

Learning representations for image-based profiling of perturbations

Authors

Nikita Moshkov, Michael Bornholdt, Santiago Benoit, Matthew Smith, Claire McQuin, Allen Goodman, Rebecca A. Senft, Yu Han, Mehrtash Babadi, Peter Horvath, Beth A. Cimini, Anne E. Carpenter, Shantanu Singh, Juan C. Caicedo*

* Corresponding author

Abstract

Measuring the phenotypic effect of treatments on cells through imaging assays is an efficient and powerful way of studying cell biology, and requires computational methods for transforming images into quantitative data that highlight phenotypic outcomes. Here, we present an optimized strategy for learning representations of treatment effects from high-throughput imaging data, which follows a causal framework for interpreting results and guiding performance improvements. We use weakly supervised learning (WSL) for modeling associations between images and treatments, and show that it encodes both confounding factors and phenotypic features in the learned representation. To facilitate their separation, we constructed a large training dataset with Cell Painting images from five different studies to maximize experimental diversity, following insights from our causal analysis. Training a WSL model with this dataset successfully improves downstream performance, and produces a reusable convolutional network for image-based profiling, which we call *Cell Painting CNN-1*. We conducted a comprehensive evaluation of our strategy on three publicly available Cell Painting datasets, discovering that representations obtained by the Cell Painting CNN-1 can improve performance in downstream analysis for biological matching up to 30% with respect to classical features, while also being more computationally efficient.

Introduction

High-throughput imaging and automated image analysis are powerful tools for studying the inner workings of cells under experimental interventions. In particular, the Cell Painting assay^{1,2} has been adopted both by academic and industrial laboratories to evaluate how perturbations alter overall cell biology. It has been successfully used for studying compound libraries^{3–5}, predicting phenotypic activity^{6–9}, and profiling human disease^{10,11}, among many other applications. To reveal the phenotypic outcome of treatments, image-based profiling transforms microscopy images into rich high-dimensional data using morphological feature extraction¹². Cell Painting datasets with thousands of experimental interventions provide a unique opportunity to use machine learning for obtaining representations of the phenotypic outcomes of treatments.

Improved feature representations of cellular morphology have the potential to increase the sensitivity and robustness of image-based profiling to support a wide range of discovery applications^{13,14}. Feature extraction has been traditionally approached with classical image processing^{15,16}, which is based on manually engineered features that may not capture all the relevant phenotypic variation. Several studies have used convolutional neural networks (CNNs) pre-trained on natural images^{17–19}, which are optimized to capture variation of macroscopic objects instead of images of cells. To recover causal representations of treatment effects, feature representations need to be sensitive to subtle changes in morphology. However, classical features and pre-trained networks may not have sufficient expressive power to realize that potential. Representation learning shows promise as a tool to learn domain-specific features from cellular images in a data-driven way^{4,20–23}, but it also brings unique challenges to avoid being dominated by confounding factors^{24,25}.

In this paper, we investigate the problem of learning representations for image-based profiling with Cell Painting. Our goal is to identify an optimal strategy for learning cellular features, and then use it for training models that recover improved representations of the phenotypic outcomes of treatments. We use a causal framework to reason about the challenges of learning representations of cell morphology (e.g., confounding factors), which naturally fits in the context of perturbation experiments^{26,27}, and serves as a tool to optimize the workflow and yield better performance (Figure 1). In addition, we adopted a quantitative evaluation of the impact of feature representations in a biological downstream task, to guide the search for an optimized workflow. The evaluation is based on querying a reference collection of treatments to find biological matches in a perturbation experiment. In each evaluation, cell morphology features change to compare different strategies, while the rest of the image-based profiling workflow remains constant. Performance is measured using metrics for the quality of a ranked list of results for each query (Figure 1G). With this evaluation framework, we conduct an extensive analysis on three publicly available Cell Painting datasets.

Within the proposed causal framework, we use weakly supervised learning (WSL)²⁰ to model the associations between images and treatments, and we found that it powerfully captures rich cellular features that simultaneously encode confounding factors and phenotypic outcomes as

latent variables. Our analysis indicates that to disentangle them, to improve the ability of models to learn the difference between the two types of variation, and to recover the causal representations of the true outcome of perturbations, it is important to train models with highly diverse data. Therefore, we constructed a training dataset that combines variation from five different experiments to maximize the diversity of treatments and confounding factors (Figure 4). As a result, we successfully trained a new, reusable single-cell feature extraction model: the Cell Painting CNN-1 (Figure 1F), which yields better performance in the evaluated benchmarks and displays sufficient generalization ability to profile other datasets effectively.

Results

Recovering features of treatment effects

We use a causal model as a conceptual framework to reason and analyze the results of representation learning strategies. Figure 1B presents the causal graph with four variables: interventions (treatments **T**), observations (images **O**), outcomes (phenotypes **Y**) and confounders (e.g. batches **C**). Some variables are observables (white circles), while others represent latent variables (shaded circles). This graph is a model of the causal assumptions we make for representation learning and for interpreting the results.

Our goal is to estimate an unbiased, multidimensional representation of treatment outcomes (**Y**), which can later be used in many downstream analysis tasks. We use WSL²⁰ (Figure 1C) for obtaining representations of the phenotypic outcome of treatments by training a classifier to distinguish all treatments from each other. In this way, the model learns associations between observed images (**O**) and treatments (**T**), while capturing unobserved variables in a latent representation (**Y** and **C**). To recover the phenotypic features of treatment effects (**Y**) from the latent representations, we employ batch correction¹⁸ to reduce the variation associated with confounders and amplify causal features of phenotypic outcomes (Figure 1D). More details of the assumptions and structure of our framework can be found in the Methods section.

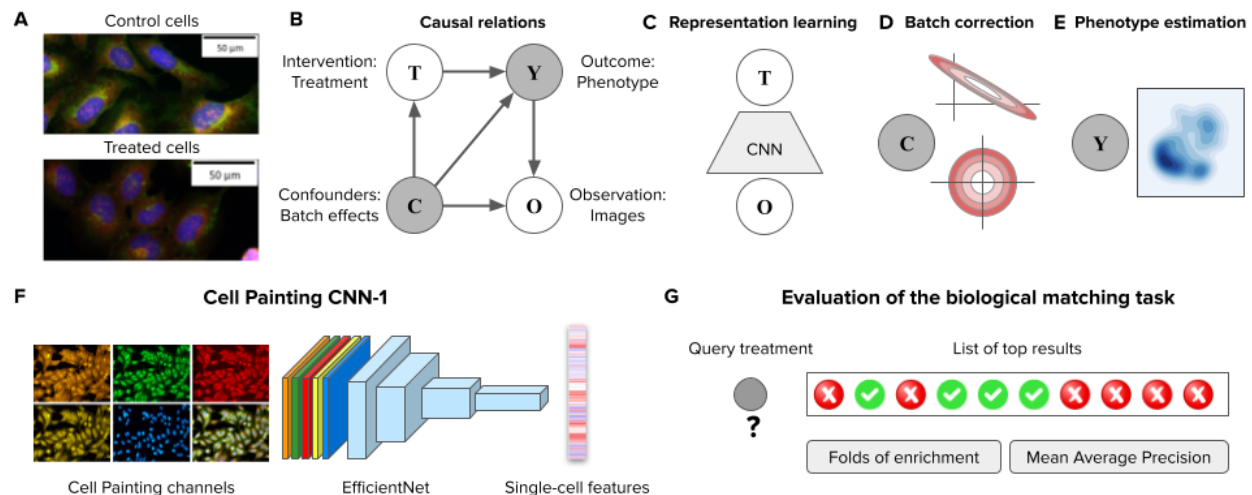


Figure 1. Framework for analyzing image-based profiling experiments. A) Example Cell Painting images from the BBBC037 dataset of control cells (empty status) and one experimental intervention (JUN wild-type overexpression) in the U2OS cell line. B) Causal graph of a conventional high-throughput Cell Painting experiment with two observables in white circles (treatments and images) and two latent variables in shaded circles (phenotypes and batch effects). The arrows indicate the direction of causation. C) Weakly supervised learning as a strategy to model associations between images (O) and treatments (T) using a convolutional neural network (CNN). The CNN captures information about the latent variables C and Y in the causal graph because both are intermediate nodes in the paths connecting images and treatments. D) Illustration of the sphering batch-correction method where control samples are a model of unwanted variation (top). After sphering, the biases of unwanted variation in control samples is reduced (bottom). E) The goal of image-based profiling is to recover the outcome of treatments by estimating a representation of the resulting phenotype, free from unwanted confounding effects. F) Illustration of the Cell Painting CNN-1, an EfficientNet model trained to extract features from single cells. G) The evaluation of performance is based on nearest neighbor queries performed in the space of phenotype representations to match treatments with the same phenotypic outcome. Performance is measured with two metrics: folds of enrichment and mean average precision (Methods).

Weakly supervised learning captures confounders and phenotypic outcomes of treatments

WSL models are trained with a classification loss to detect the treatment in images of single cells (Figure 1C, Methods), which is a pretext task to learn representations of the latent variables in the causal graph. Given that the treatment applied to cells in a well is always known, we can quantify the success rate of single-cell classifiers on this pretext task using precision and recall. We conducted single-cell classification experiments using two validation schemes to reveal how sensitive WSL is to biological and technical variation (Figure 2A). The leave-cells-out validation scheme uses single cells from all plates and treatments in the experiment for training, and leaves a random fraction out for validation. By doing so, trained CNNs have the opportunity to observe the whole distribution of phenotypic features (all treatments) as well as the whole distribution of confounding factors (all batches or plates). In contrast, the leave-plates-out validation scheme separates different technical replicates (plates)

for training and validation, resulting in a model that still observes the whole distribution of treatments, but only partially sees the distribution of confounding factors.

The results in Figure 2B show a stark contrast in performance between the two validation strategies. When leaving cells out (results in blue), a CNN can accurately learn to classify single cells in the training and validation sets with only a minor difference in performance. When leaving plates out (results in orange), the CNN learns to classify the training set well but fails to generalize correctly to the validation set. The generalization ability of the two models is further highlighted in the validation results in Figure 2C, which presents the precision and recall of each treatment.

Importantly, while these WSL models exhibit a major difference in performance in the pretext classification task, their performance in the downstream analysis task is almost the same after batch correction (Figure 2D). The large difference in performance in the pretext task followed by no difference in the downstream task reveals that WSL models leverage both phenotypes and confounders to solve the pretext task. On one hand, the validation results of leaving-cells-out are an overly optimistic estimate of how well a CNN can recognize treatments in single cells, because the models leverage batch effects to make the correct connection. On the other hand, the results of leaving-plates-out are an overly pessimistic estimate because the CNN fails to generalize to unseen replicates with unknown confounding variation (domain shift). The true estimate of performance in the pretext classification task is likely to be in the middle when accounting for confounding factors. This is indeed what we observe in the downstream analysis results: after batch correction, the representations of models trained with leave-cells-out and leave-plates-out yield similar downstream performance in the biological matching task, indicating that both models find similar phenotypic features, but capture different confounding variation.

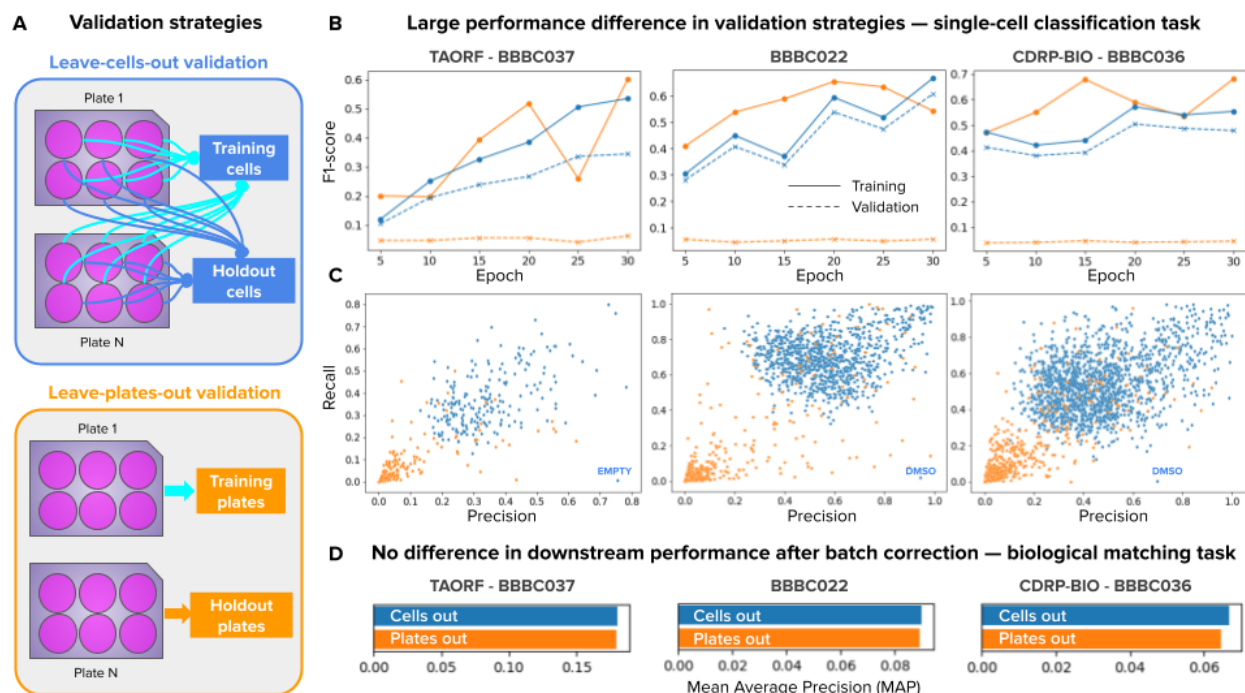


Figure 2. Validation strategies for the single-cell classification task in weakly supervised learning.

A) Illustration of the two strategies: leave-cells-out (in blue) uses cells from all plates in the dataset for training and leaves a random fraction out for validation. Leave-plates-out (in orange) uses all the cells from certain plates for training, and leaves entire plates out for validation. Any difference in performance is due to confounding factors. Note that plates-left out are selected such that all treatments have two full replicate-wells out for validation, which may or may not correspond to entire batches, depending on the experimental design. B) Learning curves of models trained with WSL for 30 epochs with all treatments from each dataset. The x-axis is the number of epochs and the y axis is the average F1-score. The color of lines indicates the validation strategy, and the style of lines indicates training (solid) or validation (dashed) data. C) Precision and recall results of each treatment in the single-cell classification task. Each point is a treatment (negative controls are labeled in blue), and the color corresponds to the validation strategy. D) Performance of models in the downstream, biological matching task after batch correction.

Treatments with strong phenotypic effect improve performance

The WSL model depicted in Figure 1C captures associations between images (O) and treatments (T) in the causal graph, while encoding technical (C) and phenotypic (Y) variation as latent variables because both are valid paths to find correlations. Given that controlling the distribution of confounding factors does not result in downstream performance changes, in this section we explore the impact of controlling the distribution of phenotypic diversity. Our reasoning is that WSL learning favors correlations between treatments and images through the path in the causal graph that makes it easier to minimize the empirical error in the pretext task. Therefore, the variation of treatments with a weak phenotypic response is overpowered by confounding factors that are stronger relative to the phenotype.

To measure the phenotypic strength of treatments we calculate the Euclidean distance between control and treatment profiles in the CellProfiler feature space after batch correction with a sphering transform (Figure 3A). We interpret this procedure as a crude approximation of the average treatment effect (ATE), a causal parameter of intervention outcomes, because the Euclidean distance calculates the difference in expected values (aggregated profiles) of the outcome variable (phenotype) between the control and treated conditions. However, since we do not observe the control and treated condition in the same cells, this remains only an approximation of the ATE, even if the cells are isogenic clones of each other. We chose distances in the CellProfiler feature space as an independent prior for estimating treatment strength because these are non-trainable, and because in our experiments CellProfiler features exhibit more robustness to confounding factors (Supplementary Figure 1).

We ranked treatments in ascending order based on the strength of the phenotypic effect and took 20% in the bottom, middle and top of the distribution (Figure 3B,C). Next, we evaluated the performance of WSL models trained on each of these three categories and we found that performance improves in the downstream biological matching task with treatments that have a stronger phenotypic effect (Figure 3D). These results were obtained by training under a leave-cells-out validation scheme, giving the CNN full access to the distribution of confounders. The trend indicates that it is possible to break the limitations of WSL for capturing useful associations between images and treatments in the latent variables as long as the phenotypic

outcome is stronger than confounding factors. Note that training with all treatments (green points in Figure 3D) may result in lower overall performance if the majority of the treatments have weak phenotypes (BBBC037), may result in marginally improved performance if the confounding effects are too strong (BBBC022), or may result in better performance when there is a balance between both latent variables (BBBC036). Training with all treatments only improves performance with respect to the CellProfiler baseline in one of the three datasets (Figure 3D).

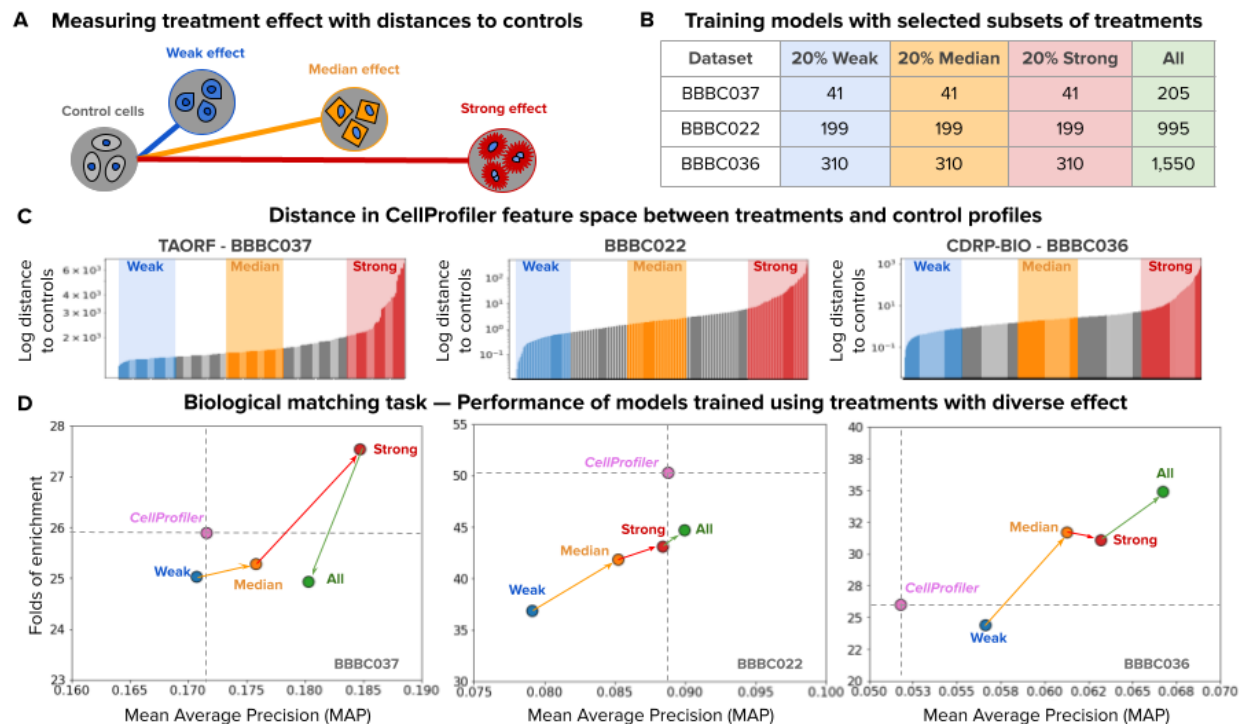


Figure 3. Effect of training models with subsets of treatments. A) Illustration of phenotypic outcomes with varied effects and their distance to controls (see Methods). B) Number of treatments used for training per dataset partition. The rows are datasets and the columns are subsets of treatments grouped by their estimated effect. C) Distribution of distances between treatments and controls as an estimation of treatment effect sorted by distance for each dataset. The x axis represents individual treatments and the y axis represents the log normalized distance to controls. From this distribution, we select 20% of treatments with the weakest (blue), median (orange), and strongest (red) treatments for experiments. D) Evaluation of performance in downstream analysis (biological matching task) for each dataset. Each point in the scatter plots represents one experiment conducted with a model trained with the corresponding subset of the data. The x axis represents performance according to mean average precision (higher is better) and the y axis represents the folds-of-enrichment metric (higher is better).

A training set with highly diverse experimental conditions

In previous sections, we presented the results of changing the distribution of confounders by training with all plates (leave-cells-out validation) or with a few plates (leave-plates-out validation), which did not change the performance in the biological matching task after batch correction (Figure 2D). We also changed the distribution of phenotypic outcomes by sampling

treatments of diverse effects, and observed that strong treatments can improve performance (Figure 3D). In addition, we found that even when training a model with the full distribution of confounders (leave-cells-out) and the full distribution of phenotypes (all treatments) in a dataset, the performance in the downstream task may not necessarily improve with respect to the CellProfiler baseline (Figure 3D green vs pink points). These results suggest that training models in individual datasets alone may limit the representation power of learned features to their local distributions of confounders and phenotypes, and therefore, we hypothesized that increasing their diversity to out of distribution examples may improve performance.

To increase the diversity of experimental conditions in the training set, we created a combined training resource by collecting strong treatments from five different dataset sources, including the three benchmarks evaluated in this work plus two additional publicly available Cell Painting datasets (Figure 4). We selected strong treatments from each source and sought to prioritize treatments shared across sources (Methods). In total, the resulting combined set has 488 strong treatments, represents two cell lines, two types of negative controls, and examples from more than 200 plates, resulting in training data with high experimental diversity with respect to the two latent variables in the causal graph: high technical variation (confounders C) and high phenotypic variation through strong treatments (outcomes Y).

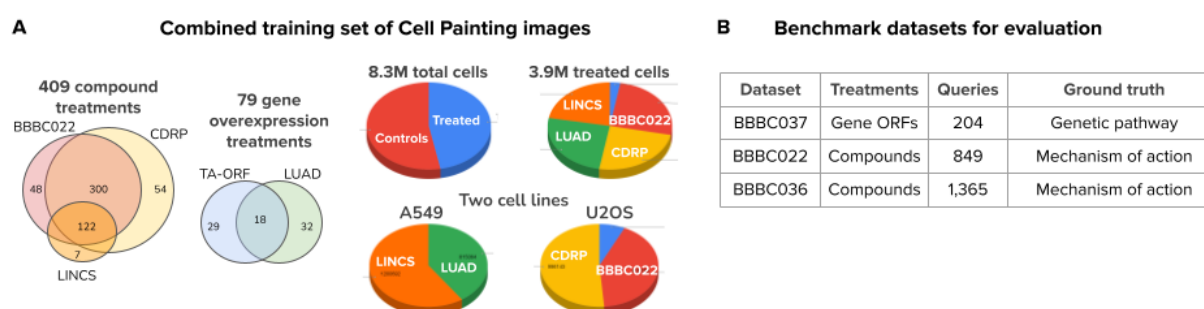


Figure 4. A combined set of Cell Painting images for training. A) Statistics of the combined Cell Painting dataset created to train a generalist model, which brings 488 treatments from 5 different publicly available sources (Methods): LINCS, LUAD, and the three datasets evaluated here; the Venn diagrams illustrate the common treatments among them. There are two types of treatments (compounds and gene overexpression), two types of controls (empty and DMSO), two cell lines (A549 and U2OS), for a total of 8.3 million single cells from 232 plates in the resulting training resource. B) Table of the three benchmark datasets used in this study with the number of treatment queries used for evaluation, and the type of ground truth annotations available for evaluation of the biological matching task. The quantitative results for each dataset are the mean across the queries listed in this table.

Cell Painting CNN-1 learns improved biological features

We found that training a model on this combined Cell Painting dataset consistently improves performance and yields better results than baseline approaches in the task of biologically matching queries against a reference annotated library of treatments, across all three benchmarks (Figure 5A). We consider two baseline strategies in our evaluation: 1) creating image-based profiles with classical features obtained with custom CellProfiler pipelines

(Methods), and 2) computing profiles with a CNN pre-trained on ImageNet, a dataset of natural images in red, green, blue (RGB) colorspace (Methods). Intuitively, we expect feature representations trained on Cell Painting images to perform better at the matching task than the baselines. In the case of CellProfiler, manually engineered features may not capture all the relevant phenotypic variation, and in the case of ImageNet pre-trained networks, they are optimized for macroscopic objects in 3-channel natural images instead of 5-channel fluorescence images of cells.

According to the MAP metric, a WSL model trained on the highly diverse combined set improves performance 7%, 8% and 25% relative to CellProfiler features on BBBC037, BBBC022 and BBBC036 respectively (difference of cyan points vs pink points in the x axis of Figure 2B). Similar improvements are observed with the Folds of Enrichment metric (y-axis of Figure 2B), obtaining 6%, 8%, and 30% relative improvement on the three benchmarks respectively. Importantly, the combined dataset allowed us to train a single model once and profile all the three benchmarks without re-training or fine-tuning on each of them, demonstrating that the model captures features of Cell Painting images relevant to distinguish more effectively the variation of the two latent variables of the causal model.

We found that ImageNet features showed variable performance compared to CellProfiler features (Figure 2B), sometimes yielding similar performance (BBBC022), sometimes lower performance (BBBC037) and sometimes slightly better performance (BBBC036). The three benchmarks used in this study are larger scale and more challenging than datasets used in previous studies^{17,18} where it was observed that ImageNet features are typically more powerful than classical features. Our results indicate that, in large scale perturbation experiments with Cell Painting, ImageNet features do not conclusively capture more cellular-specific variation than manually engineered features using classical image processing.

Figure 5B shows a UMAP projection²⁸ of the feature space obtained using our Cell Painting CNN-1 for the three datasets evaluated in this study. From a qualitative perspective, the UMAP plot of the BBBC037 dataset (a gene overexpression screen) shows groups of treatments clustered according to their corresponding genetic pathway, and recapitulates previous observations of known biology²⁹. The BBBC022³⁰ and BBBC036³¹ datasets (compound screens) likewise show many treatments grouped together according to their mechanism of action (MoA).

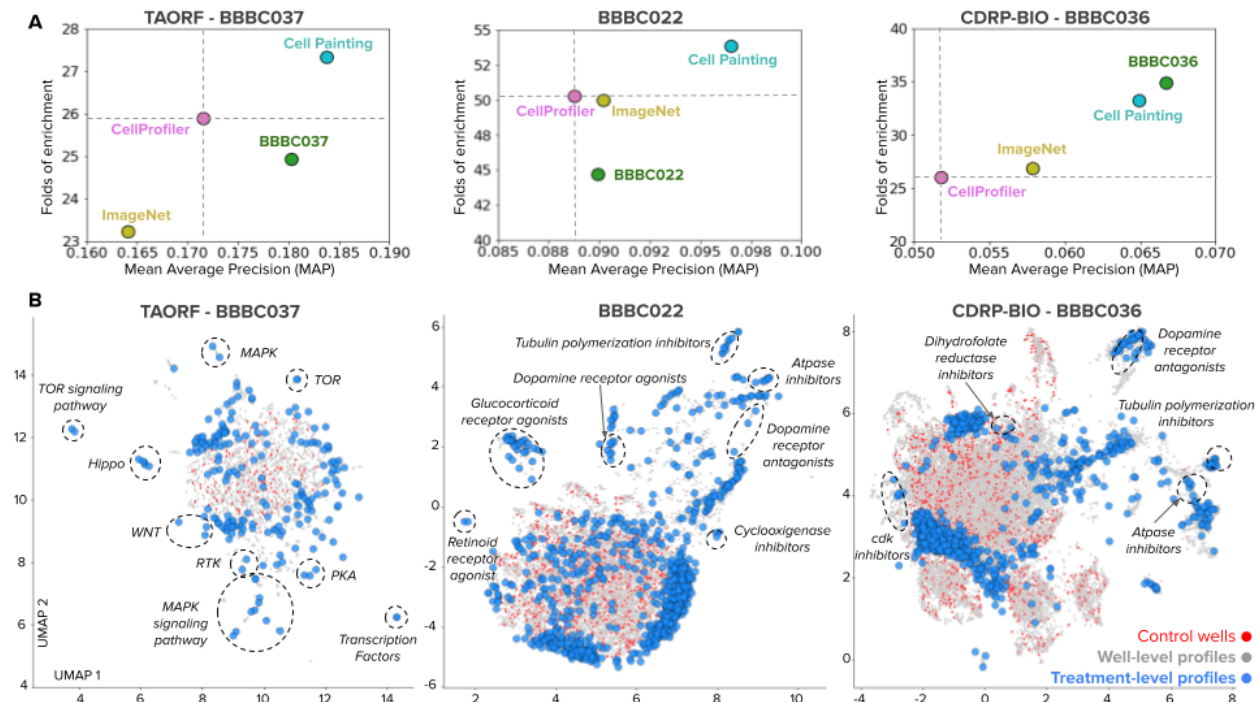


Figure 5. Quantitative evaluation of feature representations of treatment effects. The evaluation task is biological profile matching (see Figure 1G). A) Performance of feature representations for the three benchmark datasets according to two metrics: Mean Average Precision (MAP) in the x axis and Folds of Enrichment in the y axis (see Methods). Each point indicates the mean of these metrics over all queries using the following feature representations: CellProfiler (pink), a CNN pre-trained on ImageNet (yellow), a CNN trained on the combined set of Cell Painting images (cyan), and a CNN trained on Cell Painting images from the same dataset (green). B) UMAP visualizations of treatment profiles recovered with the Cell Painting CNN-1 after batch correction for the three datasets evaluated in this work. The plot includes a projection of well-level profiles (gray points), control wells (red points), and aggregated treatment-level profiles of treatments (blue points). Dashed lines indicate clusters of treatment-level profiles where all or the majority of points share the same biological annotation.

Batch correction recovers phenotype representations

Batch correction is a crucial step for image-based profiling regardless of the feature space of choice. We hypothesized that a rich representation might encode both confounders and phenotypic features in a way that facilitates separating one type of variation from the other, i.e. disentangles the sources of variation. To test this, we evaluated how representations respond to the sphering transform, a linear transformation for batch correction based on singular value decomposition SVD (Methods). Sphering first finds directions of maximal variance in the set of control samples and then reverses their importance by inverting the eigenvalues. This transform makes the assumption that large variation found in controls is associated with confounders and any variation not observed in controls is associated with phenotypes. Thus, sphering can succeed at recovering the phenotypic effects of treatments if the feature space is effective at separating the sources of variation.

We found that all the methods benefit from batch correction with the sphering transform, indicated by the upward trend of all curves from low performance with no batch correction to improved performance as batch correction increases (Figure 6B). Downstream performance in the biological matching task improves by about 50% on average when comparing against raw features without correction. The UMAP plots in Figure 6A show the Cell Painting CNN-1 feature space for well-level profiles before batch correction. When colored by Plate IDs, the data points are fragmented, and the density functions in the two UMAP axes indicate concentration of plate clusters. After sphering, the UMAP plots in Figure 6C show more integrated data points and better aligned density distributions of plates. The performance of the Cell Painting CNN-1 in the biological matching task also improves upon the baselines (Figure 6B) and displays a consistent ability to facilitate batch correction in all the three datasets, unlike the ImageNet CNN.

The sphering transform, while effective, is still far from perfect, and further research is needed to better disentangle confounding from phenotypic variation, potentially using nonlinear transformations.

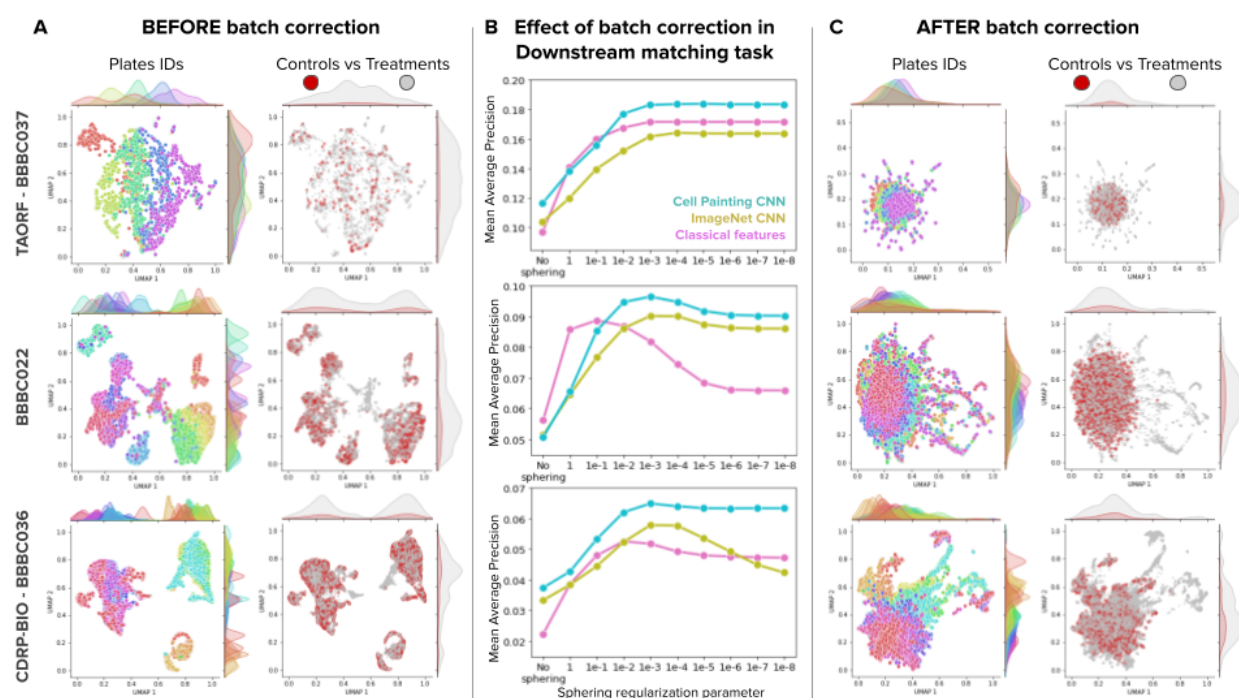


Figure 6. Effect of batch correction on feature representations. Batch correction is based on the sphering transform and applied at the well-level, before treatment-level profiling (Methods). A) UMAP plots of well-level profiles before batch correction for the three benchmark datasets (rows) colored by plate IDs (left column) and by control vs treatment status (right column). The UMAP plots display density functions on the x and y axes for each color group to highlight the spread and clustering patterns of data. B) Effect of batch correction in the biological matching task. The x axis indicates the value of the regularization parameter of the sphering transform (smaller parameter means more regularization), with no correction in the leftmost point and then in decreasing parameter order (increasing sphering effect). The y axis is Mean Average Precision in the biological matching task. C) UMAP plots of well-level profiles after batch correction for the three benchmark datasets with the same color organization as in A.

Discussion

This paper presented an optimized methodology for learning representations of phenotypes in imaging experiments, which uses weakly supervised learning, batch correction with sphering, and an evaluation framework to assess performance. We used this methodology to analyze three publicly available Cell Painting datasets with thousands of perturbations, and we found that: 1) CNNs capture confounding and phenotypic variation as latent variables, 2) the performance of CNNs can be improved by training with datasets that maximize technical and biological diversity, and 3) batch correction is necessary to recover a representation of the phenotypic outcome of treatments. The WSL approach in our methodology aims to learn *unbiased* features of cellular morphology that can be used to approach various problems and applications in cell biology. This is in contrast to supervised strategies that aim to solve one task with high accuracy, and therefore only capture features relevant to that task. Our approach can also be generalized to other imaging assays or screens, and we anticipate that the same principles will be useful for improving performance in downstream tasks.

We optimized a Cell Painting CNN-1 model that can extract single-cell features to create image-based profiles for estimating the phenotypic outcome of treatments in perturbation experiments. Following insights derived from the analysis with our methodology, we constructed a large training resource by combining five sources of Cell Painting data to maximize phenotypic and technical variation for training a reusable feature extraction model. This model successfully improved performance in all three benchmarks while also being computationally efficient (Supplementary Figure 4). The fact that the best-performing strategy involved training a single model once and profiling all the three benchmarks without re-training or fine-tuning has a remarkable implication: it indicates that generating large experimental datasets with a diversity of phenotypic impacts could be used to create a single model for the community that could be transformational in the same way that models trained on ImageNet have enabled transfer learning on natural image tasks.

We used a causal conceptual framework for analyzing the results, which we found very useful to interpret performance differences between feature extraction models. In practice, the framework was useful for guiding decision making while training new models, and it is helpful to understand and communicate the challenges of learning representations in imaging experiments. In theory, it also opens new possibilities to formulate the problem in novel ways, for instance, creating learning models that account for all four variables simultaneously. This framework is a compact way to express the causal assumptions of the underlying biological experiment, which is consistent with the experimental evidence that we observed throughout this study. We believe that this is a first step towards studying the causal relationships between disease and treatments using high-throughput imaging experiments and modern machine learning.

There are many sources of confounding factors in biological experiments, and microscopy imaging is not exempt. While batch effects have been widely studied for other data modalities (e.g. gene expression), they remain largely unexplored in imaging studies, with a few exceptions. Imaging is a powerful technology for observing cellular states, and it is sensitive to

phenotypes as well as unwanted variation. If left unaccounted for, unwanted variation can result in biased models that confound biological conclusions, and this is especially true for large capacity, deep learning models. Our experiments show that deep learning can exploit confounding factors to minimize training error, and that batch effect correction is critical to recover the biological representation of interest in conventional or deep-learning based features. More research is necessary to improve the efficiency and generalizability of batch correction methods for imaging studies.

Deep learning for high-throughput imaging promises to realize the potential of perturbation studies for decoding and understanding the phenotypic effects of treatments. The upcoming public release of the JUMP-Cell Painting Consortium's dataset of more than 100,000 chemical and genetic perturbations, collected across 12 different laboratories in academia and industry, is an excellent example of the scale and biological diversity that imaging can bring for drug discovery and functional genomics research³². Our Cell Painting CNN-1 is the first publicly available model optimized for phenotypic feature extraction in image-based profiling studies, and can generalize to new data with new perturbations. We expect that our methodology, together with larger datasets, will be useful to create new and better models for analyzing images of cells in the future.

Acknowledgements

We thank Salil Bhate for the valuable discussions and feedback provided throughout the development of this project. NM acknowledges the short-term scientific mission grants provided by eCOST action CA15124 (NEUBIAS) in 2019 and 2020. NM and PH acknowledge support from the LENDULET-BIOMAG Grant (2018-342), from the European Regional Development Funds (GINOP-2.2.1-15-2017-00072), from the H2020 and EU-Horizont (ERAPERMED-COMPASS, ERAPERMED-SYMMETRY, DiscovAIR, FAIR-CHARM) from TKP2021-EGA09, from the ELKH-Excellence grants and from the Cooperative Doctoral Programme (2020-2021) of the Ministry for Innovation and Technology. Researchers in the Carpenter–Singh lab were supported by NIH R35 GM122547 to AEC. Researchers in the Cimini lab were supported by NIH P41 GM135019. AG was supported by grant number 2018-192059 and BAC was additionally supported by grant number 2020-225720 from the Chan Zuckerberg Initiative DAF, an advised fund of the Silicon Valley Community Foundation. JCC was supported by the Schmidt Fellowship program of the Broad Institute and by the NSF DBI Award 2134695.

Methods

1. Causal framework

The causal graph in our framework includes four variables: interventions (treatments **T**), observations (images **O**), outcomes (phenotypes **Y**) and confounders (e.g. batches **C**). For simplicity, we assume that there is a single context (**X**, not in the diagram) for experimental treatments, which are clonal cells of an isogenic cell line; perturbation experiments with multiple cell lines may need a different model. We assume that images and treatments are observables (**O** and **T**) because the images are acquired as a result of the experiment, and because the treatments are chosen by researchers. We also assume that the phenotype and confounders are latent variables (**Y** and **C**) that we want to estimate and separate.

The relationships in the causal graph are interpreted as follows: the arrow from **T** to **Y** indicates that treatments are applied to cells and are the main direct biological cause of phenotypic changes in cells in the perturbation experiment. The arrow from **Y** to **O** indicates that we partially observe the phenotypic outcome through images. This observation is assumed to be noisy and incomplete, requiring hundreds of cells and multiple replicates to increase the chances of measuring the real effect of treatments. In addition, the image acquisition process and the overall experiment are influenced by technical variation. The arrow from **C** to **O** indicates that images are impacted by artifacts in image acquisition, including microscope settings and assay preparation. The arrow from **C** to **Y** indicates that phenotypes are impacted by variations in cell density and other conditions that make cells grow and respond differently. The arrow from **C** to **T** indicates that treatments are impacted by plate map designs that are not fully randomized and usually group treatments in specific plate positions.

We observe treatment outcomes (**Y**) indirectly through imaging assays, and thus, we need image analysis to recover the phenotypic effect and to separate it from unwanted variation (**C**). A representation of the phenotypic effect can be obtained with the workflow depicted in Figures 1C-E, which illustrates three major steps: 1) modeling the correlations between images and treatments using a CNN trained with weakly supervised learning (WSL), 2) using batch correction to learn a transformation of the latent representation of images obtained from intermediate layers of the CNN, and 3) generating representations of treatment effects in cellular morphology for downstream analysis.

2. Weakly Supervised Learning

Weakly supervised learning²⁰ (WSL) trains models with the auxiliary task of learning to recognize the treatment applied to single cells. Treatments are always known in a perturbation experiment, while other biological annotations, such as mechanism of action or genetic pathway may not be known for certain treatments, only reflect part of the phenotypic outcome, and is usually unknown at the single-cell level (which is the resolution used for training models in this

work). We use an EfficientNet³³ architecture with a classification loss with respect to treatment labels for training WSL models.

WSL captures the correlations between observed images and treatments and makes the following assumptions: 1) if a treatment has an observable effect then it can be seen in images, and therefore, training a CNN helps identify visual features that make it detectably different from all other treatments. 2) Treatment labels in the classification task are weak labels because they are not the final outcome of interest, they do not reflect expert biological ground truth, and there is no certainty that all treatments produce a phenotypic outcome, nor that they produce a different phenotypic outcome from each other. 3) Cells might not respond uniformly to particular treatments³⁴, which yields heterogeneous subpopulations of cells that may not be consistent with the treatment label, i.e, treatment labels do not have single-cell resolution. 4) Intermediate layers of the CNN trained with treatment labels capture all visual variation of images as latent variables, including confounders and causal phenotypic features.

3. Deep learning models

3.1 Image preprocessing

The original Cell Painting images in all the datasets used in this work are encoded and stored in 16-bit TIFF format. To facilitate image loading from disk to memory during training of deep learning models, we used image compression. This is only required for training, which requires repeated randomized loading of images for minibatch-based optimization.

The compression procedure is as follows:

- Compute one illumination correction function for each channel-plate³⁵. The illumination correction function is computed at 25% of the width/height of the original images.
- Apply the illumination correction function to images before any of the following compression steps.
- Stretch the histogram of intensities of each image by removing pixels that are too dark or too bright (below 0.05 and above 99.95 percentiles). This expands the bin ranges when changing pixel depth and prevents having dark images as a result of saturated pixels.
- Change the pixel depth from 16 bits to 8 bits. This results in 2X compression.
- Save the resulting image in PNG lossless format. This results in approximately 3X compression.
- The resulting compression factor is approximately $(2X)(3X) = 6X$.

This preprocessing pipeline is implemented in DeepProfiler and can be run with a metadata file that lists the images in the dataset that require compression, together with plate, well and site (image or field of view) information.

3.2 EfficientNet

The deep convolutional neural network architecture used in all our experiments is the EfficientNet³³. We use the base model EfficientNet-B0 to compute features on single-cell crops of 128x128 pixels. It consists of 9 stages: input, seven inverted residual convolutional blocks from MobileNetV2³⁶ (with the addition of squeeze and excitation optimization) and final layers. The usage of convolutional blocks from MobileNetV2 in combination with neural architecture search gave EfficientNet an advantage in terms of computational efficiency and accuracy compared to ResNet50. This model has only 4 million parameters and can extract features from single cells in a few milliseconds using GPU acceleration. EfficientNet has been previously used for image-based profiling, including in models trained with the CytolImageNet dataset³⁷, by top competitors in the Recursion Cellular Image Classification challenge in Kaggle, and for studying variants of unknown significance in cancer lung cells¹¹.

3.3 Training Cell Painting models

The Cell Painting models are trained on single-cell crops obtained from full images using cell locations and full-image metadata. Cropped single cells are exported to individual images together with their segmentation mask if available. In all our experiments, we used single cells cropped from a region of 128x128 pixels centered on the cell's nucleus without any resizing. We preserve the context of the single cell (background or parts of other cells), meaning that the segmentation mask is not used in our experiments.

To train a weakly supervised model we first initialize an EfficientNet with ImageNet weights and sample mini-batches of 32 examples for training with an SGD optimizer with a learning rate of 0.005. We train models for 30 epochs; each epoch makes a pass over example cells of all treatments in a balanced way. Balancing is set to draw the same number of single cells from each treatment (the median across treatments) in one epoch, and every epoch resamples new cells from the pool. This strategy leverages the variation of cells in treatments that are overrepresented (such as controls), and oversamples cells from treatments with fewer than the median. Balancing is important to optimize the categorical cross-entropy loss to compensate for rare classes among the hundreds of treatments used for training in our experiments.

All single cells go through a data augmentation process during training, which involves the following three steps:

1. Random crop and resize. This augmentation is applied with 0.5 probability. The crop region size is random, 80% to 100% of the size of the original image, then resized back to 128x128 pixels.
2. Random horizontal flips and 90-degree rotations.
3. Random brightness and contrast adjustments, each channel is augmented and renormalized separately.

We used two data-split approaches for creating training and validation subsets for the single-cell treatment classification task:

- Leave plates out: we selected the plates in a way that the data of one subset of plates is only the train data and another one is in validation.
- Leave cells out: all plates are used for training and validation. We randomly choose single cells from each well, meaning that approximately 60% of cells from each well would be in the training set and the rest in the validation set.

The training of deep learning models was performed on NVIDIA DGX with NVIDIA V100 GPUs; a single GPU was used to train each model.

3.4 Feature extraction with trained Cell Painting models

We extract features of single cells and store them in one NumPy file per field of view using an array of vectors (one per cell). The feature extraction procedure requires access to full images, metadata and location files. Since a trained model is natively trained for five-channel images, there is no need to replicate each channel separately, thus, each cell requires one inference pass and the feature vector contains representation of all channels simultaneously. The size of the feature vectors is 672 for reported results, which were extracted from EfficientNet's block6a-activation layer.

3.5 ImageNet pre-trained models

ImageNet pre-trained convolutional networks are widely used in computer vision applications and they have also been used in image-based profiling applications. ImageNet is a large collection of natural images with objects and animals of different categories³⁸. A deep learning model trained on this dataset is capable of extracting generic visual features from images for different applications, also known as transfer learning. Pre-trained models for morphological profiling have been evaluated in several studies^{11,17–19}.

We use DeepProfiler to extract features of single-cells with the pre-trained EfficientNet-B0 model available in the Keras library. The size of single-cell images in our experiments is 128x128 after being cropped from field-of-view images. An ImageNet pre-trained model expects images of higher resolution, specifically 224x224 in our case; therefore the cell crops are first resized. The pixel values are then rescaled using min-max normalization and adjusted to have values [-1:1] to match the required input range. As Cell Painting images are five-channel and ImageNet pre-trained models expect three-channel (RGB) images, we follow the well established practice of computing a pseudo RGB image for each grayscale fluorescent channel by replicating it three times before passing it through the model (thus each cell requires five inference passes). Features extracted for each channel are concatenated and the resulting feature vector size is 3,360 (672 is the size of the block6a-activation layer of the EfficientNet used in our experiments).

3. Profiling workflow

3.1 Segmentation

The cell segmentation for the benchmark datasets (BBBC037, BBBC022 and BBBC036) was performed with methods built in CellProfiler v2 based on Otsu thresholding³⁹ and propagation method⁴⁰ based on Voronoi diagrams⁴¹ or watershed from⁴². The segmentation is two-stepped: first, the images stained with Hoechst (DNA channel) were segmented using global Otsu thresholding. This prior information is then used in the second step: cell segmentation with the propagation or watershed method. The input channel for the second step depends on the dataset, as well as the other specific parameters of segmentation. The segmentation part of the pipelines is available in the published CellProfiler pipelines (see Code availability section). For the purposes of this project, we used the center of the nuclei to crop out cells in a region of 128x128 pixels. These cell crops are used in all the deep learning workflows.

3.2 Feature extraction with CellProfiler

Feature extraction for evaluated datasets was performed with CellProfiler 2. The feature extraction steps are described in the CellProfiler pipelines published together with the corresponding original datasets. These steps can be grouped in the following stages: 1) Data loading - load full image 2) Illumination correction for each channel 3) Identification of cell nuclei 4) Identification of cells using identified nuclei 4) Measurements: intensity, context, radial distribution, size and shape, texture 5) Export the features and cell outlines. Parameters of feature extraction can be found in CellProfiler pipelines which are available in the published pipelines (see Code availability section).

The features of CellProfiler are designed to be human-readable and grouped into three large groups: “Cell”, “Cytoplasm” and “Nuclei”. Each of those feature groups has several common subgroups, such as shape features, intensity-based features, texture features and context features¹². The resulting size of a feature vector is approximately 1,800 (which depends on the dataset). In our analysis, we used well-level CellProfiler features to obtain baseline results. These features were computed by the authors of the corresponding dataset and made publicly available (see Data availability section).

3.3 Feature aggregation and profiling

Image-based profiling aims to create representations of treatment effects, which is obtained by aggregating information of single cells into population-level profiles. This process follows a multi-step aggregation process. Features of single-cells are first aggregated using the median operator at field-of-view (image) level. Next, fields-of-view features are aggregated using the mean to create a well-level profile. Finally, treatment-level profiles are obtained with the average across replicate wells. The feature aggregation steps are the same for CellProfiler and deep learning features. CellProfiler well-level features with NA values were removed in the aggregation pipeline.

3.4 Batch correction with the sphering transform

To recover the phenotypic features of treatments from the latent representations of the weakly supervised CNN, we employ a batch correction model inspired by the Typical Variation Normalization (TVN) technique¹⁸. This transform aims to reduce the variation associated with confounders and amplify features caused by phenotypic outcomes (Figure 1D). The main idea of this approach is to use negative control samples as a model of unwanted variation under the assumption that their phenotypic features should be neutral, and therefore differences in control images reflect mainly confounding factors. We follow this assumption and use a sphering transformation to learn a function that projects latent features from the CNN to a corrected feature space that preserves the phenotypic features caused by treatments. We note that given how control wells are placed in plates, they may not represent all of the unwanted variation caused by plate layout effects, nevertheless, we assume it is a sufficient approximation.

In our implementation, we aim to reduce the profiles of control wells to a white noise distribution using a sphering transform, and then use the resulting transformation as a correction function for treated wells. First, the orthogonal directions of maximal variance are identified using singular value decomposition (SVD) on the matrix of control wells. Then, directions with small variation are amplified while directions with large variation are reduced by inverting their eigenvalues. We control the strength of signal amplification or reduction with a regularization parameter. The computation involves only profiles of negative controls and as a result, we obtain a linear transformation that can be applied to all well-level feature vectors in a dataset.

The sphering transformation takes n well-level profiles of negative controls with vector size d as an input matrix $X^{n \times d}$. Then, its covariance matrix is calculated as $\Sigma = \frac{X^T X}{n}$ followed by eigendecomposition $\Sigma = U \Delta U^T$, where Δ is the diagonal matrix of eigenvalues, and U is the matrix of orthonormal vectors. We renormalize the orthonormal vectors by inverting the square root of the eigenvalues in Δ together with a regularization parameter λ . The resulting ZCA-transformation⁴³ matrix is $Q = U (\Delta + \lambda)^{-\frac{1}{2}} U^T$, which can be used to compute the corrected profile of a treated well t with a matrix multiplication: $t' = Qt$. The effect of sphering and its regularization on representations and profiling performance is presented in Figure 4.

3.5 Similarity matching

To assess the similarity between treatment profiles the cosine similarity is measured between pairs of treatments. The cosine similarity is one of several similarity metrics that can be used in profiling¹² and has been used in previous studies⁴⁴.

$$\text{cosine similarity} = \frac{A \cdot B}{\|A\| \|B\|}$$

where A and B are image-based profiles, i.e., multidimensional vectors. We adopt the cosine similarity in all our similarity search and biological matching experiments.

4. Evaluation metrics

For quantitative comparison of multiple feature extraction strategies, we simulate a user searching a reference library to find a “match” to their query treatment of interest. We used a leave-one-treatment-out strategy for all annotated treatments in three benchmark datasets, following previous research in the field^{18,44,45}. In all cases, queries and reference items are aggregated treatment-level profiles matched using the cosine similarity (Methods). The result of searching the library with one treatment query is a ranked list of treatments in descending order of relevance. A result in the ranked list is considered a positive hit if it shares at least one biological annotation in common with the query; otherwise it is a negative result (Figure 1H).

There are several quantitative evaluations of feature representation quality that we use in our study. At the single-cell level, we expect neural networks to classify single cells into their corresponding treatment, and therefore use accuracy, precision and recall to evaluate performance (see Figure 3 and main text). For downstream analysis we adopted a biological matching task, which simulates a user searching for treatments that correspond to the same mechanism of action or genetic pathway (for compound and gene overexpression perturbations respectively). These queries are conducted and evaluated at the treatment-level, and the main idea is to assess how well connected treatments are in the feature space according to known biology.

We use two main metrics for evaluating the quality of the results for a given query: 1) folds of enrichment and 2) mean average precision (mAP). The folds-of-enrichment metric (see details below) is inspired by statistical analyses in biology and determines how unusual positive connections happen to be in the top 1% of the list⁴⁵. On the other hand, the mAP metric is inspired by information retrieval research, and quantifies the precision and recall trend over the entire list of results for all queries.

In order to simulate queries, we proceed as follows:

- Choose a query treatment - which belongs to an MoA or pathway that has at least two treatments in the database and, therefore, it is possible to find a match.
- Library treatments - all the others while leaving the query treatment out. Library treatments represent a database of treatments with known MoAs or pathways annotations, which can be candidate matches for a given query.

4.1 Folds of enrichment

For each query treatment we calculate the odds ratio of a one-sided Fisher’s exact test. The test is calculated using a 2x2 contingency table: the first row contains a number of treatments with the same MoAs or pathways (positive matches) and different MoAs or pathways (negative matches) at a selected threshold of the list of results. The second row is the same, but for the treatments below the threshold (the rest). Odds ratio is a sum of the first row divided by the sum of the second row. It estimates the likelihood of observing the treatment with the same MoA or pathway in the top connections.

We calculate the odds ratio of each individual query, and then obtain the average over all query treatments. The threshold we use is 1% of connections, meaning we expect the top 1% of matching results in the list to be significantly enriched for positive matches. This metric in the text is referred to as “Folds of Enrichment”. The implementation of the metric is available as a part of analysis pipelines (see Code availability section).

4.2 Mean Average Precision

For each query treatment, average precision (area under precision-recall curve) is computed following the common practice in information retrieval tasks. The evaluation starts from the most similar treatments to the query (top results) and continues until all positive pairs (response treatments with the same MoA or pathway) are found. Precision and recall are computed at each item of the result list.

$$Precision = \frac{TP}{TP + FP} \quad Recall = \frac{TP}{TP + FN}$$

where TP are the true positives, FP are the false positives, and FN are the false negatives in the list of results until the current item. This is evaluated for each query separately. As the number of treatments per MoA or pathway is not balanced, the precision-recall curve has a different number of recall points. Therefore, precision and recall are interpolated for each query to cover the maximum number of recall points possible in the dataset, and thus allow for averaging at the same recall points. The interpolated precision at each recall point is defined as follows⁴⁶:

$$p_{inter}(r) = \max_{r' \geq r} p(r')$$

Average precision for a query treatment is the mean of p_{inter} at all recall points. The reported mean average precision (mAP) is the mean average precision over all queries.

5. Datasets

5.1 Benchmarks and ground truth annotations

For this study, we used three publicly available Cell Painting datasets representing gene overexpression perturbations (BBBC037⁴⁵) and compound perturbations (BBBC022³⁰ and BBBC036³¹). The three datasets were produced at the Broad Institute using the U2OS cell-line (bone cancer) following the standardized Cell Painting protocol¹, which stains cells with six fluorescent dyes and acquires imaging samples in five channels at 20X magnification. The compound perturbation experiments used DMSO as a negative control treatment, while in the gene overexpression experiments no perturbation was used for negative control samples. All experiments were conducted using multiple 384-well plates at high-throughput with 5 replicates per treatment (except for high-replicate positive and negative controls).

We conducted quality control of images in all the three datasets by analyzing image-based features with principal component analysis. The outliers observed in the first two principal

components were flagged as candidates for exclusion, and were visually inspected to confirm rejection. We found most of these images to be noisy or empty and not suitable for training and evaluation. With this quality control, two wells were removed from BBBC037, 43 wells from BBBC022, and no wells were removed from BBBC036. If treatments had multiple concentrations in BBBC022 and BBBC036, we kept only the maximum concentration for further analysis and evaluation.

The ground-truth annotations for compounds correspond to mechanism-of-action (MoA) labels when known, and can include multiple annotations per compound. In the case of gene overexpression, the ground-truth corresponds to the genetic pathway of perturbed genes. We used annotations collected in prior work for the same three datasets ⁴⁵ and applied minor updates and corrections (see Data Availability section). Only treatments with at least two replicates left after quality control are included in the ground-truth.

5.2 Measuring treatment effect

The effect of treatments is approximated by computing the distance between the morphological features of treatments and controls. We use batch-corrected features obtained with CellProfiler (with sphering regularization parameter 1e-2) in the following way:

1. Calculate the median profile of control wells within the same plate (median control profile of the plate).
2. For each treated well, calculate the Euclidean distance between its well-level profile and the median control profile of its plate.
3. Estimate the distribution of control well distances against the median control profile per plate. Then, calculate their mean and standard deviation.
4. Using the statistics of control distances per plate, Z-score the distances of treated wells obtained in step 2.
5. Finally, we define the approximate measure of the effect of a given treatment as the average of the Z-scores of its well replicates across plates.

Intuitively, we expect treatments with stronger effects to have a high average Z-score while treatments with weaker or no detectable effect are expected to have low average Z-score. We use this measure to rank treatments and select subsets of treatments for evaluation of the impact of treatment effect during training, as well as for sampling treatments with high phenotypic effect for creating a combined training dataset.

5.3 Combined Cell Painting dataset

We combined five publicly available Cell Painting datasets to create a training resource that maximizes both phenotypic and technical variation. The five dataset sources include the three benchmarks described above (BBBC037, BBBC022, and BBBC036), as well as two additional datasets: 1) BBBC043 ¹¹, a gene overexpression experiment to study the impact of cancer variants, and 2) LINCS ⁵, a chemical screen of FDA approved compounds for drug repurposing research. Both BBBC043 and LINCS are perturbation experiments conducted with A549 cells

(lung adenocarcinoma). In total, these five sources of data have more than 6,000 treatments, in hundreds of plates, thousands of wells, and millions of images resulting in the order of hundreds of millions of imaged single cells. Our goal was to select a sample of single cells from these five sources to maximally capture phenotypic and technical variation.

Instead of sampling single cells uniformly at random, we follow the distribution of treatments to include biological diversity, and the organization of the experimental design to represent various sources of technical noise. Technical variation is organized hierarchically in experiments, starting with the five sources of Cell Painting images, continuing with batches, plate-maps, plates, and well positions. We aimed to bring samples from as many of these combinations as possible to have cells representing different types of technical variation. In terms of biological variation, three of the five data sources have U2OS cells and the other two have A549 cells, resulting in two different cellular contexts being represented. The five sources also include two types of perturbations (chemical and genetic), and multiple treatments.

To preserve as much phenotypic variation as possible, we sample treatments from both cell lines, both types of perturbations, and we identify the treatments with strongest effect in each of the five sources following the procedure described in the previous section. Several treatments overlap across data sources, and we prioritized those that can be found in two or more sources simultaneously. An example is negative controls: all compound screens use DMSO as the negative control, and we would expect their phenotype to match across data sources. The same expectation holds for the rest of treatments. Negative control wells are typically present in each plate of the experiment in several replicates.

The selection of strong treatments started with the BBBC022 and BBBC036 datasets (chemical perturbations). We selected the 500 strongest treatments (see Measuring treatment effect) from BBBC022 and searched for those in BBBC036, which resulted in 301 strong treatments in common between both datasets. We additionally selected 50 unique treatments from BBBC022 and 62 unique treatments from BBBC036. Out of those 413 treatments, 122 overlapped with the LINCS dataset and were included. We additionally selected 7 random treatments from LINCS, from top 20 (by number of associated treatments) MoAs. Treatment selection from BBBC037 and BBBC043 (gene overexpression perturbations) was similar, and we identified 28 overlapping genes. We assume that “wildtype” genes from both datasets are the same, and then we selected the 29 strongest unique perturbations from the BBBC037 dataset and the 32 strongest perturbations from BBBC043 from non-overlapping subsets.

Negative controls from compound screening datasets and negative controls from gene overexpression datasets are considered as different classes in the combined dataset (DMSO and EMPTY). Not all control wells were included from the LINCS dataset in the final sample, as these would result in extreme overrepresentation, so we randomly sampled three control wells per plate. As the final step, the treatments with less than 100 cells were filtered out. In total the dataset contains 490 classes (488 for treatments and 2 for negative controls), 8,423,455 individual single-cells (47% treatment and 53% control cells). See Venn diagrams in Figure 1C for more details.

Data availability

The Cell Painting datasets (raw images and CellProfiler profiles) are available at public S3 buckets:

BBBC037 gene overexpression dataset in U2OS cells ²⁹

s3://cytodata/datasets/TA-ORF-BBBC037-Rohban/profiles_cp/TA-ORF-BBBC037-Rohban/

BBBC022 compound screening in U2OS cells ³⁰

s3://cytodata/datasets/Bioactives-BBBC022-Gustafsdottir/profiles/Bioactives-BBBC022-Gustafsdottir/

BBBC036 compound screening in U2OS cells ³¹

s3://cytodata/datasets/CDRPBIO-BBBC036-Bray/profiles_cp/CDRPBIO-BBBC036-Bray/

BBBC043 gene overexpression dataset in A549 cells ¹¹

s3://cytodata/datasets/LUAD-BBBC043-Caicedo/profiles_cp/LUAD-BBBC043-Caicedo/

LINCS compound screening in A549 cells ⁵

s3://cellpainting-gallery/cpg0004-lincs/broad/images/2016_04_01_a549_48hr_batch1/

Code availability

A. DeepProfiler

To run all the experiments in this study, we developed DeepProfiler, a tool for learning and extracting representations from high-throughput microscopy images using convolutional neural networks (CNNs). DeepProfiler uses a standardized workflow that includes image pre-processing, training of CNNs and feature extraction, as discussed in previous sections. DeepProfiler is implemented in Tensorflow ⁴⁷ (version 2) and is publicly available on GitHub <https://github.com/cytomining/DeepProfiler>.

The DeepProfiler documentation (<https://cytomining.github.io/DeepProfiler-handbook/>) describes the steps for installing, configuring and running the software for profiling new images and for training models. In DeepProfiler we used the following EfficientNet implementation: <https://github.com/qubvel/efficientnet>.

B. DeepProfiler experiments and configurations

The processing and profiling pipelines for the three benchmarks evaluated in this work (Jupyter notebooks and Python scripts to analyze features) are available on GitHub: <https://github.com/broadinstitute/DeepProfilerExperiments>.

This repository also includes the DeepProfiler configuration files used for training the models on each dataset, as well as [the configuration for training the Cell Painting CNN-1 model](#). In addition, the ground truth files and code for evaluation of the downstream tasks are also available in this repository.

C. Models

The Cell Painting CNN-1 model (trained with leave-cells-out training-validation split) is available on Zenodo: <https://doi.org/10.5281/zenodo.7114557>.

The ImageNet pre-trained EfficientNet model used in this study can be found here:

https://github.com/Callidior/keras-applications/releases/download/efficientnet/efficientnet-b0_weights_tf_dim_ordering_tf_kernels_autoaugment.h5.

D. Others

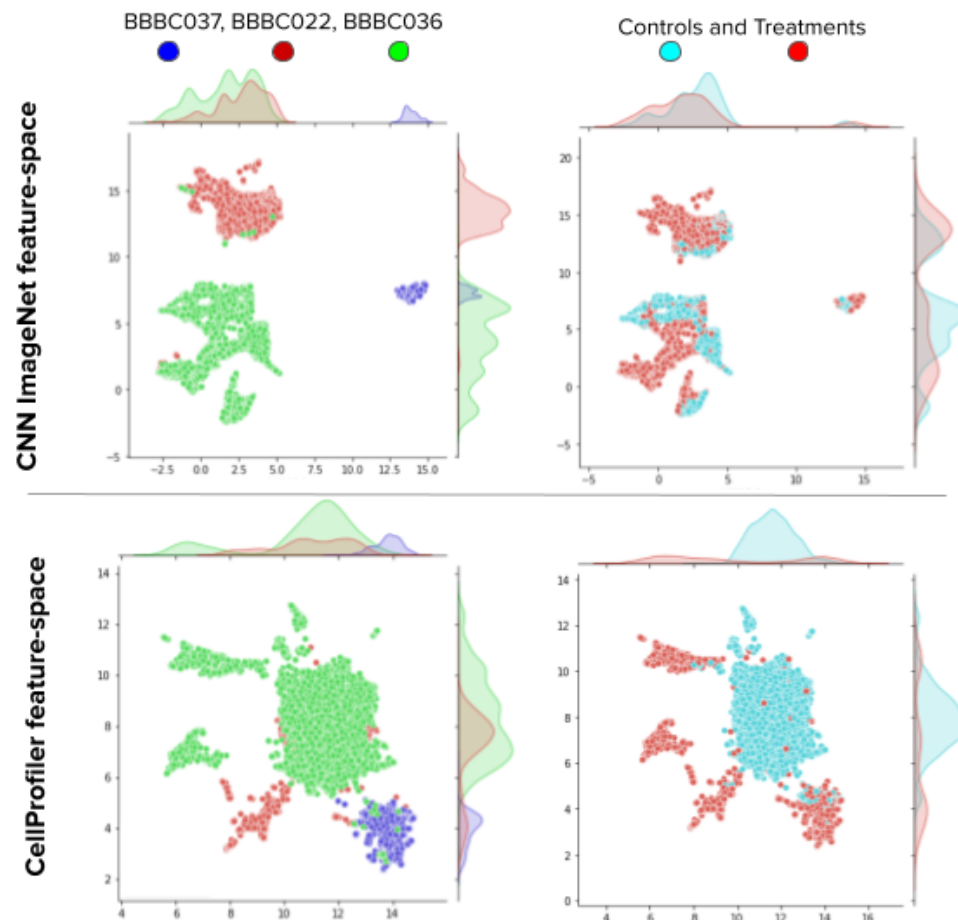
The code and CellProfiler pipelines for three evaluated datasets can be found in the associated GitHub repositories:

BBBC037: https://github.com/carpenterlab/2017_rohban_elife

BBBC036: <https://github.com/gigascience/paper-bray2017>

BBBC022: Supplementary materials in ³⁰.

Supplementary Material

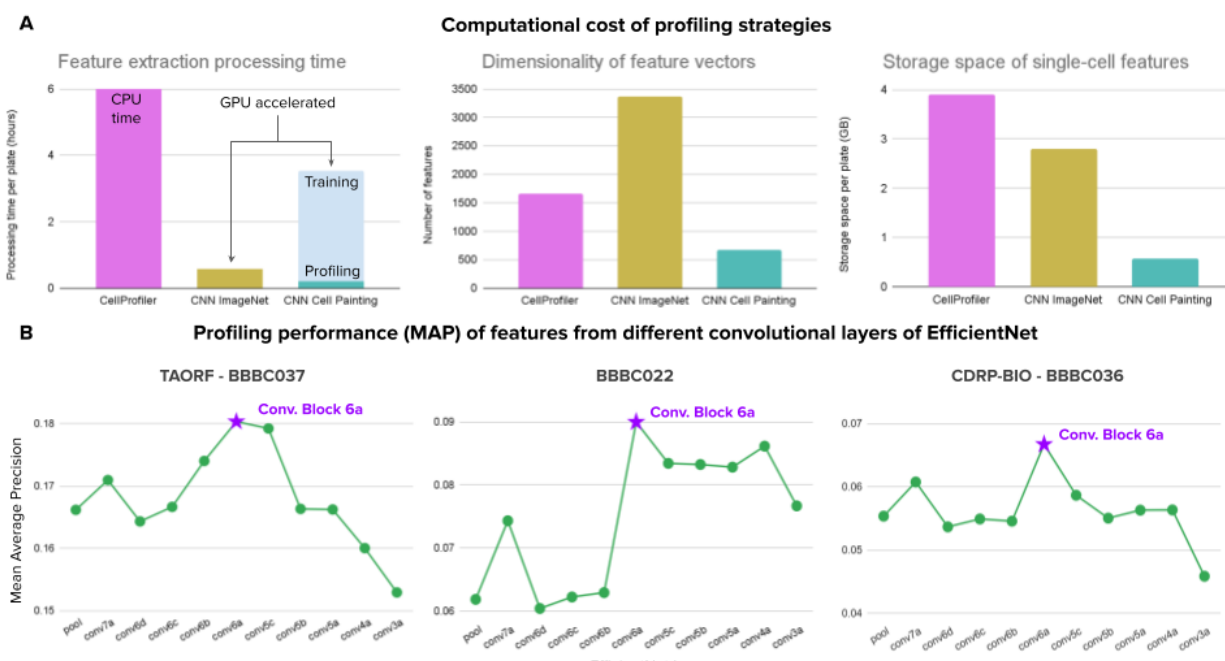


Supplementary Figure 1. UMAP visualization after combining three benchmark datasets in CNN ImageNet and CellProfiler feature-spaces. The top plots are colored by dataset (BBBC037 - blue, BBBC022 - red, BBBC036 - green). The bottom plots are colored by negative control (cyan) and treatments (red). This data was produced in preliminary experiments with dataset selection and shows the strongest treatments selected with use of corresponding features and Mahalanobis distance, instead of Euclidean in the main text. The CNN ImageNet features were selected with the penultimate layer. We see that different datasets are better integrated in the CellProfiler feature-space.

Representations learned from Cell Painting images are computationally efficient

The dimensionality of the feature space of the Cell Painting CNN-1 is also more compact than CellProfiler and the ImageNet CNN model (Supplementary Figure 2A), with the intermediate layer Conv6A of the EfficientNet B0 network being the best source of latent representations for downstream analysis in all of our experiments (Supplementary Figure 2B). This layer, after spatial average pooling, results in 672 features, compared to 1,700 of CellProfiler and 3,360 of ImageNet CNN (672 for each of the five imaging channels). The dimensionality of single-cell

features has an impact on storage space, especially for large scale experiments, making the Cell Painting CNN-1 an efficient choice too (Supplementary Figure 2A).



Supplementary Figure 2. Computational cost of profiling strategies. Beyond improved accuracy and better performance in downstream tasks, our Cell Painting CNN-1 is more computationally efficient than the baseline approaches. A) Computational cost in terms processing time per plate (in hours), dimensionality of representations (number of features), and storage space per plate (in GB), for the three representations evaluated in this work (x axes of plots). B) Downstream performance in the biological matching task from different convolutional layers of the EfficientNet model for all three datasets. We observe a consistent ability of Conv6A to yield better performance.

References

1. Bray, M.-A. *et al.* Cell Painting, a high-content image-based assay for morphological profiling using multiplexed fluorescent dyes. *Nat. Protoc.* **11**, 1757–1774 (2016).
2. Cimini, B. A. *et al.* Optimizing the Cell Painting assay for image-based profiling. *bioRxiv* 2022.07.13.499171 (2022) doi:10.1101/2022.07.13.499171.
3. Wawer, M. J. *et al.* Toward performance-diverse small-molecule libraries for cell-based phenotypic screening using multiplexed high-dimensional profiling. *Proceedings of the National Academy of Sciences* **111**, 10911–10916 (2014).
4. Cuccarese, M. F. *et al.* Functional immune mapping with deep-learning enabled phenomics applied to immunomodulatory and COVID-19 drug discovery. 2020.08.02.233064 (2020) doi:10.1101/2020.08.02.233064.
5. Way, G. P. *et al.* Morphology and gene expression profiling provide complementary information for mapping cell state. *bioRxiv* 2021.10.21.465335 (2021) doi:10.1101/2021.10.21.465335.
6. Simm, J. *et al.* Repurposing High-Throughput Image Assays Enables Biological Activity Prediction for Drug Discovery. *Cell Chem Biol* **25**, 611–618.e3 (2018).
7. Way, G. P. *et al.* Predicting cell health phenotypes using image-based morphology profiling. *Mol. Biol. Cell* mbcE20120784 (2021).
8. Moshkov, N. *et al.* Predicting compound activity from phenotypic profiles and chemical structures. *bioRxiv* 2020.12.15.422887 (2022) doi:10.1101/2020.12.15.422887.
9. Rohban, M. H. *et al.* Virtual screening for small molecule pathway regulators by image profile matching. *bioRxiv* 2021.07.29.454377 (2022) doi:10.1101/2021.07.29.454377.
10. Schiff, L. *et al.* Integrating deep learning and unbiased automated high-content screening to

- identify complex disease signatures in human fibroblasts. *bioRxiv* 2020.11.13.380576 (2021) doi:10.1101/2020.11.13.380576.
11. Caicedo, J. C., Arevalo, J. & Piccioni, F. Cell Painting predicts impact of lung cancer variants. *Mol. Biol. Cell* (2022).
12. Caicedo, J. C. *et al.* Data-analysis strategies for image-based cell profiling. *Nat. Methods* **14**, 849–863 (2017).
13. Caicedo, J. C., Singh, S. & Carpenter, A. E. Applications in image-based profiling of perturbations. *Curr. Opin. Biotechnol.* **39**, 134–142 (2016).
14. Chandrasekaran, S. N., Ceulemans, H., Boyd, J. D. & Carpenter, A. E. Image-based profiling for drug discovery: due for a machine-learning upgrade? *Nat. Rev. Drug Discov.* (2020) doi:10.1038/s41573-020-00117-w.
15. McQuin, C. *et al.* CellProfiler 3.0: Next-generation image processing for biology. *PLoS Biol.* **16**, e2005970 (2018).
16. Stirling, D. R. *et al.* CellProfiler 4: improvements in speed, utility and usability. *BMC Bioinformatics* **22**, 433 (2021).
17. Pawlowski, N., Caicedo, J. C., Singh, S., Carpenter, A. E. & Storkey, A. Automating Morphological Profiling with Generic Deep Convolutional Networks. *bioRxiv* 085118 (2016) doi:10.1101/085118.
18. Michael Ando, D., McLean, C. Y. & Berndl, M. Improving Phenotypic Measurements in High-Content Imaging Screens. *bioRxiv* 161422 (2017) doi:10.1101/161422.
19. Schiff, L. *et al.* Integrating deep learning and unbiased automated high-content screening to identify complex disease signatures in human fibroblasts. *Nat. Commun.* **13**, 1590 (2022).
20. Caicedo, J. C., McQuin, C., Goodman, A., Singh, S. & Carpenter, A. E. Weakly Supervised Learning of Single-Cell Feature Embeddings. *Proc. IEEE Comput. Soc. Conf. Comput. Vis. Pattern Recognit.* **2018**, 9309–9318 (2018).
21. Lu, A. X., Kraus, O. Z., Cooper, S. & Moses, A. M. Learning unsupervised feature

- representations for single cell microscopy images with paired cell inpainting. *PLoS Comput. Biol.* **15**, e1007348 (2019).
22. Hofmarcher, M., Rumetshofer, E. & Clevert, D. A. Accurate prediction of biological assays with high-throughput microscopy images and convolutional networks. *Journal of chemical* (2019).
23. Yang, S. J. *et al.* Applying Deep Neural Network Analysis to High-Content Image-Based Assays. *SLAS Discov* **24**, 829–841 (2019).
24. Mao, C. *et al.* Generative Interventions for Causal Learning. in *2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (IEEE, 2021). doi:10.1109/cvpr46437.2021.00394.
25. Schölkopf, B. *et al.* Toward Causal Representation Learning. *Proc. IEEE* **109**, 612–634 (2021).
26. Rubin, D. B. Estimating causal effects of treatments in randomized and nonrandomized studies. *J. Educ. Psychol.* **66**, 688–701 (1974).
27. Johansson, F., Shalit, U. & Sontag, D. Learning Representations for Counterfactual Inference. in *Proceedings of The 33rd International Conference on Machine Learning* (eds. Balcan, M. F. & Weinberger, K. Q.) vol. 48 3020–3029 (PMLR, 20–22 Jun 2016).
28. Becht, E. *et al.* Dimensionality reduction for visualizing single-cell data using UMAP. *Nat. Biotechnol.* (2018) doi:10.1038/nbt.4314.
29. Rohban, M. H. *et al.* Systematic morphological profiling of human gene and allele function via Cell Painting. *Elife* **6**, (2017).
30. Gustafsdottir, S. M. *et al.* Multiplex cytological profiling assay to measure diverse cellular states. *PLoS One* **8**, e80999 (2013).
31. Bray, M.-A. *et al.* A dataset of images and morphological profiles of 30 000 small-molecule treatments using the Cell Painting assay. *Gigascience* **6**, 1–5 (2017).
32. Chandrasekaran, S. N. *et al.* Three million images and morphological profiles of cells

- treated with matched chemical and genetic perturbations. *bioRxiv* 2022.01.05.475090 (2022) doi:10.1101/2022.01.05.475090.
33. Tan, M. & Le, Q. V. EfficientNet: Rethinking Model Scaling for Convolutional Neural Networks. *arXiv [cs.LG]* (2019).
 34. Gough, A. *et al.* Biologically Relevant Heterogeneity: Metrics and Practical Insights. *SLAS Discov* **22**, 213–237 (2017).
 35. Singh, S., Bray, M.-A., Jones, T. R. & Carpenter, A. E. Pipeline for illumination correction of images for high-throughput microscopy. *J. Microsc.* **256**, 231–236 (2014).
 36. Sandler, Howard & Zhu. Mobilenetv2: Inverted residuals and linear bottlenecks. *Proc. Estonian Acad. Sci. Biol. Ecol.*
 37. Hua, S. B. Z., Lu, A. X. & Moses, A. M. CytolImageNet: A large-scale pretraining dataset for bioimage transfer learning. *arXiv [cs.CV]* (2021).
 38. Deng, J. *et al.* ImageNet: A large-scale hierarchical image database. in *2009 IEEE Conference on Computer Vision and Pattern Recognition* 248–255 (2009).
 39. Otsu, N. A Threshold Selection Method from Gray-Level Histograms. *IEEE Trans. Syst. Man Cybern.* **9**, 62–66 (1979).
 40. Jones, T. R., Carpenter, A. & Golland, P. Voronoi-Based Segmentation of Cells on Image Manifolds. in *Computer Vision for Biomedical Image Applications* 535–543 (Springer Berlin Heidelberg, 2005).
 41. Voronoi, G. Nouvelles applications des paramètres continus à la théorie des formes quadratiques. Deuxième mémoire. Recherches sur les paralléloèdres primitifs. *J. Reine Angew. Math.* **1908**, 198–287 (1908).
 42. van der Walt, S. *et al.* scikit-image: image processing in Python. *PeerJ* **2**, e453 (2014).
 43. Kessy, A., Lewin, A. & Strimmer, K. Optimal Whitening and Decorrelation. *Am. Stat.* **72**, 309–314 (2018).
 44. Ljosa, V. *et al.* Comparison of methods for image-based profiling of cellular morphological

- responses to small-molecule treatment. *J. Biomol. Screen.* **18**, 1321–1329 (2013).
45. Rohban, M. H., Abbasi, H. S., Singh, S. & Carpenter, A. E. Capturing single-cell heterogeneity via data fusion improves image-based profiling. *Nat. Commun.* **10**, 2082 (2019).
 46. Manning, C. D. *Introduction to information retrieval*. (Syngress Publishing, 2008).
 47. Developers, T. *TensorFlow*. (Zenodo, 2021). doi:10.5281/ZENODO.4724125.