

1
2 **Deep Neural Network Improves the Estimation of Polygenic Risk Scores for**
3 **Breast Cancer**

4
5 Adrien Badré¹, Li Zhang², Wellington Muchero³, Justin C. Reynolds¹, Chongle Pan^{1,2,*}
6

7 ¹School of Computer Science, University of Oklahoma, Norman, OK

8 ²Department of Microbiology and Plant Biology, University of Oklahoma, Norman, OK

9 ³BioScience Division, Oak Ridge National Laboratory, Oak Ridge, TN
10

11 *Correspondence: cpan@ou.edu
12
13

Abstract

Polygenic risk scores (PRS) estimate the genetic risk of an individual for a complex disease based on many genetic variants across the whole genome. In this study, we compared a series of computational models for estimation of breast cancer PRS. A deep neural network (DNN) was found to outperform alternative machine learning techniques and established statistical algorithms, including BLUP, BayesA and LDpred. In the test cohort with 50% prevalence, the Area Under the receiver operating characteristic Curve (AUC) were 67.4% for DNN, 64.2% for BLUP, 64.5% for BayesA, and 62.4% for LDpred. BLUP, BayesA, and LPpred all generated PRS that followed a normal distribution in the case population. However, the PRS generated by DNN in the case population followed a bi-modal distribution composed of two normal distributions with distinctly different means. This suggests that DNN was able to separate the case population into a high-genetic-risk case sub-population with an average PRS significantly higher than the control population and a normal-genetic-risk case sub-population with an average PRS similar to the control population. This allowed DNN to achieve 18.8% recall at 90% precision in the test cohort with 50% prevalence, which can be extrapolated to 65.4% recall at 20% precision in a general population with 12% prevalence. Interpretation of the DNN model identified salient variants that were assigned insignificant p-values by association studies, but were important for DNN prediction. These variants may be associated with the phenotype through non-linear relationships.

1 Introduction

2
3 Breast cancer is the second deadliest cancer for U.S. women. Approximately
4 one in eight women in the U.S. will develop invasive breast cancer over the course of
5 their lifetime [1]. Early detection of breast cancer is an effective strategy to reduce the
6 death rate. If breast cancer is detected in the localized stage, the 5-year survival rate is
7 99% [1]. However, only ~62% of the breast cancer cases are detected in the localized
8 stage [1]. In ~30% of the cases, breast cancer is detected after it spreads to the regional
9 lymph nodes, reducing the 5-year survival rate to 85%. Furthermore, in 6% of cases, the
10 cancer is diagnosed after it has spread to a distant part of the body beyond the lymph
11 nodes and the 5-year survival rate is reduced to 27%. To detect breast cancer early, the
12 US Preventive Services Task Force (USPSTF) recommends a biennial screening
13 mammography for women over 50 years old. For women under 50 years old, the
14 decision for screening must be individualized to balance the benefit of potential early
15 detection against the risk of false positive diagnosis. False-positive mammography
16 results, which typically lead to unnecessary follow-up diagnostic testing, become
17 increasingly common for women 40 to 49 years old [2]. Nevertheless, for women with
18 high risk for breast cancer (i.e. a lifetime risk of breast cancer higher than 20%), the
19 American Cancer Society advises a yearly breast MRI and mammogram starting at 30
20 years of age [3].

21 Polygenic risk scores (PRS) assess the genetic risks of complex diseases based
22 on the aggregate statistical correlation of a disease outcome with many genetic
23 variations over the whole genome. Single-nucleotide polymorphisms (SNPs) are the
24 most commonly used genetic variations. While genome-wide association studies
25 (GWAS) report only SNPs with statistically significant associations to phenotypes [4],
26 PRS can be estimated using a greater number of SNPs with higher adjusted p-value
27 thresholds to improve prediction accuracy.

28 Previous research has developed a variety of PRS estimation models based on
29 Best Linear Unbiased Prediction (BLUP), including gBLUP [5], rr-BLUP [6], [7], and
30 other derivatives [8], [9]. These linear mixed models consider genetic variations as fixed
31 effects and use random effects to account for environmental factors and individual

variability. Furthermore, linkage disequilibrium was utilized as a basis for the LDpred [10], [11] and PRS-CS [12] algorithms

PRS estimation can also be defined as a supervised classification problem. The input features are genetic variations and the output response is the disease outcome. Thus, machine learning techniques can be used to estimate PRS based on the classification scores achieved [13]. A large-scale GWAS dataset may provide tens of thousands of individuals as training examples for model development and benchmarking. Wei et al (2019) [14] compared support vector machine and logistic regression to estimate PRS of Type-1 diabetes. The best Area Under the receiver operating characteristic Curve (AUC) was 84% in this study. More recently, neural networks have been used to estimate human height from the GWAS data, and the best R^2 scores were in the range of 0.4 to 0.5 [15]. Amyotrophic lateral sclerosis was also investigated using Convolutional Neural Networks (CNN) with 4511 cases and 6127 controls [16] and the highest accuracy was 76.9%.

Significant progress has been made for estimating PRS for breast cancer from a variety of populations. In a recent study [17], multiple large European women cohorts were combined to compare a series of PRS models. The most predictive model in this study used lasso regression with 3,820 SNPs and obtained an AUC of 65%. A PRS algorithm based on the sum of log odds ratios of important SNPs for breast cancer was used in the Singapore Chinese Health Study [18] with 46 SNPs and 56.6% AUC, the Shanghai Genome-Wide Association Studies [19] with 44 SNPs and 60.6% AUC, and a Taiwanese cohort [20] with 6 SNPs and 59.8% AUC. A pruning and thresholding method using 5,218 SNPs reached an AUC of 69% for the UK Biobank dataset [11].

In this study, deep neural network (DNN) was tested for breast cancer PRS estimation using a large cohort containing 26053 cases and 23058 controls. The performance of DNN was shown to be higher than alternative machine learning algorithms and other statistical methods in this large cohort. Furthermore, DeepLift [21] and LIME [22] were used to identify salient SNPs used by DNN for prediction.

Materials and Methods

Breast cancer GWAS data

This study used a breast cancer GWAS dataset generated by the Discovery, Biology, and Risk of Inherited Variants in Breast Cancer (DRIVE) project [23] and was obtained from the NIH dbGaP database under the accession number of phs001265.v1.p1. The DRIVE dataset was stored, processed and used on the Schooner supercomputer at the University of Oklahoma in an isolated partition with restricted access. The partition consisted of 5 computational nodes, each with 40 CPU cores (Intel Xeon Cascade Lake) and 200 GB of RAM. The DRIVE dataset in the dbGap database was composed of 49,111 subjects genotyped for 528,620 SNPs using OncoArray [23]. 55.4% of the subjects were from North America, 43.3% from Europe, and 1.3% from Africa. The disease outcome of the subjects was labeled as malignant tumor (48%), *in situ* tumor (5%), and no tumor (47%). In this study, the subjects in the malignant tumor and *in situ* tumor categories were labeled as cases and the subjects in the no tumor category were labeled as controls, resulting in 26053 (53%) cases and 23058 (47%) controls. The subjects in the case and control classes were randomly assigned to a training set (80%), a validation set (10%), and a test set (10%) (Figure 1). The association analysis was conducted on the training set using Plink 2.0 [24]. For a subject, each of the 528,620 SNPs may take the value of 0, 1 or 2, representing the genotype value on a SNP for this subject. The value of 0 meant homozygous with minor allele, 1 meant heterozygous allele, and 2 meant homozygous with the dominant allele. Such encoding of the SNP information was also used in the following machine learning and statistical approaches. The p-value for each SNP was calculated using logistic regression in Plink 2.0.

Development of deep neural network models for PRS estimation

A variety of deep neural network (DNN) architectures [25] were trained using Tensorflow 1.13. The Leaky Rectified Linear Unit (ReLU) activation function [26] was used on all hidden-layers neurons with the negative slope co-efficient set to 0.2. The output neuron used a sigmoid activation function. The training error was computed using the cross-entropy function:

$$\sum_{i=1}^n y * \log(p) + (1 - y) * \log(1 - p),$$

where $p \in [0,1]$ is the prediction probability from the model and $y \in \{0,1\}$ is the prediction target at 1 for case and 0 for control. The prediction probability was considered as the PRS from DNN.

DNNs were evaluated using different SNP feature sets. SNPs were filtered using their Plink association p-values at the thresholds of 10^{-2} , 10^{-3} , 10^{-4} and 10^{-5} . DNN was also tested using the full SNP feature set without any filtering. The DNN models were trained using mini-batches with a batch size of 512. The Adam optimizer [27], an adaptive learning rate optimization algorithm, was used to update the weights in each mini-batch. The initial learning rate was set to 10^{-4} , and the models were trained for up to 200 epochs with early stopping based on the validation AUC score. Dropout [28] was used to reduce overfitting. The dropout rates of 20%, 30%, 40%, 50%, 60%, 70%, 80%, and 90% were tested for the first hidden layer and the final dropout rate was selected based on the validation AUC score. The dropout rate was set to 50% on the other hidden layers in all architectures. Batch normalization (BN) [29] was used to accelerate the training process, and the momentum for the moving average was set to 0.9 in BN.

Development of alternative machine learning models for PRS estimation

Logistic regression, decision tree, random forest, AdaBoost, gradient boosting, support vector machine (SVM), and Gaussian naive Bayes were implemented and tested using the scikit-learn machine learning library in Python. These models were trained using the same training set as the DNNs and their hyperparameters were tuned using the same validation set based on the validation AUC (Figure 1). These models are briefly described below.

- **Decision Tree:** The gini information gain with best split was used. The maximum depth was not set to let the tree expanded until all leaves were pure or contained less than a minimum number of two examples per split (sklearn default parameters).
- **Random Forest:** classification decision trees (as configured above) were used as base learners. The optimum number of decision trees were found to be 3,000 based on a parameter sweep between 500 and 5,000 with a step size of 500. Bootstrap samples were used to build each base learner. When searching for each tree's best

split, the maximum number of considered features was set to be the square root of the number of features.

- **AdaBoost:** classification decision trees (as configured above) were used as base learners. The optimum number of decision trees were found to be 2,000 based on a parameter sweep between 500 and 5,000 with a step size of 500. The learning rate was set to 1. The algorithm used was SAMME.R [30].
- **Gradient Boosting:** regression decision trees (as configured above) were used as the base learners. The optimum number of decision trees were found to be 400 based on a parameter sweep between 100 and 1,000 with a step size of 100. Log-loss was used as the loss function. The learning rate was set to 0.1. The mean squared error with improvement score [31] was used to measure the quality of a split.
- **SVM:** The kernel was a radial basis function with $\gamma = \frac{1}{n * Var}$, where n is the number of SNPs and Var is the variance of the SNPs across individuals. The regularization parameter C was set to 1 based on a parameter sweep over 0.001, 0.01, 0.1, 1, 5, 10, 15 and 20.
- **Logistic Regression:** L2 regularization with $\alpha = 0.5$ was used based on a parameter sweep for α over 0.0001, 0.001, 0.01, 0.1, 0.2, 0.3, 0.4, 0.5, 0.6, 0.7 and 0.8. L1 regularization was tested, but not used, because it did not improve the performance.
- **Gaussian Naïve Bayes:** The likelihood of the features was assumed to be Gaussian. The classes had uninformative priors.

Development of statistical models for PRS estimation

The same training and validation sets were used to develop statistical models (Figure 1). The BLUP and BayesA models were constructed using the bWGR R package. The LDpred model was constructed as described [10].

- **BLUP:** The linear mixed model was $y = \mu + Xb + e$, where y were the response variables, μ were the intercepts, X were the input features, b were the regression coefficients, and e were the residual coefficients.

- **BayesA:** The priors were assigned from a mixture of normal distributions.
- **LDpred:** The p-values were generated by our association analysis described above. The validation set was provided as reference for LDpred data coordination. The radius of the Gibbs sampler was set to be the number of SNPs divided by 3000 as recommended by the LDpred user manual (<https://github.com/bvilhjal/ldpred/blob/master/ldpred/run.py>).

The score distributions of DNN, BayesA, BLUP and LDpred were analyzed with the Shapiro test for normality and the Bayesian Gaussian Mixture (BGM) expectation maximization algorithm. The BGM algorithm decomposed a mixture of two Gaussian distributions with weight priors at 50% over a maximum of 1000 iterations and 100 initializations.

DNN model interpretation.

LIME and DeepLift were used to interpret the DNN predictions for subjects in the test set with DNN output scores higher than 0.67, which corresponded to a precision of 90%. For LIME, the submodular pick algorithm was used, the kernel size was set to 40, and the number of explainable features was set to 41. For DeepLift, the importance of each SNPs was computed as the average across all individuals, and the reference activation value for a neuron was determined by the average value of all activations triggered across all subjects.

Results and Discussion

Development of a machine learning model for breast cancer PRS estimation

The breast cancer GWAS dataset containing 26053 cases and 23058 controls was generated by the Discovery, Biology, and Risk of Inherited Variants in Breast Cancer (DRIVE) project [23]. The DRIVE data is available from the NIH dbGaP database under the accession number of phs001265.v1.p1. The cases and controls were randomly split to a training set, a validation set, and a test set (Figure 1). The training set was used to estimate p-values of SNPs using association analysis and train

1 machine learning and statistical models. The hyperparameters of the machine learning
2 and statistical models were optimized using the validation set. The test set was used for
3 the final performance evaluation and model interpretation.

4 Statistical significance of the disease association with the 528,620 SNPs was
5 assessed with Plink using only the training set. To obtain unbiased benchmarking
6 results on the test set, it was critical not to use the test set in the association analysis
7 (Figure 1) and not to use association p-values from previous GWAS studies that
8 included subjects in the test set, as well-described in the Section 7.10.2 of Hastie et al
9 [32]. The obtained p-values for all SNPs are shown in Figure 2A as a Manhattan plot.
10 There were 1,061 SNPs with a p-value less than the critical value of $9.5 \cdot 10^{-8}$, which
11 was set using the Bonferroni correction ($9.5 \cdot 10^{-8} = 0.05/528,620$). Filtering with a
12 Bonferroni-corrected critical value may remove many informative SNPs that have small
13 effects on the phenotype, epistatic interactions with other SNPs, or non-linear
14 association with the phenotype [33]. Relaxed filtering with higher p-value cutoffs was
15 tested to find the optimal feature set for DNN (Figure 2B and Supplementary Table 1).
16 The DNN models in Figure 2B had a deep feedforward architecture consisting of an
17 input layer of variable sizes, followed by 3 successive hidden layers containing 1000,
18 250, and 50 neurons, and finally an output layer with a single neuron. As the p-value
19 cutoff increased, a greater number of SNPs were incorporated as input features, and
20 training consumed a larger amount of computational resources in terms of computing
21 time and peak memory usage. A feature set containing 5,273 SNPs above the p-value
22 cutoff of 10^{-3} provided the best prediction performance measured by the AUC and
23 accuracy on the validation set. In comparison with smaller feature sets from more
24 stringent p-value filtering, the 5,273-SNP feature set may have included many
25 informative SNPs providing additional signals to be captured by DNN for prediction. On
26 the other hand, more relaxed filtering with p-value cutoffs greater than 10^{-3} led to
27 significant overfitting as indicated by an increasing prediction performance in the
28 training set and a decreasing performance in the validation set (Figure 2B).

29 Previous studies [11], [34] have used a large number of SNPs for PRS estimation
30 on different datasets. In our study, the largest DNN model, consisting of all 528,620
31 SNPs, decreased the validation AUC score by 1.2% and the validation accuracy by

1 1.9% from the highest achieved values. This large DNN model relied an 80% dropout
2 rate to obtain strong regularization, while all the other DNN models utilized a 50%
3 dropout rate. This suggested that DNN was able to perform feature selection without
4 using association p-values, although the limited training data and the large neural
5 network size resulted in complete overfitting with a 100% training accuracy and the
6 lowest validation accuracy (Figure 2B).

7 The effects of dropout and batch normalization were tested using the 5,273-SNP
8 DNN model (Supplementary Figure 1). Without dropout, the DNN model using only
9 batch normalization had a 3.0% drop in AUC and a 4.0% drop in accuracy and its
10 training converged in only two epochs. Without batch normalization, the DNN model had
11 0.1% higher AUC and 0.3% lower accuracy but its training required a 73% increase in
12 the number of epochs to reach convergence.

13 As an alternative to filtering, autoencoding was tested to reduce SNPs to a
14 smaller set of encodings as described previously [35], [36]. An autoencoder was trained
15 to encode 5273 SNPs into 2000 features with a mean square error (MSE) of 0.053 and
16 a root mean square error (RMSE) of 0.23. The encodings from the autoencoder were
17 used as the input features to train a DNN model with the same architecture as the ones
18 shown in Figure 2B except for the number of input neurons. The autoencoder-DNN
19 model had a similar number of input neurons for DNN as the 2099-SNP DNN model, but
20 had a 1.3% higher validation AUC and a 0.2% higher validation accuracy than the 2099-
21 SNP DNN model (Figure 2B). This increased validation AUC and accuracy suggested
22 that the dimensionality reduction by the autoencoding from 5273 SNPs to 2000
23 encodings was better than the SNP filtering by the association p-values from 5273
24 SNPs to 2099 SNPs. However, the DNN models with 5,273 SNPs still had a 0.3%
25 higher validation AUC score and a 1.6% higher validation accuracy than the
26 autoencoder-DNN model.

27 The deep feedforward architecture benchmarked in Figure 2B was compared
28 with a number of alternative neural network architectures using the 5,273-SNP feature
29 set (Supplementary Table 2). A shallow neural network with only one hidden layer
30 resulted in a 0.9% lower AUC and 1.1% lower accuracy in the validation set compared
31 to the DNN. This suggested that additional hidden layers in DNN may allow additional

feature selection and transformation in the model. One-dimensional convolutional neural network (1D CNN) was previously used to estimate the PRS for bone heel mineral density, body mass index, systolic blood pressure and waist-hip ratio [15] and was also tested here for breast cancer prediction with the DRIVE dataset. The validation AUC and accuracy of 1D CNN was lower than DNN by 3.2% and 2.0%, respectively. CNN was commonly used for image analysis, because the receptive field of the convolutional layer can capture space-invariant information with shared parameters. However, the SNPs distributed across a genome may not have significant space-invariant patterns to be captured by the convolutional layer, which may explain the poor performance of CNN.

The 5,273-SNP feature set was used to test alternative machine learning approaches, including logistic regression, decision tree, naive Bayes, random forest, ADABOOST, gradient boosting, and SVM, for PRS estimation (Figure 3). These models were trained, turned, and benchmarked using the same training, validation, and test sets, respectively, as the DNN models (Figure 1). Although the decision tree had a test AUC of only 50.9%, ensemble algorithms that used decision trees as the base learner, including random forest, ADABOOST, and gradient boosting, reached test AUCs of 63.6%, 64.4%, and 65.1%, respectively. This showed the advantage of ensemble learning. SVM reached a test AUC of 65.6%. Naïve Bayes and logistic regression were both linear models with the assumption of independent features. Logistic regression had higher AUC, but lower accuracy, than SVM and gradient boosting. The test AUC and test accuracy of DNN were higher than those of logistic regression by 0.9% and 2.7%, respectively. Out of all the machine learning models, the DNN model achieved the highest test AUC at 67.4% and the highest test accuracy at 62.8% (Figure 3).

Comparison of the DNN model with statistical models for breast cancer PRS estimation

The performance of DNN was compared with three representative statistical models, including BLUP, BayesA, and LDpred (Table 1). Because the relative performance of these methods may be dependent on the number of training examples available, the original training set containing 39,289 subjects was down-sampled to create three smaller training sets containing 10000, 20000, 30000 subjects. As the

1 5,273-SNP feature set generated with a p-value cutoff of 10^{-3} may not be the most
2 appropriate for the statistical methods, a 13,890-SNP feature set (p-value cutoff = 10^{-2})
3 and a 2,099-SNP feature set (p-value cutoff = 10^{-5}) were tested for all methods.

4 Although LDpred also required training data, its prediction relied primarily on the
5 provided p-values, which were generated for all methods using all 39,289 subjects in
6 the training set. Thus, the down-sampling of the training set did not reduce the
7 performance of LDpred. LDpred reached its highest AUC score at 62.4% using the p-
8 value cutoff of 10^{-3} . A previous study [12] that applied LDpred to breast cancer
9 prediction using the UK Biobank dataset similarly obtained an AUC score of 62.4% at
10 the p-value cutoff of 10^{-3} . This showed consistent performance of LDpred in the two
11 studies. When DNN, BLUP, and BayesA used the full training set, they obtained higher
12 AUCs than LDpred at their optimum p-value cutoffs.

13 DNN, BLUP, and BayesA all gained performance with the increase in the training
14 set sizes (Table 1). The performance gain was more substantial for DNN than BLUP
15 and BayesA. The increase from 10,000 subjects to 39,258 subjects in the training set
16 resulted in a 1.9% boost to DNN's best AUC, a 0.7% boost to BLUP, and a 0.8% boost
17 to BayesA. This indicated the different variance-bias trade-offs made by DNN, BLUP,
18 and BayesA. The high variance of DNN required more training data, but could capture
19 non-linear relationships between the SNPs and the phenotype. The high bias of BLUP
20 and BayesA had lower risk for overfitting using smaller training sets, but their models
21 only considered linear relationships. The higher AUCs of DNN across all training set
22 sizes indicated that DNN had a better variance-bias balance for breast cancer PRS
23 estimation.

24 For all four training set sizes, BLUP and BayesA achieved higher AUCs using
25 more stringent p-value filtering. When using the full training set, reducing the p-value
26 cutoffs from 10^{-2} to 10^{-5} increased the AUCs of BLUP from 61.0% to 64.2% and the
27 AUCs of BayesA from 61.1% to 64.5%. This suggested that BLUP and BayesA
28 preferred a reduced number of SNPs that were significantly associated with the
29 phenotype. On the other hand, DNN produced lower AUCs using the p-value cutoff of
30 10^{-5} than the other two higher cutoffs. This suggested that DNN can perform better
31 feature selection in comparison to SNP filtering based on association p-values.

The four algorithms were compared using the PRS histograms of the case population and the control population from the test set in Figure 4. The score distributions of BLUP, BayesA and LDpred all followed normal distributions. The p-values from the Shapiro normality test of the case and control distributions were 0.46 and 0.43 for BayesA, 0.50 and 0.95 for BLUP, and 0.17 and 0.24 for LDpred, respectively. The case and control distributions were $N_{case}(\mu = 0.577, \sigma = 0.20)$ and $N_{control}(\mu = 0.479, \sigma = 0.19)$ from BayesA, $N_{cases}(\mu = 0.572, \sigma = 0.19)$ and $N_{control}(\mu = 0.483, \sigma = 0.18)$ from BLUP, and $N_{case}(\mu = -33.52, \sigma = 5.4)$ and $N_{control}(\mu = -35.86, \sigma = 4.75)$ from LDpred. The means of the case distributions were all significantly higher than the control distributions for BayesA (p-value $< 10^{-16}$), BLUP (p-value $< 10^{-16}$), and LDpred (p-value $< 10^{-16}$) and their case and control distributions had similar standard deviations.

The score histograms of DNN did not follow normal distributions based on the Shapiro normality test with a p-value of $4.1 * 10^{-34}$ for the case distribution and a p-value of $2.5 * 10^{-9}$ for the control distribution. The case distribution had the appearance of a bi-modal distribution. The Bayesian Gaussian mixture expectation maximization algorithm decomposed the case distribution to two normal distributions: $N_{case1}(\mu = 0.519, \sigma = 0.096)$ with an 86.5% weight and $N_{case2}(\mu = 0.876, \sigma = 0.065)$ with a 13.5% weight. The control distribution was resolved into two normal distributions with similar means and distinct standard deviations: $N_{control1}(\mu = 0.471, \sigma = 0.1)$ with an 85.0% weight and $N_{control2}(\mu = 0.507, \sigma = 0.03)$ with a 15.0% weight. The N_{case1} distribution had a similar mean as the $N_{control1}$ and $N_{control2}$ distributions. This suggested that the N_{case1} distribution may represent a normal-genetic-risk case sub-population, in which the subjects may have a normal level of genetic risk for breast cancer and the oncogenesis likely involved a significant environmental component. The mean of the N_{case2} distribution was higher than the means of both the N_{case1} and $N_{control1}$ distributions by more than 4 standard deviations (p-value $< 10^{-16}$). We hypothesized that the N_{case2} distribution represented a high-genetic-risk case sub-population for breast cancer, in which the subjects may have inherited many genetic variations associated with breast cancer.

Three GWAS were performed between the high-genetic-risk case sub-population with DNN PRS > 0.67, the normal-genetic-risk case sub-population with DNN PRS < 0.67, and the control population (Supplementary Table 3). The GWAS analysis of the high-genetic-risk case sub-population versus the control population identified 182 significant SNPs at the Bonferroni level of statistical significance. The GWAS analysis of the high-genetic-risk case sub-population versus the normal-genetic-risk case sub-population identified 216 significant SNPs. The two sets of significant SNPs found by these two GWAS analyses were very similar, sharing 149 significant SNPs in their intersection. Genes associated with these 149 SNPs were investigated with pathway enrichment analysis (Fisher's Exact Test; $P < 0.05$) using SNPnexus [37] (Supplementary Table 4). Many of the significant pathways were involved in DNA repair [38] signal transduction [39], and suppression of apoptosis [40]. Interestingly, the GWAS analysis of the normal-genetic-risk case sub-population and the control population identified no significant SNP. This supported our classification of the cases into the normal-genetic-risk subjects and the high-genetic-risk subjects based on their PRS scores from the DNN model.

In comparison with AUCs, it may be more relevant for practical applications of PRS to compare the recalls of different algorithms at a given precision that warrants clinical recommendations. At 90% precision, the recalls were 18.8% for DNN, 0.2% for BLUP, 1.3% for BayesA, and 1.3% for LDpred in the test set of the DRIVE cohort with a ~50% prevalence. This indicated that DNN can make a positive prediction for 18.8% of the subjects in the DRIVE cohort and these positive subjects would have an average chance of 90% to eventually develop breast cancer. American Cancer Society advises yearly breast MRI and mammogram starting at the age of 30 years for women with a lifetime risk of breast cancer greater than 20%, which meant a 20% precision for PRS. By extrapolating the performance in the DRIVE cohort, the DNN model should be able to achieve a recall of 65.4% at a precision of 20% in the general population with a 12% prevalence rate of breast cancer.

Interpretation of the DNN model

While the DNN model used 5,273 SNPs as input, we hypothesized that only a small set of these SNPs were particularly informative for identifying the subjects with high genetic risks for breast cancer. LIME and DeepLift were used to find the top-100 salient SNPs used by the DNN model to identify the subjects with PRS higher than the 0.67 cutoff at 90% precision in the test set (Figure 1). Twenty three SNPs were ranked by both algorithms to be among their top-100 salient SNPs (Supplementary Figure 2). The small overlap between their results can be attributed to their different interpretation approaches. LIME considered the DNN model as a black box and perturbed the input to estimate the importance of each variable; whereas, DeepLift analyzed the gradient information of the DNN model. 30% of LIME's salient SNPs and 49% of DeepLift's salient SNPs had p-values less than the Bonferroni significance threshold of $9.5 \cdot 10^{-8}$. This could be attributed to the non-linear relationships between the salient SNP genotype and the disease outcome, which cannot be captured by the association analysis using logistic regression. To illustrate this, four salient SNPs with significant p-values were shown in Supplementary Figure 3A, which exhibited linear relationships between their genotype values and log odds ratios as expected. Four salient SNPs with insignificant p-values were shown in Supplementary Figure 3B, which showed clear biases towards cases or controls by one of the genotype values in a non-linear fashion.

Michailidiou et al. [41] summarized a total of 172 SNPs associated with breast cancer. Out of these SNPs, 59 were not included on OncoArray, 63 had an association p-value less than 10^{-3} and were not included in the 5,273-SNP feature set for DNN, 34 were not ranked among the top-1000 SNPs by either DeepLIFT or LIME, and 16 were ranked among the top-1000 SNPs by DeepLIFT, LIME or both (Supplementary Table 5). This indicates that many SNPs with significant association may be missed by the interpretation of DNN models.

The 23 salient SNPs identified by both DeepLift and LIME in their top-100 list are shown in Table 2. Eight of the 23 SNPs had p-values higher than the Bonferroni level of significance and were missed by the association analysis using Plink. The potential oncogenesis mechanisms for some of the 8 SNPs have been investigated in previous studies. The SNP, rs139337779 at 12q24.22, is located within the gene, Nitric oxide synthase 1 (NOS1). Li et al. [42] showed that the overexpression of NOS1 can up-regulate

the expression of ATP-binding cassette, subfamily G, member 2 (ABCG2), which is a breast cancer resistant protein [43], and NOS1-induced chemo-resistance was partly mediated by the up-regulation of ABCG2 expression. Lee et al. [44] reported that NOS1 is associated with the breast cancer risk in a Korean cohort. The SNP, chr13_113796587_A_G at 13q34, is located in the F10 gene, which is the coagulation factor X. Tinholt et al [45] showed that the increased coagulation activity and genetic polymorphisms in the F10 gene are associated with breast cancer. The BNC2 gene containing the SNP, chr9_16917672_G_T at 9p22.2, is a putative tumor suppressor gene in high-grade serious ovarian carcinoma [46]. The SNP, chr2_171708059_C_T at 2q31.1, is within the GAD1 gene and the expression level of GAD1 is a significant prognostic factor in lung adenocarcinoma [47]. Thus, the interpretation of DNN models may identify novel SNPs with non-linear association with the breast cancer.

Acknowledgement

We would like to thank the OU Supercomputing Center for Education & Research (OSCER) for supercomputing technical support, the DRIVE project for the GWAS data, NIH dbGap for data access authorization, and Dr. Xu Chao for helpful discussions. The study was funded by Dr. Pan's startup funding from the University of Oklahoma and by the Oak Ridge National Laboratory (ORNL)' Directed Research Development (LDRD) Funding. Oak Ridge National Laboratory is managed by UT-Battelle, LLC for the U.S. Department of Energy under Contract Number DE-AC05-00OR22725.

References:

- [1] NIH, "Female Breast Cancer - Cancer Stat Facts." <https://seer.cancer.gov/statfacts/html/breast.html> (accessed Dec. 03, 2019).
- [2] H. D. Nelson, K. Tyne, A. Naik, C. Bougatsos, B. K. Chan, and L. Humphrey, "Screening for Breast Cancer: An Update for the U.S. Preventive Services Task Force," *Annals of Internal Medicine*, vol. 151, no. 10, pp. 727–737, Nov. 2009, doi: 10.7326/0003-4819-151-10-200911170-00009.
- [3] K. C. Oeffinger *et al.*, "Breast Cancer Screening for Women at Average Risk: 2015 Guideline Update From the American Cancer Society," *JAMA*, vol. 314, no. 15, pp. 1599–1614, Oct. 2015, doi: 10.1001/jama.2015.12783.

- [4] F. Dudbridge, "Power and Predictive Accuracy of Polygenic Risk Scores," *PLOS Genetics*, vol. 9, no. 3, p. e1003348, Mar. 2013, doi: 10.1371/journal.pgen.1003348.
- [5] S. A. Clark, B. P. Kinghorn, J. M. Hickey, and J. H. van der Werf, "The effect of genomic information on optimal contribution selection in livestock breeding programs," *Genetics Selection Evolution*, vol. 45, no. 1, p. 44, Oct. 2013, doi: 10.1186/1297-9686-45-44.
- [6] A. J. Whittaker, I. Royzman, and T. L. Orr-Weaver, "Drosophila Double parked: a conserved, essential replication protein that colocalizes with the origin recognition complex and links DNA replication with mitosis and the down-regulation of S phase transcripts," *Genes Dev*, vol. 14, no. 14, pp. 1765–1776, Jul. 2000.
- [7] T. H. E. Meuwissen, B. J. Hayes, and M. E. Goddard, "Prediction of Total Genetic Value Using Genome-Wide Dense Marker Maps," *Genetics*, vol. 157, no. 4, pp. 1819–1829, Apr. 2001.
- [8] R. Maier *et al.*, "Joint analysis of psychiatric disorders increases accuracy of risk prediction for schizophrenia, bipolar disorder, and major depressive disorder," *Am. J. Hum. Genet.*, vol. 96, no. 2, pp. 283–294, Feb. 2015, doi: 10.1016/j.ajhg.2014.12.006.
- [9] D. Speed and D. J. Balding, "MultiBLUP: improved SNP-based prediction for complex traits," *Genome Res.*, vol. 24, no. 9, pp. 1550–1557, Sep. 2014, doi: 10.1101/gr.169375.113.
- [10] B. J. Vilhjálmsson *et al.*, "Modeling Linkage Disequilibrium Increases Accuracy of Polygenic Risk Scores," *Am J Hum Genet*, vol. 97, no. 4, pp. 576–592, Oct. 2015, doi: 10.1016/j.ajhg.2015.09.001.
- [11] A. V. Khera *et al.*, "Genome-wide polygenic scores for common diseases identify individuals with risk equivalent to monogenic mutations," *Nat Genet*, vol. 50, no. 9, pp. 1219–1224, Sep. 2018, doi: 10.1038/s41588-018-0183-z.
- [12] T. Ge, C.-Y. Chen, Y. Ni, Y.-C. A. Feng, and J. W. Smoller, "Polygenic prediction via Bayesian regression and continuous shrinkage priors," *Nature Communications*, vol. 10, no. 1, pp. 1–10, Apr. 2019, doi: 10.1038/s41467-019-09718-5.
- [13] D. S. W. Ho, W. Schierding, M. Wake, R. Saffery, and J. O'Sullivan, "Machine Learning SNP Based Prediction for Precision Medicine," *Front. Genet.*, vol. 10, 2019, doi: 10.3389/fgene.2019.00267.
- [14] Z. Wei *et al.*, "From Disease Association to Risk Assessment: An Optimistic View from Genome-Wide Association Studies on Type 1 Diabetes," *PLOS Genetics*, vol. 5, no. 10, p. e1000678, Oct. 2009, doi: 10.1371/journal.pgen.1000678.
- [15] P. Bellot, G. de los Campos, and M. Pérez-Enciso, "Can Deep Learning Improve Genomic Prediction of Complex Human Traits?," *Genetics*, vol. 210, no. 3, pp. 809–819, Nov. 2018, doi: 10.1534/genetics.118.301298.
- [16] B. Yin *et al.*, "Using the structure of genome data in the design of deep neural networks for predicting amyotrophic lateral sclerosis from genotype," *bioRxiv*, p. 533679, Jan. 2019, doi: 10.1101/533679.
- [17] N. Mavaddat *et al.*, "Polygenic Risk Scores for Prediction of Breast Cancer and Breast Cancer Subtypes," *The American Journal of Human Genetics*, vol. 104, no. 1, pp. 21–34, Jan. 2019, doi: 10.1016/j.ajhg.2018.11.002.
- [18] C. H. T. Chan *et al.*, "Evaluation of three polygenic risk score models for the prediction of breast cancer risk in Singapore Chinese," *Oncotarget*, vol. 9, no. 16, pp. 12796–12804, Jan. 2018, doi: 10.18632/oncotarget.24374.

- [19] W. Wen *et al.*, “Prediction of breast cancer risk based on common genetic variants in women of East Asian ancestry,” *Breast Cancer Research*, vol. 18, no. 1, p. 124, Dec. 2016, doi: 10.1186/s13058-016-0786-1.
- [20] Y.-C. Hsieh *et al.*, “A polygenic risk score for breast cancer risk in a Taiwanese population,” *Breast Cancer Res Treat*, vol. 163, no. 1, pp. 131–138, May 2017, doi: 10.1007/s10549-017-4144-5.
- [21] A. Shrikumar, P. Greenside, and A. Kundaje, “Learning Important Features Through Propagating Activation Differences,” in *International Conference on Machine Learning*, Jul. 2017, pp. 3145–3153, Accessed: Nov. 11, 2019. [Online]. Available: <http://proceedings.mlr.press/v70/shrikumar17a.html>.
- [22] M. T. Ribeiro, S. Singh, and C. Guestrin, ““Why Should I Trust You?”: Explaining the Predictions of Any Classifier,” in *Proceedings of the 22Nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, New York, NY, USA, 2016, pp. 1135–1144, doi: 10.1145/2939672.2939778.
- [23] C. I. Amos *et al.*, “The OncoArray Consortium: A Network for Understanding the Genetic Architecture of Common Cancers,” *Cancer Epidemiol Biomarkers Prev*, vol. 26, no. 1, pp. 126–135, Jan. 2017, doi: 10.1158/1055-9965.EPI-16-0106.
- [24] C. C. Chang, C. C. Chow, L. C. Tellier, S. Vattikuti, S. M. Purcell, and J. J. Lee, “Second-generation PLINK: rising to the challenge of larger and richer datasets,” *Gigascience*, vol. 4, no. 1, Dec. 2015, doi: 10.1186/s13742-015-0047-8.
- [25] Y. Bengio, “Learning Deep Architectures for AI,” *Found. Trends Mach. Learn.*, vol. 2, no. 1, pp. 1–127, Jan. 2009, doi: 10.1561/22000000006.
- [26] B. Xu, N. Wang, T. Chen, and M. Li, “Empirical Evaluation of Rectified Activations in Convolutional Network,” *arXiv:1505.00853 [cs, stat]*, Nov. 2015, Accessed: Nov. 25, 2019. [Online]. Available: <http://arxiv.org/abs/1505.00853>.
- [27] D. P. Kingma and J. Ba, “Adam: A Method for Stochastic Optimization,” *3rd International Conference for Learning Representations*, 2015, Accessed: Nov. 26, 2019. [Online]. Available: <http://arxiv.org/abs/1412.6980>.
- [28] N. Srivastava, G. Hinton, A. Krizhevsky, I. Sutskever, and R. Salakhutdinov, “Dropout: A Simple Way to Prevent Neural Networks from Overfitting,” *Journal of Machine Learning Research*, vol. 15, pp. 1929–1958, 2014.
- [29] S. Ioffe and C. Szegedy, “Batch Normalization: Accelerating Deep Network Training by Reducing Internal Covariate Shift,” *arXiv:1502.03167 [cs]*, Mar. 2015, Accessed: Nov. 25, 2019. [Online]. Available: <http://arxiv.org/abs/1502.03167>.
- [30] T. Hastie, S. Rosset, J. Zhu, and H. Zou, “Multi-class adaboost,” *Statistics and its Interface*, vol. 2, no. 3, pp. 349–360, 2009.
- [31] J. H. Friedman, “Greedy function approximation: a gradient boosting machine,” *Annals of statistics*, pp. 1189–1232, 2001.
- [32] T. Hastie, R. Tibshirani, and J. Friedman, *The Elements of Statistical Learning: Data Mining, Inference, and Prediction, Second Edition*, 2nd ed. New York: Springer-Verlag, 2009.
- [33] R. De, W. S. Bush, and J. H. Moore, “Bioinformatics Challenges in Genome-Wide Association Studies (GWAS),” in *Clinical Bioinformatics*, R. Trent, Ed. New York, NY: Springer, 2014, pp. 63–81.

- [34] D. Gola, J. Erdmann, B. Müller-Myhsok, H. Schunkert, and I. R. König, "Polygenic risk scores outperform machine learning methods in predicting coronary artery disease status," *Genetic Epidemiology*, vol. 44, no. 2, pp. 125–138, Mar. 2020, doi: 10.1002/gepi.22279.
- [35] P. Fergus, A. Montanez, B. Abdulaimma, P. Lisboa, C. Chalmers, and B. Pineles, "Utilising Deep Learning and Genome Wide Association Studies for Epistatic-Driven Preterm Birth Classification in African-American Women," *IEEE/ACM Transactions on Computational Biology and Bioinformatics*, pp. 1–1, 2018, doi: 10.1109/TCBB.2018.2868667.
- [36] M. Cudic, H. Baweja, T. Parhar, and S. Nuske, "Prediction of Sorghum Bicolor Genotype from In-Situ Images Using Autoencoder-Identified SNPs," in *2018 17th IEEE International Conference on Machine Learning and Applications (ICMLA)*, Dec. 2018, pp. 23–31, doi: 10.1109/ICMLA.2018.00012.
- [37] A. Z. Dayem Ullah, J. Oscanoa, J. Wang, A. Nagano, N. R. Lemoine, and C. Chelala, "SNPnexus: assessing the functional relevance of genetic variation to facilitate the promise of precision medicine," *Nucleic Acids Res*, vol. 46, no. W1, pp. W109–W113, Jul. 2018, doi: 10.1093/nar/gky399.
- [38] M. J. O'Connor, "Targeting the DNA Damage Response in Cancer," *Molecular Cell*, vol. 60, no. 4, pp. 547–560, Nov. 2015, doi: 10.1016/j.molcel.2015.10.040.
- [39] W. Kolch, M. Halasz, M. Granovskaya, and B. N. Kholodenko, "The dynamic control of signal transduction networks in cancer cells," *Nature Reviews Cancer*, vol. 15, no. 9, Art. no. 9, Sep. 2015, doi: 10.1038/nrc3983.
- [40] K. Fernald and M. Kurokawa, "Evading apoptosis in cancer," *Trends in Cell Biology*, vol. 23, no. 12, pp. 620–633, Dec. 2013, doi: 10.1016/j.tcb.2013.07.006.
- [41] K. Michailidou *et al.*, "Association analysis identifies 65 new breast cancer risk loci," *Nature*, vol. 551, no. 7678, pp. 92–94, Nov. 2017, doi: 10.1038/nature24284.
- [42] X. Li *et al.*, "NOS1 upregulates ABCG2 expression contributing to DDP chemoresistance in ovarian cancer cells," *Oncology Letters*, vol. 17, no. 2, pp. 1595–1602, Feb. 2019, doi: 10.3892/ol.2018.9787.
- [43] Q. Mao and J. D. Unadkat, "Role of the Breast Cancer Resistance Protein (BCRP/ABCG2) in Drug Transport—an Update," *AAPS J*, vol. 17, no. 1, pp. 65–82, Jan. 2015, doi: 10.1208/s12248-014-9668-6.
- [44] J.-Y. Lee *et al.*, "Candidate gene approach evaluates association between innate immunity genes and breast cancer risk in Korean women," *Carcinogenesis*, vol. 30, no. 9, pp. 1528–1531, Sep. 2009, doi: 10.1093/carcin/bgp084.
- [45] M. Tinholt *et al.*, "Increased coagulation activity and genetic polymorphisms in the F5, F10 and EPCR genes are associated with breast cancer: a case-control study," *BMC Cancer*, vol. 14, Nov. 2014, doi: 10.1186/1471-2407-14-845.
- [46] L. Cesaratto *et al.*, "BNC2 is a putative tumor suppressor gene in high-grade serous ovarian carcinoma and impacts cell survival after oxidative stress," *Cell Death & Disease*, vol. 7, no. 9, Art. no. 9, Sep. 2016, doi: 10.1038/cddis.2016.278.
- [47] M. Tsuboi *et al.*, "Prognostic significance of GAD1 overexpression in patients with resected lung adenocarcinoma," *Cancer Medicine*, vol. 8, no. 9, pp. 4189–4199, 2019, doi: 10.1002/cam4.2345.
- [48] International Schizophrenia Consortium *et al.*, "Common polygenic variation contributes to risk of schizophrenia and bipolar disorder," *Nature*, vol. 460, no. 7256, pp. 748–752, Aug. 2009, doi: 10.1038/nature08185.

- 1 [49] R. A. Scott *et al.*, “An Expanded Genome-Wide Association Study of Type 2 Diabetes in
2 Europeans,” *Diabetes*, May 2017, doi: 10.2337/db16-1253.
- 3 [50] M. LeBlanc and C. Kooperberg, “Boosting predictions of treatment success,” *Proc Natl*
4 *Acad Sci U S A*, vol. 107, no. 31, pp. 13559–13560, Aug. 2010, doi:
5 10.1073/pnas.1008052107.
- 6 [51] C. Angermueller, T. Pärnamaa, L. Parts, and O. Stegle, “Deep learning for computational
7 biology,” *Molecular Systems Biology*, vol. 12, no. 7, p. 878, Jul. 2016, doi:
8 10.15252/msb.20156651.
- 9 [52] J. Schmidhuber, “Deep learning in neural networks: An overview,” *Neural Networks*, vol.
10 61, pp. 85–117, Jan. 2015, doi: 10.1016/j.neunet.2014.09.003.
- 11