

Resolving genetic linkage reveals patterns of selection in HIV-1 evolution

Muhammad S. Sohail^{1,*}, Raymond H. Y. Louie^{1,2,3,4,*}, Matthew R. McKay^{1,5,†}, and John P. Barton^{6,†}

¹Department of Electronic and Computer Engineering, Hong Kong University of Science and Technology, Hong Kong. ²Institute for Advanced Study, Hong Kong University of Science and Technology, Hong Kong. ³Kirby Institute, University of New South Wales, Australia. ⁴School of Medical Sciences, University of New South Wales, Australia. ⁵Department of Chemical and Biological Engineering, Hong Kong University of Science and Technology, Hong Kong. ⁶Department of Physics and Astronomy, University of California, Riverside, USA. *These authors contributed equally to this work. †Address correspondence to: m.mckay@ust.hk, john.barton@ucr.edu.

Identifying the genetic drivers of adaptation is a necessary step in understanding the dynamics of rapidly evolving pathogens and cancer. However, signals of selection are obscured by the complex, stochastic nature of evolution. Pervasive effects of genetic linkage, including genetic hitchhiking and clonal interference between beneficial mutants, challenge our ability to distinguish the selective effect of individual mutations. Here we describe a method to infer selection from genetic time series data that systematically resolves the confounding effects of genetic linkage. We applied our method to investigate patterns of selection in intrahost human immunodeficiency virus (HIV)-1 evolution, including a case in an individual who develops broadly neutralizing antibodies (bnAbs). Most variants that arise are observed to have negligible effects on inferred selection at other sites, but a small minority of highly influential variants have strong and far-reaching effects. In particular, we found that accounting for linkage is crucial for estimating selection due to clonal interference between escape mutants and other variants that sweep rapidly through the population. We observed only modest selection for antibody escape, in contrast with strong selection for escape from CD8⁺ T cell responses. Weak selection for escape from antibody responses may facilitate bnAb development by diversifying the viral population. Our results provide a quantitative description of the evolution of HIV-1 in response to host immunity, including selection on the viral population that accompanies bnAb development. More broadly, our analysis argues for the importance of resolving linkage effects in studies of natural selection.

Evolving populations exhibit complex dynamics. Cancers (1–6) and pathogens such as HIV-1 (7–9) and influenza (10, 11) generate multiple beneficial mutations that increase fitness or allow them to escape immunity. Subpopulations with different beneficial mutations then compete with one another for dominance, referred to as clonal interference, resulting in the loss of some mutations that increase fitness (12). Neutral or deleterious mutations can also hitchhike to high frequencies if they occur on advantageous genetic backgrounds (13). Experiments have demonstrated that these features of genetic linkage are pervasive in nature (14–16).

Linkage makes distinguishing the fitness effects of individual mutations challenging because their dynamics are contingent on the genetic background on which they appear. Lineage tracking experiments can be used to identify beneficial mutations (17), but they cannot readily be applied to evolution in natural conditions, such as HIV-1 evolution within or between hosts. Current computational methods to infer fitness from population dynamics ignore linkage or suffer from serious computational costs (18–23).

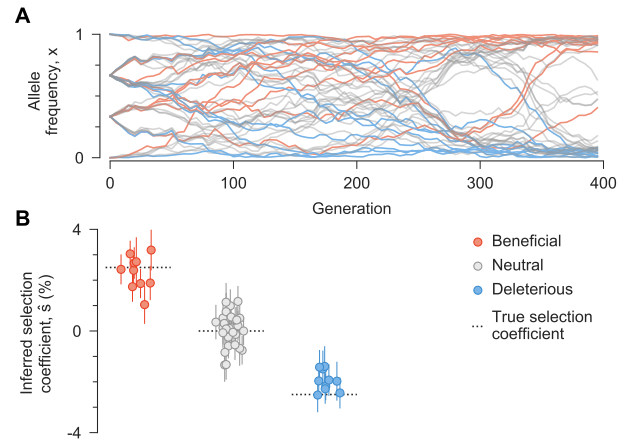


Fig. 1. MPL accurately recovers selection from complex dynamics. **A**, Simulated allele frequency trajectories in a model with 10 beneficial, 30 neutral, and 10 deleterious mutant alleles. The initial population is a mix of three subpopulations with random mutations. Selection is challenging to discern from individual trajectories alone. **B**, Selection coefficients inferred by MPL are close to their true values. Error bars denote theoretical standard deviations for inferred coefficients. Simulation parameters are the same as those defined in Supplementary Fig. 1.

Here we describe a method to infer selection from genetic time series data and demonstrate its ability to resolve linkage effects. We apply our method to reveal patterns of selection in intrahost HIV-1 evolution. Our approach is to efficiently quantify the probability of an evolutionary ‘path,’ defined by the set of all mutant allele frequencies at each time, using a path integral method derived from statistical physics (Methods). This expression can be analytically inverted to find the parameters that are most likely to have generated a path.

To define the path integral, we consider Wright-Fisher population dynamics with selection, mutation, and recombination, in the diffusion limit (24). Under an additive fitness model, the fitness of any individual is a sum of selection coefficients s_i , which quantify the selective advantage of mutant allele i relative to wild-type. The probability of an evolutionary path is then a product of probabilities of changes in mutant allele frequencies at each locus between successive generations, including the influence of selection at linked loci.

Applying Bayes’ theorem leads to an analytical expression for the maximum *a posteriori* vector of selection coefficients $\hat{s}$ corresponding to a path (Methods),

$$\hat{s} = (C_{\text{int}} + \gamma I)^{-1} (\Delta x - \mu_{\text{fl}}). \quad (1)$$

The integrated covariance matrix of mutant allele frequencies

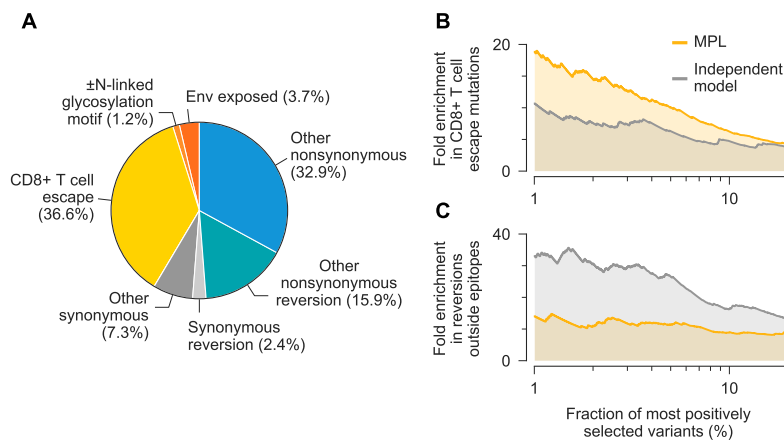


Fig. 2. Patterns of strong selection in within-host HIV evolution. **A**, Among the top 1% most beneficial variants across individuals, mutations to escape from T cell-mediated immunity are especially common. **B**, Due to clonal interference between escape mutants, MPL identifies more escape variants to be strongly beneficial than an independent model which ignores covariance. **C**, In contrast, the independent model estimates an excess in the number of strongly beneficial reversions.

C_{int} accounts for the speed of evolution and linkage effects. Here γ quantifies the width of a Gaussian prior distribution for selection coefficients, and I is the identity matrix. Intuitively, the net change in mutant allele frequencies Δx and the integrated mutational flux μ_{fl} determine whether the dynamics of a mutant allele appear to be beneficial or deleterious when linkage is ignored. Because equation (1) emerges from the likelihood of allele frequency trajectories, a subset of the full genotype distribution, we refer to it as the marginal path likelihood (MPL) estimate of the selection coefficients.

To test the ability of MPL to uncover selection, we analyzed data from simulations of a variety of evolutionary scenarios (Supplementary Information). Even in cases with strong linkage (Fig. 1A), MPL accurately recovers true selection coefficients (Fig. 1B). Further tests indicated that performance remains strong even when data is limited, an important practical consideration (Supplementary Fig. 1). Compared to existing methods of selection inference (18–23), MPL was the most accurate method in terms of both classification accuracy, measured by AUROC for classifying mutant alleles as beneficial or deleterious, and in the absolute error in inferred selection coefficients (Supplementary Fig. 2, Supplementary Information) across a range of simulated data sets. Due to the simplicity of equation (1), MPL was fastest among the methods that we compared, with a running time roughly 6 orders of magnitude faster than approaches that rely on iterative Monte Carlo methods.

Next we applied MPL to study the intrahost evolution of HIV-1 and to resolve interactions between HIV-1 and the immune system. Identifying selective pressures on HIV-1 gives insight into the evolutionary dynamics leading to HIV-1 escape from immune control and the development of bnAbs, both of which are relevant for vaccine design.

We first examined longitudinal HIV-1 half-genome sequence data from 13 individuals where early-phase CD8⁺ T cell responses were comprehensively analyzed (25). In this group, 36.6% of the top 1% most beneficial mutations reported by MPL are nonsynonymous mutations in identified (25) CD8⁺ T cell epitopes (Fig. 2A). This is a 19-fold enrichment in mutations in T cell epitopes compared to expectations by chance (Methods). Reversions to clade consensus

are also strongly beneficial. Nonsynonymous reversions outside of T cell epitopes are 14-fold enriched in this subset. Furthermore, nonsynonymous reversions within T cell epitopes are 326-fold enriched. These findings are compatible with past studies that have observed strong selection for T cell escape (8, 9, 26) and for reversions (9).

Resolving linkage leads to substantial differences in selection estimates. MPL places 1.76 times as many T cell escape mutations within the top 1% most beneficial mutations as an independent model that ignores linkage between mutant alleles (Fig. 2B). Conversely, MPL ranks 0.42 times as many nonsynonymous reversions outside of T cell epitopes to be strongly beneficial as the independent model does (Fig. 2C). These differences are explained by the collective resolution of genetic linkage effects, including clonal interference.

In order to dissect the contributions of linkage to estimates of fitness, we computed the pairwise effects $\Delta \hat{s}_{ij}$ of each variant i on the inferred selection coefficients for all other variants j (Methods). We defined $\Delta \hat{s}_{ij}$ as the difference between the estimated selection coefficient $\hat{s}_j$ for variant j using all of the data and the value of $\hat{s}_j$ when variant i is replaced by the transmitted/founder (TF) nucleotide at the same site, thereby removing the contribution to selection from linkage with variant i . Positive values of $\Delta \hat{s}_{ij}$ indicate that linkage with variant i increases the selection coefficient inferred for variant j (e.g., due to clonal interference between them). Negative values indicate that variant i decreases the selection coefficient inferred for variant j (e.g., due to hitchhiking).

Our analysis revealed that the vast majority of observed variants have essentially no effect on estimates of selection at other sites, but a small minority of highly influential variants have dramatic effects (Supplementary Figs. 3–4). Such highly influential variants are often ones that change rapidly in frequency, sweeping through the population and exerting strong effects on linked sites (Supplementary Fig. 5). Consistent with this observation, 40% of highly influential variants are putative CD8⁺ T cell escape mutations. Effects on estimated selection drop off sharply with increasing distance along the genome for most variants (Supplementary Fig. 6). However, the effects of highly influential variants routinely span across long genomic distances (Supplementary Fig. 4). Collectively,

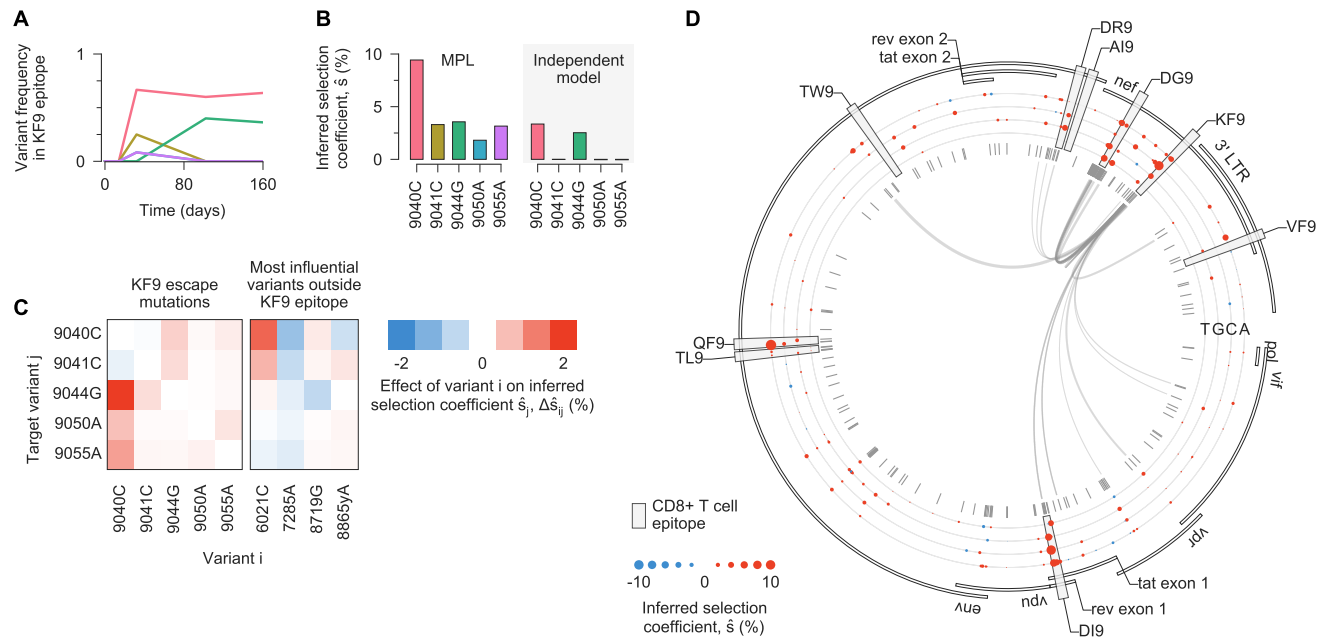


Fig. 3. Estimates of selection coefficients for viral escape mutations must account for clonal interference. **A**, Multiple escape mutations appear in the viral population in the T cell epitope KF9, targeted by individual CH77, exhibiting clonal interference. **B**, Using the full half-genome-length sequence data as input, MPL infers that all KF9 escape variants are positively selected. In contrast, estimates based solely on the trajectories of individual variants only uncover substantial positive selection for the 9040C and 9044G variants that coexist at the final time point. Furthermore, the independent model estimates of selection are attenuated because of the failure to account for competition with other beneficial mutations, including other escape mutations within the same epitope. **C**, Linkage effects on inferred selection coefficients for KF9 escape mutations. Effects shown here are due to variants within the KF9 epitope and the top four most influential variants outside the KF9 epitope, defined as the variants i for which $\sum_j |\Delta \hat{s}_{ij}|$ is the largest. All of these influential variants lie within other T cell epitopes (6021C lies in DI9, 7285A in QF9, 8719G in DR9, and 8865yA in DG9). **D**, Inferred selection in the HIV-1 half-genome sequence for CH77. Inferred selection coefficients are plotted in tracks. Coefficients of TF nucleotides are normalized to zero. Tick marks denote polymorphic sites. Inner links, shown for sites connected to the KF9 epitope, have widths proportional to matrix elements of the inverse of the integrated covariance (see Eq. (1)).

this data suggests that some highly influential variants are drivers of selective sweeps. Competition between such variants results in clonal interference.

A clear example of clonal interference is demonstrated in escape from a T cell response in individual CH77 targeting the Nef KF9 epitope (Fig. 3A). MPL infers strong positive selection for all escape variants. In contrast, when linkage is ignored escape variants that are lost are inferred to be neutral, and the magnitude of selection for 9040C is substantially decreased (Fig. 3B). Ignoring linkage thus leads to selection estimates that are qualitatively and quantitatively suspect. We observe similar instances of clonal interference in other epitopes (Supplementary Figs. 7-8). In the case of KF9, competition between the different escape variants increases the estimated selection coefficient for each of them (Fig. 3C). Inferred selection is also influenced by linkage with other mutations outside the KF9 epitope (Fig. 3C,D). For example, 9040C is inferred to be more beneficial due to its competition with the DI9 escape mutation 6021C. The selection coefficient for 9044G, in turn, is somewhat reduced due to positive linkage with 8719G, which is the dominant escape mutation in the nearby Env DR9 epitope.

Some strongly selected mutations lie in regions of Env that are exposed to antibodies, or in N-linked glycosylation motifs that affect the area of Env that is accessible to antibodies (Fig. 2A). However, these mutations are infrequent compared to others in T cell epitopes. One can also observe little posi-

tive selection in Env outside of T cell epitopes in the example of CH77 (Fig. 3D). Overall, selection for escape from antibody responses appears to be weaker or less frequent than CD8⁺ T cell-mediated selection.

We then asked whether strong antibody-mediated selection would be observed in individuals who generate bnAbs. To explore this question we studied HIV-1 evolution in individual CAP256, who eventually developed the VRC26 family of bnAbs (27, 28). This case is particularly challenging for inference because of a superinfection event 15 weeks after initial infection (Fig. 4A). This leads to strong and complex patterns of linkage as the superinfecting strain recombines and competes with the primary infecting strain (Fig. 4B, Supplementary Fig. 9). For this reason, ignoring linkage leads to poor selection inferences. Most (6 of 11) of the top 1% most beneficial mutations inferred by the independent model are from the background of the superinfecting strain and are synonymous. In contrast, none of the most beneficial mutations inferred by MPL are synonymous.

We found that selection for known VRC26 resistance mutations (27, 28) is modest (Fig. 4C). The most strongly selected mutation in the VRC26 epitope region is 6709C ($\hat{s} = 0.041$) in codon 162 in Env, a variant present in the superinfecting strain that completes an N-linked glycosylation motif which is absent from the primary infecting virus. However, this modification makes the virus more sensitive to VRC26 (27, 28). We observe selection against 6717T

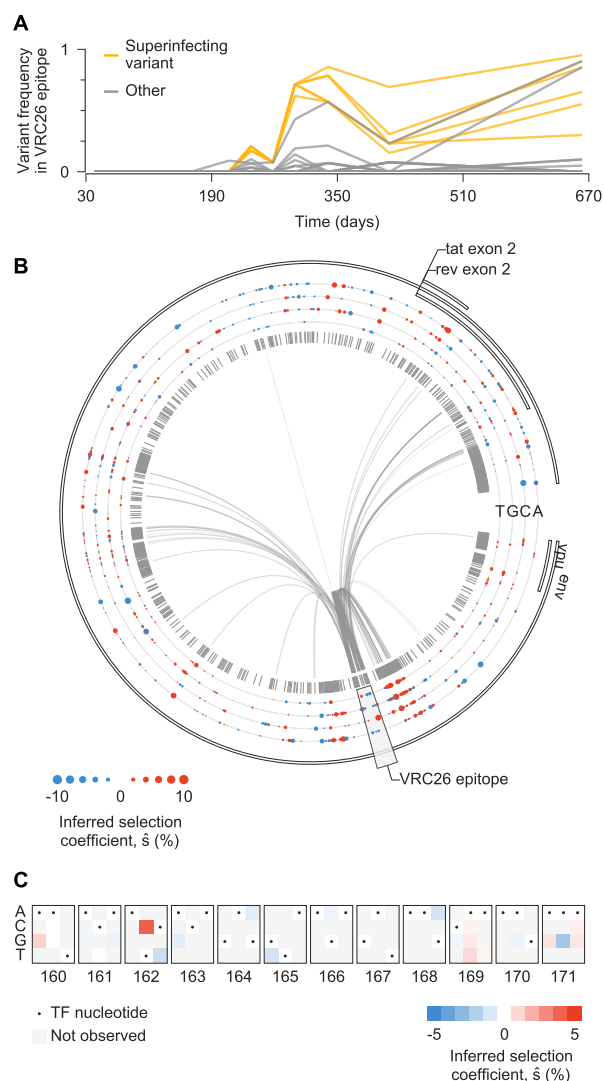


Fig. 4. Complex patterns of selection in HIV-1 Env following superinfection in an individual who develops broadly neutralizing antibodies. **A**, Multiple variants, including several from the superinfecting strain of the virus, rise and fall in frequency within the epitope targeted by the VRC26 family of antibodies. **B**, Inferred selection in CAP256 HIV-1 half-genome sequences. Inferred selection coefficients are plotted in tracks. Coefficients of TF nucleotides are normalized to zero. Tick marks denote polymorphic sites. Inner links, shown for sites connected to the VRC26 epitope, have widths proportional to matrix elements of the inverse of the integrated covariance. Linkage is extensive due to the struggle for dominance in the viral population between the TF, superinfecting, and recombinant strains. **C**, Map of inferred selection within the VRC26 epitope, consisting of codons 160–171 in Env.

($\hat{s} = -0.012$), corresponding to the Env 165L variant in the superinfecting strain. Reversion of this residue to V, the variant in the primary infecting strain, improves resistance to early VRC26 antibodies (28). We also observe modest positive selection for nonsynonymous variation at codon 169 in Env (maximum $\hat{s} = 0.010$), where mutations lead to complete resistance to VRC26 family antibodies (28). Thus, even the most strongly selected resistance mutations would fall outside of the top 5% most strongly selected mutations in the larger sample of 13 individuals.

Weak selection on the virus for antibody escape may in fact facilitate the development of bnAbs. Multiple escape vari-

ants, as well as variants that are sensitive to the antibody, can readily coexist for long times when escape is weakly selected. This coexistence increases the diversity of the viral population. Pressure on antibodies to bind to multiple variants can then select for breadth (29). Indeed, viral diversification has been observed to precede bnAb development (28, 30). In contrast, stronger pressure on the virus for escape could reduce viral diversity due to rapid fixation of beneficial escape variants and the elimination of sensitive ones.

Overall, our results reveal patterns of HIV-1 adaptation, including selection on the virus population accompanying the development of bnAbs, that were not possible to quantify using methods that cannot resolve genetic linkage. Furthermore, the scale of the data considered here is far beyond what existing methods that attempt to account for linkage can analyze. We anticipate that our method can also be widely applied to investigate selection in other evolving populations. Given the potential pitfalls of linkage-naïve inference, our results call for a greater focus on resolving linkage effects in studies of selection.

ACKNOWLEDGEMENTS

We thank A. K. Chakraborty, C. J. R. Illingworth, and J. G. Schraiber for helpful discussions and comments on the manuscript. The work of M.S.S., R.H.Y.L., and M.R.M. was supported by the Hong Kong Research Grants Council under grant number 16234716. R.H.Y.L. was also supported by Australia's National Health and Medical Research Council under grant number APP1121643. J.P.B. acknowledges support from a Regents Faculty Fellowship provided by the UCR Academic Senate.

AUTHOR CONTRIBUTIONS

All authors designed research, developed methods, analyzed data, interpreted results and wrote the paper.

References

1. Bignell, G. R. *et al.* Signatures of mutation and selection in the cancer genome. *Nature* **463**, 893–898 (2010).
2. Greaves, M. & Maley, C. C. Clonal evolution in cancer. *Nature* **481**, 306 (2012).
3. Burrell, R. A., McGranahan, N., Bartek, J. & Swanton, C. The causes and consequences of genetic heterogeneity in cancer evolution. *Nature* **501**, 338 (2013).
4. Nik-Zainal, S. *et al.* The life history of 21 breast cancers. *Cell* **149**, 994–1007 (2012).
5. Landau, D. A. *et al.* Evolution and impact of subclonal mutations in chronic lymphocytic leukemia. *Cell* **152**, 714–726 (2013).
6. Łuksza, M. *et al.* A neoantigen fitness model predicts tumour response to checkpoint blockade immunotherapy. *Nature* **551**, 517 (2017).
7. McMichael, A. J., Borrow, P., Tomaras, G. D., Goonetilleke, N. & Haynes, B. F. The immune response during acute HIV-1 infection: clues for vaccine development. *Nature Reviews Immunology* **10**, 11 (2010).
8. Allen, T. M. *et al.* Selective escape from CD8⁺ T-cell responses represents a major driving force of human immunodeficiency virus type 1 (HIV-1) sequence diversity and reveals constraints on HIV-1 evolution. *Journal of Virology* **79**, 13239–13249 (2005).
9. Zanini, F. *et al.* Population genomics of inpatient HIV-1 evolution. *eLife* **4**, e11282 (2015).
10. Strelkova, N. & Lässig, M. Clonal interference in the evolution of influenza. *Genetics* **192**, 671–682 (2012).
11. Łuksza, M. & Lässig, M. A predictive fitness model for influenza. *Nature* **507**, 57 (2014).
12. Muller, H. J. The relation of recombination to mutational advance. *Mutation Research/Fundamental and Molecular Mechanisms of Mutagenesis* **1**, 2–9 (1964).
13. Smith, J. M. & Haigh, J. The hitch-hiking effect of a favourable gene. *Genetics Research* **23**, 23–35 (1974).
14. Hegreness, M., Shores, N., Hartl, D. & Kishony, R. An equivalence principle for the incorporation of favorable mutations in asexual populations. *Science* **311**, 1615–1617 (2006).
15. Lang, G. I. *et al.* Pervasive genetic hitchhiking and clonal interference in forty evolving yeast populations. *Nature* **500**, 571 (2013).
16. Tenaillon, O. *et al.* Tempo and mode of genome evolution in a 50,000-generation experiment. *Nature* **536**, 165 (2016).
17. Levy, S. F. *et al.* Quantitative evolutionary dynamics using high-resolution lineage tracking. *Nature* **519**, 181 (2015).

18. Illingworth, C. J. & Mustonen, V. Distinguishing driver and passenger mutations in an evolutionary history categorized by interference. *Genetics* **189**, 989–1000 (2011).
19. Feder, A. F., Kryazhimskiy, S. & Plotkin, J. B. Identifying signatures of selection in genetic time series. *Genetics* **196**, 509–522 (2014).
20. Foll, M., Shim, H. & Jensen, J. D. WFABC: a Wright–Fisher ABC–based approach for inferring effective population sizes and selection coefficients from time-sampled data. *Molecular Ecology Resources* **15**, 87–98 (2015).
21. Ferrer-Admetlla, A., Leuenberger, C., Jensen, J. D. & Wegmann, D. An approximate Markov model for the Wright–Fisher diffusion and its application to time series data. *Genetics* **203**, 831–846 (2016).
22. Iranmehr, A., Akbari, A., Schlötterer, C. & Bafna, V. Clear: Composition of likelihoods for evolve and resequence experiments. *Genetics* **206**, 1011–1023 (2017).
23. Taus, T., Futschik, A. & Schlötterer, C. Quantifying selection with pool-seq time series data. *Molecular Biology and Evolution* **34**, 3023–3034 (2017).
24. Ewens, W. J. *Mathematical Population Genetics 1: Theoretical Introduction* (Springer Science & Business Media, 2012).
25. Liu, M. K. *et al.* Vertical T cell immunodominance and epitope entropy determine HIV-1 escape. *The Journal of Clinical Investigation* **123**, 380–393 (2013).
26. Liu, Y. *et al.* Selection on the human immunodeficiency virus type 1 proteome following primary infection. *Journal of Virology* **80**, 9519–9529 (2006).
27. Moore, P. L. *et al.* Multiple pathways of escape from HIV broadly cross-neutralizing V2-dependent antibodies. *Journal of Virology* **87**, 4882–4894 (2013).
28. Doria-Rose, N. A. *et al.* Developmental pathway for potent V1V2-directed HIV-neutralizing antibodies. *Nature* **509**, 55 (2014).
29. Wang, S. *et al.* Manipulating the selection forces during affinity maturation to generate cross-reactive HIV antibodies. *Cell* **160**, 785–797 (2015).
30. Liao, H.-X. *et al.* Co-evolution of a broadly neutralizing HIV-1 antibody and founder virus. *Nature* **496**, 469 (2013).

Methods

Data and code availability

All data and computer code supporting the findings of this study are available on GitHub in the repository <https://github.com/bartonlab/paper-MPL-inference>.

Evolutionary model

Our inference approach is based on the standard Wright-Fisher (WF) model of population genetics, which describes the stochastic dynamics of an evolving population of N individuals. Each individual is represented by a genetic sequence of length L . The population evolves in discrete, non-overlapping generations subject to the forces of selection, mutation, and recombination. For simplicity, we begin by describing the model with two alleles per locus, wild-type (WT) and mutant. Thus there are $M = 2^L$ unique genotypes. Later we show that our approach readily generalizes to consider multiple alleles per locus.

The state of the population at a generation t is given by the genotype frequency vector $\mathbf{z}(t) = (z_1(t), \dots, z_M(t))$, where $z_a(t)$ denotes the frequency of individuals with genotype a . Conditioned on $\mathbf{z}(t)$, the probability that the genotype frequency vector in the next generation is $\mathbf{z}(t+1)$ is multinomial (31):

$$P(\mathbf{z}(t+1) | \mathbf{z}(t)) = N! \prod_{a=1}^M \frac{(p_a(\mathbf{z}(t)))^{N z_a(t+1)}}{(N z_a(t+1))!}, \quad (2)$$

with

$$p_a(\mathbf{z}(t)) = \frac{y_a(t)f_a + \sum_{b \neq a} (\mu_{ba}y_b(t)f_b - \mu_{ab}y_a(t)f_a)}{\sum_{b=1}^M y_b(t)f_b}. \quad (3)$$

Here f_a denotes the fitness of genotype a , and μ_{ab} is the probability of genotype a mutating to genotype b . For simplicity we will assume at first that the mutation probability μ is the same at all loci, and that the probability of mutating from WT to mutant is the same as that from mutant to WT. In Eq. (3),

$$y_a(t) = (1-r)^{L-1} z_a(t) + \left(1 - (1-r)^{L-1}\right) \psi_a(t) \quad (4)$$

is the frequency of genotype a after recombination. Here r is the probability of recombination per locus per generation, and $\psi_a(t)$ is the probability that randomly recombining any two individuals in the population results in an individual of genotype a (see [Supplementary Information](#)). Though Eq. (3) and Eq. (4) appear complex, they have an intuitive interpretation. The first term in Eq. (3) reflects the fact that fitter individuals reproduce more efficiently and are therefore more likely to be observed in future generations. Mutations, captured through the second term, lead to conversions from genotype a to other genotypes, and vice versa. The denominator in Eq. (3) provides an overall normalization and indicates that relative fitness is important: in order for a particular

genotype to reliably grow in frequency, its fitness should be higher than the average fitness of all individuals in the population. The first term of Eq. (4) gives the proportion of individuals of genotype a not undergoing recombination, while the second term accounts for the net inward flow due to recombination from all other genotypes to genotype a .

For a population evolving under the WF model for T generations, the probability that the genotype frequency vector follows a particular evolutionary path $(\mathbf{z}(1), \mathbf{z}(2), \dots, \mathbf{z}(T))$, conditioned on the initial state $\mathbf{z}(0)$, is

$$P\left((\mathbf{z}(t))_{t=1}^T | \mathbf{z}(0)\right) = \prod_{t=0}^{T-1} P(\mathbf{z}(t+1) | \mathbf{z}(t)). \quad (5)$$

This expression is difficult to work with for parameter inference. This is due in part to the high dimensionality of the vector $\mathbf{z}$, which scales exponentially with the length of the genetic sequence. Thus, in the vast majority of real data sets, the sequence space is dramatically under-sampled. The functional form of Eq. (2) is also complex.

Our approach circumvents these issues by employing two approximations. First, we obtain a simplified version of Eq. (5) by using a path integral. Path integral expressions for evolutionary models have also been derived under different assumptions in past work (32–34), but they have not been widely applied for inference. We also assume that fitness is additive, such that the total fitness of each genotype a is just given by the sum of the selection coefficients s_i for mutant alleles at each locus i ,

$$f_a = 1 + \sum_{i=1}^L g_i^a s_i.$$

Here g_i^a is 1 if genotype a has a mutant allele at locus i and 0 otherwise. These assumptions will substantially simplify the expression for Eq. (5).

Path integral for mutant allele frequencies

In this section we will develop a simplified version of Eq. (5) defined at the level of allele frequencies rather than genotype frequencies. Later we will demonstrate that, if the fitness effects of mutations are additive as assumed above, this approach will lead to no loss of information for estimating the selection coefficients from data. We begin by using the WF dynamics above, which are defined for genotype frequencies, to compute the expected changes in frequency of mutant alleles. The mutant allele frequency x_i at locus i is

$$x_i(t) = \sum_{a=1}^M g_i^a z_a(t).$$

Following the assumptions above, and in the WF diffusion limit (31), one can show ([Supplementary Information](#)) that the probability density for mutant allele frequencies $\mathbf{x}(t) = (x_1(t), x_2(t), \dots, x_L(t))$ follows a Fokker-Planck equation

with a drift vector d_i and diffusion matrix C_{ij}/N given by

$$d_i(\mathbf{x}(t)) = \mu(1 - 2x_i(t)) + x_i(t)(1 - x_i(t))s_i \quad (6)$$

$$+ \sum_{j \neq i} (x_{ij}(t) - x_i(t)x_j(t))s_j$$

and

$$C_{ij}(\mathbf{x}(t)) = \begin{cases} x_i(t)(1 - x_i(t)) & i = j \\ x_{ij}(t) - x_i(t)x_j(t) & i \neq j. \end{cases} \quad (7)$$

Here x_{ij} is the frequency of individuals in the population with mutant alleles at both loci i and j . The drift vector describes the expected change in mutant allele frequencies in time. Note that the last term in Eq. (6) quantifies linked selection, i.e., how the dynamics of mutant allele frequencies are affected by the average genetic background on which they

appear. The drift vector should not be confused with *genetic drift*, the fluctuation in allele frequencies due to the inherent stochasticity of replication, which is instead described by the diffusion matrix. The diffusion matrix is simply the covariance matrix of mutant allele frequencies divided by the population size N . It therefore depends on the double mutant frequencies x_{ij} , but we will use the shortened notation $C_{ij}(\mathbf{x}(t))$ for brevity.

Applying standard methods from statistical physics (35), the Fokker-Planck equation can be converted into a path integral that quantifies the probability density for ‘paths’ of mutant allele frequencies $(\mathbf{x}(1), \mathbf{x}(2), \dots, \mathbf{x}(T))$. This expression will allow us to efficiently estimate the parameters that are most likely to have generated a specific path (see [Supplementary Information](#) for details). The probability for a path is

$$P\left((\mathbf{x}(t))_{t=1}^T | \mathbf{x}(0)\right) \approx \left[\prod_{t=0}^{T-1} \frac{1}{\sqrt{\det C(\mathbf{x}(t))}} \left(\frac{N}{2\pi}\right)^{L/2} \prod_{i=1}^L dx_i(t+1) \right] \exp\left(-\frac{N}{2} S\left((\mathbf{x}(t))_{t=0}^T\right)\right), \quad (8)$$

$$S\left((\mathbf{x}(t))_{t=0}^T\right) = \sum_{t=0}^{T-1} \sum_{i=1}^L \sum_{j=1}^L [x_i(t+1) - x_i(t) - d_i(\mathbf{x}(t))] (C^{-1}(\mathbf{x}(t)))_{ij} [x_j(t+1) - x_j(t) - d_j(\mathbf{x}(t))].$$

In the language of physics, $S\left((\mathbf{x}(t))_{t=0}^T\right)$ is referred to as the *action*. The population size N is analogous to the inverse temperature in statistical physics. The action penalizes deviation of the change in mutant frequencies between generations from the expectation given by the drift vector at the previous generation. This is normalized by the diffusion matrix, which quantifies the magnitude of typical changes in mutant frequencies due to random replication alone (i.e., genetic drift). The path integral equation in Eq. (8) follows the Itô convention.

Marginal path likelihood (MPL) estimate of the selection coefficients

Given an observed path of mutant allele frequencies, we can employ Bayesian inference to determine the maximum *a posteriori* selection coefficients $\hat{s}$ corresponding to the data, assuming that the population size N and mutation probability μ are known. In practice, our data consists of sets of genetic sequences obtained from a population at multiple time points $t_k, k \in \{0, 1, \dots, K\}$. These sequences can be used to compute the path of mutant allele frequencies $(\mathbf{x}(t_0), \mathbf{x}(t_1), \dots, \mathbf{x}(t_K))$ as well as the double mutant frequencies $x_{ij}(t_k)$, which also appear in Eq. (8). We will assume that the observed mutant allele frequencies are equal to the true ones, which simplifies the inference procedure. Our tests indicate that our results are robust to errors in the frequencies due to finite sampling (see [Supplementary Fig. 1](#)). Future work will relax this assumption.

In total, the posterior probability of the selection coefficients $\mathbf{s} = (s_1, s_2, \dots, s_L)$ given the observed path $(\mathbf{x}(t_0), \mathbf{x}(t_1), \dots, \mathbf{x}(t_K))$ is

$$P\left(\mathbf{s} | (\mathbf{x}(t_k))_{k=0}^K\right) = P\left((\mathbf{x}(t_k))_{k=0}^K | \mathbf{x}(0)\right) \times P_{\text{prior}}(\mathbf{s}), \quad (9)$$

where $P\left((\mathbf{x}(t_k))_{k=0}^K | \mathbf{x}(0)\right)$ is the probability of the path (given by Eq. (8), but extended to arbitrary sampling times as shown in [Supplementary Information](#)) and $P_{\text{prior}}(\mathbf{s})$ is the prior probability for the selection coefficients. Eq. (9) is a complicated expression of the mutant allele frequencies. However, solving for the selection coefficients that maximize the posterior probability is straightforward because Eq. (8) is a Gaussian function of the selection coefficients. Taking $P_{\text{prior}}(\mathbf{s})$ to be a Gaussian distribution with mean zero and covariance matrix $\sigma^2 I$, where I is the identity matrix, the selection coefficients that maximize Eq. (9) are given by

$$\hat{s}_i = \sum_{j=1}^L \left[\sum_{k=0}^{K-1} \Delta t_k C(\mathbf{x}(t_k)) + \gamma I \right]_{ij}^{-1} \times \left[x_j(t_K) - x_j(t_0) - \mu \sum_{k=0}^{K-1} \Delta t_k (1 - 2x_j(t_k)) \right]. \quad (10)$$

Here $\gamma = 1/N\sigma^2$, and $\Delta t_k = t_{k+1} - t_k$. We refer to Eq. (10) as the *MPL estimate* of selection coefficients. Because of the Gaussian form of Eq. (9), the maximum *a posteriori* estimates of the selection coefficients are the same as their pos-

terior means.

Eq. (10) can be readily interpreted. Let us start by considering the vector term in the “numerator” of Eq. (10) that multiplies the matrix inverse. Here the first terms quantify how the frequency of each mutant allele has changed between the initial and final generations. Naturally, alleles that increase in frequency over time are more likely to be beneficial. The remaining terms quantify the integrated mutational flux, i.e., population flow from mutant to WT (or vice versa) due to mutation. Net mutational flux from mutant to WT is also associated with higher fitness for the mutant allele. This is because this indicates that the mutant state maintained higher frequency than the WT over the trajectory, despite the force of mutation that drives the frequencies toward the same value. Together, these terms in the numerator of Eq. (10) determine whether a mutant allele is inferred to be beneficial or deleterious, at least when the off-diagonal elements of the matrix that it multiplies are zero.

While the numerator of Eq. (10) roughly determines the sign of selection, the denominator determines the strength of the inferred selection coefficient. Let us refer to $\sum_{k=0}^{K-1} \Delta t_k C(\mathbf{x}(t_k))$ as the integrated covariance matrix. From Eq. (7) we see that the entries of $C_{ij}(\mathbf{x}(t))$ are small when the mutant frequency is near the boundaries (0 or 1). Thus, the dominant contribution to the integrated covariance matrix comes from points on the path where the mutant frequency is far from the boundaries. If selection is strong, so that the mutant allele is much fitter than the WT (or vice versa), then we expect that a large portion of the path will be spent with the mutant allele frequency close to the boundary. In such cases the diagonal part of the integrated covariance will be small, and we correctly infer strong selection. The prior distribution for the selection coefficients simply adds a constant to the diagonal of the integrated covariance, which both shrinks selection estimates toward zero and ensures that the matrix is invertible. The off-diagonal terms of the integrated covariance matrix capture linkage effects, that is, how much of the change in the mutant frequency at a locus can be attributed to the average sequence background on which the mutant appears.

Equivalence of genotype- and allele-level analyses

In the preceding section we derived an estimate for the selection coefficients most likely to have generated an observed

evolutionary path. To do this we used an expression for the likelihood of a path of mutant allele frequencies that depended on the mutant frequencies $x_i(t)$ and their pairwise correlations $x_{ij}(t)$, but not on higher order correlations of the full genotype distribution. However, the WF dynamics is defined at the level of genotypes.

It can be shown that the use of Eq. (8) does not result in any loss in information beyond the approximations inherent in the WF diffusion limit. In the WF diffusion limit, the same steps as those applied to derive Eq. (8) can be performed for the genotype frequencies (see [Supplementary Information](#)). This results in a path integral expression that quantifies the probability density of genotype frequency paths. As in the allele-level analysis, the estimated selection coefficients are those that maximize

$$P(\mathbf{s} | (\mathbf{z}(t_k))_{k=0}^K) = P((\mathbf{z}(t_k))_{k=0}^K | \mathbf{z}(0)) \times P_{\text{prior}}(\mathbf{s}),$$

where $P((\mathbf{z}(t_k))_{k=0}^K | \mathbf{z}(0))$ is the probability density of the genotype frequency path. The full expression is more complicated, and less transparent, than the allele-level equivalent. Nonetheless, one can show that the expression for the selection coefficients that maximize the posterior probability above is exactly the same as Eq. (10). Full details of this derivation are given in the [Supplementary Information](#). This result is important because it shows that, following the assumptions of the WF diffusion limit and assuming that the fitness effects of mutations are additive, higher order mutational correlations contain *no further information* about the fitness effects of mutations.

Extension to multiple alleles per locus and asymmetric mutation probabilities

The MPL framework extends readily to models with ℓ alleles per locus, as well as asymmetric mutation probabilities. Let $x_{i,\alpha}(t)$ denote the frequency of allele α at locus i at generation t , and denote $\mu_{\alpha\beta}$ as the mutation probability per locus from allele α to allele β . Now, the trajectory of allele frequency vectors is $(\mathbf{x}(t_0), \mathbf{x}(t_1), \dots, \mathbf{x}(t_K))$, where $\mathbf{x}(t_k) = (x_{1,1}(t_k), x_{1,2}(t_k), \dots, x_{1,\ell}(t_k), x_{2,1}(t_k), \dots, x_{L,\ell}(t_k))$. Following parallel arguments to before (see [Supplementary Information](#)), the MPL estimate of the selection coefficient $\hat{s}_{i,\alpha}$ for allele α at locus i is

$$\hat{s}_{i,\alpha} = \sum_{j=1}^L \sum_{\beta=1}^{\ell} \left[\sum_{k=0}^{K-1} \Delta t_k C(\mathbf{x}(t_k)) + \gamma I \right]_{ij,\alpha\beta}^{-1} \times \left[x_{j,\beta}(t_K) - x_{j,\beta}(t_0) - \sum_{k=0}^{K-1} \Delta t_k \sum_{\delta=1}^{\ell} (\mu_{\delta\alpha} x_{j,\delta}(t_k) - \mu_{\alpha\delta} x_{j,\alpha}(t_k)) \right], \quad (11)$$

where $\gamma = 1/N\sigma^2$ as before. Off-diagonal entries of the covariance matrix $C(\mathbf{x}(t_k))$ are given by

$$C_{ij,\alpha\beta}(\mathbf{x}(t_k)) = x_{ij,\alpha\beta}(t_k) - x_{i,\alpha}(t_k)x_{j,\beta}(t_k),$$

where $x_{ij,\alpha\beta}(t_k)$ is the frequency of sequences with alleles α and β at loci i and j , respectively, at time t_k .

Data and code

Raw data and code used in our analysis is available in the GitHub repository located at <https://github.com/bartonlab/paper-MPL-inference>. This repository also contains Jupyter notebooks that can be run to reproduce the results presented here.

Simulation data

We implemented the Wright-Fisher model with discrete generations in Python. We used this program to record 100 evolutionary histories each for three different choices of the underlying parameters. Parameter values are detailed in Supplementary Figs. 1 and 2. For all simulations we assumed only two alleles per site. The simulation code, code for analysis, and original simulation data are contained in the GitHub repository.

Other time-series inference methods

The independent model that we compared MPL against in the main text is a single locus (SL) variant of MPL in which the off-diagonal elements of the integrated covariance matrix are set to zero.

The seven additional time series-based inference methods that we compared MPL against are FIT (36), LLS (37), CLEAR (38), EandR (39), ApproxWF (40), WFABC (41), and IM (42). Where available, we used the scripts provided by the authors. Some of these methods required preprocessing of the time-series data to obtain valid estimates of selection coefficients. See [Supplementary Information](#) for full details on implementation and data processing.

Patient cohort

We studied HIV-1 sequence data obtained from a 14 individuals recruited under the CHAVI 001 and CAPRISA 002 studies in the United States, Malawi, and South Africa. The locations of CD8⁺ T cell epitopes were experimentally (43) or computationally (44) determined in 13 of the 14 individuals. In the remaining individual, CAP256, experimental studies identified the VRC26 family of antibodies and mapped the epitope location on Env (45).

HIV-1 sequence data

Multiple sequence alignments of HIV-1 nucleotide sequences for all individuals were obtained from the Los Alamos National Laboratory HIV Sequence Database (www.hiv.lanl.gov; accessed October 19, 2018). Sequences labeled as problematic were not downloaded. For the 13 individuals with identified T cell epitopes, sets of 3' and 5' half-genome sequences were obtained, which were approximately 4500 bp in length. Only Env sequences were available for

CAP256 (approximately 2500 bp in length). All sequences were aligned with the HXB2 reference sequence (GenBank accession number K03455) for numbering, and with clade consensus sequences to determine reversions, using the Los Alamos National Laboratory HIVAalign tool (46).

In order to assure sequence quality, we 1) removed sequences with ≥ 200 gaps with respect to clade consensus, 2) removed sites with $>95\%$ gaps in the alignment, 3) imputed ambiguous nucleotides, and 4) removed time points where <4 sequences were sampled, or where the time gap between successive samples exceeded 300 days. During this process we also determined the number of reading frames in which each substitution was nonsynonymous, whether it occurred within an identified CD8⁺ T cell epitope that was actively targeted during the time for which sequence samples were available, whether it occurred within the exposed surface of Env (using surface residues as identified in ref. (47)), and whether it may have plausibly affected Env glycosylation by completing or disrupting an N-linked glycosylation motif. These analyses were performed using custom Python scripts available in the GitHub repository.

Variant indices were labeled relative to the standard HXB2 reference sequence of HIV-1. Insertions relative to HXB2 are labeled with lowercase alphabetical indices per standard conventions (48). For example, if three nucleotides were inserted relative to HXB2 after site 1, these would be labeled 1a, 1b, and 1c, respectively.

Enrichment analyses

We used fold enrichment values to determine the relative excess or lack of particular types of mutations among the HIV-1 variants that were inferred to be the most beneficial. For a given set of N_{sel} selected mutations (e.g., corresponding to the top 1% most beneficial), we computed the number n_{sel} of these mutations that have a particular property. This may represent, for example, the number of nonsynonymous mutations within identified CD8⁺ T cell epitopes, or the number of nonsynonymous reversions. The ratio $n_{\text{sel}}/N_{\text{sel}}$ then represents the fraction of the selected mutations having the specified property. This number was compared with $n_{\text{null}}/N_{\text{null}}$, where N_{null} is the total number of non-TF variants across all individuals and sequencing regions of the HIV-1 genome, and n_{null} is the number of these variants that share the specified property.

The fold enrichment of the selected set for a specified property is then naturally defined as $(n_{\text{sel}}/N_{\text{sel}})/(n_{\text{null}}/N_{\text{null}})$. A fold enrichment value greater than one indicates a larger proportion of mutants in the selected set that have the given property than expected by chance, while a value less than one indicates a smaller proportion than expected by chance.

Selection inference with MPL

We implemented the MPL method as described above in C++ and applied it to infer selection coefficients from the HIV-1 sequence data and from simulations. The original code used for inference is included in the GitHub repository. For the HIV-1 data, we assumed a regularization strength of $\gamma = 5$.

We also used mutation probabilities estimated in (49) as input. Mutation probabilities to and from gap states, representing deletions and insertions, respectively, were assumed to be very small ($\mu = 10^{-9}$). For the simulated data, we used a smaller regularization strength of $\gamma = 1$ due to the greater sampling depth. In order to increase robustness, we assumed that the true underlying allele frequency trajectories were piecewise linear and replaced the sums over time in Eq. (11) with integrals. Following the assumption of piecewise linearity, these integrals can be computed analytically. Specifically, the contribution of the mutational term to the numerator is then

$$-\sum_{k=0}^{K-1} \Delta t_k \sum_{\delta=1}^{\ell} \left(\frac{\mu_{\delta\alpha} x_{j,\delta}(t_k) + x_{j,\delta}(t_{k+1})}{2} - \mu_{\alpha\delta} \frac{x_{j,\alpha}(t_k) + x_{j,\alpha}(t_{k+1})}{2} \right),$$

the diagonal terms of the integrated covariance matrix are

$$\sum_{k=0}^{K-1} \Delta t_k \left(\frac{(3 - 2x_{i,\alpha}(t_{k+1}))(x_{i,\alpha}(t_k) + x_{i,\alpha}(t_{k+1}))}{6} - \frac{x_{i,\alpha}(t_k)x_{i,\alpha}(t_{k+1})}{3} \right),$$

and the off-diagonal terms of the integrated covariance matrix are

$$\sum_{k=0}^{K-1} \Delta t_k \frac{x_{ij,\alpha\beta}(t_k) + x_{ij,\alpha\beta}(t_{k+1})}{2} - \sum_{k=0}^{K-1} \Delta t_k \left(\frac{x_{i,\alpha}(t_k)x_{j,\beta}(t_k)}{3} + \frac{x_{i,\alpha}(t_{k+1})x_{j,\beta}(t_{k+1})}{3} + \frac{x_{i,\alpha}(t_k)x_{j,\beta}(t_{k+1})}{6} + \frac{x_{i,\alpha}(t_{k+1})x_{j,\beta}(t_k)}{6} \right).$$

After selection coefficients were inferred, we normalized them such that the transmitted/founder (HIV-1) or wild-type (simulation) allele had a selection coefficient of zero.

Calculation of effects of linkage on inferred selection

Due to the inverse of the integrated covariance matrix in Eq. (10), the selection coefficients estimated by MPL are affected by linkage. In order to quantify the effects of linkage on inferred selection during HIV-1 evolution, we computed the pairwise effects $\Delta\hat{s}_{ij}$ of each variant i on selection for other variants j , as described in the main text. Here, for ease of notation, each effective index i or j represents a single non-TF nucleotide at a particular site on the genome. That is, the indices incorporate both the label for the locus and for the allele.

To compute $\Delta\hat{s}_{ij}$, we iteratively select each nucleotide at each site, which together are represented by the index i , and generate a modified version of the sequence data in which variant i is replaced by the TF nucleotide at the same site. In

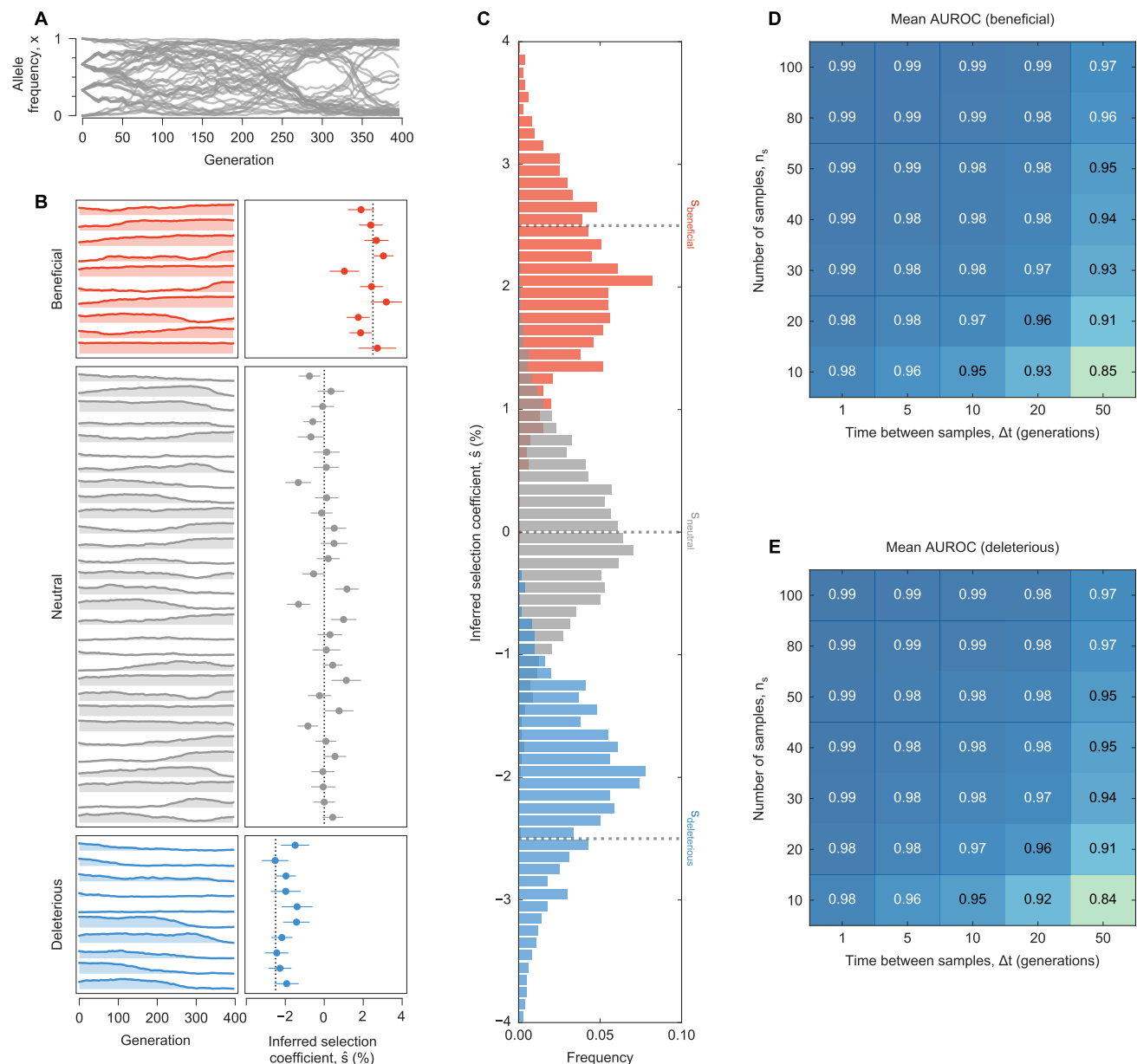
this way, linkage between the masked variant i and all other variants j is eliminated. We then infer the selection coefficients again for all variants j using the data where variant i has been replaced by TF, denoted as $\hat{s}_j^{\setminus i}$. Then we define

$$\Delta\hat{s}_{ij} = \hat{s}_j - \hat{s}_j^{\setminus i}.$$

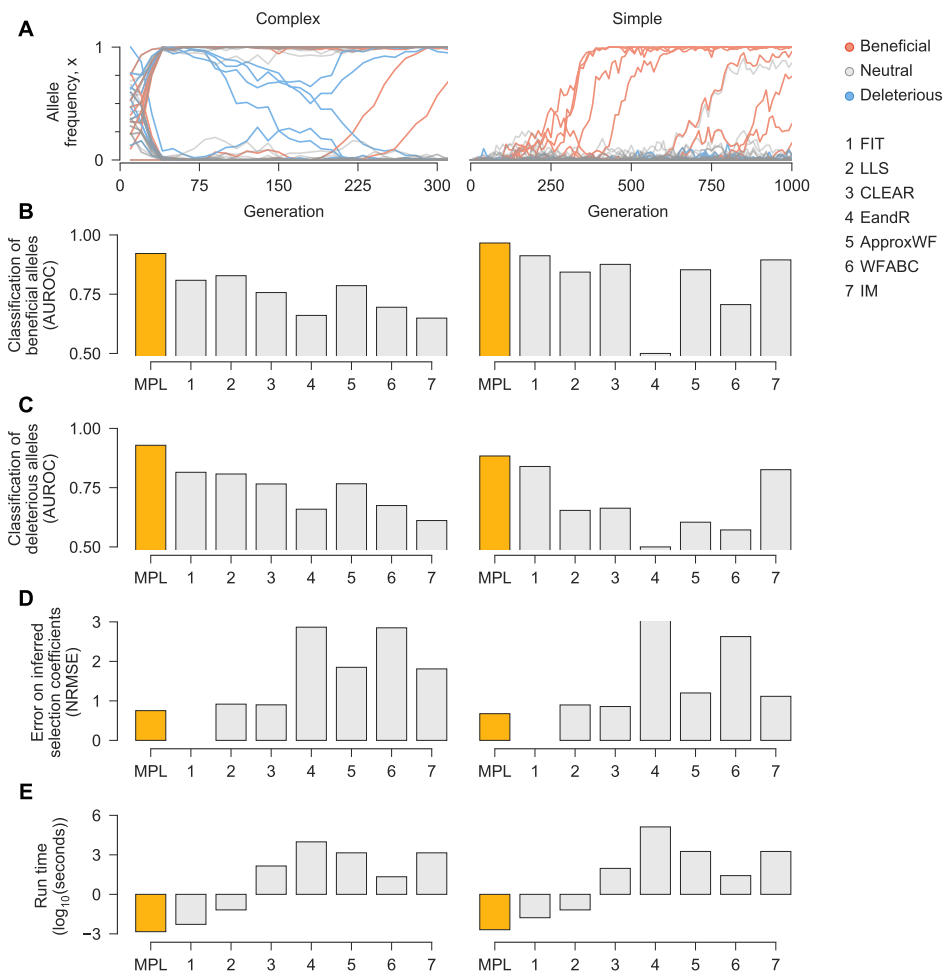
Positive values of $\Delta\hat{s}_{ij}$ thus indicate that linkage with variant i increases the selection coefficient inferred for variant j . This may be due, for example, to clonal interference between variants i and j . Negative values indicate that variant i decreases the selection coefficient inferred for variant j . This may occur, for example, if variant j hitchhikes on a beneficial genetic background that includes variant i .

References

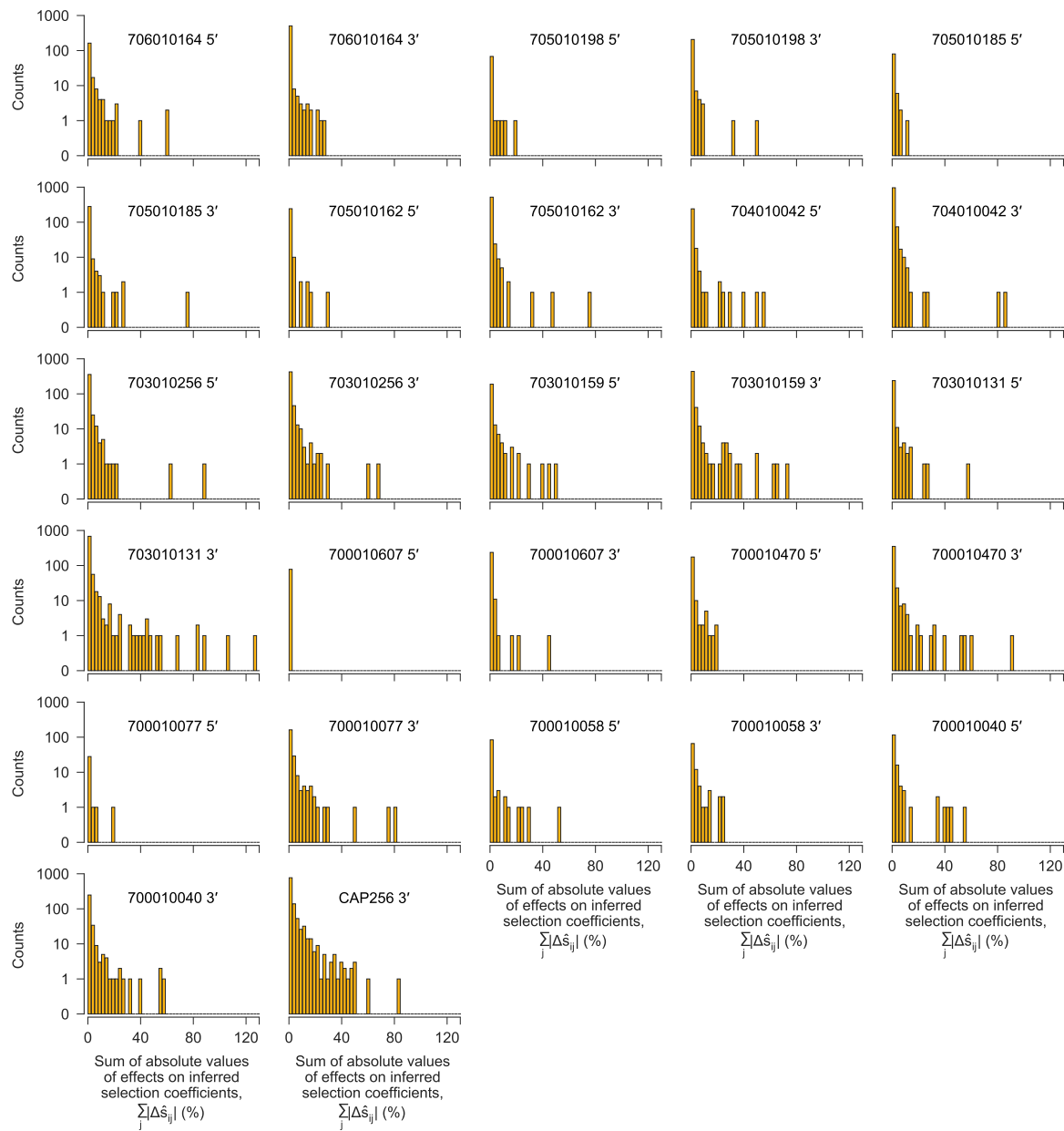
- Ewens, W. J. *Mathematical Population Genetics 1: Theoretical Introduction* (Springer Science & Business Media, 2012).
- Mustonen, V. & Lässig, M. Fitness flux and ubiquity of adaptive evolution. *Proceedings of the National Academy of Sciences* **107**, 4248–4253 (2010).
- Illingworth, C. J., Parts, L., Schiffels, S., Liti, G. & Mustonen, V. Quantifying selection acting on a complex trait using allele frequency time series data. *Molecular Biology and Evolution* **29**, 1187–1197 (2011).
- Schraiber, J. G. A path integral formulation of the Wright–Fisher process with genic selection. *Theoretical Population Biology* **92**, 30–35 (2014).
- Risken, H. *The Fokker-Planck Equation: Methods of Solution and Applications* (Springer-Verlag, 1989), 2nd edn.
- Feder, A. F., Kryazhimskiy, S. & Plotkin, J. B. Identifying signatures of selection in genetic time series. *Genetics* **196**, 509–522 (2014).
- Taus, T., Futschik, A. & Schlötterer, C. Quantifying selection with pool-seq time series data. *Molecular Biology and Evolution* **34**, 3023–3034 (2017).
- Iranmehr, A., Akbari, A., Schlötterer, C. & Bafna, V. Clear: Composition of likelihoods for evolve and resequence experiments. *Genetics* **206**, 1011–1023 (2017).
- Terhorst, J., Schlötterer, C. & Song, Y. S. Multi-locus analysis of genomic time series data from experimental evolution. *PLoS Genetics* **11**, e1005069 (2015).
- Ferrer-Admetlla, A., Leuenberger, C., Jensen, J. D. & Wegmann, D. An approximate Markov model for the Wright–Fisher diffusion and its application to time series data. *Genetics* **203**, 831–846 (2016).
- Foll, M., Shim, H. & Jensen, J. D. WFABC: a Wright–Fisher ABC-based approach for inferring effective population sizes and selection coefficients from time-sampled data. *Molecular Ecology Resources* **15**, 87–98 (2015).
- Illingworth, C. J. & Mustonen, V. Distinguishing driver and passenger mutations in an evolutionary history categorized by interference. *Genetics* **189**, 989–1000 (2011).
- Liu, M. K. *et al.* Vertical T cell immunodominance and epitope entropy determine HIV-1 escape. *The Journal of Clinical Investigation* **123**, 380–393 (2013).
- Barton, J. P. *et al.* Relative rate and location of intra-host HIV evolution to evade cellular immunity are predictable. *Nature Communications* **7**, 11660 (2016).
- Doria-Rose, N. A. *et al.* Developmental pathway for potent V1V2-directed HIV-neutralizing antibodies. *Nature* **509**, 55 (2014).
- Gaschen, B., Kuiken, C., Korber, B. & Foley, B. Retrieval and on-the-fly alignment of sequence fragments from the HIV database. *Bioinformatics* **17**, 415–418 (2001).
- Louie, R. H., Kaczorowski, K. J., Barton, J. P., Chakraborty, A. K. & McKay, M. R. Fitness landscape of the human immunodeficiency virus envelope protein that is targeted by antibodies. *Proceedings of the National Academy of Sciences* **201717765** (2018).
- Korber, B. *et al.* Numbering positions in HIV relative to HXB2CG. *Human Retroviruses and AIDS* **3**, 102–111 (1998).
- Zanini, F., Puller, V., Brodin, J., Albert, J. & Neher, R. A. In vivo mutation rates and the landscape of fitness costs of HIV-1. *Virus Evolution* **3** (2017).



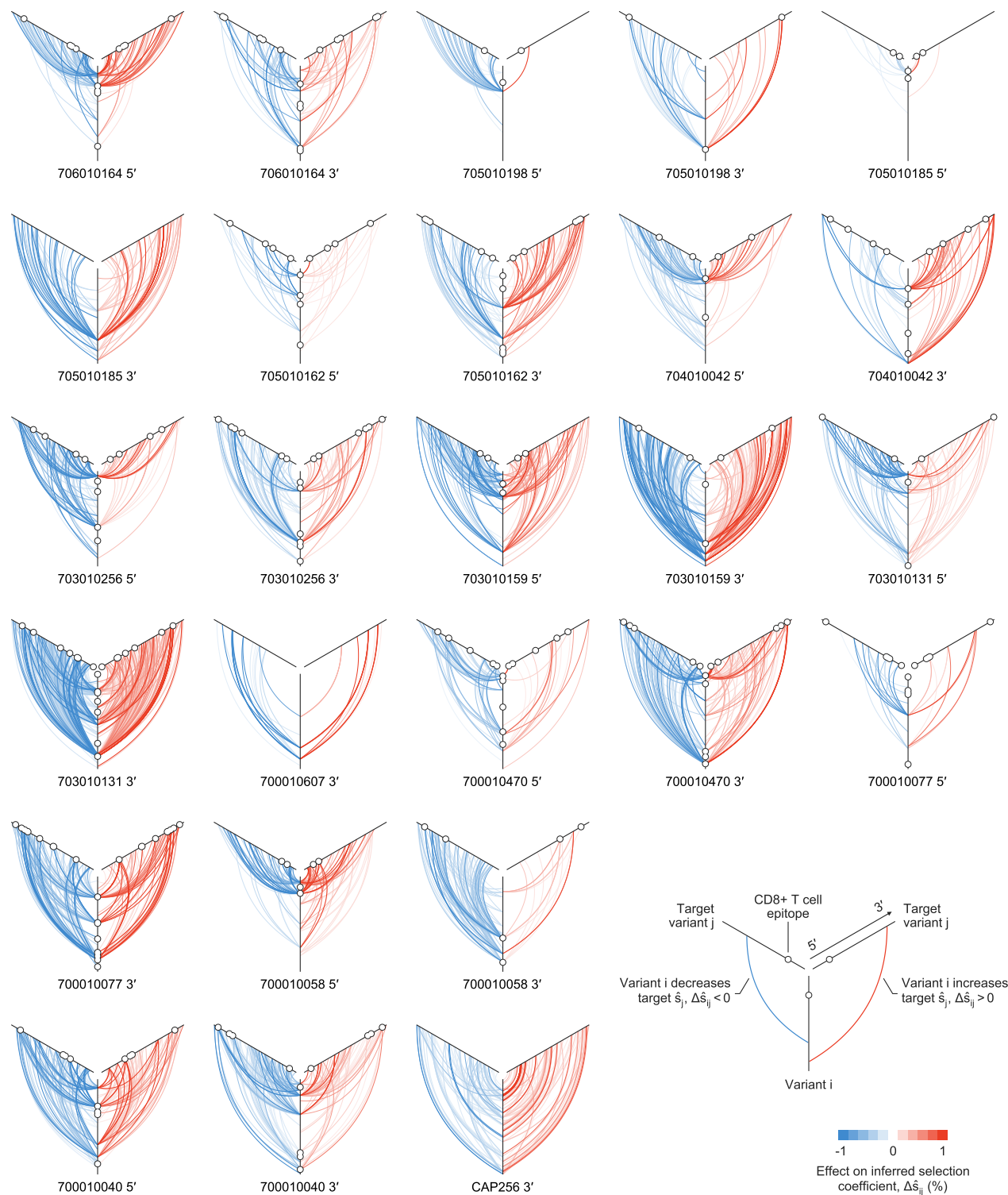
Supplementary Fig. 1. MPL accurately recovers selection coefficients from complex simulated evolutionary trajectories. **A**, Trajectories of mutant allele frequencies over time exhibit complex dynamics in a WF simulation with a simple fitness landscape. **B**, Separate views of individual trajectories for beneficial, neutral, and deleterious mutants (*left panel*) and inferred selection coefficients (*right panel*). Note that many neutral mutations exhibit temporal variation similar to beneficial or deleterious mutations. MPL estimates the underlying selection coefficients used to generate these trajectories and distinguishes between beneficial, neutral, and deleterious mutations, using equation (10). Dashed lines mark the true selection coefficients. **C**, Distributions of selection coefficient estimates across 100 replicate simulations with identical parameters in the special case of perfect sampling. MPL is also robust to finite sampling constraints, accurately classifying beneficial (**D**) and deleterious (**E**) mutants even when the number of sequences sampled per time point n_s is low, and the spacing between time samples Δt is large. *Simulation parameters.* $L = 50$ loci with two alleles at each locus (mutant and WT); 10 beneficial mutants with $s = 0.025$, 30 neutral mutants with $s = 0$, and 10 deleterious mutants with $s = -0.025$. Mutation probability $\mu = 10^{-3}$, population size $N = 10^3$. Initial population composed of approximately equal numbers of three random founder sequences, evolved over $T = 400$ generations.



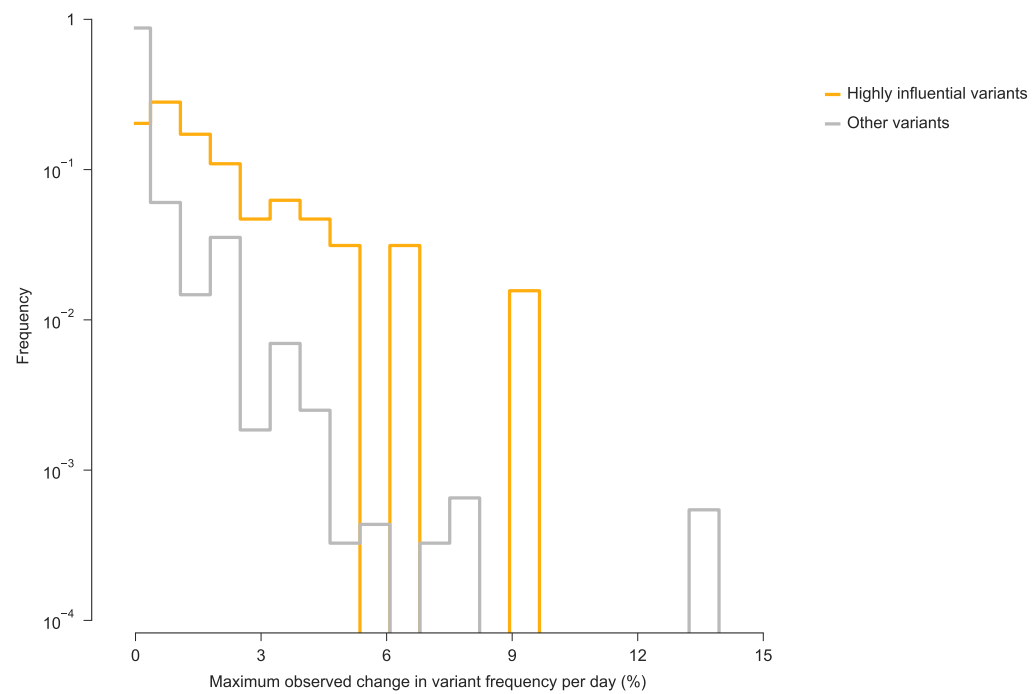
Supplementary Fig. 2. MPL compares favorably with state-of-the-art methods. We compared the ability of MPL and existing methods to infer selection from simulated test data that was rich with interference patterns and linkage, as shown in representative allele frequency trajectories (**A**). In order to evaluate robustness to finitely sampled data, we selected $n_s = 100$ sequences per time point for inference, with sampling time points separated by $\Delta t = 10$ generations. Performance was evaluated by comparing the successful classification of beneficial (**B**) and deleterious (**C**) mutations, error in the estimated selection coefficients (**D**), and run time (**E**), averaged over 100 replicate simulations with identical parameters. MPL achieves the highest performance in terms of classification and estimation accuracy, and in run time. Note that FIT does not specifically estimate selection coefficients. *Simulation parameters.* $L = 50$ loci with two alleles at each locus (mutant and WT): 10 beneficial mutants ($s = 0.1$ for Complex, $s = 0.025$ for Simple), 30 neutral mutants ($s = 0$ for both scenarios), and 10 deleterious mutants ($s = -0.1$ for Complex, $s = -0.025$ for Simple). Mutation probability $\mu = 10^{-4}$, population size $N = 10^3$. For the Complex case, the initial population is composed of equal numbers of five random founder sequences, evolved over $T = 310$ generations. Recorded trajectory used for inference begins at generation 10. For the Simple case, the initial population begins with all WT sequences, evolved over $T = 1000$ generations.



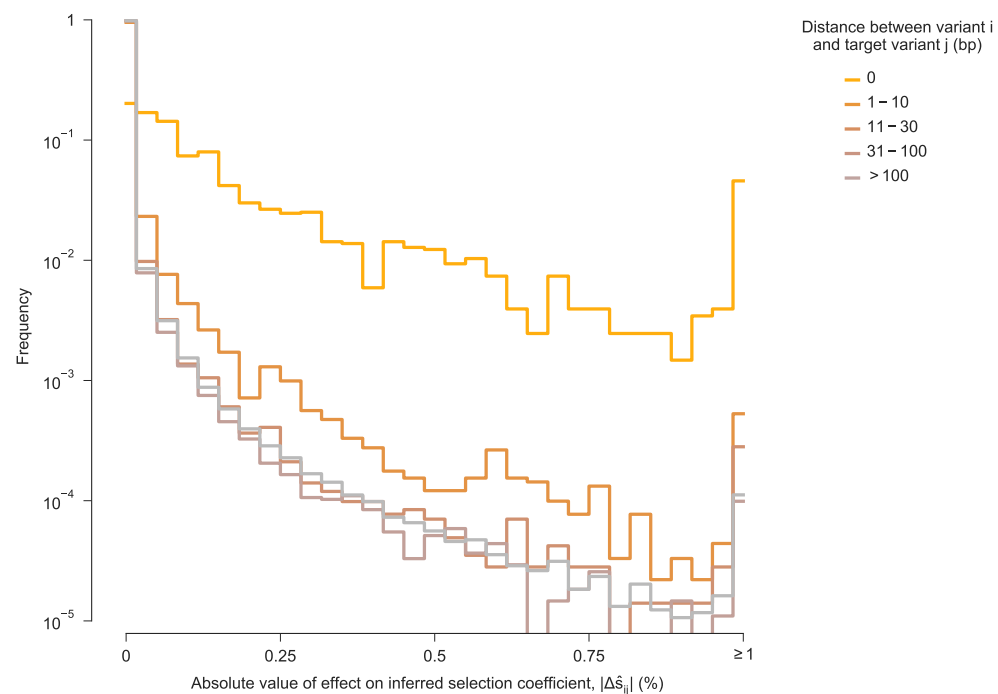
Supplementary Fig. 3. Most genetic variants have little effect on inferred selection at other sites, but a small minority have strong effects. After computing the pairwise effects $\Delta\hat{s}_{ij}$ of each variant i on the inferred selection coefficient for each other variant j , referred to as the target, we summed the absolute value of the $\Delta\hat{s}_{ij}$ values over all target variants j to quantify the influence of each variant i on selection at other sites. One histogram is shown for each sequencing region, for each individual. For the vast majority of variants, the total effect on selection at other sites is near zero. However, a small minority have strong effects. We defined a variant to be 'highly influential' if the sum of the absolute values of the $\Delta\hat{s}_{ij}$ over all targets j was larger than 0.4 (= 40%).



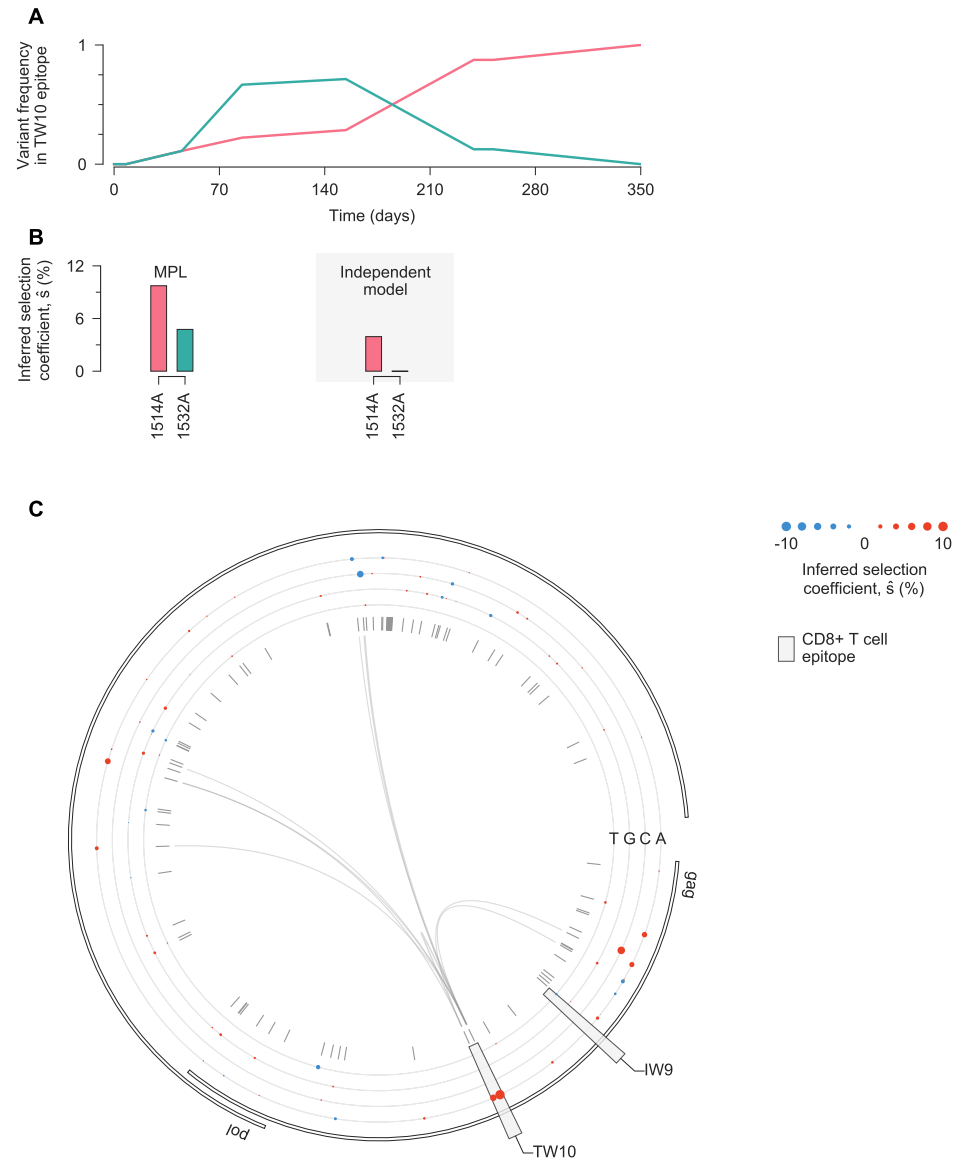
Supplementary Fig. 4. Variants that strongly influence inferred selection at other sites often act across large genomic distances. Plot of all linkage effects on inferred selection coefficients $\Delta\hat{s}_{ij}$ for which $|\Delta\hat{s}_{ij}| > 0.004$. One plot is shown for each sequencing region, for each individual. These strong effects of linkage on inferred selection coefficients can act at long range across the genome. Approximately 40% of highly influential variants, characterized by strong effects on inferred selection at other sites, lie within identified CD8⁺ T cell epitopes. The 5' region for individual 700010607 is not shown because no $\Delta\hat{s}_{ij}$ values are larger than the cutoff.



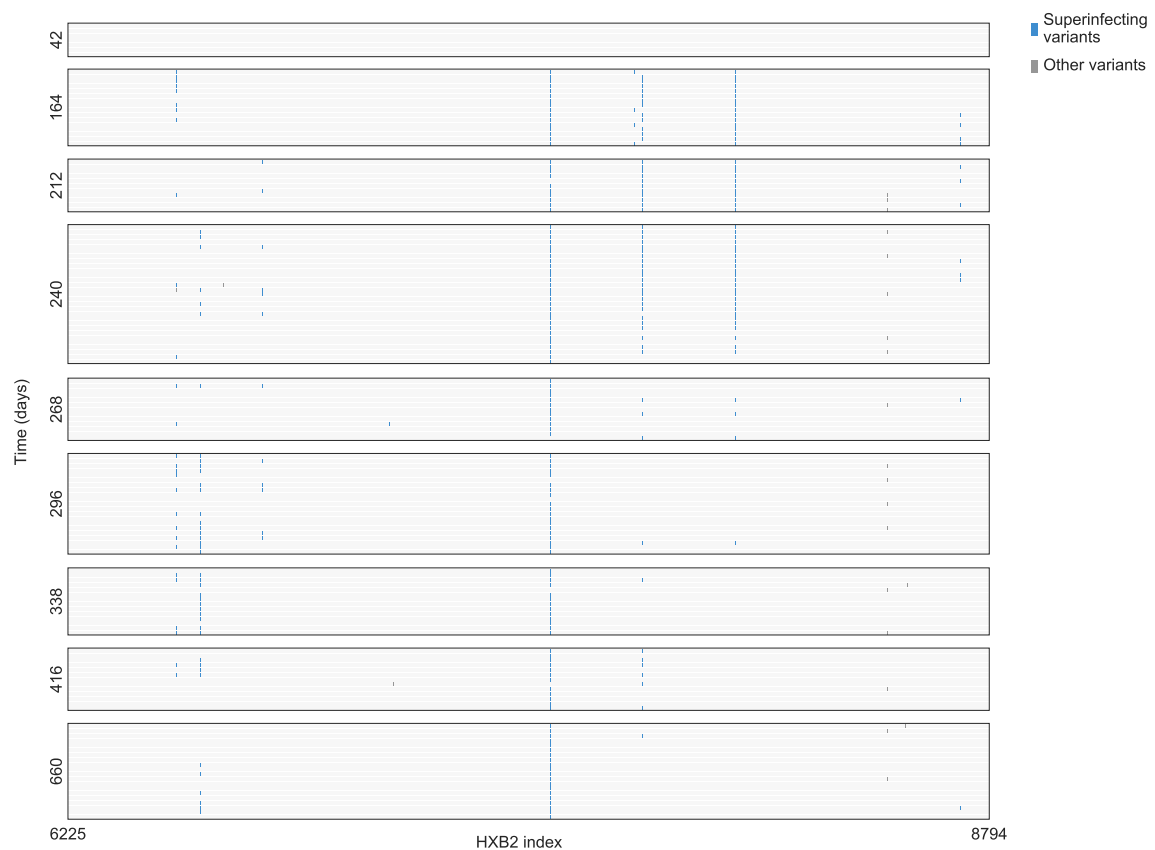
Supplementary Fig. 5. Highly influential variants are far more likely than other variants to change rapidly in frequency. For all genetic variants across all individuals and genomic regions considered in this study, we computed the maximum change in frequency per day between successive sequence samples. For most variants (those for which $\sum_j |\Delta \hat{s}_{ij}| \leq 0.4$), the maximum change in frequency per day is less than 1%. Highly influential variants, which have large effects on inferred selection coefficients at other sites ($\sum_j |\Delta \hat{s}_{ij}| > 0.4$), are much more likely than other variants to change rapidly in frequency.



Supplementary Fig. 6. For most variants, effects on inferred selection coefficients for other variants are stronger at smaller genomic distances. Histogram of the absolute value of linkage effects on inferred selection coefficients for other variants $|\Delta\hat{s}_{ij}|$, divided into subgroups based on the distance along the genome between variant i and target variant j . Consistent with intuition, the large effects on inferred selection coefficients occur most frequently for different variants that occur at the same site on the genome (i.e., distance equal to zero). "Interactions" between such variants are necessarily perfectly competitive, because only a single nucleotide is allowed at each position in the genetic sequence. For most variants, stronger linkage effects on inferred selection coefficients are more frequently observed for other variants within a distance of 10 base pairs (bp). Large linkage effects for pairs of variants within a distance of 30 bp, the approximate length of a linear T cell epitope, occur appreciably more frequently than for pairs of variants at greater genomic distances. However, there is little difference in the distribution of linkage effect sizes for pairs of variants that are between 31 bp and 100 bp apart compared to pairs of variants that are more than 100 bp apart. Nonetheless, some strong linkage effects on inferred selection are observed at long genomic distances (see Supplementary Fig. 4).



Supplementary Fig. 7. Estimates of selection coefficients in a simple example of clonal interference. **A**, Two escape mutations arise in the TW10 epitope targeted by individual CH58 and compete for dominance. **B**, MPL infers that both TW10 escape variants are positively selected. Estimates based on trajectories of individual variants only infer substantial positive selection for the 1514A variant that fixes. The magnitude of selection inferred with the independent model is also smaller than that inferred by MPL. **C**, Inferred selection in the HIV-1 5' half-genome sequence for CH58. Inferred selection coefficients are plotted in tracks. Coefficients of transmitted/founder nucleotides are normalized to zero. Tick marks denote polymorphic sites. Inner links, shown for sites connected to the TW10 epitope, have widths proportional to matrix elements of the inverse of the integrated covariance. Linked sites affect selection estimates within the epitope.



Supplementary Fig. 9. Extensive recombination between CAP256 primary and superinfecting strains. Individual CAP256 was infected by a distinct superinfecting strain of HIV-1 15 weeks after the primary infection (45). Each row represents a sequence, with nucleotide variants different from the primary infecting strain highlighted. Soon after superinfection, by 164 days after initial infection, recombinants between the primary and superinfecting strains dominate the viral population. Recombination continues throughout infection, introducing substantial variation into the VRC26 epitope region by 240 days after initial infection.

Supplementary Information

1. Derivation of the Wright-Fisher path integral

This section presents the steps used to derive the path integral approximation for the probability of the mutant allele frequencies following a path $(\mathbf{x}(1), \mathbf{x}(2), \dots, \mathbf{x}(T))$ conditioned on $\mathbf{x}(0)$, as described in Methods. While we recall that the result was presented for a biallelic model with symmetric mutation probabilities, the technique is by no means limited by this model and can be extended to incorporate other model complexities, some of which are considered subsequently.

To summarize Section 1, we first describe the WF model in Section 1.1 in more detail, followed by presenting conditional moment expansions of the mean frequency vector in Section 1.2. These expansions will then be used in Section 1.3 to derive (S15), the path integral expression for the probability of a path of *genotype* frequencies $(\mathbf{z}(1), \mathbf{z}(2), \dots, \mathbf{z}(T))$, conditioned on $\mathbf{z}(0)$. In Section 1.4, this genotype-level path integral will be used to derive (S22), the path integral expression for the probability of a path of *allele* frequencies $(\mathbf{x}(1), \mathbf{x}(2), \dots, \mathbf{x}(T))$, conditioned on $\mathbf{x}(0)$.

1.1. Wright Fisher model

We start with some brief comments on the WF model. At generation t , denote $\mathbf{Z}(t) = (Z_1(t), \dots, Z_M(t))$ as the random genotype frequency vector, and thus $\mathbf{z}(t) = (z_1(t), \dots, z_M(t))$ corresponds to an observed realization of this random vector. (These vectors, as well as all other vectors that we define, are taken as column vectors.) The stochastic dynamics of the genotype frequencies are governed by the transition probabilities reported in (2)-(4) of Methods. At generation $t+1$, $y_a(t)$ recombination occurs, and as a consequence of this action, the mean frequency of genotype a , conditioned on $\mathbf{Z}(t) = \mathbf{z}(t)$, is given by

$$y_a(t) = (1-r)^{L-1} z_a(t) + \left(1 - (1-r)^{L-1}\right) \psi_a(\mathbf{z}(t))$$

where the factor $(1-r)^{L-1}$ represents the probability of an individual not undergoing recombination, and $\psi_a(\mathbf{z}(t))$ the probability of forming genotype a after recombination from individuals in generation t . This latter quantity is given by

$$\psi_a(\mathbf{z}(t)) = \sum_{c=1}^M \sum_{d=1}^M R_{a,cd} z_c(t) z_d(t)$$

with $R_{a,cd}$ denoting the probability that genotypes c and d recombine to form genotype a . The probability $R_{a,cd}$ is a complicated function of the number of breakpoints and the particular genotypes a , c and d ; however as we will show, we do not require the exact form of $R_{a,cd}$ for our derivations. After recombination, the mean genotype frequencies at generation $t+1$ are further shaped through selection and mutation. Specifically, after recombination, selection and mutation, the mean frequency of genotype a , conditioned on $\mathbf{Z}(t) = \mathbf{z}(t)$, admits

$$\begin{aligned} p_a(\mathbf{z}(t)) &:= \mathbb{E} \left[Z_a(t+1) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] \\ &= \frac{y_a(t) f_a + \sum_{b=1, b \neq a}^M (\mu_{ba} y_b(t) f_b - \mu_{ab} y_a(t) f_a)}{\sum_{b=1}^M y_b(t) f_b}. \end{aligned}$$

Then, the WF dynamics, specified by the transition probability

$$P(\mathbf{z}(t+1) | \mathbf{z}(t)) = N! \prod_{a=1}^M \frac{(p_a(\mathbf{z}(t)))^{N z_a(t+1)}}{(N z_a(t+1))!},$$

simply results from performing random multinomial sampling (of N individuals) with these mean frequencies.

Further, we recall the assumption that $\mu_{ab} = \mu^{d_{ab}}$, with d_{ab} the Hamming distance between genotypes a and b , and the additive model for genotype fitness, such that the selection coefficient of genotype a admits

$$h_a = f_a - 1 = \sum_{j=1}^L g_j^a s_j. \quad (\text{S1})$$

We may then write

$$p_a(\mathbf{z}(t)) = \frac{(1+h_a) y_a(t) + \sum_{b=1}^M \mu^{d_{ab}} ((1+h_b) y_b(t) - (1+h_a) y_a(t))}{\sum_{b=1}^M (1+h_b) y_b(t)}. \quad (\text{S2})$$

We work under the assumption that the population size N is large, and that as $N \rightarrow \infty$,

$$s_i = \frac{\bar{s}_i}{N} + O\left(\frac{1}{N^2}\right), \quad \mu = \frac{\bar{\mu}}{N} + O\left(\frac{1}{N^2}\right), \quad r = \frac{\bar{r}}{N} + O\left(\frac{1}{N^2}\right), \quad (\text{S3})$$

and consequently

$$h_a = \frac{\bar{h}_a}{N} + O\left(\frac{1}{N^2}\right), \quad (\text{S4})$$

where $\bar{r}$, $\bar{h}_a$, $\bar{s}_i$ and $\bar{\mu}$ are constants that are independent of N .

1.2. Conditional moment expansions

The diffusion approximation, and the path integral expression, depends on the asymptotic properties of the mean frequency vector $p_a(\mathbf{z}(t))$, as well as those of higher order moments. For large N , under the parameter scalings above,

$$\begin{aligned} y_a(t) &= z_a(t) - r(L-1)(z_a(t) - \psi_a(\mathbf{z}(t))) + O\left(\frac{1}{N^2}\right) \\ &= z_a(t) + O\left(\frac{1}{N}\right) \end{aligned}$$

and thus (S2) can be expressed as

$$\begin{aligned} p_a(\mathbf{z}(t)) &= y_a(t) \left(1 + h_a - \sum_{b=1}^M h_b y_b(t)\right) + \mu \left(-Ly_a(t) + \sum_{b=1, d_{ab}=1}^M y_b(t)\right) + O\left(\frac{1}{N^2}\right) \\ &= z_a(t) \left(1 + h_a - \sum_{b=1}^M h_b z_b(t)\right) + \mu \left(-Lz_a(t) + \sum_{b=1, d_{ab}=1}^M z_b(t)\right) - r(L-1)(z_a(t) - \psi_a(\mathbf{z}(t))) + O\left(\frac{1}{N^2}\right) \\ &= z_a(t) + O\left(\frac{1}{N}\right). \end{aligned} \quad (\text{S5})$$

The conditional covariance of the genotype frequencies can also be expressed as

$$\text{Covar} \left(Z_a(t+1), Z_b(t+1) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right) = \begin{cases} \frac{p_a(\mathbf{z}(t))(1-p_a(\mathbf{z}(t)))}{N} = \frac{z_a(t)(1-z_a(t))}{N} + O\left(\frac{1}{N^2}\right) & a = b \\ -\frac{p_a(\mathbf{z}(t))p_b(\mathbf{z}(t))}{N} = -\frac{z_a(t)z_b(t)}{N} + O\left(\frac{1}{N^2}\right) & a \neq b \end{cases}. \quad (\text{S6})$$

These expansions will be used in the following subsection.

For the subsequent analysis with incomplete temporal sampling (Section 2), we will also require the following expansions:

$$\begin{aligned} \mathbb{E} \left[Z_a^2(t+1) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] &= z_a^2(t) + O\left(\frac{1}{N}\right) \\ \mathbb{E} \left[Z_a(t+1)Z_b(t+1) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] &= z_a(t)z_b(t) + O\left(\frac{1}{N}\right) \\ \text{Var} \left[Z_a(t+1)Z_b(t+1) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] &= O\left(\frac{1}{N}\right). \end{aligned} \quad (\text{S7})$$

These results are easily obtained by computing the moment generating function of $\mathbf{Z}(t+1)$ conditioned on $\mathbf{Z}(t) = \mathbf{z}(t)$,

$$\begin{aligned} M(\mathbf{w}, \mathbf{z}(t)) &= \mathbb{E} \left[\exp \left(\mathbf{w}^T \mathbf{Z}(t+1) \right) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] \\ &= \left(\sum_{k=1}^M p_k(\mathbf{z}(t)) \exp \left(\frac{w_k}{N} \right) \right)^N, \end{aligned}$$

where $\mathbf{w} = (w_1, \dots, w_M)$, and then evaluating relevant derivatives at zero in the standard way.

1.3. Genotype-level path integral

The moment expansions above give rise to a diffusion approximation and subsequently a path integral representation for the probability of genotype frequencies following a particular trajectory. The diffusion approximation is a continuous-time continuous-frequency approximation to the discrete-time discrete-frequency WF process. It is valid under the large- N parameter scalings (S3) and (S4), and corresponds to the continuous process

$$\check{\mathbf{Z}}(\tau) = (\check{Z}_1(\tau), \dots, \check{Z}_M(\tau)) := \mathbf{Z}(\lfloor N\tau \rfloor), \quad \tau \geq 0 \quad (\text{S8})$$

taken in the limit $N \rightarrow \infty$, where $\lfloor \cdot \rfloor$ denotes the floor function. Here τ is a continuous time variable with units of N generations, with one generation in discrete time (i.e., from t to $t+1$) thus taking

$$\delta\tau = \frac{1}{N} \quad (\text{S9})$$

continuous time units. The diffusion process is described by the probability density function ϕ , the solution to

$$\frac{\partial \phi}{\partial \tau} = \left[- \sum_{a=1}^M \frac{\partial}{\partial \check{z}_a} \bar{d}_a(\check{\mathbf{z}}(\tau)) + \sum_{a=1}^M \sum_{b=1}^M \frac{\partial}{\partial \check{z}_a} \frac{\partial}{\partial \check{z}_b} \bar{C}_{ab}(\check{\mathbf{z}}(\tau)) \right] \phi. \quad (\text{S10})$$

This is characterized by the drift vector $\bar{d}(\check{\mathbf{z}}(\tau))$, which describes the rate of expected changes in genotype frequencies at time τ , and the diffusion matrix $\bar{C}(\check{\mathbf{z}}(\tau))$, which describes the scaled covariance of the genotype frequency changes.

The drift vector has a th entry (see equation 4.99 of Risken(1))

$$\begin{aligned} \bar{d}_a(\check{\mathbf{z}}(\tau)) &= \lim_{\delta\tau \rightarrow 0} \frac{1}{\delta\tau} \mathbb{E} \left[\check{Z}_a(\tau + \delta\tau) - \check{Z}_a(\tau) \middle| \check{\mathbf{Z}}(\tau) = (\check{z}_1(\tau), \dots, \check{z}_M(\tau)) \right] \\ &= \lim_{N \rightarrow \infty} N \left(\check{z}_a(\tau) \left(h_a - \sum_{b=1}^M h_b \check{z}_b(\tau) \right) + \mu \left(-L \check{z}_a(\tau) + \sum_{b=1, d_{ab}=1}^M \check{z}_b(\tau) \right) - r(L-1)(\check{z}_a(\tau) - \psi_a(\check{\mathbf{z}}(\tau))) \right) \\ &= \check{z}_a(\tau) \left(\bar{h}_a - \sum_{b=1}^M \bar{h}_b \check{z}_b(\tau) \right) + \bar{\mu} \left(-L \check{z}_a(\tau) + \sum_{b=1, d_{ab}=1}^M \check{z}_b(\tau) \right) - \bar{r}(L-1)(\check{z}_a(\tau) - \psi_a(\check{\mathbf{z}}(\tau))) \end{aligned} \quad (\text{S11})$$

where the second line follows from (S9), (S8), and (S5). (As a technical point, we note that upon making the replacement $t = \lfloor N\tau \rfloor$, the variable t is considered fixed when taking limits over N . That is, we may write

$$\frac{1}{\delta\tau} \mathbb{E} \left[\check{Z}_a(\tau + \delta\tau) - \check{Z}_a(\tau) \middle| \check{\mathbf{Z}}(\tau) = (\check{z}_1(\tau), \dots, \check{z}_M(\tau)) \right] = N \mathbb{E} \left[Z_a(t+1) - Z_a(t) \middle| \mathbf{Z}(t) = (z_1(t), \dots, z_M(t)) \right]$$

and the limit of the left-hand side as $\delta\tau \rightarrow 0$ coincides with that of the right-hand side as $N \rightarrow \infty$. The same arguments will also be employed, although not explicitly stated, when taking limits in the subsequent derivations of drift vectors and diffusion matrices.) The diffusion matrix has (a, b) th entry (see equation 4.100 of Risken(1))

$$\begin{aligned} \bar{C}_{ab}(\check{\mathbf{z}}(\tau)) &= \frac{1}{2} \lim_{\delta\tau \rightarrow 0} \frac{1}{\delta\tau} \mathbb{E} \left[\left(\check{Z}_a(\tau + \delta\tau) - \check{Z}_a(\tau) \right) \left(\check{Z}_b(\tau + \delta\tau) - \check{Z}_b(\tau) \right) \middle| \check{\mathbf{Z}}(\tau) = (\check{z}_1(\tau), \dots, \check{z}_M(\tau)) \right] \\ &= \frac{1}{2} \lim_{\delta\tau \rightarrow 0} \frac{1}{\delta\tau} \text{Covar} \left(\check{Z}_a(\tau + \delta\tau), \check{Z}_b(\tau + \delta\tau) \middle| \check{\mathbf{Z}}(\tau) = (\check{z}_1(\tau), \dots, \check{z}_M(\tau)) \right) \\ &= \frac{1}{2} \begin{cases} \check{z}_a(\tau)(1 - \check{z}_a(\tau)) & a = b \\ -\check{z}_a(\tau)\check{z}_b(\tau) & a \neq b \end{cases} \end{aligned} \quad (\text{S12})$$

where the second line follows from noting that

$$\begin{aligned} &\text{Covar} \left(Z_a(t+1), Z_b(t+1) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right) \\ &= \mathbb{E} \left[\left(Z_a(t+1) - \mathbb{E} \left[Z_a(t+1) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] \right) \left(Z_b(t+1) - \mathbb{E} \left[Z_b(t+1) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] \right) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] \\ &= \mathbb{E} \left[\left(Z_a(t+1) - Z_a(t) + O\left(\frac{1}{N}\right) \right) \left(Z_b(t+1) - Z_b(t) + O\left(\frac{1}{N}\right) \right) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] \end{aligned} \quad (\text{S13})$$

along with (S8) and (S9), while the last line of (S12) follows from (S6).

The path integral (see equation 4.109 of Risken(1)) approximates the transition density of a diffusion process over a small time period. Specifically, applying this approach to the process described by (S10), for small $\delta\tau$ (equivalently large N), we obtain for the transition probability density over a single generation,

$$\phi(\tilde{z}(\tau + \delta\tau) | \tilde{z}(\tau)) \approx \frac{\exp\left(-\frac{1}{4\delta\tau}(\tilde{z}(\tau + \delta\tau) - \tilde{z}(\tau) - \bar{d}(\tilde{z}(\tau))\delta\tau)^T \bar{C}(\tilde{z}(\tau))^{-1}(\tilde{z}(\tau + \delta\tau) - \tilde{z}(\tau) - \bar{d}(\tilde{z}(\tau))\delta\tau)\right)}{(4\pi\delta\tau)^{M/2} \sqrt{\det(\bar{C}(\tilde{z}(\tau)))}}. \quad (\text{S14})$$

From this result, the transition probability for a single generation of the original discrete-time discrete-frequency WF process can (for large N) be approximated by

$$P(z(t+1) | z(t)) \approx \phi(z(t+1) | z(t)) dz(t+1)$$

where the $dz(t+1) = \prod_{a=1}^M dz_a(t+1)$ represent small frequency differences accounting for the quantization of the continuous genotype frequency space at each time point. The probability of observing a trajectory of genotype frequencies $(z(1), z(2), \dots, z(T))$ is then given by

$$\begin{aligned} P\left((z(t))_{t=1}^T | z(0)\right) &= \prod_{t=0}^{T-1} P(z(t+1) | z(t)) \\ &\approx \prod_{t=0}^{T-1} \left[\frac{1}{\sqrt{\det \bar{C}(z(t))}} \left(\frac{N}{4\pi}\right)^{M/2} dz(t+1) \right] \exp\left(-\frac{N}{4} S\left((z(t))_{t=0}^T\right)\right) \end{aligned} \quad (\text{S15})$$

where

$$S\left((z(t))_{t=0}^T\right) = \sum_{t=0}^{T-1} \sum_{a=1}^M \sum_{b=1}^M \left[z_a(t+1) - z_a(t) - \frac{\bar{d}_a(z(t))}{N} \right] (\bar{C}^{-1}(z(t)))_{ab} \left[z_b(t+1) - z_b(t) - \frac{\bar{d}_b(z(t))}{N} \right] \quad (\text{S16})$$

which is the desired path integral expression. In the language of physics, $S\left((z(t))_{t=0}^T\right)$ is referred to as the *action*.

1.4. Mutant allele-level path integral

Based on the genotype transition probability density (S14), we now provide an approximation for the mutant allele transition probability density. Let $x(t) = (x_1(t), \dots, x_L(t))$, for which

$$x_i(t) = \sum_{a=1}^M g_i^a z_a(t) \quad (\text{S17})$$

is the mutant frequency at locus i during generation t . Also, define the random mutant allele frequency vector $\mathbf{X}(t) = (X_1(t), \dots, X_L(t))$, which from (S17) is related to the random genotype frequency vector by

$$X_i(t) = \sum_{a=1}^M g_i^a Z_a(t). \quad (\text{S18})$$

The observed frequency vector $x(t)$ is thus a realization of this random vector. The continuous process which characterizes the mutant allele frequencies is

$$\check{\mathbf{X}}(\tau) = (\check{X}_1(\tau), \dots, \check{X}_L(\tau)) := \mathbf{X}(\lfloor N\tau \rfloor), \quad \tau \geq 0$$

taken as $N \rightarrow \infty$, which satisfies

$$\check{X}_i(\tau) = \sum_{a=1}^M g_i^a \check{Z}_a(\tau).$$

This allele frequency process is a diffusion process, whose probability density evolution is described by a solution to

$$\frac{\partial \phi}{\partial \tau} = \left[-\sum_{i=1}^L \frac{\partial}{\partial \check{x}_i} d_i(\check{x}(\tau)) + \sum_{i=1}^L \sum_{j=1}^L \frac{\partial}{\partial \check{x}_i} \frac{\partial}{\partial \check{x}_j} c_{ij}(\check{x}(\tau)) \right] \phi, \quad (\text{S19})$$

characterized by the drift vector $\underline{d}(\check{\mathbf{x}}(\tau))$ and diffusion matrix $\underline{C}(\check{\mathbf{x}}(\tau))$. Note that the drift and the diffusion equations given in the Methods are the same as here in the Supplementary Information, but those in the Methods use simpler notation by not explicitly distinguishing between continuous-time and discrete-time processes.

In the genotype case, the transition probability density was approximated to have a Gaussian form (S14). As the mutant allele frequencies are linear combinations of the genotype frequencies (S17), this implies that the transition probability density of mutant alleles also has a Gaussian form. Therefore, the allele-level drift and diffusion terms in (S19) are also a linear combination of the genotype drift and diffusion terms. The drift vector thus has i th entry

$$\begin{aligned} \underline{d}_i(\check{\mathbf{x}}(\tau)) &:= \sum_{a=1}^M g_i^a \bar{d}_a(\check{\mathbf{z}}(\tau)) \\ &= \sum_{a=1}^M g_i^a \left(\check{z}_a(\tau) \left(\bar{h}_a - \sum_{b=1}^M \bar{h}_b \check{z}_b(\tau) \right) + \bar{\mu} \left(-L \check{z}_a(\tau) + \sum_{b=1, d_{ab}=1}^M \check{z}_b(\tau) \right) - \bar{r}(L-1)(\check{z}_a(\tau) - \psi_a(\check{\mathbf{z}}(\tau))) \right) \\ &= \check{x}_i(\tau) (1 - \check{x}_i(\tau)) \bar{s}_i + \sum_{j=1, j \neq i}^L (\check{x}_{ij}(\tau) - \check{x}_i(\tau) \check{x}_j(\tau)) \bar{s}_j + \bar{\mu}(1 - 2\check{x}_i(\tau)) \end{aligned} \quad (\text{S20})$$

with

$$\check{x}_{ij}(\tau) := x_{ij}(\lfloor N\tau \rfloor), \quad \tau \geq 0,$$

and the last line follows by applying (S1)

$$\sum_{a=1}^M g_i^a g_j^a \check{z}_a(\tau) = \check{x}_{ij}(\tau),$$

and by noting that $\sum_{a=1}^M g_i^a (\check{z}_a(\tau) - \psi_a(\check{\mathbf{z}}(\tau))) = 0$. To see why the latter holds, first define

$$\theta_i^{cd} := \sum_{a=1}^M g_i^a R_{a,cd},$$

which is the probability that genotypes c and d recombine to form a genotype which has a mutation at locus i . Since the model is biallelic, the recombination event could involve allele pairs $(0,0)$, $(0,1)$, $(1,0)$, or $(1,1)$ at locus i of genotypes c and d . We thus have

$$\begin{aligned} \sum_{a=1}^M g_i^a \psi_a(\check{\mathbf{z}}(\tau)) &= \sum_{a=1}^M g_i^a \sum_{c=1}^M \sum_{d=1}^M R_{a,cd} \check{z}_c(\tau) \check{z}_d(\tau) \\ &= \sum_{c=1}^M \sum_{d=1}^M \theta_i^{cd} \check{z}_c(\tau) \check{z}_d(\tau) \\ &= \sum_{c=1}^M \left(\sum_{d=1}^M \theta_i^{cd} g_i^c g_i^d \check{z}_c(\tau) \check{z}_d(\tau) + \sum_{d=1}^M \theta_i^{cd} g_i^c (1 - g_i^d) \check{z}_c(\tau) \check{z}_d(\tau) \right) \\ &\quad + \sum_{c=1}^M \left(\sum_{d=1}^M \theta_i^{cd} (1 - g_i^c) g_i^d \check{z}_c(\tau) \check{z}_d(\tau) + \sum_{d=1}^M \theta_i^{cd} (1 - g_i^c) (1 - g_i^d) \check{z}_c(\tau) \check{z}_d(\tau) \right). \end{aligned}$$

Now by noting that

$$\begin{aligned} \theta_i^{cd} g_i^c g_i^d &= g_i^c g_i^d \\ \theta_i^{cd} g_i^c (1 - g_i^d) &= \frac{1}{2} g_i^c (1 - g_i^d) \\ \theta_i^{cd} (1 - g_i^c) g_i^d &= \frac{1}{2} (1 - g_i^c) g_i^d \\ \theta_i^{cd} (1 - g_i^c) (1 - g_i^d) &= 0 \end{aligned}$$

where the factor of $\frac{1}{2}$ arises because there is a 50% chance that genotype c (d) with a mutant at locus i and genotype d (c) with WT at locus i will recombine to a genotype with a mutant at locus i , we have

$$\begin{aligned}\sum_{a=1}^M g_i^a \psi_a(\tilde{\mathbf{z}}(\tau)) &= \sum_{c=1}^M \left(\sum_{d=1}^M g_i^c g_i^d \tilde{z}_c(\tau) \tilde{z}_d(\tau) + \frac{1}{2} \sum_{d=1}^M g_i^c (1 - g_i^d) \tilde{z}_c(\tau) \tilde{z}_d(\tau) \right) \\ &+ \frac{1}{2} \sum_{c=1}^M \sum_{d=1}^M (1 - g_i^c) g_i^d \tilde{z}_c(\tau) \tilde{z}_d(\tau) \\ &= \tilde{x}_i^2(\tau) + \frac{1}{2} \tilde{x}_i(\tau)(1 - x_i(\tau)) + \frac{1}{2} \tilde{x}_i(\tau)(1 - \tilde{x}_i(\tau)) \\ &= \tilde{x}_i(\tau),\end{aligned}$$

thus implying that $\sum_{a=1}^M g_i^a (\tilde{z}_a(\tau) - \psi_a(\tilde{\mathbf{z}}(\tau))) = 0$.

(Note that above and also in the subsequent derivations, since we are focused on the mutant allele frequency dynamics, we adopt notation which only explicitly demonstrates dependencies on the mutant allele frequencies. It should be recognized, however, that these dynamics also depend on the pairwise mutational frequencies. We do not explicitly show this, for the sake of notational convenience.) Considering now the diffusion matrix, this has (i, j) th entry

$$\begin{aligned}\underline{C}_{ij}(\tilde{\mathbf{x}}(\tau)) &:= \sum_{a=1}^M \sum_{b=1}^M g_i^a g_j^b \bar{C}_{ab}(\tilde{\mathbf{z}}(\tau)) \\ &= \frac{1}{2} \sum_{a=1}^M g_i^a g_j^a \tilde{z}_a(\tau)(1 - \tilde{z}_a(\tau)) - \frac{1}{2} \sum_{a=1}^M \sum_{b=1, b \neq a}^M g_i^a g_j^b \tilde{z}_a(\tau) \tilde{z}_b(\tau) \\ &= \frac{1}{2} \sum_{a=1}^M g_i^a g_j^a \tilde{z}_a(\tau) - \frac{1}{2} \left(\sum_{a=1}^M g_i^a \tilde{z}_a(\tau) \right) \left(\sum_{b=1}^M g_j^b \tilde{z}_b(\tau) \right) \\ &= \frac{1}{2} (\tilde{x}_{ij}(\tau) - \tilde{x}_i(\tau) \tilde{x}_j(\tau)).\end{aligned}$$

Note that if we concatenate the infinitesimal drift for all loci in vector form, it yields

$$\underline{d}(\tilde{\mathbf{x}}(\tau)) = 2\underline{C}(\tilde{\mathbf{x}}(\tau))\bar{\mathbf{s}} + \bar{\mu}(\mathbf{1} - 2\tilde{\mathbf{x}}(\tau)) \quad (\text{S21})$$

where $\mathbf{1}$ is the vector of ones.

As for the genotype case, equation 4.109 of Risken (1) may be applied, which approximates the transition probability density of this diffusion process over a single generation by a Gaussian distribution, given as

$$\phi(\tilde{\mathbf{x}}(\tau + \delta\tau) | \tilde{\mathbf{x}}(\tau)) \approx \frac{\exp\left(-\frac{1}{4\delta\tau} (\tilde{\mathbf{x}}(\tau + \delta\tau) - \tilde{\mathbf{x}}(\tau) - \underline{d}(\tilde{\mathbf{x}}(\tau))\delta\tau)^T [\underline{C}(\tilde{\mathbf{x}}(\tau))]^{-1} (\tilde{\mathbf{x}}(\tau + \delta\tau) - \tilde{\mathbf{x}}(\tau) - \underline{d}(\tilde{\mathbf{x}}(\tau))\delta\tau)\right)}{(4\pi\delta\tau)^{L/2} \sqrt{\det(\underline{C}(\tilde{\mathbf{x}}(\tau)))}}.$$

From this result, and recalling that $\delta\tau = 1/N$, the transition probability for a single generation of the original discrete-time discrete-frequency WF process can (for large N) be approximated by

$$\begin{aligned}P(\mathbf{x}(t+1) | \mathbf{x}(t)) &\approx \phi(\mathbf{x}(t+1) | \mathbf{x}(t)) d\mathbf{x}(t+1) \\ &= \frac{\left(\frac{N}{4\pi}\right)^{L/2} d\mathbf{x}(t+1)}{\sqrt{\det \underline{C}(\mathbf{x}(t))}} \exp\left(-\frac{N}{4} \sum_{i=1}^L \sum_{j=1}^L \left[x_i(t+1) - x_i(t) - \frac{d_i(\mathbf{x}(t))}{N} \right] (\underline{C}^{-1}(\mathbf{x}(t)))_{ij} \left[x_j(t+1) - x_j(t) - \frac{d_j(\mathbf{x}(t))}{N} \right] \right) \\ &= \frac{\left(\frac{N}{2\pi}\right)^{L/2} d\mathbf{x}(t+1)}{\sqrt{\det C(\mathbf{x}(t))}} \exp\left(-\frac{N}{2} \sum_{i=1}^L \sum_{j=1}^L \left[x_i(t+1) - x_i(t) - d_i(\mathbf{x}(t)) \right] (C^{-1}(\mathbf{x}(t)))_{ij} \left[x_j(t+1) - x_j(t) - d_j(\mathbf{x}(t)) \right] \right)\end{aligned}$$

where the $d\mathbf{x}(t+1) = \prod_{i=1}^L dx_i(t+1)$ represents small frequency differences accounting for the quantization of the continuous mutant allele frequency space, and we have made the replacements $\underline{d}_i(\mathbf{x}(t)) = Nd_i(\mathbf{x}(t))$ and $(\underline{C}(\mathbf{x}(t)))_{ij} = (1/2)C(\mathbf{x}(t))_{ij}$. The path integral expression then follows by noting that the probability of observing a trajectory of mutant allele frequencies

$(\mathbf{x}(1), \mathbf{x}(2), \dots, \mathbf{x}(T))$ is given by

$$P\left((\mathbf{x}(t))_{t=1}^T | \mathbf{x}(0)\right) = \prod_{t=0}^{T-1} P(\mathbf{x}(t+1) | \mathbf{x}(t)) \\ \approx \left(\prod_{t=0}^{T-1} \left[\frac{1}{\sqrt{\det C(\mathbf{x}(t))}} \left(\frac{N}{2\pi} \right)^{L/2} d\mathbf{x}(t+1) \right] \right) \exp\left(-\frac{N}{2} S\left((\mathbf{x}(t))_{t=0}^T\right)\right) \quad (\text{S22})$$

where

$$S\left((\mathbf{x}(t))_{t=0}^T\right) = \sum_{t=0}^{T-1} \sum_{i=1}^L \sum_{j=1}^L [x_i(t+1) - x_i(t) - d_i(\mathbf{x}(t))] (C^{-1}(\mathbf{x}(t)))_{ij} [x_j(t+1) - x_j(t) - d_j(\mathbf{x}(t))].$$

As in (S16), $S\left(\{\mathbf{x}(t)\}_{t=0}^T\right)$ is referred to as the action in physics.

2. Derivation of the MPL estimator of allele selection coefficients

This Section first presents a proof of the MPL estimate (10) of the selection coefficients. As stipulated in Methods, this is obtained as the MAP estimate of the selection coefficients, which is the solution to

$$\hat{\mathbf{s}} = \arg \max_{\mathbf{s}} \mathcal{L}\left(\mathbf{s} | N, \mu, (\mathbf{x}(t_k))_{k=0}^K\right) P_{\text{prior}}(\mathbf{s}), \quad (\text{S23})$$

where

$$\mathcal{L}\left(\mathbf{s} | N, \mu, (\mathbf{x}(t_k))_{k=0}^K\right) = P\left((\mathbf{x}(t_k))_{k=1}^K | \mathbf{x}(t_0), N, \mu, \mathbf{s}\right) \\ = \prod_{k=0}^{K-1} P(\mathbf{x}(t_{k+1}) | \mathbf{x}(t_k), N, \mu, \mathbf{s}) \quad (\text{S24})$$

is the likelihood function, while $(\mathbf{x}(t_0), \mathbf{x}(t_1), \dots, \mathbf{x}(t_K))$ is the observed trajectory of mutant allele frequencies.

The main challenge in solving (S23) is that it requires computing the likelihood (S24), which is complicated. This is simplified by adopting the path integral approach outlined in Section 1; but now extending the analysis to account for sampling times $t_0, \dots, t_K$ spanning multiple generations. Other than this difference, we may follow the same approach, starting by giving asymptotic moment expansions in Section 2.1. These expansions will then be used in Section 2.2 to derive (S32), the probability of a path of *genotype* frequencies $(\mathbf{z}(t_1), \mathbf{z}(t_2), \dots, \mathbf{z}(t_K))$, conditioned on $\mathbf{z}(t_0)$, and subsequently (S33), the probability of a path of *allele* frequencies $(\mathbf{x}(t_1), \mathbf{x}(t_2), \dots, \mathbf{x}(t_K))$, conditioned on $\mathbf{x}(t_0)$. This allele-level path integral will then be used to derive the MPL estimate (10) in Section 2.3. We then show in Section 2.4 that this MPL estimate can also be derived based on the genotype-level path integral, demonstrating no loss of optimality from derivations based on an allele-level representation. Finally, we extend the biallelic and symmetric mutation modelling assumptions to account for multiple alleles per locus and asymmetric mutation probabilities in Section 2.5. This will be required to analyze the HIV data in the main text.

2.1. Conditional moment expansions

The mean frequency of genotype a at generation $t + \Delta t$, conditioned on $\mathbf{Z}(t) = \mathbf{z}(t)$, admits

$$p_a(\mathbf{z}(t), \Delta t) := \mathbb{E}\left[Z_a(t + \Delta t) \middle| \mathbf{Z}(t) = \mathbf{z}(t)\right] \\ = \mathbb{E}\left[\mathbb{E}\left[Z_a(t + \Delta t) \middle| \mathbf{Z}(t + \Delta t - 1)\right] \middle| \mathbf{Z}(t) = \mathbf{z}(t)\right]$$

where the second line follows from the law of total expectation. Note that $p_a(\mathbf{z}(t), 1) \equiv p_a(\mathbf{z}(t))$, with $p_a(\mathbf{z}(t))$ defined in (S2). Now applying (S5), we further obtain

$$\begin{aligned} p_a(\mathbf{z}(t), \Delta t) &= \mathbb{E} \left[Z_a(t + \Delta t - 1) \left(1 + h_a - \sum_{b=1}^M h_b Z_b(t + \Delta t - 1) - \mu L \right) + \mu \sum_{b=1, d_{ab}=1}^M Z_b(t + \Delta t - 1) \right. \\ &\quad \left. - rL(Z_a(t + \Delta t - 1) - \psi_a(\mathbf{Z}(t + \Delta t - 1))) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] + O\left(\frac{1}{N^2}\right) \\ &= p_a(\mathbf{z}(t), \Delta t - 1)(1 + h_a) + \mu \left(-Lp_a(\mathbf{z}(t), \Delta t - 1) + \sum_{b=1, d_{ab}=1}^M p_b(\mathbf{z}(t), \Delta t - 1) \right) \\ &\quad - \sum_{b=1}^M h_b \mathbb{E} \left[Z_a(t + \Delta t - 1) Z_b(t + \Delta t - 1) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] \\ &\quad - rL \mathbb{E} \left[Z_a(t + \Delta t - 1) - \psi_a(\mathbf{Z}(t + \Delta t - 1)) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] + O\left(\frac{1}{N^2}\right). \end{aligned}$$

By iterating the above procedure Δt times, applying the expansions in (S5) and (S7), and recalling (S3) and (S4), we arrive at the explicit expansion

$$\begin{aligned} p_a(\mathbf{z}(t), \Delta t) &= z_a(t) + \Delta t \left(z_a(t) \left(h_a - \sum_{b=1}^M h_b z_b(t) \right) + \mu \left(-Lz_a(t) + \sum_{b=1, d_{ab}=1}^M z_b(t) \right) - rL(z_a(t) - \psi_a(\mathbf{z}(t))) \right) + O\left(\frac{1}{N^2}\right). \end{aligned} \quad (\text{S25})$$

Next, we turn to deriving corresponding expansions for the conditional variance and covariance. Using the law of total variance, the conditional variance of the frequency of genotype a at generation $t + \Delta t$, conditioned on $\mathbf{Z}(t) = \mathbf{z}(t)$, is given by

$$\begin{aligned} \text{Var} \left(Z_a(t + \Delta t) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right) &= \mathbb{E} \left[\text{Var} \left(Z_a(t + \Delta t) \middle| \mathbf{Z}(t + \Delta t - 1), \mathbf{Z}(t) = \mathbf{z}(t) \right) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] \\ &\quad + \text{Var} \left(\mathbb{E} \left[Z_a(t + \Delta t) \middle| \mathbf{Z}(t + \Delta t - 1), \mathbf{Z}(t) = \mathbf{z}(t) \right] \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right) \\ &= \mathbb{E} \left[\text{Var} \left(Z_a(t + \Delta t) \middle| \mathbf{Z}(t + \Delta t - 1) \right) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] \\ &\quad + \text{Var} \left(\mathbb{E} \left[Z_a(t + \Delta t) \middle| \mathbf{Z}(t + \Delta t - 1) \right] \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right) \end{aligned} \quad (\text{S26})$$

where the second line follows from the Markov property of the WF process. The first term in (S26) can be written as

$$\begin{aligned} \mathbb{E} \left[\text{Var} \left(Z_a(t + \Delta t) \middle| \mathbf{Z}(t + \Delta t - 1) \right) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] &= \mathbb{E} \left[\frac{p_a(\mathbf{Z}(t + \Delta t - 1))(1 - p_a(\mathbf{Z}(t + \Delta t - 1)))}{N} \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] \\ &= \frac{z_a(t)(1 - z_a(t))}{N} + O\left(\frac{1}{N^2}\right) \end{aligned}$$

which was obtained by iteratively applying the expansions (S5) and (S7). Moreover, the second term in (S26) admits

$$\begin{aligned} &\text{Var} \left(\mathbb{E} \left[Z_a(t + \Delta t) \middle| \mathbf{Z}(t + \Delta t - 1) \right] \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right) \\ &= \text{Var} \left(Z_a(t + \Delta t - 1) \left(1 + h_a - \sum_{b=1}^M h_b Z_b(t + \Delta t - 1) - \mu L \right) + \mu \sum_{b=1, d_{ab}=1}^M Z_b(t + \Delta t - 1) \right. \\ &\quad \left. - rL(Z_a(t + \Delta t - 1) - \psi_a(\mathbf{Z}(t + \Delta t - 1))) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right) + O\left(\frac{1}{N^2}\right) \\ &= \text{Var} \left(Z_a(t + \Delta t - 1) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right) + O\left(\frac{1}{N^2}\right) \end{aligned} \quad (\text{S27})$$

where the second line follows from (S5) and the last line follows from first noting that for arbitrary random variables $\{X_i\}_{i=1}^n$, we have

$$\text{Var} \left(\sum_{i=1}^n a_i X_i \right) = \sum_{i=1}^n a_i^2 \text{Var}(X_i) + 2 \sum_{1 \leq i < j \leq n} a_i a_j \text{Cov}(X_i, X_j)$$

recalling again (S3) and (S4), and applying the expansions (S6) and (S7). Substituting (2.1) and (S27) into (S26), and iterating, we arrive at

$$\text{Var} \left(Z_a(t + \Delta t) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right) = \Delta t \frac{z_a(t)(1 - z_a(t))}{N} + O \left(\frac{1}{N^2} \right). \quad (\text{S28})$$

The expansion for the conditional covariance is obtained by following a similar procedure. For $a \neq b$, it leads to

$$\text{Covar} \left(Z_a(t + \Delta t), Z_b(t + \Delta t) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right) = -\Delta t \frac{z_a(t)z_b(t)}{N} + O \left(\frac{1}{N^2} \right). \quad (\text{S29})$$

2.2. Path integral representation of the likelihood function

Based on the above conditional moment expansions, for the observed set of time points $t_0 < t_1 < \dots < t_K$, we can present a path integral expression for the likelihood function (S24). The proof follows a similar procedure to that in Sections 1.3 and 1.4, with basic modifications to the drift vector and covariance matrix to account for frequency vector observations being taken multiple generations apart.

First, we consider genotype frequency evolution, and assume that genotype frequencies are observed at generations t and $t + \Delta t$, with no observations in-between. Then, after mapping from discrete-time to continuous-time and taking $N \rightarrow \infty$ as before (i.e., in equation S8), the modified drift vector now characterizes the expected infinitesimal change in genotype frequencies between continuous time points τ and $\tau + \Delta t$, and the covariance matrix characterizes the second moment of the change in genotype frequencies between the two time points. Specifically, using (S25), the a th element of the modified drift vector can be written as (see equation 4.99 of Risken(1))

$$\begin{aligned} \bar{d}_a(\mathbf{z}(\tau), \Delta t) &= \lim_{\delta\tau\Delta t \rightarrow 0} \frac{1}{\delta\tau\Delta t} \mathbb{E} \left[\left(\check{Z}_a(\tau + \delta\tau\Delta t) - \check{Z}_a(\tau) \right) \middle| \check{\mathbf{Z}}(\tau) = (\check{z}_1(\tau), \dots, \check{z}_M(\tau)) \right] \\ &= \lim_{N \rightarrow \infty} N \left(\check{z}_a(\tau) \left(h_a - \sum_{b=1}^M h_b \check{z}_b(\tau) \right) + \mu \left(-L \check{z}_a(\tau) + \sum_{b=1, d_{ab}=1}^M \check{z}_b(\tau) \right) - r(L-1)(\check{z}_a(\tau) - \psi_a(\mathbf{z}(\tau))) \right) \end{aligned}$$

which turns out to be equivalent to (S11), and hence

$$\bar{d}_a(\mathbf{z}(\tau), \Delta t) = \bar{d}_a(\mathbf{z}(\tau)).$$

An analogous equivalence also holds for the modified diffusion matrix. Specifically, this has (a, b) th entry (see equation 4.100 of Risken (1))

$$\begin{aligned} \bar{C}_{ab}(\mathbf{z}(\tau), \Delta t) &= \frac{1}{2} \lim_{\delta\tau\Delta t \rightarrow 0} \frac{1}{\delta\tau\Delta t} \mathbb{E} \left[\left(\check{Z}_a(\tau + \delta\tau\Delta t) - \check{Z}_a(\tau) \right) \left(\check{Z}_b(\tau + \delta\tau\Delta t) - \check{Z}_b(\tau) \right) \middle| \check{\mathbf{Z}}(\tau) = (\check{z}_1(\tau), \dots, \check{z}_M(\tau)) \right] \\ &= \frac{1}{2} \lim_{\delta\tau\Delta t \rightarrow 0} \frac{1}{\delta\tau\Delta t} \text{Covar} \left(\check{Z}_a(\tau + \delta\tau\Delta t), \check{Z}_b(\tau + \delta\tau\Delta t) \middle| \check{\mathbf{Z}}(\tau) = (\check{z}_1(\tau), \dots, \check{z}_M(\tau)) \right) \\ &= \bar{C}_{ab}(\mathbf{z}(\tau)) \end{aligned} \quad (\text{S30})$$

where the second line follows via the same arguments as in (S13), while the last line was obtained by applying the expansions (S28) and (S29).

As before, the transition density can then be approximated for small $\delta\tau\Delta t$ as

$$\begin{aligned} \phi(\mathbf{z}(\tau + \delta\tau\Delta t) | \mathbf{z}(\tau), N, \mu, \mathbf{h}) &\approx \frac{\exp \left(-\frac{1}{4\delta\tau\Delta t} (\mathbf{z}(\tau + \delta\tau\Delta t) - \mathbf{z}(\tau) - \bar{\mathbf{d}}(\mathbf{z}(\tau))\delta\tau\Delta t)^T \bar{\mathbf{C}}(\mathbf{z}(\tau))^{-1} (\mathbf{z}(\tau + \delta\tau\Delta t) - \mathbf{z}(\tau) - \bar{\mathbf{d}}(\mathbf{z}(\tau))\delta\tau\Delta t) \right)}{(4\pi\delta\tau\Delta t)^{M/2} \sqrt{\det(\bar{\mathbf{C}}(\mathbf{z}(\tau)))}}. \end{aligned}$$

Here we show explicit dependence on N , μ and $\mathbf{h}$. From this result, and recalling that $\delta\tau = 1/N$, the transition probability for a single generation of the original discrete-time discrete-frequency WF process can (for large N) be approximated by

$$\begin{aligned} P(\mathbf{z}(t_{k+1})|\mathbf{z}(t_k), N, \mu, \mathbf{h}) \\ \approx \phi(\mathbf{z}(t_{k+1})|\mathbf{z}(t_k), N, \mu, \mathbf{h})d\mathbf{z}(t_{k+1}) \\ = \left(\frac{N}{2\pi\Delta t_k}\right)^{M/2} \frac{d\mathbf{z}(t_{k+1})}{\sqrt{\det C(\mathbf{z}(t_k))}} \\ \times \exp\left(-\frac{N}{2} \sum_{a=1}^M \sum_{b=1}^M \left[z_a(t_{k+1}) - z_a(t_k) - d_a(\mathbf{z}(t_k))\right] (C^{-1}(\mathbf{z}(t_k)))_{ab} \left[z_b(t_{k+1}) - z_b(t_k) - d_b(\mathbf{z}(t_k))\right]\right) \end{aligned}$$

where $\Delta t_k = t_{k+1} - t_k$, the $d\mathbf{z}(t_{k+1}) = \prod_{a=1}^M dz_a(t_{k+1})$ represent small frequency differences accounting for the quantization of the continuous genotype frequency space, and where we have defined $d_a(\mathbf{z}(t_k)) := \frac{\bar{d}_a(\mathbf{z}(t_k))}{N}$ and

$$(C(\mathbf{z}(t_k)))_{ab} := 2(\bar{C}(\mathbf{z}(t_k)))_{ab}. \quad (\text{S31})$$

The path integral expression then follows by noting that the probability of observing a trajectory of genotype frequencies $(\mathbf{z}(t_1), \mathbf{z}(t_2), \dots, \mathbf{z}(t_K))$ is given by

$$\begin{aligned} P\left((\mathbf{z}(t_k))_{k=1}^K | \mathbf{z}(t_0), N, \mu, \mathbf{h}\right) &= \prod_{k=0}^{K-1} P(\mathbf{z}(t_{k+1})|\mathbf{z}(t_k), N, \mu, \mathbf{h}) \\ &\approx \left(\prod_{k=0}^{K-1} \left[\frac{1}{\sqrt{\det C(\mathbf{z}(t_k))}} \left(\frac{N}{2\pi\Delta t_k}\right)^{M/2} d\mathbf{z}(t_{k+1})\right]\right) \exp\left(-\frac{N}{2} S\left((\mathbf{z}(t_k))_{k=0}^K\right)\right) \end{aligned} \quad (\text{S32})$$

where

$$S\left((\mathbf{z}(t_k))_{k=0}^K\right) = \sum_{k=0}^{K-1} \frac{1}{\Delta t_k} \sum_{a=1}^M \sum_{b=1}^M [z_a(t_{k+1}) - z_a(t_k) - \Delta t_k d_a(\mathbf{z}(t_k))] (C^{-1}(\mathbf{z}(t_k)))_{ab} [z_b(t_{k+1}) - z_b(t_k) - \Delta t_k d_b(\mathbf{z}(t_k))].$$

Based on these results, a corresponding path integral approximation for the probability of observing a trajectory of mutant allele frequencies $(\mathbf{x}(t_1), \mathbf{x}(t_2), \dots, \mathbf{x}(t_K))$, and hence for the likelihood function (S24), is then obtained by mirroring the steps in Section 1.4, giving

$$P\left((\mathbf{x}(t_k))_{k=1}^K | \mathbf{x}(t_0), N, \mu, \mathbf{s}\right) \approx \left(\prod_{k=0}^{K-1} \frac{1}{\sqrt{\det C(\mathbf{x}(t_k))}} \left(\frac{N}{2\pi\Delta t_k}\right)^{L/2} \prod_{i=1}^L dx_i(t_{k+1})\right) \left(\prod_{k=0}^{K-1} \exp\left(-\frac{N}{2} S\left((\mathbf{x}(t_k))_{k=0}^K\right)\right)\right) \quad (\text{S33})$$

where

$$S\left((\mathbf{x}(t_k))_{k=0}^K\right) = \sum_{k=0}^{K-1} \frac{1}{\Delta t_k} \sum_{i=1}^L \sum_{j=1}^L [x_i(t_{k+1}) - x_i(t_k) - \Delta t_k d_i(\mathbf{x}(t_k))] (C^{-1}(\mathbf{x}(t_k)))_{ij} [x_j(t_{k+1}) - x_j(t_k) - \Delta t_k d_j(\mathbf{x}(t_k))].$$

2.3. The MPL estimator solution

Returning to the MAP problem (S23), we note that this is equivalent to

$$\hat{\mathbf{s}} = \arg \max_{\mathbf{s}} \left(\log \mathcal{L}(\mathbf{s} | N, \mu, (\mathbf{x}(t_k))_{k=0}^K) + \log P_{\text{prior}}(\mathbf{s}) \right). \quad (\text{S34})$$

The likelihood $\mathcal{L}(\mathbf{s} | N, \mu, (\mathbf{x}(t_k))_{k=0}^K)$ is given by (S24) and is approximated (S33). Assuming a conjugate-prior distribution

$$P_{\text{prior}}(\mathbf{s}) = \frac{1}{(2\pi\sigma^2)^{L/2}} \exp\left(-\frac{1}{2\sigma^2} \mathbf{s}^T \mathbf{s}\right),$$

we take the vector derivative of the right-hand side of (S34) with respect to $\mathbf{s}$ (see equation 10, Chapter 10.2.1 of Lutkepohl(2)), equate to zero, and then solve for $\mathbf{s}$. This yields

$$\hat{\mathbf{s}}_i = \sum_{j=1}^L \left[\sum_{k=0}^{K-1} \Delta t_k C(\mathbf{x}(t_k)) + \gamma I \right]_{ij}^{-1} \left[x_j(t_K) - x_j(t_0) - \mu \sum_{k=0}^{K-1} \Delta t_k (1 - 2x_j(t_k)) \right] \quad (\text{S35})$$

for $i = 1, \dots, L$, where $\gamma = 1/N\sigma^2$. This is the desired *MPL estimate* of the selection coefficients, reported in (10) of Methods. Here, we note that γ can also be viewed as a parameter that regularizes the covariance matrix prior to inversion.

2.4. Equivalence of genotype- and allele-level analyses

While the MPL estimator was derived using a path integral expression for the probability of mutant allele frequency trajectories, the WF evolutionary process is defined for genotypes. Hence, one may naturally ask whether there is any loss in information in considering only the marginal frequency dynamics. As we now show, the answer is no. Specifically, we show that the estimate of the selection coefficients based on genotype frequencies is equivalent to the derived MPL estimate.

We begin by defining G as a $M \times L$ matrix with (a, j) th entry $G_{aj} = g_j^a$. Let $\mathbf{h} = (h_1, h_2, \dots, h_M)$ denote the column vector of genotype selection coefficients, which relates to the mutant allele selection coefficients via

$$\mathbf{h} = G\mathbf{s}. \quad (\text{S36})$$

Given the observed genotype frequencies $(\mathbf{z}(t_0), \mathbf{z}(t_1), \dots, \mathbf{z}(t_K))$, the MAP estimator of the selection coefficients admits

$$\hat{\mathbf{s}} = \arg \max_{\mathbf{s}} \left(\log \mathcal{L}(\mathbf{s} | \mu, N, (\mathbf{z}(t_k))_{k=0}^K) + \log P_{\text{prior}}(\mathbf{s}) \right),$$

where the likelihood function is given as

$$\mathcal{L}(\mathbf{s} | N, \mu, (\mathbf{z}(t_k))_{k=0}^K) = P((\mathbf{z}(t_k))_{k=1}^K | \mathbf{z}(t_0), N, \mu, \mathbf{h}) \quad (\text{S37})$$

Note that the right-hand side of (S37) is approximated by (S32). Expressing the genotype selection coefficients in (S32) in terms of mutant allele selection coefficients using (S36) gives the likelihood of the mutant allele selection coefficients $\mathbf{s}$ for an observed genotype frequency path and parameters N, μ . Differentiating the resulting expression with respect to $\mathbf{s}$ and equating to zero leads to

$$\mathbf{0} = \sum_{k=0}^{K-1} \left(G^T \mathbf{z}(t_{k+1}) - G^T \mathbf{z}(t_k) - G^T C(\mathbf{z}(t_k)) G \mathbf{s} - \mu G^T E \mathbf{z}(t_k) - r L \left(G^T \mathbf{z}(t_k) - G^T \boldsymbol{\psi}(\mathbf{z}(t_k)) \right) \right) + \frac{1}{N\sigma^2} \mathbf{s}, \quad (\text{S38})$$

where $\boldsymbol{\psi}(\mathbf{z}(t_k)) = \{\psi_a(\mathbf{z}(t_k))\}_{a=1}^M$ and $C(\mathbf{z}(t_k))$ is the genotype covariance matrix as given by (S30) and (S31), i.e., with elements

$$C_{ab}(\mathbf{z}(t_k)) = \begin{cases} z_a(t_k)(1 - z_a(t_k)) & a = b \\ -z_a(t_k)z_b(t_k) & a \neq b, \end{cases}$$

while matrix E has elements

$$E_{ab} = \begin{cases} -L & a = b \\ 0 & d_{ab} > 1 \\ 1 & d_{ab} = 1. \end{cases}$$

Noting that the relation between allele and genotype frequencies (S18) can be expressed in vector form as $\mathbf{x}(t_k) = G^T \mathbf{z}(t_k)$, and that $C(\mathbf{x}(t_k)) = G^T C(\mathbf{z}(t_k)) G$ and $G^T \boldsymbol{\psi}(\mathbf{z}(t_k)) = \mathbf{x}(t_k)$, we can solve (S38) to obtain the MAP estimate of the allele selection coefficients

$$\hat{\mathbf{s}} = \left[\sum_{k=0}^{K-1} \Delta t_k C(\mathbf{x}(t_k)) + \gamma I \right]^{-1} \left[\mathbf{x}(t_K) - \mathbf{x}(t_0) - \mu \sum_{k=0}^{K-1} \Delta t_k (1 - 2\mathbf{x}(t_k)) \right], \quad (\text{S39})$$

where $\gamma = 1/N\sigma^2$. This is the same as the MPL estimator (S35).

This equivalence is important. It implies that under the additive fitness model, only the single and pairwise mutational frequencies are needed for optimally estimating the selection coefficients, and higher order information (e.g., three-locus mutational frequencies) are irrelevant. The same equivalences can also be shown for extensions of the MPL estimator presented in the following section. Further extensions will be explored in future work.

2.5. Extension to multiple alleles per locus and asymmetric mutation probabilities

Here we demonstrate the extension of our inference framework to incorporate multiple alleles per locus and asymmetric mutation probabilities. The same notation and definitions introduced in Section 1 and 2 will also be used, unless stated otherwise, but will be interpreted in terms of the extended model. For example, $\tilde{Z}(\tau)$ was introduced in (S8) to refer to the binary genotype continuous process with symmetric mutation probabilities. This notation will also be used here, but it will now refer to the non-binary genotype continuous process with asymmetric mutation probabilities. We first describe the model in detail.

We assume there are ℓ alleles per locus, thus resulting in $M = \ell^L$ genotypes, with the multi-allelic sequence for genotype a denoted by $\mathbf{g}^a = (g_1^a, g_2^a, \dots, g_L^a)$, where g_i^a is the allele at locus i . Nucleotide sequences, for example, will have $g_i^a \in \{A, C, G, T\}$. For convenience, we denote each allele with an integer representation, and consider the fitness of each genotype with respect to a reference sequence comprising of allele ℓ at each locus, i.e., the reference sequence is a sequence of ℓ s. For example, if $\ell = 4$ and $L = 3$, then the reference sequence is represented by $(4, 4, 4)$. We assume, once again, an additive model of fitness, and that the reference genotype has a fitness of one. Under these assumptions, the selection coefficient for genotype a is given by

$$h_a = f_a - 1 = \sum_{i=1}^L \sum_{\alpha=1}^{\ell-1} \delta(g_i^a, \alpha) s_{i,\alpha},$$

where $s_{i,\alpha}$ is the selection coefficient of allele α at locus i , and where $\delta(\cdot, \cdot)$ is the Kronecker-delta function,

$$\delta(x, y) := \begin{cases} 1 & x = y \\ 0 & \text{otherwise.} \end{cases}$$

Note that the assumptions above imply that $s_{i,\ell} = 0$ for all $i = 1, \dots, L$.

The allele frequency vector, describing the frequency of all alleles except (the reference) allele one, is given by $\mathbf{x}(t) = (x_{1,1}(t), \dots, x_{1,\ell-1}(t), \dots, x_{L,1}(t), \dots, x_{L,\ell-1}(t))$, where $x_{i,\alpha}(t)$ denotes the observed frequency of allele α at locus i during generation t , and is related to the genotype frequencies by

$$x_{i,\alpha}(t) = \sum_{a=1}^M \delta(g_i^a, \alpha) z_a(t). \quad (\text{S40})$$

Finally, we denote $\mu_{\alpha\beta}$ as the mutation probability per generation from allele α to allele β and μ_{ab} as the mutation probability per generation from genotype a to genotype b . These are related by

$$\mu_{ab} = \prod_{i=1}^L \left(\sum_{\alpha=1}^{\ell} \sum_{\beta=1}^{\ell} \mu_{\alpha\beta} \delta(g_i^a, \alpha) \delta(g_i^b, \beta) \right).$$

Given this model, the MAP estimate of the selection coefficients is the solution to

$$\hat{\mathbf{s}} = \arg \max_{\mathbf{s}} \mathcal{L}(\mathbf{s} | \boldsymbol{\mu}, N, (\mathbf{x}(t_k))_{k=0}^T) P_{\text{prior}}(\mathbf{s}), \quad (\text{S41})$$

where

$$\begin{aligned} \mathcal{L}(\mathbf{s} | \boldsymbol{\mu}, N, (\mathbf{x}(t_k))_{k=0}^T) &= P((\mathbf{x}(t_k))_{k=1}^K | \mathbf{x}(t_0), N, \boldsymbol{\mu}, \mathbf{s}) \\ &= \prod_{k=0}^{K-1} P(\mathbf{x}(t_{k+1}) | \mathbf{x}(t_k), N, \boldsymbol{\mu}, \mathbf{s}) \end{aligned} \quad (\text{S42})$$

is the likelihood of the selection coefficients $\mathbf{s} = (s_{1,1}, \dots, s_{1,\ell-1}, \dots, s_{L,1}, \dots, s_{L,\ell-1})$ and $P_{\text{prior}}(\mathbf{s})$ their prior distribution.

As with the simpler biallelic and symmetric mutation probability scenario, the main challenge in solving (S41) is that it requires computing the likelihood (S42), which is complicated. This is simplified as before, by adopting the path integral approach outlined in Section 1; but now extending the analysis to account for sampling times spanning multiple generations, multiple alleles per locus and asymmetric mutation probabilities. Other than these differences, we may follow the same approach, starting by giving asymptotic moment expansions. By direct analogy to (S3) and (S4), we will work under the assumption that N is large, and that as $N \rightarrow \infty$,

$$s_{i,\alpha} = \frac{\bar{s}_{i,\alpha}}{N} + O\left(\frac{1}{N^2}\right), \quad h_a = \frac{\bar{h}_a}{N} + O\left(\frac{1}{N^2}\right), \quad \mu_{\beta\alpha} = \frac{\bar{\mu}_{\beta\alpha}}{N} + O\left(\frac{1}{N^2}\right), \quad \mu_{ab} = \frac{\bar{\mu}_{ab}}{N} + O\left(\frac{1}{N^2}\right), \quad r = \frac{\bar{r}}{N} + O\left(\frac{1}{N^2}\right) \quad (\text{S43})$$

where $\bar{s}_{i,\alpha}$, $\bar{h}_a$, $\bar{\mu}_{\beta\alpha}$, $\bar{\mu}_{ab}$ and $\bar{r}$ are constants independent of N .

2.5.1. Conditional moment expansions

The mean frequency of genotype a at generation $t + 1$, conditioned on $\mathbf{Z}(t) = \mathbf{z}(t)$, admits

$$\begin{aligned} p_a(\mathbf{z}(t)) &:= \mathbb{E} \left[Z_a(t+1) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right] \\ &= \frac{(1+h_a)y_a(t) + \sum_{b=1}^M (\mu_{ba}(1+h_b)y_b(t) - \mu_{ab}(1+h_a)y_a(t))}{\sum_{b=1}^M (1+h_b)y_b(t)}, \end{aligned} \quad (\text{S44})$$

where recall that

$$y_a(t) = (1-r)^L z_a(t) + \left(1 - (1-r)^L\right) \psi_a(\mathbf{z}(t)).$$

Utilizing (S44) along with (S43), and applying a similar proof to that described in Section 2.1, the mean frequency of genotype a at generation $t + \Delta t$, conditioned on $\mathbf{Z}(t) = \mathbf{z}(t)$, then admits

$$\begin{aligned} p_a(\mathbf{z}(t), \Delta t) &= z_a(t) + \Delta t \left(z_a(t) \left(h_a - \sum_{b=1}^M h_b z_b(t) \right) \right. \\ &\quad \left. + \sum_{b=1, d_{ab}=1}^M (\mu_{ba} z_b(t) - \mu_{ab} z_a(t)) - r(L-1)(z_a(t) - \psi_a(\mathbf{z}(t))) \right) + O\left(\frac{1}{N^2}\right) \\ &= z_a(t) + O\left(\frac{1}{N^2}\right), \end{aligned} \quad (\text{S45})$$

where d_{ab} here denotes the number of loci for which the allele identity at genotypes a and b differ, i.e.,

$$d_{ab} := \sum_{i=1}^L \left(1 - \delta(g_i^a, g_i^b)\right).$$

The corresponding expressions for the variance and covariance of the frequency of genotype a at generation $t + \Delta t$, conditioned on $\mathbf{Z}(t) = \mathbf{z}(t)$ are obtained analogously to (S28) and (S29) as

$$\text{Var} \left(Z_a(t + \Delta t) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right) = \Delta t \frac{z_a(t)(1-z_a(t))}{N} + O\left(\frac{1}{N^2}\right), \quad (\text{S46})$$

and

$$\text{Covar} \left(Z_a(t + \Delta t), Z_b(t + \Delta t) \middle| \mathbf{Z}(t) = \mathbf{z}(t) \right) = -\Delta t \frac{z_a(t)z_b(t)}{N} + O\left(\frac{1}{N^2}\right). \quad (\text{S47})$$

2.5.2. Path integral representation of the likelihood function

Based on the above conditional moment expansions, we can derive a path integral representation for the likelihood function (S42). Starting by considering genotype evolution, the drift vector and diffusion matrix in this case are given respectively by

$$\begin{aligned} \bar{d}_a(\check{\mathbf{z}}(\tau), \Delta t) &= \lim_{\delta\tau \Delta t \rightarrow 0} \frac{1}{\delta\tau \Delta t} \mathbb{E} \left[\check{Z}_a(\tau + \delta\tau \Delta t) - \check{Z}_a(\tau) \middle| \check{\mathbf{Z}}(\tau) = (\check{z}_1(\tau), \dots, \check{z}_M(\tau)) \right] \\ &= \check{z}_a(\tau) \left(\bar{h}_a - \sum_{b=1}^M \bar{h}_b \check{z}_b(\tau) \right) + \sum_{b=1, d_{ab}=1}^M (\bar{\mu}_{ba} \check{z}_b(\tau) - \bar{\mu}_{ab} \check{z}_a(\tau)) - \bar{r}(L-1)(\check{z}_a(\tau) - \psi_a(\check{\mathbf{z}}(\tau))) \end{aligned} \quad (\text{S48})$$

and

$$\bar{C}_{ab}(\check{\mathbf{z}}(\tau), \Delta t) = \frac{1}{2} \begin{cases} \check{z}_a(\tau)(1-\check{z}_a(\tau)) & a = b \\ -\check{z}_a(\tau)\check{z}_b(\tau) & a \neq b \end{cases}, \quad (\text{S49})$$

which follow from (S45), (S46), (S47), and by adopting the procedure described in Section 2.2. From these results, and again following the procedure in Section 2.2, the conditional probability of observing a trajectory of genotype frequencies $(z(t_1), z(t_2), \dots, z(t_K))$ is obtained as

$$P\left((z(t_k))_{k=1}^K | z(t_0)\right) = \prod_{k=0}^{K-1} P(z(t_{k+1}) | z(t_k)) \\ \approx \left(\prod_{k=0}^{K-1} \left[\frac{1}{\sqrt{\det \bar{C}(z(t_k))}} \left(\frac{N}{4\pi\Delta t_k} \right)^{M/2} \prod_{a=1}^M dz_a(t_{k+1}) \right] \right) \exp\left(-\frac{N}{4} S\left((z(t_k))_{k=0}^K\right)\right)$$

where $\Delta t_k = t_{k+1} - t_k$ and

$$S\left((z(t_k))_{k=0}^K\right) = \sum_{k=0}^{K-1} \frac{1}{\Delta t_k} \sum_{a=1}^M \sum_{b=1}^M \left[z_a(t_{k+1}) - z_a(t_k) - \frac{\Delta t_k \bar{d}_a(z(t_k))}{N} \right] \\ \times (\bar{C}^{-1}(z(t_k)))_{ab} \left[z_b(t_{k+1}) - z_b(t_k) - \frac{\Delta t_k \bar{d}_b(z(t_k))}{N} \right].$$

In this last equation we have dropped the second argument of the drift vector and diffusion matrix, as both quantities turn out to be independent of Δt , as seen from (S48) and (S49).

Now turn to mutant allele evolution. By noting the linearity of expectation and the genotype-to-allele mapping (S40), the mean frequency of allele α at locus i during time $\tau + \delta\tau\Delta t$, conditioned on $\check{X}(\tau) = \check{x}(\tau)$, is given by

$$\sum_{a=1}^M \delta(g_i^a, \alpha) \check{z}_a(\tau) + \underline{d}_{i,\alpha}(\check{x}(\tau)) \delta\tau\Delta t = \check{x}_{i,\alpha}(\tau) + \underline{d}_{i,v}(\check{x}(\tau)) \delta\tau\Delta t$$

where

$$\underline{d}_{i,\alpha}(\check{x}(\tau)) = \sum_{a=1}^M \delta(g_i^a, \alpha) \bar{d}_a(\check{z}(\tau), \Delta t) \\ = \sum_{a=1}^M \delta(g_i^a, \alpha) \left(\check{z}_a(\tau) \left(\bar{h}_a - \sum_{b=1}^M \bar{h}_b \check{z}_b(\tau) \right) + \sum_{b=1, d_{ab}=1}^M (\bar{\mu}_{ba} \check{z}_b(\tau) - \bar{\mu}_{ab} \check{z}_a(\tau)) - \bar{r}(L-1) (\check{z}_a(\tau) - \psi_a(\check{z}(\tau))) \right) \\ = I_1 + I_2 + I_3$$

with the second line following from (S48). For I_1 , we have

$$I_1 = \sum_{a=1}^M \delta(g_i^a, \alpha) \check{z}_a(\tau) \left(\bar{h}_a - \sum_{b=1}^M \bar{h}_b \check{z}_b(\tau) \right) \\ = \sum_{a=1}^M \delta(g_i^a, \alpha) \check{z}_a(\tau) \sum_{j=1}^L \sum_{\beta=1}^{\ell-1} \delta(g_j^a, \beta) \bar{s}_{j,\beta} - \sum_{a=1}^M \delta(g_i^a, \alpha) \check{z}_a(\tau) \sum_{b=1}^M \sum_{j=1}^L \sum_{\beta=1}^{\ell-1} \delta(g_j^b, \beta) \bar{s}_{j,\beta} \check{z}_b(\tau) \\ = \sum_{j=1}^L \sum_{\beta=1}^{\ell-1} (\check{x}_{ij,\alpha\beta}(\tau) - \check{x}_{i,\alpha}(\tau) \check{x}_{j,\beta}(\tau)) \bar{s}_{j,\beta}$$

where $\check{x}_{ij,\alpha\beta}(\tau)$ denotes the frequency of allele α and β occurring respectively at loci i and j during time τ , and given by

$$\check{x}_{ij,\alpha\beta}(\tau) = \sum_{a=1}^M \delta(g_i^a, \alpha) \delta(g_j^a, \beta) \check{z}_a(\tau).$$

For I_2 , we have

$$\begin{aligned} I_2 &= \sum_{a=1}^M \delta(g_i^a, \alpha) \sum_{b=1, d_{ab}=1}^M (\bar{\mu}_{ba} \check{z}_b(\tau) - \bar{\mu}_{ab} \check{z}_a(\tau)) \\ &= \sum_{a=1}^M \delta(g_i^a, \alpha) \sum_{b=1, d_{ab}=1}^M \left(\delta(g_i^b, \alpha) + 1 - \delta(g_i^b, \alpha) \right) (\bar{\mu}_{ba} \check{z}_b(\tau) - \bar{\mu}_{ab} \check{z}_a(\tau)) \\ &= I_{2a} + I_{2b} \end{aligned}$$

where

$$\begin{aligned} I_{2a} &= \sum_{a=1}^M \sum_{b=1, d_{ab}=1}^M \delta(g_i^a, \alpha) \delta(g_i^b, \alpha) (\bar{\mu}_{ba} \check{z}_b(\tau) - \bar{\mu}_{ab} \check{z}_a(\tau)) \\ I_{2b} &= \sum_{a=1}^M \sum_{b=1, d_{ab}=1}^M \delta(g_i^a, \alpha) (1 - \delta(g_i^b, \alpha)) (\bar{\mu}_{ba} \check{z}_b(\tau) - \bar{\mu}_{ab} \check{z}_a(\tau)). \end{aligned}$$

Consider now I_{2a} , where we observe that (i) $\delta(g_i^a, \alpha) \delta(g_i^b, \alpha)$ is non-zero only when genotypes a and b both have allele α at locus i , while (ii) the summation $\sum_{b=1, d_{ab}=1}^M$ is over all genotypes b which have a different allele from genotype a at only one locus. These two observations imply that $\delta(g_i^a, \alpha) \delta(g_i^b, \alpha)$ is non-zero, for $b = 1, \dots, M$, with $d_{ab} = 1$, only if a and b have a different allele at a single locus, but where the position of this locus is different from i . To illustrate this, consider a simple example with $L = 2$ and $\ell = 3$, in which case

$$\begin{aligned} \mathbf{g}^1 &= \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \quad \mathbf{g}^2 = \begin{pmatrix} 1 \\ 2 \end{pmatrix}, \quad \mathbf{g}^3 = \begin{pmatrix} 1 \\ 3 \end{pmatrix}, \\ \mathbf{g}^4 &= \begin{pmatrix} 2 \\ 1 \end{pmatrix}, \quad \mathbf{g}^5 = \begin{pmatrix} 2 \\ 2 \end{pmatrix}, \quad \mathbf{g}^6 = \begin{pmatrix} 2 \\ 3 \end{pmatrix}, \\ \mathbf{g}^7 &= \begin{pmatrix} 3 \\ 1 \end{pmatrix}, \quad \mathbf{g}^8 = \begin{pmatrix} 3 \\ 2 \end{pmatrix}, \quad \mathbf{g}^9 = \begin{pmatrix} 3 \\ 3 \end{pmatrix}. \end{aligned}$$

Then if $i = 1$ and $\alpha = 1$, the only genotype-pairs (a, b) which result in a non-zero $\delta(g_i^a, \alpha) \delta(g_i^b, \alpha)$ (subject to the condition $d_{ab} = 1$) are $(1, 2)$, $(2, 1)$, $(1, 3)$, $(3, 1)$, $(2, 3)$ and $(3, 2)$, as these genotype-pairs differ only at the second locus, while having allele one at locus $i = 1$.

Observe that every genotype-pair which result in a non-zero $\delta(g_i^a, \alpha) \delta(g_i^b, \alpha)$ (subject to the condition $d_{ab} = 1$) occur in conjugates, i.e., if $(1, 2)$ is such a genotype-pair, then so is $(2, 1)$. This implies that $I_{2a} = 0$, as for every genotype pair (a, b) resulting in a non-zero $\delta(g_i^a, \alpha) \delta(g_i^b, \alpha) \bar{\mu}_{ba} \check{z}_b(\tau)$, there always exist one other conjugate pair (b, a) which cancels this term through the $-\delta(g_i^a, \alpha) \delta(g_i^b, \alpha) \bar{\mu}_{ab} \check{z}_a(\tau)$ term.

Consider now I_{2b} , and observe that the $\delta(g_i^a, \alpha) (1 - \delta(g_i^b, \alpha))$ term is non-zero only when genotype a , but not genotype b , has allele α at locus i . Combining the above, we can thus write I_2 as

$$\begin{aligned} I_2 &= I_{2b} = \sum_{\beta=1, \beta \neq \alpha}^{\ell} (\bar{\mu}_{\beta\alpha} \check{x}_{i,\beta}(\tau) - \bar{\mu}_{\alpha\beta} \check{x}_{i,\alpha}(\tau)) \\ &= \sum_{\beta=1}^{\ell} (\bar{\mu}_{\beta\alpha} \check{x}_{i,\beta}(\tau) - \bar{\mu}_{\alpha\beta} \check{x}_{i,\alpha}(\tau)) \\ &= \sum_{\beta=1}^{\ell-1} \bar{\mu}_{\beta\alpha} \check{x}_{i,\beta}(\tau) + \bar{\mu}_{\ell\alpha} \left(1 - \sum_{\beta=1}^{\ell-1} \check{x}_{i,\beta}(\tau) \right) - \sum_{\beta=1}^{\ell} \bar{\mu}_{\alpha\beta} \check{x}_{i,\alpha}(\tau) \\ &= \bar{\mu}_{\ell\alpha} + \sum_{\beta=1}^{\ell-1} (\bar{\mu}_{\beta\alpha} - \bar{\mu}_{\ell\alpha}) \check{x}_{i,\beta}(\tau) - \check{x}_{i,\alpha}(\tau) \sum_{\beta=1}^{\ell} \bar{\mu}_{\alpha\beta}. \end{aligned}$$

Note that in the third line above we have used the fact that the frequency of the reference allele is one minus the sum of the frequency of all other alleles.

Finally, we have

$$I_3 = -\bar{r}(L-1) \sum_{a=1}^M \delta(g_i^a, \alpha) (\check{z}_a(\tau) - \psi_a(\check{\mathbf{z}}(\tau))).$$

Similar to the biallelic case, we define

$$\theta_{i,\alpha}^{cd} := \sum_{a=1}^M \delta(g_i^a, \alpha) R_{a,cd}$$

where $\theta_{i,\alpha}^{cd}$ is the probability that genotypes c and d recombine to form a genotype which has an allele α at locus i . Next we split this summation into four summations, one for each of the four possible recombination scenarios. Namely, when both genotypes c and d have allele α at locus i , when both do not have allele α at locus i and when only one of the genotypes has allele α at locus i . We thus have

$$\begin{aligned} \sum_{a=1}^M \delta(g_i^a, \alpha) \psi_a(\check{\mathbf{z}}(\tau)) &= \sum_{a=1}^M \delta(g_i^a, \alpha) \sum_{c=1}^M \sum_{d=1}^M R_{a,cd} \check{z}_c(\tau) \check{z}_d(\tau) \\ &= \sum_{c=1}^M \sum_{d=1}^M \theta_{i,\alpha}^{cd} \check{z}_c(\tau) \check{z}_d(\tau) \\ &= \sum_{c=1}^M \left(\sum_{d=1}^M \theta_{i,\alpha}^{cd} \delta(g_i^c, \alpha) \delta(g_i^d, \alpha) \check{z}_c(\tau) \check{z}_d(\tau) + \sum_{d=1}^M \theta_{i,\alpha}^{cd} \delta(g_i^c, \alpha) (1 - \delta(g_i^d, \alpha)) \check{z}_c(\tau) \check{z}_d(\tau) \right) \\ &\quad + \sum_{c=1}^M \left(\sum_{d=1}^M \theta_{i,\alpha}^{cd} (1 - \delta(g_i^c, \alpha)) \delta(g_i^d, \alpha) \check{z}_c(\tau) \check{z}_d(\tau) + \sum_{d=1}^M \theta_{i,\alpha}^{cd} (1 - \delta(g_i^c, \alpha)) (1 - \delta(g_i^d, \alpha)) \check{z}_c(\tau) \check{z}_d(\tau) \right). \end{aligned}$$

Moreover, noting that

$$\begin{aligned} \theta_{i,\alpha}^{cd} \delta(g_i^c, \alpha) \delta(g_i^d, \alpha) &= \delta(g_i^c, \alpha) \delta(g_i^d, \alpha) \\ \theta_{i,\alpha}^{cd} \delta(g_i^c, \alpha) (1 - \delta(g_i^d, \alpha)) &= \frac{1}{2} \delta(g_i^c, \alpha) (1 - \delta(g_i^d, \alpha)) \\ \theta_{i,\alpha}^{cd} (1 - \delta(g_i^c, \alpha)) \delta(g_i^d, \alpha) &= \frac{1}{2} (1 - \delta(g_i^c, \alpha)) \delta(g_i^d, \alpha) \\ \theta_{i,\alpha}^{cd} (1 - \delta(g_i^c, \alpha)) (1 - \delta(g_i^d, \alpha)) &= 0 \end{aligned}$$

where the factor of $\frac{1}{2}$ arises because there is a 50% chance that genotype c (d) with allele v at locus i and genotype d (c) which does not have allele α at locus i will recombine to a genotype with allele α at locus i , we have

$$\begin{aligned} \sum_{a=1}^M \delta(g_i^a, \alpha) \psi_a(\check{\mathbf{z}}(\tau)) &= \sum_{c=1}^M \left(\sum_{d=1}^M \delta(g_i^c, \alpha) \delta(g_i^d, \alpha) \check{z}_c(\tau) \check{z}_d(\tau) + \frac{1}{2} \sum_{d=1}^M \delta(g_i^c, \alpha) (1 - \delta(g_i^d, \alpha)) \check{z}_c(\tau) \check{z}_d(\tau) \right) \\ &\quad + \frac{1}{2} \sum_{c=1}^M \sum_{d=1}^M (1 - \delta(g_i^c, \alpha)) \delta(g_i^d, \alpha) \check{z}_c(\tau) \check{z}_d(\tau) \\ &= \check{x}_{i,\alpha}^2(\tau) + \frac{1}{2} \check{x}_{i,\alpha}(\tau) (1 - \check{x}_{i,\alpha}(\tau)) + \frac{1}{2} \check{x}_{i,\alpha}(\tau) (1 - \check{x}_{i,\alpha}(\tau)) \\ &= \check{x}_{i,\alpha}(\tau) \end{aligned}$$

thus implying that $\sum_{a=1}^M \delta(g_i^a, \alpha) (\check{z}_a(\tau) - \psi_a(\check{\mathbf{z}}(\tau))) = 0$ and hence $I_3 = 0$.

The final expression for the drift vector is thus

$$\underline{d}_{i,\alpha}(\check{\mathbf{x}}(\tau), \Delta t) = \sum_{j=1}^L \sum_{\beta=1}^{\ell-1} (\check{x}_{ij,\alpha\beta}(\tau) - \check{x}_{i,\alpha}(\tau) \check{x}_{j,\beta}(\tau)) \bar{s}_{j,\beta} + \bar{\mu}_{\ell\alpha} + \sum_{\beta=1}^{\ell-1} (\bar{\mu}_{\beta\alpha} - \bar{\mu}_{\ell\alpha}) \check{x}_{i,\beta}(\tau) - \check{x}_{i,\alpha}(\tau) \sum_{\beta=1}^{\ell} \bar{\mu}_{\alpha\beta}. \quad (\text{S50})$$

As expected, we observe that when there are two alleles, the mutation probability is symmetric and the reference allele is the WT, the expression in (S50) reduces to the expression in (S20).

We now move on to the diffusion matrix. This can be partitioned into L^2 blocks each of size $(\ell - 1) \times (\ell - 1)$, where the (i, j) th block has (α, β) th entry

$$\begin{aligned}\underline{C}_{ij, \alpha \beta}(\mathbf{x}(\tau)) &= \sum_{a=1}^M \sum_{b=1}^M \delta(g_i^a, \alpha) \delta(g_j^b, \beta) \bar{C}_{ab}(\mathbf{z}(\tau)) \\ &= \frac{1}{2} \sum_{a=1}^M \delta(g_i^a, \alpha) \delta(g_j^a, \beta) \tilde{z}_a(\tau) (1 - \tilde{z}_a(\tau)) - \frac{1}{2} \sum_{a=1}^M \sum_{b=1, b \neq a}^M \delta(g_i^a, \alpha) \delta(g_j^b, \beta) \tilde{z}_a(\tau) \tilde{z}_b(\tau) \\ &= \frac{1}{2} \sum_{a=1}^M \delta(g_i^a, \alpha) \delta(g_j^a, \beta) \tilde{z}_a(\tau) - \frac{1}{2} \left(\sum_{a=1}^M \delta(g_i^a, \alpha) \tilde{z}_a(\tau) \right) \left(\sum_{b=1}^M \delta(g_j^b, \beta) \tilde{z}_b(\tau) \right) \\ &= \frac{1}{2} (\tilde{x}_{ij, \alpha \beta}(\tau) - \tilde{x}_{i, \alpha}(\tau) \tilde{x}_{j, \beta}(\tau)).\end{aligned}$$

Following along the lines of the proof in Section 1.4, the probability of observing mutant allele frequencies $(\mathbf{x}(t_1), \mathbf{x}(t_2), \dots, \mathbf{x}(t_K))$, and hence the likelihood function (S42), is approximated by the path integral representation

$$\mathcal{L}(\mathbf{s} | \boldsymbol{\mu}, N, (\mathbf{x}(t_k))_{k=0}^K) \approx \left[\prod_{k=0}^{K-1} \frac{1}{\sqrt{\det C(\mathbf{x}(t_k))}} \left(\frac{N}{2\pi \Delta t_k} \right)^{L(\ell-1)/2} d\mathbf{x}(t_{k+1}) \right] \exp \left(-\frac{N}{2} S((\mathbf{x}(t_k))_{k=0}^K) \right) \quad (\text{S51})$$

where $d\mathbf{x}(t_{k+1}) = \prod_{i=1}^L \prod_{\alpha=1}^{\ell-1} dx_{i, \alpha}(t_{k+1})$ and

$$\begin{aligned}S((\mathbf{x}(t_k))_{k=0}^K) &= \sum_1 \frac{[x_{i, \alpha}(t_{k+1}) - x_{i, \alpha}(t_k) - \Delta t_k d_{i, \alpha}(\mathbf{x}(t_k))] (C^{-1}(\mathbf{x}(t_k)))_{ij, \alpha \beta} [x_{j, \beta}(t_{k+1}) - x_{j, \beta}(t_k) - \Delta t_k d_{j, \beta}(\mathbf{x}(t_k))]}{\Delta t_k}.\end{aligned}$$

Here we have adopted the shorthand notation $\sum_1 = \sum_{k=0}^{K-1} \sum_{i=1}^L \sum_{\alpha=1}^{\ell-1} \sum_{j=1}^L \sum_{\beta=1}^{\ell-1}$, while $(\mathbf{X})_{ij, \alpha \beta}$ indicates the (α, β) th entry of the (i, j) th block, with each block having dimension $(\ell - 1) \times (\ell - 1)$. Moreover, we have

$$d_{i, \alpha}(\mathbf{x}(t_k)) = \sum_{j=1}^L \sum_{\beta=1}^{\ell-1} (x_{ij, \alpha \beta}(t_k) - x_{i, \alpha}(t_k) x_{j, \beta}(t_k)) s_{j, \beta} + \mu_{\ell \alpha} + \sum_{\beta=1}^{\ell-1} (\mu_{\beta \alpha} - \mu_{\ell \alpha}) x_{i, \beta}(t_k) - x_{i, \alpha}(t_k) \sum_{\beta=1}^{\ell} \mu_{\alpha \beta}$$

and

$$C_{ij, \alpha \beta}(\mathbf{x}(t_k)) = x_{ij, \alpha \beta}(t_k) - x_{i, \alpha}(t_k) x_{j, \beta}(t_k).$$

2.5.3. The MPL estimator solution

Substituting the likelihood approximation (S51) and the prior

$$P_{\text{prior}}(\mathbf{s}) = \frac{1}{(2\pi\sigma^2)^{L(\ell-1)/2}} \exp \left(-\frac{1}{2\sigma^2} \mathbf{s}^T \mathbf{s} \right)$$

into the equivalent MAP problem

$$\hat{\mathbf{s}} = \arg \max_{\mathbf{s}} \left(\log \mathcal{L}(\mathbf{s} | \boldsymbol{\mu}, N, (\mathbf{x}(t_k))_{k=0}^K) + \log P_{\text{prior}}(\mathbf{s}) \right),$$

we take the vector derivative with respect to $\mathbf{s}$ and equate to zero. This yields the MPL estimate

$$\begin{aligned}\hat{s}_{i, \alpha} &= \sum_{j=1}^L \sum_{\beta=1}^{\ell-1} \left[\sum_{k=0}^{K-1} \Delta t_k C(\mathbf{x}(t_k)) + \gamma I \right]_{ij, \alpha \beta}^{-1} \\ &\quad \times \left[x_{j, \beta}(t_K) - x_{j, \beta}(t_0) - \sum_{k=1}^K \Delta t_k \left(\mu_{\ell \alpha} + \sum_{l=1}^{\ell-1} (\mu_{l \alpha} - \mu_{\ell \alpha}) x_{j, l}(t_k) - x_{j, \alpha}(t_k) \sum_{l=1}^{\ell} \mu_{\alpha l} \right) \right]\end{aligned}$$

for $i = 1, \dots, L$ and $\alpha = 1, \dots, \ell - 1$, which is the result quoted in (11) in Methods.

3. Data analysis

This section provides details about the synthetic data and its analysis used to obtain results presented in the paper.

3.1. Simulated data generation and performance analysis

The results in Figure 1, Supplementary Figures 1 and 2 were obtained for a 50-locus biallelic system. We generated a population of $N = 1000$ sequences, each of length $L = 50$, and evolved the population according to the WF model with selection and mutation. We assumed loci to be biallelic with 0 representing the WT and 1 representing the mutant allele. Parameters for the simulations are included in the captions of Supplementary Figures 1 and 2. Scripts for generating and analyzing these data are located in the GitHub repository.

We compared MPL with seven methods from the literature: WFABC (3), FIT (4), ApproxWF (5), LLS (6), CLEAR (7), EandR (8), and IM (9). We compared the accuracy of these methods by measuring the AUROC for classification of beneficial and deleterious selection coefficients, as well as the normalized root-mean-squared error (NRMSE) of the inferred selection coefficients,

$$\text{NRMSE} := \sqrt{\frac{\sum_{i=1}^L (\hat{s}_i - s_i)^2}{\sum_{i=1}^L s_i^2}}.$$

We also recorded the run time of all algorithms in the same computational environment. All values were averaged over samples from 100 WF simulations with the same underlying parameters.

3.2. Implementation and data preprocessing specifications

We wrote custom codes for the FIT and IM methods, while for the remaining methods we used software implementations provided by the respective authors. Scripts used to generate the comparison results in the paper are available at the GitHub repository.

Preprocessing: Some methods required the input sequence or allele frequency data to be modified in order to produce reasonable results. This *pre-processing* was performed according to one or two rules, as indicated below. In describing these, we define a *trajectory* as a set of mutant allele frequencies at consecutively sampled time points. The start of a trajectory is marked by the first observation of a polymorphism, while the end is marked by either the fixation or loss of the mutant allele, or by the last sampled time point. Thus, over an entire observation period, it is possible to observe multiple “trajectories” at the same locus.

1. *Pre-filtering:* Trajectories were filtered such that (a) the minimum length of a trajectory was two time points and (b) the maximum (minimum) frequency was at least 0.05 (0.95) for a trajectory that started from a frequency of zero (one).
2. *Bit-flipping:* The definitions of WT and mutant were reversed at loci for which trajectories had initial frequency greater than 0.95.

MPL and SL. MPL and its single locus variant, the independent model described in the main text, were implemented in C++. No preprocessing of the sequence data was required for MPL and SL.

WFABC. Trajectories were converted to WFABC’s format using custom Matlab scripts and then analyzed using the code provided at <http://jjensenlab.org/software> (accessed on September 16, 2017). Initial tests revealed that we had to pre-process the raw sampled trajectories before applying the WFABC method in order to get meaningful results. All trajectories were pre-filtered except the ones that were rendered monomorphic due to pre-filtering. Bit-flipping was also applied. These processing steps removed spurious noisy blips from the trajectories and redefined WT/mutant, resulting in improved performance of the WFABC method.

ApproxWF. Trajectories were converted to an input format compatible with ApproxWF using custom Matlab scripts and then analyzed using the code provided at <https://bitbucket.org/wegmannlab/approxwf/wiki/Home> (commit fcc7964 dated: 2016-10-04 accessed on September 26, 2017). All trajectories were pre-filtered as outlined above, except the ones that were rendered monomorphic due to filtering. Bit-flipping was not used.

IM. We developed a custom MATLAB script to implement the IM method (9). Trajectories were pre-filtered, but not flipped. We used the filtering threshold of 0.1 (0.9) as specified by the authors of this method. The method was implemented according to the description given in ref. (9), which applies a simulated annealing algorithm to estimate the selection coefficients. We used the simulated annealing algorithm implementation provided in the Global Optimization Toolbox of MATLAB 2017a. We let the algorithm run for a maximum of 100,000 iterations and stopped the algorithm when the average change in value of the objective function in the previous 10,000 iterations was less than a tolerance value (default value of 10^{-8}).

In all simulation runs, the optimization stopped in less than 100,000 iterations, indicating that the algorithm had converged. To further validate our in-house implementation of IM, we applied it to data sets generated using the system parameters as in ref.(9) and under the same simulation setup; e.g., continuous long runs of several hundred thousand generations. The results were qualitatively similar to those reported in ref. (9). We found that IM’s procedure of calling trajectories (described in ref. (9)),

was sensitive to the evolutionary parameters used, e.g., population size, mutation probability ($N = 10,000$ and $\mu = 5 \times 10^{-7}$ were used in ref. (9)), as well as to the simulation condition of not allowing a locus to mutate after it has reached fixation (in the simulation setup of IM, a locus that reaches fixation remains at fixation for the next 3200 generations). For the simulation results presented in this paper, where $N = 1000$, $\mu = 10^{-4}$ and there are no restrictions on mutations at a locus reaching fixation, we had to change the parameters of the procedure for calling trajectories so that the algorithm returned meaningful results. Using the default parameters of ref. (9) for calling trajectories resulted in the IM method missing the start/end of trajectories. The modified parameters are specified in Matlab scripts in the GitHub repository.

LLS. Trajectories were fed directly into the LLS code provided by the authors at <https://github.com/ThomasTaus/poolSeq>. A small number of loci resulted in a “N/A” selection coefficient estimate, which occurred when the loci had a frequency of zero or one for the vast majority of time points, and having a frequency close to zero or one for the other (small number of) time points. These loci were excluded from the analysis.

FIT. No additional preprocessing was required. The trajectories were analyzed using custom Matlab scripts.

CLEAR. Simulation data was converted into a format readable by CLEAR using custom Python scripts. No preprocessing of the data was necessary.

EandR. Simulation data was converted into a format readable by EandR using custom Python scripts. No preprocessing of the data was required.

Supplementary References

1. Risken, H. *The Fokker-Planck Equation: Methods of Solution and Applications* (Springer-Verlag, 1989), 2nd edn.
2. Lutkepohl, H. Handbook of matrices. *Computational Statistics and Data Analysis* **2**, 243 (1997).
3. Foll, M., Shim, H. & Jensen, J. D. WFABC: a Wright–Fisher ABC–based approach for inferring effective population sizes and selection coefficients from time-sampled data. *Molecular Ecology Resources* **15**, 87–98 (2015).
4. Feder, A. F., Kryazhimskiy, S. & Plotkin, J. B. Identifying signatures of selection in genetic time series. *Genetics* **196**, 509–522 (2014).
5. Ferrer-Admetlla, A., Leuenberger, C., Jensen, J. D. & Wegmann, D. An approximate Markov model for the Wright–Fisher diffusion and its application to time series data. *Genetics* **203**, 831–846 (2016).
6. Taus, T., Futschik, A. & Schlötterer, C. Quantifying selection with pool-seq time series data. *Molecular Biology and Evolution* **34**, 3023–3034 (2017).
7. Iranmehr, A., Akbari, A., Schlötterer, C. & Bafna, V. Clear: Composition of likelihoods for evolve and resequence experiments. *Genetics* **206**, 1011–1023 (2017).
8. Terhorst, J., Schlötterer, C. & Song, Y. S. Multi-locus analysis of genomic time series data from experimental evolution. *PLoS Genetics* **11**, e1005069 (2015).
9. Illingworth, C. J. & Mustonen, V. Distinguishing driver and passenger mutations in an evolutionary history categorized by interference. *Genetics* **189**, 989–1000 (2011).