

On the limits of fitting complex models of population history to genetic data

Robert Maier^{1,*}, Pavel Flegontov^{1,2,*}, Olga Flegontova², Piya Changmai² and David Reich^{1,3,4,5,§}

¹ *Department of Human Evolutionary Biology, Harvard University, Cambridge, MA, USA*

² *Department of Biology and Ecology, Faculty of Science, University of Ostrava, Ostrava, Czechia*

³ *Howard Hughes Medical Institute, Harvard Medical School, Boston, MA, USA*

⁴ *Broad Institute of Harvard and MIT, Cambridge, MA, USA*

⁵ *Department of Genetics, Harvard Medical School, Boston, MA, USA*

* These authors contributed equally

§ corresponding author's email address: рмаier@broadinstitute.org, pavel.flegontov@osu.cz, reich@genetics.med.harvard.edu

Abstract

Our understanding of human population history in deep time has been assisted by fitting “admixture graphs” to data: models that specify the ordering of population splits and mixtures which is the only information needed to capture the patterns of allele frequency correlation among populations. Not needing to specify population size changes, split times, or whether admixture events were sudden or drawn out simplifies the space of models that need to be searched. However, the space of possible admixture graphs relating populations is vast and cannot be sampled fully, and thus most published studies have identified fitting admixture graphs through a manual process driven by prior hypotheses, leaving the vast majority of alternative models unexplored. Here, we develop a method for systematically searching the space of all admixture graphs that can incorporate non-genetic information in the form of topology constraints. We implement this *findGraphs* tool within a software package, *ADMIXTOOLS 2*, which is a reimplement of the *ADMIXTOOLS* software with new features and large performance gains. We apply this methodology to identify alternative models to admixture graphs that played key roles in eight published studies and find that graphs modeling more than six populations and two or three admixture events are often not unique, with many alternative models fitting nominally or significantly better than the published one. Our results suggest that strong claims about population history from admixture graphs should only be made when all well-fitting and temporally plausible models share common topological features. Our re-evaluation of published data also provides insight into the population histories of humans, dogs, and horses, identifying features that are stable across the models we explored, as well as scenarios of populations relationships that differ in important ways from models that have been highlighted in the literature, that fit the allele frequency correlation data, and that are not obviously wrong.

24 Introduction

25

26 Admixture graph models provide a powerful intellectual framework for describing the relationships
 27 among populations that allows not only branching of populations from a common ancestor but also
 28 mixture events. Admixture graphs can precisely summarize important features of population history
 29 without requiring specification of all parameters such as population sizes, split times, mixture times,
 30 and distinguishing between sudden splits or drawn-out separations. All these parameters describe
 31 important features of demographic history and are considered by several tools for fitting
 32 demographic models (Gutenkunst et al. 2009, Gronau et al. 2011, Schiffels et al. 2016, Flegontov et
 33 al. 2019, Kamm et al. 2019, Rogers 2019, Hubisz et al. 2020). However, the fact that it is possible to
 34 first infer important aspects of the topology (admixture graphs), and then fit these additional
 35 parameters, simplifies demographic inference (Patterson et al. 2012, Pickrell and Pritchard 2012,
 36 Lipson et al. 2013, Leppälä et al. 2017, Lipson 2020, Molloy et al. 2021, Yan et al. 2021). Admixture
 37 graphs thus serve both as conceptual frameworks that allow us to think about the relationships of
 38 populations deep in time, and as mathematical models we can fit to genetic data.

39

40 A challenge for fitting admixture graph models is that they are often not uniquely constrained by the
 41 data, with many providing equally good fits to the f_2 -, f_3 -, and f_4 -statistics used to constrain them
 42 within the limits of statistical resolution. Previously published methods for finding fitting admixture
 43 graphs were not well-equipped to handle the large range of equally well-fitting models for three
 44 reasons: (1) They did not reliably provide information on whether there is a uniquely fitting
 45 parsimonious model or alternatively whether there are many models that fit equally well to within
 46 the limits of statistical resolution, (2) they did not provide formal goodness-of-fit tests, and related
 47 to this, (3) they not provide tests for whether the difference between the fits of any two models is
 48 statistically significant. As a consequence and as we demonstrate in what follows, many published
 49 admixture graph models have been interpreted as providing more confidence than is merited about
 50 the extent to which genetic data allows us to disentangle ancestral relationships.

51

52 There are two main approaches to studying demographic history with admixture graphs.

53

54 The first approach is to identify admixture graphs automatically, either without human intervention
 55 or with guidance. It is possible in theory to exhaustively test all possible graphs for a given set of
 56 populations and pre-specified number of admixture events, as implemented, for example, in the
 57 *admixturegraph* R package (Leppälä et al. 2017). An exhaustive approach can provide a complete
 58 and unbiased view on the kinds of models that are consistent with the data for a specified level of
 59 parsimony (total number of admixture events allowed in the graph), but this approach is limited to
 60 small graphs (typically up to 6 groups, 2 admixture events) due to the rapid increase in the number
 61 of possible admixture graphs as the number of populations and admixture events grows. In addition,
 62 as we show in our discussion of case studies, the assumption of parsimony that needs to be made in
 63 order to use an exhaustive approach can lead to misleading models of population history because
 64 not including additional populations can blind users to additional mixture events that occurred (and
 65 whose existence is revealed by examining data from additional populations). Specifically, models
 66 with additional admixture events that are qualitatively profoundly different to the best fitting
 67 parsimonious graph and that capture the true history, will sometimes be completely missed when
 68 applying the parsimony assumption. Alternatively, the programs *TreeMix* (Pickrell and Pritchard
 69 2012, Molloy et al. 2021), *MixMapper* (Lipson et al. 2013), and *migoGraph* (Yan et al. 2021) all
 70 address the problem of how to rapidly explore the vast space of admixture graphs relating a set of
 71 populations by applying algorithmic ideas or heuristics; all of these methods speed up model search
 72 by orders of magnitude. The new method we introduce here, *findGraphs*, belongs to this class of
 73 algorithms. Algorithmic innovations in *findGraphs* enable us to get around some of the limitations
 74 associated with parsimony assumptions by more efficiently exploring a larger proportion of plausible
 75 of admixture graph space, and by increasing the speed of evaluation of individual graphs.

76

The second approach to fitting admixture graphs is to manually build them up by grafting additional populations onto simpler smaller graphs that fit the data. This approach involves stepwise addition of populations in an order that is chosen based on the best judgment of the user, and for each newly added population involves adding admixture events or tweaks in the graph until a fit is obtained; the user then moves on to adding the next population (see Reich et al. Nature 2009, Reich et al. AJHG 2011, Reich et al. Nature 2012, Lazaridis et al. 2014, Seguin-Orlando et al. 2014, Fu et al. 2016, Skoglund et al. 2016, Yang et al. 2017, McColl et al. 2018, Moreno-Mayar et al. 2018, Tambets et al. 2018, van de Loosdrecht et al. 2018, Flegontov et al. 2019, Sikora et al. 2019, Wang et al. 2019, Lipson et al. 2020, Shinde et al. 2019, Yang et al. 2020, Hajdinjak et al. 2021, Wang et al. 2021 for examples). The program *qpGraph* in the *ADMIXTOOLS* package (Patterson et al. 2012) has been the most common computational method used for testing fits of individual admixture graphs. Most published admixture graphs have been constructed manually, often acknowledging the existence of alternative models by presenting plausible models side-by-side, and this approach has been the basis for many claims about population history (Lazaridis et al. 2014, Yang et al. 2017, Posth et al. 2018, Sikora et al. 2019, Shinde et al. 2019, Bergström et al. 2020, Lipson et al. 2020, Hajdinjak et al. 2021, Wang et al. 2021). A strength of this approach is that it takes advantage of human judgment and outside knowledge about what graphs best fit the history of the human populations being analyzed. This external information is powerful as it can incorporate non-genetic evidence such as geographic plausibility and temporal ordering of populations or linguistic similarity, or other genetic data such as estimates of population split times or shared Y chromosomes or rejection of proposed scenarios based on joint analysis of much larger numbers of populations than can reasonably be analyzed within a single admixture graph. Thus, while manual approaches explore many orders of magnitude fewer topologies than automatic approaches often do, they still may provide inferences about population history that are more useful than those provided by automatic approaches. These methods' strength is also their weakness: by relying on intuition, following a manual approach has the potential to validate the biases users have as to what types of histories are most plausible (these may be the only types of histories that will be carefully explored). This can blind users to surprises: to profoundly different topologies that may correspond more closely to the true history, and we discuss examples of this in the Results section.

The *findGraphs* method combines the advantages of automated and manual topology exploration by allowing users to encode various sources of information as constraints on the space of admixture graphs, which is then explored automatically.

Regardless of the approach used to search through the space of possibly fitting admixture graphs, a challenge in the effort to find a uniquely well-fitting admixture graph is that it is difficult to quantify the absolute quality of the fit of a model, as well as the relative quality of the fits of multiple models. Performance gains relative to the original implementation of *qpGraph* allow us to address this problem by obtaining bootstrap confidence intervals and *p*-values for estimated parameters of single models, as well as for the difference in fit quality of two models. Existing methods for comparing admixture graph models (for example based on Akaike information criterion (AIC) or Bayesian information criterion (BIC), see Flegontov et al. 2019, Shinde et al. 2019) do not take the variability across SNPs into account appropriately since they do not rely on dataset resampling approaches such as bootstrap, and thus they tend to overestimate our ability to differentiate competing models. In combination with the previously described approach to automating the search of admixture graphs, this leads to a situation where we are able to find and test a large number of models, many of which fit equally well despite often having very different topological features.

The methods for automated graph topology inference and model comparison relying on bootstrap resampling described above are implemented in *ADMIXTOOLS 2*, a comprehensive platform for learning about population history from *f*-statistics. It is built to provide a stand-alone workspace for research in this area and is implemented as an R package. For all computations, *ADMIXTOOLS 2* exhibits large speedups relative to previously published platforms for *f*-statistic analysis (e.g. *popstats* and *ADMIXTOOLS* version 6.0 which we call "Classic *ADMIXTOOLS*" in what follows to

distinguish it from updated *ADMIXTOOLS* version 7.0.2 which implements some of the speed-up ideas also implemented in *ADMIXTOOLS* 2). This is achieved by deploying a series of algorithmic improvements, most notably storage of pre-computed f -statistics in random access memory, which avoids having to rely on reading in extremely large genotype matrices to perform most computations. In addition to the new algorithmic ideas allowing efficient searching through the space of admixture graphs and comparing the fits of two admixture graphs, *ADMIXTOOLS* 2 also provides a solution to the question of which parameters of an admixture graph are identifiable in the limit of infinite data. The most important algorithmic improvements presented in *ADMIXTOOLS* 2 and its philosophical differences relative to classic *ADMIXTOOLS* are described in the next section, while the full methodological details are presented in the Methods and Supplement.

Results

New ideas implemented in ADMIXTOOLS 2

We present a new implementation of the popular *ADMIXTOOLS* software (called “Classic *ADMIXTOOLS*”) (Patterson et al. 2012, Haak et al. 2015). Our implementation (*ADMIXTOOLS* 2, see documentation at <https://uqarmaie1.github.io/admixtools>) enhances performance by greatly reducing runtime and memory requirements across a wide range of different methods, relative to Classic *ADMIXTOOLS* (Figure 1a). We note that some of these improvements have now been implemented in version 7.0.2 of *ADMIXTOOLS*. The present study focuses not on the performance differences between Classic *ADMIXTOOLS* and *ADMIXTOOLS* 2, but on the description of new ideas implemented in one or both of these tools.

Computation and use of f-statistics

A key idea that facilitates the performance increases shared by *ADMIXTOOLS* 2 and *ADMIXTOOLS* v. 7.0.2 is that any f -statistic (which form the basis of almost all *ADMIXTOOLS* programs as well as other toolkits for studying population history such as *popstats*) can be computed from a small number of f_2 -statistics. For most f -statistic-based analyses (for example *qpWave*, *qpAdm*, and *qpGraph*; Figure 1c), the time required to process f -statistics is trivial compared to the time required to compute f -statistics from genotype data. These f_2 -statistics can be stored and re-used to compute f_3 - and f_4 -statistics, thus reducing the size of the input data, runtime, and memory requirements by orders of magnitude (Figure 1a, 1d).

Using precomputed f_2 -statistics is not always the best solution. In data sets with large amounts of missing data, computing f_3 - and f_4 -statistics from pre-computed f_2 -statistics may introduce bias. In this case, it is necessary to compute f_3 - and f_4 -statistics directly, using different SNPs for each f -statistic (all available SNPs in each population triplet or quadruplet). However, even without the use of pre-computed f_2 -statistics, *ADMIXTOOLS* 2 often achieves large performance gains (Figure 1a).

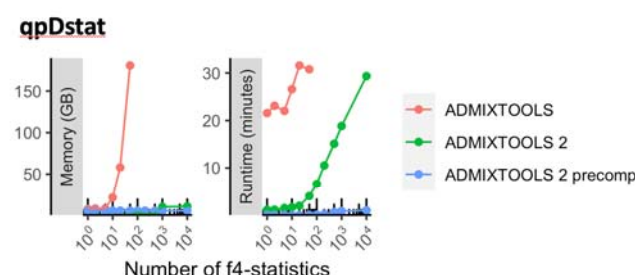
The program *qpfstats* in Classic *ADMIXTOOLS* implements an idea which strikes a balance between these two extremes. It increases the accuracy of estimation of f -statistics by using a regression approach to jointly estimate the values of all f_2 -, f_3 - and f_4 -statistics relating a set of populations, taking advantage of the algebraic relationships of the expected values of these statistics. This approach makes it possible to obtain more precise estimates of the values of these statistics than can be obtained by inferring them only using SNPs that are covered in each of the populations being compared. This feature is available in *ADMIXTOOLS* 2 through the *qpfstats* option in the *extract_f2* function.

Another improvement introduced in *ADMIXTOOLS* 2 relates to accurate evaluation of the match between observed and expected f_3 -statistics when fitting admixture graphs where at least one

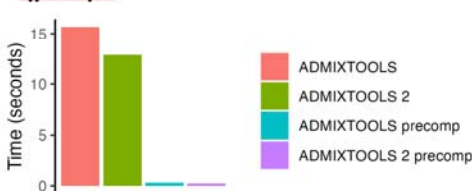
population is represented by a single individual with genotypes derived from randomly selected sequencing reads (“pseudo-diploid” or “pseudo-haploid” data). f_3 -statistic computations need to be modified when analyzing pseudo-diploid data, because heterozygosity cannot be computed using comparisons of sequences within the same individual; however, computation of heterozygosity is essential to calculate “admixture” f_3 -statistics where negative values provide proof of the mixed nature of the target population. When a population is represented by multiple individuals, unbiased estimation of admixture f_3 -statistics can be carried out even for pseudo-diploid data by analyzing positions covered by sequences from at least two individuals and only computing variation rates across individuals. This approach is implemented in Classic *ADMIXTOOLS* with the “inbreed: YES” option. However, no admixture f_3 -statistic can be computed when the “inbreed: YES” option is turned on and the target population is represented by a single individual (as no variation across individuals within a population can be detected in this case). Classic *ADMIXTOOLS* deals with this case by failing to run if any population in an analysis is represented by a single individual and the “inbreed: YES” option is turned on. Because the datasets from all the admixture graphs revisited here included at least one population represented by a single individual (**Table S1**), the “inbreed: YES” option could not be used in the original studies (the program failed with this option, by design). Thus, admixture graph fitting in those studies relied on the incorrect algorithm for calculating f_3 -statistics (except for Librado *et al.* 2021, who used *TreeMix* instead of *qpGraph*) and, as a result, some f_3 -statistics that are negative and could provide important constraints for admixture graph fitting were evaluated as positive. These concerns are relevant for the Shinde *et al.* (2019), Lipson *et al.* (2020), and Wang *et al.* (2021) studies we revisit below (see Table S1 for datasets where negative f_3 -statistics were encountered). To be able to detect negative f_3 -statistics and thus take advantage of their power for constraining the space of possibly fitting historical models, in *ADMIXTOOLS 2* we introduced an option which makes it possible to compute negative f_3 -statistics on pseudo-diploid data, at a cost of removing sites with only one chromosome genotyped in any population that is represented by at least two individuals (so that it is possible in theory to compute heterozygosity in these populations). Admixture f_3 -statistics continue to be incorrectly computed using *ADMIXTOOLS 2* for targets that are singleton populations represented by pseudo-diploid data, as there is no avoiding this particular problem. See Methods for a description of the new algorithm for calculating f_3 -statistics.

Figure 1

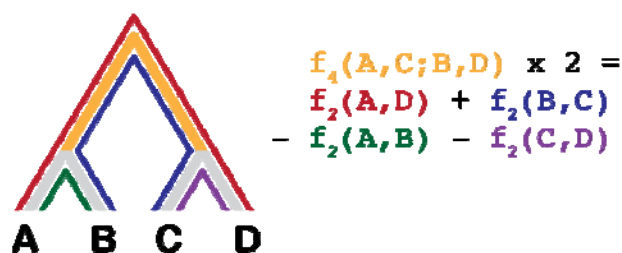
a



qpGraph



b



c

Program	f-statistic	Primary use
qp3pop	f_3	Test for admixture
qpDstat	f_4	Test for admixture
qpF4ratio	f_4	Estimating admixture proportions
qpWave	f_4	Finding how many gene flows connect two sets of populations
qpAdm	f_4	Estimating admixture proportions
qpGraph	f_3	Fitting admixture graph models of history

d

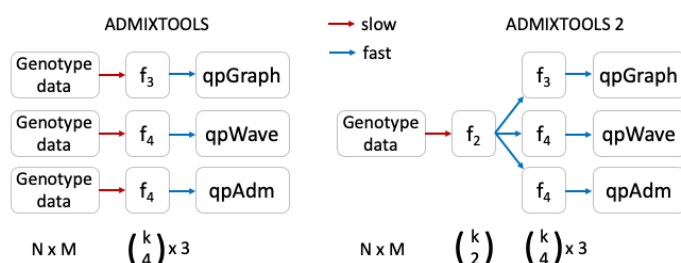


Figure 1:

a Performance comparison of f -statistic computation and admixture graph fitting. Top: Memory usage and runtime for computing f -statistics using (1) the *qpDstat* program in *ADMIXTOOLS* v7.0.2 released in 06/2021 (2) the function in *ADMIXTOOLS* 2 without precomputing f -statistics, and (3) the function in *ADMIXTOOLS* 2 with precomputed f -statistics. (1) and (2) give identical results, whereas (3) only gives identical results in the absence of missing data, which limits usefulness beyond a moderate number of populations. Bottom: Runtime comparison of *qpGraph* with and without precomputed f -statistics.

b Illustration of f_2 - and f_4 -statistics. f_2 measures the amount of drift separating any two populations, while f_4 measures the amount of drift shared between two population pairs. Every f_4 -statistic is a linear combination of four f_2 -statistics.

c Overview of the major *ADMIXTOOLS* programs, their primary use cases, and their associated f -statistics.

d Schematic representation of the computations behind the *ADMIXTOOLS* programs *qpGraph*, *qpWave*, and *qpAdm*. *ADMIXTOOLS* 2 separates the computation of f -statistics from the later steps in the pipeline. Shown below are the number of data points for individuals, SNPs, and populations. The exact number of all possible non-redundant f_2 , f_3 , and f_4 -statistics for k populations are $\binom{k}{2}$, $\binom{k}{3}$, and $\binom{k}{4}$. A small number of f_2 -statistics can be used to obtain a much larger number of f_3 - and f_4 -statistics and requires much less space than the raw genotype data.

Admixture graph fitting, model comparison, and interpretation

There are several challenges that arise when modeling the ancestral relationships among populations with admixture graphs, and *ADMIXTOOLS* 2 implements solutions to several key problems that were not adequately addressed with previous approaches: (1) Automated admixture graph inference; (2) Estimating confidence intervals for admixture graph parameters; (3) Comparing fits of different admixture graphs; (4) Determining identifiability of admixture graph parameters; and (5) Drawing conclusions from a large number of fitting graphs.

Here, we describe these challenges, how we address them, and how our approaches compare to other approaches, while the Methods section gives detailed descriptions.

(1) Automated admixture graph inference

Constructing admixture graphs manually runs the risk of overlooking models that challenge conventional hypotheses. On the other hand, current methods for inferring admixture graphs automatically (Leppälä et al. 2017, Molloy et al. 2021, Pickrell and Pritchard 2012, Yan et al. 2021) do not allow external information to be integrated into the analysis, and often result in models that may fit the genetic data but can be rejected on other grounds. In addition, *TreeMix* (Pickrell and Pritchard 2012), as well as *OrientAGraph* (Molloy et al. 2021), an improved version of *TreeMix*, can miss admixture graph topologies that exist on parts of the non-convex likelihood surface that are bypassed by these algorithms for exploring admixture graphs (for example, topology M7 in Figure 4 of (Molloy et al. 2021)). *MixMapper* (Lipson et al. 2013) and *miqoGraph* (Yan et al. 2021) have a different limitation: exploring topologies with more than one admixture event in the history of any group is not possible. Due to these limitations, many published findings are based on manual proposal of topologies and evaluation of fit, and the great majority of studies using this manual approach (see, for example, Reich et al. 2011, Reich et al. 2012, Lazaridis et al. 2014, Fu et al. 2016, Skoglund et al. 2016, Yang et al. 2017, McColl et al. 2018, Moreno-Mayar et al. 2018, Tambets et al. 2018, van de Loosdrecht et al. 2018, Flegontov et al. 2019, Sikora et al. 2019, Wang et al. 2019, Lipson et al. 2020, Shinde et al. 2019, Yang et al. 2020, Hajdinjak et al. 2021, Wang et al. 2021) rely on the software *qpGraph*. We introduce an approach for finding well-fitting admixture graphs automatically that can integrate external information, and that recovers graph topologies more accurately than *TreeMix* (Figure 2). External information can be integrated by specifying a set of constraints that admixture graphs must satisfy. This not only ensures that resulting models are temporally plausible, but also cleanly separates prior assumptions from the independent constraints provided by genetic data. Our strategy implemented in the function “*findGraphs*”, differs from *TreeMix/OrientAGraph* in several deep ways, most notably in that it optimizes graphs directly, rather than optimizing trees first and adding admixture events later. This makes it less prone to getting stuck in local optima: our simulation results show that *findGraphs* is more accurate for random graphs (Figure 2), and that it can recover specific topologies that pose problems for *TreeMix* and *OrientAGraph*.

(2) Estimating confidence intervals for admixture graph parameters

Since our new implementation of *qpGraph* can evaluate models much more rapidly, it becomes feasible to evaluate the same model multiple times on different SNP sets. This allows us to derive bootstrap confidence intervals (Boos 2003) for all parameters estimated by *qpGraph*, including drift lengths, admixture weights, log-likelihood (LL) scores, and f_4 -statistic residuals. Parameters with extremely wide confidence intervals can thus be immediately shown to be poorly determined. It should also be noted that the estimated confidence intervals do not take into account uncertainty about the graph topology.

(3) Comparing the fits of different admixture graphs

Using the bootstrap method for evaluating a graph multiple times on different SNP sets not only allows us to obtain confidence intervals for single graphs, but also allows us to test whether the fit of one graph is significantly better than the fit of another graph, by obtaining confidence intervals for the log-likelihood score difference or worst-residual (WR) difference of two graphs. When we apply this approach to a range of data sets, we find that models with modest log-likelihood differences are often not distinguishable after accounting for the variability across SNPs, even if one might expect them to be distinguishable based on the magnitude of the likelihood difference (Figure 3a,b). Thus, previous methods relying on AIC or BIC (such as Shinde et al. 2019, Flegontov et al. 2019) that used specified likelihood difference thresholds to reject some models over others, were over-aggressive.

a

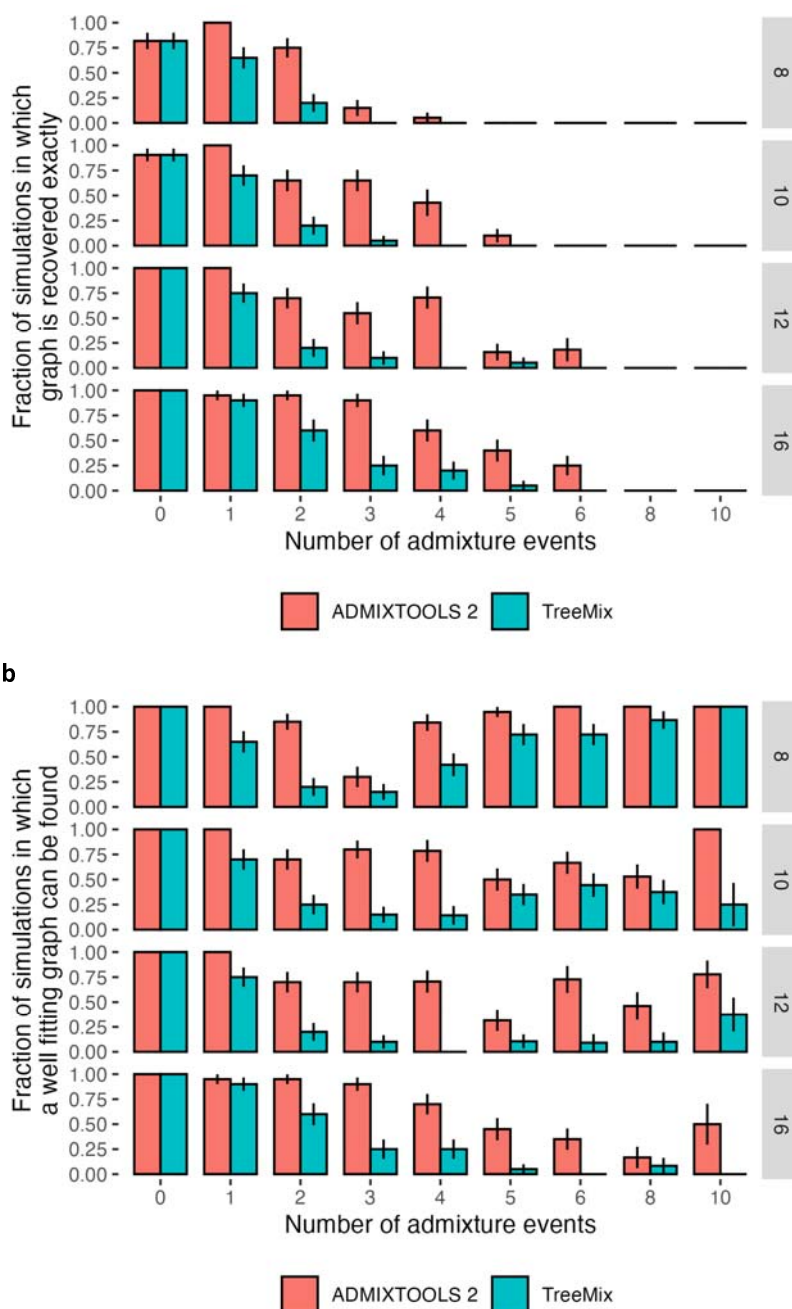


Figure 2:

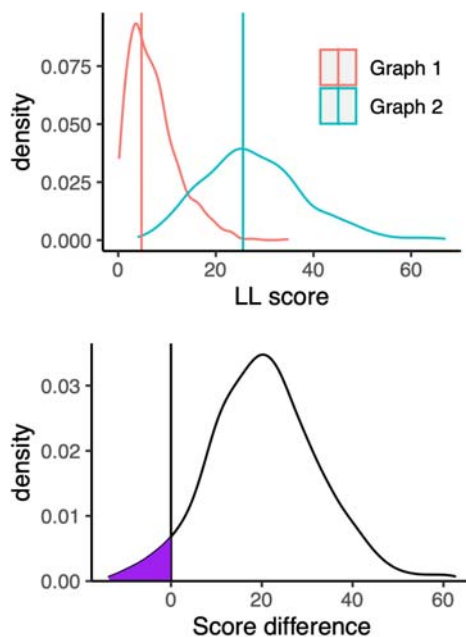
Comparison of accuracy of automated search for optimal topology in the *findGraphs* function of *ADMIXTOOLS* 2 and *TreeMix* using simulated graphs with 8, 10, 12, and 16 populations, and 0 to 10 admixture events. Error bars show standard errors calculated as $SE^2 = p(1 - p)/n$ where p is the fraction on the y-axis and n is the number of simulations in each group (typically 20). In the case of *ADMIXTOOLS* 2, we applied *findGraphs* three times on each simulated data set and picked a result with the best fit score. More details are provided in the Methods section. **a** Fraction of simulations where the simulated graph is recovered exactly. **b** Fraction of simulations where the simulated graph is either recovered exactly, or the score is at least as good as the score of the simulated graph, when both graphs are evaluated by *ADMIXTOOLS* 2. More admixture edges greatly increase the search space and make it more difficult to recover the simulated graph, but they do make it easier to find alternative graphs with good fits.

A second challenge in comparing different admixture graph models arises when comparing models of different complexity (i.e., with a different number of admixture events). Established methods such as AIC and BIC are applicable and also can account for different model complexity if the

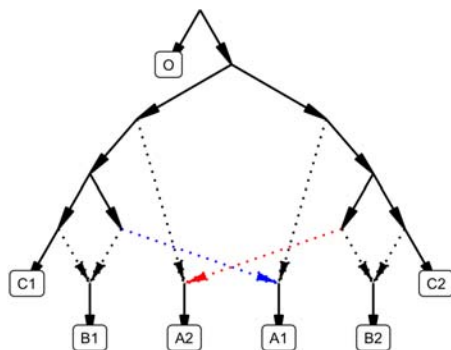
number of independent parameters in a model is known. However, the number of independent parameters estimated in an admixture graph is not simply determined by the number of groups, drift edges, or admixture events, as it also depends on the graph topology in a complex way. We implement a method to compare admixture graph models of different complexity by using a new scoring function, which uses different blocks of SNPs for deriving fitted and estimated f -statistics. This ensures that our model comparison test does not favor more complex models by allowing them to overfit the data. This cross-validation approach can also be used to rank alternative models of the same complexity and deal with overfitting. We note that the calculation of cross-validated log-likelihood scores is not turned on in *findGraphs* by default, and to make our results more comparable to those of the published studies we revisited, we relied on standard log-likelihood scores below.

To test if our method is well calibrated, we simulated 100,000 SNPs under the same graph in 1000 replicates. We then created two new topologies by removing one out of two symmetric edges from the first graph (Figure 3c). These new incorrect models are symmetrically related to the first graph and can be used to test the null hypothesis that the true difference in log-likelihood scores of these two graphs is zero. The uniform distribution of p -values confirms that our method is well calibrated (Figure 3d). A caveat is that only one symmetric topology was explored in this way.

a, b



c



d

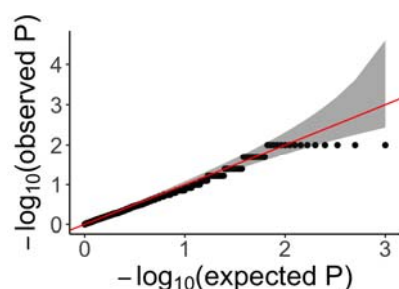


Figure 3: Calibrating the bootstrap model comparison approach

a Bootstrap sampling distributions of the log-likelihood scores for two admixture graphs (shown in **Figure S1**) of the same populations fitted using real data. Vertical lines show the log-likelihood scores computed on all SNP blocks.

b Distribution of differences of the bootstrap log-likelihood scores for both graphs (same data as in **a**). The purple area shows the proportion of resamplings in which the first graph has a higher score than the second graph. The two-sided p -value for the null hypothesis of no difference is equivalent to twice that area (or one over the number of bootstrap iterations if all values fall on one side of zero). In this case it is 0.078.

c The admixture graph which was used to evaluate our method for testing the significance of the difference of two graph fits on simulated data. We simulated under the full graph and fitted two graphs that result from deleting either the red admixture edge or the blue admixture edge. These two graphs have the same expected *qpGraph* fit score but can have different scores in any one simulation iteration.

d QQ plot of p -values testing for a score difference between the two graphs (on simulated data) under the null hypothesis of no difference, confirming that the method is well calibrated.

(4) Determining identifiability of admixture graph parameters

Fitting admixture graphs results in an estimate of the overall model fit, as well as in estimates of branch lengths and admixture weights. However, even with infinite data some of these parameters cannot be estimated, as they are not identifiable from the system of equations that corresponds to the admixture graph. Issues like this have been well described for simple topological features of a graph. For example, the lengths of the two branches connected to the root node cannot be estimated without either fixing one of them or forcing the lengths to be evenly distributed. Furthermore, in a graph with populations and admixture events, at least one parameter will not be identifiable unless the inequality – is satisfied (Lipson 2020). However, even in graphs that meet this criterion, some parameters are not identifiable, and until *findGraphs*, there was no method for testing whether any given parameter in an admixture graph is identifiable. We introduce a method for testing which parameters in an admixture graph are identifiable, and which are not, based on the Jacobi matrix of the graph's system of f_2 equations. Like our method for deriving confidence intervals for admixture graph parameters, this can improve the interpretability of admixture graph analyses.

Our methods for automated topology inference, for bootstrapping log-likelihood scores or worst residuals for comparing model fits, for cross-validation of admixture graphs of different complexities, for estimating confidence intervals and determining identifiability of admixture graph parameters can greatly improve the interpretability of admixture graph analyses. We implement all these methods in *findGraphs* to assist the user in building a series of models that can explain admixture processes of ancient populations in a manner that surpasses all other programs in accuracy.

(5) Drawing conclusions from a large number of fitted models

As discussed in the Results section, when we apply our methods for finding optimal graphs and comparing admixture graphs to a number of previously published models, we find that there often exists a much larger number of fitting models than has previously been appreciated. In these cases, we are unable to prioritize a single model, or even a small number of models, based on the evidence

we have. However, we are still able to reject the vast majority of all tested models. This suggests that insight can be gained by identifying common features among the well-fitting models. We therefore introduce methods for summarizing collections of well-fitting admixture graphs to determine which features they share. In practice, we find that these methods can aid the manual inspection of *findGraphs* results, but the high diversity of well-fitting topologies we see in most case studies and the importance of fitted parameters (especially admixture proportions) for historical interpretation of topologies makes it difficult to reliably automatize the process of interpreting fitted admixture graph models.

Revisiting published admixture graphs

We studied admixture graphs from eight publications (Lazaridis et al. 2014, Shinde et al. 2019, Sikora et al. 2019, Bergström et al. 2020, Lipson et al. 2020, Hajdinjak et al. 2021, Librado et al. 2021, Wang et al. 2021) with the goal of comparing published models to models identified by our algorithm for automatically inferring optimal (best-fitting) admixture graphs (Table 1). In all but one study *qpGraph* or its automated reimplementation (*admixturegraph*, Leppälä et al. 2017) was used for fitting topologies to genetic data, while Librado et al. (2021) relied on the automated *OrientAGraph* method. The main question we were interested in is whether we can find alternative models which (1) fit as well as, or better than the published graph, (2) differ in important ways from the published graph, and (3) cannot immediately be rejected based on other evidence such as temporal plausibility. The studies were selected according to the criterion that an admixture graph model inferred in the study is used as primary evidence for at least one statement about population history in the main text of the study. In other words, the admixture graph method was used in the original studies to generate new insights into population history, and not simply to show that there is a model that exists (even if it is not unique) that does not contradict results of other genetic analyses, an approach that is a valid use of admixture graphs and has been taken in some studies (e.g., Seguin-Orlando et al. 2014, Narasimhan et al. 2019, Wang et al. 2019). There are many published studies that could have been included in our re-evaluation exercise as they meet our key criterion (e.g., Yang et al. 2017, McColl et al. 2018, Posth et al. 2018, Flegontov et al. 2019, Calhough et al. 2021, Lipson et al. 2022, Vallini et al. 2022). However, critical re-evaluation of each published graph was an intensive process, and the sample of studies we revisited is diverse enough to identify some general patterns and to support strong conclusions.

Publication	Figure in the original publication	Groups (populations)	Admixture events	SNPs used	Publ. model: worst residual, SE	Distinct alternative topologies found	Significantly better fitting topologies, %	Non-significantly better fitting topologies, %	Non-significantly worse fitting topologies, %	Significantly worse fitting topologies, %
Bergström et al. 2020	1e	7	3	312,282	2.1	221	0.5	2.3	16.7	80.5
Lazaridis et al. 2014	3	7	4	642,247	2.2	306	1.0	12.1	80.7	6.2
Shinde et al. 2019	3	8*	3	249,009	2.6	143	0.0	2.8	3.5	93.7
Librado et al. 2021	3b		3		23.9	324	6.8	15.7	24.1	53.4
	Ext 5d	10*	4	1,767,419	14.1	535	0.0	0.0	4.5	95.5
	Ext 5e		5		6.9	784	0.0	0.3	28.4	71.3
Hajdinjak et al. 2021	2d	12	8	263,698	4.8	1,988	15.7	55.7	6.6	22.0
Lipson et al. 2020	Ext 4	12	11**	211,738	2.3	2,000	0.0	11.9	77.1	10.4
Wang et al. 2021	Ext 6	12	8	203,753	3.8	1,778	12.6	84.3	3.1	0.0
Sikora et al. 2019	3f (left)	13	6**	344,903	3.8	894	0.3	17.1	34.6	48.0
	3f (right)	14	6**	613,509	4.2	2,785	0.1	0.9	9.8	89.2

* the population composition was modified, see Table S1 and the text

** certain gene flows were removed from the published model for simplicity, see Table S1 and the text

Table 1: Published graphs in the context of automatically found graphs.

We compared graphs from 8 publications to alternative graphs inferred on the same or very similar data (see **Table S1** for details).

Publication: Last name of the first author and year of the relevant publication.

Figure in the original publication: Figure number in the original paper where the admixture graph is presented.

Groups (populations): The number of populations in each graph.

Admixture events: The number of admixture events in each graph.

SNPs used: The number of SNPs (with no missing data at the group level) used for fitting the admixture graphs. For all case studies, we tested the original data (SNPs, population composition, and the published graph topology) and obtained model fits very similar to the published ones. However, for the purpose of efficient topology search we adjusted settings for f_3 -statistic calculation, population composition, or graph complexity as noted in the footnotes, in **Table S1**, and discussed in the text.

Publ. model: worst residual, SE: The worst f -statistic residual of the published graph fitted to the SNP set shown in the “SNPs used” column, measured in standard errors (SE).

Distinct alternative topologies found: The number of distinct newly found topologies differing from the published one.

Significantly better fitting topologies, %: The percentage of distinct alternative topologies that fit significantly better than the published graph according to the bootstrap model comparison test (two-tailed empirical p -value < 0.05). If the number of distinct topologies was very large, a representative sample of models (1/20 to 1/3 of models evenly distributed along the log-likelihood spectrum) was compared to the published one instead, and the percentages in this and following columns were calculated on this sample.

Non-significantly better fitting topologies, %: The percentage of distinct topologies that fit non-significantly (nominally) better than the published graph according to the bootstrap model comparison test (two-tailed empirical p -value ≥ 0.05).

Non-significantly worse fitting topologies, %: The percentage of distinct topologies that fit non-significantly (nominally) worse than the published graph according to the bootstrap model comparison test (two-tailed empirical p -value ≥ 0.05).

Significantly worse fitting topologies, %: The percentage of distinct topologies that fit significantly worse than the published graph according to the bootstrap model comparison test (two-tailed empirical p -value < 0.05).

Here we present a high-level summary of these analyses. Discussions of individual graphs, as well an overview of the methodology, can be found in the next section.

For 19 out of 22 published graphs we examined we were able to find at least one, but usually many, graphs of the same complexity (number of groups and admixture events) and with a log-likelihood score that was nominally better than that of the published graph (see results for 11 selected graphs in **Table 1** and full results for all 22 in **Table S1**). The 22 graphs were drawn from the 8 publications as there were multiple final graphs presented in some of the publications (Shinde et al. 2019, Sikora et al. 2019, Librado et al. 2021), or we examined selected intermediates in the model construction process (Bergström et al. 2020, Lazaridis et al. 2014, Lipson et al. 2020, Wang et al. 2021), or we introduced an outgroup not used in the original study (Hajdinjak et al. 2021, Sikora et al. 2019), or we tested additional graph complexity classes dropping “unnecessary” admixture events (Lipson et al. 2020, Sikora et al. 2019).

Oftentimes these alternative graphs are not significantly better than the published one after taking into account variability across SNPs via bootstrapping. In the following cases, at least one model that fits significantly better than the published one according to our bootstrap model comparison method was found: the Bergström et al. and Lazaridis et al. seven-population graphs; the Librado et al. graph with 3 admixture events; the Hajdinjak et al. graphs with or without adding a chimpanzee outgroup; the Lipson et al. intermediate graphs with 7 groups and 4 admixture events and with 10 groups and 8 admixture events; the Wang et al. 12-population graph; and the Sikora et al. graphs for West Eurasians and for East Eurasians with 10 or 6 admixture events (**Table S1**). In nearly all cases (except for the Lazaridis et al. six-population graph, Shinde et al. graph with 8 populations and 3 admixture events, and the Librado et al. graph with 4 admixture events), we also identified a large number of graphs that fit the data not significantly worse than the published ones. In every example,

some of these graphs have topologies that are qualitatively different in important ways from those of the published graphs. Features such as which populations are admixed or unadmixed, direction of gene flow, or the order of split events, if not constrained *a priori*, are generally not the same between alternative fitting models for the same populations. While some of these graphs can be rejected since their topologies appear highly unlikely because of non-genetic or unrelated genetic evidence, for all of the publications except one (Shinde et al. 2019), there are alternative equally-well-or-better-fitting graphs we identified and examined manually that differ in qualitatively important ways with regard to the implications about history, are temporally plausible (for instance, very ancient populations do not receive gene flows from sources closely related to much less ancient groups), and not obviously wrong based on other lines of evidence. These findings suggest the possibility that complex admixture graph models, even with a very good fit to the data, may differ in important ways from true population histories.

The previous statements are valid if the original parsimony constraints are applied, i.e., if the graph complexity (the number of admixture events) is not altered. Below in selected case studies (Shinde et al., Librado et al.) we explore the effect of relaxing the parsimony constraint.

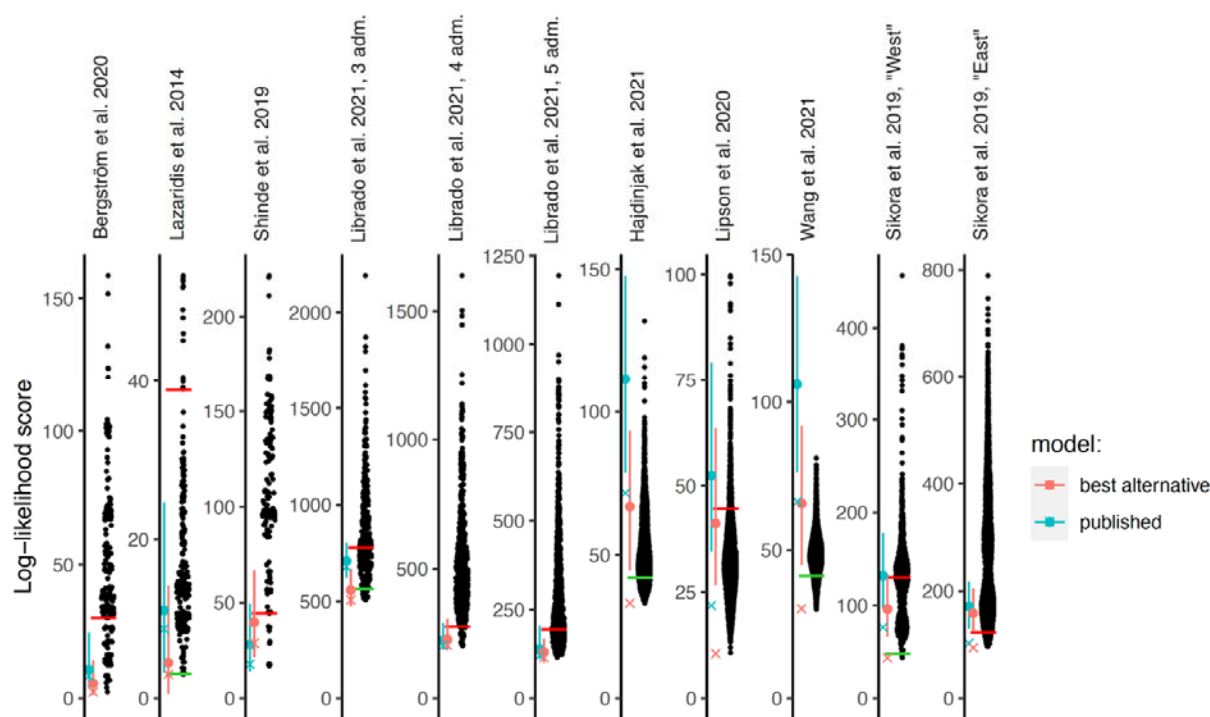


Figure 4:

Log-likelihood scores of published graphs (those shown in **Table 1**) and automatically inferred graphs. Each dot represents the log likelihood score of a best-fitting graph from one *findGraphs* iteration (low values of the score indicate a better fit); only topologically distinct graphs are shown. Log-likelihood scores for the published models and best-fitting alternative models found are shown by blue and pink x's, respectively. Bootstrap distributions of log-likelihood scores for these models (vertical lines, 90% CI) and their medians (solid dots) are also shown. Lower scores of the fits obtained using all SNPs, relative to the bootstrap distribution, indicate overfitting to the full data set. Green and red horizontal lines show the approximate locations where newly found models consistently have fits significantly better or worse, respectively, than those of the published model. In the case of the Bergström et al., Lazaridis et al. and Hajdinjak et al. studies, one or more worst-fitting models were removed for improving the visualization. The setups shown here (population composition, number of groups and admixture events, topology search constraints) match those shown in **Table 1**.

Detailed reconsideration of eight published admixture graphs

We investigated admixture graphs from eight publications. We usually focused on one final graph from each publication, and in some cases we also discuss simpler intermediates in the published model-building process, apply various topological constraints to the model inference process, or decrease/increase the number of admixture events to explore the influence of parsimony constraints. **Table 1** and **Figure 4** summarize these results for one or a few graphs from each publication, while **Table S1** contains the full results for all studied graphs and setups. **Table 2** summarizes our assessment of inferences in the original publications that were supported by the published graphs.

To identify alternative models, we ran many iterations of *findGraphs* for each set of input populations, constraints, and the number of admixture events we investigated, and we selected the best-fitting graph in each iteration, that is, the graph with the lowest log-likelihood (LL) score. Each algorithm iteration was initiated from a random graph, and the algorithm is non-deterministic so that in each iteration it takes a different trajectory through graph space, possibly terminating in a different final best graph. The number of admixture events in the initial random graphs and in the output graphs was always kept equal to that of the published graph. For each example, we counted how many distinct topologies were found with significantly or non-significantly better or worse LL scores than that of the published graph (**Table 1**, **Table S1**). To obtain a formally correct comparison of model fit, the published graph and each alternative model were fitted to resampled replicates of the dataset and the resulting LL score distributions were compared (see Methods). As shown in **Figure 4**, for 4 of the 8 publications we re-analyzed, the LL score of the published graph run on the full data is better than almost all the bootstrap replicates on the same data (it falls below the 5th percentile), which is a sign of overfitting, and underscores the importance of applying bootstrap to assess the robustness of fitted models and conclusions drawn from them.

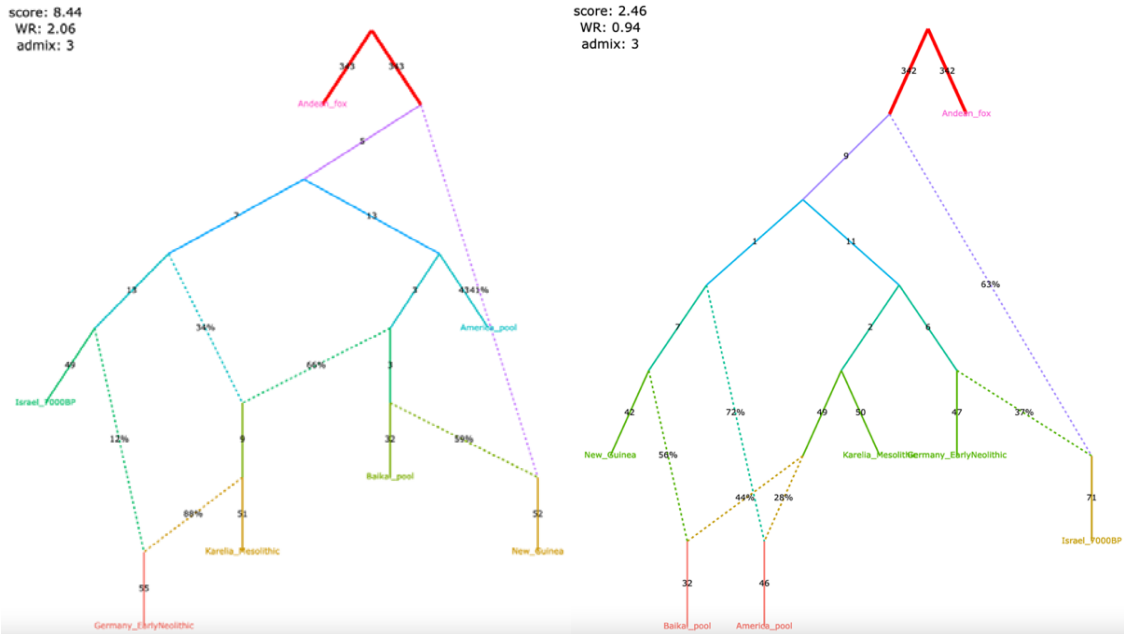
The fraction of graphs with scores better than the score of the published graph should not be overinterpreted, as it is influenced by the *findGraphs* algorithm which does not guarantee ergodic sampling from the space of well-fitting admixture graphs. In particular, it is possible that despite *findGraphs*'s strategies for efficiently identifying classes of well-fitting admixture graphs (see Methods), it has a bias toward missing certain classes of graphs. However, even a single alternative graph which is no worse-fitting than the published graph suggests that we are not able to identify a single best-fitting model. Many of these alternatives, despite providing a good fit to the data, appear unlikely, for example, because they suggest that Paleolithic era humans are mixed between different lineages of present-day humans. We were mainly interested in alternative models which are also plausible, and so we constrained the space of allowed topologies in *findGraphs* to those we considered plausible *a priori*, in cases where this was necessary for reducing the search space size. Constraints were either integrated into the topology search itself, or were applied to outcomes of unconstrained searches, as detailed below.

Figure 5:

Published graphs from six studies for which we explored alternative admixture graph fits. In all cases, we selected a temporally plausible alternative model that fits nominally or significantly better than the published model and has important qualitative differences compared to the published model with respect to the interpretation about population relationships. In all but one case, the model has the same complexity as the published model shown on the left with respect to the number of admixture events; the exception is the reanalysis of Librado *et al.* 2021 horse dataset since the published model with 3 admixture events is a poor fit (worst Z-score comparing the observed and expected *f*-statistics has an absolute value of 23.9 even when changing the composition of the population groups to increase their homogeneity and improve the fit relative to the composition used in the published study). For this case, we show an alternative model with 8 admixture events that fits well and has important qualitative differences from the point of view of population history interpretation relative to the published model. The existence of well-fitting admixture graph models does not mean that the alternative models are the correct models; however, their identification is important because they prove that alternative reasonable scenarios exist to published models.

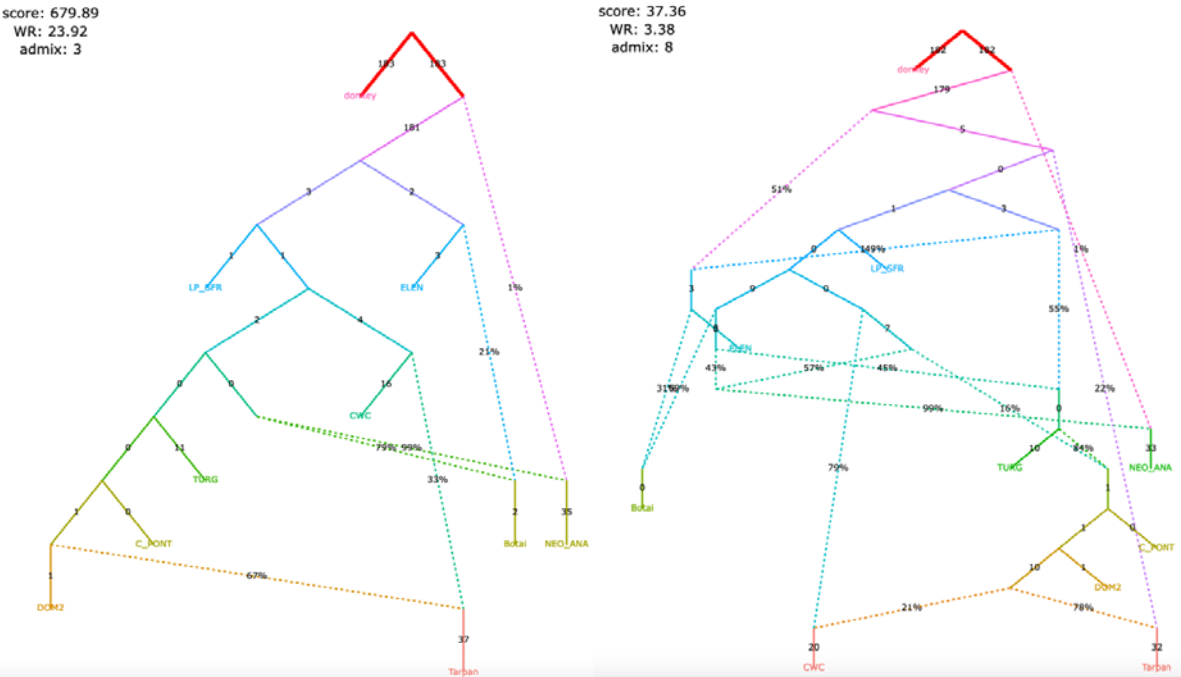
a, Bergström *et al.* 2020 published

Nominally better fitting graph for dogs that is more congruent to human history. For both species, Baikal and Native American groups are mixed between European- and East Asian-related lineages, and a “Basal Eurasian” lineage contributes to Near Eastern groups; these features are all characteristic of human history but absent in the published dog graph.



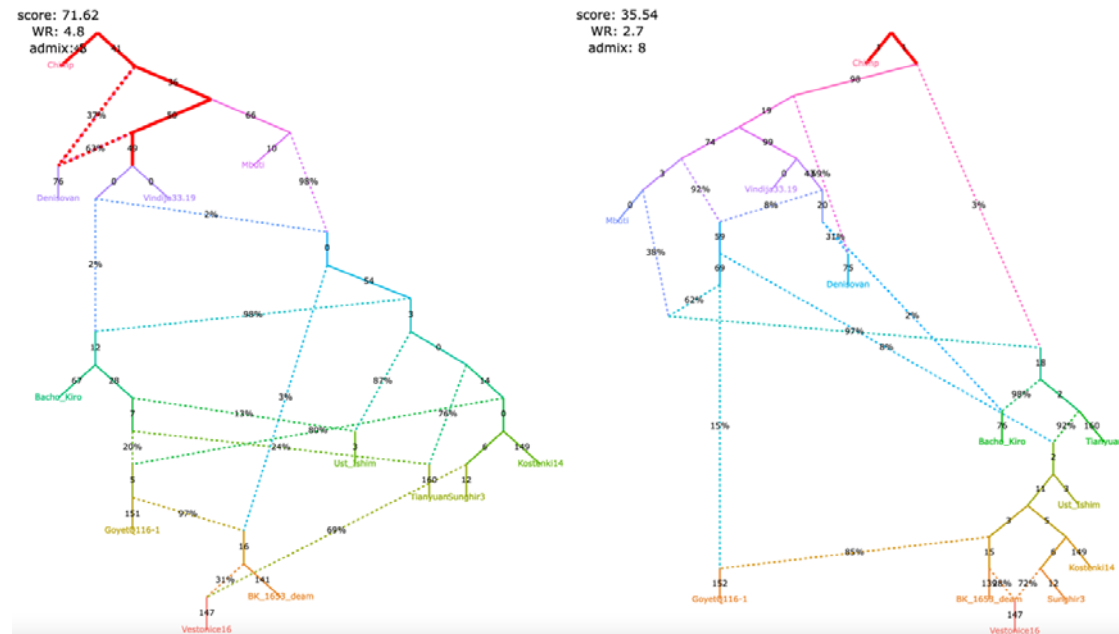
b, Librado *et al.* 2021 (modified population composition) published

Significantly better fitting, temporally/geographically plausible. In contrast to the published graph, in this graph with 8 mixture events (the minimum necessary to obtain an acceptable statistical fit to the data), a lineage maximized in horses associated with Yamnaya steppe pastoralists or their Sintashta descendants (C-PONT, TURG, or DOM2) contributes a substantial amount of ancestry to the horses from the Corded Ware archaeological context (CWC). Thus, in this model both CWC humans and horses are mixtures of Yamnaya and European farmer-associated lineages. This is qualitatively different from the suggestion that there was no Yamnaya-associated contribution to CWC horses which was a possibility raised in the paper. The eight-admixture graph also is different from the published model in that it shows a fitting model where the Tarpan horse does not have the history claimed in the study (as an admixture of the CWC and DOM2 horses).



c, Hajdinjak *et al.* 2021 published

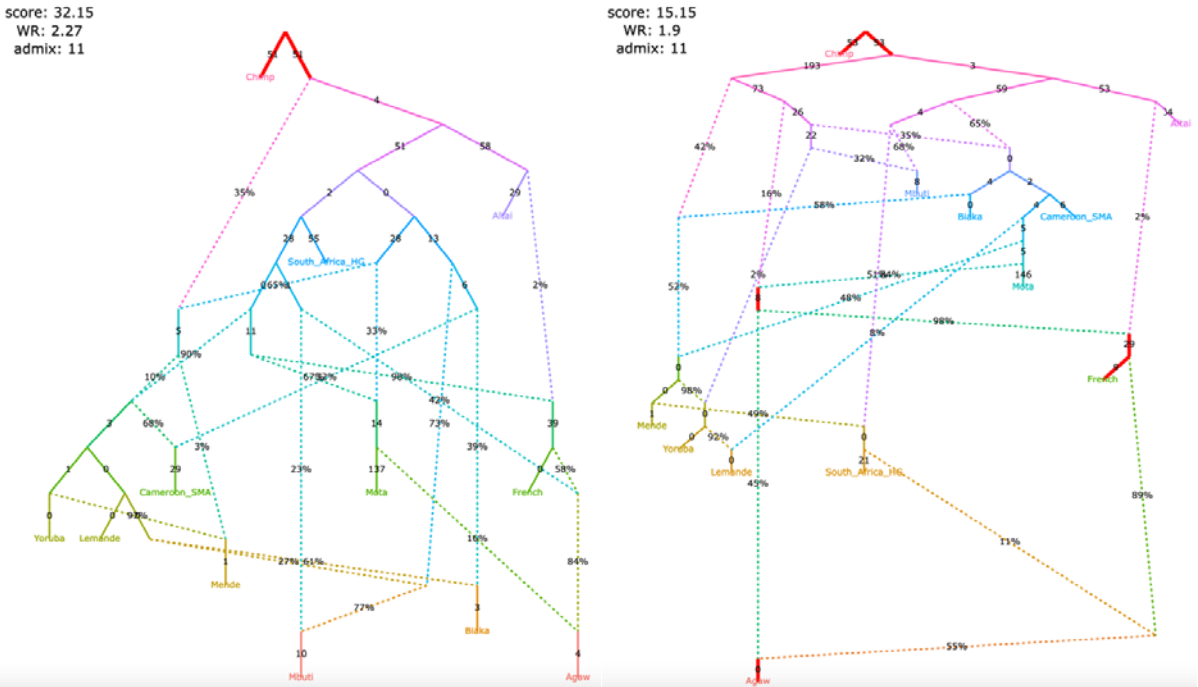
Significantly better fitting, but without a specific lineage shared between Bacho Kiro Initial Upper Paleolithic and East Asians. In this model, all the lineages shared between Bacho Kiro IUP and East Asians contributed a large fraction of the ancestry of later European hunter-gatherers as well, and thus this graph does not imply distinctive shared ancestry between the earliest modern humans in Europe and later people in East Asia, and instead could be explained by a quite different and also archaeologically plausible scenario of a primary modern human expansion out of the Near East contributing serially to the major lineages leading to Bacho Kiro, then later East Asians, then Ust'-Ishim, then the primary ancestry in later European hunter-gatherers.



618

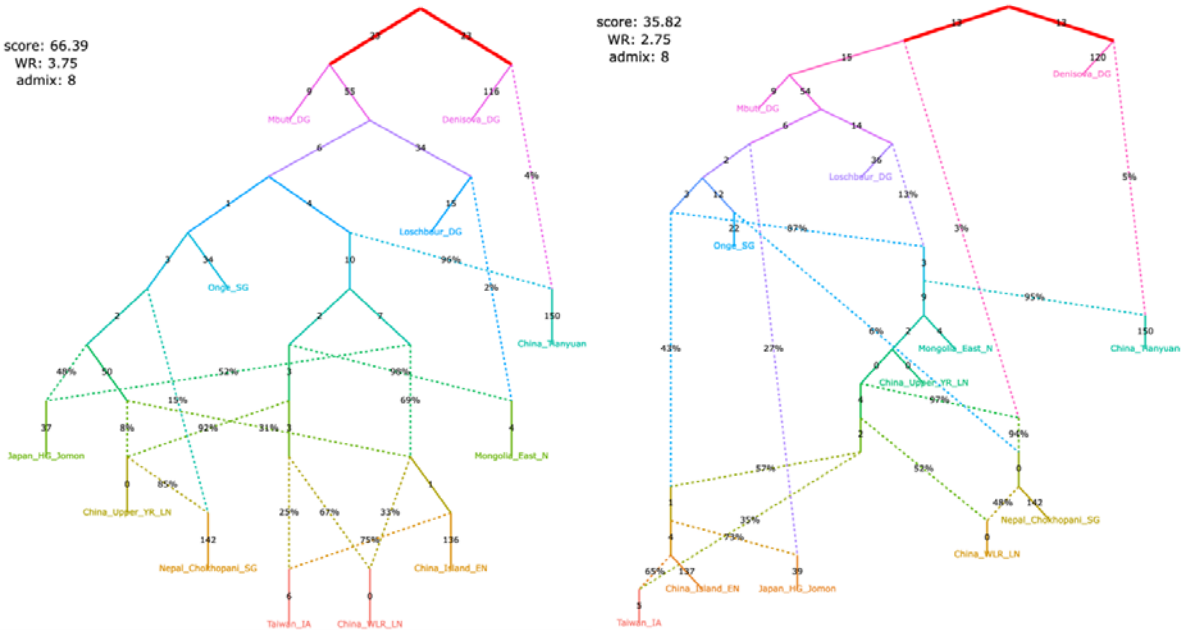
d, Lipson *et al.* 2020 published

Nominally better fitting. In contrast to the published graph, there is no single lineage specific to modern rainforest hunter-gatherers (Biaka and Mbuti) and Shum Laka (Cameroon_SMA). Rather, the primary ancestries in each group are separate deep-branching lineages (the deeper lineage they all share is also the source of the majority of ancestry in all anatomically modern humans modeled here). In contrast to the graph in the published paper, there is no West African-maximized ancestry present in mixed form in Biaka, Mbuti, and Shum Laka; archaic admixture is not limited to a subset of Africans, but is present in all anatomically modern humans in various proportions; and there is no ghost modern human ancestry in Agaw, Biaka, Lemande, Mbuti, Mende, Mota, Shum Laka, and Yoruba.



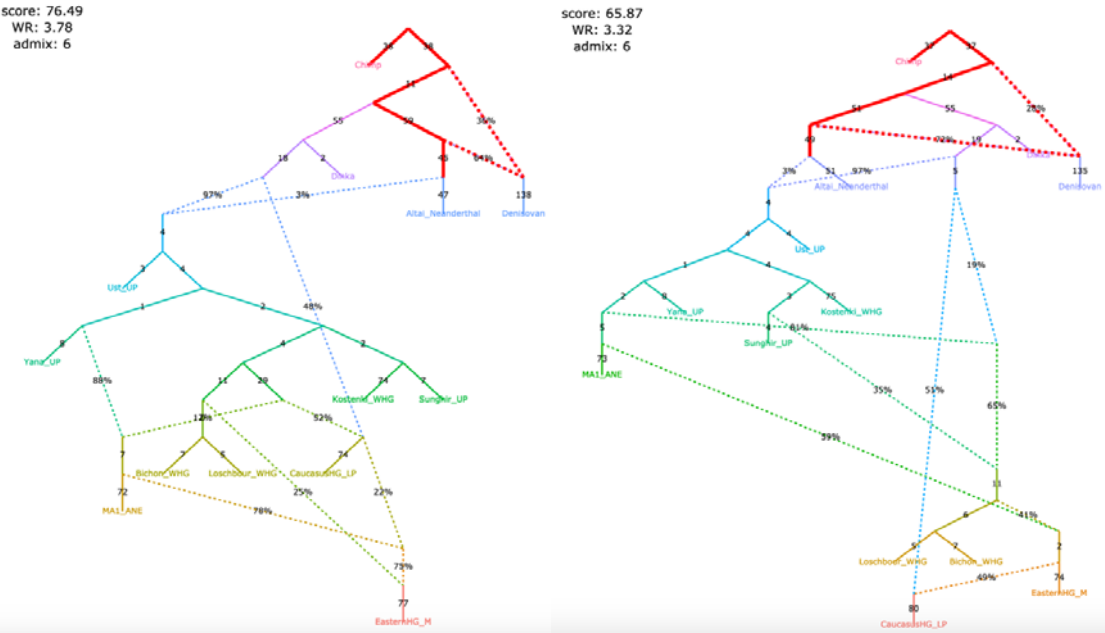
e, Wang *et al.* 2021 published

Significantly better fitting while meeting the constraints used to inform model-building in the published paper. The finding of Onge-related admixture that is widespread in East Asia suggesting an early peopling via a coastal route is not a feature of this model.



f, Sikora *et al.* 2019 (simplified Western graph)

Nominally better fitting. The striking feature of the admixture graph suggested in the paper whereby Mal'ta (MA1_ANE) derives some ancestry from the CHG-associated lineage is not a feature of this alternative model.



Bergström et al. 2020. The admixture graph for dogs in Figure 1e of Bergström et al. 2020 was inferred by exhaustively evaluating all graphs with two admixture events and outgroup ‘Andean fox’ for the six populations that remain in the graph after excluding an Early Neolithic dog from Germany. The only six-population graph with a worst residual (WR) below 3 standard errors (SE) was then chosen as a scaffold onto which the Early Neolithic dog genome from Germany was mapped, allowing for one more admixture event, and a seven-population graph with the lowest LL score was shown as the final model in the paper. Alternative six-population scaffolds were not explored in the original publication, although two six-population graphs with fits very similar to the best one were found. LL was not used as a ranking metric for alternative models in the original study; instead, the number of f_4 -statistics having residuals above 3 SE was considered. Since no f_3 -statistics were negative when all sites available for each population triplet were used (the “useallsnps: YES” option), we did not use the upgraded algorithm for calculating f_3 -statistics on pseudo-diploid data.

Our *findGraphs* results confirm the published six-population graph in that no graph with lower LL score is identified, but 3 of 14 unique alternative graphs found fit not significantly worse than the published graph (the published graph was also recovered by *findGraphs*) (**Table 1, Table S1**). When we used *findGraphs* to infer seven-population graphs with three admixture events (again fixing Andean fox as the outgroup), we identified 5 graphs with log-likelihood score nominally better than that of the published graph and one with a score that is slightly lower than that of the published graph but actually significantly better according to our model comparison methodology (this model is very similar to the published graph, **Figure S2**). In the newly found seven-population graph with the best log-likelihood score (**Figure S2**), the Siberian (Baikal), American, and Levantine dogs are admixed, and the West European, East European (Karelia), and dogs of Southeast Asian origin (New Guinea singing dog) are unadmixed, while the opposite pattern is found in the published graph (**Table 2**). The best-fitting graph does not fit the data significantly better than the published graph (two-tailed empirical p -value = 0.332), but it bears a closer resemblance to the human population history (see the third-best graph found by *findGraphs* on human data from Bergström et al. 2020 in **Figure S3**) than the published seven-population graph (**Figure 5a, Figure S2**).

In this new seven-population model (**Figure 5a**), both American and Siberian dog lineages represent a mixture between groups related to the Asian and East European dog lineages, and robust genetic results suggest that in the time horizon investigated in the original publication (after ca. 10,900 years ago) nearly all Siberian (Jeong et al. 2019, Sikora et al. 2019) and all American (Raghavan et al. 2014, Raghavan et al. 2015, Moreno-Mayar et al. 2018) human populations were admixed between groups most closely related to Europeans and Asians. According to this model, Levantine dogs are modeled as a mixture of a basal branch (splitting deeper than the divergence of the Asian and European dogs) and West European dogs, again in agreement with current models of genetic history of Middle Eastern human populations who are modeled as a mixture of “basal Eurasians” and West European hunter-gatherers (Lazaridis et al. 2016, Lipson et al. 2017). Although greater congruence with human history increases the plausibility of *findGraphs*’s newly identified model relative to the published model, to make unbiased comparisons between the history of the two species, model selection should be done strictly independently for each species, and so the genetic data alone does not favor one model more than another. Our results provide a specific alternative hypothesis that differs in qualitatively important ways from the published model and can be tested against new genetic data as it becomes available as well as other lines of genetic analysis of existing data.

To explain why the original paper on the population history of dogs missed the model that *findGraphs* identified that is plausibly a closer match to the true history, we observe that the Bergström et al. 2020 admixture graph search was exhaustive under the parsimony constraint (no more than 2 admixture events for 6 populations), and thus missed the potentially true topology including 3 admixture events for these 6 populations. This case study also illustrates that even in a relatively low complexity context (7 groups and 3 admixture events) applying manual approaches for finding optimal models is risky. When any new group such as an Early Neolithic dog from Germany is added to the model, it may introduce crucial information into the system, and re-exploring the

whole graph space in an automated way is advisable. In contrast, mapping a newly added group on a simple skeleton graph (even when that skeleton is a uniquely best-fitting model) may yield a topology that is at odds with the true history. As the original Bergström *et al.* paper noted (Fig. 3C of that study), no congruent six-population graph models were found for humans and dogs under the parsimony assumption: the three most congruent graphs for dogs resulted in poorly fitting models for the corresponding human populations (WR above 10 SE), and the three most congruent graphs for humans resulted in poorly fitting models for the corresponding dog populations (WR between 5 and 10 SE). We added a West European hunter-gatherer group to the set of human groups from the original publication and using *findGraphs* on the original set of 77K transversion SNPs we found that the third best-fitting model for humans (**Figure S3**) (which is not significantly different in fit from the first one) is topologically identical to the newly found dog graph.

Even though *findGraphs* identified an admixture graph topology that fits the data as well as the seven-population graph in Bergström *et al.* and is qualitatively quite different with respect to which populations were admixed, the new topology continues to support another of the key inferences of that study: that many of the early divergences among domesticated dog lineages occurred prior to the date of the Karelian dog (~10,900 ya). Thus, both graphs concur in providing strong evidence that the radiation of domesticated dog lineages occurred by the early Holocene, prior to the domestication of other animals. We further emphasize that the Bergström *et al.* 2020 graph is the best-case scenario (along with Lazaridis *et al.* 2014 discussed below) for published admixture graphs. Most published graphs are far less stable even than this.

Table 2: Features of the published admixture graphs that support inferences in the original studies and the level of their support in our re-analysis.

The table lists key features whose support we assessed in sets of alternative well-fitting and temporally plausible models generated by *findGraphs*. Since this assessment had to be performed manually, only in two cases (marked by asterisks) all models fitting better and non-significantly worse than the published one were scrutinized; in other cases only a subset of best-fitting models was examined (see the sections for details).

Study	Groups / admixture events	Features of the published model supported by temporally plausible alternative models generated by <i>findGraphs</i>	Features of the published model <u>not</u> supported by temporally plausible alternative models generated by <i>findGraphs</i>
Bergström et al. 2020	7 / 3 *	Early divergence of domesticated dog lineages (prior to the date of the Karelian dog, 10,900 ya).	Siberian (Baikal), American, and Levantine dog lineages are unadmixed, and the West European (Germany Early Neolithic), East European (Karelia), and dogs of Southeast Asian origin (New Guinea singing dog) are admixed.
Lazaridis et al. 2014	7 / 4	Present-day Europeans represent a mixture of three ancestral sources related to the following groups: Mal'ta (MA1), West European hunter-gatherers, and early European farmers.	N/A
Shinde et al. 2019	8 / 3 *	(1) Iranian farmer-related ancestry in the Indus Periphery group is not derived from the Hajji Firuz Neolithic or Tepe Hissar Chalcolithic groups. (2) There is Asian-related ancestry in the Indus Periphery group.	N/A

	8 / 4	(2) There is Asian ancestry in the Indus Periphery group.	(1) Iranian farmer-related ancestry in the Indus Periphery group is not derived from the Hajji Firuz Neolithic or Tepe Hissar Chalcolithic groups.
Librado et al. 2021	10 / 8 or 9	(2) DOM2 and C-PONT are sister groups (they form a clade); (4) there was gene flow from a deep-branching ghost group to the NEO-ANA group.	(1) NEO-ANA-related admixture is absent in the DOM2 group; (3) there is no gene flow connecting the CWC group and the cluster associated with Yamnaya horses and horses of the later Sintashta culture whose ancestry is maximized in the Western Steppe (DOM2, C-PONT, TURG); (5) Tarpan is a mixture of a CWC-related and a DOM2-related lineage.
Hajdinjak et al. 2021	12 / 8	(3) the Vestonice16 lineage is a mixture of a Sunghir-related and a BK1653-related lineage.	(1) there are gene flows from the lineage found in the ~45,000-43,000-years-old Bacho Kiro Initial Upper Paleolithic (IUP) associated lineage to the Ust'-Ishim, Tianyuan, and GoyetQ116-1 lineages; (2) the ~35,000-years-old Bacho Kiro Cave individual BK1653 belonged to a population that was related, but not identical, to that of the GoyetQ116-1 individual.
Lipson et al. 2020	12/ 11	N/A	(1) A lineage maximized in present-day West African groups (Lemane, Mende, and Yoruba) also contributed some ancestry to the ancient Shum Laka individual and to present-day Biaka and Mbuti; (2) another ancestry component in Shum Laka is a deep-branching lineage maximized in the rainforest hunter-gatherers Biaka and Mbuti; (3) "super-archaic" ancestry (i.e., diverging at the modern human/Neanderthal split point or deeper) contributed to Biaka, Mbuti, Shum Laka, Lemane, Mende, and Yoruba; (4) a ghost modern human lineage (or lineages) contributed to Agaw, Mota, Biaka, Mbuti, Shum Laka, Lemane, Mende,

			and Yoruba.
Wang et al. 2021	12 / 8	N/A	Admixture from a source related to Andamanese hunter-gatherers is almost universal in East Asians, occurring in the Jomon, Tibetan, Upper Yellow River Late Neolithic, West Liao River Late Neolithic, Taiwan Iron Age, and China Island Early Neolithic (Liangdao) groups.
Sikora et al. 2019 “West”	13 / 6	N/A	The Mal’ta (MA1_ANE) lineage received a gene flow from the Caucasus hunter-gatherer (CaucasusHG_LP or CHG) lineage.
Sikora et al. 2019 “East”	14 / 6	(2) European-related ancestry in the Kolyma, USR1, and Clovis lineages is closer to Mal’ta than to Yana.	(1) the Mal’ta (MA1_ANE) and Yana (Yana_UP) lineages received gene flow from a common East Asian-associated source diverging before the ones contributing to the Devil’s Cave (DevilsCave_N), Kolyma (Kolyma_M), USR1 (Alaska_LP), and Clovis (Clovis_LP) lineages; (3) the Devil’s Cave lineage received no European-related gene flows, and Kolyma has less European-related ancestry than ancient Americans (USR1 and Clovis).

Lazaridis et al. 2014. The graph in Figure 3 (Lazaridis et al. 2014) was inferred in the following manner. First, a phylogenetic tree without admixture was constructed which was the best fit for all f_4 -statistics among the populations “Mbuti”, “WHG Loschbour” (Lazaridis et al. 2014), “LBK Stuttgart” (Lazaridis et al. 2014), “Onge”, and “Karitiana”, with “Mbuti” fixed as an outgroup. Next, all possible admixture graphs were considered that result from adding a single admixture edge to this tree. After it was found that each of them had a $WR > 3$ SE, several graphs with two admixture events were considered, and some of them had $WR < 3$ SE. The “MA1” genome (Raghavan et al. 2014) was added to these graphs in several different ways, and only one of these configurations was found to have $WR < 3$ SE. This was then used as a skeleton graph onto which a European population (represented by different present-day groups) was added. No fitting graph was found in which present-day Europeans could be modeled as a two-way mixture (adding one admixture event to the graph). After inspecting the non-fitting f -statistics of one of these graphs, it was found that modeling modern Europeans as a three-way mixture (adding two admixture events to the graph) is consistent with all f -statistics. Six f_3 -statistics were negative when all sites available for each population triplet were used (the “useallsnps: YES” option, **Table S1**), but the upgraded algorithm for calculating f_3 -statistics on pseudo-diploid data had no effect since the only pseudo-diploid group in the dataset (MA1) was a singleton population, and the algorithm removes sites with only one chromosome genotyped in any non-singleton population. Thus, below we show results generated using the standard algorithm for calculating f -statistics.

First, we considered the published skeleton graph onto which a European population was later added (Lazaridis et al. 2014). As in the (Bergström et al. 2020) example, the best six-population graph with two admixture events found by *findGraphs* is identical to the published six-population

graph, which has a LL score of 3.0 (**Table S1**). The second-best graph found has a LL score of 31.8. When computing the bootstrap p -value for the difference between these two graphs, we find that in 1.6% of all SNP resamplings the second-best graph has a better score than the published graph, resulting in a two-tailed empirical p -value of 0.032 for a difference in fits between these two graphs. All 14 alternative graphs found by our algorithm fit significantly worse than the published graph (**Table S1**).

When we add the European population (French) and consider seven-population graphs with four admixture events, we find 40 out of 306 distinct graphs with a score better than that of the published graph (10 of those graphs are shown in **Figure S4**). The best-fitting newly found model and two other models fit the data significantly better than the published model (**Table S1**), but their topology is qualitatively very similar to that of the published graph (**Figure S4**). In the best-fitting newly found model, French and Karitiana share some drift to the exclusion of MA1, while in the published model the source of MA1-related ancestry in French is closer to MA1 than to Karitiana.

It is important to point out that not all of the 40 alternative graphs that fit nominally or significantly better than the published one are consistent with the conclusion that modern European populations are admixed between three different ancestral populations (**Figure S4**). For example, the fifth alternative graph in **Figure S4** that is fitting nominally better than the published model (p -value = 0.464) includes no basal Eurasian ancestry in early European farmers (LBK Stuttgart), and instead models Onge as having ~50% West Eurasian-related ancestry and MA1 as having ~25% Asian ancestry. According to that graph, the present-day European population was formed by admixture of a MA1-related lineage and a European Neolithic-related lineage, with no West European hunter-gatherer (WHG) contribution. Of course, other lines of evidence make it clear that LBK Stuttgart is a mixture of Anatolian farmer-related ancestry and WHG Loschbour-related ancestry (Lazaridis et al. 2016, Lipson et al. 2017), thus providing external information in favor of the Lazaridis *et al.* 2014 model, and the use of such ancillary information in concert with graph exploration is important in order to obtain more confident inferences about population history taking advantage of admixture graphs.

The second alternative graph in **Figure S4** that fits just negligibly worse than the highest-ranking model has another distinctive feature: LBK Stuttgart is modeled as a mixture of a WHG-related and a basal Eurasian lineage, but modern Europeans receive a gene flow not from the LBK-related lineage, but from its basal Eurasian source. Although temporally plausible, this model is much less plausible from the archaeological point of view than the published model, and thus in this case too we can reject it as unlikely based on non-genetic evidence. We note, however, that a large group of newly found models (247 graphs) fits not significantly worse than the published one (**Table S1**), and those are topologically diverse. Thus, strictly speaking, the admixture graph method on the given dataset cannot be used to prove that the published model is the only one fitting the data.

Shinde et al. 2019. The skeleton admixture graph in the original study (Shinde et al. 2019) was constructed manually on the basis of a SNP set derived from the 1240K enrichment panel, and subsequently all possible branching orders (105) within the five-population Iranian farmer-related clade were tested. The published model (Fig. 3 in that study) included 9 groups and 3 admixture events, but one group (Belt Cave Mesolithic) had a very high missing data rate, and as a result model fitting relied not just on the merged dataset which included 19,000 polymorphic sites without missing data across groups, but also on a dataset with approximately 470,000 sites that excluded the Belt Cave individual. The topological inferences were consistent for both analyses (Table S3 of that study). Following the approach of the published paper, we repeated *findGraphs* analysis both with and without the Belt Cave individual. Thus, we initially explored the following topology classes: 9 groups with 3 admixture events on ca. 19,000 polymorphic sites and 8 groups with 3 admixture events on ca. 470,000 sites (**Table S1**). The sample composition of the groups and the SNP dataset

matched that in the original study. We summarize results across 4,000 independent iterations of the *findGraphs* algorithm for each topology class.

For the nine-population graph we found 89 models with LL nominally better than that of the published model (**Table S1**). For the eight-population graph, we found 61 nominally and 4 significantly better fitting models (**Table S1**), and their topological diversity was high (**Figure S5**). We note that the following groups were admixed by default in the graph models compared in the original study: Hajji Firuz Neolithic (labeled “Chalcolithic” in that study but the dates are Neolithic) and Tepe Hissar Chalcolithic were considered as mixtures of an Anatolian farmer-related lineage and an Iranian farmer-related lineage; Indus Periphery was considered as a mixture of an Andamanese-related lineage representing ancient South Indians (ASI) and an Iranian farmer-related lineage. However, calculation of negative “admixture” f_3 -statistics for these target groups is impossible using the original dataset and the original model fitting algorithm for several reasons. First, the Indus Periphery group was represented by a single pseudo-diploid individual (I8726) from the “Indus Valley cline” for whom the best-quality data were available. But direct calculation of “admixture” f_3 -statistics for such a group as a target is impossible since its heterozygosity cannot be estimated. Second, as discussed above, in Classic *ADMIXTOOLS* it is impossible to apply a correction intended for accurate calculation of f_3 -statistics on pseudo-diploid data (the “inbreed: YES” option) if there is at least one population composed of one individual only (a singleton population). Third, the original Hajji Firuz Neolithic group composed of five individuals included a family of three 2nd or 3rd-degree relatives, and that artificially inflated the drift on the Hajji Firuz branch and made detecting a negative statistic f_3 (Hajji Firuz Neolithic; X, Y), even if present, highly unlikely. Indeed, no f_3 -statistic turned out to be nominally negative for the groups on the eight-population graph when statistics were calculated according to the original settings (we used settings equivalent to “useallsnps: NO” and “inbreed: NO” in classic *ADMIXTOOLS*, 470,389 polymorphic sites were available). Considering this fact, it is not surprising that our automated topology space search is not well constrained. The original study differed from ours since the constraints were introduced manually, but we wanted our topology search to be automatic and to explore a wider range of parameter space.

In order to provide power to detect negative f_3 -statistics useful for constraining the model search, we 1) removed two members of the family from the Hajji Firuz Neolithic group, 2) extended the number of individuals and sites available for the Indus Periphery group by generating new shotgun-sequencing data for a previously published library (Narasimhan et al. 2019) derived from individual I8726 (see **Table S2**) and by adding published data for three other individuals from the Indus Valley cline (from Gonur in Turkmenistan and Shahr-i-Sokhta in Iran; Narasimhan et al. 2019, Shinde et al. 2019), 3) removed from other groups two individuals based on 2nd-3rd degree relatedness, and 4) removed two individuals from other groups based on evidence of contamination with modern human DNA. All the changes to the dataset are shown in **Table S3**. In addition to these dataset adjustments, the new algorithm for calculating f -statistics makes it possible to compute negative f_3 -statistics on pseudo-diploid data, but at a cost of removing sites with only one chromosome genotyped in any non-singleton population (see Methods). We eventually detected significantly negative “admixture” f_3 -statistics f_3 (Tepe Hissar Chalcolithic; Ganj Dareh Neolithic, Anatolia Neolithic), f_3 (Indus Periphery; Ganj Dareh Neolithic, Onge), and other similar statistics for the same target groups. We also observed a nominally negative (Z-score = -0.6) statistic f_3 (Hajji Firuz Neolithic; Ganj Dareh Neolithic, Anatolia Neolithic), which is suggestive but does not by itself prove admixture in Hajji Firuz Neolithic. For this updated analysis, 249,009 variable sites without missing data at the group level were available for the eight populations.

We repeated topology search with this set of f -statistics providing additional constraints, performing 4,000 runs of the *findGraphs* algorithm. The Mota ancient African individual was set as an outgroup and 3 admixture events were allowed in the eight-population graph. Among 4,000 resulting graphs (one from each *findGraphs* run), 144 were distinct topologically, and the published model was recovered in 13 runs of 4,000 (**Table S1**). Only 4 distinct topologies fitting nominally better than the published one were found, and those had LL scores almost identical to that of the published eight-

population model (16.97 and 17.66 vs. 17.85). These four alternative models (**Figure S6b**) shared all topologically important features of the published model (**Figure S6a**). Five other topologies differed in important ways from the published one and emerged as fitting the data worse, but not significantly worse, than the published one (**Figure S6c**): two-tailed empirical *p*-values reported by our bootstrap model comparison method ranged between 0.060 and 0.112. Three of these topologies included a trifurcation of Iranian farmer-related lineages leading to the Indus Periphery, Hajji Firuz Neolithic, and Ganj Dareh Neolithic groups. The other two topologies included Hajji Firuz Neolithic as an unadmixed Anatolian-related lineage. In both cases, Indus Periphery was modeled as receiving a gene flow from either the Onge lineage (a proxy for ASI) or a deep Asian lineage.

The finding that the predominant ancestry component of the Indus Periphery group was the most basal branch in the Iranian farmer clade was a prominent claim of the original study (Shinde et al. 2019); for example, the abstract stated: “The Iranian-related ancestry in the IVC derives from a lineage leading to early Iranian farmers, herders, and hunter gatherers before their ancestors separated”. Our finding that the Hajji Firuz Neolithic lineage may be as deep within the Iranian clade as the Indus Periphery lineage or may even diverge from the Anatolian branch shows that this statement cannot be confidently made based on admixture graph analysis alone.

However, the findings we have described up to this point do not invalidate the broader conclusion that the admixture graph modeling in Shinde *et al.* was used to support; namely (using the phrasing from the abstract) that the genetic data “contradict... the hypothesis that the shared ancestry between early Iranians and South Asians reflects a large-scale spread of western Iranian farmers east.” This finding if correct is important, since it implies that the Iranian-related ancestry in the IVC (Indus Valley Civilization genetic grouping, which is the same group as IP), split from the Iranian-related ancestry in the first Iranian plateau farmers before the date of the Hajji Firuz farmers, who at ~8000 years ago are among the earliest people living on the Iranian plateau known to have grown West Asian crops. The ancient DNA record combined with radiocarbon dating evidence suggests that beginning around the time of the Hajji Firuz farmers, both West Asian domesticated plants such as wheat and barley, and Anatolian farmer-related admixture, began spreading eastward across the Iranian-plateau. If the Iranian-related ancestry in IP was spread eastward into the Indus Valley across the Iranian-plateau as part of the same agriculturally-associated expansion—perhaps brought by people speaking Indo-European languages as well as introducing West Asian crops—then we would expect to see at least some of the Iranian-related ancestry in IP being a clade with that in Hajji Firuz relative to Ganj Dareh. The fact that we do not find any models compatible with this scenario is thus a potentially important finding. In summary, there are two reasons the genetic analyses we have reported up to this point continue to support the finding that the Iranian-related ancestry in IP is not a clade with the Iranian-related ancestry in Hajji Firuz (and Tepe Hissar) and thus is unlikely to reflect the same eastward movement of agriculturalists. First, in *findGraphs* analysis, all models specifying IP and Tepe Hissar and/or Hajji Firuz as a clade relative to Ganj Dareh were significantly worse-fitting than the published one. Instead, either the Iranian-related ancestry in IP definitively splits off first (the topology from Shinde *et al.*), or the branching order of IP, Ganj Dareh, and the Hajji Firuz / Tepe Hissar lineages cannot be determined, or IP, Ganj Dareh, and Tepe Hissar are a clade relative to Hajji Firuz. In all these fitting topologies, the ~10,000-year-old radiocarbon date of the Ganj Dareh individuals sets a lower bound on the split time between IP and Hajji Firuz / Tepe Hissar, which is pre-agriculturalist. This suggests that the Iranian-related ancestry in IP is not due to an eastward agriculturalist expansion.

But in fact, the admixture graph analysis reported above is not an adequate exploration of the problem. Although absolute fits of the best models found are good (WR = 2.5 SE), the parsimony constraint allowing only 3 admixture events precluded correct modeling of basal Eurasian ancestry shared by all Middle Eastern groups (Lazaridis et al. 2016) or of the Indus Periphery group itself, for which a more complex 3-component admixture model was proposed (Narasimhan et al. 2019). Concerned that this oversimplification could be causing our search to miss important classes of models, we explored *qpAdm* models for the Indus Periphery group further, following the “distal”

protocol with “rotating” outgroups outlined by Narasimhan *et al.* (2019) and using the dataset and outgroups (“right” populations) from that study. All sites available for analyses were used, following Narasimhan *et al.* (the “useallsnps: YES” option). The combined Indus Periphery group we analyzed included 7 individuals from Shahr-i-Sokhta and 3 individuals from Gonur (3 individuals were removed from the Narasimhan *et al.* 2019 dataset due to potential contamination with modern human DNA and low coverage). We removed one individual from the Ganj Dareh Neolithic group as potentially contaminated, and one 2nd or 3rd degree relative was removed from the Anatolia Neolithic group, see the dataset composition in **Table S4**. We note that no “distal” *qpAdm* models were tested for the combined Indus Periphery group by Narasimhan *et al.* (2019), and individuals from this group were modeled one by one (Table S82 from Narasimhan *et al.* 2019), which potentially reduced the sensitivity of the method.

A model “Indus Periphery = Ganj Dareh Neolithic + Onge (ASI)” was strongly rejected for the Indus Periphery group of 10 individuals with a p -value = 2×10^{-15} , and a model that was shown to be fitting for all Indus Periphery individuals modeled one by one by Narasimhan *et al.* (Ganj Dareh Neolithic + Onge (ASI) + West Siberian hunter-gatherers (WSHG)) was rejected for the grouped individuals with a p -value = 0.0044. In contrast, a model “Indus Periphery = Ganj Dareh Neolithic + Onge (ASI) + WSHG + Anatolia Neolithic” was not rejected based on the $p > 0.01$ threshold used in Narasimhan *et al.* (p -value was marginal but passing at 0.03) and produced plausible admixture proportions for all four sources that are confidently above zero: $53.2 \pm 5.3\%$, $28.7 \pm 2.1\%$, $10.5 \pm 1.3\%$, $7.7 \pm 2.9\%$, respectively (**Table S5**). The same “distal” model albeit with Anatolian Neolithic always in higher proportion was found as one of the simplest models (or the only simplest model) fitting the data for many other groups from Iran and Central Asia explored by Narasimhan *et al.* (2019): Aligrama2_IA (13% Anatolia Neolithic), Barikot_H (21%), BMAC (26%), Bustan_BA_o2 (15%), Butkara_H (24%), Saidu_Sharif_H_o (12%), Shahr_I_Sokhta_BA1 (19%), SPGT (23%) (**Table S5** gives a compendium of distal modeling results by Narasimhan *et al.*). When we modeled Indus Periphery individuals separately, as in Narasimhan *et al.* (2019), the simplest two-component model “Ganj Dareh Neolithic + Onge (ASI)” was rejected for 5 of 10 individuals (at least ~315,000 sites were genotyped per individual), including the individual I8726 used for the admixture graph analysis in Shinde *et al.* (**Table S6**). The model “Ganj Dareh Neolithic + Onge (ASI)” was not rejected only for individuals with fewer than 141,000 sites genotyped, suggesting that this result is attributed not to population heterogeneity, but to lack of power.

These *qpAdm* results show that the parsimony assumption that was made when constructing the admixture graph analysis in Shinde *et al.* (2019) is contradicted by f -statistic evidence, and indeed Narasimhan *et al.* themselves showed this when they presented a distal *qpAdm* model that was more complex (Ganj Dareh Neolithic + Onge (ASI) + WSHG) than the one used for constraining the admixture graph model comparison (Ganj Dareh Neolithic + Onge (ASI)). Another line of evidence used to support the principal historical conclusion by Shinde *et al.* was a series of f_4 -statistic cladality tests following correction of allele frequencies using an admixture model “target group = Iranian farmer + Anatolia Neolithic + Onge (ASI)”, with a great majority of tests supporting the deepest position of the Iranian farmer ancestry component in the Indus Periphery group within the Iranian farmer clade (Shinde *et al.* 2019). However, the model used for allele frequency correction (Ganj Dareh Neolithic + Onge (ASI) + Anatolia Neolithic) was simpler than the 4-component model and different from the 3-component model for the Indus Periphery group suggested by Narasimhan *et al.* (2019), which is a weakness of that analysis. A valuable direction for future work would be to repeat this analysis with a 4-component allele frequency correction model (Ganj Dareh Neolithic + Onge (ASI) + WSHG + Anatolia Neolithic), although that is beyond the scope of the present study, which simply aims to re-visit the reported analyses and test if they fully support their inferences by ruling out alternative explanations.

To explore how the parsimony constraint influences results, we allowed 4 admixture events in the eight-population graph (**Table S1**). Among 4,000 resulting graphs (one from each *findGraphs* run), 443 were distinct topologically, and 270 had WRs between 2 and 3 SE, i.e., fitted the data well. We

explored 35 topologies with LL scores in a narrow range between 9.3 (the best value) and 13.3. In **Figure S7b** we show four graphs with four admixture events that model the Indus Periphery group as a mixture of three or four sources, with a significant fraction of its ancestry derived from the Hajji Firuz Neolithic or Tepe Hissar Chalcolithic lineages including both Iranian and Anatolian ancestries. The fits of these models are just slightly different (e.g., LL = 11.7 vs. 9.3, both WRs = 2.4 SE) from that of the best-fitting model (**Figure S7a**), and similar to that of the published graph. Besides these four illustrative graphs, dozens of topologies with very different models for the Indus Periphery group fit the data approximately equally well, suggesting that there is no useful signal in this type of admixture graph analysis when the parsimony constraint is relaxed (this finding is similar to that in our reanalysis of the dog admixture graph in Bergström *et al.* 2020, where relaxation of the parsimony constraint identified equally well fitting admixture graphs that were very different with regard to their inferences about population history). These results show that at least with regard to the admixture graph analysis, a key historical conclusion of the study (that the predominant genetic component in the Indus Periphery lineage diverged from the Iranian clade prior to the date of the Ganj Dareh Neolithic group at ca. 10 kya and thus prior to the arrival of West Asian crops and Anatolian genetics in Iran) depends on the parsimony assumption, but the preference for three admixture events instead of four is hard to justify based on archaeological or other arguments.

Why did the Shinde *et al.* 2019 admixture graph analysis find support for the IP Iranian-related lineage being the first to split, while our *findGraphs* analysis did not? The Shinde *et al.* 2019 study sought to carry out a systematic exploration of the admixture graph space in the same spirit as *findGraphs*—one of only a few papers in the literature where there has been an attempt to do so—and thus this qualitative difference in findings is notable. We hypothesize that the inconsistency reflects the fact that the deeply-diverging WSHG-related ancestry (Narasimhan *et al.* 2019) present in the IVC (Indus Valley Civilization genetic grouping, which is the same group as Indus Periphery) at a level of ca. 10% was not taken into account explicitly neither in the admixture graph analysis nor in the admixture-corrected f_4 -symmetry tests also reported in Shinde *et al.* (2019). The difference in qualitative conclusions may also reflect the fact that the Shinde *et al.* study was distinguishing between fitting models relying on a LL difference threshold of 4 units (based on the AIC). As discussed above, AIC is not applicable to admixture graphs where the number of independent model parameters is topology-dependent even if the numbers of groups and admixture events are fixed, and models compared with AIC should have the same number of parameters. Moreover, model comparisons with AIC do not account for the variability across SNPs, unlike the bootstrap model-comparison method we use. Thus, the analysis by Shinde *et al.* was over-optimistic about being able to reject models that were in fact plausible using its admixture graph fitting setup.

The archaeological and linguistic implications of the Shinde *et al.* study are important, and there are several avenues available for further attempting to distinguish historical scenarios using f -statistics that are outside the scope of a methodological study like this one. Some of our observations that are most challenging for the conclusions of Shinde *et al.* are those related to the graphs with four admixture events in **Figure S7b** that fit the Iranian farmer-related ancestry in the Indus Periphery group as deriving partially from the Hajji Firuz Neolithic or Tepe Hissar Chalcolithic-related lineages. *qpAdm* is able to use information from distal outgroups (such as WSHG) not included in the admixture graph modeling exercise revisited here. Leveraging this information might be able to obtain constraints that would further test the key historical conclusions from Shinde *et al.* Non- f -statistic-based methods could also be informative. Finally, we emphasize that the f_4 -statistic cladality tests correcting for the Anatolian farmer-related and Onge-related admixture in the Indus Periphery grouping do continue to provide support for the historical conclusion of Shinde *et al.* (these analyses reject models where the Tepe Hissar or Hajji Firuz groups share genetic drift with the Indus Periphery individuals), with the caveat that they do not correct for the WSHG admixture.

Librado *et al.* 2021. In contrast to the other studies revisited in our work, the admixture graph published by Librado *et al.* (2021) was inferred automatically using *OrientAGraph*. Models with three

(Fig. 3b in that study) and zero to five (Ext. Data Fig. 5a-d) admixture events were shown. The dataset included 10 populations (9 horse populations and donkey as an outgroup) and was based on 7.4 million polymorphic transversion sites with no missing data at the group level. We observed that some groups used for the *OrientAGraph* and *qpAdm* analyses were very broad geographically and temporally (see Table S1 in the original study), and thus we tested two alternative group compositions: the original one and a streamlined one. In the latter case we included individuals from one archaeological site and one archaeological period per group: the Botai, C-PONT, DOM2, ELEN, and NEO-ANA groups were modified in this way, and the CWC, LP-SFR, Tarpan, and TURG groups were left with the composition used in the original paper (Table S7). In addition, 7 individuals with missing data proportion exceeding 80% were removed from the analysis, affecting the donkey outgroup, DOM2, and NEO-ANA groups (Table S7). Since among all possible f_3 -statistics for the 10 populations three were negative (using all sites available for each population triplet, “useallsnps: YES”), we applied the upgraded algorithm for calculating f -statistics, which removed sites with only one chromosome genotyped in any non-singleton population, resulting in the following site counts for the original and modified population compositions: 11,092 and 1,767,419 sites, respectively. The very low number of sites available in the former case is due to the fact that all individuals are pseudohaploid, and that two groups (the donkey outgroup and NEO-ANA) are composed of two individuals, a high-coverage one and a low-coverage one. Thus, just sites genotyped in both donkey individuals and in both NEO-ANA individuals were kept. Considering this problem, we focused on the modified group composition only. We tested a range of model complexities (from 3 to 9 gene flows) and performed 1,000 *findGraphs* topology search iterations per model complexity class.

Unlike all the other admixture graphs we re-evaluate in this study whose fits to the data were evaluated in the published studies using *qpGraph*, the topologies published in Librado *et al.* 2021 (with 3 to 5 admixture events) were not evaluated for statistical goodness-of-fit, and in fact fit the f -statistic data so poorly that even simple statistics show they cannot be correct (Figure 5b, Figure S8a,c,e, Table S1). In this case, the approach of using *findGraphs* to identify alternative topologies with the same number of admixture events that fit the data better is meaningless, as both the published models and the alternative models do not have enough degrees of freedom to accommodate the complexity present in the real data; all models are guaranteed to be wrong. In particular, we found that WR of the published model with 3 admixture events is 23.9 SE (Figure S8a). In this complexity class *findGraphs* found 22 topologically diverse models that fit significantly better than the published one (Table 1, Table S1), but nevertheless have extremely poor absolute fits (from 16.2 to 21.3 SE, see a temporally plausible example in Figure S8b). In the complexity class with 4 admixture events, no model fitting better than the published one was found; however, five alternative models fitting not significantly worse than the published one had lower WR (10 or 12 SE vs. 14.1 SE, Figure S8d). The WR of the published model with 5 admixture events was 6.9 SE (Figure S8e); just two models fitting nominally better and 223 models fitting non-significantly worse than the published model and having similar or higher WR were found (Table 1, Table S1). These results suggest that while *OrientAGraph* was often (but not always) able to find the same tentative global likelihood optimum as *findGraphs*, neither 3 nor 5 admixture events are enough to explain the data since nearly all the groups are admixed.

For this reason, we moved to topology searches in more complex model spaces incorporating 6 to 9 admixture events. Temporally plausible models with even a modest fit (WR between 3 and 4 SE) were encountered only among models with 8 and 9 admixture events (Figure S8j-r). In the complexity class with 8 admixture events, 5 such temporally plausible fitting models were found, with WRs ranging from 3.4 to 3.9 SE (all these models are shown in Figure S8j-l). In the complexity class with 9 admixture events, 11 such models were found, with WRs ranging from 3.4 to 4.0 SE (all these models are shown in Figure S8m-r).

Librado *et al.* 2021 discussed the following inferences relying fully or partially on their published admixture graphs reported in that study (Table 2): 1) NEO-ANA-related admixture is absent in DOM2; 2) DOM2 and C-PONT are sister groups (they form a clade); 3) there is no gene flow

connecting the CWC and the cluster associated with Yamnaya horses and horses of the later Sintashta culture whose ancestry is maximized in the Western Steppe (DOM2, C-PONT, TURG); 4) there was gene flow from a deep-branching ghost group to NEO-ANA; and 5) Tarpan is a mixture of a CWC-related and a DOM2-related lineage.

The simplest temporally plausible and best-fitting (WR = 3.4 SE) model we found (modified group composition, 8 admixture events, **Figure 5b**, **Figure S8j**, upper panels) supports inferences 2 and 4, and is incompatible with inferences 1, 3, and 5 (**Table 2**). This newly found model can be interpreted as follows. There is a trifurcation of three deep lineages: a lineage maximized in Western and Central Europe (up to 100% of ancestry in a Late Paleolithic group from France, LP_SFR), a Western Steppe-specific lineage (up to 55% in TURG), and a Tarpan-specific lineage (22% in Tarpan). Western and Central European horses, represented by LP-SFR, by the majority ancestry in horses found in the Corded Ware culture context (CWC), and by the majority ancestry in wild Neolithic Anatolian horses (NEO-ANA), contributed about half of the ancestry in the Western Steppe groups TURG, C-PONT, and DOM2. The other half of ancestry in the Western Steppe groups is represented by the Western Steppe-specific lineage. That lineage also contributed about 50% of ancestry in wild horses from the Yana Upper Paleolithic site (ELEN), and the other half of ELEN's ancestry is derived from an even deeper lineage. The Botai group is modeled as a mixture of European horses (69%) and Siberian horses (31% ELEN-related ancestry). In contrast to Librado *et al.* (2021), Tarpan is modeled as a mixture of its specific lineage (22%) and a DOM2-related group (78%), and CWC also received ancestry (21%) from a DOM2-related group. All the populations included in the model except for LP_SFR are admixed, and there is evidence of substantial genetic influence from a lineage that was eventually maximized in the Western Steppe (although it did not necessarily originate there) in the ELEN and Botai groups. We consider this model to be plausible from both temporal and geographical perspectives.

We are not arguing here that our 8-admixture-event model represents the true history; in fact, it is highly unlikely to be the truth, given how large the space of all possible admixture events is and how much admixture evidently occurred relating all these groups (which makes finding the unique truly fitting model extremely unlikely based on *f*-statistic fitting, see the results on simulated data in **Figure 2b**). We have also not attempted in any way to replicate the admixture graph exploration procedure performed in the Librado *et al.* study; the graph fitting procedure was quite different from ours, based on *OrientAGraph* optimization rather than *findGraphs* optimization, and a Block Jackknife procedure with a different genome block size for determining standard errors (4 Mbp in our protocol and ca. 500 kbp in the Librado *et al.* study). Regardless of how the graph was obtained, it is valuable for providing readers with guidance about which topological features of the graphs are meaningful and stable, and which are less certain, especially—as in the case of the admixture graph presented in the paper—when some features of the presented model cannot be right, as evident by the WR of 6.9 in the published model for 5 admixture events. Our set of 16 temporally plausible and fitting (WR < 4 SE) models with 8 or 9 admixture events (**Figure S8j-r**) is consistent with some features of the published graph being stable: the features (2) that DOM2 and C-PONT are sister groups, and (4) that there was a gene flow from a deep-branching ghost group to NEO-ANA (**Table 2**).

Equally important, however, is our finding that there are plausible models that are inconsistent with other inferences in Librado *et al.* (**Table 2**). For example, 13 of these 16 models are inconsistent with the suggestion that there was no gene flow connecting the CWC and the cluster maximized in the Western steppe (DOM2, C-PONT, TURG) (**Figure S8j-r**). In the 8-admixture-event best-fitting plausible model (**Figure 5b**, **Figure S8j**, upper panels), CWC actually derives appreciable ancestry from the early domestic horse lineage DOM2 associated with the Sintashta culture to the exclusion of the more distant Yamnaya-associated TURG and C_PONT horses. This scenario presents a parallel to the one observed in humans, with individuals associated with the CWC receiving admixture from Steppe pastoralists albeit in different proportions: ~75% for humans, versus ~20% in horses. These models specifying a substantial Steppe horse contribution to CWC horses would weaken support for

the inference in Librado *et al.* that “Our results reject the commonly held association between horseback riding and the massive expansion of Yamnaya steppe pastoralists into Europe around 3000 BC”. We are not aware of other lines of evidence in the paper (apart from the fitted admixture graph) that support the claim of no Yamnaya horse impact on CWC horses.

Another example of a feature of the published graph that turned out to be unstable is the model for the Tarpan horse. Only 8 of 16 temporally plausible and fitting models (**Figure S8j-r**) support the conclusion by Librado *et al.* that the Tarpan is a mixture of a DOM2-related and a CWC-related lineage. The other 8 models suggest that Tarpan is a mixture of a deep lineage and a DOM2-related lineage (**Figure 5b, Figure S8j**, upper panels), echoing a hypothesis that Tarpan may be a hybrid with Przewalski’s horses not represented in the admixture graph (Librado *et al.* 2021).

Again, we are not arguing here that our fitting alternative model is right—indeed we are nearly certain it is wrong in important aspects—but we are merely pointing out that the complexity of the admixture graph space means that qualitatively quite different conclusions are compatible with the genetic data. Other aspects of the Librado *et al.* study, most notably the dramatic geographic expansion of the DOM2 modern domestic horse lineage after 4000 years ago in association with the Sintashta culture which is the most extraordinary finding of Librado *et al.*, are in no way challenged by our results.

Hajdinjak *et al.* 2021. The admixture graph inferred by Hajdinjak *et al.* was constructed manually on the basis of a SNP set derived from in-solution enrichment of two SNP panels (1240K + a further million transversion polymorphisms discovered as polymorphic within one or two sub-Saharan African individuals or among archaic humans) and incorporated 11 groups and 8 admixture events (Figure 2d in the original study). The published graph has no clear outgroup since the deepest branch (Denisovan) is admixed. This property of the graph makes automated graph space exploration difficult. We explored two topology classes: 1) 11 groups with 8 admixture events, the original SNP set, Denisovan assigned as an outgroup only at the stage of generating random starting graphs (gene flows to/from the Denisovan branch were allowed at the topology optimization step); and 2) 12 groups with 8 admixture events, chimpanzee added and the original SNP set changed due to the zero missing rate condition, and chimpanzee assigned as an outgroup at both algorithm stages (**Table S1**). For both graph complexity classes, two topology search settings were tested: 1) either no additional constraints were applied beyond the outgroup constraints described above, or 2) the Vindija Neanderthal and Mbuti were allowed to have no admixture events in their history, and the Denisovan lineage was allowed to have up to one admixture event in its history (these constraints were in line with the model in the original study and with literature on the genetic history of archaic humans, e.g., Prüfer *et al.* 2014). The composition of the groups matched that in the original study, as did the parameter settings for *qpGraph*, with the exception of “least squares mode”, which was used in the original study, but not in our analysis. “Least squares mode” computes LL scores without taking into account the f -statistic covariance matrix, and we confirmed that changing this parameter does not qualitatively change our results. Since no f_3 -statistics were negative when all sites available for each population triplet were used (the “useallsnps: YES” option), we did not use the upgraded algorithm for calculating f_3 -statistics on pseudo-diploid data. We summarize results across 2,000 to 4,000 independent iterations of the *findGraphs* algorithm (**Table S1**).

When chimpanzee was not included into the analysis and no topology constraints were applied, nearly all newly found models turned out to be distinct (3,996 of 4,000), nearly all (96.8%) fit nominally better and 15.9% fit significantly better than the published model (**Table S1**), and absolute fits of 91.3% of novel models are good ($WR < 3 SE$). Similar results were obtained when the topology search algorithm was constrained: nearly all (89.5%) of 1,999 newly found models fit nominally better and 26% fit significantly better than the published model (**Table S1**).

When chimpanzee was set as an outgroup and no topology constraints were applied, the picture remained similar. Nearly all newly found models turned out to be distinct (1,996 of 2,000), and a very large fraction of them (56.8%) fit significantly better than the published model (**Table S1**); 16.4% of novel models demonstrated $WR < 3$ SE. Similar results were obtained when the topology search algorithm was constrained: most (71.4%) newly found models fit nominally better and 15.7% fit significantly better than the published model (**Table 1, Table S1, Figure 4**), which has a poor absolute fit on this set of sites and groups ($WR = 4.8$ SE, **Figure 5c, Figure S9**). The statistics described above and the fact that LL scores on all sites lie outside of the LL distribution on resampled datasets (**Figure 4**) suggest that models in this complexity class are overfitted, but the published topology emerged as fitting relatively poorly.

Overfitting arises naturally during manual graph construction as performed in many studies (not only in Hajdinjak *et al.* (2021), but also in Fu *et al.* 2016, Skoglund *et al.* 2016, Yang *et al.* 2017, Posth *et al.* 2018, McColl *et al.* 2018, Moreno-Mayar *et al.* 2018, Tambets *et al.* 2018, van de Loosdrecht *et al.* 2018, Flegontov *et al.* 2019, Sikora *et al.* 2019, Wang *et al.* 2019, Lipson *et al.* 2020, Shinde *et al.* 2019, Yang *et al.* 2020, and Wang *et al.* 2021). The graph grew one group at a time, and each newly added group was mapped on to the pre-existing skeleton graph as unadmixed or as a 2-way mixture. This imposed constraints on the model-building process. Another constraint imposed was the requirement that all intermediate graphs have good absolute fits (WR below 3 or 4 SE). When the model-building process is constrained in a particular path and fits of all intermediates are required to be good, unnecessary admixture events are often added along the way, and the resulting graph belongs to a complexity class in which models are overfitted and many alternative models fit equally well. There is no single obviously correct order of adding branches to a growing graph. For example, the Kostenki and Sunghir lineages were included into the initial graph (Fig. S6.1 in the original study) as unadmixed lineages, and their admixture status was not revisited at subsequent steps (unlike that of Tianyuan and Ust'-Ishim), except for adding the archaic gene flow common for non-Africans. For that reason, the published graph differs from many alternative better-fitting and temporally plausible graphs where the Kostenki and Sunghir lineages are modeled as more complex mixtures (**Figure S9**).

Hajdinjak *et al.* (2021)'s published graph had the following notable features that were interpreted by the authors and used to support some conclusions of the study (**Table 2**): 1) there are gene flows from the lineage found in the ~45,000-43,000-years-old Bacho Kiro Initial Upper Paleolithic (IUP) individuals to the Ust'-Ishim, Tianyuan, and GoyetQ116-1 lineages; 2) the ~35,000-years-old Bacho Kiro Cave individual BK1653 belonged to a population that was related, but not identical, to that of the GoyetQ116-1 individual; and 3) the Vestonice16 lineage is a mixture of a Sunghir-related and a BK1653-related lineage. To assess if these features are supported by our re-analysis, we focused on our most constrained *findGraphs* run: with chimpanzee set as an outgroup and with the topology constraints applied at the topology search step. We identified 1,421 topologies fitting nominally or significantly better than the published model and satisfying the constraints and moved on to inspect 50 best-fitting topologies for temporal plausibility (all of them fitting significantly better than the published model). All non-African individuals included in the model are Upper Paleolithic and their dates are not drastically different in relative terms: from ca. 45 kya (thousand years before present) for some Bacho Kiro IUP individuals (Hajdinjak *et al.* 2021) to ca. 30 kya for the Vestonice16 individual (Fu *et al.* 2016). Nevertheless, we considered most gene flows from later-attested to earlier-attested lineages as temporally implausible (for instance, GoyetQ116-1 (~35 kya) → Ust'-Ishim (~44 kya), GoyetQ116-1 (35 kya) → Bacho Kiro IUP (45-43 kya), Kostenki14 (38 kya) → Ust'-Ishim (44 kya), Sunghir III (34.5 kya) → Tianyuan (40 kya), Vestonice16 (30 kya) → Tianyuan (40 kya)) since they imply great antiquity of the later-attested lineages, e.g., >40 kya for the Vestonice16 lineage, and even greater antiquity for the related lineages such as Sunghir III and Kostenki14.

Of the 50 topologies inspected, 32 were considered temporally plausible. Of those topologies, none supported feature 1 of the published admixture graph (there is no replication of the finding of gene

flows from the Bacho Kiro IUP lineage specifically into all three of the Ust'-Ishim, Tianyuan, and GoyetQ116-1 lineages). One topology supported features 2 and 3, and partially supported feature 1 (there was Bacho Kiro → GoyetQ116-1 gene flow, but no Bacho Kiro → Tianyuan and Bacho Kiro → Ust'-Ishim gene flows). A total of 17 topologies supported features 2 and 3 but were inconsistent with feature 1; and 14 topologies supported feature 3 only (**Table 2**). Best-fitting representatives of each of these topology classes are shown along with the published model in **Figure S9**. Considering topological diversity among models that are temporally plausible, conform to current knowledge about relationships between modern and archaic humans, and fit significantly better than the published model, we conclude that feature 3 is probably robust but other details of the fitted admixture graph in Hajdinjak *et al.* (Figure 2d of that study)—for example, gene flows to the Ust'-Ishim, Tianyuan and Goyet Q116-1 lineages from sources sharing drift exclusively with the Upper Paleolithic Bacho Kiro lineage—should not be interpreted as providing meaningful inferences about population history of Upper Paleolithic modern humans. For example, the upper right-hand alternative model plotted in **Figure S9c** supports features 2 and 3 but includes no gene flows from the Bacho Kiro IUP lineage.

A central finding of Hajdinjak *et al.* is that the Bacho Kiro IUP group shares more alleles with present-day East Asians than with Upper Paleolithic Holocene Europeans despite coming from Europe. Specifically, the study documents significantly positive statistics of the form $D(\text{an Asian group, Kostenki14; Bacho Kiro IUP, Mbuti})$ (Fig. 2b and Extended Data Fig. 5 in the original study). For example, $D(\text{Tianyuan, Kostenki14; Bacho Kiro IUP, Mbuti})$ is significantly positive ($D = 0.0032$, $SE = 0.0010$, $Z = 3.2$) on the dataset used for testing the twelve-population graphs (263,698 sites without missing data across all 12 groups). The same statistic is also significantly positive ($D = 0.0029$, $SE = 0.0006$, $Z = 4.4$) when all 1,312,292 non-missing sites in the population quadruplet are analyzed. Hajdinjak *et al.*'s interpretation of this observation, using the language from the abstract, is that “there was at least some continuity between the earliest modern humans in Europe [Bacho Kiro IUP] and later people in Eurasia [East Asians]”.

However, a significant D -statistic can have multiple explanations. The statistic $f_4(\text{Tianyuan, Kostenki14; Bacho Kiro IUP, Mbuti})$ is fitted equally well by the published twelve-population admixture graph (Z -score for the difference between the observed and fitted statistics = 0.64) and by, for example, the lower left-hand graph in **Figure S9c** (Z -score = 0.94) reproduced in **Figure 5c**. Under the latter model that fits the data significantly better than the published model (p -value = 0.02), the Bacho Kiro IUP and Tianyuan branches are not connected by a gene flow and do not receive gene flows from a third common source, but the common ancestor of Ust'-Ishim and all European Paleolithic lineages receives an 8% gene flow from a divergent modern human lineage splitting deeper than Bacho Kiro IUP and Tianyuan (**Figure 5c**, **Figure S9c**). This scenario or some version of it seems archaeologically and geographically plausible and is not disproven by any other line of genetic or non-genetic evidence of which we are aware. It could correspond to a scenario where a primary modern human expansion out of the Near East contributed serially to the major lineages leading to Bacho Kiro, then later East Asians, then Ust'-Ishim, and finally the primary ancestry in later European hunter-gatherers. This has a very different interpretation from the scenario of distinctive shared ancestry between the earliest modern humans in Europe such as Bacho Kiro IUP and later people in East Asia—to the exclusion of later European hunter-gatherers—that is suggested by the Hajdinjak *et al.* published graph.

We are not claiming that this specific alternative model is correct—indeed, it is almost certainly not the correct one given the topological complexity of the set of all admixture graphs consistent with the data—but the existence of it and many other models that fit the data makes it clear that we do not yet have a unique historical explanation for the excess sharing of alleles that has been documented between some Upper Paleolithic European groups (Bacho Kiro IUP, Hajdinjak *et al.* 2021, and GoyetQ116-1, Yang *et al.* 2017 and Hajdinjak *et al.* 2021) and all East Asians.

Lipson et al. 2020. The admixture graph in the original study (Lipson et al. 2020) was constructed manually based on a SNP set derived from the 1240K enrichment panel, and the final model was alternatively tested on the combined HumanOrigins sub-panels 4 and 5 (each ascertained on one African individual) or on sites ascertained as polymorphic in archaic humans. The final published model (Extended Data Fig. 4 in that study) is very complex (12 groups and 12 admixture events): it exists in a space of $\sim 10^{44}$ topologies of this complexity. We note that one admixture event was added by Lipson et al. (2020) to account for potential modern DNA contamination in ancient Shum Laka individuals, and removing it caused a negligible difference in the fit of the published model (**Table S1**). Thus, to decrease the complexity of the graph search space, we considered graphs with 12 groups and 11 admixture events. Twenty-two f_3 -statistics for these 12 groups turned out to be negative (when the “useallsnps: YES” setting was used), and thus for exploring this graph complexity class we had to remove sites with only one chromosome genotyped in any non-singleton population (**Table S1**). The following constraints were applied during the topology search: chimpanzee was assigned as an outgroup at both stages of the process (while generating random starting graphs and while searching the topology space); Altai Neanderthal was required to be unadmixed; and non-Africans (French) were required to have at least one admixture event in their history. The composition of the groups we analyzed matched that in the original study. We summarize results across 2,000 independent iterations of the *findGraphs* algorithm.

All newly found models turn out to be distinct (2,000), and 11.9% fit nominally (but not significantly) better than the published model (**Table 1, Table S1, Figure 4**). Absolute fits of 36.7% of novel models are good ($WR < 3$ SE). Fits of the highest-ranking model and the published model are not significantly different according to the bootstrap model comparison method (p -value = 0.176). These metrics, along with the fact that LL scores on all sites lie outside of the LL distribution on resampled datasets (**Figure 4**), suggest that models in this complexity class, including the published model, are overfitted. Of the admixture graphs we re-evaluate in this study, Lipson et al. 2020 shares with Hajdinjak et al. 2021, Sikora et al. 2019, and Wang et al. 2021, evidence of being overfitted (**Figure 4**).

We also wanted to check if overfitting would be found in the graph complexity classes corresponding to two simpler intermediate graphs from the original study (**Table S1**): 7 groups and 4 admixture events (Figure S3.24 in that study) and 10 groups and 8 admixture events (Figure S3.25 in that study). The population composition of the dataset we used for this analysis was slightly different from the dataset used by Lipson et al.: the ancient South African hunter-gatherer group was replaced by a related group (present-day Juǀ'hoan North), and instead of the Shum Laka ancient group, only one high-coverage individual from the same group (I10871) was used. We summarize results across 2,000 or 10,000 independent iterations for each SNP set, for the small and large graphs, respectively. For 7 groups, we found 201 novel topologies fitting better than the published one, and for 10 groups we found nearly 9,000 such topologies (**Table S1**). In the latter case 6.8% of newly found topologies fit significantly better than the published topology. For the more complex graph class with 10 groups and 8 admixture events we also found evidence of overfitting: the LL score of the published graph run on the full data is better than almost all the bootstrap replicates on the same data (it falls below the 5th percentile).

Below we discuss selected prominent features of the admixture graph published in the original study (that were interpreted by the authors and used to support some conclusions of the study) and the extent to which these features consistently replicate across the large number of fitting 12-population graphs with 11 admixture events (**Table 2**): 1) A lineage maximized in present-day West African groups (Lemane, Mende, and Yoruba) also contributed some ancestry to the ancient Shum Laka individual and to present-day Biaka and Mbuti; 2) another ancestry component in Shum Laka is a deep-branching lineage maximized in the rainforest hunter-gatherers Biaka and Mbuti; 3) “super-archaic” ancestry (i.e., diverging at the modern human/Neanderthal split point or deeper) contributed to Biaka, Mbuti, Shum Laka, Lemane, Mende, and Yoruba; and 4) a ghost modern human lineage (or lineages) contributed to Agaw, Mota, Biaka, Mbuti, Shum Laka, Lemane, Mende,

and Yoruba. We identified 232 twelve-population topologies that fit nominally better than the published one, 34 best-fitting topologies (of 232) were manually assessed for temporal plausibility, and we focus on 30 topologies identified as temporally plausible and including a low-level Neanderthal contribution ($\leq 10\%$) in non-Africans (French). These 30 topologies are shown along with the published model in **Figure S10**.

In this set of alternative models, high topological diversity is observed (see an example in **Figure 5d** and further topologies in **Figure S10**). We classified the topologies as follows. If an ancestral lineage defined above (for example, a deep-branching lineage maximized in rainforest hunter-gatherers Biaka and Mbuti) exists in the graph, we compared the sets of populations where it is found in the published model and in the model examined. If there was no more than one population where the ancestry is expected according to the published model but not present, or present but not expected, we considered this feature of the published graph supported by the alternative graph. If no ancestral lineage meets the definition above, the feature of the published graph was considered not supported. In all other cases partial support for the feature was declared (**Figure S10**). Considering extreme cases, two alternative graphs completely lacked support for three features of the published graph (**Figure 5d**, **Figure S10c**), and one graph supported all four features of the published graph fully (**Figure S10q**, bottom panels). There are some graphs where defining two distinct ancestral lineages maximized in West Africans and in Mbuti and Biaka (features 1 and 2) is essentially impossible since all or nearly all Africans are modeled as a mixture of at least two deep lineages (see graph no. 4, **Figure S10d**). In some graphs there is no single lineage specific to rainforest hunter-gatherers (Biaka, Mbuti, and Shum Laka) since the primary ancestries in these groups form independent deep branches in the African graph (see graph no. 2 in **Figure 5d** and graph no. 16 in **Figure S10j**, bottom panels). The ghost modern and super-archaic gene flows to Africans also had no universal support in the set of alternative graphs we examined (see, for example, **Figure 5d** and **Figure S10c**).

Considering the high degree of topological diversity among models that are temporally plausible, conform to known findings about relationships between modern and archaic humans, and fit nominally better than the published model, we conclude that the features from the original study are not supported by our re-analysis (**Table 2**). As in the case study above, the published manually constructed model is a representative of a large class of models that are equally well fitting to the limits of our resolution. This situation may be attributed to 1) overfitting and/or to 2) the lack of information in the dataset (in the combination of groups and SNP sites) and/or to 3) inherent limitations of *f*-statistics, when distinct topologies predict identical *f*-statistics.

In reconsidering the findings of Lipson *et al.* 2020 it is important to keep in mind that analysis of allele frequency correlation statistics is not the only type of information that can be used to make inferences about population relationships in deep time. Other methodologies have provided important insights into deep African population history, and the model-building in Lipson *et al.* 2020 was guided in an informal way by these other lines of evidence. For example, unknown archaic lineages admixing into some African populations were hypothesized through identification of deeply splitting haplotypes that are too long to have been freely mixing with other haplotypes in present-day populations for all of their history (Hammer *et al.* 2011, Lachance *et al.* 2012, Speidel *et al.* 2019). Similarly, analysis of haplotype divergence times of pairs of populations has been used to provide evidence of an early radiation of modern human lineages maximized today in southern African hunter-gatherers, Mbuti rainforest hunter-gatherers, and the great majority of other present-day populations; and a later split of lineages related to East African hunter-gatherers, West African agriculturalists, and non-Africans, which is a feature of the Lipson *et al.* model (Campbell and Tishkoff 2008, Mallick *et al.* 2016). Notably, some alternative models we found do not contradict the above-mentioned results and are profoundly different from the published model at the same time (see, for example, **Figure 5d**). These constraints are not enough, however, to provide evidence for all the topological details of the Lipson *et al.* 2020 admixture graph highlighted in this section, or for other features of Lipson *et al.* 2020 that were not invoked in the previous literature and newly

proposed in that study, such as the “ghost modern” lineage splitting around the same time as the lineages leading to southern African hunter-gatherers and central African rainforest hunter-gatherers and mixing in highest proportion to Ethiopian hunter gatherers and to a lesser proportion to West Africans, and the “basal West African” lineage that contributes uniquely to Shum Laka. Many of the models that emerged as good fits in our admixture graph building exercise as the published one did not share some of these features (**Figure S10**).

The high diversity of well-fitting admixture graph models that satisfy known constraints relating diverse African populations highlights the need for further research based on multiple lines of genetic analysis (in addition to allele frequency correlation patterns) to obtain further insights into deep African history. Our results particularly highlight the mystery around the highly distinctive genetic ancestry of the Shum Laka individuals themselves, who represent the newly reported data in the Lipson *et al.* 2020 study, and a highly important set of genetic datapoints that was not available prior to the study. The ancestral relationships of these four individuals to both rainforest hunter-gatherers, and to the primary lineage in present-day West Africans, remains an open question, one whose resolution promises meaningful new insights into modern human population history.

Wang et al. 2021. The admixture graph inferred by Wang *et al.* 2021 was constructed manually on the basis of a SNP set derived from the 1240K enrichment panel. We focused our analysis on the final graph (Extended Data Figure 6 in Wang et al. 2021, 12 groups and 8 admixture events) and on two simpler intermediates in the model building process (Figures S13-9 and S13-10a in Wang et al. 2021). To simplify the latter two models further, we removed a low-level gene flow (1%) from a West European hunter-gatherer-related lineage (Loschbour) to the Mongolia Neolithic group, which resulted in negligible LL differences (0.5 and 2.4 log-units, respectively). Thus, using *findGraphs* we explored the following topology classes: 9 groups with 4 admixture events, 10 groups with 5 admixture events, and 12 groups with 8 admixture events (**Table S1**). The composition of the groups matched that in the original study. We summarized results across 2,000 independent iterations of the *findGraphs* algorithm for each topology class. In the case of the most extensive population set (12 groups), three f_3 -statistics turned out to be negative (when the “useallsnps: YES” setting was used), and thus for exploring this graph complexity class we had to remove sites with only one chromosome genotyped in any non-singleton population (**Table S1**). For this complexity class, we also applied several constraints on the graph space exploration process all of which were shared with the Wang *et al.* graphs: the Denisovan genome was assigned as an outgroup in the random starting graphs, but not at the topology search stage; up to one admixture event was allowed in the history of the Denisovan group; no admixture events were allowed in the history of Mbuti, Loschbour, and Onge; and the (Denisovan, (Mbuti, (Loschbour, Onge))) branching order was required.

For each topology class we found hundreds to thousands of topologically unique graphs fitting nominally better than the published models (**Table 1, Table S1**). For both simple topology classes, no model fitting significantly better than the published one was found (**Table S1**). However, the final published model fits the data significantly worse than 12.6% of newly found models of the same complexity (**Table 1, Table S1**). The fact that many topologically diverse models had good absolute fits (65%, 55%, and 15% of distinct newly found graphs with 9, 10, and 12 groups, respectively, had $WR < 3$ SE) suggests that admixture graph models in these complexity classes are overfitted. Further evidence of overfitting comes from the poor fits of the published model on bootstrap-resampled datasets as compared to their fits on all sites (**Figure 4**).

An important feature of the published graphs in Wang *et al.* 2021 that was remarked upon in the study is admixture from a source related to Andamanese hunter-gatherers that is almost universal in East Asians, occurring in the Jomon, Tibetan, Upper Yellow River Late Neolithic, West Liao River Late Neolithic, Taiwan Iron Age, and China Island Early Neolithic (Liangdao) groups (**Table 2**). For example, the abstract states “Hunter-gatherers from Japan, the Amur River Basin, and people of

Neolithic and Iron Age Taiwan and the Tibetan Plateau are linked by a deeply splitting lineage that probably reflects a coastal migration during the Late Pleistocene epoch.” We performed 2,000 *findGraphs* iterations and obtained 1,778 distinct topologies satisfying all the constraints, nearly all of them (1,724) fitting nominally better than the published model, and 12.6% fitting significantly better (**Table S1**). The models were ranked by LL, and 56 highest-ranking topologies, all of them fitting significantly better than the published one, were assessed for temporal plausibility (models with gene flows from a later group to Tianyuan dated to 40 kya were removed), and 20 topologies were considered temporally plausible (all of them are shown in **Figure S11**). According to these topologies, 0 to 2 East Asian groups had a fraction of their ancestry derived from a source specifically related to Onge, and 19 topologies included gene flows from the European (Loschbour)-related branch to all 8 East Asian groups (**Figure S11**). The inferred topological relationships among East Asians are variable in this group of 20 models, and we decided to apply further constraints that guided model ranking and elimination by Wang *et al.*, based on considerations from archaeological evidence, Y chromosome haplogroup divergence patterns, and population split time estimation.

The constraints that are not based on correlation of allele frequencies across populations that Wang *et al.* applied and that we applied in our re-examination are as follows. First, combined evidence from archaeology, linguistics, and genetics (a closely shared Y chromosome haplogroup) suggests that the present-day Tibetan Plateau population harbors a substantial proportion of ancestry from a large-scale migration from the Neolithic farming groups from the Upper and Middle Yellow River (Chen *et al.* 2015, Lu *et al.* 2016, Zhang *et al.* 2019). These arguments and radiocarbon dates favor the following branching order of predominant ancestry components: (Mongolia East N, (China Upper YR LN, Nepal Chokhopani)). Second, evidence from archaeology, linguistics and genetics suggests that the expansion of Austronesian speakers and the peopling of Taiwan was from southeast coastal China to Taiwan and Southeast Asia, but not from Taiwan to mainland China (Bellwood 2011, Gray and Jordan 2000, Ko *et al.* 2011). These arguments make a China Island EN → Taiwan IA gene flow direction plausible and make the opposite direction of flow less likely. Third, in the original study MSMC cross-coalescence rates were computed for a few pairs of present-day proxies for the ancient groups, and it was argued that they impose constraints on the graph topology. The inferred coalescence date for the Tibetan and Ulchi groups was slightly younger than the Tibetan-Ami and Tibetan-Atayal dates (Fig. S13-1 in the original study), suggesting that the Nepal Chokhopani and Mongolia East N group may share ancestral source populations more recently than these two groups and Taiwan IA. We note that it was not clear in the original paper if the difference in coalescence dates is statistically significant, the finding was clearer in MSMC than in MSMC2 analysis, and there was no attempt to calculate expected cross-coalescence profiles using these methods from models incorporating many gene flows. Nevertheless, we applied this constraint as well in an attempt to understand whether, if we used a constraint system similar to that in Wang *et al.*, we would obtain results that agreed with respect to the finding of Onge-related admixture ubiquitous among East Asian groups.

Applying these three additional constraints, we identified two models (among the 56 ones subjected to manual inspection) that satisfied all of them. The highest-ranking of those models is shown in **Figure 5e** and **Figure S11c** (lower panels), and it includes a 13% (deeply) European-related gene flow to the common ancestor of all East Asians, and gene flows from the Onge-related branch to just two East Asian groups: Nepal Chokhopani and China WLR LN. This model fits the data significantly better than the published model (p -value = 0.028). We do not claim that this is the correct model (indeed we are almost certain that it is not given the high degeneracy of fitting models), but it is not obviously wrong and differs in qualitatively important ways from the published one.

The Wang *et al.* 2021 admixture graph provides an illuminating example that helps us to understand the value added by admixture graph construction. The admixture graph construction process in Wang *et al.* followed a philosophy of not relying entirely on the allele frequency correlation data (not treating the genetic data as independent to explore how much new insight could come from genetic data alone). Instead, the study integrated other lines of genetic evidence as well as linguistic

and archaeological insights explicitly into the admixture graph construction process, with the goal of identifying models consistent with multiple lines of evidence. The fact that after this procedure a fitting graph was obtained is not of great interest, as it is essentially always possible to obtain a fit to allele frequency correlation data when enough admixture events are added. The important question is whether any of the emergent features of the graph that were not applied as constraints in the construction process—for example the evidence of ubiquitous Andamanese-related gene flow throughout East Asia suggesting a coastal route expansion that admixed with an interior route expansion proxied by Tianyuan—were stably inferred. Our analysis does not come to this finding consistently among well-fitting and plausible admixture graphs. We conclude that an important feature of the published graph, i.e. variable levels of Andamanese-related ancestry found in all East Asians except for Siberians (Mongolia Neolithic) and the Upper Paleolithic Tianyuan (Fig. 2 in (Wang et al. 2021)), is not supported by f -statistic analysis alone (**Table 2**), and indeed we are not aware of a single feature of the Wang *et al.* 2021 admixture graph that is stably inferred beyond the constraints applied to build it.

Sikora et al. 2019. Two admixture graphs inferred by Sikora *et al.* (2019) were constructed manually based on a SNP set derived from whole-genome shotgun data and incorporated 12 or 13 groups and 10 admixture events (Extended Data Figure 3f in the original study). One graph was focused on West Eurasians, and the other one on East Eurasians, and both included a Neanderthal, a Denisovan, and an African group (Dinka). Although the chimpanzee outgroup was not included in the original graphs, we added it as it drastically constrains the topology search space. The following additional constraints were applied at the *findGraphs* model optimization stage: the Neanderthal and African groups were unadmixed and the Denisovan group had no more than one admixture event in its history. These three constraints match the features of the published graph. We also repeated topology searches without constraining the admixture status of the Neanderthal, Denisovan, and Dinka. Since no f_3 -statistics were negative when all sites available for each population triplet were used (the “useallsnps: YES” option), we did not apply the algorithm that allows unbiased calculation of f_3 -statistics on pseudo-diploid data at the expense of loss of analyzed SNPs.

In contrast to most other published graphs discussed above, gene flows in the graphs inferred by Sikora *et al.* do not have equal standing: four low-level gene flows (0-1%) connect the Neanderthal lineage to Upper Paleolithic lineages (Kostenki, Sungir, Yana, Ust'-Ishim in the ‘Western’ graph and Sungir, Yana, Mal'ta, Ust'-Ishim in the ‘Eastern’ graph). We repeated each topology search under two alternative settings: either keeping the number of admixture events at 10 to match the published graphs, or at 6 to match simplified versions of the published graphs lacking these low-level Neanderthal gene flows. We performed that modification to simplify the search space and to alleviate the overfitting problem which becomes severe if 10 gene flows across the graph are allowed (**Table S1**). Here we compare LL and WR for the original published models and their simplified versions: the Western graph including chimpanzee (LL = 65.7, WR = 3.32 SE) vs. its simplified version (LL = 76.5, WR = 3.78 SE) and the Eastern graph including chimpanzee (LL = 85.3, WR = 3.11 SE) vs. its simplified version (LL = 102.4, WR = 4.16 SE). In both cases we found no statistically significant differences in model fits (relying on the bootstrap model comparison method). In summary, topology search was repeated under 8 settings: for the Western or Eastern graphs, with no constraints on the admixture status or with the constraints specified above, and with 10 or 6 gene flows (**Table S1**). Below we focus on results for constrained models with 6 admixture events. In contrast, **Figure 4** and **Table 1** show results for constrained Western graphs with 10 admixture events.

In the case of the constrained Western graphs with 6 admixture events, 1,000 topology search iterations were performed, 894 distinct topologies were found, 4 models fit significantly better, and 151 models fit nominally better than the published one (**Table 1**, **Table S1**). We inspected those 155 topologies and identified 29 topologies (**Figure S12**) that are temporally plausible and include no non-canonical gene flows from archaic groups such as Denisovan or gene flows ghost archaic →

non-Africans. Sikora *et al.* came to the following striking conclusion relying on the Western admixture graph (**Table 2**): the Mal'ta (MA1_ANE) lineage received a gene flow from the Caucasus hunter-gatherer (CaucasusHG_LP or CHG) lineage. However, in our *findGraphs* exploration this direction of gene flow (CHG → Mal'ta) was supported by two of the 29 topologies, and the opposite gene flow direction (from the Mal'ta and East European hunter-gatherer clade to CHG) was supported by the remaining 27 plausible topologies (**Figure S12**). The highest-ranking plausible topology (**Figure 5f**) has a fit that is not significantly different from that of the simplified published model (p -value = 0.392). We note that the gene flow direction contradicting the graph by Sikora *et al.* was supported by a published *qpAdm* analyses (Lazaridis *et al.* 2016, Narasimhan *et al.* 2019), and *qpAdm* is not affected by the same model degeneracy issues that are the focus of this study. Considering the topological diversity among models that are temporally plausible, conform to robust findings about relationships between modern and archaic humans, and fit nominally better than the published model, we conclude that the direction of the Mal'ta-CHG gene flow cannot be resolved by admixture graph analysis (**Table 2**).

Some important conclusions based on the Eastern graph also do not replicate across all plausible admixture graphs (**Table 2**). In the case of the constrained Eastern graphs with 6 admixture events, 4,446 topology search iterations were performed, and 2,785 distinct topologies were found. Only 3 topologies fit significantly and 13 nominally better than the published one (p -value for the highest-ranking newly found model vs. the simplified published model = 0.112), and 9.8% of topologies fit not significantly worse than the published one (**Table 1, Table S1**). Of the topologies belonging to these groups, we inspected 116 best-fitting ones and identified 97 topologies that are temporally plausible and include no gene flows from archaic groups such as Denisovan or ghost archaic → non-Africans that are qualitatively different from the gene flows that are currently widely accepted. The Sikora *et al.* Eastern admixture graph had the following distinctive features that were used to support some conclusions of the study (**Table 2**): 1) the Mal'ta (MA1_ANE) and Yana (Yana_UP) lineages receive a gene flow from a common East Asian-associated source diverging before the ones contributing to the Devil's Cave (DevilsCave_N), Kolyma (Kolyma_M), USR1 (Alaska_LP), and Clovis (Clovis_LP) lineages; 2) European-related ancestry in the Kolyma, USR1, and Clovis lineages is closer to Mal'ta than to Yana; 3) the Devil's Cave lineage received no European-related gene flows, and Kolyma has less European-related ancestry than ancient Americans (USR1 and Clovis). Only feature 2 was universally supported by all the 97 plausible alternative models fitting significantly better, nominally better, or not significantly worse than the simplified published model, while feature 3 was supported by 83 of 97 plausible models, and feature 1 was supported by 28 of 97 plausible models (**Table 2**). We plotted 14 plausible graphs as examples of topologies supporting all three features, two features, or one feature of the published graph (**Figure S13**). We note that all the Eastern graphs discussed here, both the published and alternative ones, have relatively poor absolute fits with WR above 4 or 5 SE. Increasing the number of gene flows to 10 allowed us to reach much better absolute fits (with WR as low as 2.42 SE), but that resulted in high topological diversity (on a par with some other case studies discussed above). In the case of the constrained Eastern graphs with 10 admixture events, 1,000 topology search iterations were performed, and 1,000 distinct topologies were found. Of these topologies, 13.2% fit significantly better, 30% nominally better, and 17.6% non-significantly worse than the published model (p -value for the highest-ranking newly found model vs. the published model < 0.002) (**Table S1**).

A Proposed Protocol for Using Admixture Graph Fitting in Genetic Studies

Admixture graphs represent a conceptually powerful framework for thinking about demographic history, but the practice of manually constructing a small number of complex models without exploring admixture graph space in an automated way can lead to overconfidence in the validity of these models. An ideal outcome of an admixture graph model exploration exercise would be the identification of a model or a group of topologically very similar models which fits the data well and significantly better than all alternative models with the same number of admixture events; however,

this is almost never achieved for graphs with more than eight populations and three admixture events in our experience, and even this approach can lead to potentially unstable results as relaxing the assumption of parsimony (that fewer admixture events is more likely) can lead to qualitatively quite different equally well fitting topologies as in our reanalysis of the Bergström *et al.* and Shinde *et al.* datasets. Most of the examples of admixture graphs in eight recently published studies we revisited do not fit this ideal pattern, as we were able to identify many topologically different alternative models that could not easily be rejected based on temporal plausibility or other constraints (Figures S5-S13). In particular, for all studies except Shinde *et al.* 2019 (under a strict parsimony assumption however), we identified admixture graphs that were not significantly worse fitting than the published ones, and with topological features that were different in qualitatively important ways. There were also some more encouraging findings of the exercise we performed to re-evaluate published models. For example, at least one of the key inferences about population history relying on the admixture graph modeling were stable for all analyzed models for the Lazaridis *et al.*, Librado *et al.*, Hajdinjak *et al.*, Shinde *et al.* 2019 (under the parsimony assumption), and Sikora *et al.* (simplified Eastern graph) studies. The existence of some stable features in these graphs helps to point the way toward a protocol that we believe should be applied in all future studies that use admixture graph fitting exercises to support claims about population history.

We propose the following tentative protocol to identify features of fitting admixture graphs that are stable enough to be used to make inferences about population history.

1. For a given combination of populations, carry out an initial scan using *findGraphs* to identify reasonable parameter values for the number of allowed admixture events (the graph complexity class). For example, run *findGraphs* allowing between zero and eight admixture events (100 algorithm iterations per graph complexity class), each iteration saving one or a few best-fitting outcomes. The smallest number of admixture events that yields models where the (negative) LL score or the worst *f*-statistic residual is lower than some threshold can then be explored more deeply by running more iterations of *findGraphs*.
2. Run *findGraphs* on the determined complexity class, where some of the resulting graphs should be inspected manually to determine whether they could in principle be historically plausible models. Implausible models (for example, models where a very ancient population appears to be admixed between two modern populations) can be filtered out by imposing topological constraints. If no or only a few graphs remain, *findGraphs* can be run again under these constraints. This can be repeated until one or more graphs with an acceptable LL score or worst residual has been identified. At this stage we apply the bootstrap method to determine whether the best-fitting graph is significantly better than the next best-fitting graph. If it is not, we identify a set of graphs which are not clearly worse than the best-fitting graph by performing the bootstrap model comparison for many model pairs.
3. Researchers should compare the resulting graphs to each other with the goal of identifying common features. Although *ADMIXTOOLS 2* includes automated tools for cataloguing common topological features (Suppl. Methods), we found a manual approach to be valuable, as the fitted parameters (especially admixture proportions) are as important for this task as graph topology.
4. Once a set of fitting graphs and stable topological features shared between them is identified, researchers should carry out a *findGraphs* exploration of the space of graphs with one additional admixture event. If inferences are stable even when fitting graphs with one more level of complexity than the graphs with the minimal number of admixture events needed to fit the data, this increases confidence in the inferences. Furthermore, the addition of a new population may introduce crucial information to an existing set of populations, which can change the space of fitting topologies in a profound way, as in our reanalysis of the data from Bergström *et al.* 2020 (Figure 5a, Figure S2). Thus, it is advisable that the topology optimization procedure is repeated on several alternative

population sets, in addition to considering models that allow an additional admixture event beyond the minimum required for parsimony, to explore if inferences about topology change qualitatively.

5. Admixture graphs fitted with f -statistics do not distinguish between time and population size as the two sources of genetic drift, and many different complex genetic histories for a set of populations can result in the exact same expected f -statistics. This provides an important opportunity to further constrain a model fitting procedure. Methods that take advantage of information from the site frequency spectra (*mom2*, *fastsimcoal*, Kamm et al. 2019, Excoffier et al. 2013) or derived site patterns, a special case of site frequency spectra (*Legofit*, Rogers 2019), can supply alternative information not captured by f -statistics (further information can come from methods that fit haplotype divergence patterns such as *MSMC* (Schiffels and Durbin 2014) and *SMC++* (Terhorst et al. 2017), or inferences based on fitted gene trees such as *RELATE* (Speidel et al. 2019), and *ARGweaver* (Hubisz et al. 2020, Hubisz and Siepel 2020)). These tools are too computationally intensive to explore a large number of models, but the advantages of the different approaches can be combined by first identifying a set of candidate models using *findGraphs*, and then testing these candidate models with other methods. This approach is also expected to deal with overfitting since different data types almost always include different variable site sets.

We believe that researchers should only begin to make strong claims about population history with admixture graphs once a protocol such as we propose is applied.

We see the guidelines above as analogous in spirit to the protocols that were introduced in medical genetics at a time when a reproducibility crisis was found in the field of candidate gene association studies. Many studies looking for risk factors for common complex diseases resulted in publications with marginally significant p -values without correcting for the multiple hypothesis testing that was implicitly performed due to many candidate genes being tested and only those with significant findings being published. Unsurprisingly, most of these claims failed to replicate in follow-up studies in independent sets of samples (Ioannidis 2005, Border et al. 2019, Collins et al. 2012, Duncan et al. 2019). The human medical genetic community addressed this challenge by coming together to support a rigorous set of commonly accepted standards for declaring genome-wide statistical significance, such as the requirement that p -values be corrected for the effective number of independent common variants in the genome and requiring correcting for the known confounders of population structure and undocumented relatedness among individuals (Hirschhorn and Daly 2005).

Conclusion

Sampling admixture graph space is a useful method for modeling population histories, but finding robust and accurate models can be challenging. As we demonstrated by revisiting a handful of published admixture graphs and re-analyzing the same datasets used to fit them, f -statistic and, more generally, allele frequency data alone are usually insufficient for building accurate graph models, making it necessary to incorporate other sources of evidence. This provides a challenge to previous approaches for automated model building. We investigated several published admixture graph models and, in nearly all cases, found many alternative models, some of which are historically and geographically plausible but contradict conclusions that were derived from the published models. To conduct these analyses, we developed a method for automated admixture graph topology optimization which can incorporate external sources of information as topological constraints. This method is developed in a framework called *ADMIXTOOLS 2*, which aside from admixture graph modeling, implements many other methods for population history inference based on f -statistics. In the process of revisiting published admixture graphs we found a well-fitting model for the history of dogs that is substantially different from the published model (Bergström et al. 2020) and is strikingly congruent with the known history of relevant human lineages. We also found a novel admixture graph for domestic and wild ancient horses (Librado et al. 2021) that is

substantially different from the published model, fits the data significantly better, and is geographically and historically plausible. These alternative graphs, however, have some of the same challenges as the published ones: they are almost certainly oversimplified relative to the true histories, and they exist in a large space of admixture graphs with meaningful topological differences that fit the allele frequency correlation data equally well. An important topic for future work should be to test these new alternative models as well as the previously published models as hypotheses with newly reported ancient samples and additional lines of genetic, archaeological, and other forms of analysis to obtain further clarity about population history.

It is important to recognize that the key concern we have highlighted in this study—the fact that there can often be multiple different topologies that are equally good fits to the allele frequency correlation patterns relating a set of populations—does not invalidate the use of allele frequency correlation testing in many other contexts in which it has been applied to make inferences about population history. For example, negative f_3 -statistics (“admixture” f_3 -statistics) continue to provide unambiguous evidence for a history of mixture in tested populations, and f_4 and D symmetry statistics remain a powerful way of evaluating whether a tested pair of populations is consistent with descending from a common ancestral population since separation from the ancestors of two groups used for comparison. The *qpWave* methodology remains a fully valid generalization of f_4 -statistics, making it possible to test whether a set of populations is consistent with descending from a specified number of ancestral populations (which separated at earlier times from a comparison set of populations). In addition, the *qpAdm* extension of *qpWave*—which allows for estimating proportions of mixtures for the tested population under the assumption that we have data from the source populations for the mixture—remains a valid approach, unaffected by the concerns identified here. Instead of relying on a specific model of deep population relationships, *qpAdm* relies on an empirically measured covariance matrix of f_4 -statistics for the analyzed populations, which is highly constraining with respect to estimation of mixture proportions but can be consistent with a wide range of deep history models. All these methods are implemented in *ADMIXTOOLS 2*.

Finally, approaches that use admixture graphs to adjust for the covariance structure relating a set of populations without insisting that the particular admixture graph model that is proposed is true with can be useful, for example for the purpose of analyzing shared genetic drift patterns of a group of populations that derive from similar mixtures. One example was a study that attempted to test for different source populations for Neolithic migrations into the Balkans after controlling for different proportions of hunter-gatherer admixture (Mathieson et al. 2018). Another example was a study that attempted to study shared ancestry between different East African forager populations after controlling for different proportions of deeply divergent source populations (Lipson et al. 2022). However, with respect to the inferences about deep history produced by admixture graphs themselves, our results highlight the importance of caution in proposing specific models of population history that relate a set of groups.

Methods

Technical presentation of *ADMIXTOOLS 2* in the context of f -statistic modeling methods

Much of the content that follows recapitulates theory presented in previous work, notably Reich et al. 2009, Green et al. 2010 and Patterson et al. 2012, but we summarize it here for coherence.

f-statistics

All *ADMIXTOOLS* programs are based on the statistics f_2 , f_3 , and f_4 , for population pairs, triplets, and quadruples, respectively.

f_2 quantifies the genetic drift separating two populations A and B . For a single SNP, it is given by $f_2(A, B) = \frac{1}{M} \sum_j (a_j - b_j)^2$, where a_j and b_j are the allele frequencies for SNP j in populations A and B . When allele frequencies are estimated using a small number of samples, this estimator of f_2 will be biased upwards. An unbiased estimator of f_2 is given by $f_2 = \frac{1}{M} \sum_j (a_j - b_j)^2 - \frac{a_j(1-a_j)}{n_{A,j}-1} - \frac{b_j(1-b_j)}{n_{B,j}-1}$, where $n_{A,j}$ and $n_{B,j}$ are the observed allele counts in populations A and B .

$f_3(A; B, C) = \frac{1}{M} \sum_j (a_j - b_j)(a_j - c_j)$ is the covariance of the allele frequency differences between populations A and B , and the allele frequency differences between populations A and C (assuming that alleles are coded randomly, so that $a - b$ and $a - c$ are both 0 in expectation). Significantly negative values of $f_3(A; B, C)$ suggest that A is a mixture of sources related to B and C (although the converse does not hold: A might be admixed between B and C even if f_3 is positive).

$f_4(A, B; C, D) = \frac{1}{M} \sum_j (a_j - b_j)(c_j - d_j)$ is the covariance of the allele frequency differences between A and B , and the allele frequency differences between C and D . Significantly positive values of $f_4(A, B; C, D)$ (or equivalently significantly negative values of $f_4(A, B; D, C)$) reveal that A and B do not form a clade with respect to C and D , and that some of the drift separating A from C is shared with the drift separating B from D .

f_3 and f_4 can be written as linear combinations of f_2 statistics:

$$f_3(A; B, C) = \frac{1}{2} (f_2(A, B) + f_2(A, C) - f_2(B, C)) \quad (\text{Eq. 1})$$

$$f_4(A, B; C, D) = \frac{1}{2} (f_2(A, D) + f_2(B, C) - f_2(A, C) - f_2(B, D)) \quad (\text{Eq. 2})$$

This implies that all f_3 - and f_4 -statistics can be computed from f_2 -statistics as long as they are defined on the same SNPs.

For revisiting published studies, we used the “*extract_f2*” function with the “*maxmiss*” argument set at 0, which corresponds to the “*useallsnps: NO*” setting in classic *ADMIXTOOLS*. It means that no missing data are allowed (at the level of populations) in the specified set of populations for which pairwise f_2 -statistics are calculated. For the values of the “*blgsiz*”, “*adjust_pseudohaploid*”, and “*minac2*” arguments we use in our analyses, see **Table S1**. The “*blgsiz*” argument sets the SNP block size in Morgans, and we used either the default value of 0.05 (5 cM), or 4,000,000 bp when a genetic map was not available. Genotypes of pseudo-haploid samples are usually coded as 0 or 2 (i.e., they are, strictly speaking, pseudo-diploid), even though only one allele is observed. The “*adjust_pseudohaploid*” argument ensures that the observed allele count increases only by 1 for each pseudo-haploid sample. If “*TRUE*” (default), samples that do not have any genotypes coded as 1 among the first 1,000 SNPs are automatically identified as pseudo-haploid. This leads to slightly more accurate estimates of f -statistics. Setting this parameter to “*FALSE*” treats all samples as diploid.

Another important argument (“*minac2=2*”) of the “*extract_f2*” function removes sites with only one chromosome genotyped in any non-singleton population and is needed for unbiased estimation of negative f_3 -statistics in non-singleton pseudo-haploid populations. In the absence of negative f_3 -statistics or pseudo-haploid populations, this argument has no influence on admixture graph log-likelihood scores. This algorithm for calculation of f -statistics triggered by the “*minac2=2*” argument is described below.

For $f_3(a; b, c)$, we compute the uncorrected numerator for each SNP, $(a - b) \times (a - c)$. We then subtract a bias correction factor at each SNP, $p(1 - p) / (ac - 1)$, which we only need for population a (because the other factors cancel out); p is the allele frequency, and ac is the observed allele count. In pseudo-haploid samples, $(ac - 1)$ would be zero and produce an error in any sites with only one observed allele. With the “*inbreed: NO*” setting in Classic *ADMIXTOOLS*, the smallest non-zero value

for ac is 2, so the division by 0 problem is avoided, but the correction factor is slightly smaller than it should be. *ADMIXTOOLS2* adds only an allele count of 1 for each site in a pseudo-haploid sample (with the default option “*adjust_pseudohaploid = TRUE*”), so there can be cases where $ac = 1$. To imitate what the setting “*inbreed: NO*” in Classic *ADMIXTOOLS* is doing, ac is set to 2 at those sites (or the denominator is set to 1). There is still a small difference between Classic *ADMIXTOOLS* and *ADMIXTOOLS2* at other sites because each observed site adds 2 alleles in *ADMIXTOOLS* with the default setting “*inbreed: NO*”, but only 1 allele in *ADMIXTOOLS2* with the default setting “*adjust_pseudohaploid = TRUE*”, but for admixture graph fitting that does not matter. One solution to avoid biased correction factors is to only consider sites with ac of at least two, which is what the “*inbreed: YES*” setting in Classic *ADMIXTOOLS* does. The problem with this is that we cannot use populations with a single pseudo-haploid sample, which is often useful, and would only give misleading results if that population is admixed. The new option “*minac2=2*” in *ADMIXTOOLS2* is different from the “*inbreed: YES*” setting in Classic *ADMIXTOOLS* since it makes an exception for populations consisting of a single pseudo-haploid sample in that it sets ac to 2 at each site (denominator is set to 1) when computing the correction factor of those populations.

Fitting admixture graphs

An admixture graph is a directed acyclic graph specifying the topology of the ancestral relationships among a set of populations. Each node in this graph represents a (present-day or ancient) population. Terminal nodes (also called leaf nodes) represent observed populations, while internal nodes represent unobserved ancestral populations. Modeling all observed populations as leaf nodes confers some robustness to drift specific to single populations and to genotyping errors. The edges connecting the populations are weighted and correspond either to the magnitude of genetic drift that has occurred along that branch (drift edges), or to the admixture proportions (admixture edges, where two edges point to the same node).

The goal of *qpGraph* is to test how well a given graph topology fits the observed f -statistics. This is achieved by varying the edge weights until the maximum likelihood fit is obtained. The following section describes the graph fitting in more detail.

First, for k populations, all $\frac{k(k+1)}{2}$ f_3 -statistics of the form $f_3(O; X_1, X_2)$ are computed, where O is one of the k populations (typically an outgroup), and X_1 and X_2 are all pairs formed from the other populations (including pairs where $X_1 = X_2$). These f_3 -statistics can then be used to fit the graph and to compute the likelihood. The likelihood score of a graph is the dot product of the differences between the expected and observed f_3 -statistics, weighted by the inverse covariance matrix of f_3 -statistics:

$$L(g) = -\frac{1}{2} (f_{3,obs} - f_{3,fit})' Q^{-1} (f_{3,obs} - f_{3,fit}) \quad (\text{Eq. 3})$$

Here, $f_{3,obs}$ are the observed f_3 -statistics and $f_{3,fit}$ are the fitted f_3 -statistics. Both are vectors of length $q = \frac{k(k+1)}{2}$ for k populations excluding the outgroup. Q is the $q \times q$ covariance matrix of f_3 -statistics, where the diagonal entries are the f_3 -statistic variances, and the off-diagonal entries are the covariances for all pairs of f_3 -statistics. Just like the variances (the squared standard errors), the covariances are estimated from the jackknife leave-one-block-out f_3 -statistics.

Finding the edge weights which maximize the likelihood score involves two nested optimization steps. The inner optimization finds the drift weights which maximize the likelihood score while fixing the admixture weights. The outer optimization finds the admixture weights which maximize the likelihood score, while optimizing the drift weights for each set of admixture weights. The inner

optimization uses a quadratic programming solver to find the optimal drift weights, while the outer optimization uses a general purpose optimization algorithm to find the optimal admixture weights. While the gradient function in the outer optimization adjusts the admixture weights, the objective function iterates over the following steps:

1. Optimization of drift weights conditional on admixture weights
2. Estimation of fitted f_3 -statistics
3. Calculation of the graph likelihood using observed and fitted f_3 -statistics

These steps are repeated until convergence is reached and the likelihood score can no longer be improved by adjusting the admixture weights.

Step 1 optimizes the drift edge weights, while holding the admixture weights constant. All drift edge weights are required to be non-negative, which makes this a constrained quadratic programming problem (hence *qpGraph*). Additional upper and lower bounds can be specified for individual graph edges.

Step 2 turns the edge weights into fitted f_3 -statistics. To see how edge weights in an admixture graph translate to f_3 -statistics, it helps to first consider how they translate into f_2 -statistics for a pair of populations. Without any admixture events, there is exactly one path p connecting any two populations. The fitted f_2 -statistic ($f_{2,fit}$) is the sum of edge weights w_e along this path p connecting two populations. The fitted f_2 -statistic is the sum of edge weights w_e along this path:

$$f_{2,fit} = \sum_{e \in p} w_e$$

In the presence of admixture events, two populations may be connected via multiple paths. Each admixture node that lies between the two populations increases the number of possible paths. The fitted f_2 -statistic for the two populations now becomes the weighted sum of all these paths, where the weight of each path is given by the product of all estimated admixture proportions w_a along this path ($\prod_{a \in p} w_a$):

$$f_{2,fit} = \sum_{p \in P} \prod_{a \in p} w_a \sum_{e \in p} w_e \quad (\text{Eq. 4})$$

The fitted f_2 -statistics are then used to obtain fitted f_3 -statistics using Eq. 1.

Step 3 uses the fitted and observed f_3 -statistics to estimate the likelihood score using Eq. 3.

Prior to these three steps, initial admixture weights are drawn randomly. To ensure that the end results do not depend on the random initialization, the whole optimization process is repeated multiple times with different random initial values. The original *ADMIXTOOLS* implementation retains only the results from the initial values resulting in the lowest absolute likelihood score. The new *ADMIXTOOLS* implementation provides an option to retrieve the results for all random initializations. This can be useful, as large fluctuations between different random initializations can be an indicator of an overparameterized or otherwise poorly fitting model.

Automated admixture graph inference

To find graph topologies that could conceivably have given rise to the observed f -statistics, we start with a randomly generated graph with a fixed number of admixture events, apply a number of modifications to this graph, and evaluate each of the resulting graphs. We then pick the best-fitting graph and repeat this procedure until graph modifications no longer lead to improved scores. We use a number of random graph modifications, as well as targeted modifications which are informed by parameters obtained during the fitting of the current graph.

For the targeted modifications we change the optimization of a single graph from a constrained optimization problem, in which drift edges are constrained to be positive and admixture weights are constrained to be between zero and one, to an unconstrained optimization problem in which both types of parameters can take any real values. Rearranging the nodes adjacent to edges which were estimated to be negative results in an improved fit at a much higher rate than random graph adjustments.

The random modifications include (1) pruning and randomly re-grafting leaf nodes, (2) pruning and randomly re-grafting a set of connected nodes in the graph, (3) swapping the orientation of admixture edges, (4) shifting admixture edges, (5) re-rooting the graph, (6) combinations of two or more of any of these modifications.

The number of admixture events is not affected by the graph modifications described so far. A significant score improvement can often be achieved by adding a single admixture edge to several random positions in a graph. This is unsurprising since it increases the degrees of freedom of the original graph. However, picking the best fitting graph with one admixture edge added, and testing all graphs that result from removing a single admixture edge from that graph often results in a graph with the same number of admixture events and a better fit than the original graph. We employ this strategy whenever the regular graph modifications described above do not lead to any further improvements.

We keep track of the search tree of all previously evaluated graphs and their scores in order to not evaluate any graph more than once, and so that backtracking in the search space is possible in cases where no more local improvements can be identified. Nevertheless, multiple iterations with different random starting graphs are usually necessary to find graphs with good fits. The number of iterations needed to approach a global optimum depends on the size of the search space, but the optimal number of iterations is hard to estimate in practice.

For revisiting published studies, we used the following settings of the *findGraphs* algorithm:

- *mutfuns* = *namedList(spr_leaves, spr_all, swap_leaves, move_admixededge_once, flipadmix_random, place_root_random, mutate_n)*, a list of functions used to modify graphs.
- *numgraphs* = 10, number of alternative graphs produced by randomly applying the mutation functions at the start of each generation.
- *stop_gen* = 10000, total number of generations after which to stop.
- *stop_gen2* = 30, number of generations without LL score improvement after which to stop.
- *plusminus_generations* = 10. If the best score does not improve after *plusminus_generations* generations, another approach to improving the score is attempted: A number of graphs with an additional admixture edge is generated and evaluated. The resulting graph with the best score is picked, and new graphs are created by removing any one admixture edge (bringing the number back to what it was originally). The graph with the lowest score is then selected. This approach often makes it possible to break out of local optima.
- *opt_worst_residual* = *FALSE*. Optimize for lowest worst residual instead of best score. “*FALSE*” by default, because the LL score is generally a better indicator of the quality of the model fit, and because optimizing for the lowest worst residual is much slower since f_4 -statistics need to be computed.
- *reject_f4z* = 0. If this is a number greater than zero, all f_4 -statistics with $|Z\text{-score}| > \text{reject_f4z}$ will be used to constrain the search space of admixture graphs: Any graphs in which f_4 -statistics greater than *reject_f4z* are expected to be zero will not be evaluated.
- *diag* = *1e-04*. This argument is passed to the *qpgraph* function and determines the regularization term added to the diagonal elements of the covariance matrix of fitted branch lengths (after scaling by the matrix trace). Default is 0.0001.

- *numstart* = 10. This argument is passed to the *qpgraph* function and determines the number of random initializations of starting weights (defaults to 10). Increasing this number will make the optimization slower but reduce the risk of not finding the optimal weights.
- *lsqmode* = FALSE. This argument is passed to the *qpgraph* function. If set to "FALSE", the inverse f_3 -statistic covariance matrix is not discarded by the algorithm.

The arguments "*admix_constraints*" (constraints on the number of admixture events in the history of a given population), "*event_constraints*" (constraints on the branching order of specified lineages), and "*outpop*" (the population assigned as an outgroup) were set according to **Table S1**. Each *findGraphs* run was initiated by a random graph with a specified number of admixture events. Usually, the same topology constraints were applied at the stage of random graph generation and the topology search stage, for exceptions see **Table S1**.

Evaluating automated admixture graph inference through simulations

We evaluated the performance of *findGraphs* by simulating genetic data under a large number of different admixture graph models, applying *findGraphs* to each simulated data set in three independent iterations, and comparing the resulting best graph across three iterations to the simulated graph. We applied *TreeMix* to the same simulated data for comparison. We simulated between 8 and 16 populations per graph, and between 0 and 10 admixture events. For each parameter combination, we simulated 20 different admixture graphs generated by the *random_admixturegraph* function. We counted both the fraction of random simulated graphs where the best inferred graph was identical to the simulated graph, as well as the fraction of random simulated graphs where the best inferred graph was either identical to the simulated graph or had a better score than the simulated graph. For models with a large number of admixture events the number of possible models is so large that it becomes increasingly likely that there will be some alternative models which fit the data better than the model under which the data were simulated.

We used *msprime* and the *msprime_sim* wrapper function in *ADMIXTOOLS* 2 to simulate data for 100,000 unlinked SNPs and 100 diploid samples per population for each admixture graph. The simulation parameters we chose were aimed at facilitating fast simulations of large numbers of informative SNPs rather than at being as realistic as possible. We therefore expect that the simulation results allow us to make comparisons across groups, but not that they are informative about the rate at which "true" models can be recovered in empirical data. We simulated under a constant mutation rate of 0.001 per site per generation, a constant haploid effective population size of 1000, with neighboring nodes in the graph separated by 1000 or more generations, and all admixture events occurring in discrete pulses of 50/50 proportions.

To allow for a fair comparison between *findGraphs* and *TreeMix*, we made sure that small differences in the way admixture graphs are modeled in *findGraphs* and in *TreeMix* were accounted for before testing graphs for identical topology. For example, *TreeMix* admixture graphs can have lineages terminating at an admixture node, whereas in *findGraphs* lineages always end at a 'leaf' node with a single ancestor.

Comparing the fits of different admixture graphs

We are interested in determining whether one admixture graph fits the data significantly better than another admixture graph, or whether an observed score difference $\Delta = S_1 - S_2$ can be attributed to variability across independent SNPs.

We first consider two admixture graphs with the same number of admixture events, where we can ignore the problem of comparing two models with different complexity. As in other bootstrap standard error calculations, we divide the genome into n blocks indexed by i , and we draw b sets of

n blocks with replacement, indexed by j . We fit both graphs b times - once for each bootstrap set of SNP blocks. This results in a set of b score differences Δ_j . The bootstrap confidence interval for the difference in scores is given by the quantiles of the distribution of Δ_j . We also compute an empirical bootstrap p -value, testing the null hypothesis that two different graphs fit the data equally well. It is computed as $p = \max(\frac{1}{b}, 2\delta)$ (Boos 2003), where δ is either the fraction of $\Delta_j > 0$, or the fraction of $\Delta_j < 0$, whichever is smaller. The reason for applying bootstrap resampling, as opposed to jackknife resampling in this case, is that the distribution of score differences tends to have a high kurtosis, which can make jackknife estimates inaccurate. Simulating data under the null hypothesis is not straightforward in this case, because it involves finding two non-identical graphs which in expectation fit the data equally well. We decided to simulate under one graph and compare two graphs which are symmetrically related to the simulated graph (Figure 3). This confirmed that the p -value follows a uniform distribution under the null hypothesis.

Next, we consider comparisons of two graphs of different complexity. The problem here is that more complex graphs have more degrees of freedom which allow them to overfit the data better, without necessarily being any closer to the truth. To solve this problem, we introduce an out-of-sample likelihood score. The regular likelihood score is given by:

$L(g) = -\frac{1}{2} (f_{3,obs} - f_{3,fit})' Q^{-1} (f_{3,obs} - f_{3,fit})$, with $f_{3,obs}$ and $f_{3,fit}$ defined on the same set of SNPs. The out-of-sample likelihood score is defined in the same way, except that $f_{3,obs}$ and $f_{3,fit}$ are defined on mutually exclusive sets of SNP blocks, thereby preventing any overfitting. The covariance matrix Q is defined on the same set of SNP blocks as $f_{3,fit}$. As described earlier, we use block-bootstrap to fit both graphs multiple times on different SNP blocks. In each bootstrap iteration, we use all SNP blocks which are not used in fitting the graph for estimating $f_{3,obs}$.

Admixture graph identifiability

An edge in an admixture graph is unidentifiable, if small changes to the weight of this edge (admixture proportions in the case of an admixture edge, drift length in the case of a drift edge) do not necessarily lead to changes in expected f -statistics. This is the case if the small change in weight can be offset by small changes in other graph edges, leading to a situation where observed f -statistics can be explained by more than one weight estimate for that edge. To find unidentifiable edges, we derive the Jacobi matrix of the graph's system of f_2 equations (Eq. 4 applied to each population pair). In principle, whether or not a parameter is identifiable can depend on the values of all other parameters. However, in practice this is rarely the case, and so we draw values for all parameters from a uniform distribution, which gives us a Jacobi matrix with numeric values. We then determine the rank of the Jacobi matrix, along with the rank of all matrices that result from dropping a single column (a parameter corresponding to a graph edge). For identifiable edges, the rank of the full matrix will be greater than the rank of the reduced matrix, and for unidentifiable edges, the ranks will be the same.

Drawing conclusions from a large number of fitting models

We developed several methods that aim to summarize a collection of graphs which all fit the data similarly well. By highlighting features which are observed repeatedly across graphs, it becomes possible to extract interpretable conclusions from an otherwise hard to interpret collection of possible models. These graph summaries identify features in each graph that can be compared to different graphs describing the same populations. We summarize each graph in several ways:

(1) Admixture status of each population

For each population, we count the total number of admixture events that is encountered along all paths to the root.

- (2) Order of population split events
For each pair of population pairs, we determine if the most recent split of the first pair has occurred before or after the most recent split of the second pair, or whether the graph does not specify the order in which those splits occurred.
 - (3) Proxy populations
For each admixed population in a graph, we attempt to identify proxy sources: populations closest to the admixing populations. In contrast to the other approaches to summarizing graphs which are based only on the topology of each graph, this can also rely on information about the estimated graph parameters.
 - (4) Cladality
For each group of four populations, we test whether the graph implies that any f_4 -statistic describing the relationship between the four populations is expected to be zero.
 - (5) Node descendants
Each internal node in an admixture graph is an ancestor to a specific set of leaf populations. An admixture graph can be characterized by the sets of leaf populations formed by the internal nodes. Multiple admixture graphs may be compared by counting the number of overlapping sets. This also makes it possible to quantify for each internal node in a single graph, how often a matching internal node can be found across a collection of alternative graphs, which is conceptually similar to bootstrap support values in phylogenetic trees.
- While these methods provide some help in comparing features across many graphs, they are not able to reliably answer the question whether the fitting graphs are relatively similar or dissimilar from each other, and whether they are similar to any particular graph. This is in part due to the fact that small topological changes involving populations of interest may be more relevant than similar topological changes involving only populations that are not the focus of the study.

Acknowledgements

We thank Anders Bergström, Esther Brielle, Mateja Hajdinjak, Iosif Lazaridis, Pablo Librado, Mark Lipson, Vagheesh Narasimhan, Ludovic Orlando, Nick Patterson, Mary Prendergast, Jakob Sedig, Kendra Sirak, Pontus Skoglund, and Chuanchao Wang, for suggestions for how to improve specific analyses, and for conversations and critical comments. We thank Matthew Mah, Shop Mallick, Adam Micco, Nadin Rohland, Ron Pinhasi for help in generating additional data from an ancient DNA library from individual I8726 for which 1.24 million SNP capture data was generated and published in Narasimhan *et al.* 2019 and for which we report 2.6-fold shotgun data here (**Table S2**). P.F., P.C., and O.F. were supported by the Czech Ministry of Education, Youth and Sports (program ERC CZ, project no. LL2103) and by the Czech Science Foundation (project no. 21-27624S). P.F. and P.C. were also supported by the Czech Ministry of Education, Youth and Sports (Large Infrastructures for Research, Experimental Development and Innovations project "IT4Innovations National Supercomputing Center – LM2015070"; Inter-Excellence program, project no. LTAUSA18153). R.M. and D.R. were supported by grants from the National Institutes of Health (GM100233 and HG012287), the John Templeton Foundation (grant 61220), and the Allen Discovery Center program, a Paul G. Allen Frontiers Group advised program of the Paul G. Allen Family Foundation. D.R. was also supported by a private gift from J.-F. Clin and is an Investigator of the Howard Hughes Medical Institute.

References

1. Bellwood P. The checkered prehistory of rice movement southwards as a domesticated cereal—from the Yangzi to the equator. *Rice*. 2011; 4:93–103. doi: 10.1007/s12284-011-9068-9
2. Bergström A, Frantz L, Schmidt R, Ersmark E, Lebrasseur O, Girdland-Flink L, Lin AT, Storå J, Sjögren KG, Anthony D, Antipina E, Amiri S, Bar-Oz G, Bazaliiskii VI, Bulatović J, Brown D, Carmagnini A, Davy T, Fedorov S, Fiore I, Fulton D, Germonpré M, Haile J, Irving-Pease EK, Jamieson A, Janssens L, Kirillova I, Horwitz LK, Kuzmanovic-Cvetković J, Kuzmin Y, Losey RJ,

- 2144 Dizdar DL, Mashkour M, Novak M, Onar V, Orton D, Pasarić M, Radivojević M, Rajković D,
2145 Roberts B, Ryan H, Sablin M, Shidlovskiy F, Stojanović I, Tagliacozzo A, Trantalidou K, Ullén I,
2146 Villaluenga A, Wapnish P, Dobney K, Götherström A, Linderholm A, Dalén L, Pinhasi R, Larson G,
2147 Skoglund P. Origins and genetic legacy of prehistoric dogs. *Science*. 2020 Oct 30;370(6516):557-
2148 564. doi: 10.1126/science.aba9572.
- 2149 3. Boos DD. Introduction to the Bootstrap World. *Statist. Sci.* 18(2): 168-174 (May 2003). DOI:
2150 10.1214/ss/1063994971.
- 2151 4. Border R, Johnson EC, Evans LM, Smolen A, Berley N, Sullivan PF, Keller MC (2019) No support
2152 for historical candidate gene or candidate gene-by-interaction hypotheses for major depression
2153 across multiple large samples. *American Journal of Psychiatry* 176: 376-387.
- 2154 5. Chen FH, Dong GH, Zhang DJ, Liu XY, Jia X, An CB, Ma MM, Xie YW, Barton L, Ren XY, Zhao ZJ, Wu
2155 XH, Jones MK. Agriculture facilitated permanent human occupation of the Tibetan Plateau after
2156 3600 B.P. *Science*. 2015 Jan 16;347(6219):248-50. doi: 10.1126/science.1259172.
- 2157 6. Collins AL, Kim Y, Sklar P; International Schizophrenia Consortium, O'Donovan MC, Sullivan PF
2158 (2012) Hypothesis-driven candidate genes for schizophrenia compared to genome-wide
2159 association results. *Psychological Medicine* 42: 607-616.
- 2160 7. Duncan LE, Ostacher M, Ballon J (2019) How genome-wide association studies (GWAS) made
2161 traditional candidate gene studies obsolete. *Neuropsychopharmacology* 44: 1518-
2162 1523. Campbell MC, Tishkoff SA. African genetic diversity: implications for human demographic
2163 history, modern human origins, and complex disease mapping. *Annu Rev Genomics Hum Genet.*
2164 2008;9:403-33. doi: 10.1146/annurev.genom.9.081307.164258.
- 2165 8. Durvasula A, Sankararaman S. Recovering signals of ghost archaic introgression in African
2166 populations. *Sci Adv*. 2020 Feb 12;6(7):eaax5097. doi: 10.1126/sciadv.aax5097.
- 2167 9. Flegontov P, Altınışık NE, Changmai P, Rohland N, Mallick S, Adamski N, Bolnick DA,
2168 Broomandkhoshbacht N, Candilio F, Culleton BJ, Flegontova O, Friesen TM, Jeong C, Harper TK,
2169 Keating D, Kennett DJ, Kim AM, Lamnidis TC, Lawson AM, Olalde I, Oppenheimer J, Potter BA,
2170 Raff J, Sattler RA, Skoglund P, Stewardson K, Vajda EJ, Vasilyev S, Veselovskaya E, Hayes MG,
2171 O'Rourke DH, Krause J, Pinhasi R, Reich D, Schiffels S. Palaeo-Eskimo genetic ancestry and the
2172 peopling of Chukotka and North America. *Nature*. 2019 Jun;570(7760):236-240. doi:
2173 10.1038/s41586-019-1251-y.
- 2174 10. Fu Q, Posth C, Hajdinjak M, Petr M, Mallick S, Fernandes D, Furtwängler A, Haak W, Meyer M,
2175 Mittnik A, Nickel B, Peltzer A, Rohland N, Slon V, Talamo S, Lazaridis I, Lipson M, Mathieson I,
2176 Schiffels S, Skoglund P, Derevianko AP, Drozdov N, Slavinsky V, Tsybankov A, Cremonesi RG,
2177 Mallegni F, Gély B, Vacca E, Morales MR, Straus LG, Neugebauer-Maresch C, Teschler-Nicola M,
2178 Constantin S, Moldovan OT, Benazzi S, Peresani M, Coppola D, Lari M, Ricci S, Ronchitelli A,
2179 Valentin F, Thevenet C, Wehrberger K, Grigorescu D, Rougier H, Crevecoeur I, Flas D, Semal P,
2180 Mannino MA, Cupillard C, Bocherens H, Conard NJ, Harvati K, Moiseyev V, Drucker DG, Svoboda
2181 J, Richards MP, Caramelli D, Pinhasi R, Kelso J, Patterson N, Krause J, Pääbo S, Reich D. The
2182 genetic history of Ice Age Europe. *Nature*. 2016 Jun 9;534(7606):200-5. doi:
2183 10.1038/nature17993.
- 2184 11. Gray RD, Jordan FM. Language trees support the express-train sequence of Austronesian
2185 expansion. *Nature*. 2000 Jun 29;405(6790):1052-5. doi: 10.1038/35016575.
- 2186 12. Green RE, Krause J, Briggs AW, Maricic T, Stenzel U, Kircher M, Patterson N, Li H, Zhai W, Fritz
2187 MH, Hansen NF, Durand EY, Malaspina AS, Jensen JD, Marques-Bonet T, Alkan C, Prüfer K,
2188 Meyer M, Burbano HA, Good JM, Schultz R, Aximu-Petri A, Butthof A, Höber B, Höffner B,
2189 Siegemund M, Weihmann A, Nusbaum C, Lander ES, Russ C, Novod N, Affourtit J, Egholm M,
2190 Verna C, Rudan P, Brajkovic D, Kucan Ž, Gušić I, Doronichev VB, Golovanova LV, Lalueza-Fox C, de
2191 la Rasilla M, Fortea J, Rosas A, Schmitz RW, Johnson PLF, Eichler EE, Falush D, Birney E, Mullikin
2192 JC, Slatkin M, Nielsen R, Kelso J, Lachmann M, Reich D, Pääbo S. A draft sequence of the
2193 Neandertal genome. *Science*. 2010 May 7;328(5979):710-722. doi: 10.1126/science.1188021.
- 2194 13. Gronau I, Hubisz MJ, Gulko B, Danko CG, Siepel A. Bayesian inference of ancient human
2195 demography from individual genome sequences. *Nat Genet*. 2011 Sep 18;43(10):1031-4. doi:
2196 10.1038/ng.937.
- 2197 14. Gutenkunst RN, Hernandez RD, Williamson SH, Bustamante CD. Inferring the joint demographic

- 2198 history of multiple populations from multidimensional SNP frequency data. *PLoS Genet.* 2009
2199 Oct;5(10):e1000695. doi: 10.1371/journal.pgen.1000695.
- 2200 15. Haak W, Lazaridis I, Patterson N, Rohland N, Mallick S, Llamas B, Brandt G, Nordenfelt S, Harney
2201 E, Stewardson K, Fu Q, Mittnik A, Bánffy E, Economou C, Francken M, Friederich S, Pena RG,
2202 Hallgren F, Khartanovich V, Khokhlov A, Kunst M, Kuznetsov P, Meller H, Mochalov O, Moiseyev
2203 V, Nicklisch N, Pichler SL, Risch R, Rojo Guerra MA, Roth C, Szécsényi-Nagy A, Wahl J, Meyer M,
2204 Krause J, Brown D, Anthony D, Cooper A, Alt KW, Reich D. Massive migration from the steppe
2205 was a source for Indo-European languages in Europe. *Nature.* 2015 Jun 11;522(7555):207-11.
2206 doi: 10.1038/nature14317.
- 2207 16. Hajdinjak M, Mafessoni F, Skov L, Vernot B, Hübner A, Fu Q, Essel E, Nagel S, Nickel B, Richter J,
2208 Moldovan OT, Constantin S, Endarova E, Zahariev N, Spasov R, Welker F, Smith GM, Sinet-
2209 Mathiot V, Paskulin L, Fewlass H, Talamo S, Rezek Z, Sirakova S, Sirakov N, McPherron SP,
2210 Tsanova T, Hublin JJ, Peter BM, Meyer M, Skoglund P, Kelso J, Pääbo S. Initial Upper Palaeolithic
2211 humans in Europe had recent Neanderthal ancestry. *Nature.* 2021 Apr;592(7853):253-257. doi:
2212 10.1038/s41586-021-03335-3.
- 2213 17. Hammer MF, Woerner AE, Mendez FL, Watkins JC, Wall JD. Genetic evidence for archaic
2214 admixture in Africa. *Proc Natl Acad Sci U S A.* 2011 Sep 13;108(37):15123-8. doi:
2215 10.1073/pnas.1109300108.
- 2216 18. Hirschhorn JN, Daly MJ. Genome-wide association studies for common diseases and complex
2217 traits. *Nat Rev Genet.* 2005 Feb;6(2):95-108. doi: 10.1038/nrg1521.
- 2218 19. Hubisz MJ, Williams AL, Siepel A. Mapping gene flow between ancient hominins through
2219 demography-aware inference of the ancestral recombination graph. *PLoS Genet.* 2020 Aug
2220 6;16(8):e1008895. doi: 10.1371/journal.pgen.1008895.
- 2221 20. Hubisz M, Siepel A. Inference of Ancestral Recombination Graphs Using ARGweaver. *Methods*
2222 *Mol Biol.* 2020;2090:231-266. doi: 10.1007/978-1-0716-0199-0_10.
- 2223 21. Ioannidis JP (2005) Why most published research findings are false. *PLoS Med* 2: e124.
- 2224 22. Jeong C, Balanovsky O, Lukianova E, Kahbatkyy N, Flegontov P, Zaporozhchenko V, Immel A,
2225 Wang CC, Ixan O, Khussainova E, Bekmanov B, Zaibert V, Lavryashina M, Pocheshkhova E,
2226 Yusupov Y, Agdzhoyan A, Koshel S, Bukin A, Nymadawa P, Turdikulova S, Dalimova D, Churnosov
2227 M, Skhalyakho R, Daragan D, Bogunov Y, Bogunova A, Shtrunov A, Dubova N, Zhabagin M,
2228 Yepiskoposyan L, Churakov V, Pislegin N, Damba L, Saroyants L, Dibirova K, Atramentova L,
2229 Utevska O, Idrisov E, Kamenshchikova E, Evseeva I, Metspalu M, Outram AK, Robbeets M,
2230 Djansugurova L, Balanovska E, Schiffels S, Haak W, Reich D, Krause J. The genetic history of
2231 admixture across inner Eurasia. *Nat Ecol Evol.* 2019 Jun;3(6):966-976. doi: 10.1038/s41559-019-
2232 0878-2.
- 2233 23. Kamm J, Terhorst J, Durbin R, Song YS. Efficiently inferring the demographic history of many
2234 populations with allele count data. *J Am Stat Assoc.* 2020;115(531):1472-1487. doi:
2235 10.1080/01621459.2019.1635482.
- 2236 24. Ko AM, Chen CY, Fu Q, Delfin F, Li M, Chiu HL, Stoneking M, Ko YC. Early Austronesians: into and
2237 out of Taiwan. *Am J Hum Genet.* 2014 Mar 6;94(3):426-36. doi: 10.1016/j.ajhg.2014.02.003.
- 2238 25. Lachance J, Vernot B, Elbers CC, Ferwerda B, Froment A, Bodo JM, Lema G, Fu W, Nyambo TB,
2239 Rebbeck TR, Zhang K, Akey JM, Tishkoff SA. Evolutionary history and adaptation from high-
2240 coverage whole-genome sequences of diverse African hunter-gatherers. *Cell.* 2012 Aug
2241 3;150(3):457-69. doi: 10.1016/j.cell.2012.07.009.
- 2242 26. Lazaridis I, Patterson N, Mittnik A, Renaud G, Mallick S, Kirsanow K, Sudmant PH, Schraiber JG,
2243 Castellano S, Lipson M, Berger B, Economou C, Bollongino R, Fu Q, Bos KI, Nordenfelt S, Li H, de
2244 Filippo C, Prüfer K, Sawyer S, Posth C, Haak W, Hallgren F, Fornander E, Rohland N, Delsate D,
2245 Francken M, Guinet JM, Wahl J, Ayodo G, Babiker HA, Bailliet G, Balanovska E, Balanovsky O,
2246 Barrantes R, Bedoya G, Ben-Ami H, Bene J, Berrada F, Bravi CM, Brisighelli F, Busby GB, Cali F,
2247 Churnosov M, Cole DE, Corach D, Damba L, van Driem G, Dryomov S, Dugoujon JM, Fedorova SA,
2248 Gallego Romero I, Gubina M, Hammer M, Henn BM, Hervig T, Hodoglugil U, Jha AR, Karachanak-
2249 Yankova S, Khusainova R, Khusnutdinova E, Kittles R, Kivisild T, Klitz W, Kučinskas V,
2250 Kushniarevich A, Laredj L, Litvinov S, Loukidis T, Mahley RW, Melegh B, Metspalu E, Molina J,
2251 Mountain J, Näkkäläjärvi K, Nesheva D, Nyambo T, Osipova L, Parik J, Platonov F, Posukh O,

- 2252 Romano V, Rothhammer F, Rudan I, Ruizbakiev R, Sahakyan H, Sajantila A, Salas A, Starikovskaya
2253 EB, Tarekegn A, Toncheva D, Turdikulova S, Uktveryte I, Utevska O, Vasquez R, Villena M,
2254 Voevoda M, Winkler CA, Yepiskoposyan L, Zalloua P, Zemunik T, Cooper A, Capelli C, Thomas
2255 MG, Ruiz-Linares A, Tishkoff SA, Singh L, Thangaraj K, Vilems R, Comas D, Sukernik R, Metspalu
2256 M, Meyer M, Eichler EE, Burger J, Slatkin M, Pääbo S, Kelso J, Reich D, Krause J. Ancient human
2257 genomes suggest three ancestral populations for present-day Europeans. *Nature*. 2014 Sep
2258 18;513(7518):409-13. doi: 10.1038/nature13673.
- 2259 27. Lazaridis I, Nadel D, Rollefson G, Merrett DC, Rohland N, Mallick S, Fernandes D, Novak M,
2260 Gamarra B, Sirak K, Connell S, Stewardson K, Harney E, Fu Q, Gonzalez-Fortes G, Jones ER,
2261 Roodenberg SA, Lengyel G, Bocquentin F, Gasparian B, Monge JM, Gregg M, Eshed V, Mizrahi AS,
2262 Meiklejohn C, Gerritsen F, Bejenaru L, Blüher M, Campbell A, Cavalleri G, Comas D, Froguel P,
2263 Gilbert E, Kerr SM, Kovacs P, Krause J, McGettigan D, Merrigan M, Merriwether DA, O'Reilly S,
2264 Richards MB, Semino O, Shamoon-Pour M, Stefanescu G, Stumvoll M, Tönjes A, Torroni A,
2265 Wilson JF, Yengo L, Hovhannisyan NA, Patterson N, Pinhasi R, Reich D. Genomic insights into the
2266 origin of farming in the ancient Near East. *Nature*. 2016 Aug 25;536(7617):419-24. doi:
2267 10.1038/nature19310.
- 2268 28. Leppälä K, Nielsen SV, Mailund T. admixturegraph: an R package for admixture graph
2269 manipulation and fitting. *Bioinformatics*. 2017 Jun 1;33(11):1738-1740. doi:
2270 10.1093/bioinformatics/btx048.
- 2271 29. Librado P, Khan N, Fages A, Kusliy MA, Suchan T, Tonasso-Calvière L, Schiavinato S, Alioglu D,
2272 Fromentier A, Perdereau A, Aury JM, Gaunitz C, Chauvey L, Seguin-Orlando A, Der Sarkissian C,
2273 Southon J, Shapiro B, Tishkin AA, Kovalev AA, Alquraishi S, Alfarhan AH, Al-Rasheid KAS, Seregély
2274 T, Klassen L, Iversen R, Bignon-Lau O, Bodu P, Olive M, Castel JC, Boudadi-Maligne M, Alvarez N,
2275 Germonpré M, Moskal-Del Hoyo M, Wilczyński J, Pospuła S, Lasota-Kuś A, Tunia K, Nowak M,
2276 Rannamäe E, Saarma U, Boeskorov G, Lõugas L, Kyselý R, Peške L, Bălăşescu A, Dumitraşcu V,
2277 Dobrescu R, Gerber D, Kiss V, Szécsényi-Nagy A, Mende BG, Gallina Z, Somogyi K, Kulcsár G, Gál
2278 E, Bendrey R, Allentoft ME, Sirbu G, Dergachev V, Shephard H, Tomadini N, Grouard S, Kasparov
2279 A, Basilyan AE, Anisimov MA, Nikolskiy PA, Pavlova EY, Pitulko V, Brem G, Wallner B, Schwall C,
2280 Keller M, Kitagawa K, Bessudnov AN, Bessudnov A, Taylor W, Magail J, Gantulga JO,
2281 Bayarsaikhan J, Erdenebaatar D, Tabaldiev K, Mijiddorj E, Boldgiv B, Tsagaan T, Pruvost M, Olsen
2282 S, Makarewicz CA, Valenzuela Lamas S, Albizuri Canadell S, Nieto Espinet A, Iborra MP, Lira
2283 Garrido J, Rodríguez González E, Celestino S, Olària C, Arsuaga JL, Kotova N, Pryor A, Crabtree P,
2284 Zhumatayev R, Toleubaev A, Morgunova NL, Kuznetsova T, Lordkipanize D, Marzullo M, Prato O,
2285 Bagnasco Gianni G, Tecchiati U, Clavel B, Lepetz S, Davoudi H, Mashkour M, Berezina NY,
2286 Stockhammer PW, Krause J, Haak W, Morales-Muñiz A, Benecke N, Hofreiter M, Ludwig A,
2287 Graphodatsky AS, Peters J, Kiryushin KY, Iderkhangai TO, Bokovenko NA, Vasiliev SK, Seregin NN,
2288 Chugunov KV, Plasteeva NA, Baryshnikov GF, Petrova E, Sablin M, Ananyevskaya E, Logvin A,
2289 Shevnina I, Logvin V, Kalieva S, Loman V, Kukushkin I, Merz I, Merz V, Sakenov S, Varfolomeyev
2290 V, Usmanova E, Zaibert V, Arbuckle B, Belinskiy AB, Kalmykov A, Reinhold S, Hansen S, Yudin AI,
2291 Vybornov AA, Epimakhov A, Berezina NS, Roslyakova N, Kosintsev PA, Kuznetsov PF, Anthony D,
2292 Kroonen GJ, Kristiansen K, Wincker P, Outram A, Orlando L. The origins and spread of domestic
2293 horses from the Western Eurasian steppes. *Nature*. 2021 Oct;598(7882):634-640. doi:
2294 10.1038/s41586-021-04018-9.
- 2295 30. Lipson M, Loh PR, Levin A, Reich D, Patterson N, Berger B. Efficient moment-based inference of
2296 admixture parameters and sources of gene flow. *Mol Biol Evol*. 2013 Aug;30(8):1788-802. doi:
2297 10.1093/molbev/mst099.
- 2298 31. Lipson M, Szécsényi-Nagy A, Mallick S, Pósa A, Stégmár B, Keerl V, Rohland N, Stewardson K,
2299 Ferry M, Michel M, Oppenheimer J, Broomandkhoshbacht N, Harney E, Nordenfelt S, Llamas B,
2300 Gusztáv Mende B, Köhler K, Oross K, Bondár M, Marton T, Oszás A, Jakucs J, Paluch T, Horváth
2301 F, Csengeri P, Koós J, Sebők K, Anders A, Raczyk P, Regénye J, Barna JP, Fábíán S, Serlegi G, Toldi
2302 Z, Gyöngyvér Nagy E, Dani J, Molnár E, Pálfi G, Márk L, Melegh B, Bánfai Z, Domboróczi L,
2303 Fernández-Eraso J, Antonio Mujika-Alustiza J, Alonso Fernández C, Jiménez Echevarría J,
2304 Bollongino R, Orschiedt J, Schierhold K, Meller H, Cooper A, Burger J, Bánffy E, Alt KW, Lalueza-
2305 Fox C, Haak W, Reich D. Parallel palaeogenomic transects reveal complex genetic history of early

- European farmers. *Nature*. 2017 Nov 16;551(7680):368-372. doi: 10.1038/nature24476.
32. Lipson M, Ribot I, Mallick S, Rohland N, Olalde I, Adamski N, Broomandkhoshbacht N, Lawson AM, López S, Oppenheimer J, Stewardson K, Asombang RN, Bocherens H, Bradman N, Culleton BJ, Cornelissen E, Crevecoeur I, de Maret P, Fomine FLM, Lavachery P, Mindzie CM, Orban R, Sawchuk E, Semal P, Thomas MG, Van Neer W, Veeramah KR, Kennett DJ, Patterson N, Hellenthal G, Lalueza-Fox C, MacEachern S, Prendergast ME, Reich D. Ancient West African foragers in the context of African population history. *Nature*. 2020 Jan;577(7792):665-670. doi: 10.1038/s41586-020-1929-1.
33. Lipson M. Applying f4-statistics and admixture graphs: Theory and examples. *Mol Ecol Resour*. 2020 Nov;20(6):1658-1667. doi: 10.1111/1755-0998.13230.
34. Lipson M, Sawchuk EA, Thompson JC, Oppenheimer J, Tryon CA, Ranhorn KL, de Luna KM, Sirak KA, Olalde I, Ambrose SH, Arthur JW, Arthur KJW, Ayodo G, Bertacchi A, Cerezo-Román JJ, Culleton BJ, Curtis MC, Davis J, Gidna AO, Hanson A, Kaliba P, Katongo M, Kwekason A, Laird MF, Lewis J, Mabulla AZP, Mapemba F, Morris A, Mudenda G, Mwafurirwa R, Mwangomba D, Ndiema E, Ogola C, Schilt F, Willoughby PR, Wright DK, Zipkin A, Pinhasi R, Kennett DJ, Manthi FK, Rohland N, Patterson N, Reich D, Prendergast ME. Ancient DNA and deep population structure in sub-Saharan African foragers. *Nature*. 2022 Mar;603(7900):290-296. doi: 10.1038/s41586-022-04430-9.
35. Lu D, Lou H, Yuan K, Wang X, Wang Y, Zhang C, Lu Y, Yang X, Deng L, Zhou Y, Feng Q, Hu Y, Ding Q, Yang Y, Li S, Jin L, Guan Y, Su B, Kang L, Xu S. Ancestral Origins and Genetic History of Tibetan Highlanders. *Am J Hum Genet*. 2016 Sep 1;99(3):580-594. doi: 10.1016/j.ajhg.2016.07.002.
36. Mallick S, Li H, Lipson M, Mathieson I, Gymrek M, Racimo F, Zhao M, Chennagiri N, Nordenfelt S, Tandon A, Skoglund P, Lazaridis I, Sankararaman S, Fu Q, Rohland N, Renaud G, Erlich Y, Willems T, Gallo C, Spence JP, Song YS, Poletti G, Balloux F, van Driem G, de Knijff P, Romero IG, Jha AR, Behar DM, Bravi CM, Capelli C, Hervig T, Moreno-Estrada A, Posukh OL, Balanovska E, Balanovsky O, Karachanak-Yankova S, Sahakyan H, Toncheva D, Yepiskoposyan L, Tyler-Smith C, Xue Y, Abdullah MS, Ruiz-Linares A, Beall CM, Di Rienzo A, Jeong C, Starikovskaya EB, Metspalu E, Parik J, Vilems R, Henn BM, Hodoglugil U, Mahley R, Sajantila A, Stamatoyannopoulos G, Wee JT, Khusainova R, Khusnutdinova E, Litvinov S, Ayodo G, Comas D, Hammer MF, Kivisild T, Klitz W, Winkler CA, Labuda D, Bamshad M, Jorde LB, Tishkoff SA, Watkins WS, Metspalu M, Dryomov S, Sukernik R, Singh L, Thangaraj K, Pääbo S, Kelso J, Patterson N, Reich D. The Simons Genome Diversity Project: 300 genomes from 142 diverse populations. *Nature*. 2016 Oct 13;538(7624):201-206. Doi: 10.1038/nature18964.
37. McColl H, Racimo F, Vinner L, Demeter F, Gakuhari T, Moreno-Mayar JV, van Driem G, Gram Wilken U, Seguin-Orlando A, de la Fuente Castro C, Wasef S, Shoocongdej R, Souksavatdy V, Sayavongkhamdy T, Saidin MM, Allentoft ME, Sato T, Malaspinas AS, Aghakhanian FA, Korneliussen T, Prohaska A, Margaryan A, de Barros Damgaard P, Kaewsutthi S, Lertrit P, Nguyen TMH, Hung HC, Minh Tran T, Nghia Truong H, Nguyen GH, Shahidan S, Wiradnyana K, Matsumae H, Shigehara N, Yoneda M, Ishida H, Masuyama T, Yamada Y, Tajima A, Shibata H, Toyoda A, Hanihara T, Nakagome S, Deviese T, Bacon AM, Durringer P, Ponche JL, Shackelford L, Patole-Edoumba E, Nguyen AT, Bellina-Pryce B, Galipaud JC, Kinaston R, Buckley H, Pottier C, Rasmussen S, Higham T, Foley RA, Lahr MM, Orlando L, Sikora M, Phipps ME, Oota H, Higham C, Lambert DM, Willerslev E. The prehistoric peopling of Southeast Asia. *Science*. 2018 Jul 6;361(6397):88-92. Doi: 10.1126/science.aat3628.
38. McVean G. A genealogical interpretation of principal components analysis. *PloS Genet*. 2009 Oct;5(10):e1000686. Doi: 10.1371/journal.pgen.1000686. Epub 2009 Oct 16.
39. Molloy EK, Durvasula A, Sankararaman S. Advancing admixture graph estimation via maximum likelihood network orientation. *Bioinformatics*. 2021 Jul 12;37(Suppl_1):i142-i150. Doi: 10.1093/bioinformatics/btab267.
40. Moreno-Mayar JV, Vinner L, de Barros Damgaard P, de la Fuente C, Chan J, Spence JP, Allentoft ME, Vimala T, Racimo F, Pinotti T, Rasmussen S, Margaryan A, Iraeta Orbegozo M, Mylopotamitaki D, Wooler M, Bataille C, Becerra-Valdivia L, Chivall D, Comeskey D, Deviese T, Grayson DK, George L, Harry H, Alexandersen V, Primeau C, Erlandson J, Rodrigues-Carvalho C, Reis S, Bastos MQR, Cybulski J, Vullo C, Morello F, Vilar M, Wells S, Gregersen K, Hansen KL,

- 2360 Lynnerup N, Mirazón Lahr M, Kjær K, Strauss A, Alfonso-Durruty M, Salas A, Schroeder H,
2361 Higham T, Malhi RS, Rasic JT, Souza L, Santos FR, Malaspinas AS, Sikora M, Nielsen R, Song YS,
2362 Meltzer DJ, Willerslev E. Early human dispersals within the Americas. *Science*. 2018 Dec
2363 7;362(6419):eaav2621. Doi: 10.1126/science.aav2621.
- 2364 41. Narasimhan VM, Patterson N, Moorjani P, Rohland N, Bernardos R, Mallick S, Lazaridis I,
2365 Nakatsuka N, Olalde I, Lipson M, Kim AM, Olivieri LM, Coppa A, Vidale M, Mallory J, Moiseyev V,
2366 Kitov E, Monge J, Adamski N, Alex N, Broomandkhoshbacht N, Candilio F, Callan K, Cheronet O,
2367 Culleton BJ, Ferry M, Fernandes D, Freilich S, Gamarra B, Gaudio D, Hajdinjak M, Harney É,
2368 Harper TK, Keating D, Lawson AM, Mah M, Mandl K, Michel M, Novak M, Oppenheimer J, Rai N,
2369 Sirak K, Slon V, Stewardson K, Zalzal F, Zhang Z, Akhatov G, Bagashev AN, Bagnera A, Baitanayev
2370 B, Bendezu-Sarmiento J, Bissembaev AA, Bonora GL, Chavgynov TT, Chikisheva T, Dashkovskiy
2371 PK, Derevianko A, Dobeš M, Douka K, Dubova N, Duisengali MN, Enshin D, Epimakhov A, Fribus
2372 AV, Fuller D, Goryachev A, Gromov A, Grushin SP, Hanks B, Judd M, Kazizov E, Khokhlov A, Krygin
2373 AP, Kupriyanova E, Kuznetsov P, Luiselli D, Maksudov F, Mamedov AM, Mamirov TB, Meiklejohn
2374 C, Merrett DC, Micheli R, Mochalov O, Mustafokulov S, Nayak A, Pettener D, Potts R, Razhev D,
2375 Rykun M, Sarno S, Savenkova TM, Sikhymbaeva K, Slepchenko SM, Soltobaev OA, Stepanova N,
2376 Svyatko S, Tabaldiev K, Teschler-Nicola M, Tishkin AA, Tkachev VV, Vasilyev S, Velemínský P,
2377 Voyakin D, Yermolayeva A, Zahir M, Zubkov VS, Zubova A, Shinde VS, Lalueza-Fox C, Meyer M,
2378 Anthony D, Boivin N, Thangaraj K, Kennett DJ, Frachetti M, Pinhasi R, Reich D. The formation of
2379 human populations in South and Central Asia. *Science*. 2019 Sep 6;365(6457):eaat7487. Doi:
2380 10.1126/science.aat7487.
- 2381 42. Patterson N, Moorjani P, Luo Y, Mallick S, Rohland N, Zhan Y, Genschoreck T, Webster T, Reich D.
2382 Ancient admixture in human history. *Genetics*. 2012 Nov;192(3):1065-93. Doi:
2383 10.1534/genetics.112.145037.
- 2384 43. Pickrell JK, Pritchard JK. Inference of population splits and mixtures from genome-wide allele
2385 frequency data. *PLoS Genet*. 2012;8(11):e1002967. Doi: 10.1371/journal.pgen.1002967.
- 2386 44. Posth C, Nakatsuka N, Lazaridis I, Skoglund P, Mallick S, Lamnidis TC, Rohland N, Nägele K,
2387 Adamski N, Bertolini E, Broomandkhoshbacht N, Cooper A, Culleton BJ, Ferraz T, Ferry M,
2388 Furtwängler A, Haak W, Harkins K, Harper TK, Hünemeier T, Lawson AM, Llamas B, Michel M,
2389 Nelson E, Oppenheimer J, Patterson N, Schiffels S, Sedig J, Stewardson K, Talamo S, Wang CC,
2390 Hublin JJ, Hubbe M, Harvati K, Nuevo Delaunay A, Beier J, Francken M, Kaulicke P, Reyes-
2391 Centeno H, Rademaker K, Trask WR, Robinson M, Gutierrez SM, Prufer KM, Salazar-García DC,
2392 Chim EN, Müller Plumm Gomes L, Alves ML, Liryo A, Ingles M, Oliveira RE, Bernardo DV, Barioni
2393 A, Wesolowski V, Scheifler NA, Rivera MA, Plens CR, Messineo PG, Figuti L, Corach D, Scabuzzo C,
2394 Eggers S, DeBlasis P, Reindel M, Méndez C, Politis G, Tomasto-Cagigao E, Kennett DJ, Strauss A,
2395 Fehren-Schmitz L, Krause J, Reich D. Reconstructing the Deep Population History of Central and
2396 South America. *Cell*. 2018 Nov 15;175(5):1185-1197.e22. doi: 10.1016/j.cell.2018.10.027.
- 2397 45. Prüfer K, Racimo F, Patterson N, Jay F, Sankararaman S, Sawyer S, Heinze A, Renaud G, Sudmant
2398 PH, de Filippo C, Li H, Mallick S, Dannemann M, Fu Q, Kircher M, Kuhlwilm M, Lachmann M,
2399 Meyer M, Ongyerth M, Siebauer M, Theunert C, Tandon A, Moorjani P, Pickrell J, Mullikin JC,
2400 Vohr SH, Green RE, Hellmann I, Johnson PL, Blanche H, Cann H, Kitzman JO, Shendure J, Eichler
2401 EE, Lein ES, Bakken TE, Golovanova LV, Doronichev VB, Shunkov MV, Derevianko AP, Viola B,
2402 Slatkin M, Reich D, Kelso J, Pääbo S. The complete genome sequence of a Neanderthal from the
2403 Altai Mountains. *Nature*. 2014 Jan 2;505(7481):43-9. Doi: 10.1038/nature12886.
- 2404 46. Raghavan M, Skoglund P, Graf KE, Metspalu M, Albrechtsen A, Moltke I, Rasmussen S, Stafford
2405 TW Jr, Orlando L, Metspalu E, Karmin M, Tambets K, Rootsi S, Mägi R, Campos PF, Balanovska E,
2406 Balanovsky O, Khusnutdinova E, Litvinov S, Osipova LP, Fedorova SA, Voevoda MI, DeGiorgio M,
2407 Sicheritz-Ponten T, Brunak S, Demeshchenko S, Kivisild T, Villems R, Nielsen R, Jakobsson M,
2408 Willerslev E. Upper Palaeolithic Siberian genome reveals dual ancestry of Native Americans.
2409 *Nature*. 2014 Jan 2;505(7481):87-91. Doi: 10.1038/nature12736.
- 2410 47. Raghavan M, Steinrücken M, Harris K, Schiffels S, Rasmussen S, DeGiorgio M, Albrechtsen A,
2411 Valdiosera C, Ávila-Arcos MC, Malaspinas AS, Eriksson A, Moltke I, Metspalu M, Homburger JR,
2412 Wall J, Cornejo OE, Moreno-Mayar JV, Korneliussen TS, Pierre T, Rasmussen M, Campos PF, de
2413 Barros Damgaard P, Allentoft ME, Lindo J, Metspalu E, Rodríguez-Varela R, Mansilla J,

- 2414 Henrickson C, Seguin-Orlando A, Malmström H, Stafford T Jr, Shringarpure SS, Moreno-Estrada
2415 A, Karmin M, Tambets K, Bergström A, Xue Y, Warmuth V, Friend AD, Singarayer J, Valdes P,
2416 Balloux F, Lebreiro I, Vera JL, Rangel-Villalobos H, Pettener D, Luiselli D, Davis LG, Heyer E,
2417 Zollikofer CPE, Ponce de León MS, Smith CI, Grimes V, Pike KA, Deal M, Fuller BT, Arriaza B,
2418 Standen V, Luz MF, Ricaut F, Guidon N, Osipova L, Voevoda MI, Posukh OL, Balanovsky O,
2419 Lavryashina M, Bogunov Y, Khusnutdinova E, Gubina M, Balanovska E, Fedorova S, Litvinov S,
2420 Malyarchuk B, Derenko M, Mosher MJ, Archer D, Cybulski J, Petzelt B, Mitchell J, Worl R,
2421 Norman PJ, Parham P, Kemp BM, Kivisild T, Tyler-Smith C, Sandhu MS, Crawford M, Villemes R,
2422 Smith DG, Waters MR, Goebel T, Johnson JR, Malhi RS, Jakobsson M, Meltzer DJ, Manica A,
2423 Durbin R, Bustamante CD, Song YS, Nielsen R, Willerslev E. Genomic evidence for the Pleistocene
2424 and recent population history of Native Americans. *Science*. 2015 Aug 21;349(6250):aab3884.
2425 Doi: 10.1126/science.aab3884.
- 2426 48. Reich D, Thangaraj K, Patterson N, Price AL, Singh L. Reconstructing Indian population history.
2427 *Nature*. 2009 Sep 24;461(7263):489-94. Doi: 10.1038/nature08365.
- 2428 49. Reich D, Patterson N, Kircher M, Delfin F, Nandineni MR, Pugach I, Ko AM, Ko YC, Jinam TA,
2429 Phipps ME, Saitou N, Wollstein A, Kayser M, Pääbo S, Stoneking M. Denisova admixture and the
2430 first modern human dispersals into Southeast Asia and Oceania. *Am J Hum Genet*. 2011 Oct
2431 7;89(4):516-28. Doi: 10.1016/j.ajhg.2011.09.005.
- 2432 50. Reich D, Patterson N, Campbell D, Tandon A, Mazieres S, Ray N, Parra MV, Rojas W, Duque C,
2433 Mesa N, García LF, Triana O, Blair S, Maestre A, Dib JC, Bravi CM, Bailliet G, Corach D, Hünemeier
2434 T, Bortolini MC, Salzano FM, Petzl-Erler ML, Acuña-Alonzo V, Aguilar-Salinas C, Canizales-
2435 Quinteros S, Tusié-Luna T, Riba L, Rodríguez-Cruz M, Lopez-Alarcón M, Coral-Vazquez R, Canto-
2436 Cetina T, Silva-Zolezzi I, Fernandez-Lopez JC, Contreras AV, Jimenez-Sanchez G, Gómez-Vázquez
2437 MJ, Molina J, Carracedo A, Salas A, Gallo C, Poletti G, Witonsky DB, Alkorta-Aranburu G, Sukernik
2438 RI, Osipova L, Fedorova SA, Vasquez R, Villena M, Moreau C, Barrantes R, Pauls D, Excoffier L,
2439 Bedoya G, Rothhammer F, Dugoujon JM, Larrouy G, Klitz W, Labuda D, Kidd J, Kidd K, Di Rienzo
2440 A, Freimer NB, Price AL, Ruiz-Linares A. Reconstructing Native American population history.
2441 *Nature*. 2012 Aug 16;488(7411):370-4. Doi: 10.1038/nature11258.
- 2442 51. Rogers AR. Legofit: estimating population history from genetic data. *BMC Bioinformatics*. 2019
2443 Oct 28;20(1):526. Doi: 10.1186/s12859-019-3154-1.
- 2444 52. Schiffels S, Durbin R. Inferring human population size and separation history from multiple
2445 genome sequences. *Nat Genet*. 2014 Aug;46(8):919-25. Doi: 10.1038/ng.3015.
- 2446 53. Schiffels S, Haak W, Paajanen P, Llamas B, Popescu E, Loe L, Clarke R, Lyons A, Mortimer R, Sayer
2447 D, Tyler-Smith C, Cooper A, Durbin R. Iron Age and Anglo-Saxon genomes from East England
2448 reveal British migration history. *Nat Commun*. 2016 Jan 19;7:10408. Doi:
2449 10.1038/ncomms10408.
- 2450 54. Seguin-Orlando A, Korneliussen TS, Sikora M, Malaspina AS, Manica A, Moltke I, Albrechtsen A,
2451 Ko A, Margaryan A, Moiseyev V, Goebel T, Westaway M, Lambert D, Khartanovich V, Wall JD,
2452 Nigst PR, Foley RA, Lahr MM, Nielsen R, Orlando L, Willerslev E. Paleogenomics. Genomic
2453 structure in Europeans dating back at least 36,200 years. *Science*. 2014 Nov 28;346(6213):1113-
2454 8. Doi: 10.1126/science.aaa0114.
- 2455 55. Shinde V, Narasimhan VM, Rohland N, Mallick S, Mah M, Lipson M, Nakatsuka N, Adamski N,
2456 Broomandkhoshbacht N, Ferry M, Lawson AM, Michel M, Oppenheimer J, Stewardson K, Jadhav
2457 N, Kim YJ, Chatterjee M, Munshi A, Panyam A, Waghmare P, Yadav Y, Patel H, Kaushik A,
2458 Thangaraj K, Meyer M, Patterson N, Rai N, Reich D. An Ancient Harappan Genome Lacks
2459 Ancestry from Steppe Pastoralists or Iranian Farmers. *Cell*. 2019 Oct 17;179(3):729-735.e10. doi:
2460 10.1016/j.cell.2019.08.048.
- 2461 56. Sikora M, Pitulko VV, Sousa VC, Allentoft ME, Vinner L, Rasmussen S, Margaryan A, de Barros
2462 Damgaard P, de la Fuente C, Renaud G, Yang MA, Fu Q, Dupanloup I, Giampoudakis K, Nogués-
2463 Bravo D, Rahbek C, Kroonen G, Peyrot M, McColl H, Vasilyev SV, Veselovskaya E, Gerasimova M,
2464 Pavlova EY, Chasnyk VG, Nikolskiy PA, Gromov AV, Khartanovich VI, Moiseyev V, Grebenyuk PS,
2465 Fedorchenko AY, Lebedintsev AI, Slobodin SB, Malyarchuk BA, Martiniano R, Meldgaard M,
2466 Arppe L, Palo JU, Sundell T, Mannermaa K, Putkonen M, Alexandersen V, Primeau C,
2467 Baimukhanov N, Malhi RS, Sjögren KG, Kristiansen K, Wessman A, Sajantila A, Lahr MM, Durbin

- 2468 R, Nielsen R, Meltzer DJ, Excoffier L, Willerslev E. The population history of northeastern Siberia
2469 since the Pleistocene. *Nature*. 2019 Jun;570(7760):182-188. Doi: 10.1038/s41586-019-1279-z.
- 2470 57. Skoglund P, Posth C, Sirak K, Spriggs M, Valentin F, Bedford S, Clark GR, Reepmeyer C, Petchey F,
2471 Fernandes D, Fu Q, Harney E, Lipson M, Mallick S, Novak M, Rohland N, Stewardson K, Abdullah
2472 S, Cox MP, Friedlaender FR, Friedlaender JS, Kivisild T, Koki G, Kusuma P, Merriwether DA, Ricaut
2473 FX, Wee JT, Patterson N, Krause J, Pinhasi R, Reich D. Genomic insights into the peopling of the
2474 Southwest Pacific. *Nature*. 2016 Oct 27;538(7626):510-513. Doi: 10.1038/nature19844.
- 2475 58. Speidel L, Forest M, Shi S, Myers SR. A method for genome-wide genealogy estimation for
2476 thousands of samples. *Nat Genet*. 2019 Sep;51(9):1321-1329. Doi: 10.1038/s41588-019-0484-x.
- 2477 59. Tambets K, Yunusbayev B, Hudjashov G, Ilumäe AM, Rootsi S, Honkola T, Vesakoski O, Atkinson
2478 Q, Skoglund P, Kushniarevich A, Litvinov S, Reidla M, Metspalu E, Saag L, Rantanen T, Karmin M,
2479 Parik J, Zhadanov SI, Gubina M, Damba LD, Bermisheva M, Reisberg T, Dibirova K, Evseeva I,
2480 Nelis M, Klovins J, Metspalu A, Esko T, Balanovsky O, Balanovska E, Khusnutdinova EK, Osipova
2481 LP, Voevoda M, Villems R, Kivisild T, Metspalu M. Genes reveal traces of common recent
2482 demographic history for most of the Uralic-speaking populations. *Genome Biol*. 2018 Sep
2483 21;19(1):139. Doi: 10.1186/s13059-018-1522-1.
- 2484 60. Terhorst J, Kamm JA, Song YS. Robust and scalable inference of population history from
2485 hundreds of unphased whole genomes. *Nat Genet*. 2017 Feb;49(2):303-309. Doi:
2486 10.1038/ng.3748.
- 2487 61. Vallini L, Marciani G, Aneli S, Bortolini E, Benazzi S, Pievani T, Pagani L. Genetics and material
2488 culture support repeated expansions into Paleolithic Eurasia from a population hub out of Africa.
2489 *Genome Biol Evol*. 2022 Apr 10;14(4):evac045. Doi: 10.1093/gbe/evac045.
- 2490 62. van de Loosdrecht M, Bouzouggar A, Humphrey L, Posth C, Barton N, Aximu-Petri A, Nickel B,
2491 Nagel S, Talbi EH, El Hajraoui MA, Amzazi S, Hublin JJ, Pääbo S, Schiffels S, Meyer M, Haak W,
2492 Jeong C, Krause J. Pleistocene North African genomes link Near Eastern and sub-Saharan African
2493 human populations. *Science*. 2018 May 4;360(6388):548-552. Doi: 10.1126/science.aar8380.
- 2494 63. Wang CC, Reinhold S, Kalmykov A, Wissgott A, Brandt G, Jeong C, Cheronet O, Ferry M, Harney E,
2495 Keating D, Mallick S, Rohland N, Stewardson K, Kantorovich AR, Maslov VE, Petrenko VG, Erlikh
2496 VR, Atabiev BC, Magomedov RG, Kohl PL, Alt KW, Pichler SL, Gerling C, Meller H, Vardanyan B,
2497 Yeganyan L, Rezepkin AD, Mariaschk D, Berezina N, Gresky J, Fuchs K, Knipper C, Schiffels S,
2498 Balanovska E, Balanovsky O, Mathieson I, Higham T, Berezin YB, Buzhilova A, Trifonov V, Pinhasi
2499 R, Belinskij AB, Reich D, Hansen S, Krause J, Haak W. Ancient human genome-wide data from a
2500 3000-year interval in the Caucasus corresponds with eco-geographic regions. *Nat Commun*. 2019
2501 Feb 4;10(1):590. Doi: 10.1038/s41467-018-08220-8.
- 2502 64. Wang CC, Yeh HY, Popov AN, Zhang HQ, Matsumura H, Sirak K, Cheronet O, Kovalev A, Rohland
2503 N, Kim AM, Mallick S, Bernardos R, Tumen D, Zhao J, Liu YC, Liu JY, Mah M, Wang K, Zhang Z,
2504 Adamski N, Broomandkhoshbacht N, Callan K, Candilio F, Carlson KSD, Culleton BJ, Eccles L,
2505 Freilich S, Keating D, Lawson AM, Mandl K, Michel M, Oppenheimer J, Özdoğan KT, Stewardson
2506 K, Wen S, Yan S, Zalzal F, Chuang R, Huang CJ, Looch H, Shiung CC, Nikitin YG, Tabarev AV,
2507 Tishkin AA, Lin S, Sun ZY, Wu XM, Yang TL, Hu X, Chen L, Du H, Bayarsaikhan J, Mijiddorj E,
2508 Erdenebaatar D, Iderkhangai TO, Myagmar E, Kanzawa-Kiriyama H, Nishino M, Shinoda KI,
2509 Shubina OA, Guo J, Cai W, Deng Q, Kang L, Li D, Li D, Lin R, Nini, Shrestha R, Wang LX, Wei L, Xie
2510 G, Yao H, Zhang M, He G, Yang X, Hu R, Robbeets M, Schiffels S, Kennett DJ, Jin L, Li H, Krause J,
2511 Pinhasi R, Reich D. Genomic insights into the formation of human populations in East Asia.
2512 *Nature*. 2021 Mar;591(7850):413-419. Doi: 10.1038/s41586-021-03336-2.
- 2513 65. Yan J, Patterson N, Narasimhan VM. miqoGraph: fitting admixture graphs using mixed-integer
2514 quadratic optimization. *Bioinformatics*. 2021 Aug 25;37(16):2488-2490. Doi:
2515 10.1093/bioinformatics/btaa988.
- 2516 66. Yang MA, Gao X, Theunert C, Tong H, Aximu-Petri A, Nickel B, Slatkin M, Meyer M, Pääbo S,
2517 Kelso J, Fu Q. 40,000-Year-Old Individual from Asia Provides Insight into Early Population
2518 Structure in Eurasia. *Curr Biol*. 2017 Oct 23;27(20):3202-3208.e9. doi:
2519 10.1016/j.cub.2017.09.030. Epub 2017 Oct 12.
- 2520 67. Yang MA, Fan X, Sun B, Chen C, Lang J, Ko YC, Tsang CH, Chiu H, Wang T, Bao Q, Wu X, Hajdinjak
2521 M, Ko AM, Ding M, Cao P, Yang R, Liu F, Nickel B, Dai Q, Feng X, Zhang L, Sun C, Ning C, Zeng W,

- 2522 Zhao Y, Zhang M, Gao X, Cui Y, Reich D, Stoneking M, Fu Q. Ancient DNA indicates human
- 2523 population shifts and admixture in northern and southern China. *Science*. 2020 Jul
- 2524 17;369(6501):282-288. Doi: 10.1126/science.aba0909.
- 2525 68. Zhang M, Yan S, Pan W, Jin L. Phylogenetic evidence for Sino-Tibetan origin in northern China in
- 2526 the Late Neolithic. *Nature*. 2019 May;569(7754):112-115. doi: 10.1038/s41586-019-1153-z.

Supplementary Tables

Publication	Figure in the original publication	Groups (populations)	Singleton pseudo-haploid populations	No. of negative F3-stats (all pop = YES)	Admixture events	Pd1 model: log likelihood (LL)	Pd1 model: LL, median of bootstrap dist.	Pd1 model: worst fit actual (WBF) SE	SNPs used	Settings for calculating f2-statistics	Topology search constraints and population modifications	Iterations	Iterations confirming published graph	Distinct alternative topologies found ^(a)	Significantly better fitting topologies ^(a)	Non-significantly better fitting topologies ^(a)	Significantly worse fitting topologies ^(a)	Non-significantly worse fitting topologies ^(a)	Significantly better fitting topologies ^(a)	Non-significantly better fitting topologies ^(a)	P-value best alternative vs. publ.	Used in table 1		
Bergström et al. 2020 1e	6	4	0	2	4.7	4.9			312,282	max miss=0, b1g_size=4000000, adjust_pseudohaploid=T	Andean Fox OG	100	38	14	0	0	3	11	0.0	0.0	21.4	78.6	0.060*	
	7	5	0	3	8.4	11.8	2.1					10,000	33.7	221	1	5	37	178	0.5	2.3	16.7	80.5	0.332	yes
Lazaridis et al. 2014 3	6	1	6	2	3.0	4.2			624,583	max miss=0, b1g_size=0.05, adjust_pseudohaploid=T	Mbuti OG	100	32	14	0	0	0	14	0.0	0.0	0.0	100.0	0.032*	
	7	1	6	4	8.7	11.1	2.2		642,247	minac2=2		1,000	0	306	3	37	247	19	1.0	12.1	80.7	6.2	0.032	yes
	9	3	6	3	22.1	40.9	2.8		19,017	max miss=0, b1g_size=0.05, adjust_pseudohaploid=T	Mota OG [4]	4,000	0	398	0	89	266	43	0.0	22.4	66.8	10.8	0.136	
	8	2	2	3	29	40.8	3.2		470,389	max miss=0, b1g_size=0.05, adjust_pseudohaploid=T	Mota OG; Indus_Periphery group expanded to 4 Ind., relatives and contaminated Ind. removed	4,000	0	216	4	61	79	72	1.9	28.2	36.6	33.3	0.072	
	8	1	10	4	N/A	N/A	2.6		249,009	minac2=2		4,000	13	143	0	4	5	134	0.0	2.8	3.5	93.7	0.088*	yes
Shinde et al. 2019 3	3b				3	907.3	942	28.3		max miss=0, b1g_size=4000000, adjust_pseudohaploid=T		1,000	25	335									0.002	
	Ext 5d				4	363.3	386.2	15.7				1,000	10	531								0.016		
	Ext 5e				5	267.7	286.5	10.4				1,000	0	747								0.002		
	N/A	10	3	3	6				6,343,116			1,000	894											
					7	N/A	N/A	N/A				1,000	986											
					8							1,000	999											
					9							1,000	1,000											
	3b				3	679.9	711.1	23.9				1,000	1	324	22	51	78	173	6.8	15.7	24.1	53.4	0.040	yes
	Ext 5d				4	202.8	224.4	14.1				1,000	30	535	0	0	24	511	0.0	0.0	4.5	95.5	0.788*	yes
	Ext 5e				5	121	139.6	6.9				1,000	3	784	0	2	223	559	0.0	0.3	28.4	71.3	0.720	yes
Librado et al. 2021		10	3	3	6				1,767,419	max miss=0, b1g_size=4000000, adjust_pseudohaploid=T, minac2=2		1,000	954											
	N/A				7	N/A	N/A	N/A				1,000	995											
					8							1,000	1,000											
					9							1,000	1,000											
	3b				3	1133	1163	33.8				1,000	0	363									0.002	
	Ext 5d				4	477.1	500.5	16.3				1,000	3	567									0.002	
	Ext 5e				5	380.5	395	13.2				1,000	0	729									0.002	
	N/A	10	1	1	6				7,403,037	max miss=0, b1g_size=4000000, adjust_pseudohaploid=T	Donkey OG	1,000	906											
					7	N/A	N/A	N/A				1,000	981											
					8							1,000	999											
					9							1,000	998											
Hajdinjak et al. 2021 2d	11	5	0	8	35.8	65.1	2.8		283,890	max miss=0, b1g_size=0.05, adjust_pseudohaploid=T	Denisovan OG [4]	4,000	0	3,996				15.9	80.8	3.2	0.0	0.008		
												2,000	0	1,999					26.0	63.5	6.6	3.9	0.002	
	12	6	0	8	71.6	112.4	4.8		263,698	max miss=0, b1g_size=0.05, adjust_pseudohaploid=T	Vindija, Mbuti unadmixed, Denisovan max. 1 admix. event	2,000	0	1,995				56.8	41.4	1.8	0.0	0.002		
												2,000	0	1,988					15.7	55.7	6.6	22.0	0.002	yes
	S3.24	7	1	0	4	5.4	8.7	1.9	801,362	max miss=0, b1g_size=4000000, adjust_pseudohaploid=T, minac2=2	Chimp OG [4]	2,000	0	778	2	199	489	89	0.3	25.6	62.9	11.4	0.260	
	S3.25	10	1	15	8	33.6	51.2	2.7	838,910	max miss=0, b1g_size=0.05, adjust_pseudohaploid=T	Chimp OG	10,000	0	9,927				6.8	83.5	9.1	0.6	0.032		
	Ext 4	12	2	22	12	29.5	61.6	2.2	363,131	max miss=0, b1g_size=0.05, adjust_pseudohaploid=T	Chimp OG	2,000	0	2,000								0.064		
					12	21.8	52.3	2.3	211,738	max miss=0, b1g_size=0.05, adjust_pseudohaploid=T	Chimp OG, Altai unadmixed, French min. 1 admix. event	2,000	0	2,000								0.100		
					11	21.8	52.5	2.3	543,124	max miss=0, b1g_size=0.05, adjust_pseudohaploid=T		2,000	0	906				0.0	28.5	68.5	3.0	0.100	yes	
	S13-9 [2]	9	2	0	4	16	29.6	2.5	544,068	max miss=0, b1g_size=0.05, adjust_pseudohaploid=T	Denisovan OG [4]	2,000	0	1,524				0.0	11.9	77.1	10.4	0.176		
	S13-10a [3]	10	3	0	5	23.3	41.8	2.5		max miss=0, b1g_size=0.05, adjust_pseudohaploid=T	Denisovan OG [4]	2,000	0	1,895				0.0	38.8	60.8	0.4	0.524		
												2,000	0	1,993					0.0	44.0	56.0	0.0	0.288	
	Ext 6	12	3	3	8	62.1	99.2	3.1	496,233	max miss=0, b1g_size=0.05, adjust_pseudohaploid=T	Denisovan OG [4]	1,900	0	1,895								0.004		
					66.4	106	3.8		203,753	same as above + minac2=2	Deniskovan OG [4]	2,000	0	1,993								0.004		
												2,000	0	1,778					12.6	84.3	3.1	0.0	0.008	yes
					10	65.7	116.4	3.2				1,000	0	1,000					33.2	13.6	50.0	3.2	0.002	
	3f [left]	13	4	0	6	76.5	132	3.8				1,000	0	996					28.5	68.7	2.8	0.0	0.002	
												1,000	0	894					0.3	17.1	34.6	48.0	0.016	yes
												1,000	0	966					0.1	37.0	52.3	10.6	0.024	
					10	85.3	147.8	3.1				1,000	0	1,000					13.2	30.0	17.6	39.2	0.002	
												1,000	0	1,000					26.0	44.4	15.6	14.0	0.002	
					6	102.4	171.9	4.2				4,446	0	2,785					0.1	0.9	9.8	89.2	0.112	yes
												3,505	0	2,894					0.1	6.7	40.3	52.8	0.056	
												published model was recovered												
⁽¹⁾ Counting at most one topology per iteration, not double counting topologies found in multiple iterations																								
⁽²⁾ a 1% gene flow from Loschbour to the Mongolia Neolithic group was dropped from the published model to decrease the complexity of the baseline model. That increased model LL slightly (by 0.5 log-units).																								
⁽³⁾ a 1% gene flow from Loschbour to the Mongolia Neolithic group was dropped from the published model to decrease the complexity of the baseline model. That increased model LL slightly (by 2.4 log-units).																								
⁽⁴⁾ OG was set for generating random starting graphs, but not for the "find_graphs" algorithm itself																								
[*] here the direction of comparison is reversed since the published model is the highest-ranking among all models found																								

[1] Counting at most one topology per iteration, not double counting topologies found in multiple iterations

[2] a 1% gene flow from Loschbour to the Mongolia Neolithic group was dropped from the published model to decrease the complexity of the baseline model. That increased model LL slightly (by 0.5 log-units).

[3] a 1% gene flow from Loschbour to the Mongolia Neolithic group was dropped from the published model to decrease the complexity of the baseline model. That increased model LL slightly (by 2.4 log-units).

[4] OG was set for generating random starting graphs, but not for the "find_graphs" algorithm itself

* here the direction of comparison is reversed since the published model is the highest-ranking among all models found

Table S1: Published graphs in the context of automatically found graphs. We compared 22 different graphs from 8 publications to alternative graphs inferred on the same or very similar data; these *findGraphs* runs are highlighted in blue in the “iterations” column. In total, 51 *findGraphs* runs are summarized here since in some cases models more complex or less complex than the published one were explored and/or different population compositions were tested (see the “Topology search constraints and population modifications” column and footnotes for details). The columns with names in blue show various information on the published graphs or their modified versions and some properties of the published population sets. The columns with names in magenta show settings used for calculating f -statistics and for exploring the admixture graph space, and the number of SNPs used that depends on them. The columns with names in black summarize the outcomes of *findGraphs* runs, i.e., the properties of alternative model sets found.

Publication: Last name of the first author and year of the relevant publication.

Figure in the original publication: Figure number in the original paper where the admixture graph is presented.

Groups (populations): The number of populations in each graph.

Singleton pseudo-diploid populations: The number of populations in the graph composed of a single pseudo-diploid individual. Calculation of negative “admixture” f_3 -statistics is impossible for such populations since their heterozygosity cannot be estimated (see the text for details).

No. of negative f_3 -stats (allsnps: YES): The number of negative f_3 -statistics among all possible f_3 -statistics for a given set of populations when all available sites are used for each statistic. If no negative f_3 -statistics exist for a set of populations, admixture graph fits are not affected by the “minac2=2” setting intended for accurate calculation of f -statistics for non-singleton pseudo-diploid groups.

Admixture events: The number of admixture events in each graph.

Publ. model: log-likelihood (LL): Log-likelihood score of the published graph fitted to the SNP set shown in the “SNPs used” column.

Publ. model: LL, median of bootstrap distr.: Median of the log-likelihood scores of 100 or 500 fits of the published graph using bootstrap resampled SNPs.

Publ. model: worst residual (WR), SE: The worst f -statistic residual of the published graph fitted to the SNP set shown in the “SNPs used” column, measured in standard errors (SE).

SNPs used: The number of SNPs (with no missing data at the group level) used for fitting the admixture graph. For all case studies, we tested the original data (SNPs, population composition, and the published graph topology) and obtained model fits very similar to the published ones. However, for the purpose of efficient topology search we adjusted settings for f_3 -statistic calculation, population composition, or graph complexity as shown here and discussed in the text.

Settings for calculating f_2 -statistics: Arguments of the *extract_f2* function used for calculating all possible f_2 -statistics for a set of groups, which were then used by *findGraphs* for calculating f_3 -statistics needed for fitting admixture graph models. See Methods for descriptions of each argument.

Topology search constraints and population modifications: Constraints applied when generating random starting graphs and/or when searching the topology space.

Modifications of the original population composition are also described in this column, where applicable.

Iterations: The number of *findGraphs* iterations, each started from a random graph of a certain complexity. For each case study, *findGraphs* setups that were considered optimal are highlighted in blue in this column.

Iterations confirming published graph: The number of iterations in which the resulting graph was topologically identical to the published graph. In the cases when the published model was irrelevant since more complex graphs were explored, “N/A” appears in this and subsequent columns. If less complex models were explored, the published model was still relevant since its version without selected admixture edges was tested.

Distinct alternative topologies found: The number of distinct newly found topologies. If graph complexity was equal to (or less than) that of the published graph, the published topology (or its simplified version) is not counted here. If graph complexity exceeded that of the published graph, all newly found topologies are counted. If the published topology was recovered by *findGraphs*, the numbers in this column are shown in bold.

2566 **Significantly better fitting topologies:** The number of distinct topologies that fit significantly better than the published graph according to the bootstrap model comparison test
2567 (two-tailed empirical p -value <0.05). If the number of distinct topologies was very large, a representative sample of models (1/20 to 1/3 of models evenly distributed along the
2568 log-likelihood spectrum) was compared to the published one instead. These cases are marked as “a fraction of models tested” in this column. If model complexity was higher than
2569 that of the published model, model comparison was irrelevant and was not performed.
2570 **Non-significantly better fitting topologies:** The number of distinct topologies that fit non-significantly (nominally) better than the published graph according to the bootstrap
2571 model comparison test (two-tailed empirical p -value ≥ 0.05).
2572 **Non-significantly worse fitting topologies:** The number of distinct topologies that fit non-significantly (nominally) worse than the published graph according to the bootstrap
2573 model comparison test (two-tailed empirical p -value ≥ 0.05).
2574 **Significantly worse fitting topologies:** The number of distinct topologies that fit significantly worse than the published graph according to the bootstrap model comparison test
2575 (two-tailed empirical p -value ≥ 0.05).
2576 **Significantly better fitting topologies, %:** The percentage of distinct topologies that fit significantly better than the published graph according to the bootstrap model comparison
2577 test (two-tailed empirical p -value <0.05). If the number of distinct topologies was very large, a representative sample of models (1/20 to 1/3 of models evenly distributed along
2578 the log-likelihood spectrum) was compared to the published one instead, and the percentages in this and following columns were calculated on this sample.
2579 **Non-significantly better fitting topologies, %:** The percentage of distinct topologies that fit non-significantly (nominally) better than the published graph according to the
2580 bootstrap model comparison test (two-tailed empirical p -value ≥ 0.05).
2581 **Non-significantly worse fitting topologies, %:** The percentage of distinct topologies that fit non-significantly (nominally) worse than the published graph according to the
2582 bootstrap model comparison test (two-tailed empirical p -value ≥ 0.05).
2583 **Significantly worse fitting topologies, %:** The percentage of distinct topologies that fit significantly worse than the published graph according to the bootstrap model comparison
2584 test (two-tailed empirical p -value ≥ 0.05).
2585 **P-value best alternative vs. publ.:** An empirical two-tailed p -value of a test comparing log-likelihood distributions across bootstrap replicates for two topologies, the highest-
2586 ranking newly found topology and the published topology. In some cases, the highest ranking newly found topology (according to LL) has a fit that is not significantly better than
2587 that of the published model, but other newly found models fit significantly better despite having higher LL. P -values below 0.05 are highlighted in green.
2588 **Used in Table 1:** Here the *findGraphs* runs featured in **Table 1** are marked.

Table S2: Statistics for shotgun sequencing of individual I8726

Individual ID	I8726
Archaeological IDs	SHAR_201 (Grave 201)
Skeletal element	Petrous bone
Location	Seistan, Shahr-i-Sokhta, Iran
Archaeological context date	3100-3000 BCE
Latitude	30.649857
Longitude	61.400311
First publication of library	Narasimhan, Patterson et al. Science 2019
Library ID	S8726.E1.L1
HiSeqX10 lanes for shotgun sequencing	3
Molecular sex	Male
Mean coverage measured on 1240k autosomal targets	2.60306

2592 **Supplementary Items Available as Separate Files**

2593

2594 **5 supplementary tables are available as separate files (Tables S3-S7)**

2595

2596 **13 supplementary figures are available as separate files (Figures S1-S13)**