

Simulating Language

Lecture 13: Evolution of learning bias

Simon Kirby

simon@ling.ed.ac.uk



Linguistic nativism

- Is language innate?
 - Not really a useful question...
 - Language is a product of biology and environment, like everything else
- A better question: does our biology provide a **domain-specific** learning device which imposes **strong constraints** on the form that language can take?

What is domain-specificity?

- A general definition
 - A learning device that only applies to a specific domain (e.g. language, causal relationships, social relationships, ...)
 - Domain-general: I use the same mechanism to learn language, causal relationships, social relationships, ...
- An evolutionary definition
 - Evolved under selection for a specific function (e.g. language learning mechanism evolved for language learning)
 - Domain-general: mechanism did not evolve under selection solely for the function it is currently used for (e.g. general-purpose learning mechanism evolved for learning language, causal relationships, ...)

A classic nativist argument

- Pinker & Bloom (1990): **Yes**, our biology provides a domain-specific learning device which imposes strong constraints on the form that language can take?
- Domain-specific:
 - *“we have argued ... that human language, like other specialized biological systems, evolved by natural selection. Our conclusion is based on two facts ...: language shows signs of complex design for the communication of propositional structures, and the only explanation for the origin of organs with complex design is the process of natural selection”*

A classic nativist argument

- Pinker & Bloom (1990): **Yes**, our biology provides a domain-specific learning device which imposes strong constraints on the form that language can take?
- Strong constraints:
 - *“Children are fluent speakers of complex grammatical sentences by the age of three, without benefit of formal instruction. They are capable of inventing languages that are more systematic than those they hear, showing resemblances to languages that they have never heard, and they obey subtle grammatical principles for which there is no evidence in their environments.”*

Wait a minute...

- *“language shows signs of complex design for the communication of propositional structures”*
- *“the only explanation for the origin of **organs** with complex design is the process of natural selection”*
- Language isn't an organ, it's a socially-learnt behaviour
 - The language organ / faculty is a device for learning a language from data
 - What's the relationship between an evolving learning device and an evolving socially-transmitted language? What kind of language faculties evolve?

What's the relationship between an evolving bias and an evolving socially-transmitted language?

- Damned if I know - let's model it and find out
- We already have a lovely model of cultural evolution
 - Iterated Bayesian Learning
- We already have a lovely model of biological evolution
 - Genetic algorithms
- Let's just evolve the prior biases of Bayesian learners
- What happens to the priors?

A co-evolutionary model

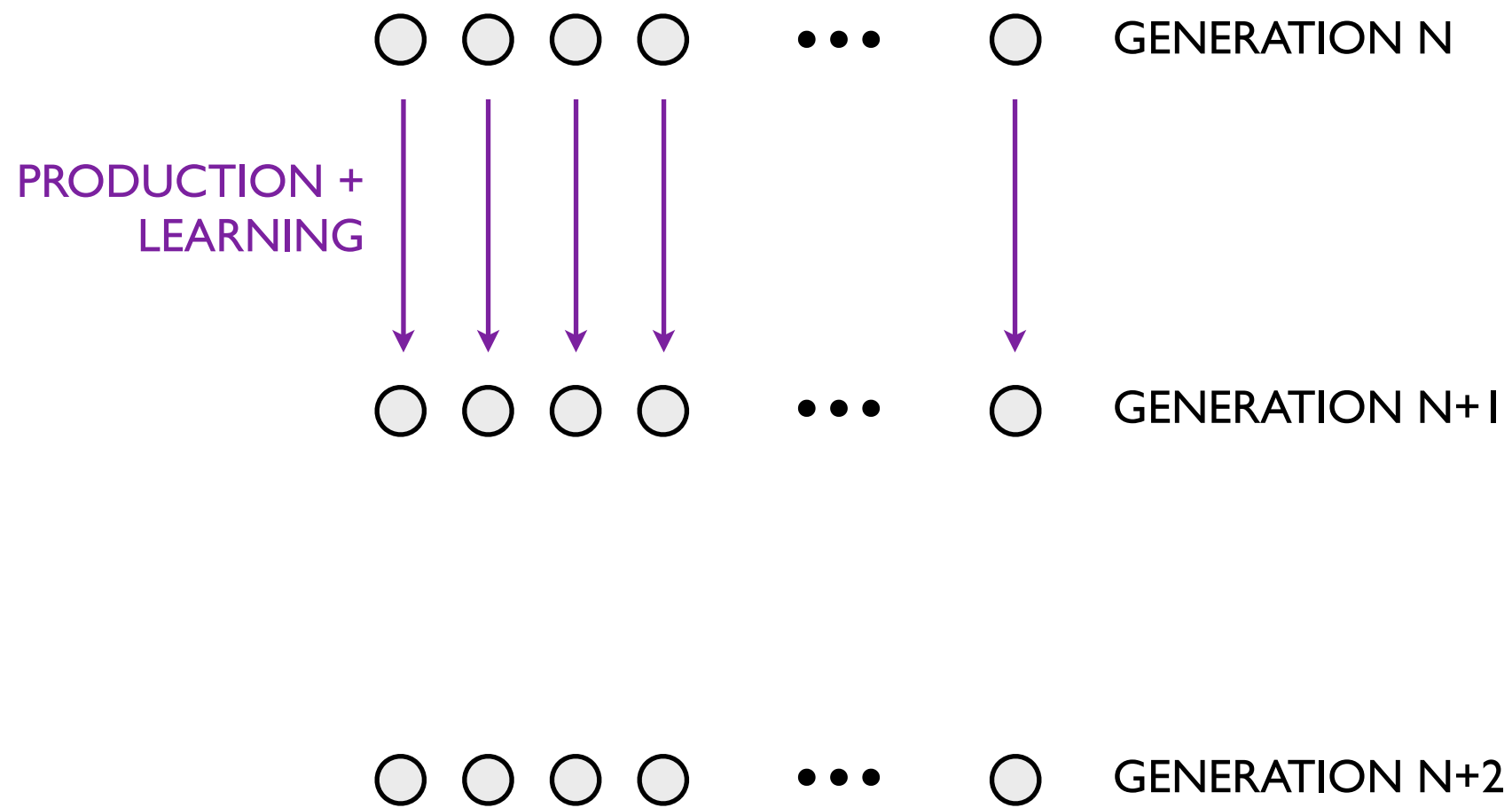
- A population is a series of generations, multiple individuals per generation
- Each agent learns a language from data produced by the previous generation
 - Single teacher, for now
- Prior encoded genetically
 - Initially: uninformative (neutral) prior
- Fitness determined by how closely your language matches the rest of your generation
- Fittest individuals pass on their genes to next generation, with some small probability of mutation
 - Evolving domain-specific priors, since they're really **for** language

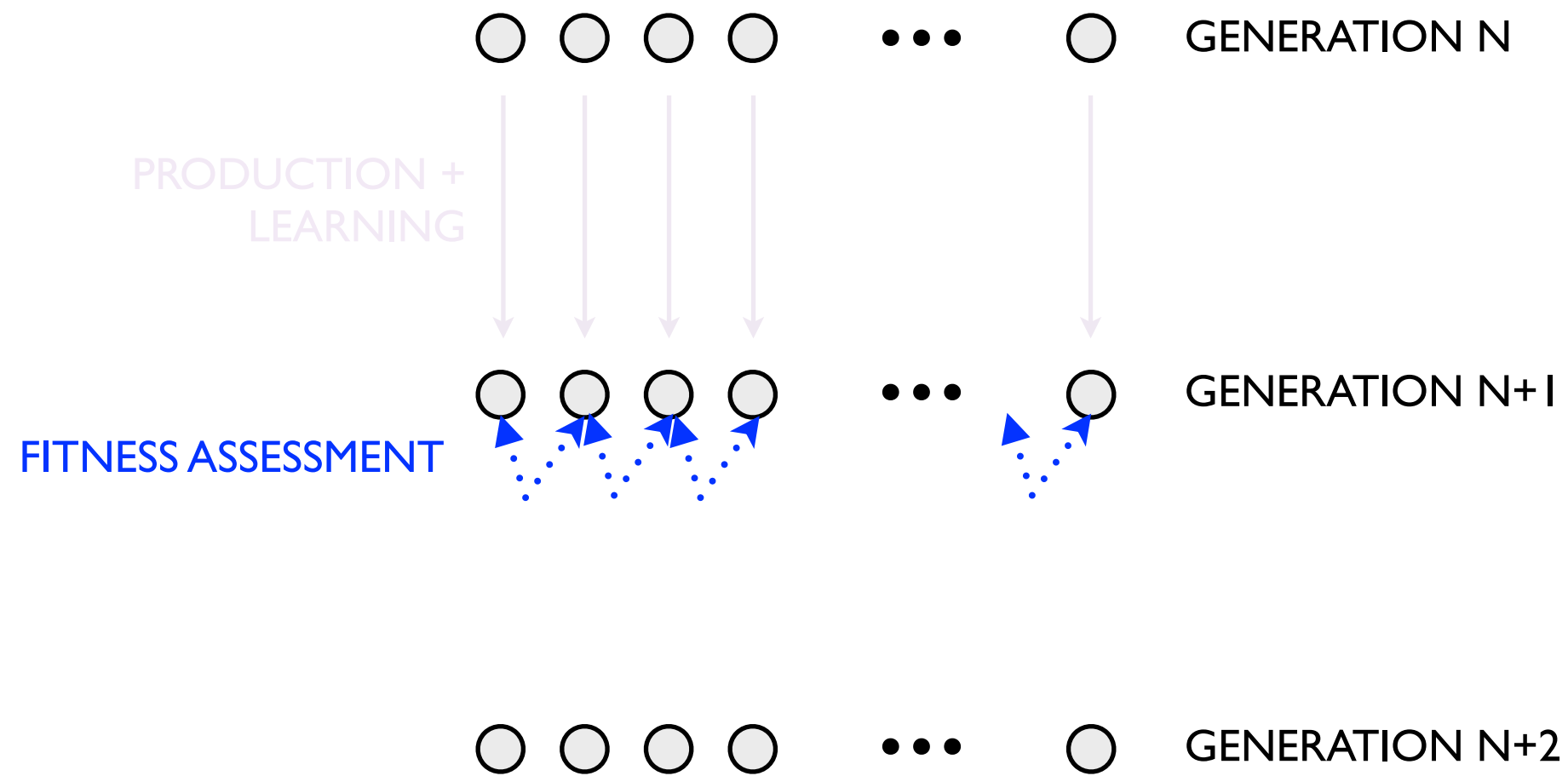
Details

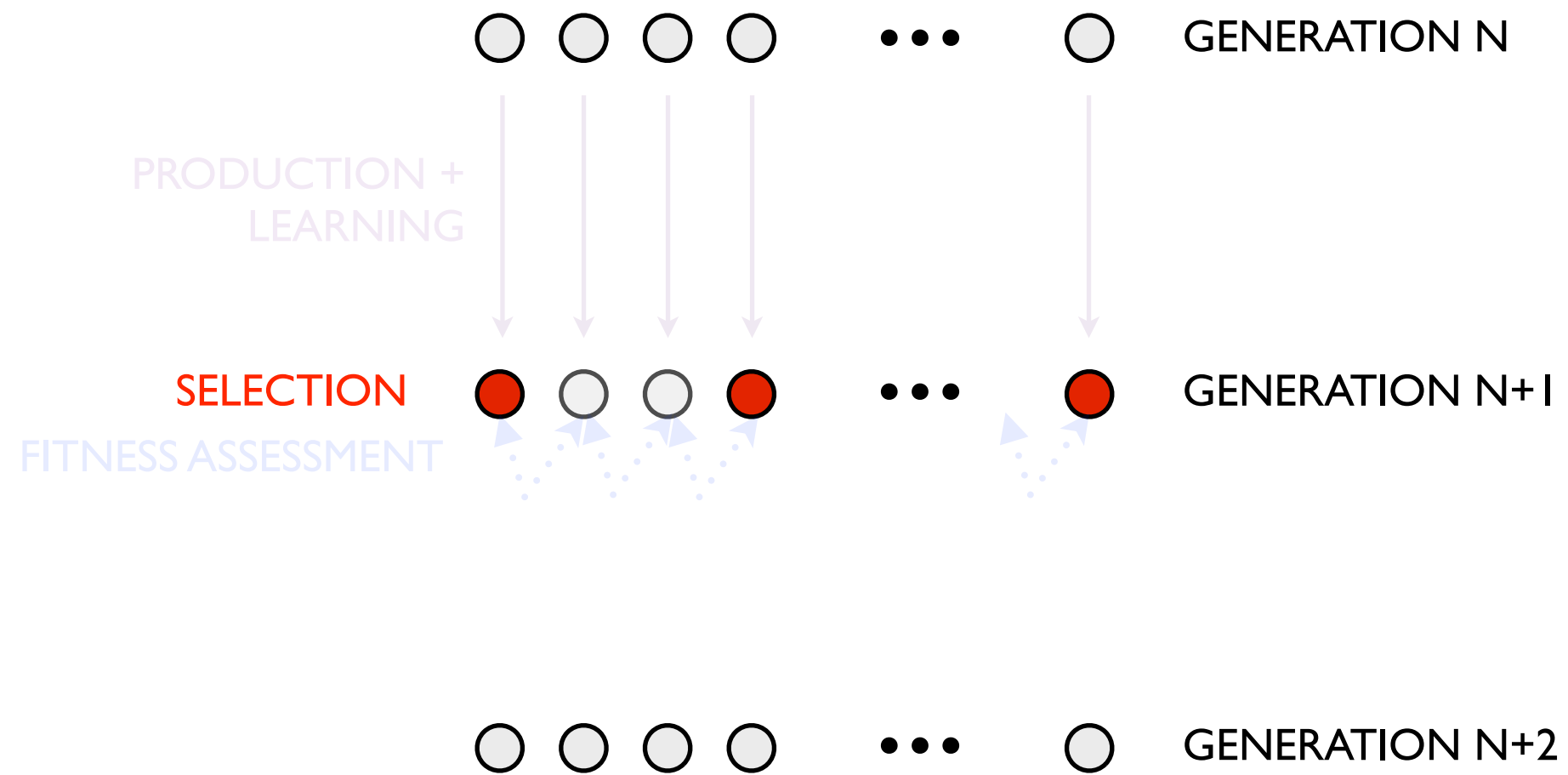
- The language model
 - Two possible languages, 0 and 1 (as per bayes1.py)
- The bias
 - > 0.5 , biased in favour of language 1
- Learning
 - MAP or sampling

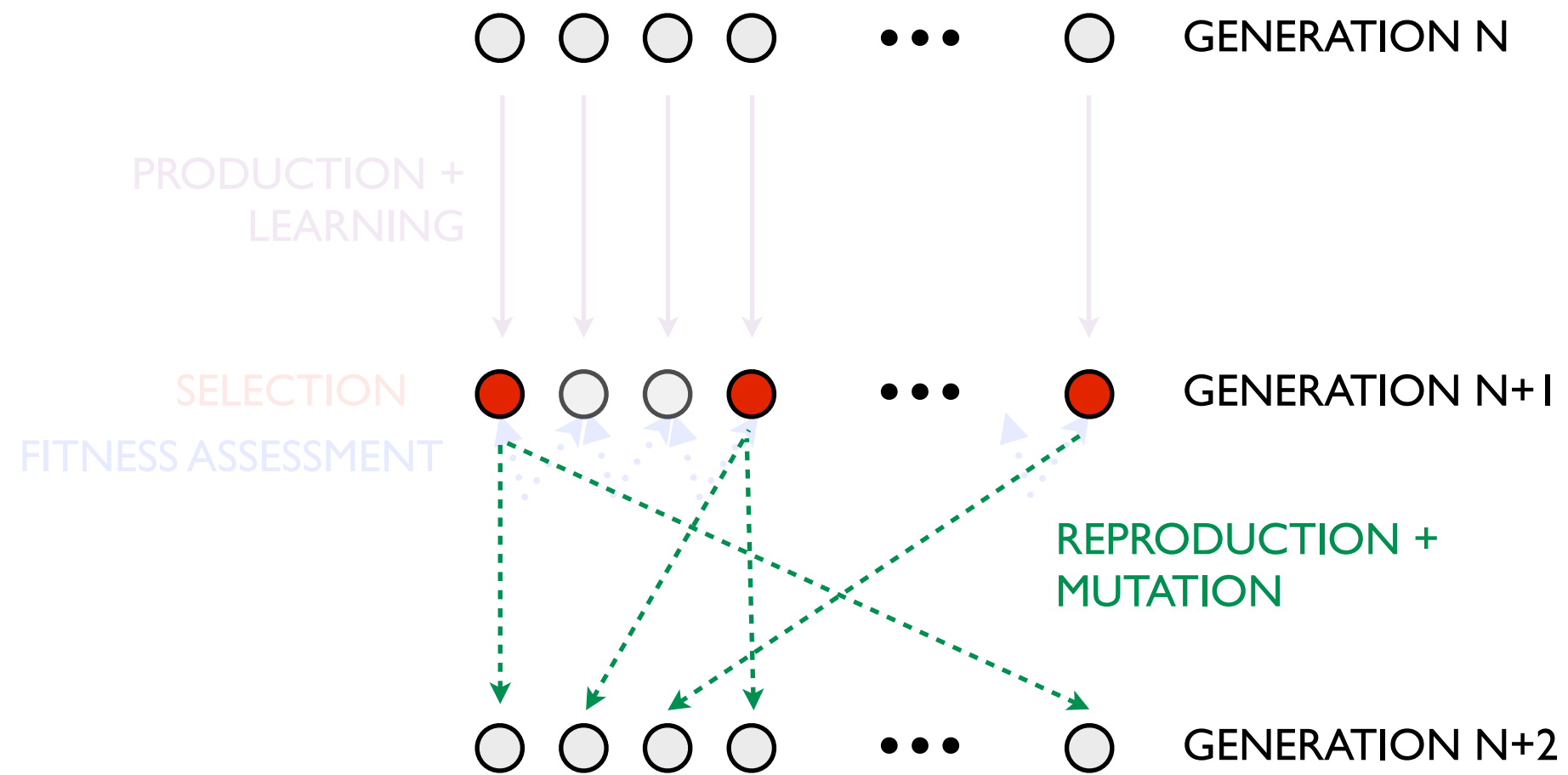
Genes for prior bias

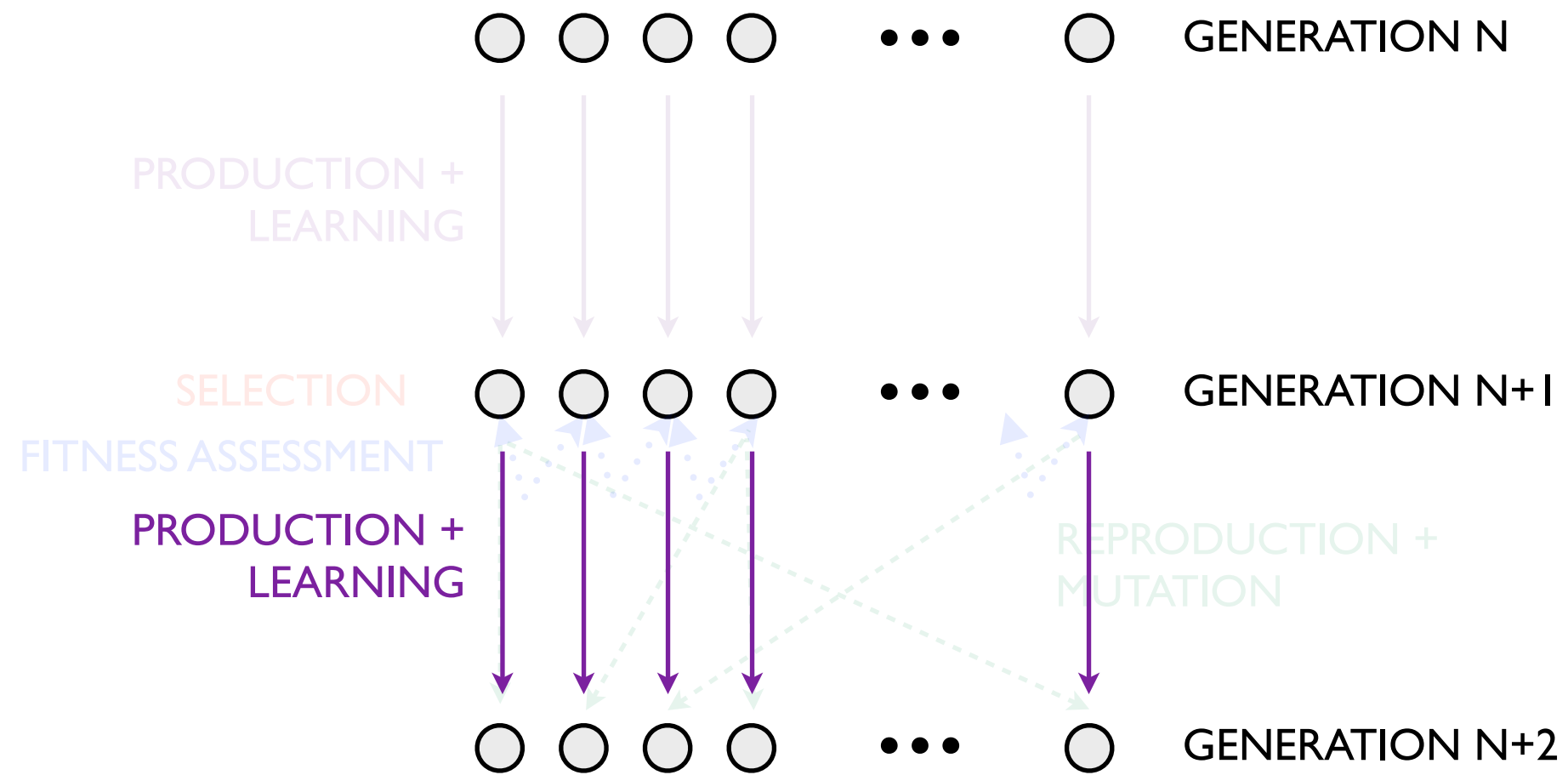
- How can we represent a genome that encodes a bias?
- One solution: additive coding of bias
 - Multiple genes
 - Each contribute a small amount to bias
 - $[0,0,0,0,0,0,0,0,0,0]$: bias = 0
 - $[1,1,1,1,1,1,1,1,1,1]$: bias = 1
 - $[1,1,1,0,0,1,0,0,1,0]$: bias = 0.5
- Any bias possible, but maintaining a strong bias against mutation requires selection for that bias

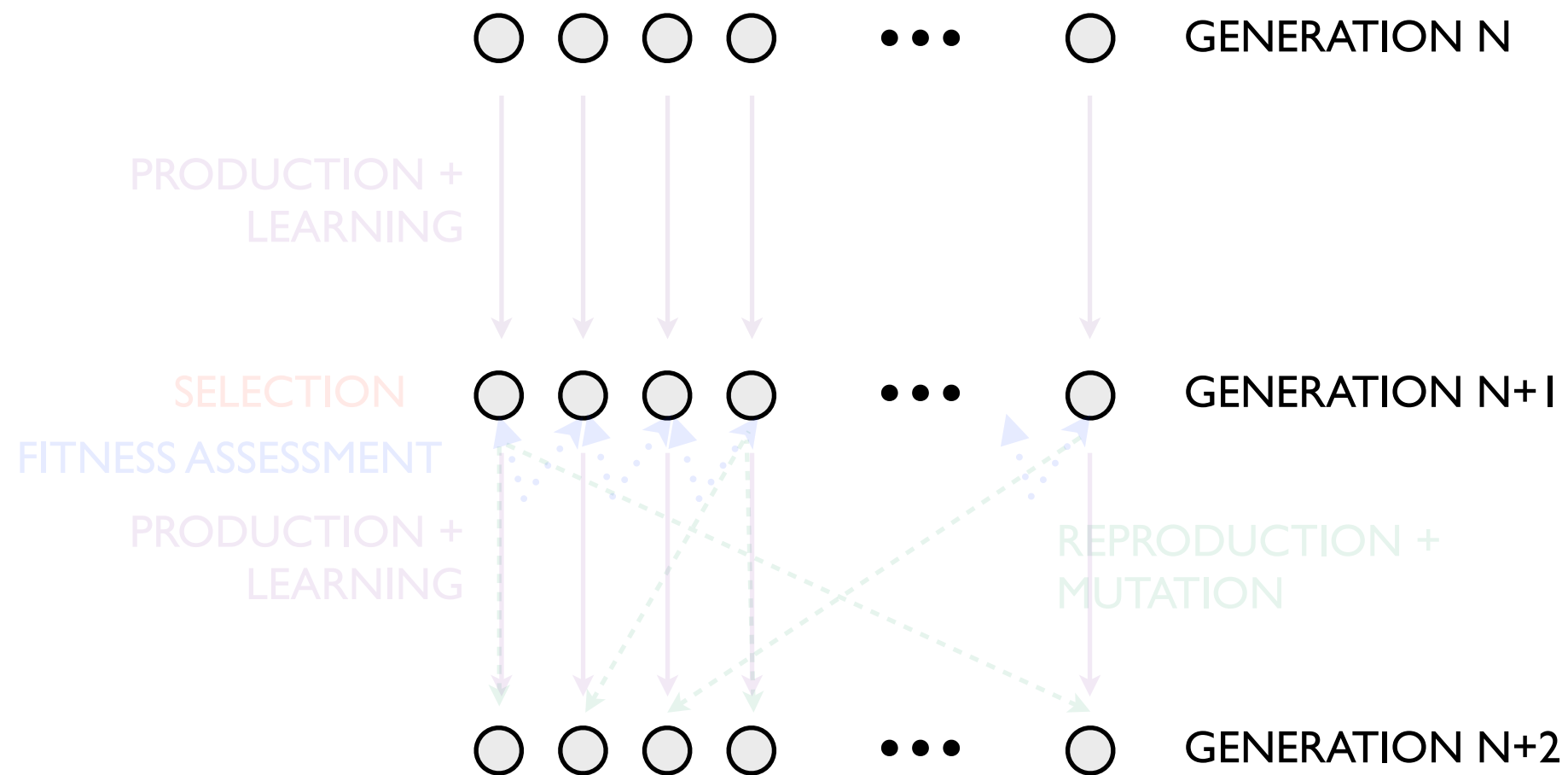












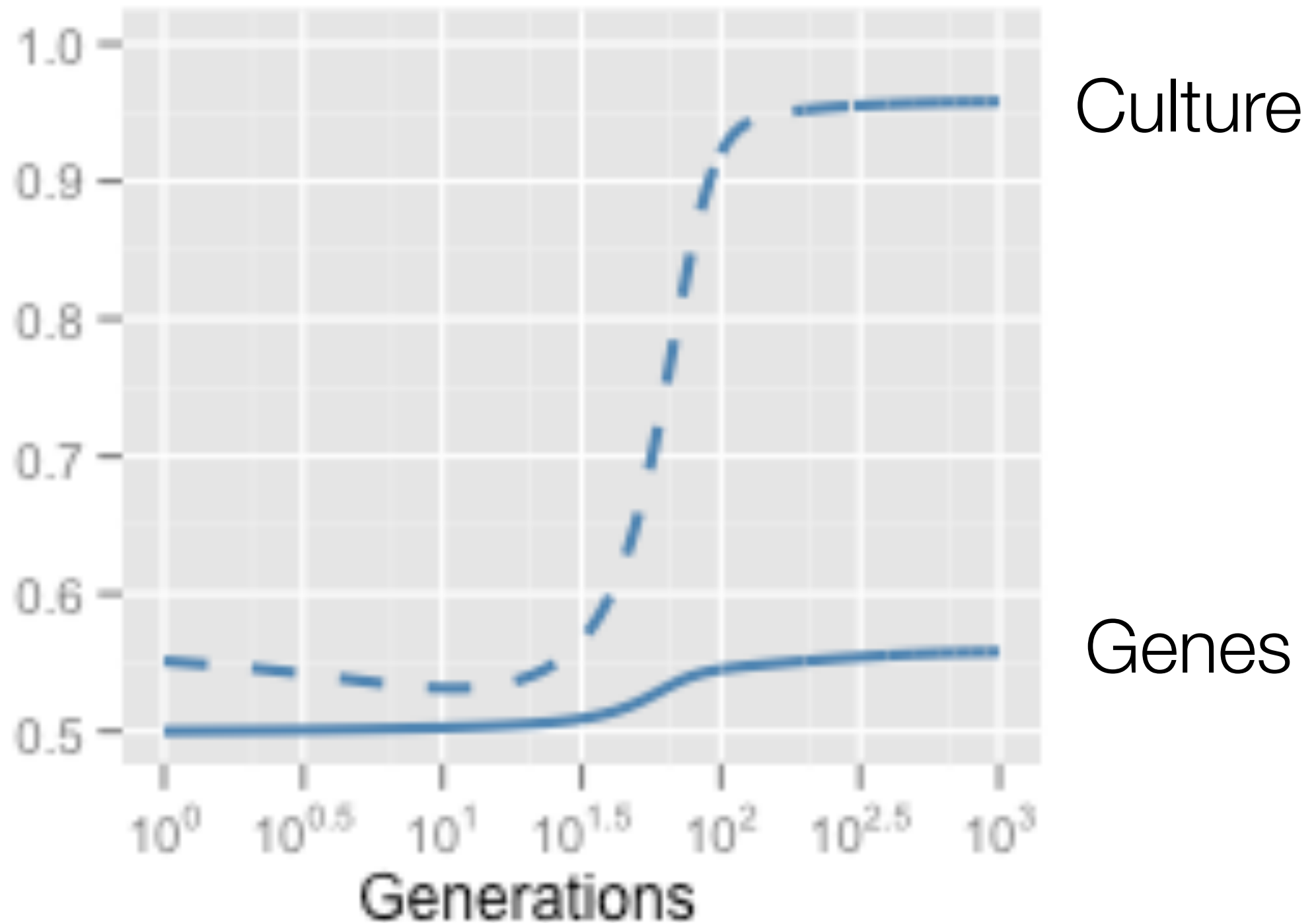
Note two kinds of inheritance - iterated learning and genetic transmission. Evolution due to all of: misconvergence in learning, natural selection, and mutation.

Before we run the model, can we work out what to expect?

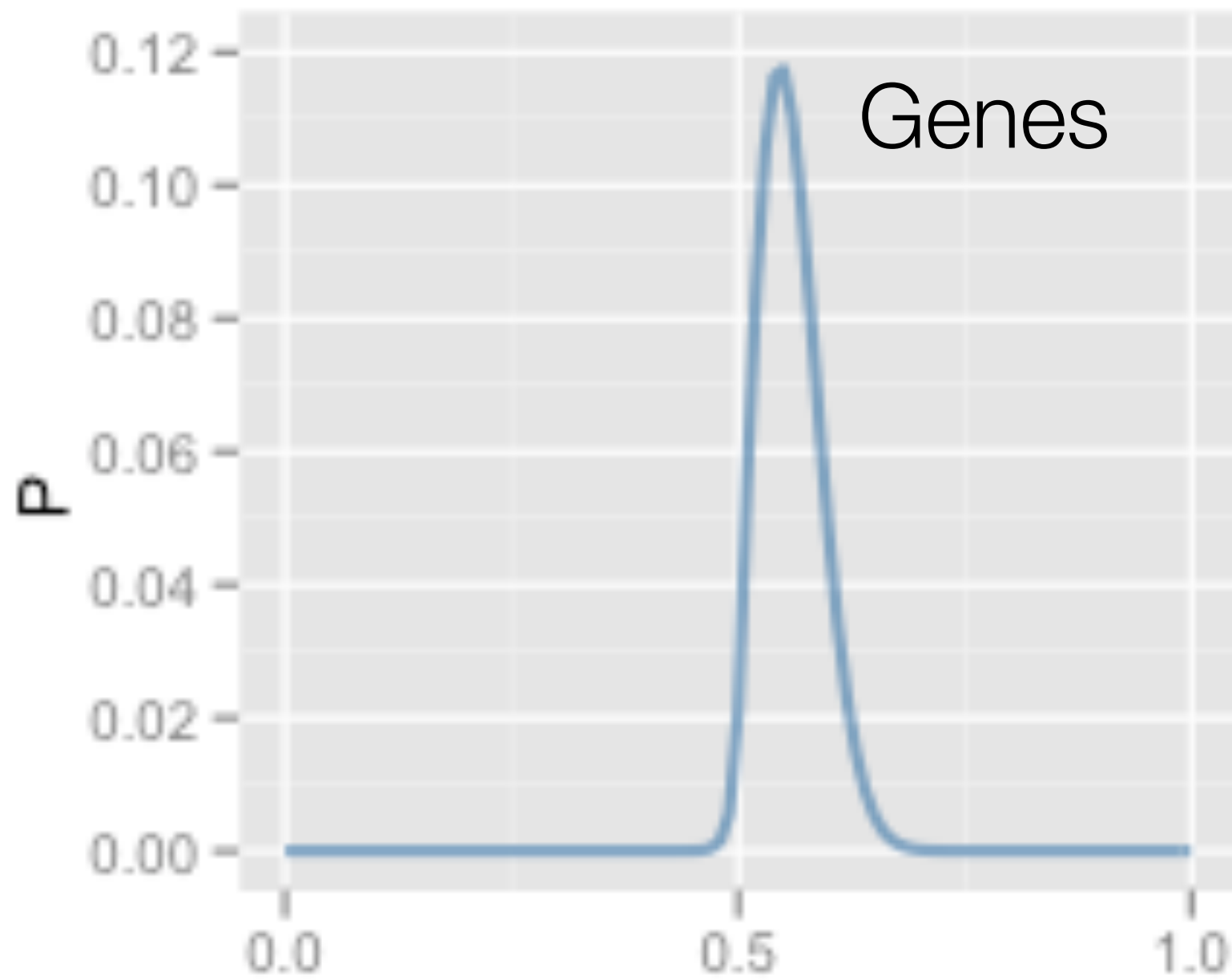
- When does it pay to have a bias in favour of a particular language?
- For samplers, when does one language become very common?
- For MAP learners, when does one language become very common?
- What's going to happen from our starting point of neutral priors?

N.B. From this point on, results are in preparation for publication. Please don't re-use or distribute any graphs without permission!

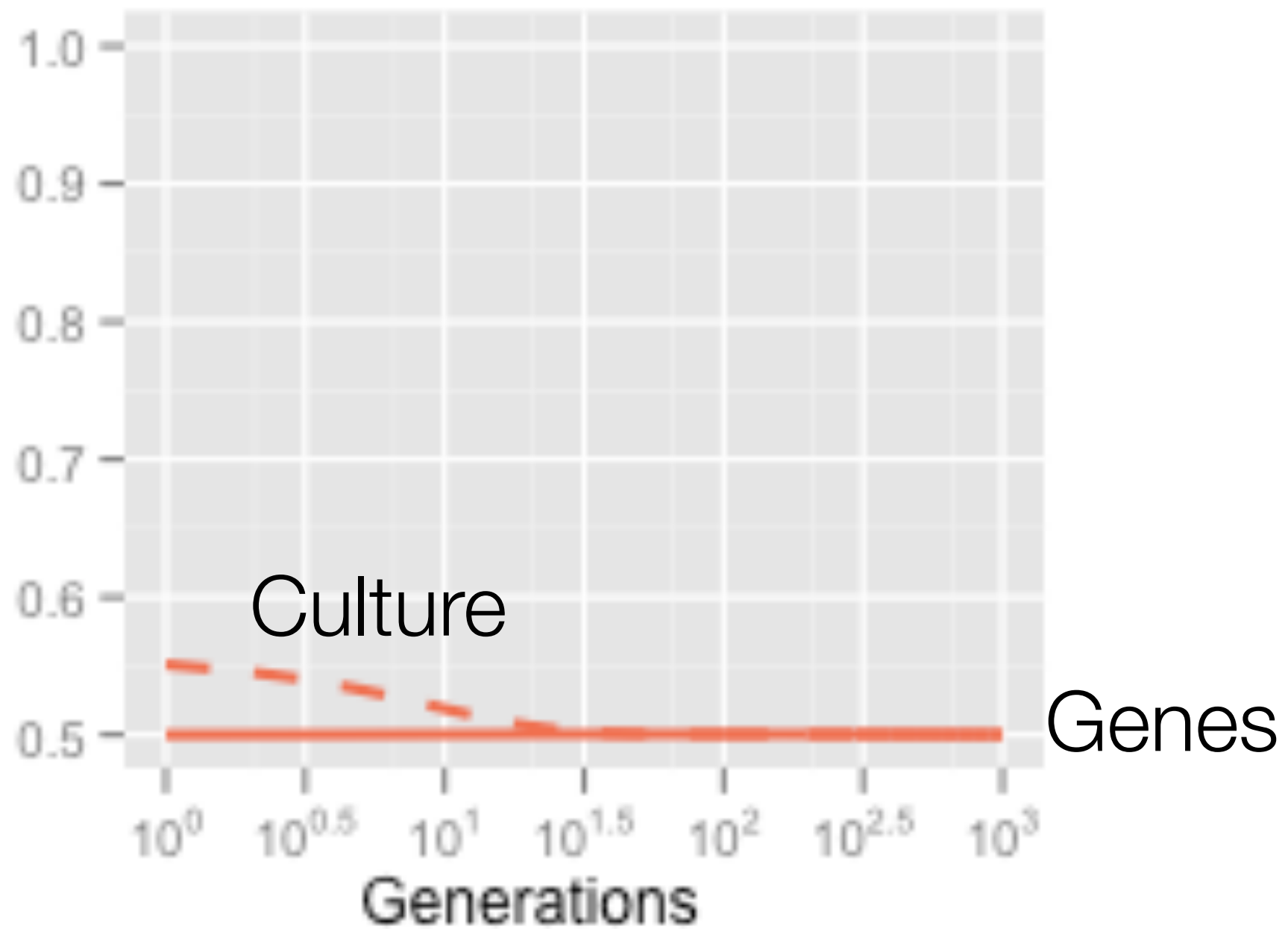
MAP coevolution



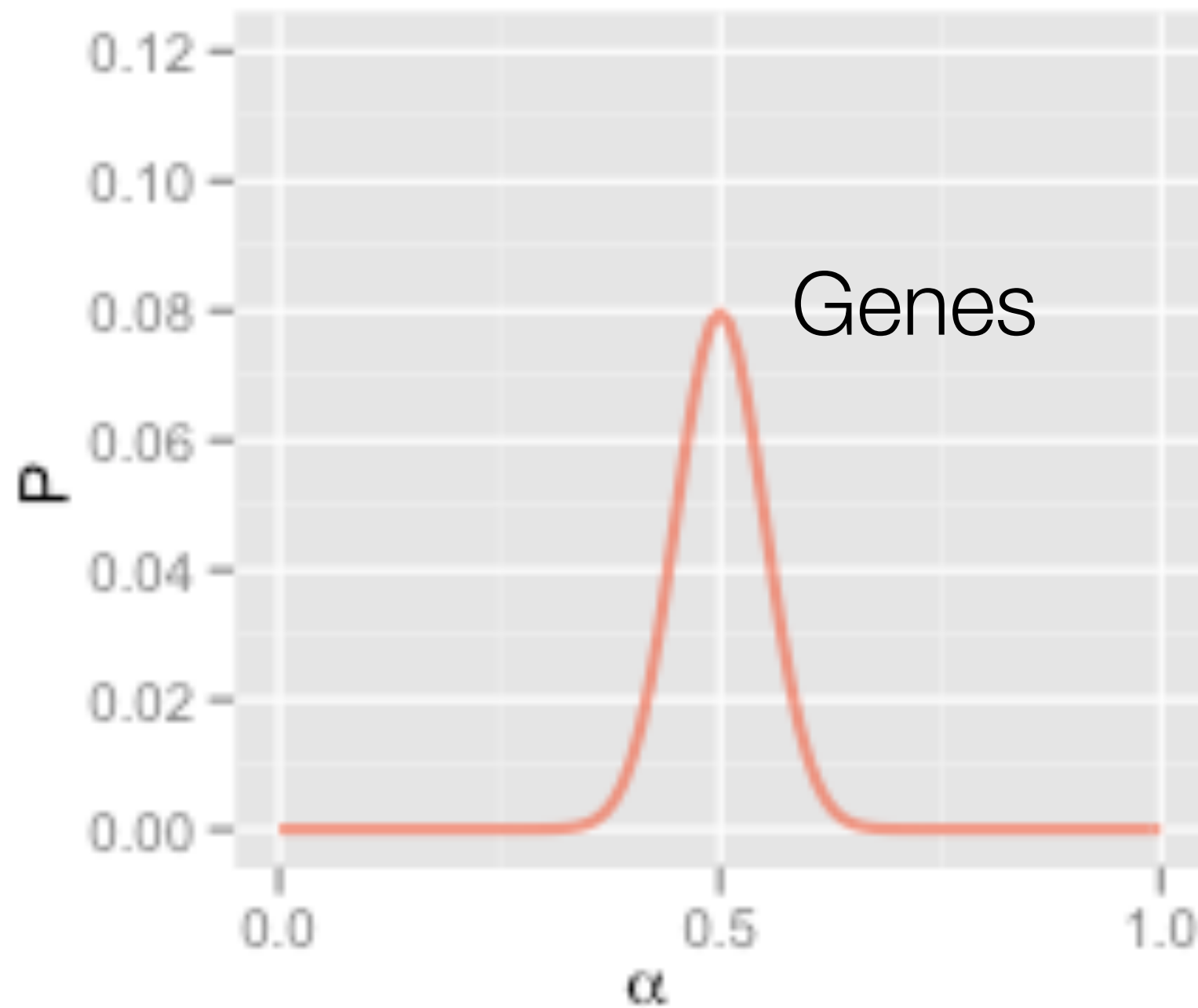
MAP coevolution



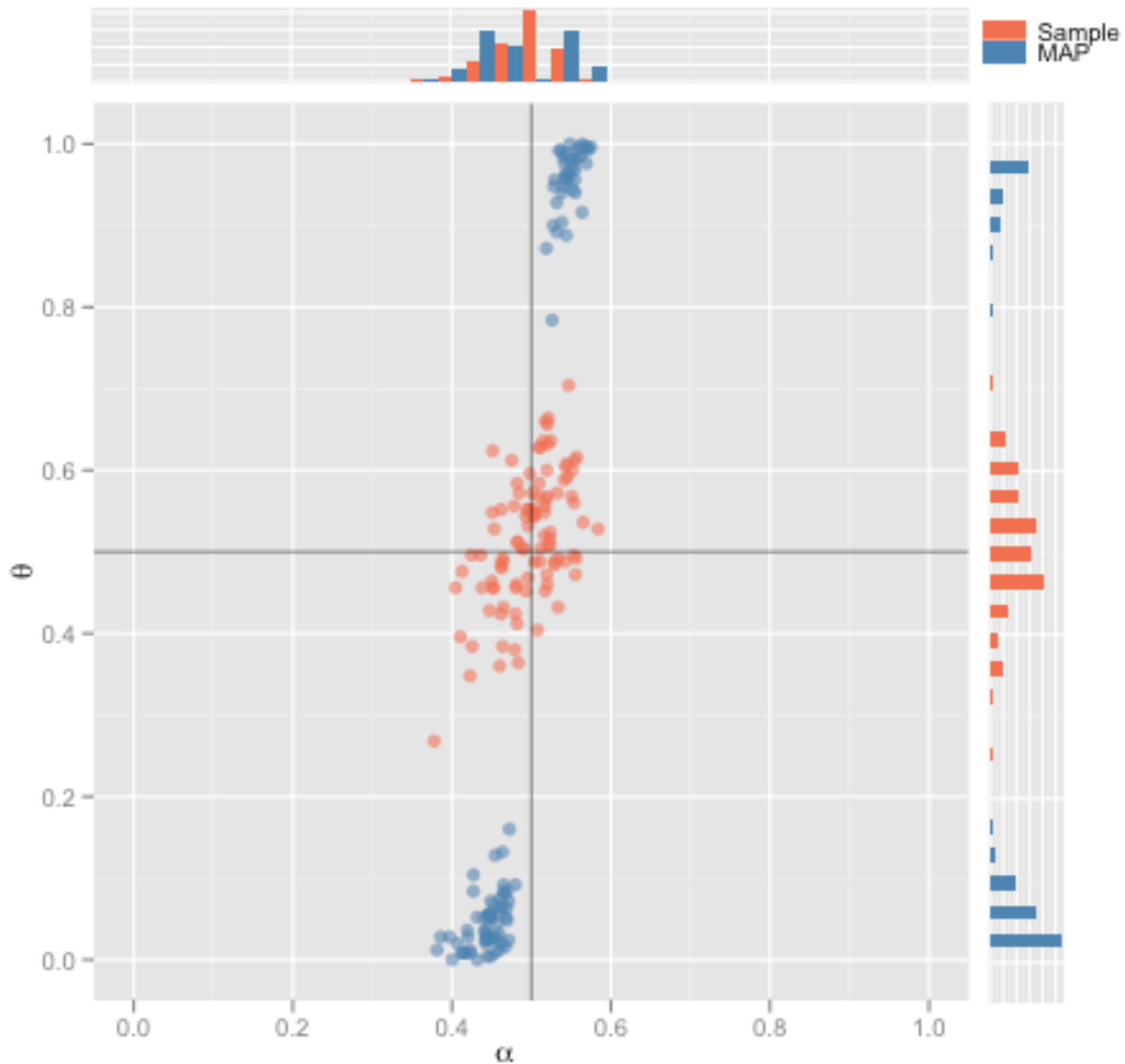
Sampler co-evolution



Sampler co-evolution



Culture

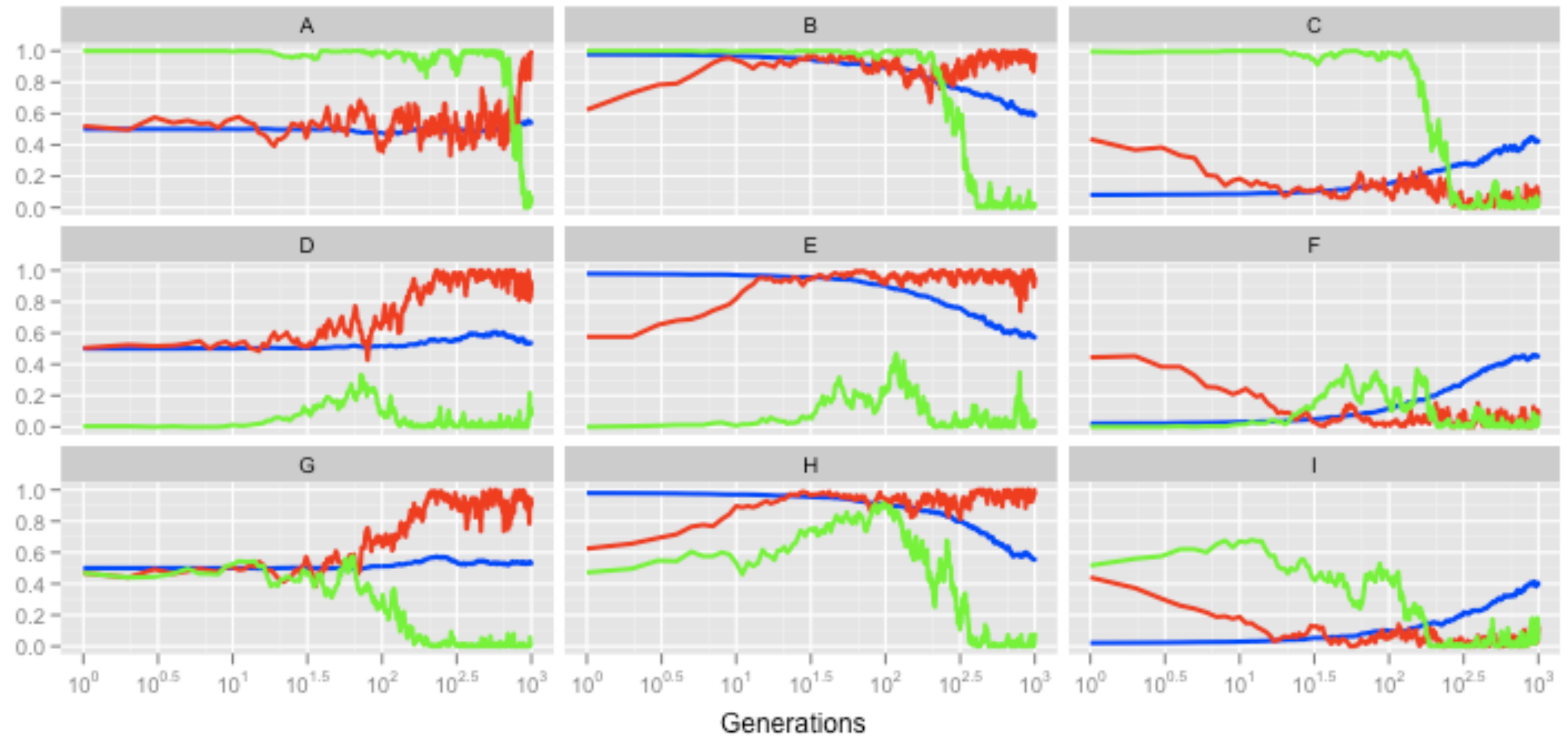


Genes

Sampling vs. MAP

- If you have to pick the same language as someone else trained on similar data to yourself, would you pick the MAP language or sample?
- Smith & Kirby (2008): MAP learning is **always** selected for over sampling, for coordination problems (Reading for this week!)
- Suggests that evolution might have given us a specialised strategy for learning coordinated tasks
 - We can imagine an evolutionary transition from sampling to MAP for language

Evolution of MAP



PROPORTION OF SAMPLERS (GENES)

BIAS (GENES)

LANGUAGE (CULTURE)

Why do we get these results?

- Think about evolution in terms of *masking* and *unmasking*
- MAP learning rapidly selected for.
- Subsequently:
 - Non-neutrality is *unmasked*
 - Bias strength is *masked*
- Tiny changes in the genes from the starting condition have big fitness effects
 - But there is no pressure to make these into strong constraints

Conclusions

- Recall: linguistic nativism proposes domain-specific strong constraints
- Model's predictions:
 - Samplers drift randomly leading to no strong constraints or universals (and sampling is selected against anyway)
 - MAP learners lead to domain specific biases that are *as weak as possible*
- If we do find a strong innate constraints in language learning, they are likely to have come from selection for something else (i.e. be domain-general)
- You can get *either* domain-specific weak biases, *or* domain-general strong biases
 - **But not linguistic nativism**