

The evolution of sex for artificial intelligence

A population-genetic framework for multigenerational model populations

Giorgio F. Gilestro

Department of Life Sciences, Imperial College London
giorgio@gilest.ro · lab.gilest.ro

Draft

Significance statement

Artificial intelligence increasingly consists of populations of models. Models are fine-tuned from common ancestors, trained on data that earlier models generated, and combined by weight merging. These practices couple model generations the way reproduction couples biological generations, and they raise the same question: how does a population keep and accumulate abilities over time? I transfer the population genetics of sexual reproduction to this setting and test it in simulations, small neural networks, and language models. The framework recasts continual learning at the population scale and yields design rules: how much real data retraining needs, when to combine models, when to keep them separate, when to stop combining them, and how to anticipate a failed combination before making it.

Abstract

AI development increasingly resembles a population process. Models are specialised, retrained on model output, and recombined by weight merging, in evolutionary vocabulary with little evolutionary theory. I treat multigenerational model populations as systems whose inheritance, diversity, and compatibility must be managed, and transfer to them the population genetics of sexual reproduction. That training on model output is genetic drift, with model collapse its signature, is established; here I develop what follows. A minimal inheritance model is exactly Wright–Fisher, and trained networks depart from it by a measurable, architecture-specific bias. In this model grounding is immigration: a real-data fraction far below one retained most equilibrium diversity, and protecting a rare capability costs the inverse of its frequency. Refitting a child to the average of its parents’ outputs cancels the gain of having several parents, to first order for rare items, whereas operators that keep each parent’s strongest contribution realise it. Merged language-model specialists exceeded every parent across seeds. In a six-generation language-model population, lineages obliged to merge collapsed once partners stopped knowing different things; lineages allowed to refuse a merge, or made to stop after three generations, finished level with never merging, with or without selection between lineages, and merging with one’s own ancestor was safer than merging with a contemporary. Blind recombination fails on rugged task landscapes; screening candidate offspring restores the gain. I introduce model speciation: the merge barrier remaining after permutation-and-rescaling alignment tracks functional conflict, isolation did not emerge from specialisation alone, and pre-merge functional disagreement predicted merge damage where weight geometry did not.

Introduction

Machine learning has become a population-scale phenomenon. Public repositories host millions of models (Hugging Face passed three million by 2026), most of them fine-tunes, distillations, or merges of a few foundation models, forming family trees already mapped by phylogenetic methods (1–3). *Model merging*, the combination of trained parents into a new model by averaging their weights, is mainstream practice with standard tooling and thousands of hybrid checkpoints, some topping leaderboards (4–7), and its literature already speaks of “crossover,” “mutation,” and “mate choice” in populations of merging models that climb benchmarks (5, 8–10) and stagnate as their members grow alike (11).

Generations are coupled through data as well as weights. Models increasingly learn from model output: frontier alignment pipelines are predominantly synthetic (over 98% in documented cases; 12, 13), self-generated instruction data seeds whole lineages (14), much of the public web is machine-generated or machine-translated (15, 16), and the stock of human text is projected to run out within this decade (17). Multi-agent systems and agent economies put many models into sustained contact (18–21). A population whose members inherit from one another, recombine, and retransmit is an evolving population in the technical sense, and I transfer to it the branch of biology built for that situation, the population genetics of the evolution of sex (a transfer anticipated by the reading of sex as an algorithm for mixability; 22).

Training each generation on the previous generation’s output degrades it (*model collapse*). Rare capabilities vanish first and the lineage drifts toward its own most common behaviour (23). That degradation is *genetic drift*, the loss of rare variants in any finite population when each generation is a finite sample of the last (the accident by which rare surnames vanish from small villages, with nothing selecting against them). The identification has been made repeatedly and independently, for sequential inference chains before deep learning (24), for language-model text ecosystems (25), as a first-extinction law (26), and in quantitative-genetic form for self-consuming diffusion models (27). Drift is only the entry point, because population genetics is above all a theory of what keeps a finite population from decaying (immigration, recombination, selection, population structure) and of where each of those fails, and every one of them has a counterpart that the operator of a model population can switch on: real data entering each generation, merging, selection against a verifier, and the choice of which models merge with which.

An operator of a model population faces recurring decisions with no principled guidance. How much verified real data does retraining need? Will combining two models compose their abilities or damage them? Can incompatibility be detected before a failed merge is paid for? When should specialists be kept separate? These are machine learning’s oldest problem, *continual learning* (acquiring new abilities without losing old ones; 28, 29), transposed from a single network to a population whose members inherit from one another, and each has a population-genetic answer with a number attached (how many real samples per generation, how far the average sits below the best parent, how much the parents disagree on shared inputs). Table 1 gives the correspondences the argument runs on. Fig. 1A maps the programme across three tiers (an inheritance model in simulation, trained neural networks, language models). Fig. 1B draws the change of viewpoint the transfer rests on. Models are usually pictured as a society in space, contemporaries exchanging messages, but the couplings that matter here (training on model output, merging, real data entering each generation) run between generations, and a society coupled in time is what population genetics describes.

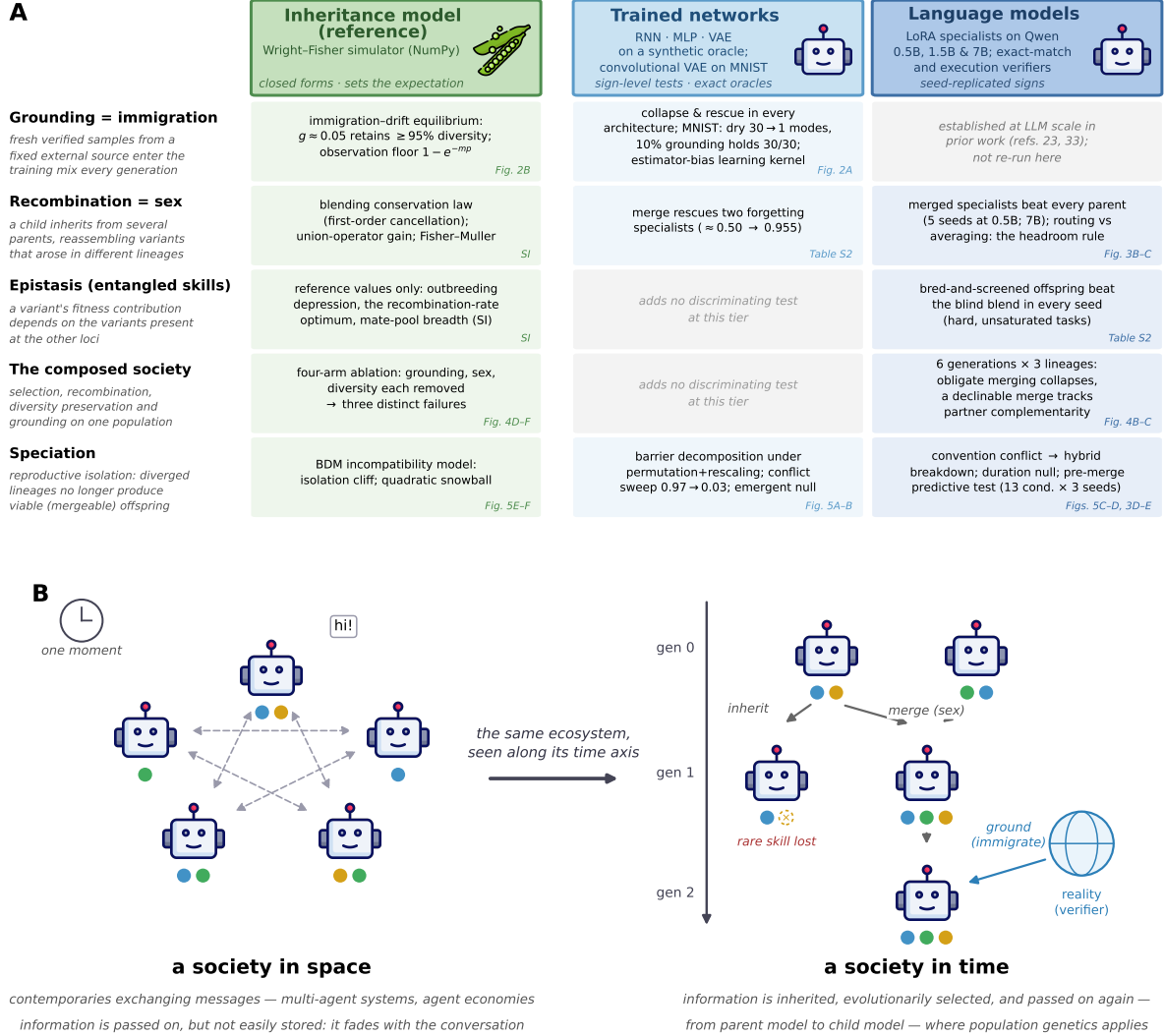


Figure 1: A map of the study. (A) Each row is a biological mechanism the paper borrows, each column a level of realism at which it is tested: an inheritance model (an exact simulation of knowledge transmission, green), trained neural networks measured against exact oracles (blue), and language models (blue). Filled cells name the experiments run at each level and, in the corner, the figure or table reporting them; grey cells were not run, either because the result is established in prior work (23, 34) or because that level adds no new test for that question. The inheritance model is the reference column: it sets the expectation the real-model experiments are read against. (B) The change of viewpoint the transfer rests on. A group of models is usually pictured as a society in space, contemporaries exchanging messages. The couplings studied here run between generations: training on model output (inheritance), weight-space merging (recombination), and verified real data entering each generation (immigration from reality). That is a society in time, which is what population genetics describes. Dots are capabilities: the rare one (gold) is lost under single-parent inheritance, reassembled by merging complementary parents, and re-supplied by grounding.

Results

The inheritance model and its calibration against trained networks

Knowledge is modelled as a distribution $\mathbf{p}_t$ over K discrete *items*, each standing for a capability, a fact or a mode of behaviour. An item is the counterpart of an allele, and a *capability* is what an item stands for. A fixed true distribution $\mathbf{p}^*$ gives each item its true frequency, and its rare tail (the items of lowest frequency) carries the knowledge most at risk. Following population genetics I call an item’s frequency $\mathbf{p}_i$ its *mass*, the probability that one sample drawn from the distribution is that item (the allele frequency of Table 1), and the mass of a set of items is the sum of their frequencies. One generation has a single parent and a single child (several parents are the subject of the merging section) and consists of three steps: draw $\mathbf{n}$ samples from the parent’s distribution; optionally add $\mathbf{m}$ samples drawn from $\mathbf{p}^*$ itself, standing for real data that has passed a verifier (*grounding*, with grounding fraction $\mathbf{g} = \mathbf{m}/(\mathbf{n}+\mathbf{m})$); and fit the child’s distribution to the pooled $\mathbf{n} + \mathbf{m}$ samples (the *refit*, which in the minimal model is simply the observed frequencies). The resampling step is the Wright–Fisher process, population genetics’ canonical model of neutral evolution, in which each generation is a random sample of size $\mathbf{n}$ from the last. In this *inheritance model* the Wright–Fisher “population” is the sample a child is trained on and its “individuals” are the $\mathbf{n} + \mathbf{m}$ samples, so it is a model of a learner. Diversity throughout is *heterozygosity*, $\mathbf{H} = 1 - \sum \mathbf{p}_i^2$, the probability that two items sampled independently from the distribution differ (high when the mass is spread over many items, zero when one item holds it all). The simulator reproduces three closed forms of the process to within 0.5% of the analytic value (Methods): the heterozygosity decay under drift alone, $\mathbf{E}[\mathbf{H}_t] = \mathbf{H}_0(1 - 1/\mathbf{n})^t$; the stationary diversity under real data, written in the next subsection; and, for K parents that each hold a given rare item with probability $\mathbf{q}$ and whose holdings are correlated by ρ (0 fully complementary, 1 identical), the expected fraction of rare items held by at least one parent, $\rho\mathbf{q} + (1 - \rho)(1 - (1 - \mathbf{q})^K)$, used in the merging section.

Trained networks are not exact copiers, because they add approximation error, optimisation noise and their own inductive bias to the resampling step, so before using Wright–Fisher as a reference I measured how far real learners depart from it. Run through the same generational loop against an exact oracle, they departed in opposite directions (Fig. S2). The sequence generators (a recurrent and a feedforward network) *smooth*, spreading probability onto items they have never seen, and so collapse more slowly than drift predicts while keeping spurious variants alive. The image autoencoder *sharpens*, concentrating probability on its commonest modes, and so collapses faster (Fig. 2A; the comparison with drift in Fig. S2). Both departures are reproduced by adding one knob to the copying step, a mutation rate toward a prior for smoothing or a temperature for sharpening (Fig. S2). A real learner is therefore treated throughout as Wright–Fisher plus a signed, measurable bias, and the two predictions that matter here (rare items are lost first, and real data arrests the loss) held in every architecture tested (Figs. 2 and S1).

In biological terms, retraining a child on a single parent is *asexual reproduction*. In a population that never recombines, a loss that happens to reach every individual can never be undone, because no individual retains the copy from which it could be rebuilt. Each such loss clicks the population one notch down, and the notches turn only one way. This is *Muller’s ratchet* (30), and model collapse has the same irreversible arm. Once every copy of a rare item is gone from all parents and all sources nothing can rebuild it, and the inheritance model shows the trap in its commonest form: a population that adopts its own collapsed output as its new reference never recovers the items it had lost, whatever real data it is fed afterwards (Fig. S3). Remedies must therefore act while copies still survive somewhere in the population.

Table 1. The dictionary. Each biological term is introduced in the section that develops it. The support column names where the evidence comes from: a figure panel or Supplementary figure or text of this paper, a reference number for the literature, or both. “Closed form” means derived in the inheritance model and verified against simulation; “empirical” means measured in a trained system; “hypothesis” means stated with a falsifier and untested.

Population genetics	Model populations	Support
Genetic drift in a finite population	Training on finite samples of model output	Closed form (Fig. 2B); collapse measured (Fig. 2A); the identification is prior work (23–27)
Immigration from a fixed source	Grounding with verified real data	Closed-form equilibrium and per-item floor (Fig. 2B); sign confirmed in trained nets (Fig. 2A); stationarity and stability under fresh data (31, 32); comparable fractions reported (23, 33, 34); conservation analogue (35)
Muller’s ratchet (asexual decay)	Irreversible arm of model collapse	The irreversibility is reproduced in the inheritance model (Fig. S3); the mutational mechanism of the ratchet is not modelled (30)
Recombination / sexual reproduction	Model merging	Fig. 3B–C: merging beats blending wherever the weight-average scores well below the best parent, and blending suffices where it does not; that merges can beat parents is established (4, 36)
Fisher–Muller effect	Merged specialists exceed every parent	Fig. 3B; inheritance-model expectation (Fig. S9); classical theory (37, 38)
Outbreeding depression under epistasis	Merging entangled skills harms offspring	Inheritance model only (Fig. S10), reproducing (39, 40); hypothesis at LLM scale
Mating systems / population structure	Who merges with whom (breadth of the parent pool)	Inheritance model only (Fig. S13), reproducing (41); hypothesis for real populations
Reproductive isolation (Bateson–Dobzhansky–Muller incompatibilities)	Merge failure from functional conflict	Fig. 5A–D and SI Text S1, Proposition S2; emergent form not observed; classical theory (42, 43); alignment tools and known residuals (44–47)
Seed bank (mating with a stored earlier generation)	Merging with one’s own ancestor	Six-generation population (Results; SI Table S2): own-ancestor merge beat a contemporary in every seed; checkpoint averaging as a stabiliser (48, 49)
Recombination modifier (a gene that sets how often other genes are shuffled)	A declinable merge: keeping the parent unchanged is scored as one candidate offspring	Fig. 4B–C (six generations, 3 seeds): a fixed early stop matched it, and declines tracked generation, not complementarity, once the two were decoupled. Modifier theory (50–52) is the motivating frame; its reduction-principle reading was not supported; gated and early-stopped merging in continual settings (53, 54)
Selection on a fitness function	Verifier-anchored selection (“reality that can say no”)	Fig. 4D–F; diversity-preserving selection from (55), inheritance-model reference (Fig. S12)

The real-data fraction required to arrest collapse

Grounding, the mixing of verified real data into each generation’s training sample, plays in the inheritance model the role that immigration plays in population genetics. A fixed external

source (p^*) supplies a fraction g of each generation’s sample, and a population that would otherwise drift to fixation settles instead at a stationary diversity (33, 34, 56). I swept g from 0 to 0.4 across 100 independent lineages (Fig. 2B and Fig. S4) to separate two questions: how much real data holds aggregate diversity, and what happens to an individual rare item.

Part of the aggregate answer exists already: that a self-consuming loop fed fresh real data settles at a stationary state instead of collapsing was shown for generative models (31), a sufficient condition on the real fraction for stability has been proved (32), the same loop with any non-vanishing synthetic fraction never recovers the real-data scaling law (57, 58), and in the first collapse study retaining 10% of the original data held perplexity steady over ten generations (23). These results establish that a grounded lineage stabilises below the real data without saying where, and the inheritance model gives the level in closed form: with m real samples added to n inherited ones each generation, diversity settles at $H_{eq} = H^* \cdot m(2n+m-1)/(n+2nm+m^2)$, where H^* is the diversity of the source, and the simulator matches this to within 0.5% (Fig. 2B). Two consequences follow that the earlier results could not show. The first is that what holds diversity is the *count* of real samples per generation, not their share of the training set. Whenever real samples are a minority ($m \ll n$) the formula reduces to $H_{eq} \approx H^* \cdot 2m/(2m+1)$ and n drops out: one real sample per generation keeps two thirds of the source’s diversity and ten keep 95%, however large the inherited sample is. The expression is Wright’s island model in haploid form: the shortfall $1/(2m+1)$ is its fixation index F_{ST} for a population receiving m migrants a generation, and the rule of thumb of conservation genetics is stated as *one migrant per generation* (35), a count and not a fraction, because of the same cancellation. The size of the receiving population drops out, and how much of the source’s diversity an island keeps is set by how many migrants reach it. In the tested setting ($K = 1000$ items, $n = 200$ inherited samples per generation, and a true distribution whose item frequencies fall off as a power law, a *Zipf* distribution, the standard model of the long tail of natural data) 95% of the source’s diversity was kept from $g \approx 0.05$ upward (Fig. S4), but that fraction is ten real samples divided by a training set of 200, and it shrinks as the training set grows. The second is that the curve is smooth. Diversity rises gradually with m , there is no value at which a lineage switches from collapsing to safe, and the lineage never reaches the source (the shortfall is about $1/(2m+1)$ at any budget, as the scaling-law results require; 57, 58). Any threshold quoted for real data is therefore a retention target one chooses and reads off the curve, not a property of the system. Comparable fractions are reported for accumulating real data in language models (34) and for the replay ratios of continual learning. Optimal mixing ratios derived for squared-error regression are far higher (about 0.6; 59), because that objective weighs every sample equally where the question here is which items survive at all.

Aggregate diversity cannot say whether one particular rare item survives, and for that the answer is elementary. Call the number m of verified real samples added per generation the *real-data budget*. Under unstratified sampling an item of frequency p appears in a batch of m real samples with probability $1 - e^{-mp}$, so a budget of $m \approx 1/p$ gives only a 63% chance of seeing the item once per generation; an item that appears in one real sample in ten thousand needs a budget of about ten thousand real samples every generation. The budget is therefore set by the rarest item one refuses to lose, and it is a lower bound, because a single copy that does arrive enters a pool of $n + m$ samples and can still be lost when the child is resampled from it (Fig. S4D, where the rarest items recover last). The rule is the immigration counterpart of the per-item extinction laws derived for closed loops (25, 60). It also explains an observation reported by others and left unexplained, that the absolute count of real samples predicts collapse better than their proportion (61): the aggregate closed form and the per-item rule both depend on m , not on g . The same arithmetic has been observed on the acquisition side, in pretraining itself: about 250 documents install a rare behaviour in models from 600 million to 13 billion parameters, although the larger models see twenty times more data, so the documents’ share of the corpus falls twentyfold while their effect does not (62). One migrant per generation, 250

poisoned documents and $\mathfrak{m} \cdot \mathfrak{p} \gtrsim 1$ are one rule read three times: what a population keeps, or acquires, of a rare item is set by the number of copies that reach it each generation, not by the size of everything else it is trained on. A fixed budget stretches further in two ways. Real data protects only the topics it covers, since when the 1,000 items are split into ten topics and the same budget is spent either on one topic or evenly over all ten, real data aimed at the topic keeps about half of its rare items alive and real data spread over all topics keeps 7% (Fig. S5), so a capability is protected by real data about that capability, not by real data in general. And an item lost from one lineage can be recovered from another lineage that still holds it, which is the subject of the next section.

In the trained networks (the recurrent and feedforward generators on the synthetic universe, Fig. S6, and the convolutional VAE on MNIST, Fig. 2A and Fig. S7) grounding reduced collapse in every case, as prior work at language-model scale had found (23, 34). Compared against the exact model, the trained networks depart in two ways, both consequences of the estimator bias measured above. The threshold softens: in the recurrent network the distance from the truth falls gradually over the whole range of $\mathfrak{g}$ tested (Fig. S6B), where the inheritance model’s diversity saturates within a few percent. And the usual measure of collapse fails for a smoothing learner. Such a network keeps assigning probability to items it was never trained on, so counting how many rare modes survive overstates its health; in the recurrent network that count is not even monotone in $\mathfrak{g}$ (Fig. S6D), while a network can retain every mode and still hold the mass in the wrong proportions. For smoothing learners I therefore measure collapse by the forward Kullback–Leibler divergence from truth to model, the standard measure of how well a model covers a distribution, which penalises every region where the truth has mass and the model has little. On real images (Fig. 2A) ungrounded self-training collapsed a convolutional VAE from thirty modes to one within fifteen generations, while about 10% grounding held all thirty (Fig. S7). The autoencoder needed about 10% real data where the inheritance model needed 5%, and the difference is what its sharpening bias costs: a learner that concentrates mass on its commonest modes loses rare ones faster than sampling alone would, and needs more real copies to hold them.

Merging operators and the retention of rare capabilities

Refitting a child on the average of its parents’ output distributions is *blending inheritance*, the pre-Mendelian view of heredity in which offspring are an average of their parents. Fleeming Jenkin’s objection to Darwin (63, 64) was that under blending a rare favourable variant is halved at every cross and swamped within a few generations, so selection could never establish it; particulate (Mendelian) inheritance, in which an allele passes intact or not at all, answered the objection, and blending was abandoned as a theory of heredity. Averaging does to a rare capability exactly what Jenkin said blending would do to a rare variant, and blending inheritance is therefore the right null model of merging. The same dilution has been reported in machine learning under three different names, without being recognised as one phenomenon: distilling onto an ensemble mean discards the members’ diversity (65), averaging expert weights loses to routing among the same experts (66), and an update held by one of $\mathfrak{N}$ parents is scaled by $1/\mathfrak{N}$ in their soup (67). In the inheritance model the dilution is a conservation law: the expected mass of a rare item in the child is $\mathfrak{q} \cdot \mathfrak{p}$ (its mass $\mathfrak{p}$ in a parent that holds it, times the probability $\mathfrak{q}$ that a parent holds it) whatever the number of parents, so averaging over more parents neither helps nor harms a rare item’s expected share, and the proposition below says exactly when the same holds for its survival.

Proposition (blending inheritance, rare-item regime). Let each of K parents independently retain a rare item, which has mass $\mathfrak{p}$ in a parent that retains it, and let the child draw $\mathfrak{n}$ samples either from one parent chosen at random or from the mean of the K parents’

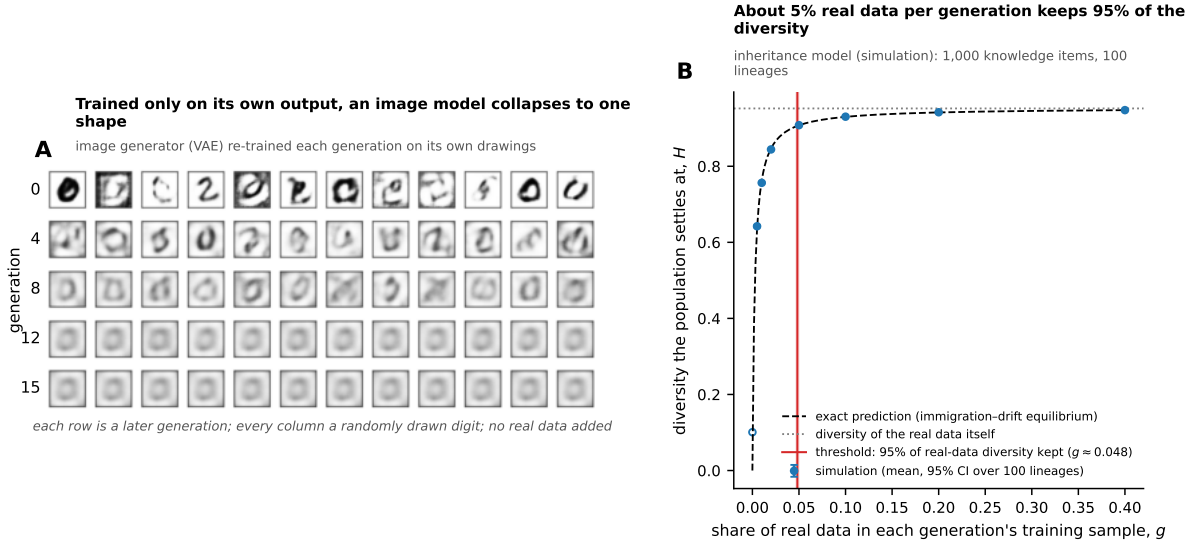


Figure 2: How much real data stops model collapse. (A) An image-generating network (a variational autoencoder) is trained on handwritten digits, then a fresh copy is trained only on the digits the previous one drew, for fifteen generations, with no real data added. Each row is a later generation (0, 4, 8, 12, 15) and each column a randomly chosen drawing. The thirty kinds of digit (ten digits $\times$ three stroke thicknesses, some kinds rare) collapse to one blurred shape; an independent classifier confirms that the number of kinds still drawn falls from 30 to 1, while adding 10% real digits each generation keeps all 30 (Fig. S7; 4 replicates). (B) The same question in the inheritance model, the exact simulation: 1,000 knowledge items, 200 samples drawn per generation, and a fraction g of fresh real samples mixed in. Points are the diversity the population settles at after 500 generations (mean and 95% CI over 100 lineages), the dashed line the exact prediction (the immigration–drift equilibrium), the dotted line the diversity of the real data itself. The curve is smooth, so any threshold is a choice: the red line marks the g at which 95% of the real data’s diversity is kept, about 0.05 (bootstrap CI shaded). The hollow point at $g = 0$ has not yet reached its equilibrium of zero. The trained image model needed about twice this fraction, because a trained network is not the exact copier the simulation assumes (Fig. S2).

distributions. The expected mass of the item in the child’s sample is the same under both schemes. When the item is rare enough that even a parent holding it rarely contributes more than one copy to the child’s sample ($n \cdot p \ll 1$), the probability that the item survives into the child is the same too: averaging over K parents makes the item K times more likely to be present in the mixture, and K times less frequent when it is, and the two factors cancel (proof in SI Text S4).

The proposition fixes the baseline against which any merging operator is judged, and it has two boundaries. For items common enough that the child usually sees several copies, averaging is safer than inheriting from one random parent, because the probability of losing an item is a convex function of its mass and averaging evens out which parent happened to hold it; the cancellation is a statement about rare items, which are the ones at risk. A *union* operator, which keeps for each item the mass it has in the parent holding it most strongly (and therefore needs a verifier to say which parent that is), raises expected retention with every additional parent at every rarity tested (Fig. S8).

Neither scheme is what model merging does in practice. The two operators in use are *weight averaging*, which averages the parents’ parameters (a network is nonlinear in its weights, so averaging weights does not average outputs and the proposition applies only by analogy; but an update held by one of N parents is still scaled by $1/N$ in the average (67), which is the dilution the proposition describes), and *routing*, which keeps every specialist intact and sends each input to the specialist that owns it (68), the practical form of the union. I compared the two at two model sizes (0.5B and 7B parameters) on easy and on deliberately hard task families (Fig. 3C for 7B on the hard families; the other size and difficulty combinations in Supplementary Information, Table S2). Routing wins by the amount averaging loses to dilution, and two things set that loss. On the easy families a 7B base has nothing to lose: after averaging it scores at ceiling on two of the three families (1.00 on both), so routing has nothing to recover and the two are equivalent. On the hard families the average falls to the level of the best single specialist (0.41 for both, over three 7B seeds), because it dilutes each specialist’s own skill, and routing among the intact specialists wins by a wide margin (0.50, ahead in every seed). A weak base (0.5B) shows the same gap even on the easy families. The operative variable is the *headroom*, the distance between what the weight-average scores and what the specialists would jointly score if every input reached the right one: it is large wherever there is room to lose to dilution (a weak base, or hard tasks at a strong one), and neither model size nor task difficulty alone predicts it. On the second base lineage the ordering is the same and the margin larger (routing 0.33 against soup 0.17 on the hard families at 1.7B, ahead in every seed, with the soup below the best specialist in every seed; Fig. S16). Whether the gain scales quantitatively with the headroom is untested.

Merging complementary specialists can also yield a model better than any of them, the *Fisher–Muller effect* (37, 38). In an asexual population two useful variants that arise in different individuals can never meet in one descendant; the lineages carrying them compete, and one is lost. Recombination puts both into one offspring, which is why sexual populations adapt faster. In the multi-locus inheritance model, merged decorrelated specialists reach a combination of variants (a *genotype*) that no parent held, while the best parent and the blended average plateau below (Fig. S9). Merges of three LoRA (69) specialists reproduced the signature, beating every parent overall (0.65 against 0.59 over five seeds at 0.5B; 0.87 against 0.81 over three seeds at 7B, in every seed), and on worst-family accuracy they were the only models competent everywhere, in every seed (Fig. 3B). The same protocol on an unrelated base lineage (SmolLM2-1.7B-Instruct: a different laboratory, architecture family and pretraining corpus) gave the same result in every one of five seeds (merge 0.66 against best specialist 0.61 overall; worst family 0.32 against 0.13; Fig. S16). That merges can exceed their parents is established for adapters (4, 36, 70); the model contributes the condition under which it happens and the operator that

realises it.

Blind recombination is not always safe. On rugged (*epistatic*) landscapes, where a variant’s contribution depends on the variants around it (71), recombining two adapted parents yields offspring below both, and the optimal recombination rate falls as entanglement grows. Both results are long established in population genetics (39) and evolutionary computation (41) and are reproduced here only to fix reference values (Fig. S10). An engineered population has an option a natural one lacks: breed many candidate offspring and keep whichever a verifier scores highest. In the inheritance model this *directed* recombination recovers the gain on every landscape where blind recombination loses it (Fig. S11), and in language models it beat the a-priori blend in every seed on hard tasks, including one seed where the blend failed catastrophically and selection was unaffected (Supplementary Information, Table S2).

Ablation of a composed population

Grounding enters a population at two points. In the inheritance model it is *grounded inheritance*, real samples added to the pooled sample the child is fit to. In a selecting population it is *grounded evaluation*: an agent is scored partly against reality and partly against the population’s own consensus ($g \cdot \text{true-fitness} + (1-g) \cdot \text{conformity}$). The consensus term stands for what a population does when it has no verifier, which is to learn from its own outputs, so $g = 0$ is a population that rewards agreement with itself. To ask whether grounding, recombination and diversity contribute separately, I ran a four-arm ablation in the multi-locus inheritance model: a population of 60 agents, each a genotype of 12 loci, adapting on a rugged (NK) landscape for 80 generations (SI Methods M3), with one operator removed per arm (Fig. 4D–F). The full system (grounded evaluation, directed recombination, and diversity-preserving selection (54; its inheritance-model reference in Fig. S12)) approached the global optimum while keeping its specialists. Removing grounded evaluation converged the population confidently on an unfit consensus, the self-consumption failure. Removing recombination stranded it on local optima, and removing diversity converged it prematurely on a worse answer. The arm without grounding fails by construction, since a rule that scores agreement will converge on agreement, but the other two removals fail in ways of their own, so under these conditions recombination and diversity are not substitutes for grounding or for each other. Magnitudes depend on the mutation, restart and selection schemes, which were not varied.

A six-generation language-model population

Merging has been iterated before, in two forms. Evolutionary merging holds a pool of parents fixed and recombines it repeatedly (5, 8, 9), and over several generations the pool stagnates as its members grow alike (11). Continual merging folds a stream of independently trained experts into one running model (53, 54, 72, 73), and in long streams it degrades unless merging is gated by similarity or stopped early (53, 54). In neither form does a lineage learn a new skill by training between merges, so what happens to a composed capability when it is inherited, extended and recombined has not been measured. I ran inheritance, recombination and immigration together as a population of language models across six generations on real datasets.

Three lineages start from one frozen base model (Qwen2.5, 1.5 billion parameters, untrained on the tasks). Each generation, every lineage acquires one new skill from six public datasets (natural-language inference (MNLI; 74), science questions (ARC-Easy; 75), commonsense completion (HellaSwag; 76), reading-comprehension spans (SQuAD; 77), yes/no questions (BoolQ; 78), pronoun resolution (WinoGrande; 79)), each scored by its own verifier, a program that marks an answer right or wrong. A skill lives in a *LoRA adapter*, a small set of trainable weights added to the frozen base (the base a shared textbook, the adapter one specialist’s mar-

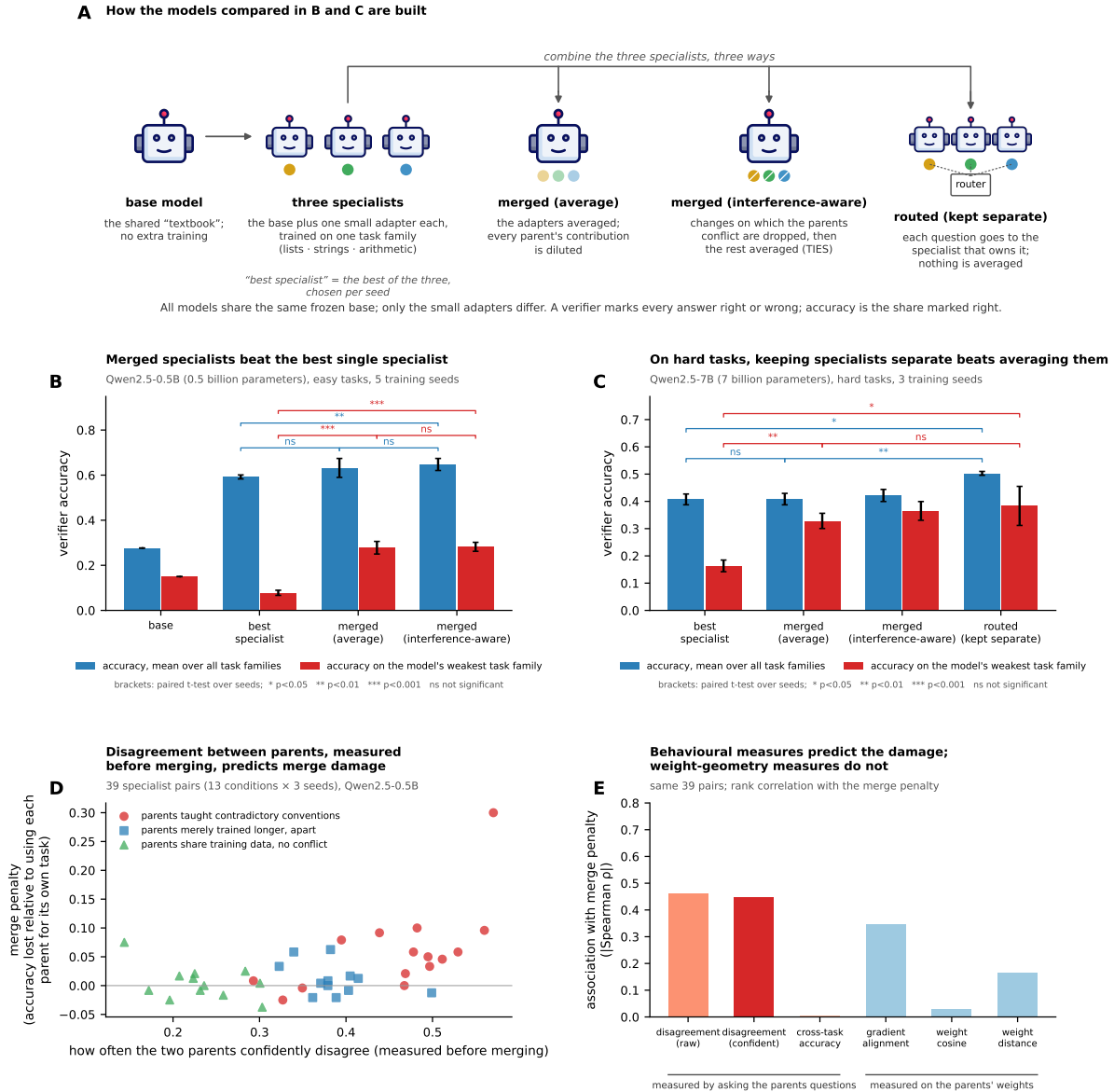


Figure 3: Merging language-model specialists: when it helps, and predicting when it will hurt. All models are built from one frozen base (Qwen2.5) plus a LoRA adapter, a small set of extra weights trained on one family of tasks (list puzzles, string puzzles or arithmetic); a verifier marks every answer right or wrong, and accuracy is the share marked right on held-out questions. (A) The models compared: the base alone; three specialists (one adapter each); their merge by averaging the adapters; their merge after dropping the changes on which the parents conflict (TIES); and routing, which keeps the specialists separate and sends each question to the one that owns it. (B) Easy tasks, 0.5-billion-parameter base, five training seeds (fixed test sets; mean and 95% CI). Both merges beat the best single specialist on the weakest task family (paired t-test over seeds, $p < 10^{-4}$), and the interference-aware merge beats it overall ($p = 0.006$; the plain average $p = 0.09$, ahead in 4 of 5 seeds); the two merges do not differ from each other. Only merged models are competent on every family. (C) Deliberately hard tasks, 7-billion-parameter base, three seeds. Averaging only matches the best specialist overall ($p = 0.96$) although it lifts the weakest family ($p = 0.009$); routing beats averaging overall ($p = 0.007$, ahead in every seed) and beats the best specialist on both measures ($p = 0.018$ and 0.014). With three seeds, some comparisons that hold in every seed are not significant (ns). (D) Predicting merge damage before merging: 39 pairs of specialists built along three axes, parents taught contradictory conventions (red), parents merely trained longer on different tasks (blue), and parents sharing training data without conflict (green). The horizontal axis is how often the two parents confidently disagree when asked the same questions before merging; the vertical axis is the merge penalty, the accuracy the merged model loses relative to answering each task with the parent that owns it. Damage concentrates in the conflicting pairs. (E) Six pre-merge measures ranked by how strongly they track the penalty (absolute Spearman correlation): measures taken by asking the

gin notes). A child inherits by continuing to train its parent’s adapter, so what the parent learned in its lifetime passes to the child (the inheritance of acquired characters that Lamarck proposed and biology rejected, and that a weight file makes trivial). Each child’s training set also contains a fixed number of examples from the skills its lineage learned in earlier generations (150, beside 300 new), so that new training does not overwrite old skills; this *replay* is the standard remedy for forgetting in continual learning (28, 29).

The curriculum is a Latin square: the lineages take the same six skills in rotated orders, like three students working through one syllabus in different sequences. A partner therefore knows things a lineage lacks early (*complementarity*, the share of the partner’s skills one lacks, is 1.0 at the first two generations) and nothing it lacks by the end (0.0 at the sixth). Complementarity is thus a swept variable, but it is also collinear with generation number, so any effect that grows with an adapter’s training age shares its signature; a second curriculum, below, breaks the collinearity. Merging averages two adapters at a weight chosen on validation data and reported on held-out tests. The arms are: never merge; always merge with a contemporary from another lineage (with verified or with self-generated replay); merge with one’s own ancestor three generations back; and a *declinable* merge, in which keeping the parent unchanged is scored as a candidate beside every merge and wins if none beats it. A control arm merges obligately through generation 2 and never afterwards (a *forced stop*), the fixed schedule the declinable arm must be compared against. Lineages are never culled, so the population has inheritance, recombination and immigration of new skills but no differential reproduction. Three training seeds; the outcome is a lineage’s accuracy over all six families.

Obligate recombination collapsed (Fig. 4B): the always-merge arm tracked the never-merge arm for three generations, then fell from 0.65 to 0.27, beginning when partner complementarity dropped below 0.8; its self-replay variant did the same (0.31), so replay was not what failed. The declinable arm neither collapsed nor won. It led at the start (0.68 against 0.60), was overtaken, and finished level with never merging (0.792 against 0.796; per-seed -0.03 , $+0.01$, $+0.01$), while one model taught the curriculum alone reached 0.80 (with replay, forgetting was not a pressure recombination could relieve). In both non-obligate arms accuracy on the skills a lineage had been taught held near 0.78 and the first skill learned never eroded ($0.85 \rightarrow 0.88$); the obligate arm fell to 0.24 on those same skills.

The choice of partner mattered more than whether to merge. Merging with one’s own ancestor three generations back, a partner that lacks the lineage’s three most recent skills but shares every convention it holds, beat merging with a contemporary in every seed (0.66 against 0.27). The ancestor supplies complementarity in time: what it lacks is exactly what the lineage has since learned, and nothing it holds was learned differently. A *seed bank* plays this role in population genetics, letting a population mate with its own stored past. Averaging a model with its own earlier checkpoint is a known stabiliser in continual learning and in self-improvement loops (48, 49); the comparison against a contemporary partner under matched conditions is what this population adds. In the declinable arm the fraction of proposed merges that were declined rose from 0.44 to 1.00 across the six generations (Fig. 4C), until every lineage declined every merge and the population had become the never-merge arm by its own choice. A control arm that merges through generation 2 and never afterwards (the forced stop) finished level with the declinable arm in every seed (0.793 against 0.792; per-seed differences -0.008 , -0.006 , $+0.011$), so the declinable arm’s outcome is explained by when it stopped and not by which merges it chose. A second curriculum, in which every lineage starts with the same skill so that complementarity is zero at the first generation, peaks at the third (0.70) and returns to zero, produced the same rise in declines with generation ($0.44 \rightarrow 0.89$). Pooled over both curricula with generation controlled, declines did not track complementarity (partial Spearman $\rho = -0.07$, 95% CI -0.21 to 0.09 , $n = 36$) but did track generation (partial $\rho = 0.31$).

Three things rise with generation in both curricula: the adapters’ training age, the number

of skills each holds, and the arrival in every lineage of the two families whose answer conventions conflict (yes/no against 1/2). Two further curricula moved only the third. In one the conflicting pair arrives in generations 1–2 of every lineage, in the other in generations 5–6, with the four compatible families filling the rest in rotated orders, so age and skill count rise identically in both (Fig. S14). Neither the decline curve nor the collapse moved with the conflict. Declines rose with generation on the same schedule in both ($0.56 \rightarrow 0.78$ and $0.44 \rightarrow 0.89$), and with generation controlled they did not track the presence of conflict (partial $\rho = -0.09$, 95% CI -0.45 to 0.15 , $n = 36$) but did track generation (partial $\rho = 0.45$). The obligate arm collapsed in both (final accuracy 0.28 and 0.39 against 0.80 and 0.78 for never merging, in every seed): the conflict-early population dipped when the pair arrived, recovered to the others’ level by generation 3, and collapsed from generation 5, while the conflict-late population collapsed from generation 4 with its conflicting pair still to come. What the four curricula leave confounded is adapter age with skill count, which rise together by construction.

A skill whose answer convention conflicts with nothing a lineage holds occupies a *new locus*, a new position in the genome filled without displacing anything, and lineages accumulate loci freely (six here; half a million facts in a lifelong-editing benchmark that averages a fresh adapter per period into the accumulated one; 80). Two skills demanding different conventions for the same kind of question (“yes/no” against “1/2” for a two-way choice) are *alternative alleles at one locus*, and a model, like a chromosome, carries one. Where conventions disagree a merged child must err against at least one parent (SI Text S1, Proposition S2). A lineage obliged to merge pays that error every generation on every pair of conflicting conventions, and the errors accumulate into collapse. In the Latin-square curriculum the collapse began at the generation when partners stopped bringing skills a lineage lacked and started bringing conventions that clashed with the ones it held, but the conflict-arrival curricula above show that moving the clash by four generations does not move the collapse: conflicting conventions set the size of each merge’s error, and something that grows with generation sets when the errors stop being repaired. Single models show the same divide: non-contradictory updates integrate safely while contradictory ones corrupt unrelated knowledge (81), and disjoint tasks make forgetting eliminable where conflicting overlap imposes a floor (82). The collapse is the second kind of knowledge arriving in a population obliged to merge.

The declinable merge was designed as a *recombination modifier*, in genetics a gene that sets how often other genes are shuffled between parents. Modifier theory holds that recombination is favoured when it assembles complementary alleles from different parents and disfavoured when it breaks combinations that already work (39, 50, 51), and that when shuffling gains nothing the *reduction principle* drives its rate to zero (52), turning the lineage asexual; on that reading the declinable merge should have switched itself off as partners stopped being complementary. The controls do not support that reading here. Acceptance fell with generation whether or not partners were complementary, and a fixed schedule reproduced the outcome. What the population establishes is narrower: one bit of selection on each recombination event, or a fixed early stop, avoids the collapse of obligate merging at no cost against never merging, and the declinable version does so without knowing in advance when to stop. The result was obtained under six generations, a single base model, and replay throughout, none of which was varied. The population also had no differential reproduction, and the Fisher–Muller argument predicts that selection is what turns recombination’s early lead into a level advantage, because a lineage that assembles the skills first leaves more descendants. Adding truncation selection (after every generation the lowest-scoring lineage is re-founded from the highest, keeping its own place in the curriculum) did not bear this out (Fig. S15). Selection acted every generation and lifted the population mean early, but the final levels converged: with selection, never merging reached 0.804 and the declinable merge 0.793 (below in every seed, by 0.011 ± 0.003), against 0.796 and 0.792 without it. Recombination’s early lead was the same with and without selection and gone by generation 5 in both. Under a curriculum that delivers every skill to every lineage the

ceiling is what one adapter can hold (0.80 for the single model taught the whole syllabus), and sex and selection each reach it sooner without raising it.

Merge failure and its dependence on functional conflict

Recombination presupposes compatible parents. In biology, lineages pushed far enough apart become separate species (*reproductive isolation*) through Bateson–Dobzhansky–Muller incompatibilities (42, 43), changes harmless on their own genetic background but deleterious in combination. This is the mechanism behind the mule’s sterility, in which two genomes that each work cannot run in the same cell. A merged model is that exposed hybrid. In the inheritance model of the process (Fig. 5 E and F) hybrid fitness stays at the parents’ level while the lineages remain compatible and then falls below the ancestor, sooner the more incompatibilities the genomes carry, and Orr showed that the number of such incompatibilities grows with the square of divergence (43). Whether a growing number of conflicts produces a fall in performance in a trained network is the question the simulation cannot answer.

In trained networks the claim must survive a known alternative. Two networks trained separately can differ in their weights for a trivial reason: the hidden units of a network can be renumbered, and in a ReLU network each unit’s incoming weights can be scaled up and its outgoing weights scaled down by the same factor, without changing what the network computes. Two networks that compute similar functions can therefore lie far apart in weight space, and averaging them gives a poor model, a *coordinate barrier*. Merge barriers between independently trained networks are famously of this kind, removable by re-aligning hidden units (44) and renormalising their activations (46) before averaging, and richer symmetry groups remove more (83). A residual that alignment does not remove is also known: networks trained on different tasks keep a barrier after permutation (47), and experts diverged far from a shared base keep one with symmetries accounted for (45). What has not been asked is what the residual measures, divergence as such or conflict in what the networks compute. To separate the two I aligned pairs of networks under permutation matching combined with exact per-unit rescaling (the complete unit symmetry group of plain ReLU MLPs; 44, 46) and measured the barrier before and after (Fig. 5 A and B). Two networks trained from different initialisations on the *same* task have a barrier the alignment removes almost entirely (residual ≈ 0.001 , the aligned merge performing at parent level): their barrier was coordinate mismatch. Two networks trained on *conflicting* label maps (the same inputs, with a fraction of the classes relabelled) have a barrier the alignment leaves unchanged ($0.502 \rightarrow 0.497$), and the merged model is functionally dead. The aligner is validated only on a special case (exact recovery of a permuted-and-rescaled copy of a network), so the share of the barrier it removes is a lower bound on the removable share, and the residual an upper bound. Sweeping the fraction of classes in conflict traces the fall in hybrid fitness from 0.97 to 0.03. That no single model can answer one prompt two ways is a matter of information, not of training (SI Text S1, Proposition S2). What the population view adds is where the cliff sits: it moves with the share of shared inputs on which the parents’ conventions contradict (Fig. 5B), and in a population that share grows whenever lineages adopt conventions independently.

The sharpest test is whether isolation emerges with no conflicting signal anywhere, as a true Bateson–Dobzhansky–Muller incompatibility would (each lineage’s changes are harmless alone). Children were diverged with no conflicting signal anywhere, using complementary class specialists and divergent input conventions, to $6.4\times$ the base training. No isolation emerged (residual 0.000 throughout). Instead the merge rescued the two specialists: each had forgotten the other’s classes and scored about 0.50 alone, and their weight-average scored 0.955 at every divergence tested. Divergence six times the base training produced the strongest Fisher–Muller effect in the paper, and no incompatibility. The language-model tier gave the same double result in each of three training seeds (Fig. 5 C and D): conflicting conventions produce function-

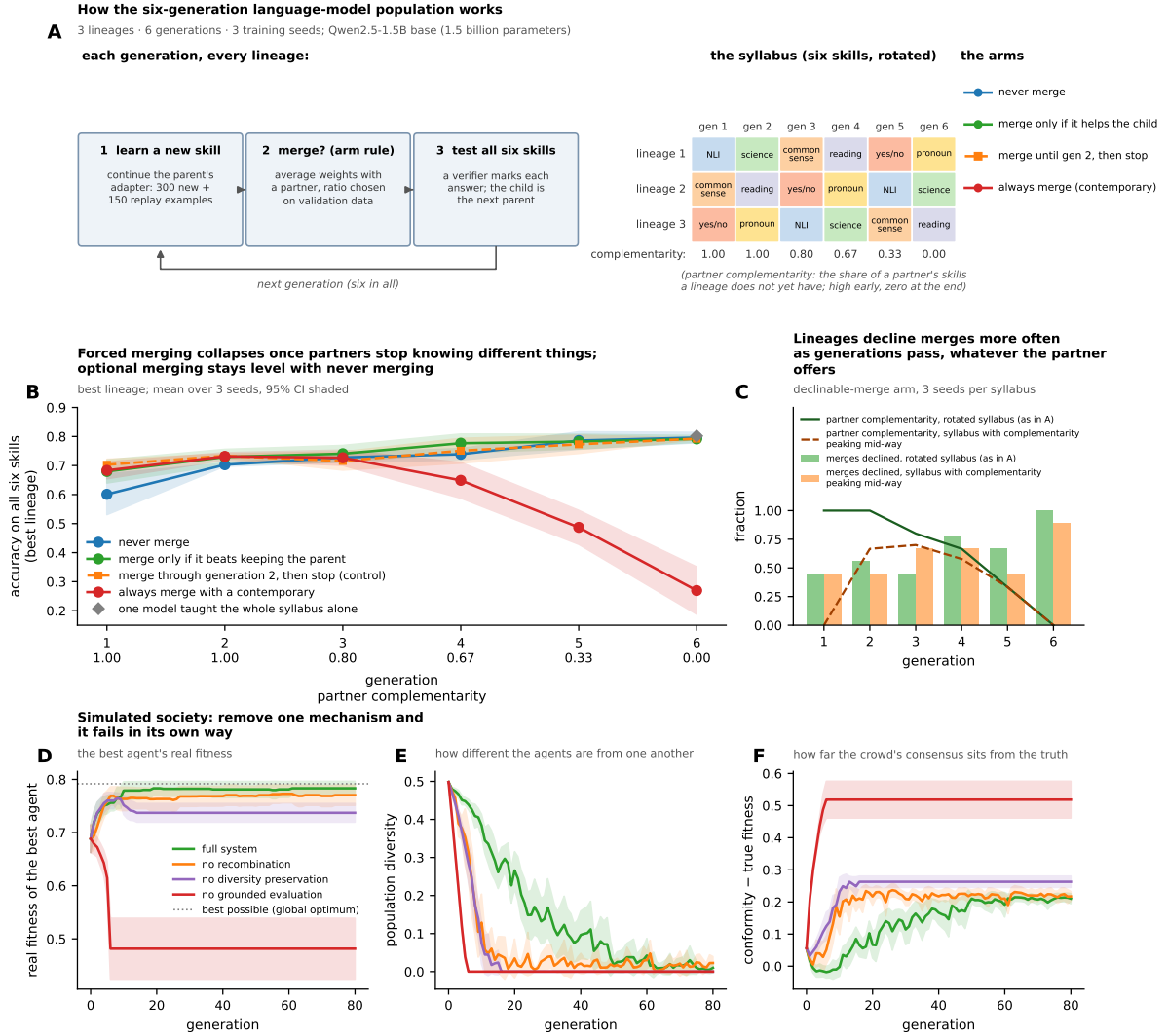


Figure 4: A population of language models over six generations. (A) The set-up. Three lineages start from one frozen 1.5-billion-parameter base (Qwen2.5-1.5B). Each generation, every lineage learns one new skill from a public dataset by continuing to train its parent’s adapter (300 new examples plus 150 replayed from earlier skills), may merge with a partner according to its arm’s rule (weights averaged at a ratio chosen on validation data), and is tested on all six skills by a verifier; the child becomes the next parent. The six skills are taken in rotated order, so a partner knows things a lineage lacks early on (complementarity 1.0) and nothing it lacks by the end (0.0). Three training seeds. (B) Accuracy over all six skills of the best lineage (mean and 95% CI). Never merging and merging only when it beats keeping the parent finish level (0.80 and 0.79); merging with a contemporary every generation collapses to 0.27, beginning when partners stop being complementary; a control that merges through generation 2 and then stops (dashed) matches the declinable arm in every seed, and a single model taught the whole syllabus alone (diamond) matches the population. (C) How often the declinable lineages refused a merge (bars) against partner complementarity (lines), under the rotated syllabus and under a second syllabus in which complementarity is zero at the start, peaks mid-way and returns to zero. Refusals rise with generation under both; with generation held fixed they do not track complementarity (partial Spearman $\rho = -0.07$, 95% CI -0.21 to 0.09 , $n = 36$). (D–F) The simulation that motivated the design: 60 agents evolving on a rugged fitness landscape with all four mechanisms (grounded evaluation, recombination, diversity preservation, mutation) and one removed per arm (12 replicates; mean and 95% CI). Removing grounded evaluation, so that agents are scored on agreement with the crowd instead of on the truth, collapses the population onto a confident but wrong consensus (D, F); removing recombination or diversity preservation strands it below the optimum (D) and drains diversity fastest (E). Each removal fails in its own way.

specific breakdown (at full conflict the merge scores 0.02, 0.12 and 0.16 on the conflicted function against 0.23–0.25 for either parent, while a budget-controlled design shows the disjoint skills merge unharmed), and over-training disjoint specialists from 1 to 12 epochs (cf. the expert-duration effect; 84) produces no isolation, the merge improving instead in every seed (0.76 $\rightarrow$ 0.95 on the parents’ private tasks). Longer expert training is reported to harm merging (84, 85) and deepening specialisation to lower feature similarity between experts (68); in the regimes tested here neither produced isolation without conflict (a complementary-class merge rescued by alignment had been seen before on label-skewed splits; 44). In every tier tested, isolation had to be provoked by functional conflict; specialisation alone did not speciate. What breaks merging is conflicting conventions on shared circuitry, not divergence as such, and this is the cost the obligate-merge arm of the six-generation population paid from its fourth generation onward, once its partners held skills it had already learned under conventions of its own (Fig. 4B).

Predicting merge damage before merging

If functional conflict is what breaks a merge, measuring it on the parents should forecast the damage before any merge is made. I tested this on thirty-nine pairs of LoRA specialists (13 training conditions $\times$ 3 seeds), built so that three properties of a pair vary independently of one another (Fig. 3D): *conflict* (the parents answer the same prompts under contradictory conventions, with their private training budgets held fixed), *compatible overlap* (the parents are trained on the same prompts under the same convention, so they share data and volume without conflict), and *duration* (the parents are trained longer on disjoint tasks, so their weights diverge with no conflict at all).

Six quantities were computed on each pair before merging. Two are functional, obtained by putting the same probe questions to both parents (probes drawn without knowledge of where the conflict lies): the fraction of probes on which the parents answer differently (*raw disagreement*), and the fraction on which they answer differently and both confidently (*confidence-weighted conflict*, proposed here as the better proxy for merge-relevant interaction, because raw disagreement also counts the harmless case in which one parent is merely ignorant). Three describe the geometry of the parents’ weight changes: the cosine similarity and the distance between the two LoRA updates, and the alignment of the two tasks’ gradients at the shared base (86). The sixth is a baseline, each parent’s accuracy on the other’s task. The pre-registered outcome is the *merge penalty*: how far the merged model falls short of the accuracy the pair would reach if each task were answered by the parent that owns it. In population genetics that shortfall is *hybrid load*, the fitness a hybrid loses relative to what its parents’ genes could jointly supply.

Functional disagreement measured before merging predicted the merge penalty (Fig. 3 D and E). Its rank correlation with the penalty was $\rho = +0.45$ (+0.46 for the confidence-weighted variant), with a 95% confidence interval excluding zero (bootstrapped over conditions, because the three seeds of one condition are not independent), and it kept $\rho \approx 0.35$ –0.40 when each condition in turn was held out and predicted from the rest. The cosine and the distance between LoRA updates showed no detectable association, and gradient alignment carried intermediate signal. The direction agrees with three recent reports: hidden-state distance between parents tracks merging loss where four parameter-space metrics, cosine among them, do not (87); global cosine, sign conflict and subspace overlap miss functional interference between task vectors (88); and gradient distance outpredicts task-vector cosine in vision (86). Those studies are correlational or in-sample; the design here holds conditions out and adds the control below. At this sample size the differences between predictors are not individually significant, only these baselines were tested, and three seeds leave substantial uncertainty about generalisation, though the functional measures led within every seed taken alone (Supplementary Information, Table

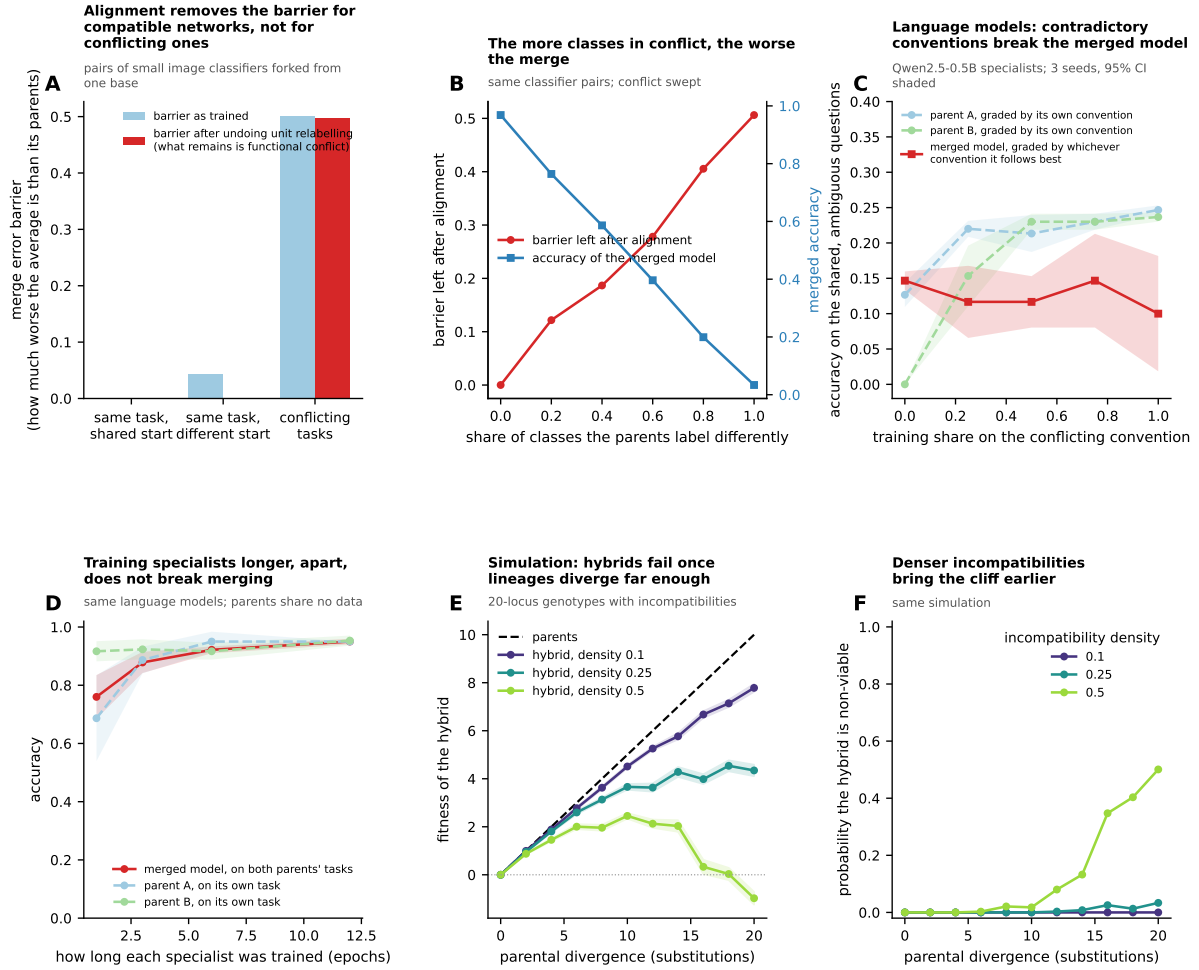


Figure 5: Model speciation: when two lineages can no longer merge. (A, B) Small image classifiers (multilayer perceptrons) forked from one trained base. Two networks that compute the same function can still differ in their weights, because hidden units can be renumbered and rescaled without changing the output; alignment undoes this before averaging. The merge error barrier is how much worse the average of two networks is than the networks themselves. (A) Two copies trained from different random starts on the same task have a barrier that alignment removes almost entirely (0.04 to 0.001); two trained on conflicting labels (the same images, some classes relabelled) keep theirs (0.50), and their average is useless (3 replicates). (B) Sweeping the share of classes in conflict moves the merged model’s accuracy from 0.97 to 0.03. (C) Language models: two specialists share a set of ambiguous questions (“sort this list”, direction unstated) and are taught opposite conventions. As the share of conflicting training grows, each parent stays good under its own convention while the merged model falls below both, in all three seeds (95% CI shaded). (D) The control: specialists trained longer and longer on different tasks, with no conflict, merge better, not worse, in every seed. (E, F) The simulation: 20-position genotypes carrying incompatibilities of the Bateson–Dobzhansky–Muller kind. Hybrid fitness tracks the parents while lineages are compatible, then crashes, sooner the denser the incompatibilities (E), and the probability of a non-viable hybrid rises with divergence (F). What breaks merging is conflicting conventions on shared machinery, not distance or specialisation as such.

S2).

The compatible-overlap control produced a finding of its own. In an initial grid that varied only conflict and duration, the best predictor was the cosine between LoRA updates ($\rho = +0.60$). Parents trained on the same prompts have aligned weight changes and also merge worse, so the cosine was reading shared training data, not incompatibility: adding pairs that share prompts without conflicting collapsed its correlation to $+0.03$. Any merge predictor validated on a grid in which conflict and shared data vary together inherits this artefact. I know of no study that has controlled for it, and it bears on the merge-prediction literature (86–88) independently of the biology. One pre-registered prediction failed: confidence weighting did not beat raw disagreement as a rank predictor, so the evidence supports functional disagreement in general and not the incompatibility-specific refinement. Headline quantitative results, with sample sizes and uncertainty, are collected in Supplementary Information, Table S2.

Discussion

Design rules. *Ground every generation* in verified reality. A few percent of real data kept most of the diversity here, but what protects a capability is the number of real examples of it that arrive each generation, not their share of the training set (the one-migrant-per-generation rule, 35; the few hundred documents that poison a model of any size, 62). The rarest capabilities therefore need a budget of about $1/p$ real examples per generation, real data aimed at them, or a parent that still holds them. *Route or screen rather than average whenever the average falls short of the best parent on any task.* On the hard families routing (sending each input to the specialist that owns it) beat weight averaging by 0.09 in every seed and screening candidate merges beat it by 0.07 (Fig. 3C), and the plain average lost nothing only where the base already answered at ceiling. *Stop recombining early, by rule or by test.* A fixed early stop, or scoring the unchanged parent beside every candidate merge, avoided the collapse of obligate merging at no cost against never merging. *When a partner must be found, prefer a stored ancestor to a divergent contemporary,* which shares every convention and beat a contemporary in every seed. *Preserve diversity as an objective in itself,* since selection can only keep what exists. *Before merging, measure functional conflict* (whether the parents answer the same prompts differently), which was cheap and predictive where weight distance was not; divergence or specialisation alone is no evidence of incompatibility, since what broke merging in every regime was conflicting conventions. The inheritance model adds one untested rule: merge sparingly, and with offspring selection, when skills are entangled (40; Fig. S13).

Continual learning at the population scale. Continual learning, the machine-learning field that teaches one network new things without erasing old ones, has found remedies for forgetting that are this framework’s operators applied to a single lineage. Rehearsal of stored real data (28, 29) is grounding, and the replay fractions the field has settled on (about 1% in instruction tuning, 89; 5% to 25% in continual pretraining, 90) look inconsistent only as fractions: at typical batch sizes each delivers tens to thousands of replayed examples of a skill per step, far more than the ten copies per generation that hold 95% of diversity. Pseudo-rehearsal, replaying the network’s own generated samples (91, 92), is grounding with no real data at all, harmless over one step and compounding over generations (Fig. 2) unless the samples are verified (33, 93). Adapters on a frozen base (94, 95) keep lineages decorrelated, consolidating them into the base is the slow store of complementary-learning-systems models (96–98), and merging as a continual-learning mechanism (72, 73, 80, 99, 100) accumulates new skills and breaks on contradictory conventions (81, 82), as the six-generation population did. That rare knowledge is forgotten first (101–103) is tail extinction observed one model at a time: forgetting and collapse differ in mechanism (interference against sampling drift) but lose the same items to the same remedies.

Two results carry over directly. A pre-merge test, disagreement between the parents on shared probes, predicts interference where weight distance does not, with the control for shared training data that earlier regression (86) and distance (87, 88) studies lacked. Weight distance fails because two adapters that learned the same skill in different runs are nearly orthogonal (cosine 0.006) yet merge with no penalty: most of a weight difference is neutral, like most DNA substitutions (Supplementary Information, Text S3). Whether to consolidate specialists or keep them modular (72, 73, 98–100) follows the same rule: route while the plain average falls short of the best parent, average once it does not. Since drift removes rare items first and a lost item is recoverable only while some parent or source still holds a copy (Fig. S3), the number to watch is accuracy on the rarest items, not the mean. Apparent forgetting can also be task misrecognition rather than lost capability (104), which the oracle excludes at the small tiers only.

Three theories of heredity. A model population runs on all three historical accounts of inheritance at once. A child continues training its parent’s adapter, so what the parent learned in its lifetime passes on (Lamarck); weight averaging blends the parents (Jenkin); and a verifier selects among variants (Darwin). Biology discarded the first for want of a mechanism and the second because blending would swamp any new variant. Here Lamarckian transmission is what lets a lineage accumulate skills (the never-merge arm reached 0.80 without any recombination). Blending dilutes whichever parent’s skill is rarest, so routing and offspring screening pay only where the plain average falls short of the best parent (Fig. 3B–C). Grounded selection is the only operator that looks outside the population, and removing it is the one ablation that fails outright: a population selected on agreement with its own consensus settles at 0.48 against 0.78 for the full society (Fig. 4D–F), confident and wrong.

Recombination’s speed advantage. In the six-generation population recombination bought speed and not level: an early lead, then parity with never merging once every skill had reached every lineage. The Fisher–Muller argument (that sex speeds adaptation by combining beneficial variants that arose in different individuals) predicts parity in exactly this case, since the curriculum guaranteed every lineage every skill, and that letting the faster lineages leave more descendants should break the parity, which it did not: selected populations reached the same ceiling, recombination’s lead again gone by generation 5. The ceiling is what one adapter can carry, and sex and selection only reach it sooner. The inheritance-model society climbs under the same operators (Fig. 4D–F) because no curriculum delivers its skills; a language-model population in which some skills come only by merging would separate the two regimes. Three refinements the framework proposed were not supported: weighting disagreement by confidence did not improve the merge predictor, the declinable merge did not track complementarity as a recombination modifier (a gene that sets how often other genes are shuffled) would, and selection did not turn recombination’s speed advantage into a level advantage. What population genetics supplied was the questions, the nulls and the controls, not a mechanism only it can explain.

Open problems. The hardest is the fitness function. Selection optimises what is measured, and for knowledge the persuasive and the true compete; a reality that can refuse is the only anchor, and building it into institutions (verification, replication, challenge among models) is a problem this paper poses and does not solve. Whether speciation emerges at scale is the second: here isolation had to be provoked by conflicting conventions, and whether long specialisation supplies such conflict on its own (84, 85) needs a population diverged far longer than any here. Collapse also reaches style: models trained on model output lose lexical and syntactic diversity (105) and model-assisted writing is individually better but collectively less diverse (106, 107), because a voice is a distribution over rare variants, exactly what drift erases first and blending averages away; whether the remedies transfer is untested.

Outlook. Language-model development is consolidating around the operators studied here: synthetic-data flywheels (inheritance), merging and routing of specialist fine-tunes (recombi-

nation and population structure), verifier-gated pipelines (grounded selection), and periodic consolidation of adapters into new bases. The forecast is a population that recombines early and consolidates late, until conflicting conventions split it into lineages connected by routing instead of merging, and the pre-merge conflict test can measure which way it goes. Biology receives in return a model system in which every genotype, environment and mating decision is observable and manipulable, and the evolution of sex can be studied with interventions (unbounded parents, offspring preview, directed mating) no living system permits.

Materials and Methods

Full procedures, parameters, and replicate counts are in Supplementary Information, Methods. Appendix 1 (*The figures explained*) restates every main and supplementary figure with a legend that explains the machine-learning experiment behind it for readers from biology.

Inheritance-model tier. A NumPy/SciPy Wright–Fisher simulator over K-item distributions (knowledge as `p_t`, Zipf-tailed truth `p*`, and drift–grounding–refit generations), extended with a learning kernel (a smoothing and a sharpening knob on the refit), multi-locus genotypes on additive and Kauffman NK landscapes, n-parent crossover, and finite-population loops. All parameters live in per-experiment YAML configs. Every run derives its randomness from one master seed (`SeedSequence.spawn`) and is bitwise reproducible. Scientific-validation tests assert the closed forms to within 0.5% and run in CI alongside 151 further correctness tests.

Neural tier. Trained-network experiments realise the same abstractions against an exact oracle. Histogram, RNN, MLP and VAE generators run on a synthetic mode universe, where the histogram model reduces the harness exactly to the inheritance model (the bridge gate), and a convolutional VAE runs on MNIST with a frozen CNN oracle at 98.5% mode accuracy (its confusion matrix is recorded as the measurement floor). Speciation experiments fork no-BatchNorm MLPs (784–512–512–10) from a shared base, weight-average them, and measure the error barrier along the straight line between the two weight vectors (the linear-mode-connectivity barrier) before and after alignment. Alignment composes deterministic Git Re-Basin permutation matching with exact per-unit scale canonicalisation, the unit symmetry group of this architecture class taken as the search space, and is gated by exact recovery of a permuted-and-rescaled copy. Control recovery does not establish global optimality.

Language-model tier. LoRA specialists (rank 16) on procedurally generated task families with an exact-match verifier, on frozen Qwen2.5-Instruct bases (0.5B on one 16 GB GPU; 7B on one L40S). Operators: weight-space merges via adapter arithmetic (the plain weight average, or soup, and TIES, which reconciles the sign of each parameter change across parents before averaging; 4), per-input routing, and Dirichlet-sampled offspring populations screened on held-out validation splits. The six-generation population uses the Qwen2.5-1.5B base model, six public datasets with per-family exact-match or execution verifiers, and rank-16 adapters continued from the parent adapter each generation (300 new and 150 replay examples, 3 epochs), merged over the weight grid 0.5/0.5, 0.3/0.7, 0.7/0.3 chosen on 20 validation items per family and reported on 60 held-out test items, with the unchanged parent as a further candidate in the declinable arm; three training seeds. Multi-seed protocols fix the test sets and vary the training seed. The predictive test computes all predictors pre-merge (generation confidence from token log-probabilities, base-model gradient cosines, and LoRA-delta geometry computed exactly in the adapters’ low-rank factor space) and evaluates merges on held-out tests. Its rows are not independent, because parents share task-data seeds across conditions, so inference is condition-clustered and per-seed and leave-one-seed-out sensitivity are reported alongside; a committed script produces these statistics. Statistical, per-seed reproducibility is documented for the GPU tiers.

Data and code availability. All code, configs, seeds, results artifacts (with content hashes), figures, and a one-command reproduction script will be deposited openly (repository + archived DOI) on publication; every figure in this paper regenerates from committed artifacts without re-simulation.

Acknowledgements

This work was done in close collaboration with Claude Opus 5 and Claude Fable 5.1 (Anthropic). I conceived the framework and the population-genetic reading, chose the questions and the experiments, set the pre-registered predictions and falsifiers, directed every stage, judged the results and edited the text; the models contributed to the experimental design, wrote the code and ran the experiments under my direction, performed the analyses and drafted the text. I take full responsibility for the content. I thank Imperial College London for funding.

References

1. B. Laufer, H. Oderinwale, J. Kleinberg, Anatomy of a machine learning ecosystem: 2 million models on Hugging Face. *arXiv [Preprint]* (2025). <https://doi.org/10.48550/arXiv.2508.06811>.
2. E. Horwitz, A. Shul, Y. Hoshen, Unsupervised model tree heritage recovery. *Int. Conf. Learn. Represent.* (2025). <https://doi.org/10.48550/arXiv.2405.18432>.
3. W. Jiang, et al., PeaTMOSS: A dataset and initial analysis of pre-trained models in open-source software. *Proc. Int. Conf. Min. Softw. Repos.* (2024). <https://doi.org/10.48550/arXiv.2402.0069>.
4. P. Yadav, D. Tam, L. Choshen, C. Raffel, M. Bansal, TIES-Merging: Resolving interference when merging models. *Adv. Neural Inf. Process. Syst.* **36** (2023). <https://doi.org/10.48550/arXiv.2405.18432>.
5. T. Akiba, M. Shing, Y. Tang, Q. Sun, D. Ha, Evolutionary optimization of model merging recipes. *Nat. Mach. Intell.* **7**, 195–204 (2025).
6. C. Goddard, et al., Arcee’s MergeKit: A toolkit for merging large language models. *Proc. Conf. Empir. Methods Nat. Lang. Process. (Industry Track)*, 477–485 (2024). <https://doi.org/10.48550/arXiv.2403.13257>.
7. E. Yang, et al., Model merging in LLMs, MLLMs, and beyond: Methods, theories, applications and opportunities. *arXiv [Preprint]* (2024). <https://doi.org/10.48550/arXiv.2408.07666>.
8. Y. Zhang, et al., Nature-inspired population-based evolution of large language models (GENOME). *arXiv [Preprint]* (2025). <https://doi.org/10.48550/arXiv.2503.01155>.
9. J. Abrantes, et al., Competition and attraction improve model fusion (M2N2). *Proc. Genet. Evol. Comput. Conf.* (2025). <https://doi.org/10.48550/arXiv.2508.16204>.
10. V. Subramaniam, et al., Multiagent finetuning: Self-improvement with diverse reasoning chains. *arXiv [Preprint]* (2025). <https://doi.org/10.48550/arXiv.2501.05707>.
11. Y. Hu, Y. Yao, N. Zhang, H. Chen, S. Deng, Exploring model kinship for merging large language models. *arXiv [Preprint]* (2024). <https://doi.org/10.48550/arXiv.2410.12613>.
12. NVIDIA (B. Adler, et al.), Nemotron-4 340B technical report. *arXiv [Preprint]* (2024). <https://doi.org/10.48550/arXiv.2406.11704>.
13. M. Abdin, et al., Phi-4 technical report. *arXiv [Preprint]* (2024). <https://doi.org/10.48550/arXiv.2412.089>.

14. Y. Wang, et al., Self-Instruct: Aligning language models with self-generated instructions. *Proc. Annu. Meet. Assoc. Comput. Linguist.* (2023). <https://doi.org/10.48550/arXiv.2212.10560>.
15. B. Thompson, et al., A shocking amount of the web is machine translated: Insights from multi-way parallelism. *Findings Assoc. Comput. Linguist.: ACL* (2024). <https://doi.org/10.48550/arXiv>
16. W. Liang, et al., Monitoring AI-modified content at scale: A case study on the impact of ChatGPT on AI conference peer reviews. *Proc. Int. Conf. Mach. Learn.* (2024). <https://doi.org/10.48550/arXiv.2403.07183>.
17. P. Villalobos, et al., Position: Will we run out of data? Limits of LLM scaling based on human-generated data. *Proc. Int. Conf. Mach. Learn.* (2024). <https://doi.org/10.48550/arXiv.2211.043>
18. L. Brinkmann, et al., Machine culture. *Nat. Hum. Behav.* **7**, 1855–1868 (2023).
19. J. S. Park, et al., Generative agents: Interactive simulacra of human behavior. *Proc. ACM Symp. User Interface Softw. Technol.* (2023). <https://doi.org/10.48550/arXiv.2304.03442>.
20. T. Guo, et al., Large language model based multi-agents: A survey of progress and challenges. *Proc. Int. Joint Conf. Artif. Intell.* (2024). <https://doi.org/10.48550/arXiv.2402.01680>.
21. N. Tomasev, et al., Virtual agent economies. *arXiv [Preprint]* (2025). <https://doi.org/10.48550/arXiv.250>
22. A. Livnat, C. Papadimitriou, Sex as an algorithm: The theory of evolution under the lens of computation. *Commun. ACM* **59**, 84–93 (2016).
23. I. Shumailov, et al., AI models collapse when trained on recursively generated data. *Nature* **631**, 755–759 (2024).
24. J. P. Crutchfield, S. Whalen, Structural drift: The population dynamics of sequential learning. *PLOS Comput. Biol.* **8**, e1002510 (2012).
25. S. Riis, Drift and selection in LLM text ecosystems. *arXiv [Preprint]* (2026). <https://doi.org/10.48550/arXiv>
26. M. Benati, A. Londei, D. Lanzieri, V. Loreto, First-extinction law for resampling processes. *arXiv [Preprint]* (2025). <https://doi.org/10.48550/arXiv.2509.20101>.
27. Y. Yoon, D. Hu, I. Weissburg, Y. Qin, H. Jeong, Model collapse in the self-consuming chain of diffusion finetuning: A novel perspective from quantitative trait modeling. *Int. Conf. Learn. Represent.* (2025). <https://doi.org/10.48550/arXiv.2407.17493>.
28. M. McCloskey, N. J. Cohen, Catastrophic interference in connectionist networks: The sequential learning problem. *Psychol. Learn. Motiv.* **24**, 109–165 (1989).
29. R. M. French, Catastrophic forgetting in connectionist networks. *Trends Cogn. Sci.* **3**, 128–135 (1999).
30. H. J. Muller, The relation of recombination to mutational advance. *Mutat. Res.* **1**, 2–9 (1964).
31. S. Alemohammad, et al., Self-consuming generative models go MAD. *Int. Conf. Learn. Represent.* (2024). <https://doi.org/10.48550/arXiv.2307.01850>.
32. Q. Bertrand, A. J. Bose, A. Duplessis, M. Jiralerspong, G. Gidel, On the stability of iterative retraining of generative models on their own data. *Int. Conf. Learn. Represent.* (2024). <https://doi.org/10.48550/arXiv.2310.00429>.
33. B. Yi, Q. Liu, Y. Cheng, H. Xu, Escaping model collapse via synthetic data verification. *arXiv [Preprint]* (2025). <https://doi.org/10.48550/arXiv.2510.16657>.

34. M. Gerstgrasser, et al., Is model collapse inevitable? Breaking the curse of recursion by accumulating real and synthetic data. *Conf. Lang. Model.* (2024). <https://doi.org/10.48550/arXiv.2404.01>
35. L. S. Mills, F. W. Allendorf, The one-migrant-per-generation rule in conservation and management. *Conserv. Biol.* **10**, 1509–1518 (1996).
36. M. Wortsman, et al., Model soups: Averaging weights of multiple fine-tuned models improves accuracy without increasing inference time. *Proc. Int. Conf. Mach. Learn.* (2022). <https://doi.org/10.48550/arXiv.2203.05482>.
37. R. A. Fisher, *The Genetical Theory of Natural Selection* (Clarendon Press, 1930).
38. H. J. Muller, Some genetic aspects of sex. *Am. Nat.* **66**, 118–138 (1932).
39. S. P. Otto, M. W. Feldman, Deleterious mutations, variable epistatic interactions, and the evolution of recombination. *Theor. Popul. Biol.* **51**, 134–147 (1997).
40. A. R. Templeton, “Coadaptation and outbreeding depression” in *Conservation Biology: The Science of Scarcity and Diversity*, M. E. Soulé, Ed. (Sinauer, 1986), pp. 105–116.
41. M. Tomassini, *Spatially Structured Evolutionary Algorithms: Artificial Evolution in Space and Time* (Springer, 2005).
42. H. A. Orr, The population genetics of speciation: The evolution of hybrid incompatibilities. *Genetics* **139**, 1805–1813 (1995).
43. H. A. Orr, M. Turelli, The evolution of postzygotic isolation: Accumulating Dobzhansky–Muller incompatibilities. *Evolution* **55**, 1085–1094 (2001).
44. S. K. Ainsworth, J. Hayase, S. Srinivasa, Git Re-Basin: Merging models modulo permutation symmetries. *Int. Conf. Learn. Represent.* (2023). <https://doi.org/10.48550/arXiv.2209.04836>.
45. E. Sharma, D. M. Roy, G. K. Dziugaite, The non-local model merging problem: Permutation symmetries and variance collapse. arXiv [Preprint] (2024). <https://doi.org/10.48550/arXiv.2410.1270>
46. K. Jordan, H. Sedghi, O. Saukh, R. Entezari, B. Neyshabur, REPAIR: REnormalizing permuted activations for interpolation repair. *Int. Conf. Learn. Represent.* (2023). <https://doi.org/10.48550/arXiv.2211.08403>.
47. G. Stoica, et al., ZipIt! Merging models from different tasks without training. *Int. Conf. Learn. Represent.* (2024). <https://doi.org/10.48550/arXiv.2305.03053>.
48. A. Kleiman, G. K. Dziugaite, J. Frankle, S. Kakade, M. Paul, Soup to go: Mitigating forgetting during continual learning with model averaging. arXiv [Preprint] (2025). <https://doi.org/10.48550/arXiv.2501.05559>.
49. X. Yuan, et al., Superficial self-improved reasoners benefit from model merging. arXiv [Preprint] (2025). <https://doi.org/10.48550/arXiv.2503.02103>.
50. N. H. Barton, A general model for the evolution of recombination. *Genet. Res.* **65**, 123–144 (1995).
51. S. P. Otto, T. Lenormand, Resolving the paradox of sex and recombination. *Nat. Rev. Genet.* **3**, 252–261 (2002).
52. L. Altenberg, M. W. Feldman, Selection, generalized transmission and the evolution of modifier genes. I. The reduction principle. *Genetics* **117**, 559–572 (1987).

53. T. Fukuda, H. Kera, K. Kawamoto, Adapter merging with centroid prototype mapping for scalable class-incremental learning. *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit.* (2025). <https://doi.org/10.48550/arXiv.2412.18219>.
54. D. Shenaj, O. Bohdal, T. Ceritli, M. Ozay, P. Zanuttigh, U. Michieli, K-Merge: Online continual merging of adapters for on-device large language models. *arXiv [Preprint]* (2025). <https://doi.org/10.48550/arXiv.2510.13537>.
55. J. Lehman, K. O. Stanley, Abandoning objectives: Evolution through the search for novelty alone. *Evol. Comput.* **19**, 189–223 (2011).
56. S. Wright, Evolution in Mendelian populations. *Genetics* **16**, 97–159 (1931).
57. E. Dohmatob, Y. Feng, P. Yang, F. Charton, J. Kempe, A tale of tails: Model collapse as a change of scaling laws. *Proc. Int. Conf. Mach. Learn.* (2024). <https://doi.org/10.48550/arXiv.2402.0704>.
58. E. Dohmatob, Y. Feng, A. Subramonian, J. Kempe, Strong model collapse. *arXiv [Preprint]* (2024). <https://doi.org/10.48550/arXiv.2410.04840>.
59. A. Garg, S. Bhattacharya, P. Sur, Preventing model collapse under overparametrization: Optimal mixing ratios for interpolation learning and ridge regression. *arXiv [Preprint]* (2025). <https://doi.org/10.48550/arXiv.2509.22341>.
60. A. T. Suresh, A. Thangaraj, A. N. K. Khandavally, Rate of model collapse in recursive training. *arXiv [Preprint]* (2024). <https://doi.org/10.48550/arXiv.2412.17646>.
61. J. Kazdan, et al., Collapse or thrive? Perils and promises of synthetic data in a self-generating world. *arXiv [Preprint]* (2024). <https://doi.org/10.48550/arXiv.2410.16713>.
62. A. Souly, et al., Poisoning attacks on LLMs require a near-constant number of poison samples. *arXiv [Preprint]* (2025). <https://doi.org/10.48550/arXiv.2510.07192>.
63. F. Jenkin, The origin of species [review]. *North Br. Rev.* **46**, 277–318 (1867).
64. M. Bulmer, Did Jenkin’s swamping argument invalidate Darwin’s theory of natural selection? *Br. J. Hist. Sci.* **37**, 281–297 (2004).
65. A. Malinin, B. Mlodozieniec, M. Gales, Ensemble distribution distillation. *Int. Conf. Learn. Represent.* (2020). <https://doi.org/10.48550/arXiv.1905.00076>.
66. M. Li, et al., Branch-Train-Merge: Embarrassingly parallel training of expert language models. *arXiv [Preprint]* (2022). <https://doi.org/10.48550/arXiv.2208.03306>.
67. X. Yuan, et al., Behavior knowledge merge in reinforced agentic models. *arXiv [Preprint]* (2026). <https://doi.org/10.48550/arXiv.2601.13572>.
68. J. Pari, S. Jelassi, P. Agrawal, Collective model intelligence requires compatible specialization. *arXiv [Preprint]* (2024). <https://doi.org/10.48550/arXiv.2411.02207>.
69. E. J. Hu, et al., LoRA: Low-rank adaptation of large language models. *Int. Conf. Learn. Represent.* (2022). <https://doi.org/10.48550/arXiv.2106.09685>.
70. L. Yu, B. Yu, H. Yu, F. Huang, Y. Li, Language models are super Mario: Absorbing abilities from homologous models as a free lunch. *Proc. Int. Conf. Mach. Learn.* (2024). <https://doi.org/10.48550/arXiv.2311.03099>.
71. S. A. Kauffman, S. Levin, Towards a general theory of adaptive walks on rugged landscapes. *J. Theor. Biol.* **128**, 11–45 (1987).

72. D. Marczak, B. Twardowski, T. Trzciński, S. Cygert, MagMax: Leveraging model merging for seamless continual learning. *Proc. Eur. Conf. Comput. Vis.* (2024). <https://doi.org/10.48550/arXiv>.
73. S. Dziadzio, et al., How to merge your multimodal models over time? *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit.* (2025). <https://doi.org/10.48550/arXiv.2412.06712>.
74. A. Williams, N. Nangia, S. R. Bowman, A broad-coverage challenge corpus for sentence understanding through inference. *Proc. Conf. North Am. Chapter Assoc. Comput. Linguist. Hum. Lang. Technol.*, 1112–1122 (2018).
75. P. Clark, et al., Think you have solved question answering? Try ARC, the AI2 Reasoning Challenge. *arXiv [Preprint]* (2018). <https://doi.org/10.48550/arXiv.1803.05457>.
76. R. Zellers, A. Holtzman, Y. Bisk, A. Farhadi, Y. Choi, HellaSwag: Can a machine really finish your sentence? *Proc. Annu. Meet. Assoc. Comput. Linguist.*, 4791–4800 (2019).
77. P. Rajpurkar, J. Zhang, K. Lopyrev, P. Liang, SQuAD: 100,000+ questions for machine comprehension of text. *Proc. Conf. Empir. Methods Nat. Lang. Process.*, 2383–2392 (2016).
78. C. Clark, et al., BoolQ: Exploring the surprising difficulty of natural yes/no questions. *Proc. Conf. North Am. Chapter Assoc. Comput. Linguist. Hum. Lang. Technol.*, 2924–2936 (2019).
79. K. Sakaguchi, R. Le Bras, C. Bhagavatula, Y. Choi, WinoGrande: An adversarial Winograd schema challenge at scale. *Proc. AAAI Conf. Artif. Intell.* **34**, 8732–8740 (2020).
80. L. Thede, K. Roth, M. Bethge, Z. Akata, T. Hartvigsen, WikiBigEdit: Understanding the limits of lifelong knowledge editing in LLMs. *Proc. Int. Conf. Mach. Learn.* (2025). <https://doi.org/10.48550/arXiv.2503.05683>.
81. S. Clemente, et al., In praise of stubbornness: An empirical case for cognitive-dissonance aware continual update of knowledge in LLMs. *arXiv [Preprint]* (2025). <https://doi.org/10.48550/arXiv.2503.05683>.
82. J. Störk, Interference and retention in continual learning. *arXiv [Preprint]* (2026). <https://doi.org/10.48550/arXiv.2601.22285>.
83. T. Li, Z. Shen, Scaling linear mode connectivity and merging to billion-parameter pre-trained transformers. *arXiv [Preprint]* (2026). <https://doi.org/10.48550/arXiv.2606.23607>.
84. N. Kozodoi, Z. Afolabi, J. Butler, Are we merging the right models? Impact of expert training duration on model merging for LLMs. *arXiv [Preprint]* (2026). <https://doi.org/10.48550/arXiv.2607.14126>.
85. S. Horoi, G. Wolf, E. Belilovsky, G. K. Dziugaite, From memorization to parameter interference: How overtraining experts harms model merging. *Proc. Int. Conf. Mach. Learn.* (2026). <https://doi.org/10.48550/arXiv.2506.14126>.
86. L. Zhou, B. Zhao, R. Yu, E. Rodolà, Demystifying mergeability: Interpretable properties to predict model merging success. *arXiv [Preprint]* (2026). <https://doi.org/10.48550/arXiv.2601.22285>.
87. Y. Cao, et al., An empirical study and theoretical explanation on task-level model-merging collapse. *arXiv [Preprint]* (2026). <https://doi.org/10.48550/arXiv.2603.09463>.
88. C. Zhu, X. Li, T. Cai, When do task vectors interfere? Mapping the validity boundaries of weight-space composition. *arXiv [Preprint]* (2026). <https://doi.org/10.48550/arXiv.2608.09490>.
89. T. Scialom, T. Chakrabarty, S. Muresan, Fine-tuned language models are continual learners. *Proc. Conf. Empir. Methods Nat. Lang. Process.* (2022). <https://doi.org/10.48550/arXiv.2205.12390>.

90. A. Ibrahim, et al., Simple and scalable strategies to continually pre-train large language models. *Trans. Mach. Learn. Res.* (2024). <https://doi.org/10.48550/arXiv.2403.08763>.
91. A. Robins, Catastrophic forgetting, rehearsal and pseudorehearsal. *Connect. Sci.* **7**, 123–146 (1995).
92. H. Shin, J. K. Lee, J. Kim, J. Kim, Continual learning with deep generative replay. *Adv. Neural Inf. Process. Syst.* **30** (2017). <https://doi.org/10.48550/arXiv.1705.08690>.
93. Y. Feng, et al., Beyond model collapse: Scaling up with synthesized data requires verification. *arXiv [Preprint]* (2024). <https://doi.org/10.48550/arXiv.2406.07515>.
94. A. A. Rusu, et al., Progressive neural networks. *arXiv [Preprint]* (2016). <https://doi.org/10.48550/arXiv.1605.06461>.
95. D. Biderman, et al., LoRA learns less and forgets less. *Trans. Mach. Learn. Res.* (2024). <https://doi.org/10.48550/arXiv.2405.09673>.
96. J. L. McClelland, B. L. McNaughton, R. C. O'Reilly, Why there are complementary learning systems in the hippocampus and neocortex: Insights from the successes and failures of connectionist models of learning and memory. *Psychol. Rev.* **102**, 419–457 (1995).
97. D. Kumaran, D. Hassabis, J. L. McClelland, What learning systems do intelligent agents need? Complementary learning systems theory updated. *Trends Cogn. Sci.* **20**, 512–534 (2016).
98. J. Schwarz, et al., Progress & Compress: A scalable framework for continual learning. *Proc. Int. Conf. Mach. Learn.* (2018).
99. G. Ilharco, et al., Editing models with task arithmetic. *Int. Conf. Learn. Represent.* (2023). <https://doi.org/10.48550/arXiv.2212.04089>.
100. A. Alexandrov, et al., Mitigating catastrophic forgetting in language transfer via model merging. *Findings Assoc. Comput. Linguist.: EMNLP* (2024). <https://doi.org/10.48550/arXiv.2407.08690>.
101. M. Toneva, et al., An empirical study of example forgetting during deep neural network learning. *Int. Conf. Learn. Represent.* (2019). <https://doi.org/10.48550/arXiv.1812.05159>.
102. N. Kandpal, H. Deng, A. Roberts, E. Wallace, C. Raffel, Large language models struggle to learn long-tail knowledge. *Proc. Int. Conf. Mach. Learn.* (2023). <https://doi.org/10.48550/arXiv.2211.00000>.
103. X. Liu, et al., Long-tailed class incremental learning. *Proc. Eur. Conf. Comput. Vis.* (2022). <https://doi.org/10.48550/arXiv.2210.00266>.
104. S. Kotha, J. M. Springer, A. Raghunathan, Understanding catastrophic forgetting in language models via implicit inference. *Int. Conf. Learn. Represent.* (2024). <https://doi.org/10.48550/arXiv.2310.00000>.
105. Y. Guo, G. Shang, M. Vazirgiannis, C. Clavel, The curious decline of linguistic diversity: Training language models on synthetic text. *Findings Assoc. Comput. Linguist.: NAACL* (2024). <https://doi.org/10.48550/arXiv.2311.09807>.
106. V. Padmakumar, H. He, Does writing with language models reduce content diversity? *Int. Conf. Learn. Represent.* (2024). <https://doi.org/10.48550/arXiv.2309.05196>.
107. A. R. Doshi, O. P. Hauser, Generative AI enhances individual creativity but reduces the collective diversity of novel content. *Sci. Adv.* **10**, eadn5290 (2024).