

METHODOLOGY

Open Access



MOTL: enhancing multi-omics matrix factorization with transfer learning

David P. Hirst^{1*}, Morgane Térézol¹, Laura Cantini², Paul Villoutreix¹, Matthieu Vignes³ and Anaïs Baudot^{1,4,5*}

*Correspondence:
david.hirst@univ-amu.fr; anais.baudot@univ-amu.fr

¹ Aix Marseille Univ, INSERM, MMG, Centuri, Marseille, France

² Institut Pasteur, Université Paris Cité, CNRS UMR 3738, Machine Learning for Integrative Genomics Group, Paris F-75015, France

³ School of Mathematical and Computational Sciences, College of Science, Massey University, Palmerston North, New Zealand

⁴ CNRS, Marseille, France

⁵ Barcelona Supercomputing Center, Barcelona, Spain

Abstract

Joint matrix factorization is popular for extracting lower dimensional representations of multi-omics data but loses effectiveness with limited samples. Addressing this limitation, we introduce MOTL (Multi-Omics Transfer Learning), a framework that enhances MOFA (Multi-Omics Factor Analysis) by inferring latent factors for small multi-omics target datasets with respect to those inferred from a large heterogeneous learning dataset. We evaluate MOTL by designing simulated and real data protocols and demonstrate that MOTL improves the factorization of limited-sample multi-omics datasets when compared to factorization without transfer learning. When applied to actual glioblastoma samples, MOTL enhances delineation of cancer status and subtype.

Keywords: Matrix factorization, Dimensionality reduction, Multi-omics, Data integration, Transfer learning, MOFA

Background

Omics data have transformed the study of biology and medicine by enabling high-throughput measurements of the activity and abundance of biological molecules and processes [1–3]. In recent years, the fields of biology and medicine have been revolutionized by the increased availability of multi-omics datasets [1, 4, 5]. A multi-omics dataset is comprised of multiple data matrices, each containing a different type of omics data (e.g., mRNA transcript counts, genomic mutations, DNA methylation prevalence). The integrative analysis of multi-omics data can provide a better understanding of a biological system than that obtained from the analysis of a single omics data matrix, as the complementary information contained in different omics enables a more comprehensive overview of the underlying biological system [6–11]. Additionally, using multiple omics can reveal insights into relationships between the different biological layers they represent [4, 5, 12]. Combining omics is also expected to reduce the impact of noise [6, 9, 12]. However, multi-omics data poses further analysis challenges beyond those encountered in single omics data analysis. These challenges include increased dimensionality, the presence of multiple data types, diverse sources of technological noise, and diverse



© The Author(s) 2025. **Open Access** This article is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License, which permits any non-commercial use, sharing, 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 licence, and indicate if you modified the licensed material. You do not have permission under this licence to share adapted material derived from this article or parts of it. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence 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. To view a copy of this licence, visit <http://creativecommons.org/licenses/by-nc-nd/4.0/>.

ranges of variability. In this context, there has been an increased need for methods able to carry out integrative analysis of multiple omics.

The development of multi-omics analysis tools is an active area of research and a large variety of strategies have been proposed. These strategies encompass broad and overlapping categories, including Bayesian methods, network-based approaches, or dimensionality reduction techniques [4, 9], alongside more recent deep learning strategies [10, 11, 13]. A category of multi-omics analysis tools that is widely used is dimensionality reduction with matrix factorization. Matrix factorization infers a lower dimensional representation of the observed data, in which a sufficiently informative proportion of the original signal is retained [14]. It has proven to be computationally efficient, but also interpretable, and effective for the analysis of large datasets [6, 7, 9, 12, 15].

Most classical matrix factorization approaches were designed for the analysis of a single data matrix. Applying matrix factorization to a single omics matrix produces a score matrix and a weight matrix, both of which contain values for latent factors that are potentially associated with different sources of underlying biological signal. The values in the weight matrix ideally represent signal across the assayed biological features, and the values in the score matrix represent the signal across the samples. For a multi-omics dataset, one of the strategies is to jointly factorize multiple omics data matrices. Various methods are now available for this purpose [9]. Multi-omics joint matrix factorization methods typically produce a weight matrix for each omics, and either a shared score matrix or a combination of shared and omics specific score matrices. Many multi-omics matrix factorization methods are extensions of classical methods. For example, intNMF [16] extends non-negative matrix factorization to the multi-omics setting, and allows the user to determine the relative contribution of each omics to the extraction of joint signal. JIVE [17] extends principal component analysis to model both joint and omics specific signal. moCluster [18] extends canonical correlation analysis to produce a shared score matrix based on score matrices produced for each of two or more omics. MOFA [19], which is an extension of Factor Analysis, uses a Bayesian framework to account for the presence of multiple data types and to distinguish between joint and omics specific signal. Overall, the factors inferred by multi-omics matrix factorization can be used for clustering samples to reveal disease sub-types, for identifying molecular profiles and biomarkers associated with diseases, as well as for prediction of outcomes such as drug response and survival [6, 7, 9, 20, 21]. A challenge for matrix factorization is that it requires a large amount of observed data to produce a meaningful representation. However, there are cases where omics are measured from only a small number of samples, due to the rareness or cost of obtaining the data, and so there is a need for methods which help mitigate this challenge [14, 21, 22].

For a dataset generated from a small number of samples, transfer learning is a potential solution to the limited effectiveness of matrix factorization. Transfer learning is a machine learning approach in which information extracted from a large learning domain is used to improve the performance of a task applied to a smaller target domain [20–23]. It is assumed that the two domains share an overlapping latent space, allowing knowledge from the learning domain to be transferred to the application of the task to the target domain. Transfer learning has been successfully used in various machine learning applications, including image classification, text sentiment classification and

recommendation systems [21, 22, 24, 25]. In a transfer learning approach to omics matrix factorization, information inferred from the prior factorization of a learning dataset, comprised of a large number of samples from a heterogeneous set of biological conditions, is incorporated into the factorization of a small target dataset [21, 26]. It is assumed that if the latent factors inferred from the learning dataset represent common underlying biological processes, they should help improve the factorization of the target dataset [23].

The usefulness of transfer learning approaches to matrix factorization, for omics data analysis, has been demonstrated in contexts in which both the target and learning datasets were comprised of single matrices of omics data. In these cases, transfer learning was used to infer a score matrix for the target dataset by projecting it onto a weight matrix inferred from a learning dataset. In one study, Stein-O'Brien et al. [23] factorized a mouse single cell RNA-seq learning dataset with the Bayesian non-negative matrix factorization algorithm CoGAPS [27]. Then, they used the transfer learning tool projectR [28] to infer a score matrix for a human time course bulk RNA-seq dataset. The resulting factors were associated with known spatiotemporal differences across the samples. In another example, Davis-Marcisak et al. [29] factorized a mouse single cell RNA-seq learning dataset with CoGAPS, and then used projectR to infer a score matrix for bulk RNA-seq data from human cancer samples. They observed an association between a particular projectR factor and outcomes in metastatic melanoma. Taroni et al. [20] developed MultiPLIER, a transfer learning framework, which they demonstrated by firstly applying the non-negative matrix factorization algorithm PLIER [30] to a subset of Recount2 to infer a weight matrix. Recount2 is a compendium of RNA-seq data obtained from 70,000+ human samples taken across more than 2000 studies [31]. Taroni et al. [20] then used a blood cell compendium of microarray gene expression data as the target dataset, for which they inferred a score matrix with MultiPLIER, as well as factorizing the target dataset directly with PLIER. For counts of a cell type of interest, MultiPLIER inferred a more highly correlated factor than was inferred by direct factorization of the target dataset. They also used microarray gene expression data for 79 samples from a rare disease group called antineutrophil cytoplasmic autoantibody associated vasculitis (AAV) as a target dataset. There are no AAV samples in the Recount2 compendium, yet the MultiPLIER factors were positively correlated with their best match from factors inferred by direct factorization of the target dataset.

It has thus been demonstrated that the application of matrix factorization to a large, heterogeneous learning dataset can yield factors containing transferable information, that are biologically relevant to target datasets from different organisms, diseases, cell types and omics platforms. However, existing transfer learning approaches to matrix factorization have been designed for, and demonstrated on, datasets comprised of single omics data only. To the best of our knowledge, transfer learning approaches to joint multi-omics matrix factorization are currently lacking.

We here introduce MOTL (Multi-Omics Transfer Learning), a novel Bayesian transfer learning algorithm for multi-omics matrix factorization. MOTL is based on MOFA, a popular tool for integrative multi-omics analysis [19]. We first present the statistical framework and implementation of MOTL. Next, we propose two protocols, that we designed based on simulated and real multi-omics datasets, for evaluating the

performance of transfer learning approaches. We used these protocols to evaluate MOTL, and observed that, for a target multi-omics dataset comprised of a small number of samples, our transfer learning approach to matrix factorization is more effective than matrix factorization without transfer learning. Lastly, we showcase a practical use case of MOTL on a limited glioblastoma sample set, revealing an enhanced delineation of cancer status and subtype thanks to transfer learning.

Results

MOTL: a new transfer learning framework for multi-omics matrix factorization

We propose MOTL, a transfer learning approach to multi-omics matrix factorization. MOTL is based on MOFA [19], which uses variational Bayesian inference [32]. Consider a multi-omics target dataset, T , consisting of omics matrices, $T^{(m)}$, $m = 1, \dots, M$. Each $T^{(m)} = \begin{bmatrix} t_{nd}^{(m)} \end{bmatrix} \in \mathbb{R}^{N_t \times D_m}$ contains data for N_t samples (rows) and D_m features (columns), where $t_{nd}^{(m)}$ is the value for the n th sample and the d th feature from the m th matrix (see the [Methods “Mathematical Notation”](#) section for a summary of the mathematical notation used in this document). The features depend on which molecules were assayed to generate a given omics matrix; for example the features for mRNA counts are genes, while those for DNA methylation are CpG sites.

We wish to jointly factorize $T^{(m)}$ into a matrix of sample scores, $Z = [z_{nk}] \in \mathbb{R}^{N_t \times K}$, and an omics specific matrix of feature weights, $W^{(m)} = \begin{bmatrix} w_{kd}^{(m)} \end{bmatrix} \in \mathbb{R}^{K \times D_m}$. The resulting lower dimensional representation is based on K factors, which ideally represent underlying biological signals associated with some biological condition(s) of interest. z_{nk} is the score for the n th sample and the k th factor, while $w_{kd}^{(m)}$ is the weight for the k th factor and the d th feature from the m th matrix. The k th column vector of Z , denoted by $z_{:k}$, contains scores for factor k , while the n th row vector, z_n , contains scores for sample n . The k th row vector of $W^{(m)}$, denoted by $w_{k:}^{(m)}$, contains weights for factor k , for the m th matrix, while the d th column vector, $w_{:d}^{(m)}$, contains weights for feature d from the m th matrix.

We are concerned with the situation in which N_t is small, exacerbating the curse of dimensionality, and therefore, we expect to improve the factorization of T by employing a transfer learning approach (see Fig. 1). We do this transfer learning by incorporating values that have already been inferred from the prior factorization of a learning dataset, L , and we assume that the Bayesian matrix factorization algorithm MOFA [19] was used for factorizing L .

The learning dataset consists of omics matrices $L^{(m)}$, $m = 1, \dots, M$. Each $L^{(m)} = \begin{bmatrix} l_{nd}^{(m)} \end{bmatrix} \in \mathbb{R}^{N_l \times D_m}$ contains data for the same D_m features as $T^{(m)}$, but for a different set of $N_l > N_t$ samples. We hypothesize that if L is comprised of samples from a heterogeneous set of biological conditions, then the factorization of L will yield information that is relevant for the factorization of T .

MOTL is based on the variational Bayesian inference methodology used by MOFA (the [Methods “The MOFA model”](#) section). We have modified the MOFA algorithm to enable us to supplement the factorization of T by incorporating values already inferred from the prior factorization of L . For MOTL, we assume that each observed $t_{nd}^{(m)}$ is a random variable, with a likelihood that is conditional on vectors z_n and $w_{:d}^{(m)}$. We model

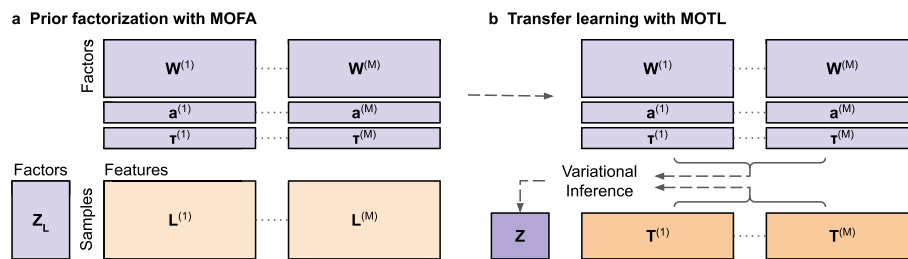


Fig. 1 Overview of MOTL, our transfer learning approach to joint multi-omics matrix factorization based on variational Bayesian inference. **a** A multi-omics learning dataset, L , consisting of M omics matrices, $L^{(m)}$, $m = 1, \dots, M$, is factorized with MOFA to infer a matrix of feature weights, $W^{(m)}$, vector of feature-wise intercepts, $a^{(m)}$, and a vector of feature-wise precision parameter values, $\tau^{(m)}$, for each $L^{(m)}$. **b** The feature weight, intercept, and precision parameter values, inferred from the factorization of L , are incorporated into the factorization of a multi-omics target dataset, T , for which MOTL infers a matrix of sample scores, Z , with variational inference

continuous, counts and binary data with the same likelihoods and link functions that MOFA uses. For observed continuous data, we thus assume a Gaussian likelihood, into which we include a feature-wise precision parameter, $\tau_d^{(m)}$, for each feature d from matrix m . For observed binary data, we assume a Bernoulli likelihood, and for observed counts data we assume a Poisson likelihood. In contrast to MOFA, MOTL does not center the input data during factorization fitting, as we want to incorporate an intercept that is compatible with the factorization of L . We therefore replace $z_n \cdot w_{:d}^{(m)}$ with $a_d^{(m)} + z_n \cdot w_{:d}^{(m)}$ in the likelihood, where $a_d^{(m)}$ is the feature-wise intercept for feature d , from matrix m . We infer $a_d^{(m)}$ values based on the MOFA factorization of L (the [Methods “Application of MOTL to simulated, TCGA and glioblastoma multi-omics datasets”](#) section). MOTL accepts missing $t_{nd}^{(m)}$ values; therefore it is not necessary to remove features with missing values, or perform imputation, before using MOTL.

In order to carry out a transfer learning approach to matrix factorization, MOTL uses the matrix of feature weights, $W^{(m)}$, vector of feature-wise intercepts, $a^{(m)} = [a_d^{(m)}] \in \mathbb{R}^{D_m}$, and vector of feature-wise precision parameter values, $\tau^{(m)} = [\tau_d^{(m)}] \in \mathbb{R}^{D_m}$, inferred for each $L^{(m)}$ with a prior MOFA factorization of L . Instead of modeling these as random variables, we treat them as constants. We aim to obtain point estimates of z_{nk} values, for which, we assume the same joint prior distribution as MOFA does,

$$p(Z) = \prod_{n=1}^{N_t} \prod_{k=1}^K \text{Normal}(z_{nk} | 0, 1) \quad (1)$$

MOTL obtains point estimates of z_{nk} values by approximating the joint posterior distribution $p(Z|T)$ with a variational distribution:

$$q(Z) = \prod_{n=1}^{N_t} \prod_{k=1}^K q(z_{nk}) = \prod_{n=1}^{N_t} \prod_{k=1}^K \text{Normal}(z_{nk} | \mu_{nk}, \sigma_{nk}) \quad (2)$$

MOTL infers $q(Z)$ iteratively. At each iteration, the value of each parameter is updated while all other parameter values are held fixed. MOTL optimizes the joint

variational distribution by iterating until convergence. For each z_{nk} , the expected value, $\mathbb{E}_q[z_{nk}] = \mu_{nk}$, is used as the point estimate throughout and after model fitting. MOTL uses the same update equations for the parameters of $q(z_{nk})$ as MOFA, but with the inclusion of intercepts:

$$\sigma_{nk}^2 = \left(\sum_{m=1}^M \sum_{d=1}^{D_m} \tau_{nd}^{(m)} \left(w_{kd}^{(m)} \right)^2 + 1 \right)^{-1} \quad (3)$$

$$\mu_{nk} = \sigma_{nk}^2 \sum_{m=1}^M \sum_{d=1}^{D_m} \tau_{nd}^{(m)} w_{kd}^{(m)} \left(\hat{t}_{nd}^{(m)} - a_d^{(m)} - \sum_{j \neq k} z_{nj} w_{jd}^{(m)} \right) \quad (4)$$

where $\tau_{nd}^{(m)}$ is the precision for the n th sample and d th feature from the m th matrix, and $\hat{t}_{nd}^{(m)}$ denotes a (possibly) transformed observed data point (the [Methods “The MOFA model”](#) section). For observed data with a Gaussian assumed likelihood, a feature-wise precision, $\tau_d^{(m)}$, is used instead of $\tau_{nd}^{(m)}$, and there is no transformation, meaning $\hat{t}_{nd}^{(m)} = t_{nd}^{(m)}$. For observed data with a non-Gaussian assumed likelihood, MOTL transforms the data to yield Gaussian pseudo-data values, which it does not center. The transformation to Gaussian pseudo-data allows updates of $q(\mathbf{Z})$ to be based on the assumption of Gaussian observed data. When MOFA transforms observed data with a Bernoulli assumed likelihood, it derives and uses a precision parameter, $\tau_{nd}^{(m)}$, for each sample and feature. For observed data with a Poisson assumed likelihood, it derives and uses a feature-wise precision, $\tau_d^{(m)}$. Thus for Bernoulli observed data, MOTL initializes $\tau_{nd}^{(m)}$ values with $\tau_d^{(m)}$ values, which are averages of the $\tau_{nd}^{(m)}$ values returned by the factorization of $\mathbf{L}$, and these are subsequently updated at each iteration of the algorithm. For Poisson observed data, MOTL uses the $\tau_d^{(m)}$ values obtained from the prior factorization of $\mathbf{L}$, and holds them fixed.

To monitor convergence we calculate the evidence lower bound (ELBO), which can be used to evaluate how well a variational distribution approximates a posterior distribution of interest. We calculate the ELBO with respect to $\mathbf{Z}$:

$$\text{ELBO}(\mathbf{Z}) = \mathbb{E}_q[\log p(\mathbf{T}|\mathbf{Z})] + \mathbb{E}_q[\log p(\mathbf{Z})] - \mathbb{E}_q[\log q(\mathbf{Z})] \quad (5)$$

For $\mathbf{T}^{(m)}$ with a non-Gaussian assumed likelihood, we use the same lower bound for $\log p(t_{nd}^{(m)}|\mathbf{Z})$ as MOFA does. Maximizing this lower bound, coupled with the use of $\hat{t}_{nd}^{(m)}$ values, allows updates of $q(\mathbf{Z})$ based on the assumption of Gaussian observed data [33, 34]. We calculate the ELBO at regular intervals, and the number of iterations between each calculation is a user defined parameter. We check for convergence based on the absolute change in the ELBO (from the previous check) as a percentage of the initial ELBO. The algorithm is deemed to have converged when a specified number of changes in the ELBO are consecutively below a threshold. Both the threshold, and the required number of consecutive changes falling below this threshold, are user defined parameters.

We allow factors to be dropped during training, based on the fraction of variance explained:

$$R_{mk}^2 = 1 - \frac{\sum_{n=1}^{N_t} \sum_{d=1}^{D_m} \left(\hat{t}_{nd}^{(m)} - a_d^{(m)} - z_{nk} w_{kd}^{(m)} \right)^2}{\sum_{n=1}^{N_t} \sum_{d=1}^{D_m} \left(\hat{t}_{nd}^{(m)} - a_d^{(m)} \right)^2} \quad (6)$$

We drop the factor with the lowest R_{mk}^2 that does not have any R_{mk}^2 above the threshold. We assess factors in this way after each round of updates. After convergence the algorithm returns $\mathbf{Z}$ and $\mathbf{W}^{(m)}$ matrices for the factors that have not been dropped.

MOTL is available as an open source R implementation [35].

Evaluation protocol using simulated multi-omics data

We first designed and implemented a transfer learning evaluation protocol based on simulated multi-omics datasets, which we generated from groundtruth factors (the [Methods “Multi-omics data simulated with groundtruth factors”](#) section). In each simulation instance, we generated a multi-omics dataset, $\mathbf{Y}$, which we subsequently split into a target dataset, $\mathbf{T}$, and a learning dataset, $\mathbf{L}$. $\mathbf{Y}$ consisted of matrices of counts, continuous, and binary data. We generated each matrix, $\mathbf{Y}^{(m)}$, from a statistical distribution conditional on random matrices $\mathbf{Z}$ and $\mathbf{W}^{(m)}$, which each contained values for K groundtruth factors. The k th column vector of $\mathbf{Z}$ contained sample scores for the k th groundtruth factor. The k th row vector of $\mathbf{W}^{(m)}$ contained feature weights for that same factor. We varied the number of groundtruth factors across configurations, using $K \in \{20, 30\}$. We generated $\mathbf{Z}$ based on the group membership of samples. In each instance, we created two groups of five samples for the target dataset. The learning dataset samples belonged to either 20 or 40 differently sized groups of randomly selected sizes. For each groundtruth factor and group, the sample scores were generated using a mean parameter value that was common to all samples in the group. We induced heterogeneity by allowing the means to vary across groups and factors, randomly selecting each group mean, for a given groundtruth factor, from a pool of three possible values. We split each $\mathbf{Y}^{(m)}$ into $\mathbf{T}^{(m)}$ and $\mathbf{L}^{(m)}$, based on the sample groups used to generate $\mathbf{Z}$. In each instance $\mathbf{T}$ contained data for 10 samples, while the expected number of samples for $\mathbf{L}$ was $\in \{400, 1000\}$.

For each simulation instance, we factorized $\mathbf{L}$ with MOFA (the [Methods “Application of MOFA to simulated, TCGA and glioblastoma multi-omics datasets”](#) section). We then factorized $\mathbf{T}$ with our transfer learning method MOTL (the [Methods “Application of MOTL to simulated, TCGA and glioblastoma multi-omics datasets”](#) section), incorporating output from the factorization of $\mathbf{L}$. To benchmark the performance of MOTL, we also performed direct MOFA factorizations (i.e., factorization without transfer learning) of $\mathbf{T}$ datasets. We evaluated both the MOTL and direct MOFA factorization of each $\mathbf{T}$, and compared the overall performance of each approach. We evaluated factorizations of each $\mathbf{T}$ by calculating an $F1$ score (the [Methods “Evaluation methods”](#) section), to measure how well the factorization allowed us to uncover differentially active groundtruth factors underlying $\mathbf{T}$. The k th groundtruth factor was differentially active for $\mathbf{T}$ if the mean parameter values used to simulate the sample scores, for that factor, differed between the two groups of target dataset samples. Factorization with MOTL led to higher $F1$ scores than direct MOFA factorization, indicating that the MOTL

factorizations were more effective in uncovering differentially active latent signal from T datasets (Fig. 2a). This was observed across all simulation configurations, and the overall uplift in mean $F1$ score for MOTL, when compared to direct MOFA factorization, was 0.21 (p -value < 0.01 , the [Methods “Evaluation methods”](#) section). We thus observed that transfer learning with MOTL was more effective in uncovering differentially active latent signal, when compared to direct MOFA factorization (without transfer learning) of T . Of note, MOTL was also more effective than direct factorization with the alternative multi-omics matrix factorization approaches intNMF and moCluster (Additional file 1: Fig. S1).

We next wanted to evaluate the robustness of MOTL when there is a decline in the overlap between the latent spaces of L and T (the [Methods “Application of MOTL to simulated, TCGA and glioblastoma multi-omics datasets”](#) section). We forced this decline in overlap by permuting feature vectors in the $W^{(m)}$ matrices inferred from L datasets, based on a range of permutation proportions between 0 and 1. For each simulation instance, and permutation proportion, p , we created new $W^{(m)}$ matrices by permuting the values in $(p \times 100)\%$ of the feature vectors in the $W^{(m)}$ matrices inferred from L . We then factorized T with MOTL, using the permuted $W^{(m)}$ matrices, and calculated the $F1$ score. We observed that MOTL outperformed direct MOFA factorization even when there were large declines in the overlap between the latent spaces of L and T , and that the performance of MOTL tended to drop below that of direct MOFA factorization when the values in 80% or more of the feature vectors were permuted (Fig. 2b).

Evaluation protocol using TCGA multi-omics data

We next designed, and implemented, a second transfer learning evaluation protocol, based on TCGA multi-omics data (the [Methods “TCGA multi-omics data acquisition and pre-processing”](#) section). We used four types of omics data: log2 transformed mRNA counts, log2 transformed miRNA counts, DNA methylation M -values, and simple nucleotide variation (SNV) binary data, which we obtained for 32 different cancer types. We created target datasets using data from three cancer types; acute myeloid leukemia (LAML), pancreatic adenocarcinoma (PAAD) and skin cutaneous melanoma (SKCM). We created these target datasets by firstly creating four reference datasets. Each reference dataset, R , contained multi-omics data for all samples from either two, or all three of the cancer types. We then randomly split every R into non-overlapping target datasets which each contained only five samples per cancer type (Fig. 3a). We merged data from the remaining 29 cancer types into a learning dataset, L , which contained multi-omics data for 7217 samples.

We factorized L with MOFA (the [Methods “Application of MOFA to simulated, TCGA and glioblastoma multi-omics datasets”](#) section), based on which we used MOTL to factorize each T (the [Methods “Application of MOTL to simulated, TCGA and glioblastoma multi-omics datasets”](#) section). To benchmark the performance of MOTL, we also performed direct MOFA factorizations (without transfer learning) of T datasets. In order to evaluate the factorizations of T datasets, we factorized R datasets with MOFA and treated the resulting score, Z , and weight, $W^{(m)}$, matrices as groundtruth factor matrices (the [Methods “Evaluation methods”](#) section). We were interested in how well the

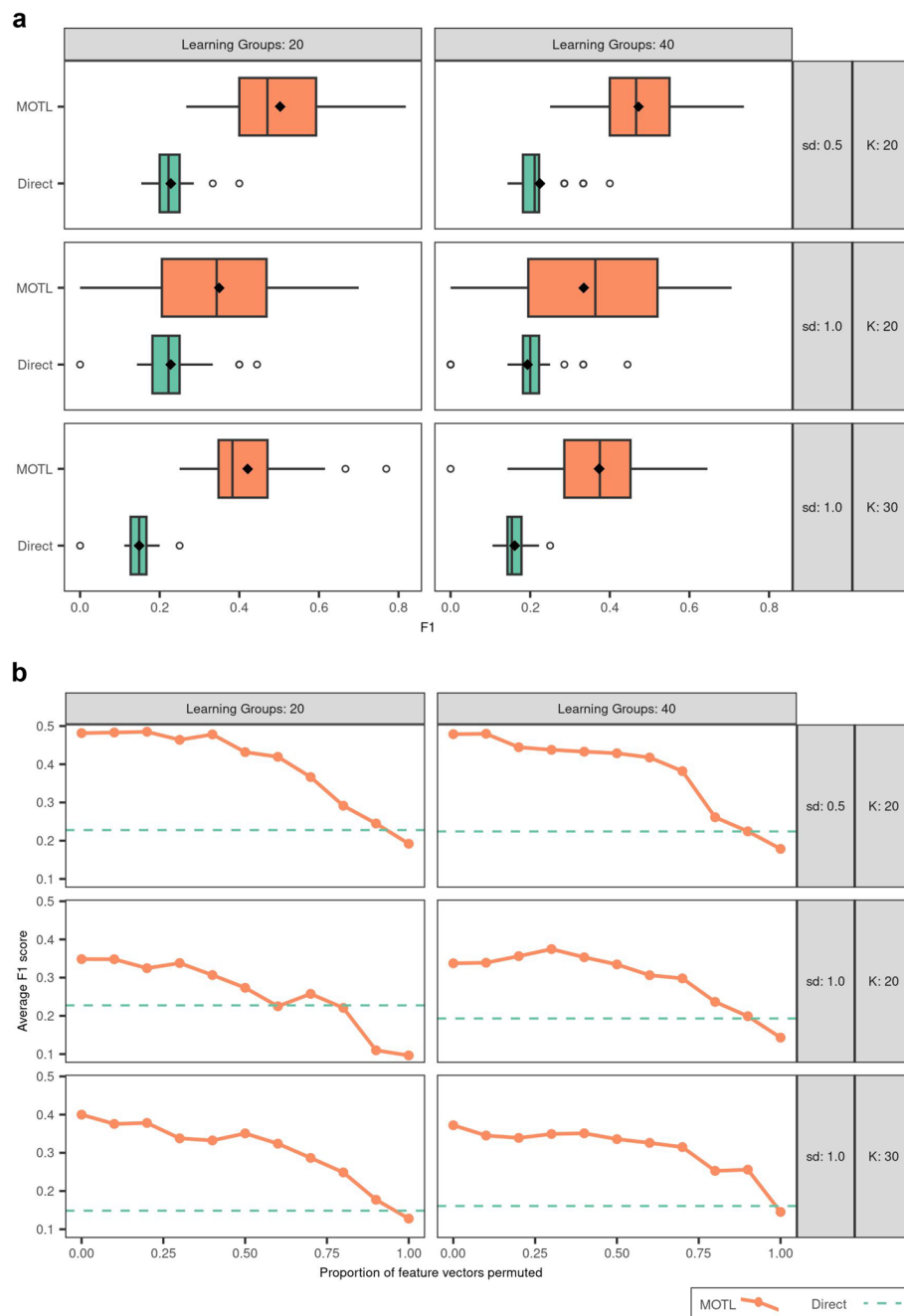


Fig. 2 Evaluation of factorizations of small simulated multi-omics target datasets. **a** The boxplots represent the F1 scores obtained for factorizations with and without MOTL transfer learning, for different simulation configuration settings. Simulation configurations varied in the number of groups of samples used for the learning dataset (*Learning Groups*), the number of groundtruth factors (*K*), or the standard deviation used to simulate z_{nk} values (*sd*). F1 scores take a value between 0 and 1, and higher values indicate better factorizations. Each boxplot is based on 30 F1 scores. The hinges of the boxes are the 25th and 75th percentiles, the middle lines are medians, the diamonds are the mean values, and the whiskers are either extreme values or extend 1.5 times the inter-quartile range from the hinge. **b** The line plots represent the average F1 score obtained from MOTL factorizations, for different simulation configuration settings, after permuting the values in proportions of the feature vectors in the $\mathbf{W}^{(m)}$ matrices obtained from prior factorization of $\mathbf{L}$. The dashed line is the average F1 score obtained with direct MOFA factorization for the simulation configuration

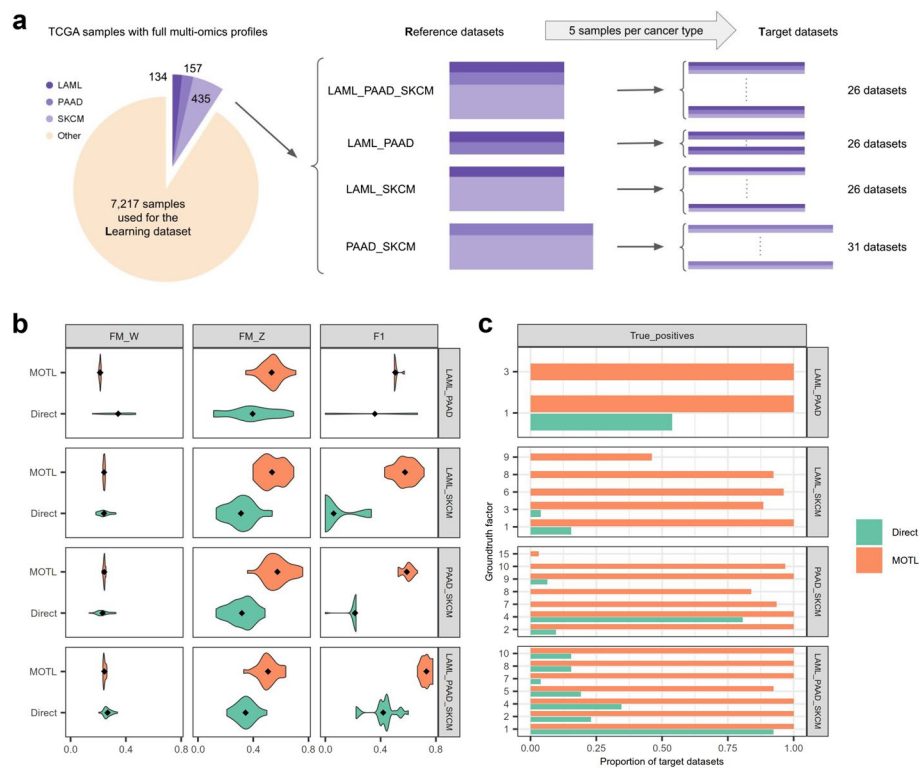


Fig. 3 Evaluations using TCGA multi-omics data. **a** We created TCGA target, T , datasets from two or three cancer types: For each combination of cancer types, we created a reference, R , dataset containing multi-omics data for all samples from the selected cancer types. We then randomly split R into non-overlapping T datasets, containing multi-omics data for five samples per cancer type. We did this for subsets of the set of cancer types {LAML, PAAD, SKCM}. In total we created, and split, four reference datasets, each of which contained multi-omics data for all samples from either two (LAML and PAAD, LAML and SKCM, PAAD and SKCM), or three (LAML, PAAD, and SKCM) cancer types. For datasets containing PAAD and SKCM samples only, we included mRNA, miRNA, DNA methylation, and SNV data. We did not include SNV data in datasets containing LAML samples, due to the sparsity of SNV data. **b** Comparison of factorization approaches applied to TCGA multi-omics datasets. Violin plots of F -measure values for weight matrix factors (FM_W), F -measure values for score matrix factors (FM_Z), and $F1$ scores ($F1$). For each evaluation score, higher values indicate better factorizations. Scores are plotted by factorization method and by the cancer types characterizing the target dataset samples. **c** Frequency with which differentially active groundtruth TCGA factors were true positives. A differentially active groundtruth factor was a true positive if it was predicted as being differentially active based on a factorization of a target dataset. Each bar represents the proportion of target datasets, for which the factorization led to the differentially active groundtruth factor being a true positive. Proportions are plotted by factorization method, and by the cancer types characterizing the reference and target dataset samples

factorizations of T datasets uncovered the groundtruth, and we used F -measure values and $F1$ scores to evaluate this.

We calculated F -measure values to assess the correlation between factors inferred from each T , and the groundtruth factors obtained from the factorization of the corresponding R dataset (the [Methods “Evaluation methods”](#) section). We calculated F -measure values for weight matrices (FM_W), as well as for score matrices (FM_Z). The overall mean FM_W for MOTL was slightly lower (0.03 reduction, p -value < 0.01) than for direct MOFA factorizations of T datasets (Fig. 3b, column 1), which is the result of lower average relevance counterbalancing higher average recovery. We concluded from this that groundtruth $W^{(m)}$ factors were more easily uncovered with those transferred from

L than with direct MOFA factorization. However, despite factor trimming during MOTL factorization, some remaining transferred factors were less associated with groundtruth factors than those obtained with direct MOFA factorization. It is of note that the difference in average FM_W is attributable to the datasets containing LAML and PAAD samples only. If we exclude these, there is no difference in FM_W (p -value 0.59). The overall mean FM_Z for MOTL was 0.20 higher (p -value < 0.01, the [Methods “Evaluation methods”](#) section) than for direct MOFA factorizations (Fig. 3b, column 2). We thus observed that the Z factors obtained with MOTL, from T datasets, were more correlated with groundtruth factors, overall, than those obtained with direct MOFA factorization.

We also calculated $F1$ scores to measure how well the factorizations of T datasets uncovered differentially active groundtruth factors (the [Methods “Evaluation methods”](#) section). For each T dataset, the groundtruth factors were the factors obtained from factorization of the corresponding R dataset. We considered the k th groundtruth factor to be differentially active if the distribution of scores in the k th column vector, of groundtruth Z , differed between the cancer types. We can simultaneously evaluate the Z and $W^{(m)}$ factors, and assess the overall quality of factorizations, by checking for an uplift in $F1$ scores (Fig. 3b, column 3). MOTL (0.34 uplift, p -value < 0.01) yielded higher $F1$ scores than direct MOFA factorization, meaning it was more effective in uncovering latent activity that varied across cancer types. Similarly to the evaluation using simulated data, MOTL was also more effective than direct factorization with the alternative multi-omics matrix factorization approaches intNMF and moCluster (Additional file 1: Fig. S2).

We next examined differentially active groundtruth factors, with an initial focus on the frequency with which these factors were true positives (Fig. 3c). For each factorization of a T dataset, a differentially active groundtruth factor was a true positive if it was predicted as being differentially active based on the factorization of T . The unique count of true positives was a component of each $F1$ score value (the [Methods “Evaluation methods”](#) section). We further performed a gene set enrichment analysis to identify the pathways and processes associated with differentially active groundtruth factors that were true positives (the [Methods “Evaluation methods”](#) section, Additional file 2: Table S1), and that explained at least 1% of the mRNA variance in R (Additional file 3: Table S2).

The factorization of the R dataset containing all LAML and PAAD samples yielded six groundtruth factors, of which two were differentially active; Factor 1 and Factor 3. Both of these factors were true positives for 100% of MOTL factorizations of T datasets containing subsets of five LAML and PAAD samples (Fig. 3c). In contrast, only one of these factors was a true positive for direct MOFA factorizations, and for just over half of the same T datasets (Fig. 3c). Factor 1 is significantly associated with developmental processes, cell communication and immunity signaling. Factor 3 displays similar enrichments, with an additional specific enrichment related to the regulation of gene expression in beta cells (Additional file 2: Table S1).

The factorization of the R dataset containing all LAML and SKCM samples yielded 12 groundtruth factors, of which five were differentially active. Four out these five factors were true positives for MOTL factorizations for more than 80% of the T datasets; one factor was a true positive for just under half of the T datasets. Importantly, only two of these five groundtruth factors were true positives for direct MOFA factorizations,

and only for a small proportion of the T datasets. Four out of these five differentially active groundtruth factors, that were true positives, explained at least 1% of the mRNA variance in R : Factor 1, Factor 3, Factor 8, and Factor 9. We performed gene set enrichment analysis on these four factors. Factor 1 is significantly associated with extracellular organization, developmental processes, cell communication signaling, and Fc Receptor mediated immune processes (Additional file 2: Table S1). Factor 3 is significantly associated with hematopoietic cell lineage, Pi3K/AKT and G protein-coupled signaling, and chemokine, interleukin, interferon signaling. Both Factors 1 and 3 are associated with keratinization and formation of the cornified envelope. Factors 8 and 9 do not present significant pathway enrichments beyond keratinization and processes already associated with the first two factors.

The factorization of the R dataset containing all PAAD and SKCM samples yielded 17 groundtruth factors, of which eight were differentially active. Seven of these eight factors were true positives for MOTL factorizations, six with high frequency (i.e., identified for more than 80% of the T target datasets). Only one differentially active groundtruth factor is frequently a true positive for direct MOFA factorization of the T target datasets. Factor 2 is related to B cell receptor signaling and Fc Receptor mediated immune processes, drug metabolism by cytochrome p450 and other metabolism-related processes. This Factor 2 is rarely uncovered by direct MOFA factorization. Contrarily, Factor 4 is a true positive for MOTL for 100% of the target datasets and for direct MOFA factorizations for more than 75% of the target datasets. This factor is associated with developmental processes and cytokine-cytokine receptor interactions. Factor 7 is associated with keratinization and formation of the cornified envelope, Factor 9 with cell adhesion and migration, and Factor 10 with cytokine and chemokine signaling.

The factorization of the R dataset containing all LAML, PAAD and SKCM samples yielded 13 groundtruth factors, of which seven were differentially active. All seven of these differentially active groundtruth factors were true positives for MOTL factorization with high frequency, compared to one factor for direct MOFA factorization. These groundtruth factors, differentially active between all three cancer types, are associated with the same cellular processes and pathways identified when comparing the cancer types pairwise (Additional file 2: Table S1).

Overall, the factors that were differentially active when comparing two or the three cancer types reflect the different embryonic origins of the cancerous tissues, and highlight the importance of immunity and microenvironment in cancer pathophysiology and response to treatments. In conclusion, matrix factorization with transfer learning using MOTL better uncovers differentially active groundtruth factors from target datasets containing only a small number of samples.

Application of MOTL to glioblastoma

Glioblastoma is a rare, heterogeneous, and aggressive cancer type. Multi-omics datasets offer an important opportunity to better characterize glioblastoma subtypes, identify biomarkers, and propose novel therapeutic options [36]. However, large collections of glioblastoma tissue samples are difficult to obtain due to the relative scarcity of the disease and the challenges involved in acquiring samples via invasive biopsies.

In [36], the authors conducted a multi-omics profiling (mRNA expression, DNA methylation) for four normal brain samples and nine patient-derived glioblastoma stem cell (pd-GBSC) cultures. The nine cancer samples had been previously classified into three subtypes thanks to transcriptome-based signatures: classical (CL), proneural (PN), and mesenchymal (MS). Given the small number of samples, the authors devised a strategy based on analyzing this dataset in parallel with datasets gathered from the literature. We illustrate here how MOTL could help in analyzing such a dataset comprised of a limited number of samples.

We first applied a direct MOFA factorization (the [Methods “Application of MOFA to simulated, TCGA and glioblastoma multi-omics datasets”](#) section) to a target dataset comprised of the four normal and nine pd-GBSC samples (the [Methods “Glioblastoma target dataset acquisition and pre-processing”](#) section), revealing eight factors. Heatmap clustering of the samples, based on these factors, does not demonstrate clear grouping with respect to either cancer status or subtype (Fig. 4a). Next, we applied MOTL to the same target dataset (the [Methods “Application of MOTL to simulated, TCGA and glioblastoma multi-omics datasets”](#) section). In this case, we first created a TCGA learning dataset containing mRNA expression, miRNA expression, DNA methylation, and SNV data for samples from all 32 cancer types (the [Methods “TCGA multi-omics data acquisition and pre-processing”](#) section). It is noteworthy that this learning dataset did not

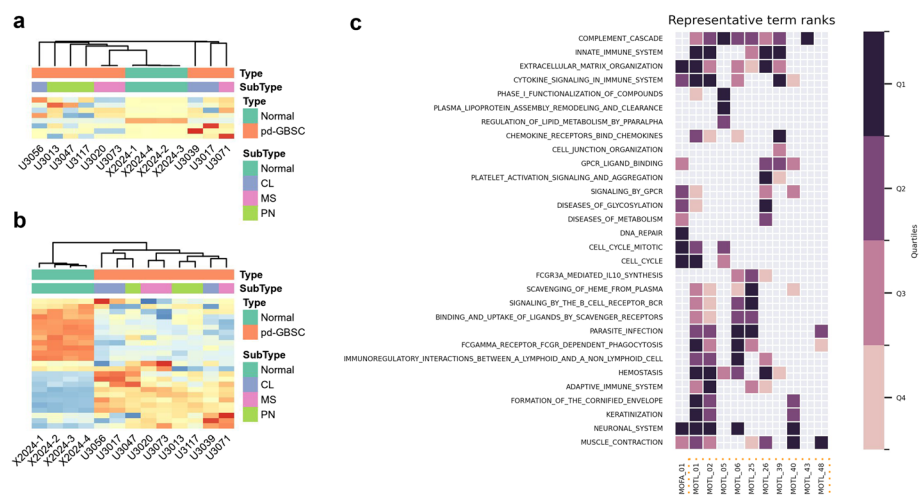


Fig. 4 Heatmaps of factorizations and gene set enrichment analysis of glioblastoma data: The heatmaps are based on factorizations and subsequent gene set enrichment analysis of the target dataset comprised of multi-omics data (mRNA expression, DNA methylation) for all normal and patient-derived glioblastoma stem cell (pd-GBSC) samples. The multi-omics data was obtained from Santamarina-Ojeda et al. [36], who also provided subtypes for the cancer samples, previously defined from transcriptomic signatures: CN (classical), PN (proneural), MS (mesenchymal). **a** Heatmap of the score matrix, Z , inferred with Direct MOFA factorization of the target dataset. The rows are the factors, the columns are the samples, and the cells are the row-wise centered and scaled factor values. The rows and columns were ordered with hierarchical clustering (complete-linkage). **b** Heatmap of Z matrix inferred with MOTL factorization of the target dataset. **c** Reactome enrichment results for direct MOFA and MOTL factors. We performed gene set enrichment analysis on the direct MOFA and the 11 MOTL factors that were differentially active between normal and cancer samples, and that explained at least 1% of mRNA variance. This yielded a set of Reactome processes and pathways significantly associated with either one direct MOFA factor or one or more of 10 different MOTL factors. The Reactome processes and pathways were filtered, clustered, and plotted with orsum. The colors represent the quartiles of enrichment significance for each factor (darker means more significant)

contain data for glioblastoma, as there were no TCGA glioblastoma samples fulfilling our selection criteria (i.e., with complete 4-layer multi-omics profiles). We factorized this learning dataset with MOFA (the [Methods “Application of MOFA to simulated, TCGA and glioblastoma multi-omics datasets”](#) section), based on which we applied MOTL to the target dataset. MOTL transfer learning factorization revealed 25 factors. In this case, the heatmap clustering of the MOTL factors separates cancer and normal samples and also displays subgroups partially matching subtypes previously defined based on transcriptome signatures (Fig. 4b).

Further statistical tests revealed that only one of the eight factors identified by direct MOFA factorization was differentially active between normal and cancer samples, whereas 19 of the 25 factors obtained by MOTL transfer learning factorization were differentially active (the [Methods “Evaluation methods”](#) section). Gene set enrichment analysis using Reactome, GO:CC, GO:BP and KEGG (the [Methods “Evaluation methods”](#) section), focusing on differentially active factors that explained at least 1% of mRNA variance, revealed 715 processes and pathways associated with the direct MOFA factor, and 1061 processes and pathways associated with 11 MOTL factors (Additional file 4: Table S3). The overlap between the two sets of associated processes and pathways was 318, which is statistically significant (hypergeometric test, one-sided p -value < 0.01). We filtered and integrated these enrichment results (Fig. 4c; Additional file 1: Fig. S3–S5) using orsum [37] (the [Methods “Evaluation methods”](#) section). This analysis revealed that MOTL factors capture a broader spectrum of biological pathways than the direct MOFA approach, including for instance immune/inflammatory processes and lipid metabolism pathways.

We also applied both direct MOFA, and MOTL factorizations, to target datasets comprised of normal samples and samples from just a single cancer subtype. In all cases, the direct MOFA factorization yielded only one differentially active factor, whereas MOTL yielded 11 (CL vs normal), 16 (MS vs normal), and 12 (PN vs normal) differentially active factors. Focusing on the subset of differentially active factors that explained at least 1% of mRNA variance, we performed gene set enrichment analyses. We identified processes and pathways that were associated with only a single cancer subtype, such as fatty acid metabolism enrichment for the MS subtype and clathrin-mediated endocytosis for the PN subtype (Additional file 5: Table S4).

Discussion

We presented MOTL, which factorizes a multi-omics target dataset by incorporating latent factor values already inferred from the factorization of a multi-omics learning dataset. In the application of MOTL, we are concerned with the situation in which we analyze a target dataset comprised of a limited number of samples. In our evaluations, the target datasets never contained data for more than 15 samples, as our primary concern was with target datasets considered too small for useful factor analysis. It would be insightful to extend the evaluations by using a larger range of sizes for target datasets, in order to identify a crossover point at which transfer learning no longer enhances matrix factorization.

An assumption underlying our application of transfer learning is that there is an overlapping latent space stemming from shared biological signal. With target datasets

containing few samples, the learning dataset is likely to represent a set of biological conditions that neither is entirely specific to, nor fully overlaps with the set of biological conditions represented by the target dataset. In our evaluation based on simulated data, we observed that MOTL performed well even when there were large declines in the overlap between the latent spaces underlying the target and learning datasets, except for the more extreme situation where there was an almost total absence of overlap. It would be relevant to perform a similar analysis using real data, investigating how similar the learning and target datasets need to be in order for informative shared factors to exist. For instance, it would be beneficial to use a measure of similarity, between a given learning and target dataset, which would predict the effectiveness of using a transfer learning approach to apply matrix factorization to the target dataset. The similarity between the learning and target datasets could potentially be data driven, based on measures such as optimal transport [38] or maximum mean discrepancy [39]. Alternatively, a relevant ontology [40, 41] could be used to assess the overall similarity between the biological conditions characterizing the learning and target datasets. It could also be interesting to quantify how heterogeneous (i.e., representing a large diversity of tissues, diseases, experimental conditions...) a learning dataset needs to be, in order to yield factors which can be relevant for a given target dataset.

To evaluate MOTL, we designed two evaluation protocols, based on simulated and real data. Importantly, these protocols can be reused to evaluate other transfer learning strategies for multi-omics data integration. For the evaluation protocol based on real data, we used TCGA, a public repository of multi-omics data with a large number of samples, representing various cancer and tissue types. We selected three different cancer types as references, from which to build target datasets. We did not include samples from these cancer types in the learning dataset. In this setting, it was unknown, prior to evaluation, whether there were latent factors common to the learning and target datasets. Yet, MOTL was effective in uncovering differentially active latent factors, demonstrating that latent factors can be shared across different cancer types. We applied MOTL to a glioblastoma use-case as a proof-of-concept, but we envisage MOTL as being a helpful tool in the study of rare diseases in general. Therefore a future extension of our work would be to evaluate the application of MOTL to target datasets with non-cancer rare disease samples, using factors inferred from the TCGA learning dataset. We foresee the results of such an evaluation being accompanied by measures of how similar the target datasets are to the TCGA learning dataset, allowing for guidance on when MOTL is a suitable tool for the analysis of other rare disease datasets.

With MOTL, we have designed a transfer learning framework that is compatible with a prior learning dataset factorization, as carried out with the MOFA Bayesian approach [19]. The appeal of a Bayesian framework is the flexibility with regard to the incorporation of prior information, and variational inference serves as a fast alternative to sampling methods. In the future, MOTL could be extended to allow information to be incorporated at other levels of the assumed hierarchy of latent variables. For example, instead of fixing the feature weight values, they could be treated as random variables by MOTL, with priors informed by the factorization of the learning dataset.

In addition to MOFA, there are numerous methods available for multi-omics matrix factorization [16, 17, 42, 43]. A future extension of our work could be a transfer learning

framework matched to different multi-omics matrix factorization methods. Similarly, we foresee value in extending an existing transfer learning method that has been designed for single omics data [20, 28], so it can be applied in the multi-omics context. The evaluation of such a method, based on the factorization of a learning dataset with a variety of matrix factorization methods, would be informative.

A limitation with MOTL is that we are restricted to a factorization based on features that were retained for factorization of the learning dataset. A consequence is that some features which are highly variable in the target dataset may not contribute to the MOTL factorization. Therefore a future extension could be to add flexibility into the MOTL workflow, so that all features that are highly variable in the target dataset contribute to the factorization, even if they were not retained for the factorization of the learning dataset. Adding flexibility in this way may enhance the weight matrices that are used by MOTL. In the evaluation based on TCGA data, we observed that factors from the groundtruth weight matrices were more easily uncovered using those transferred from the learning dataset than by using those from direct MOFA factorization. However, the transferred factors were slightly less correlated with the factors from the groundtruth weight matrices than the direct MOFA factors were. Despite this, MOTL factorizations were more effective than direct MOFA factorizations in uncovering differentially active groundtruth latent factors. We attribute this superior performance to the fact that the factors in the score matrices inferred by MOTL showed higher correlations with factors from the groundtruth score matrices than the direct MOFA factors. We thus expect that incorporating greater flexibility in the MOTL workflow to retain all of the features that are highly variable in the target dataset would further enhance the weight matrices and produce even more informative factors.

Recently, deep learning, generative, and foundation models have been tested for multi-omics data integration [10, 11, 13]. We however did not identify benchmarks comparing linear methods based on MF with deep learning methods in the context of bulk multi-omics data integration, in particular for small target datasets. It will be interesting to see the developments in this context with more multi-omics data becoming available.

Conclusions

We presented MOTL, an approach for multi-omics matrix factorization with transfer learning, which infers latent factor values for a multi-omics target dataset comprised of a small number of samples. MOTL factorizes the target dataset by incorporating latent factor values already inferred from the factorization of a learning dataset. We designed two protocols, based on simulated and real multi-omics datasets, for evaluating the performance of multi-omics matrix factorization with transfer learning. We implemented these protocols to evaluate MOTL, and observed that MOTL was more effective in uncovering differentially active groundtruth latent factors than direct matrix factorization without transfer learning. This result is observed from comparison to direct factorization with MOFA as well as with moCluster and intNMF.

Finally, the application of MOTL to a glioblastoma dataset, comprised of a small number of samples, revealed an enhanced delineation of cancer status and subtype thanks to transfer learning. We thus demonstrated, in the case of a multi-omics dataset comprised of a small number of samples, that MOTL can enhance the discovery of biological processes

and pathways associated with a biological condition of interest. MOTL is accessible as an open source R implementation, as are the evaluation protocols we used in this study.

Methods

Mathematical notation

- We denote matrices and datasets with bold capital letters: $\mathbf{Y}$
- If $\mathbf{Y}$ denotes a matrix, we introduce it as $\mathbf{Y} = [y_{nd}] \in \mathbb{R}^{N \times D}$ for which:
 - there are N rows and D columns
 - y_{nd} denotes the value in the n th row and the d th column
 - $\mathbf{y}_n$ denotes the n th row vector, and $\mathbf{y}_{:d}$ denotes the d th column vector
- If $\mathbf{Y}$ denotes a dataset comprised of multiple matrices, we specify this, and denote each of the matrices as $\mathbf{Y}^{(m)} = [y_{nd}^{(m)}] \in \mathbb{R}^{N \times D_m}$
- We denote parameters for statistical distributions as non-bold, lower case letters. If the parameter is for a random variable stored in a matrix, we add indices. For example, $\tau_{nd}^{(m)}$ is a parameter for a random variable in the n th row and d th column of the m th matrix of some dataset. If $\tau_d^{(m)}$ is a parameter for the same matrix, then it is used for all values in the d th column.

The MOFA model

Consider a multi-omics dataset $\mathbf{Y}$ consisting of omics matrices $\mathbf{Y}^{(m)}$, $m = 1, \dots, M$. Each $\mathbf{Y}^{(m)} = [y_{nd}^{(m)}] \in \mathbb{R}^{N \times D_m}$ contains data for N samples (rows) and D_m features (columns), where $y_{nd}^{(m)}$ is the value for the n th sample and the d th feature from the m th matrix. MOFA [44], assumes the existence of latent factors, and jointly factorizes each $\mathbf{Y}^{(m)}$ into a shared matrix of sample scores $\mathbf{Z} = [z_{nk}] \in \mathbb{R}^{N \times K}$, and an omics specific matrix of feature weights $\mathbf{W}^{(m)} = [w_{kd}^{(m)}] \in \mathbb{R}^{K \times D_m}$. The k th column of $\mathbf{Z}$ contains scores for the k th factor, and the k th row of $\mathbf{W}^{(m)}$ contains corresponding weights for that factor.

MOFA assumes that each observed $y_{nd}^{(m)}$ is a random variable, characterized by a probability distribution conditional on a set of latent random variables $\boldsymbol{\beta}$. It is assumed that the joint likelihood $p(\mathbf{Y}|\boldsymbol{\beta})$ is equal to $\prod_{m=1}^M \prod_{n=1}^N \prod_{d=1}^{D_m} p(y_{nd}^{(m)} | \boldsymbol{\beta})$, and the choice of probability distribution depends on the type of observed data. A Gaussian likelihood is assumed for observed continuous data:

$$p(y_{nd}^{(m)} | \boldsymbol{\beta}) = \text{Normal}(y_{nd}^{(m)} | \mathbf{z}_n: \mathbf{w}_{:d}^{(m)}, 1/\tau_d^{(m)}) \quad (7)$$

where $\mathbf{z}_n$ is the vector of scores for the n th sample, $\mathbf{w}_{:d}^{(m)}$ is the vector of weights the d th feature from the m th matrix, and $\tau_d^{(m)}$ is the precision for that feature. A Bernoulli likelihood is assumed for observed binary data and the logistic link function $\pi(x) = (1 + e^{-x})^{-1}$ is used:

$$p(y_{nd}^{(m)} | \beta) = \text{Bernoulli}\left(y_{nd}^{(m)} | \pi\left(\mathbf{z}_n: \mathbf{w}_{:d}^{(m)}\right)\right) \quad (8)$$

A Poisson likelihood is assumed for observed counts data and the link function $\lambda(x) = \log(1 + e^x)$ is used:

$$p(y_{nd}^{(m)} | \beta) = \text{Poisson}\left(y_{nd}^{(m)} | \lambda\left(\mathbf{z}_n: \mathbf{w}_{:d}^{(m)}\right)\right) \quad (9)$$

The assumed joint prior distribution, $p(\beta)$, is comprised of independent priors: $z_{nk} \sim \text{Normal}(0, 1)$, $w_{kd}^{(m)} = \hat{w}_{kd}^{(m)} s_{kd}^{(m)}$, $\hat{w}_{kd}^{(m)} \sim \text{Normal}(0, 1/\alpha_k^{(m)})$, $\alpha_k^{(m)} \sim \text{Gamma}(1e^{-14}, 1e^{-14})$, $s_{kd}^{(m)} \sim \text{Bernoulli}(\theta_k^{(m)})$, $\theta_k^{(m)} \sim \text{Beta}(1, 1)$, $\tau_d^{(m)} \sim \text{Gamma}(1e^{-14}, 1e^{-14})$.

MOFA uses mean-field variational inference [32, 45, 46] to approximate the joint posterior distribution, $p(\beta|Y)$, with a joint variational distribution factorized over J disjoint groups of variables:

$$\begin{aligned} q(\beta) &= \prod_{j=1}^J q(\beta_j) \\ &= \prod_{n=1}^N \prod_{k=1}^K q(z_{nk}) \prod_{m=1}^M \prod_{k=1}^K q(\alpha_k^{(m)}) q(\theta_k^{(m)}) \prod_{m=1}^M \prod_{d=1}^{D_m} q(\tau_d^{(m)}) \prod_{m=1}^M \prod_{d=1}^{D_m} \prod_{k=1}^K q(\hat{w}_{kd}^{(m)}, s_{kd}^{(m)}) \end{aligned} \quad (10)$$

MOFA infers $q(\beta)$ iteratively until convergence. At each iteration, each $q(\beta_j)$ is updated as

$$q(\beta_j) \propto \exp\{\mathbb{E}_{q_{-j}}[\log p(\beta, \hat{Y})]\} \quad (11)$$

where $\mathbb{E}_{q_{-j}}$ denotes an expectation with respect to the joint variational distribution, after removing $q(\beta_j)$. The dataset $\hat{Y}$ is derived by transformation of Y . Observed data with a Gaussian assumed likelihood are transformed with feature-wise centering, which avoids the need to estimate intercepts. Observed data with a non-Gaussian assumed likelihood are transformed to derive Gaussian pseudo-data. The derivation of Gaussian pseudo-data occurs at each iteration, and is based on a new parameter, $\zeta_{nd}^{(m)}$, which is derived for each sample n and feature d from matrix m . For observed data with a Bernoulli assumed likelihood, a precision parameter, $\tau_{nd}^{(m)}$, is introduced for each sample and feature as part of the transformation:

$$\hat{y}_{nd}^{(m)} = \frac{2y_{nd}^{(m)} - 1}{2\tau_{nd}^{(m)}} \quad (12)$$

$$\tau_{nd}^{(m)} = 2\lambda\left(\zeta_{nd}^{(m)}\right) \quad (13)$$

$$\left(\zeta_{nd}^{(m)}\right)^2 = \mathbb{E}_q\left[\left(\mathbf{z}_n: \mathbf{w}_{:d}^{(m)}\right)^2\right] \quad (14)$$

For observed data with a Poisson assumed likelihood, a precision parameter, $\tau_d^{(m)}$, is introduced for each feature as part of the transformation:

$$\hat{y}_{nd}^{(m)} = \zeta_{nd}^{(m)} - \frac{\pi\left(\zeta_{nd}^{(m)}\right)\left(1 - y_{nd}^{(m)} / \lambda\left(\zeta_{nd}^{(m)}\right)\right)}{\tau_d^{(m)}} \quad (15)$$

$$\zeta_{nd}^{(m)} = \mathbb{E}_q\left[\mathbf{z}_n \cdot \mathbf{w}_{:d}^{(m)}\right] \quad (16)$$

$$\tau_d^{(m)} = 0.25 + 0.17 \times \max\left(\mathbf{y}_{:d}^{(m)}\right) \quad (17)$$

where $\mathbf{y}_{:d}^{(m)}$ is the vector of observed values for the d th feature from the m th matrix. For both Bernoulli and Poisson observed data, the $\hat{y}_{nd}^{(m)}$ values are centered at each iteration, and $\zeta_{nd}^{(m)}$ values are derived using the factorization fit from the preceding iteration. MOFA monitors convergence with the evidence lower bound (ELBO), which is used to evaluate how well the variational distribution approximates the posterior distribution. The ELBO is calculated as:

$$\text{ELBO}(\boldsymbol{\beta}) = \mathbb{E}_q[\log p(\mathbf{Y}|\boldsymbol{\beta})] + \mathbb{E}_q[\log p(\boldsymbol{\beta})] - \mathbb{E}_q[\log q(\boldsymbol{\beta})] \quad (18)$$

For $\mathbf{Y}^{(m)}$ with non-Gaussian assumed likelihood, MOFA uses a lower bound for each $\log p\left(y_{nd}^{(m)}|\boldsymbol{\beta}\right)$. Maximizing this lower bound, coupled with the use of $\hat{y}_{nd}^{(m)}$ values, allows updates of $q(\boldsymbol{\beta})$ based on the assumption of Gaussian observed data [33, 34]. MOFA assesses convergence at regular intervals, based on the percentage change in ELBO after. MOFA allows factors to be dropped during training, based on the fraction of variance explained for each matrix. After each iteration, MOFA identifies factors that do not explain a fraction of variance, for any omics matrix, over a threshold. MOFA then drops one of the identified factors.

Multi-omics data simulated with groundtruth factors

We simulated multi-omics datasets, from groundtruth factors, with various configurations. For each simulation configuration, we generated 30 instances of a multi-omics dataset, $\mathbf{Y}$, consisting of matrices, $\mathbf{Y}^{(m)}$, $m = 1, 2, 3$. We split each $\mathbf{Y}$ into a target dataset, $\mathbf{T}$, and a learning dataset, $\mathbf{L}$. Each $\mathbf{Y}^{(m)} = \left[y_{nd}^{(m)}\right] \in \mathbb{R}^{N \times D_m}$ contained data for $N = N_t + N_l$ samples (rows) and $D_m = 2000$ features (columns), where $y_{nd}^{(m)}$ is the value for the n th sample and the d th feature from the m th matrix. N_t is the number of samples subsequently belonging to $\mathbf{T}$, and N_l is the number of samples belonging to $\mathbf{L}$. We generated each $\mathbf{Y}^{(m)}$ from a different statistical distribution, conditional on a random matrix of sample scores, $\mathbf{Z} = [z_{nk}] \in \mathbb{R}^{N \times K}$, and a random matrix of feature weights, $\mathbf{W}^{(m)} = [w_{kd}^{(m)}] \in \mathbb{R}^{K \times D_m}$. The k th column vector of $\mathbf{Z}$ contained sample scores for the k th groundtruth factor. The k th row vector of $\mathbf{W}^{(m)}$ contained feature weights for that same factor. We varied the number of groundtruth factors across configurations, using $K \in \{20, 30\}$.

We generated $\mathbf{Z}$ based on each sample being a member of a group. In each instance we created two groups of five samples belonging to $\mathbf{T}$, meaning N_t was always equal to 10 samples. We allowed N_t to vary across instances, with samples belonging to $\mathbf{L}$ being in differently sized groups of randomly selected sizes. We used either 20 learning groups of size $\in \{10, 20, 30\}$, or 40 groups of size $\in \{10, 25, 40\}$. For the n th sample and k th groundtruth factor, we generated the score as $z_{nk} \sim \text{Normal}(\mu_{g(n)k}, \sigma_z)$, where $\mu_{g(n)k}$ is the mean parameter for groundtruth factor k , for the group that sample n belonged to, $g(n)$. In each instance we selected $\mu_{g(n)k}$ randomly for each group and groundtruth factor, with probabilities $Pr(3) = 1/8$, $Pr(5) = 3/4$, $Pr(7) = 1/8$. The k th groundtruth factor was differentially active for $\mathbf{T}$ if $\mu_{g(n)k}$ differed between the two target dataset groups. For all instances of a given simulation configuration, the same standard deviation parameter, σ_z , was shared by all groups and groundtruth factors. We varied the latent noise-to-signal ratio across our simulation configurations by using $\sigma_z \in \{0.5, 1.0\}$.

For the k th groundtruth factor, and the d th feature from the m th matrix, we generated the weight as $w_{kd}^{(m)} = \hat{w}_{kd}^{(m)} s_{kd}^{(m)}$. As such, each $w_{kd}^{(m)}$ was the product of a continuous random variable, $\hat{w}_{kd}^{(m)} \sim \text{Normal}(\mu^{(m)}, \sigma_k^{(m)})$, and a binary random variable, $s_{kd}^{(m)} \sim \text{Bernoulli}(\theta_k^{(m)})$. We specified $\mu^{(m)}$, the mean parameter for the m th matrix, with $\mu^{(1)} = 5$; $\mu^{(2)} = 0$; $\mu^{(3)} = 0$. We generated, $\sigma_k^{(m)}$, the standard deviation parameter for the k th groundtruth factor and the m th matrix, with $\sigma_k^{(1)} \sim \text{Uniform}(0.5, 1.5)$; $\sigma_k^{(2)} \sim \text{Uniform}(0.5, 1.5)$; $\sigma_k^{(3)} \sim \text{Uniform}(0.1, 0.2)$. We generated the sparsity for the k th groundtruth factor and the m th matrix, $1 - \theta_k^{(m)}$, with $\theta_k^{(m)} \sim \text{Uniform}(0.15, 0.25)$.

We generated the values in each $\mathbf{Y}^{(m)}$ as:

$$\begin{aligned} y_{nd}^{(1)} &\sim \text{Poisson}\left(\log\left(1 + \exp\left(\mathbf{z}_n \cdot \mathbf{w}_{:d}^{(1)}\right)\right)\right) \\ y_{nd}^{(2)} &\sim \text{Normal}\left(\mathbf{z}_n \cdot \mathbf{w}_{:d}^{(2)}, \sigma_d \sim \text{Uniform}(0.25, 0.75)\right) \\ y_{nd}^{(3)} &\sim \text{Bernoulli}\left(1 / \left(1 + \exp\left(-\mathbf{z}_n \cdot \mathbf{w}_{:d}^{(3)}\right)\right)\right) \end{aligned} \quad (19)$$

where $\mathbf{z}_n$ is the vector of scores for the n th sample, $\mathbf{w}_{:d}^{(m)}$ is the vector of weights the d th feature from the m th matrix, and σ_d is the standard deviation for the d th feature.

We split each $\mathbf{Y}^{(m)}$ into $\mathbf{T}^{(m)} = \begin{bmatrix} \mathbf{t}_{nd}^{(m)} \end{bmatrix} \in \mathbb{R}^{N_t \times D_m}$, which contained values for the target group samples, and $\mathbf{L}^{(m)} = \begin{bmatrix} \mathbf{l}_{nd}^{(m)} \end{bmatrix} \in \mathbb{R}^{N_l \times D_m}$, which contained values for the learning group samples.

Before direct factorization with MOFA we pre-processed simulated $\mathbf{T}$ and $\mathbf{L}$ datasets by removing features with 0 variance across samples. Before factorization with transfer learning with MOTL, we pre-processed simulated $\mathbf{T}$ datasets by removing features that had 0 variance across samples or that had been removed from the corresponding $\mathbf{L}$ datasets.

TCGA multi-omics data acquisition and pre-processing

We used the R packages *TCGAbiolinks* (v.2.25.3) and *SummarizedExperiment* (v.1.28.0) to download and save TCGA mRNA expression, miRNA expression, DNA methylation, and simple nucleotide variation (SNV) data [47–49]. The mRNA and miRNA expression data consisted of raw counts. The DNA methylation data consisted of CpG site β -values,

which had been derived from HM450 array intensities with R package *SeSAmE* (v.1.16.0) [50]. The SNV data consisted of masked somatic mutation files.

We created four reference datasets, using data from three cancer types; acute myeloid leukemia (LAML), pancreatic adenocarcinoma (PAAD) and skin cutaneous melanoma (SKCM). Each reference dataset, $\mathbf{R}$, contained multi-omics data for all samples from either two, or all three of the cancer types. We did not include SNV data in $\mathbf{R}$ datasets containing LAML samples, due to the sparsity of SNV data for LAML. We only used samples that had data for all omics of interest, and only included one sample per study participant. We thus had multi-omics data for 134 LAML samples, 157 PAAD samples, and 435 SKCM samples. We then randomly split each $\mathbf{R}$ into non-overlapping target datasets. Each resulting target dataset, $\mathbf{T}$, contained multi-omics data for five samples per cancer type.

For the evaluation protocol based on TCGA multi-omics data, we merged data from the remaining 29 cancer types into a learning dataset, $\mathbf{L}$. For this $\mathbf{L}$, we only used samples that had data for all four omics, and only included one sample per study participant. This $\mathbf{L}$ contained multi-omics data (mRNA, miRNA, DNA methylation, and SNV) for 7217 samples.

For the application of MOTL to the pd-GBSC target datasets, we created a new learning dataset by merging data from all 32 cancer types. This new learning dataset contained multi-omics data (mRNA, miRNA, DNA methylation, and SNV) for 7866 samples.

Before direct factorization with MOFA, we pre-processed $\mathbf{R}$, $\mathbf{T}$, and $\mathbf{L}$ datasets in the same way. For mRNA data, we removed genes that map to the Y chromosome. For both mRNA and miRNA, we removed genes if they had a count of zero in $\geq 90\%$ of samples, or had zero variance across samples. We normalized mRNA and miRNA counts with the *DESeq2* (v.1.38.0) R package [51], and $\log_2(x + 1)$ transformed the normalized counts. For DNA methylation data, we removed CpG sites that map to the X or Y chromosome, were masked during *SeSAmE* quality control, had missing values in $\geq 20\%$ of samples, or had zero variance across samples. We converted DNA methylation β -values to M -values [52]. We included SNV records whose variant classification was either *Frame_Shift_Del*, *Frame_Shift_Ins*, *In_Frame_Del*, *In_Frame_Ins*, *Missense_Mutation*, *Nonsense_Mutation*, *Nonstop_Mutation*, *Splice_Site*, or *Translation_Start_Site*. We then created binary SNV matrices aggregated by gene and sample. We removed genes from SNV matrices if the mutation rate across samples was $\leq 1\%$. We filtered all omics to include only the 5000 most variable features. We did not perform any batch effect correction on $\mathbf{L}$ datasets in order to preserve biological signal [53]. We checked each $\mathbf{R}$ for batch effects with visualizations of UMAP coordinates [54]. We used the R package *uwot* (v.0.1.14) to derive UMAP coordinates from MOFA factorizations, and we did not observe the need to correct $\mathbf{R}$ datasets for batch effects.

Before factorization with transfer learning with MOTL, we pre-processed $\mathbf{T}$ datasets by removing all omics features that had zero variance across samples, or that had been removed from $\mathbf{L}$ during pre-processing. We used *DESeq2* to normalize mRNA and miRNA counts with the geometric means from $\mathbf{L}$, and then $\log_2(x + 1)$ transformed the normalized counts. We converted DNA methylation β -values to M -values,

and converted SNV data to binary matrices after filtering on variant classification, as described previously.

Glioblastoma target dataset acquisition and pre-processing

We created four pd-GBSC target datasets, based on multi-omics profiling conducted by Santamarina-Ojeda et al. [36] for four normal brain samples and nine patient-derived glioblastoma stem cell (pd-GBSC) cultures. The nine cancer samples had been previously classified into three subtypes thanks to transcriptome-based signatures: classical (CL), proneural (PN), and mesenchymal (MS). Each pd-GBSC target dataset contained mRNA expression and DNA methylation data for the four normal brain cortex samples, as well as either all nine cancer samples or just the samples from a subtype.

For factorization with transfer learning with MOTL, the pd-GBSC target datasets initially consisted of mRNA expression raw counts and DNA methylation β -values. Before factorization with MOTL, we pre-processed a pd-GBSC target dataset by removing all omics features that had zero variance across samples, or that had been removed from L during pre-processing. We used *DESeq2* to normalize mRNA counts with the geometric means from L , and then $\log_2(x + 1)$ transformed the normalized counts. We converted DNA methylation β -values to M -values.

For direct MOFA factorization, without transfer learning, the pd-GBSC target datasets initially consisted of the same mRNA expression data, but already normalized and transformed by Santamarina-Ojeda et al. [36], and DNA methylation β -values. Before direct MOFA factorization, we pre-processed mRNA data by removing genes that map to the Y chromosome, if they had a count of zero in $\geq 90\%$ of samples, or had zero variance across samples. We pre-processed DNA methylation data by removing CpG sites that had missing values in $\geq 20\%$ of samples, or had zero variance across samples. We converted DNA methylation β -values to M -values. We filtered both omics to include only the 5000 most variable features.

Application of MOFA to simulated, TCGA, and glioblastoma multi-omics datasets

We factorized simulated target, T , and learning, L , datasets with the MOFA Python implementation *mofapy2* (v.0.6.4). The number of factors we used for each MOFA factorization was equal to the lesser of the number of samples and the number of groundtruth factors that were differentially active when simulating the dataset. The k th groundtruth factor was differentially active for a dataset if the mean parameter, $\mu_{g(n)k}$, for the sample scores for that factor, was not the same for all groups of samples in the dataset. We specified observed data likelihoods corresponding to those used for simulating $Y^{(m)}$ matrices. We set the maximum number of iterations to 10,000 to ensure convergence. For the remaining settings, we used the *mofapy2* defaults, meaning that all datasets were feature-wise centered during factorization fitting.

We factorized pre-processed TCGA reference, R , target, T , and learning, L , datasets with the MOFA Python implementation *mofapy2* (v.0.7.0). We specified Gaussian as the observed data likelihood for mRNA, miRNA, and DNA methylation data, and specified Bernoulli as the likelihood for SNV data. For the L datasets, we started the factorization with 100 factors and allowed factors to be dropped based on the fraction of variance

explained, for which we set the threshold to 0.001. We set the threshold so low in order to retain factors that explained little of the variance in L , yet could be potentially relevant for transfer learning. For R datasets, we also started with 100 factors. For T datasets, we started with the maximum number of factors allowed by MOFA, which was either 10 factors (two cancer types) or 15 factors (three cancer types). For R and T datasets, we dropped factors based on a threshold of 0.01, in order to only retain relevant factors. For all TCGA datasets, we set the maximum number of iterations to 10,000, to ensure convergence, and the frequency of convergence checking to five, to ensure that the algorithm had stopped dropping factors before converging.

When saving the factorizations of simulated and TCGA L datasets, we set the expectations argument to *all*. We did this to ensure that the point estimate for each precision parameter was saved in addition to those that are saved by default.

We factorized the pre-processed pd-GBSC target datasets with the MOFA Python implementation *mofapy2* (v.0.7.0). We specified Gaussian as the observed data likelihood for the mRNA and the DNA methylation data. We started with the maximum allowable number of factors and dropped factors based on a threshold of 0.01.

Application of MOTL to simulated, TCGA, and glioblastoma multi-omics datasets

We applied MOTL to simulated, TCGA and pd-GBSC multi-omics target datasets. For each target dataset, we used point estimates of feature weight and precision values saved from the MOFA factorization of the corresponding learning, L , dataset. For observed data with a Gaussian or Poisson assumed likelihood, the transferred value of the precision for each feature, $\tau_d^{(m)}$, was held fixed throughout iterations of MOTL updates. For observed data with a Bernoulli assumed likelihood, we initialized the value of the precision for each sample and feature, $\tau_{nd}^{(m)}$, with a feature-wise average, $\tau_d^{(m)}$, of the $\tau_{nd}^{(m)}$ values from the factorization of L . The precisions for Bernoulli observed data were then iteratively updated by MOTL. We estimated intercepts using likelihoods assumed for L , combined with outputs from the MOFA factorization of L . For Gaussian observed data we calculated the intercept for the d th feature, from the m th matrix, as $a_d^{(m)} = \frac{1}{N_l} \sum_{n=1}^{N_l} l_{nd}^{(m)}$, where $l_{nd}^{(m)}$ denotes an uncentered learning dataset value after pre-processing. For Poisson and Bernoulli observed data we obtained maximum likelihood estimates of $a_d^{(m)}$ values, for which we used the *mle* function from the R package *stats4* (v.4.2.0). For Poisson observed data we initialized each estimate with $a_d^{(m)} = \log\left(-1 + \exp\left(\frac{1}{N_l} \sum_{n=1}^{N_l} l_{nd}^{(m)}\right)\right)$ and for Bernoulli observed data we initialized it with $a_d^{(m)} = \log\left(\left(\frac{1}{N_l} \sum_{n=1}^{N_l} l_{nd}^{(m)}\right)\left(1 - \frac{1}{N_l} \sum_{n=1}^{N_l} l_{nd}^{(m)}\right)^{-1}\right)$.

To evaluate robustness, we also applied MOTL to each simulated target dataset after permuting the values in feature vectors of the weight matrices inferred from L . We used a range of proportions for each simulation instance, with the proportions of feature weight vectors permuted in each instance being:

$$\{0.0, 0.1, 0.2, 0.3, 0.4, 0.5, 0.6, 0.7, 0.8, 0.9, 1.0\}$$

When checking the ELBO for convergence, we used 0.0005% as the threshold, which is the default for MOFA. The algorithm was stopped when the absolute change in ELBO was under this threshold for two consecutive checks, and we set the maximum number of iterations to 10,000 to be consistent with our application of MOFA. For TCGA and pd-GBSC target datasets, and for simulated target datasets factorized after permuting feature weight values, we allowed factors to be dropped based on a threshold of 0.01 for the fraction of variance explained. We checked the ELBO after every five iterations, to ensure that the algorithm had stopped dropping factors before converging.

Evaluation methods

Groundtruth factors: For each simulated T (the [Methods “Multi-omics data simulated with groundtruth factors”](#) section), the groundtruth factor values were contained in the corresponding simulated Z and $W^{(m)}$ matrices. The sample scores for the k th groundtruth factor were contained in $z_{:k}$, the k th column vector of simulated Z . The feature weights for the m th matrix, for that same groundtruth factor, were contained in $w_{k:}^{(m)}$, the k th row vector of simulated $W^{(m)}$. For each TCGA T (the [Methods “TCGA multi-omics data acquisition and pre-processing”](#) section), the groundtruth factors were based on the R dataset which we had split to create T . We factorized each R with MOFA, and treated the inferred $z_{:k}$ and $w_{k:}^{(m)}$ vectors as groundtruth factors for each T that had been created by splitting R .

Differentially active groundtruth factors: For each simulated and TCGA T , groundtruth factor k was differentially active if the group means for groundtruth $z_{:k}$ differed between the target dataset groups. For each simulated T , this was the group mean, $\mu_{g(n)k}$, used to simulate groundtruth $z_{:k}$. For each TCGA T , the factorization of corresponding R provided groundtruth $z_{:k}$ and $w_{k:}^{(m)}$ factor vectors. We performed either the Wilcoxon rank sum test (two cancer types), or the Kruskal-Wallis test (three cancer types), on each groundtruth $z_{:k}$ to determine if there was a statistically significant difference between the cancer types. We classed a groundtruth factor as differentially active if its BH-adjusted p -value was below 0.05.

Post-processing: We post-processed inferred and groundtruth $W^{(m)}$ matrices before evaluation. We scaled each feature vector, $w_{d:}^{(m)}$, by its Frobenius norm. We then centered each factor vector, $w_{k:}^{(m)}$, of scaled values separately for each m . We then concatenated $w_{k:}^{(m)}$ vectors to produce a single vector, $w_{k:}$, of centered and scaled feature weights for each factor k .

Best hits: For each factorization of each simulated and TCGA T , we identified the best hits between the factor vectors inferred with the factorization of T , and the groundtruth factor vectors. For two sets of vectors $\{v_1, \dots, v_{K_v}\}$ and $\{x_1, \dots, x_{K_x}\}$, we define the best hit for vector v_{k_v} as

$$\text{BestHit}(v_{k_v}) = \arg \max_{x_{k_x}} \text{cor}(v_{k_v}, x_{k_x}) \quad (20)$$

where $\text{cor}(v, x)$ is the Pearson correlation coefficient between vectors v and x . We define the best hit for vector x_{k_x} as

$$\text{BestHit}(\mathbf{x}_{k_x}) = \arg \max_{\mathbf{v}_{k_v}} \text{cor}(\mathbf{x}_{k_x}, \mathbf{v}_{k_v}) \quad (21)$$

For each simulated T , we identified best hits between inferred and groundtruth $\mathbf{w}_{k:}$ vectors. For each TCGA T , we identified best hits between inferred and groundtruth $\mathbf{w}_{k:}$ vectors, as well as between inferred and groundtruth $\mathbf{z}_{:k}$ vectors. We used shared features when calculating correlations for $\mathbf{w}_{k:}$ vectors, and we used shared samples for $\mathbf{z}_{:k}$ vectors. We calculated p -values for the correlations, and only considered correlations with a p -value < 0.05 (two-sided alternative hypothesis) when identifying best hits.

F -measure values: For each factorization of each TCGA T , we calculated an F -measure value to assess the overall correlation between factor vectors inferred with the factorization of T , and groundtruth factor vectors. We based this on the F -measure presented by Saelens et al. [55], which we adapted in order to assess correlations. For a given set of inferred factor vectors, $\{\mathbf{v}_1, \dots, \mathbf{v}_{K_v}\}$, and a set of groundtruth factor vectors, $\{\mathbf{x}_1, \dots, \mathbf{x}_{K_x}\}$, we calculated the F -measure as

$$FM = 2 / ((1/Relevance) + (1/Recovery)) \quad (22)$$

where

$$Relevance = \frac{1}{K_v} \sum_{k_v=1}^{K_v} \text{cor}(\mathbf{v}_{k_v}, \text{BestHit}(\mathbf{v}_{k_v})) \quad (23)$$

and

$$Recovery = \frac{1}{K_x} \sum_{k_x=1}^{K_x} \text{cor}(\mathbf{x}_{k_x}, \text{BestHit}(\mathbf{x}_{k_x})) \quad (24)$$

Here, $\text{cor}(\mathbf{v}, \mathbf{x})$ is the Pearson correlation coefficient between vectors $\mathbf{v}$ and $\mathbf{x}$. We calculated F -measure values for sets of inferred and groundtruth $\mathbf{w}_{k:}$ vectors, as well as for $\mathbf{z}_{:k}$ vectors.

$F1$ scores: We calculated $F1$ scores to evaluate the factorizations of simulated and TCGA T datasets:

$$\begin{aligned} F1 &= (2 \times \text{Precision} \times \text{Recall}) / (\text{Precision} + \text{Recall}) \\ \text{Precision} &= \text{True Positives} / \text{Predicted Positives} \\ \text{Recall} &= \text{True Positives} / \text{Actual Positives} \end{aligned} \quad (25)$$

Actual Positives were the groundtruth factors of T , that were differentially active.

Predicted Positives were the groundtruth factors that were predicted as being differentially active, based on the factorization of T . We firstly performed either the Wilcoxon rank sum test (two groups), or the Kruskal-Wallis test (three groups), on each $\mathbf{z}_{:k}$ vector inferred with the factorization of T , and classed factors with a p -value < 0.05 as differentially active. For inferred factors classed as differentially active, we identified the best hits for their corresponding inferred $\mathbf{w}_{k:}$ vectors. We selected these best hits from the set of groundtruth $\mathbf{w}_{k:}$ vectors for T . If groundtruth $\mathbf{w}_{k:}$ was selected as a best hit for a differentially active inferred factor, then groundtruth factor k was predicted as being differentially active.

True Positives were the differentially active groundtruth factors that were predicted as being differentially active, based on the factorization of T .

Statistical testing of differences between factorization methods: We calculated the differences in evaluation measures between factorization methods, and tested the statistical significance of these differences. To do this, we fit generalized least squares regressions with the R package *nlme* (v.3.1.157) [56]. We fit a single regression to model the $F1$ scores for simulated data. For TCGA data, we fit a separate regression for each evaluation measure. For each regression, we modeled $y_i = \beta_0 + \mathbf{d}_i \boldsymbol{\beta}_d + \mathbf{f}_i \boldsymbol{\beta}_f + \epsilon_i$. The vector $\mathbf{d}_i = (d_{i1}, \dots, d_{iT})$ indicates the simulation configuration, or cancer type, that y_i relates to, and vector $\mathbf{f}_i = (f_{i1}, \dots, f_{iM})$ indicates the factorization method. Vectors $\boldsymbol{\beta}_d = (\beta_{d1}, \dots, \beta_{dT})^\top$ and $\boldsymbol{\beta}_f = (\beta_{f1}, \dots, \beta_{fM})^\top$ are estimated fixed effects and ϵ_i is the residual. We incorporated correlations between residuals from the same target dataset using the compound symmetry structure method. We calculated contrasts for the factorization method effects in $\boldsymbol{\beta}_f$ using the R package *emmeans* (v.1.8.7) [57], and used Tukey-adjusted p -values for assessing statistical significance.

Differentially active factors from glioblastoma target datasets: We identified differentially active factors from the MOTL factorization of each pd-GBSC target dataset, as well as from the direct MOFA factorization (without transfer learning), of each pd-GBSCs target dataset. We performed the Wilcoxon Rank Sum test on each $\mathbf{z}_{:,k}$ vector inferred with the factorization of a pd-GBSC target dataset. We classed factors with a BH-adjusted p -value < 0.05 as differentially active between the normal samples and the cancer samples.

Gene set enrichment analysis: We used R package *fgsea* (v.1.24.0) [58] to perform gene set enrichment analysis on differentially active groundtruth factors that were true positives for factorizations of TCGA T datasets. For each differentially active groundtruth TCGA factor k , we analyzed vector $\mathbf{w}_k^{(m)}$ if the fraction of mRNA variance explained by k was > 0.01 , and where m corresponded to the mRNA matrix. We tested KEGG, Reactome, GO:BP, and GO:CC gene sets that have a size of between 15 and 500 genes, obtained using the R package *msigdb* (v.7.5.1). We used an BH-adjusted p -value cutoff of 0.01 for selecting enriched gene sets. We also performed gene set enrichment analysis on differentially active factors from the pd-GBSC target datasets, and used the same criteria as outlined above for differentially active groundtruth TCGA factors. We filtered, clustered, and plotted enrichment analysis results with the Python package *orsum* (v.1.8.0) [37]. We ran *orsum* with `maxRepSize = 2000`; `maxTermSize = 3000`; `minTermSize = 15`; `numberOfTermsToPlot = 30`.

Processing time

We used a Dell computer with 20 cores at 3GHz, and 64 GB of RAM, to perform factorizations. To pre-process and factorize the L used in the TCGA evaluation protocol, it took 26,405 seconds (over seven hours). Hence, we have made the factorization of a large TCGA L dataset publicly available for transfer learning. It took an average of 37 seconds to pre-process a T dataset, comprised of four omics, and factorize it directly with MOFA. The average time increased to 134 seconds for MOTL.

Supplementary Information

The online version contains supplementary material available at <https://doi.org/10.1186/s13059-025-03675-7>.

Additional file 1. Supplementary methods and supplementary Fig. S1–S5.

Additional file 2. Table S1: gene sets associated with differentially active groundtruth TCGA factors.

Additional file 3. Table S2: the percentage of variance explained, by groundtruth TCGA factors, for each omics.

Additional file 4. Table S3: gene sets associated with factors differentially active between all pd-GBSC samples and normal samples.

Additional file 5. Table S4: gene sets associated with factors differentially active between all pd-GBSC samples and normal samples, as well as those associated with factors differentially active between pd-GBSC subtype samples and normal samples.

Acknowledgements

We would like to thank Carl Herrmann, Céline Chevalier, Kim-Anh Lê Cao, Lionel Spinelli, Olivia Angelin-Bonnet, and two anonymous reviewers for helpful feedback and discussions.

For the purpose of open access, the authors have applied a CC BY-NC-ND public copyright licence to any Author Accepted Manuscript (AAM) version arising from this submission.

Peer review information

Andrew Cosgrove was the primary editor of this article and managed its editorial process and peer review in collaboration with the rest of the editorial team. The peer-review history is available in the online version of this article.

Authors' contributions

A.B., L.C., P.V., M.V., and D.P.H. conceived the project. D.P.H. conceived the model and evaluation protocols with input from all authors. D.P.H. and M.T. contributed to the code and implemented the model. D.P.H. generated the figures and results. D.P.H., A.B., and M.V. wrote the manuscript with feedback from all authors. A.B. and M.V. supervised the project. All authors reviewed and approved the final manuscript.

Funding

This research was supported by l'Agence Nationale de la Recherche (ANR), project ANR-21-CE45-0001-01, by France 2030 PEPR Digital Health managed by ANR, project ANR-22-PESN-0013, by the Royal Society of New Zealand, Catalyst project CSG-MAU1902 and by the Association Française contre les Myopathies (AFM).

Data availability

An open source (GPL-3.0 license) R implementation of MOTL is available at [35] <https://github.com/david-hirst/MOTL> along with the code for the evaluation protocols. The version of the code that can be used to reproduce the analyses in this study has been archived on Zenodo [59] (<https://doi.org/10.5281/zenodo.15486697>). We obtained the TCGA [48] data used in this study from the GDC data portal [60] (<https://portal.gdc.cancer.gov/>) with the TCGA biolinks R package [47]. We assembled the glioblastoma target dataset with data archived on Zenodo [61] (<https://doi.org/10.5281/zenodo.7380252>). The factorization fit of the full TCGA learning dataset, used for the application of MOTL to the glioblastoma target dataset, is available at Zenodo [62] (<https://doi.org/10.5281/zenodo.10848217>).

Declarations

Ethics approval and consent to participate

Not applicable.

Consent for publication

Not applicable.

Competing interests

The authors declare no competing interests.

Received: 30 May 2024 Accepted: 30 June 2025

Published online: 25 July 2025

References

1. Conesa A, Beck S. Making multi-omics data accessible to researchers. *Sci Data*. 2019;6(1):251. <https://doi.org/10.1038/s41597-019-0258-4>.
2. Dermitzakis ET. From gene expression to disease risk. *Nat Genet*. 2008;40(5):492–3. <https://doi.org/10.1038/ng0508-492>.
3. Manzoni C, Kia DA, Vandrovicova J, Hardy J, Wood NW, Lewis PA, et al. Genome, transcriptome and proteome: the rise of omics data and their integration in biomedical sciences. *Brief Bioinforma*. 2018;19(2):286–302. <https://doi.org/10.1093/bib/bbw114>.
4. Subramanian I, Verma S, Kumar S, Jere A, Anamika K. Multi-omics Data Integration, Interpretation, and Its Application. *Bioinforma Biol Insights*. 2020;14. <https://doi.org/10.1177/1177932219899051>.

5. Huang S, Chaudhary K, Garmire LX. More Is Better: Recent Progress in Multi-Omics Data Integration Methods. *Front Genet.* 2017;8:84. <https://doi.org/10.3389/fgene.2017.00084>.
6. Rappoport N, Shamir R. Multi-omic and multi-view clustering algorithms: review and cancer benchmark. *Nucleic Acids Res.* 2018;46(20):10546–62. <https://doi.org/10.1093/nar/gky889>.
7. Pierre-Jean M, Deleuze JF, Le Floch E, Mauger F. Clustering and variable selection evaluation of 13 unsupervised methods for multi-omics data integration. *Brief Bioinforma.* 2020;21(6):2011–30. <https://doi.org/10.1093/bib/bbz138>.
8. Huang Y, Du C, Xue Z, Chen X, Zhao H, Huang L. What Makes Multi-Modal Learning Better than Single (Provably). In: *Advances in Neural Information Processing Systems*. vol. 34. Curran Associates, Inc.; 2021. pp. 10944–56. https://papers.nips.cc/paper_files/paper/2021/file/5aa3405a3f865c10f420a4a7b55cbff3-Paper.pdf.
9. Cantini L, Zakeri P, Hernandez C, Naldi A, Thieffry D, Remy E, et al. Benchmarking joint multi-omics dimensionality reduction approaches for the study of cancer. *Nat Commun.* 2021;12(1):124. <https://doi.org/10.1038/s41467-020-20430-7>.
10. Ballard J, Wang Z, Li W, Shen L, Long Q. Deep learning-based approaches for multi-omics data integration and analysis. *BioData Min.* 2024;17(1):38. <https://doi.org/10.1186/s13040-024-00391-z>.
11. Baião AR, Cai Z, Poulos RC, Robinson PJ, Reddel RR, Zhong Q, et al. A technical review of multi-omics data integration methods: from classical statistical to deep generative approaches. *arXiv.* 2025. <https://doi.org/10.48550/arXiv.2501.17729>.
12. Chauvel C, Novoloaca A, Veyre P, Reynier F, Becker J. Evaluation of integrative clustering methods for the analysis of multi-omics data. *Brief Bioinforma.* 2020;21(2):541–52. <https://doi.org/10.1093/bib/bbz015>.
13. Wen Y, Zheng L, Leng D, Dai C, Lu J, Zhang Z, et al. Deep Learning-Based Multiomics Data Integration Methods for Biomedical Application. *Adv Intell Syst.* 2023;5(5):2200247. <https://doi.org/10.1002/aisy.202200247>.
14. Stein-O'Brien GL, Arora R, Culhane AC, Favorov AV, Garmire LX, Greene CS, et al. Enter the Matrix: Factorization Uncovers Knowledge from Omics. *Trends Genet.* 2018;34(10):790–805. <https://doi.org/10.1016/j.tig.2018.07.003>.
15. Tini G, Marchetti L, Priami C, Scott-Boyer MP. Multi-omics integration—a comparison of unsupervised clustering methodologies. *Brief Bioinforma.* 2019;20(4):1269–79. <https://doi.org/10.1093/bib/bbx167>.
16. Chalise P, Fridley BL. Integrative clustering of multi-level 'omic data based on non-negative matrix factorization algorithm. *PLoS ONE.* 2017;12(5):e0176278. <https://doi.org/10.1371/journal.pone.0176278>.
17. Lock EF, Hoadley KA, Marron JS, Nobel AB. Joint and individual variation explained (JIVE) for integrated analysis of multiple data types. *Ann Appl Stat.* 2013;7(1):523–42. <https://doi.org/10.1214/12-AOAS597>.
18. Meng C, Helm D, Frejno M, Kuster B. moCluster: Identifying Joint Patterns Across Multiple Omics Data Sets. *J Proteome Res.* 2016;15(3):755–65. <https://doi.org/10.1021/acs.jproteome.5b00824>.
19. Argelaguet R, Velten B, Arnol D, Dietrich S, Zenz T, Marioni JC, et al. Multi-Omics Factor Analysis—a framework for unsupervised integration of multi-omics data sets. *Mol Syst Biol.* 2018;14(6):e8124. <https://doi.org/10.15252/msb.20178124>.
20. Taroni JN, Grayson PC, Hu Q, Eddy S, Kretzler M, Merkel PA, et al. MultiPLIER: A Transfer Learning Framework for Transcriptomics Reveals Systemic Features of Rare Disease. *Cell Syst.* 2019;8(5):380–394.e4. <https://doi.org/10.1016/j.cels.2019.04.003>.
21. Banerjee J, Taroni JN, Allaway RJ, Prasad DV, Guinney J, Greene C. Machine learning in rare disease. *Nat Methods.* 2023;1–12. <https://doi.org/10.1038/s41592-023-01886-z>.
22. Weiss K, Khoshgoftaar TM, Wang D. A survey of transfer learning. *J Big Data.* 2016;3(1):9. <https://doi.org/10.1186/s40537-016-0043-6>.
23. Stein-O'Brien GL, Clark BS, Sherman T, Zibetti C, Hu Q, Sealfon R, et al. Decomposing Cell Identity for Transfer Learning across Cellular Measurements, Platforms, Tissues, and Species. *Cell Syst.* 2019;8(5):395–411.e8. <https://doi.org/10.1016/j.cels.2019.04.004>.
24. Veeramachaneni SD, Pujari AK, Padmanabhan V, Kumar V. A Maximum Margin Matrix Factorization based Transfer Learning Approach for Cross-Domain Recommendation. *Appl Soft Comput.* 2019;85:105751. <https://doi.org/10.1016/j.asoc.2019.105751>.
25. Dong A, Li Z, Zheng Q. Transferred Subspace Learning Based on Non-negative Matrix Factorization for EEG Signal Classification. *Front Neurosci.* 2021;15. <https://doi.org/10.3389/fnins.2021.647393>.
26. Peng M, Li Y, Wamsley B, Wei Y, Roeder K. Integration and transfer learning of single-cell transcriptomes via cFIT. *Proc Natl Acad Sci.* 2021;118(10):e2024383118. <https://doi.org/10.1073/pnas.2024383118>.
27. Fertig EJ, Ding J, Favorov AV, Parmigiani G, Ochs MF. CoGAPS: an R/C++ package to identify patterns and biological process activity in transcriptomic data. *Bioinformatics.* 2010;26(21):2792–3. <https://doi.org/10.1093/bioinformatics/btq503>.
28. Sharma G, Colantuoni C, Goff LA, Fertig EJ, Stein-O'Brien G. projectR: an R/Bioconductor package for transfer learning via PCA, NMF, correlation and clustering. *Bioinformatics.* 2020;36(11):3592–3. <https://doi.org/10.1093/bioinformatics/btaa183>.
29. Davis-Marcisak EF, Fitzgerald AA, Kessler MD, Danilova L, Jaffee EM, Zaidi N, et al. Transfer learning between preclinical models and human tumors identifies a conserved NK cell activation signature in anti-CTLA-4 responsive tumors. *Genome Med.* 2021;13(1). <https://doi.org/10.1186/s13073-021-00944-5>.
30. Mao W, Zaslavsky E, Hartmann BM, Sealfon SC, Chikina M. Pathway-level information extractor (PLIER) for gene expression data. *Nat Methods.* 2019;16(7):607–10. <https://doi.org/10.1038/s41592-019-0456-1>.
31. Collado-Torres L, Nellore A, Kammers K, Ellis SE, Taub MA, Hansen KD, et al. Reproducible RNA-seq analysis using recount2. *Nat Biotechnol.* 2017;35(4):319–21. <https://doi.org/10.1038/nbt.3838>.
32. Blei DM, Kucukelbir A, McAuliffe JD. Variational Inference: A Review for Statisticians. *arXiv.* 2017. <https://doi.org/10.48550/arXiv.1601.00670>.
33. Jaakkola TS, Jordan MI. Bayesian parameter estimation via variational methods. *Stat Comput.* 2000;10(1):25–37. <https://doi.org/10.1023/A:1008932416310>.

34. Seeger M, Bouchard G. Fast Variational Bayesian Inference for Non-Conjugate Matrix Factorization Models. In: Proceedings of the Fifteenth International Conference on Artificial Intelligence and Statistics. PMLR; 2012. pp. 1012–8. <https://proceedings.mlr.press/v22/seeger12.html>.
35. Hirst DP, Térézol M, Cantini L, Villoutreix P, Vignes M, Baudot A. MOTL: enhancing multi-omics matrix factorization with transfer learning. Github. 2025. <https://github.com/david-hirst/MOTL>.
36. Ozisik O, Térézol M, Baudot A. orsum: a Python package for filtering and comparing enrichment analyses using a simple principle. *BMC Bioinformatics*. 2022;23(1):293. <https://doi.org/10.1186/s12859-022-04828-2>.
37. Bunne C, Schiebinger G, Krause A, Regev A, Cuturi M. Optimal transport for single-cell and spatial omics. *Nat Rev Methods Prim*. 2024;4(1):58. <https://doi.org/10.1038/s43586-024-00334-2>.
38. Pan SJ, Tsang IW, Kwok JT, Yang Q. Domain Adaptation via Transfer Component Analysis. *IEEE Trans Neural Netw*. 2011;22(2):199–210. <https://doi.org/10.1109/TNN.2010.2091281>.
39. Gargano MA, Matentzoglou N, Coleman B, Addo-Lartey EB, Anagnostopoulos A, Anderton J, et al. The Human Phenotype Ontology in 2024: phenotypes around the world. *Nucleic Acids Res*. 2024;52(D1):D1333–46. <https://doi.org/10.1093/nar/gkad1005>.
40. Schriml LM, Arze C, Nadendla S, Chang YWW, Mazaitis M, Felix V, et al. Disease Ontology: a backbone for disease semantic integration. *Nucleic Acids Res*. 2012;40(D1):D940–6. <https://doi.org/10.1093/nar/gkr972>.
41. Rohart F, Gautier B, Singh A, Cao KAL. mixOmics: An R package for 'omics feature selection and multiple data integration. *PLoS Comput Biol*. 2017;13(11):e1005752. <https://doi.org/10.1371/journal.pcbi.1005752>.
42. Mo Q, Shen R, Guo C, Vannucci M, Chan KS, Hilsenbeck SG. A fully Bayesian latent variable model for integrative clustering analysis of multi-type omics data. *Biostatistics (Oxford, England)*. 2018;19(1):71–86. <https://doi.org/10.1093/biostatistics/kxx017>.
43. Argelaguet R, Arnol D, Bredikhin D, Deloro Y, Velten B, Marioni JC, et al. MOFA+: a statistical framework for comprehensive integration of multi-modal single-cell data. *Genome Biol*. 2020;21(1):111. <https://doi.org/10.1186/s13059-020-02015-1>.
44. Grimmer J. An Introduction to Bayesian Inference via Variational Approximations. *Political Anal*. 2011;19(1):32–47.
45. Fox CW, Roberts SJ. A tutorial on variational Bayesian inference. *Artif Intell Rev*. 2012;38(2):85–95. <https://doi.org/10.1007/s10462-011-9236-8>.
46. Silva TC, Colaprico A, Olsen C, D'Angelo F, Bontempi G, Ceccarelli M, et al. TCGA Workflow: Analyze cancer genomics and epigenomics data using Bioconductor packages. *F1000Research*. 2016;5:1542.
47. Hutter C, Zenklusen JC. The Cancer Genome Atlas: Creating Lasting Value beyond Its Data. *Cell*. 2018;173(2):283–5. <https://doi.org/10.1016/j.cell.2018.03.042>.
48. Mounir M, Lucchetta M, Silva TC, Olsen C, Bontempi G, Chen X, et al. New functionalities in the TCGAAbilinks package for the study and integration of cancer data from GDC and GTEx. *PLoS Comput Biol*. 2019;15(3):e1006701. <https://doi.org/10.1371/journal.pcbi.1006701>.
49. Zhou W, Triche TJ Jr, Laird PW, Shen H. SeSAMe: reducing artifactual detection of DNA methylation by Infinium BeadChips in genomic deletions. *Nucleic Acids Res*. 2018;46(20):e123. <https://doi.org/10.1093/nar/gky691>.
50. Love MI, Huber W, Anders S. Moderated estimation of fold change and dispersion for RNA-seq data with DESeq2. *Genome Biol*. 2014;15(12):550. <https://doi.org/10.1186/s13059-014-0550-8>.
51. Du P, Zhang X, Huang CC, Jafari N, Kibbe WA, Hou L, et al. Comparison of Beta-value and M-value methods for quantifying methylation levels by microarray analysis. *BMC Bioinformatics*. 2010;11(1):587. <https://doi.org/10.1186/1471-2105-11-587>.
52. Lee AJ, Park Y, Doing G, Hogan DA, Greene CS. Correcting for experiment-specific variability in expression compendia can remove underlying signals. *GigaScience*. 2020;9(11). <https://doi.org/10.1093/gigascience/giaa117>.
53. McInnes L, Healy J, Melville J. UMAP: Uniform Manifold Approximation and Projection for Dimension Reduction. *arXiv*. 2020. <https://doi.org/10.48550/arXiv.1802.03426>.
54. Saelens W, Cannoodt R, Saeys Y. A comprehensive evaluation of module detection methods for gene expression data. *Nat Commun*. 2018;9(1):1090. <https://doi.org/10.1038/s41467-018-03424-4>.
55. Pinheiro JC, Bates DM. *Mixed-Effects Models in S and S-PLUS*. Statistics and Computing. New York: Springer; 2000.
56. Searle SR, Speed FM, Milliken GA. Population Marginal Means in the Linear Model: An Alternative to Least Squares Means. *Am Stat*. 1980;34(4):216–21.
57. Sergushichev AA. An algorithm for fast preranked gene set enrichment analysis using cumulative statistic calculation. *bioRxiv*. 2016. <https://doi.org/10.1101/060012>. Accessed 30 Mar 2021.
58. Santamarina-Ojeda P, Tejedor JR, Pérez RF, López V, Roberti A, Mangas C, et al. Multi-omic integration of DNA methylation and gene expression data reveals molecular vulnerabilities in glioblastoma. *Mol Oncol*. 2023;17(9):1726–43. <https://doi.org/10.1002/1878-0261.13479>.
59. Hirst DP, Térézol M, Cantini L, Villoutreix P, Vignes M, Baudot A. david-hirst/MOTL: MOTL: enhancing multi-omics matrix factorization with transfer learning (v1.0.0). Zenodo. 2025. <https://doi.org/10.5281/zenodo.15486697>.
60. Grossman RL, Heath AP, Ferretti V, Varmus HE, Lowy DR, Kibbe WA, et al. Toward a Shared Vision for Cancer Genomic Data. *N Engl J Med*. 2016;375(12):1109–12. <https://doi.org/10.1056/nejmp1607591>.
61. Santamarina-Ojeda P, Tejedor JR, Pérez RF, López V, Roberti A, Mangas C, et al. Multi-omic integration of DNA methylation and gene expression data reveals molecular vulnerabilities in glioblastoma (processed data) (v1.0) [Data set]. Zenodo. 2023. <https://doi.org/10.5281/zenodo.7380252>.
62. Hirst DP, Térézol M, Cantini L, Villoutreix P, Vignes M, Baudot A. MOTL: enhancing multi-omics matrix factorization with transfer learning [Data set]. Zenodo. 2024. <https://doi.org/10.5281/zenodo.10848217>.

Publisher's Note

Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.