

**Constructing genotype and phenotype network helps reveal disease
heritability and phenome-wide association studies**

Xuewei Cao^{1,#}, Lirong Zhu^{1,#}, Xiaoyu Liang², Shuanglin Zhang¹, Qiuying Sha^{1,*}

¹Department of Mathematical Sciences, Michigan Technological University, Houghton,
Michigan, USA, 49931

²Department of Epidemiology and Biostatistics, Michigan State University, East Lansing,
MI, USA, 48824

Both authors contributed equally

*Corresponding author: QIUYING SHA, Department of Mathematical Sciences,
Michigan Technological University, Houghton, Michigan 49931, USA. E-mail:
qsha@mtu.edu

Text S1. Details of other network properties by comparing with random network

In this study, we also consider two network properties, degree entropy, and cross entropy, in comparison between the constructed sparse representation of GPN, $\mathcal{G}_{GPN}^\tau$, and the corresponding random network, $\mathcal{G}_{GPN}^{random}$, for a specific $\tau \in (0,1)$.

Degree entropy. The Shannon entropy of the degree can be used to measure the diversity of associations between genetic variants and phenotypes through their degrees(1). For a specific threshold τ , we define the Shannon entropy for the degrees

of genetic variants and phenotypes as $H_\tau^G = -\sum_{m=1}^M \bar{d}_m^G \log(\bar{d}_m^G)$ and $H_\tau^P = -\sum_{k=1}^K \bar{d}_k^P \log(\bar{d}_k^P)$,

where the min-max standardized degrees are given by

$\bar{d}_m^G = (d_m^G - \min_m \{d_m^G\}) / (\max_m \{d_m^G\} - \min_m \{d_m^G\})$ for the m^{th} genetic variant and

$\bar{d}_k^P = (d_k^P - \min_k \{d_k^P\}) / (\max_k \{d_k^P\} - \min_k \{d_k^P\})$ for the k^{th} phenotype. The global degree

entropy of a bipartite network is given by $H_\tau = H_\tau^G + H_\tau^P$. For the corresponding random

network, we use the same way to calculate the degree entropies,

$H_\tau^{G,random}, H_\tau^{P,random}$, and H_τ^{random} .

Cross Entropy. We define the cross-entropy of weighted or unweighted degree of genetic variants and phenotypes between $\mathcal{G}_{GPN}^\tau$ and $\mathcal{G}_{GPN}^{random}$ to determine the diversity between a bipartite GPN and a random bipartite network.

$$H_{cross}^G = -\sum_{m=1}^M \left[\bar{d}_m^G \log(\bar{d}_m^{G,random}) + (1 - \bar{d}_m^G) \log(1 - \bar{d}_m^{G,random}) \right],$$

$$H_{cross}^P = -\sum_{k=1}^K \left[\bar{d}_k^P \log(\bar{d}_k^{P,random}) + (1 - \bar{d}_k^P) \log(1 - \bar{d}_k^{P,random}) \right],$$

32 where, H_{cross}^G and H_{cross}^P are used to measure the difference between degree
 33 distributions of genetic variants and phenotypes in $\mathcal{G}_{GPN}^\tau$ and $\mathcal{G}_{GPN}^{random}$, respectively.
 34 H_{cross}^G and H_{cross}^P are always positively valued and it will increase if the degree
 35 distributions tend to be more different. Same as degree entropy, we also define the
 36 global cross entropy of a bipartite network as $H_{cross} = H_{cross}^G + H_{cross}^P$. Without the loss of
 37 the generality, the optimal threshold τ should be selected by maximizing H_{cross}^G and
 38 H_{cross}^P . Meanwhile, in the case of equivalent numbers and weights of edges, the greater
 39 the difference of network topologies between $\mathcal{G}_{GPN}^\tau$ and $\mathcal{G}_{GPN}^{random}$, the more information
 40 the $\mathcal{G}_{GPN}^\tau$ includes. Therefore, we also assess the difference of degree entropy between
 41 $\mathcal{G}_{GPN}^\tau$ and $\mathcal{G}_{GPN}^{random}$ for $\tau \in [0,1]$, which are defined as $\Delta H^G = H_\tau^G - H_\tau^{G,random}$ and
 42 $\Delta H^P = H_\tau^P - H_\tau^{P,random}$. To investigate the significance of differences and the stability of
 43 the cross entropies, we construct 1,000 random networks corresponding to $\mathcal{G}_{GPN}^\tau$. For
 44 network degree entropy, we evaluate their distributions of the random network, then
 45 compare $H_\tau^G(H_\tau^P)$ and the upper bound of the 95% confidence interval (CI) of
 46 $H_\tau^{G,random}(H_\tau^{P,random})$. For cross-entropy, we can estimate the standard errors of them and
 47 then obtain the stability by computing their 95% CIs.

48

Text S2. Details of the four multiple phenotype association tests

In this study, we apply four powerful GWAS summary-based multiple phenotype association tests to identify the association between phenotypes in each network module and a genetic variant, including minP(2), ACAT(3), MTAG(4), SHom(5). To simplify the notation, we assume that the tests are applied to test the association between K phenotypes and a genetic variant.

minP. Consider the z-score vector is $\mathbf{Z} = (Z_1, \dots, Z_K)^T$, where Z_k is the z-score for testing the association between the k^{th} phenotype and a genetic variant. Assume that $\mathbf{Z}$ is asymptotically multivariate normal $MVN(\mathbf{0}, \mathbf{R})$ under the null hypothesis that K phenotypes and a genetic variant have no association, where $\mathbf{R}$ is the correlation matrix of phenotypes. The minP test statistic is given by $T_{minP} = \max_{k=1, \dots, K} \{|Z_k|\}$ and the corresponding p-value can be calculated by

$$p_{minP} = 1 - \int_{-T_{minP}}^{T_{minP}} \dots \int_{-T_{minP}}^{T_{minP}} f(x_1, \dots, x_K; \mathbf{0}, \hat{\mathbf{R}}) dx_1 \dots dx_K, \text{ where } f(x_1, \dots, x_K; \mathbf{0}, \hat{\mathbf{R}}) \text{ is the density function for } MVN(\mathbf{0}, \hat{\mathbf{R}}) \text{ and } \hat{\mathbf{R}} \text{ can be estimated by using 'estcov' function in aSPU package.}$$

ACAT. Let $p_1, \dots, p_K$ be the p-values of $Z_1, \dots, Z_K$, respectively. The ACAT test statistic is $T_{ACAT} = \sum_{k=1}^K \tan\{(.5 - p_k)\pi\}/K$ and the p-value of T_{ACAT} is approximated by

$$p_{ACAT} \approx 0.5 - \arctan\{T_{ACAT}\}/\pi.$$

MTAG. Let $\hat{\beta}_{MTAG} = (\hat{\beta}_{MTAG,1}, \dots, \hat{\beta}_{MTAG,K})^T$ be the vector of the MTAG estimator correcting for both genetic correlation $\mathbf{\Omega}$ and phenotypic correlation $\mathbf{\Sigma}$ among K phenotypes. $\hat{\mathbf{\Omega}}$

69 can be estimated by the method of moments using the moment condition, and $\hat{\Sigma}$ can
70 be estimated by LD score regression (LDSC)(6). Let $p_{MTAG,1}, \dots, p_{MTAG,K}$ be the
71 corresponding p-value of MTAG estimator. Then, a Bonferroni correction is used to
72 adjust for multiple testing for K phenotypes.

73 **SHom**. The SHom test statistic is given by $T_{SHom} = (\mathbf{1}^T \hat{\Sigma}^{-1} \mathbf{Z})(\mathbf{1}^T \hat{\Sigma}^{-1} \mathbf{Z})^T / (\mathbf{1}^T \hat{\Sigma}^{-1} \mathbf{1})$, where
74 $\mathbf{1} = (1, \dots, 1)^T$ is a $K \times 1$ vector that contains all 1s and $\hat{\Sigma}$ can be estimated by LDSC. The
75 p-value of T_{SHom} is calculated by assuming T_{SHom} follows a chi-square distribution with 1
76 degree of freedom.

77

Text S3. Simulation studies for PheWAS

In this study, we expand the application of constructing the bipartite GPN and unipartite PPN to phenome-wide association studies (PheWAS). In PheWAS, correcting for multiple testing is crucial to reduce the risk of false positives and ensure the reliability of the results. Therefore, applying the community detection method for GPN and PPN can obtain a prior grouping of phenotypes based on the shared genetic architecture. Then, jointly testing multiple phenotypes in each network module and one genetic variant can discover the cross-phenotype associations and pleiotropy. Finally, significance thresholds for PheWAS are adjusted for multiple testing by applying the refined false discovery rate (FDR) control approach. We conduct comprehensive simulations to evaluate the FDR of PheWAS based on network modules detected by GPN.

We directly generate a z-score matrix, $\mathbf{Z}$, for M genetic variants and K phenotypes in the whole phenome ($K = 500$ and $1,000$ in our simulation studies). Suppose there are L phenotypic categories and $k = K/L$ in each category. Let $\mathbf{\Sigma}$ be the phenotypic correlation matrix, where $\mathbf{\Sigma} = Bdiag(\mathbf{\Sigma}_1, \dots, \mathbf{\Sigma}_L)$ is a block diagonal matrix, that is, the phenotypes within a category are correlated and between categories are uncorrelated. We consider two scenarios of $\mathbf{\Sigma}$ (7), SAME and DIFF. In the SAME scenario, there is the same correlation coefficient of each pair of phenotypes within the category, that is, the off-diagonal elements of $\mathbf{\Sigma}_l$ equal ρ . In the DIFF scenario, the correlation coefficients are different and $\mathbf{\Sigma}_l$ is generated by using an autoregressive (AR(1)) model, that is, $\Sigma_l = \rho^{|k-l|}$. We use $\rho = 0.3$ in the simulation studies.

Assume S_{causal} is the set of M_{causal} true causal variants. Then, we generate a z-score vector for the m^{th} ($m = 1, \dots, M$) genetic variant from

$$\mathbf{Z}_m = (Z_{m1}, \dots, Z_{mK})^T \sim \begin{cases} MVN(\mathbf{0}, \mathbf{\Sigma}), & \text{for } m \notin S_{causal} \\ MVN(\boldsymbol{\mu}_m, \mathbf{\Sigma}), & \text{for } m \in S_{causal} \end{cases},$$

where $\boldsymbol{\mu}_m = (\mu_{m1}, \dots, \mu_{mK})^T$ is the K dimensional vector of the true effects for causal variants. In the simulation studies, we consider a total of $M = 10^6$ genetic variants and define the first $M_{causal} = 100$ as the true causal variants. Based on the different numbers of phenotypic categories $L = 50$ and 100 , we consider the following two models to define $\boldsymbol{\mu}_m$. Let $\boldsymbol{\delta}_1 = \delta(1, \dots, 1)^T$ and $\boldsymbol{\delta}_2 = \frac{2\delta}{k/2} \left(1, \dots, \frac{k}{2}, \frac{k}{2}, \dots, 1\right)^T$ be two k dimensional vectors of true effect sizes.

Model 1. Only the phenotypes in the first two categories are associated with at least one causal variant with the same effect sizes but in different directions. That is, the first 50 causal variants impact the phenotypes in the first category with $\boldsymbol{\delta}_1$; the second 50 causal variants impact the phenotypes in the second category with $-\boldsymbol{\delta}_1$.

Model 2. Only the phenotypes in the first four categories are associated with at least one causal variant with different effect sizes and different directions. That is, the first 25 causal variants impact the phenotypes in the first category with $\boldsymbol{\delta}_1$; the second 25 causal variants impact the phenotypes in the second category with $-\boldsymbol{\delta}_1$; the third 25 causal variants impact the phenotypes in the third category with $\boldsymbol{\delta}_2$; the fourth 25 causal variants impact the phenotypes in the fourth category with $-\boldsymbol{\delta}_2$.

For each simulation model, we run B Monte-Carlo (MC) runs and use the following steps for the b^{th} MC run: i) generate a z-score matrix using the above simulation models; ii) construct the bipartite GPN using the method introduced in section 2.2.1; iii) detect $L^{(b)}$ network modules of K phenotypes using the method introduced in section 2.2.4; iv) test the association between phenotypes in each of $L^{(b)}$ network modules and each and each of M genetic variants using one of the multiple phenotype association tests (**Text B.2**), obtaining $p_{ml}^{(b)}$; v) calculate the optimal threshold, $\hat{p}_m^{(b)}$, by applying the refined FDR controlling approach.

Let $D_m^{(b)} = \sum_{l=1}^{L^{(b)}} \mathbf{I}(p_{ml}^{(b)} \leq \hat{p}_m^{(b)})$ be the total number of discoveries. Then, we define the true discoveries and false discoveries as $TD_m^{(b)} = \sum_{l \in L_a} \mathbf{I}(p_{ml}^{(b)} \leq \hat{p}_m^{(b)})$ and $FD_m^{(b)} = D_m^{(b)} - TD_m^{(b)}$, respectively. L_a is the set of network modules containing at least one phenotype that is associated with the m^{th} genetic variant. Therefore, the average FDR can be computed by

$$FDR = \frac{1}{B \times M_{causal}} \sum_{b=1}^B \sum_{m=1}^{M_{causal}} \frac{FD_m^{(b)}}{\max\{D_m^{(b)}, 1\}}.$$

Note that we do not generate linkage disequilibrium (LD) of genetic variants in our simulation studies, therefore, we can not use LDSC(6) to estimate the phenotypic correlation matrix Σ in applications of SHom. We use the same estimation method introduced in minP and Chisq(2) to approximately estimate Σ in the simulation studies. Meanwhile, we directly generate Z-scores instead of effect sizes of genetic variants on phenotypes, therefore, we do not consider MTAG in our simulation studies.

Figure S1. Network connectance of GPN with different thresholds for (a) 12 genetically correlated phenotypes and (b) 588 EHR-derived phenotypes in the UK Biobank.

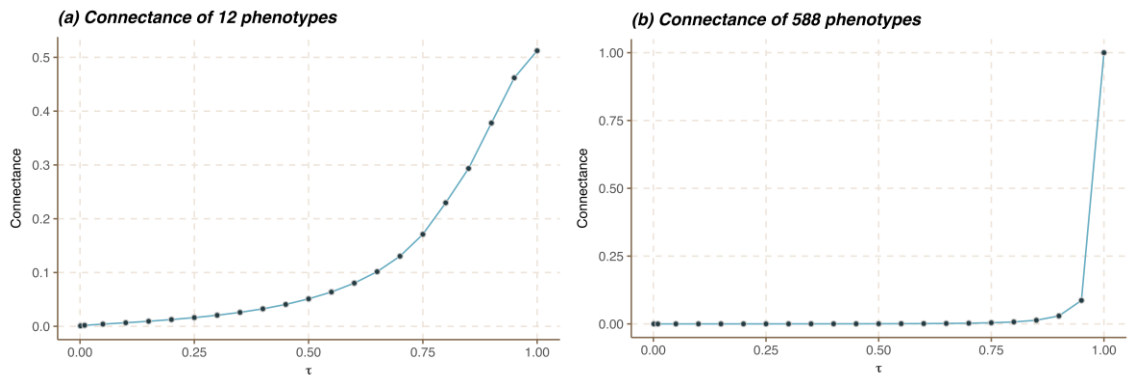


Figure S2. Network properties of the unweighted bipartite GPNs for 12 genetically correlated phenotypes. (a) KL divergence for genetic variants. The blue line is the mean of KL divergencies across 1,000 random network comparisons. The boxplots show the scaled distribution of KL divergence for each threshold. (b) Cross entropy for genetic variants. Blue lines are the means of cross entropy across 1,000 random network comparisons. The boxplots show the scaled distribution of cross entropy for each threshold. Red lines represent the degree entropy for the original network. The boxplots show the distribution of degree entropy for each threshold across 1,000 random networks. The blue line represents the difference between the original and random networks. (c) Unweighted degree entropy for genetic variants. (d) plot of the unweighted degree distribution of genetic variants for three GPNs on the log-log scale, denser representation ($\tau = 1$), well-defined sparse representation ($\tau = 0.45$), and an arbitrary threshold sparse representation ($\tau = 0.1$).

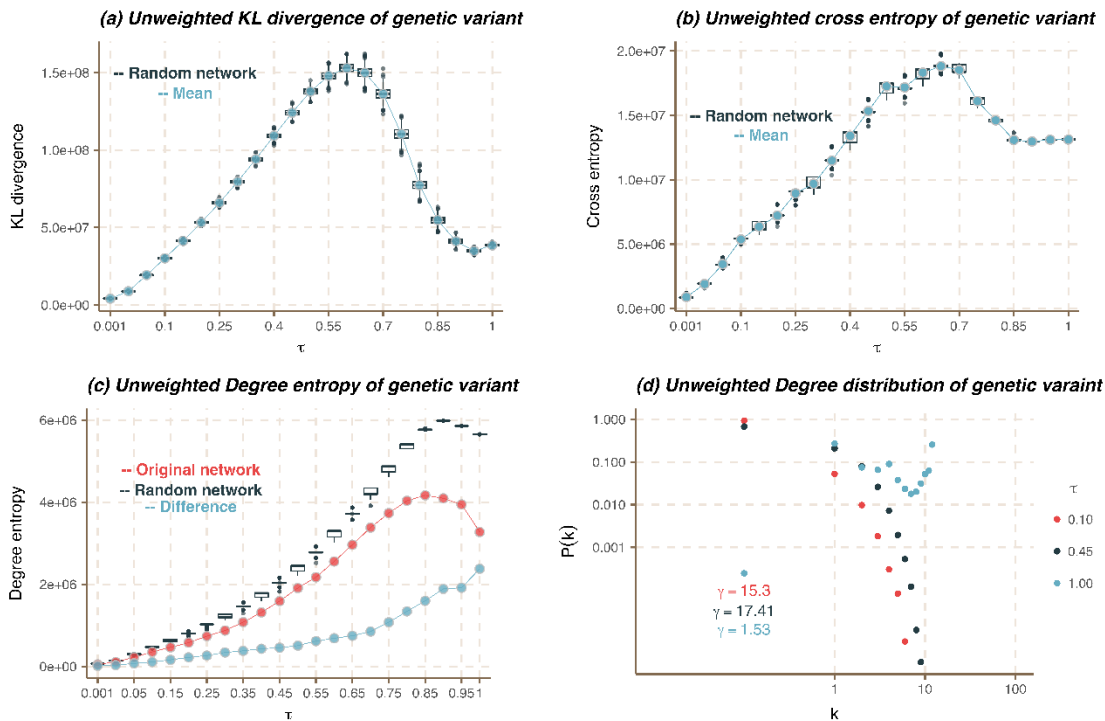


Figure S3. The correlation of 12 highly correlated phenotypes calculated by different methods: (a) Adjacency matrix of Phenotype and Phenotype Network (PPN) projected from the denser representation of the bipartite GPN; (b) Adjacency matrix of PPN from the well-defined sparse representation of GPN; (c) Genetic correlation matrix estimated by LDSC; (d) Global genetic correlation estimated by SUPERGENOVA.

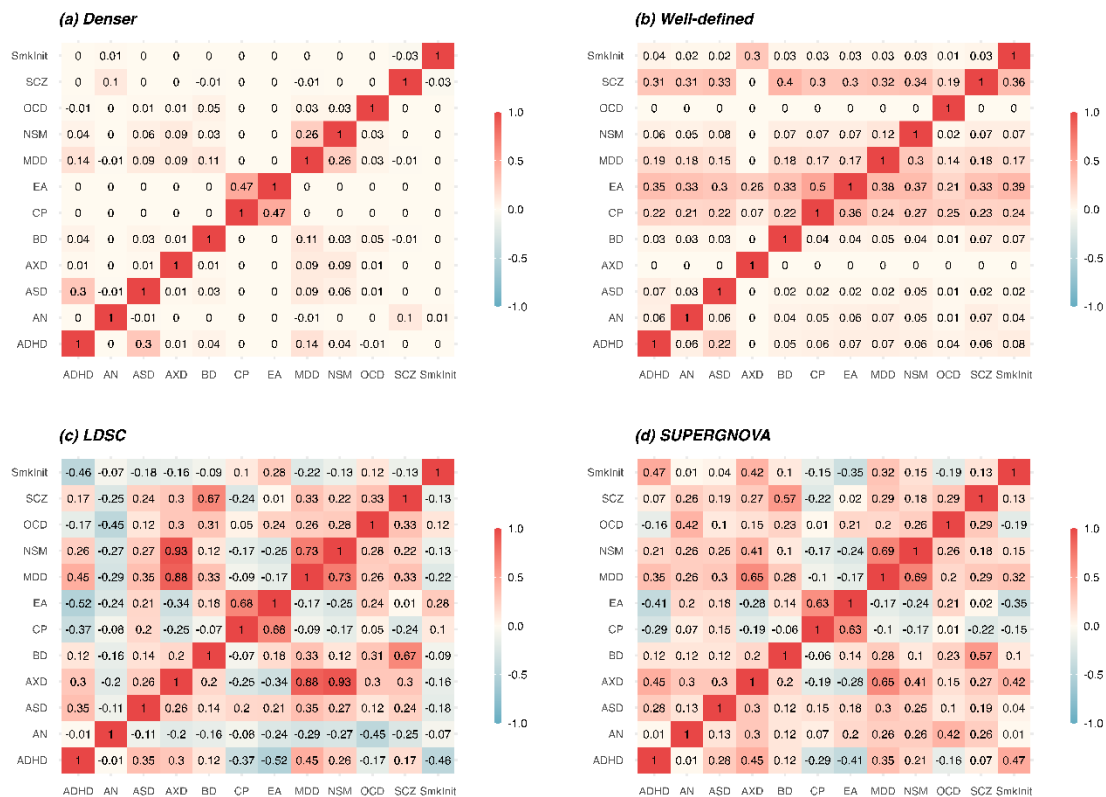
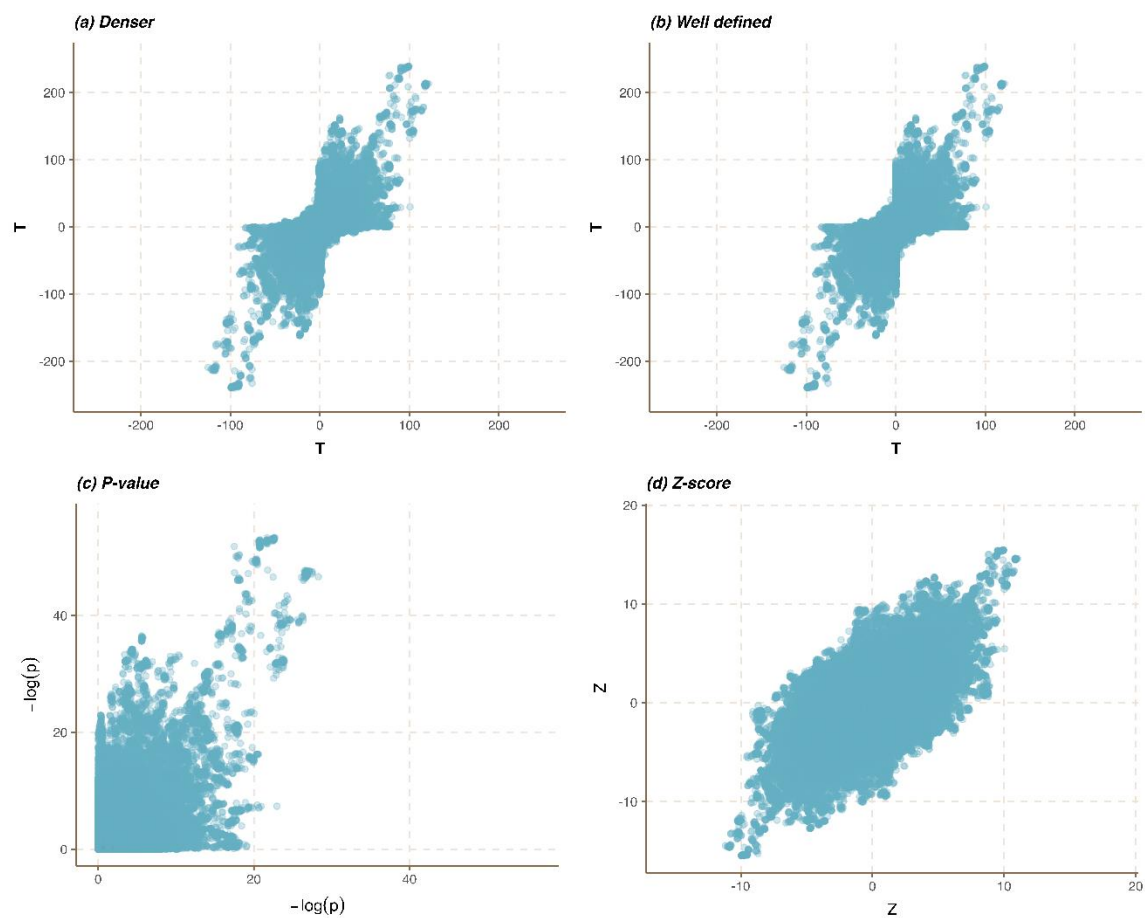


Figure S4. The qq-plot of EA versus CP based on the weight of (a) the denser representation of GPN and (b) the well-defined sparse representation of GPN. The qq-plot of EA versus CP based on (c) $-\log_{10}(\text{p-values})$ and (d) z-scores from GWAS summaries.



170 **Figure S5.** Heatmap of edge weights in the well-defined sparse representation of GPN
171 for (a) the top 100 and (b) the top 1000 genetic variants with the highest degree
172 centrality.

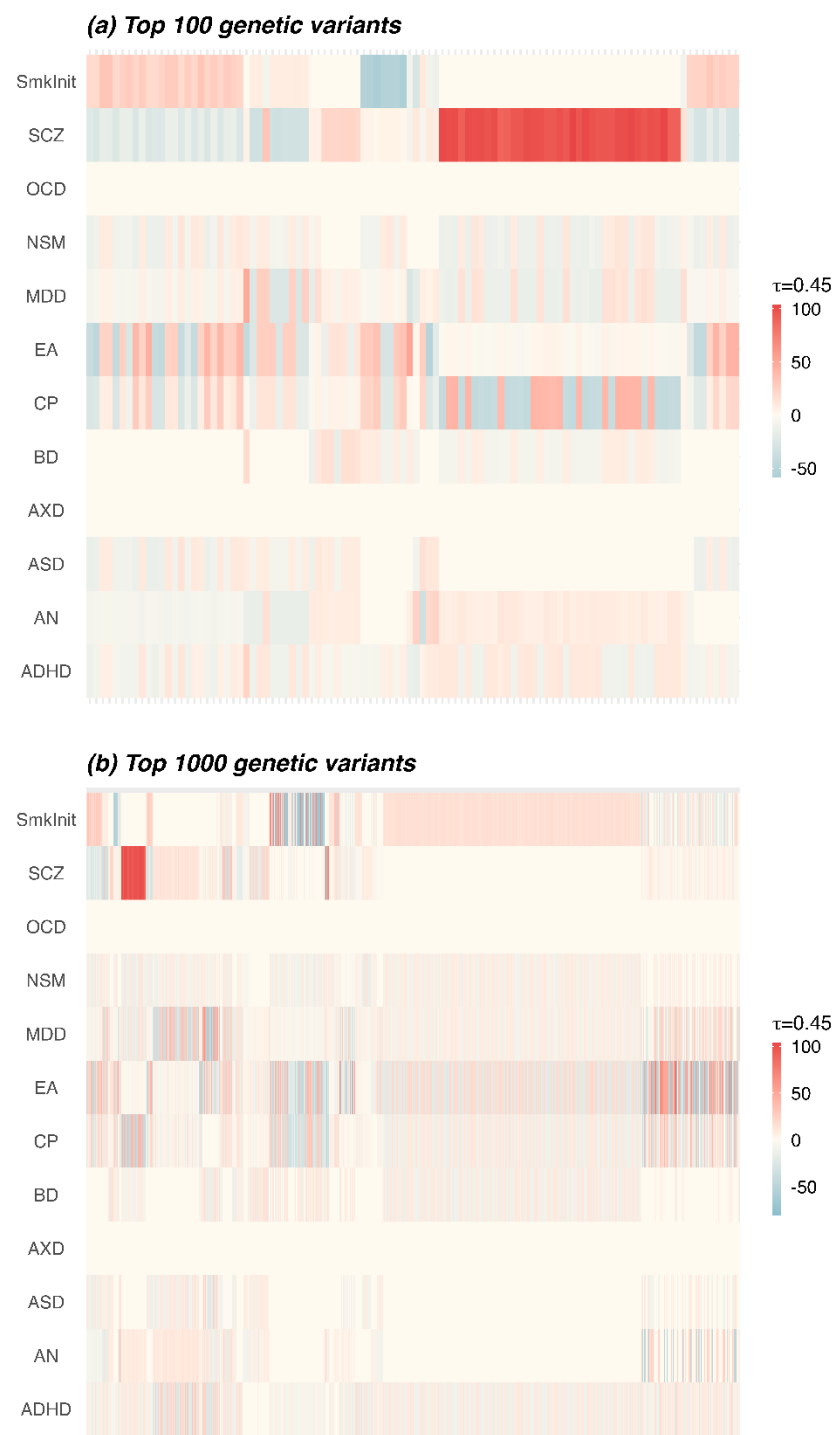
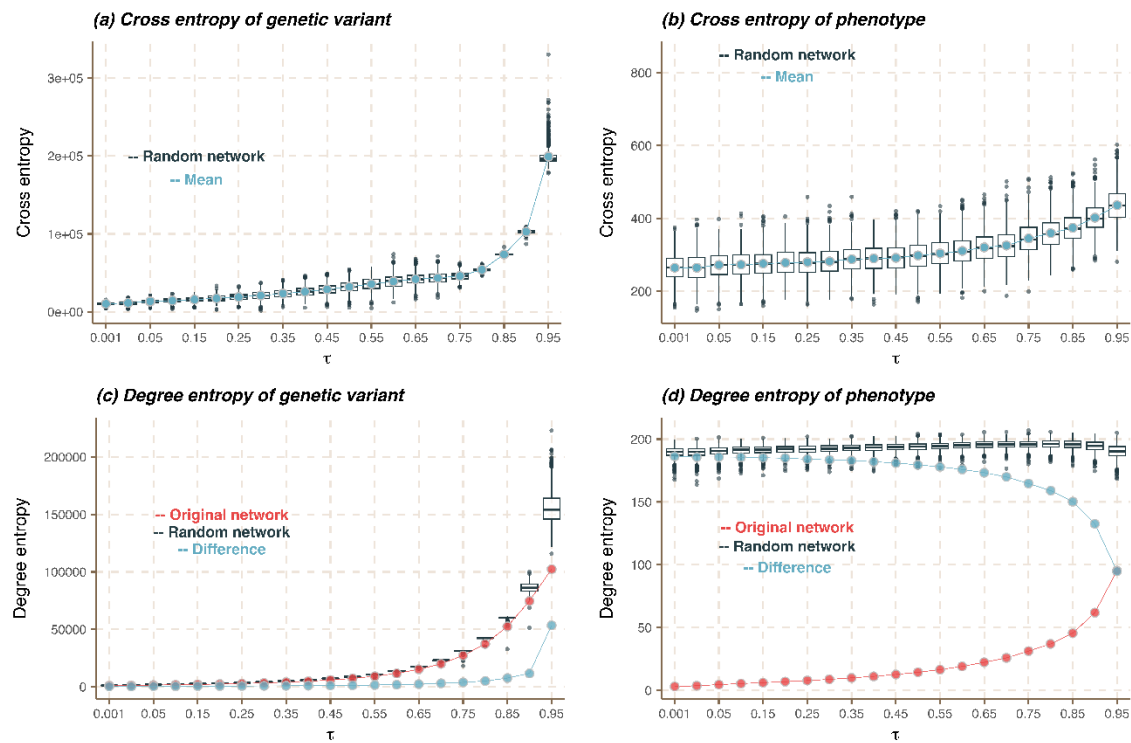
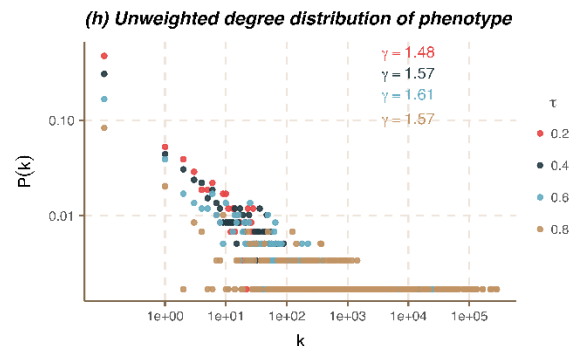
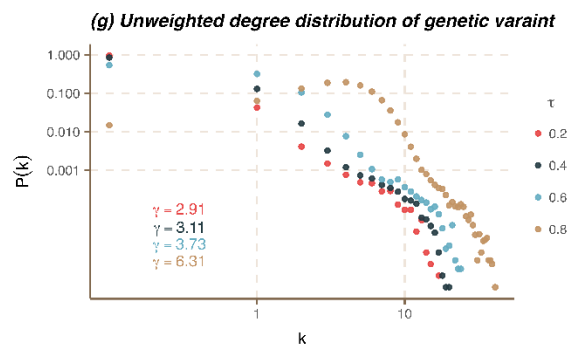
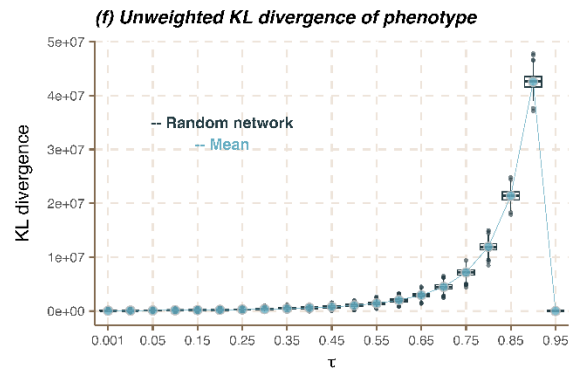
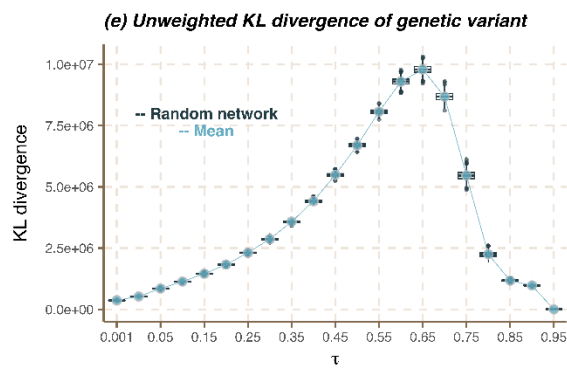
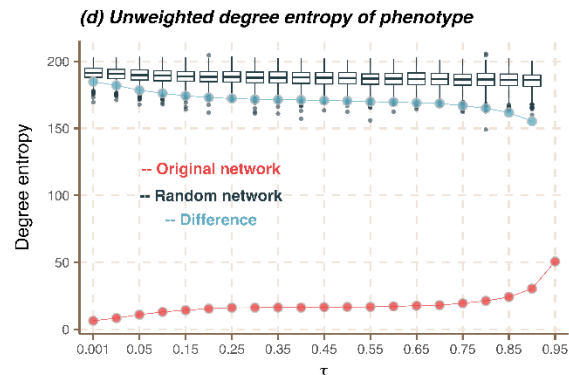
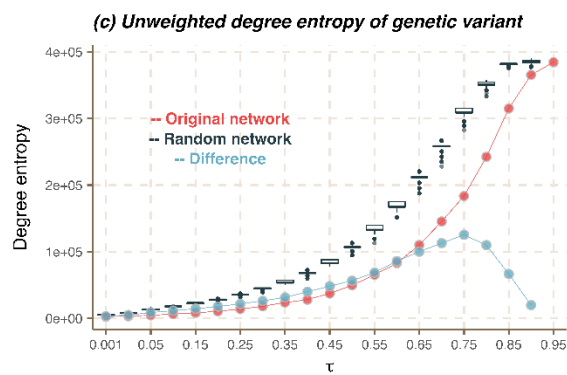
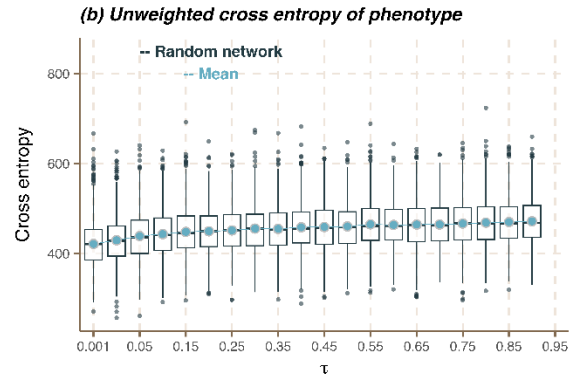
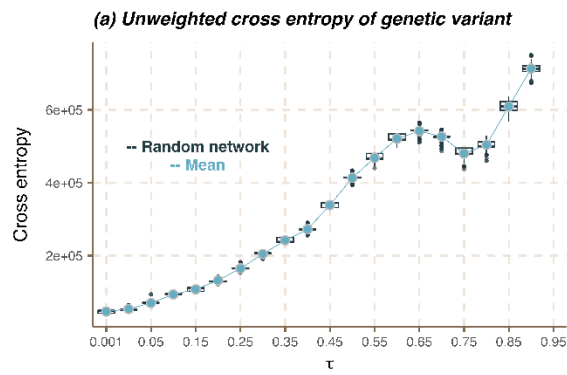


Figure S6. Cross entropy and degree entropy the unweighted bipartite GPNs for 588 EHR-derived phenotypes in the UK Biobank. Cross entropy for (a) genetic variants and (b) phenotypes. The blue line is the mean of cross entropy across 1,000 random network comparisons. The boxplots show the scaled distribution of cross entropy for each threshold. Degree entropy for (a) genetic variants and (b) phenotypes. Red lines represent the degree entropy for the original network. The boxplots show the distribution of degree entropy for each threshold across 1,000 random networks. The blue line represents the difference between the original and random networks.



184 **Figure S7.** Network properties of the unweighted bipartite GPNs for 588 EHR-derived
185 phenotypes in the UK Biobank. (a) and (b) KL divergence for genetic variants and
186 phenotypes. Blue line is the mean of KL divergencies across 1,000 random network
187 comparisons. The boxplots show the scaled distribution of KL divergence for each
188 threshold. (c) and (d) Cross entropy for genetic variants and phenotypes. Blue line is the
189 mean of cross entropy across 1,000 random network comparisons. Boxplots show the
190 scaled distribution of cross entropy for each threshold. (e) and (f) Unweighted degree
191 entropy for genetic variants and phenotypes. The red line represents the degree entropy
192 for the original network. The boxplots show the distribution of degree entropy for each
193 threshold across 1,000 random networks. The blue line represents the difference
194 between the original and random networks. (g) and (h) Unweighted degree distribution
195 of genetic variants and phenotypes for the four GPNs, much more denser
196 representation ($\tau = 0.8$), well-defined sparse representation ($\tau = 0.6$), and two arbitrary
197 threshold sparse representations ($\tau = 0.2$ and $\tau = 0.4$).

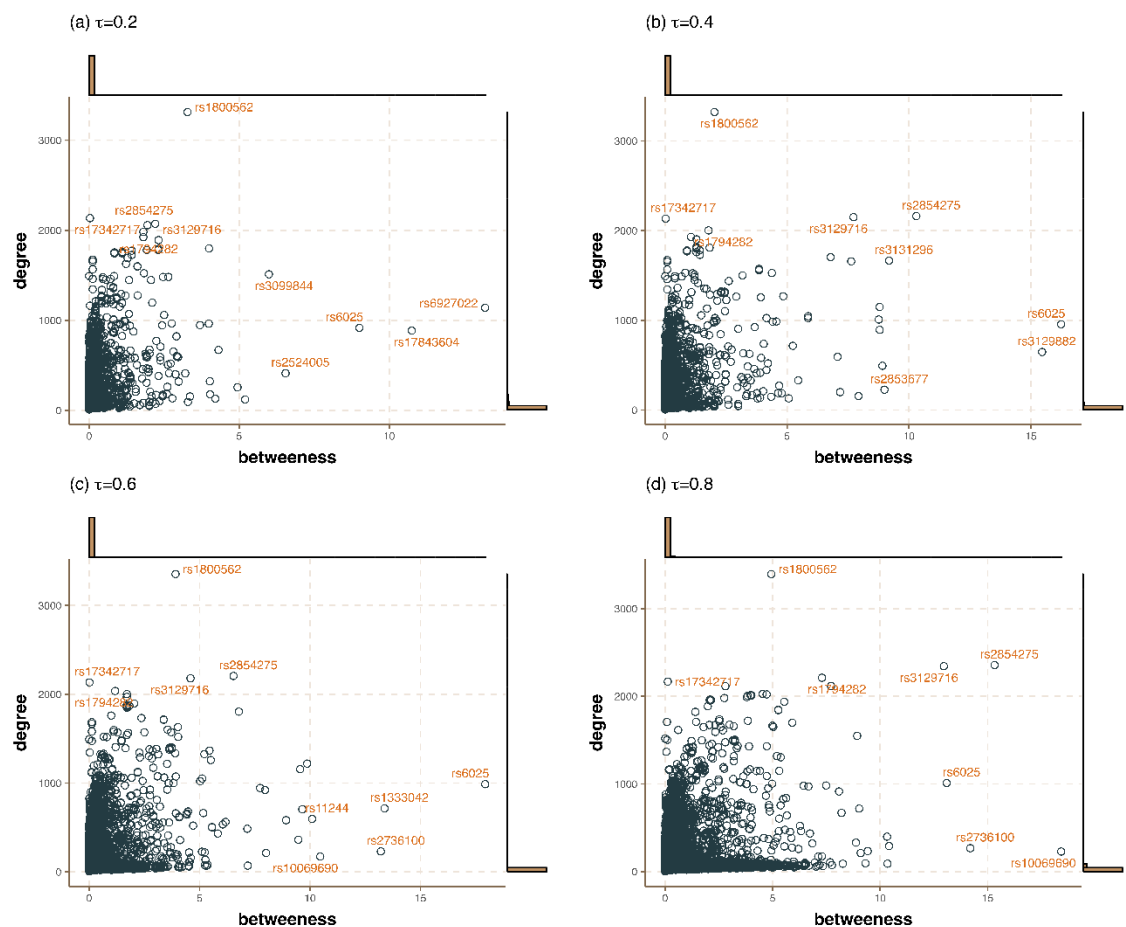


198

199

200

Figure S8. Degree centrality and betweenness centrality of genetic variants of the weighted bipartite GPNs for 588 EHR-derived phenotypes in the UK Biobank.



205 **Figure S9.** Degree centrality and betweenness centrality of genetic variants of the
206 unweighted bipartite GPNs for 588 EHR-derived phenotypes in the UK Biobank.

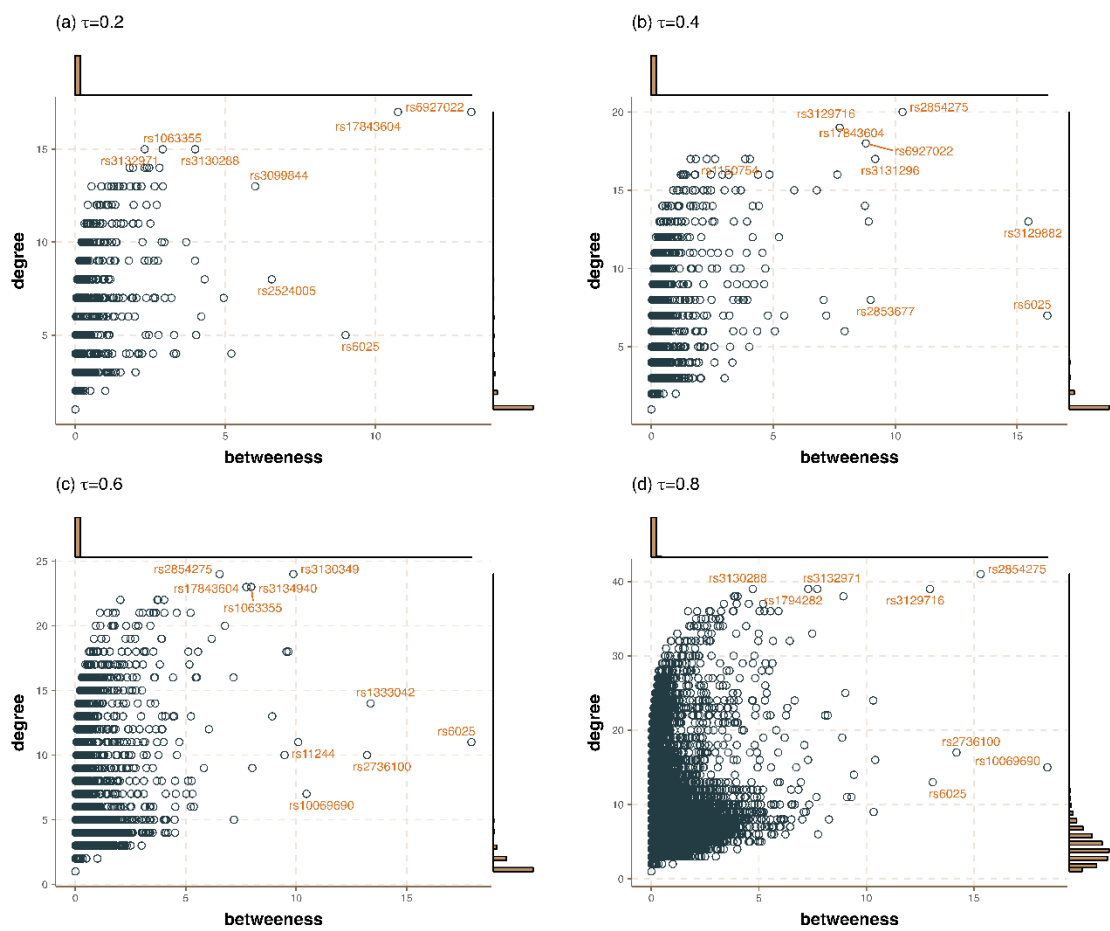


Figure S10. Heatmap of $-\log_{10}(\text{p-value})$ of 100 genetic variants from GWAS summary datasets, which are uniquely identified by ACAT based on LDSC compared with ACAT based on the UK Biobank. Only the p-values smaller than the GWAS significance level (5×10^{-8}) shown in the heatmap.

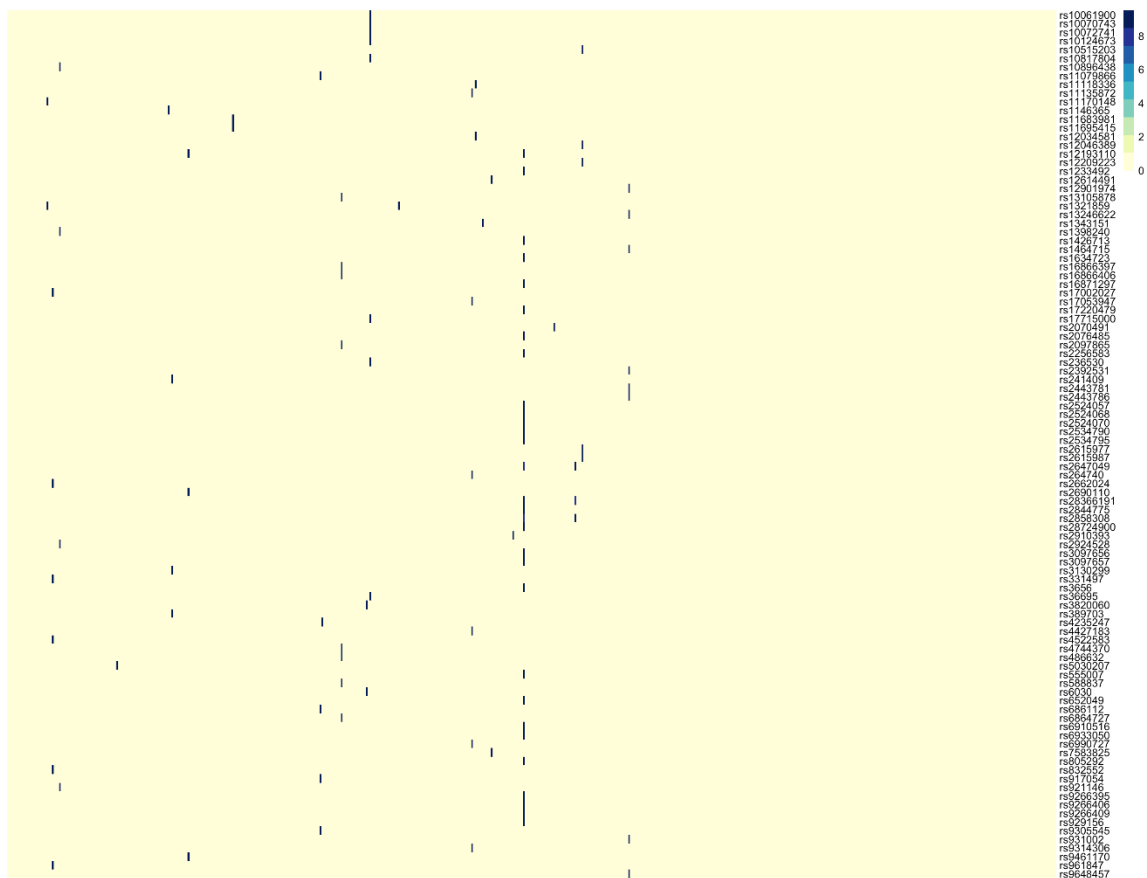


Table S1. Phenotypes, abbreviations, samples sizes, disease heritability, and GWAS resources used in heritability enrichment analyses.

Phenotype	Abbreviation	Sample size	Heritability	Reference
Attention deficit/hyperactivity disorder	ADHD	53,293	0.2354 (0.0153)	Demontis et al.(8)
Smoking initiation	SmkInit	632,802	0.0724 (0.0068)	Liu et al.(9)
Autism spectrum disorder	ASD	46,351	0.1941 (0.0168)	Grove et al.(10)
Neuroticism	NSM	170,911	0.0877 (0.0067)	Okbay et al.(11)
Anxiety disorder	AXD	31,890	0.0417 (0.0156)	Meier et al.(12)
Major depressive disorder	MDD	500,199	0.0599 (0.0023)	Howard et al.(13)
Obsessive-compulsive disorder	OCD	9,725	0.3217 (0.0496)	Arnold et al.(14)
Anorexia nervosa	AN	72,517	0.1773 (0.0116)	Watson et al.(15)
Bipolar disorder	BD	51,710	0.3469 (0.0174)	Stahl et al.(16)
Schizophrenia	SCZ	105,318	0.4101 (0.0113)	Pardinas et al.(17)
Educational attainment	EA	766,345	0.1066 (0.0026)	Lee et al.(18)
Cognitive performance	CP	257,828	0.192 (0.0062)	Lee et al.(18)

Notes: Heritability is calculated by LD score regression(6): heritability (standard error of heritability).

Table S2. Global genetic correlations (right upper triangle) and proportions of correlated regions (left lower triangle) estimated by SUPERGNOVA(19).

	ADHD	Smklnit	ASD	NSM	AXD	MDD	OCD	AN	BD	SCZ	EA	CP
ADHD		0.47*	0.28*	0.21*	0.45*	0.35*	-0.16	0.008	0.12*	0.07	-0.41*	-0.29*
Smklnit	81%*		0.04	0.15*	0.42*	0.32*	-0.19*	0.01	0.10*	0.13*	-0.35*	-0.15*
ASD	32%*	1.33%		0.25*	0.30*	0.30*	0.10	0.13	0.12	0.19*	0.18*	0.15*
NSM	41%	46%*	15%*		0.41*	0.69*	0.26*	0.26*	0.10	0.18*	-0.24*	-0.17*
AXD	52%	63%*	41%	69%		0.65*	0.15	0.30*	0.20*	0.27*	-0.28*	-0.19*
MDD	69%	65%*	51%*	89%*	77%*		0.20*	0.26*	0.28*	0.29*	-0.17*	-0.10
OCD	0.25%	15%*	0.32%	2.3%*	0.11%*	6.4%*		0.42*	0.23*	0.29*	0.21*	0.01
AN	0.26%	1.52%*	0.19%	54%*	30%*	53%*	37%		0.12*	0.26*	0.20*	0.07
BD	1.37%	2.72%*	4.53%*	1.11%	35%	55%*	2%	21%		0.57*	0.14*	-0.06
SCZ	12%*	39%*	35%*	35%*	42%*	60%*	33%*	58%*	83%*		0.02	-0.22*
EA	76%*	79%*	33%*	61%	51%*	38%*	31%	47%*	32%*	7.5%*		0.63*
CP	67%*	39%*	15%*	38%	20%*	27%*	0.5%	1.9%*	3.4%*	57%*	93%*	

Notes: * indicates the significance of genetic correlations and proportions of correlated regions between two phenotypes.

228 **Table S3.** Phenotypic correlations (right upper triangle) and genetic correlations (left
229 lower triangle) estimated by LDSC(6).

	ADHD	SmkInit	ASD	NSM	AXD	MDD	OCD	AN	BD	SCZ	EA	CP
ADHD		0.0034	0.3626	0.0092	0.0040	0.0860	0.0082	-0.1235	0.0398	0.0289	-0.0165	-0.0040
SmkInit	-0.4628*		0.0033	-0.0125	-0.0039	-0.0167	-0.0080	-0.0081	0.0058	0.0055	0.0332	0.0115
ASD	0.3459*	-0.1819*		-0.0125	-0.0024	0.0598	-0.0003	-0.1226	0.0134	0.0170	-0.0003	-0.0013
NSM	0.2642*	-0.1342*	0.2723*		0.0514	0.1211	0.0049	-0.0026	0.0069	0.0036	-0.0403	-0.0212
AXD	0.3008	-0.1579	0.2607	0.9339*		0.0342	0.0027	-0.0108	0.0078	0.0144	-0.0067	-0.0030
MDD	0.4537*	-0.2187*	0.3526*	0.7313*	0.8790*		0.0062	-0.0627	0.0594	0.0375	-0.0146	-0.0111
OCD	-0.1695	0.1185	0.1185	0.2821*	0.2989	0.2591*		-0.0127	0.0376	0.0331	-0.0080	-0.0070
AN	-0.0082	-0.0668	-0.1057	-0.2671*	-0.1976	-0.2870*	-0.4490*		-0.0536	-0.0321	0.0147	0.0116
BD	0.1205	-0.0873	0.1373	0.1213	0.2000	0.3320*	0.3106*	-0.1592*		0.1898	0.0021	0.0119
SCZ	0.1679*	-0.1335*	0.2379*	0.2164*	0.3044*	0.3301*	0.3318*	-0.2527*	0.6667*		-0.0054	-0.0066
EA	-0.5159*	0.2825*	0.2081*	-0.2512*	-0.3416*	-0.1734*	0.2390*	-0.2380*	0.1820*	0.0106		0.1681
CP	-0.3677*	0.0992*	0.2002*	-0.1677*	-0.2459	-0.0866*	0.0456	-0.0819	-0.0701	-0.2438*	0.6840*	

230
231 Notes: * indicates the significance genetic correlations between two phenotypes.

232

Table S4. Heritability enrichment analyses of network topology annotation (betweenness centrality) calculated from denser and sparse representations of bipartite GPN for each of 12 phenotypes.

Trait	Denser		Sparse ($\tau = 0.45$)		Sparse ($\tau = 0.1$)	
	Enrichment (Standard error) <i>p</i> -value	Effect τ^* (<i>se</i> (τ^*)) z-score	Enrichment (Standard error) <i>p</i> -value	Effect τ^* (<i>se</i> (τ^*)) z-score	Enrichment (Standard error) <i>p</i> -value	Effect τ^* (<i>se</i> (τ^*)) z-score
ADHD	1.1065 (0.0397) 0.0088	19.9545 (14.6723) 1.3600	1197.52 (134.831) 3.18e-23	10.0116 (0.8860) 11.3003	498.39 (76.248) 1.10e-10	6.9712 (1.0241) 6.8073
AN	1.1342 (0.0450) 0.0019	13.4534 (10.1444) 1.3262	670.61 (106.029) 5.40e-11	4.1649 (0.6006) 6.9350	292.19 (57.763) 4.27e-07	3.1944 (0.6120) 5.2198
ASD	1.0499 (0.1275) 0.7047	-10.6735 (34.4418) -0.3101	1284.50 (216.018) 1.47e-16	7.4179 (0.8219) 9.0256	129.45 (28.518) 5.34e-07	3.9185 (0.7609) 5.1501
AXD	1.3694 (0.3546) 0.0827	12.5210 (10.8825) 1.1506	85.3863 (83.1066) 0.1246	0.3191 (0.2555) 1.5317	187.30 (158.517) 0.0387	0.4038 (0.1944) 2.0767
BD	1.1634 (0.0258) 5.94e-14	40.1394 (6.5713) 6.1082	1005.12 (109.477) 3.47e-20	10.8984 (1.0510) 10.2816	500.87 (65.368) 3.70e-14	10.0078 (1.2369) 8.1479
CP	1.0888 (0.0143) 1.94e-09	9.6528 (4.1834) 2.3074	863.223 (55.5744) 7.68e-27	8.3341 (0.6674) 12.4864	744.55 (101.854) 3.04e-08	7.1091 (1.2339) 5.7613
EA	1.0876 (0.0096) 2.20e-17	6.4808 (1.5985) 4.0543	991.484 (58.1901) 5.39e-29	5.6591 (0.4291) 13.1869	738.32 (102.370) 1.31e-07	4.1130 (0.7518) 5.4712
MDD	1.1198 (0.0144) 7.81e-16	5.2452 (1.0707) 4.8990	1345.72 (93.838) 6.24e-25	2.6991 (0.2275) 11.8632	624.41 (73.576) 2.85e-12	1.9156 (0.2575) 7.4392
NSM	1.0392 (0.0692) 0.5804	-2.0694 (11.3005) -0.1832	1030.86 (110.868) 4.37e-16	3.0760 (0.3472) 8.8586	730.19 (88.155) 1.65e-11	3.4041 (0.4767) 7.1413
OCD	1.1530 (0.1135) 0.0942	40.7436 (37.4199) 1.0888	236.746 (118.689) 0.0141	3.6566 (1.4789) 2.4725	52.57 (46.183) 0.2173	1.0209 (0.8355) 1.2218
SCZ	1.1640 (0.0198) 7.51e-16	60.5467 (10.4817) 5.7764	1275.95 (86.122) 3.33e-27	18.4038 (1.4660) 12.6051	624.72 (85.240) 1.95e-09	13.0659 (2.0790) 6.2848
Smklnit	1.0719 (0.0221) 9.78e-05	3.7899 (1.6834) 2.2514	568.13 (88.322) 5.35e-12	1.8560 (0.2193) 8.4633	205.23 (45.469) 1.61e-07	1.5744 (0.2906) 5.4185

Notes: The betweenness are scaled by multiplying the number of phenotypes and the number of genetic variants due to it being much smaller than the baseline LD annotations. The estimated effect size and its estimated standard error, τ^* and *se*(τ^*), are scaled by dividing 10^{-12} . Z-score of the effect size is reported to test the null hypothesis that either $\tau \leq 0$ (one-sided) or $\tau = 0$ (two-sided). P-value of enrichment is reported to test the null hypothesis that *Enrichment* > 1. The bold-faced p-values indicate the annotation is significantly enriched in the disease heritability after accounting for multiple testing (p-value < 0.05/12 $\approx$ 0.0041).

Table S5. The number of associated phenotypes for the top 5 genetic variants with the highest approximate betweenness centrality and the weighted degree centrality, respectively, is calculated when $\tau = 0.6$.

	Genetic variants	Location	Mapped Gene	Associated phenotypes
Top 5 genetic variants with the highest weighted degree centrality	rs1800562	6:26092913	HFE	100
	rs17342717	6:25821542	SLC17A1	3
	rs1794282	6:32698749	HLA-DQB1, MTCO3P1	1
	rs3129716	6:32689659	HLA-DQB1, MTCO3P1	1
	rs2854275	6:32660651	HLA-DQB1, HLA-DQB1-AS1	2
Top 5 genetic variants with the highest approximate betweenness centrality	rs6025	1:169549811	F5	25
	rs1333042	9:22103814	CDKN2B-AS1	9
	rs10069690	5:1279675	TERT	54
	rs2736100	5:1286401	TERT	35
	rs11244	6:32812947	HLA-DOB	1

Notes: For instance, the genetic variant rs1800562, which is located on chromosome 6, has been associated with 100 phenotypes as reported in the GWAS Catalog.

Table S6. The number of associated phenotypes for the top 5 genetic variants with the highest approximate betweenness centrality and the unweighted degree centrality, respectively, is calculated when $\tau = 0.6$.

	Genetic variants	Location	Mapped Gene	Associated phenotypes
Top 5 genetic variants with the highest unweighted degree centrality	rs3130349	6:32179919	RNF5	3
	rs2854275	6:32660651	HLA-DQB1, HLA-DQB1-AS1	2
	rs1063355	6:32659937	HLA-DQB1, HLA-DQB1-AS1	2
	rs17843604	6:32652506	HLA-DQB1, HLA-DQA1	1
	rs3134940	6:32182039	AGER	1
Top 5 genetic variants with the highest approximate betweenness centrality	rs6025	1:169549811	F5	25
	rs1333042	9:22103814	CDKN2B-AS1	9
	rs10069690	5:1279675	TERT	54
	rs2736100	5:1286401	TERT	35
	rs11244	6:32812947	HLA-DOB	1

Table S7. The average FDR in the simulation studies for 500 phenotypes and 50 phenotypic categories.

Ce=SAME, Model 1				Ce=SAME, Model 2			
μ	minP	ACAT	SHom	μ	minP	ACAT	SHom
2.0	0.0846	0.0522	0.0415	1.5	0.0860	0.0462	0.0493
2.2	0.1203	0.0632	0.0407	2.0	0.1111	0.0583	0.0528
2.5	0.1040	0.0630	0.0527	2.5	0.1001	0.0512	0.0451
2.8	0.0959	0.0532	0.0493	3.0	0.1017	0.0521	0.0518
Ce=DIFF, Model 1				Ce=DIFF, Model 2			
μ	minP	ACAT	SHom	μ	minP	ACAT	SHom
1.3	0.1030	0.0527	0.0462	1.3	0.0947	0.0427	0.0505
1.5	0.0993	0.0483	0.0517	1.5	0.0844	0.0423	0.0522
1.7	0.1109	0.0622	0.0491	1.7	0.1080	0.0645	0.0453
1.9	0.0928	0.0477	0.0482	1.9	0.1011	0.0491	0.0427

Notes: FDR is evaluated using 10 MC runs, equivalent to 1,000 replications at a nominal FDR level of 5%. The 95% confidence interval (CI) is [0.0365, 0.0635] and bold-faced values indicate that the values are beyond the upper bounds of the 95% CI.

Table S8. The average FDR in the simulation studies for 1,000 phenotypes and 100 phenotypic categories.

Ce=SAME, Model 1				Ce=SAME, Model 2			
μ	minP	ACAT	SHom	μ	minP	ACAT	SHom
2.0	0.0982	0.0535	0.0496	1.5	0.0809	0.0452	0.0473
2.2	0.1143	0.0556	0.0482	2.0	0.1036	0.0475	0.0500
2.5	0.1102	0.0578	0.0501	2.5	0.1110	0.0576	0.0556
2.8	0.1024	0.0535	0.0481	3.0	0.1097	0.0535	0.0526
Ce=DIFF, Model 1				Ce=DIFF, Model 2			
μ	minP	ACAT	SHom	μ	minP	ACAT	SHom
1.3	0.1068	0.0535	0.0556	1.3	0.0997	0.0425	0.0396
1.5	0.0925	0.0503	0.0491	1.5	0.1002	0.0560	0.0431
1.7	0.0942	0.0460	0.0483	1.7	0.0914	0.0412	0.0483
1.9	0.0977	0.0470	0.0483	1.9	0.0998	0.0555	0.0515

Notes: FDR is evaluated using 10 MC runs, equivalent to 1,000 replications at a nominal FDR level of 5%. The 95% confidence interval (CI) is [0.0365, 0.0635] and bold-faced values indicate that the values are beyond the upper bounds of the 95% CI.

References

1. Freitas CG, Aquino AL, Ramos HS, Frery AC, Rosso OA. A detailed characterization of complex networks using Information Theory. *Scientific reports*. 2019;9(1):16689.
2. Kim J, Bai Y, Pan W. An adaptive association test for multiple phenotypes with GWAS summary statistics. *Genetic epidemiology*. 2015;39(8):651-63.
3. Liu Y, Chen S, Li Z, Morrison AC, Boerwinkle E, Lin X. ACAT: a fast and powerful p value combination method for rare-variant analysis in sequencing studies. *The American Journal of Human Genetics*. 2019;104(3):410-21.
4. Turley P, Walters RK, Maghzian O, Okbay A, Lee JJ, Fontana MA, et al. Multi-trait analysis of genome-wide association summary statistics using MTAG. *Nature genetics*. 2018;50(2):229-37.
5. Zhu X, Feng T, Tayo BO, Liang J, Young JH, Franceschini N, et al. Meta-analysis of correlated traits via summary statistics from GWASs with an application in hypertension. *The American Journal of Human Genetics*. 2015;96(1):21-36.
6. Bulik-Sullivan B, Finucane HK, Anttila V, Gusev A, Day FR, Loh P-R, et al. An atlas of genetic correlations across human diseases and traits. *Nature genetics*. 2015;47(11):1236-41.
7. Lee CH, Shi H, Pasaniuc B, Eskin E, Han B. PLEIO: a method to map and interpret pleiotropic loci with GWAS summary statistics. *The American Journal of Human Genetics*. 2021;108(1):36-48.
8. Demontis D, Walters RK, Martin J, Mattheisen M, Als TD, Agerbo E, et al. Discovery of the first genome-wide significant risk loci for attention deficit/hyperactivity disorder. *Nature genetics*. 2019;51(1):63-75.

- 295 9. Liu M, Jiang Y, Wedow R, Li Y, Brazel DM, Chen F, et al. Association studies of
296 up to 1.2 million individuals yield new insights into the genetic etiology of tobacco and
297 alcohol use. *Nature genetics*. 2019;51(2):237-44.
- 298 10. Grove J, Ripke S, Als TD, Mattheisen M, Walters R, Won H, et al. Common risk
299 variants identified in autism spectrum disorder. *bioRxiv*. 2017:224774.
- 300 11. Okbay A, Baselmans BM, De Neve J-E, Turley P, Nivard MG, Fontana MA, et al.
301 Genetic variants associated with subjective well-being, depressive symptoms, and
302 neuroticism identified through genome-wide analyses. *Nature genetics*. 2016;48(6):624-
303 33.
- 304 12. Meier SM, Trontti K, Purves KL, Als TD, Grove J, Laine M, et al. Genetic variants
305 associated with anxiety and stress-related disorders: a genome-wide association study
306 and mouse-model study. *JAMA psychiatry*. 2019;76(9):924-32.
- 307 13. Howard DM, Adams MJ, Clarke T-K, Hafferty JD, Gibson J, Shiri M, et al.
308 Genome-wide meta-analysis of depression identifies 102 independent variants and
309 highlights the importance of the prefrontal brain regions. *Nature neuroscience*.
310 2019;22(3):343-52.
- 311 14. Arnold PD, Askland KD, Barlassina C, Bellodi L, Bienvenu O, Black D, et al.
312 Revealing the complex genetic architecture of obsessive-compulsive disorder using
313 meta-analysis. *Molecular psychiatry*. 2018;23(5):1181-.
- 314 15. Watson HJ, Yilmaz Z, Thornton LM, Hübel C, Coleman JR, Gaspar HA, et al.
315 Genome-wide association study identifies eight risk loci and implicates metabo-
316 psychiatric origins for anorexia nervosa. *Nature genetics*. 2019;51(8):1207-14.

- 317 16. Stahl EA, Breen G, Forstner AJ, McQuillin A, Ripke S, Trubetskoy V, et al.
318 Genome-wide association study identifies 30 loci associated with bipolar disorder.
319 Nature genetics. 2019;51(5):793-803.
- 320 17. Pardiñas AF, Holmans P, Pocklington AJ, Escott-Price V, Ripke S, Carrera N, et
321 al. Common schizophrenia alleles are enriched in mutation-intolerant genes and in
322 regions under strong background selection. Nature genetics. 2018;50(3):381-9.
- 323 18. Lee JJ, Wedow R, Okbay A, Kong E, Maghzian O, Zacher M, et al. Gene
324 discovery and polygenic prediction from a 1.1-million-person GWAS of educational
325 attainment. Nature genetics. 2018;50(8):1112.
- 326 19. Zhang Y, Lu Q, Ye Y, Huang K, Liu W, Wu Y, et al. SUPERGNOVA: local genetic
327 correlation analysis reveals heterogeneous etiologic sharing of complex traits. Genome
328 biology. 2021;22:1-30.
- 329