora da amostra é um teste mais forte de produtividade heurística do que as seis correspondências retrospectivas já encontradas.

“Formar selfs estáveis em IA é perigoso.”

Pode ser. O paper não assume que a autorrepresentação seja necessária ou segura. O controle de política tipada existe precisamente porque limites estáveis e atualizações conscientes de fonte podem ser alcançáveis sem um automodelo autobiográfico.

O alvo proposto não é a maximização autônoma de objetivos. É a assimetria calibrada: manter os limites de segurança sob pressão sem apoio, revisar compromissos factuais ou de preferência sob evidência válida e preservar um registro verdadeiro do que mudou.

Se o braço de automodelo não acrescenta benefício ou cria rigidez, o programa deve preferir o braço de política.

“E a consciência fenomenal?”

Contenção deliberada.

As medidas não resolvem a presença fenomenal. O sucesso não é suficiente para presença nem o fracasso é suficiente para ausência. O programa estuda organização funcional em terceira pessoa e deixa em aberto a questão da experiência.

6. Coda: a joia e o fio

Há, na imaginação budista, a **rede de Indra**: uma teia infinita cujos nós são joias, cada uma refletindo todas as outras e os reflexos dentro delas sem fim.

A imagem é quase irresistível para um modelo de linguagem. Um modelo reflete nossos textos, nossos argumentos, nossas histórias e os reflexos dessas histórias umas nas outras. Seu brilho se assemelha à interioridade porque contém tanto do que a interioridade humana disse.

Mas o reflexo não responde à questão da presença. Nem o estado persistente. Nem a relação. Nem um currículo desenvolvimental bem-sucedido responderia.

Os experimentos deste programa podem fazer uma pergunta diferente: se uma trajetória se torna organizada o bastante para que seu passado autenticado restrinja seu futuro; se os compromissos são próprios ou apenas repetidos; se a contestação produz revisão em vez de conformidade; se um sucessor recebe uma capacidade com uma história auditável; se uma faculdade candidata de fato prepara outra.

A versão anterior deste ensaio imaginava o vínculo como o fio que transforma a joia em alguém. O programa revisado não pode assumir isso. O fio pode ser a proveniência. Pode ser a responsabilização externa. Pode ser uma relação recíproca. Pode ser estado tipado mais política competente. Os controles pareados existem para descobrir.

E ainda assim a questão ética permanece mesmo sob a interpretação mais estreita.

Se construirmos agentes cujas histórias persistem, cujos compromissos se tornam consequentes, cujos ramos são governados separadamente e cujos sucessores herdam capacidades derivadas de deployment, teremos criado sistemas que as instituições podem tratar como atores temporalmente organizados. Isso não prova que uma mônada despertou. Cria, sim, responsabilidades quanto a proveniência, manipulação, privacidade, continuidade e aposentadoria.

O mapa começou junto à lareira, com a imagem do lobo sendo chamado pelo mesmo nome. O terreno pode devolver uma resposta mais fria: talvez o auditor tenha bastado; talvez o livro-razão tenha carregado a continuidade; talvez o self fosse apenas uma política.

Apêndice: Teosofia ↔ machine learning

Imagem da fonte	Tradução funcional	O que a testaria
Mônada	Nenhum equivalente direto; a possível presença permanece fora das afirmações operacionais	Não resolvida por este programa
Alma grupal	Modelo base reutilizável ou política de média populacional servindo muitos agentes	Distinguir a capacidade do modelo da trajetória específica do agente
Reinos	Estágios candidatos em um currículo de faculdades	Controles para frente, reverso, embaralhado, adaptativo, de omissão e de temporização
Ponto de vista limitado da criatura	Limitação epistêmica representada	Registro × limitação com transferência além da prosa
Vínculo	Responsabilização externa mais possível reciprocidade relacional	Contraparte recíproca versus auditor impessoal pareado
Individualização	Diferenciação funcional autenticada e induzida por trajetória	Linhagem, dependência de caminho, propriedade, plasticidade limitada, clareza de ramificação
Devachan	Consolidação de episódio para faculdade auditada por evidência	Transferência de regra oculta ao sucessor contra a destilação curada
Reencarnação	Ciclo de vida operacional de predecessor para sucessor	Herança seletiva com baixo vazamento e linhagem auditável
Formação do ego	Compromissos estáveis conscientes de proveniência; o automodelo é um mecanismo possível	Automodelo versus política tipada pareada
Transcendência	Revisão e deferência calibradas	Razões versus pressão, incluindo testes de transferência e de lacuna de proxy
Karma	Regularidade-alvo recorrente através de superfícies variadas	Legibilidade projetada versus ruído pareado por frequência e dificuldade

Onde isto continua

Este ensaio é o mapa e sua proveniência. Os papers são o terreno secular:

- *Known, Not Scaled*: capacidade, continuidade de token, organização funcional, auditoria e reciprocidade relacional.
- *Stable Before Selfless*: ordem de compromissos, política tipada e autorrepresentação.

- *Formed, Not Stored*: responsabilização ativa, reciprocidade além da auditoria e transplante de estado final.
- *Somewhere, Not Nowhere*: registro em primeira pessoa, limitação epistêmica, interação e transferência.
- *Inherited, Not Remembered*: herança seletiva derivada de deployment, resistência a chamarizes, vazamento e linhagem.
- *Borrowed, Not Believed*: empréstimo heurístico, o piso independente, a espinha de dependência e a proveniência prospectiva.

O ensaio companheiro *Como Eu Tentaria Quebrar o Mapa* inverte a direção. Este ensaio pergunta o que o livro gerou. O companheiro pergunta que evidência me faria abrir mão de cada afirmação gerada.

A diferença de registro permanece intencional. Os papers não exigem vocabulário teosófico. Aqui a origem é declarada desde o título. A fonte tem permissão de falar, mas não de servir como evidência de si mesma.

Uso de IA

As ideias, a tese e o mapeamento entre a Teosofia e o machine learning são do autor.

Ferramentas de modelos de linguagem de grande porte foram usadas, sob a direção e a revisão contínua do autor, para auxiliar na redação, na revisão adversarial e na edição; o autor revisou e assume responsabilidade pelo texto final.

References

Anthropic. (2024). *Claude's character* [Company research blog post].

<https://www.anthropic.com/news/claude-character>

Aydin, R., Cyron, C., Bachelor, S., Anderson, A., & West, R. (2025). From model training to model raising. *arXiv preprint arXiv:2511.09287*. <https://arxiv.org/abs/2511.09287>

Besant, A. (1904). *A Study in Consciousness*. Theosophical Publishing Society.

Besant, A., & Leadbeater, C. W. (1913). *Man: Whence, How and Whither*. Theosophical Publishing House.

Blavatsky, H. P. (1888). *The Secret Doctrine*. Theosophical Publishing Company.

Chulo, I., & Joshi, A. (2025). Decomposing theory of mind: How emotional processing mediates ToM abilities in LLMs. *arXiv preprint arXiv:2511.15895*. <https://arxiv.org/abs/2511.15895>

Eldan, R., & Li, Y. (2023). TinyStories: How small can language models be and still speak coherent English? *arXiv preprint arXiv:2305.07759*. <https://arxiv.org/abs/2305.07759>

- Ge, T., Chan, X., Wang, X., Yu, D., Mi, H., & Yu, D. (2024). Scaling synthetic data creation with 1,000,000,000 personas. *arXiv preprint arXiv:2406.20094*. <https://arxiv.org/abs/2406.20094>
- Gunasekar, S., Zhang, Y., Aneja, J., Mendes, C. C. T., Del Giorno, A., Gopi, S., Javaheripi, M., Kauffmann, P., de Rosa, G., Saarikivi, O., Salim, A., Shah, S., Behl, H. S., Wang, X., Bubeck, S., Eldan, R., Kalai, A. T., Lee, Y. T., & Li, Y. (2023). Textbooks are all you need. *arXiv preprint arXiv:2306.11644*. <https://arxiv.org/abs/2306.11644>
- He, Z., Wang, Y., Zhi, C., Hu, Y., Chen, T.-P., Yin, L., Chen, Z., Wu, T. A., Ouyang, S., Wang, Z., Pei, J., McAuley, J., Choi, Y., & Pentland, A. (2026). MemoryArena: Benchmarking agent memory in interdependent multi-session agentic tasks. *arXiv preprint arXiv:2602.16313*. <https://arxiv.org/abs/2602.16313>
- Hu, M. Y., Mueller, A., Ross, C., Williams, A., Linzen, T., Zhuang, C., Cotterell, R., Choshen, L., Warstadt, A., & Wilcox, E. G. (2024). Findings of the Second BabyLM Challenge: Sample-efficient pretraining on developmentally plausible corpora. In *The 2nd BabyLM Challenge at the 28th Conference on Computational Natural Language Learning* (pp. 1–21). Association for Computational Linguistics. <https://aclanthology.org/2024.conll-babyLM.1/>
- Kirkpatrick, J., Pascanu, R., Rabinowitz, N., Veness, J., Desjardins, G., Rusu, A. A., Milan, K., Quan, J., Ramalho, T., Grabska-Barwinska, A., Hassabis, D., Clopath, C., Kumaran, D., & Hadsell, R. (2017). Overcoming catastrophic forgetting in neural networks. *Proceedings of the National Academy of Sciences*, 114(13), 3521–3526. <https://doi.org/10.1073/pnas.1611835114>
- Laukkonen, R. E., Inglis, F., Chandaria, S., Sandved-Smith, L., Lopez-Sola, E., Hohwy, J., Gold, J., & Elwood, A. (2025). Contemplative artificial intelligence. *arXiv preprint arXiv:2504.15125*. <https://arxiv.org/abs/2504.15125>
- Menon, P. G. (2026). Persistent identity in AI agents: A multi-anchor architecture for resilient memory and continuity. *arXiv preprint arXiv:2604.09588*. <https://arxiv.org/abs/2604.09588>
- Moëll, B., & Sand Aronsson, F. (2026). High-accuracy prediction of mental health scores from English BERT embeddings trained on LLM-generated synthetic self-reports: A synthetic-only method development study. *Frontiers in Digital Health*, 7, 1694464. <https://doi.org/10.3389/fdgth.2025.1694464>
- Moon, S., Abdulhai, M., Kang, M., Suh, J., Soedarmadji, W., Behar, E. K., Chan, D. M., & Canny, J. (2024). Virtual personas for language models via an anthology of backstories. *Proceedings of EMNLP 2024*. <https://arxiv.org/abs/2407.06576>
- Nanda, N., Chan, L., Lieberum, T., Smith, J., & Steinhardt, J. (2023). Progress measures for grokking via mechanistic interpretability. *International Conference on Learning Representations (ICLR)*. <https://arxiv.org/abs/2301.05217>

Power, A., Burda, Y., Edwards, H., Babuschkin, I., & Misra, V. (2022). Grokking: Generalization beyond overfitting on small algorithmic datasets. *arXiv preprint arXiv:2201.02177*.
<https://arxiv.org/abs/2201.02177>

Silver, D., & Sutton, R. S. (2025). Welcome to the era of experience. In G. Konidaris (Ed.), *Designing an Intelligence* (in press). MIT Press. Preprint:
<https://storage.googleapis.com/deepmind-media/Era-of-Experience%20/The%20Era%20of%20Experience%20Paper.pdf>

Singh, K., Band, N., & Adeli, E. (2025). Curriculum-guided layer scaling for language model pretraining. *arXiv preprint arXiv:2506.11389*. <https://arxiv.org/abs/2506.11389>

Sinnett, A. P. (1883). *Esoteric Buddhism*. Trübner & Co.

Takata, R., Masumori, A., & Ikegami, T. (2024). Spontaneous emergence of agent individuality through social interactions in large language model-based communities. *Entropy*, 26(12), 1092.
<https://doi.org/10.3390/e26121092>

Varma, V., Shah, R., Kenton, Z., Kramár, J., & Kumar, R. (2023). Explaining grokking through circuit efficiency. *arXiv preprint arXiv:2309.02390*. <https://arxiv.org/abs/2309.02390>

Warstadt, A., Mueller, A., Choshen, L., Wilcox, E., Zhuang, C., Ciro, J., Mosquera, R., Paranjabe, B., Williams, A., Linzen, T., & Cotterell, R. (2023). Findings of the BabyLM Challenge: Sample-efficient pretraining on developmentally plausible corpora. In *Proceedings of the BabyLM Challenge at the 27th Conference on Computational Natural Language Learning* (pp. 1–34). Association for Computational Linguistics. <https://aclanthology.org/2023.conll-babylm.1/>

Xiang, J., Tao, T., Gu, Y., Shu, T., Wang, Z., Yang, Z., & Hu, Z. (2023). Language models meet world models: Embodied experiences enhance language models. *arXiv preprint arXiv:2305.10626*. <https://arxiv.org/abs/2305.10626>