



天津科技大学学报
Journal of Tianjin University of Science & Technology
ISSN 1672-6510, CN 12-1355/N

《天津科技大学学报》网络首发论文

题目: 基于迁移学习框架的单细胞药物反应预测
作者: 李玉田, 葛凤雅, 王林
DOI: 10.13364/j.issn.1672-6510.20240222
收稿日期: 2024-11-06
网络首发日期: 2025-04-30
引用格式: 李玉田, 葛凤雅, 王林. 基于迁移学习框架的单细胞药物反应预测[J/OL]. 天津科技大学学报. <https://doi.org/10.13364/j.issn.1672-6510.20240222>



网络首发: 在编辑部工作流程中, 稿件从录用到出版要经历录用定稿、排版定稿、整期汇编定稿等阶段。录用定稿指内容已经确定, 且通过同行评议、主编终审同意刊用的稿件。排版定稿指录用定稿按照期刊特定版式(包括网络呈现版式)排版后的稿件, 可暂不确定出版年、卷、期和页码。整期汇编定稿指出版年、卷、期、页码均已确定的印刷或数字出版的整期汇编稿件。录用定稿网络首发稿件内容必须符合《出版管理条例》和《期刊出版管理规定》的有关规定; 学术研究成果具有创新性、科学性和先进性, 符合编辑部对刊文的录用要求, 不存在学术不端行为及其他侵权行为; 稿件内容应基本符合国家有关书刊编辑、出版的技术标准, 正确使用和统一规范语言文字、符号、数字、外文字母、法定计量单位及地图标注等。为确保录用定稿网络首发的严肃性, 录用定稿一经发布, 不得修改论文题目、作者、机构名称和学术内容, 只可基于编辑规范进行少量文字的修改。

出版确认: 纸质期刊编辑部通过与《中国学术期刊(光盘版)》电子杂志社有限公司签约, 在《中国学术期刊(网络版)》出版传播平台上创办与纸质期刊内容一致的网络版, 以单篇或整期出版形式, 在印刷出版之前刊发论文的录用定稿、排版定稿、整期汇编定稿。因为《中国学术期刊(网络版)》是国家新闻出版广电总局批准的网络连续型出版物(ISSN 2096-4188, CN 11-6037/Z), 所以签约期刊的网络版上网络首发论文视为正式出版。



基于迁移学习框架的单细胞药物反应预测

李玉田, 葛凤雅, 王 林
(天津科技大学人工智能学院, 天津 300457)

摘要: 肿瘤组织中存在多种类型的细胞, 这些细胞在形态、功能和生物学特性上存在差异, 诸如不同细胞类型对同一药物的敏感性反应不同, 而这也是导致癌症治疗高复发率和低效力的重要原因。单细胞 RNA 测序技术是针对这一细胞类型异质性的研究工具, 本文提出了一种基于变分自动编码器的迁移学习框架, 并用于单细胞药物反应预测。该框架整合了混合测序的细胞系药物基因组数据和单细胞 RNA 测序数据, 使用最大均值差异和最优传输, 在低维潜在空间中对细胞系和单细胞进行特征对齐, 从而在单细胞水平上准确预测药物反应。模型和最新的单细胞药物反应预测方法做基准比较, 在 8 个数据集上进行训练验证, 模型取得了较高的预测性能。消融实验和鲁棒性实验验证了模型各模块的有效性以及模型的稳定性。本研究提供了新的迁移学习思路, 为个性化和精确的癌症治疗提供了一种新的策略, 并有助于精准医学的发展。

关键词: 迁移学习; 特征对齐; 单细胞; 药物反应预测

中图分类号: TP3-05

文献标志码: A

文章编号: 1672-6510 (0000)00-0000-00

Single-Cell Drug Response Prediction Based on Transfer Learning Framework

LI Yutian, GE Fengya, WANG Lin

(College of Artificial Intelligence, Tianjin University of Science & Technology, Tianjin 300457, China)

Abstract: There are many types of cells in tumor tissue, and these cells have differences in morphology, function and biological characteristics, such as different cell types have different sensitivity to the same drug, which is also an important reason for the high recurrence rate and low efficacy of cancer treatment. Single-cell RNA sequencing technology is an important research tool for the cell type heterogeneity, and this paper proposes a transfer learning framework based on variational auto-encoder, which is used for single-cell drug response prediction. The framework integrates bulk-sequenced cell line pharmacogenomic data and single-cell RNA sequencing data, and adopts maximum mean discrepancy and optimal transport to align cell lines and single cells in a low-dimensional latent space to accurately predict drug response at the single-cell level. The model was compared with the state-of-the-art single-cell drug response prediction methods, i.e., these methods were trained and verified on eight data sets, and our model achieved higher predictive performance. Ablation experiments and robustness experiments verified the validity of each module of the model and the stability of the model. This study provides a new idea for transfer learning, and a new strategy for personalized and precise cancer treatment, which contributes to the development of precision medicine.

Key words: transfer learning; feature alignment; single cell; drug response prediction

精准医学在揭示癌症基因组复杂性方面取得了显著进展^[1]。随着基因测序技术的快速发展, 对癌症的治疗研究越来越全面, 混合细胞系测序技术可以帮助研究人员在组织水平上研究癌症对药物的反

应, 这有力促进了精准医疗的发展。然而, 由于混合细胞系测序数据中不同细胞状态之间的异质性, 癌症药物治疗常常面临高复发率和低效力的问题。肿瘤内异质性是癌症治疗失败的重要原因^[2], 这种异

质性导致个别细胞对药物的反应不同, 最终导致少量残留癌细胞并引发癌症复发^[3]。近年来, 单细胞 RNA 测序技术得到了快速发展。单细胞测序技术可以在细胞水平上研究癌细胞异质性, 使研究者能够以更高的分辨率分析肿瘤内的异质性^[4]。但是, 目前受制于单细胞测序技术的成本以及数据量较少的缘故, 使用单细胞测序技术支撑肿瘤药物研发仍较为受限。另一方面, 得益于前人研究, 体外药物筛选研究已经为各种癌症细胞系产生了药物反应数据^[5-6], 积累了大规模的细胞系混合测序数据。因此, 采用创新的方法, 将大规模的细胞系混合 RNA 测序数据与单细胞 RNA 测序数据整合起来, 可以充分利用两者的优势。

迁移学习旨在将一个数据域(源域)中获得的知识转移到另一个数据域(目标域), 以改善目标域的预测性能。在此背景下, 迁移学习可以用于将细胞系混合 RNA 测序数据中的药物反应信息映射到单细胞 RNA 测序数据, 从而实现单细胞水平的药物反应预测。一些研究方法已经展示了迁移学习在单细胞药物反应预测研究中的潜力^[1,7], 本研究提出了一种新的基于迁移学习的单细胞药物反应预测框架(single-cell drug response prediction based on transfer learning, SCTL), 该框架基于最优传输和最大均值差异, 在细胞系测序数据和单细胞测序数据之间构建了迁移桥梁, 并在单细胞水平上实现了准确的药物反应预测。

1 数据介绍

共收集整理 8 个数据集。对于源域数据, 使用了来自公开数据库 GDSC^[8]和 CCLE^[9]的混合细胞系测序数据。Genomics of Drug Sensitivity in Cancer (GDSC) 数据库的在线访问途径为, <https://www.cancerrxgene.org/>。Cancer Cell Line Encyclopedia (CCLE) 数据库的在线访问途径为, <https://sites.broadinstitute.org/ccle>。数据 1~5 的源域数据收集了 1280 个细胞系, 其中有 804 个细胞系来自 GDSC 数据库, 剩下的来自 CCLE 数据库。数据 6~8 的源域数据收集了 1001 个细胞系, 全部来自 GDSC

数据库。对于目标域数据, 使用了来自口腔鳞状细胞癌、肺癌、前列腺癌、急性髓性白血病、头颈部鳞状细胞癌的单细胞测序数据, 数据 1~5 目标域数据的 GEO 登记号在表格 2 中, 数据 6~8 的目标域数据来自 Zheng 等^[7]的研究。详细的源域数据见表 1, 目标域数据见表 2。

表 1 细胞系混合 RNA 测序数据(源域数据)

Tab. 1 Bulk RNA-seq data of cell lines (source data).

序号	药物	耐药细胞系 个数	敏感细胞系 个数	基因 数
Data1	Cisplatin	791	489	12966
Data2	Gefitinib	799	481	12499
Data3	Docetaxel	788	492	13476
Data4	LBET.762	724	556	9566
Data5	Bortezomib	770	510	12715
Data6	Gefitinib	714	115	10610
Data7	Vorinostat	774	60	10610
Data8	AR-42	811	80	10610

表 2 单细胞 RNA 测序数据(目标域数据)

Tab. 2 Single-cell RNA-seq data (target data).

序号	药物	耐药细 胞个数	敏感细 胞个数	基因 数	来源
Data1	Cisplatin	187	361	12966	GSE117872
Data2	Gefitinib	470	37	12499	GSE112274
Data3	Docetaxel	162	162	13476	GSE140440
Data4	LBET.762	685	734	9566	GSE110894
Data5	Bortezomib	4565	2875	12715	GSE114461
Data6	Gefitinib	33	33	10610	JHU006
Data7	Vorinostat	33	33	10610	JHU006
Data8	AR-42	33	33	10610	JHU006

使用 Scanpy 包^[10]对单细胞 RNA 测序数据进行预处理, 过滤掉了检测到的基因少于 200 个的细胞, 以及仅在 3 个以下细胞中检测到的基因。仅保留了在细胞系和单细胞中都共有的基因。最后, 使用 MinMaxScaler 函数将源域和目标域数据归一化。

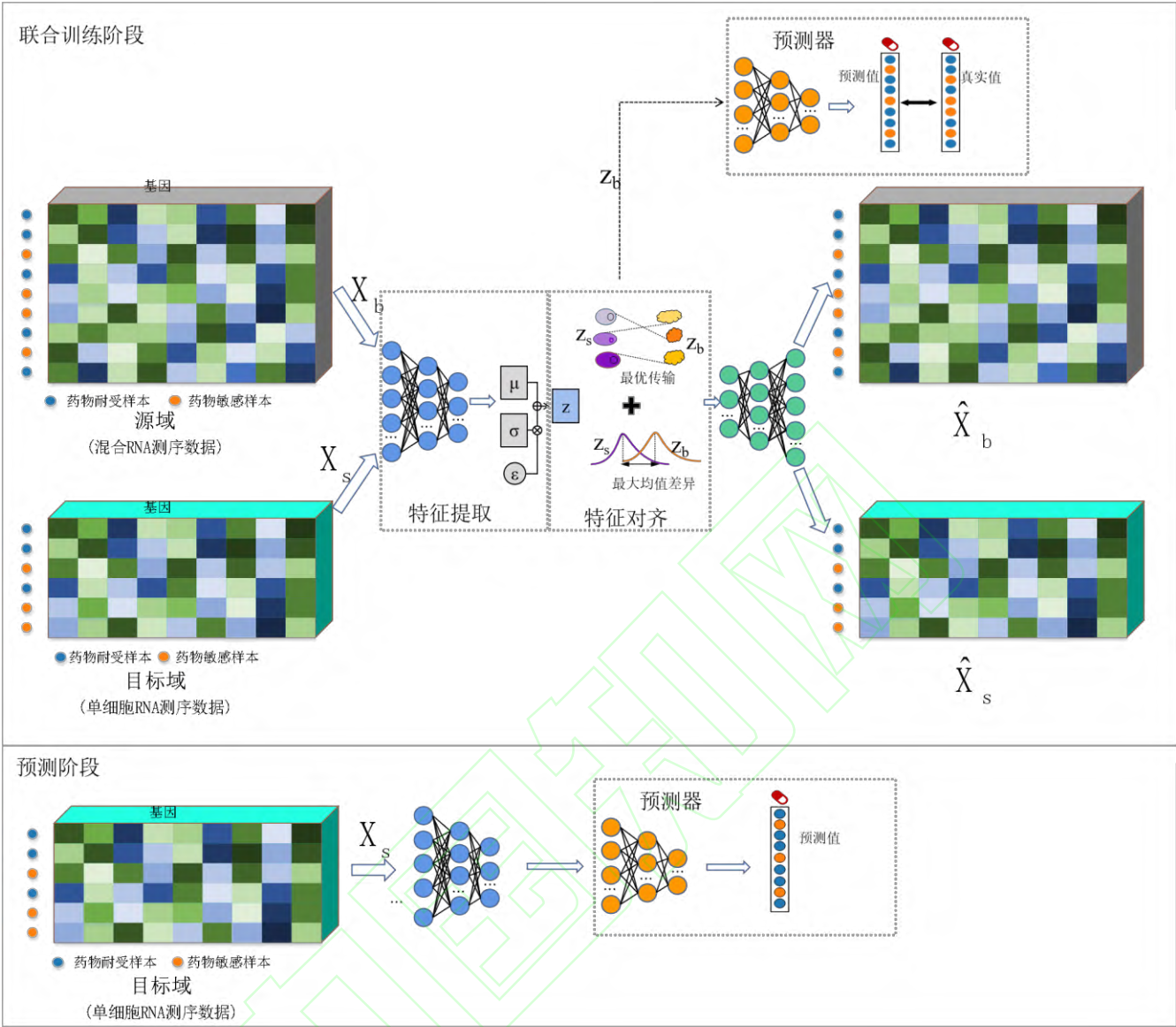


图 1 SCTL 模型框架图

Fig. 1 Framework of the SCTL model

2 计算方法

2.1 迁移学习框架的构建

本研究针对细胞系混合 RNA 测序数据和单细胞 RNA 测序数据之间的差异, 构建深度学习网络提取跨域数据的特征, 通过特征对齐的计算方法, 实现从细胞系水平到单细胞水平的药物反应预测。模型框架见图 1。

对于源域数据和目标域数据, 构建了一个变分自动编码器^[11], 包括一个用于数据特征提取的编码器, 和一个相同层数的解码器, 解码器用于将特征重构回原数据。通过模型的训练迭代, 编码器会将源域数据和目标域数据映射到共享的隐空间中, 解

码器会在每一轮训练中从隐空间中重构数据, 规范隐空间数据特征。训练结束后, 可以获取在同一隐空间的源域和目标域数据的特征。

在特征对齐的计算方法上, 采用了多种手段弥合跨域数据的特征, 包括最大均值差异^[12]的计算和最优传输^[11]的方法。

针对单细胞药物反应预测问题, 构建了一个用于药物反应预测二分类任务的多层感知机, 其中输出神经元 1 表示敏感, 0 表示耐药。在训练药物反应预测器的过程中, 使用自动编码器提取的细胞系特征和真实的细胞系药物反应信息训练药物反应预测器。当模型训练结束后, 将编码器提取的单细胞隐空间特征输入到药物反应预测器, 获得模型对单细胞

胞的药物反应预测值。

2.2 跨域特征的提取

来自源域的细胞系数据是一个包含 m 个细胞系的混合 RNA 测序数据 $X_b = \{x_b^i\}_{i=1}^m$ ；来自目标域的单细胞数据是一个包含 n 个单细胞的 RNA 测序数据 $X_s = \{x_s^j\}_{j=1}^n$ 。在这里，将两个数据 X_b 和 X_s 纵向拼接作为联合输入 $X = [X_b, X_s]$ ，利用变分自动编码器将其映射到共享潜在空间，学习嵌入特征。构建的变分自动编码器由一个跨域共享的编码器 $q_\phi(z|x)$ ，和一个共享的解码器 $p_\theta(x|z)$ 组成，其中 ϕ 和 θ 分别是编码器和解码器的权重参数。对于数据 x ，编码器输出均值向量 $\mu_{z|x}$ 以及标准差向量 $\sigma_{z|x}$ ，潜在向量 z 通过重参数化计算得到，即 $z = \mu_{z|x} + \sigma_{z|x} \square v$ ，其中 $\square$ 表示元素逐个相乘， v 从 $N(0, I)$ 中采样。通过这种方式，可以分别得到细胞系 x_b^i ($i=1, \dots, m$) 和单细胞 x_s^j ($j=1, \dots, n$) 的 H 维潜在向量 z_b^i 和 z_s^j 。

对于重构部分，基于潜在向量 z_b^i 和 z_s^j 作为输入，解码器会输出源域重构数据 $\hat{x}_b^i$ 以及目标域重构数据 $\hat{x}_s^j$ 。重构的损失函数如下：

$$\begin{aligned} L_{rec} &= L_{rec_b} + L_{rec_s} \\ &= \frac{1}{B_b} \sum_{i=1}^{B_b} MSE(x_b^i, \hat{x}_b^i) + \frac{1}{B_s} \sum_{j=1}^{B_s} MSE(x_s^j, \hat{x}_s^j) \\ &= \frac{1}{B_b g} \sum_{i=1}^{B_b} \sum_{k=1}^g (x_b^{ik} - \hat{x}_b^{ik})^2 \\ &\quad + \frac{1}{B_s g} \sum_{j=1}^{B_s} \sum_{k=1}^g (x_s^{jk} - \hat{x}_s^{jk})^2 \end{aligned} \quad (1)$$

其中： B_b 和 B_s 分别表示数据集 X_b 和 X_s 的批量大小， $MSE()$ 表示均方误差， g 是基因个数。此外，还考虑了 H 维潜在向量 z 的后验分布与多元高斯分布之间的 Kullback-Leibler (KL) 散度，表示为：

$$\begin{aligned} D_{KL}(q_\phi(z|x) || N(0, I)) &= \\ -\frac{1}{2} \sum_{k=1}^H (1 + \log((\sigma_{z|x}^k)^2) - (\mu_{z|x}^k)^2 - (\sigma_{z|x}^k)^2) \end{aligned} \quad (2)$$

最后，结合重构损失和 KL 散度，定义变分自动编码器损失函数如下：

$$\begin{aligned} L_{vae} &= \lambda_1 L_{rec} + \lambda_2 L_{KL} \\ &= \frac{1}{B_b} \sum_{i=1}^{B_b} MSE(x_b^i, \hat{x}_b^i) + \frac{1}{B_s} \sum_{j=1}^{B_s} MSE(x_s^j, \hat{x}_s^j) \\ &\quad + \frac{1}{B_b + B_s} \sum_{i=1}^{B_b} D_{KL}(q_\phi(z_b^i | x_b^i) || N(0, I)) \\ &\quad + \frac{1}{B_b + B_s} \sum_{j=1}^{B_s} D_{KL}(q_\phi(z_s^j | x_s^j) || N(0, I)) \end{aligned} \quad (3)$$

其中 $\lambda_1 = 1$, $\lambda_2 = 1$ 。

2.3 特征对齐的计算

对模型提取的源域低维特征 z_b 和目标域低维特征 z_s ，采用最优均值差异和最优传输的计算方法，在隐空间上，缩小两种数据域特征分布的距离，实现对齐两种数据域。

最大均值差异 (maximum mean discrepancy, MMD) 是一种基于核的统计检验方法，用于判断两个给定分布是否相同^[12]。MMD 是通过比较两个分布在一个再生核希尔伯特空间中的嵌入完成的。在本研究中，对于输入的细胞系 RNA 测序数据 X_b 和单细胞 RNA 测序数据 X_s ，数据经过编码器映射到潜在空间时，引入最大均值差异估算源域嵌入 z_b 和目标域嵌入 z_s 之间的相似度。损失函数定义如下：

$$L_{mmd} = || \frac{1}{m} \sum_{i=1}^m \phi(x_b^i) - \frac{1}{n} \sum_{j=1}^n \phi(x_s^j) ||_H \quad (4)$$

其中， $\phi(\cdot)$ 表示是将数据映射到再生核希尔伯特空间 (RKHS)。RKHS 范数 $||\cdot||_H$ 用于测量两个数据向量之间的距离。

最优传输是一种求解两个分布之间最小距离的优化算法^[13]。其核心思想是将一个概率分布转化为另一个概率分布，使得转化后的分布与目标分布之间的距离最小。最优传输算法在机器学习领域受到了越来越多的关注，特别是 Wasserstein-GAN^[14] 的出现，引起了研究者对最优传输的关注。在本研究中，引入最优传输作为正则化项调整两种数据域的嵌入。对于嵌入 z_b^i 和 z_s^j ，其后验分布分别为高斯分布 $N(\mu_{z_b^i|x_b^i}, \sigma_{z_b^i|x_b^i}^2 I)$ 和 $N(\mu_{z_s^j|x_s^j}, \sigma_{z_s^j|x_s^j}^2 I)$ ，嵌入之间的传输成本为：

$$C_{ij} = || \mu_{z_b^i|x_b^i} - \mu_{z_s^j|x_s^j} ||^2 + || \sigma_{z_b^i|x_b^i} - \sigma_{z_s^j|x_s^j} ||^2$$

(5)

然后, 批量大小为 B_b 和 B_s 的最优传输方案 T^* 计算如下:

$$T^* = \underset{T \in \square_{+}^{B_b \times B_s}}{\operatorname{argmin}} \langle C, T \rangle + \lambda_3 H(T) + \lambda_4 D_{KL} \left(T 1_{B_s} \parallel \frac{1}{B_b} 1_{B_b} \right) + \lambda_4 D_{KL} \left(T^T 1_{B_b} \parallel \frac{1}{B_s} 1_{B_s} \right) \quad (6)$$

其中: $\lambda_3 = 0.1$ 和 $\lambda_4 = 1$ 是平衡参数。 T 是传输方案,

$\langle C, T \rangle = \sum_{i,j} C_{ij} T_{ij}$ 表示传输成本。 $H(T) = \sum_{i,j} T_{ij} \log T_{ij}$ 为熵正则

化项, 可以平衡传输成本和分布均匀性之间的关系。

$D_{KL} \left(T 1_{B_s} \parallel \frac{1}{B_b} 1_{B_b} \right)$ 和 $D_{KL} \left(T^T 1_{B_b} \parallel \frac{1}{B_s} 1_{B_s} \right)$ 两个 KL 散度项分别表示当源域和目标域的经验分布为均匀分布时的软边缘约束, 1_{B_s} 和 1_{B_b} 表示元素全为 1 的向量, 其

维度分别为 B_s 和 B_b 。公式(6)表示一个严格的凸优化问题, 可以通过采用一种高效的不精确近端点方法^[15]

计算最优传输方案 T^* , 其迭代公式如下所示:

$$\alpha^{(l+1)} = \left(\frac{\frac{1}{B_b} 1_{B_b}}{G \beta^{(l)}} \right)^{\frac{\lambda_5}{\lambda_5 + \lambda_4}}, \beta^{(l+1)} = \left(\frac{\frac{1}{B_s} 1_{B_s}}{G^T \alpha^{(l+1)}} \right)^{\frac{\lambda_5}{\lambda_5 + \lambda_4}} \quad (7)$$

初始化时, $\beta^{(0)} = \frac{1}{B_s} 1_{B_s}$, 其中 $G = [G_{ij}]$ 且

$G_{ij} = T_{ij}^{(l)} e^{-\frac{C_{ij}}{\lambda_4}}$ 。可以从最后一次近端点迭代中获得 α 和 β 的最终值。最优传输方案的计算就转换成了 $T_{ij}^* = \alpha_i G_{ij} \beta_j$ 。因此, 最优传输损失确定如下:

$$L_{ot} = \sum_{i,j} C_{ij} \times T_{ij}^* \quad (8)$$

2.4 药物反应预测器的构建

针对来自细胞系的药物反应数据 $\{y_b^i\}_{i=1}^m$, 即 m

个细胞系对给定药物的二值化反应信息, 其中 0 表示药物耐受, 1 表示药物敏感; 以及来自变分自动编码器提取的细胞系特征 z_b^i ($i=1, \dots, m$), 构建了深度为 5 的多层感知机, 用作药物反应预测这一二分类任务。对于训练数据 $\{z_b^i, y_b^i\}_{i=1}^m$, 药物反应预测器的损失函数定义为交叉熵损失函数 BCE(), 即:

$$BCE(y_b^i, \hat{y}_b^i) = -y_b^i \log(\hat{y}_b^i) - (1 - y_b^i) \log(1 - \hat{y}_b^i) \quad (9)$$

其中, y_b^i 是细胞系对于给定药物的真实二值化反应值, $\hat{y}_b^i$ 是预测器输出的预测值。对于批量大小为 B_b 的细胞系, 药物反应预测器的损失函数如下:

$$L_{pred} = \sum_{i=1}^{B_b} BCE(y_b^i, \hat{y}_b^i) \quad (10)$$

最后, SCTL 模型总的损失函数定义为:

$$L_{all} = \rho_1 L_{vae} + \rho_2 L_{ot} + L_{mmd} + \rho_3 L_{pred} \quad (11)$$

其中 ρ_1 、 ρ_2 和 ρ_3 是平衡参数。 ρ_1 可选为 0.1、0.5、

1、2, ρ_2 可选为 0.5、1、10, ρ_3 可选为 0.5、1、5、

15。对于每个数据集, 可以通过调整平衡参数, 使模型的性能处在最佳状态。

当模型训练完成后, 将经过跨域特征提取和特征对齐后的单细胞嵌入 z_s 输入到药物反应预测器得到单细胞对于给定药物的预测反应信息。

2.5 模型评价指标

在本研究中, 单细胞药物反应预测是一种二分类预测任务。二分类预测问题通常采用 ROC(receiver operating characteristic) 曲线下面积 (AUC), 以及 PR (precision-recall) 曲线下面积 (AUPR) 作为评价指标。其中 ROC 曲线描绘了在不同阈值下模型预测的真阳性率 (TPR) 和假阳性率 (FPR), AUC 的值范围在 0~1 之间, 值越大表示模型性能越好。AUPR 表示模型预测在所有阈值下的精确率与召回率的表现, 值越高表示模型性能越好。

在模型的训练过程中, 参考了 SCAD 工作^[7]中关于环境随机种子的设置, SCAD 选取了 seed=42、seed=50、seed=60、seed=70、seed=80 环境随机种子, 在此基础上, 又增加一个种子 seed=124, 共六个随机种子, 用于数据批次的划分, 然后计算了模型在这六个随机种子下的平均性能。

2.6 模型的可解释性

模型在训练过程中,通过最优传输的计算,可以获得当 m 个细胞系和 n 个单细胞都是均匀分布时,最小化传输费用的最优传输矩阵,即 $T_{m \times n}^*$,其中行表示细胞系,列表示细胞,矩阵值表示细胞系分布和单细胞分布间的数据传输量,即细胞系和单细胞间的映射匹配关系。接着对模型取得较好性能的数据集展开关键基因的分析。对模型预测的细胞药物响应值二值化,将细胞划分为预测敏感细胞簇和预测耐药细胞簇。然后进行 Wilcoxon 秩和检验,选择 Bonferroni 校正后的 p 值小于 0.05、对数倍数变化大于 0.7 的基因作为高表达的基因。

3 实验

3.1 性能和结果

一共在 8 个数据集上进行了模型测试,并与当前最新的单细胞癌症药物敏感性预测方法进行了比较,这些方法包括 scDEAL、SCAD、Beyondcell^[16]和 DREEP^[17]。

scDEAL 模型通过对细胞系混合 RNA 测序数据和单细胞 RNA 测序数据建模,构建了两个去噪自编码器和一个药物反应预测器,并用最大均值差异来对齐细胞系和单细胞的表征,从而实现将细胞系的真实药物反应信息迁移学习到单细胞水平上。

SCAD 模型采用了对抗生成网络 GAN^[18]框架,其构建了一个特征提取器,用于学习细胞系和单细胞的表征;一个药物反应预测器,用于学习对药物敏感性推断有用的信息特征;以及一个域判别器,使特征提取器能够从两个域(细胞系和单细胞)中学习域不变特征。

Beyondcell 方法先从细胞系的基因表达数据中识别出药物反应的上调和下调基因集,然后基于单细胞基因表达数据,计算单细胞在这些基因集的 BCS (Beyondcell Score) 得分,该得分可以用来对单细胞依据药物敏感反应的强弱进行排序。

DREEP 方法利用皮尔逊相关系数将混合测序基因表达与药物基因组筛选的药物反应数据联系起来,构建了药物敏感性基因组图谱(GPDS)的排序列表,GPDS 排序列表中药物的耐药性预测生物标志物位于顶部,敏感性预测生物标志物位于底部。为了在单细胞水平上预测药物反应,该方法基于单细胞 RNA 测序数据,使用 GF-ICF (gene frequency-inverse cell frequency) 归一化获取每个细胞的最相关基因集,然后计算该基因集在 GPDS 排序列表中的富集得分,并用来衡量每个细胞的药物敏感性。值得注意的是,Beyondcell 和 DREEP 都是基于药物敏感性的基因标志物的方法。

表 3 模型在 8 个数据集上的性能结果

Tab. 3 Performance results of models on eight datasets.

序号	AUC					AUPR				
	SCTL	scDEAL	SCAD	Beyondcell	DREEP	SCTL	scDEAL	SCAD	Beyondcell	DREEP
Data1	0.928	0.823	0.809	0.386	0.497	0.966	0.901	0.885	0.595	0.657
Data2	0.996	0.880	0.936	0.698	0.590	0.961	0.850	0.498	0.141	0.088
Data3	0.916	0.950	0.651	0.488	0.500	0.861	0.913	0.749	0.500	0.500
Data4	0.944	0.947	0.918	0.269	0.521	0.931	0.965	0.856	0.382	0.528
Data5	0.916	0.648	0.762	0.344	0.500	0.859	0.650	0.633	0.330	0.398
Data6	0.964	0.904	0.967	0.977	0.788	0.966	0.890	0.973	0.982	0.788
Data7	0.952	0.613	0.942	0.731	0.667	0.967	0.552	0.958	0.642	0.600
Data8	0.962	0.600	0.968	0.733	0.621	0.954	0.616	0.970	0.663	0.569

表 3 显示了本文模型 SCTL 与四个基准方法在 8 个数据集上关于 AUC 和 AUPR 两个预测性能指标的结果。SCTL 在 4 个数据集上取得了最优的预测性能,在其他 4 个数据集上的预测性能排名第二,而且仅略低于最优预测性能。

与 4 种基线方法相比, SCTL 在 8 个数据集上的

平均 AUC 为 0.949, 平均 AUPR 为 0.933 (图 2)。关于平均 AUC, SCTL 比 SCAD 提高了 9.21%, 比 scDEAL 提高了 19.22%。关于平均 AUPR, SCTL 比 SCAD 提高了 0.145%, 比 scDEAL 提高了 0.178%。另外, 深度学习模型 SCTL、scDEAL 以及 SCAD 比 Beyondcell 和 DREEP 这两种基于药物敏感性的基因

标志物的方法的性能都要好,这也反映出了深度学习在提取深层次有效特征方面的能力。

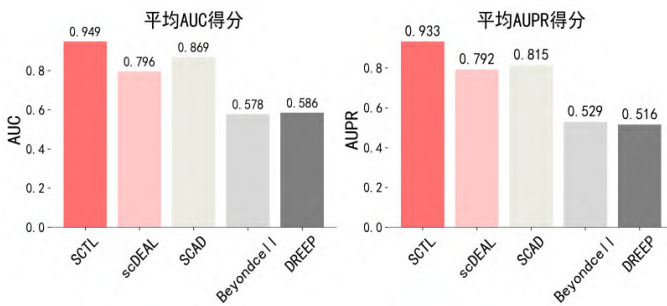


图 2 5 种方法的平均 AUC 值和平均 AUPR 值
Fig. 2 Average AUC and average AUPR results of five methods.

3.2 消融实验

为了评估模型的各个功能模块对预测性能的影响,在 8 个数据集上进行了详细的消融实验分析。

针对特征对齐的计算环节,分别移除最优传输的特征对齐和最大均值差异的特征对齐计算,构建了两个对比模型。保持每个数据集的最优参数组合不变,然后重新计算了两个对比模型在所有数据集上的 AUC 和 AUPR。详细结果见表 4。当把最大均值差异的特征对齐计算模块移除时,其平均 AUC 比原先模型下降了 0.038,平均 AUPR 下降了 0.036。当移除最优传输特征对齐计算时,模型的预测性能出现了明显的下降,平均 AUC 比原先模型下降了 0.264,平均 AUPR 下降了 0.323。这指示着在特征对齐的计算里,最优传输计算起了重要的影响。

表 4 特征对齐的消融实验结果

Tab. 4 Ablation experiment results for feature alignment.

数据	SCTL		SCTL_1		SCTL_2	
	AUC	AUPR	AUC	AUPR	AUC	AUPR
Data1	0.928	0.966	0.846	0.911	0.635	0.742

表 6 不同药物反应预测器的性能对比结果

Tab. 6 Performance comparison results for different drug response predictors.

序号	AUC					AUPR				
	predictor1	predictor2	predictor3	predictor4	predictor5	predictor1	predictor2	predictor3	predictor4	predictor5
Data1	0.711	0.797	0.786	0.928	0.907	0.875	0.908	0.882	0.966	0.958
Data2	0.966	0.977	0.989	0.996	0.917	0.936	0.940	0.941	0.961	0.619
Data3	0.878	0.889	0.891	0.916	0.862	0.817	0.831	0.852	0.861	0.820
Data4	0.935	0.937	0.943	0.944	0.938	0.921	0.918	0.931	0.931	0.917
Data5	0.805	0.832	0.801	0.916	0.908	0.821	0.771	0.816	0.859	0.894
Data6	0.879	0.907	0.895	0.964	0.884	0.915	0.911	0.919	0.966	0.898
Data7	0.906	0.919	0.924	0.965	0.776	0.925	0.932	0.939	0.967	0.852
Data8	0.919	0.921	0.937	0.962	0.896	0.912	0.914	0.926	0.954	0.881

Data2	0.996	0.961	0.886	0.828	0.828	0.328
Data3	0.916	0.861	0.913	0.86	0.708	0.737
Data4	0.944	0.931	0.932	0.926	0.517	0.516
Data5	0.916	0.859	0.847	0.841	0.459	0.372
Data6	0.964	0.966	0.943	0.942	0.64	0.674
Data7	0.952	0.966	0.895	0.913	0.845	0.826
Data8	0.962	0.967	0.889	0.896	0.617	0.62
Mean	0.949	0.933	0.894	0.89	0.656	0.603
Std	0.0281	0.0466	0.0356	0.0417	0.1354	0.1808

注: SCTL_1 表示去除最大均值差异; SCTL_2 表示去除最优传输。

针对药物反应预测器的模型结构,对多层感知机的隐藏层层数进行了详细的消融分析。构建了 6 个网络层深度各不相同的药物反应预测器,分别是具有 2 层、3 层、4 层、5 层和 6 层隐藏层层数的神经网络。设置每层隐藏层的节点数为 128 或 64,激活函数为 ReLU。详细的药物反应预测器的网络结构设置如表 5 所示。

表 5 药物反应预测器的网络结构

Tab. 5 Network architecture of drug response predictor.

网络	网络层及节点数	激活函数
predictor1	128,64	ReLU
predictor2	128,128,64	ReLU
predictor3	128,128,128,64	ReLU
predictor4	128,128,128,128,64	ReLU
Predictor5	128,128,128,128,128,64	ReLU

在 8 个数据集上分别对 5 个药物反应预测器进行了实验,结果见表 6。初始预测器网络层数为 2,随着预测器隐藏层层数增加,AUC 和 AUPR 得分增加,在隐藏层深度增加到 5 时,模型整理的性能表现最好;当多层感知机深度继续增加到 6 层时,模型在大部分数据集上的性能出现了明显下降。

在跨域特征提取部分,对细胞系混合测序数据有 3 种数据采样方式,分别为不采样、加权随机采样、合成少数类过采样 (Synthetic Minority Oversampling Technique, SMOTE)。加权随机采样是根据数据中正负样本数的比例,对占比小的数据,提高其被采样的概率。SMOTE 是通过在特征空间中合成新的少数类样本平衡数据集,从而提高分类模型对少数类的识别能力^[19]。在 8 个数据集上,分别展示了 3 种数据采样方式对结果的影响,见表 7。

表 7 3 种采样的性能对比结果

Tab. 7 Performance comparison results for different sampling methods.

数据	不采样		加权随机采样		SMOTE	
	AUC	AUPR	AUC	AUPR	AUC	AUPR
Data1	0.865	0.941	0.928	0.966	0.828	0.895
Data2	0.996	0.961	0.932	0.458	0.897	0.430
Data3	0.916	0.861	0.710	0.738	0.719	0.743
Data4	0.944	0.931	0.927	0.884	0.928	0.891
Data5	0.851	0.846	0.828	0.767	0.916	0.859
Data6	0.657	0.640	0.964	0.966	0.740	0.752
Data7	0.582	0.587	0.965	0.967	0.839	0.900
Data8	0.541	0.571	0.952	0.932	0.962	0.954

此外,还对 8 个数据集的源域(混合 RNA 测序数据)的正负样本数进行了比率计算,数据 1 源域正负样本数比值为 0.618、数据 2 为 0.602、数据 3 为 0.624、数据 4 为 0.768、数据 5 为 0.662、数据 6 为 0.161、数据 7 为 0.078、数据 8 为 0.099。按照正负样本比率和 3 种采样方式的 AUC、AUPR 结果进行了可视化分析,见图 3。当源域数据正负样本处于基本平衡时,不进行采样就可以获得较好的性能;而当正负样本数差别很大时,如数据 6、7、8,使用采样的方法可以很好地提高模型的性能。

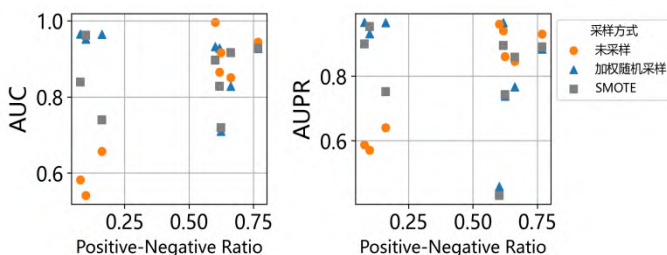


图 3 正负样本比率和采样方法的关系

Fig. 3 Relationship between positive and negative sample ratio and sampling method.

3.3 鲁棒性实验

为了说明模型预测结果的稳定性,还做了鲁棒性实验。随机的环境种子会影响数据集批次的划分。固定模型在每个数据集上训练的最优参数,然后在随机的环境种子下重复实验 20 次,结果见图 4。模型在 8 个数据集的波动性都很小,AUC 和 AUPR 的平均标准差分别为 0.0504 和 0.0482,这表明模型具有很好的稳定性。

分别对每个单细胞 RNA 测序数据进行 80% 的随机采样,然后对采样得到的单细胞数据进行模型训练预测,并重复实验 20 次得到性能结果。模型每一次实验都会不一样的目标域数据上进行跨域特征对齐和药物反应预测,结果见图 5。模型在 8 个单细胞 RNA 测序数据上的结果均保持了良好的稳定性,这说明模型对不同的目标域数据均进行了有效的特征学习和特征对齐。

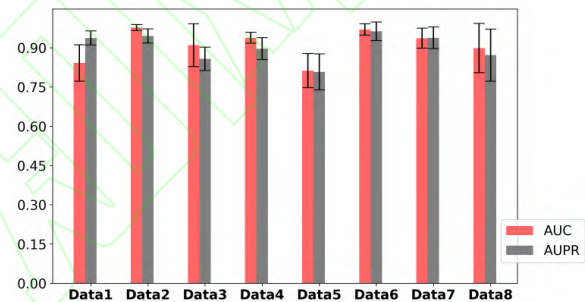


图 4 不固定随机种子重复 20 次的鲁棒性实验

Fig. 4 Robustness experiment repeated 20 times with unfixed random seed.

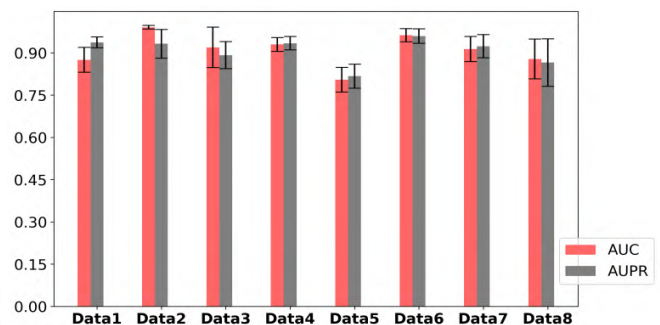


图 5 随机采样 80% 单细胞 RNA 测序数据重复 20 次的鲁棒性实验

Fig. 5 Robustness experiment repeated 20 times with 80% random sampling of single-cell RNA-seq data.

3.4 模型可解释性的案例分析

以 Data6 为例子,其是关于头颈部鳞状细胞癌在吉非替尼药物作用下的数据。该单细胞数据来自头颈部鳞状细胞癌模型细胞系^[20-21],由两簇细胞组成: JHU006_Gefitinib_Sen 和 JHU006_Gefitinib_Res。其

中, JHU006_Gefitinib_Sen 细胞簇对吉非替尼敏感, JHU006_Gefitinib_Res 细胞簇对吉非替尼耐药。

模型训练完成后输出细胞系和单细胞间映射匹配的最优传输矩阵 T^* 。使用热图可视化细胞系和单细胞间的最优传输矩阵值 (见图 6), 为了方便展示, 这里分别随机选择了 90 个敏感细胞系以及 90 个耐药细胞系。在热图中, 可以观察到对吉非替尼敏感的细胞与对吉非替尼敏感的细胞系之间的匹配关系更紧密, 类似的, 对吉非替尼耐药细胞与对吉非替尼耐药细胞系之间的联系也更强。这表明在特征提取和特征对齐的环节, 模型获得了细胞系和单细胞间关于药物反应分层的准确的匹配关系, 从而模型可以获得准确的单细胞药物反应预测结果。

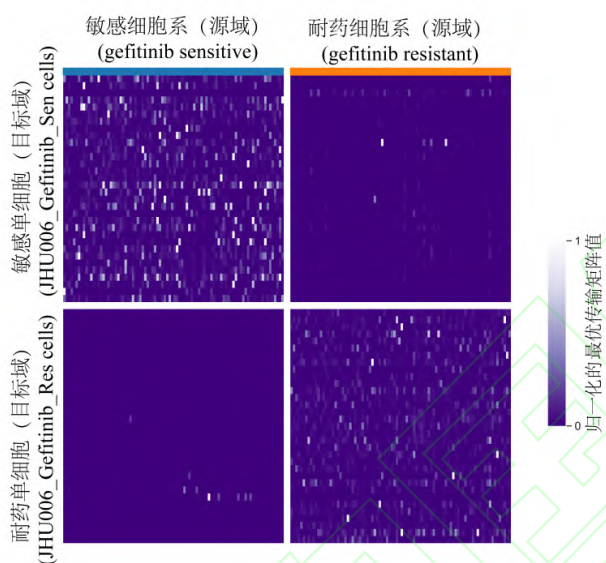


图 6 数据 6 中细胞系和单细胞间最优传输矩阵的热图

Fig. 6 Heatmap of optimal matrix between source domain and target domain for data6.

分析在模型输出药物响应预测值的过程中起到关键作用的基因的影响。对模型输出的药物响应预测连续值二值化, 阈值选择为这组预测连续值的中位数, 将单细胞划分为预测敏感细胞簇 (值为 1) 和预测耐药细胞簇 (值为 0), 成功预测分类了 90.9% 的细胞的药物响应状态。接着, 对两簇细胞采用 Wilcoxon 秩和检验。在两簇细胞上分别筛选了 top 10 个敏感高表达基因和 top 10 个耐药高表达基因, 并在预测敏感细胞簇和预测耐药细胞簇上对这些基因的表达值用热图可视化, 以及在真实敏感和真实耐药细胞簇上对这些基因的表达值可视化, 见图 7。可以明显看到, 筛选得到的基因在模型预测的细胞簇上有明显的表达差异, 也在真实的细胞簇上有明显

的表达差异。这些高表达基因在模型预测细胞药物响应状态过程中, 可被认为起到了关键作用。总的来说, 本模型在预测任务上取得优异性能的同时, 还在预测的细胞簇上筛选到高表达基因, 这些基因作为癌细胞的药物敏感标志物和耐药标志物, 对后续的治疗分析起到了潜在支持。

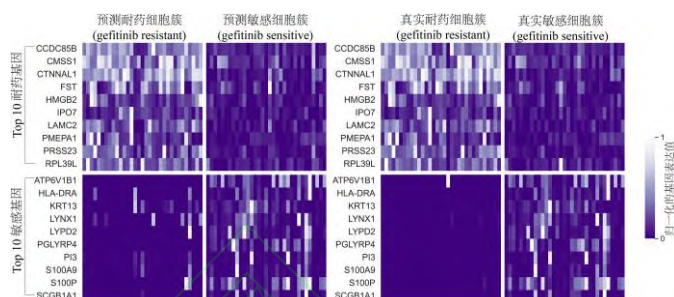


图 7 数据 6 中 top 10 高表达基因在预测细胞簇和真实细胞簇上的基因表达热图

Fig. 7 Heatmap of gene expression for the top 10 highly expressed genes in data6 on predicted cell clusters and actual cell clusters.

4 结语

本研究提出了一种用于单细胞抗癌药物反应预测的迁移学习模型 SCTL, 其利用现有的大规模细胞系混合 RNA 测序数据, 将细胞系的药物基因组信息迁移到单细胞 RNA 测序数据上。通过与基准方法对比, 验证了 SCTL 具有较高的预测性能。消融实验的结果展示了不同模块对于模型准确预测的有效性。鲁棒性实验显示了模型预测结果在不同的环境随机种子以及单细胞数据随机采样等情况下具有稳定性。

针对单细胞药物响应预测研究领域, 本研究提出使用迁移学习计算方法, 有效地弥合了单细胞测序数据和细胞系测序数据之间的差异。然而, 的方法也存在着局限性。一个主要挑战是如何在不损失特征信息的情况下, 实现源域与目标域数据特征的对齐和信息的跨域迁移。的方法由于使用共享编码器学习源域和目标域特征, 这一过程可能会损失不同数据域的特异性信息, 继而在特征对齐的环节损害特征信息的迁移。本研究中, 虽然通过调整模型超参数在不同数据集上表现较好的性能, 但模型的泛化能力仍有待提高。为了应对这些问题, 计划在接下来的研究中引入注意力机制, 以增强模型对源域和目标域数据特征的学习能力, 提高模型的泛化性和预测性能。

本文中用到的单细胞数据主要集中在临床前模型上, 缺乏临床水平的患者验证。随着临床单细胞组学数据的日益广泛, 在单细胞水平上预测药物反应的 SCTL 模型可以更深入地揭示细胞间的异质性和功能差异, 为个性化和精确的癌症治疗提供一种新的策略。

参考文献:

- [1] CHEN J, WANG X, MA A, et al. Deep transfer learning of cancer drug responses by integrating bulk and single-cell RNA-seq data[J]. *Nature communications*, 2022, 13(1): 6494.
- [2] CARREIRA S, ROMANEL A, GOODALL J, et al. Tumor clone dynamics in lethal prostate cancer[J]. *Science translational medicine*, 2014, 6(254): 254ra125.
- [3] WONG C H, SIAH K W, LO A W. Estimation of clinical trial success rates and related parameters[J]. *Biostatistics*, 2019, 20(2): 273-286.
- [4] RAMBOW F, ROGIERS A, MARIN-BEJAR O, et al. Toward minimal residual disease-directed therapy in melanoma[J]. *Cell*, 2018, 174(4): 843-855.
- [5] VERJANS E T, DOIJEN J, LUYTEN W, et al. Three-dimensional cell culture models for anticancer drug screening: worth the effort?[J]. *Journal of cellular physiology*, 2018, 233(4): 2993-3003.
- [6] SCHIRLE M, JENKINS J L. Identifying compound efficacy targets in phenotypic drug discovery[J]. *Drug discovery today*, 2016, 21(1): 82-89.
- [7] ZHENG Z, CHEN J, CHEN X, et al. Enabling single-cell drug response annotations from bulk RNA-Seq using SCAD[J]. *Advanced science*, 2023, 10(11): e2204113.
- [8] YANG W, SOARES J, GRENINGER P, et al. Genomics of drug sensitivity in cancer (GDSC): a resource for therapeutic biomarker discovery in cancer cells[J]. *Nucleic acids research*, 2013, 41(D1): 955-961.
- [9] BARRETINA J, CAPONIGRO G, STRANSKY N, et al. The cancer cell line encyclopedia enables predictive modelling of anticancer drug sensitivity[J]. *Nature*, 2012, 483(7391): 603-607.
- [10] WOLF F A, ANGERER P, THEIS F J. SCANPY: large-scale single-cell gene expression data analysis[J]. *Genome biology*, 2018, 19(1): 15.
- [11] CAO K, GONG Q, HONG Y, et al. A unified computational framework for single-cell data integration with optimal transport[J]. *Nature communications*, 2022, 13(1): 7419.
- [12] ZHANG W, WU D. Discriminative joint probability maximum mean discrepancy (DJP-MMD) for domain adaptation[C]//IEEE. 2020 international joint conference on neural networks (IJCNN). New York: IEEE, 2020: 1-8.
- [13] CUTURI M. Sinkhorn distances: lightspeed computation of optimal transport[J]. *Advances in neural information processing systems*, 2013, 26: 15966283.
- [14] ARJOVSKY M, CHINTALA S, BOTTOU L. Wasserstein generative adversarial networks[C]//PMLR. International conference on machine learning. New York: PMLR, 2017: 214-223.
- [15] XIE Y, WANG X, WANG R, et al. A fast proximal point method for computing exact Wasserstein distance[C]//PMLR. Uncertainty in artificial intelligence. New York: PMLR, 2020: 433-453.
- [16] FUSTERO-TORRE C, JIMÉNEZ-SANTOS M J, GARCÍA-MARTÍN S, et al. Beyondcell: targeting cancer therapeutic heterogeneity in single-cell RNA-seq data[J]. *Genome medicine*, 2021, 13(1): 187.
- [17] PELLECCIA S, VISCIDO G, FRANCHINI M, et al. Predicting drug response from single-cell expression profiles of tumours[J]. *BMC Medicine*, 2023, 21(1): 476.
- [18] GOODFELLOW I, POUGET-ABADIE J, MIRZA M, et al. Generative adversarial networks[J]. *Communications of the ACM*, 2020, 63(11): 139-144.
- [19] CHAWLA N V, BOWYER K W, HALL L O, et al. SMOTE: synthetic minority over-sampling technique[J]. *Journal of artificial intelligence research*, 2002, 16: 321-357.
- [20] KINKER G S, GREENWALD A C, TAL R, et al. Pan-cancer single-cell RNA-seq identifies recurring programs of cellular heterogeneity[J]. *Nature genetics*, 2020, 52(1): 1208-1218.
- [21] SINHA S, VEGESNA R, MUKHERJEE S, et al. PERCEPTION predicts patient response and resistance to treatment using single-cell transcriptomics of their tumors[J]. *Nature cancer*, 2024, 5(6): 938-952.