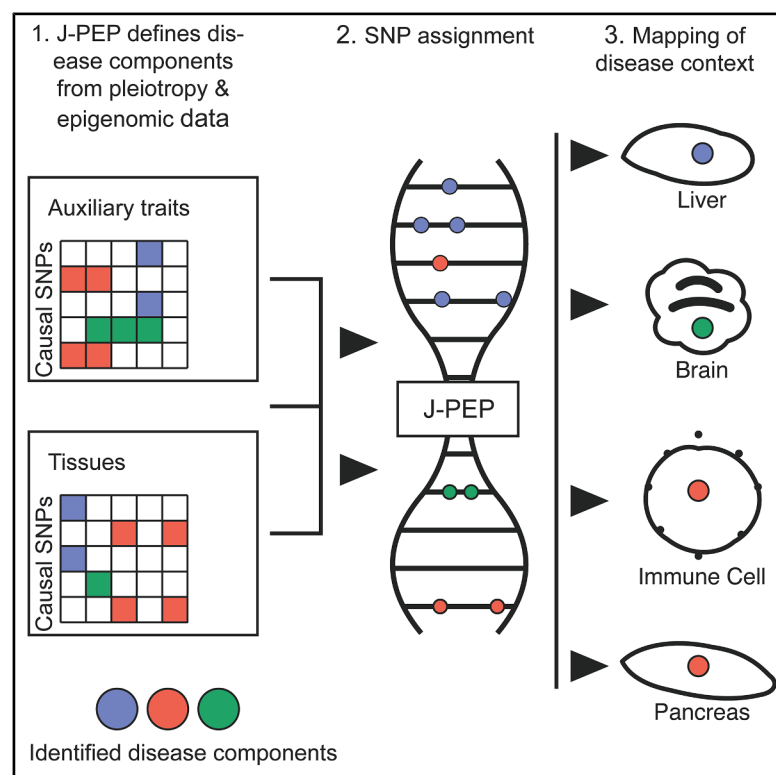


Mapping disease loci to biological processes via joint pleiotropic and epigenomic partitioning

Graphical abstract



Authors

Gaspard Kerner, Nolan Kamitaki, Benjamin Strober, Alkes L. Price

Correspondence

gkerner@hsph.harvard.edu (G.K.),
aprice@hsph.harvard.edu (A.L.P.)

In brief

Kerner et al. introduce J-PEP, a framework that integrates pleiotropic and epigenomic data to partition GWAS loci into biologically meaningful clusters. Using a new cross-modal prediction metric and single-cell epigenomics, they reveal tissue- and cell-type-specific processes underlying diverse complex diseases.

Highlights

- J-PEP integrates pleiotropic and epigenomic data to cluster GWAS loci
- Cross-modal PEPA metric quantifies trait-tissue prediction accuracy
- J-PEP improves biological interpretability across 165 diseases/traits
- Single-cell epigenomics refines clusters to cell-type resolution



Article

Mapping disease loci to biological processes via joint pleiotropic and epigenomic partitioning

Gaspard Kerner,^{1,2,6,*} Nolan Kamitaki,³ Benjamin Strober,^{1,4} and Alkes L. Price^{1,2,5,*}¹Department of Epidemiology, Harvard T.H. Chan School of Public Health, Boston, MA, USA²Broad Institute of MIT and Harvard, Boston, MA, USA³Department of Biomedical Informatics, Harvard Medical School, Boston, MA, USA⁴Computational Health Informatics Program, Boston Children's Hospital, Boston, MA, USA⁵Department of Biostatistics, Harvard T.H. Chan School of Public Health, Boston, MA, USA⁶Lead contact*Correspondence: gkerner@hsph.harvard.edu (G.K.), aprice@hsph.harvard.edu (A.L.P.)<https://doi.org/10.1016/j.xgen.2025.101138>

SUMMARY

Genome-wide association studies have identified thousands of disease-associated loci, yet their biological interpretation remains limited. We propose joint pleiotropic and epigenomic partitioning (J-PEP), a clustering framework that integrates pleiotropic SNP effects on auxiliary traits and tissue-specific epigenomic data to partition disease-associated loci into biologically distinct clusters. We introduce a metric—pleiotropic and epigenomic prediction accuracy (PEPA)—that evaluates how well the clusters predict SNP-to-trait and SNP-to-tissue associations in off-chromosome data. Analyzing summary statistics for 165 diseases/traits (average $N = 290,000$), J-PEP attained 16%–30% higher PEPA than pleiotropic or epigenomic partitioning approaches, with larger improvements for well-powered traits, consistent with simulations; these gains arise from J-PEP's tendency to upweight signals present in both auxiliary trait and tissue data, emphasizing shared components. Notably, integrating single-cell chromatin accessibility data refined bulk-based clusters, enhancing cell-type resolution and specificity. For type 2 diabetes, hypertension, and other diseases/traits, J-PEP clusters recapitulated known pathways while revealing underexplored biological processes.

INTRODUCTION

A central challenge in complex disease genomics is to decipher the biological processes underlying the thousands of loci identified by genome-wide association studies (GWASs).^{1–6} While interpretation of individual loci has yielded important insights into disease biology and functionally validated putative causal genes,^{7–11} findings at individual loci remain limited,^{1,12} motivating genome-wide approaches to map disease loci to biological processes. Extensive pleiotropy between diseases and complex traits^{13–22} motivates pleiotropic partitioning^{23–25} (see also Zhang et al.,^{26–31} Kim et al.,^{26–31} Reeve et al.,^{26–31} Qi et al.,^{26–31} Ghatan et al.,^{26–31} and Li et al.^{26–31})—leveraging pleiotropic associations with auxiliary traits to group disease loci into clusters that often align with disease-critical cell types and biological pathways. However, interpreting clusters based solely on pleiotropy can be challenging when biological processes are partially overlapping or produce similar multi-trait signatures. Epigenomic annotations provide valuable information about disease biology^{32–37} and have been used to retrospectively interpret clustering results^{23–25} but have not been directly integrated into clustering models. Furthermore, no quantitative metric has been proposed to evaluate clustering performance.

Here, we introduce joint pleiotropic and epigenomic partitioning (J-PEP), a method that jointly partitions disease loci into clus-

ters using both pleiotropic associations with auxiliary traits and tissue-specific epigenomic profiles under a tissue sparsity constraint that improves algorithmic stability and interpretability; J-PEP clusters capture both SNP-to-trait and SNP-to-tissue associations, enhancing cluster interpretability. We also introduce a new evaluation metric, pleiotropic and epigenomic prediction accuracy (PEPA), which evaluates how well the clusters predict SNP-to-trait and SNP-to-tissue associations. We compare J-PEP to pleiotropic and epigenomic partitioning approaches via simulations and analyses of 165 GWAS diseases/traits using the PEPA metric. We find that J-PEP outperforms other approaches using the PEPA metric, and we highlight novel biological insights for type 2 diabetes (T2D), hypertension (HTN), and neutrophil count (NC). Finally, we show that by integrating single-cell chromatin accessibility data, J-PEP enables the association of clusters—i.e., subsets of loci—with fine-grained cell types, refining broader bulk-derived associations.

RESULTS

Overview of methods

J-PEP partitions disease loci into biologically interpretable clusters by jointly leveraging pleiotropic data (across auxiliary traits) and epigenomic data (across tissues) (Figure 1). For a given focal disease/trait, J-PEP inputs a SNP-to-trait matrix (V_{trait}) derived



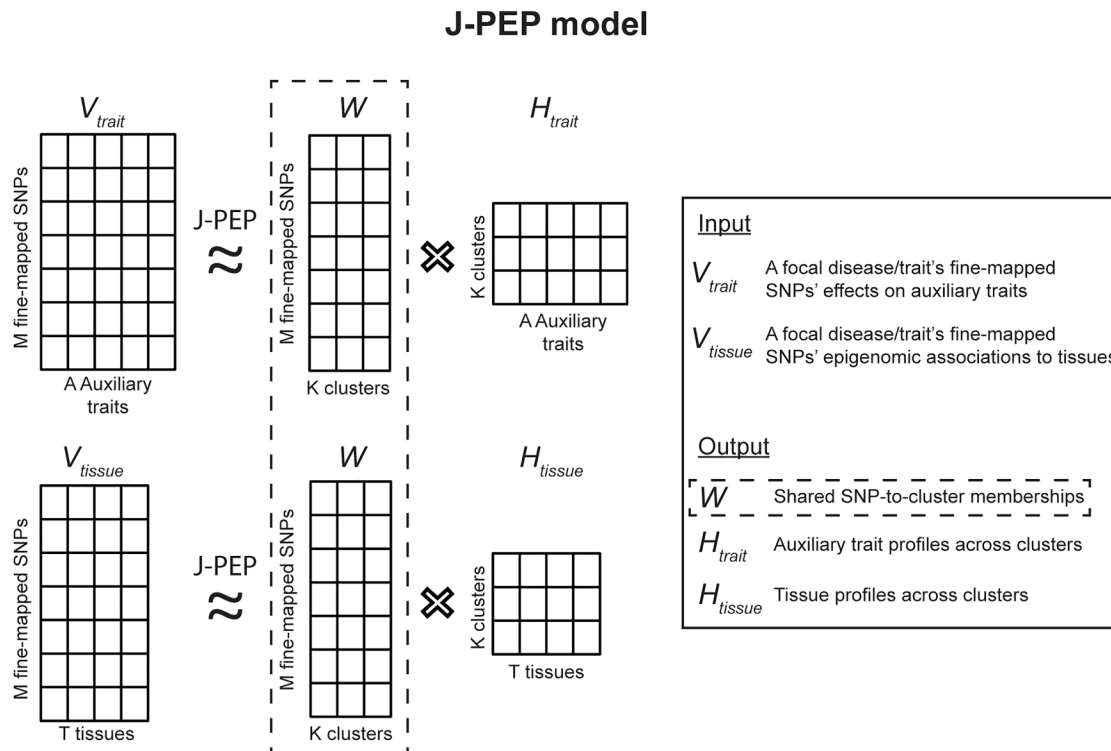


Figure 1. Schematic representation of the J-PEP model

The J-PEP model performs an extended version of joint Bayesian non-negative matrix factorization (bNMF) to fine-mapped SNPs from a given focal disease/trait. It jointly factorizes two input matrices: a SNP-to-trait matrix (V_{trait}) and a SNP-to-tissue matrix (V_{tissue}). J-PEP infers a shared SNP-to-cluster membership matrix (W), an auxiliary trait-to-cluster profile matrix (H_{trait}) and a tissue-to-cluster profile matrix (H_{tissue}).

from fine-mapping results^{38,39} across auxiliary traits, with each entry defined as the product of posterior inclusion probabilities (PIPs) for a given fine-mapped SNP (focal trait PIP > 0.01) across the focal trait and a given auxiliary trait (STAR Methods). To accommodate the sign of pleiotropic effects, each auxiliary trait is split into two components: one representing concordant (or positive) pleiotropic effects and the other representing discordant (or negative) pleiotropic effects (STAR Methods). J-PEP also inputs a SNP-to-tissue matrix (V_{tissue}) derived from epigenomic data across tissues, where each entry reflects the normalized strength of association between a fine-mapped SNP and a tissue, computed using an expectation maximization (EM) algorithm (STAR Methods). J-PEP then jointly factorizes V_{trait} and V_{tissue} to produce estimates of a SNP-to-cluster membership matrix (W), an auxiliary trait-to-cluster profile matrix (H_{trait}), and a tissue-to-cluster profile matrix (H_{tissue}), using an extended version of joint Bayesian non-negative matrix factorization (bNMF).^{40,41} Standard bNMF decomposes a single non-negative data matrix into lower-dimensional non-negative factors under Bayesian priors; J-PEP extends this framework to two matrices, encouraging cross-data alignment across clusters (see below). To enhance cluster interpretability, J-PEP restricts each tissue to at most a single cluster, whereas auxiliary traits may span multiple clusters. This iterative optimization process continues until convergence, yielding clusters that integrate pleiotropic and epigenomic information. To enhance stability, we perform ten opti-

mization runs with different random seeds and select the run with the highest average projection correlation between modalities (i.e., traits and tissue) among those yielding the most consistent number of identified clusters (STAR Methods). We note that of all input/output matrices, only V_{tissue} has normalized rows that specifically represent probabilities.

In detail, J-PEP solves the following optimization problem:

$$\{W, H_{\text{trait}}, H_{\text{tissue}}\} = \underset{W \geq 0, H_{\text{trait}} \geq 0, H_{\text{tissue}} \geq 0}{\operatorname{argmin}} \left\{ \frac{1}{2} \|V_{\text{trait}} - WH_{\text{trait}}\|^2 + \frac{1}{2} \|V_{\text{tissue}} - WH_{\text{tissue}}\|^2 + \sum_k \lambda \left((V_{\text{trait}} H_{\text{trait}}^T)_k, (V_{\text{tissue}} H_{\text{tissue}}^T)_k \right) \right\},$$

(Equation 1)

where W , H_{trait} , H_{tissue} , V_{trait} , and V_{tissue} are as defined above, and $(V_{\text{trait}} H_{\text{trait}}^T)_k$ and $(V_{\text{tissue}} H_{\text{tissue}}^T)_k$ are vectors of length M (the number of focal disease/trait fine-mapped SNPs with PIP > 0.01) representing the projections of fine-mapped SNPs onto the k^{th} cluster as defined by auxiliary traits and tissues, respectively. To guide the factorization process, the first two terms of the optimization function minimize the reconstruction error of the input matrices, while the third term is a penalization

term that promotes alignment between pleiotropic and epigenomic data and suppresses irrelevant clusters. Specifically, $\lambda(A, B) = -\log_{10}(p_{corr}(A, B))^{-1}$, where p_{corr} indicates the one-sided Pearson correlation p value for positive correlation. Additionally, J-PEP guarantees sparsity in tissue profiles by setting all non-maximal tissue-to-cluster memberships (columns of H_{tissue}) to zero, leaving each tissue assigned to at most a single cluster. Conceptually, J-PEP optimizes the objective function of Equation 1 under the constraint that $H_{tissue} \in C$, where C is the set of matrices whose columns each have at most one non-zero entry. The initial number of clusters (set to 15 following prior work²³) serves as an upper bound and may be reduced during optimization as uninformative clusters are pruned.

We systematically compared J-PEP to alternative methods for partitioning disease-associated loci. First, we considered two simpler single-modality approaches: pleiotropic partitioning,^{23–25} which factorizes a SNP-to-trait matrix (e.g., V_{trait}) without incorporating tissue information; and epigenomic partitioning, which factorizes a SNP-to-tissue matrix (e.g., V_{tissue}) without incorporating pleiotropic information (STAR Methods). In both approaches, partitioning is performed using standard bNMF. For pleiotropic partitioning, we learn H_{tissue} from V_{tissue} and W using gradient descent by solving $V_{tissue} = W \times H_{tissue}$. For epigenomic partitioning, we learn H_{trait} from V_{trait} and W using gradient descent by solving $V_{trait} = W \times H_{trait}$. To support comparisons between methods, the sparsity constraint on tissue profiles applied to J-PEP was also imposed on the two simpler approaches (STAR Methods). Second, we benchmarked J-PEP against two independently developed software packages that perform pleiotropic partitioning: FactorGO²⁶ and Flashier.⁴² FactorGO is a variational Bayesian factor analysis model that decomposes GWAS Z-score matrices into latent pleiotropic factors by modeling true genetic effects with uncertainty, using priors to prune uninformative factors. Flashier is an empirical Bayes matrix factorization framework that assumes the observed data can be approximated by the product of two low-rank latent matrices (with Gaussian noise) and places priors on factors; we benchmarked three versions of Flashier that differ in their prior specifications, regularization strategies, and input data.

We applied J-PEP, pleiotropic partitioning, and epigenomic partitioning to GWAS summary statistics for 165 focal diseases/traits (average of 290,000 samples and 77 fine-mapped loci; a fine-mapped locus is defined as a locus containing at least one SNP with $PIP > 0.5$), including 57 UK Biobank diseases/traits⁴³ and 108 non-UK Biobank diseases/traits (Table S1; see data and code availability). We integrated EpiMap data³⁵ comprising 1,595 chromatin tracks, each defined by a unique combination of one of six histone or accessibility marks (DNase HS, H3K4me1, H3K4me2, H3K4me3, H3K9ac, and H3K27ac) and a cell-type label. These tracks are grouped into 32 broad tissue categories defined by EpiMap³⁵ (STAR Methods; see data and code availability); we refined our interpretation of J-PEP clusters using single-cell assay for transposase-accessible chromatin using sequencing (scATAC-seq),⁴⁴ profiling 615,998 nuclei across 30 adult human tissues, ultimately resolving 111 fine-grained cell types. We restricted most analyses to 55 approximately independent focal diseases/traits (pairwise genetic correlation $[r_g] < 0.5$), prioritizing traits with greater statisti-

cal power (defined by the number of fine-mapped loci); focal traits yielding fewer than two clusters for any of the three partitioning methods were excluded from most comparisons, yielding 38 focal diseases/traits (STAR Methods). Auxiliary traits for each focal trait were selected from the full set of 165 diseases/traits based on two criteria: (1) a genetic correlation of < 0.5 with the focal trait (other thresholds were also tested); and (2) the existence of at least one shared causal SNP ($PIP > 0.5$) with the focal trait. On average, a focal disease/trait had 58 auxiliary traits. To improve cluster interpretability in selected real trait examples, we also conducted Gene Ontology (GO) enrichment⁴⁵ informed by SNP-to-gene linking scores from Gazal et al.⁴⁶ (STAR Methods).

To assess the performance of J-PEP and compare it to other clustering approaches, we developed PEPA, a new validation metric that evaluates how well a method predicts pleiotropic information (SNP-to-trait associations; V_{trait}) from epigenomic information (SNP-to-tissue associations; V_{tissue}) and vice versa. PEPA employs an off-chromosome prediction framework,^{47–50} where entries in V_{trait} and V_{tissue} , respectively, are predicted from the corresponding entries in V_{tissue} and V_{trait} , respectively, using cluster-specific profiles (H_{trait} and H_{tissue}) estimated using off-chromosome data. We note that the off-chromosome prediction framework ensures that this metric is robust to overfitting. PEPA is defined as the geometric mean of two constituent metrics: pleiotropic prediction accuracy (PPA), which measures how well SNP-to-trait associations (V_{trait}) can be predicted from SNP-to-tissue associations (V_{tissue}); and epigenomic prediction accuracy (EPA), which measures how well SNP-to-tissue associations (V_{tissue}) can be predicted from SNP-to-trait associations (V_{trait}). We use the geometric mean of PPA and EPA, rather than the arithmetic mean, because it prioritizes clusters with balanced support from both auxiliary traits and tissues while penalizing imbalance between modalities, reflecting the overarching goal of J-PEP. We also note that even though different inference approaches can differ in the number of discovered clusters, the PEPA metric ensures comparability across methods because it evaluates prediction accuracy in the observed data space, independent of the number of inferred clusters. Standard errors on each metric are computed using a genomic block jackknife, similarly to Finucane et al.³⁴ Further details are provided in STAR Methods.

Simulations

To evaluate the performance of different clustering methods, we performed simulations by simulating GWAS summary statistics at real SNPs, incorporating genetic architectures informed by predefined cluster profiles (STAR Methods). Our primary simulations included 32 tissues and 20 auxiliary traits, all of which were genetically uncorrelated with one another. In brief, we simulated data using $K_{causal} = 5$ predefined causal clusters and sampled 100–400 causal SNPs from chromosomes 1–2 of the 1000 Genomes Project⁵¹ (1KG) according to their aggregated cluster memberships, generating correlated SNP-to-trait (V_{trait}) and SNP-to-tissue (V_{tissue}) profiles. Using 1KG linkage disequilibrium (LD), we then simulated GWAS summary statistics and applied single-causal-variant fine mapping to match our real-data analyses (STAR Methods; see Figure S1 for an illustrative example).

Our primary simulations also included uncorrelated structure, i.e., structure within V_{tissue} and V_{trait} , respectively, that is not correlated to any structure in V_{trait} and V_{tissue} , respectively, by injecting random subsets of s non-causal tissues and s non-causal traits with permuted profiles from causal ones (STAR Methods); we also performed simulations with other parameter configurations (STAR Methods). We note that uncorrelated structure is likely to arise in real data, e.g., if the relevant tissues or cell types driving pleiotropic associations to auxiliary traits are missing or under-represented in epigenomic annotations, or analogously if a cell-type-specific disease process is not captured by the available auxiliary traits.

We first evaluated the accuracy of pleiotropic partitioning, epigenomic partitioning, and J-PEP in reconstructing the auxiliary trait-to-cluster profile matrix (H_{trait}) and tissue-to-cluster profile matrix (H_{tissue}) by computing the Frobenius norm error (defined as the root square of the sum of squared element-wise differences). These errors were averaged across 100 simulation replicates under four simulation settings, corresponding to $m = 100, 200, 300$, or 400 causal SNPs for the focal disease/trait. We determined that J-PEP attained the lowest average Frobenius norm error for both H_{tissue} and H_{trait} , with improvements becoming more pronounced as the number of causal SNPs increased (Figure 2A and Table S2). This indicates that J-PEP's enhanced predictive performance is not solely driven by the constraint of identifying shared structures between traits and tissues, which is enforced by the third term of the optimization function in Equation 1, but also benefits from the first two terms of Equation 1, which ensure that the inferred cluster memberships (W) remain closely aligned with the auxiliary trait and tissue data.

Next, we evaluated the accuracy of pleiotropic partitioning, epigenomic partitioning, and J-PEP using our new metric (PEPA) and its two constituent metrics (PPA and EPA). The average PEPA achieved by J-PEP increased with larger values of m , showing progressively greater and more significant improvements over both pleiotropic and epigenomic partitioning as m increased (Figure 2B and Table S2). At $m = 200$, J-PEP outperformed pleiotropic and epigenomic partitioning by 18% (paired t test $p = 1.2 \times 10^{-4}$) and 8% ($p = 0.04$), respectively; at $m = 400$, these gains rose to 19% ($p = 2.0 \times 10^{-10}$) and 38% ($p = 7.8 \times 10^{-20}$). Among the two constituent methods, pleiotropic partitioning performed better at predicting SNP-to-tissue associations (V_{tissue}) from SNP-to-trait associations (V_{trait}) (EPA), while epigenomic partitioning performed better at predicting SNP-to-trait associations (V_{trait}) from SNP-to-tissue associations (V_{tissue}) (PPA); this is consistent with pleiotropic partitioning and epigenomic partitioning yielding lower error when reconstructing the SNP-to-cluster membership matrix (W) from SNP-to-trait and SNP-to-tissue association data, respectively, as evaluated in the respective EPA and PPA metrics—compared to reconstructing W from SNP-to-tissue and SNP-to-trait data as evaluated in PPA and EPA, respectively.

To investigate the impact of LD and PIP thresholds on model performance, we conducted two additional analyses. First, we downsampled non-causal SNPs, thereby reducing the number of tagging SNPs contributing LD structure (Figure 2C and Table S2). As expected, performance of all three methods

improved as downsampling increased, and the data more closely reflected the true causal signals; importantly, J-PEP consistently maintained its relative advantage over pleiotropic and epigenomic partitioning across all settings. Second, we varied the PIP threshold used to define the fine-mapped SNPs included as input to each method, testing both stricter (≥ 0.05) and more lenient (down to 0.001) thresholds (Figure 2D and Table S2). This analysis revealed that applying stricter thresholds reduced prediction accuracy while including additional low-PIP SNPs improved prediction accuracy across all three methods (except that accuracy plateaued for epigenomic partitioning). This suggests that retaining a larger set of SNPs, weighted by their posterior probabilities, allows the models to capture additional true information and attain improved accuracy (consistent with expectations under well-calibrated fine mapping⁵²). J-PEP continued to outperform pleiotropic and epigenomic partitioning at more lenient PIP thresholds, but not at stricter PIP thresholds, perhaps because J-PEP benefits from a sufficiently large set of causal SNPs to reveal concordant structure between pleiotropic effects and epigenomic annotations; when too many informative SNPs are removed, the shared signal across the two data sources becomes too sparse to provide additional gains over single-modality approaches. Based on these results, we used a PIP threshold of 0.01 for all real trait analyses in this study (limiting computational cost).

Finally, we benchmarked FactorGO²⁶ and three versions of Flashier⁴² against J-PEP (Figure 2E and Table S2). With respect to reconstruction error (Frobenius norm), FactorGo performed worse than J-PEP, but Flashier (particularly the semi-non-negative version, Flashier_snn) performed almost as well as J-PEP in some cases, consistent with its design to optimize reconstruction of the single modality being partitioned. However, FactorGO and Flashier both performed far worse than J-PEP with respect to our cross-modal PEPA metric, consistent with the fact that FactorGO and Flashier do not consider tissue data. This suggests that unlike FactorGO and Flashier, J-PEP recovers clusters that align traits with tissues, making them more biologically meaningful by capturing disease processes in a tissue context.

We performed several secondary analyses. First, we formally evaluated receiver operating characteristic area under the curve (ROC AUC), sensitivity, and specificity using a co-clustering framework: for each pair of SNPs, we assessed the extent to which they co-cluster, taking into account their probabilistic cluster assignments (Figure S2). J-PEP attained a mean ROC AUC of 0.5237, closely followed by pleiotropic partitioning (0.5236, two-sided paired t test $p = 0.38$ for the difference with J-PEP across replicates) and well above epigenomic partitioning (0.5017, $p < 10^{-12}$). In terms of sensitivity and specificity, J-PEP consistently attained higher sensitivity than pleiotropic partitioning and higher specificity than epigenomic partitioning. Pleiotropic partitioning attained high specificity but low sensitivity, consistent with the fact that shared genetic effects across traits are sparse, hence stringent; epigenomic partitioning attained high sensitivity (at low PIP thresholds) but low specificity, as SNPs tend to overlap many regulatory annotations, capturing true co-clustering while also introducing false co-clustering.

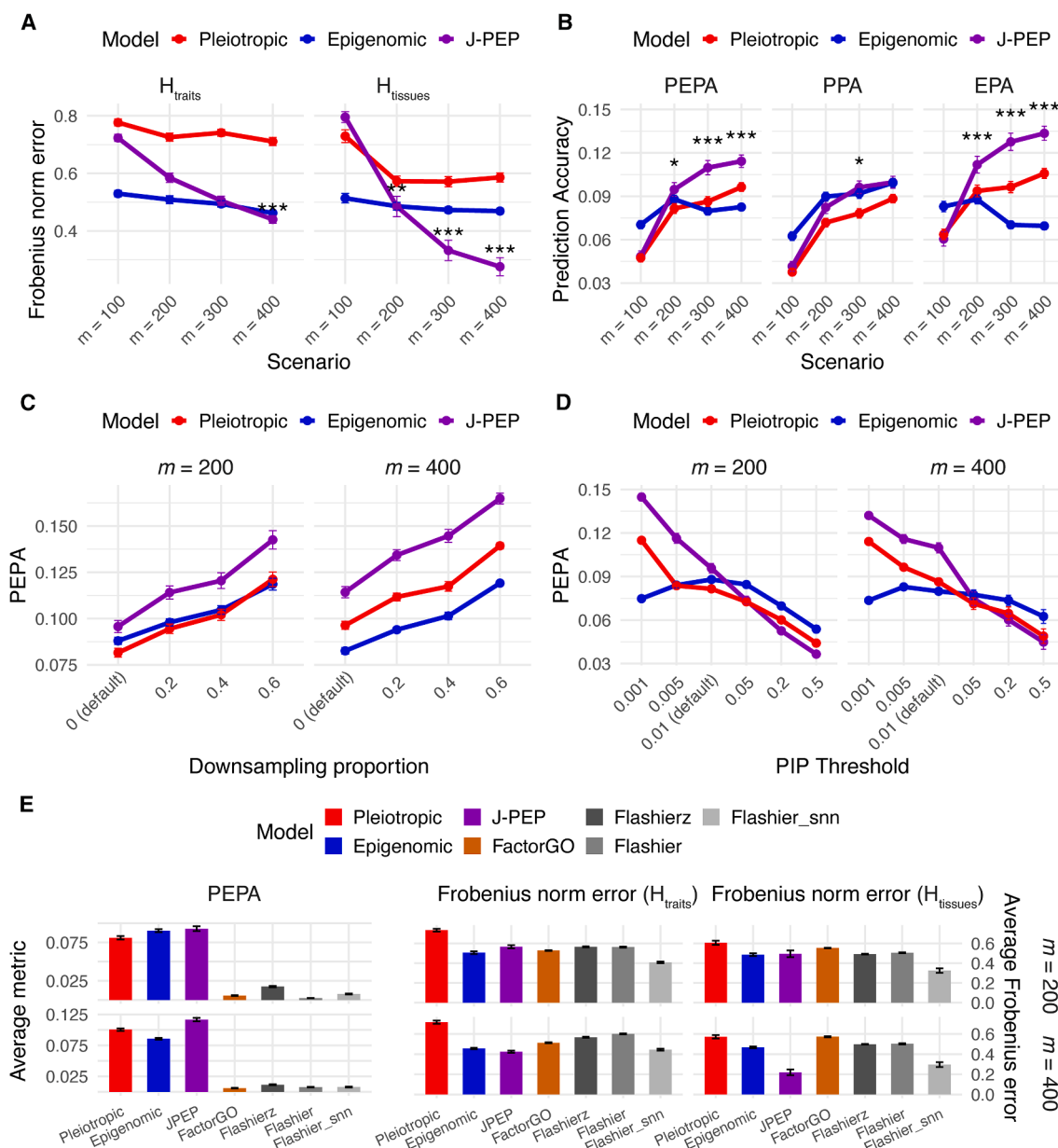


Figure 2. J-PEP outperforms pleiotropic partitioning and epigenomic partitioning in simulations

(A) Frobenius norm errors between the true vs. inferred auxiliary trait-to-cluster profile matrix (H_{trait}) and tissue-to-cluster profile matrix (H_{tissue}).

(B) PEPA, PPA, and EPA prediction accuracy metrics.

(C) PEPA metric across downsampling proportions of non-causal SNPs.

(D) PEPA metric across various PIP thresholds. In (A)–(D), results represent averages across 100 simulation replicates under four simulation settings with 100, 200, 300, or 400 causal SNPs for the focal disease/trait (denoted along the x axis as $m = 100, 200, 300$, or 400, respectively).

(E) Comparison between pleiotropic partitioning, epigenomic partitioning, J-PEP, FactorGO, and three versions of Flashier across three metrics: PEPA (left), Frobenius norm error for H_{trait} (center), and Frobenius norm error for H_{tissue} (right).

Error bars denote standard errors across replicates. p values were computed using paired two-sided Wilcoxon signed-rank tests. $*p < 0.05$, $**p < 0.01$, $***p < 0.001$. Numerical results are reported in Table S2.

Across six additional simulation analyses, we found that J-PEP's advantages were robust across generative settings (Figures S3 and S4). Performance gains increased when uncorrelated structure was present and persisted across variations in the number of causal clusters, maximum allowed K , and the pro-

portion of missing causal tissues; reducing K improved PEPA, and when half of the causal tissues were absent but auxiliary traits were available, J-PEP achieved even larger improvements over pleiotropic partitioning. We then considered a version of J-PEP in which the tissue sparsity constraint is directly included

in the objective function, encouraging but not enforcing sparsity in H_{tissue} (J-PEP_L1; [STAR Methods](#)), instead of the default “hard max sparsity constraint” of J-PEP. J-PEP_L1 produced very similar results to J-PEP when convergence was achieved ([Figure S5A](#)). However, J-PEP_L1 failed to converge or did not identify two distinct clusters in 37% of simulation replicates, whereas J-PEP converged in all replicates ([Figure S5B](#)). Finally, across simulations varying dimensionality (numbers of auxiliary traits and tissues), SNP-inclusion thresholds, and partitioning methods ([Figures S6, S7A, and S7B](#)), J-PEP showed stable PEPA behavior and remained computationally tractable, except at very low PIP thresholds (<0.01) where computational cost became prohibitive.

Application of J-PEP to 165 diseases and complex traits

We applied J-PEP, pleiotropic partitioning, and epigenomic partitioning to single causal variant fine-mapping results of GWAS summary statistics for 165 focal diseases/traits (average of 290,000 samples and 77 fine-mapped loci; [Table S1](#)), focusing our comparisons on 38 approximately independent diseases/traits for which all three methods identified at least two clusters. We performed single causal variant fine mapping instead of multiple causal variant fine mapping, as this is the recommended approach when in-sample LD is not available.⁵² Across these 38 traits, J-PEP attained significantly higher and lower) PEPA than pleiotropic partitioning (genomic block jackknife $p < 0.05$) for nine traits and two traits, respectively, and significantly higher and lower PEPA than epigenomic partitioning for six traits and no traits, respectively ([Figure 3A](#) and [Table S3](#)). J-PEP attained the most significant improvements for immune-related traits (e.g., blood cell counts; [Figure 3A](#)) and metabolic traits (e.g., total cholesterol; [Table S3](#)). In contrast, pleiotropic partitioning outperformed J-PEP for one brain-related trait (fractional anisotropy; [Table S3](#)) and one lung-related trait (FEV₁) ([Table S3](#)); in each case, J-PEP’s inferred cluster tissue profiles (H_{tissue}) lacked associations with brain and lung annotations, respectively, despite significant enrichment of these and other tissues in the corresponding SNP-to-tissue association matrix (V_{tissue}) matrices.

Averaging across the 38 focal diseases/traits, J-PEP attained 16% higher PEPA than pleiotropic partitioning (genomic jackknife on the difference $p = 2.4 \times 10^{-2}$; p from inverse-variance weighting [p_w] = 1.2×10^{-3} ; genomic block jackknife on the average across 38 traits; [STAR Methods](#)) and 22% higher PEPA than epigenomic partitioning ($p = 2.6 \times 10^{-4}$; $p_w = 3.2 \times 10^{-5}$) ([Figure 3B](#) and [Table S3](#)). These improvements are consistent with our primary simulations, suggesting that real disease/trait data harbor uncorrelated structure, i.e., structure within V_{tissue} and V_{trait} that is not correlated to any structure in V_{trait} and V_{tissue} , respectively. Notably, J-PEP’s performance advantage grew more pronounced in analyses of well-powered traits with more fine-mapped loci, also consistent with simulation results. For example, J-PEP outperformed pleiotropic partitioning by 27% ($p = 1.6 \times 10^{-3}$; $p_w = 2.0 \times 10^{-3}$) and epigenomic partitioning by 30% ($p = 1.4 \times 10^{-8}$; $p_w = 3.1 \times 10^{-8}$) across 16 focal traits with at least 100 fine-mapped loci ([Figure 3B](#) and [Table S3](#)). We generally observed stronger concordance between J-PEP and epigenomic partitioning than between J-PEP

and pleiotropic partitioning (average maximum SNP-level correlation of 0.86 between J-PEP and epigenomic partitioning clusters and 0.25 between J-PEP and pleiotropic partitioning clusters across all 38 focal diseases/traits; [Table S4](#); see below). Pleiotropic partitioning performed better at predicting SNP-to-tissue associations (V_{tissue}) from SNP-to-trait associations (V_{trait}) (EPA), while epigenomic partitioning performed better at predicting SNP-to-trait associations (V_{trait}) from SNP-to-tissue associations (V_{tissue}) (PPA), consistent with our simulations. As expected, a given constituent method tends to perform well when J-PEP significantly outperforms the other constituent method ([Figure S8](#)). These results demonstrate that J-PEP consistently improves tissue-trait prediction accuracy across a broad set of diseases/traits, particularly in analyses of well-powered traits.

The above analyses restricted auxiliary traits to those with $r_g < 0.5$ with the focal trait (average of 58 auxiliary traits per focal trait). To assess the impact of this choice, we repeated our analyses using a more stringent threshold ($r_g < 0.2$; average of 55.3 auxiliary traits) or a more lenient threshold ($r_g < 0.9$; average of 62.5 auxiliary traits) relative to the default threshold ($r_g < 0.5$; average of 61.3 auxiliary traits), restricting all comparisons to the 32 focal traits for which all tested thresholds yielded at least two clusters ([Figures 3C and 3D](#); [Table S3](#)). We determined that the more stringent threshold ($r_g < 0.2$) led to a reduction in accuracy as quantified by the PEPA metric (e.g., J-PEP PEPA = 0.040 vs. 0.050), with J-PEP continuing to outperform pleiotropic partitioning and epigenomic partitioning by a similar margin. The more lenient threshold ($r_g < 0.9$) slightly increased prediction accuracy (e.g., J-PEP PEPA = 0.056 vs. 0.050), with J-PEP continuing to outperform pleiotropic partitioning and epigenomic partitioning by a similar margin. However, six of the 38 focal traits that produced at least two clusters at $r_g < 0.5$ failed to do so at $r_g < 0.9$, all of which had gained new auxiliary traits. Conversely, no focal trait yielded at least two clusters at $r_g < 0.9$ and not at $r_g < 0.5$. This suggests that inclusion of highly correlated traits may destabilize the model. For this reason, we decided not to increase the default threshold to $r_g < 0.9$.

We performed several secondary analyses. First, we investigated whether the stronger concordance between J-PEP and epigenomic partitioning than between J-PEP and pleiotropic partitioning ([Table S4](#)) could arise from differences in the scaling of the two input matrices. In the default version of J-PEP, entries of the SNP-to-trait matrix (V_{trait}) are constructed from products of (focal trait PIP) $\times$ (auxiliary trait PIP), whereas entries of the SNP-to-tissue matrix (V_{tissue}) are individual probabilities that are not weighted by (focal trait PIP). We performed a new analysis of real traits in which we weighted V_{tissue} by (focal trait PIP), so that V_{trait} and V_{tissue} are on the same scale with respect to focal trait PIPs. In this analysis, we observed stronger concordance between J-PEP and pleiotropic partitioning than between J-PEP and epigenomic partitioning ([Figure S5C](#)). We elected to continue to not weight V_{tissue} by (focal trait PIP) in the default version of J-PEP because J-PEP is explicitly designed to integrate tissue and pleiotropic signals indirectly through joint factorization, and weighting V_{tissue} by (focal trait PIP) may suppress much of the cross-tissue structure captured by the epigenomic data, likely contributing to the frequent failures to converge or

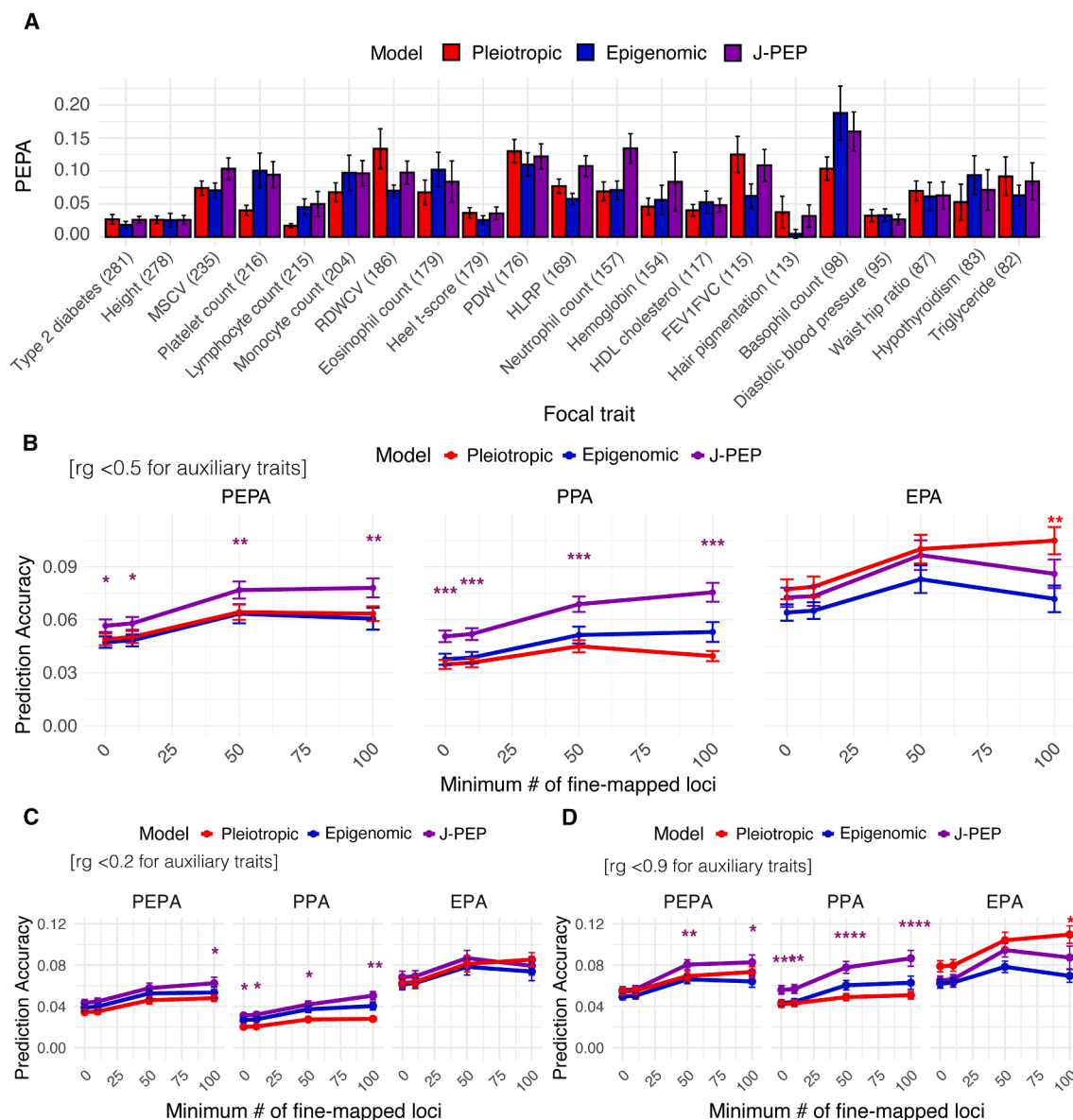


Figure 3. J-PEP outperforms pleiotropic partitioning and epigenomic partitioning in analyses of real diseases/traits

(A) PEPA prediction accuracy metric for 20 focal diseases/traits with the largest number of fine-mapped loci (denoted in parentheses). Error bars denote standard errors (via genomic block jackknife).

(B) PEPA, PPA, and EPA prediction accuracy metrics averaged across 38, 33, 24, or 16 focal traits (of a total of 38) with at least 0, 10, 50, or 100 fine-mapped loci, respectively.

(C and D) Analogous to (B), with the genetic correlation (r_g) threshold for auxiliary traits changed from $r_g < 0.5$ to $r_g < 0.2$ and $r_g < 0.9$, respectively. Error bars denote standard errors (via genomic block jackknife). p values were computed using two-sided genomic block jackknife tests on the difference between J-PEP and each other method. * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$ (vs. each other method). Numerical results for all 165 focal diseases/traits are reported in Table S3. MSCV, mean spheroid corpuscular volume; RDWCV, red cell distribution width—coefficient of variation; PDW, platelet distribution width; HLRP, high light reticulocyte proportion; FEV1/FVC, forced expiratory volume in 1 s/forced vital capacity.

recover clusters. We also investigated whether the stronger concordance between J-PEP and epigenomic partitioning than between J-PEP and pleiotropic partitioning could arise from the sparsity constraint on tissue-cluster membership but concluded that this was not the case (Figures S5C and S5D). Across an additional comprehensive set of secondary analyses

(Figures S7–S14; Tables S5, S6, and S7), J-PEP substantially outperformed FactorGO²⁶ and Flashier⁴² with respect to the cross-modal PEPA metric, consistent with simulations and with the fact that these methods do not use tissue data. J-PEP also behaved robustly across traits and parameter settings: cluster resolution increased with GWAS power; SNPs and loci were

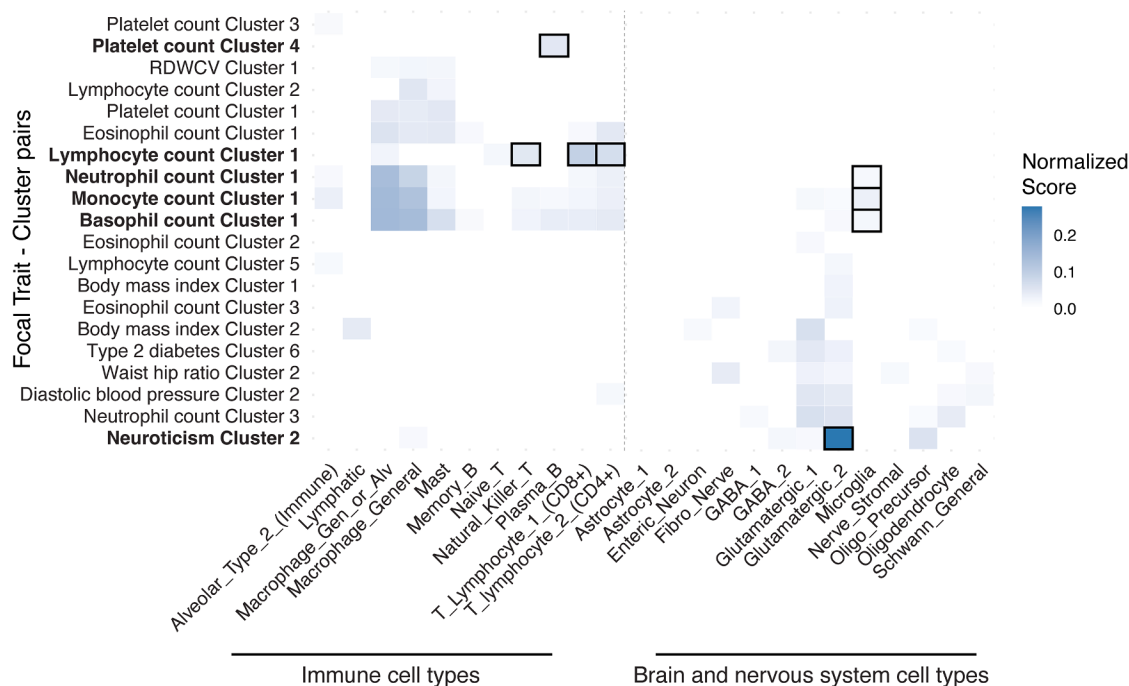


Figure 4. Single-cell data improve resolution and specificity of cell-type associations

We report normalized cell-type scores for each of 20 selected focal trait-cluster pairs (selected as described below; rows) across 24 single-cell-derived immune cell types (left) and brain/nervous system cell types (right). Normalized scores are defined as the proportion of each cluster assigned to a given cell type, normalized across all cell types for that cluster. We selected focal trait-cluster pairs with at least one normalized score $>1\%$ and sorted them by identifying the tissue category (immune or brain) with the highest aggregate score for each pair, then sorting first by tissue category and second by the magnitude of the maximum aggregate score within that category. Black borders denote focal trait/cluster/cell type triplets discussed in the main text, with corresponding focal trait-cluster pairs in bold font. Numerical results for 38 diseases/traits are reported in Table S6. RDWCV, red cell distribution width—coefficient of variation.

well differentiated; aggregated tissue profiles recapitulated expected biology; reducing the maximum number of clusters modestly improved PEPA without changing method rankings; tissue-to-cluster scores reflected annotation density; running times remained low; and all runs converged.

Single-cell accessibility profiles enhance cell-type resolution and specificity

We sought to assess whether single-cell chromatin accessibility profiles could help interpret the J-PEP clusters for 38 focal disease/traits identified by J-PEP using bulk epigenomic data. We analyzed a single-cell atlas of chromatin accessibility comprising 615,998 cells across 111 fine-grained adult human cell types.⁴⁴ To derive cell-type-to-cluster profile matrices ($H_{\text{cell-type}}$), we used a two-step procedure (J-PEP-CT): (1) apply J-PEP to bulk data in order to infer the shared SNP-to-cluster membership matrix W ; and (2) solve $V_{\text{cell-type}} = W \times H_{\text{cell-type}}$ using a non-negative least-squares approach with multiplicative update rules analogous to those used in bNMF, where $V_{\text{cell-type}}$ denotes the SNP-to-cell-type association matrix obtained using the EM algorithm on scATAC-seq data across the 111 cell types (analogous to V_{tissue} as defined above; Table S8 and STAR Methods). We compared $H_{\text{cell-type}}$ to H_{tissue} by collapsing both to 24 unified tissue categories (Table S9); we observed high concordance (mean $r = 0.55$; Figure S15), confirming high concordance between the two modalities. We did not include

direct application of J-PEP to single-cell data (i.e., applying J-PEP to V_{trait} and $V_{\text{cell-type}}$ instead of the two-step J-PEP-CT procedure) in our primary analyses; due to increased sparsity, cell-type complexity, and technical variability of single-cell assays, this led to lower prediction accuracy than application of J-PEP to bulk data (although J-PEP still substantially outperformed pleiotropic partitioning and epigenomic partitioning in direct application to single-cell data; Figure S13), such that we do not recommend direct application of J-PEP to single-cell data.

We evaluated whether scATAC-seq data could improve the resolution and specificity of cell-type associations. Restricting to the 11 of 24 unified tissue categories with at least four annotated cell types, we determined that focal trait-cluster pairs positively associated with these tissue categories were associated with only 21% of the corresponding cell types on average (Figures 4 and S15; Table S8), indicating a high degree of specificity. For example, in the immune category, cluster 1 of lymphocyte count was predominantly associated with T cell subtypes, whereas cluster 4 of platelet count was exclusively linked to plasma B cells. In the brain and nervous system category, cluster 2 of neuroticism was specifically associated with a glutamatergic neuron subtype—excitatory cells critical for brain signaling—whereas cluster 1 of each white blood cell trait (NC, monocyte count, and basophil count) showed distinct associations with microglia, the brain's resident immune cells. In the skin category,

which comprises four annotated cell types, cluster 1 of hair pigmentation was uniquely associated with melanocytes, the pigment-producing cells of the epidermis (Figure S15). These findings demonstrate that single-cell epigenomic data can resolve the cellular basis of trait-tissue associations, enhancing the interpretability of J-PEP-derived clusters.

J-PEP identifies non-canonical axes of T2D pathogenesis

T2D is a multifactorial disorder involving heterogeneous pathological mechanisms including pancreatic β cell dysfunction, impaired insulin production by excess hepatic fat accumulation, and obesity-related insulin resistance.^{23–25,53,54} Prior efforts to partition T2D loci—often via pleiotropic partitioning—have consistently identified clusters reflecting these distinct processes.^{23–25,27,30,31,55} However, these approaches have often lacked resolution in the specific cellular contexts in which they occur, a limitation that J-PEP is designed to address. To enable direct comparison with previous work, all results in this section are based on the 109 auxiliary traits from the most comprehensive pleiotropic partitioning study to date²⁵ (Table S10), which employed a method closely aligned with pleiotropic partitioning as defined in the current study.

We applied J-PEP and J-PEP-CT to T2D summary statistics from Suzuki et al.²⁴ (428,452 cases and 2,107,149 controls). J-PEP identified seven clusters, capturing both established and underexplored biological processes. These clusters are reported in Figure 5, sorted in decreasing order of their joint proportion of variance explained across both the trait and tissue matrices, computed as the squared magnitude of each cluster's reconstruction relative to the total squared magnitude of the observed data (see STAR Methods). With respect to the PEPA metric, J-PEP (PEPA = 0.03) did not significantly differ from epigenomic partitioning (PEPA = 0.04, genomic jackknife on the difference $p = 0.28$; five clusters) and was outperformed by pleiotropic partitioning (PEPA = 0.06, $p = 5.1 \times 10^{-4}$; five clusters) (Table S3), consistent with the trade-off between the biological information in additional clusters vs. optimizing PEPA (Figure S13). We detail J-PEP's results below, as its clusters integrate elements from both pleiotropic and epigenomic partitions, offering a broad perspective on disease processes.

For each J-PEP cluster, we computed SNP-level correlations between each cluster produced by a different method (defined as the correlation between the corresponding columns of the respective SNP-to-cluster membership matrices W) and assessed the maximum correlation. For clusters 2, 4, and 5, SNP-level correlations were strongest between J-PEP and pleiotropy-based methods (including Smith et al.²⁵ and Suzuki et al.²⁴), yet overall correlations remained modest (mean max $r = 0.34$; Figure 5); for clusters 1, 3, 6, and 7, correlations were strongest between J-PEP and epigenomic partitioning, with high correlations (mean max $r = 0.96$; Figure 5), underscoring the informativeness of tissue data.

Clusters 2–4 align with canonical T2D pathways^{23–25,27,30,31,55} (Figure 5). Cluster 2 is associated with insulin secretion traits and endocrine/pancreatic tissues and is enriched for GO terms related to insulin signaling and glucose regulation (Table S11 and STAR Methods). Notably, single-cell data refined this asso-

ciation to pancreatic islet β cells, supporting a β cell dysfunction mechanism contributing to insulin deficiency.^{23,24,56,57} Cluster 3 exhibits a consistent liver signature across bulk and single-cell profiles with the single-cell association refined to hepatocytes—the main liver cell type and the only one represented in the single-cell data—and GO enrichment for cholesterol efflux and triglyceride metabolism, indicative of hepatic lipid dysregulation contributing to insulin resistance.^{23–25,27} Cluster 4 is associated with liver enzyme and adiposity traits and maps to digestive tissues. Single-cell data refined this digestive association to multiple gastrointestinal cell types, and GO term enrichment for hydrogen peroxide catabolism points to hepatic oxidative stress linked to lipid overload and steatosis—a known mechanism of insulin resistance.⁷

Clusters 1 and 5–7 reflect non-canonical components of T2D biology (Figure 5). Cluster 1 is associated with red blood cell traits and epithelial tissues. Single-cell data refined this association to thyroid follicular cells, and GO enrichment revealed epithelial-to-mesenchymal transition and oxidative stress, suggesting a stress-related process potentially affecting insulin signaling.⁵⁸ Cluster 5 lacks a clear pathological interpretation but is enriched for chromatin organization and associated with embryonic stem cells. Refinement using single-cell data was not feasible, as the single-cell atlas that we employed lacks embryonic cell types. This cluster may reflect early developmental or transcriptional regulatory influences on T2D risk. Cluster 6 is associated with obesity-related traits and immune tissues. Single-cell data refined this association to macrophages, with GO enrichment for inflammatory signaling and thermogenesis, suggesting a role for immune-metabolic regulation of energy balance in T2D.^{59–61} Finally, cluster 7 is characterized by associations with liver enzymes and triglycerides, GO enrichment for lipid homeostasis pathways, and a stromal bulk tissue profile refined by single-cell data to a mix of adipocyte and neuronal cell types, suggesting a metabolic mechanism.

J-PEP identifies stromal and endocrine axes of HTN pathogenesis

HTN is a highly polygenic disease with contributions from multiple biological pathways.^{62,63} Prior clustering efforts have often focused on shared metabolic comorbidities such as obesity, lipid levels, or T2D, providing limited insight into tissue-specific regulatory contributions to blood pressure control.^{64,65}

We applied J-PEP and J-PEP-CT to HTN summary statistics from Bycroft et al.⁴³ (3,174 cases and 455,380 controls) using 68 auxiliary traits (Table S10). J-PEP identified two clusters described in detail in Figure 6A, sorted in decreasing order of their joint proportion of variance explained across both the trait and tissue matrices (see above). With respect to the PEPA metric, J-PEP (PEPA = 0.07) outperformed both pleiotropic partitioning (PEPA = 0.04, genomic jackknife on the difference $p = 0.07$; three clusters) and epigenomic partitioning (PEPA = 0.03, $p = 2.1 \times 10^{-3}$; five clusters) (Figure 3A and Table S3).

Cluster 1 is associated with bulk stromal tissue, refined by single-cell data to mesothelial cells, with GO enrichment for Wnt signaling and cytoskeletal remodeling (Table S11), suggesting mesothelial-driven extracellular matrix (ECM) remodeling and vascular stiffening, a known contributor to elevated vascular

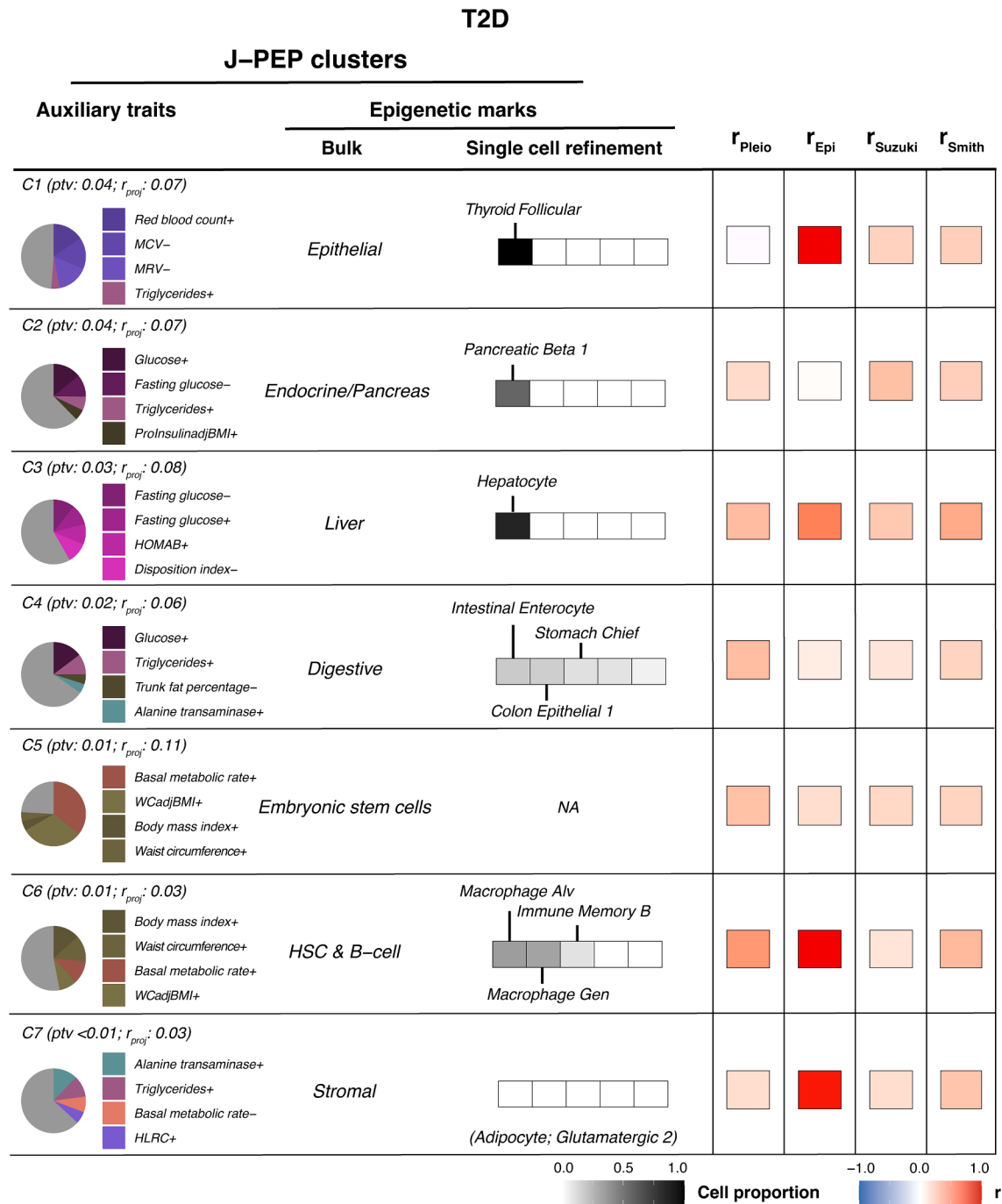


Figure 5. Comparison of T2D clusters identified by J-PEP and J-PEP-CT with other clustering methods

For each T2D cluster identified by J-PEP and J-PEP-CT, the pie chart (left) shows the pleiotropic auxiliary trait profile (top four traits), and the bulk and single-cell columns show the associated epigenomic tissue and cell-type profiles, respectively. The “+” or “-” symbol appended to auxiliary trait names denotes pleiotropic associations that are concordant or discordant, respectively, with the focal trait (STAR Methods). Single-cell refinement results show the top five contributing cell types within the most enriched bulk tissue category, shaded by relative cell proportion. NA indicates that the single-cell atlas that we employed lacks cell types from the implicated bulk tissue. Cell types that do not correspond to a refinement of the implicated bulk tissue are denoted in parentheses. Cluster labels indicate the projection correlation (r_{proj}), quantifying internal concordance between trait and tissue profiles and the proportion of total variance (ptv) they explain jointly across traits and tissues (STAR Methods). Clusters are displayed in decreasing order of ptv. The heatmap shows the maximum correlation between each J-PEP cluster and clusters obtained from pleiotropic partitioning, epigenomic partitioning, or two previous T2D clustering studies (Smith et al.²⁵ and Suzuki et al.²⁴). See data and code availability for numerical results. MCV, mean corpuscular volume; MRV, mean reticulocyte volume; ProInsulinadjBMI, proinsulin adjusted for body mass index; HOMAB, homeostasis model assessment of β cell function; WCadjBMI, waist circumference adjusted for body mass index; HLRC, high light reticulocyte count.

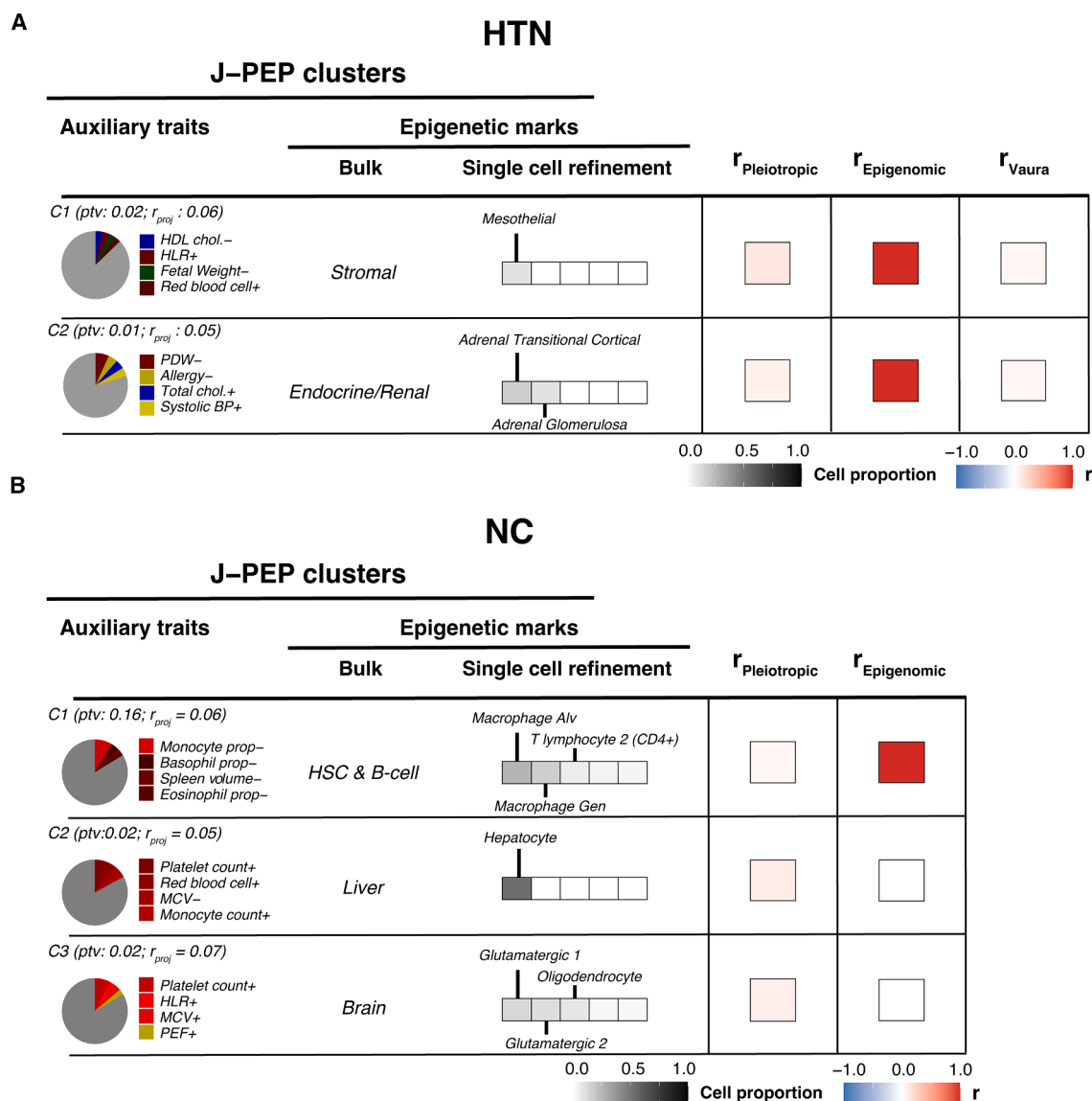


Figure 6. Comparison of HTN and NC clusters identified by J-PEP and J-PEP-CT with other clustering methods

For each hypertension (HTN) cluster (A) and neutrophil count (NC) cluster (B) identified by J-PEP and J-PEP-CT, the pie chart (left) shows the pleiotropic auxiliary trait profile (top four traits), and the bulk and single-cell columns show the associated epigenomic tissue and cell-type profiles, respectively. Single-cell refinement results show the top five contributing cell types within the most enriched bulk tissue category, shaded by relative cell proportion. The “+” or “-” symbol appended to auxiliary trait names denotes pleiotropic associations that are concordant or discordant, respectively, with the focal trait (STAR Methods). Cluster labels indicate the projection correlation (r_{proj}), quantifying internal consistency between tissue and trait profiles and the proportion of total variance (ptv) they explain jointly across traits and tissues (STAR Methods). Clusters are displayed in decreasing order of ptv. The heatmap shows the maximum correlation between each J-PEP cluster and clusters obtained from pleiotropic partitioning or epigenomic partitioning (or Vaura et al.⁶⁴ for HTN). See data and code availability for numerical results. HLR, high light reticulocyte count; PDW, platelet distribution width; BP, blood pressure; MCV, mean corpuscular volume; PEF, peak expiratory flow.

resistance in HTN.⁶⁶ While ECM remodeling is a recognized mechanism in HTN pathophysiology,^{67,68} the specific role of mesothelial cells in this process remains speculative.⁶⁹ Cluster 2 is associated with endocrine bulk tissue, refined by single-cell data to adrenal cortex cells, with GO terms highlighting neutrophil-mediated immunity, indicating a potential feedback between aldosterone signaling and local immune activity.⁷⁰

Although aldosterone is known to modulate various immune pathways involved in inflammation, and immune-targeted therapies have shown benefit in managing HTN, the precise mechanisms underlying its crosstalk with the immune system remain unknown.⁷¹ For both cluster 1 and cluster 2, SNP-level correlations were strongest between J-PEP and epigenomic partitioning (Figure 6A), highlighting the benefit of incorporating

epigenomic data. Compared to the four metabolically themed clusters identified by Vaura et al.,⁶⁴ who applied a method similar to pleiotropic partitioning to FinnGen data⁷²—comprising obesity, two lipid-related, and one stature-associated clusters—J-PEP produced a distinct set of clusters with minimal overlap, likely reflecting its integrative design that incorporates epigenomic information.

J-PEP identifies immune, hepatic, and neuroinflammatory axes of NC

NC is a key hematological trait central to immune and inflammatory processes,^{58,59} yet to our knowledge it has not previously been analyzed using clustering methods that integrate either pleiotropic or epigenomic data. We applied J-PEP and J-PEP-CT to NC summary statistics from Vuckovic et al.⁷³ (563,085 samples) using 83 auxiliary traits (Table S10). J-PEP identified three clusters reflecting immune, hepatic, and neuroinflammatory axes (Figure 6B, sorted in decreasing order of their joint proportion of variance explained across both the trait and tissue matrices; see above). With respect to the PEPA metric, J-PEP (PEPA = 0.14) significantly outperformed pleiotropic partitioning (PEPA = 0.05, genomic jackknife on the difference $p = 3.8 \times 10^{-9}$; five clusters) and epigenomic partitioning (PEPA = 0.06, $p = 6.4 \times 10^{-6}$; three clusters) (Figure 3A and Table S3).

Cluster 1 is associated with bulk hematopoietic stem and progenitor cells, refined by single-cell data to alveolar and general macrophages, with GO enrichment for chemotaxis, immune activation, and antimicrobial responses (Table S11), highlighting canonical immune regulation of neutrophils.⁷⁴ Cluster 2 exhibits a consistent liver signature across bulk and single-cell profiles (refined to hepatocytes in single-cell data) and is linked to platelet-related traits, suggesting a speculative hepatic-inflammatory connection to neutrophil regulation, consistent with neutrophil activation in metabolic disorders.⁷⁵ Cluster 3 exhibits a bulk brain profile refined by single-cell data to glutamatergic neurons and oligodendrocytes, along with associations to hematocrit and platelet count, pointing to emerging but unconfirmed links between neural regulation and immune tone.^{76,77} SNP-level correlations indicate that cluster 1 aligns closely with epigenomic partitioning, whereas clusters 2 and 3 are only weakly recovered by pleiotropic partitioning—despite their most distinctive features (compared to cluster 1 and to each other) being reflected in their tissue and cell-type profiles, which notably do not implicate immune-related components.

DISCUSSION

Summary of findings

We have developed J-PEP, a method that integrates pleiotropic and epigenomic information to partition disease loci into biologically interpretable clusters. By leveraging shared structure across auxiliary traits and tissues, J-PEP improves the accuracy of predicting SNP-to-trait and SNP-to-tissue associations, which we quantify using our new PEPA metric. These gains are most pronounced for focal diseases/traits whose underlying biology is well captured by available epigenomic datasets, such as immune- and liver-related traits. In our applications to GWAS summary statistics for a broad set of diseases and traits,

J-PEP recapitulates canonical disease pathways, e.g., endocrine/pancreatic dysfunction and hepatic-driven insulin resistance for T2D,^{23–25} but also reveals underexplored biological processes, e.g., immune-metabolic and developmental components for T2D. Notably, incorporating single-cell epigenomic profiles refines tissue associations to cell-type resolution, e.g., resolving an endocrine/pancreatic cluster for T2D to pancreatic β cells, and resolving an endocrine/renal cluster for HTN to specific adrenal cortex cell types.

Alternative approaches

J-PEP shares conceptual goals with several recent approaches aimed at disentangling shared genetic architecture across complex traits. DeGAs¹⁷ applies truncated singular value decomposition (TSVD) to GWAS summary statistics to extract latent components of genetic variation. While scalable and effective for broad surveys, it lacks structured priors to promote sparsity, usually requiring post hoc enrichment analyses for interpretation. FactorGo²⁶ addresses this by modeling uncertainty in GWAS effect estimates within a variational Bayesian framework and promoting sparsity through automatic relevance determination. However, it operates on LD-pruned data, which limits compatibility with regulatory annotations that do not depend on LD structure, such as epigenomic marks. The V2G2P framework³⁷ takes a more mechanistic route, linking GWAS variants to genes through enhancer maps and to regulatory programs derived from single-cell perturbation screens. This strategy enables interpretable, cell-type-specific disease insights, but it requires substantial experimental resources and is currently limited to well-characterized cellular systems, limiting its scalability.

Downstream implications

Our results have several downstream implications. From a translational perspective, J-PEP provides a scalable means to distill complex polygenic signals into a small number of interpretable axes of biological activity. This may aid in prioritizing genes, tissues, and pathways for functional follow-up and therapeutic development.^{4,78,79} Also, by uncovering shared axes across diseases/traits, J-PEP may help identify opportunities for drug repurposing—for example, targeting pathways implicated in both T2D and cardiovascular disease.^{80,81} In the context of precision medicine, clusters inferred by J-PEP could aid the construction of partitioned polygenic risk scores, previously shown to support disease subtyping.^{24,25} Finally, our PEPA metric provides a principled framework for benchmarking future methods that integrate pleiotropic and epigenomic data.

Limitations of the study

We note six limitations of J-PEP and this study. First, although J-PEP outperforms single-modality methods, overall PEPA scores remain modest, reflecting inherent biological complexity, measurement noise, and incomplete regulatory and phenotypic coverage—gaps that are likely to narrow as new genetic and epigenomic resources continue to emerge.^{4,82–85} Second, while J-PEP captures trait-tissue correlations, it does not explicitly infer causality between them; nonetheless, clusters derived from joint partitioning are less prone to spurious associations compared to single-modality approaches, as they better reflect

disease-relevant biological context. Third, in bulk epigenomic data, the results of J-PEP (and other methods) are biased toward tissues with denser annotations; in single-cell epigenomic data, direct application of J-PEP (instead of J-PEP-CT) is limited by sparsity and technical variability, and we do not recommend this approach. Ongoing expansion of tissue-specific and single-cell epigenomic datasets will help mitigate both issues.^{4,82–85} Fourth, the algorithm can be computationally demanding for large-scale datasets, particularly when incorporating high-dimensional modalities like single-cell multi-omics; however, the method is computationally tractable (Figure S14), and efficiency can be further enhanced by restricting inputs to critical tissues or a curated set of auxiliary traits. Fifth, J-PEP employs a non-standard procedure to enforce sparsity in tissue-to-cluster scores at each iteration. However, our comparative analyses demonstrate that this sparsity constraint is essential for ensuring convergence, computational efficiency, and biologically meaningful clusters, and it outperforms alternative strategies (Figures S5A and S5B), such that we recommend it as our default approach. Finally, we did not evaluate J-PEP under multiple-causal variant fine mapping with in-sample LD; because such methods can be miscalibrated with out-of-sample LD,⁵² and our empirical analyses relied primarily on out-of-sample LD, we used single-causal variant fine mapping for both simulations and real data to ensure calibration and consistency. Despite these limitations, J-PEP provides a generalizable and interpretable approach for partitioning GWAS loci into biologically distinct clusters.

RESOURCE AVAILABILITY

Lead contact

Requests for further information and resources should be directed to and will be fulfilled by the lead contact, Gaspard Kerner (gkerner@hsph.harvard.edu).

Materials availability

This study did not generate new unique reagents.

Data and code availability

All data used in this study and the corresponding access information are listed in the [key resources table](#). This includes the `jpepR` package, developed as part of this work, which provides an environment to run J-PEP, as well as GWAS summary statistics for 165 traits and the corresponding J-PEP results. Links to the public software and algorithms used in the present study are listed in the [key resources table](#).

ACKNOWLEDGMENTS

We thank Kushal Dey and all members of Alkes Price's lab for helpful discussions and data sharing. This research was funded by NIH grants R01 HG006399, U01 HG012009, R01 MH101244, R01 HG013083, R37 MH107649, and T32 HG002295.

AUTHOR CONTRIBUTIONS

G.K. and A.L.P. conceived and designed the study. G.K. was the lead analyst, with important contributions from A.L.P., B.S., and N.K. N.K. and B.S. conceived and designed the SNP-to-tissue scoring algorithm using the EM algorithm. G.K. wrote the manuscript with substantial contributions from A.L.P.

DECLARATION OF INTERESTS

The authors declare no competing interests.

STAR★METHODS

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

- [KEY RESOURCES TABLE](#)
- [METHOD DETAILS](#)
 - J-PEP method: Input and output
 - J-PEP method: Objective function
 - J-PEP method: Optimization
 - Pleiotropic partitioning method
 - Epigenomic partitioning method
 - Prediction accuracy metrics
 - Methodological considerations and limitations of J-PEP
 - Simulation framework
 - Additional simulation analyses
 - FactorGO and Flashier implementations
 - GWAS summary statistics
 - EpiMap dataset
 - Single-cell dataset
 - J-PEP-CT: Inference of cell-type-derived cluster profiles
 - Gene ontology (GO) enrichment analysis
- [QUANTIFICATION AND STATISTICAL ANALYSIS](#)

SUPPLEMENTAL INFORMATION

Supplemental information can be found online at <https://doi.org/10.1016/j.xgen.2025.101138>.

Received: June 4, 2025

Revised: October 9, 2025

Accepted: December 30, 2025

Published: January 26, 2026

REFERENCES

1. Tam, V., Patel, N., Turcotte, M., Bossé, Y., Paré, G., and Meyre, D. (2019). Benefits and limitations of genome-wide association studies. *Nat. Rev. Genet.* 20, 467–484. <https://doi.org/10.1038/s41576-019-0127-1>.
2. Shendure, J., Findlay, G.M., and Snyder, M.W. (2019). Genomic Medicine—Progress, Pitfalls, and Promise. *Cell* 177, 45–57. <https://doi.org/10.1016/j.cell.2019.02.003>.
3. Claussnitzer, M., Cho, J.H., Collins, R., Cox, N.J., Dermitzakis, E.T., Hurles, M.E., Kathiresan, S., Kenny, E.E., Lindgren, C.M., MacArthur, D.G., et al. (2020). A brief history of human disease genetics. *Nature* 577, 179–189. <https://doi.org/10.1038/s41586-019-1879-7>.
4. Abdellaoui, A., Yengo, L., Verweij, K.J.H., and Visscher, P.M. (2023). 15 years of GWAS discovery: Realizing the promise. *Am. J. Hum. Genet.* 110, 179–194. <https://doi.org/10.1016/j.ajhg.2022.12.011>.
5. Sollis, E., Mosaku, A., Abid, A., Buniello, A., Cerezo, M., Gil, L., Groza, T., Güneş, O., Hall, P., Hayhurst, J., et al. (2023). The NHGRI-EBI GWAS Catalog: knowledgebase and deposition resource. *Nucleic Acids Res.* 51, D977–D985. <https://doi.org/10.1093/nar/gkac1010>.
6. Southam, L., and Zeggini, E. (2025). Twenty years of genome-wide association studies. *Nature* 641, 47–49. <https://doi.org/10.1038/d41586-025-01128-6>.
7. Musunuru, K., Strong, A., Frank-Kamenetsky, M., Lee, N.E., Ahfeldt, T., Sachs, K.V., Li, X., Li, H., Kuperwasser, N., Ruda, V.M., et al. (2010). From noncoding variant to phenotype via SORT1 at the 1p13 cholesterol locus. *Nature* 466, 714–719. <https://doi.org/10.1038/nature09266>.

8. Bauer, D.E., Kamran, S.C., Lessard, S., Xu, J., Fujiwara, Y., Lin, C., Shao, Z., Canver, M.C., Smith, E.C., Pinello, L., et al. (2013). An erythroid enhancer of BCL11A subject to genetic variation determines fetal hemoglobin level. *Science* 342, 253–257. <https://doi.org/10.1126/science.1242088>.
9. Claussnitzer, M., Dankel, S.N., Kim, K.-H., Quon, G., Meuleman, W., Haugen, C., Glunk, V., Sousa, I.S., Beaudry, J.L., Puviindran, V., et al. (2015). FTO Obesity Variant Circuitry and Adipocyte Browning in Humans. *N. Engl. J. Med.* 373, 895–907. <https://doi.org/10.1056/NEJMoa1502214>.
10. Downes, D.J., Cross, A.R., Hua, P., Roberts, N., Schwessinger, R., Cutler, A.J., Munis, A.M., Brown, J., Mielczarek, O., de Andrea, C.E., et al. (2021). Identification of LZTFL1 as a candidate effector gene at a COVID-19 risk locus. *Nat. Genet.* 53, 1606–1615. <https://doi.org/10.1038/s41588-021-00955-3>.
11. Morris, J.A., Caragine, C., Daniloski, Z., Domingo, J., Barry, T., Lu, L., Davis, K., Ziosi, M., Glinos, D.A., Hao, S., et al. (2023). Discovery of target genes and pathways at GWAS loci by pooled single-cell CRISPR screens. *Science* 380, eadh7699. <https://doi.org/10.1126/science.adh7699>.
12. Derks, E.M., Thorp, J.G., and Gerring, Z.F. (2022). Ten challenges for clinical translation in psychiatric genetics. *Nat. Genet.* 54, 1457–1465. <https://doi.org/10.1038/s41588-022-01174-0>.
13. Cotsapas, C., Voight, B.F., Rossin, E., Lage, K., Neale, B.M., Wallace, C., Abecasis, G.R., Barrett, J.C., Behrens, T., Cho, J., et al. (2011). Pervasive sharing of genetic effects in autoimmune disease. *PLoS Genet.* 7, e1002254. <https://doi.org/10.1371/journal.pgen.1002254>.
14. Solovieff, N., Cotsapas, C., Lee, P.H., Purcell, S.M., and Smoller, J.W. (2013). Pleiotropy in complex traits: challenges and strategies. *Nat. Rev. Genet.* 14, 483–495. <https://doi.org/10.1038/nrg3461>.
15. Bulik-Sullivan, B., Finucane, H.K., Anttila, V., Gusev, A., Day, F.R., Loh, P.-R., and ReproGen Consortium; Psychiatric Genomics Consortium; Genetic Consortium for Anorexia Nervosa of the Wellcome Trust Case Control Consortium 3; and Duncan, L. (2015). An atlas of genetic correlations across human diseases and traits. *Nat. Genet.* 47, 1236–1241. <https://doi.org/10.1038/ng.3406>.
16. Pickrell, J.K., Berisa, T., Liu, J.Z., Séguérel, L., Tung, J.Y., and Hinds, D.A. (2016). Detection and interpretation of shared genetic influences on 42 human traits. *Nat. Genet.* 48, 709–717. <https://doi.org/10.1038/ng.3570>.
17. Tanigawa, Y., Li, J., Justesen, J.M., Horn, H., Aguirre, M., DeBoever, C., Chang, C., Narasimhan, B., Lage, K., Hastie, T., et al. (2019). Components of genetic associations across 2,138 phenotypes in the UK Biobank high-light adipocyte biology. *Nat. Commun.* 10, 4064. <https://doi.org/10.1038/s41467-019-11953-9>.
18. Grotzinger, A.D., Rhemtulla, M., de Vlaming, R., Ritchie, S.J., Mallard, T.T., Hill, W.D., Ip, H.F., Marioni, R.E., McIntosh, A.M., Deary, I.J., et al. (2019). Genomic structural equation modelling provides insights into the multivariate genetic architecture of complex traits. *Nat. Hum. Behav.* 3, 513–525. <https://doi.org/10.1038/s41562-019-0566-x>.
19. Ballard, J.L., and O'Connor, L.J. (2022). Shared components of heritability across genetically correlated traits. *Am. J. Hum. Genet.* 109, 989–1006. <https://doi.org/10.1016/j.ajhg.2022.04.003>.
20. Grotzinger, A.D., Mallard, T.T., Akingbuwa, W.A., Ip, H.F., Adams, M.J., Lewis, C.M., McIntosh, A.M., Grove, J., Dalsgaard, S., Lesch, K.-P., et al. (2022). Genetic architecture of 11 major psychiatric disorders at biobehavioral, functional genomic and molecular genetic levels of analysis. *Nat. Genet.* 54, 548–559. <https://doi.org/10.1038/s41588-022-01057-4>.
21. Romero, C., Werme, J., Jansen, P.R., Gelernter, J., Stein, M.B., Levey, D., Polimanti, R., de Leeuw, C., Posthuma, D., Nagel, M., and van der Sluis, S. (2022). Exploring the genetic overlap between twelve psychiatric disorders. *Nat. Genet.* 54, 1795–1802. <https://doi.org/10.1038/s41588-022-01245-2>.
22. Mackay, T.F.C., and Anholt, R.R.H. (2024). Pleiotropy, epistasis and the genetic architecture of quantitative traits. *Nat. Rev. Genet.* 25, 639–657. <https://doi.org/10.1038/s41576-024-00711-3>.
23. Udler, M.S., Kim, J., von Grotthuss, M., Bonàs-Guarch, S., Cole, J.B., Chiou, J., Anderson, C.D., on behalf of METASTROKE and the ISGC; Boehnke, M., Laakso, M., and Atzmon, G. (2018). Type 2 diabetes genetic loci informed by multi-trait associations point to disease mechanisms and subtypes: A soft clustering analysis. *PLoS Med.* 15, e1002654. <https://doi.org/10.1371/journal.pmed.1002654>.
24. Suzuki, K., Hatzikotoulas, K., Southam, L., Taylor, H.J., Yin, X., Lorenz, K.M., Mandla, R., Huerta-Chagoya, A., Melloni, G.E.M., Kanoni, S., et al. (2024). Genetic drivers of heterogeneity in type 2 diabetes pathophysiology. *Nature* 627, 347–357. <https://doi.org/10.1038/s41586-024-07019-6>.
25. Smith, K., Deutsch, A.J., McGrail, C., Kim, H., Hsu, S., Huerta-Chagoya, A., Mandla, R., Schroeder, P.H., Westerman, K.E., Szczerbinski, L., et al. (2024). Multi-ancestry polygenic mechanisms of type 2 diabetes. *Nat. Med.* 30, 1065–1074. <https://doi.org/10.1038/s41591-024-02865-3>.
26. Zhang, Z., Jung, J., Kim, A., Suboc, N., Gazal, S., and Mancuso, N. (2023). A scalable approach to characterize pleiotropy across thousands of human diseases and complex traits using GWAS summary statistics. *Am. J. Hum. Genet.* 110, 1863–1874. <https://doi.org/10.1016/j.ajhg.2023.09.015>.
27. Kim, H., Westerman, K.E., Smith, K., Chiou, J., Cole, J.B., Majarian, T., von Grotthuss, M., Kwak, S.H., Kim, J., Mercader, J.M., et al. (2023). High-throughput genetic clustering of type 2 diabetes loci reveals heterogeneous mechanistic pathways of metabolic disease. *Diabetologia* 66, 495–507. <https://doi.org/10.1007/s00125-022-05848-6>.
28. Reeve, M., Kanai, M., Graham, D., Karjalainen, J., Luo, S., Kolosov, N., Adams, C., Ritari, J., Karczewski, K., Kiiskinen, T., et al. (2024). Autoimmune hypothyroidism GWAS reveals independent autoimmune and thyroid-specific contributions and an inverse relation with cancer risk. Preprint at Res. Sq. <https://doi.org/10.21203/rs.3.rs-4626646/v1>.
29. Qi, G., Chhetri, S.B., Ray, D., Dutta, D., Battle, A., Bhattacharjee, S., and Chatterjee, N. (2024). Genome-wide large-scale multi-trait analysis characterizes global patterns of pleiotropy and unique trait-specific variants. *Nat. Commun.* 15, 6985. <https://doi.org/10.1038/s41467-024-51075-5>.
30. Ghatan, S., van Rooij, J., van Hoek, M., Boer, C.G., Felix, J.F., Kavousi, M., Jaddoe, V.W., Sijbrands, E.J.G., Medina-Gomez, C., Rivadeneira, F., and Oei, L. (2024). Defining type 2 diabetes polygenic risk scores through colocalization and network-based clustering of metabolic trait genetic associations. *Genome Med.* 16, 10. <https://doi.org/10.1186/s13073-023-01255-7>.
31. Li, Y., Chen, G.-C., Moon, J.-Y., Arthur, R., Sotres-Alvarez, D., Daviglus, M.L., Pirzada, A., Mattei, J., Ferreira, K.M., Rotter, J.I., et al. (2024). Genetic Subtypes of Prediabetes, Healthy Lifestyle, and Risk of Type 2 Diabetes. *Diabetes* 73, 1178–1187. <https://doi.org/10.2337/db23-0699>.
32. Roadmap Epigenomics Consortium; Kundaje, A., Meuleman, W., Ernst, J., Bilenky, M., Yen, A., Heravi-Moussavi, A., Kheradpour, P., Zhang, Z., and Wang, J. (2015). Integrative analysis of 111 reference human epigenomes. *Nature* 518, 317–330. <https://doi.org/10.1038/nature14248>.
33. Farh, K.K.-H., Marson, A., Zhu, J., Kleinewietfeld, M., Housley, W.J., Beik, S., Shores, N., Whitton, H., Ryan, R.J.H., Shishkin, A.A., et al. (2015). Genetic and epigenetic fine mapping of causal autoimmune disease variants. *Nature* 518, 337–343. <https://doi.org/10.1038/nature13835>.
34. Finucane, H.K., Bulik-Sullivan, B., Gusev, A., Trynka, G., Reshef, Y., Loh, P.-R., Anttila, V., Xu, H., Zang, C., Farh, K., et al. (2015). Partitioning heritability by functional annotation using genome-wide association summary statistics. *Nat. Genet.* 47, 1228–1235. <https://doi.org/10.1038/ng.3404>.
35. Boix, C.A., James, B.T., Park, Y.P., Meuleman, W., and Kellis, M. (2021). Regulatory genomic circuitry of human disease loci by integrative epigenomics. *Nature* 590, 300–307. <https://doi.org/10.1038/s41586-020-03145-z>.
36. Gaulton, K.J., Preissl, S., and Ren, B. (2023). Interpreting non-coding disease-associated human variants using single-cell epigenomics. *Nat. Rev. Genet.* 24, 516–534. <https://doi.org/10.1038/s41576-023-00598-6>.

37. Schnitzler, G.R., Kang, H., Fang, S., Angom, R.S., Lee-Kim, V.S., Ma, X.R., Zhou, R., Zeng, T., Guo, K., Taylor, M.S., et al. (2024). Convergence of coronary artery disease genes onto endothelial cell programs. *Nature* 626, 799–807. <https://doi.org/10.1038/s41586-024-07022-x>.
38. Wellcome Trust Case Control Consortium; Maller, J.B., McVean, G., Byrnes, J., Vukcevic, D., Palin, K., Su, Z., Howson, J.M.M., Auton, A., and Myers, S. (2012). Bayesian refinement of association signals for 14 loci in 3 common diseases. *Nat. Genet.* 44, 1294–1301. <https://doi.org/10.1038/ng.2435>.
39. Schaid, D.J., Chen, W., and Larson, N.B. (2018). From genome-wide associations to candidate causal variants by statistical fine-mapping. *Nat. Rev. Genet.* 19, 491–504. <https://doi.org/10.1038/s41576-018-0016-z>.
40. Schmidt, M.N., Winther, O., and Hansen, L.K. (2009). Bayesian Non-negative Matrix Factorization. In *Independent Component Analysis and Signal Separation*, T. Adali, C. Jutten, J.M.T. Romano, and A.K. Barros, eds. (Springer), pp. 540–547. https://doi.org/10.1007/978-3-642-00599-2_68.
41. Tan, V.Y.F., and Févotte, C. (2013). Automatic Relevance Determination in Nonnegative Matrix Factorization with the/spl beta/-Divergence. *IEEE Trans. Pattern Anal. Mach. Intell.* 35, 1592–1605. <https://doi.org/10.1109/TPAMI.2012.240>.
42. Willwerscheid, J., Carbonetto, P., and Stephens, M. (2024). ebnm: An R Package for Solving the Empirical Bayes Normal Means Problem Using a Variety of Prior Families. Preprint at arXiv. <https://doi.org/10.48550/arXiv.2110.00152>.
43. 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. <https://doi.org/10.1038/s41586-018-0579-z>.
44. Zhang, K., Hocker, J.D., Miller, M., Hou, X., Chiou, J., Poirion, O.B., Qiu, Y., Li, Y.E., Gaulton, K.J., Wang, A., et al. (2021). A single-cell atlas of chromatin accessibility in the human genome. *Cell* 184, 5985–6001.e19. <https://doi.org/10.1016/j.cell.2021.10.024>.
45. Young, M.D., Wakefield, M.J., Smyth, G.K., and Oshlack, A. (2010). Gene ontology analysis for RNA-seq: accounting for selection bias. *Genome Biol.* 11, R14. <https://doi.org/10.1186/gb-2010-11-2-r14>.
46. Gazal, S., Weissbrod, O., Hormozdiani, F., Dey, K.K., Nasser, J., Jagadeesh, K.A., Weiner, D.J., Shi, H., Fulco, C.P., O'Connor, L., et al. (2022). Combining SNP-to-gene linking strategies to identify disease genes and assess disease omnigenicity. *Nat. Genet.* 54, 827–836. <https://doi.org/10.1038/s41588-022-01087-y>.
47. Chen, K.M., Wong, A.K., Troyanskaya, O.G., and Zhou, J. (2022). A sequence-based global map of regulatory activity for deciphering human genetics. *Nat. Genet.* 54, 940–949. <https://doi.org/10.1038/s41588-022-01102-2>.
48. Gao, V.R., Yang, R., Das, A., Luo, R., Luo, H., McNally, D.R., Karagiannis, I., Rivas, M.A., Wang, Z.-M., Barisic, D., et al. (2024). ChromaFold predicts the 3D contact map from single-cell chromatin accessibility. *Nat. Commun.* 15, 9432. <https://doi.org/10.1038/s41467-024-53628-0>.
49. Fu, X., Mo, S., Buendia, A., Laurent, A.P., Shao, A., Alvarez-Torres, M.D.M., Yu, T., Tan, J., Su, J., Sagatelian, R., et al. (2025). A foundation model of transcription across human cell types. *Nature* 637, 965–973. <https://doi.org/10.1038/s41586-024-08391-z>.
50. Dorans, E., Jagadeesh, K., Dey, K., and Price, A.L. (2025). Linking regulatory variants to target genes by integrating single-cell multiome methods and genomic distance. *Nat. Genet.* 57, 1649–1658. <https://doi.org/10.1038/s41588-025-02220-3>.
51. 1000 Genomes Project Consortium; Auton, A., Brooks, L.D., Durbin, R.M., Garrison, E.P., Kang, H.M., Korbel, J.O., Marchini, J.L., McCarthy, S., McVean, G.A., and Abecasis, G.R. (2015). A global reference for human genetic variation. *Nature* 526, 68–74. <https://doi.org/10.1038/nature15393>.
52. Weissbrod, O., Hormozdiani, F., Benner, C., Cui, R., Ulirsch, J., Gazal, S., Schoech, A.P., van de Geijn, B., Reshef, Y., Márquez-Luna, C., et al. (2020). Functionally informed fine-mapping and polygenic localization of complex trait heritability. *Nat. Genet.* 52, 1355–1363. <https://doi.org/10.1038/s41588-020-00735-5>.
53. Tobias, D.K., Merino, J., Ahmad, A., Aiken, C., Benham, J.L., Bodhini, D., Clark, A.L., Colclough, K., Corcoy, R., Cromer, S.J., et al. (2023). Second international consensus report on gaps and opportunities for the clinical translation of precision diabetes medicine. *Nat. Med.* 29, 2438–2457. <https://doi.org/10.1038/s41591-023-02502-5>.
54. Misra, S., Wagner, R., Ozkan, B., Schön, M., Sevilla-Gonzalez, M., Prys-tupa, K., Wang, C.C., Kreienkamp, R.J., Cromer, S.J., Rooney, M.R., et al. (2023). Precision subclassification of type 2 diabetes: a systematic review. *Commun. Med.* 3, 138. <https://doi.org/10.1038/s43856-023-00360-3>.
55. DiCorpo, D., LeClair, J., Cole, J.B., Sarnowski, C., Ahmadizar, F., Bielak, L.F., Blokstra, A., Bottinger, E.P., Chaker, L., Chen, Y.-D.I., et al. (2022). Type 2 Diabetes Partitioned Polygenic Scores Associate With Disease Outcomes in 454,193 Individuals Across 13 Cohorts. *Diabetes Care* 45, 674–683. <https://doi.org/10.2337/dc21-1395>.
56. Leibiger, I.B., Leibiger, B., and Berggren, P.-O. (2008). Insulin Signaling in the Pancreatic β -Cell. *Annu. Rev. Nutr.* 28, 233–251. <https://doi.org/10.1146/annurev.nutr.28.061807.155530>.
57. Dooley, J., Tian, L., Schonefeldt, S., Delghingaro-Augusto, V., Garcia-Perez, J.E., Pasciuto, E., Di Marino, D., Carr, E.J., Oskolkov, N., Lyssenko, V., et al. (2016). Genetic predisposition for beta cell fragility underlies type 1 and type 2 diabetes. *Nat. Genet.* 48, 519–527. <https://doi.org/10.1038/ng.3531>.
58. Ghionescu, A.-V., Sorop, A., Linioudaki, E., Coman, C., Savu, L., Fogarasi, M., Lixandru, D., and Dima, S.O. (2024). A predicted epithelial-to-mesenchymal transition-associated mRNA/miRNA axis contributes to the progression of diabetic liver disease. *Sci. Rep.* 14, 27678. <https://doi.org/10.1038/s41598-024-77416-4>.
59. Donath, M.Y., and Shoelson, S.E. (2011). Type 2 diabetes as an inflammatory disease. *Nat. Rev. Immunol.* 11, 98–107. <https://doi.org/10.1038/nri2925>.
60. Luchsinger, J.A. (2012). Type 2 diabetes and cognitive impairment: linking mechanisms. *J. Alzheimers Dis.* 3, S185–S198. <https://doi.org/10.3233/JAD-2012-111433>.
61. Roh, E., Song, D.K., and Kim, M.-S. (2016). Emerging role of the brain in the homeostatic regulation of energy and glucose metabolism. *Exp. Mol. Med.* 48, e216. <https://doi.org/10.1038/emm.2016.4>.
62. Evangelou, E., Warren, H.R., Mosen-Ansorena, D., Mifsud, B., Pazoki, R., Gao, H., Ntritsos, G., Dimou, N., Cabrera, C.P., Karaman, I., et al. (2018). Genetic analysis of over 1 million people identifies 535 new loci associated with blood pressure traits. *Nat. Genet.* 50, 1412–1425. <https://doi.org/10.1038/s41588-018-0205-x>.
63. Keaton, J.M., Kamali, Z., Xie, T., Vaez, A., Williams, A., Goleva, S.B., Ani, A., Evangelou, E., Hellwege, J.N., Yengo, L., et al. (2024). Genome-wide analysis in over 1 million individuals of European ancestry yields improved polygenic risk scores for blood pressure traits. *Nat. Genet.* 56, 778–791. <https://doi.org/10.1038/s41588-024-01714-w>.
64. Vaura, F., Kim, H., Udler, M.S., Niiranen, T., FinnGen; Salomaa, V., and Lahti, L. (2022). Multi-Trait Genetic Analysis Reveals Clinically Interpretable Hypertension Subtypes. *Circ. Genom. Precis. Med.* 15, e003583. <https://doi.org/10.1161/CIRCGEN.121.003583>.
65. Pascat, V., Zudina, L., Maurin, L., Ulrich, A., Maina, J.G., Demirkan, A., Balkhiyarova, Z., Pupko, I., Sharhorodska, Y., Pattou, F., et al. (2025). Mechanistic Heterogeneity in Type 2 Diabetes and Hypertension Comorbidity Revealed with Partitioned Polygenic Scores. Preprint at medRxiv. <https://doi.org/10.1101/2025.03.02.25323190>.
66. Jelinic, M., Dona, M.S., Hughes, T.A.G., Farrugia, G., Harper, R., Hsu, I., Haslem, A., Tran, V., Galvao, H.B.F., Diep, H., et al. (2025). High-resolution transcriptomic profiling of the aortic cellular landscape during hypertension reveals novel drivers of vascular fibrosis. Preprint at bioRxiv. <https://doi.org/10.1101/2025.03.13.642936>.

67. Wagenseil, J.E., and Mecham, R.P. (2009). Vascular extracellular matrix and arterial mechanics. *Physiol. Rev.* 89, 957–989. <https://doi.org/10.1152/physrev.00041.2008>.
68. Humphrey, J.D., and Schwartz, M.A. (2021). Vascular Mechanobiology: Homeostasis, Adaptation, and Disease. *Annu. Rev. Biomed. Eng.* 23, 1–27. <https://doi.org/10.1146/annurev-bioeng-092419-060810>.
69. Rampino, T., and Dal Canton, A. (2003). Peritoneal Dialysis and Epithelial-to-Mesenchymal Transition (2003). *N. Engl. J. Med.* 348, 2037–2039. <https://doi.org/10.1056/NEJM200305153482020>.
70. Wang, J., Fan, C.-N., Wei, G.-X., Zhao, X.-J., Liu, K., Wang, M.-L., Li, L., Liu, M., and Zhao, H.-Y. (2025). Recurrence of primary aldosteronism 20 years after surgery: a case report. *J. Hum. Hypertens.* 39, 237–239. <https://doi.org/10.1038/s41371-025-00994-x>.
71. Ferreira, N.S., Tostes, R.C., Paradis, P., and Schiffrin, E.L. (2021). Aldosterone, Inflammation, Immune System, and Hypertension. *Am. J. Hypertens.* 34, 15–27. <https://doi.org/10.1093/ajh/hpaa137>.
72. Kurki, M.I., Karjalainen, J., Palta, P., Sipilä, T.P., Kristiansson, K., Donner, K.M., Reeve, M.P., Laivuori, H., Aavikko, M., Kaunisto, M.A., et al. (2023). FinnGen provides genetic insights from a well-phenotyped isolated population. *Nature* 613, 508–518. <https://doi.org/10.1038/s41586-022-05473-8>.
73. Vuckovic, D., Bao, E.L., Akbari, P., Lareau, C.A., Mousas, A., Jiang, T., Chen, M.-H., Raffield, L.M., Tardaguila, M., Huffman, J.E., et al. (2020). The Polygenic and Monogenic Basis of Blood Traits and Diseases. *Cell* 182, 1214–1231.e11. <https://doi.org/10.1016/j.cell.2020.08.008>.
74. Zhang, F., Xia, Y., Su, J., Quan, F., Zhou, H., Li, Q., Feng, Q., Lin, C., Wang, D., and Jiang, Z. (2024). Neutrophil diversity and function in health and disease. *Signal Transduct. Target. Ther.* 9, 343–349. <https://doi.org/10.1038/s41392-024-02049-y>.
75. Herrero-Cervera, A., Soehnlein, O., and Kenne, E. (2022). Neutrophils in chronic inflammatory diseases. *Cell. Mol. Immunol.* 19, 177–191. <https://doi.org/10.1038/s41423-021-00832-3>.
76. Vaibhav, K., Braun, M., Alverson, K., Khodadadi, H., Kutiyawalla, A., Ward, A., Banerjee, C., Sparks, T., Malik, A., Rashid, M.H., et al. (2020). Neutrophil extracellular traps exacerbate neurological deficits after traumatic brain injury. *Sci. Adv.* 6, eaax8847. <https://doi.org/10.1126/sciadv.aax8847>.
77. Zenaro, E., Pietronigro, E., Della Bianca, V., Piacentino, G., Marongiu, L., Budui, S., Turano, E., Rossi, B., Angiari, S., Dusi, S., et al. (2015). Neutrophils promote Alzheimer's disease-like pathology and cognitive decline via LFA-1 integrin. *Nat. Med.* 21, 880–886. <https://doi.org/10.1038/nm.3913>.
78. Trajanoska, K., Bhéer, C., Taliun, D., Zhou, S., Richards, J.B., and Mooser, V. (2023). From target discovery to clinical drug development with human genetics. *Nature* 620, 737–745. <https://doi.org/10.1038/s41586-023-06388-8>.
79. Minikel, E.V., Painter, J.L., Dong, C.C., and Nelson, M.R. (2024). Refining the impact of genetic evidence on clinical success. *Nature* 629, 624–629. <https://doi.org/10.1038/s41586-024-07316-0>.
80. Lumeng, C.N., Bodzin, J.L., and Saltiel, A.R. (2007). Obesity induces a phenotypic switch in adipose tissue macrophage polarization. *J. Clin. Invest.* 117, 175–184. <https://doi.org/10.1172/JCI29881>.
81. Jin, L., Deng, Z., Zhang, J., Yang, C., Liu, J., Han, W., Ye, P., Si, Y., and Chen, G. (2019). Mesenchymal stem cells promote type 2 macrophage polarization to ameliorate the myocardial injury caused by diabetic cardiomyopathy. *J. Transl. Med.* 17, 251. <https://doi.org/10.1186/s12967-019-1999-8>.
82. Cuomo, A.S.E., Nathan, A., Raychaudhuri, S., MacArthur, D.G., and Powell, J.E. (2023). Single-cell genomics meets human genetics. *Nat. Rev. Genet.* 24, 535–549. <https://doi.org/10.1038/s41576-023-00599-5>.
83. Kanamaru, K., Cranley, J., Muraro, D., Miranda, A.M.A., Ho, S.Y., Wilbrey-Clark, A., Patrick Pett, J., Polanski, K., Richardson, L., Litvinukova, M., et al. (2023). Spatially resolved multiomics of human cardiac niches. *Nature* 619, 801–810. <https://doi.org/10.1038/s41586-023-06311-1>.
84. Emani, P.S., Liu, J.J., Clarke, D., Jensen, M., Warrell, J., Gupta, C., Meng, R., Lee, C.Y., Xu, S., Dursun, C., et al. (2024). Single-cell genomics and regulatory networks for 388 human brains. *Science* 384, eadi5199. <https://doi.org/10.1126/science.adi5199>.
85. Richard, D., Muthurulan, P., Young, M., Yengo, L., Vedantam, S., Marouli, E., Bartell, E., GIANT Consortium; Hirschhorn, J., and Capellini, T.D. (2025). Functional genomics of human skeletal development and the patterning of height heritability. *Cell* 188, 15–32.e24. <https://doi.org/10.1016/j.cell.2024.10.040>.
86. Lee, D.D., and Seung, H.S. (1999). Learning the parts of objects by non-negative matrix factorization. *Nature* 401, 788–791. <https://doi.org/10.1038/44565>.
87. Wang, W., and Stephens, M. (2021). Empirical Bayes Matrix Factorization. *J. Mach. Learn. Res.* 22, 120.
88. Hou, L., Xiong, X., Park, Y., Boix, C., James, B., Sun, N., He, L., Patel, A., Zhang, Z., Molin, B., et al. (2023). Multitissue H3K27ac profiling of GTEx samples links epigenomic variation to disease. *Nat. Genet.* 55, 1665–1676. <https://doi.org/10.1038/s41588-023-01509-5>.
89. Hill, W.G., Goddard, M.E., and Visscher, P.M. (2008). Data and Theory Point to Mainly Additive Genetic Variance for Complex Traits. *PLoS Genet.* 4, e1000008. <https://doi.org/10.1371/journal.pgen.1000008>.
90. Mäki-Tanila, A., and Hill, W.G. (2014). Influence of gene interaction on complex trait variation with multilocus models. *Genetics* 198, 355–367. <https://doi.org/10.1534/genetics.114.165282>.
91. Zeng, J., de Vlaming, R., Wu, Y., Robinson, M.R., Lloyd-Jones, L.R., Yengo, L., Yap, C.X., Xue, A., Sidorenko, J., McRae, A.F., et al. (2018). Signatures of negative selection in the genetic architecture of human complex traits. *Nat. Genet.* 50, 746–753. <https://doi.org/10.1038/s41588-018-0101-4>.
92. Gazal, S., Loh, P.-R., Finucane, H.K., Ganna, A., Schoech, A., Sunyaev, S., and Price, A.L. (2018). Functional architecture of low-frequency variants highlights strength of negative selection across coding and non-coding annotations. *Nat. Genet.* 50, 1600–1607. <https://doi.org/10.1038/s41588-018-0231-8>.
93. Schoech, A.P., Jordan, D.M., Loh, P.-R., Gazal, S., O'Connor, L.J., Balick, D.J., Palamara, P.F., Finucane, H.K., Sunyaev, S.R., and Price, A.L. (2019). Quantification of frequency-dependent genetic architectures in 25 UK Biobank traits reveals action of negative selection. *Nat. Commun.* 10, 790. <https://doi.org/10.1038/s41467-019-08424-6>.
94. O'Connor, L.J., Schoech, A.P., Hormozdiari, F., Gazal, S., Patterson, N., and Price, A.L. (2019). Extreme Polygenicity of Complex Traits Is Explained by Negative Selection. *Am. J. Hum. Genet.* 105, 456–476. <https://doi.org/10.1016/j.ajhg.2019.07.003>.
95. Zhu, X., and Stephens, M. (2017). Bayesian large-scale multiple regression with summary statistics from genome-wide association studies. *Ann. Appl. Stat.* 11, 1561–1592. <https://doi.org/10.1214/17-aos1046>.

STAR★METHODS

KEY RESOURCES TABLE

REAGENT or RESOURCE	SOURCE	IDENTIFIER
Deposited data		
EpiMap bulk epigenomic tracks	Boix et al. ³⁵	https://compbio.mit.edu/epimap/
Single-cell ATAC-seq	Zhang et al. ⁴⁴	GEO: GSE184462
GWAS summary statistics for 165 diseases/traits	This paper	https://alkesgroup.broadinstitute.org/J-PEP
Auxiliary traits (109 traits from ref. 25)	This paper	https://alkesgroup.broadinstitute.org/J-PEP
cS2G (SNP-to-gene linking annotations)	Gazal et al. ⁴⁶	
J-PEP input/output matrices and PEPA/PPA/EPA	This paper	https://alkesgroup.broadinstitute.org/J-PEP
Software and algorithms		
jpepR (J-PEP, J-PEP-CT, PEPA and bNMF code)	This paper	https://github.com/kernerg/jpepR (https://doi.org/10.5281/zenodo.17704042)
Flashier	Willwerscheid et al. ⁴²	https://doi.org/10.48550/arXiv.2110.00152
FactorGO	Zhang et al. ²⁶	https://github.com/mancusolab/FactorGo
MACS2 (peak calling)	Github	https://macs3-project.github.io/MACS/
UCSC bigWigToBedGraph	UCSC Genome Browser	https://genome.ucsc.edu/goldenpath/help/bigWig.html
R (statistical computing environment)	R Core Team	Version 4.1.2

METHOD DETAILS

J-PEP method: Input and output

J-PEP inputs V_{trait} and V_{tissue} and outputs W , H_{trait} and H_{tissue} .

$V_{\text{trait}} \in \mathbb{R}^{M \times A}$ is a SNP-to-auxiliary trait matrix defined for a given focal trait. M is the number of fine-mapped SNPs (PIP > 0.01) for the focal trait and A is the number of auxiliary traits. The entries of this matrix are defined as:

$$(V_{\text{trait}})_{m,a} = \text{PIP}_m^f \times \text{PIP}_m^a \quad (\text{Equation 2})$$

where PIP_m^f and PIP_m^a denote the posterior inclusion probabilities (PIPs) for SNP m in the focal trait f and auxiliary trait a , respectively. Because J-PEP relies on non-negative matrix factorization, each auxiliary trait is split into two components: a concordant pleiotropic component (denoted by a “+” sign in figures) and a discordant pleiotropic component (denoted by a “−” sign).

$V_{\text{tissue}} \in \mathbb{R}^{M \times T}$ is a SNP-to-tissue matrix defined for a given focal trait. M is the number of fine-mapped SNPs (PIP > 0.01) for the focal trait and T is the number of tissues. The entries of this matrix are defined as:

$$(V_{\text{tissue}})_{m,t} = \hat{a}_{m,t} \quad (\text{Equation 3})$$

where $\hat{a}_{m,t}$ is a normalized score quantifying the strength of association between SNP m and tissue t derived from tissue-specific epigenomic profiles. We generate this using the E-M algorithm. Details are below (see **EpiMap dataset**).

$W \in \mathbb{R}^{M \times K}$ is a SNP-to-cluster membership matrix defined for a given focal trait. M is the number of fine-mapped SNPs (PIP > 0.01) for the focal trait and K is the number of clusters identified by the J-PEP method.

$H_{\text{trait}} \in \mathbb{R}^{K \times A}$ is a cluster-specific profile matrix for auxiliary traits. K is the number of clusters identified by the J-PEP method and A is the number of auxiliary traits, as in V_{trait} .

$H_{\text{tissue}} \in \mathbb{R}^{K \times T}$ is a cluster-specific profile matrix for tissues. K is the number of clusters identified by the J-PEP method and T is the number of tissues, as in V_{tissue} .

We note that neither the columns of W nor the rows of H matrices are constrained to sum to 1, as is the case in other factorization approaches,^{26,42} since the partitioning approach preserves relative contributions of clusters to the overall partition. Reconstructed V matrices are not used for benchmarking but only inferred W and H matrices, since their reconstruction would only reproduce observed data.

J-PEP method: Objective function

J-PEP performs joint non-negative matrix factorization on V_{trait} and V_{tissue} , enforcing a shared SNP-to-cluster membership matrix W . The objective function is:

$$\{W, H_{\text{trait}}, H_{\text{tissue}}\} = \underset{W \geq 0, H_{\text{trait}} \geq 0, H_{\text{tissue}} \geq 0}{\operatorname{argmin}} \left\{ \frac{1}{2} \|V_{\text{trait}} - WH_{\text{trait}}\|^2 + \frac{1}{2} \|V_{\text{tissue}} - WH_{\text{tissue}}\|^2 + \sum_k \lambda ((V_{\text{trait}} H_{\text{trait}}^T)_k, (V_{\text{tissue}} H_{\text{tissue}}^T)_k) \right\} \quad (\text{Equation 4})$$

where k is taken across the K clusters (columns of $V_{\text{trait}} H_{\text{trait}}^T$ and $V_{\text{tissue}} H_{\text{tissue}}^T$) from the current iteration of the algorithm and the penalty function λ promotes alignment between pleiotropic and epigenomic data:

$$\lambda(A, B) = -\log_{10}(p_{\text{corr}}(A, B))^{-1} \quad (\text{Equation 5})$$

where p_{corr} indicates the one-sided Pearson's correlation p value for positive correlation.

We note that since $V_{\text{trait}} H_{\text{trait}}^T$ and $V_{\text{tissue}} H_{\text{tissue}}^T$ approximate the shared cluster membership matrix W (and are equal to it when H_{trait} and H_{tissue} are orthogonal matrices, which is rare in practice, and otherwise approximate W by projecting the observed data into the cluster membership space spanned by H , as is standard in matrix factorization approaches^{86,87}), the penalty function λ encourages clusters where SNP memberships are consistent across modalities (i.e., trait and tissue data). We refer to λ as the penalty function and the entries of the vector λ as the penalty weights.

We also note that λ depends on the number of fine-mapped SNPs, so the method uses a scaling parameter to adjust the strength of the penalty and can be tuned by the user.

J-PEP method: Optimization

We optimized the J-PEP objective function using an iterative algorithm alternating updates of the shared SNP-to-cluster membership matrix W , the trait-specific cluster profile matrix H_{trait} , and the tissue-specific profile matrix H_{tissue} . All matrices were initialized with small positive random values scaled to the input range. Updates were performed using multiplicative update rules derived from the gradient of the objective function⁴¹ (also referred to as gradient descent updates), as done in bNMF.⁴⁰ Specifically, these updates can be derived from the gradient of the objective function with respect to the H and W matrices, and the multiplicative update rules correspond to a projected gradient descent step under non-negativity constraints. Updates also incorporate the penalization function λ , whose length is equal to the inferred number of clusters K . Entries of the vector λ (penalty weights) enforce consistency between trait-based and tissue-based cluster assignments by penalizing clusters whose SNP memberships are poorly correlated across the two modalities. To maintain numerical stability, a small constant $\varepsilon = 10^{-50}$ was added to all matrix approximations.

Convergence was assessed by monitoring relative changes in the penalization weights, i.e., the entries of the vector λ , applied to each cluster. The algorithm terminated when the maximum relative change of the penalty weights across iterations fell below a pre-defined threshold (default: 10^{-7}) or when the maximum number of iterations (default: 5,000) was reached (very rare in practice, see Figure S7C). The iteration cap served as a fail-safe rather than a convergence condition; if the cap is hit without meeting the tolerance, the routine stops with an explicit “maximum iterations without convergence” error and the run is not used. We discarded solutions in which fewer than two non-redundant clusters remained active (defined by row sums of H_{trait} and H_{tissue} exceeding a minimal threshold, defined by default to 10^{-10}). Additionally, we discarded solutions in which at least one method failed to identify at least two clusters in a minimum of 18 out of the 22 model runs required for computing prediction accuracy metrics (one per chromosome removed; see below). Applying these two criteria retained 38 of the original 55 focal traits from the full set of 165 defined as approximately genetically independent (pairwise $r_g < 0.5$).

Sparsity constraint

To improve stability and interpretability, J-PEP imposes tissue-to-cluster memberships (columns of H_{tissue}) to be either zero or positive for at most one cluster, as previously proposed,⁸⁸ whereas auxiliary traits remain free to associate with multiple clusters. Importantly, a tissue remains unassigned if its highest tissue-to-cluster value is zero (or for practical purposes, very close to 0). Conceptually, J-PEP optimizes the objective function of Equation 1 under the constraint that $H_{\text{tissue}} \in C$, where C is the set of matrices whose columns each have at most one non-zero entry. For robustness, we perform ten optimization runs, as previously proposed,²³ with different random seeds and retain the solution that yielded the highest average projection correlation between modalities, i.e., the highest average correlation between $V_{\text{trait}} H_{\text{trait}}^T$ and $V_{\text{tissue}} H_{\text{tissue}}^T$ across clusters, among those yielding the most consistent number of clusters; in case of ties, we prioritize the run with fewer clusters.

Cluster ordering

In Figures 5 and 6, we ordered clusters based on the proportion of total variance they jointly explained across both trait and tissue association matrices. Specifically, for each cluster k , we computed its contribution to the SNP-to-trait matrix V_{trait} and the SNP-to-tissue matrix V_{tissue} using rank-1 reconstructions: $V_{\text{trait},k} = W_{\cdot,k} H_{\text{trait}}[k, \cdot]$ and $V_{\text{tissue},k} = W_{\cdot,k} H_{\text{tissue}}[k, \cdot]$. The total variance explained by cluster k was defined as the sum of squared entries in $V_{\text{trait},k}$ and $V_{\text{tissue},k}$, normalized by the total variance of the original matrices. Clusters were then ranked in decreasing order of this joint variance contribution.

Pleiotropic partitioning method

The pleiotropic partitioning method factorizes V_{trait} using bNMF, that is, by solving the optimization problem:

$$\{W, H_{trait}\} = \underset{W \geq 0, H_{trait} \geq 0}{\operatorname{argmin}} \left\{ \|V_{trait} - WH_{trait}\|^2 + \sum_k \lambda_{single}(W_k, H_{trait,k}) \right\} \quad (\text{Equation 6})$$

where W and H_{trait} are defined as above, and λ_{single} is a penalty term that reduces overfitting by pruning irrelevant clusters.⁴⁰ The maximum number of clusters K is predefined (typically $K = 15$), but the actual number of identified clusters by the partitioning method is often reduced as irrelevant components are penalized to zero. M includes all fine-mapped SNPs with PIP values ≥ 0.01 for the focal trait. The code to run this step was directly taken from ref.²³ without modification. On the other hand, the following steps to obtain H_{tissue} (which is not produced by pleiotropic partitioning methods but is required here to enable comparison to other methods using our new PEPA metric) are original to this work.

We generate cluster-specific tissue profiles (H_{tissue}) while ensuring that each tissue associates with a single cluster, as follows.

- 1) We estimate H_{tissue} from V_{tissue} and W using gradient descent by solving $V_{tissue} = W \times H_{tissue}$, without applying any constraint on H_{tissue} .
- 2) Clusters with highly similar tissue profiles are merged by averaging their respective cluster memberships. Similarity is assessed using the *simil* function from the R v.4.2 proxy package, with “high similarity” defined as a similarity score > 0.5 .
- 3) We repeat step (1) but now enforce the sparsity condition, as implemented in the J-PEP method.

Epigenomic partitioning method

The epigenomic partitioning method factorizes V_{tissue} using bNMF, that is, by solving the analogous optimization problem to the pleiotropic partitioning method, using as input V_{tissue} instead of V_{trait} :

$$\{W, H_{tissue}\} = \underset{W \geq 0, H_{tissue} \geq 0}{\operatorname{argmin}} \left\{ \|V_{tissue} - WH_{tissue}\|^2 + \sum_k \lambda_{single}(W_k, H_{tissue,k}) \right\} \quad (\text{Equation 7})$$

The code to run this step was directly taken from ref.²³ without modification.

We learn cluster-specific auxiliary trait profiles (H_{trait}) from V_{trait} and W using gradient descent by solving $V_{trait} = W \times H_{trait}$.

Prediction accuracy metrics

PPA. Let $V_{trait} \in \mathbb{R}^{M \times A}$ and $V_{tissue} \in \mathbb{R}^{M \times T}$ represent the SNP-to-trait and SNP-to-tissue matrices, respectively, where M is the number of fine-mapped SNPs (PIP > 0.01) for the focal trait, A is the number of auxiliary traits, and T is the number of tissues. The prediction process aims to estimate V_{trait} using the SNP-to-cluster membership matrix W and the tissue cluster profiles H_{tissue} and H_{trait} . Essentially, PPA follows a two-step approach.

- 1) For a given cluster k and a held-out chromosome c , we predict the cluster membership $W_k^{(c)}$ based on $V_{tissue}^{(c)}$ and $H_{tissue_k}^{-(c)}$, the SNP-to-tissue associations for chromosome c and the tissue profiles inferred from the remaining 21 chromosomes, using gradient descent.
- 2) Then, the predicted SNP-to-trait associations for cluster k and chromosome c , i.e., $\hat{V}_{trait_k}^{(c)}$, are calculated as:

$$\hat{V}_{trait_k}^{(c)} = W_k^{(c)} H_{trait_k}^{-(c)} \quad (\text{Equation 8})$$

The Pearson correlation between the predicted values ($\hat{V}_{trait_k}^{(c)}$) and the true values ($V_{trait}^{(c)}$) for each cluster is computed as:

$$r_k^{(c)} = \operatorname{cor}(\hat{V}_{trait_k}^{(c)}, V_{trait}^{(c)}) \quad (\text{Equation 9})$$

The weighted mean correlation across all clusters in chromosome c is computed as:

$$r_{trait}^{(c)} = \frac{\sum_k r_k^{(c)} w_k^{(c)}}{\sum_k w_k^{(c)}} \quad (\text{Equation 10})$$

where $w_k^{(c)}$ is defined as:

$$w_k^{(c)} = \frac{\operatorname{Var}(W_k^{(c)})}{\operatorname{Var}(\sum_j W_j^{(c)})} \quad (\text{Equation 11})$$

EPA

Computation for EPA is analogous to PPA, inverting the roles of trait and tissue data (e.g., V_{tissue} for V_{trait} and V_{trait} for V_{tissue}). The Pearson correlation between the predicted values ($\hat{V}_{tissue_k}^{(c)}$) and the true values ($V_{tissue}^{(c)}$) for each cluster is computed as:

$$r_k^{(c)} = \text{cor}(\hat{V}_{tissue_k}^{(c)}, V_{tissue}^{(c)}) \quad (\text{Equation 12})$$

The weighted mean correlation across all clusters in chromosome c is computed as:

$$r_{tissue}^{(c)} = \frac{\sum_k r_k^{(c)} w_k^{(c)}}{\sum_k w_k^{(c)}} \quad (\text{Equation 13})$$

PEPA

The overall PEPA score is the sum of the weighted average correlations for both prediction tasks across all chromosomes:

$$PEPA = \sum_{c \in C} (\sqrt{r_{tissue}^{(c)} \times r_{trait}^{(c)}} \times \text{pip}^{(c)}) \quad (\text{Equation 14})$$

where $\text{pip}^{(c)}$ denotes the proportion of fine-mapping PIPs on chromosome c relative to the total across all held-out chromosomes in the set C , where C indicates the group of chromosomes containing at least five GWAS loci. PEPA was only computed in cases in which $|C| > 18$.

Although inference approaches can differ in the number of discovered clusters, the PEPA metric ensures comparability across methods because it evaluates prediction accuracy in the observed data space, independent of the number of inferred clusters.

Projection correlation

To validate trait-tissue coherence across identified clusters, we computed a projection correlation metric beyond the PEPA metric, defined as $r_{proj} = \text{cor}(V_{trait} \times H_{trait}^t, V_{tissue} \times H_{tissue}^t)$.

Methodological considerations and limitations of J-PEP

We highlight three methodological considerations. First, J-PEP assumes a linear shared structure between modalities, which may overlook non-linear interactions between pleiotropy and epigenomic data; however, the linearity assumption is likely to be a reasonable approximation given the largely additive nature of genetic effects.^{22,89,90} Second, the requirement for concordant clustering across traits and tissues may suppress modality-specific associations (e.g., clusters detectable only through pleiotropy or epigenomics alone); this constraint also underlies the definition of the PEPA metric. However, this restriction likely improves robustness by focusing on clusters supported by multiple lines of evidence. Similarly, the geometric mean used in the PEPA metric penalizes imbalance between modalities (traits and tissues), thus failing to prioritize solutions where one modality performs well while the other performs poorly. Nevertheless, we chose the geometric mean specifically because J-PEP aims to identify clusters with concordant support across modalities. Third, the different statistical properties of V_{trait} and V_{tissue} arising from their different scaling with respect to focal trait PIPs is a primary factor impacting the relative concordance of J-PEP with pleiotropic partitioning vs. epigenomic partitioning; however, the alternative that rescales V_{tissue} by focal trait PIPs performed worse than the default approach (Figure S5C; Table S3).

Simulation framework

Generating simulated data consists of 8 steps: Step 1: simulating cluster-specific profiles, Step 2: estimating the SNP-to-cluster membership matrix, Step 3: computing SNP-to-trait associations, Step 4: sampling causal SNPs, Step 5: sampling causal effect sizes, Step 6: sampling GWAS summary statistics, Step 7: fine-mapping, Step 8: generating uncorrelated structure.

Step 1. Simulating cluster-specific profiles

We initialized cluster-specific trait and tissue profiles as matrices $H_{trait} \in \mathbb{R}^{A \times K}$ and $H_{tissue} \in \mathbb{R}^{T \times K}$, where A is the number of auxiliary traits, T the number of tissues, and K the number of causal clusters. Each matrix was initialized with standard uniform values, and a subset of traits and tissues was randomly assigned as causal for each cluster (one trait and one tissue per cluster) by retaining non-zero values only in the corresponding columns:

$$H_{trait}[i, j] = \begin{cases} U[0.1, 1], & \text{if } i \in \text{causal trait for cluster } j \\ 0, & \text{otherwise} \end{cases} \quad (\text{Equation 15})$$

$$H_{tissue}[i, j] = \begin{cases} U[0.1, 1], & \text{if } i \in \text{causal tissue for cluster } j \\ 0, & \text{otherwise} \end{cases} \quad (\text{Equation 16})$$

Each row was normalized to sum to one. We used $K = 5$ by default and tested scenarios with $K = 3$ and $K = 8$.

Step 2. Estimating SNP-to-cluster membership matrix

The SNP-to-cluster membership W , which probabilistically maps SNPs to clusters, was inferred using gradient descent, a SNP-to-tissue matrix from real data analysis and simulated H_{tissue} , by solving the factorization

problem: $\{\hat{W}\} = \underset{W \geq 0}{\operatorname{argmin}} \left\{ \frac{1}{2} \|\tilde{V}_{tissue} - WH_{tissue}\|^2 \right\}$ where $\tilde{V}_{tissue}$ is a SNP-to-tissue matrix constructed by intersecting SNPs from chromosomes 1 and 2 of 1000 Genomes Phase 3⁴³ with EpiMap³⁵ annotations (see **EpiMap dataset**). Only SNPs with MAF > 0.05 were retained.

Step 3. Computing SNP-to-trait associations

The SNP-to-trait matrix V_{trait} was obtained by matrix multiplication: $W \times H_{trait}$.

Step 4. Sampling causal SNPs

Causal SNPs for the focal and auxiliary traits were sampled from distinct causal loci, ensuring independence between all causal SNPs. These were sampled with probability proportional to:

$$P(\text{causal SNP } i \text{ for trait } j) \propto V_{trait}[i, j] \times MAF_i \quad (\text{Equation 17})$$

where MAF_i is the MAF of SNP i . We assumed the first column of V_{trait} represents the focal trait. This sampling scheme has the consequence that common variants explain more trait variance (analogous to real traits, due to negative selection^{91–93}) and also the consequence that common variant architectures are more polygenic (analogous to real traits, due to negative selection⁹⁴).

For auxiliary traits contributing to clusters, we enforced at least 25% overlap in causal SNPs with the focal trait. We varied the number of total causal SNPs $m \in \{100; 200; 300; 400\}$, sampling each from different LD blocks on chromosomes 1 and 2.

Step 5. Sampling causal effect sizes

Effect sizes were sampled from an independent standard normal distribution. Each trait's total SNP-heritability ($h_j^2 = 0.5$) was distributed equally across its causal SNPs by rescaling effect sizes such that their squared sum matched h_j^2 .

Step 6. Simulating GWAS summary statistics

We then simulated summary statistics using the RSS likelihood,⁹⁵ $\hat{\beta} \sim N(SRS^{-1}\beta, SRS)$, where β is the vector of true causal effect sizes, R is the LD matrix from chromosomes 1 and 2 of the 1000 Genomes Project Phase 3, S is a diagonal matrix with elements $S_{ii} = \frac{1}{N2p_i(1-p_i)}$, N is the sample size and p_i the minor allele frequency of SNP i . We used a regularized LD matrix, $R \times (1-\lambda) + \lambda I$, where $\lambda = 0.001$.

Step 7. Fine-mapping

We performed single causal variant fine-mapping^{38,39} on all simulated GWAS summary statistics as in ref.²⁰. We note that multiple causal variant fine-mapping is not recommended when in-sample LD is not available.⁵²

Step 8. Generating uncorrelated structure

We define uncorrelated structure as a reproducible, modality-specific component of variation, i.e., a pattern present in one modality that has no counterpart in the other. To simulate trait- or tissue-specific variation unrelated to the predefined cluster structure, we added uncorrelated, yet “important”, features (i.e., auxiliary traits for V_{trait} and tissues for V_{tissue}) to both V_{trait} and V_{tissue} . We first ranked auxiliary traits and tissues by their total contribution to the simulated clusters, then randomly selected a subset of low-contributing (“non-important”) features, defined as features that were not associated to the predefined clusters and thus had little or no contribution to the simulated clusters, and an equally sized subset of high-contributing (“important”) features. For each selected non-important feature, we replaced its profile (i.e., the SNP-to-feature scores) by resampling values from one of the sampled top-contributing feature of the same modality—maintaining realistic signal strength while disrupting alignment with the cluster structure. The number of such substituted features (traits and tissues) defines the level of uncorrelated structure simulated (s). This is an approach original to this paper.

Representative simulation example

We selected a representative simulation scenario with $K_{causal} = 5$ and $m = 200$ for visualization in [Figure S1](#). To improve clarity, only 4 of the 5 clusters are shown, excluding the cluster with the weakest trait–tissue correlation based on r_{proj} . To align colors between true and inferred clusters across methods, we computed correlations between the true and inferred trait and tissue profiles, assigning each inferred cluster the color of its best-matching true cluster.

Additional simulation analyses

Downsampling

To evaluate the impact of LD among non-causal SNPs on model performance, we performed a downsampling procedure using the same simulation parameters as in the main simulations presented in the main text (e.g., $K_{causal} = 5$ and $s = 2$). In each replicate, all true causal SNPs were retained, while the set of non-causal SNPs was progressively reduced. Specifically, we randomly removed 0% (default), 20%, 40%, or 60% of the remaining non-causal SNPs. We then applied J-PEP, pleiotropic partitioning, and epigenomic partitioning to the downsampled data and evaluated performance using the PEPA metric. We present results for scenarios with $m = 200$ or $m = 400$ causal SNPs.

Varying PIP threshold

To assess how the choice of posterior inclusion probability (PIP) threshold influences model performance, we conducted additional simulations under the same parameter settings as in the main simulations presented in the main text (e.g., $K_{causal} = 5$ and $s = 2$). For each replicate, we varied the minimum PIP threshold used to retain SNPs in the input matrices. In addition to the default threshold of 0.01, we tested stricter thresholds (0.05, 0.1, 0.2, 0.5) as well as more lenient thresholds (0.005 and 0.001). For each threshold, we

constructed V_{trait} and V_{tissue} accordingly, ran J-PEP, pleiotropic partitioning, and epigenomic partitioning, and evaluated prediction accuracy using the PEPA metric. We present results for scenarios with $m = 200$ or $m = 400$ causal SNPs.

FactorGO and Flashier implementations

To benchmark J-PEP against independently developed approaches, we conducted additional analyses with FactorGO²⁶ and three variants of Flashier⁴² (Flashierz, Flashier, and Flashier_snn), applied both in simulations (under the default primary simulations procedure) and in real data.

FactorGO

FactorGO was run on a z-score-based V_{trait} matrix, where entries are z-scores of auxiliary traits and the focal trait for fine-mapped SNPs of the focal trait. We used default settings, with a maximum of $K = 15$ initial factors, consistent with the maximum cluster number allowed for J-PEP, pleiotropic, and epigenomic partitioning.

Flashierz

Flashier applied to the same z-score-based V_{trait} matrix with $K = 15$.

Flashier

Flashier directly applied to the original PIP-based V_{trait} .

Flashier_snn

A semi-non-negative version of Flashier, enforcing non-negativity on one of the inferred factor matrices, applied to the PIP-based V_{trait} . For this configuration we additionally applied non-negative gradient descent to obtain H_{tissue} . This is the only Flashier variant used in real data analysis as it was the best performing in simulations.

For all Flashier runs, we set backfit = TRUE. Since these methods only factorize the trait matrix, they do not provide H_{tissue} . To enable comparison with J-PEP, we projected tissues onto the inferred SNP-to-cluster membership matrix W using gradient descent, applying a non-negativity constraint only for the Flashier_snn configuration (as also done in pleiotropic partitioning).

Performance was evaluated using two criteria: (i) the PEPA metric (in both simulations and real data) and (ii) matrix reconstruction error, assessed via the Frobenius norm error of the reconstructed H_{trait} and H_{tissue} matrices, and only assessed in simulations.

GWAS summary statistics

We applied partitioning methods to GWAS summary statistics for 165 focal diseases/traits (average of 290K samples and 77 fine-mapped loci), including 57 UK Biobank diseases/traits⁴³ and 108 non-UK Biobank diseases/traits (Table S1; see Data Availability).

Fine mapping

We conducted single causal variant fine-mapping on all tested GWAS traits as in ref.²⁰

Combining results across GWAS diseases/traits

To evaluate significant differences between J-PEP's PEPA average across GWAS diseases/traits with respect to other partitioning strategies, we used two approaches. First, we applied a genomic block jackknife to the unweighted average PEPA difference across traits (denoted p in the main text). Second, we performed a jackknife using trait-specific weights based on the inverse variance of each trait's PEPA difference (denoted p_w in the main text).

EpiMap dataset

Track calling

We obtained EpiMap³⁵ p value signal tracks from the Washington University Epigenome for downstream analysis. We restricted our selection to tracks that met two criteria: (1) the tracks were directly observed rather than imputed, and (2) the tracks corresponded to one of six epigenomic annotations — DNase hypersensitive sites (DNase HS), H3K4me1, H3K4me2, H3K4me3, H3K9ac, or H3K27ac. The selected tracks were initially downloaded in bigWig format and converted to bedGraph format using the UCSC command-line utility bigWigToBedGraph. Following this conversion, we performed peak calling using MACS2 to identify significant regions of epigenomic activity. Peaks were called using the macs2 bdgpeakcall command with a significance threshold of $-\log_{10}(p) > 2$, ensuring that only high-confidence peaks were included in the subsequent analyses.

SNP-to-tissue scoring using the EM algorithm

To compute SNP-to-tissue scores from epigenomic annotations, we applied an expectation maximization (EM) algorithm to fine-mapped variants overlapping epigenomic peaks from the EpiMap project. This approach is original to this paper.

Briefly, the EM algorithm is used to infer a latent, unobserved variable a_{pt} that indicates if SNP p is active (i.e., overlaps a causal regulatory element) in annotation t .

Let $V \in \{0, 1\}^{M \times t}$ denote the binary overlap matrix of M SNPs and t epigenomic annotations. Then, formally, the algorithm maximizes the following likelihood:

$$L(\theta) = \prod_{p=1}^M \sum_{t=1}^T \theta_t \frac{V_{pt} w_p}{b_t} \quad (\text{Equation 18})$$

where V_{pt} indicates the observed overlap of SNP p and annotation t , w_p is the PIP weight for SNP p , b_t indicates the background expectation for annotation t , and θ_t is the annotation-level enrichment weight.

In practice, we initialize posterior weights a_{pt} across annotations for each SNP p by normalizing overlap rows:

$$a_{pt}^{(0)} = \frac{V_{pt}}{\sum_{t'} V_{pt'}} \quad (\text{Equation 19})$$

At each iteration z , we re-estimate annotation weights θ_t as a weighted sum over SNPs (M-step):

$$\theta_t^{(z)} = \frac{1}{Z} \sum_{p=1}^M \frac{a_{pt}^{(z-1)} w_p}{b_t} \quad (\text{Equation 20})$$

where w_p is the SNP-specific weight based on its PIP, b_t is the background expectation for annotation t , and Z is a normalization constant ensuring $\sum_t \theta_t = 1$.

Posterior weights a_{pt} are then updated as (E-step):

$$a_{pt}^{(z)} = \frac{V_{pt} \theta_t^{(z)}}{\sum_{t'} V_{pt'} \theta_{t'}^{(z)}} \quad (\text{Equation 21})$$

Iterations proceed until convergence is achieved (mean squared difference between $a_{pt}^{(z)}$ and $a_{pt}^{(z-1)}$ below 10^{-8} , or a maximum of 200 iterations). SNPs were retained if their posterior inclusion probability (PIP) exceeded a default PIP threshold (default: 0.01). Functional annotations were filtered to exclude cancer samples and tracks with extreme background enrichment values (computed on chromosome 1 of the 1000 Genomes Project⁵¹).

Track-level scores were aggregated into tissue-level scores by summing the normalized posterior probabilities from the EM algorithm across all tracks assigned to predefined tissue groups (as defined in EpiMap metadata³⁵), yielding a final SNP-to-tissue matrix $V_{\text{tissue}} \in \mathbb{R}^{M \times T}$, where T is the number of grouped tissues.

Single-cell dataset

We analyzed a single-cell atlas of chromatin accessibility comprising 615,998 cells across 111 fine-grained adult human cell types.⁴⁴ Cell-type-specific peak annotations were downloaded using the following command: `wget -r -level=1 -np -R index.html http://catlas.org/catlas_downloads/humantissues/Peaks/`.

We retained the original cell-type labels as provided by the authors. The SNP-to-cell-type association matrix $V_{\text{cell-type}}$ was constructed using the same procedure applied to the bulk tissue annotations (see **EpiMap dataset**) but substituting in the single-cell-derived peak tracks.

J-PEP-CT: Inference of cell-type-derived cluster profiles

We refer to the extension of J-PEP to single-cell resolution as J-PEP-CT. J-PEP-CT consists of two steps. (i) First, we applied J-PEP to bulk data in order to infer the shared SNP-to-cluster membership matrix W ; (ii) then, we derived cell-type-specific cluster profiles, denoted $H_{\text{cell-type}}$, by solving

$$V_{\text{cell-type}} \approx W \times H_{\text{cell-type}},$$

where $V_{\text{cell-type}}$ denotes the SNP-to-cell-type association matrix obtained from the EM algorithm on scATAC-seq data across 111 cell types (analogous to V_{tissue}).

Estimation of $H_{\text{cell-type}}$ was done using a non-negative least squares procedure with multiplicative update rules analogous to those used in bNMF. Importantly, we maintained the same tissue-sparsity constraint as in the bulk analysis, restricting each cell type to be assigned to at most one cluster, while allowing the possibility of remaining unassigned. This approach enables projection of cell-type data onto bulk-derived clusters, ensuring consistency between bulk- and cell-type-level analyses.

Gene ontology (GO) enrichment analysis

To investigate the biological processes associated with each cluster, we conducted Gene Ontology (GO) term enrichment analysis using the “Enrichr” R package. SNP-to-cluster scores were mapped to genes using pre-annotated SNP-to-gene mappings,⁴⁶ and gene-level scores were calculated by summing the scores of all SNPs assigned to each gene. To prioritize the most informative signals, genes were ranked by their cluster scores, and only the top genes cumulatively contributing to 80% of the total cluster weight were retained. GO terms were considered significant if they included at least four overlapping genes and had a Benjamini-Hochberg adjusted p -value below 0.05.

QUANTIFICATION AND STATISTICAL ANALYSIS

The quantitative and statistical analyses are described thoroughly in the relevant sections of the [method details](#).