

Genetics and population analysis

BTS: a scalable Bayesian Tissue Score for prioritizing GWAS variants and their functional contexts across >1000s of omics datasets

Pavel P. Kuksa^{1,†}, Matei Ionita^{1,†}, Luke Carter¹, Jeffrey Cifello¹, Prabhakaran Gangadharan¹, Kaylyn Clark¹, Otto Valladares¹, Yuk Yee Leung^{1,*}, Li-San Wang^{1,*}

¹Department of Pathology and Laboratory Medicine, Penn Neurodegeneration Genomics Center, University of Pennsylvania, Philadelphia, PA 19104, United States

*Corresponding authors. Yuk Yee Leung, 3700 Hamilton Walk, D101 Richards Medical Research Building, University of Pennsylvania, Philadelphia, PA 19104, United States. E-mail: yyee@penmedicine.upenn.edu; Li-San Wang, 3700 Hamilton Walk, D102 Richards Medical Research Building, University of Pennsylvania, Philadelphia, PA 19104, United States. E-mail: lswang@penmedicine.upenn.edu.

[†]The authors wish it to be known that, in their opinions, the first two authors should be regarded as Joint First Authors.

Associate Editor: Russell Schwartz

Abstract

Motivation: Summary statistics from genome-wide association studies (GWAS) are widely used in fine-mapping and colocalization analyses to identify causal variants and their enrichment in functional contexts, such as affected cell types and genomic features. With the expansion of functional genomic (FG) datasets, which now include hundreds of thousands of tracks across various cell and tissue types, it is critical to establish scalable algorithms integrating thousands of diverse FG annotations with GWAS results.

Results: We propose BTS (Bayesian Tissue Score), a novel, highly efficient algorithm uniquely designed for (i) identifying affected cell types and functional elements (context-mapping) and (ii) fine-mapping potentially causal variants in a context-specific manner using large collections of cell type-specific FG annotation tracks. BTS leverages GWAS summary statistics and annotation-specific Bayesian models to analyze genome-wide annotation tracks, including enhancers, open chromatin, and histone marks. We evaluated BTS on GWAS summary statistics for immune and cardiovascular traits, such as Inflammatory Bowel Disease (IBD), Rheumatoid Arthritis (RA), Systemic Lupus Erythematosus (SLE), and Coronary Artery Disease (CAD). Our results demonstrate that BTS is over 100x more efficient in estimating functional annotation effects and context-specific variant fine-mapping compared to existing methods. Importantly, this large-scale Bayesian approach prioritizes both known and novel annotations, cell types, genomic regions, and variants and provides valuable biological insights into the functional contexts of these diseases.

Availability and implementation: Docker image is available at <https://hub.docker.com/r/wanglab/bts> with preinstalled BTS R package (<https://bitbucket.org/wanglab-upenn/BTS-R>) and BTS GWAS summary statistics analysis pipeline (<https://bitbucket.org/wanglab-upenn/bts-pipeline>).

1 Introduction

Typical single-marker-based genome-wide association studies (GWASs) do not consider correlation between variants (linkage disequilibrium, LD) and the potential cellular or epigenetic context of the genetic variants. Interpretation and prioritization of GWAS results and variants thus require downstream analyses using various types of functional genomic (FG) data, e.g. enhancer, open chromatin, histone modification, and other epigenetic markers.

Many existing methods can perform fine-mapping of the GWAS signals within a region, based on the summary statistics alone (approximate Bayes factor), or summary statistics with LD information. Examples of these methods include conditional analysis (Galarneau *et al.* 2010, Knight *et al.* 2012, Yang *et al.* 2012, Uffelmann *et al.* 2021), CAVIAR/CAVIARBF (Hormozdiari *et al.* 2014, Chen *et al.* 2015, 2016), DAP-G (Wen *et al.* 2016), FINEMAP (Benner *et al.* 2016), and SuSiE (Wang *et al.* 2020, Zou *et al.* 2022). Other methods prioritize variants based on their functional annotations, such as coding, promoter, or enhancer regions, e.

g. RegulomeDB (Boyle *et al.* 2012, Dong *et al.* 2023). Colocalization methods such as coloc (Giambartolomei *et al.* 2014, Wallace 2021), HyPrColoc (Foley *et al.* 2021), eCaviar (Hormozdiari *et al.* 2016), ENLOC (Wen *et al.* 2017), and their variants jointly analyze summary statistics from GWAS and molecular traits such as expression quantitative trait loci (eQTL). There are also tools for integrative analysis [e.g. INFERNO (Amlie-Wolf *et al.* 2018), SparkINFERNO (Kuksa *et al.* 2020), FUMA (Watanabe *et al.* 2017)] which combine data from multiple sources: for example, prioritizing variants based on both colocalization posterior probabilities and overlaps with functional annotations. Additionally, some methods [e.g. PAINTOR (Kichaev *et al.* 2014), fGWAS (Pickrell 2014), BFGWAS_QUANT (Chen *et al.* 2022), PolyFun (Weissbrod *et al.* 2020), CARMA (Yang *et al.* 2023)] further formalize overlap with functional annotations and incorporate functional annotations as priors for inferring causal variant status.

The recent increase in availability of large functional annotation databases such as ENCODE (Encode Project Consortium 2012, Encode Project Consortium *et al.* 2020),

Received: 6 March 2025; Revised: 25 July 2025; Accepted: 30 August 2025

© The Author(s) 2025. Published by Oxford University Press.

This is an Open Access article distributed under the terms of the Creative Commons Attribution License (<https://creativecommons.org/licenses/by/4.0/>), which permits unrestricted reuse, distribution, and reproduction in any medium, provided the original work is properly cited.

ROADMAP (Roadmap Epigenomics Consortium et al. 2015, Encode Project Consortium et al. 2020), and FILER (Kuksa et al. 2022) provides an opportunity to carry out unbiased functional analyses of GWAS results across thousands of cell types and genome-wide annotations. However, such large-scale, systematic analyses can come at a significant computational cost, as they need to build and evaluate different statistical models for each annotation or combination of annotations and cell types of interest. In addition, each model evaluation often requires running time exponential in the number of potentially causal variants (Kichaev et al. 2014, Asimit et al. 2019).

Moreover, most methods that model LD are susceptible to errors due to a mismatch between GWAS summary statistics and LD. Such mismatches can occur when the GWAS and the reference genotype panel used to compute LD differ demographically. Even more importantly, mismatch is all but guaranteed in the case of GWAS meta-analyses, as reviewed recently in the SLALOM publication (Kanai et al. 2022). Briefly, if two variants belong to the same haplotype, most models expect them to have similar summary statistics and behave erratically if this expectation is violated. But different studies in a meta-analysis can incorporate one variant and not the other, which can cause large differences in the amount of evidence supporting each variant. To accommodate this common situation, fine-mapping methods must be robust to LD mismatch (Chen et al. 2021, Kanai et al. 2022, Yang et al. 2023).

To address these large-scale analysis issues, we present BTS, a Bayesian Tissue Score model and an efficient implementation of this model. BTS performs joint fine-mapping of variants and their context-mapping based on FG annotations and provides easy-to-interpret summaries of the results. Table 1, available as supplementary data at *Bioinformatics* online summarizes the main features of the BTS pipeline

compared to existing methods and tools. The main features of the proposed BTS framework are:

- End-to-end GWAS summary statistics analysis pipeline (Fig. 1; see Section 2.1; [Supplementary Methods, available as supplementary data at Bioinformatics online](#)): users can provide their own functional annotations for running with BTS, or use the FILER FG database (Kuksa et al. 2022) to obtain cell type-specific annotations from data sources such as ENCODE (Encode Project Consortium 2012, Encode Project Consortium et al. 2020), EPIMAP (Boix et al. 2021), and GTEx (GTEx Consortium 2020). By default, the user only needs to provide GWAS summary statistics as input.
- Robustness to mismatch between GWAS summary statistics and LD estimates (Fig. 2). The BTS model introduces a single model parameter (the prior on variance of the true effect sizes) to control most of the sensitivity to LD mismatch. We also provide guidelines on how to choose this parameter (see Section 2.4).
- BTS can perform joint context-mapping (inference of cell types, genomic features) and context-specific fine-mapping of variants (Figs 1, 3, 4, and 5; see Section 2).
- Scalability: BTS can conduct a systematic and exhaustive search through thousands of genome-wide annotation tracks and has running times two orders-of-magnitude faster per genome-wide track (see Section 3.4; Fig. 6) using a novel, more efficient factored Bayesian model (see Section 2; Section 2.2) for FG annotations, LD, and GWAS summary statistics.

We applied BTS to GWAS datasets from four different diseases: Coronary Artery Disease (CAD) (van der Harst and

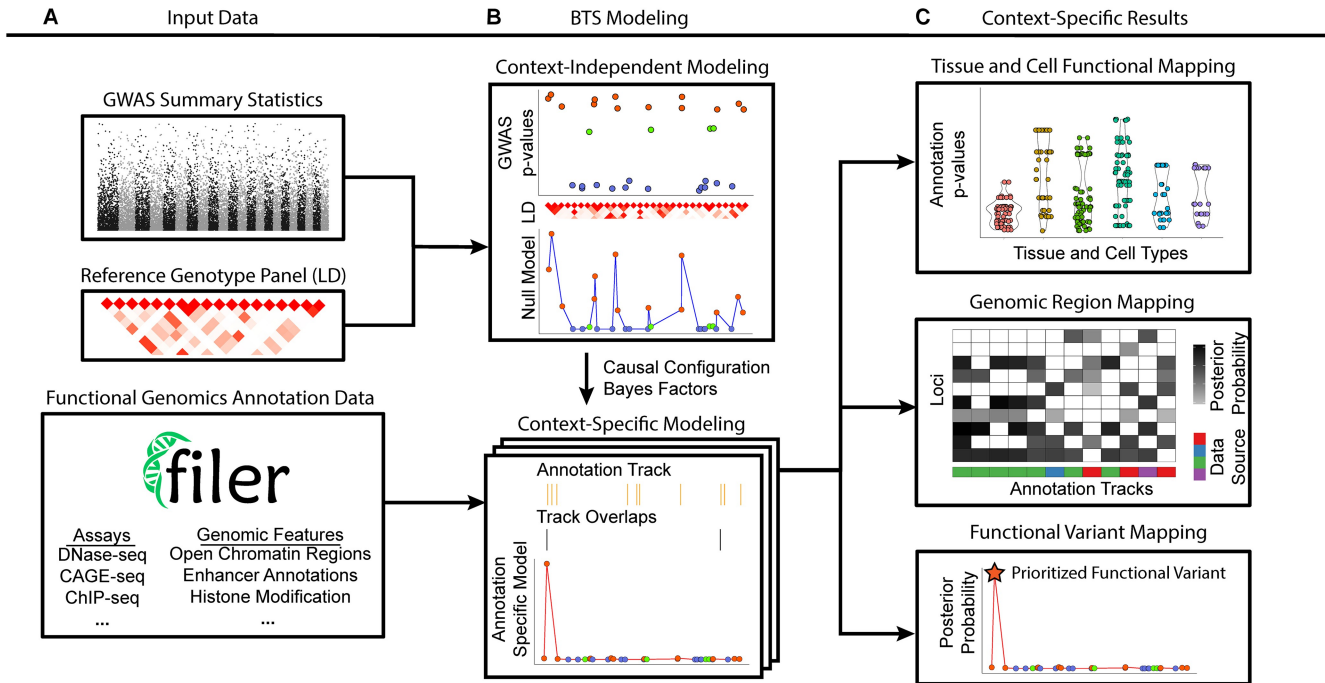


Figure 1. Overview of the BTS framework. (A) Input GWAS summary statistics and LD data (top subpanel) are analyzed across cell type and tissue-level FG annotations (bottom subpanel) to jointly perform context-mapping and identify context-specific sets of genomic regions and potentially causal variants for each context and region. (B) The two-factor structure of BTS model (see Section 2.2) allows it to be evaluated efficiently across thousands of genome-wide annotations (Fig. 6; see Section 3.4). BTS estimates context-independent GWAS+LD null model (top subpanel) and many annotation-specific models (bottom subpanel). (C) BTS outputs annotation relevance (enrichment P -values) (top subpanel), functional genomic regions within each of the prioritized contexts (middle subpanel), and annotation-specific causal variant posterior probabilities (bottom subpanel) (Fig. 3; see Sections 3.2 and 3.3). Figure 1, available as supplementary data at *Bioinformatics* online and Section 2.1 detail BTS analysis workflow and its main steps.

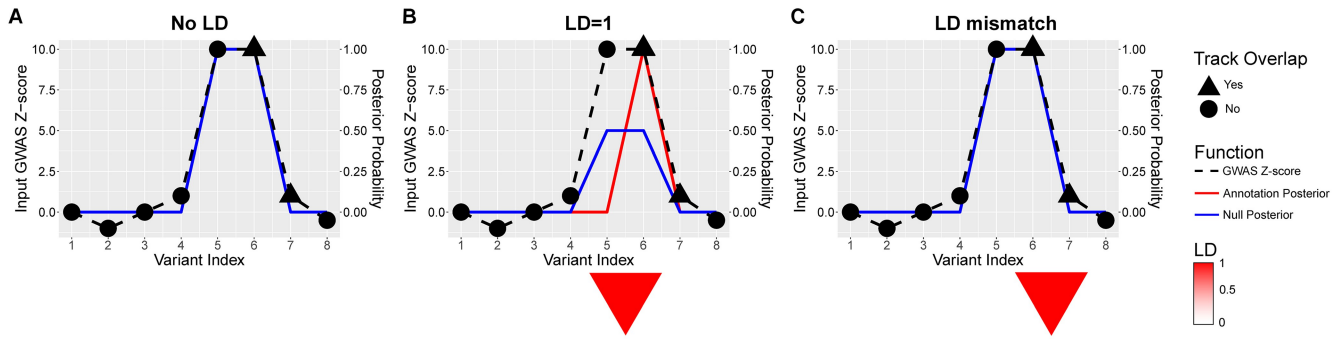


Figure 2. Variant prioritization with BTS using summary statistics, LD, and annotation information. (A) In the case of two independent signals (numbered 5 and 6 on the x-axis), both variants are prioritized by BTS (posterior probability for both variants is 1). (B) If the two variants (5 and 6) are in LD, variant 6 is prioritized (Annotation Posterior) based on its overlap with a relevant annotation track (overlap signified by triangle). Note that the null model (Null Posterior) will not be able to prioritize between these variants (posterior = 0.5) since both variants have similar Z-scores and are in LD. (C) In case of LD and GWAS mismatch, two variants (6 and 7) are in LD but have very different GWAS summary statistics (Z-scores). As shown here, BTS is robust with respect to the LD mismatch and prioritizes variant 6 with a more significant association statistic (Z-score) and avoids prioritizing false-positive variant 7 by default model. Table 6, available as supplementary data at *Bioinformatics* online illustrates the BTS model behavior in the LD mismatch scenario using matching and mismatched genotype panels from 1000 Genomes (Genomes Project Consortium *et al.* 2015).

Verweij 2018), Inflammatory Bowel Disease (IBD) (Liu *et al.* 2015), Rheumatoid Arthritis (RA) (Stahl *et al.* 2010), and Systemic Lupus Erythematosus (SLE) (Bentham *et al.* 2015). In each case, BTS took under one hour (Fig. 6) to prioritize cell types, tissues, genomic regions, and variants relevant to the disease across a variety of FG annotation tracks (>900; see Table 2, available as supplementary data at *Bioinformatics* online for details on annotation tracks used). This shows BTS can serve as a tool for systematic and unbiased functional context mapping and context-specific variant mapping.

The BTS GWAS summary statistics analysis pipeline, including scripts for obtaining relevant LD and functional annotations starting from the GWAS summary statistics, is freely available at <https://bitbucket.org/wanglab-upenn/BTS-pipeline>. BTS model estimation is implemented as an R package and is freely available at <https://bitbucket.org/wanglab-upenn/BTS-R>. BTS Docker image including preinstalled GWAS summary statistics pipeline is also available at <https://hub.docker.com/r/wanglab/bts>.

2 Methods

2.1 BTS GWAS summary statistics analysis workflow

We provide an end-to-end functional variant fine-mapping and context-mapping pipeline for analysis of GWAS summary statistics based on user-supplied full GWAS summary statistics as input (Fig. 1; Fig. 1 and Supplementary Methods, available as supplementary data at *Bioinformatics* online). The BTS GWAS summary statistics pipeline aims to automate genome-wide post-GWAS functional analysis and provide a systematic and more complete report of all potentially causal variants, genomic loci, and likely functional contexts including cell type and tissue-specific regulatory mechanisms underlying observed GWAS signals.

To do this, input GWAS summary statistics are first pre-processed to resolve reference and nonreference (alternative) alleles and normalize effect sizes to consistently reflect effects of the alternative alleles. Using the normalized GWAS summary statistics, a set of genomic regions for further downstream analyses and all potential candidate variants are then identified. This is achieved through identification of pairwise-independent GWAS signals (tag variants) by performing linkage disequilibrium-based pruning ($LD\ r^2 > 0.7$) of all genome-wide significant ($P < 5e-8$) GWAS variants. Based on the identified tag variants, the

candidate set of potentially causal variants is then formed as a set of variants in LD with the tagging variants ($LD\ r^2 > 0.7$), including the tag variants themselves. LD-based genomic regions for analysis are constructed by defining the LD region for each of the tag variants as a genomic region with the left and right boundaries corresponding to the leftmost and rightmost variants linked with the tag variant. The leftmost and rightmost variants are restricted to be within 1 Mbp from the tag variant and have no more than 1000 variants between them and the tag variant. The final set of nonoverlapping genomic regions for downstream analyses is obtained by merging any overlapping LD-based regions into larger regions.

For each such identified genomic region, the pipeline will then generate all the information required to fit the BTS model (see Section 2.2) including the pairwise LD-based variant correlation matrix L ($n \times n$, where n is the number of variants in the region), functional annotation matrix A ($n \times NA$), and a vector of GWAS summary statistics Z (Z-scores) ($n \times 1$). Pairwise LD calculation for all variants located in the locus is conducted based on the reference genotype panel (Genomes Project Consortium *et al.* 2015, Byrka-Bishop *et al.* 2022).

Functional annotation matrix A is obtained by querying FILER FG database (Kuksa *et al.* 2022) for each of the NA genomic annotations and FG data tracks of interest and noting annotation overlaps for each of the variants ($A_{i,j}$ will be set to 1 if variant i overlaps annotation j). The summary of included annotations and detailed list of annotation tracks used are provided in Tables 2 and 4, available as supplementary data at *Bioinformatics* online. Summary statistics (Z-scores) are extracted from the input GWAS summary statistics after reference and alternative allele resolution and effect normalization to consistently reflect the effect of the nonreference (alternative) allele.

BTS algorithm (Fig. 1; see Section 2; Section 2.3) will then be applied to fit the model and estimate variant and locus posteriors and functional annotation enrichment for each of the target annotations (Fig. 3). To find potentially causal variants within each locus, BTS uses variant LD matrix and Z-scores from GWAS summary statistics to precompute and store Bayes factors for each possible causal variant configuration in every locus. BTS then uses an EM-based algorithm to iteratively estimate annotation enrichment coefficients and compute annotation-specific causal priors for each of the analyzed variants. These functional annotation-specific priors

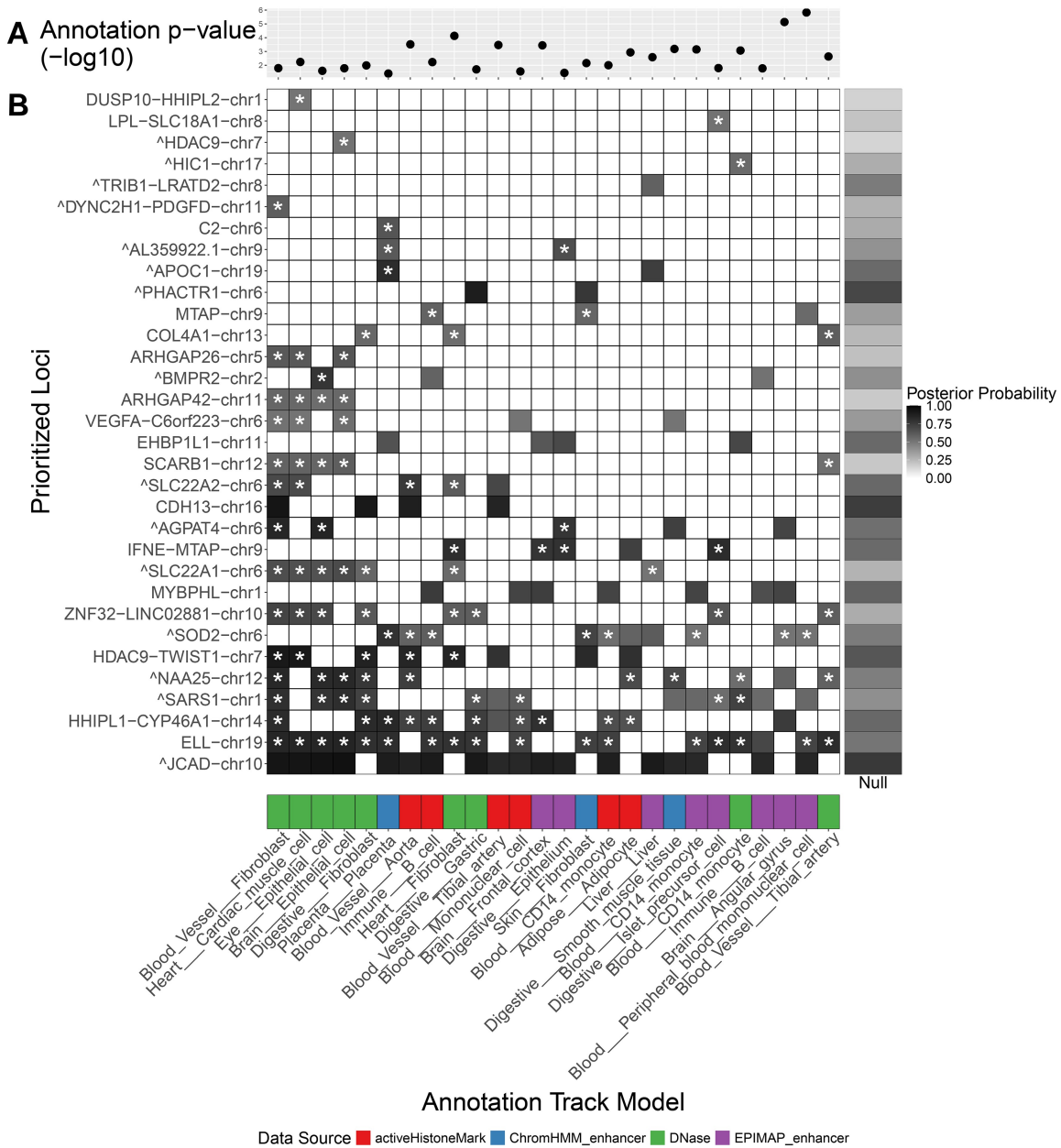


Figure 3. BTS prioritization for functional contexts (X-axis) and genomic regions (Y-axis) in CAD GWAS. (A) Genome-wide enrichment of the functional annotations across all identified regions of interest and using all variants located within these regions. Shown are the annotation significance (P -values) as given by the likelihood ratio test of the BTS model with and without annotation (see Section 2.3; [Supplementary Methods, available as supplementary data at Bioinformatics online](#), for details on estimating annotation significance). Out of 103 significant annotations ($P < .01$), shown are 25 annotations overlapping with at least five prioritized genomic regions (variant posterior > 0.5). In general, annotations with greater enrichment in top GWAS variants have lower, more significant P -values. (B) BTS-prioritized regions of interest (Y-axis) and their potential functional contexts (X-axis). Shown are regions with annotation overlaps and top variant posterior > 0.5 . Asterisks (*) mark regions with increased causal variant posterior compared to the null model (GWAS+LD only, without annotation) in one or more functional contexts. Darker colors correspond to regions and annotations with greater top variant posteriors. Merged regions (i.e. regions containing more than one overlapping LD blocks) are labeled with a caret (^) before the region name. For CAD, BTS prioritizes blood vessel, monocyte, cardiac muscle, and immune cell types across active histone marks, open chromatin, and enhancer genomic features.

are then combined with precomputed configuration Bayes factors to obtain context-specific causal variant posteriors.

Evaluation results reported in Sections 3.2 and 3.3 were generated by applying this pipeline to CAD ([van der Harst and Verweij 2018](#)), IBD ([Liu *et al.* 2015](#)), RA ([Bentham *et al.* 2015](#)), and SLE ([Stahl *et al.* 2010](#)) GWAS summary statistics.

2.2 BTS statistical model

To accommodate for correlation between variants (linkage disequilibrium, LD) as well as the tissue and cell-type-specific

functions of variants and genomic regions, we adopt the Bayesian probabilistic framework first proposed in [Kichaev *et al.* \(2014\)](#). We consider the following information for all n variants in a genomic region of interest: (i) a vector $\mathbf{Z}$ of GWAS Z-scores (standardized regression coefficients; observed), (ii) a vector $\mathbf{A}$ of variant tissue and cell type-specific annotations (observed), (iii) LD (linkage disequilibrium) correlation matrix $\mathbf{\Sigma}$, and a vector $\mathbf{C}$ of variant causal status (unobserved; latent binary variable: set to 1 for causal variants, and 0 otherwise), a vector $\mathbf{\Lambda}$ of unknown true effect

sizes for each variant. To model each genomic region, we use a Bayesian model in which the likelihood of observing $\mathbf{Z}$ is a multivariate normal parametrized by causal configuration $\mathbf{C}$, true effect sizes $\mathbf{\Lambda}$, and LD correlation matrix $\mathbf{\Sigma}$:

$$P(\mathbf{Z}|\mathbf{C}, \mathbf{\Lambda}, \mathbf{\Sigma}) = N(\mathbf{Z}; \mathbf{\Sigma}(\mathbf{\Lambda}^\circ \mathbf{C}), \mathbf{\Sigma}) \quad (1)$$

where $\mathbf{\Lambda}^\circ \mathbf{C}$ is an element-wise vector product. The true effect size $\mathbf{\Lambda}$ is also modeled as normal, with mean 0 and diagonal variance, using the scalar W as a model parameter:

$$P(\mathbf{\Lambda}|\mathbf{C}) = N(\mathbf{\Lambda}; \mathbf{0}, W\mathbf{I}_C) \quad (2)$$

where $W\mathbf{I}_C$ is a scaled diagonal matrix, with diagonal elements of $\mathbf{I}_C$ set to 0 and 1 according to the causal configuration $\mathbf{C}$. Integrating out $\mathbf{\Lambda}$ gives the following formula for the full likelihood as proved in [Kichaev et al. \(2014\)](#):

$$P(\mathbf{Z}|\mathbf{C}, \mathbf{\Sigma}) = N(\mathbf{Z}; \mathbf{0}, \mathbf{\Sigma} + W\mathbf{\Sigma}\mathbf{I}_C\mathbf{\Sigma}) \quad (3)$$

Bayes factor (BF) for a causal variant configuration $\mathbf{C}$ in any particular genomic region:

$$\begin{aligned} BF_C &= \frac{P(\mathbf{Z}|\mathbf{C})}{P(\mathbf{Z}|\mathbf{C}=\mathbf{0})} \\ &= \frac{N(\mathbf{Z}; \mathbf{0}, \mathbf{\Sigma} + W\mathbf{\Sigma}\mathbf{I}_C\mathbf{\Sigma})}{N(\mathbf{Z}; \mathbf{0}, \mathbf{\Sigma})} \\ [Lemma1] &= \frac{N(\mathbf{Z}_1; \mathbf{0}, \mathbf{\Sigma}_{11} + W\mathbf{\Sigma}_{11}^2)}{N(\mathbf{Z}_1; \mathbf{0}, \mathbf{\Sigma}_{11})} \end{aligned} \quad (4)$$

$$[Lemma2] = \det(\mathbf{I} + W\mathbf{\Sigma}_{11})^{-1/2} \exp\left(\frac{W}{2} \mathbf{Z}_1^T (\mathbf{I} + W\mathbf{\Sigma}_{11})^{-1} \mathbf{Z}_1\right)$$

where the first simplification (Lemma 1) reduces BF computation from full vectors and matrices to the much smaller $\mathbf{Z}_1 = \mathbf{Z}_{i:C_i=1}$, and $\mathbf{\Sigma}_{11} = \mathbf{\Sigma}_{i:C_i=1}$ corresponding to the Z -scores and correlation between the causal variants ($C_i = 1$), and the second simplification (Lemma 2) further reduces BF computation to a single matrix inversion of a positive semidefinite matrix (see [Supplementary Methods, available as supplementary data at Bioinformatics online](#)).

The prior probability of causality for each variant i is modeled as a logistic function of its annotation A :

$$\text{Prior}(C_i) = P(C_i = 1|A, E) = 1/(1 + \exp(-EA)) \quad (5)$$

where the annotation effect size coefficient E for annotation A is estimated genome-wide and is shared by all variants in all regions. Note that for variants overlapping annotation A with positive effect E , their prior probability of causality will be greater than prior probabilities for variants located outside of annotation A . Prior for causal configuration $\mathbf{C}_j$ for region j is then a product of all n individual variant priors $P(C_{ij})$

$$\begin{aligned} \text{Prior}(\mathbf{C}_j) &= \prod_i P(C_{ij}|A, E) = \\ &= \prod_i P(C_{ij} = 1|A, E)^{C_{ij}} P(C_{ij} = 0|A, E)^{1-C_{ij}} \end{aligned} \quad (6)$$

The full data likelihood across all genomic regions j is a product of individual region data likelihoods:

$$L(\mathbf{Z}; E, A) = \prod_j \sum_{\mathbf{C}_j} P(\mathbf{Z}_j|\mathbf{C}_j) P(\mathbf{C}_j|E, A) \quad (7)$$

The computational complexity for each region j is then consists of prior computation and data likelihood computation [Equation (3)] for every possible causal variant configuration $\mathbf{C}$, $O(|\mathbf{C}| \times (L + d^3))$.

To improve computational efficiency, we first note that $P(C_i = 0|A, E) = 1 - P(C_i = 1|A, E) = (1 + \exp(EA))^{-1}$ and a full variant configuration probability can be computed in $O(d)$ time (where d is the maximum number of independent causal variants) as an update to the null configuration probability:

$$P(\mathbf{C}|\mathbf{E}, A) = P(\mathbf{C}_0|\mathbf{E}) \prod_{i:C_i=1} P(C_i = 1|\mathbf{E}, A) / P(C_i = 0|\mathbf{E}) \quad (8)$$

where $\mathbf{C}_0$ is a null configuration (all variants are noncausal) and the $P(\mathbf{C}_0)$ term is only computed once per locus.

Marginalizing over $\mathbf{C}$, and restricting to configurations that have at most d causal variants, we obtain the posterior probability that a variant i is causal in a particular genomic region:

$$\begin{aligned} P(C_i = 1|\mathbf{Z}, \mathbf{\Sigma}, A, E) &= \sum_{\mathbf{C}: C_i=1} P(\mathbf{C}|\mathbf{Z}, \mathbf{\Sigma}, A, E) \\ &= \frac{\sum_{\mathbf{C}: C_i=1} P(\mathbf{C}|\mathbf{Z}, \mathbf{\Sigma}) P(\mathbf{C}|\mathbf{E}, A)}{\sum_{\text{all } \mathbf{C}} P(\mathbf{C}|\mathbf{Z}, \mathbf{\Sigma}) P(\mathbf{C}|\mathbf{E}, A)} \\ &= \frac{\sum_{\mathbf{C}: C_i=1} BF_C P(\mathbf{C}|\mathbf{E}, A)}{\sum_{\text{all } \mathbf{C}} BF_C P(\mathbf{C}|\mathbf{E}, A)} \end{aligned} \quad (9)$$

expressed in terms of the Bayes factors $BF_C = P(\mathbf{Z}|\mathbf{C})/P(\mathbf{Z}|\mathbf{C}_0)$ [Equation (4)] and configuration priors $P(\mathbf{C})$ [Equation (6)].

More conceptually, the variant posterior in Equation (9) is a dot product between a vector of annotation-independent Bayes factors BF_C and a vector of annotation-dependent variant priors $P(\mathbf{C})$, where each vector is indexed by causal variant configurations $\mathbf{C}$ with up to d causal variants (i.e. each vector is of $|\mathbf{C}| = \sum_{l=0}^d \binom{n}{l}$ dimensionality, where n is the number of variants in the locus). These vectors can be computed independently from each other, as Bayes factors (BF) only depend on GWAS summary statistics and LD matrix [Equation (4)], while the variant priors only depend on annotations and their enrichment coefficients [Equations (5) and (6)].

Given access to the precomputed Bayes factors for each possible configuration $\mathbf{C}$, the overall complexity of computing variant posteriors [using Equations (8) and (9)] is then linear $O(|\mathbf{C}| \times d)$ for any given genome-wide annotation A , which is a $O(L + d^2)$ improvement compared to nonfactorized model [Equation (7)] $O(|\mathbf{C}| \times (L + d^3))$ with the on-the-fly prior and Bayes factor computation.

The model outputs the variant causal posterior probabilities [Equation (9)] for each of analyzed variants and the estimated annotation effect size coefficients E_A for each tested annotation A .

2.3 BTS algorithm

We use an expectation-maximization (EM) algorithm ([Kichaev et al. 2014](#)) to fit the statistical model in Equation

(9). Intuitively, this is an iterative algorithm which optimizes overall likelihood and takes turns updating the posterior probabilities and enrichment coefficients, until a convergence criterion is reached.

Given multiple annotations, with a possibly complex correlation structure, it is standard practice (Kichaev et al. 2014, Pickrell 2014) to perform feature selection, by first fitting a separate model for each annotation, and then selecting a few high-ranking annotations for a final model. Therefore, our systematic approach to annotations requires that thousands of models be fitted for all annotation, tissue, and cell types.

Our key observation is that, in Equation (9), the posterior probability decomposes into a factor which only involves GWAS data, and one which only involves annotations. Furthermore, the factor involving GWAS data is the same for all iterations of the EM algorithm. Because of this, BTS computes it only once, and distributes it to all models and EM iterations as necessary. Since likelihood computation is the most time-intensive part of fitting the model, the choice to only perform it once is responsible for the bulk of our computational improvements.

This algorithm design choice is a trade-off between speed and memory: to compute the likelihood only once, BTS must store it until all models have been processed. The amount of time and memory spent on likelihood computations is proportional to the number of allowed causal configurations:

- If there are N variants, and any subset of them could be the causal set, then there are 2^N causal sets to be enumerated. Due to the nature of exponential growth, this is impractical even for moderately large regions ($N > 30$) and impossible for large ones ($N > 80$).
- BTS implements a common solution to this problem, which is to only consider configurations of size smaller or equal to d (Kichaev et al. 2014, Asimit et al. 2019), a user-provided parameter, with default value 2. Then storing the likelihood of such configurations requires $O(N^d)$ memory for a region with N variants.
- If $d = 2$, the necessary memory is equal to that of storing the LD matrix, so BTS gains two orders of magnitude in speed, at the cost of only doubling its memory use.
- For $d > 2$, we mitigate the memory use by only storing those likelihoods which are at most t orders of magnitude smaller than the largest, where t is a user-provided parameter, with default value 12. In our experiments with $d = 3, 4, 5$, this optimization reduces memory use by 100-fold and does not change the final results within the first 5 significant digits. Our experiments suggest that BTS can accommodate values of d up to 5 without significant memory issues, while for $d > 5$ runtime increases severely.

We obtain further computational improvements by using the matrix inversion lemma (Lemma 2) to compute Bayes factors (Supplementary Methods, available as supplementary data at *Bioinformatics* online) and more efficiently computing variant configuration probabilities [Equation (8)]. When computing likelihood ratios, BTS needs to evaluate the ratio of multivariate normal densities with different variance matrices. Naively, this involves inverting each variance matrix. The matrix inversion lemma provides an equivalent expression in which a single matrix needs to be inverted and has the following benefits:

- Decreased computation time, since matrix inversion is the most time-consuming part of likelihood computation.
- The covariance matrices are singular in the case of variants in perfect LD. In our formulation, the matrix to be inverted is strictly positive definite, which improves numerical stability and removes the need for regularization.

Figure 1 summarizes the BTS algorithm:

- A module for computing Bayesian factors and likelihoods.
- A loop which distributes annotations and likelihoods to each model to be fit.
- Aggregation of results, and prioritization of annotations, loci, and variants.

BTS GWAS summary statistics pipeline using core BTS algorithm is outlined in Fig. 1, available as supplementary data at *Bioinformatics* online.

2.4 Mitigation of LD mismatch

We investigated the effects of LD mismatch, which can occur whenever in-sample LD for the GWAS cohort is unavailable, and a reference genotype panel is used instead. There are existing methods to flag regions with high suspicion of LD mismatch [DENTIST (Chen et al. 2021), SLALOM (Kanai et al. 2022)], but they do not address the problem after identifying it. Moreover, SLALOM only flags a region if the mismatch involves the variant with highest association, which is not completely general.

Our approach is to quantify the spurious effect of LD mismatch on likelihood computations and provide guidance in choosing algorithm parameters so that this effect is minimized. In the supplementary material (see Supplementary Methods, available as supplementary data at *Bioinformatics* online), we show that, for two variants in perfect LD, with Z-scores a, b , the likelihood of the configuration where both are causal is:

$$(1 + 2W)^{-1/2} \exp \left(\frac{W}{2(1 + 2W)} [(a^2 + b^2) + W(a - b)^2] \right) \quad (10)$$

where W is the prior variance from Equation (2). Since the variants are perfectly correlated, it should be the case that $a = b$ in the absence of LD mismatch. In practice, we often observe $LD = 1$ but $a \neq b$; one Z-score could be large while the other is close to zero. In this case, the term $W(a - b)^2$ in the exponent is the spurious effect which should be minimized. If W is much larger than 1, then the spurious second term can end up dominating the first one. However, if W is much smaller than 1, then the null configuration can end up dominating all others, leading to posterior probabilities close to 0. To balance these requirements, BTS uses $W = 1$.

3 Results

3.1 BTS overview

To jointly estimate variant posterior probabilities and enrichment of annotations in causal variants, we build on a well-known Bayesian model [PAINTOR (Kichaev et al. 2014), CAVIARBF (Chen et al. 2015), see also Section 2; Section 2.2; Supplementary Methods, available as supplementary

data at *Bioinformatics* online]. In this framework, for each possible configuration of causal variants within a locus, its Bayes factor only depends on the GWAS summary statistics and the correlation between variants (linkage disequilibrium, LD). On the other hand, causal variant configuration priors depend only on the functional annotations of the variants, as well as the annotation importance (e.g. enrichments of the annotations in causal variants). We use an Expectation-Maximization (EM) algorithm to maximize overall data likelihood by iteratively updating the annotation enrichments and evaluating configuration posterior probabilities, until a convergence criterion is reached (Fig. 1; see Section 2; [Supplementary Methods, available as supplementary data at *Bioinformatics* online](#)).

A key insight underlying the BTS method is that the posterior probability of each variant configuration can be decomposed into two components: one that depends only on the GWAS data and LD, and the other that depends solely on functional annotations [see Section 2.2; Equation (9)]. Crucially, the GWAS-based factor is the same across all iterations of the EM algorithm. This allows BTS to compute Bayes factors just once and reuse them across annotation evaluations, significantly improving computational efficiency (Figs 1 and 6).

We also implemented two key computational improvements to further improve performance:

- (1) Efficient Bayes factor computation using the matrix inversion lemma [see Section 2.2; [Supplementary Methods, available as supplementary data at *Bioinformatics* online; Lemma 2, Equation \(4\)](#)];
- (2) More efficient evaluation of variant configuration probabilities [see Section 2.2; Equation (8)].

These observations, alongside computational improvements detailed in Section 2 and [Supplementary Methods, available as supplementary data at *Bioinformatics* online](#), are the main reason for the improved runtime of BTS (Fig. 6).

To run the BTS algorithm, there are two options. (A) The user supplies three types of information for each genomic region to be analyzed: per-locus GWAS summary statistics, pre-computed LD matrix, and user-prepared functional annotations. This option skips the data preparation steps 1 and 2 referenced in the list below, which identify regions of interest and generate their LD and functional annotations. (B) The user can provide full GWAS summary statistics and utilize the BTS end-to-end analysis pipeline (Fig. 1) which automates the rest of the workflow, including the generation of regions of interest and their functional annotations. This end-to-end pipeline:

- (1) Systematically identifies genomic regions of interest by performing LD pruning to select pairwise-independent (tag) genome-wide significant variants, followed by LD expansion to define regions consisting of one or more overlapping LD blocks (see Section 2; Section 2.1).
- (2) Computes variant-level LD and functional annotation matrices for each region using the 1000 Genomes genotype panel and the FILER FG annotation data, respectively (see Section 2; Section 2.1).

Performs BTS model estimation using the assembled input data (Steps 3, 4 in Fig. 1, [available as supplementary data at *Bioinformatics* online](#)).

Aggregates across estimated models to prioritize functional annotations and provide context-specific prioritization for the genomic regions and variants.

This GWAS summary statistics analysis pipeline leverages the FILER FG data repository ([Kuksa et al. 2022](#)) to obtain tissue- and cell type-specific functional annotations, and the 1000 Genomes project genotype reference panel ([Genomes Project Consortium et al. 2015, Byrsk-Bishop et al. 2022](#)) to estimate LD between variants. The pipeline outputs: (i) cell type and functional context-specific posterior probabilities for individual variants within each of the identified genomic regions of interest, (ii) annotation and functional context importance (enrichment scores, causal variant prior odds), as well as (iii) cell type, region, and variant mapping summary plots (e.g. Figs 3 and 4).

As illustrated in Fig. 2, given the same summary statistics as input, BTS discovers two independent signals if the variants are uncorrelated (Fig. 2, first column, LD = 0), and one independent signal shared between two variants, if these are correlated (Fig. 2, second column, LD = 1). Crucially, in the baseline model with no functional annotations, the posterior probability mass is shared equally between the two variants, but in the model that uses an annotation the probability is assigned to the variant which overlaps trait-relevant annotation.

Finally, BTS is robust to mismatch between LD and summary statistics (Fig. 2, third column, LD mismatch; [Table 6, available as supplementary data at *Bioinformatics* online](#)): when two variants with LD close to 1 have wildly different summary statistics, the one with low phenotype association is not prioritized. We emphasize that this common-sense result was difficult to achieve: the same statistical model with different parameter settings has a propensity for prioritizing variants with low phenotype association in such cases of LD mismatch. One of our contributions in this article is to identify parameter values that make the model robust to LD mismatch ([Supplementary Methods, available as supplementary data at *Bioinformatics* online; Section 2 and Section 2.4](#)).

3.2 Prioritizing regions, variants, and their contexts

To illustrate the performance of BTS, we started with summary statistics from a Coronary Artery Disease (CAD) GWAS ([van der Harst and Verweij 2018](#)). Our end-to-end pipeline first performed LD pruning of the 4298 genome-wide significant GWAS variants ($P\text{-value} < 5 \times 10^{-8}$) to identify 389 tag variants (pairwise-independent, $r^2 < 0.7$). We further filtered out 2 loci belonging to the HLA region on chromosome 6, where complicated linkage patterns make fine-mapping extremely difficult. LD expansion of the remaining tag variants yielded 6513 candidate variants. Using the LD-expanded candidate set, we then defined the LD block for each tag variant using the left-most and right-most variants in LD with the tag variant (i) within a 1 megabase window and (ii) with no more than 1000 variants separating the boundary variants and the tag variant. Any of the overlapping LD blocks were then merged to obtain 167 nonoverlapping regions of interest. By extracting variants within these regions from the initial GWAS summary statistics, we obtained a total of 24 727 variants for further analysis. We note that each of the regions will not only contain variants that are highly associated with the CAD phenotype

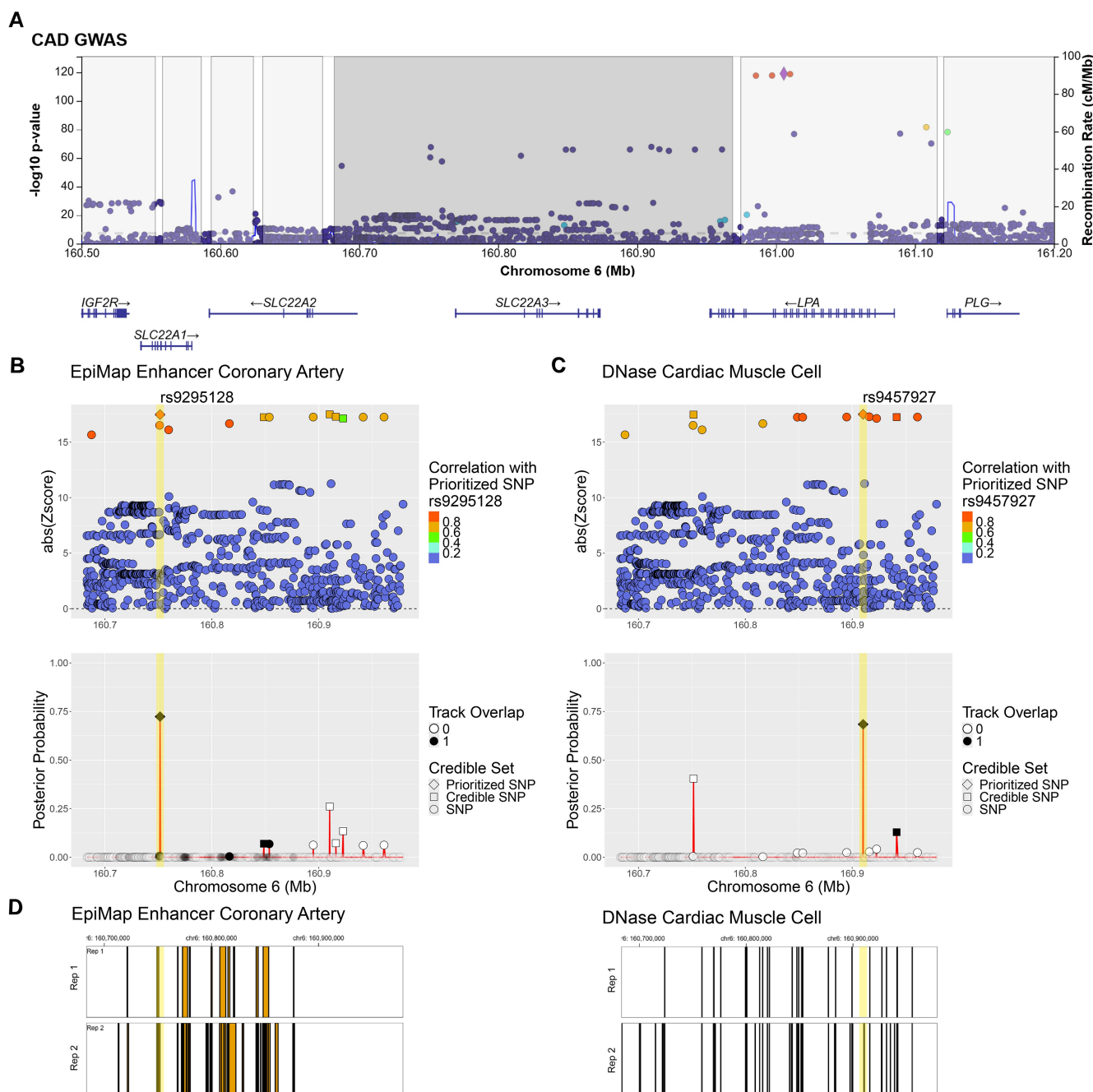


Figure 4. BTS within-region context-specific variant fine-mapping. Shown is an example of a context-dependent variant prioritization in *SLC22A* region, where two variants are co-prioritized at the same time in different functional contexts: rs9295128 is prioritized in the context of coronary artery enhancers (left) and rs9457927 is prioritized in the context of open chromatin regions in cardiac muscle cells (right). The prioritized variant rs9457927 shows an increased posterior of 0.68 compared to 0.3581 null (GWAS+LD) posterior. (A) GWAS results for genomic region encompassing *SLC22A* region (GRCh37/hg19 chr6:160683145–160978997), with discrete LD blocks shaded separately. Subfigures B, C, D show BTS analysis results for the middle (dark grey) LD block. (B) BTS prioritizes rs9295128 in the coronary artery context as it overlaps a CAD-related annotated enhancer (chr6:160750600–160752000) from EpiMAP (BTS estimated annotation prior odds = 2.85; P -value = 1.7×10^{-3}). (C) In a different context (heart), BTS prioritizes rs9457927 in the cardiac muscle cell based on its location within another CAD-relevant cardiac muscle cell open chromatin region (chr6:160910195–160910345) (BTS estimated prior odds of being causal for cardiac muscle chromatin region annotation = 5.51; P -value = 7.1×10^{-4}). (D) Annotated enhancer regions for coronary artery (left) and open chromatin regions for cardiac muscle cell (right) across two biological replicates (Rep1, Rep2) each.

or are in LD with them but also any unrelated or unassociated variants in between. For each of the identified regions of interest, the BTS pipeline used PLINK (Purcell *et al.* 2007, Chang *et al.* 2015) to compute LD matrices based on the 1000 Genomes dataset (Genomes Project Consortium *et al.* 2015, Byrsk-Bishop *et al.* 2022), using genotypes from samples belonging to the European super-population as reference.

BTS used FILER (Kuksa *et al.* 2022) to obtain a discovery set of 943 various regulatory annotations for testing including FANTOM5 enhancers (Andersson *et al.* 2014), Roadmap Epigenomics ChromHMM enhancers (Roadmap Epigenomics Consortium *et al.* 2015), ENCODE DNase hypersensitivity sites (Encode Project Consortium 2012, Encode Project Consortium *et al.* 2020), EpiMap enhancers (Boix

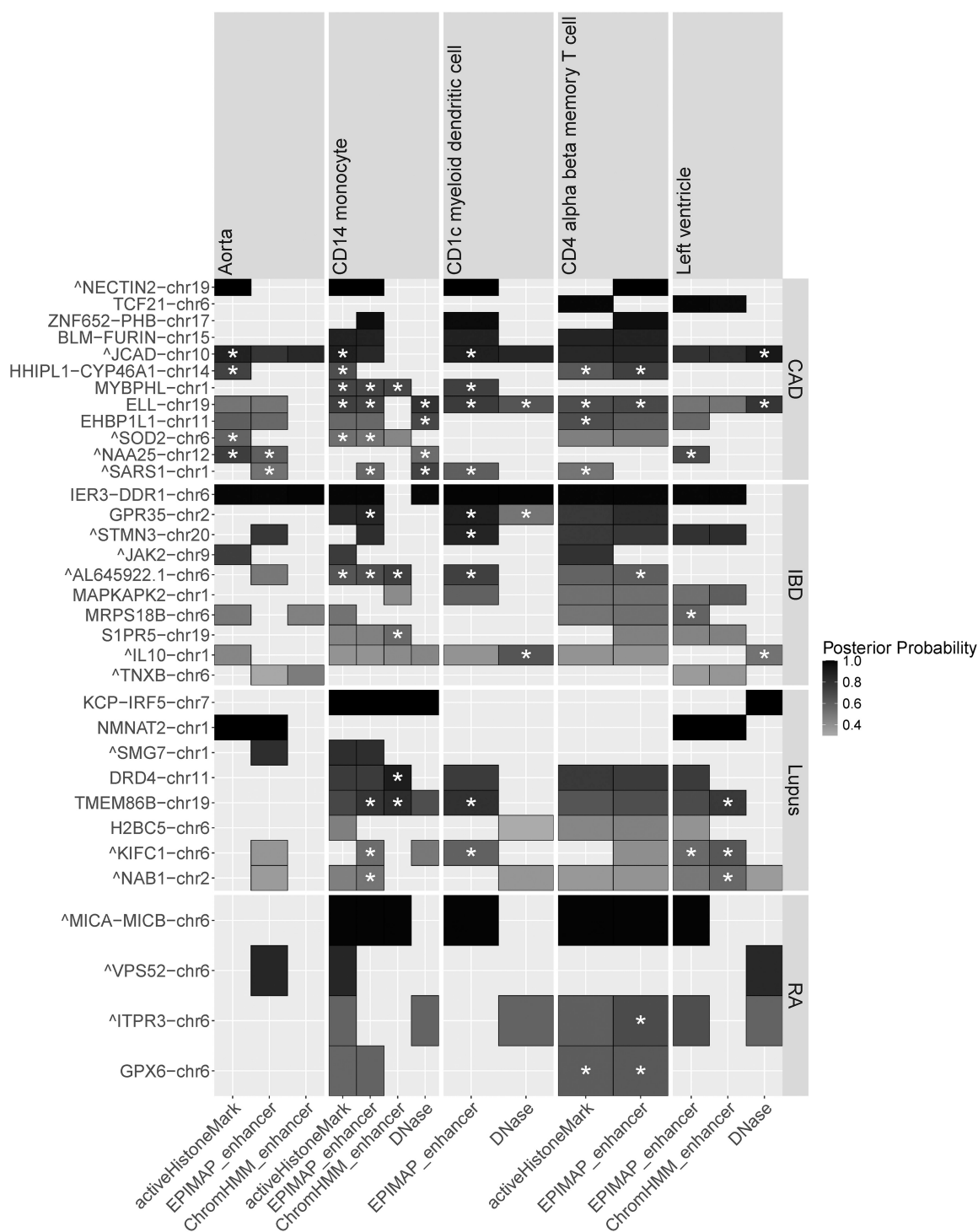


Figure 5. BTS results across four immune-related and cardiovascular GWAS traits. Shown are identified functional contexts (X-axis) and genomic regions (Y-axis) at the cell-type level for each of the traits (CAD, IBD, SLE, and RA panels on the right). For each region and functional context, the top variant posterior is shown (shades of gray) with a star (*) indicating posterior increase of at least 0.2 in that context compared to the null model without annotation. BTS identifies specific cell types and genomic feature types (open chromatin, histone marks, enhancers on the X-axis) for each of the traits (e.g. CD4 T cell for RA; myeloid dendritic cells for IBD; aorta blood vessel and CD4 alpha-beta T cells for CAD; heart left ventricle for SLE).

et al. 2021), and all available tissues and cell types therein (see [Tables 2 and 4](#), available as supplementary data at *Bioinformatics* online). After excluding annotations which did not overlap any of the regions of interest, 532 genome-wide annotation tracks remained across all data sources and tissues. Then BTS was run on the 167 genomic regions, with the default setting of $d=2$, i.e. looking for at most two

distinct causal signals in each region (we note that even though $d=2$, the credible variant set may include more than d variants).

BTS prioritized open chromatin annotations for blood vessel, monocyte and B cells, and other relevant cell types and tissues ([Fig. 3](#)) consistent with the tissues and cell types involved in the disease ([Libby and Theroux 2005](#), [Ghaffar *et al.*](#)

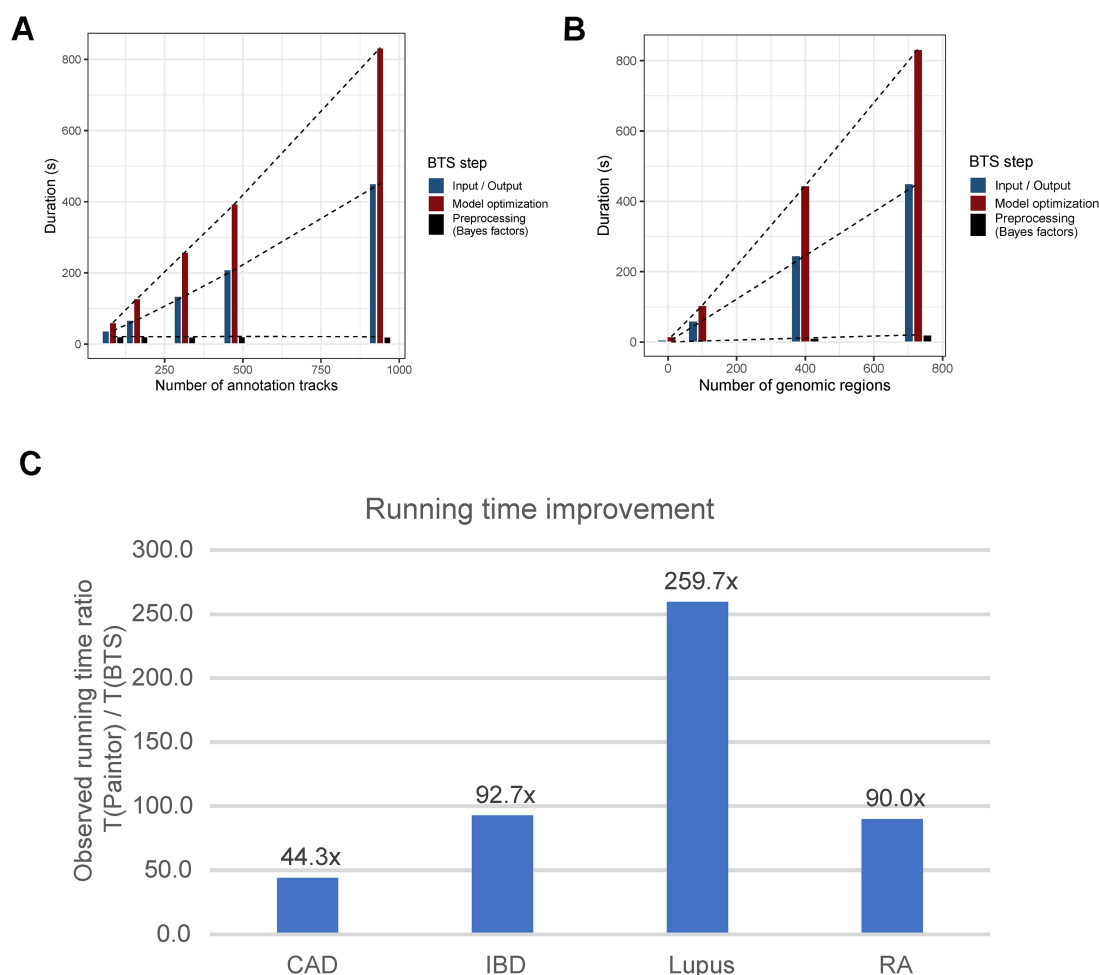


Figure 6. BTS running time. (A) BTS running time as a function of the number of annotation tracks used for functional evaluation. BTS scales linearly with the number of input annotation tracks. Note the constant running time for precomputing annotation-independent Bayes factors. (B) BTS running time as a function of the number of genomic regions. BTS scales linearly with the number of genomic regions of interest. (C) Comparison of the per-model BTS running time and the running time using the reference implementation [fastPainter 3.0 (Kichaev *et al.* 2017)]. BTS improves running time by a factor of 44–259x (average 120x) across the four GWAS datasets analyzed.

2013, Srikakulapu and McNamara 2017). Blood vessel, blood, and immune-related annotations accounted for 23 out of 103 (22%) of the significant annotations (annotation statistical significance assessed with a likelihood ratio test, $P < .01$).

Using cell type and tissue-specific annotations with BTS helped identify a more precise set of potentially causal variants across genomic regions for CAD. Compared to the GWAS+LD null model without annotation, we observed a significant reduction in average credible set sizes for BTS models that use annotations. For example, across 46 genomic regions prioritized by BTS and 200 annotations, we observed the average 90% credible set size of 2.44 ± 2.7 , while the null model had a significantly larger credible set size of 6.5 ± 7.77 across the same genomic regions. Importantly, using annotations allowed BTS to identify 44 additional potentially causal variants that were otherwise not prioritized by the null model. We also note 38 out of 72 variants identified by GWAS+LD-only model were de-prioritized as they were not found to be located within any of the prioritized annotations or functional elements. Overall, using annotations with BTS allowed for the identification of a total of 78 potentially causal CAD variants across 46 loci (see Table 5a, available as supplementary data at *Bioinformatics* online).

Importantly, BTS provides context-specific variant fine-mapping as shown in Fig. 4, where BTS prioritized two variants (rs9295128 and rs9457927) within the same SLC22A2 region, with one variant (rs9295128) prioritized in coronary artery context, and another variant (rs9457927) prioritized in the cardiac muscle cell. This context-dependent prioritization of the variants in LD with each other is enabled using functional annotation information in the BTS model. Table 7, available as supplementary data at *Bioinformatics* online provides a summary of additional annotations supporting prioritization of these two variants using Hi-C, QTL, ATAC-seq, and other annotations from the FILER (Kuksa *et al.* 2022) database.

We note that the variants prioritized by BTS model using functional annotations showed a much higher posterior probability of being causal compared to GWAS+LD-only model without annotations. For example, for the CAD GWAS, across all prioritized contexts/annotations, we have observed a 46% median increase in the posterior probabilities of being causal compared to GWAS+LD-only model posteriors.

We have also validated the prioritized variants *in silico* using FG data from external data sources (see Table 8, available as supplementary data at *Bioinformatics* online). For example, BTS-prioritized variants for CAD were found to be

involved in promoter chromatin interactions in aorta, right atrium, right ventricle (enrichment of 1.84 $\times$, 2.02 $\times$, and 2.9 $\times$, respectively). The variants were also found in ABC (activity-by-contact) interaction sets for heart ventricle, cardiac muscle cell, blood vessel smooth muscle cell. Additionally, 25/58 and 24/58 of prioritized variants for CAD were found to be significant eQTLs in aorta and blood categories from GTEx v8 (GTEx Consortium 2020).

3.3 Cross-trait BTS evaluation

We further tested performance of BTS on several other GWAS datasets for immune-related traits (IBD, RA, SLE) using the same set of 943 annotation tracks as input. As shown in Fig. 5, BTS prioritizes disease and cell type associations consistent with previous studies, including myeloid dendritic cells for IBD (Baumgart *et al.* 2005), T cells for RA (Weyand *et al.* 2000), and monocytes for SLE (Hirose *et al.* 2019) (ChromHMM enhancers) and CAD (Ghaffar *et al.* 2013) (active histone marks, enhancer, and DNase-hypersensitive regions). See Tables 5a–d, available as supplementary data at *Bioinformatics* online for a detailed list of all prioritized annotations, genomic regions, and variants for each tested GWAS dataset.

Moreover, BTS prioritized regions which have been implicated in IBD (Fig. 2, available as supplementary data at *Bioinformatics* online). Among these are an intergenic region near the gene coding for the transcription factor CEBPA (Zhou *et al.* 2019), or genes coding for interleukins and their receptors, such as IL10 (Ip *et al.* 2017) and ILR1 (Dosh *et al.* 2019). Importantly, while recent functional work on IBD (Stankey *et al.* 2024) identified an intergenic region directing macrophage inflammation through ETS2, the ETS2-PSMG1 region was prioritized by BTS in granulocytes (Fig. 2, available as supplementary data at *Bioinformatics* online), belonging to the same myeloid family with macrophage.

3.4 Runtime improvement

BTS only took minutes to process genome-wide GWAS summary statistics data, synthesizing tens of thousands of variants and hundreds of annotation tracks (Fig. 6). Compared to reference implementation fastPaintor/Paintor v3.0 (Kichaev *et al.* 2017), BTS is faster by two orders of magnitude (Fig. 6C; see Section 2.2), being able to compute an annotation-specific model in under one second on average (Fig. 6A and B). This comes at the cost of a preprocessing step that takes a few seconds but is only carried out once for the entire analysis. Moreover, BTS scales linearly (Fig. 6A and B) with both the number of genomic regions and the number of annotation tracks, which guarantees reasonable running times for even larger datasets.

4 Discussion

In this paper, we report BTS, a new algorithm that performs joint fine-mapping of variants and context-mapping using genome-wide functional annotations. The algorithm has multiple statistical and algorithmic innovations to allow one to characterize genetic association signals across the genome against thousands of functional assay experiments across different tissues and cell types. The algorithm provides easily interpretable results that highlight important cellular and tissue context of genetic trait associations with genome-wide enrichment and statistical confidence. The BTS algorithm

implementation is highly scalable and can allow multiple (up to $d = 5$) independent association signals in each locus.

We applied BTS to summary statistics from four immune-related and cardiovascular GWASs, leveraging publicly available functional genomic annotations from multiple data sources, accounting for >200 cell types and tissues. BTS successfully prioritized relevant tissue and cellular types known to be associated with disease biology, validating the underlying statistical model and showing the value of our approach to translate genetic findings to biological mechanisms of disease.

BTS has certain limitations. First, currently available functional annotations are still limited in their specificity in cell and tissue types, and this remains to be addressed in the future as the research community continues to generate highly specific functional experiments. Additionally, users can provide their own specialized functional annotations when running BTS. Second, the number of independent causal variants (parameter d) per locus needs to be chosen in advance and the memory use is still exponential, which prohibits analysis of high-order genetic interactions. In practice, this might not be a serious limitation as the algorithm is still scalable up to $d = 5$. Third, as each functional annotation is tested independently of other annotations, all tissues or cell types that are biologically similar, e.g. in terms of their genomic features such as open chromatin regions or enhancers, will likely be prioritized by the current BTS model thus limiting tissue/cell type resolution. This remains to be addressed in the future, e.g. by explicitly modeling tissue/cell type correlation in the multi-annotation/multi-context model. Fourth, as the model assumes independent causal variants, more complex, epistatic multi-SNP relationships are not explicitly modeled. Fifth, as the current BTS implementation is using a single reference genotype panel (e.g. EUR 1000 genome), more accurate LD estimates could be obtained using GWAS in-sample LD matrices or using, e.g. ancestry-weighted LD [Summix (Arriaga-MacKenzie *et al.* 2021, Stoneman *et al.* 2025), GAUSS (Lee and Bacanu 2024)].

There are several useful usage scenarios for BTS by biologists, geneticists, and other researchers. As complex diseases involve multiple cell types and their interactions, the choice of input tracks and genomic features can reflect multiple hypotheses involving suspected or potentially causal tissues, cell types, and disease mechanisms. BTS can then be used to test these hypotheses and select and prioritize tissues, cell types, genomic mechanisms, functional genomic regions, and variants.

We note that compared to enrichment-based frameworks such as LDSC (Bulik-Sullivan *et al.* 2015) and S-LDSC (Finucane *et al.* 2015), which are commonly used to gain insight into potentially relevant cell types and tissues, BTS not only identifies relevant functional context(s) but also allows users to identify and prioritize functional genomic regions for such contexts. Additionally, BTS provides context-specific fine-mapping and prioritization for variants within each of the analyzed genomic regions and for each of the analyzed contexts.

While multiple overlapping tracks are typically used to show or prioritize functional regions and individual variants (e.g. in the genome browser), BTS can also output per-annotation or per-locus variant posteriors which can be used to select variants with high causal posteriors across annotations. Overall, BTS provides data-driven discovery and prioritization by combining GWAS signals and FG annotation data.

There are many directions for future development of this work. First, although BTS ranks region and annotation pairs in its output, it does not rank regions by themselves, nor does the algorithm provide affected target genes [e.g. as in Activity-by-contact (Fulco *et al.* 2019), Effector Index (Forgetta *et al.* 2022) models] regulated by the causal variants. Incorporation of expression QTL data or chromatin interaction data could address the second problem, and this idea is explored in many other papers including the INFERNO algorithm we reported previously [INFERNO (Amlie-Wolf *et al.* 2018), FUMA (Watanabe *et al.* 2017), SparkINFERNO (Kuksa *et al.* 2020)]. BTS may also be expanded to accommodate for other types of annotations such as experimental functional assays (MPRAs), at the gene or transcript level, and predicted functional activity or pathogenicity of variants [e.g. GENO-NET (He *et al.* 2018), CADD (Rentzsch *et al.* 2019, 2021); JARVIS (Vitsios *et al.* 2021); PoPS (Weeks *et al.* 2023)], but new statistical models need to be developed to integrate these types of data.

A key advantage of BTS is its scalability, allowing us to leverage thousands of genome-wide functional annotations and systematically explore the biological mechanism and cell or tissue context without bias. One may extend our algorithm to analyze multiple traits and tissues jointly or carry out combinatorial analysis of multiple cell type and tissue functional surveys across many genomic features. Applications of such an approach to the growing body of GWAS and WGS-based studies, along with further experimental validations, could lead to significant insights into disease-underlying variants, loci, and molecular mechanisms.

Acknowledgements

We thank Wang Lab Members, Mingyao Li, and Struan Grant for their valuable feedback and comments.

Author contributions

Pavel P. Kuksa (Conceptualization [equal], Data curation [equal], Formal analysis [equal], Investigation [equal], Methodology [equal], Resources [equal], Software [equal], Writing—original draft [equal], Writing—review & editing [equal]), Matei Ionita (Conceptualization [equal], Data curation [equal], Formal analysis [equal], Investigation [equal], Methodology [equal], Resources [equal], Software [equal], Writing—original draft [equal], Writing—review & editing [equal]), Luke Carter (Investigation [supporting], Resources [supporting], Software [supporting], Visualization [equal], Writing—review & editing [supporting]), Jeffrey Cifello (Data curation [supporting], Resources [supporting], Writing—review & editing [supporting]), Prabhakaran Gangadharan (Data curation [supporting], Resources [supporting]), Kaylyn Clark (Data curation [supporting], Resources [supporting], Writing—review & editing [supporting]), Otto Valladares (Resources [supporting], Software [supporting]), Yuk Yee Leung (Conceptualization [equal], Data curation [equal], Formal analysis [equal], Methodology [equal], Writing—original draft [equal], Writing—review & editing [equal]), and Li-San Wang (Conceptualization [equal], Funding acquisition [lead], Investigation [equal], Methodology [equal], Resources [equal], Supervision [lead], Writing—original draft [equal], Writing—review & editing [equal])

Supplementary data

Supplementary data are available at *Bioinformatics* online.

Conflict of interest: none declared.

Funding

This work has been supported by the National Institute on Aging [U24-AG041689, U54-AG052427, U01-AG032984].

Data availability

GWAS summary statistics used in this study were obtained from EMBL-EBI GWAS catalog for IBD (GCST003043), RA (GCST000679), and SLE (GCST003156), and from MRC IEU OpenGWAS database for CAD (ebi-a-GCST005195). ENCODE (<https://encodeproject.org>), Roadmap Epigenomics (https://egg2.wustl.edu/roadmap/web_portal), and EpiMap (<https://compbio.mit.edu/epimap/>) functional annotations used in BTS evaluations and 1000 Genomes (<https://www.internationalgenome.org/>) reference genotype panel data for LD computation were obtained from the FILER database (<https://lisanwanglab.org/FILER>). BTS code used in the analyses and results on tested GWASs are available from <https://doi.org/10.5281/zenodo.13286305>. BTS runs on HPC clusters, single servers, or cloud-based instances and is available at <https://hub.docker.com/r/wanglab/bts> (Docker container) and <https://bitbucket.org/wanglab-upenn/bts-pipeline> (BTS pipeline source code). BTS core R package is also available at <https://bitbucket.com/wanglab-upenn/BTS-R>.

References

- Amlie-Wolf A, Tang M, Mlynarski EE *et al.* INFERNO: inferring the molecular mechanisms of noncoding genetic variants. *Nucleic Acids Res* 2018;**46**:8740–53. <https://doi.org/10.1093/nar/gky686>
- Andersson R, Gebhard C, Miguel-Escalada I *et al.* An atlas of active enhancers across human cell types and tissues. *Nature* 2014;**507**: 455–61. <https://doi.org/10.1038/nature12787>
- Arriaga-Mackenzie IS, Matesi G, Chen S *et al.* Summix: a method for detecting and adjusting for population structure in genetic summary data. *Am J Hum Genet* 2021;**108**:1270–82. <https://doi.org/10.1016/j.ajhg.2021.05.016>
- Asimit JL, Rainbow DB, Fortune MD *et al.* Stochastic search and joint fine-mapping increases accuracy and identifies previously unreported associations in immune-mediated diseases. *Nat Commun* 2019;**10**:3216. <https://doi.org/10.1038/s41467-019-11271-0>
- Baumgart DC, Metzke D, Schmitz J *et al.* Patients with active inflammatory bowel disease lack immature peripheral blood plasmacytoid and myeloid dendritic cells. *Gut* 2005;**54**:228–36. <https://doi.org/10.1136/gut.2004.040360>
- Benner C, Spencer CC, Havulinna AS *et al.* FINEMAP: efficient variable selection using summary data from genome-wide association studies. *Bioinformatics* 2016;**32**:1493–501. <https://doi.org/10.1093/bioinformatics/btw018>
- Bentham J, Morris DL, Graham DSC *et al.* Genetic association analyses implicate aberrant regulation of innate and adaptive immunity genes in the pathogenesis of systemic lupus erythematosus. *Nat Genet* 2015;**47**:1457–64. <https://doi.org/10.1038/ng.3434>
- Boix CA, James BT, Park YP *et al.* Regulatory genomic circuitry of human disease loci by integrative epigenomics. *Nature* 2021;**590**: 300–7. <https://doi.org/10.1038/s41586-020-03145-z>
- Boyle AP, Hong EL, Hariharan M *et al.* Annotation of functional variation in personal genomes using RegulomeDB. *Genome Res* 2012;**22**:1790–7. <https://doi.org/10.1101/gr.137323.112>

- Bulik-Sullivan BK, Loh PR, Finucane HK *et al.*; Schizophrenia Working Group of the Psychiatric Genomics Consortium. LD score regression distinguishes confounding from polygenicity in genome-wide association studies. *Nat Genet* 2015;47:291–5. <https://doi.org/10.1038/ng.3211>
- Byrska-Bishop M, Evani US, Zhao X *et al.*; Human Genome Structural Variation Consortium. High-coverage whole-genome sequencing of the expanded 1000 genomes project cohort including 602 trios. *Cell* 2022;185:3426–40.e19. <https://doi.org/10.1016/j.cell.2022.08.004>
- Chang CC, Chow CC, Tellier LC *et al.* Second-generation PLINK: rising to the challenge of larger and richer datasets. *Gigascience* 2015; 4:7. <https://doi.org/10.1186/s13742-015-0047-8>
- Chen J, Wang L, De Jager PL *et al.* A scalable Bayesian functional GWAS method accounting for multivariate quantitative functional annotations with applications for studying Alzheimer disease. *HGG Adv* 2022;3:100143. <https://doi.org/10.1016/j.xhgg.2022.100143>
- Chen W, Larrabee BR, Ovsyannikova IG *et al.* Fine mapping causal variants with an approximate Bayesian method using marginal test statistics. *Genetics* 2015;200:719–36. <https://doi.org/10.1534/genetics.115.176107>
- Chen W, McDonnell SK, Thibodeau SN *et al.* Incorporating functional annotations for fine-mapping causal variants in a Bayesian framework using summary statistics. *Genetics* 2016;204:933–58. <https://doi.org/10.1534/genetics.116.188953>
- Chen W, Wu Y, Zheng Z *et al.* Improved analyses of GWAS summary statistics by reducing data heterogeneity and errors. *Nat Commun* 2021;12:7117. <https://doi.org/10.1038/s41467-021-27438-7>
- Dong S, Zhao N, Spragins E *et al.* Annotating and prioritizing human non-coding variants with RegulomeDB v.2. *Nat Genet* 2023;55: 724–6. <https://doi.org/10.1038/s41588-023-01365-3>
- Dosh RH, Jordan-Mahy N, Sammon C *et al.* Interleukin 1 is a key driver of inflammatory bowel disease-demonstration in a murine IL-1Ra knockout model. *Oncotarget* 2019;10:3559–75. <https://doi.org/10.18632/oncotarget.26894>
- Encode Project Consortium. An integrated encyclopedia of DNA elements in the human genome. *Nature* 2012;489:57–74. <https://doi.org/10.1038/nature11247>
- Encode Project Consortium; Moore JE, Purcaro MJ, Pratt HE *et al.* Expanded encyclopaedias of DNA elements in the human and mouse genomes. *Nature* 2020;583:699–710. <https://doi.org/10.1038/s41586-020-2493-4>
- Finucane HK, Bulik-Sullivan B, Gusev A *et al.*; RACI Consortium. Partitioning heritability by functional annotation using genome-wide association summary statistics. *Nat Genet* 2015;47:1228–35. <https://doi.org/10.1038/ng.3404>
- Foley CN, Staley JR, Breen PG *et al.* A fast and efficient colocalization algorithm for identifying shared genetic risk factors across multiple traits. *Nat Commun* 2021;12:764. <https://doi.org/10.1038/s41467-020-20885-8>
- Forgetta V, Jiang L, Vulpescu NA *et al.* An effector index to predict target genes at GWAS loci. *Hum Genet* 2022;141:1431–47. <https://doi.org/10.1007/s00439-022-02434-z>
- Fulco CP, Nasser J, Jones TR *et al.* Activity-by-contact model of enhancer-promoter regulation from thousands of CRISPR perturbations. *Nat Genet* 2019;51:1664–9. <https://doi.org/10.1038/s41588-019-0538-0>
- Galarneau G, Palmer CD, Sankaran VG *et al.* Fine-mapping at three loci known to affect fetal hemoglobin levels explains additional genetic variation. *Nat Genet* 2010;42:1049–51. <https://doi.org/10.1038/ng.707>
- Genomes Project Consortium; Auton A, Brooks LD, Durbin RM *et al.* A global reference for human genetic variation. *Nature* 2015;526: 68–74. <https://doi.org/10.1038/nature15393>
- Ghaffar A, Griffiths HR, Devitt A *et al.* Monocytes in coronary artery disease and atherosclerosis: where are we now? *J Am Coll Cardiol* 2013;62:1541–51. <https://doi.org/10.1016/j.jacc.2013.07.043>
- Giambartolomei C, Vukcevic D, Schadt EE *et al.* Bayesian test for colocalisation between pairs of genetic association studies using summary statistics. *PLoS Genet* 2014;10:e1004383. <https://doi.org/10.1371/journal.pgen.1004383>
- GTEx Consortium. The GTEx Consortium Atlas of genetic regulatory effects across human tissues. *Science* 2020;369:1318–30. <https://doi.org/10.1126/science.aaz1776>
- He Z, Liu L, Wang K *et al.* A semi-supervised approach for predicting cell-type specific functional consequences of non-coding variation using MPRA. *Nat Commun* 2018;9:5199. <https://doi.org/10.1038/s41467-018-07349-w>
- Hirose S, Lin Q, Ohtsui M *et al.* Monocyte subsets involved in the development of systemic lupus erythematosus and rheumatoid arthritis. *Int Immunol* 2019;31:687–96. <https://doi.org/10.1093/intimm/dxz036>
- Hormozdiari F, Kostem E, Kang EY *et al.* Identifying causal variants at loci with multiple signals of association. *Genetics* 2014;198: 497–508. <https://doi.org/10.1534/genetics.114.167908>
- Hormozdiari F, van de Bunt M, Segre AV *et al.* Colocalization of GWAS and eQTL signals detects target genes. *Am J Hum Genet* 2016;99:1245–60. <https://doi.org/10.1016/j.ajhg.2016.10.003>
- Ip WKE, Hoshi N, Shouval DS *et al.* Anti-inflammatory effect of IL-10 mediated by metabolic reprogramming of macrophages. *Science* 2017;356:513–9. <https://doi.org/10.1126/science.aal3535>
- Kanai M, Elzur R, Zhou W *et al.*; Global Biobank Meta-Analysis Initiative. Meta-analysis fine-mapping is often miscalibrated at single-variant resolution. *Cell Genom* 2022;2:100210. <https://doi.org/10.1016/j.xgen.2022.100210>
- Kichaev G, Roytman M, Johnson R *et al.* Improved methods for multi-trait fine mapping of pleiotropic risk loci. *Bioinformatics* 2017;33: 248–55. <https://doi.org/10.1093/bioinformatics/btw615>
- Kichaev G, Yang WY, Lindstrom S *et al.* Integrating functional data to prioritize causal variants in statistical fine-mapping studies. *PLoS Genet* 2014;10:e1004722. <https://doi.org/10.1371/journal.pgen.1004722>
- Knight J, Spain SL, Capon F *et al.*; I-chip for Psoriasis Consortium. Conditional analysis identifies three novel major histocompatibility complex loci associated with psoriasis. *Hum Mol Genet* 2012;21: 5185–92. <https://doi.org/10.1093/hmg/dds344>
- Kuksa PP, Lee CY, Amlie-Wolf A *et al.* SparkINFERNO: a scalable high-throughput pipeline for inferring molecular mechanisms of non-coding genetic variants. *Bioinformatics* 2020;36:3879–81. <https://doi.org/10.1093/bioinformatics/btaa246>
- Kuksa PP, Leung YY, Gangadharan P *et al.* FILER: a framework for harmonizing and querying large-scale functional genomics knowledge. *NAR Genom Bioinform* 2022;4:lqab123. <https://doi.org/10.1093/nargab/lqab123>
- Lee D, Bacanu SA. GAUSS: a summary-statistics-based R package for accurate estimation of linkage disequilibrium for variants, gaussian imputation, and TWAS analysis of cosmopolitan cohorts. *Bioinformatics* 2024;40:btac203. <https://doi.org/10.1093/bioinformatics/btac203>
- Libby P, Theroux P. Pathophysiology of coronary artery disease. *Circulation* 2005;111:3481–8. <https://doi.org/10.1161/CIRCULATIONAHA.105.537878>
- Liu JZ, van Sommeren S, Huang H *et al.*; International IBD Genetics Consortium. Association analyses identify 38 susceptibility loci for inflammatory bowel disease and highlight shared genetic risk across populations. *Nat Genet* 2015;47:979–86. <https://doi.org/10.1038/ng.3359>
- Pickrell JK. Joint analysis of functional genomic data and genome-wide association studies of 18 human traits. *Am J Hum Genet* 2014;94: 559–73. <https://doi.org/10.1016/j.ajhg.2014.03.004>
- Purcell S, Neale B, Todd-Brown K *et al.* PLINK: a tool set for whole-genome association and population-based linkage analyses. *Am J Hum Genet* 2007;81:559–75. <https://doi.org/10.1086/519795>
- Rentzsch P, Schubach M, Shendure J *et al.* CADD-Splice-improving genome-wide variant effect prediction using deep learning-derived splice scores. *Genome Med* 2021;13:31. <https://doi.org/10.1186/s13073-021-00835-9>
- Rentzsch P, Witten D, Cooper GM *et al.* CADD: predicting the deleteriousness of variants throughout the human genome. *Nucleic Acids Res* 2019;47:D886–94. <https://doi.org/10.1093/nar/gky1016>

- Roadmap Epigenomics Consortium; Kundaje A, Meuleman W, Ernst J *et al.* Integrative analysis of 111 reference human epigenomes. *Nature* 2015;518:317–30. <https://doi.org/10.1038/nature14248>
- Srikakulapu P, McNamara CA. B cells and atherosclerosis. *Am J Physiol Heart Circ Physiol* 2017;312:H1060–7. <https://doi.org/10.1152/ajpheart.00859.2016>
- Stahl EA, Raychaudhuri S, Remmers EF *et al.*; YEAR Consortium. Genome-wide association study meta-analysis identifies seven new rheumatoid arthritis risk loci. *Nat Genet* 2010;42:508–14. <https://doi.org/10.1038/ng.582>
- Stankey CT, Bourges C, Haag LM *et al.* A disease-associated gene desert directs macrophage inflammation through ETS2. *Nature* 2024; 630:447–56. <https://doi.org/10.1038/s41586-024-07501-1>
- Stoneman HR, Price AM, Trout NS *et al.*; Colorado Center for Personalized Medicine. Characterizing substructure via mixture modeling in large-scale genetic summary statistics. *Am J Hum Genet* 2025;112:235–53. <https://doi.org/10.1016/j.ajhg.2024.12.007>
- Uffelmann E, Huang QQ, Munung NS *et al.* Genome-wide association studies. *Nat Rev Methods Primers* 2021;1:59. <https://doi.org/10.1038/s43586-021-00056-9>
- van der Harst P, Verweij N. Identification of 64 novel genetic loci provides an expanded view on the genetic architecture of coronary artery disease. *Circ Res* 2018;122:433–43. <https://doi.org/10.1161/CIRCRESAHA.117.312086>
- Vitsios D, Dhindsa RS, Middleton L *et al.* Prioritizing non-coding regions based on human genomic constraint and sequence context with deep learning. *Nat Commun* 2021;12:1504. <https://doi.org/10.1038/s41467-021-21790-4>
- Wallace C. A more accurate method for colocalisation analysis allowing for multiple causal variants. *PLoS Genet* 2021;17:e1009440. <https://doi.org/10.1371/journal.pgen.1009440>
- Wang G, Sarkar A, Carbonetto P *et al.* A simple new approach to variable selection in regression, with application to genetic fine mapping. *J R Stat Soc Series B Stat Methodol* 2020;82:1273–300. <https://doi.org/10.1111/rssb.12388>
- Watanabe K, Taskesen E, van Bochoven A *et al.* Functional mapping and annotation of genetic associations with FUMA. *Nat Commun* 2017;8:1826. <https://doi.org/10.1038/s41467-017-01261-5>
- Weeks EM, Ulirsch JC, Cheng NY *et al.* Leveraging polygenic enrichments of gene features to predict genes underlying complex traits and diseases. *Nat Genet* 2023;55:1267–76. <https://doi.org/10.1038/s41588-023-01443-6>
- Weissbrod O, Hormozdiari F, Benner C *et al.* Functionally informed fine-mapping and polygenic localization of complex trait heritability. *Nat Genet* 2020;52:1355–63. <https://doi.org/10.1038/s41588-020-00735-5>
- Wen X, Lee Y, Luca F *et al.* Efficient integrative Multi-SNP association analysis via deterministic approximation of posteriors. *Am J Hum Genet* 2016;98:1114–29. <https://doi.org/10.1016/j.ajhg.2016.03.029>
- Wen X, Pique-Regi R, Luca F. Integrating molecular QTL data into genome-wide genetic association analysis: probabilistic assessment of enrichment and colocalization. *PLoS Genet* 2017;13:e1006646. <https://doi.org/10.1371/journal.pgen.1006646>
- Weyand CM, Bryl E, Goronzy JJ. The role of T cells in rheumatoid arthritis. *Arch Immunol Ther Exp (Warsz)* 2000;48:429–35. <https://www.ncbi.nlm.nih.gov/pubmed/11140470>
- Yang J, Ferreira T, Morris AP *et al.*; DIABetes Genetics Replication And Meta-analysis (DIAGRAM) Consortium. Conditional and joint multiple-SNP analysis of GWAS summary statistics identifies additional variants influencing complex traits. *Nat Genet* 2012;44: 369–75, S1–3. <https://doi.org/10.1038/ng.2213>
- Yang Z, Wang C, Liu L *et al.* CARMA is a new Bayesian model for fine-mapping in genome-wide association meta-analyses. *Nat Genet* 2023;55:1057–65. <https://doi.org/10.1038/s41588-023-01392-0>
- Zhou J, Li H, Xia X *et al.* Anti-inflammatory activity of MTL-CEBPA, a small activating RNA drug, in LPS-Stimulated monocytes and humanized mice. *Mol Ther* 2019;27:999–1016. <https://doi.org/10.1016/j.ymthe.2019.02.018>
- Zou Y, Carbonetto P, Wang G *et al.* Fine-mapping from summary data with the “sum of single effects” model. *PLoS Genet* 2022;18: e1010299. <https://doi.org/10.1371/journal.pgen.1010299>