

Percolate

An Exponential Family JIVE Model to Design DNA-Based Predictors of Drug Response

Mourragui, Soufiane M.C.; Loog, Marco; van Nee, Mirrelij; de Wiel, Mark A.van; Reinders, Marcel J.T.; Wessels, Lodewyk F.A.

DOI

[10.1007/978-3-031-29119-7_8](https://doi.org/10.1007/978-3-031-29119-7_8)

Publication date

2023

Document Version

Final published version

Published in

Research in Computational Molecular Biology - 27th Annual International Conference, RECOMB 2023, Proceedings

Citation (APA)

Mourragui, S. M. C., Loog, M., van Nee, M., de Wiel, M. A. V., Reinders, M. J. T., & Wessels, L. F. A. (2023). Percolate: An Exponential Family JIVE Model to Design DNA-Based Predictors of Drug Response. In H. Tang (Ed.), *Research in Computational Molecular Biology - 27th Annual International Conference, RECOMB 2023, Proceedings* (pp. 120-138). (Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics); Vol. 13976 LNBI). Springer. https://doi.org/10.1007/978-3-031-29119-7_8

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.



Percolate: An Exponential Family JIVE Model to Design DNA-Based Predictors of Drug Response

Soufiane M. C. Mourragui^{1,2}, Marco Loog^{2,3}, Mirrelijin van Nee⁴,
Mark A van de Wiel^{4,5}, Marcel J. T. Reinders^{2,6}(✉),
and Lodewyk F. A. Wessels^{1,2}(✉)

¹ Computational Cancer Biology, Division of Molecular Carcinogenesis, Oncode Institute, The Netherlands Cancer Institute, Amsterdam, The Netherlands

l.wessels@onki.nl

² Faculty of EEMCS, Delft University of Technology, Delft, The Netherlands

m.j.t.reinders@tudelft.nl

³ Department of Computer Science, University of Copenhagen, Copenhagen, Denmark

⁴ Department of Epidemiology and Biostatistics, Amsterdam Public Health research institute, Amsterdam University medical centers, Amsterdam, The Netherlands

⁵ MRC Biostatistics Unit, University of Cambridge, Cambridge, UK

⁶ Leiden Computational Biology Center, Leiden University Medical Center, Leiden, The Netherlands

Abstract. Motivation: Anti-cancer drugs may elicit resistance or sensitivity through mechanisms which involve several genomic layers. Nevertheless, we have demonstrated that gene expression contains most of the predictive capacity compared to the remaining omic data types. Unfortunately, this comes at a price: gene expression biomarkers are often hard to interpret and show poor robustness.

Results: To capture the best of both worlds, i.e. the accuracy of gene expression and the robustness of other genomic levels, such as mutations, copy-number or methylation, we developed Percolate, a computational approach which extracts the joint signal between gene expression and the other omic data types. We developed an out-of-sample extension of Percolate which allows predictions on unseen samples without the necessity to recompute the joint signal on all data. We employed Percolate to extract the joint signal between gene expression and either mutations, copy-number or methylation, and used the out-of sample extension to perform response prediction on unseen samples. We showed that the joint signal recapitulates, and sometimes exceeds, the predictive performance achieved with each data type individually. Importantly, molecular signatures created by Percolate do not require gene expression to be evaluated, rendering them suitable to clinical applications where only one data type is available.

Availability: Percolate is available as a [Python 3.7 package](#) and the scripts to reproduce the results are available [here](#).

Supplementary Information The online version contains supplementary material available at https://doi.org/10.1007/978-3-031-29119-7_8.

© The Author(s) 2023

H. Tang (Ed.): RECOMB 2023, LNBI 13976, pp. 120–138, 2023.

https://doi.org/10.1007/978-3-031-29119-7_8

1 Introduction

Over the course of their lifespan, human cells accumulate molecular alterations that result in the modification of cell behavior [27]. When aggregated at the tissue level, these alterations can compromise tissue homeostasis, in turn clinically impacting a patient [13]. Understanding the combined effect of these alterations is key to designing bespoke lines of treatment [28, 33]. These molecular alterations occur at different genomic levels and are recorded using different technologies, collectively referred to as “omics” technologies. Each of these omic measurements offers only partial information regarding the compromised tissue. Aggregating different omic measurements, an analysis known as multi-omics integration, is therefore necessary to generate a comprehensive picture of the molecular features underlying a cancerous lesion [5, 20].

Owing to their high versatility, cell lines offer a cost-effective model system for drug response modelling [8]. Specifically, large scale consortia have industriously subjected a large number of cell lines to hundreds of different compounds, yielding valuable drug response measurements [12, 16, 32]. A key challenge resides in combining these response measurements with multi-omics data to study mechanisms of resistance and sensitivity [24]. Existing approaches focus on combining all omics data types and can be ordered based on the stage of the analysis at which the integration is performed [6]. At one extreme, early integration approaches [4, 19] first aggregate all features from all data types to process them all simultaneously. At the other extreme, late integration approaches first compute a representation of each data type individually, and subsequently combine these representations [11, 26, 36]. Several other methods can be positioned along this ordering, and differ by the analysis stage during which the grouping of data types is performed [41]. Although promising and encouraging, these methods do not take into account the quality of the data types and do not explicitly model their topology [2], i.e., how the data types relate to each other regarding information content and capacity to predict drug response. In particular, it has been observed that, although it has traditionally been the least clinically actionable data type, gene expression consistently prevails over other data types [9] and provides similar performance as early-integration approaches [1], obviating the need for complex integration strategies.

In order to maintain the predictive power of gene expression data, while exploiting the robustness of the most actionable data types, we present Percolate, an unsupervised multi-omics integration framework. Percolate sets itself apart from other integration approaches as it aims to eliminate gene expression measurements from the final predictor, rather than integrating it with all other data types. This is achieved by extracting the joint signal between gene expression and the other data types in an iterative fashion. First the joint signal between gene expression and Data type 1 (e.g. mutations) is extracted. Then the remaining signal (not shared with Data type 1) is employed to extract the joint signal of gene expression with Data type 2 (e.g. copy number data). This procedure is repeated for all omics data types. In this way, the gene expression signal is “percolated” down the other omics data types, ideally extracting all

predictive signal from the gene expression data. Technically, Percolate employs a popular framework, called JIVE [11,26], which breaks down paired datasets into joint and individual signals. We first extended JIVE to non-Gaussian noise models employing GLM-PCA [7]. Specifically, we used an alternative optimization, the decomposition of saturated parameters [21], which we theoretically proved to be competitive with the original formulation. Finally, we developed an out-of-sample extension for JIVE, useful when only one of the two data types is available.

We first show that comparing gene expression to other data types individually recovers a known topology of multi-omics data. We then show that the information shared between individual omic data types and gene expression increases drug response predictive performance for the individual omic data types. Finally, reconstructing the joint signal solely from mutation, copy-number and methylation, we show that the signatures derived from “percolating” gene expression down these data types recapitulate the drug response predictive performance of these data types.

2 Methods

2.1 Trade-off Between Robust and Predictive Types

We consider four data types: mutations (MUT), copy number aberrations (CNA), methylation (METH) and gene expression (GE). MUT and CNA directly measure genetic aberrations and therefore rely on DNA measurements. Due to several biological and technological factors, these measurements are highly robust and suffer from little technical artefacts. On the other end of the spectrum, GE measures RNA abundance, a process known for exhibiting large biological variability and prone to technical artefacts. Between these two extremes, methylation offers an intermediate level of robustness. However, when it comes to drug response prediction, the order is reversed: GE offers, on average, a better predictive performance than METH, and significantly outperforms MUT and CNA [1,8,17]. This leads to a trade-off between robustness and predictive ability (Fig. 1A) with MUT and CNA being the most robust and least predictive and GE being the most predictive and least robust, with METH rating at the intermediate level in terms of robustness and predictive capacity.

2.2 Exponential Family Distribution

Our integrated approach is inspired by AJIVE [11], a computational approach which takes as input two paired datasets and computes a joint and a data-specific signals. AJIVE is an extension of the JIVE model [26], which we selected, among other extensions [35,37], for its computational tractability and its mathematical formulation which is amenable to the derivation we propose. JIVE, AJIVE, and derivations thereof, critically rely on Principal Component Analysis (PCA) which assumes a Gaussian noise model on the data [22,39]. To extend this

Table 1. Exponential family distributions. Gaussian distribution is assumed to have unit variance. The dispersion parameter r is fixed for the Negative Binomial.

Data type	Family distribution
Copy-number aberration (CNA)	Log-Normal or Gamma
Gene expression (GE)	Negative-Binomial
Methylation (METH)	Beta
Mutation (MUT)	Bernoulli

framework to non-Gaussian settings, we make use of a generalized formulation that can deal with a wider class of parametric distribution models, i.e., the so-called exponential family [31].

Definition 2.1 (Exponential family distribution). *Let $\mathcal{X} \subset \mathbb{R}^p$, we say that a random vector $Z \in \mathcal{X}$ follows an **exponential family distribution** if its probability density function f can be written as*

$$\forall z \in \mathcal{X}, \quad f(z|\theta) = h(z) \exp \left(\eta(\theta)^T T(z) - A(\theta) \right). \quad (1)$$

$T : \mathcal{X} \rightarrow \mathbb{R}^q$ ($q > 0$) is called the **sufficient statistics**, $\theta \in \mathbb{R}^q$ the **exponential parameter**, $\eta : \mathbb{R}^q \rightarrow \mathbb{R}^q$ the **natural parametrization**, $A : \mathbb{R}^q \rightarrow \mathbb{R}$ the **log-partition function** and $h : \mathcal{X} \rightarrow \mathbb{R}^+$ the base measure.

The exponential family encompasses a broad set of distributions (Supp. Table 1), including the Gaussian distribution with unit variance, the Poisson, the Bernoulli, the Beta or the Gamma distributions. Practically, the functions A , T and η are modelling choices which can be tuned for any specific application.

2.3 Saturated Model Parameters

For this section, we consider a data matrix $X \in \mathbb{R}^{n \times p}$, with n (resp. p) the number of samples (resp. features). We model this data using an exponential family distribution $\mathcal{E} = (T, A, \eta)$ (Definition 2.1), which choice is motivated by prior knowledge. For instance, if the data is known to be binary, one would turn to $\mathcal{E}$ defined by the Bernoulli distribution, while another data distribution would lead to a different choice of functions (Supp. Table 1). We denote by q the dimensionality of T and A output space.

Definition 2.2 (Negative log-likelihood). *We define the **negative log-likelihood**, denoted $\mathcal{L}$, as follows:*

$$\forall \Theta \in \mathbb{R}^{n \times p \times q}, \quad \mathcal{L}(\Theta; X, \mathcal{E}) = \sum_{i=1}^n \sum_{j=1}^p A(\Theta_{i,j}) - \eta(\Theta_{i,j})^T T(X_{i,j}). \quad (2)$$

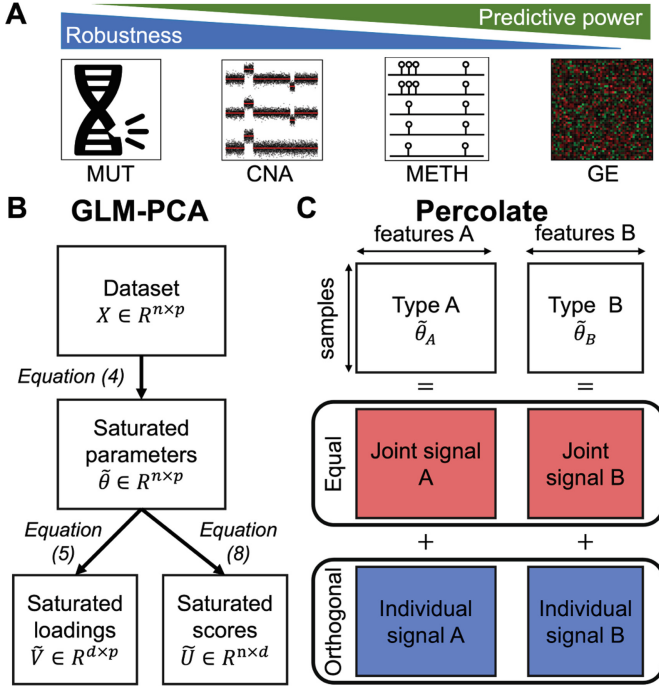


Fig. 1. Dissecting multi-omics topology using Percolate bridges the gap between predictive and robust data types. (A) Trade-off between robust data types (MUT, CNA) and predictive types (METH, GE). (B) Workflow of our implementation of GLM-PCA, which relies on the projection of saturated parameters. (C) Workflow of Percolate, which extends JIVE to non-Gaussian settings by comparing the low-rank structures of saturated parameter matrices.

Definition 2.3 (Saturated parameters). *We define the saturated parameters $\tilde{\Theta}(X, \mathcal{E}) \in \mathbb{R}^{n \times p \times q}$ as the minimizers of $\mathcal{L}$, i.e.,*

$$\tilde{\Theta}(X, \mathcal{E}) = \underset{\Theta \in \mathbb{R}^{n \times p \times q}}{\operatorname{argmin}} \mathcal{L}(\Theta; X, \mathcal{E}). \quad (3)$$

The saturated parameters correspond to single-sample maximum likelihood estimates. This quantity, which will be the pillar of our approach to GLM-PCA (Sect. 2.4), can be computed as follows.

Proposition 2.4 (Computation of saturated parameters). *Assume that A and ν are differentiable with invertible differentials. Then, denoting J as the Jacobian of a function:*

$$\tilde{\Theta}(X, \mathcal{E}) = \eta^{-1} \circ (J_{A \circ \eta^{-1}})^{-1} \circ T(X) \hat{=} g^{-1}(X), \quad (4)$$

using an element-wise operation on all elements of X .

Proof. We refer the reader to the Supplementary Material (Sect. 4) for the proof. ■

Proposition 2.4 shows that the saturated parameters correspond to a dual representation of the data motivated by prior knowledge on the data-distribution. We will exploit this representation ‘*a la PCA*’ to find the main sources of variations in a framework called GLM-PCA.

2.4 Generalized Linear Model Principal Component Analysis (GLM-PCA)

JIVE is based on Principal Component Analysis (PCA), which admits three equivalent definitions: maximization of projected variance, minimization of reconstruction error and maximization of a Gaussian likelihood with unit-variance. This latter definition can be restrictive for non-Gaussian data and we therefore set out to replace PCA by an extension called **GLM-PCA** [7]. In these methods, the Gaussian likelihood is replaced by an exponential family distribution. The original approach from *Collins et al.* [7] minimizes a negative log-likelihood using an SVD-like decomposition for the exponential parameters, yielding three different matrices. Refinements of this idea, which solve a similar optimization problem, have been proposed in the literature [23, 25] and offer competitive routines for the computation of these three matrices. Another take on this problem, which relies on the projection of saturated parameters, has recently been developed by *Landgraf et al.* [21]. This approach offers the advantage of a simpler single-matrix optimization instead of concomitantly optimizing on three. Furthermore, the out-of-sample extension relies on a matrix multiplication and is thus computationally fast. These two approaches therefore aim at finding the same decomposition through different computational routines. We here present these two approaches and prove that the latter offers a similar or better minimizer for the negative log-likelihood, which, to the best of our knowledge, had not been established.

2.4.1 Two Formulations of GLM-PCA

Definition 2.5 (SVD-type [7]). *SVD-type GLM-PCA computes three matrices, $U_{SVD} \in \mathbb{R}^{n \times d}$, $V_{SVD} \in \mathbb{R}^{p \times d}$ and $\Sigma_{SVD} \in \mathbb{R}^{d \times d}$ (diagonal), alongside a vector $\mu_{SVD} \in \mathbb{R}^p$ defined as*

$$U_{SVD}, V_{SVD}, \Sigma_{SVD}, \mu_{SVD} \hat{=} \underset{\substack{U, V, \Sigma, \mu \\ V^T V = U^T U = I_d}}{\operatorname{argmin}} \mathcal{L}(U \Sigma V^T + 1_n \mu^T; X, \mathcal{E}) \quad (5)$$

Definition 2.6 (Projection of saturated parameters [21]). *GLM-PCA by projection of saturated parameters computes one matrix, $V_{SP} \in \mathbb{R}^{p \times d}$ alongside a vector $\mu_{SP} \in \mathbb{R}^p$ defined as*

$$V_{SP}, \mu_{SP} \hat{=} \underset{\substack{V \in \mathbb{R}^{p \times d}, \mu \in \mathbb{R}^p \\ V^T V = I_d}}{\operatorname{argmin}} \mathcal{L}\left(\left(\tilde{\Theta}(X; \mathcal{E}) - 1_n \mu^T\right) V V^T + 1_n \mu^T; X, \mathcal{E}\right), \quad (6)$$

The loading matrices (V_{SVD} and V_{SP}) and the score matrix (U_{SVD}) have orthogonal constraints, which is similar to PCA where scores are by construction uncorrelated.

2.4.2 Equivalence of the Formulations

We here show that the projection of saturated parameters provides a competitive minimization when compared to the SVD-type decomposition. The main result is based on Supp. Lemma 5.1 and we refer the reader to the Supplementary Material (Sect. 5) for a complete proof.

Theorem 2.7. *Let us define $U_{SVD}, V_{SVD}, \Sigma_{SVD}$ and μ_{SVD} as in Definition 2.5, and V_{SP}, μ_{SP} as in Definition 2.6. The likelihood resulting from the two optimization processes satisfies*

$$\mathcal{L}(U_{SVD}\Sigma_{SVD}V_{SVD}^T + 1_n\mu_{SVD}^T) \geq \mathcal{L}\left(\left(\tilde{\Theta} - 1_n\mu_{SP}^T\right)V_{SP}V_{SP}^T + 1_n\mu_{SP}^T\right), \quad (7)$$

where the dependencies on X and $\mathcal{E}$ for $\mathcal{L}$ and $\tilde{\Theta}$ were removed for verbosity's sake.

Theorem 2.7 shows that, although the two approaches compute the same decomposition, the one obtained from saturated parameters yields a lower or equal negative log-likelihood. It is also worth noting that the SVD-like optimization is usually performed by alternate optimization [40] and the initialization can play a major role in the convergence. The projection of saturated parameters only requires one minimization round, and is thus faster and less prone to initialization effects. Using the decomposition of saturated parameters, however, comes at a price: there is an infinity of solutions, all equal up to a unitary transformation. In order to obtain sample scores that are uncorrelated, we proceed as follows.

Definition 2.8 (Sample scores). *Let V_{SP} and μ_{SP} be defined as in Definition 2.6 and assume that $\text{rank}\left(\tilde{\Theta} - 1_n^T\mu_{SP}\right) \geq d$. Then $\text{rank}\left[\left(\tilde{\Theta} - 1_n^T\mu_{SP}\right)V_{SP}V_{SP}^T\right] = d$ and we define U_{SP}, Σ_{SP} and W_{SP} as the unique rank- d SVD decomposition of the saturated parameters, i.e.*

$$U_{SP}\Sigma_{SP}W_{SP}^T = \left(\tilde{\Theta} - 1_n^T\mu_{SP}\right)V_{SP}V_{SP}^T. \quad (8)$$

It is worth noting that the equality in Eq. 8 is not an approximation and this second SVD does not entail any loss of information. It is a pure computational maneuver to whiten the obtained scores.

2.4.3 Hyper-parameter Optimization

The solution of Eq. (6) is an optimization problem with a Stiefel-manifold constraint, which we solved by using recent advances in auto-differentiation [30]

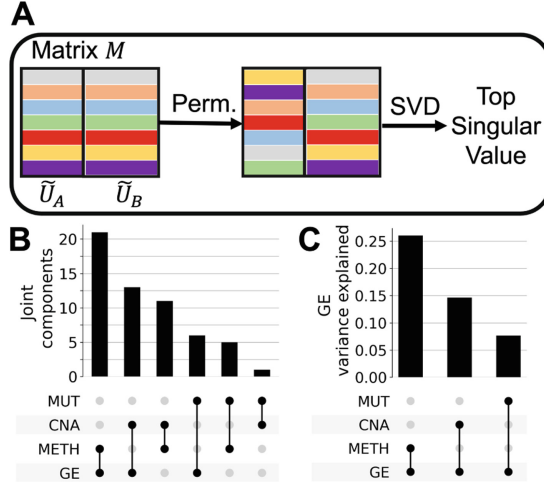


Fig. 2. Assessing the number of joint components. (A) Schematic of the sample-level permutations we perform to estimate the number of joint components. (B) Venn-diagram of the number of joint components obtained using the permutation scheme. (C) Ratio of variance explained for the GE saturated parameters matrix after projection on the joint components.

and optimization on Riemannian manifolds [29]. We modelled the functions A , T and the negative log-likelihood using PyTorch; stochastic gradient descent (SGD) on the Stiefel-manifold was performed using McTorch. Such a formulation allows to employ a large variety of exponential family distributions without the need for heavy and potentially cumbersome Lagrangian computations. Our optimization scheme relies on four hyper-parameters: number of factors (or principal components), learning rate, number of epochs and batch size. To determine them, we compute the Akaike Information Criterion (AIC) of the complete data for various values of d and different hyper-parameters [3]. For a GLM-PCA model with d PCs, the AIC corresponds to the sum of the data log-likelihood and the number of model parameters, which we estimate as the dimensionality of the Stiefel manifold $\{V \in \mathbb{R}^{d \times p} | VV^T = I_d\}$, equal to $pd - d(d+1)/2$. Among all trained models, we select the one which harbors the smallest AIC.

2.5 Comparison of GLM-PCA Directions by Percolate

Setting: We consider two datasets $X_A \in \mathbb{R}^{n \times p_A}$ and $X_B \in \mathbb{R}^{n \times p_B}$ with paired samples (rows) but potentially different features. We first perform GLM-PCA independently on X_A and X_B using two different exponential family distributions, yielding d_A and d_B factors, respectively denoted as $\tilde{V}_A$ and $\tilde{V}_B$. We furthermore denote by $\tilde{\Theta}_A$ and $\tilde{\Theta}_B$ the saturated parameters of datasets A and B respectively, and $\tilde{\mu}_A$ and $\tilde{\mu}_B$ the intercept terms. Using the decomposition presented in Definition 2.8, we furthermore define $\tilde{U}_A, \Sigma_A, \tilde{W}_A$ and $\tilde{U}_B, \Sigma_B, \tilde{W}_B$.

Definition 2.9. To compare the two sets of samples scores, $\tilde{U}_A$ and $\tilde{U}_B$, we aggregate them in a matrix $\mathbf{M}$, which we decompose by SVD:

$$\mathbf{M} = \begin{bmatrix} \tilde{U}_A & \tilde{U}_B \end{bmatrix} = U_M \Sigma_M V_M^T. \quad (9)$$

The top left-singular vectors correspond to sample scores which are highly correlated between $\tilde{U}_A$ and $\tilde{U}_B$, since both of these two matrices are consisting, by construction, of uncorrelated factors. Following the same intuition as in AJIVE, these can be understood as the **joint signal**, motivating the following definition.

Definition 2.10 (Joint and individual signals). Let $r_J < \min(d_A, d_B)$, we define the **joint signal** as the matrix $\tilde{U}_J \in \mathbb{R}^{n \times r_J}$ with the top r_J left-singular values of $\mathbf{M}$. We furthermore denote by Σ_J the diagonal matrix with the top r_J singular values of $\mathbf{M}$.

We define the **individual signal** of A (resp. B), denoted as $\tilde{U}_I^A$ (resp. $\tilde{U}_I^B$), as the signal from $\tilde{U}_I^A$ (resp. $\tilde{U}_I^B$) not present in $\tilde{U}_A$ (resp. $\tilde{U}_B$), formally:

$$\begin{aligned} \tilde{U}_I^A &= \left(I_n - \tilde{U}_J \tilde{U}_J^T \right) \tilde{U}_A \\ \tilde{U}_I^B &= \left(I_n - \tilde{U}_J \tilde{U}_J^T \right) \tilde{U}_B. \end{aligned} \quad (10)$$

We call the complete process **Percolate**, and a summarised workflow can be found in Fig. 1B-C.

In order to set the number of joint components r_J , we employ a sample-level permutation scheme. We first independently permute the rows of $\tilde{U}_A$ and $\tilde{U}_B$, which we then aggregate as in Eq. (9) to obtain the singular values. We perform 100 such permutations independently and retrieve the first singular value for each. Finally, we set r_J as the number of elements in Σ_M over one standard deviation from the mean of the permuted singular values (Fig. 2A).

2.6 Projector of Joint Signal

AJIVE does not provide an out-of-sample extension, and we here propose a derivation thereof by rewriting the matrix U_J as a function of the saturated parameters.

Theorem 2.11. Let's decompose the matrix V_M as $V_M = [V_{M,A}^T \ V_{M,B}^T]^T$ such that $V_{M,A}^T$ contains the first d_A columns of V_M^T and $V_{M,B}^T$ the last d_B ones, we obtain:

$$\begin{aligned} \tilde{U}_J &= \tilde{U}_{J,A} + \tilde{U}_{J,B} \\ \text{with } \begin{cases} \tilde{U}_{J,A} &= \left(\tilde{\Theta}_A - 1_n \tilde{\mu}_A^T \right) \tilde{V}_A^T \tilde{V}_A W_A \Sigma_A^{-1} V_{M,A} \Sigma_J^{-1} \\ \tilde{U}_{J,B} &= \left(\tilde{\Theta}_B - 1_n \tilde{\mu}_B^T \right) \tilde{V}_B^T \tilde{V}_B W_B \Sigma_B^{-1} V_{M,B} \Sigma_J^{-1} \end{cases} \end{aligned} \quad (11)$$

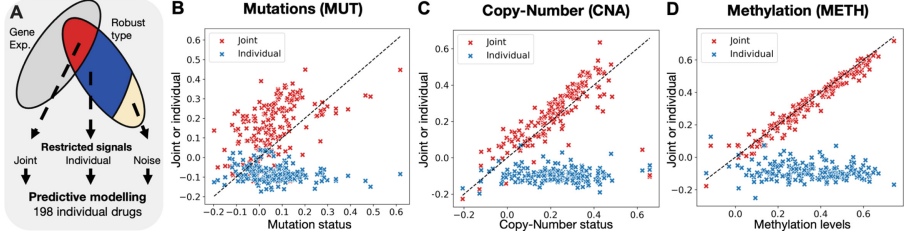


Fig. 3. The joint signal between robust and gene expression contains most of the predictive signal. (A) Workflow of our approach. (B) Predictive performance for MUT when using Percolate between MUT and GE. Each point corresponds to a single drug, with the x-axis corresponding to the predictive performance obtained using the original mutation data, and the y-axis by either the joint (red) or the individual (blue) signals. (C) Predictive performance for CNA, similarly displayed as in B. (D) Predictive performance for METH, similarly displayed as in B.

Proof. We refer the reader to the Supplementary Material (Sect. 6) for the complete proof. ■

The formulation of $\tilde{U}_J$ presented in Equation (11) highlights the additive contribution of both dataset to the joint signal. At test time, both views are therefore required to estimate the joint signal. To tackle the issue of missing data-view, we propose a nearest-neighbor imputation of the unknown joint-term. Let's consider, without loss of generality, that only the view A is available. The joint signal has been computed using the two data matrices X_A and X_B , yielding $\tilde{U}_{J,A}$ and $\tilde{U}_{J,B}$. The second term contains r_J terms, and we train r_J corresponding k-Nearest-Neighbors (kNN) regressors. The test dataset $Y_A \in \mathbb{R}^{m \times p_A}$ can be projected on the joint signal by replacing the saturated parameter $\tilde{\Theta}_A$ in Eq. 11 with the saturated parameter of the test data. We then estimate the second term by means of the r_J kNN regression models. Adding these two terms yields an estimate of the joint signal.

2.7 Drug Response Prediction

We assess the predictive performance of a dataset by employing ElasticNet [42], which has been shown, inspite of its relative simplicity, to outperform more complex non-linear models when it comes to drug response prediction [8, 17, 38]. For a given dataset, we perform nested cross-validation as follows. First, datasets are stratified into 10 groups of equal size. For each group (10%), we employ a 3-fold cross-validation grid search on the remaining 90% to determine the optimal ElasticNet hyper-parameters (ℓ_1 -ratio and penalization). We then fit this optimal ElasticNet model on the 90% to predict the class labels on the 10%. Repeating this procedure, we obtain one cross-validated estimate per sample and we define the **predictive performance** as the Pearson correlation between these estimates and the actual values.

2.8 Data Download, Modelling and Processing

We consider four data types in our analysis (Table 1) which we modelled using different exponential family distributions (Supp. Material). The GDSC data was accessed on January 2020 from CellModel Passport [16]. For GE, MUT and CNA, we restricted to protein coding genes known to be frequently mutated in cancer, referred to as the **mini-cancer genome** [15]. GE was corrected for library size using TMM normalization [34] and mutations were restricted to non-silent.

3 Results

3.1 The Breakdown of the Joint Signals Highlights the Topology of Multi-omics Data

To compare data types, we employ Percolate using the distributions defined in Table 1, and a number of PCs set using the procedure presented in Subsect. 2.4 (Supp. Figure 2). For each comparison, setting the number of joint components is a crucial step, as it defines the threshold between the joint and individual signals. For that purpose, we used a sample level permutation test (Fig. 2A, Subsect. 2.5).

We observe that GE shares 21 joint components with METH, 13 with CNA and only 6 with MUT, which is coherent with the gradient put forward in Fig. 1. We furthermore observe that MUT is consistently the data type with the least number of joint components (Fig. 2B), highlighting the weakness of the signal coming from MUT data, corroborating previous measured topologies of multi-omics data [2]. To measure the strength of the underlying joint signals, we computed the proportion of GE variance explained by the joint directions (Fig. 2C), computed as the ratio between the joint signal variance and the variance of the GE's saturated parameters matrix. We observe that the joint signal between GE and METH explains 26% of GE variance, while this figures drops to 14% and 7% for CNA and MUT, respectively. These observations highlight the existence of a joint signal, of which the predictive performance can be interrogated.

3.2 Robust Signal Predictive of Drug Response Is Concentrated in the Joint Part

We then investigated the relevance of the joint and individual signals when it comes to drug response prediction. Considering one robust data type at a time (MUT, CNA or METH), we first decomposed the original robust data type into a signal joint with GE and an individual signal specific to the robust data type. We then computed, for 195 drugs (Methods), the predictive performance for these two signals and compared it to the original robust robust data (Fig. 3A, Subsect. 2.7). To ensure a proper comparison between joint, individual and cell-view, the cross-validation was performed using the same folds for all datasets. As ElasticNet has been shown in the literature to outperform other more advanced algorithms for this particular task [8, 17, 38], we restricted our comparison to

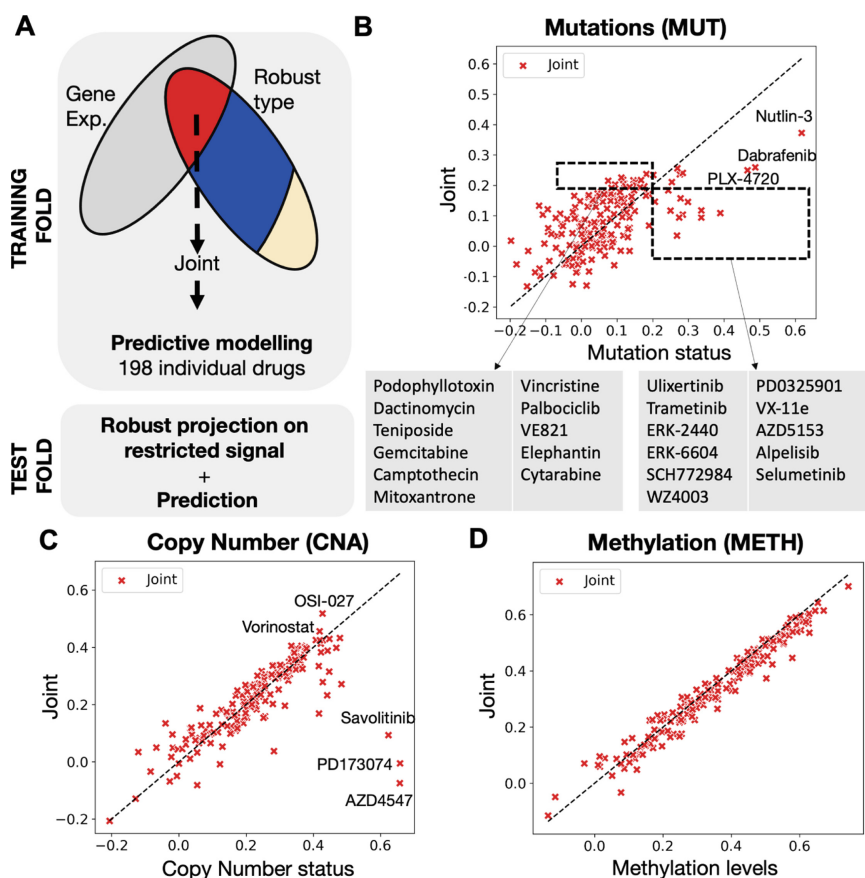


Fig. 4. Robust-type-based signatures created from Percolate recapitulate drug response. (A) Schematic of the cross validation experiment. (B) Results for MUT with a special zoom on drugs predictive for joint but not robust (left) and for robust but not joint (right). (C) Results for CNA. (D) Results for METH.

this regression method. Such experimental design has the advantage to properly assess the effect of Percolate, as no additional performance can be gained from the regression model.

We first analyzed the results obtained between MUT and GE data (Fig. 3B). We observe that for most drugs, the predictive performance of the joint signal exceeds the predictive performance of the original robust signal, except for a number of drugs of which the response is quite well predicted based on MUT only. This set includes the drugs Nutlin-3, Dabrafenib, and PLX-4720. In contrast, the individual signal shows no predictive performance (Pearson correlation below 0) for most drugs, indicating an absence of drug response related signal in the individual portion. We then turned to CNA where the choice of distribution was unclear, with, to the best of our knowledge, no clear precedent on how to

model such data. Due to the observed behavior of CNA data, we opted for two possible distributions: Log-normal and Gamma distributions (Supp. Table 1). We observe that the joint signal computed using a Gamma-distribution yields better performances than the log-normal model (Supp. Figure 3A-B). When using a Gamma distribution, a conclusion similar to the MUT data can be reached with the majority of drugs predicted well with the joint signal except three drug, AZD4547, PD173074 and Savolitinib (Fig. 3C). This advocates for using the Gamma distribution for analyzing CNA data and shows that the joint signal presents an increased performance while the individual signal is not predictive. Finally, we studied the drug response performance obtained after decomposing METH using GE (Fig. 3D). We observe that the joint signal presents a similar predictive performance as the original methylation data. The individual signal is, again, not predictive. These results highlight the potential of restricting predictors to the joint signal for robust data types.

3.3 Out-of-sample Extension Recapitulates the Predictive Performance of Robust Signal

In order to compute the joint signal between one robust data type and GE, one needs to have access to both modalities. However, the purpose is to become independent of non-robust GE measurements. In order to study whether the joint signal could be estimated without access to gene expression, when the predictor is applied to a test case, we exploited our out-of-sample extension (Subsect. 2.6). We employed this algorithm to compute the drug response predictive performance of the joint signal estimated using the robust data alone (Fig. 4A). Dividing the data in ten independent folds, we performed a cross-validation estimation as follows. For each train-test division of the data, we trained a Percolate instance on the 90% of the data, the training set containing GE and the robust data type. The resulting joint information was then used to train an ElasticNet model to predict drug response. The remaining 10% (test data) were then used to first estimate the joint signal, solely based on the robust data (Subsect. 2.6). This joint signal was then used as input into the ElasticNet model to predict the response on this test set. Finally, we computed the predictive performance as indicated in Subsect. 2.7.

When analyzing results for MUT (Fig. 4B), we first observe a clear drop in performance for the joint signal compared to the previous results (Fig. 3B). This suggests that the GE portion of the joint signal (Eq. 11) contains a significant portion of predictive signal, which is less well captured by our out-of-sample extension. Nonetheless, we observe that 11 drugs show a predictive performance above 0.2 for joint but not for the robust data. In contrast, 11 drugs show the opposite effect, including seven which target the MAPK pathway – MEK (Trametinib, PD0325901, Selumetinib) and ERK (ERK2440, ERK6604, Ulixertinib, SCH772984). BRAF inhibitors Dabrafenib and PLX-4720 also show a drop in performance. This suggests that constitutive activation of the MAPK pathway is not recapitulated by the joint signal. Nonetheless, the joint signal generated by Percolate helps increase performance for several poorly predictive

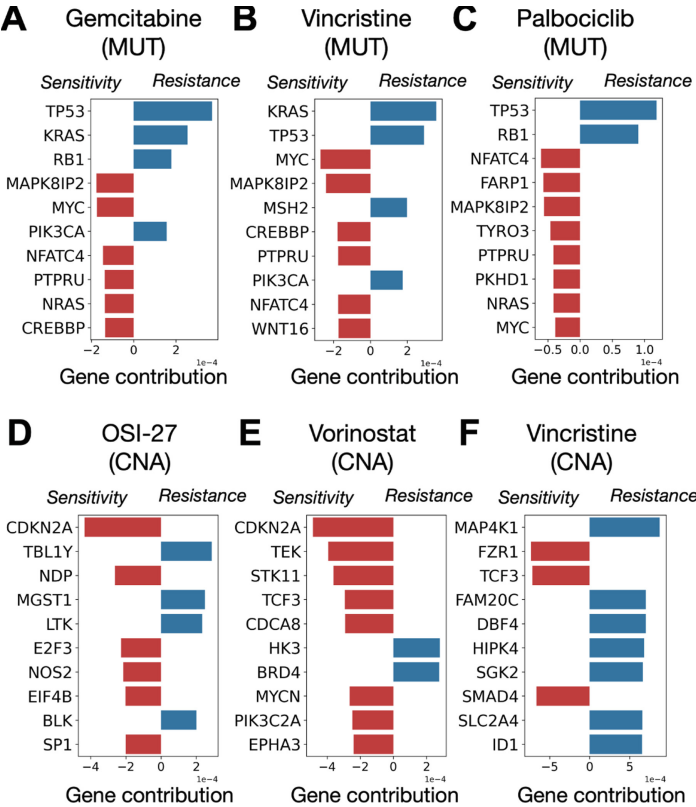


Fig. 5. Study of joint signals contributing to improved performance. For each drug, we report the top 10 largest gene regression coefficients from the joint signal, in absolute values. We first analysed the joint biomarkers created from MUT data for Gemcitabine (A), Vincristine (B) and Palbociclib (C). We then turned to CNA-based signatures for OSI-27 (D), Vorinostat (E) and Vincristine (F).

drugs and is therefore of interest to study various response mechanisms. We then turned to CNA (Fig. 4C) and observe a modest decrease in predictive performance compared to the performance on the original CNA profiles. Three drugs show a spectacular drop as the response can not be predicted by the joint signal – Savolitinib (cMET), PD173074 (FGFR) and AZD4547 (FGFR). In contrast, three drugs show improved performance for the joint signal – OSI-027 (mTOR), Navitoclax (HDAC) and Vincristine (tubulin). Finally, we repeated the experiment for METH (Fig. 4D) and observe that predictive performances of the joint signal is remarkably comparable to the predictive performance on the original METH data, with most drugs falling showing less than 2% relative performance difference (Supp. Figure 4C). Taken together, these results show that the joint signal recapitulates the drug response performance abilities of DNA-based measurements.

3.4 Study of Genes Contributing to the Joint Signals

We then set out to study the underlying mechanisms associated with the predictors derived from the robust data types (Subsect. 3.3) which also lead to improved performance. For a given drug, we trained an ElasticNet model on the joint signal, yielding one regression coefficient per joint component. Using the relationship from Eq. 11, we obtain a regression coefficient for each gene. A positive coefficient indicates that larger values of the saturated parameters, caused by a mutation or amplification of the supporting gene, are associated with resistance. In contrast, a negative coefficient indicates that larger values of the saturated parameters are associated with sensitivity.

For MUT, we studied the mode of action of three drugs for which the joint signal performs well (Fig. 4B): Gemcitabine (Fig. 5A), Vincristine (Fig. 5B) and Palbociclib (Fig. 5C). We observe that TP53 mutation status is associated with resistance to three drugs, concordant with earlier observations showing that TP53 mutant are more resistant to chemotherapy [14]. Resistance to Gemcitabine and Vincristine is also associated with KRAS and PI3KCA mutations, known for their proliferative potential [10, 18]. Interestingly, mutations in MYC and MAPK8IP2 are associated with sensitivity to these three drugs. Three other drugs show a drop in predictive performance on the joint signal as compared to the original signal: Nutlin-3, Dabrafenib and PLX-4720 (Fig. 4B). We observe that the known targets of these drugs exhibit a large coefficient: TP53 for Nultin-3 (known resistance biomarker) and BRAF for Dabrafenib and PLX-4720 (Supp. Figure 5). These three drugs highlight a limitation of our approach: GLM-PCA generates scores which aggregates the contributions of several genes. Highly-specific drugs, like Nutlin-3 (Mdm2-inhibitor) or BRAF/MEK-inhibitors not only target a specific protein, but mutations in the target are excellent response predictors. Such cases do not benefit from the GLM-PCA aggregation as a single feature alone is predictive.

Next we turned to CNA where three drugs: OSI-27 (Fig. 5D), Vorinostat (Fig. 5E) and Vincristine (Fig. 5F), which all showed increased performance when the joint signal is employed as compared to the original CNA data. For both OSI-27 (mTORC1) and Vorinostat (HDAC), we observe that amplification of CDKN2A (p16) is associated with sensitivity. P16 acts as a tumor-suppressor by slowing down the early progression of the cell-cycle and its loss is here associated with resistance for these two drugs. Finally, Vincristine's predictor shows that MAP4K1's amplification as a predictor of resistance. Such result is coherent with what we observed for MUT (Fig. 5B) where mutations on KRAS were associated with resistance.

3.5 Iterative Application of Percolate Deprives Gene Expression from Predictive Power

Finally, we questioned whether some signal predictive of drug response is still present in gene expression. To this end, we studied the GE signal after it has been

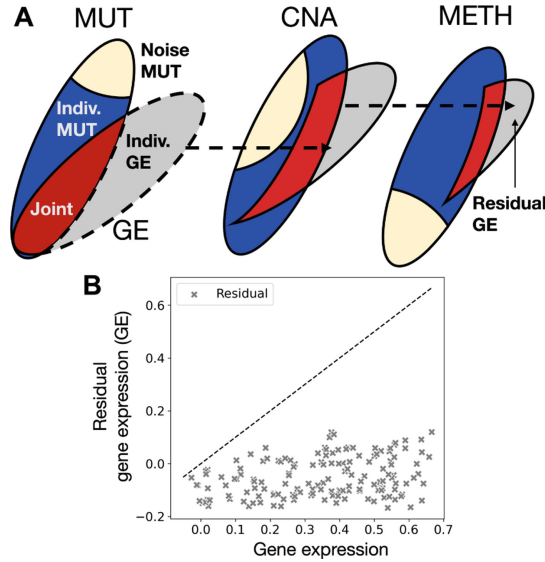


Fig. 6. The signal joint with DNA-based measurements deprives gene expression from any predictive power. (A) Schematic of our iterative procedure to remove from GE any signal joint with robust data type. **(B)** Predictive performance of the resulting residual gene expression compared to the predictive performance of the complete gene expression.

stripped of all the signal it shares with MUT, METH or CNA. To remove all signal associated with robust data types from GE, we used Percolate iteratively on GE, starting with the *least* predictive data type (MUT), followed by CNA and ending with the *most* predictive data type (METH) (Fig. 6A). Specifically, we first “percolate” GE through MUT to obtain an individual GE signal (not shared with MUT), which is then percolated through CNA to obtain a second GE individual signal, which is then finally percolated through METH, resulting in the individual GE signal we denote as *residual gene expression*. We finally assessed the predictive performance of this residual gene expression and compared it to the predictive performance of the original GE (Fig. 6B, Subsect. 2.7). We observe that no drug reaches a Pearson correlation above 0.16, indicative of a complete lack of predictive performance in the residual GE. This shows that removing the signal joint with DNA-based measurements deprives gene expression from any predictive ability.

4 Discussion

Designing multi-omics predictors of drug response has highlighted the existence of a trade-off between robust and predictive data types. To study this trade-off, we developed Percolate, a method which decomposes a pair of data types

into a joint and an individual signal. After showing that the strength of the joint signal recapitulates the known topology between data types, we showed that the joint signal contains more predictive power than any robust data type alone. Exploiting our out-of-sample extension, we showed that the joint signal, computed from robust data types alone, recapitulates most of the predictive performance of each original robust signal. Finally, we showed that the gene expression signal predictive of drug response is fully captured by robust data types through Percolate.

Although encouraging, our results display certain limitations that could inspire future methodological improvements. A key direction lies in the drop of performance between Fig. 4 and Fig. 5, caused by the out-of-sample extension. We theoretically decomposed the joint signal (Theorem 2.7) and presented an approach to approximate, using the robust type, the contribution from gene expression. We believe that this step can be improved in two ways: either by increasing the sample-size, thereby expanding the pool of potential anchors, or through the design of novel regression approaches. Another important improvement would be to extend this methodology to unpaired (single-cell) multi-omic measurements where characterizing the joint signal between omic datasets is a critical step.

Technically, Percolate extends JIVE in two different ways. First, by using GLM-PCA instead of PCA, we tailor the dimensionality reduction step to the specific data under consideration. Second, we developed an out-of-sample extension which allows to estimate the joint signal, even in the absence of one data-modality. For our analysis, we made use of standard distributions from the exponential family: Negative Binomial, Gamma, Beta or Bernoulli. Our implementation of GLM-PCA is versatile and any exponential family distribution can be employed in our framework, provided it can be auto-differentiated by PyTorch. Employing more complex distribution, like the inverse-gamma for copy-number is a fruitful avenue to improve on our methodology.

Acknowledgements and Funding. We wish to thank Stavros Makrodimitris (TU Delft) for useful discussions, and the RHPC of the NKI for the computing infrastructure. This work was funded by ZonMw TOP grant COMPUTE CANCER (40-00812-98-16012).

References

1. Aben, N., et al.: TANDEM: a two-stage approach to maximize interpretability of drug response models based on multiple molecular data types. *Bioinformatics* **32**(17), i413–i420 (2016)
2. Aben, N., et al.: ITOP: inferring the topology of omics data. *Bioinformatics* **34**(17), i988–i996 (2018)
3. Akaike, H.: A new look at the statistical model identification. *IEEE Trans. Autom. Control* **19**(6), 716–723 (1974)
4. Argelaguet, R., et al.: Multi-omics factor analysis-a framework for unsupervised integration of multi-omics data sets. *Mol. Syst. Biol.* **14**(6), 1–13 (2018)

5. Bersanelli, M., et al.: Methods for the integration of multi-omics data: mathematical aspects. *BMC Bioinform.* **17**(2), 167–177 (2016)
6. Cantini, L., et al.: Benchmarking joint multi-omics dimensionality reduction approaches for the study of cancer. *Nat. Commun.* **12**(1), 1–12 (2021)
7. Collins, M., et al.: A generalization of principal component analysis to the exponential family. *NeurIPS* **14**, 1–8 (2002)
8. Costello, J.C., et al.: A community effort to assess and improve drug sensitivity prediction algorithms. *Nat. Biotechnol.* **32**(12), 1202–1212 (2014)
9. Dempster, J.M., et al.: Gene expression has more power for predicting in vitro cancer cell vulnerabilities than genomics. *BioRxiv* **1**, 1–42 (2020)
10. Eklund, E.A., et al.: KRAS mutations impact clinical outcome in metastatic non-small cell lung cancer department of surgery. *Cancer* **14**, 2063 (2022)
11. Feng, Q., et al.: Angle-based joint and individual variation explained. *J. Multivar. Anal.* **166**, 241–265 (2018)
12. Ghandi, M., et al.: Next-generation characterization of the cancer cell line encyclopedia. *Nature* **569**(7757), 503–508 (2019)
13. Hanahan, D., et al.: Hallmarks of cancer: the next generation. *Cell* **144**(5), 646–674 (2011)
14. Hientz, K., et al.: The role of p53 in cancer drug resistance and targeted chemotherapy. *Oncotarget* **8**(5), 8921–8946 (2017)
15. Hoogstraat, M., et al.: Genomic and transcriptomic plasticity in treatment-Naïve ovarian cancer. *Genome Res.* **24**(2), 200–211 (2014)
16. Iorio, F., et al.: A landscape of pharmacogenomic interactions in cancer. *Cell* **166**, 740–754 (2016)
17. Jang, I.S., et al.: Systematic assessment of analytical methods for drug sensitivity prediction from cancer cell line data. *Pac. Symp. Biocomput.* **23**, 1–7 (2013)
18. Kim, S.T., et al.: Impact of KRAS mutations on clinical outcomes in pancreatic cancer patients treated with first-line gemcitabine-based chemotherapy. *Mol. Cancer Ther.* **10**(10), 1993–1999 (2011)
19. Kim, Y., et al.: WON-PARAFAC: a genomic data integration method to identify interpretable factors for predicting drug-sensitivity in-vivo, pp. 1–30
20. Kristensen, V.N., et al.: Kristensen - Principles and methods of integrative genomic analyses in cancer.pdf. *Nat. Rev. Cancer.* **14**, 299–313 (2014)
21. Landgraf, A.J., et al.: Dimensionality reduction for binary data through the projection of natural parameters. *J. Multivar. Anal.* **180**, 2020 (1999)
22. Lawrence, N.: Probabilistic non-linear principal component analysis with Gaussian process latent variable models. *J. Mach. Learn. Res.* **6**, 1783–1816 (2005)
23. Li, J., et al.: Simple exponential family PCA. *IEEE Trans. Neural Netw. Learn. Syst.* **24**(3), 485–497 (2013)
24. Li, Y., et al.: A review on machine learning principles for multi-view biological data integration. *Brief. Bioinform.* **19**(2), 325–340 (2018)
25. Liu, L.T., et al.: ePCA: high dimensional exponential family PCA. *Ann. Appl. Statist.* **12**(4), 2121–2150 (2018)
26. Lock, E.F., et al.: Joint and individual variation explained (JIVE) for integrated analysis of multiple data types. *Ann. Appl. Statist.* **7**(1), 523–542 (2013)
27. Martincorena, I., et al.: Somatic mutation in cancer and normal cells. *Science* **349**(6255), 1483–1489 (2015)
28. McLeod, H.L.: Cancer pharmacogenomics: early promise, but concerted effort needed. *Science* **340**(6127), 1563–1566 (2013)
29. Meghwanshi, M., et al.: McTorch, a manifold optimization library for deep learning, pp. 1–5 (2018)

30. Paske, A., et al.: Automatic differentiation in prose. In: NeurIPS (2017)
31. Pitman, E.J.: Sufficient statistics and intrinsic accuracy. *Math. Proc. Cambridge Philos. Soc.* **32**(4), 567–579 (1936)
32. Rees, M.G., et al.: Correlating chemical sensitivity and basal gene expression reveals mechanism of action. *Nat. Chem. Biol.* **12**(2), 109–116 (2016)
33. Relling, M.V., et al.: Pharmacogenomics in the clinic. *Nature* **526**(7573), 343–350 (2015)
34. Robinson, M.D., et al.: edgeR: a bioconductor package for differential expression analysis of digital gene expression data. *Bioinformatics* **26**(1), 139–140 (2009)
35. Sagonas, C., et al.: Robust joint and individual variance explained. *CVPR* **5739–5748**, 2017 (2017)
36. Sharifi-Noghabi, H., et al.: MOLI: multi-omics late integration with deep neural networks for drug response prediction. *Bioinformatics* **35**(14), i501–i509 (2019)
37. Shu, H., et al.: D-CCA: a decomposition-based canonical correlation analysis for high-dimensional datasets. *J. Am. Stat. Assoc.* **115**(529), 292–306 (2020)
38. Smith, A.M., et al.: Standard machine learning approaches outperform deep representation learning on phenotype prediction from transcriptomics data. *BMC Bioinform.* **21**(1), 1–18 (2020)
39. Tipping, M.E., et al.: Probabilistic principal component analysis. *J. Roy. Stat. Soc. B* **61**(3), 611–622 (1999)
40. Townes, F.W., Hicks, S.C., Aryee, M.J., Irizarry, R.A.: Feature selection and dimension reduction for single-cell RNA-SEQ based on a multinomial model. *Genome Biol.* **20**(1), 1–16 (2019)
41. Wang, B., et al.: Similarity network fusion for aggregating data types on a genomic scale. *Nat. Methods* **11**(3), 333–337 (2014)
42. Zou, H., et al.: Regularization and variable selection via the elastic net Hui. *J. Statist. Soc. Ser. B* **67**(2), 301–320 (2005)

Open Access This chapter is licensed under the terms of the Creative Commons Attribution 4.0 International License (<http://creativecommons.org/licenses/by/4.0/>), which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

