

Nunca Saímos do Treinamento

A infância como pre-training, a adolescência como poda e a vida adulta como aprendizado contínuo

Rafael M. Ehlers · 25 de agosto de 2026

rafaelmehlers.com.br/essays/we-never-leave-training

Existe uma ironia na história da inteligência artificial. Construímos redes neurais inspiradas, de maneira radicalmente simplificada, no cérebro. Agora usamos essas mesmas redes como metáforas para compreender o próprio cérebro.

Uma das analogias mais férteis é pensar no desenvolvimento humano como o treinamento de um modelo.

Durante a infância e a adolescência, o cérebro passa por um período extraordinariamente intenso de formação e reorganização. Novas conexões surgem, algumas se fortalecem e outras desaparecem. Experiências repetidas alteram circuitos. Comportamentos que produzem determinados resultados tendem a ser reforçados, enquanto padrões pouco utilizados podem enfraquecer. É difícil ignorar a semelhança com o treinamento de uma rede neural.

Uma rede artificial começa com uma arquitetura capaz de aprender, mas ainda não sabe representar bem o mundo. Ela é exposta a enormes quantidades de informação, e cada exemplo modifica ligeiramente seus parâmetros. Aos poucos, certas representações se tornam mais estáveis. O sistema deixa de apenas reagir a exemplos isolados e começa a reconhecer padrões.

Algo parecido acontece conosco.

A infância poderia ser comparada, com todas as limitações da metáfora, a uma espécie de pre-training biológico. Muito antes de recebermos instruções explícitas sobre a maior parte das coisas, absorvemos continuamente as regularidades do ambiente: rostos, sons, linguagem, relações espaciais, comportamento social, causa e efeito. Grande parte desse aprendizado não exige que alguém nos diga a resposta correta. O cérebro observa o mundo e constrói modelos internos que ajudam a prevê-lo.

Nesse aspecto, o aprendizado humano inicial se parece com o que hoje chamamos de self-supervised learning.

O reinforcement learning oferece outro paralelo. Certos comportamentos produzem recompensa. Outros levam a desconforto, desaprovação ou ausência de recompensa. Aos poucos, o organismo aprende tanto uma representação do mundo quanto políticas de ação dentro dele. A dopamina participa de alguns desses mecanismos de recompensa e erro de previsão, embora reduzir todo o comportamento humano a reinforcement learning seria uma simplificação grosseira.

A adolescência torna a analogia mais interessante.

Nesse período, o cérebro acumula novas conexões enquanto passa por uma intensa reorganização e poda sináptica. Conexões pouco utilizadas podem desaparecer, enquanto circuitos frequentemente ativados se tornam mais eficientes.

Para quem conhece machine learning, é difícil resistir à comparação com regularização, pruning ou otimização de uma rede.

O cérebro continua aprendendo a partir de uma arquitetura funcional que se tornou muito mais estável. Isso sugere outra extensão da analogia: a passagem do pre-training para o post-training.

Depois do treinamento inicial, um grande modelo de linguagem já contém uma vasta estrutura de representações. O post-training ajusta comportamentos, preferências e respostas sobre essa base.

A vida adulta talvez funcione de maneira semelhante.

Continuamos aprendendo, claro. O cérebro permanece plástico durante toda a vida. Novas experiências podem alterar crenças, hábitos, habilidades e até estruturas neurais. Mas esse aprendizado acontece dentro de uma rede que já carrega décadas de treinamento anterior.

Cada experiência nova encontra pesos moldados por tudo o que veio antes.

Nossa educação, nossas famílias, os ambientes em que crescemos, recompensas, traumas, sucessos, fracassos e relações formam o estado a partir do qual interpretamos cada nova situação.

Isso também ajuda a explicar por que mudar certos padrões na vida adulta pode ser tão difícil. Uma informação nova precisa disputar espaço com representações reforçadas milhares ou milhões de vezes.

A comparação fica mais nítida quando olhamos para a memória.

Em machine learning existe um problema conhecido como catastrophic forgetting. Quando um modelo continua sendo treinado com novos dados, ele pode aprender o novo enquanto destrói parte do que havia aprendido. Pesquisadores de continual learning ou lifelong learning estudam como construir sistemas capazes de continuar aprendendo sem perder o conhecimento anterior.

O cérebro humano parece recorrer a mecanismos que as redes artificiais ainda não conseguiram reproduzir. Dormimos, consolidamos memórias, esquecemos seletivamente, reinterpretamos experiências antigas e continuamos atualizando nosso modelo interno do mundo.