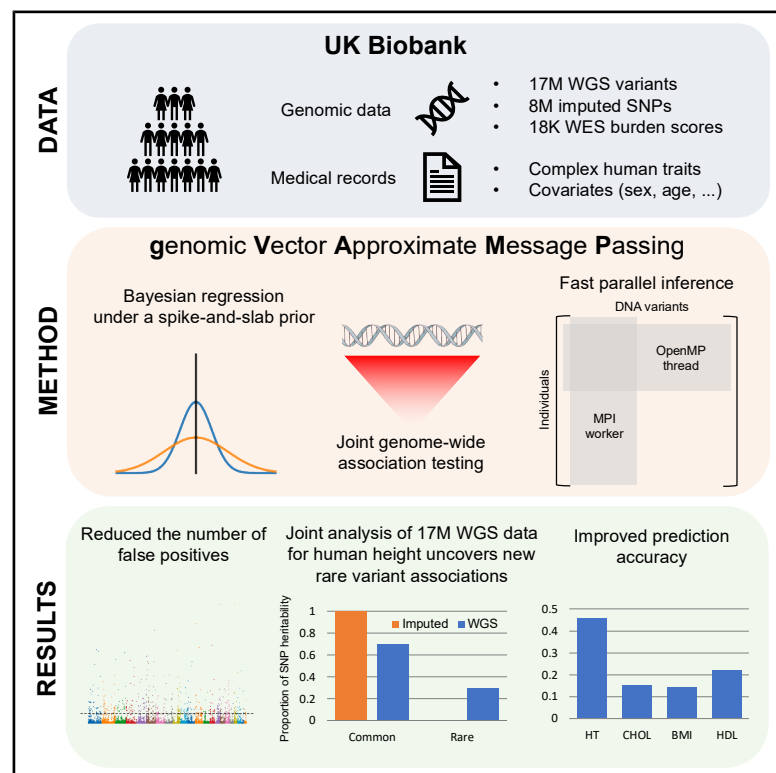


Joint modeling of whole-genome sequencing data for human height via approximate message passing

Graphical abstract



Authors

Al Depope, Jakub Bajzik, Marco Mondelli, Matthew R. Robinson

Correspondence

al.depope@ist.ac.at (A.D.), marco.mondelli@ist.ac.at (M.M.), mrobinso@ist.ac.at (M.R.R.)

In brief

Depope et al. develop a computationally efficient statistical algorithm for genomics that, when applied to half a million whole-genome sequences, reduces false discoveries and provides high reported accuracy for the prediction of human traits from the DNA.

Highlights

- Computationally efficient application of the AMP framework for genomics
- gVAMP gives the highest reported prediction accuracy for human height
- gVAMP predicts complex traits and identifies associated genetic variants
- 17 M WGS variants were analyzed jointly for human height



Technology

Joint modeling of whole-genome sequencing data for human height via approximate message passing

Al Depope,^{1,*} Jakub Bajzik,¹ Marco Mondelli,^{1,2,*} and Matthew R. Robinson^{1,2,3,*}¹Institute of Science and Technology Austria, Klosterneuburg, Austria²These authors contributed equally³Lead contact*Correspondence: al.depope@ist.ac.at (A.D.), marco.mondelli@ist.ac.at (M.M.), mrobinso@ist.ac.at (M.R.R.)<https://doi.org/10.1016/j.xgen.2026.101162>

SUMMARY

Human height is a model for the genetic analysis of complex traits, and recent studies suggest the presence of thousands of common genetic variant associations and hundreds of low-frequency/rare variants. Here, we develop a new algorithmic paradigm based on approximate message passing (genomic vector approximate message passing [gVAMP]) for identifying DNA sequence variants associated with complex traits and common diseases in large-scale whole-genome sequencing (WGS) data. We show that gVAMP accurately localizes associations to variants with the correct frequency and position in the DNA, outperforming existing fine-mapping methods in selecting the appropriate genetic variants within WGS data. We then apply gVAMP to jointly model the relationship of tens of millions of WGS variants with human height in hundreds of thousands of UK Biobank individuals. We identify 59 rare variants and gene burden scores alongside many hundreds of DNA regions containing common variant associations and show that understanding the genetic basis of complex traits will require the joint analysis of hundreds of millions of variables measured on millions of people. The polygenic risk scores obtained from gVAMP have high accuracy (including a prediction accuracy of ~46% for human height) and outperform current methods for downstream tasks such as mixed linear model association testing across 13 UK Biobank traits. In conclusion, gVAMP offers a scalable foundation for a wider range of analyses in WGS data.

INTRODUCTION

Efficient utilization of large-scale biobank data is crucial for inferring the genetic basis of disease and predicting health outcomes from DNA. By mapping associations to regions of the DNA, we seek to understand the underlying genetic architecture of phenotypes. The common approach is to use single-marker or single-gene burden score regression,^{1–4} which gives marginal associations that do not account for linkage disequilibrium (LD) among loci. Fine-mapping methods are then employed with the aim of controlling for LD when conducting variable selection and resolving the associations to a set of likely causal variants.⁵

Broadly speaking, inference for association testing,^{1–4} fine-mapping,⁵ and polygenic risk score^{6,7} tasks are performed using either Markov chain Monte Carlo (MCMC) or variational inference (VI) approaches. From a practical perspective, current software implementations of these algorithms limit either the number of markers or the number of individuals. Mixed linear model association (MLMA) approaches are typically restricted to using less than 1 M SNPs to control for the polygenic background,^{1,2} resulting in a loss of power⁷ and the potential for inadequate control for fine-scale confounding factors.⁸ Polygenic risk score algorithms are limited to a few million SNPs, and they lose power by modeling only blocks of genetic markers.^{9,10} Likewise, fine-

mapping methods are generally limited to focal segments of the DNA,⁵ and they are unable to fit all genome-wide DNA variants together. Thus, no existing approach can truly estimate all genetic effects jointly in large-scale whole-genome sequencing (WGS) data.

Here, we propose a new inference approach, genomic vector approximate message passing (gVAMP), a form of Bayesian whole-genome regression, capable of fitting tens of millions of WGS variants jointly for hundreds of thousands of individuals (see [STAR Methods](#) and [Methods S1](#)). We extensively benchmark gVAMP in simulations on WGS data and show that it accurately localizes associations to variants with the correct frequency and position in the DNA, outperforming marginal testing and marginal testing plus fine-mapping. We then describe an application of gVAMP to WGS data and measurements of human height for hundreds of thousands of UK Biobank¹¹ participants. We also validate the application of gVAMP for polygenic risk score prediction on a number of datasets by benchmarking against the state of the art. Here, gVAMP outperforms widely used summary statistic methods such as LDpred2⁹ and SBayesR¹⁰; its performance matches that of MCMC sampling schemes⁷ but with a dramatic speedup in time (analyzing 8.4 M SNPs jointly in hours as opposed to weeks); and it provides some improvement over a recently proposed individual-level VI



approach.¹² When polygenic risk scores are used for association testing in a mixed linear model framework, gVAMP also outperforms both the REGENIE¹ and LDAC-KVIK¹² software. In summary, our approach lays the foundation for a wider range of analyses in large WGS datasets.

DESIGN

Our focus is on joint modeling of all genetic variants genome-wide, controlling for local and long-range LD. To do this, we consider a general form of whole-genome Bayesian variable selection regression (BVSR), common to genome-wide association studies (GWASs),^{7,10} estimating the effects vector $\beta \in \mathbb{R}^P$ from a vector of normalized phenotype measurements $\mathbf{y} = (y_1, \dots, y_N) \in \mathbb{R}^N$ given by

$$y_i = \langle \mathbf{x}_i, \beta \rangle + \epsilon_i, \text{ for } i \in \{1, \dots, N\}. \quad (\text{Equation 1})$$

Here, $\mathbf{x}_i$ is the row of the normalized genotype matrix $\mathbf{X}$ corresponding to the i -th individual, $\langle \mathbf{x}_i, \beta \rangle = \mathbf{x}_i^T \beta$ denotes the inner product, and $\epsilon = (\epsilon_1, \dots, \epsilon_N)$ is an unknown noise vector with multivariate normal distribution $\mathcal{N}(0, \gamma_\epsilon^{-1} \cdot \mathbf{I})$ and unknown noise precision γ_ϵ^{-1} . We note that the normalized genotype design matrix was used in its original form, without correcting for sample covariances, because the analysis performed in this work targets a genetically homogeneous population. To allow for a range of genetic effects, we select the prior on β to be of an adaptive spike-and-slab form:

$$\beta_j \sim (1 - \lambda) \cdot \delta_0(\cdot) + \lambda \cdot \sum_{i=1}^L \pi_i \cdot \mathcal{N}(\cdot, 0, \sigma_i^2), \text{ for } j \in \{1, \dots, P\}, \quad (\text{Equation 2})$$

where we learn from the data the DNA variant inclusion rate $\lambda \in [0, 1]$, the unknown number of Gaussian mixtures L , the mixture probabilities $(\pi_i)_{i=1}^L$, and the variances for the slab components $(\sigma_i^2)_{i=1}^L$.

We apply the statistical model of Equations 1 and 2 by developing a new approach for GWAS inference, dubbed gVAMP. AMP^{13–15} refers to a family of iterative algorithms with several attractive properties: (1) AMP allows the usage of a wide range of Bayesian priors; (2) the AMP performance for high-dimensional data can be precisely characterized by a simple recursion called state evolution¹⁶; (3) using state evolution, joint association test statistics can be obtained¹⁷; and (4) AMP achieves Bayes-optimal performance in several settings^{17–19} (for an overview of AMP, see Methods S2). However, we find that existing AMP algorithms proposed for various applications^{20–23} cannot be transferred to biobank analyses as (1) they are entirely infeasible at scale, requiring expensive singular value decompositions, and (2) they give diverging estimates of the signal in either simulated genomic data or the UK Biobank data. To address the problem, we combine a number of principled approaches to produce an algorithm tailored to whole-genome regression, as described in the [method details](#) section of [STAR Methods](#) (Algorithm 1) and [Methods S3–S5](#).

gVAMP provides both (1) a point estimate of the posterior mean $\mathbb{E}[\beta | \mathbf{X}, \mathbf{y}]$, which can be directly used to create a polygenic risk score, and (2) a statistical framework, via state evolution, to

test whether each marker makes a non-trivial contribution to the outcome. Specifically, state evolution yields posterior inclusion probabilities (PIPs) for each SNP (see the [quantification and statistical analysis](#) section of [STAR Methods](#)). Finally, we learn all unknown parameters in an adaptive Bayes expectation maximization (EM) framework,^{25,26} which avoids expensive cross-validation and yields biologically informative inference of the phenotypic variance attributable to the genomic data (SNP heritability h_{SNP}^2), allowing for a full genome-wide characterization of the genetic architecture of human complex traits in WGS data.

RESULTS

Variable selection in WGS data

We begin by modeling the UK Biobank WGS data, curating a set of 16,854,878 WGS variants observed for 426,462 individuals of European genetic ancestry. To create this set, we remove individuals of first-degree ancestry to avoid close-familial environmental sharing, and because of cloud computing restrictions, we group highly correlated common variants with an absolute correlation of ≥ 0.6 , keeping all rare variation within the data with a minor-allele frequency (MAF) of ≥ 0.0001 (see the [method details](#) section of [STAR Methods](#)).

We first assess the ability of gVAMP to estimate genetic effects within this curated UK Biobank WGS dataset by simulating 12 different scenarios on top of data from chromosomes 11–15 (3,285,117 SNPs). The 12 scenarios differ in (1) the relationship of effect size and MAF, (2) the ratio of common to rare causal variants, and (3) the number of causal variants (500 or 1,000) that contribute 7.25% to the phenotypic variance (see the [method details](#) section of [STAR Methods](#) and [Table S1](#) for an overview of the scenarios). Chromosomes 11–15 represent $\sim 18.8\%$ of the total length of the DNA, and if we extrapolate genome-wide, then our simulated phenotypes would have $\sim 39\%$ phenotypic variance attributable to between approximately 2,700 and 5,300 underlying causal variants genome-wide. This gives an expected average variance contributed per DNA variant of $39\% / 5,300 = 0.0074\%$, comparable to that of the human height of $70\% / 10,000 = 0.007\%$ suggested by recent results.²⁷ These settings are chosen to give adequate power to select variables and to enable clear comparisons of fine-mapping methods. We also note that in [Note S1](#), we conduct extensive benchmarks of gVAMP in imputed SNP data for traits of highly polygenic architecture with different effect size distributions, different relationships of marker frequency and effect size, different numbers of causal variants, and different variances attributable to the SNPs, across 12 additional settings each with 4 sets of data, for a total of hundreds of models trained with a sample size of ~ 400 k and up to >8 M SNPs (see the [method details](#) section of [STAR Methods](#) and [Note S1](#)).

We analyze the WGS data with either (1) our proposed gVAMP or (2) the FINEMAP²⁸ and SuSiE variable selection models,⁵ run on 3 Mb regions surrounding each significant variant discovered via GWAS marginal testing one SNP at a time, following standard practice. With access to 2,240 virtual CPUs and just under 9 TB (35 parallel instances of 256 GB and 64 CPUs), running gVAMP for all replicates of the 12 scenarios requires 5 days,

Algorithm 1. gVAMP

```

1: Input: preprocessed normalized genotype matrix  $\mathbf{X} \in \mathbb{R}^{N \times P}$ , max number of iterations  $N_{it}$ , initial estimate of effect sizes  $\mathbf{r}_{1,0} = \mathbf{0}_P \in \mathbb{R}^P$ , initial estimate of effect sizes precision  $\gamma_{1,0} = 10^{-6} > 0$ , initial estimate of noise precision  $\gamma_{\epsilon,0} = 2$ , initial set of parameters defining the prior distribution  $\Theta_0 = \left\{ \lambda, \left( \pi_i^{(0)} \right)_{i=1}^L, \left( \sigma_i^{(0)} \right)_{i=1}^L \right\}$ , max number of variance auto-tuning steps  $N_{var\_tune} = 5 \in \mathbb{N}$ , threshold for stopping criterion  $\varepsilon = 10^{-4} > 0$ , damping factor  $\rho \in (0,1), \alpha_{2,-1} = 0$ .
2: for  $t = 0, 1, \dots, N_{it}$  do
3:   Denoising step
4:   for  $k = 0, 1, \dots, N_{var\_tune}$  do
5:      $\hat{\beta}_{1,t} = \mathbb{E}[\beta | \mathbf{r}_{1,t} = \beta + \mathcal{N}(0, \gamma_{1,t}^{-1} \mathbf{I}), \gamma_{1,t}, \Theta_t]$ 
6:     if  $t > 0$  then
7:       Variance auto-tuning step of estimation error for  $\beta$  in the denoising step, called  $\gamma_{1,t}$ 
8:       EM update of the prior distribution parameters  $\Theta$ , called  $\Theta_t$ 
9:       if  $|\gamma_{1,t} - \gamma_{1,t}^{(previous)}| < 10^{-3}$  then
10:         break
11:       end if
12:     end if
13:   end for
14:   if  $t \geq 0$  then
15:      $\hat{\beta}_{1,t} = \rho \cdot \hat{\beta}_{1,t} + (1 - \rho) \cdot \hat{\beta}_{1,t-1}$ 
16:   end if
17:    $\alpha_{1,t} = \gamma_{1,t} \cdot \text{Var}[\beta | \mathbf{r}_{1,t} = \beta + \mathcal{N}(0, \gamma_{1,t}^{-1} \mathbf{I}), \gamma_{1,t}, \Theta_t]$ 
18:    $\gamma_{2,t} = \gamma_{1,t} \cdot (1 - \alpha_{1,t}) / \alpha_{1,t}$ 
19:    $\mathbf{r}_{2,t} = (\hat{\beta}_{1,t} - \alpha_{1,t} \mathbf{r}_{1,t}) / (1 - \alpha_{1,t})$ 
20:    $\xi = \min(2 \cdot \min(\alpha_{1,t}, \alpha_{2,t-1}), 1.0)$  // Inspired by24
21:    $\rho = \max(\rho, \xi)$ 
22:   LMMSE step
23:    $\hat{\beta}_{2,t} = (\gamma_{\epsilon,t} \mathbf{X}^T \mathbf{X} + \gamma_{2,t} \mathbf{I})^{-1} (\gamma_{\epsilon,t} \mathbf{X}^T \mathbf{y} + \gamma_{2,t} \mathbf{r}_{2,t})$ 
24:    $\alpha_{2,t} = \gamma_{2,t} \cdot \text{Tr}[(\gamma_{\epsilon,t} \mathbf{X}^T \mathbf{X} + \gamma_{2,t} \mathbf{I})^{-1}] / P$ 
25:    $\gamma_{1,t+1} = \gamma_{2,t} \cdot (1 - \alpha_{2,t}) / \alpha_{2,t}$ 
26:   if  $t > 1$  then
27:     Variance auto-tuning step of estimation error for  $\beta$  in the LMMSE step, called  $\gamma_{2,t}$ 
28:   end if
29:    $\mathbf{r}_{1,t+1} = (\hat{\beta}_{2,t} - \alpha_{2,t} \mathbf{r}_{2,t}) / (1 - \alpha_{2,t})$ 
30:   EM update of the estimate of  $\gamma_{\epsilon}$ , called  $\gamma_{\epsilon,t}$ 
31:   if  $t \geq 1$  and  $\|\hat{\beta}_{1,t} - \hat{\beta}_{1,t-1}\|_2 / \|\hat{\beta}_{1,t-1}\|_2 < \varepsilon$  then
32:     break
33:   end if
34: end for
35: return  $\hat{\beta}_{1,t}$ 

```

with FINEMAP requiring 8 days. However, running SuSiE required 30 days to complete 4 scenarios (with standard LD and eigenvalue calculation in LDstore2; see the [method details](#) section of [STAR Methods](#)), and given its consistently comparable performance to FINEMAP, we restrict our comparisons with SuSiE to only the settings presented in [Figure 1](#). We first calculate the true positive rate (TPR = number of identified variants that are causal/number of causal variants) and the false discovery rate (FDR = number of identified variants that are not causal/number of identified variants) at a given threshold. gVAMP shows improved power (TPR) for identifying the correct underlying causal variants and a well-controlled FDR (FDR $\leq 5\%$ corresponding to a selection of variants with PIP, $\geq 95\%$), as compared to both the methods of FINEMAP and SuSiE across most scenarios ([Figures 1A and S1–S3](#)). FINEMAP shows a well-controlled FDR in only five of the 12 settings (settings $f, h,$

$j, k,$ and l in [Figures S1–S3](#)), where half of the simulated causal variants are common, and the power relationship of MAF and effect size is weaker. In contrast, we find only two settings where the FDR of gVAMP rises slightly above 5% (settings i and k in [Figures S1–S3](#)). In these two settings, the relationship between effect size and MAF is weakly negative, and variants are distributed randomly across the WGS data. As a result, most causal variants are rare, and each makes a very small contribution to the phenotypic variance, making them difficult to detect. Consequently, these are the lowest-powered settings, where the TPRs of both gVAMP and FINEMAP are the smallest, and the fewest variants overall are recovered ([Figure S1](#)). gVAMP outperforms FINEMAP in terms of TPR in all but one scenario, namely the challenging setting k , where the TPR of both methods is low ([Figures S1–S3](#)). SuSiE performs similarly to FINEMAP, as it is typically unable

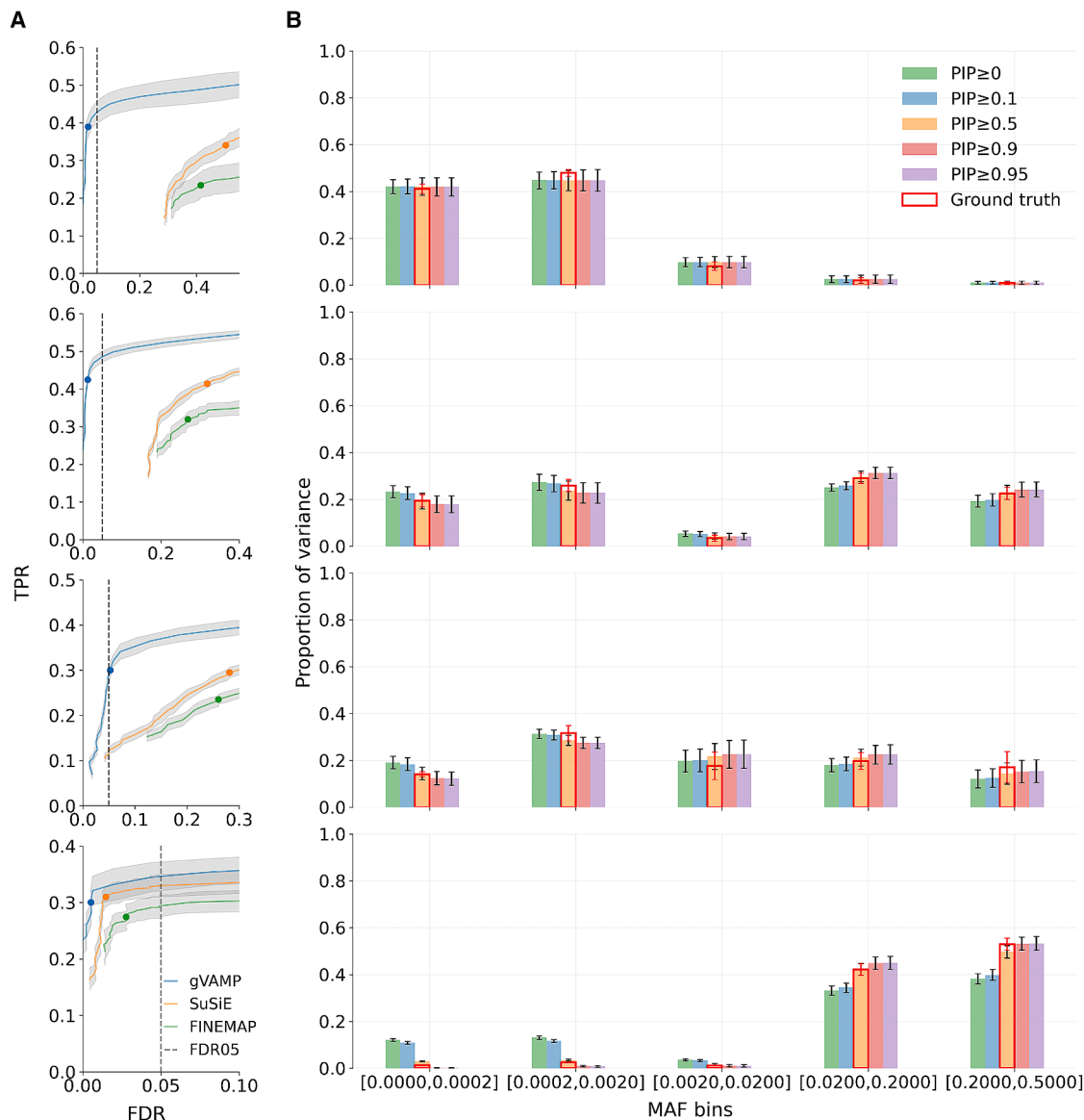


Figure 1. Simulation study of variable selection performance using whole-genome sequencing data of chromosomes 11–15 in the UK Biobank

(A) Comparison of the true positive rate (TPR; y axis) versus the false discovery rate (FDR; x axis) between gVAMP (blue) and marginal GWAS association testing plus fine-mapping, implemented in the FINEMAP (green) and SuSiE (orange) algorithms across four simulation scenarios. The rows of both (A) and (B) correspond to settings a, b, e, and f, which are described in the [method details](#) section of [STAR Methods](#), [Table S1](#), and [Figures S1–S3](#), alongside additional simulation settings. For gVAMP, variable selection is conducted from posterior inclusion probability (PIP) values obtained from a single model run genome-wide. For FINEMAP and SuSiE, we use a Bonferroni-corrected *p* value significance threshold for marginal GWAS testing, and then, having identified associated regions, we use the algorithms to obtain the PIP values to conduct variable selection. Error bands give the standard deviation across 5 simulation replicates. The blue, green, and orange dots give the TPR and FDR at a PIP ≥ 0.95, which should correspond to an FDR of ≤ 0.05 if the model correctly conducts variable selection based on PIP values.

(B) The proportion of variance attributable to DNA loci of different frequency. We plot the proportion of the sum of the squared regression coefficients attributable to variants of different frequency as estimated from gVAMP across different PIP thresholds. The ground truth (red frame) is calculated as the sum of squared true effects across different-frequency groups. The variance attributable to SNPs of different frequency matches the ground truth when selecting variants at a PIP ≥ 0.5. Error bars give the standard deviation across simulation replicates.

to control the FDR (Figure 1A). SuSiE has a higher TPR than FINEMAP, but it is still consistently outperformed by gVAMP, setting *f* being the only one where all methods perform similarly (Figure 1A).

The advantage of the gVAMP framework is confirmed in additional simulation comparisons. In particular, we find that infinitesimal versions of SuSiE and FINEMAP²⁹ that have been proposed in the literature also show only slight improvements on

the standard versions in all WGS settings we explored (Figure S7). Additionally, we find that variants selected by gVAMP at a PIP of ≥ 0.5 have a frequency distribution that reflects the simulated ground truth across scenarios (Figures 1B and S4). Furthermore, we validate the part of our analysis that combines WGS variants and whole-exome sequencing (WES) burden scores by performing an additional set of simulations on top of WGS variants from chromosomes 11–15 and WES burden scores corresponding to that region of the DNA. We simulated signals with a total heritability of $h^2 = 7.25\%$ by allocating 75% of 500 causal markers to WGS variants and 25% to WES burden scores, with the latter exhibiting effect sizes with variance 3-fold larger than those of individual WGS variants. This means that both groups should explain half the variance in settings where the variance of effect sizes is independent of the MAF. We find that we can correctly recover the percentage of variance recovered by each of the groups (Figure S6). Finally, when we simulate causal effects in both coding and non-coding regions, with 0%, 50%, or 100% of the causal variants located in coding (exonic) regions, we find that as the percentage of causal variants allocated to coding regions increases, so does the enrichment of the effects (Figure S7). Taken together, these findings suggest that existing methods fall short of the performance required for variable selection in WGS data, while gVAMP both reliably localizes the signal to the correct genomic regions and selects a variant of appropriate frequency.

The genetic architecture of human height in WGS data

We then apply gVAMP to model the genetic architecture of human height in our 16,854,878 WGS variant dataset. We find significant amounts of common genetic variation detected and a considerable contribution of rare genetic variants (Figure 2). Generally, the regions identified are replicated in the most recent GWAS of human height²⁷ within 100–500 kb (Figure 2A). We find that for variants of a frequency < 0.05 identified in the WGS analysis, variants discovered by Yengo et al.²⁷ of the strongest correlation and most similar MAF, significant at $p \leq 5 \times 10^{-8}$ within a 500 kb region, are more common in frequency (Figure 2B). Variants of a frequency < 0.005 discovered by gVAMP are likely not captured at all by Yengo et al. because these variants have lower average marginal chi-squared test statistic values than the other variants we identify (Figure 2C), and Yengo et al. set a MAF threshold of 1% for their meta-analysis. We estimate the maximum correlation within the UK Biobank WGS data of each gVAMP discovery with Yengo et al. SNPs significant at $p \leq 5 \times 10^{-8}$ within a 1 Mb window. We find that none of the variants of a frequency < 0.005 identified by gVAMP have a strong correlation with those identified by Yengo et al. (Figure S8). Thus, very rare gVAMP-identified variants are likely independent of the colocating common variants from Yengo et al., being simply a randomly selected set with a higher frequency (Figure 2B). In contrast, all gVAMP-identified variants of a frequency ≥ 0.05 are highly correlated to variants identified by Yengo et al. (Figure S8) with similar frequencies (Figure 2B). Thus, consistent with the literature,³⁰ common variant-tagged rare variant associations likely make up a small fraction of common-variant associations in Yengo et al., and taken together, our results suggest that discovered WGS regions are broadly the same as those previously identified.

19 out of the 21 genome-wide-significant gene burden scores are for genes that have previous GWAS height associations linked to them in Open Targets, but our analysis suggests that the effect is attributable to a rare protein-coding variant rather than the common variants suggested by current GWASs. This includes insulin-like growth factor 2 genes IGF1R and IGF2BP2 and known height genes ACAN, ADAMTS10 and ADAMTS17, CRISPLD1, DDR2, FLNB, HMG20B, LTBP2, LCORL, NF1, and NPR2. Novel genes are COL27A1, which encodes a collagen type XXVII alpha 1 chain, and HECTD1. The top genes, where gVAMP attributes $\geq 0.04\%$ of height variation in addition to those listed above, are EFEMP1, ZBTB38, and ZFAT. We find that many genome-wide-significant height-associated rare variants discovered in the WGS data were not found in previous UK Biobank analyses in Open Targets, including rs116467226, an intronic variant by TPRG1; rs766919361, an intronic variant by FGF18; rs141168133, an intergenic variant 19 kb from ID4; rs150556786, an upstream gene variant for GRM4; rs574843917, a non-coding transcript exon variant in GPR21; rs532230290, an intergenic variant 42 kb from SCYL1; rs543038394, intronic in OVOL1; rs1247942912, a non-coding transcript exon variant in AC024257.3; rs577630729, a regulatory region variant for ISG20; and rs140846043, a non-coding transcript exon variant of MIRLET7BHG. 10 out of our top 38 height-associated WGS variants of $\leq 1\%$ MAF were not previously discovered and only became height associated when conditioning on the entire polygenic background captured by WGS data within our analysis.

Generally, marginal estimates made in WGS data for height have similar joint effect sizes when included in the gVAMP model (Figure S9). We find high replication for variants directly observed within the All of Us study (Figure S10) when we compare the jointly fitted gVAMP estimates to simple marginal ordinary least squares (OLS) estimates made within All of Us, controlling for sex and 16 available genetic principal components (PCs) in Figure S10. The slope of the regression coefficients is 0.8471, which is a previously used metric for replication.³¹

We then compare our gVAMP WGS analysis to gVAMP run on two curated sets of 8,430,446 and 2,174,071 imputed SNPs for the same individuals (see the [method details](#) section of [STAR Methods](#)), but with an elevated MAF threshold, to contrast the findings obtained when setting different MAF thresholds and to explore the benefit of using WGS data over imputed SNP data. gVAMP estimates the proportion of phenotypic variance in human height attributable to WGS data as 0.652, a value that is comparable to 0.63 for imputed SNP data, to the previously published REML estimate in a different WGS dataset,³² and to family-based estimates.³³ This is expected, as the imputed data are a subset of the WGS data and very-low-frequency variants are expected to make little contribution to the population-level variance, so the addition of large numbers of rare variants will have little impact, especially if they capture a signal that was previously weakly tagged by more common variation.

gVAMP identifies 2,070 sets of variants when we apply the model to a set of 2,174,071 imputed SNPs at a PIP $\geq 95\%$, which is consistent with findings from MLMA testing described in the next section within imputed data (Table 2). We note that in the WGS data, however, gVAMP identifies

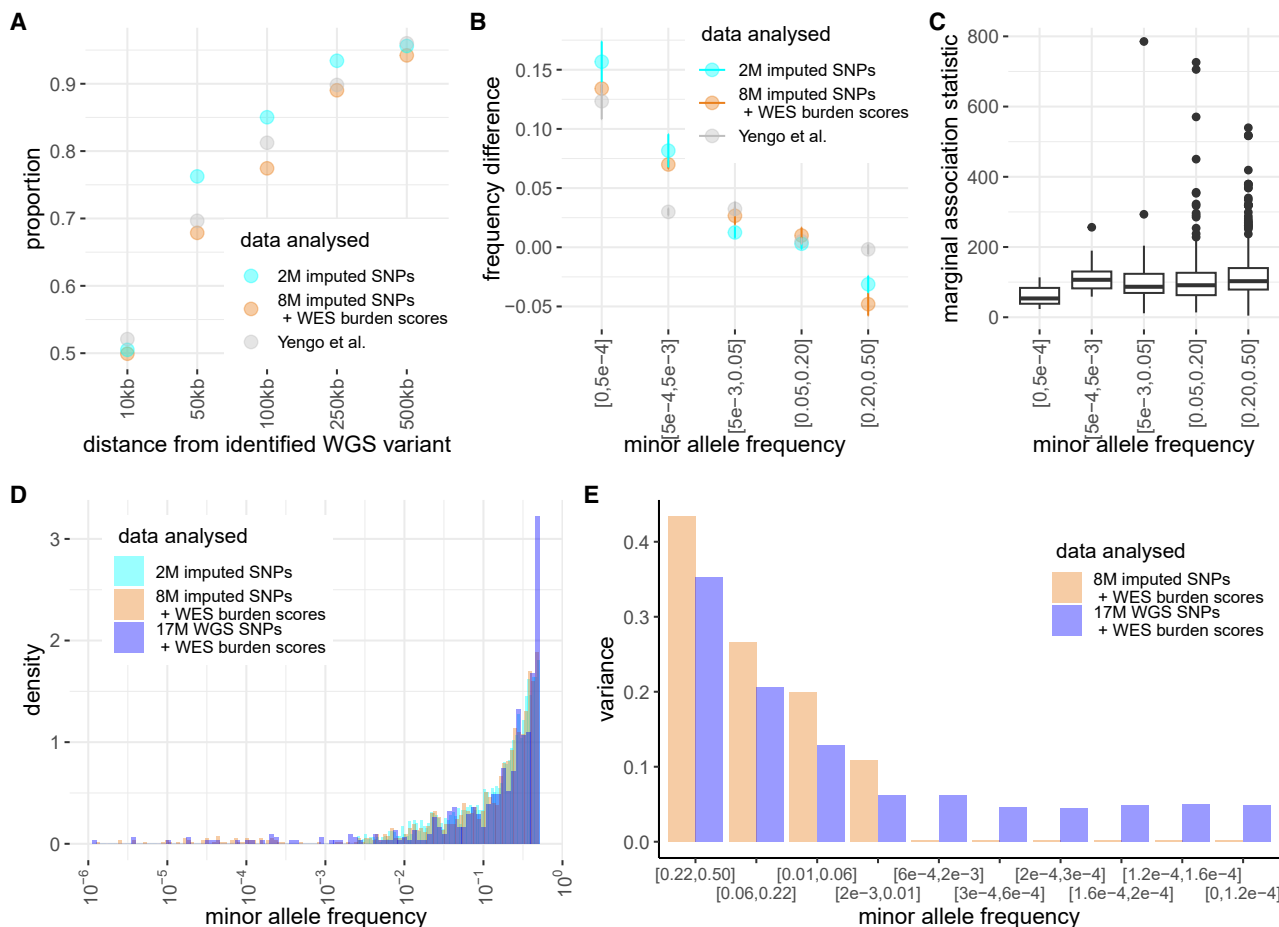


Figure 2. The genetic architecture of human height inferred by gVAMP from 16,854,878 whole-genome sequencing variants in the UK Biobank

(A) For each whole-genome sequencing (WGS) variant discovered at a PIP $\geq 95\%$, we determine whether we also identify a variant within a given number of base pairs (distance, x axis) from (1) the latest height GWAS by Yengo et al.,²⁷ selecting variants at significance $p \leq 5 \times 10^{-8}$ from the marginal summary statistics for European individuals excluding the UK Biobank, or (2) an analysis of imputed SNP data with gVAMP, selecting variants with a PIP $\geq 95\%$. The proportion overlap (y axis) is calculated as the proportion of the identified WGS variants that are discovered within a certain base-pair distance.

(B) For each WGS variant discovered at a PIP $\geq 95\%$, we determine (1) the variant discovered by Yengo et al. of the most similar minor-allele frequency (MAF) significant at $p \leq 5 \times 10^{-8}$ within a 500 kb region or (2) the imputed variant of most similar MAF discovered at a PIP $\geq 95\%$ within a 500 kb region. We then calculate the mean difference in MAF of the imputed/GWAS variant and the corresponding WGS variant, with error bars showing twice the standard error.

(C) For each WGS variant discovered at a PIP $\geq 95\%$, we determine the maximum marginal association test statistic of any variant with an absolute correlation of ≥ 0.6 . Boxplots of these maximum chi-squared test statistic values are then given for different MAF groups, showing a drop in average power to detect rare variants.

(D) Density plot of the MAF distribution of all discovered variants at a PIP $\geq 95\%$ for 2,174,071 imputed SNPs ("2 M imputed SNPs"), 8,430,446 imputed SNPs fit jointly alongside 17,852 whole-exome sequencing (WES) gene burden scores ("8 M imputed SNPs + WES burden scores"), and 16,854,878 WGS SNPs fit jointly alongside 17,852 WES gene burden scores ("17 M imputed SNPs + WES burden scores"). Irrespective of the data analyzed, the frequency distributions are similar, with the exception that the WGS analysis identifies a greater number of variants (both SNPs and gene burden scores) at lower frequencies.

(E) We calculate the proportional contribution to phenotypic variance across different MAF groups for 8,430,446 imputed SNPs fit jointly alongside 17,852 WES gene burden scores ("8 M imputed SNPs + WES burden scores") and 16,854,878 WGS SNPs fit jointly alongside 17,852 WES gene burden scores ("17 M imputed SNPs + WES burden scores"). A larger proportion of variance is attributable to rare variation at the expense of more common variants in WGS data.

fewer genome-wide significant effects, with only 501 found at a PIP $\geq 95\%$. There are two potential explanations for this. The first is that some of the height associations are re-allocated to DNA variants of lower MAF, and we then have lower power to detect them above a PIP $\geq 95\%$ threshold. The second explanation is that two variants can have very low correlation with each other in the imputed data while still being

moderately correlated with a third SNP. If the third SNP is not in the imputed panel but is height associated, this would result in two quasi-independent associations in the imputed data, but only one in the WGS data. We find that our 501 findings are within 50 kb of 1,042 of the 2,070 sets of variants identified by gVAMP in the 2,174,071 imputed SNP data and are within 500 kb of 1,989 of the 2,070 variant sets. Thus,

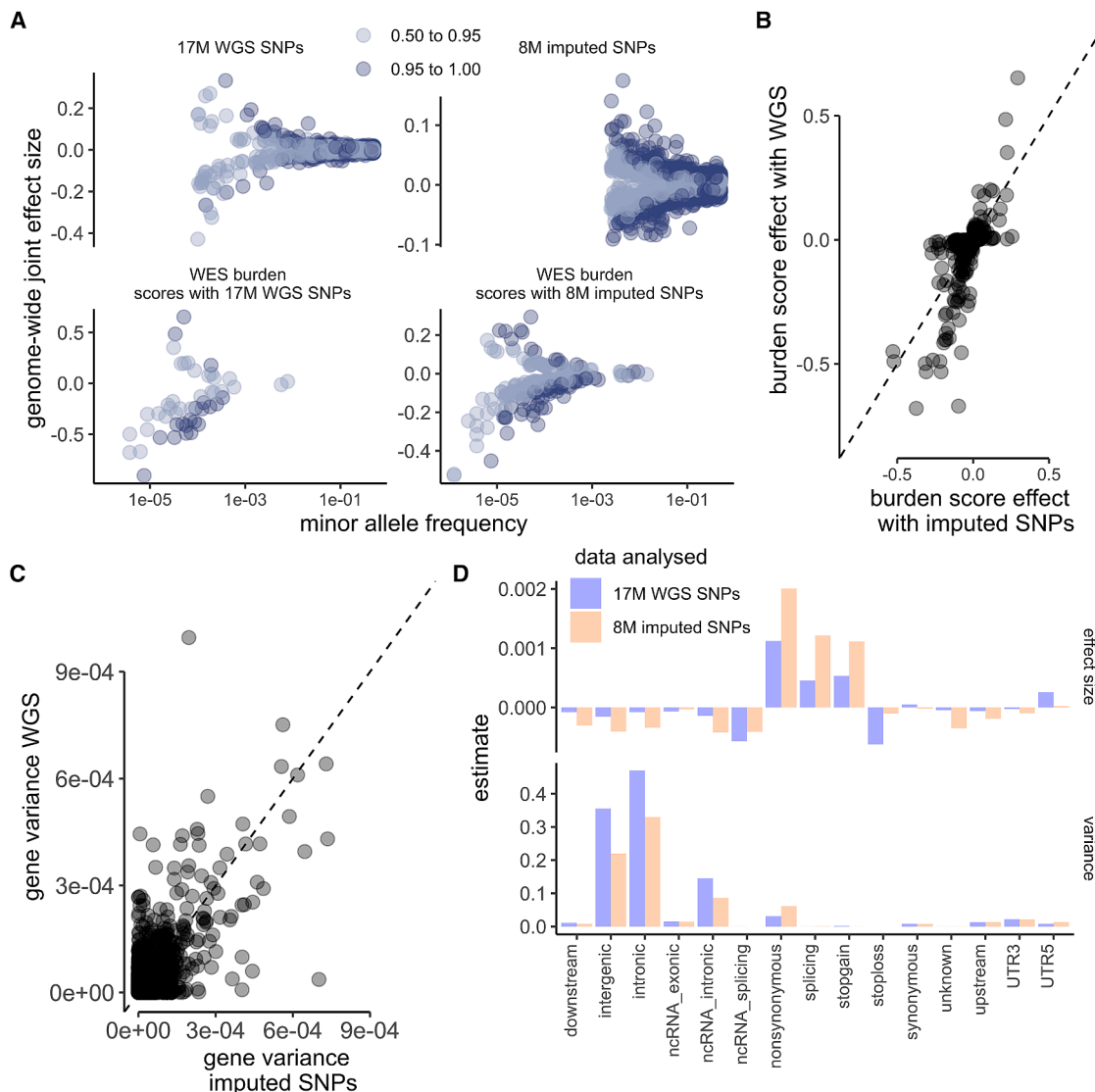


Figure 3. The relationship between effect size and minor-allele frequency, burden score effects, and the genetic architecture of human height inferred from 16,854,878 whole-genome sequencing variants in the UK Biobank

(A) For different posterior inclusion probability thresholds (0.5–0.95 and ≥ 0.95), we plot the relationship between joint effect size and minor-allele frequency (MAF) for 16,854,878 WGS variants (“17 M WGS SNPs”), 8,430,446 imputed SNP variants (“8 M imputed SNPs”), and 17,852 whole-exome sequencing (WES) gene burden scores fit alongside either the WGS or imputed variants.

(B) Across 17,852 WES gene burden scores, we find general concordance of the estimated effect sizes when fit alongside either imputed or WGS data, for most but not all genes, with a squared correlation of 0.532 and a regression slope of 0.993 (0.007 SE). The dashed line gives $y = x$.

(C) Likewise, we also find general concordance of the phenotypic variance attributable to markers annotated to each of the 17,852 genes, when fit conditionally on either imputed or WGS variants, with a squared correlation of 0.494 and a regression slope of 0.698 (0.004 SE). The dashed line gives $y = x$.

(D) Finally, we show similar patterns of enrichment when annotating markers to functional annotations in either the average effect size relative to the genome-wide average effect size (labeled “effect size”) or the proportion of variance attributable to each group (labeled “variance”), from joint estimation of either imputed or WGS data.

the WGS gVAMP findings are often neighbored by multiple variants discovered by gVAMP in the imputed SNP data. The WGS analysis identifies a greater proportion of variants (both SNPs and gene burden scores) at lower frequencies (Figure 2D) and finds a larger proportion of variance attributable to rare variation at the expense of slightly less variation attributable to common variants (Figure 2E).

We show that rare effects are larger than the effects of common variants when both are fit jointly (Figure 3A). We estimate the effects of 17,852 WES gene burden scores when fit alongside the WGS data, finding again that conditioning on rare variation enhances joint effect sizes (Figures 3A and 3B). We find a general concordance of the estimated effect sizes for most but not all WES gene burden scores (Figure 3B). WES gene burden

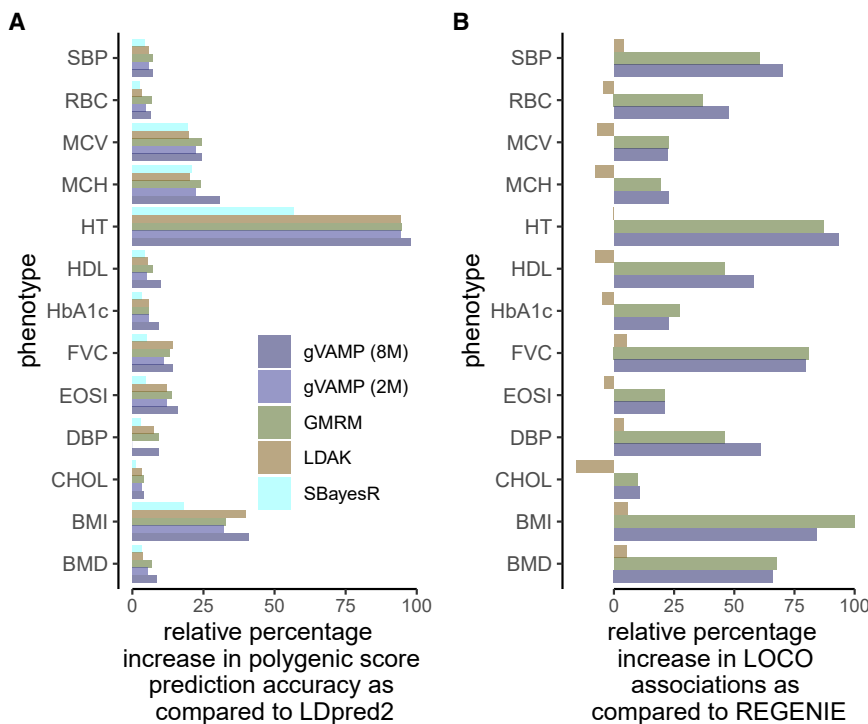


Figure 4. Validating gVAMP through polygenic risk score accuracy and association testing benchmarks in the UK Biobank within imputed SNP data

(A) Relative prediction accuracy of gVAMP in a hold-out set of the UK Biobank across 13 traits as compared to other approaches.

(B) Relative number of leave-one-chromosome-out (LOCO) associations of gVAMP across 13 UK Biobank traits as compared to other approaches at 8,430,446 markers. Trait codes are given in [Table S2](#).

incorporate more rare variation so that effects can be appropriately assigned.

Creating genetic predictors via gVAMP

We validate and benchmark gVAMP against state-of-the-art polygenic risk score prediction approaches in a number of alternative datasets. Applying our model to 13 traits in the UK Biobank imputed data (training data sample size and trait codes are given in [Table S2](#)), we compare the prediction accuracy of gVAMP to the widely used summary statistic methods LDpred2⁹ and SBayesR¹⁰ and to the individual-level methods GMRM⁷ and LDAK.¹² gVAMP matches or outperforms all methods for most phenotypes ([Figure 4A](#); [Table 1](#), with the best prediction accuracy highlighted in bold). Specifically, for human height (HT), our prediction accuracy is the highest, and individual-level approaches clearly outperform summary statistic ones: we obtain an accuracy of 45.7%, which is a 97.8% relative increase over LDpred2 and a 26.2% relative increase over SBayesR ([Figure 4A](#); [Table 1](#)). We also note that our polygenic risk score has higher accuracy than that obtained from SBayesR using the latest height GWAS of 3.5 M people of 44.7%,²⁷ despite our sample size of only 414,055.

The polygenic risk score accuracy presented may represent the limit of what is achievable for these traits at the UK Biobank sample size. While modeling 2.17 M SNPs improves accuracy over modeling 880,000 variants, we find rather small differences when going from 2.17 M to 8 M. This is because the 2.17 M dataset is an LD-grouped subset of the 8 M data, and the addition of highly correlated variants captures little additional phenotypic variance. gVAMP estimates the h^2_{SNP} of each of the 13 traits at an average of 3.4% less than GMRM when using 2,174,071 SNP markers but at an average of only 3.9% greater than GMRM when using 8,430,446 SNP markers. This implies that only slightly more of the phenotypic variance is captured by the SNPs when the full imputed SNP data are used ([Figure S12](#)). In summary, gVAMP (an adaptive mixture) performs similarly to LDAK (an elastic net) and GMRM (an MCMC sampling algorithm), improving over both when analyzing the full set of 8,430,446 imputed SNPs ([Table 1](#); [Figure 4A](#)).

score effects can differ in the WGS analysis by being either (1) more strongly associated because of improved conditioning on the genetic background or (2) attenuated because the variants that comprise the burden score are also included as variants in the analysis. We sum up the joint effects to determine the variation in height attributable to each gene, and we find a general concordance across markers annotated to each of the 17,852 genes, conditional on all other markers, across the imputed SNP and WGS analysis ([Figure 3C](#)). We see little relationship between the variance attributable to the burden score of a given gene and the SNPs annotated to the gene ([Figure S10](#)), with some notable exceptions being ACAN, ADAMTS17, ADAMTS10, and LCORL.

Additionally, across annotations, we find that the phenotypic variance attributable to different DNA regions is higher for intergenic and intronic variants ([Figure 3D](#)). However, when adjusting for the number of SNPs contained by each category by using the average effect size of the group relative to the average effect size genome-wide, we find that exonic variants that are nonsynonymous, splicing, and stop-gain have the largest average effects of the WGS variants included within our model ([Figure 3D](#)).

Finally, across annotations, we find similar patterns across both WGS and imputed SNP analyses ([Figure 3D](#)). Taken together, we find strong concordance across WGS and imputed data analyses performed on the same people, but we highlight that MAF cutoffs influence the variants to which effects are assigned and thus where phenotypic variance is allocated. Therefore, we expect WGS data to provide a more accurate characterization of the genetic architecture of human traits, but very large sample sizes are required to

Table 1. Polygenic risk score prediction accuracy R^2 for 13 different traits from statistical models trained in the UK Biobank data and tested in a UK Biobank hold-out set

Phen	LDpred	SBayesR	LDAK 2.17 M	GMRM 2.17 M	gVAMP 880,000	gVAMP 2.17 M	gVAMP 8 M
CHOL	0.147	0.149	0.152	0.153	0.140	0.152	0.153
EOSI	0.107	0.112	0.120	0.122	0.114	0.120	0.124
HbA1c	0.087	0.090	0.092	0.092	0.085	0.092	0.095
HDL	0.199	0.208	0.210	0.213	0.192	0.209	0.219
MCH	0.178	0.215	0.214	0.221	0.203	0.218	0.223
MCV	0.196	0.234	0.235	0.244	0.222	0.240	0.244
RBC	0.186	0.191	0.192	0.199	0.182	0.195	0.198
BMI	0.100	0.118	0.140	0.133	0.107	0.132	0.141
DBP	0.065	0.067	0.070	0.071	0.058	0.065	0.071
FVC	0.098	0.103	0.112	0.111	0.097	0.109	0.112
BMD	0.188	0.194	0.195	0.201	0.183	0.198	0.204
HT	0.231	0.362	0.449	0.450	0.419	0.449	0.457
SBP	0.068	0.071	0.072	0.073	0.061	0.072	0.073

Training data sample size and trait codes are given in Table S2 for each trait. The sample size of the hold-out test set is 15,000 for all phenotypes. LDpred2 and SBayesR give estimates obtained from the LDpred2 and SBayesR software, respectively, using summary statistic data of 8,430,446 SNPs obtained from the REGENIE software. LDAK gives the estimates obtained by the LDAK software using the elastic net predictor and individual-level data for 2,174,071 SNP markers. GMRM denotes estimates obtained from a Bayesian mixture model at 2,174,071 SNP markers ("GMRM 2 M"). gVAMP denotes estimates obtained from the adaptive EM Bayesian mixture model within a vector approximate message passing (VAMP) framework presented in this work, using either 887,060 ("gVAMP 880,000"), 2,174,071 ("gVAMP 2 M"), or 8,430,446 ("gVAMP 8 M") SNP markers. The best prediction accuracy is in bold.

Next, we compare gVAMP to REGENIE, GMRM, and LDAK-KVIK for association testing of the 13 traits in the UK Biobank imputed data. gVAMP performs similarly to MLMA leave-one-chromosome-out (MLMA-LOCO) testing conducted using a predictor from GMRM, with the use of the full 8,430,446 imputed SNP markers improving performance (Table 2; Figure 4B). REGENIE yields far fewer associations than either GMRM or gVAMP for all traits (Table 2; Figure 4B), with LDAK-KVIK being equivalent to REGENIE. We additionally verify that placing the predictors obtained by gVAMP into the second step of the LDAK-KVIK algorithm also yields improved performance as compared to the LDAK model (Table S3). Finally, using gVAMP to conduct variable selection, we find that hundreds of the LOCO marker associations for each trait can be localized (Table 2), with the obvious caveat that these results are for imputed SNP data, and re-analysis of WGS data may yield different results, as we show above for height. Using height as an example, we find that 3,367 of the 5,242 MLMA-LOCO associations can be localized directly to the 2,070 imputed variant sets discovered by gVAMP via the PIP testing within 50 kb. Thus, the gVAMP findings obtained by performing variable selection via the state evolution framework are often neighbored by multiple variants discovered by MLMA-LOCO in the imputed SNP data.

The increase in MLMA-LOCO findings for gVAMP over LDAK and REGENIE within the UK Biobank is consistent with the extensive simulation study results presented in Note S1. There, we show that MLMA-LOCO testing power scales with out-of-sample prediction accuracy only for gVAMP and not for the other methods (Figures S13 and S14). Including more predictors in gVAMP to generate the polygenic risk scores used in the first stage of MLMA-LOCO generally improves out-of-sam-

ple prediction accuracy (Figure S13) and association testing power while maintaining controlled type I error (Figures S14–S16). Additional simulations varying the number of markers included in the model lead to similar findings (Figures S17–S19). We also find that, in practice, our framework has similar performance to an MCMC approach, GMRM,⁷ that has a similar mixture prior (Figure S20).

In terms of compute time and resource use, gVAMP completes in less than 1/2 of the time given the same data and compute resources as compared to a single-trait analysis with REGENIE (Figure S21A), it is dramatically faster than GMRM (12.5× speedup; Figure S20C), and it is marginally faster than LDAK when run on our simulated data (Figure S21A). Across the 13 UK Biobank traits, gVAMP is faster than LDAK-KVIK, but it requires more random-access memory (RAM) (roughly the PLINK file size on disk; Figure S21B). If run on cloud computing at the current DNAnexus rates, step 1 of REGENIE can accommodate all 13 UK Biobank traits together, requiring 75.6 h with 30 threads, 15 GB RAM, and 421 GB disk space, which costs 75 GBP for the whole analysis or 5.77 GBP per trait. LDAK-KVIK analysis of 13 UK Biobank traits with 2.17 M SNPs using 15 threads requires up to 6 GB of RAM and 210 GB of storage, with an average runtime across the 13 traits of 26 h, giving an average cost of 11.47 GBP per trait. Increasing the CPU given to LDAK to 60, as shown in Figure S21B, reduces the average run time to 14.5 h but raises costs to 28.57 GBP per trait. GMRM analysis of 2.17 M SNP markers requires 82 h using 60 threads and 210 GB RAM and disk space, which costs 163.59 GBP per trait. In comparison, gVAMP analysis of 2.17 M SNP markers requires 16 h on average with 32 threads and 210 GB RAM and disk space, which costs

Table 2. Genome-wide significant associations for 13 UK Biobank traits from the software GMRM, gVAMP, LDAK-KVIK, and REGENIE at 8,430,446 genetic variants from either mixed linear model association leave-one-chromosome-out testing or gVAMP variable selection

Trait	REGENIE 880,000 LOCO	GMRM 2M LOCO	LDAK 2M LOCO	gVAMP 2M LOCO	gVAMP 8M LOCO	gVAMP PIP
CHOL	571	627	482	603	632	381
EOSI	572	693	549	630	693	527
HbA1c	337	429	321	385	413	227
HDL	692	1,009	639	812	1,092	536
MCH	746	889	690	810	916	691
MCV	970	1,190	903	1,062	1,185	898
RBC	897	1,229	856	1,079	1,325	651
BMI	688	1,376	726	1,175	1,266	261
DBP	291	425	303	419	468	315
FVC	549	994	577	721	986	318
BMD	522	874	549	668	867	491
HT	2,712	5,070	2,703	4,452	5,242	2,070
SBP	311	499	323	388	529	245

“REGENIE 880,000 LOCO” denotes results obtained from MLMA-LOCO testing using the REGENIE software, with 882,727 SNP markers used for step 1 and 8,430,446 markers used for the LOCO testing of step 2. “GMRM 2 M LOCO” refers to MLMA-LOCO testing at 8,430,446 SNPs, using a Bayesian MCMC mixture model in step 1, with 2,174,071 SNP markers. “LDAK 2 M LOCO” refers to MLMA-LOCO testing at 8,430,446 SNPs, using the LDAK-KVIK software in step 1, with 2,174,071 SNP markers. “gVAMP 2 M LOCO” refers to MLMA-LOCO testing at 8,430,446 SNPs, using the framework presented here, where in step 1, 2,174,071 markers are used, and the second step is identical to that of REGENIE. We then extend that to using all 8,430,446 SNP markers (“gVAMP 8 M LOCO”) in step 1. We also present gVAMP posterior inclusion probability (PIP) testing when applying a correlation threshold on the markers prior to analysis of 0.9 to the SNP markers (2,174,071 markers). For LOCO testing, the values give the number of genome-wide significant linkage disequilibrium independent associations selected based upon a standard p value threshold of less than 5×10^{-8} and R^2 between SNPs in a 1 Mb genomic segment of less than 1%. For the PIP testing, values give the number of genome-wide significant associations selected based on a ≥ 0.95 threshold. Bold indicates the best-performing method.

20.59 GBP per trait. Increasing the CPU to 60, as shown in Figure S21B, reduces the average run time to 11 h but slightly raises the cost to 21.89 GBP per trait.

DISCUSSION

Here, we develop an adaptive form of BVSR capable of fitting tens of millions of WGS variants jointly for hundreds of thousands of individuals. For human height, a joint analysis of a set of 17 M WGS variants localizes associations to numerous rare variants and gene burden scores, as well as hundreds of regions containing common variation. gVAMP provides a point estimate from the posterior distribution, and as such, its performance is expected to match that of standard MCMC Gibbs sampling algorithms. We show that this is the case in imputed SNP data: gVAMP and GMRM have comparable out-of-sample prediction accuracy, comparable estimates of the proportion of variance attributable to the DNA, comparable numbers of MLMA-LOCO findings, and comparable performance in our simulation study. However, unlike Gibbs sampling algorithms, which cannot scale to current data sizes, gVAMP exhibits a remarkable speedup, which allows high-dimensional biobank data with full WGS observations to be analyzed. Taken together, gVAMP offers a flexible framework suitable for (1) conducting variable selection in WGS data, where SuSiE, FINEMAP, and their variants are miscalibrated, and (2) creating polygenic risk scores with high accuracy that can also be used for downstream tasks such as MLMA-LOCO testing, with improved performance over LDAK-KVIK and REGENIE.

Limitations of the study

There are a number of remaining limitations and possible ways to expand and improve the analyses presented here. Our results suggest that the prioritization of relevant regions, genes, and gene sets is feasible at current sample sizes, but the detection of specific height-associated mutations is expected to be a challenging task. There is a very large number of rare variants within the human population that are missing from our analysis of 16,854,878 WGS variants. While conducting variable selection genome-wide for these will require millions of WGS samples, it may be beneficial to include as many of these rare variants as possible. Our algorithm requires RAM that is the same as the PLINK data file on disk, and although we evidence that this is not prohibitive to running the algorithm on the UK Biobank cloud computing, two alternative solutions could be to (1) apply gVAMP region by region or (2) make slight modifications in the implementation, so that gVAMP streams data at the cost of a slightly increased run time. Our ongoing work focuses on exploring these options.

A further limitation is that we have exclusively tested and analyzed individuals of European ancestry. Combining data of different ancestries or utilizing admixed populations is of great importance,³⁴ and previous work suggests that better modeling within a single large biobank can facilitate improved association testing in other global biobanks.³⁵ However, tackling this problem requires the development of a modeling framework that explicitly allows for differences in MAF, LD, and effect size across the human population, which we leave to future work.

Additionally, we removed first-degree relatives from all of our analyses to avoid confounding shared environmental effects. While some relatedness still exists even after removing first-degree relatives, we show in simulations that gVAMP performs well in both WGS and imputed data when planting a signal on real observed genotype data. Our motivation is that genotype-phenotype data of close relatives are often utilized for follow-up studies, and there are unique models (to which our framework could be adapted) employed for such tasks. Additionally, removing first-degree relatives loses 39,592 individuals from the UK Biobank, which does not alter the sample size in a significant way. We note that methods such as fastGWA⁴ or the equivalent implementation in LDK could be used, where a first step calculates the parameters needed to control for environmental confounding before the model is run. Furthermore, our method currently requires LD pruning of highly correlated common variants to improve numerical stability and convergence. While this limits the ability of gVAMP to fine-map regions with extremely high correlations, we note that when rare variants are included in the analysis, gVAMP demonstrates superior performance compared to FINEMAP, SuSiE, and their respective variants.

Finally, we also now seek to extend the model presented here to different outcome distributions (binary outcomes, time to event, count data, etc.), to have a grouped prior for annotations, to model multiple outcomes jointly, and to do all of this using summary statistics as well as individual-level data across different biobanks. Progress in these areas has been made in the AMP literature, see, e.g., Schnitter et al.³⁶ and Fletcher et al.³⁷ Combining datasets together in a principled way without a loss of accuracy, whilst preserving privacy, is key to obtaining the sample sizes that are likely required to fully explore the genetic basis of complex traits.

Summary

In summary, gVAMP is a different way to create genetic predictors and conduct variable selection. With increasing sample sizes reducing standard errors, a vast number of genomic regions are being identified as significantly associated with trait outcomes by one-SNP-at-a-time association testing. Such large numbers of findings will make it increasingly difficult to determine the relative importance of a given mutation, especially in WGS data with dense, highly correlated variants. Thus, it is crucial to develop statistical approaches fitting all variants jointly and asking whether, given the LD structure of the data, there is evidence for an effect at each locus.

RESOURCE AVAILABILITY

Lead contact

Requests for further information and resources should be directed to and will be fulfilled by the lead contact, Matthew R. Robinson (mrobinso@ist.ac.at).

Materials availability

This study did not generate new materials.

Data and code availability

- This project uses the UK Biobank data under project number 35520. UK Biobank genotypic and phenotypic data are available through a formal

request at <http://www.ukbiobank.ac.uk>. It also uses genotypic and phenotypic data from the All of Us study, which are also available through a formal request at <https://www.researchallofus.org/data-tools/data-access/>. All summary statistic estimates are released publicly on Dryad: <https://doi.org/10.5061/dryad.cz8w9gjjc>.

- The gVAMP code developed in this work is open source and has been deposited on GitHub, where it is publicly available at <https://github.com/medical-genomics-group/gVAMP>, and the code used to generate the data in the manuscript are available from Zenodo <https://doi.org/10.5281/zenodo.17935521>. The URLs of other software used are listed in the [key resources table](#).

ACKNOWLEDGMENTS

We thank Malgorzata Borczyk for creating the gene burden scores. We thank Robin Beaumont, Amedeo Roberto Esposito, Gareth Hawkes, Philip Schnitter, Matthew Stephens, Pragya Sur, Peter Visscher, Michael Weedon, and Harry Wright for providing valuable suggestions and comments on earlier versions of the work. This project was funded by a Lopez-Loreta Prize to M.M., an SNSF Eccellenza Grant to M.R.R. (PCEGP3-181181), an ERC Starting Grant to M.M. (INF², project number 101161364), and core funding from ISTA. High-performance computing was supported by the Scientific Service Units (SSU) of ISTA through resources provided by Scientific Computing (SciComp). We would like to acknowledge the participants and investigators of the UK Biobank study. We gratefully acknowledge the All of Us participants for their contributions, without whom this research would not have been possible. We also thank the National Institutes of Health All of Us Research Program for making available the participant data (and/or samples and/or cohort) examined in this study.

AUTHOR CONTRIBUTIONS

M.M. and M.R.R. conceived the study. A.D., M.M., and M.R.R. designed the study. A.D. derived the model and the algorithm, with input from M.M. and M.R.R. A.D. wrote the software with input from M.M. and M.R.R. A.D., J.B., M.M., and M.R.R. conducted the analysis and wrote the paper. All authors approved the final manuscript prior to submission.

DECLARATION OF INTERESTS

M.R.R. receives research funding from Boehringer Ingelheim for work that is unrelated to that presented here.

STAR★METHODS

Detailed methods are provided in the online version of this paper and include the following:

- [KEY RESOURCES TABLE](#)
- [EXPERIMENTAL MODEL AND STUDY PARTICIPANT DETAILS](#)
 - Participant inclusion in UK Biobank
 - Whole-genome sequencing (WGS) data
 - Imputed SNP data
 - Whole-exome sequencing (WES) data burden scores
 - Phenotypic records
- [METHOD DETAILS](#)
 - gVAMP installation
 - gVAMP algorithm overview
 - gVAMP implementation
 - gVAMP algorithm stability
 - gVAMP model parameters for the analysis on imputed data
 - gVAMP model parameters for the analysis of human height in WGS data
 - Polygenic risk scores and SNP heritability
 - Comparison to polygenic risk scores obtained by other methods
 - Replication of results in All of Us
 - Simulation study for WGS data

- QUANTIFICATION AND STATISTICAL ANALYSIS
 - gVAMP association testing
 - Mixed linear model association testing
- ADDITIONAL RESOURCES

SUPPLEMENTAL INFORMATION

Supplemental information can be found online at <https://doi.org/10.1016/j.xgen.2026.101162>.

Received: August 11, 2025

Revised: November 7, 2025

Accepted: January 13, 2026

Published: February 18, 2026

REFERENCES

1. Mbatchou, J., Barnard, L., Backman, J., Marcketta, A., Kosmicki, J.A., Ziyatdinov, A., Benner, C., O'Dushlaine, C., Barber, M., Boutkov, B., et al. (2021). Computationally efficient whole-genome regression for quantitative and binary traits. *Nat. Genet.* 53, 1097–1103.
2. Loh, P.-R., Tucker, G., Bulik-Sullivan, B.K., Vilhjálmsson, B.J., Finucane, H.K., Salem, R.M., Chasman, D.I., Ridker, P.M., Neale, B.M., Berger, B., et al. (2015). Efficient bayesian mixed-model analysis increases association power in large cohorts. *Nat. Genet.* 47, 284–290.
3. Zhou, W., Nielsen, J.B., Fritsche, L.G., Dey, R., Gabrielsen, M.E., Wolford, B.N., LeFaive, J., VandeHaar, P., Gagliano, S.A., Gifford, A., et al. (2018). Efficiently controlling for case-control imbalance and sample relatedness in large-scale genetic association studies. *Nat. Genet.* 50, 1335–1341.
4. Jiang, L., Zheng, Z., Fang, H., and Yang, J. (2021). A generalized linear mixed model association tool for biobank-scale data. *Nat. Genet.* 53, 1616–1621.
5. Zou, Y., Carbonetto, P., Wang, G., and Stephens, M. (2022). Fine-mapping from summary data with the “sum of single effects” model. *PLoS Genet.* 18, e1010299.
6. Spence, J.P., Sinnott-Armstrong, N., Assimes, T.L., and Pritchard, J.K. (2022). A flexible modeling and inference framework for estimating variant effect sizes from gwas summary statistics. Preprint at bioRxiv. <https://doi.org/10.1101/2022.04.18.488696>.
7. Orlic, E.J., Banos, D.T., Ojavee, S.E., Läll, K., Mägi, R., Visscher, P.M., and Robinson, M.R. (2022). Improving Gwas Discovery and Genomic Prediction Accuracy in Biobank Data. *Proc. Natl. Acad. Sci. USA* 119, e2121279119.
8. Lawson, D.J., Davies, N.M., Haworth, S., Ashraf, B., Howe, L., Crawford, A., Hemani, G., Davey Smith, G., and Timpson, N.J. (2020). Is population structure in the genetic biobank era irrelevant, a challenge, or an opportunity? *Hum. Genet.* 139, 23–41.
9. Privé, F., Arbel, J., and Vilhjálmsson, B.J. (2020). LDpred2: better, faster, stronger. *Bioinformatics* 36, 5424–5431.
10. Lloyd-Jones, L.R., Zeng, J., Sidorenko, J., Yengo, L., Moser, G., Kemper, K.E., Wang, H., Zheng, Z., Magi, R., Esko, T., et al. (2019). Improved polygenic prediction by bayesian multiple regression on summary statistics. *Nat. Commun.* 10, 5086.
11. Bycroft, C., Freeman, C., Petkova, D., Band, G., Elliott, L.T., Sharp, K., Motyer, A., Vukcevic, D., Delaneau, O., O'Connell, J., et al. (2018). The UK Biobank resource with deep phenotyping and genomic data. *Nature* 562, 203–209.
12. Hof, J.P., and Speed, D. (2025). LDK-KVIK performs fast and powerful mixed-model association analysis of quantitative and binary phenotypes. *Nat. Genet.* 57, 2116–2123.
13. Donoho, D.L., Maleki, A., and Montanari, A. (2009). Message Passing Algorithms for Compressed Sensing. *Proc. Natl. Acad. Sci. USA* 106, 18914–18919.
14. Rangan, S., Schniter, P., and Fletcher, A.K. (2019). Vector approximate message passing. *IEEE Trans. Inf. Theor.* 65, 6664–6684.
15. Feng, O.Y., Venkataramanan, R., Rush, C., and Samworth, R.J. (2022). A unifying tutorial on approximate message passing. *Foundations and Trends® in Machine Learning* 15, 335–536.
16. Bayati, M., and Montanari, A. (2011). The dynamics of message passing on dense graphs, with applications to compressed sensing. *IEEE Trans. Inf. Theor.* 57, 764–785.
17. Montanari, A., and Venkataramanan, R. (2021). Estimation of low-rank matrices via approximate message passing. *Ann. Stat.* 45, 321–345.
18. Barbier, J., Krzakala, F., Macris, N., Miolane, L., and Zdeborová, L. (2019). Optimal errors and phase transitions in high-dimensional generalized linear models. *Proc. Natl. Acad. Sci. USA* 116, 5451–5460.
19. Barbier, J., Camilli, F., Mondelli, M., and Sáenz, M. (2023). Fundamental limits in structured principal component analysis and how to reach them. *Proc. Natl. Acad. Sci. USA* 120, e2302028120.
20. Jeon, C., Ghods, R., Maleki, A., and Studer, C. (2015). Optimality of large MIMO detection via approximate message passing. In *IEEE International Symposium on Information Theory*, pp. 1227–1231.
21. Metzler, C.A., Maleki, A., and Baraniuk, R.G. (2016). From denoising to compressed sensing. *IEEE Trans. Inf. Theor.* 62, 5117–5144.
22. Eksioğlu, E.M., and Tanc, A.K. (2018). Denoising AMP for MRI reconstruction: BM3D-AMP-MRI. *SIAM J. Imag. Sci.* 11, 2090–2109.
23. Zhong, X., Su, C., and Fan, Z. (2022). Empirical bayes pca in high dimensions. *J. Roy. Stat. Soc. B Stat. Methodol.* 84, 853–878.
24. Sarkar, S., Potter, L., and Ahmad, R. (2020). Solving Linear and Bilinear Inverse Problems Using Approximate Message Passing Methods (Ph.D. thesis). AAI28433395.
25. Fletcher, A.K., and Schniter, P. (2017). Learning and free energies for vector approximate message passing. In *2017 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pp. 4247–4251.
26. Fletcher, A.K., Sahraee-Ardakan, M., Rangan, S., and Schniter, P. (2017). Rigorous dynamics and consistent estimation in arbitrarily conditioned linear systems. In *Advances in Neural Information Processing Systems*.
27. Yengo, L., Vedantam, S., Marouli, E., Sidorenko, J., Bartell, E., Sakaue, S., Graff, M., Eliassen, A.U., Jiang, Y., Raghavan, S., et al. (2022). A saturated map of common genetic variants associated with human height. *Nature* 610, 704–712.
28. Benner, C., Spencer, C.C.A., Havulinna, A.S., Salomaa, V., Ripatti, S., and Pirinen, M. (2016). Finemap: efficient variable selection using summary data from genome-wide association studies. *Bioinformatics* 32, 1493–1501.
29. Cui, R., Elzur, R.A., Kanai, M., Ulirsch, J.C., Weissbrod, O., Daly, M.J., Neale, B.M., Fan, Z., and Finucane, H.K. (2024). Improving fine-mapping by modeling infinitesimal effects. *Nat. Genet.* 56, 162–169.
30. Wray, N.R., Purcell, S.M., and Visscher, P.M. (2011). Synthetic associations created by rare variants do not explain most gwas results. *PLoS Biol.* 9, e1000579.
31. Yengo, L., Sidorenko, J., Kemper, K.E., Zheng, Z., Wood, A.R., Weedon, M.N., Frayling, T.M., Hirschhorn, J., Yang, J., and Visscher, P.M.; GIANT Consortium (2018). Meta-analysis of genome-wide association studies for height and body mass index in 700,000 individuals of european ancestry. *Hum. Mol. Genet.* 27, 3641–3649.
32. Wainschtein, P., Jain, D., Zheng, Z., Cupples, L.A., Shadyab, A.H., McKnight, B., Shoemaker, B.M., Mitchell, B.D., et al.; TOPMed Anthropometry Working Group; NHLBI Trans-Omics for Precision Medicine TOPMed Consortium (2022). Assessing the contribution of rare variants to complex trait heritability from whole-genome sequence data. *Nat. Genet.* 54, 263–273.
33. Robinson, M.R., English, G., Moser, G., Lloyd-Jones, L.R., Triplett, M.A., Zhu, Z., Nolte, I.M., van Vliet-Ostapchouk, J.V., Snieder, H., et al.; LifeLines Cohort Study (2017). Genotype-covariate interaction effects and the heritability of adult body mass index. *Nat. Genet.* 49, 1174–1181.

34. Hindorf, L.A., Bonham, V.L., Brody, L.C., Ginoza, M.E.C., Hutter, C.M., Manolio, T.A., and Green, E.D. (2018). Prioritizing diversity in human genomics research. *Nat. Rev. Genet.* **19**, 175–185.
35. Campos, A.I., Namba, S., Lin, S.-C., Nam, K., Sidorenko, J., Wang, H., Kamatani, Y., Biobank Japan Project; Wang, L.H., Lee, S., et al. (2023). Boosting the power of genome-wide association studies within and across ancestries by using polygenic scores. *Nat. Genet.* **55**, 1769–1776.
36. Schniter, P., Rangan, S., and Fletcher, A.K. (2016). Vector approximate message passing for the generalized linear model. In 2016 50th Asilomar conference on signals, systems and computers (IEEE), pp. 1525–1529.
37. Fletcher, A.K., Pandit, P., Rangan, S., Sarkar, S., and Schniter, P. (2019). Plug in estimation in high dimensional linear inverse problems: A rigorous analysis. *J. Stat. Mech.* **2019**, 124021.
38. Benner, C., Havulinna, A.S., Järvelin, M.-R., Salomaa, V., Ripatti, S., and Pirinen, M. (2017). Prospects of fine-mapping trait-associated genomic regions by using summary statistics from genome-wide association studies. *Am. J. Hum. Genet.* **101**, 539–551.
39. Danecek, P., Bonfield, J.K., Liddle, J., Marshall, J., Ohan, V., Pollard, M.O., Whitwham, A., Keane, T., McCarthy, S.A., Davies, R.M., and Li, H. (2021). Twelve years of samtools and bcftools. *GigaScience* **10**, giab008.
40. Band, G., and Marchini, J. (2018). BGEN: a binary file format for imputed genotype and haplotype data. Preprint at bioRxiv 308296. <https://doi.org/10.1101/308296>.
41. Speed, D., Cai, N., UCLEB Consortium; Johnson, M.R., Nejentsev, S., and Balding, D.J. (2017). Reevaluation of snp heritability in complex human traits. *Nat. Genet.* **49**, 986–992.
42. Li, X., Wood, A.R., Yuan, Y., Zhang, M., Huang, Y., Hawkes, G., Beaumont, R.N., Weedon, M.N., Li, W., Li, X., et al. (2025). Streamlining large-scale genomic data management: Insights from the UK Biobank whole-genome sequencing data. *Cell Genomics* **5**, 101009.
43. Skuratovs, N., and Davies, M.E. (2022). Warm-starting in message passing algorithms. In IEEE International Symposium on Information Theory (ISIT), pp. 1187–1192.
44. Hutchinson, M.F. (1990). A stochastic estimator of the trace of the influence matrix for laplacian smoothing splines. *Commun. Stat. Simulat. Comput.* **19**, 433–450.
45. Vila, J., and Schniter, P. (2012). Expectation-maximization gaussian-mixture approximate message passing. *IEEE Trans. Signal Process.* **61**.
46. Vila, J., Schniter, P., Rangan, S., Krzakala, F., and Zdeborová, L. (2015). Adaptive damping and mean removal for the generalized approximate message passing algorithm. In IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), pp. 2021–2025.
47. Wen, X., and Stephens, M. (2010). Using linear predictors to impute allele frequencies from summary or pooled genotype data. *Ann. Appl. Stat.* **4**, 1158–1182.
48. Bick, A.G., Metcalf, G.A., Mayo, K.R., Lichtenstein, L., Rura, S., Carroll, R.J., Musick, A., Linder, J.E., Jordan, I.K., Nagar, S.D., et al. (2024). Genomic data in the all of us research program. *Nature* **627**, 340–346.
49. Chang, C.C., Chow, C.C., Tellier, L.C., Vattikuti, S., Purcell, S.M., and Lee, J.J. (2015). Second-generation PLINK: rising to the challenge of larger and richer datasets. *GigaScience* **4**, 7.

STAR★METHODS

KEY RESOURCES TABLE

REAGENT or RESOURCE	SOURCE	IDENTIFIER
Deposited data		
UK Biobank dataset	UK Biobank	http://www.ukbiobank.ac.uk
All of Us Research Program dataset	National Institutes of Health (NIH)	https://allofus.nih.gov
Software and algorithms		
Python version 3.11	Python Software Foundation	https://www.python.org
R version 4.2.1	The R Foundation for Statistical Computing	https://www.r-project.org
PLINK version 1.9	Chang et al. ³⁸	https://www.cog-genomics.org/plink/1.9/
PLINK version 2.0	Christopher Chang	https://www.cog-genomics.org/plink/2.0/
gVAMP	This paper	https://github.com/medical-genomics-group/gVAMP
REGENIE	Mbatchou et al. ¹	https://github.com/rcggithub/regenie
LDpred2	Privé et al. ⁹	https://privefl.github.io/bigsnpR
SBayesR	Lloyd-Jones et al. ¹⁰	https://gctbhub.cloud.edu.au/software/gctb
LDAC-KVIK	Hof et al. ¹²	https://dougsspeed.com/
GMRM	Orliac et al. ⁷	https://github.com/medical-genomics-group/gmrM
FINEMAP	Benner et al. ²⁸	http://www.christianbenner.com
SuSiE	Zou et al. ⁵	https://stephenslab.github.io/susieR
SuSie-Inf	Cui et al. ²⁹	https://github.com/FinucaneLab/fine-mapping-inf
FINEMAP-Inf	Cui et al. ²⁹	https://github.com/FinucaneLab/fine-mapping-inf
bcftools	Danecek et al. ³⁹	https://samtools.github.io/bcftools/
LDstore	Benner et al. ⁴⁰	http://www.christianbenner.com
bgenix	Band and Marchini ⁴¹	https://enkre.net/cgi-bin/code/bgen
VCF Trimmer	Li et al. ⁴²	https://github.com/drarwood/vcf_trimmer

EXPERIMENTAL MODEL AND STUDY PARTICIPANT DETAILS

In this section, we discuss the data preparation setup. A more detailed description that also motivates our choices further is available in [supplemental methods 1](#).

Participant inclusion in UK Biobank

The study utilized data from the UK Biobank, and all procedures were covered by ethical approval from the North-West Multicenter Research Ethics Committee, and all participants provided written informed consent. The analysis focused on a sample of European-ancestry UK Biobank individuals. Ancestry inference was based on two criteria: self-reported ethnic background (UK Biobank field 21000–0, selecting coding 1) and genetic ethnicity (UK Biobank field 22006–0, selecting coding 1). Participants were projected onto the first two genotypic principal components (PCs) derived from 2,504 individuals in the 1,000 Genomes project. Individuals were retained if their projections were within $PC1 \leq \text{absolute value } 4$ and $PC2 \leq \text{absolute value } 3$. Samples were also excluded based on several UK Biobank quality control procedures.

- Individuals identified as extreme heterozygosity or missing genotype outliers,
- Individuals whose genetically inferred gender did not match their self-reported gender,
- Individuals with putative sex chromosome aneuploidy,
- Individuals excluded from kinship inference,
- Individuals who had withdrawn consent.

Furthermore, 39,592 individuals identified as first-degree relatives were removed from the analyses to avoid potential confounding from shared environmental effects, leaving 419,155 individuals for the whole genome sequence analysis.

Whole-genome sequencing (WGS) data

The study used population-level WGS variants from the UK Biobank DNAnexus platform. Thousands of pVCF files were processed using bcftools³⁹ and VCF Trimmer.⁴² A series of elementary filters were applied to the WGS data.

- Genotype Quality (GQ) ≤ 10 .
- Local allele depth (smpL_sum LAD) < 8 .
- Missing genotype (F_MISSING) > 0.1 .
- Minor Allele Frequency (MAF) < 0.0001 .

Additionally, indels were normalized to the most recent reference, redundant data fields were removed, and any duplicate variants sharing the same base pair position were excluded. Due to RAM limitations of 2TBs on the compute nodes, the variant set was reduced. Variants were ranked by MAF, and a PLINK clumping approach was used to remove variants in high linkage disequilibrium (LD) with the most common variants. This was set with a 1000 kb radius and an R^2 threshold of 0.36. The final WGS dataset consisted of 16,854,878 variants.

Imputed SNP data

We use version 3 of the imputed autosomal genotype data from the UK Biobank. Genotype probabilities were hard-called for variants with an imputation quality score above 0.3, using a hard-call threshold of 0.1 (genotypes with probability ≤ 0.9 were set to missing). Markers were filtered for missingness ($< 5\%$) and Hardy-Weinberg test p -value ($> 10^{-6}$), and only those with an rs identifier were kept, resulting in 23,609,048 SNPs. A primary subset of 8,430,446 autosomal markers was created by selecting markers with MAF ≥ 0.002 . Two smaller subsets were also created from the 8.4M set for comparative analyses.

- 2,174,071 SNPs by taking one marker from variants with LD $R^2 \geq 0.8$ within a 1MB window,
- 882,727 SNPs by further subsetting markers based on $R^2 \geq 0.5$ criteria.

Whole-exome sequencing (WES) data burden scores

The UK Biobank final release of population-level exome variant call files was processed. High-quality, biallelic sites were retained based on: individual and variant missingness $< 10\%$, Hardy-Weinberg Equilibrium p -value $> 10^{-15}$, minimum read coverage depth of 7, and at least one sample per site passing the allele balance threshold > 0.15 . Genomic variants in canonical, protein-coding transcripts were annotated using the Ensembl Variant Effect Predictor (VEP) tool. The LOFTEE plugin was used to identify high-confidence (HC) loss-of-function (LoF) variants. For each of the 17,852 genes, a burden score was calculated.

- Score of 2: Individuals homozygous or multiple heterozygous for LoF variants.
- Score of 1: Individuals with a single heterozygous LoF variant.
- Score of 0: All other individuals.

Phenotypic records

The prepared DNA datasets were linked to phenotypic measurements from the UK Biobank. For the imputed SNP data analyses, 13 quantitative traits (including 7 blood-based biomarkers and 6 other quantitative measures) were selected. These traits were chosen because they showed $\geq 15\%$ SNP heritability and $\geq 5\%$ out-of-sample prediction accuracy in previous work. The sample was divided into a training set and a testing set. The test set consisted of 15,000 individuals who were unrelated (SNP marker relatedness < 0.05) to the individuals in the training set. Prior to analysis, all phenotypic values were adjusted for several covariates: participant's recruitment age, genetic sex, north-south and east-west location coordinates, measurement center, genotype batch, and the leading 20 genetic principal components.

METHOD DETAILS

gVAMP installation

The primary method used for analysis in this paper is gVAMP. Its source code is available at <https://github.com/medical-genomics-group/gVAMP> and can be downloaded using the `git clone` command. Compilation in a High-Performance Computing (HPC) environment requires `gcc`, `openmpi` and `boost` libraries. Upon loading the required modules (e.g., `module load gcc openmpi boost`), the executable (e.g., `main_real.exe`) is built from the C++ source files using an `mpic++` command with relevant flags (such as `-march = native`, `-Ofast`, and `-fopenmp`).

To execute an analysis, gVAMP requires specific inputs passed as command-line arguments, primarily the genotype data in a standard PLINK `.bed` file (specified with `-bed-file`), the phenotype data in a `.phen` file (`-phen-files`), and other model parameters like the number of samples (`-N`) and total markers (`-Mt`). A typical execution command integrating these inputs is shown below.

```
# Example with 30 Processes
time mpirun -np 30 /path/to/main_real.exe \
--bed-file /path/to/genotypes.bed \
--phen-files /path/to/phenotypes.phen \
--N 401452 \
--Mt 2174071 \
... [other arguments]
```

gVAMP algorithm overview

gVAMP extends EM-VAMP, introduced in,^{14,25,26} in which the prior parameters are adaptively learned from the data via EM, and it is an iterative procedure consisting of two steps: (i) denoising, and (ii) linear minimum mean square error estimation (LMMSE). The denoising step accounts for the prior structure given a noisy estimate of the signal β , while the LMMSE step utilizes phenotype values to further refine the estimate by accounting for the LD structure of the data. Further information on historical development and an informal description of AMP-based approaches for Bayesian regressions are presented in [supplemental methods 2](#).

A key feature of the algorithm is the so called *Onsager correction*: this is added to ensure the asymptotic normality of the noise corrupting the estimates of β at every iteration. Here, in contrast to MCMC or other iterative approaches, the normality is guaranteed under mild assumptions on the normalized genotype matrix. This property allows a precise performance analysis via state evolution and, consequently, the optimization of the method.

In particular, the quantity $\gamma_{1,t}$ in line 7 of [Algorithm 1](#) is the state evolution parameter tracking the error incurred by $r_{1,t}$ in estimating β at iteration t . The state evolution result gives that $r_{1,t}$ is asymptotically Gaussian, i.e., for sufficiently large N and P , $r_{1,t}$ is approximately distributed as $\mathcal{N}(\beta, \gamma_{1,t}^{-1} \mathbf{I})$. Here, β represents the signal to be estimated, with the prior learned via EM steps at iteration t :

$$\beta_i \sim (1 - \lambda_t) \cdot \delta_0(\cdot) + \lambda_t \cdot \sum_{l=1}^L \pi_{t,l} \cdot \mathcal{N}(\cdot, 0, \sigma_{t,l}^2), \forall i = 1, \dots, P.$$

Compared to [Equation 2](#), the subscript t in $\lambda_t, \pi_{t,l}, \sigma_{t,l}$ indicates that these parameters change through iterations, as they are adaptively learned by the algorithm. Similarly, $r_{2,t}$ is approximately distributed as $\mathcal{N}(\beta, \gamma_{2,t}^{-1} \mathbf{I})$. The Gaussianity of $r_{1,t}, r_{2,t}$ is enforced by the presence of the Onsager coefficients $\alpha_{1,t}$ and $\alpha_{2,t}$, see lines 17 and 22 of [Algorithm 1](#), respectively. We also note that $\alpha_{1,t}$ (resp. $\alpha_{2,t}$) is the state evolution parameter linked to the error incurred by $\hat{\beta}_{1,t}$ (resp. $\hat{\beta}_{2,t}$).

The vectors $r_{1,t}, r_{2,t}$ are obtained after the LMMSE step, and they are further improved via the denoising step, which respectively gives $\hat{\beta}_{1,t}, \hat{\beta}_{2,t}$. In the denoising step, we exploit our estimate of the approximated posterior by computing the conditional expectation of β with respect to $r_{1,t}, r_{2,t}$ in order to minimize the mean square error of the estimated effects. For example, let us focus on the pair $(r_{1,t}, \hat{\beta}_{1,t})$ (analogous considerations hold for $(r_{2,t}, \hat{\beta}_{2,t})$). Then, we have that

$$\hat{\beta}_{1,t} = f_t(r_{1,t}) = \mathbb{E}[\beta | r_{1,t} = \beta + \mathcal{N}(0, \gamma_{1,t}^{-1} \mathbf{I}), \lambda_t, \{\pi_{t,l}\}_{l=1}^L, \{\sigma_{t,l}^2\}_{l=1}^L]. \quad (\text{Equation 3})$$

Here, $f_t: \mathbb{R} \rightarrow \mathbb{R}$ denotes the denoiser at iteration t and the notation $f_t(r_{1,t})$ assumes that the denoiser f_t is applied component-wise to elements of $r_{1,t}$. We give explicit formulas for denoisers f_t as well as the Onsager corrections in the [supplemental methods 3](#). Moreover, note that, in line 15 of [Algorithm 1](#), we take this approach one step further by performing an additional step of damping, see “gVAMP algorithm stability” below.

If one has access to the singular value decomposition (SVD) of the data matrix $\mathbf{X}$, the per-iteration complexity is of order $\mathcal{O}(NP)$. However, at biobank scales, performing the SVD is computationally infeasible. Thus, the linear system $(\gamma_{\epsilon,t} \mathbf{X}^T \mathbf{X} + \gamma_{2,t} \mathbf{I})^{-1} (\gamma_{\epsilon,t} \mathbf{X}^T \mathbf{y} + \gamma_{2,t} \mathbf{r}_{2,t})$ (see line 21 of [Algorithm 1](#)) needs to be solved using an iterative method, in contrast to having an analytic solution in terms of the elements of the singular value decomposition of $\mathbf{X}$. We approximate the solution of the linear system $(\gamma_{\epsilon,t} \mathbf{X}^T \mathbf{X} + \gamma_{2,t} \mathbf{I})^{-1} (\gamma_{\epsilon,t} \mathbf{X}^T \mathbf{y} + \gamma_{2,t} \mathbf{r}_{2,t})$ with a symmetric and positive-definite matrix via the *conjugate gradient method* (CG), see [Algorithm 2](#) in [supplemental methods 4](#). If κ is the condition number of $\gamma_{\epsilon,t} \mathbf{X}^T \mathbf{X} + \gamma_{2,t} \mathbf{I}$, the method requires $\mathcal{O}(\sqrt{\kappa})$ iterations to return a reliable approximation.

Additionally, inspired by,⁴³ we initialize the CG iteration with an estimate of the signal from the previous iteration of gVAMP. This warm-starting technique leads to a reduced number of CG steps that need to be performed and, therefore, to a computational speed-up. However, this comes at the expense of potentially introducing spurious correlations between the signal estimate and the Gaussian error from the state evolution. Such spurious correlations may lead to algorithm instability when run for a large number of iterations (also extensively discussed below). This effect is prevented by simply stopping the algorithm as soon as the R^2 measure on the training data starts decreasing.

In order to calculate the Onsager correction in the LMMSE step of gVAMP (see line 22 of [Algorithm 1](#)), we use the Hutchinson estimator⁴⁴ to estimate the quantity $\text{Tr}[(\gamma_{\epsilon,t} \mathbf{X}^T \mathbf{X} + \gamma_{2,t} \mathbf{I})^{-1}] / P$. We recall that this estimator is unbiased, in the sense that, if $\mathbf{u}$ has i.i.d. entries equal to -1 and $+1$ with the same probability, then

$$\mathbb{E}[\mathbf{u}^T (\gamma_{\epsilon,t} \mathbf{X}^T \mathbf{X} + \gamma_{2,t} \mathbf{I})^{-1} \mathbf{u} / P] = \text{Tr}[(\gamma_{\epsilon,t} \mathbf{X}^T \mathbf{X} + \gamma_{2,t} \mathbf{I})^{-1}] / P.$$

Furthermore, in order to perform an EM update for the noise precision γ_{ϵ} , one has to calculate the trace of a matrix which is closely connected to the one we have seen in the previous paragraph. In order to do so efficiently, i.e., to avoid solving another large-dimensional linear system, we store the inverted vector $(\gamma_{\epsilon,t} \mathbf{X}^T \mathbf{X} + \gamma_{2,t} \mathbf{I})^{-1} \mathbf{u}$ and reuse it again in the EM update step.

The VAMP approach in¹⁴ assumes exact knowledge of the prior on the signal β , which deviates from the setting in which genome-wide association studies are performed. Hence, we adaptively learn the signal prior from the data using expectation maximization (EM) steps, see lines 8 and 28 of [Algorithm 1](#). This leverages the variational characterization of EM-VAMP,²⁵ and its rigorous theoretical analysis presented in.²⁶ In the [supplemental methods 5](#), we summarize the hyperparameter estimation results derived based upon⁴⁵ in the context of our model.

gVAMP implementation

Our open-source gVAMP software (<https://github.com/medical-genomics-group/gVAMP>) is implemented in C++, and it incorporates parallelization using the OpenMP and MPI libraries. MPI parallelization is implemented in a way that the columns of the normalized genotype matrix are approximately equally split between the workers. OpenMP parallelization is done on top of that and used to further boost performance within each worker by simultaneously performing operations such as summations within matrix vector product calculations. Moreover, data streaming is employed using a lookup table, enabling byte-by-byte processing of the genotype matrix stored in PLINK format with entries encoded to a set $\{0, 1, 2\}$.

$$(\begin{matrix} 0 & 1 & 0 & 0 & 1 & 1 & 1 & 0 \end{matrix}) \mapsto (\begin{matrix} \text{NaN} & 2 & 0 & 1 \end{matrix})$$

The lookup table enables streaming in the data in bytes, where every byte (8 bits) encodes the information of 4 individuals. This reduces the amount of memory needed to load the genotype matrix. In addition, given a suitable computer architecture, our implementation supports SIMD instructions which allow handling four consecutive entries of the genotype matrix simultaneously. To make the comparisons between different methods fair, the results presented in the paper do not assume usage of SIMD instructions. Additionally, we emphasize that all calculations take un-standardized values of the genotype matrix in the form of standard PLINK binary files, but are conducted in a manner that yields the parameter estimates one would obtain if each column of the genotype matrix was standardized.

gVAMP algorithm stability

We find that the application of existing EM-VAMP algorithms to the UK Biobank dataset leads to diverging estimates of the signal. This is due to the fact that the data matrix (the SNP data) might not conform to the properties required in,¹⁴ especially that of right-rotational invariance. Furthermore, incorrect estimation of the noise precision in line 28 of [Algorithm 1](#) may also lead to instability of the algorithm, as previous applications of EM-VAMP generally do not leave many hyperparameters to estimate.

To mitigate these issues, different approaches have been proposed including row or/and column normalization, damping (i.e., doing convex combinations of new and previous estimates),⁴⁶ and variance auto-tuning.²⁶ In particular, to prevent EM-VAMP from diverging and ensure it follows its state evolution, we empirically observe that the combination of the following techniques is crucial.

1. We perform *damping* in the space of denoised signals. Thus, line 15 of [Algorithm 1](#) reads as

$$\hat{\beta}_{1,t} = \rho \cdot \mathbb{E}[\beta | \mathbf{r}_{1,t}, \Theta_t] + (1 - \rho) \cdot \hat{\beta}_{1,t-1},$$

in place of $\hat{\beta}_{1,t} = \mathbb{E}[\beta | \mathbf{r}_{1,t}, \Theta_t]$. Here, $\rho \in (0, 1)$ denotes the damping factor. This ensures that the algorithm is making smaller steps when updating a signal estimate.

2. We perform *auto-tuning* of $\gamma_{1,t}$ via the approach from.²⁶ Namely, in the auto-tuning step, one refines the estimate of $\gamma_{1,t}$ and the prior distribution of the effect size vector β by jointly re-estimating them. If we denote the previous estimates of $\gamma_{1,t}$ and Θ with $\gamma_{1,t}^{(\text{previous})}$ and $\Theta^{(\text{previous})}$, then this is achieved by setting up an expectation maximization procedure whose aim is to maximize

$$\mathbb{E}[\log p(\beta, \mathbf{r}_{1,t} | \gamma_{1,t}, \Theta) | \mathbf{r}_{1,t}, \gamma_{1,t}^{(\text{previous})}, \Theta^{(\text{previous})}]$$

with respect to $\gamma_{1,t}$ and Θ .

3. We *filter* the design matrix for first-degree relatives to reduce the correlation between rows, which has the additional advantage of avoiding potential confounding of shared-environmental effects among relatives.

gVAMP model parameters for the analysis on imputed data

For the analyses of the 13 UK Biobank phenotypes using the 887,060 and 2,174,071 imputed SNP sets, a damping factor ρ of 0.1 was used. The prior was initialized with a spike at zero and 22 slab mixtures. The variances of these mixtures followed a geometric progression in the interval $[10^{-6}, 10]$, and their probabilities followed a geometric progression with a factor of 1/2. The initial probability of being assigned to the zero mixture was 97%.

This configuration works well for all phenotypes. We also note that our inference of the number of mixtures, their probabilities, their variances and the SNP marker effects is not dependent upon specific starting parameters for the analyses of the 2,174,071 and 882,727 SNP datasets, and the algorithm is rather stable for a range of initialization choices. Similarly, the algorithm is stable for different choices of the damping ρ , as long as said value is not too large.

Generally, appropriate starting parameters are not known in advance and this is why we learn them from the data within the EM steps of our algorithm. However, it is known that EM can be sensitive to the starting values given and, thus, we recommend initialising a series of models at different values to check that this is not the case (similar to starting multiple Monte Carlo Markov chains in standard Bayesian methods). The feasibility of this recommendation is guaranteed by the significant speed-up of our algorithm compared to existing approaches (see [Note S1](#), [Figure S21](#)).

Given the same data, gVAMP yields estimates that more closely match GMRM when convergence in the training R^2 , SNP heritability, residual variance, and out-of-sample test R^2 are smoothly monotonic within around 10–40 iterations. Following this, training R^2 , SNP heritability, residual variance, and out-of-sample test R^2 may then begin to slightly decrease as the number of iterations becomes large. Thus, as a stopping criterion for the 2,174,071 and 882,727 SNP datasets, we choose the iteration that maximizes the training R^2 .

We highlight the iterative nature of our method. Thus, improved computational speed and more rapid convergence is achieved by providing better starting values for the SNP marker effects. Specifically, when moving from 2,174,071 to 8,430,446 SNPs, only columns with correlation $R^2 \geq 0.8$ are being added back into the data. Thus, for the 8,430,446 SNP set, we initialise the model with the converged SNP marker and prior estimates obtained from the 2,174,071 SNP runs, setting to 0 the missing markers. Furthermore, we lower the value of the damping factor ρ , with typical values being 0.05 and 0.01. We experiment both with using the noise precision from the initial 2,174,071 SNP runs and with setting it to 2. We then choose the model that leads to a smoothly monotonic curve in the training R^2 . We observe that SNP heritability, residual variance, and out-of-sample test R^2 are also smoothly monotonic within 25 iterations. Thus, as a stopping criterion for the 8,430,446 SNP dataset, we choose the estimates obtained after 25 iterations for all the 13 traits. We follow the same process when extending the analyses to include the WES rare burden gene scores.

gVAMP model parameters for the analysis of human height in WGS data

We apply gVAMP to the WGS data to analyze human height using the largest computational instance currently available on the DNAnexus platform, employing 128 cores and 1921.4 GB total memory. Efficient C++ gVAMP implementation allows for parallel computing, utilizing OpenMP and MPI libraries. Here, we split the memory requirements and computational workload between 2 OpenMP threads and 64 MPI workers. The prior initialization of gVAMP is similar to that used in the simulation study, except that we opt for a sparser prior with probability 99.5% of variants being assigned to the 0 mixture. As before, the SNP marker effect sizes are initialised with 0. Based on the experiments in the UK Biobank imputed dataset, in WGS we set the initial damping factor ρ to 0.1, and adjust it to 0.05 from iteration 4 onward, which stabilizes the algorithm and leads us to stop it at iteration 18. We also note that, after the first few iterations, the marker inclusion probabilities are stable.

In the simulation study on WGS data from chromosomes 11–15, gVAMP was run with a damping factor ρ of 0.05. The prior was initialized with a spike at zero and 22 slab mixtures. The mixture variances followed a geometric progression in the interval $[10^{-7}, 1]$, and the mixture probabilities followed a geometric progression with a factor of 1/2. The initial probability of a variant being assigned to the zero mixture was 99%. All variant effect sizes were initialized to 0, and the method was run until iteration 30.

Polygenic risk scores and SNP heritability

gVAMP produces SNP effect estimates that can be directly used to create polygenic risk scores. The estimated effect sizes are on the scale of normalised SNP values, i.e., $(\mathbf{X}_j - \mu_{X_j})/SD(\mathbf{X}_j)$, with μ_{X_j} the column mean and $SD(\mathbf{X}_j)$ the standard deviation, and thus SNPs in the out-of-sample prediction data must also be normalized. We provide an option within the gVAMP software to do phenotypic prediction, returning the adjusted prediction R^2 value when given input data of a PLINK file and a corresponding file of phenotypic values. gVAMP estimates the SNP heritability as the phenotypic variance (equal to 1 due to normalization) minus 1 divided by the estimate of the noise precision, i.e., $h_{SNP}^2 = 1 - 1/\gamma_e$.

Comparison to polygenic risk scores obtained by other methods

We compare gVAMP to an MCMC sampler approach (GMRM) with a similar prior (the same number of starting mixtures) as presented in.⁷ We select this comparison as the MCMC sampler was demonstrated to exhibit the highest genomic prediction accuracy up to

date.⁷ We run GMRM for 2000 iterations, taking the last 1800 iterations as the posterior. We calculate the posterior means for the SNP effects and the posterior inclusion probabilities of the SNPs belonging to the non-zero mixture group. GMRM estimates the SNP heritability in each iteration by sampling from an inverse χ^2 distribution using the sum of the squared regression coefficient estimates. We also compare to the LDAK-KVIK elastic net software,¹² run using their suggested parameters and default options.

Finally, we also compare gVAMP to the summary statistics prediction methods LDpred2⁹ and SBayesR¹⁰ using summary statistic data of 8,430,446 SNPs obtained from the REGENIE software. For SBayesR, following the recommendation on the software webpage (<https://cnsngenomics.com/software/gctb/#SummaryBayesianAlphabet>), after splitting the genomic data per chromosomes, we calculate the so-called *shrunk* LD matrix, which use the method proposed by⁴⁷ to shrink the off-diagonal entries of the sample LD matrix toward zero based on a provided genetic map. We make use of all the default values: `--genmap-n 183`, `--ne 11400` and `--shrunk-cutoff 10-5`. Following that, we run the SBayesR software using summary statistics generated via the REGENIE software (see “Mixed linear association testing” below) by grouping several chromosomes in one run. Namely, we run the inference jointly on the following groups of chromosomes: {1},{2},{3},{4},{5.6},{7.8},{9,10,11},{12,13,14} and {15:22}. This allows to have locally joint inference, while keeping the memory requirements reasonable. All the traits except for Blood cholesterol (CHOL) and Heel bone mineral density T-score (BMD) give non-negative R^2 ; CHOL and BMD are then re-run using the option to remove SNPs based on their GWAS p -values (threshold set to 0.4) and the option to filter SNPs based on LD R-Squared (threshold set to 0.64). For more details on why one would take such an approach, one can check <https://cnsngenomics.com/software/gctb/#FAQ>. As the obtained test R^2 values are still similar, as a final remedy, we run standard linear regression over the per-group predictors obtained from SBayesR on the training dataset. Following that, using the learned parameters, we make a linear combination of the per-group predictors in the test dataset to obtain the prediction accuracy given in the table.

Replication of results in All of Us

The analysis used whole genome sequence (WGS) data from the All of Us cohort.⁴⁸ Informed consent for all participants is conducted in person or through an eConsent platform that includes primary consent, HIPAA Authorization for Research use of electronic health records and other external health data, and Consent for Return of Genomic Results. The protocol was reviewed by the Institutional Review Board (IRB) of the All of Us Research Program. The All of Us IRB follows the regulations and guidance of the NIH Office for Human Research Protections for all studies, ensuring that the rights and welfare of research participants are overseen and protected uniformly.

A cohort of 226,147 individuals of European genetic ancestry with height data was selected for the replication analysis. The phenotypic values for height were adjusted for sex and 16 available principal components (PCs) within the All of Us dataset. A marginal regression was performed on the standardized WGS dosage data. This analysis was conducted using the Hail Python library. The results of this marginal regression (estimating one variant at a time) from All of Us were then compared to the joint-effect estimates generated by gVAMP from the WGS runs in the UK Biobank data for standing height as a trait.

Simulation study for WGS data

To demonstrate the usefulness of gVAMP as a signal localization tool, we investigate a series of 12 simulation scenarios in which the signal is planted on top of the observed whole genome sequence data of chromosomes 11–15, which consists of 3,285,117 variants and 411,536 individuals. At a scale of over 3 million WGS variants and over 400,000 individuals, cloud computing costs naturally restrict the number of scenarios we can explore, especially as we have to benchmark against a range of other fine-mapping approaches. The simulation details for each of the 12 scenarios we considered are reported in Table S1 and described below.

The signal itself is simulated from a Bernoulli-Gaussian prior in which the variance of the slab part is determined as the ratio of the phenotypic variance explained by the causal variants and the number of causal variants. We fix the phenotypic variance explained by the causal variants to be 7.25% and vary the number of causal variants as either 500 or 1000. Given that the length of the DNA of chromosomes 11–15 is ~18.8% of the total autosomal DNA length, the settings we use are equivalent to a trait with a proportion of variance explained of ~39% spread over ~2700 to ~5300 causal variants genome-wide. Our objective in the simulation study is to access the performance of our model at scale in settings similar to our WGS analysis of human height, but where power is sufficient to accurately compare variable selection methods. $39\%/5300 = 0.0074\%$ is a value similar to that expected for human height $70\%/10000 = 0.007\%$ suggested by recent results.²⁷ In the settings we explore, we vary γ , the percentage of non-zero effects attributed to common variants (those with minor allele frequency ≥ 0.05). In simulation scenarios with $\gamma = 0\%$, we spread the signals uniformly across variants maintaining our minimal distance. $\gamma = 50\%$ allows common variants in the interval [0.05,0.5] to be causal with the probability 50%. When selecting causal variants within these two settings, we do so at random, but to spread variants across the DNA we enforce a minimal distance between consecutive causal variants of size 100,000 base pairs. This minimal distance facilitates a simple direct comparison of variable selection methods to localize a single signal within a given genomic region, but restricts us to a smaller absolute number of causal variants than may be expected for some traits genome-wide. However, as noted above, the average expected contribution to variance is in-line with that expected for human height and, thus, our simulation represents a realistic potential architecture.

We vary α , the power relationship between effect size and minor allele frequency. For each SNP position j , upon randomly choosing a mixture group l it belongs to, according to the true effect probability distribution, we draw its effect size from

$\mathcal{N}(0, \sigma_j^2 \cdot (2 \cdot \text{MAF}_j \cdot [1 - \text{MAF}_j])^{1+\alpha})$. Given a desired heritability h^2 of the simulated trait, we normalize the estimates so that the sum of squares of all simulated effects is equal to h^2 . This model conforms exactly to that described in⁴¹ and we set values of -1 , -0.5 , and -0.25 to cover those reported for complex traits.⁴¹ As described in,⁴¹ when $\alpha = 0$ the unscaled effect-size variance is invariant to MAF, indicative of selective neutrality. The uniform heritability model ($\alpha = -1$) implies that the unscaled effect-size variance increases as MAF decreases, which is interpreted as a signal of negative (or purifying) selection.

All simulated traits are adjusted by the first 20 principal components following standard practice. For running SuSiE and SuSiE-inf (v1.4), we use the standard fine-mapping with infinitesimal effects v1.2 Python framework provided by the authors.²⁹ The framework allows to initialize FINEMAP-inf using tau-squared and sigma-squared parameters generated by SuSiE-inf to improve accuracy. When running FINEMAP, we use a shotgun stochastic search subprogram. We perform genome-wide marginal testing and define fine-mapping blocks of size 3Mbp around a marginal discovery, passing the Bonferroni-adjusted p -value threshold of 0.05. When the 3Mbp contained or neighbored multiple causal variants, we algorithmically allow for merging the overlapping regions, allowing FINEMAP and SuSiE to run on region sizes of 6Mbp. In fact, we found that this gave favorable performance to minimize the probability of counting an association as false, when it is simply correlated to a neighboring true causal variant. For each region, we create an index of variants using bgenix,⁴⁰ compute LD matrices using LDstore v2.0,³⁸ and marginal estimates using PLINK 1.9.⁴⁹

We run the gVAMP method using as prior initialization a spike at zero and 22 slab mixtures. We let the variance of those mixtures follow a geometric progression in the interval $[10^{-7}, 1]$, and we let the probabilities follow a geometric progression with factor 1/2. The prior probability $1-\lambda$ of variants being assigned to the 0 mixture is initialized to 99%. Variant effect sizes are initialised with 0. Based on the experiments in the UK Biobank imputed dataset, in WGS we use damping factor $\rho = 0.05$ and run the method until iteration 30 when the hyper-parameters stabilize.

We then conduct an additional simulation study focusing on the inclusion of WES burden scores alongside WGS variants. For the settings of $(\alpha, \gamma) \in \{(-1.0), (-1.50), (-0.5.0), (-0.5.50)\}$, we used the same whole genome sequence data of chromosomes 11–15 described above. We simulated 500 underlying causal variants with a total heritability of $h^2 = 7.25\%$, allocating 75% of the causal markers (CM) to randomly selected WGS variants and 25% to randomly selected WES burden scores from the same chromosomes. The WES burden scores exhibit effect sizes drawn from the normal distribution with variance 3-fold larger than those of individual WGS variants. This gives the expectation that the variance attributable to WGS variants should equal that of the WES burden scores. We then ran gVAMP using the same initialization described above.

Finally, we conduct a further simulation study where we vary effects sizes across genomic annotations. Specifically, we simulated 500 causal variants with a total heritability of $h^2 = 7.25\%$, stratifying them into coding and non-coding regions of the genome. We vary the parameter $\gamma_{\text{annot}} \in \{0, 50, 100\}$ which specifies the percentage of causal variants assigned to coding (exonic) regions. This simulation was performed only for settings a , b , e , and f , described in Table S1. We then ran gVAMP using the same initialization described above.

QUANTIFICATION AND STATISTICAL ANALYSIS

gVAMP association testing

In line with the standard outputs of MCMC methods, the AMP framework offers the possibility of calculating posterior inclusion probabilities (PIPs) jointly for each of the SNPs. Once the prior information in a particular iteration is updated, we use Equation 4 below to approximate the posterior inclusion probability of each SNP. To conduct variable selection, we apply a threshold of $\text{PIP} \geq 0.95$.

The gVAMP framework provides posterior inclusion probability testing for detecting associations. The calculation uses the output of the state evolution, specifically the noisy estimate of the effects $\mathbf{r}_{1,t}$ and its associated precision $\gamma_{1,t}$ at a given iteration t . The posterior inclusion probability for a single SNP j , denoted as ζ_j , is calculated using the following equation:

$$\zeta_j = \frac{\lambda_t \cdot \sum_{i=1}^L \pi_{i,t} \cdot \mathcal{N}(\mathbf{r}_{1,t,j}; 0, \sigma_{i,t}^2 + \gamma_{1,t}^{-1})}{\lambda_t \cdot \sum_{i=1}^L \pi_{i,t} \cdot \mathcal{N}(\mathbf{r}_{1,t,j}; 0, \sigma_{i,t}^2 + \gamma_{1,t}^{-1}) + (1 - \lambda_t) \cdot \mathcal{N}(\mathbf{r}_{1,t,j}; 0, \gamma_{1,t}^{-1})}. \quad (\text{Equation 4})$$

In this formula, the terms represent the following theoretical quantities:

- ζ_j : The posterior probability that SNP j has a non-zero effect,
- $(\mathbf{r}_{1,t})_j$: The noisy estimate of the effect for SNP j at iteration t ,
- $\gamma_{1,t}$: The precision of the estimate $(\mathbf{r}_{1,t})_j$, tracked by state evolution,
- λ_t : The adaptively learned sparsity rate, which is the prior probability of a non-zero effect,
- L : The number of Gaussian components in the “slab” part of the prior.
- $\pi_{i,t}$: The prior probability of being in the i -th Gaussian mixture, given that the effect is non-zero,
- $\sigma_{i,t}^2$: The variance of the i -th Gaussian mixture component.
- $\mathcal{N}(\cdot; \mu, \nu)$: The probability density function of a Gaussian distribution with mean μ and variance ν .

Mixed linear model association testing

We conduct mixed linear model association testing using a leave-one-chromosome-out (LOCO) estimation approach on the 8,430,446 and 2,174,071 imputed SNP markers. LOCO association testing approaches have become the field standard and they are two-stage: a subset of markers is selected for the first stage to create genetic predictors; then, statistical testing is conducted in the second stage for all markers one-at-a-time. We consider REGENIE,¹ as it is a recent commonly applied approach and LDAK-KVIK.¹² We also compare to GMRM,⁷ a Bayesian linear mixture of regressions model that has been shown to outperform REGENIE for LOCO testing. For the first stage of LOCO, REGENIE is given 887,060 markers to create the LOCO genetic predictors, even if it is recommended to use 0.5 million genetic markers. We compare the number of significant loci obtained from REGENIE to those obtained if one were to replace the LOCO predictors with: (i) those obtained from GMRM using the LD pruned sets of 2,174,071 and 887,060 markers; and (ii) those obtained from gVAMP at all 8,430,446 markers and the LD pruned sets of 2,174,071 and 887,060 markers. We note that obtaining predictors from GMRM at all 8,430,446 markers is computationally infeasible, as using the LD pruned set of 2,174,071 markers already takes GMRM several days. In contrast, gVAMP is able to use all 8,430,446 markers and still be faster than GMRM with the LD pruned set of 2,174,071 markers. LDAK-KVIK is given 2,174,071 markers to facilitate a direct comparison with gVAMP and we compare (i) placing the predictors of gVAMP directly within the LDAK-KVIK step 2; (ii) using gVAMP step 2 procedure which is identical to that of REGENIE.

LOCO testing does not control for linkage disequilibrium within a chromosome. Thus, to facilitate a simple, fair comparison across methods, we clump the LOCO results obtained with the following PLINK commands: `--clump-kb 1000 --clump-r2 0.01 --clump-p1 0.00000005`. Therefore, within 1Mb windows of the DNA, we calculate the number of independent associations (squared correlation ≤ 0.01) identified by each approach that pass the genome-wide significance testing threshold of $5 \cdot 10^{-8}$. As LOCO can only detect regions of the DNA associated with the phenotype and not specific SNPs, given that it does not control for the surrounding linkage disequilibrium, a comparison of the number of uncorrelated genome-wide significance findings is conservative. A full benchmark of approaches in a simulation study is also provided in [Note S1](#), where we consider a wider range of metrics, scenarios and comparisons.

ADDITIONAL RESOURCES

Code url: <https://github.com/medical-genomics-group/gVAMP>.