

Journal Club 15/05/18

Joseph Boyd JOSEPH.BOYD@MINES-PARISTECH.FR



May 27, 2018



A Multi-Scale Convolutional Neural Network for Phenotyping High-Content Cellular Images

William J. Godinez^{1,*}, Imtiaz Hossain¹, Stanley E. Lazic^{1,3}, John W. Davies² and Xian Zhang^{1,*}

¹Novartis Institutes for BioMedical Research Inc., Basel, Switzerland.

²Novartis Institutes for BioMedical Research Inc., Cambridge, Massachusetts, USA.

³Current address: Quantitative Biology, AstraZeneca, Cambridge, CB4 0WG, UK.

*To whom correspondence should be addressed.

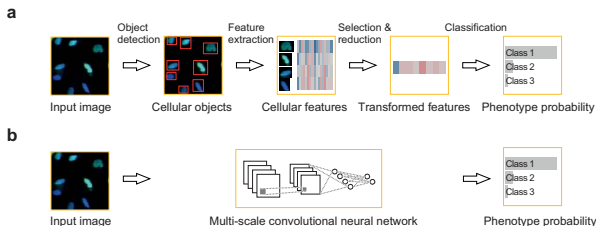
Associate Editor: Prof. Robert Murphy

Received on XXXXX; revised on XXXXX; accepted on XXXXX

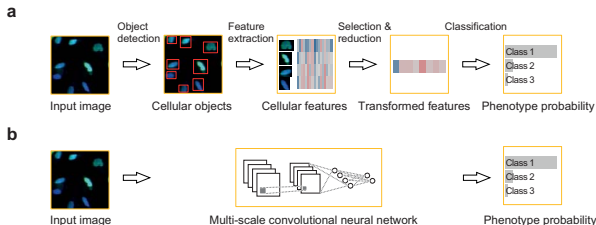
Abstract

Motivation: Identifying phenotypes based on high-content cellular images is challenging. Conventional image analysis pipelines for phenotype identification comprise multiple independent steps, with each step requiring method customization and adjustment of multiple parameters.

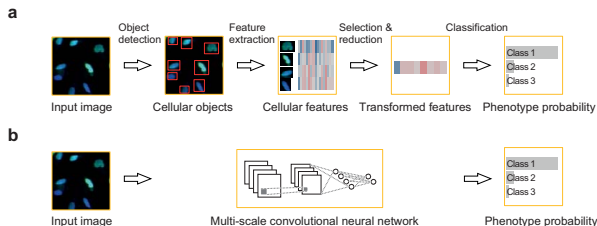
Results: Here we present an approach based on a multi-scale convolutional neural network (M-CNN) that classifies, in a single cohesive step, cellular images into phenotypes by using directly and solely the images' pixel intensity values. The only parameters in the approach are the weights of the neural network, which are automatically optimized based on training images. The approach requires no a priori knowledge or manual customization, and is applicable to single- or multi-channel images displaying single or multiple cells. We evaluated the classification performance of the approach on eight diverse benchmark datasets. The approach yielded overall a higher classification accuracy compared to state-of-the-art results, including those of other deep CNN architectures. In addition to using the network to simply obtain a yes-or-no prediction for a given phenotype, we use the probability outputs calculated by the network to quantitatively describe the phenotypes. Our study shows that these probability values correlate with chemical treatment concentrations. This finding validates further our approach and enables chemical treatment potency estimation via convolutional neural networks.



- The joint optimisation of all parameters across the entire [classical] analysis pipeline remains challenging (page 2)



- The joint optimisation of all parameters across the entire [classical] analysis pipeline remains challenging (page 2)
- A CNN-based approach that [...] in one cohesive step, classifies cellular images into phenotypes (page 2)



- The joint optimisation of all parameters across the entire [classical] analysis pipeline remains challenging (page 2)
- A CNN-based approach that [...] in one cohesive step, classifies cellular images into phenotypes (page 2)
- See also Orlov et al. (2008) or Uhlmann and Singh (2016)



- Network tested on 8 datasets spanning 6 cell lines (versatility)



- Network tested on 8 datasets spanning 6 cell lines (versatility)
- Tasks differ by dataset, for example CHO cell dataset used in Boland et al. 2001 for organelle classification



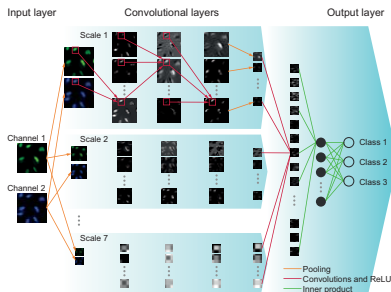
- Network tested on 8 datasets spanning 6 cell lines (versatility)
- Tasks differ by dataset, for example CHO cell dataset used in Boland et al. 2001 for organelle classification
- Of particular interest is the image set **BBBC21v2**¹ (MCF-7 breast cancer cell line)

¹Broad Institute Bioimage Benchmark Collection

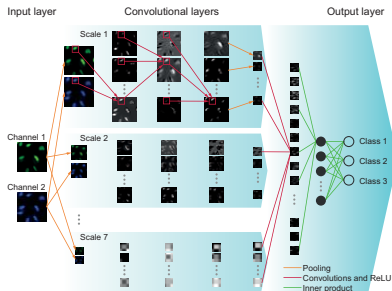


- Network tested on 8 datasets spanning 6 cell lines (versatility)
- Tasks differ by dataset, for example CHO cell dataset used in Boland et al. 2001 for organelle classification
- Of particular interest is the image set **BBBC21v2**¹ (MCF-7 breast cancer cell line)
- 38 compounds at 8 concentrations annotated into 13 mechanisms of action (MOA), including negative (DMSO)

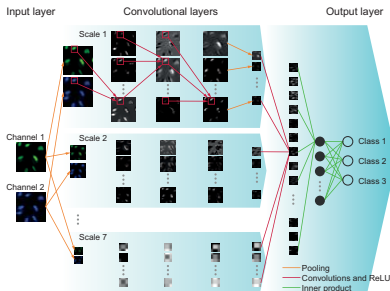
¹Broad Institute Bioimage Benchmark Collection



- MCNN consists of 7 paths of three 5×5 convolutional layers



- MCNN consists of 7 paths of three 5×5 convolutional layers
- Each path processes a downsampled version of the input $(1, \frac{1}{2}, \frac{1}{4}, \dots, \frac{1}{2^6})$ before aggregation



- MCNN consists of 7 paths of three 5×5 convolutional layers
- Each path processes a downsampled version of the input $(1, \frac{1}{2}, \frac{1}{4}, \dots, \frac{1}{2^6})$ before aggregation
- Network designed to be robust to cell line, magnification, etc.



Network produces $N_p = 13$ probabilities per replicate / drug / concentration / field



Network produces $N_p = 13$ probabilities per replicate / drug / concentration / field

- **Hard classification:** Predict MOA in a leave-one-out cross validation scheme

$$\hat{y} = \arg \max_k p(y = k | \mathbf{x})$$



Network produces $N_p = 13$ probabilities per replicate / drug / concentration / field

- **Hard classification:** Predict MOA in a leave-one-out cross validation scheme

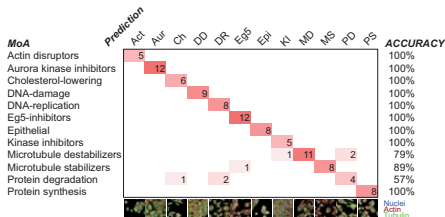
$$\hat{y} = \arg \max_k p(y = k | \mathbf{x})$$

- **Soft classification:** Track probabilities $p_{rf}(y = k)$ over titration series as median over replicate r of median over field f of view

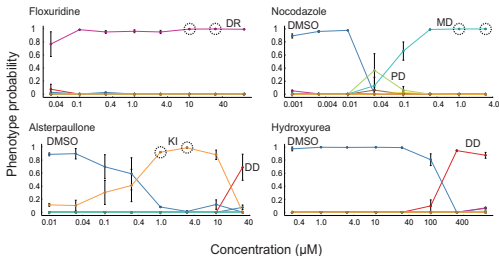
$$\rho_{rfk} = \text{med}\{\text{med}\{p_{rf}(y = k)\}_f\}_r$$



a



b





 Original Report

Compound Functional Prediction Using Multiple Unrelated Morphological Profiling Assays

SLAS Technology
1–9
© 2017 Society for Laboratory
Automation and Screening
DOI: 10.1177/2472630317740831
journals.sagepub.com/home/jla


France Rose^{1*}, Sreetama Basu^{1*}, Elton Rexhepaj^{1,2}, Anne Chauchereau³,
Elaine Del Nery², and Auguste Genovesio¹

Abstract

Phenotypic cell-based assays have proven to be efficient at discovering first-in-class therapeutic drugs mainly because they allow for scanning a wide spectrum of possible targets at once. However, despite compelling methodological advances, posterior identification of a compound's mechanism of action (MOA) has remained difficult and highly refractory to automated analyses. Methods such as the cell painting assay and multiplexing fluorescent dyes to reveal broadly relevant cellular components were recently suggested for MOA prediction. We demonstrated that adding fluorescent dyes to a single assay has limited impact on MOA prediction accuracy, as monitoring only the nuclei stain could reach compelling levels of accuracy. This observation suggested that multiplexed measurements are highly correlated and nuclei stain could possibly reflect the general state of the cell. We then hypothesized that combining unrelated and possibly simple cell-based assays could bring a solution that would be biologically and technically more relevant to predict a drug target than using a single assay multiplexing dyes. We show that such a combination of past screen data could rationally be reused in screening facilities to train an ensemble classifier to predict drug targets and prioritize a possibly large list of unknown compound hits at once.

Keywords

target prediction, high-content screening, mechanism of action, ensemble classifier



- The contribution of phenotypic screening to the discovery of first-in-class small-molecule drugs exceeded that of target-based approaches [in recent years] (page 1, paraphrasing Swinney & Anthony (2011))



- The contribution of phenotypic screening to the discovery of first-in-class small-molecule drugs exceeded that of target-based approaches [in recent years] (page 1, paraphrasing Swinney & Anthony (2011))
- While a complex painting assay on an optimised cell line represents a compelling approach [...] simple image-based assays on several cell lines may be more relevant biologically. (page 2)



- 9 (distinct) pleural mesothelioma + 1 prostate cancer cell lines



- 9 (distinct) pleural mesothelioma + 1 prostate cancer cell lines
- Prestwick dataset (1200 compounds), of which 614 have an MOA (DrugBank)



- 9 (distinct) pleural mesothelioma + 1 prostate cancer cell lines
- Prestwick dataset (1200 compounds), of which 614 have an MOA (DrugBank)
- Predict MOA (from 113 target classes) in leave-one-out cross validation scheme from Ljosa et al. (2013)



Following a classical pipeline:

- Average per-cell features over well to create *phenotypic profile* for each drug k :

$$profile_k = \left[\frac{1}{N} \sum_{n=1}^N x_{n1}^k, \dots, \frac{1}{N} \sum_{n=1}^N x_{iD}^k \right]$$

where x_{nd}^k is feature d of D for cell n of N for the k th drug.



Following a classical pipeline:

- Average per-cell features over well to create *phenotypic profile* for each drug k :

$$profile_k = \left[\frac{1}{N} \sum_{n=1}^N x_{n1}^k, \dots, \frac{1}{N} \sum_{n=1}^N x_{iD}^k \right]$$

where x_{nd}^k is feature d of D for cell n of N for the k th drug.

- Train one random forest (30 trees) per cell line (10 forests)



Following a classical pipeline:

- Average per-cell features over well to create *phenotypic profile* for each drug k :

$$profile_k = \left[\frac{1}{N} \sum_{n=1}^N x_{n1}^k, \dots, \frac{1}{N} \sum_{n=1}^N x_{iD}^k \right]$$

where x_{nd}^k is feature d of D for cell n of N for the k th drug.

- Train one random forest (30 trees) per cell line (10 forests)
- Predict MOA based on combining probabilities (ensemble prediction)

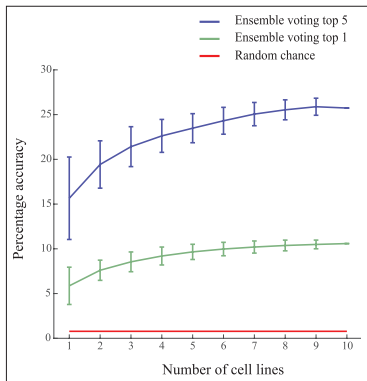


Figure: Top-5 and top-1 classification accuracy improves with the addition of cell lines



Improving drug discovery with high-content phenotypic screens by systematic selection of reporter cell lines

Jungseog Kang^{1,2,6}, Chien-Hsiang Hsu^{1,3,4,6}, Qi Wu¹, Shanshan Liu¹, Adam D Coster¹, Bruce A Posner⁵, Steven J Altschuler^{1,4} & Lani F Wu^{1,4}

High-content, image-based screens enable the identification of compounds that induce cellular responses similar to those of known drugs but through different chemical structures or targets. A central challenge in designing phenotypic screens is choosing suitable imaging biomarkers. Here we present a method for systematically identifying optimal reporter cell lines for annotating compound libraries (ORACLs), whose phenotypic profiles most accurately classify a training set of known drugs. We generate a library of fluorescently tagged reporter cell lines, and let analytical criteria determine which among them—the ORACL—best classifies compounds into multiple, diverse drug classes. We demonstrate that an ORACL can functionally annotate large compound libraries across diverse drug classes in a single-pass screen and confirm high prediction accuracy by means of orthogonal, secondary validation assays. Our approach will increase the efficiency, scale and accuracy of phenotypic screens by maximizing their discriminatory power.



- Molecular screens expensive for large-scale drug screening



- Molecular screens expensive for large-scale drug screening
- Image-based high-content screening (HCS) a promising alternative



- Molecular screens expensive for large-scale drug screening
- Image-based high-content screening (HCS) a promising alternative
- Main objective of *identifying lead compounds in a single-pass screen* (page 1) while respecting scalability



In a *live-cell* assay, an optimal cell line is chosen by brute force to be used in a large-scale screen:



In a *live-cell* assay, an optimal cell line is chosen by brute force to be used in a large-scale screen:

- 1 Construct library of reporter cell lines



In a *live-cell assay*, an optimal cell line is chosen by brute force to be used in a large-scale screen:

- 1 Construct library of reporter cell lines
- 2 Identify optimal reporter cell line for compound libraries (ORACL)



In a *live-cell assay*, an optimal cell line is chosen by brute force to be used in a large-scale screen:

- 1 Construct library of reporter cell lines
- 2 Identify optimal reporter cell line for compound libraries (ORACL)
- 3 Use ORACL to identify lead compounds in a single-pass screen



A549 non-small cell lung cancer *parent* cell line. 93 reporter clones were produced by:

- randomly labeling proteins with yellow fluorescent protein as an extra exon via viral transfection (*central dogma tagging*).



A549 non-small cell lung cancer *parent* cell line. 93 reporter clones were produced by:

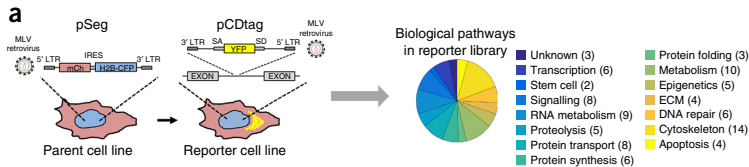
- randomly labeling proteins with yellow fluorescent protein as an extra exon via viral transfection (*central dogma tagging*).
- The genes can then be identified using 3'RACE (duplicates discarded)

Step 1: library of reporter cell lines



A549 non-small cell lung cancer *parent* cell line. 93 reporter clones were produced by:

- randomly labeling proteins with yellow fluorescent protein as an extra exon via viral transfection (*central dogma tagging*).
- The genes can then be identified using 3'RACE (duplicates discarded)





- The 93 reporter cell lines were treated by a small 6 classes $\times$ 5 drugs + DMSO *reference set*



- The 93 reporter cell lines were treated by a small 6 classes $\times$ 5 drugs + DMSO *reference set*
- Phenotypic profile created as in Perlman et al. (2004):

$$profile_k = \left[KS(X_{:,1}^k, X_{:,1}^-), \dots, KS(X_{:,D}^k, X_{:,D}^-) \right]$$

where $X_{:,d}^k$ and $X_{:,d}^-$ is feature d of D of all cells for drug k and the negative control respectively.



- The 93 reporter cell lines were treated by a small 6 classes $\times$ 5 drugs + DMSO *reference set*
- Phenotypic profile created as in Perlman et al. (2004):

$$profile_k = \left[KS(X_{:,1}^k, X_{:,1}^-), \dots, KS(X_{:,D}^k, X_{:,D}^-) \right]$$

where $X_{:,d}^k$ and $X_{:,d}^-$ is feature d of D of all cells for drug k and the negative control respectively.

- One drug was removed at random from each class and assigned a class as the nearest centroid of the remaining four drugs per class

Step 2: identification of the ORACL

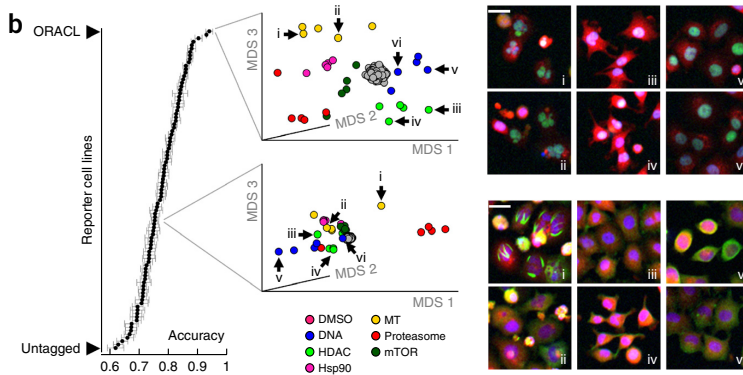


Figure: The reporter cell line maximising prediction accuracy (94% over 100 repetitions) was chosen as the ORACL (XRCC5 DSB-repair protein)



- The selected ORACL is then tested on a large suite of drug libraries (10483 drugs total)



- The selected ORACL is then tested on a large suite of drug libraries (10483 drugs total)
- This yields ~ 62000 live-cell images after 24h and 48h and ~ 60 million cells with ~ 230 morphological features apiece



- The selected ORACL is then tested on a large suite of drug libraries (10483 drugs total)
- This yields ~ 62000 live-cell images after 24h and 48h and ~ 60 million cells with ~ 230 morphological features apiece
- Linear discriminant analysis to reduce dimensionality, in which reference drug clusters are used as nearest neighbour classifier



- The selected ORACL is then tested on a large suite of drug libraries (10483 drugs total)
- This yields ~ 62000 live-cell images after 24h and 48h and ~ 60 million cells with ~ 230 morphological features apiece
- Linear discriminant analysis to reduce dimensionality, in which reference drug clusters are used as nearest neighbour classifier
- Confidence scores (based on distance, calibrated using reference set) calculated for each classification. Compounds left unclassified when confidence low ($p < 0.1$)

Step 3: identification of multi-class hits

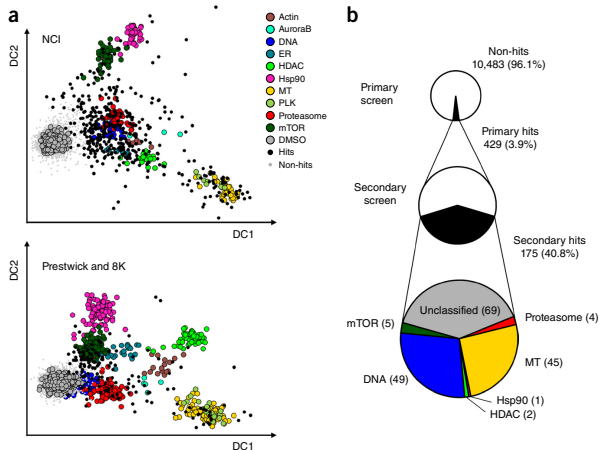


Figure: *Primary hits* (429 total) defined as non-DMSO classified (whether to a known class or not)



- A secondary screen on the primary hits was performed to eliminate weak phenotypes (precision over recall) leaving 175 *secondary hits*



- A secondary screen on the primary hits was performed to eliminate weak phenotypes (precision over recall) leaving 175 *secondary hits*
- They are then able to validate HDAC directly through literature (gratifyingly)



- A secondary screen on the primary hits was performed to eliminate weak phenotypes (precision over recall) leaving 175 *secondary hits*
- They are then able to validate HDAC directly through literature (gratifyingly)
- The others are validated through a series of secondary experiments