

Learning sequence-function relationships with scalable, interpretable Gaussian processes

Reviewed Preprint

v1 • November 25, 2025

Not revised

Juannan Zhou, Carlos Martí-Gómez, Samantha Petti, David M McCandlish 

Department of Biology, University of Florida, Gainesville, United States • Simons Center for Quantitative Biology, Cold Spring Harbor Laboratory, Cold Spring Harbor, United States • Department of Mathematics, Tufts University, Medford, United States

 https://en.wikipedia.org/wiki/Open_access
 Copyright information

eLife Assessment

This **important** work introduces a family of interpretable Gaussian process models that allows us to learn and model sequence-function relationships in biomolecules. These models are applied to three recent empirical fitness landscapes, providing **convincing** evidence of their predictive power. The findings should be of interest to the community working on the sequence-function relationship, on epistasis, and on fitness landscapes.

<https://doi.org/10.7554/eLife.108964.1.sa3>

Abstract

Understanding the relationship between biological sequences, such as DNA, RNA or protein sequences, and their resulting phenotypes is one of the central goals of genetics. This task is complicated by epistasis, i.e., the context dependence of mutational effects. Advances in high-throughput phenotyping now make it possible to study these relationships at unprecedented scale, generating large datasets that measure phenotypes for tens or hundreds of thousands of sequences. However, standard regression models for analyzing such datasets often make unrealistic assumptions about the generalizability of mutational effects and epistatic coefficients across genetic backgrounds. Deep neural networks offer greater flexibility but suffer from limited interpretability and lack uncertainty quantification. Here, we introduce a family of interpretable Gaussian process models for sequence-function relationships that capture epistasis through flexible prior distributions that generalize classical theoretical models from the fitness landscape literature. In particular, these priors are parameterized by interpretable site-, allele-, and mutation-specific factors controlling the degree to which specific mutations decrease the predictability of the effects of other mutations. Using GPU acceleration to scale to large protein, RNA, and genome-wide SNP datasets, our models consistently deliver superior predictive performance while yielding interpretable parameters that both recover known features and uncover novel epistatic interactions. Overall, our methods provide new insights into the structure of the genotype-phenotype map and offer scalable, interpretable approaches for exploring complex genetic interactions across diverse biological systems.

Introduction

Understanding how genetic sequences (DNA, RNA, and proteins) encode biological function remains one of the major open challenges in biology. Deciphering sequence-function relationships is crucial for understanding and predicting evolution [1–8] as well as for uncovering the genetic factors underlying human genetic and infectious diseases and cancer [9–11]. It is also essential for designing and engineering sequences with desired properties, such as maximizing the expression or activity of a protein of interest [12–14] or breeding crops and livestock with improved yield or tolerance to stress and pathogens [15–18]. Recent advances in high-throughput phenotyping technologies, such as deep mutational scanning and massively parallel reporter assays, have led to the rapid proliferation of datasets measuring functionality across large libraries of sequence variants [19]. These approaches have been applied to proteins [9, 20–37], regulatory sequences [38–48] and even genome-wide genetic variation [49–52], providing valuable insights into protein function and gene regulation. However, building models that capture the structure of these empirical sequence-function relationships and accurately predict phenotypes for novel genotypes remains challenging, in part because mutational effects often change due to the presence of other mutations [3, 5, 7, 8, 53–56].

Existing methods for modeling sequence-function relationships often make strong assumptions about the structure of this background dependence of mutational effects. For example, an additive model assumes that the effects of mutations are constant, such that the observed mutational effects in the data can be generalized to any novel background [57]. Pairwise interaction models relax this assumption by allowing mutational effects to vary, but still assume that the epistatic interaction between any pair of mutations is itself background-independent [58]. Perhaps surprisingly, this background independence of double-mutant epistatic coefficients implies that for any pairwise interaction model, the mean predicted phenotype changes quadratically as a function of the number of mutations from the wild type (or any arbitrarily chosen focal sequence)—a strong restriction on the types of geometric features such models can represent [55]. As a result, simple regression models often lack the flexibility to account for the full complexity of epistasis in empirical sequence-function relationships [55] and are frequently outperformed in prediction tasks by more expressive methods, such as neural networks [34, 35, 46, 59–66]. However, while neural networks are able to make more accurate predictions than traditional regression models, they provide limited capabilities for interpretability and uncertainty quantification, which reduces their utility.

Gaussian process models offer an alternative class of highly flexible models [55, 58, 67–74] that provide better interpretability and uncertainty quantification. These models work by defining a prior probability distribution over the space of possible sequence-function relationships and then computing the posterior distribution given the data [75], where in particular the prior distribution can be defined by interpretable parameters that control the predictability of mutational effects across different genetic backgrounds [55] and the posterior distribution provides a natural means to quantify uncertainty. While our previously described prior distributions for sequence-function relationships are isotropic, this is, they assume that every mutation has the same effect on the predictability of other mutations [55, 58, 70], empirical evidence suggests substantial heterogeneity. Different mutations can vary widely in how they affect predictability at other positions [28, 76] and in the total fraction of genetic variance explained by their interactions [77], with key mutations dramatically altering the effects of mutations elsewhere in the sequence [30, 77–79].

In this paper, we introduce a new family of anisotropic Gaussian process priors that capture heterogeneity in how mutations influence the predictability of other mutations based on two key principles. First, each position, allele, or mutation is assigned a decay factor that controls how much it affects the predictability of mutations at other positions. Second, when several mutations occur together, their effects on the predictability of other mutations are assumed to combine multiplicatively. These design choices enable us to generalize previously proposed theoretical models of fitness landscapes in several key directions [55, 77, 80], particularly by allowing us to generate families of complex random fitness landscapes using a moderately sized set of interpretable parameters, each tied directly to the behavior of individual mutations. We combine these new priors with recent advances in Gaussian process inference and modern GPU infrastructure [81–83] to analyze diverse high-throughput phenotyping experiments, including nearly complete genotype-phenotype maps for protein GB1 [84] and human 5' splice sites [42], a dataset for the longer AAV2 capsid protein [35], and a genome-wide dataset for yeast [50]. We then use our Gaussian process framework for a range of data-analytic tasks, including: (i) identifying sites, alleles, and mutations with strong effects on the predictability of other mutations, (ii) making phenotypic predictions for specific genotypes of interest and quantifying the uncertainty in these predictions, (iii) predicting the effects of all single point mutations across different genetic backgrounds and assessing our confidence in whether those effects have changed, and (iv) reconstructing combinatorially complete landscapes for moderately sized subsets of mutations. Overall, by studying these diverse systems, we find that mutations exhibit highly heterogeneous effects on the predictability of other mutations. By modeling this heterogeneity explicitly, we capture the background dependence of mutations with greater granularity and precision, thereby substantially enhancing the predictive power of our models.

Results

Gaussian processes for learning sequence-function relationships

To model sequence-function relationships, we use Gaussian process regression, which places a Gaussian prior distribution over the space of functions mapping genotypes to their phenotypes. This prior is fully specified by its mean and a covariance function, or kernel, $k(x, x')$, which defines the covariance of phenotypes between each possible pair of sequences x and x' under the prior. Gaussian process regression makes predictions based on the data y by computing the posterior mean for any genotype x

$$\mathbb{E}[f(x)|y] = k_x^T (K_{BB} + E)^{-1} y, \quad (1)$$

where K_{BB} and E are square matrices modeling the prior covariance and noise variance for the set of measured genotypes B , and k_x is the vector encoding the covariance between the genotype x whose phenotype we wish to predict and the set B of measured genotypes. A key advantage of Gaussian process regression is that it also provides a full characterization of the prediction uncertainty via the posterior covariance matrix (see Supplement 1.1).

Because Eq. 1 is a generic result applicable to all Gaussian process regression problems [75], the key determinant of model behavior and performance is the choice of the prior defined by the kernel function. In previous work, we introduced priors parameterized by biologically interpretable quantities such as the average squared epistatic coefficient [58, 70] or the variance associated with genetic interactions of each possible order [55]. These priors also control how the predictability of mutational effects decays with genetic distance [55].

In this section, we review our previously developed method, empirical variance component regression (VC regression) [55], propose three new methods that model how specific mutations influence the predictability of mutational effects at other sites, and analyze the flexibility in the set

of expected variance components that can be under each prior.

Empirical variance component regression

In [55], we introduced a family of flexible priors in which the kernel function for a pair of genotypes x and x' separated by d mutations is given by

$$k(x, x') = k(d) = \sum_{k=1}^{\ell} \lambda_k W_k(d), \quad (2)$$

where $W_k(d)$ is a kernel that captures the covariance structure of pure k -th order interactions (given by a family of orthogonal polynomials known as the Krawtchouk polynomials, see [55], [85] for details) and $\lambda_k \geq 0$ is the expected variance of k -th order interactions under the prior. Importantly, this family of priors coincides with a well-studied class of theoretical random fitness landscapes where the covariance in phenotypes and mutational effects is completely determined by the Hamming distance between genotypes or genetic backgrounds [85–88]. These isotropic random landscapes are parameterized by the expected variance components of additive effects, pairwise, and each possible type of higher-order interaction between mutations. These variance components then control how fast the correlation in mutational effects between genetic backgrounds decays with the genetic distance between them.

Traditionally, the scalability of Gaussian process regression has been limited to fitting tens of thousands of training data points because computation of Eq. 1 and hyperparameter optimization scale cubically with dataset size [75, 81]. In the original proposal of VC regression [55], we inferred the model hyperparameters from the data by minimizing the discrepancy between the prior covariance and the empirical phenotypic covariance using weighted least squares [89] and exploited a sparse representation of the kernel matrix for the full sequence space to compute the posterior mean and covariances via iterative methods [55, 74]. This strategy allowed us to scale our methods to sequence spaces with up to 10 million sequences, equivalent to 12 nucleotides for DNAs/RNAs, 24 biallelic loci, or 5 sites for proteins.

Here, we alleviate these constraints on sequence length and alphabet size by leveraging recent developments in Gaussian process inference techniques and modern GPU computation using GPyTorch [81] and KeOps [83] to instead work with the dense matrix K_{BB} (Eq. 1) whose dimensions depend on the size of the training data but not on the size of sequence space. These modern machine learning libraries implement algorithms that combine the power of GPUs and iterative methods to compute posterior means and covariances more efficiently, allowing computations for datasets with up to roughly 1 million observations [82]. Moreover, these techniques also enable a more principled approach for inferring hyperparameters, such as the λ_k , via evidence maximization [75]. In other words, we find the combination of kernel hyperparameter values that maximize the likelihood of observing the empirical data (see Supplement 1.1). These improvements make our methods scalable to large datasets without constraints on the size of sequence space.

Connectedness model

By tuning the relative contributions of different orders of epistatic interaction, VC regression allows priors that can reproduce any distance-dependent correlation function for mutational effects. However, the isotropic assumption that all mutations have equal effects on the predictability of other mutations still limits the expressivity of these priors. In this section, we introduce our first new model, which relaxes this isotropic assumption by allowing mutations at each site to exert a site-specific characteristic effect on the predictability of the phenotypic effects of mutations at other sites.

To construct the model, we introduce a site-specific decay factor δ_p for each of the ℓ positions ($p = 1, \dots, \ell$). Biologically, δ_p captures how much a mutation at site p disrupts the predictability of mutational effects across genetic backgrounds. For example, if two genotypes differ only at position p , then the correlation under the prior in the effects of mutations at all other positions is precisely $1 - \delta_p$. A small δ_p suggests that mutations at p have minimal effects on the effects other mutations, while a large δ_p indicates strong epistatic influence. To generalize this idea to genetic backgrounds differing at multiple sites, we assume that the effects of different mutations on the predictability of other mutations combine multiplicatively. That is, mutating an additional position q further reduces the original correlation by a factor of $1 - \delta_q$, producing an expected correlation of $(1 - \delta_p)(1 - \delta_q)$. This results in a biologically interpretable kernel function, where the covariance between two sequences decreases according to the cumulative effect of their mutational differences:

$$k(x, x') = \sigma^2 \prod_{p: x_p \neq x'_p} (1 - \delta_p), \quad (3)$$

where σ^2 represents the shared variance across all genotypes. Importantly, the decay factors δ_p not only control the phenotypic correlation between different sequences, but also the correlation in mutational effects in genetic backgrounds that differ at specific positions p , given by $p(1 - \delta_p)$ (see Supplement 1.2).

Our prior is closely related to the ‘connectedness model’ [77], a statistical model of Gaussian fitness landscapes first proposed for studying evolution of the distribution of fitness effects for adapting populations. The connectedness model is parameterized by ℓ site-specific probabilities $0 \leq \mu_p \leq 1$, corresponding to the tendency for each site p to be involved in epistatic interactions. In particular, the strength of interactions among any set of positions is proportional to the product of their corresponding μ_p factors.

In Supplement 1.4, we show that Eq. 3 provides a slight generalization of the connectedness model and that the two site-specific factors μ_p and δ_p are related through the following equation:

$$\delta_p = \frac{\alpha \mu_p}{1 + (\alpha - 1) \mu_p}. \quad (4)$$

Interestingly, we show that our connectedness kernel in Eq. 3 remains a valid (i.e. positive-definite) kernel for some $\delta_p > 1$, which would correspond to probabilities $\mu_p > 1$ in the original connectedness model. In such cases, introducing a mutation with a decay factor δ_p larger than one leads to anticorrelated phenotypes and mutational effects. This added flexibility enables the model to capture a broader range of epistatic patterns than the original connectedness model, although empirical support for this so-called egg box [28, 76] correlation structure remains limited. In addition, while the original connectedness model was proposed for biallelic genotypic spaces, our connectedness kernel is compatible with an arbitrary number of possible alleles per site.

The relationship between the connectedness kernel and the VC regression kernel can be clarified by examining how each generalizes a simple kernel where the correlation between genotypes decays geometrically with Hamming distance $d(x, x')$, i.e. $k(x, x') = \sigma^2 \beta^{-d(x, x')}$. This geometric decay kernel includes the exponential kernel as a special case ($\beta > 0$). The exponential kernel itself arises when applying the standard radial basis function (RBF) kernel to one-hot encoded sequences [67, 87], when the magnitude of epistatic interactions decays exponentially with interaction order (i.e. $\lambda_k = e^{-b \times k}$) [55, 73, 87] and as the diffusion kernel on the Hamming graph [69, 90]. Our VC regression kernel can then be viewed as a generalization of the geometric kernel to allow the λ_k to deviate from this strict multiplicative scaling while retaining the isotropic property of the geometric kernel. In contrast, the connectedness kernel offers a different generalization: whereas the geometric decay kernel treats all mutations the same ($\beta = (1 - \delta)$ and $\delta_p = \delta$), the

connectedness kernel introduces a detailed dependence on the mutated sites through site-specific decay factors (or equivalently, site-specific epistatic strengths). Interestingly, both the VC regression and the connectedness priors have the same number of hyperparameters so that comparing their performance provides an indication as to the relative importance of precisely matching the contribution of higher-order epistasis versus incorporating site-to-site heterogeneity in the extent of epistasis.

Jenga model

The connectedness model introduces flexibility to the prior by allowing the extent to which mutational effects generalize across genetic backgrounds to depend on which specific sites are mutated. One limitation of this approach is that it treats all mutations at a given site equally. In practice, however, some alleles may be more likely than others to participate in genetic interactions and to influence the predictability of mutations at other sites. In this section, we generalize the connectedness model by allowing the decay in the correlation of mutational effects to depend not only on the mutated sites, but also on which specific alleles have been altered.

In the connectedness model, a mutation at site p reduces the correlation in mutational effects at other sites by a factor of $1 - \delta_p$. Here, instead of assigning a uniform decay factor δ_p to all mutations at site p , we assign an allele-specific decay factor δ^a , such that mutating allele a to a new allele a' results in an overall reduction in correlation by $100 \times (1 - \delta^a)(1 - \delta^{a'})$ percent. Assuming that the effects of mutations on phenotypic correlations across positions combine multiplicatively in this manner, we obtain a new distribution over the space of all sequence-function relationships, with covariance characterized by the following kernel function:

$$k(x, x') = \sigma^2 \prod_{p: x_p \neq x'_p} (1 - \delta_p^a)(1 - \delta_p^{a'}). \quad (5)$$

Importantly, this expression also generalizes to describe the decay in the predictability of mutational effects, such that the correlation in these local effects is $(1 - \delta_a)(1 - \delta_{a'})$ for all positions p segregating between a and a' . It is easy to see that this model contains the connectedness model as a special case when all alleles at a given position have a common decay factor. Furthermore, in the special case $\delta^a = 1$, according to Eq. 5, mutating allele a results in zero correlation in phenotypes or mutational effects. Indeed, if all alleles in Eq. 5 have decay factors equal to 1, we recover the classical ‘House-of-Cards’ (HoC) model [80], in which mutating any allele completely eliminates predictability in phenotypes or mutational effects—akin to how removing any component from a house made of playing cards causes the entire structure to collapse. Thus, by allowing alleles to have a continuum of effects on the overall correlational structure, our model can be viewed as a generalization of the HoC model where different alleles at a site have different effects on the correlation structure. We therefore refer to this new model of random fitness landscapes as the ‘Jenga model’, reflecting the idea that some mutations may have negligible effects, while others can disrupt the entire correlational structure, much like how specific blocks in the game Jenga can either be easily removed or cause the whole tower to fall.

The Jenga kernel is also closely related to the classical automatic relevance determination (ARD) kernel [75, 91] applied to the one-hot encoding of genotypes (see Supplement 1.5). However, one important difference is that the allele-specific decay factors δ^a in the Jenga model are allowed to exceed one under conditions that ensure the positive definiteness of the kernel on sequence space as a whole (see Supplement 1.6). As a result, whereas the ARD kernel requires all correlations to be strictly positive, the Jenga kernel permits negative correlations between phenotypes and mutational effects across adjacent genetic backgrounds, thereby offering the flexibility to model eggbox-like patterns that could, in principle, be induced by mutations at some positions.

General product model

In the Jenga model, the effect of a mutation on the correlation between genotypes is parameterized by multiplying the two allele-specific factors, meaning that each allele makes the same contribution to the decay in correlations for all mutations involving that allele. In this section, we relax this assumption to further generalize the Jenga model. We do so by assigning a mutation-specific decay factor $\delta^{a,a'}$ to each pair of alleles a, a' , such that mutating from a to a' (or vice versa) reduces the overall correlation by a factor of $(1 - \delta^{a,a'})$.

This yields the following general product kernel:

$$k(x, x') = \sigma^2 \prod_{p: x_p \neq x'_p} (1 - \delta_p^{a,a'}). \quad (6)$$

Similar to the Jenga model and connectedness model, the correlation in mutational effects is given by the product of the mutation-specific decay factors for sites segregating between the two genetic backgrounds: IT $(1 - \delta_{a,a'})$.

It is easy to see that [Eq. 6](#) provides greater flexibility than the Jenga model, as the general product kernel has α hyperparameters per site, whereas the Jenga model has only α hyperparameters. In fact, [Eq. 6](#) represents the most general form of a homoscedastic product kernel such that any product kernel with uniform variance can be expressed in this form and can be guaranteed to remain positive definite through an appropriate parameterization [\[92\]](#).

Variance components of the connectedness, Jenga, and general product kernels

We can understand certain aspects of the geometry of random fitness landscapes such as the connectedness model, the Jenga model, and the general product model by quantifying how much of the overall phenotypic variance is expected to be explained by epistatic interactions of different orders under these priors. While VC priors are directly parameterized in terms of these expected variance components, our new random fitness landscape models are parameterized by decay factors, which have a less straightforward relationship to variance components.

In the Supplement 1.7, using a general result for deriving variance components from arbitrary product kernels, we show that all three models are capable of capturing epistatic interactions of every possible order between any number of positions. Furthermore, we find that a single decay factor affects all expected variance components, such that increasing any decay factor systematically shifts the variance toward higher-order components.

We also find that there are certain combinations of expected variance components that cannot be obtained by tuning the decay factors of our three new priors, for example, priors with interactions of a single order alone (e.g., additive, pairwise). This behavior is more limited than in the VC kernel, which can match any valid combination of expected variance components. Furthermore, we show that the set of expected variance components under a prior from the connectedness model is in fact the same as that of based on the general product model (see Supplement 1.7, Theorem 3). Therefore, although the connectedness model has the fewest hyperparameters, it is just as expressive as the general product and Jenga models in terms of realizable combinations of expected variance components when used as prior distributions. Finally, although these new priors can express only limited set of variance components in expectation, it is important to realize that the posteriors can still capture arbitrary patterns of epistasis given sufficient data, since these priors produce a positive density for any possible landscape.

Applications to experimental genotype-phenotype datasets

In the previous section, we introduced new models of random sequence-function relationships in which the decay in the predictability of mutational effects between genetic backgrounds is determined by the positions, alleles, or mutations at which sequences differ. Importantly, the connectedness model, the Jenga model, and the general product model form a hierarchical family, with each representing a generalization of the preceding model. These kernels equip us with the flexibility to model complex epistasis in the data while allowing for a balance between model expressivity, interpretability, and computational efficiency.

In this section, we implement these models as prior distributions in a Gaussian process framework for modeling empirical genotype-phenotype datasets. We infer the site-, allele-, or mutation-specific decay factors from the data by fitting the kernels in [Eq. 3](#), [Eq. 5](#), and [Eq. 6](#) through evidence maximization, and we use them for two main purposes. First, the inferred decay factors provide important broad-scale insights into the structure of epistasis in the empirical sequence-function relationship by highlighting mutations that reduce the predictability of other mutations. Second, we use the inferred decay rates as hyperparameters of a prior distribution for Gaussian process regression to predict phenotypes for unobserved genotypes and to assess our confidence in these predictions.

Application to protein GB1

We first applied our methods to the protein GB1 dataset [\[84\]](#), which contains relative enrichment values of nearly all $20^4 = 160,000$ genotypes across four highly epistatic amino acid sites on the IgG-binding domain of streptococcal protein G, where selection was imposed for IgG-binding.

To assess the predictive performance of our methods, we first fit several standard models for studying sequence-function relationships to this dataset, including a simple additive model, a pairwise interaction model, and a global epistasis model, which maps the sequence first to an underlying additive phenotype and then onto the observation scale through a nonlinear monotonic function [\[61, 94, 95\]](#). We also fit Gaussian process regression models with the VC kernel and the geometric kernel, with the model hyperparameters optimized via evidence maximization. The out-of-sample performance of these models under different proportions of training data is summarized in [Figure 1A](#). We first found that the additive model achieves an R^2 of approximately 0.5 on the test data, and that adding a global nonlinear function to the additive phenotype increased the test R^2 to slightly over 0.65. Adding pairwise interaction terms to the additive model without a global epistasis nonlinearity resulted in a more substantial improvement, with test R^2 values close to 0.75 when the training data proportion exceeded 10%. Next, we found that Gaussian process regression using the VC kernel and geometric kernel show similar performance, achieving a test R^2 over 0.8 at high data density, corresponding to a 5% improvement over the pairwise model.

With the performance of these baseline methods established, we proceeded to fit our three new methods to quantify site-, allele-, or mutation-specific effects on the predictability of other mutations and examine how accounting for heterogeneity at various levels of the correlational structure improves predictive performance. First, we found that the connectedness model indeed reveals heterogeneous effects across the four amino acid positions. Specifically, we found the site-specific decay factors inferred using evidence maximization showed a nearly seven-fold difference ($\delta_{39} = 0.23$, $\delta_{40} = 0.08$, $\delta_{41} = 0.54$, $\delta_{54} = 0.38$; [Figure 1B](#)). This suggests that mutating sites 41 and 54 have the strongest effect on the predictability of phenotypes and mutational effects, such that the same mutations have an average correlation of 0.53 across backgrounds segregating only at site 41 or 54, and an average correlation of 0.28 when both sites segregate (see [Eq. 3](#)). In contrast, mutations at site 40 have a much smaller effect on the predictability of other mutations,

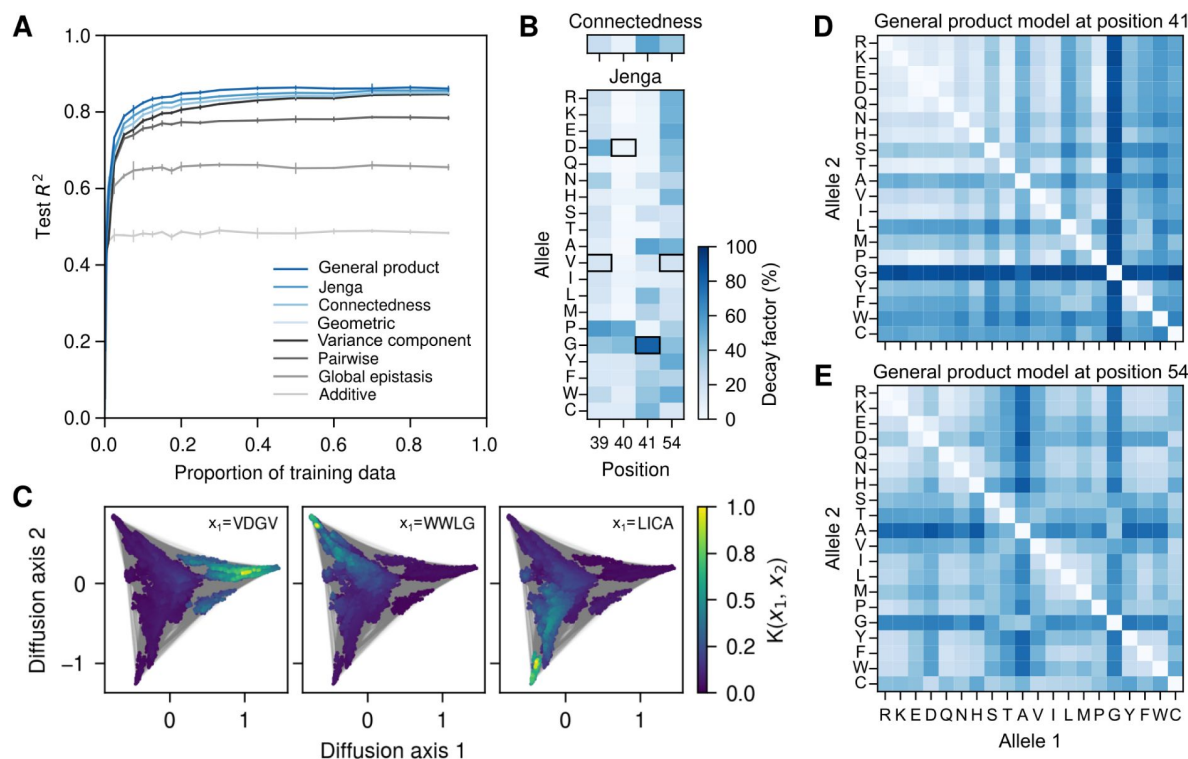


Figure 1

Application to the genotype-phenotype dataset for protein GB1.

(A) Model predictive performance, measured by the R^2 on held-out data, as a function of the fraction of data used for training. Error bars represent one standard deviation across three independent subsets of sequences used for training at each proportion. (B) Inferred position- and allele-specific decay factors under the connectedness model (top) and Jenga model (bottom), respectively. Black squares highlight the allele of the wild-type sequence at each position. (C) The Jenga prior is able to capture information about the three main fitness peaks of the GB1 landscape [58]. In each panel, dots represent genotypes and lines denote single point mutations between genotypes. Squared distances between dots optimally approximate the commute times between genotypes under a weak-mutation evolutionary model under selection for high phenotypic values [74, 93] (see Methods). Genotypes are colored by their correlations with three genotypes VDG, WWLG and LICA, each representing a local fitness peak. Correlations were calculated using the Jenga kernel with hyperparameters inferred using evidence maximization. The complete GB1 fitness landscape for the evolutionary model was constructed using the maximum a posteriori estimate under the Jenga model. (D, E) Inferred mutation-specific decay factors under the general product model for positions 41 and 54. These matrices show great heterogeneity in the influence of mutations at a given site on the predictability of mutational effects at other sites both in terms of the degree of influence conferred by mutations between different pairs of alleles and in the effect of mutations between the same pair of alleles at different sites.

such that for genetic backgrounds segregating only at site 40 the same mutations have an average correlation of 0.92. This rank order between sites is consistent with previous analysis showing that higher-order interactions in this dataset are enriched among site 39, 41 and 54, due to steric interactions [23, 84]. We found that using these decay factors to parameterize the connectedness kernel for Gaussian process regression yields a slight improvement in predictive accuracy over both VC regression and geometric kernel regression (**Figure 1A**).

We next fit the Jenga model to the data and examined how specific alleles at each position affect the correlation in phenotype and mutational effects. We first present the inferred allele-specific decay factors as a heatmap in **Figure 1B**. The decay factors exhibit substantial heterogeneity across alleles. For example, at position 40, mutations involving Gly and Pro result in more than a 50% reduction in correlation, whereas mutations involving other amino acids have minimal or negligible effects. The Jenga model heatmap is also consistent with the aforementioned steric interaction model [23, 84]. This is most evident at position 41, where amino acids with the highest decay factors tend to have either small side chains (e.g., Gly and Ala) or bulky ones (e.g., Leu, Trp, and Phe). These mutations likely cause the largest disruption to the optimal pattern of steric interactions, possibly resulting in nonfunctional proteins and consequently a breakdown in the correlation of mutational effects. Finally, we observed a modest improvement in R^2 for the Jenga model over the connectedness model, which was comparable in magnitude to the improvement of the connectedness model over VC regression (**Figure 1A**).

To gain some intuition as to how Jenga regression works, we examined the qualitative structure of the Jenga kernel using a visualization technique to understand how the correlations incorporated into the Jenga kernel relate to the structure of the genotype-phenotype map. In **Figure 1C**, we first represented all 160,000 genotypes in the GB1 fitness landscape using a 2-dimensional embedding generated using a method similar to diffusion maps [74, 93]. The method is based on modeling how a population would evolve across this fitness landscape, and results in a triangular structure with three high-fitness clusters, i.e. fitness peaks, occupying the corners and separated by a central mass of low-fitness genotypes. Importantly, we previously found that the local structure within each cluster is mainly additive such that the effects of mutations are well-correlated across backgrounds within clusters but not between clusters [58, 96]. We next chose a focal sequence from each of the three clusters and colored each sequence by its correlation with this focal sequence under the Jenga prior (**Figure 1C**). We observe that all three high-fitness genotypes exhibit strong correlation with neighboring genotypes within the same cluster, and that this correlation decays to near zero outside each cluster. This indicates that the model predominantly uses mutational effects observed within the same cluster to make predictions for a focal genotype, and thereby in essence is fitting multiple independent models across different regions of the fitness landscape. For comparison, we also present the same visualization for all models fit to the GB1 data in Figure S1. Unlike the highly localized correlation structure of the Jenga prior, the isotropic priors (which only depend on the number of mutations separating two sequences and not their identity, e.g. the VC and geometric priors) show strong to moderate correlation between each focal genotype and genotypes across the whole landscape. The key reason for this qualitative difference is that the geometric kernel can only capture the typical or average effect of mutations on the predictability of other mutations. Thus, two genotypes that are mutationally adjacent will still have a high correlation even if the mutation that separates them is particularly epistatic. In contrast, the Jenga model's ability to account for variation in epistatic effects across alleles enables it to model the lack of correlation between adjacent genotypes when separated by strong-effect mutations, while allowing distant genotypes to maintain high correlation if separated only by mutations with small decay factors. As a result, the Jenga kernel more closely reflects the structure of the GB1 landscape, which likely contributes to its improved predictive accuracy.

Finally, we fit our most flexible model, i.e., the general product model. Unlike the Jenga model, which assigns decay factors to individual alleles at each position, the general product model directly characterizes the effects of all mutations. This added flexibility yields another modest improvement in predictive accuracy, comparable in magnitude to the improvement achieved by the Jenga model over VC regression (**Figure 1A**). To investigate the cause of this improvement, we visualized the mutation-specific decay factors as 20×20 matrices. We first examined the matrices for positions 41 and 54 in **Figures 1D** and **1E**, and then considered side-by-side comparisons of the matrices inferred by the general product model with those constructed by multiplying the corresponding allele-specific decay factors from the Jenga model (Figure S2). We found that the effects of some mutations in the general product prior can sometimes be well approximated by the Jenga prior; for example, mutating Gly at position 41 uniformly induces strong epistasis in both the Jenga model and the general product priors. However, there are also mutations with mutation specific decay factors that deviate from the Jenga prior's expectation. For example, according to **Figure 1E**, the three aromatic residues Tyr, Phe, and Trp at position 54 are largely interchangeable, but under the Jenga prior, mutations among these residues are expected to have moderately large epistatic effects on other mutations (Figure S2). This difference arises because the Jenga prior can only encode the average epistatic effects of each allele, whereas the actual effects of mutations at position 54 have a more subtle dependence on the physical properties of the specific alleles involved.

Application to human 5' splice site

Having demonstrated the performance of our methods on the protein GB1 dataset, we now apply our methods to modeling an RNA sequence-function relationship. The experimental data consists of the activity of nearly all variants covering 8 nucleotides on the 5' splice site of exon 7 for the survival of motor neuron 1 (SMN1) gene in a minigene construct [42], with splicing activity measured as the proportion of correctly spliced transcripts (percent spliced in, PSI).

To generate performance baselines, we first fit the additive, pairwise epistatic, and global epistasis models to the data. In **Figure 2A**, we can see that the additive model exhibits poor predictive accuracy, with test R^2 around 0.15. The pairwise model provides a substantial improvement, achieving an R^2 of 0.45 when the training data proportion exceeds 20%. However, the pairwise model is outperformed by the global epistasis model, which typically achieved an R^2 of approximately 0.6. This rank order is consistent with our previous understanding of the SMN1 splicing landscape, which at a coarse scale can be approximated as a single-peaked fitness landscape resulting from an additive fitness landscape transformed by a steep sigmoidal nonlinearity on the measurement scale (Figure S3A) [55].

The limited predictive performance of these basic models suggest that the data contain a more complex pattern of genetic interactions. In fact, the empirical distance-correlation function (**Figure 2B**, black dots) shows that correlation between pairs of sequences in the data cannot be accurately modeled by a linear or quadratic function as expected under an additive or pairwise interaction model [55]. Instead, each mutation, in average, seemed to decrease the correlation by about half, as expected under a geometric kernel. In line with this, we found that Gaussian process regression with the VC kernel and the geometric kernel showed similar performance throughout the range of training data density, with test R^2 increasing nearly linearly once the training data proportion exceeded 20%, and surpassing 0.8 with dense training data. This improvement over the global epistasis model can be attributed to their ability to fit higher-order epistasis, which can provide a good approximation of the global nonlinearity while also capturing mutation-specific interaction patterns, often referred to as specific epistasis [97], that the global epistasis model cannot capture.

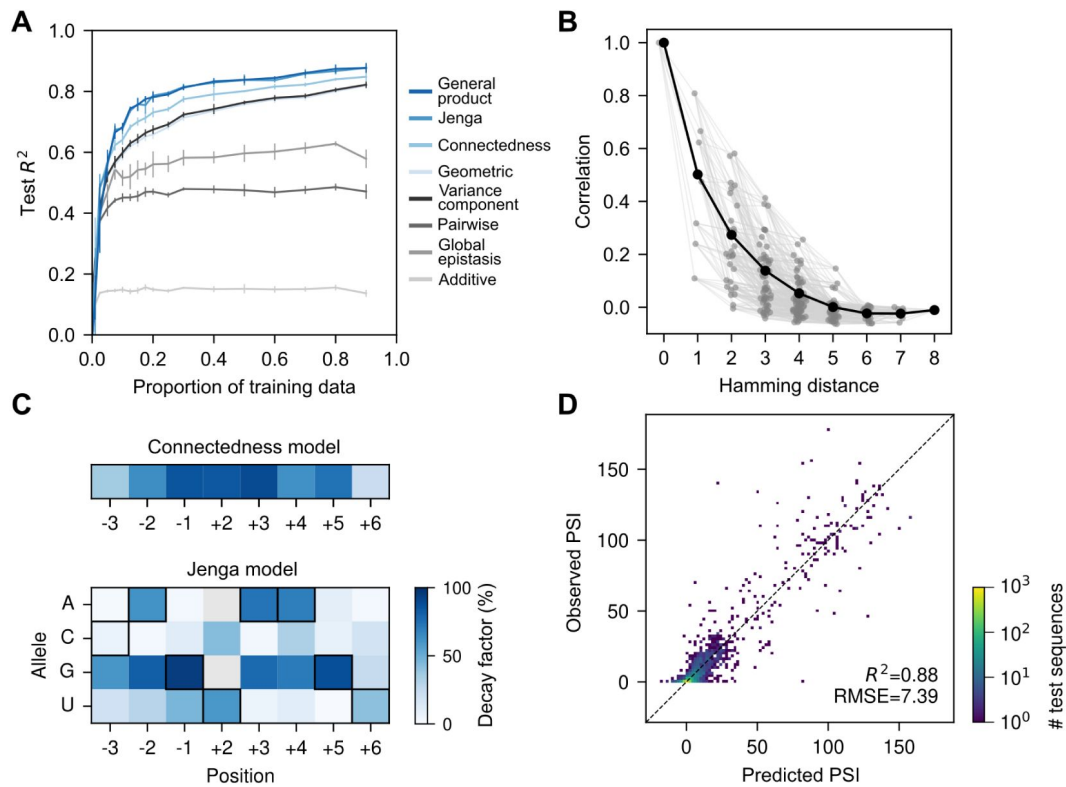


Figure 2

Application to the high-throughput dataset for 5' splice site of SMN1 exon 7.

(A) Model predictive performance, measured by the R^2 on held-out data, as a function of the fraction of data used for training. Error bars represent one standard deviation across three independent subsets of sequences used for training at each proportion. (B) Empirical correlations between pairs of sequences. Black dots show the average correlation for all pairs of sequences at a given Hamming distance. Gray dots show the correlation between genotypes segregating at specific sets of positions for each distance class. Values on the x-axis were jittered to facilitate visualization. (C) Inferred position- and allele-specific decay factors under the connectedness model (top) and Jenga model (bottom), respectively. Black squares highlight the canonical 5' splice site nucleotides complementary to the U1 snRNA template. (D) Two-dimensional histogram comparing the observed Percent Spliced In (PSI) values against predictions of the Jenga model on test 5' splice sites, comprising 10% of the measured genotypes.

However, when stratifying the empirical distance-correlation function depending on the sites at which sequences differ (**Figure 2B** [↗](#), gray dots), we can see that the empirical correlations exhibit large variation within each distance class, to the extent that two genotypes differing at up to five positions can be more correlated than pairs separated by only one position. This observation suggests that mutations at different sites in the SMN1 dataset have very heterogeneous effects on the predictability of other mutations. Consistent with this observation, the connectedness model yielded an improvement in predictive performance, typically increasing test R^2 by 3–4% over the VC and geometric kernels (**Figure 2A** [↗](#)). In **Figure 2C** [↗](#), we show the decay factors of the connectedness prior inferred using evidence maximization. We can see that some positions when mutated can reduce the correlation under the prior by over 80% (e.g. $\delta_{-1} = 0.86$, $\delta_{+2} = 0.85$, $\delta_{+3} = 0.89$), while others have more moderate effects on the correlation ($\delta_{-3} = 0.36$, $\delta_{+6} = 0.23$). This allows the connectedness prior to better capture the empirical correlation structure observed in **Figure 2B** [↗](#) than the VC regression prior, which depends on mutational distance but not the specific sites at which sequences differ.

Next, we turn to the Jenga model to examine how changing specific alleles affects the predictability of other mutations. We find that the allele-specific decay factors of the Jenga model inferred using evidence maximization show substantial variation among alleles within each site (**Figure 2C** [↗](#)). Comparing the decay factors of the connectedness model and the Jenga model, we find that the former provides a coarse approximation of the latter, such that sites with large allele-specific decay factors also tend to have large site-specific decay factors under the inferred connectedness prior. On the other hand, the additional flexibility of the Jenga model is reflected in a roughly 3% improvement in R^2 over the connectedness model at most training data proportions.

Interestingly, we find that the alleles with the largest decay factors under the Jenga model often coincide with the most preferred nucleotides under the global epistasis model (Figure S3B), which infers the presence of a sharp, threshold-like effect, consistent with the biophysical intuition that PSI can be expressed as a sigmoid function of the binding affinity between the 5' splice site and the U1 snRNA [43 [↗](#), 55 [↗](#), 98 [↗](#)]. Because mutating critical nucleotides (e.g., U at position +2) typically results in nonfunctional splice sites, subsequent mutations have little to no effect since the splice site has already been rendered inactive. Consequently, these critical alleles are inferred to have large decay factors in the Jenga model, as mutational effects in backgrounds with and without the allele tend to be poorly correlated. This ability to incorporate different phenotypic correlations for different mutations appears to result in a qualitatively better fit of the global nonlinearity, so that while VC regression produces a systematic nonlinear pattern of residuals (Figure S3C), inference under the Jenga model produces a pattern of residuals for out-of-sample predictions centered on the observed values (**Figure 2D** [↗](#)). Overall, these results show that beyond capturing specific epistatic interactions, our new models continue to perform well in the presence of a strong global nonlinearity.

Finally, we applied our most flexible model, general product regression, to the data. Unlike in the protein GB1 example, allowing mutation-specific decay factors for SMN1 does not improve predictive performance relative to the Jenga model (**Figure 2A** [↗](#)). Examination of the inferred mutation-specific decay factors shows strong agreement with those implied by the Jenga model (Figure S3D), which accounts for their nearly identical performance. A contributing factor to the differing behavior of these models across datasets may be the alphabet size, as the difference in parameter count and expressivity between the Jenga and general product models is much smaller in nucleotide space (4 alleles versus 6 allele pairs) than in amino acid space (20 alleles versus 190 allele pairs).

Application to AAV2 capsid protein

So far, we have applied our methods to relatively small sequence spaces. To demonstrate their utility in larger sequence spaces, we now apply them to a sequence–function dataset for the capsid protein of adeno-associated virus 2 (AAV2). This dataset encompasses 42,328 variants targeting 28

amino acid sites (positions 561–588), spanning buried, surface, and interface regions [35, 99], with corresponding deep mutational scanning (DMS) scores measuring viral production efficiency. Importantly, the sequence space for this example contains $20^{28} \approx 2.7 \times 10^{36}$ sequences, far larger than could be accommodated by our previous approaches [55, 58, 70, 74] which could only handle sequence spaces containing low millions of sequences.

As in previous sections, we first fit the baseline additive and global epistasis models across increasing proportions of training data (Figure 3A). Both models exhibit strong predictive performance, achieving test R^2 values close to 0.6 and 0.8, respectively, as the training data proportion increases. The pairwise interaction model performs comparably to the global epistasis model at higher sampling densities but is inferior at lower sampling densities. We next fit the geometric, connectedness, Jenga, and general product models to the AAV2 dataset, selecting model hyperparameters via evidence maximization. We did not explicitly fit the VC regression to this dataset because the dataset lacks pairs of sequences at all relevant distances, preventing reliable estimation of the empirical autocorrelation function. The more expressive priors (Jenga and general product models) consistently outperformed the geometric and connectedness priors when trained with over 30% of the data. In contrast, the simpler priors performed better when less training data was available, a pattern that differs from our earlier results using nearly complete combinatorial data for short sequences. This difference is expected, as this dataset samples only a small fraction of the possible sequence space, and the number of hyperparameters to learn varies by an order of magnitude across models. For example, the general product prior has 5,322 hyperparameters, compared to 562 for the Jenga prior and only 28 for the connectedness prior.

In the previous examples, we found that the site-, allele-, and mutation-specific decay factors are often related to the qualitative properties of complex sequence–function relationships and can help identify regions of sequence space where mutational effects differ systematically. We next examine decay factors inferred under different priors to guide exploration and interpretation of the AAV2 landscape. First, we found that the site-specific decay factors δ_p inferred by the connectedness prior show substantial variation across positions, highlighting sites with little to no influence on the effects of other mutations ($\delta_{575} = 0.08$, $\delta_{561} = 0.10$) and sites with strong influence ($\delta_{568} = 0.37$, $\delta_{569} = 0.52$, Figure 3B). The Jenga prior exhibits even greater heterogeneity in the inferred decay factors, revealing that the identity of the mutated allele can matter more than the site itself. For example, while position 569 has the highest site-specific decay factor under the connectedness prior, the Jenga model shows that this effect is driven primarily by the Asn allele ($\delta^{\text{Asn}} = 0.43$), with other substitutions at the same position showing minimal influence on predictability ($\delta^{a \neq \text{Asn}} < 0.08$). To explore this further, we used the Gaussian process framework to compare posterior distributions of mutational effects in the wild-type background and in the presence of an Asn569Gln substitution (Figure S5A). In the wild-type background, mutations had widespread effects, but in the Asn569Gln background, these effects became largely neutral. This pattern reflects an essential role for Asn569 in capsid assembly, as all measured sequences lacking Asn at position 569 were non-functional (Figure 3C), with little apparent influence from genetic background (Figure S5B). Structurally, Asn569 forms hydrogen bonds with nearby residues, including 522 and 507 in AAV2 (Figure S5C), and these interactions are conserved across multiple AAV serotypes (Figure S5D, E). Even the conservative Asn569Gln substitution is predicted to introduce steric clashes that likely destabilize the capsid (Figure S5F).

Next, rather than focusing on a single position, we examined a region spanning positions 579–588, located in a solvent-exposed loop (Figure S6A), where the positively charged residues Arg and Lys exhibit consistently high decay rates (Figure 3B). In the wild-type background, mutational effect estimates reveal that Arg and Lys are generally deleterious throughout this region, whereas substitutions from positively charged to neutral or negatively charged residues, such as Glu and Asp, tend to be beneficial (Figure S6B). These patterns led us to hypothesize that the functional constraint in this region is charge-dependent. Supporting this hypothesis, mutations that reduce the overall charge, such as R588Q, are estimated to become neutral when introduced in an

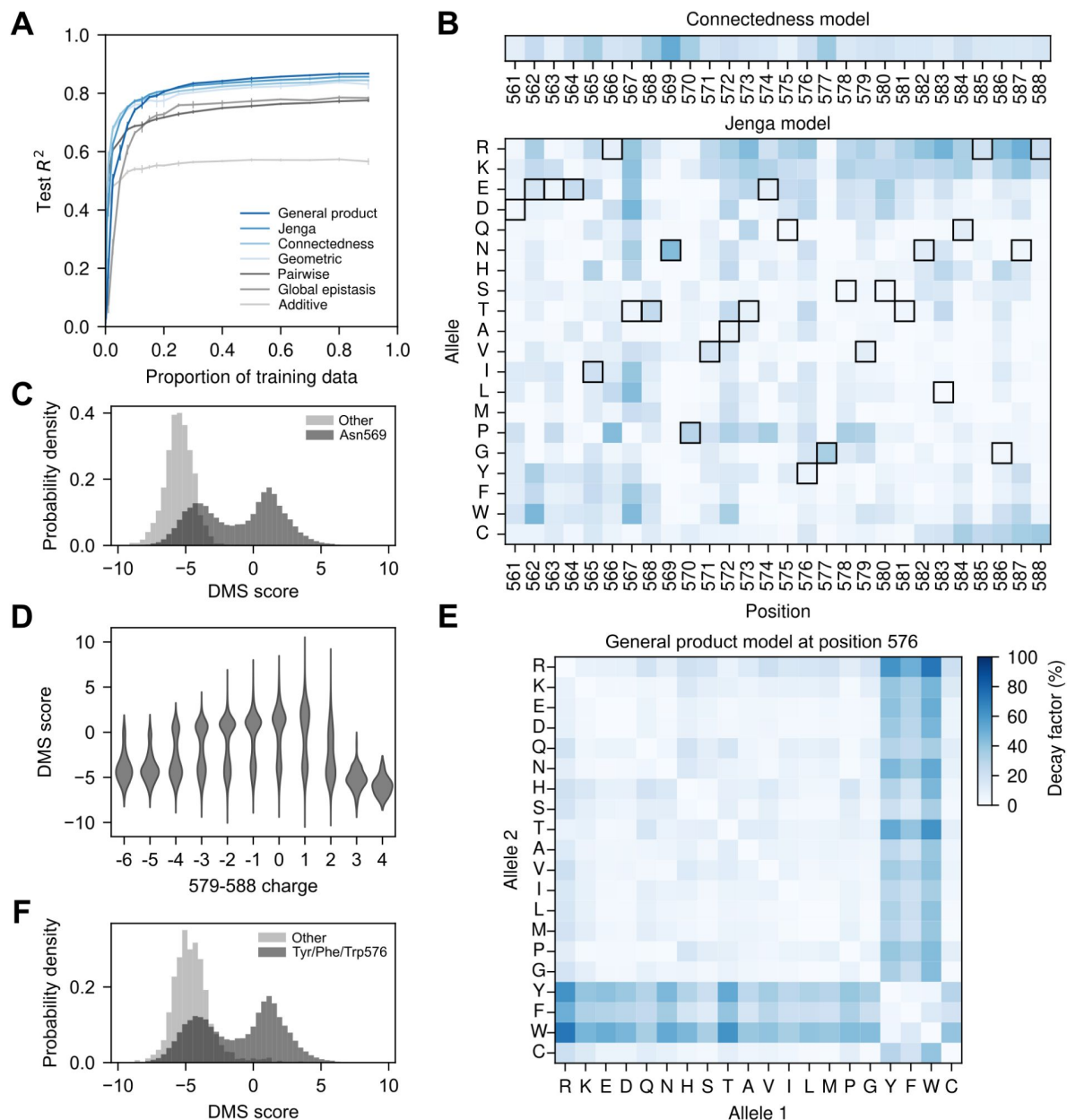


Figure 3

Application to the AAV2 capsid protein.

(A) Model predictive performance measured by test R^2 on held-out data, shown as a function of the proportion of training data. Error bars represent one standard deviation across three random subsets of training sequences at each proportion. (B) Inferred decay factors for each of the 28 positions under the connectedness model (top) and for 20 amino acids at each position under the Jenga model (bottom). Black squares indicate the wild-type allele at each position. (C) Distribution of raw DMS scores for sequences with Asn at position 569 (dark gray) compared to sequences with other amino acids at that site (light gray). (D) Raw DMS scores as a function of net charge across amino acids 579–588, indicating that intermediate charge levels are associated with higher viral production. (E) Inferred mutation-specific decay factors from the general product model for position 576. (F) Distribution of raw DMS scores for sequences with an aromatic residue (Tyr, Phe, or Trp) at position 576 (dark gray) compared to sequences with other amino acids at that site (light gray).

Asn587Glu background (Figure S6C) or even deleterious when combined with an additional Arg585Glu substitution (Figure S6D). To directly test this charge-balance hypothesis, we stratified the experimental data by the net charge in the 579–588 region. Consistent with our expectations, only sequences with a net charge between -4 and +2 could be functional, while those with excessive positive or negative charge were invariably non-functional. This suggests that an intermediate charge in this region is required, but not sufficient, for capsid assembly (Figure 3D).

We also examined the mutation-specific decay factors across all positions under the general product model (Figure S4). These values highlight the extensive variability and idiosyncrasy in amino acid exchangeability across different positions and biochemical contexts. Several positions exhibit mutation-specific decay factors that are incompatible with those expected under the Jenga prior, where mutations within groups are not expected to alter the effects of other mutations, whereas mutations across groups are. This pattern resembles position 54 in the GB1 dataset (Figure 1E). For instance, mutation-specific decay factors at position 576 exhibit two distinct groups: aromatic and non-aromatic residues (Figure 3E). Consistent with this hypothesis, the estimated mutational effects in the wild-type context were largely preserved in the presence of a Tyr576Phe substitution (Figure S5G), whereas other substitutions, like Tyr576Cys, systematically reduced the effects of other deleterious mutations (Figure S5H). Moreover, measured sequences without an aromatic at position 576 were generally non-functional (Figure 3F), thereby limiting the ability of other deleterious substitutions to further contribute to loss of function. This pattern is consistent with structural evidence: the wild-type Tyr576 occupies a hydrophobic pocket at the interface between two monomers, a location likely incompatible with non-aromatic residues (Figure S5I).

Application to a genome-wide genotype-phenotype map

So far, we have applied our methods to genotype–phenotype maps for proteins and regulatory RNA sequences. Here, we extend our approach to a genome-wide genotype–phenotype map. We analyze a dataset derived from a large barcoded segregant library for haploid yeast [50]. Specifically, we modeled the relative fitness under lithium exposure for nearly 20,000 segregants with high-quality genotype data, focusing on the 83 previously identified quantitative trait loci (QTLs) spread across 15 chromosomes [50].

We fit the baseline additive, pairwise interaction, and geometric models to different proportions of the training data. All three models exhibit highly similar prediction performance across both low and high training data densities and converge to a test R^2 of approximately 0.69 (Figure 4A). The failure of these epistatic models to substantially outperform the additive model suggests a lack of pervasive genetic interactions across genomic loci. We next fit connectedness regression, which in bi-allelic datasets such as this one is equivalent to both the Jenga and general product kernel regression models. Interestingly, we observed a consistent improvement in predictive performance relative to the baseline additive and epistatic models across all proportions of training data, achieving a test R^2 of 0.76 on held-out data when including at least 50% of the data (Figure 4A). This result suggests that loci likely have heterogeneous epistatic contributions and differ in their influence on the predictability of mutations at other positions, highlighting the importance of accounting for this heterogeneity to achieve high predictive performance.

To further investigate the heterogeneous contribution of loci to the predictability of mutational effects, we examined locus-specific decay factors across the 83 focal QTLs (Figure 4B). Our analysis revealed that fitness variation is dominated by a single QTL located near the gene *ENA1* ($\delta_{ENA1} = 0.43$), which encodes a Na⁺/Li⁺ ATPase pump involved in Na⁺ and Li⁺ efflux. As expected, this *ENA1* QTL also has a strong main effect, with genotypes carrying the beneficial BY or deleterious RM allele showing strongly shifted phenotypic distributions (Figure 4C).

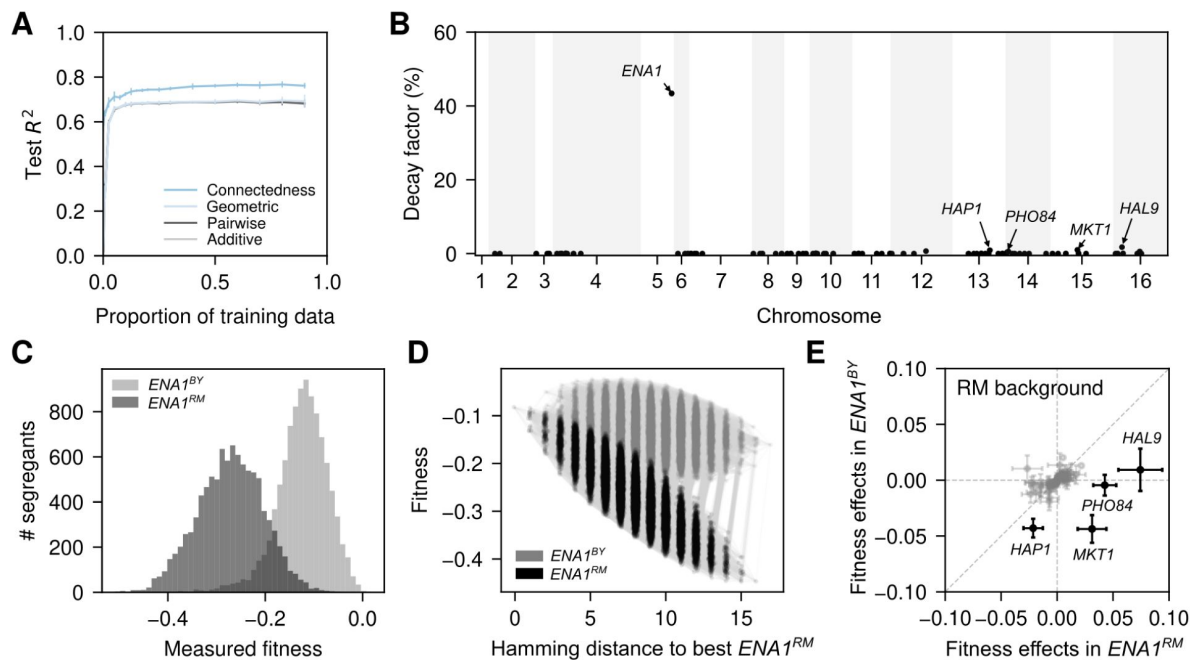


Figure 4

Connectedness regression applied to genome-wide genotype-phenotype data measuring relative fitness of *Saccharomyces cerevisiae* under lithium exposure.

(A) Model predictive performance, measured by the R^2 on held-out data, as a function of the fraction of data used for training. Error bars represent one standard deviation over three independent subsets of sequences used for training at each proportion. (B) Inferred decay factors for the 83 QTLs across the genome under the connectedness model. QTLs are named by the gene at which they are located or their closest gene in the reference genome S288C. (C) Fitness distribution of segregants stratified by their allele at the *ENA1* locus. (D) Representation of the complete genotype-phenotype map across the 16 loci with the highest inferred decay factors (mapped to the closest genes: *ENA1*, *HAL9*, *MKT1*, *PHO84*, *HAP1*, *HAL5*, *TAO3*, *BUL2*, *PTR2*, *NRT1*, *SUP45*, *DPH5*, *MLF3*, *SUS1*, *IRA2*, *VIP1*) and an additional pseudo-locus representing the genetic background across all other loci (RM vs. BY). The posterior mean fitnesses of the $2^{17} = 131,072$ possible genotypes are plotted against their Hamming distances to the genotype with the highest predicted fitness when combined with the *ENA1*^{RM} variant. Nodes represent genotypic combinations and are colored according to the allele at the *ENA1* locus. Edges connect genotypes separated by single point mutations. Values on the x-axis were jittered to facilitate visualization. (E) Comparison of the estimated mutational effects in the RM background in the presence of RM vs. BY alleles at the *ENA1* locus. Labeled QTLs highlighted in black correspond to loci with the largest decay factors after *ENA1*, as shown in (B). Error bars represent one standard deviation of the posterior distribution.

To examine how the *ENA1* QTL reshapes the genotype–phenotype map, we computed the posterior distribution for the complete combinatorial landscape involving all possible genotypes at the 16 loci with the largest decay factors, with the background (BY or RM) encoded as an additional locus (**Figure 4D** [↗](#)). In contrast to the AAV2 example, where mutations at Asn569 caused complete loss of function and eliminated the effects of all other mutations, the deleterious *ENA1*^{RM} allele modulates the effects of many mutations without uniformly suppressing them. Indeed, in this altered genetic context, previously neutral or deleterious variants can become beneficial, allowing combinations that achieve a relative fitness of −0.08, higher than the fitness of both the RM (−0.30) and BY (−0.18) wild-type strains.

To investigate how the *ENA1* QTL influences the effects of other QTLs in more detail, we computed the posterior distribution of the mutational effect for each of the other QTLs in the presence of *ENA1*^{RM} or *ENA1*^{BY} in an otherwise RM genetic background (**Figure 4E** [↗](#)). Whereas the effect of some QTLs with large decay factors, such as *HAP1*, remain deleterious in the presence of both *ENA1* alleles, we find 3 QTLs whose effects substantially change depending on the *ENA1* locus (**Figures 4E** [↗](#) and S7).

The first of these QTLs is near *HAL9*, a known regulator of *ENA1* expression [[100](#) [↗](#)]. We found that the *HAL9*^{BY} allele is neutral in the *ENA1*^{BY} background, but beneficial in the presence of the deleterious *ENA1*^{RM} allele, suggesting that *HAL9*^{BY} likely compensates for the deleterious *ENA1*^{RM} via up-regulation of *ENA1*. Similarly, we found that for the second QTL, which mapped to *PHO84*, the BY allele is beneficial only in the *ENA1*^{RM} background (**Figure 4E** [↗](#)). As *PHO84* is a high-affinity proton-phosphate symporter [[101](#) [↗](#)] and uses the proton gradient to import inorganic phosphate in the cell, it may influence the efficiency of a secondary detoxification mechanism involving the Na⁺, Li⁺-proton antiporter *NHA1*, which is known to become relevant in *ENA1* - deficient cells [[100](#) [↗](#)]. Finally, we found that the BY allele of a previously reported large-effect QTL near *MKT1* [[50](#) [↗](#), [102](#) [↗](#)] is detrimental in the presence of *ENA1*^{BY}, but becomes beneficial when combined with the deleterious *ENA1*^{RM} allele (**Figure 4E** [↗](#)). *MKT1* is a pleiotropic gene with known effects across multiple environments [[50](#) [↗](#), [102](#) [↗](#)], potentially by stabilizing specific mRNAs in a *PBP1* - and *PUF3* -dependent manner [[103](#) [↗](#)], although its interaction with genes involved in Li⁺ homeostasis, has not been previously reported to our knowledge. Importantly, the effects of these QTLs do not only depend on *ENA1*, but also vary across a wider range of genetic backgrounds (Figure S7). For instance, the beneficial effect of *HAL9*^{BY} in the presence of *ENA1*^{RM} is stronger in an RM rather than a BY genetic background across all other loci (Figure S7, upper left).

In summary, this yeast example demonstrates that our proposed models can infer genome-wide genotype- phenotype maps using data not only from engineered sequence libraries but also from high-throughput phenotyping of segregating variation in large populations. The high predictive performance of connectedness regression on held-out genotypes provides evidence of an important contribution of genetic interactions and coordinated epistasis to genome-scale genotype-phenotype maps, allowing us to detect and quantify novel genetic interactions and to exhaustively reconstruct small portions of even astronomically large genotype- phenotype maps (e.g., **Figure 4D** [↗](#), showing a non-trivial sub-landscape consisting of 131,072 genotypes out of the $2^{83} \approx 10^{25}$ possible genotypes in this system).

Discussion

In this study, we introduced a family of models motivated by a simple yet informative aspect of epistasis: some mutations change the predictability of mutational effects more than others. Our methods provide a comprehensive framework for studying empirical sequence-function relationships by both allowing accurate phenotypic prediction for genotype-phenotype maps containing all orders of genetic interaction and providing biologically interpretable summaries of

the structure of epistasis. The strong performance of these models across four empirical datasets, encompassing sequence-function relationships for proteins, RNAs, and complex traits, provides further evidence that heterogeneity across mutations in the degree to which they alter the effects of other mutations is likely a general feature of empirical sequence-function relationships, corresponding to genotype-phenotype maps that are rugged in some directions but smooth in others.

Our contribution also addresses a longstanding limitation of Gaussian process approaches: scalability to long sequences and large datasets. Classically, Gaussian processes could only be applied exactly to datasets containing fewer than low tens of thousands of observations [75], and while our previous methods exploited symmetries in biological sequence space to accommodate hundreds of thousands to low millions of observations [55, 58, 70, 74], they remained applicable only to short sequences. By leveraging recent advances in GPU acceleration [81–83], we can now overcome these previous constraints on sequence length, allowing us to provide software that can scale to longer protein sequences and genome-scale QTLs. Despite working in these astronomically large sequence spaces, our modeling framework maintains a high degree of interpretability because our hyperparameters have a definite genetic meaning in terms of how much each mutation reduces the predictability of other mutations. Moreover, by displaying these decay factors in heatmaps, we can at a glance identify the key alleles at each site most involved in epistatic interactions (e.g. **Figure 3B**) or subsets of amino acids that behave equivalently at a given site (e.g. **Figure 3C**, Figure S3), both of which can help guide the choice of mutations for follow-up experiments, and as we show, motivate biological hypotheses to explain the observed pattern of epistasis.

Our work also advances the ongoing process of unification and integration between theoretical and empirical approaches to fitness landscapes [5, 7, 8, 104]. In particular, an important advantage of Gaussian processes is that we can use well-characterized families of random fitness landscapes as priors for empirical data analysis [77, 85–88, 105]. By learning the hyperparameters for these priors via evidence maximization, our methods effectively identify the ensemble of theoretical landscapes whose patterns of epistasis are most similar to those in the observed data. At the same time, our proposal of the Jenga and general product kernel models provides an opportunity to analyze the dynamics of adaptive evolution and the accessibility of mutational paths under these more realistic forms of anisotropy [77]. Our work is also closely connected to the existing fitness landscape literature via our decay factors, which are conceptually related to the generalized γ statistics [28, 76] that measure the correlation in mutational effects across pairs of genetic backgrounds that differ by a specific mutation. We derive the mathematical relationship between the γ statistics and our decay factors in Supplement 1.3.

This work is also situated within a broader literature on Gaussian process approaches to modeling genotype-phenotype relationships [67–69, 106–111]. Previously, anisotropy has typically been introduced by incorporating external information, e.g. by building Gaussian processes over sequence embeddings [106, 107], by incorporating structural information [67], or by using information from sequence alignments [110]. By contrast, the models introduced here learn anisotropy directly from the data. In terms of expressiveness, our models are also similar in their motivations to deep neural networks that can act as general function approximators [34, 35, 46, 59–66], but Gaussian processes offer advantages in terms of interpretability and uncertainty quantification. One key advantage of DNNs is their ability to accommodate nonlinear transformations of model outputs, enabling them to capture global or nonspecific forms of epistasis [60, 61, 65, 112, 113]. In contrast, adding similar nonlinearities to Gaussian process models introduces non-Gaussian likelihoods requiring additional approximations, such as Laplace approximation [75] or variational inference [114–116]. Nonetheless, our new anisotropic Gaussian process models appear to consistently outperform global epistasis models except at very low training data densities. These new models also performed particularly well on the SMN1 dataset, which includes a strong global nonlinearity, and captured this nonlinear

structure much better than our previous isotropic priors [55]. While our models achieved excellent predictive performance, they also reveal several promising directions for future research. The connectedness, Jenga, and general product kernel priors by construction encode a form of “coordinated epistasis,” where the mutations with the largest effects also tend to be the most epistatic, a pattern consistent with many empirical observations [51, 117–120]. However, more work is needed to develop priors that capture not just how epistatic each mutation tends to be, but also provide more fine-grained summaries of which sites or mutations tend to interact with each other [121]. Another limitation is that while our priors allow anisotropy in the predictability of mutational effects, they are nonetheless homoskedastic, assigning uniform phenotypic variance across functionally distinct regions of sequence space. Developing heteroskedastic priors, where phenotypic variance varies across sequence space would better capture the expectation that functionally inert regions are less variable than those containing functionally active sequences. However, it is important to note that these are limitations on the inductive biases expressible through the prior, rather than limitations of the Bayesian regression models themselves. In particular, these models are capable of learning arbitrary genotype-phenotype maps, and any additional structure will still be reflected in the posterior, even if not directly encoded in the prior.

Methods

Data processing

GB1 protein data was processed as previously described [55, 58]. Briefly, we used the number of sequencing reads for each sequence x in the input sample (c_x^{input}) and in the selected sample (c_x^{sel}) reported in [84] to estimate the log-enrichment ratio relative to the wild-type sequence y_x as a measure of the binding strength. Moreover, we estimated the error variance σ^2 of this estimate [122]:

$$y_x = \log \left(\frac{c_x^{sel} + 0.5}{c_x^{input} + 0.5} \right) - \log \left(\frac{c_{wt}^{sel} + 0.5}{c_{wt}^{input} + 0.5} \right) \quad (7)$$

$$\sigma_x^2 = \frac{1}{c_x^{input} + 0.5} + \frac{1}{c_x^{sel} + 0.5} + \frac{1}{c_{wt}^{input} + 0.5} + \frac{1}{c_{wt}^{sel} + 0.5}. \quad (8)$$

SMN1 5' splice site data was downloaded from <https://github.com/davidmccandlish/vcregression>, which used data from the original publication [42]. Briefly, assuming a log-normal distribution in 1-7 replicates, for each different 5' splice site sequence x , the bias corrected geometric mean of the enrichment ratio was used as an estimate of the median enrichment ratio when the enrichment ratio was strictly positive for all replicates. Otherwise, the median of enrichment scores was used to estimate the phenotype y_x . Sequence-specific variance σ^2 was estimated as indicated below, where s^2 is the sample variance of the log-enrichment ratios if all replicates were strictly positive and were measured in at least two samples or the median of all s^2 for sequences x with at least two replicate measurements:

$$\sigma_x^2 = \left(e^{s_x^2} - 1 \right) e^{2y_x + s_x^2}, \quad (9)$$

see [55] for more details.

AAV2 capsid data was downloaded from the ProteinGym database [123] and used as provided with- out external estimates of the experimental variance (the Gaussian process fit still included an experimental noise term with a learned uniform variance, see below).

Yeast fitness data was downloaded from the Dryad repository (<https://doi.org/10.5061/dryad.1rn8pk0vd>). This dataset provided estimates of the mean y_x and variance σ^2 of the relative fitness for each barcoded segregant x . We used the reported genotype probabilities $p_{x,i}$ for each position i across all ℓ loci to compute a segregant-level genome uncertainty U_i as previously reported [50].

$$U_x = \frac{4}{\ell} \sum_m p_{x,i} (1 - p_{x,i}). \quad (10)$$

We kept only the 19833 (20%) segregants with lowest genome uncertainty for further analysis. We modeled the relative fitness in high Li^+ concentration as a function of the 83 previously reported Quantitative Trait Loci (QTL) [50]. QTLs were labeled by associating them with the closest protein-coding gene within a 10kb window from the BY4741 genome annotation downloaded from the *Saccharomyces* Genome Database (SGD) [124].

Global epistasis models

Global epistasis models for each data set were fit using MAVE-NN v1.0.2 [61] with an unregularized additive genotype-phenotype map, a monotonic global epistasis non-linearity, and Gaussian homoskedastic noise. Models were trained by minimizing the variational information with a batch size equal to the size of the dataset for 5000 epochs with a learning rate of 0.01 to ensure model convergence. All other parameters were set to their default values.

Gaussian process models

Gaussian process models for the different datasets were fit using our newly developed library EpiK. EpiK implements the proposed prior distributions for performing Bayesian inference using GPyTorch [81] and KeOps [83] for GPU-accelerated inference of Gaussian process models. Specifically, we use a zero mean Gaussian process prior and a Gaussian likelihood with known experimental variance σ^2 when available and fit an additional noise term σ^2 for additional sequence-independent variance (see Supplement 1.1). Hyperparameters of each kernel were optimized by maximizing the Bayesian evidence using the Adam optimizer with a learning rate of 0.02 for 500 iterations. For accurate estimation of the evidence, we use a maximum number of 600 Lanczos vectors with 100 random vectors for trace estimation when approximating the log-determinant term (see Supplement 1.1). Site-specific product kernels (geometric, connectedness, Jenga and general product kernels) were initialized to have a prior variance equal to the empirical phenotypic variance and uniform decay factors across mutations and sites, which decay exponentially with the number of mutations down to 0.1 at the maximal distance between sequences. As the evidence is a non-convex function with potential local optima, to obtain the final estimates for the decay factors when fitting the complete dataset, we run 5 independent optimizations with different random initializations and kept the hyperparameter values that minimized the evidence along the 5 different optimizations. However, in practice, we see convergence to the same optima independent of the random initialization [75].

Visualization of GB1 sequence-function map

We used the maximum a posteriori (MAP) estimate for the complete sequence-function map of protein GB1 inferred with the complete dataset under the Jenga model to generate a low-dimensional representation using gpmmap-tools [74], in which squared distances optimally approximate the time to evolve between pairs of sequences under a model of molecular evolution where the population evolves under selection for higher log relative enrichment ratios. The different diffusion axes capture directions in which evolution is the slowest and typically separate qualitatively different sets of sequences with high phenotypic values that are separated by at least

partial fitness valleys [58 [↗](#), 70 [↗](#), 93 [↗](#)]. Selection strength was set such that the stationary distribution of the resulting Markov chain had a phenotypic mean that matches the wild-type phenotypic value.

Data and code availability

The EpiK software package allows general Gaussian Process regression for sequence-function relationships under the kernels presented in this work. EpiK is freely available to the community at <https://github.com/cmarti/epik> [↗](#) under MIT license, with documentation available at <https://epik.readthedocs.io> [↗](#). Data and code to reproduce the analysis and figures from this work are available at https://github.com/cmarti/epik_analyses [↗](#).

Acknowledgements

This work was supported by a Burroughs-Wellcome Careers at the Scientific Interface (CASI) award (SP), NIH grant R35 GM133613 (CMG, DMM), NIH grant R01 HG011787 (DMM), and NIH grant R35 GM154908 (JZ), as well as additional funding from the Simons Center for Quantitative Biology at CSHL (DMM) and the College of Liberal Arts and Sciences at the University of Florida (JZ).

Additional files

Supplementary Information [↗](#)

Additional information

Funding

National Institutes of Health (GM133613)

National Institutes of Health (HG011787)

National Institutes of Health (GM154908)

Burroughs Wellcome Fund (Burroughs-Wellcome Careers at the Scientific Interface (CASI) award)

Simons Center for Quantitative Biology at CSHL

College of Liberal Arts and Sciences at the University of Florida

References

1. Wright S (1932) **The roles of mutation, inbreeding, crossbreeding and selection in evolution** In: *Proceedings of the Sixth International Congress of Genetics* pp. 356–366 [Google Scholar](#)

2. de Visser JA, Krug J (2014) **Empirical fitness landscapes and the predictability of evolution** *Nat Rev Genet* **15**:480 <https://doi.org/10.1038/nrg3744> | Google Scholar
- [3] Phillips Patrick C (2008) **Epistasis—the essential role of gene interactions in the structure and evolution of genetic systems** *Nat. Rev. Genet* **9**:855–867 <https://doi.org/10.1038/nrg2452> | Google Scholar
- [4] Kondrashov Alexey S, Sunyaev Shamil, Kondrashov Fyodor A (2002) **Dobzhansky–Muller incompatibilities in protein evolution** *Proc. Natl. Acad. Sci. U.S.A* **99**:14878–14883 <https://doi.org/10.1073/pnas.232565499> | Google Scholar
- [5] Weinreich Daniel M, Lan Yinghong, Scott Wylie C, Heckendorn Robert B (2013) **Should evolutionary geneticists worry about higher-order epistasis?** *Curr. Opin. Genet. Dev* **23**:700–707 <https://doi.org/10.1016/j.gde.2013.10.007> | Google Scholar
- [6] Sailer Zachary R, Harms Michael J (2017) **High-order epistasis shapes evolutionary trajectories** *PLoS Comput Biol* **13**:e1005541 <https://doi.org/10.1371/journal.pcbi.1005541> | Google Scholar
- [7] Bank Claudia (2022) **Epistasis and Adaptation on Fitness Landscapes** *Annual Review of Ecology, Evolution, and Systematics* **53**:457–479 <https://doi.org/10.1146/annurev-ecolsys-102320-112153> | Google Scholar
- [8] Johnson Milo S., Reddy Gautam, Desai Michael M. (2023) **Epistasis and evolution: recent advances and an outlook for prediction** *BMC Biology* **21**:120 <https://doi.org/10.1186/s12915-023-01585-3> | Google Scholar
- [9] Moulana Alief, Dupic Thomas, Phillips Angela M, Chang Jeffrey, Roffler Anne A, Greaney Allison J, Starr Tyler N, Bloom Jesse D, Desai Michael M (2023) **The landscape of antibody binding affinity in SARS-CoV-2 Omicron BA.1 evolution** *eLife* **12**:e83442 <https://doi.org/10.7554/eLife.83442> | Google Scholar
- [10] Moore Jason H., Williams Scott M. (2009) **Epistasis and Its Implications for Personal Genetics** *The American Journal of Human Genetics* **85**:309–320 <https://doi.org/10.1016/j.ajhg.2009.08.006> | Google Scholar
- [11] Dasari Krishna, Somarelli Jason A., Kumar Sudhir, Townsend Jeffrey P. (2021) **The somatic molecular evolution of cancer: Mutation, selection, and epistasis** *Progress in Biophysics and Molecular Biology* **165**:56–65 <https://doi.org/10.1016/j.pbiomolbio.2021.08.003> | Google Scholar
- [12] Freschlin Chase R, Fahlberg Sarah A, Romero Philip A (2022) **Machine learning to navigate fitness landscapes for protein engineering** *Current Opinion in Biotechnology* **75**:102713 <https://doi.org/10.1016/j.copbio.2022.102713> | Google Scholar
- [13] Yang G, et al (2019) **Higher-order epistasis shapes the fitness landscape of a xenobiotic-degrading enzyme** *Nature Chemical Biology* **15**:1120–1128 <https://doi.org/10.1038/s41589-019-0386-3> | Google Scholar
- [14] Lipsh-Sokolik Rosalie, Fleishman Sarel J. (2024) **Addressing epistasis in the design of protein function** *Proceedings of the National Academy of Sciences* **121** <https://doi.org/10.1073/pnas.2314999121> | Google Scholar
- [15] Sackton Timothy B., Hartl Daniel L. (2016) **Genotypic Context and Epistasis in Individuals and Populations** *Cell* **166**:279–287 <https://doi.org/10.1016/j.cell.2016.06.047> | Google Scholar

- [16] De Los Campos Gustavo, Hickey John M, Pong-Wong Ricardo, Daetwyler Hans D, Calus Mario P L (2013) **Whole-Genome Regression and Prediction Methods Applied to Plant and Animal Breeding** *Genetics* **193**:327–345 <https://doi.org/10.1534/genetics.112.143313> | Google Scholar
- [17] Soyk Sebastian, Benoit Matthias, Lippman Zachary B. (2020) **New Horizons for Dissecting Epistasis in Crop Quantitative Trait Variation** *Annual Review of Genetics* **54**:287–307 <https://doi.org/10.1146/annurev-genet-050720-122916> | Google Scholar
- [18] Dwivedi Sangam L., Heslop-Harrison Pat, Amas Junrey, Ortiz Rodomiro, Edwards David (2024) **Epistasis and pleiotropy-induced variation for plant breeding** *Plant Biotechnology Journal* **22**:2788–2807 <https://doi.org/10.1111/pbi.14405> | Google Scholar
19. Kinney JB, McCandlish DM (2019) **Massively Parallel Assays and Quantitative Sequence-Function Relationships** *Annu. Rev. Genomics Hum. Genet* **20** | Google Scholar
- [20] Fowler Douglas M, Araya Carlos L, Fleishman Sarel J, Kellogg Elizabeth H, Stephany Jason J, Baker David, Fields Stanley (2010) **High-resolution mapping of protein sequence-function relationships** *Nat Methods* **7**:741–746 <https://doi.org/10.1038/nmeth.1492> | Google Scholar
- [21] Starita Lea M, Pruneda Jonathan N, Lo Russell S, Fowler Douglas M, Kim Helen J, Hiatt Joseph B, Shendure Jay, Brzovic Peter S, Fields Stanley, Klevit Rachel E (2013) **Activity-enhancing mutations in an E3 ubiquitin ligase identified by high-throughput mutagenesis** *Proc. Natl. Acad. Sci. U.S.A* **110**:E1263–E1272 <https://doi.org/10.1073/pnas.1303309110> | Google Scholar
- [22] Melamed Daniel, Young David L, Gamble Caitlin E, Miller Christina R, Fields Stanley (2013) **Deep mutational scanning of an RRM domain of the *Saccharomyces cerevisiae* poly (A)-binding protein** *RNA* **19**:1537–1551 <https://doi.org/10.1261/rna.040709.113> | Google Scholar
- [23] Anders Olson C, Wu Nicholas C, Sun Ren (2014) **A comprehensive biophysical description of pairwise epistasis throughout an entire protein domain** *Curr Biol* **24**:2643–2651 <https://doi.org/10.1016/j.cub.2014.09.072> | Google Scholar
- [24] Doud Michael B, Ashenberg Orr, Bloom Jesse D (2015) **Site-specific amino acid preferences are mostly conserved in two closely related protein homologs** *Mol. Biol. Evol* **32**:2944–2960 <https://doi.org/10.1093/molbev/msv167> | Google Scholar
- [25] Podgornaia Anna I, Laub Michael T (2015) **Pervasive degeneracy and epistasis in a protein-protein interface** *Science* **347**:673–677 <https://doi.org/10.1126/science.1257360> | Google Scholar
- [26] Sarkisyan Karen S, Bolotin Dmitry A, Meer Margarita V, Usmanova Dinara R, Mishin Alexander S, Sharonov George V, Ivankov Dmitry N, Bozhanova Nina G, Baranov Mikhail S, Soylemez Onuralp, et al. (2016) **Local fitness landscape of the green fluorescent protein** *Nature* **533**:397 <https://doi.org/10.1038/nature17995> | Google Scholar
- [27] Steinberg Barrett, Ostermeier Marc (2016) **Shifting fitness and epistatic landscapes reflect trade-offs along an evolutionary pathway** *J Mol. Biol* **428**:2730–2743 <https://doi.org/10.1016/j.jmb.2016.04.033> | Google Scholar
- [28] Bank Claudia, Matuszewski Sebastian, Hietpas Ryan T, Jensen Jeffrey D (2016) **On the (un)predictability of a large intragenic fitness landscape** *Proc. Natl. Acad. Sci. U.S.A* **113**:14085–14090 <https://doi.org/10.1073/pnas.1612676113> | Google Scholar

- [29] Starr Tyler N, Picton Lora K, Thornton Joseph W (2017) **Alternative evolutionary histories in the sequence space of an ancient protein** *Nature* **549**:409–413 <https://doi.org/10.1038/nature23902> | [Google Scholar](#)
- [30] Pokusaeva Victoria O, Usmanova Dinara R, Putintseva Ekaterina V, Espinar Lorena, Sarkisyan Karen S, Mishin Alexander S, Bogatyreva Natalya S, Ivankov Dmitry N, Akopyan Arseniy V, Avvakumov Sergey Ya, et al. (2019) **An experimental assay of the interactions of amino acids from orthologous sequences shaping a complex fitness landscape** *PLoS Genet* **15**:e1008079 <https://doi.org/10.1371/journal.pgen.1008079> | [Google Scholar](#)
- [31] Plesa C, et al (2018) **Multiplexed gene synthesis in emulsions for exploring protein functional landscapes** *Science* **359**:343–347 <https://doi.org/10.1126/science.aao5167> | [Google Scholar](#)
- [32] Tack Drew S, Tonner Peter D, Pressman Abe, Olson Nathan D, Levy Sasha F, Romantseva Eugenia F, Alperovich Nina, Vasilyeva Olga, Ross David (2021) **The genotype-phenotype landscape of an allosteric protein** *Molecular systems biology* **17**:e10179 <https://doi.org/10.15252/msb.202010179> | [Google Scholar](#)
- [33] Starr Tyler N, Greaney Allison J, Hilton Sarah K, Ellis Daniel, Crawford Katharine HD, Dingens Adam S, Navarro Mary Jane, Bowen John E, Alejandra Tortorici M, Walls Alexandra C, et al. (2020) **Deep mutational scanning of SARS-CoV-2 receptor binding domain reveals constraints on folding and ACE2 binding** *Cell* **182**:1295–1310 <https://doi.org/10.1016/j.cell.2020.08.012> | [Google Scholar](#)
- [34] Somermeyer Louisa Gonzalez, Fleiss Aubin, Mishin Alexander S, Bozhanova Nina G, Igoikina Anna A, Meiler Jens, Pujol Maria-Elisenda Alaball, Putintseva Ekaterina V, Sarkisyan Karen S, Kondrashov Fyodor A (2022) **Heterogeneity of the GFP fitness landscape and data-driven protein design** *eLife* **11**:e75842 <https://doi.org/10.7554/eLife.75842> | [Google Scholar](#)
- [35] Bryant Drew H, Bashir Ali, Sinai Sam, Jain Nina K, Ogden Pierce J, Riley Patrick F, Church George M, Colwell Lucy J, Kelsic Eric D (2021) **Deep diversification of an AAV capsid protein by machine learning** *Nature Biotechnology* **39**:691–696 <https://doi.org/10.1038/s41587-020-00793-4> | [Google Scholar](#)
- [36] Faure Andre J., Martí-Aranda Aina, Hidalgo-Carcedo Cristina, Beltran Antoni, Schmiedel Jörn M., Lehner Ben (2024) **The genetic architecture of protein stability** *Nature* **634**:995–1003 <https://doi.org/10.1038/s41586-024-07966-0> | [Google Scholar](#)
- [37] Beltran A, et al (2025) **Site-saturation mutagenesis of 500 human protein domains** *Nature* :1–10 <https://doi.org/10.1038/s41586-024-08370-4> | [Google Scholar](#)
- [38] Kinney Justin B, Murugan Anand, Callan Curtis G, Cox Edward C (2010) **Using deep sequencing to characterize the biophysical mechanism of a transcriptional regulatory sequence** *Proc. Natl. Acad. Sci. U.S.A* **107**:9158–9163 <https://doi.org/10.1073/pnas.1004290107> | [Google Scholar](#)
- [39] Rosenberg Alexander B. B., Patwardhan Rupali P. P., Shendure Jay, Seelig Georg (2015) **Learning the Sequence Determinants of Alternative Splicing from Millions of Random Sequences** *Cell* **163**:698–711 <https://doi.org/10.1016/j.cell.2015.09.054> | [Google Scholar](#)
- [40] Rabani Michal, Pieper Lindsey, Chew Guo-Liang, Schier Alexander F. (2017) **A Massively Parallel Reporter Assay of 3 UTR Sequences Identifies In Vivo Rules for mRNA Degradation** *Molecular Cell* **68**:1083–1094 <https://doi.org/10.1016/j.molcel.2017.11.014> | [Google Scholar](#)

- [41] Evfratov Sergey A., Osterman Ilya A., Komarova Ekaterina S., Pogorelskaya Alexandra M., Rubtsova Maria P., Zatsepin Timofei S., Semashko Tatiana A., Kostryukova Elena S., Mironov Andrey A., Burnaev Evgeny, Krymova Ekaterina, Gelfand Mikhail S., Govorun Vadim M., Bogdanov Alexey A., Sergiev Petr V., Dontsova Olga A. (2017) **Application of sorting and next generation sequencing to study 5'-UTR influence on translation efficiency in Escherichia coli** *Nucleic Acids Research* **45**:3487–3502 <https://doi.org/10.1093/nar/gkw1141> | Google Scholar
- [42] Wong Mandy S, Kinney Justin B, Krainer Adrian R (2018) **Quantitative Activity Profile and Context Dependence of All Human 5' Splice Sites** *Mol Cell* <https://doi.org/10.1016/j.molcel.2018.07.033> | Google Scholar
- [43] Baeza-Centurion Pablo, Miñana Belén, Schmiedel Jörn M., Valcárcel Juan, Lehner Ben (2019) **Com- binatorial Genetics Reveals a Scaling Law for the Effects of Mutations on Splicing** *Cell* **176**:549–563 <https://doi.org/10.1016/j.cell.2018.12.010> | Google Scholar
- [44] de Boer Carl G., Vaishnav Eeshit Dhaval, Sadeh Ronen, Abeyta Esteban Luis, Friedman Nir, Regav Aviv (2020) **Deciphering eukaryotic gene-regulatory logic with 100 million random promoters** *Nature Biotechnology* **38**:56–65 <https://doi.org/10.1038/s41587-019-0315-8> | Google Scholar
- [45] Kuo Syue-Ting, Jahn Ruey-Lin, Cheng Yuan-Ju, Chen Yi-Lan, Lee Yun-Ju, Hollfelder Florian, Wen Jin-Der, Chou Hsin-Hung David (2020) **Global fitness landscapes of the Shine-Dalgarno sequence** *Genome Research* **30**:711–723 <https://doi.org/10.1101/gr.260182.119> | Google Scholar
- [46] Vaishnav Eeshit Dhaval, de Boer Carl G., Molinet Jennifer, Yassour Moran, Fan Lin, Adiconis Xian, Thompson Dawn A., Levin Joshua Z., Cubillos Francisco A., Regav Aviv (2022) **The evolution, evolvability and engineering of gene regulatory DNA** *Nature* <https://doi.org/10.1038/s41586-022-04506-6> | Google Scholar
- [47] Liao Susan E., Sudarshan Mukund, Regav Oded (2023) **Deciphering RNA splicing logic with interpretable machine learning** *Proceedings of the National Academy of Sciences* **120**:e2221165120 <https://doi.org/10.1073/pnas.2221165120> | Google Scholar
- [48] Agarwal Vikram, Inoue Fumitaka, Schubach Max, Penzar Dmitry, Martin Beth K., Dash Pyaree Mohan, Keukeleire Pia, Zhang Zicong, Sohota Ajuni, Zhao Jingjing, Georgakopoulos-Soares Ilias, Noble William S., Gürkan Yardımcı Galip, Kulakovskiy Ivan V., Kircher Martin, Shendure Jay, Ahituv Nadav (2025) **Massively parallel characterization of transcriptional regulatory elements** *Nature* :1–10 <https://doi.org/10.1038/s41586-024-08430-9> | Google Scholar
- [49] Bakerlee CW, Nguyen Ba AN, Shulgina Y, Rojas Echenique JI, Desai MM (2022) **Idiosyncratic epistasis leads to global fitness-correlated trends** *Science* [Google Scholar](https://doi.org/10.1126/science.abc)
- 50. Nguyen Ba AN, et al (2022) **Barcoded bulk QTL mapping reveals highly polygenic and epistatic architecture of complex traits in yeast** *eLife* **11**:e73983 [Google Scholar](https://doi.org/10.7554/eLife.73983)
- 51. Matsui T, et al (2022) **The interplay of additivity, dominance, and epistasis on fitness in a diploid yeast cross** *Nature Communications* **13**:1463 [Google Scholar](https://doi.org/10.1038/s41467-022-28000-0)
- 52. N'Guessan A, Tong WY, Heydari H, Nguyen Ba AN (2025) **Refining the resolution of the yeast genotype-phenotype map using single-cell RNA-sequencing** *eLife* **13** [Google Scholar](https://doi.org/10.7554/eLife.80000)

53. Kondrashov DA, Kondrashov FA (2015) **Topological features of rugged fitness landscapes in sequence space** *Trends Genet* **31**:24–33 [Google Scholar](#)
54. Domingo J, et al (2019) **The Causes and Consequences of Genetic Interactions (Epistasis)** *Annu. Rev. Genomics Hum. Genet* **20** [Google Scholar](#)
55. Zhou J, et al (2022) **Higher-order epistasis and phenotypic prediction** *Proceedings of the National Academy of Sciences* **119**:e2204233119 [Google Scholar](#)
56. Park Y, et al (2022) **Epistatic drift causes gradual decay of predictability in protein evolution** *Science* :823–830 [Google Scholar](#)
57. Fisher RA (1918) **The Correlation Between Relatives on the Supposition of Mendelian Inheritance** *Trans. R. Soc. Edinburgh* :399–433 [Google Scholar](#)
58. Zhou J, McCandlish DM (2020) **Minimum epistasis interpolation for sequence-function relationships** *Nature Communications* **11**:1–14 [Google Scholar](#)
59. Aghazadeh A, et al (2021) **Epistatic Net allows the sparse spectral regularization of deep neural networks for inferring fitness functions** *Nature Communications* **12**:1–10 [Google Scholar](#)
60. Sarkisyan KS, et al (2016) **Local fitness landscape of the green fluorescent protein** *Nature* **533**:397–401 [Google Scholar](#)
61. Tareen A, et al (2022) **MAVE-NN: learning genotype-phenotype maps from multiplex assays of variant effect** *Genome biology* **23**:98 [Google Scholar](#)
62. Gelman S, et al (2021) **Neural networks to learn protein sequence–function relationships from deep mutational scanning data** *Proceedings of the National Academy of Sciences* **118**:e2104878118 [Google Scholar](#)
- [63] de Almeida Bernardo P., Reiter Franziska, Pagani Michaela, Stark Alexander (2022) **DeepSTARR predicts enhancer activity from DNA sequence and enables the de novo design of synthetic enhancers** *Nature Genetics* **54**:613–624 <https://doi.org/10.1038/s41588-022-01048-5> | [Google Scholar](#)
64. Freschlin CR, et al (2024) **Neural network extrapolation to distant regions of the protein fitness landscape** *Nature Communications* **15**:6405 [Google Scholar](#)
- [65] Sethi Palash, Zhou Juannan (2024) **Importance of higher-order epistasis in large protein sequence-function relationships** *bioRxiv* [Google Scholar](#)
- [66] Thompson Mike, Martín Mariano, Olmo Trinidad Sanmartín, Rajesh Chandana, Koo Peter K, Bolognesi Benedetta, Lehner Ben (2025) **Massive experimental quantification allows interpretable deep learning of protein aggregation** *Science Advances* [Google Scholar](#)
67. Romero PA, et al (2013) **Navigating the protein fitness landscape with Gaussian processes** *Proc. Natl. Acad. Sci. U.S.A* **110**:E193–E201 [Google Scholar](#)
68. Gianola D, van Kaam JB (2008) **Reproducing kernel Hilbert spaces regression methods for genomic assisted prediction of quantitative traits** *Genetics* **178**:2289–2303 [Google Scholar](#)

69. Morota G, et al (2013) **Predicting complex traits using a diffusion kernel on genetic markers with an application to dairy cattle and wheat data** *Genetics Selection Evolution* :1–15 [Google Scholar](#)
70. Chen WC, et al (2021) **Field- theoretic density estimation for biological sequence space with applications to 5 splice site diversity and aneuploidy in cancer** *Proceedings of the National Academy of Sciences* **118** [Google Scholar](#)
71. Chen WC, et al (2024) **Density estimation for ordinal biological sequences and its applications** *Physical Review E* **110**:044408 [Google Scholar](#)
- [72] Rapp Jacob T., Bremer Bennett J., Romero Philip A. (2024) **Self-driving laboratories to autonomously navigate the protein fitness landscape** *Nature Chemical Engineering* **1**:97–107 <https://doi.org/10.1038/s44286-023-00002-4> | [Google Scholar](#)
73. Petti S, et al (2025) **On learning functions over biological sequence space: relating Gaussian process priors, regularization, and gauge fixing** *arXiv* [Google Scholar](#)
74. Martí-Gómez C, et al (2025) **Inference and visualization of complex genotype-phenotype maps with gpmmap-tools** *bioRxiv* [Google Scholar](#)
75. Rasmussen CE, Williams CKI (2006) **Gaussian processes for machine learning** MIT Press [Google Scholar](#)
76. Ferretti L, et al (2016) **Measuring epistasis in fitness landscapes: The correlation of fitness effects of mutations** *J. Theor. Biol* :132–143 [Google Scholar](#)
77. Reddy G, Desai MM (2021) **Global epistasis emerges from a generic model of a complex trait** *eLife* **10**:e64740 [Google Scholar](#)
78. Weinreich DM, et al (2005) **Perspective: sign epistasis and genetic constraint on evolutionary trajectories** *Evolution* :1165–1174 [Google Scholar](#)
79. Kvitek DJ, Sherlock G (2011) **Reciprocal sign epistasis between frequently experimentally evolved adaptive mutations causes a rugged fitness landscape** *PLoS genetics* **7**:e1002056 [Google Scholar](#)
- [80] Kingman John FC (1978) **A simple model for the balance between selection and mutation** *Journal of Applied Probability* **15**:1–12 [Google Scholar](#)
81. Gardner J, et al (2018) **Gpytorch: Blackbox matrix-matrix gaussian process inference with gpu acceleration** *Advances in neural information processing systems* [Google Scholar](#)
- [82] Wang Ke, Pleiss Geoff, Gardner Jacob, Tyree Stephen, Weinberger Kilian Q, Wilson Andrew Gordon (2019) **Exact Gaussian Processes on a Million Data Points** In: *NeurIPS* [Google Scholar](#)
- [83] Charlier Benjamin, Feydy Jean, Glaunes Joan Alexis, Collin François-David, Durif Ghislain (2021) **Kernel operations on the GPU, with autodiff, without memory overflows** *Journal of Machine Learning Research* **22**:1–6 [Google Scholar](#)
- [84] Wu Nicholas C., Anders Olson C., Sun Ren (2016) **High-throughput identification of protein mutant stability computed from a double mutant fitness landscape** *Protein Sci* **25**:530–539 <https://doi.org/10.1002/pro.2840> | [Google Scholar](#)

- [85] Stadler Peter F, Happel Robert, et al. (1994) **Canonical approximation of landscapes** Santa Fe Institute [Google Scholar](#)
86. Stadler PF, Happel R (1999) **Random field models for fitness landscapes** *J. Math. Biol* **38**:435–478 [Google Scholar](#)
- [87] Neidhart Johannes, Szendro Ivan G, Krug Joachim (2013) **Exact results for amplitude spectra of fitness landscapes** *Journal of Theoretical Biology* **332**:218–227 [Google Scholar](#)
88. Agarwala A, Fisher DS (2019) **Adaptive walks on high-dimensional fitness landscapes and seascapes with distance-dependent statistics** *Theoretical Population Biology* **130**:13–49 [Google Scholar](#)
89. Wang T, et al (2015) **An overview of kernel alignment and its applications** *Artificial Intelligence Review* :179–192 [Google Scholar](#)
90. Kondor RI, Lafferty JD (2002) **Diffusion Kernels on Graphs and Other Discrete Input Spaces** In: *Proceedings of the Nineteenth International Conference on Machine Learning* pp. 315–322 [Google Scholar](#)
- [91] Neal Radford M. (1996) **Bayesian Learning for Neural Networks** New York, NY: Springer <https://doi.org/10.1007/978-1-4612-0745-0> | [Google Scholar](#)
- [92] Lewandowski Daniel, Kurowicka Dorota, Joe Harry (2009) **Generating random correlation matrices based on vines and extended onion method** *Journal of Multivariate Analysis* **100**:1989–2001 <https://doi.org/10.1016/j.jmva.2009.04.008> | [Google Scholar](#)
93. McCandlish DM (2011) **Visualizing fitness landscapes** *Evolution* :1544–1558 [Google Scholar](#)
94. Otwinowski J, et al (2018) **Inferring the shape of global epistasis** *Proceedings of the National Academy of Sciences* **115**:E7550–E7558 [Google Scholar](#)
95. Sailer ZR, Harms MJ (2017) **Detecting high-order epistasis in nonlinear genotype-phenotype maps** *Genetics* :1079–1088 [Google Scholar](#)
96. Posfai A, et al (2025) **Gauge fixing for sequence- function relationships** *PLoS Computational Biology* **21**:e1012818 [Google Scholar](#)
97. Starr TN, Thornton JW (2016) **Epistasis in protein evolution** *Protein Sci* **25**:1204–1218 [Google Scholar](#)
- [98] Ishigami Yuma, Wong Mandy S., Martí-Gómez Carlos, Ayaz Andalus, Kooshkbaghi Mahdi, Hanson Sonya M., McCandlish David M., Krainer Adrian R., Kinney Justin B. (2024) **Specificity, synergy, and mechanisms of splice-modifying drugs** *Nature Communications* **15**:1880 <https://doi.org/10.1038/s41467-024-46090-5> | [Google Scholar](#)
99. Sinai S, et al (2021) **Generative AAV capsid diversification by latent interpolation** *bioRxiv* <https://doi.org/10.1101/2021.04.16.440236> | [Google Scholar](#)
100. Ruiz A, Ariño J (2007) **Function and Regulation of the *Saccharomyces cerevisiae* ENA Sodium ATPase System** *Eukaryotic Cell* **6**:2175 [Google Scholar](#)
101. Eskes E, et al (2018) **pH homeostasis in yeast; the phosphate perspective** *Current Genetics* **64**:155–161 [Google Scholar](#)

- [102] Deutschbauer Adam M., Davis Ronald W. (2005) **Quantitative trait loci mapped to single-nucleotide resolution in yeast** *Nature Genetics* **37**:1333–1340 <https://doi.org/10.1038/ng1674> | [Google Scholar](#)
- 103. Chaithanya KV, Sinha H (2023) **MKT1 alleles regulate stress responses through posttranscriptional modulation of Puf3 targets in budding yeast** *Yeast* **40**:616–627 <https://doi.org/10.1002/yea.3908> | [Google Scholar](#)
- [104] Szendro IG, et al (2013) **Quantitative analyses of empirical fitness landscapes** *J Stat Mech Theory Exp* <https://doi.org/10.1088/1742-5468/2013/01/P01005> | [Google Scholar](#)
- 105. Neidhart J, et al (2014) **Adaptation in Tunably Rugged Fitness Landscapes: The Rough Mount Fuji Model** *Genetics* **198**:699–721 [Google Scholar](#)
- 106. Yang KK, et al (2019) **Machine-learning-guided directed evolution for protein engineering** *Nat. Methods* **16**:687–694 [Google Scholar](#)
- 107. Groth PM, et al (2024) **Kermut: Composite kernel regression for protein variant effects** *arXiv* [Google Scholar](#)
- 108. Leslie CS, et al (2002) **The spectrum kernel: a string kernel for SVM protein classification** *Pac Symp Biocomput* :564–575 [Google Scholar](#)
- 109. Leslie CS, et al (2004) **Mismatch string kernels for discriminative protein classification** *Bioinformatics* **20**:467–476 [Google Scholar](#)
- 110. Haussler D (1999) **Convolution Kernels on Discrete Structures** [Google Scholar](#)
- [111] Amin Alan Nawzad, Weinstein Eli Nathan, Marks Debora Susan (2023) **Biological Sequence Kernels with Guaranteed Flexibility** *arXiv* [Google Scholar](#)
- 112. Gonzalez Somermeyer L, et al (2022) **Heterogeneity of the GFP fitness landscape and data-driven protein design** *eLife* **11**:e75842 [Google Scholar](#)
- 113. Faure AJ, Lehner B (2024) **MoCHI: neural networks to fit interpretable models and quantify energies, energetic couplings, epistasis, and allostery from deep mutational scanning data** *Genome Biology* **25**:303 [Google Scholar](#)
- 114. Hensman J, Matthews A, Ghahramani Z (2014) **Scalable Variational Gaussian Process Classification** *arXiv* [Google Scholar](#)
- [115] Kucukelbir Alp, Tran Dustin, Ranganath Rajesh, Gelman Andrew, Blei David M. (2016) **Automatic Differentiation Variational Inference** *arXiv* [Google Scholar](#)
- 116. Tonner PD, et al (2022) **Interpretable modeling of genotype–phenotype landscapes with state-of-the-art predictive power** *Proceedings of the National Academy of Sciences* **119**:e2114021119 [Google Scholar](#)
- 117. Costanzo M, et al (2010) **The genetic landscape of a cell** *Science* **327**:425–431 [Google Scholar](#)
- 118. Bloom JS, et al (2015) **Genetic interactions contribute less than additive effects to quantitative trait variation in yeast** *Nature Communications* **6**:8712 [Google Scholar](#)

119. Sheppard B, et al (2021) **A model and test for coordinated polygenic epistasis in complex traits** *Proceedings of the National Academy of Sciences* **118**:e1922305118 [Google Scholar](#)
- [120] Tang David, Freudenberg Jerome, Dahl Andy (2023) **Factorizing polygenic epistasis improves prediction and uncovers biological pathways in complex traits** *The American Journal of Human Genetics* **110**:1875–1887 [Google Scholar](#)
- [121] Hwang Sungmin, Schmiegel Benjamin, Ferretti Luca, Krug Joachim (2018) **Universality Classes of Interaction Structures for NK Fitness Landscapes** *Journal of Statistical Physics* **172**:226–278 [Google Scholar](#)
122. Rubin F, et al (2017) **A statistical framework for analyzing deep mutational scanning data** *Genome Biol* **18**:150 [Google Scholar](#)
123. Notin1 P, et al (2023) **ProteinGym: Large-Scale Benchmarks for Protein Design and Fitness Prediction** *bioRxiv* <https://doi.org/10.1101/2023.12.07.570727> | [Google Scholar](#)
- [124] Engel SR, et al (2022) **New data and collaborations at the Saccharomyces Genome Database: updated reference genome, alleles, and the Alliance of Genome Resources** *Genetics* **220**:iyab224 <https://doi.org/10.1093/genetics/iyab224> | [Google Scholar](#)

Author information

Juannan Zhou*

Department of Biology, University of Florida, Gainesville, United States

*These authors contributed equally to this work

Carlos Martí-Gómez*

Simons Center for Quantitative Biology, Cold Spring Harbor Laboratory, Cold Spring Harbor, United States

ORCID iD: [0000-0002-2042-843X](https://orcid.org/0000-0002-2042-843X)

*These authors contributed equally to this work

Samantha Petti

Department of Mathematics, Tufts University, Medford, United States

David M McCandlish

Simons Center for Quantitative Biology, Cold Spring Harbor Laboratory, Cold Spring Harbor, United States

For correspondence: mccandlish@cshl.edu

Editors

Reviewing Editor

Anne-Florence Bitbol

Ecole Polytechnique Federale de Lausanne (EPFL), Lausanne, Switzerland

Senior Editor

Alan Moses

University of Toronto, Toronto, Canada

Reviewer #1 (Public review):

Summary:

Zhou and colleagues introduce a series of generalized Gaussian process models for genotype-phenotype mapping. The goal was to develop models that were more powerful than standard linear models, while retaining explanatory power as opposed to neural network approaches. The novelty stems from choices of prior distributions (and I suppose fitted posteriors) that model epistasis based on some form of site/allele-specific modifier effect and genotype distance. The authors then apply their models to three empirical datasets, the GB1 antibody-binding dataset, the human 5' splice set dataset, and a yeast meiotic cross dataset, and find substantially improved variance explained while retaining strong explanatory power when compared to linear models.

Strengths:

The main strength of the manuscript lies in the development of the modeling approaches, as well as the evidence from the empirical dataset that the variance explained is improved.

Weaknesses:

The main weakness of the paper is that none of the models were tested on an *in silico* dataset where the ground truth is known. Therefore, it is unclear if their model actually retains any explanatory power.

Impact:

Genotype-phenotype mapping is a central point of genetics. However, the function is complex and unknown. Simple linear models can uncover some functional link between genes and their effects, but do so through severe oversimplification of the system. On the other hand, neural networks can, in principle, model the function perfectly, but it does so without easy interpretation. Gaussian regression is another approach that improves on linear regression, allowing better fitting of the data while allowing interpretation of the underlying alleles and their effects. This approach, now computable with state-of-the-art algorithms, will advance the field of genotype-to-phenotype associations.

<https://doi.org/10.7554/eLife.108964.1.sa2>

Reviewer #2 (Public review):

This paper builds on prior work by some of the same authors on how to model fitness landscapes in the presence of epistasis. They have previously shown how simply writing general expansions of fitness in terms of one-body plus two-body plus three-body, etc., terms often fails to generalize to good predictions. They have also previously introduced a Gaussian process regression approach regarding how much epistasis there should be of each order.

This paper contains several main advances:

- (1) They implement a more efficient form of the Gaussian process model fitting that uses GPUs and related algorithmic advances to enable better fitting of these models to datasets for larger sequences.
- (2) They provide a software package implementing the above.

(3) They generalize the models to allow the extent of epistasis associated with changes in sequence to depend on specific sites, alleles, and mutations.

(4) They show modest improvements in prediction and substantial improvements in interpretability with the more generalized models above.

Overall, while this paper is quite technical, my assessment is that it represents a substantial conceptual and algorithmic advance for the above reasons, and I would recommend only modest revisions. The paper seems well-written and clear, given the inherent complexity of this topic.

<https://doi.org/10.7554/eLife.108964.1.sa1>

Reviewer #3 (Public review):

Summary:

The authors propose three types of Gaussian process kernels that extend and generalize standard kernels used for sequence-function prediction tasks, giving rise to the connectedness, Jenga, and general product models. The associated hyperparameters are interpretable and represent epistatic effects of varying complexity. The proposed models significantly outperform the simpler baselines, including the additive model, pairwise interaction model, and Gaussian process with a geometric kernel, in terms of R^2 .

Strengths:

- (1) The demonstrated performance boost and improved scaling with increasing training data are compelling.
- (2) The hyperparameter selection step using the marginal likelihood, as implemented by the authors, seems to yield a reasonable hyperparameter combination that lends itself to biologically plausible interpretations.
- (3) The proposed kernels generalize existing kernels in domain-interpretable ways, and can correspond to cases that would not be "physical" in the original models (e.g., $\mu_p > 1$ in the original connectedness model that allows modeling of anticorrelated phenotypes).

Weaknesses:

- (1) While enabling uncertainty quantification is a key advantage of Gaussian processes, the authors do not present metrics specific to the predicted uncertainties; all metrics seem to concern the mean predictions only. It would be helpful to evaluate coverage metrics and maybe include an application of the uncertainties, such as in active learning or Bayesian optimization.
- (2) The more complex models, like the general product model, place a heavier burden on the hyperparameter selection step. Explicitly discussing the optimization routine used here would be helpful to potential users of the method and code.

<https://doi.org/10.7554/eLife.108964.1.sa0>