

doi:10.19920/j.cnki.jmsc.2026.07.009

大样本中识别标杆：社会网络数据包络分析^①

安庆贤^{1,2}, 王 萍³, 文 瑶^{4*}

(1. 中南大学商学院, 长沙 410083; 2. 合肥工业大学经济学院, 合肥 230601;
3. 合肥工业大学管理学院, 合肥 230009; 4. 福州大学经济与管理学院, 福州 350108)

摘要: 为传播积极的意识形态、体现良好的社会风气, 各行各业推出一些行业模范人物(即标杆)的评选活动, 如“最美医生”、“最美教师”等。这些模范人物(即标杆)往往是大量候选人中的少数, 然而传统方法难以对大样本中数以万计的候选人进行绩效评价, 且尚未考虑候选人之间的学习关系和具体优势, 较难实现大样本中标杆的识别。为减少标杆识别的工作量, 本研究结合数据包络分析(data envelopment analysis, DEA)和社会网络分析(social network analysis, SNA), 提出一种社会网络数据包络分析方法, 用于从大样本中识别标杆。首先通过 DEA 方法构建一个两两评价过程, 以探索每个决策单元(decision making unit, DMU)相对于另一个 DMU 的效率状态, 得到一个能反映任意两个 DMU 之间参考关系的两两评价矩阵。接着, 基于该矩阵建立社会网络, 并结合社会网络分析中的入度中心性指标来识别大样本中的标杆 DMU。然后, 根据 DEA 模型中权重的含义, 分析标杆 DMU 的具体优势。最后, 利用春雨医生平台上 55 个科室中 10 418 名医生的在线诊断数据进行实验, 对本研究提出的社会网络数据包络分析方法进行验证。

关键词: 数据包络分析; 大样本; 标杆; 社会网络分析; 两两评价

中图分类号: C934 **文献标识码:** A **文章编号:** 1007-9807(2026)07-0125-16

0 引 言

“最美医生”、“最美教师”等楷模, 往往具有很多值得推广和学习的优点, 是多数人学习的榜样, 即标杆。这些标杆往往是多数人中的少数, 要科学客观地识别出来是比较困难的。根据《中国统计年鉴》, 2020 年我国有超过 1 793.7 万名教师和 408.5 万名医生。如何从成千上万人员中选出教师或医生界的标杆? 这更是一个极具挑战性的问题。为回答该问题, 本研究提供了一种从大样本中识别标杆的新方法。

社会个体成员之间相互作用形成的稳定的关系系统称为社会网络, 其中每个节点代表一个组

织或实体, 每两个节点之间的边表示他们的学习关系。从一组样本中识别标杆与从社会网络中识别重要节点的思路是一致的。社会网络分析(social network analysis, SNA)^[1,2] 是识别社会网络中重要节点的经典方法, 具体包括三类社会网络中心性^[3]: 1) 基于邻居节点的中心性是一种较简单的确定重要节点的办法, 反映了与节点相连的邻居个数。如度中心性、局部排名、聚类排名、核中心性、H 指数都是基于邻居节点的中心性。其中, 度中心性指一个节点邻居的数量, 邻居越多该节点越重要^[1]。2) 基于路径的中心性描述了节点控制信息流的能力, 即, 有些节点的度虽然很小, 但却发挥了重要作用, 如离心中心性^[4]、Katz 中心

① 收稿日期: 2022-05-20; 修订日期: 2024-06-06。

基金项目: 国家自然科学基金资助项目(72171238; 72188101; 72401067); 国家资助博士后研究人员计划项目(GZC20250532)。

通讯作者: 文 瑶(1992—), 女, 四川眉山人, 博士, 讲师, 硕士生导师。Email: wenyaphd@163.com

性^[5]、中介中心性^[6]。3) 基于向量的中心性也同时考虑了邻居节点的质量^[7],如累计提名法^[8],PageRank 算法^[9],LeaderRank 算法^[10],Hits 算法。以上确定社会网络重要节点的方法主要是基于节点特征和路径的,这些重要节点主要有以下特征:1) 邻居节点较多;2) 连接的节点质量较高;3) 处于重要的位置。不同的中心性测量方法适用于不同的社会网络。本研究所构建的社会网络是一种考虑每两个样本之间学习关系的网络,任何表现差的样本都可以向表现好的样本学习,与样本所处的位置和邻居无关,因此本研究选择第 1) 类中的度中心性来识别标杆。此外,以往文献中,鲜有研究能客观地给出每两个样本之间的学习关系。要客观确定两个样本之间的学习关系,就需要对两者进行客观评价。现实中,对医生或教师的评价都需综合多个方面的,而非单一的评价指标,这便为任意两个样本之间学习关系的识别增加了困难。

为评价一组具有多投入多产出指标的决策单元(decision making unit, DMU)的相对绩效,Charnes 等^[11]提出了数据包络分析(data envelopment analysis, DEA)方法。DEA 是一种数据驱动的非参数评价方法^[12],不需要定义生产函数就可以提供客观的效率评价结果。通过 DEA 方法得到的 DMU 效率值反映的是被评价 DMU 相对于有效 DMU 所构成的生产前沿的距离^[13],是一种相对效率^[14]。此外,DEA 方法还可以为效率较低的决策单元提供改进方案^[15],识别标杆 DMU^[16],进而揭示决策单元之间的参考关系。在所有待评价 DMU 构成的生产可能集中,一个 DMU 以标杆 DMU 为参考对象,向标杆 DMU 学习以提高自身效率^[17],因此 DMU 之间的参考关系也可以看作 DMU 之间的学习关系。由此可见,DEA 方法可用于识别 DMU 之间的学习关系。

然而,大样本中有成千上万,甚至上百万个 DMU,传统的 DEA 模型难以实现其效率评价和标杆识别。传统 DEA 模型以最大化被评价 DMU 的效率为目标函数,以所有待评价 DMU 效率值小于等于 1 为约束,因此每运行一次只能得到一个被评价 DMU 的相对效率值及其所要参考的标杆 DMU 和改进方案,使得所有 DMU 效率值的计算

时间会随样本量的增加而急剧增加^[18]。已有不少学者将 DEA 和大样本结合^[19],针对 5 000 及以上 DMU 效率评价问题,提出大规模或大样本 DEA,力求减缓传统 DEA 模型在大样本评价时模型计算的难度,缩短计算时间。起初,Ali^[20]提出了一种称为“the restricted basis entry”的方法,通过删除 DEA 模型中“强度”系数为 0 的 DMU 来减小 DEA 模型的规模。此后,许多研究者以减少 DEA 模型的规模为目标,对大样本的相对效率评价问题展开研究。比如,Barr 和 Durchholz^[21]提出了层次分解过程,将待评价的 DMU 分为几组,然后用所有组的有效 DMU 来评价其他 DMU。后来,Zhu 等^[22]和 Jie^[23]对该分解过程进行了一些改进和应用。此外,另一种思路是识别构成生产前沿面的最少 DMU。首先,Dulá^[24]提出了“build-hull”算法,通过 DEA 模型对少量 DMU 进行评价以获得初始的“hull”,然后基于该“hull”,对剩余的 DMU 逐个进行有效和非有效的判断,不断更新“hull”并最终得到生产前沿面。其次,Khezrimotlagh 等^[18]在以往研究的基础上,提出了一个先验程序,用于识别构成生产前沿面的最少 DMU,使得计算时间远小于层次分解过程和“build-hull”算法。除以上经典的算法外,Chen 和 Lai^[25]以及 Dellnitz^[26]也对大样本的评价提供了思路。根据这些研究,本研究将包含 5 000 及以上 DMU 的样本界定为大样本。虽然以上研究中学者们提出了各种方法来减缓传统 DEA 模型对大样本进行效率评价的难度,不同程度地减少了 DEA 模型的计算时间,但是这些方法只提供了被评价 DMU 参考整个生产前沿面的相对效率值,尚未考虑任意 DMU 之间的学习关系。综上,已有的 DEA 方法较难识别大样本中任意两个 DMU 之间的学习关系。

目前已有较多研究者考虑到 DMU 之间的参考关系,针对传统 DEA 模型无法区分多个有效 DMU 相对效率这一不足,将 SNA 与 DEA 结合对 DMU 进行评价,以提高评价结果的区分度。比如,Liu 等^[27]提出了一种基于 SNA 的方法,首先用 DEA 模型用于计算决策单元的效率值,然后使用“强度”变量值来构建社会网络的邻接矩阵,最后在构建加权有向网络的基础上使用特征向量中心

性实现对有效 DMU 的排序. 随后, Liu 和 Lu^[28] 对 Liu 等^[27] 的方法进行了改进, 对“强度”变量值进行归一化, 并给出了一种基于 SNA 来识别 DMU 优缺点的方法. De Blas 等^[29] 结合 SNA、DEA 及加权 HITS 算法, 提出了一种对有效 DMU 进行排序的方法. 此外, Aydın 等^[30] 根据 DEA 模型中“强度”变量的值构建社会网络, 并使用特征向量中心性值作为超效率的权重, 结合 SNA 与超效率 DEA 模型来实现对有效 DMU 的排序, 并确定 DMU 在各维度上的优势和劣势. 钟韬和彭勤科^[31] 建立了金融市场的社会网络模型, 并利用其中的连通补网进行投资组合选择. Ang 等^[32] 结合 DEA 和 SNA 为不同组的 DMU 选择合适的标杆, 并结合 HITS 算法对所有 DMU 进行排序. 以上这些研究多基于 DEA 模型中的“强度”变量构建社会网络, 主要目的是提高 DEA 评价结果的区分度, 较少涉及大样本中标杆识别的问题. 如上所述, 在社会网络中, 识别标杆的首要任务是确定任意两个 DMU 之间的参考关系. Kang 等^[33] 提出一种新的算法, 该算法基于自助抽样算法生成可以反映 DMU 之间原始参考关系的社会网络和标杆, 进而对 DMU 进行排序, 增加了排名的区分度. 以往研究者结合 DEA 与 SNA 方法为有效 DMU 的排序提供了思路, 可以为有效 DMU 和无效 DMU 之间的参考关系提供依据. 然而, 一方面, 这些研究需要用 DEA 模型依次对每一个 DMU 相对效率进行评价, 若用于大样本 DMU 效率评价将因模型复杂度和模型数量较多而耗时很长; 另一方面, 以往 DEA 与 SNA 的结合研究多集中于判断 DMU 是否有效或对多个有效 DMU 进行排序, 尚未考虑任意两个 DMU 之间的参考关系, 难以判断多个无效 DMU 或多个有效 DMU 之间的学习关系.

另外, 如果直接用传统的 DEA 模型对任意两个 DMU 的相对效率进行计算, 进而生成任意两个 DMU 之间的参考关系, 从理论上是可以实现的, 但难以适用于大样本. 这是因为当样本中包含 n 个 DMU, 则需构建 C_n^2 个 DEA 模型并运行 $2C_n^2$ 次, 于是当样本量很大时, 这样的评价过程是非常耗时甚至无法实现的. 基于上述分析, 现有 DEA 方法, 包括 SNA 和 DEA 结合的方法, 难以判断任意两个 DMU 之间的参考关系, 因此尚无法用于大

样本中标杆识别问题.

为了从大样本中识别标杆, 本研究基于 SNA 和 DEA 方法, 提出一种社会网络 DEA 方法. 具体来说, 首先基于 DEA 有效性与多目标规划的 Pareto 最优解的等价性^[34], 构建一个两两评价过程. 通过简单的数学表达式对任意两个 DMU 的相对效率状态进行评价, 使得在识别两两 DMU 之间参考关系时不需要计算大规模线性或非线性规划模型, 有助于节省计算时间. 基于该过程, 可得到样本中所有两两 DMU 之间的比较结果, 记为两两评价矩阵. 该矩阵可揭示任意两个 DMU 之间的学习关系. 然后, 基于该矩阵构建一个有向社会网络. 通过计算该社会网络中所有 DMU 的入度中心性, 找到大样本中的标杆 DMU. 最后, 通过评价标杆 DMU 相对于向他学习的 DMU 的效率, 以 DEA 模型中投入产出指标的权重值为依据, 客观分析他们被选为标杆的原因, 剖析标杆 DMU 的优势指标.

与现有研究相比, 本研究的贡献主要有三方面, 一是新提出的社会网络 DEA 方法可识别大样本中任意两个 DMU 之间的学习关系, 并从客观的角度识别其中值得大家学习的标杆; 二是本研究提供的方法不要求解线性规划或非线性规划, 节省了大量计算时间, 可较快识别大样本中的标杆 DMU; 三是本研究的方法不仅可以识别大样本中的标杆, 还可以分析被选为标杆的原因, 即从客观的角度明晰标杆的优势.

1 基本理论与方法介绍

1.1 DEA 模型

假设一共有 n 个 DMU 属于集合 N , 其中每个 DMU 都使用 m 个投入, 生成 s 个产出. 令 $X_j = (x_{1j}, \dots, x_{mj})^T$ 和 $Y_j = (y_{1j}, \dots, y_{sj})^T$ 为 DMU _{j} 的投入产出向量. 下标 j 表示第 j 个 DMU. 根据 Charnes 等^[11], DMU _{k} 的效率值可由下面的模型计算得到

$$\begin{aligned} \max E_k &= \sum_{r=1}^s u_r y_{rk} \\ \text{s. t. } \sum_{r=1}^s u_r y_{rj} - \sum_{i=1}^m v_i x_{ij} &\leq 0, j = 1, 2, \dots, n \end{aligned} \quad (1)$$

$$\sum_{i=1}^m v_i x_{ik} = 1, v_i, u_r \geq 0,$$

$$\forall i = 1, \dots, m; r = 1, \dots, s,$$

其中 v_i, u_r 分别表示第 i 个投入和第 r 个产出的权重. 通过计算模型 (1), 可得到一个被评价 DMU (即 DMU_k) 的效率值, 而 $k = 1, 2, \dots, n$, 所以为了得到所有 DMU 的效率值, 需构建和计算 n 个类似于模型 (1) 的 DEA 模型, 每个模型中包含了 $n+1$ 个约束条件, $m+s$ 个决策变量.

模型 (1) 称为乘数形式的 DEA 模型, 其对偶形式如下

$$\begin{aligned} \min \quad & \theta_k \\ \text{s. t.} \quad & \sum_{j=1}^n \lambda_j x_{ij} \leq \theta_k x_{ik}, i = 1, \dots, m \\ & \sum_{j=1}^n \lambda_j y_{rj} \geq y_{r0}, r = 1, \dots, s \\ & \lambda_j \geq 0, j = 1, \dots, n \end{aligned} \quad (2)$$

模型 (2) 称为包络形式的 DEA 模型, 其中 $\lambda_j, j = 1, \dots, n$ 为“结构”或“强度”系数, 表示 n 个 DMU 的线性组合. 通过模型 (1) 或模型 (2) 的计算, 可得到被评价 DMU 的效率值. 若 $E_k = \sum_{r=1}^s u_r y_{rk} = \theta_k = 1$, 则 DMU_k 为有效, 否则为非有效. 在 n 个待评价 DMU 中, 无效 DMU 可参考有效 DMU, 通过减少投入或增加产出来提高自身效率.

1.2 社会网络概念

社会网络指的是社会个体之间相互作用形成的稳定关系系统. 社会网络分析 (social network analysis, SNA) 主要关注的是个体之间的互动和联系. 社会网络图由节点和连线构成, 节点可以表示

个体或组织, 连线表示各种各样的关系, 如学习关系、朋友关系、商业联系、信任关系、引文关系等.

社会网络通常由图或邻接矩阵表示, 见表 1. 社会网络图包括节点和连接, 通常记为 $G(E, \Lambda)$, 其中 Λ 表示节点集合, $\Lambda = \{\Lambda_1, \dots, \Lambda_i, \dots, \Lambda_n\}$, E 表示连接的集合, $(\Lambda_i, \Lambda_j) \in E$ 表示节点 Λ_i 和 Λ_j 间存在连接. 社会网络图分为有向图和无向图, 如表 1 所示, 有向图和无向图定义如下.

定义 1 在有向图中, 所有连接都是有向的.

定义 2 在无向图中, 所有连接都是无向的.

邻接矩阵将节点间的连接关系表示成一个二维向量. 对于给定的有向图或无向图 $G(E, \Lambda)$, 它的邻接矩阵用 $C = (c_{ij})_{n \times n}$ 表示, 其中元素 $c_{ij} = 1$ 表示节点 Λ_i 和 Λ_j 间存在连接, $c_{ij} = 0$ 表示节点 i 和 j 之间没有连接. 基于此, 邻接矩阵可以分为二值邻接矩阵和加权邻接矩阵, 定义如下.

定义 3 图网络 $G(E, \Lambda)$ 的二值矩阵为 $C = (c_{ij})_{n \times n}$, 其中 $c_{ij} \in \{0, 1\}$, $c_{ij} = 1$ 表示节点 Λ_i 和 Λ_j 间有联系, $c_{ij} = 0$ 表示节点 Λ_i 和 Λ_j 间没有联系.

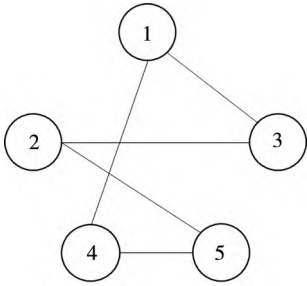
在某些情况下, 二值邻接矩阵可能会丢失信息, 因为二值只能反映是否存在联系, 不能反映联系的强弱, 为了避免该现象, 学者们又提出了加权邻接矩阵.

定义 4 网络图 $G(E, \Lambda)$ 的加权邻接矩阵为 $C = (c_{ij})_{n \times n}$, 其中 c_{ij} 为任意值, 表示节点 Λ_i 和 Λ_j 间联系的强度.

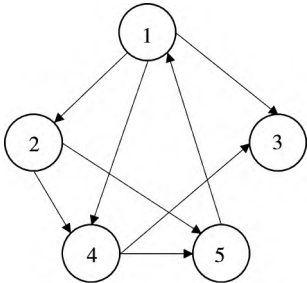
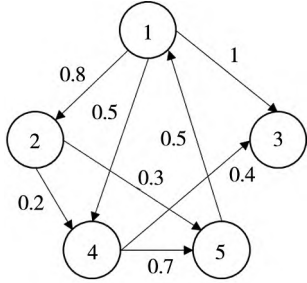
表 1 展示了不同种类的网络图和对应的邻接矩阵, 包括无向图、有向图、加权有向图、二值邻接矩阵和加权邻接矩阵.

表 1 社会网络的表达

Table 1 The expressions of social network

社会网络图	邻接矩阵
 无向图	$C_1 = \begin{bmatrix} 0 & 0 & 1 & 1 & 0 \\ 0 & 0 & 1 & 0 & 1 \\ 1 & 1 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 1 \\ 0 & 1 & 0 & 1 & 0 \end{bmatrix}$

续表 1
Table 1 Continues

社会网络图	邻接矩阵
 有向图	$C_2 = \begin{bmatrix} 0 & 1 & 1 & 1 & 0 \\ 0 & 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 1 \\ 1 & 0 & 0 & 0 & 0 \end{bmatrix}$
 加权有向图	$C_3 = \begin{bmatrix} 0.0 & 0.8 & 1.0 & 0.5 & 0.0 \\ 0.0 & 0.0 & 0.0 & 0.2 & 0.3 \\ 0.0 & 0.0 & 0.0 & 0.0 & 0.0 \\ 0.0 & 0.0 & 0.4 & 0.0 & 0.7 \\ 0.5 & 0.0 & 0.0 & 0.0 & 0.0 \end{bmatrix}$

1.3 度中心性

在社会网络中,衡量节点的重要性是一个非常
重要的研究话题,即确定哪个节点最关键、最
具有影响力.在本研究所构建的社会网络中,节
点的度中心性越高,向该节点学习的邻居节点
越多,该节点越重要^[1].衡量节点重要性的方
法有多种,如向量中心性、中介中心性、度中
心性,在有向网络图中,节点的度中心性又可
分为入度中心性和出度中心性^[35].以下分
别介绍度中心性、入度中心性和出度中心性
的定义.

定义 5 度中心性 在无向图中,度中心性
是节点连边(权重)的加和.令 d_i 表示节点 i
的度,则

$$d_i = \sum_{j=1}^n c_{ij}.$$

定义 6 入度中心性 在有向图中,入度表
示所有指向该节点的连线的加和.令 d_i^{in} 表
示节点 i

的入度,则
$$d_i^{\text{in}} = \sum_{j=1}^n c_{ji}.$$

定义 7 出度中心性 在有向图中,出度表
示

由该节点指向其它节点的连线的和.令 d_i^{out}
表示节点 i 的出度,则
$$d_i^{\text{out}} = \sum_{j=1}^n c_{ij}.$$

本研究是基于两个节点之间的学习关系来
识别重要节点,需表明谁向谁学习,因此将构
建有向网络图,利用二值矩阵表示各节点之
间的联系.其中,如果节点 A_i 和 A_j 之间存在
学习关系,且节点 A_i 向节点 A_j 学习,则有
向连接由节点 A_i 指向节点 A_j .如果有较
多的节点指向节点 A_j ,代表该节点被较多
的节点学习.由于本研究所构建的社会网络
是一种基于学习关系的网络,任何表现差的
节点都可以向表现好的节点学习,与样本所
处的位置和邻居无关,因此本研究选择入度
中心性识别标杆(具体方法见 2.2 节).

2 基于社会网络 DEA 方法识别标杆

2.1 构建社会网络

如前文所述,要通过构建社会网络来识别

杆,首先需识别出两两节点之间的学习关系.为了识别大样本中两两 DMU 之间的学习关系,本节将基于 DEA 有效性和 Pareto 最优解之间的等价关系,提出一个两两评价过程.在该过程中,每次仅两个 DMU 进行比较,而不是整个样本中所有 DMU 放在一起进行比较.在传统 DEA 模型中,被评价的 DMU 通常被分为两类:有效 DMU 和非有效 DMU.非有效 DMU 参考有效 DMU,通过向有效 DMU 学习来提升自己的效率.

令 (X_A, Y_A) 和 (X_B, Y_B) 分别表示 DMU_A 和 DMU_B 的投入产出指标,其中 $A, B \in N, N = \{1, \dots, n\}$ 表示所有待评价 DMU 的集合.另外, $X_A = (x_{1A}, \dots, x_{iA}, \dots, x_{mA})^T$, $Y_A = (y_{1A}, \dots, y_{rA}, \dots, y_{sA})^T$, $X_B = (x_{1B}, \dots, x_{iB}, \dots, x_{mB})^T$, $Y_B = (y_{1B}, \dots, y_{rB}, \dots, y_{sB})^T$. 根据第 1.1 节的介绍,在 DMU_A 和 DMU_B 两个 DMU 构成的生产可能集 (production possibility set, PPS) 中,评价 DMU_B 时,传统的 CCR 模型包络形式如下所示

$$\begin{aligned} \min \theta_B \\ \text{s. t. } X_A \lambda_A + X_B \lambda_B \leq \theta_B X_B \\ Y_A \lambda_A + Y_B \lambda_B \geq Y_B \\ \lambda_A, \lambda_B \geq 0 \end{aligned} \quad (3)$$

模型(3) 或其对偶模型(即,类似于模型(1) 的乘数形式的 DEA 模型) 可用于计算 DMU_A 和 DMU_B 的效率值.然而,这需求解两个规划模型,所以在 DMU 数量较大时,该过程是极其耗时的(见图 2). 以下将介绍一个简单的两两评价过程来大量缩短任意两个 DMU 之间学习关系的判断时间,为识别大样本中的标杆奠定基础.

定理 1 令 $\delta = \max_{r=1, \dots, s} \frac{y_{Br}}{y_{Ar}}$. 如果 $X'_A = \delta X_A = \left(\max_{r=1, \dots, s} \frac{y_{Br}}{y_{Ar}} \right) X_A < X_B$, 则 DMU_B 相对于 DMU_A 来说是非有效的; 否则, DMU_B 相对于 DMU_A 来说是有效的.

其中 DMU_B 相对于 DMU_A 来说是(非)有效的是指 DMU_B 在 A 和 B 两个 DMU 构成的 PPS 中是相对(非)有效的.

证明 首先证明 DMU_B 相对于 DMU_A 来说是

非有效的, 如果 $X'_A = \delta X_A = \left(\max_{r=1, \dots, s} \frac{y_{Br}}{y_{Ar}} \right) X_A < X_B$.

将 $\lambda_A = \delta$ 和 $\lambda_B = 0$ 代入模型(3) 可得

$$\begin{aligned} \min \theta_B \\ \text{s. t. } X_A \delta \leq \theta_B X_B \\ Y_A \delta \geq Y_B \\ \delta \geq 0 \end{aligned} \quad (4)$$

模型(4) 中当 $\delta = \max_{r=1, \dots, s} \frac{y_{Br}}{y_{Ar}}$ 时, $Y_A \delta \geq Y_B$ 必然是

成立的. 另外, 有 $X'_A = \delta X_A = \left(\max_{r=1, \dots, s} \frac{y_{Br}}{y_{Ar}} \right) X_A < X_B$. 因此, 存在 $\theta_B < 1$ 使得 $\lambda_A = \delta$ 和 $\lambda_B = 0$ 时 $X_A \delta \leq \theta_B X_B$ 成立. 于是, 模型(4) 的最优值 $\theta_B^* = \min \theta_B$ 必然是小于 1 的. 由此可见, 被评价的 DMU_B 在 DMU_A 和 DMU_B 构成的生产可能集中是相对非有效的.

然后, 证明如果 $X'_A = \delta X_A = \left(\max_{r=1, \dots, s} \frac{y_{Br}}{y_{Ar}} \right) X_A < X_B$ 不成立, 则 DMU_B 在 DMU_A 和 DMU_B 构成的生产可能集中是有效的(包括强有效和弱有效). 为此, 假设 $X'_A = \delta X_A = \left(\max_{r=1, \dots, s} \frac{y_{Br}}{y_{Ar}} \right) X_A \geq X_B$ 时, DMU_B 在 DMU_A 和 DMU_B 构成的生产可能集中是非有效的.

根据 1.1 节的介绍, 在 DMU_A 和 DMU_B 构成的生产可能集中, DMU_B 被评价时, 投入导向的 CCR 模型的乘数形式如下所示

$$\begin{aligned} \max \mu Y_B \\ \text{s. t. } \omega X_j - \mu Y_j \geq 0, j = A, B \\ \omega X_B = 1 \\ \omega, \mu \geq 0 \end{aligned} \quad (5)$$

由于模型(3) 和模型(5) 互为对偶模型, 因此两者的最优值是相等的. 令 ω^* , μ^* 和 λ_A^* , λ_B^* , θ_B^* 分别是模型(5) 和模型(3) 的最优解. 若 DMU_B 是无效的, 模型(5) 的最优值 $\mu^* Y_B < 1$. 于是, $\omega^* X_B - \mu^* Y_B > 0$ 成立. 根据对偶理论可知, $\lambda_B^* =$

$$0. \text{ 因此, } \begin{cases} X_A \lambda_A^* \leq \theta_B^* X_B \\ Y_A \lambda_A^* \geq Y_B \\ \lambda_A^* \geq 0 \end{cases} \text{ 成立, 并且至少有一个 } i$$

使得 $x_{iA}\lambda_A^* = \theta_B^* x_{iB}$ 成立. 据此,可得

$$1 > \theta_B^* \geq \frac{X_A}{X_B} \lambda_A^* \geq \frac{X_A}{X_B} \times \frac{Y_B}{Y_A} \quad (6)$$

因此,至少有一个 i 使得 $1 > \theta_B^* = \frac{x_{iA}}{x_{iB}} \lambda_A^* \geq$

$\frac{x_{iA}}{x_{iB}} \times \frac{y_{Br}}{y_{Ar}}$ 成立, 其中 $r = 1, \dots, s$. 显然, $X_B >$

$X_A \left(\max_{r=1, \dots, s} \frac{y_{Br}}{y_{Ar}} \right) = X'_A$ 成立, 与假设不一致. 进而得

出 $X'_A = \delta X_A = \left(\max_{r=1, \dots, s} \frac{y_{Br}}{y_{Ar}} \right) X_A \geq X_B$ 时, DMU_B 在

DMU_A 和 DMU_B 构成的生产可能集中不是非有效的, 即 DMU_B 相对于 DMU_A 来说是有效的 (包括强有效和弱有效). 证毕

对于一个大样本, 任意一个 DMU 相对于另一个 DMU 来说, 其相对有效或相对非有效的状态可由定理 1 得出. 从而可以得到样本中任意两个 DMU 之间的学习关系. DMU_B 相对 DMU_A 来说是有效的, 则说明 DMU_A 应向 DMU_B 学习; DMU_B 相对 DMU_A 来说是非有效的, 则说明 DMU_B 应向 DMU_A 学习. 本研究将该过程称为两两评价过程. 该过程与传统 DEA 评价过程 (即模型 (1) 和模型 (2)) 的区别在于该过程中仅由两个 DMU 构成生产可能集, 可通过定理 1 直接进行数值上的比较, 进而得到两个 DMU 之间的相对效率状态; 而传统 DEA 评价过程用样本中所有 DMU 构建生产可能集, 这将极大地增加模型 (1) 和模型 (2) 的计算复杂度, 使得计算费时甚至无法计算, 且得不到两两 DMU 之间的学习关系. 为了区分, 本研究将两两评价过程中相对于另一个 DMU 来说有效的 DMU 称为局部有效 DMU , 相对非有效的 DMU 称为局部非有效 DMU ; 将整个样本中所有 DMU 构成的生产可能集中相对有效的 DMU 称为全局有效 DMU , 非有效的 DMU 称为全局非有效 DMU .

引理 1 若 DMU_B 相对于 DMU_A 来说是有效的, 同时 DMU_A 相对于 DMU_B 来说是有效的, 则两个 DMU 均为局部有效.

引理 2 局部有效 DMU 可能是全局有效的,

也可能是全局非有效的. 即两两评价过程中相对有效的 DMU , 在整个样本中所有 DMU 构成的生产可能集中可能有效, 也可能非有效.

引理 3 局部非有效 DMU 一定是全局非有效的, 即, 两两评价过程中相对非有效 DMU 在整个样本中一定是非有效的.

除以上引理外, 两两评价过程中, 任意 DMU_A 和 DMU_B 之间存在如下一些性质, 其中 $A, B \in \{1, 2, \dots, n\}$.

性质 1 DMU_A 是局部有效的, 如果 DMU_B 是局部非有效的.

性质 2 如果 DMU_A 相对于 DMU_B 来说是有效, 且 DMU_B 相对于 DMU_F 有效, 则 DMU_A 相对于 DMU_F 有效.

通过两两评价过程, 可得到每个 DMU 相对于另一个 DMU 的局部效率状态, 进而得到任意两个 DMU 之间的学习关系. 用矩阵 C 来记录两个 DMU 之间的学习关系, 并将其称为两两评价矩阵 $C = c_{AB}, A, B \in \{1, 2, \dots, n\}$. 假设大样本的集合 $\{1, 2, \dots, n\}$ 中任意两个 DMU 为 DMU_A 和 DMU_B , DMU_B 为被评价 DMU . 于是在两两评价矩阵中, 用 c_{AB} 来表示 DMU_B 相对于 DMU_A 的局部效率状态. 令 $c_{AB} = 0$ 表示 DMU_B 相对于 DMU_A 来说是局部非有效的, $c_{AB} = 1$ 表示 DMU_B 相对于 DMU_A 来说是局部有效的. 基于此, 若 $c_{AB} = c_{BA} = 1$, 则说明 DMU_A 和 DMU_B 均为局部有效, 于是两个 DMU 无法区分谁向谁学习. 为便于构建社会网络, 将所有 $c_{AB} = c_{BA} = 1$ 替换为 $c_{AB} = c_{BA} = 0$. 另外, 当 A 等于 B , 即两个相同 DMU 进行两两评价时, c_{AB} 恒等于 1, 也就是说, 两个相同 DMU 之间不存在向谁学习的关系. 此外, 根据性质 1, 如果 DMU_B 在两两评价过程中相对于 DMU_A 来说是局部非有效的, 那么 DMU_A 一定是局部有效的, 于是, 由 $c_{AB} = 0$ 可推导出 $c_{BA} = 1$. 基于以上分析, 令 $P = ((A, B) | A = 1, \dots, n; B = 1, \dots, n)$ 表示 n 个 DMU 中任意两个 DMU 的组合, 则两两评价过程如图 1 所示.

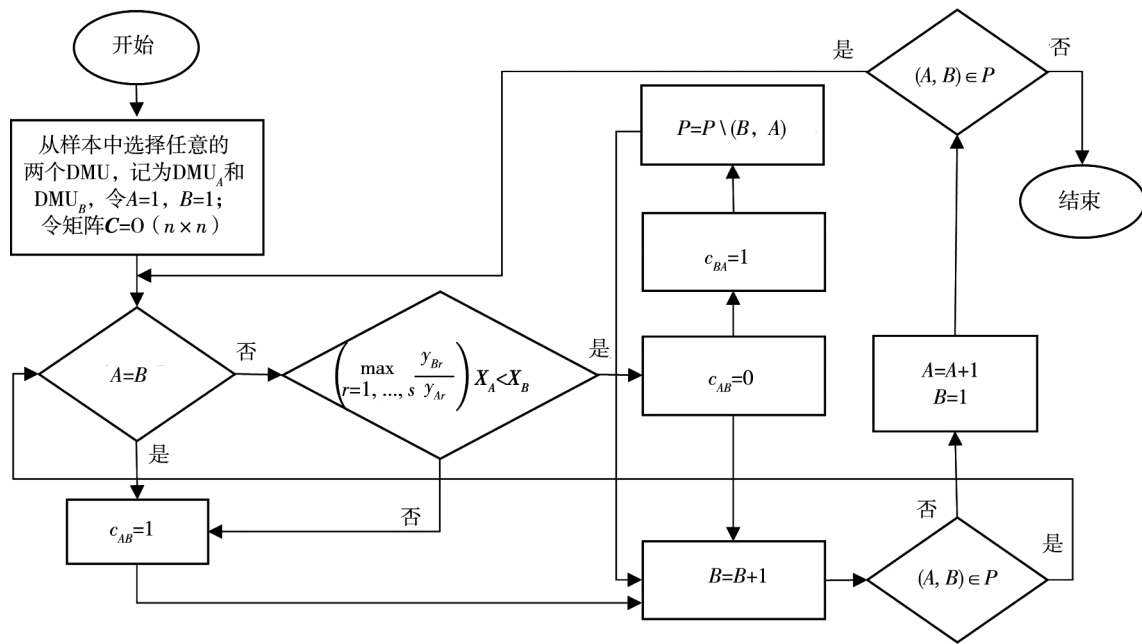


图 1 两两评价过程

Fig. 1 Pairwise evaluation process

根据图 1, 可通过如下伪代码来得到两两评价矩阵中所有元素的值.

```

{ 如果  $(A, B) \in P$ 
  如果  $(A, B) \in P$ 
    如果  $A \neq B$ 
      如果  $\max_{r=1, \dots, s} \left( \frac{y_{Br}}{y_{Ar}} \right) X_A$ 
        小于  $X_B$ 
           $c_{AB} = 0$ 
           $c_{BA} = 1$ 
           $P = P \setminus \{A, B\}$ 
        否则
           $c_{AB} = 1$ 
        结束
      否则
         $c_{AB} = 1$ 
    结束
  否则
     $c_{AB} = 1$ 
  结束
   $B = B + 1$ 
  结束
   $A = A + 1$ 
结束
}

```

以上两两评价过程无需对任何线性规划或非线性规划模型进行求解, 因此, 不同于传统 DEA 模型, 即使待评价 DMU 数量极大, 定理 1 中所示

的数值比较都是可以简便且快速完成的. 如图 2 所示, 上述两两评价过程与用模型 (1) 对任意两个 DMU 之间相对效率进行评价的过程相比, 在计算时间上具有明显的优势.

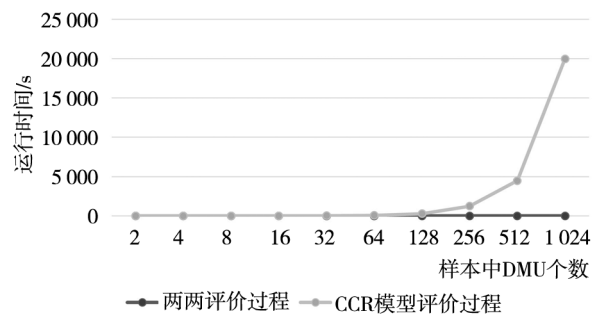


图 2 计算时间对比

Fig. 2 Comparison of computation time

并且两两评价矩阵可反映出大样本中任意两个 DMU 之间的学习关系, 而不是传统 DEA 效率值反映的一个 DMU 相对于多个 DMU 的学习关系. 另外, 由引理 2 和引理 3 可知, 局部非有效 DMU 必然也是全局非有效 DMU, 因此不可能被其它 DMU 学习, 也就是说, 大样本中的标杆 DMU 一定是在两两评价过程中局部有效 DMU 中产生. 因此, 上述两两评价矩阵揭示了大样本中任意两个 DMU 之间的学习关系, 从而可以避免求解规划问题的情况下从大样本中找到标杆 DMU.

2.2 识别社会网络中的标杆

本节主要提出确定社会网络中标杆的方法. 根据2.1节可以获得两两评价矩阵 C , 据此可以构建一个二值有向网络图, 记为 $G(E, A, C)$. 如果 $c_{AB} = 1$, 说明 DMU_A 可以向 DMU_B 学习. 对于一个差的DMU, 可以向比它稍微好一点的DMU学习, 也可以向比它好很多的DMU学习; 对于一个DMU, 如果被很多DMU学习, 说明这个DMU在社会网络中较重要, 因此将其作为样本中的标杆. 据此, 可以用1.3节介绍的入度中心性来衡量DMU重要性, 进而从样本中识别标杆.

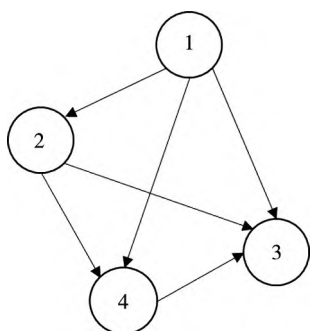


图3 4个DMUs构成的社会网络

Fig. 3 The social network graph of four DMUs

以4个DMU为例, DMU_1 可以向 DMU_2 、 DMU_3 、 DMU_4 学习; DMU_2 可以向 DMU_3 、 DMU_4 学习; DMU_4 可以向 DMU_3 学习. 这些学习关系可用图3所示的社会网络图来表示. 于是各节点的入度中心性如表2所示.

表2 各节点的入度中心性

Table 2 The in-degree centrality of nodes

节点	入度中心性	排名
1	0	4
2	1	3
3	3	1
4	2	2

由表2所示的入度中心性可见, DMU_3 的入度中心性为3, 排名第一, 是其他DMU的学习标杆. 根据该结果, 本研究发现 DMU_1 向 DMU_2 学习, DMU_1 的入度中心性小于 DMU_2 ; DMU_2 向 DMU_3 学习, DMU_2 的入度中心性小于3等. 这种现象说明入度中心性之间存在一种序关系. 该性质和证明如下.

性质3 如果一个DMU被另一个DMU学

习, 那么前者的入度中心性不小于后者. 这种性质称为“序关系”.

证明 如果 DMU_1 向 DMU_2 学习, 在社会网络图中的连接由 DMU_1 指向 DMU_2 . 此外, 由性质2可得, 如果 DMU_1 向 DMU_2 学习, DMU_2 向 DMU_3 学习, 则 DMU_1 也会向 DMU_3 学习. 因此, 所有向 DMU_1 学习的DMU也会向 DMU_2 学习, DMU_2 的入度中心性至少比 DMU_1 多1. 即一个DMU向另一个DMU学习, 那么前者的入度中心性小于后者. 证毕.

当网络图中DMU数量较多时, 可以通过计算和对比各节点的入度中心性来识别标杆.

2.3 分析标杆的优势

在识别标杆后, 揭示标杆的内在优势也很重要, 因为这可以为其他DMU提供明确的学习方向. 结合DEA和SNA, 本节按如下步骤分析标杆DMU的优势.

首先, 将所有投入产出数据进行归一化, 如式(8)将所有DMU的第 i ($i = 1, \dots, m$) 个投入以及第 r ($r = 1, \dots, s$) 个产出归一化到区间 $(0, 1]$

$$\begin{aligned} \tilde{x}_{ij} &= \frac{x_{ij}}{x_i^{\max}}, j = 1, \dots, n, \forall i \\ \tilde{y}_{rj} &= \frac{y_{rj}}{y_r^{\max}}, j = 1, \dots, n, \forall r \end{aligned} \quad (8)$$

式(8)中 $x_i^{\max}$, $y_r^{\max}$ 分别是所有DMU中第 i 个投入和第 r 个产出的最大值, x_{ij} , y_{rj} 分别是第 j 个DMU中第 i 个投入和第 r 个产出的原始值, $\tilde{x}_{ij}$ 和 $\tilde{y}_{rj}$ 分别表示标准化之后的投入产出指标值.

然后, 利用权重分析标杆DMU的优势.

在DEA中, 投入产出的权重可反映其重要性, 如果 DMU_A 的第 r 个产出的权重 u_r 较大, 说明它在产出指标 y_r 上表现较好, 且被评价DMU的目标是使自身效率最大, 往往会在占优势的指标上赋予更大的权重. 如果针对某个特定的标杆DMU, 用任意向他学习的DMU来评价该标杆DMU, 得到的最优权重表示标杆DMU相对于向它学习的DMU来说, 对各指标的重视程度, 也可以称为优势所在.

为获取标杆DMU的优势, 本研究假设社会网络 $G(E, A)$ 中有 k 个标杆DMU, 每个标杆DMU

被其他 DMU 学习. 令 P_j 表示向标杆 DMU_j 学习的集合, $P_j = \{DMU_1, DMU_2, \dots, DMU_{l_j}\}$, 集合 P_j 中有 l_j 个元素, 表示向 DMU_j 学习的 DMU 个数. 假设 DMU_A 被集合 $P_A = \{DMU_1, DMU_2, \dots, DMU_{l_A}\}$ 中的 DMU 学习, 将 DMU_A 和 P_A 中的每一个 DMU 比较, 获得 DMU_A 的优势. 以 DMU_A 和 DMU_k 为参考集, DMU_k 为集合 P_A 中的任意 DMU, 利用传统 CCR 模型对 DMU_A 进行评价, 模型如下

$$\begin{aligned} \max E_A &= \sum_{r=1}^s u_r \tilde{y}_{rA} \\ \text{s. t. } &\sum_{r=1}^s u_r \tilde{y}_{rk} - \sum_{i=1}^m v_i \tilde{x}_{ik} \leq 0, \forall k, k \in P_A \\ &\sum_{i=1}^m v_i \tilde{x}_{iA} = 1 \\ &v_i, u_r \geq 0, \forall i=1, \dots, m; r=1, \dots, s \quad (9) \end{aligned}$$

由于 DMU_A 需要依次和 P_A 中的每一个 DMU 形成参考集, 然后对自身效率进行评价, 对于 DMU_A 而言, 共有 $|P_A|$ 组最优权重 ($|P_A|$ 表示集合 P_A 中的 DMU 个数). 令 W_A 表示 DMU_A 的最优权重矩阵, 可表示如下

$$W_A = \begin{bmatrix} u_1^1 & u_2^1 & u_3^1 & u_4^1 & u_5^1 & v_1^1 \\ u_1^2 & u_2^2 & u_3^2 & u_4^2 & u_5^2 & v_1^2 \\ \dots & \dots & \dots & \dots & \dots & \dots \\ u_1^k & u_2^k & u_3^k & u_4^k & u_5^k & v_1^k \\ \dots & \dots & \dots & \dots & \dots & \dots \\ u_1^{|P_A|} & u_2^{|P_A|} & u_3^{|P_A|} & u_4^{|P_A|} & u_5^{|P_A|} & v_1^{|P_A|} \end{bmatrix}$$

W_A 是一个 $l_A \times (m+s)$ 维的矩阵. 为了衡量 DMU_A 相较于所有 DMU_k ($DMU_k \in P_A$) 的优

势, 对所有 v_i, u_r 进行平均, 即 $\bar{u}_r = \frac{\sum_{k=1}^{|P_A|} u_r^k}{|P_A|}$, $\bar{v}_i = \frac{\sum_{k=1}^{|P_A|} v_i^k}{|P_A|}$, 以此来判断 DMU_A 的优势. 如果 $\bar{u}_r$ 或 $\bar{v}_i$ 较高, 那么可以认为该标杆 (DMU_A) 在产出 r 或投入 i 上具有优势.

3 应用

春雨医生平台创立于 2011 年, 已服务 1.4 亿

注册用户和 66 万以上注册医生, 是较早进行问诊的医疗平台. 该平台涵盖 17 个一级科室, 每次问诊 3 min 之内响应, 可为用户建立真实的个人健康档案, 减缓了普通人看病难的问题. 基于此, 本文以春雨医生平台为例, 爬取该平台 55 个科室的医生诊断数据, 并利用上述方法, 从该平台的众多医生中识别出标杆医生, 并确定他们的优势.

3.1 数据描述

目前, 人们越来越注重身体健康, 尤其是新冠疫情爆发之后, 这使得一些在线医疗平台发挥了重要作用. 在线医疗平台不仅可以方便病人问诊, 提高寻医问药效率, 还可以避免线下问诊引起的交叉感染. 本研究以识别春雨医生平台上的标杆医生为例, 利用 Python 爬虫代码爬取了该网站上的医生信息, 包括医生的咨询价格, 服务人次, 患者对医生在服务态度、病情讲解是否清楚、回复是否及时、建议是否有帮助等方面的评价. 爬取的原始诊断数据一共 16 642 条, 但为了保证结果的可靠性, 对原始数据进行了预处理, 删除了部分缺失数据和不合理数据, 最后还剩 10 418 条. 于是, 根据以往关于大规模或大样本 DEA 的研究, 该样本可视为一个大样本. 本研究从中选取了 6 个指标来进行验证, 分别记为

X1 (图文咨询价格): 医生提供单次医疗服务的价格;

Y1 (服务人次): 医生共服务过的人次数;

Y2 (态度非常好): 患者对医生的评价中认为医生态度好的评价数;

Y3 (讲解很清楚): 患者对医生的评价中认为医生讲解很清楚的评价数;

Y4 (回复很及时): 患者对医生的评价中认为医生回复信息及时的评价数;

Y5 (建议很有帮助): 患者对医生的评价中认为医生提供的建议很有帮助的评价数.

将每个医生视为决策单元 (DMU), X1 表示投入指标, Y1, Y2, Y3, Y4, Y5 表示产出指标. 理论上而言, 医生的图文咨询价格越低, 服务人次越多, 好评数越多 (Y2, Y3, Y4, Y5), 则表现越好. 表 3 给出了所有投入产出数据的统计信息, 包括均值、标准差、最小值、最大值和四分位数.

表 3 10 418 名医生诊断数据统计信息
Table 3 Statistical information of 10 418 doctors' diagnostic data

指标	X1/元	Y1/(人次)	Y2/条	Y3/条	Y4/条	Y5/条
均值	22.37	3 133.48	262.76	222.20	254.18	227.36
标准差	18.54	7 684.86	630.18	524.41	621.97	542.14
最小值	1.00	1.00	1.00	1.00	1.00	1.00
Q1(25%)	10.00	245.00	19.00	17.00	19.00	17.00
Q2(50%)	18.00	785.50	64.00	56.00	61.00	56.00
Q3(75%)	30.00	2 779.75	223.00	192.00	211.00	193.00
最大值	100.00	235 128.00	20 119.00	16 364.00	21 200.00	17 350.00

从以上表中可知,图文咨询价格(X1)的平均值为 22.37,四分位数分别为 10,18,30,最大值为 100,说明各个医生的图文咨询价格差异较大,大部分保持在 10~30 之间. Y1 的标准差较大,说明各个医生服务的人次差异大. Y2,Y3,Y4,Y5 的各个指标较近,

但 Y2 的标准差最大,说明在医生的评价中,患者认为各个医生的服务态度好的评价数波动相对较大.

3.2 实验步骤和结果

利用以上数据,对本研究所提出的方法进行验证,具体的实验步骤如以下流程图所示.

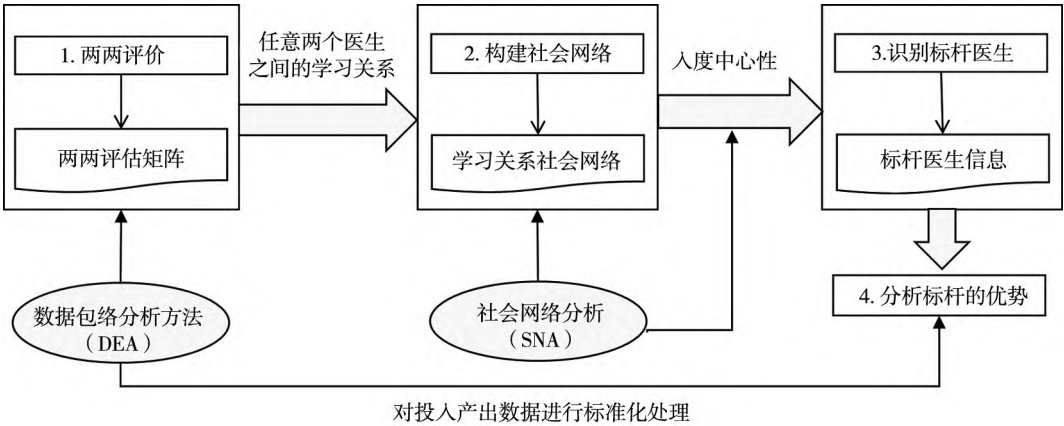


图 4 实验步骤

Fig. 4 Process of experiment

经过步骤 1 中的两两评价过程后,可以得到一个两两评价矩阵,该过程计算时间仅 657.258 s. 然后,根据该矩阵可绘制医生之间的学习网络图. 由于节点(医生)数量较多,不能够完全清晰的展示医生之间的学习关系,因此过滤了入度较低的节点,只在图中展示了入度较高节点之间的关系. 在图 5 中,节点代表医生,连线代表学习关系,箭头指向的节点是被学习的对象. 节点大小与入度中心性有关,节点越大代表该节点入度越高,是其他节点学习的榜样;

相反,节点越小表示该节点入度越小,需要向箭头指向的节点学习. 如图 5,深灰色且形状较大的节点入度高,是被学习的对象;左下方与右侧形状较小的节点入度小,需要向其他节点学习. 本研究选取了几个节点作为标杆,此处用字母 G^{leader} 表示标杆构成的集合, $G^{leader} = \{3605, 4110, 4663, 3589, 3541, 3770, 776, 3660, 9580, 6817\}$.

为了详细分析实验结果,表 4 列出了春雨医生平台标杆医生的具体数据.

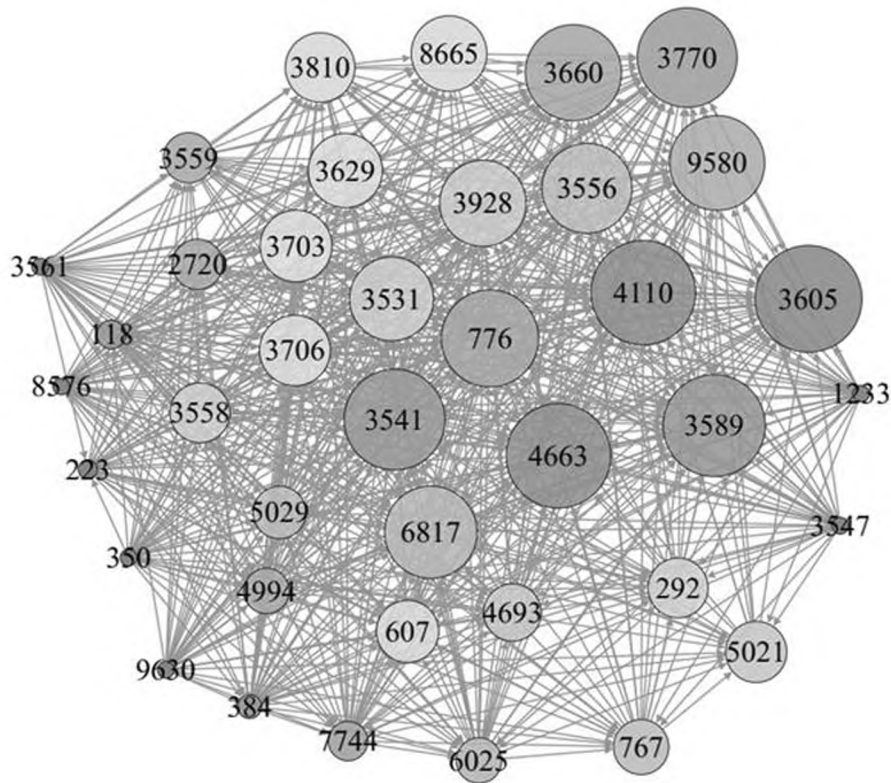


图 5 在线医生学习网络

Fig.5 The learning network graph of online doctors

表 4 从 10 418 名医生中识别出的标杆医生评价指标

Table 4 The evaluation indicators of benchmark doctors selected from 10 418 doctors

标杆医生	Y1/(人次)	Y2/条	Y3/条	Y4/条	Y5/条	X1/元
DMU3605	76 707	7 203	5 558	7 119	5 872	2
DMU4110	21 651	3 205	2 443	3 033	2 398	1
DMU4663	43 321	5 971	4 925	6 190	5 109	2
DMU3589	112 002	7 010	5 723	7 138	6 550	6
DMU3541	235 128	20 119	16 364	21 200	17 350	15
DMU3770	50 539	2 654	2 100	2 423	2 296	3
DMU776	37 160	3 031	2 498	3 026	2 630	3
DMU3660	91 160	4 120	3 275	4 283	3 562	6
DMU9580	40 402	1 878	1 541	1 923	1 638	3
DMU6817	14 765	1 407	1 113	1 368	1 154	2

在识别标杆医生后,了解这些医生之所以成为标杆的原因也很重要,这就需要探索这些医生的优势.根据 2.3 节中的方法,用权重反映标杆医

生的优势,权重结果如表 5 所示. $\bar{u}_1$, $\bar{u}_2$, $\bar{u}_3$, $\bar{u}_4$, $\bar{u}_5$ 分别代表产出的平均权重, $\bar{v}_1$ 代表投入的平均权重.表中展示了 10 个标杆医生的权重,对

比各个标杆医生的权重,本研究发现权重 $\bar{v}_1$ 的值比其他权重的值高,说明图文资询价(X1)在标杆医生选取过程中发挥了重要作用.每个标杆医生的优势是不同的,以序号为 3605 的医生(DMU3605)为例, $\bar{u}_1$, $\bar{v}_1$ 的值比其他权重大,说明该医生在服务人次和咨询价格上表现较好.因此,标杆医生 3605 的优势是 X1, Y1,在集合 P_{3605} 中的医生要以医生 3605 为标杆,增加自身在春雨医生平台的服务人次,同时降低图文资询价.对于

医生 4110 和医生 4663,权重 $\bar{v}_1$ 和 $\bar{u}_5$ 的值相对其他权重较高,表示这两个医生在 X1 和 Y5 上具有优势,因此,集合 P_{4110} 和 P_{4663} 中的医生应该降低图文资询价,并向病人提供较多的有利于身体健康的建议.剩余的其他标杆医生大多数在 X1 和 Y1 上具有优势.总体而言,这些标杆医生通常具有一些相同的特点,即以较低的图文资询价提供较高质量的服务,这些特征可以为春雨医生平台选取标杆医生提供参考.

表 5 标杆医生权重结果
Table 5 The weight results of benchmark doctors

榜样医生	$\bar{u}_1$	$\bar{u}_2$	$\bar{u}_3$	$\bar{u}_4$	$\bar{u}_5$	$\bar{v}_1$
DMU3605	1.036 4	0.751 1	0.123 1	0.435 0	0.606 0	50.000 0
DMU4110	1.007 2	1.815 9	0.451 3	0.768 5	3.188 2	100.000 0
DMU4663	0.509 6	0.702 4	0.220 7	0.413 1	1.734 0	50.000 0
DMU3589	1.245 7	0.272 5	0.119 0	0.387 6	0.369 6	16.666 7
DMU3541	0.393 0	0.189 4	0.071 3	0.171 6	0.174 6	6.666 7
DMU3770	3.504 6	0.681 9	0.218 7	0.381 8	0.642 7	33.333 3
DMU776	2.697 0	1.039 8	0.426 9	0.901 6	1.472 8	33.333 3
DMU3660	2.026 5	0.275 0	0.084 5	0.487 6	0.207 4	16.666 7
DMU9580	4.598 2	0.576 6	0.291 9	0.837 0	0.557 9	33.333 3
DMU6817	5.249 6	3.468 6	0.863 4	1.796 2	3.806 0	0

4 结束语

本研究主要提出社会网络 DEA 方法来识别大样本中的标杆.首先利用 DEA 方法对任意两个 DMU 进行评价,将评价结果用两两评价矩阵表示;接着,基于两两评价矩阵构建 DMU 之间的社会网络图,通过网络图展示两两 DMU 之间的学习关系;然后,计算入度中心性并据此识别社会网络中的重要节点,以此作为大样本中的标杆;随后,利用 DEA 模型中投入产出的权重来分析标杆的优势,揭示 DMU 被选为标杆的原因;最后,为了验证该方法的可行性,利用春雨医生平台上的诊断数据进行实验.

本研究所提出的方法对于管理者或决策者选择某个领域内标杆具有重要的参考价值.其一,该方法可以避免求解多个线性/非线性规划模型,仅需通过简单的数值比较,实施起来较以往的方法更便捷.其二,该方法以一种客观的方式确定社会网络中任意节点间的学习关系,可避免主观因素带来的影响.其三,该方法可以分析标杆的优势,进而使表现较差者知道哪些方面需要改进,进而提高整体的效率水平.

在未来的研究中,可以基于本研究进行以下拓展:首先,结合社会网络和 DEA 进行研究是一项有趣和有挑战性的任务,未来可以结合社会网络和 DEA,探讨 DMU 之间其他关系;然后,各行业中经常会评选某一方面表现较好的单位和个

体,如道德模范,最佳会议酒店,最佳服务酒店等,研究在指定的某一方面表现较好的标杆也非常重要;最后,大数据具有规模性、高速性、多样性和价

值性四个特征,本研究只研究了大数据的规模性这一特征,未来可以研究如何结合其他三个特征识别标杆.

参 考 文 献:

- [1]Freeman L C. Centrality in social networks conceptual clarification[J]. *Social Networks*, 1978, 1(3): 215-239.
- [2]Newman M E J. The structure and function of complex networks[J]. *SIAM Review*, 2003, 45(2): 167-256.
- [3]Lü L, Chen D, Ren X L, et al. Vital nodes identification in complex networks[J]. *Physics Reports*, 2016, (650): 1-63.
- [4]Hage P, Harary F. Eccentricity and centrality in networks[J]. *Social Networks*, 1995, 17(1): 57-63.
- [5]Katz L. A new status index derived from sociometric analysis[J]. *Psychometrika*, 1953, 18(1): 39-43.
- [6]Bavelas A. A mathematical model for group structures[J]. *Human Organization*, 1948, 7(3): 16-30.
- [7]Bonacich P. Factoring and weighting approaches to status scores and clique identification[J]. *Journal of Mathematical Sociology*, 1972, 2(1): 113-120.
- [8]Poulin R, Boily M C, Masse B R. Dynamical systems to define centrality in social networks[J]. *Social Networks*, 2000, 22(3): 187-220.
- [9]Brin S, Page L. The anatomy of a large-scale hypertextual web search engine[J]. *Computer Networks and ISDN Systems*, 1998, 30(1-7): 107-117.
- [10]Lü L, Zhang Y C, Yeung C H, et al. Leaders in social networks, the delicious case [J]. *PloS One*, 2011, 6(6): e21202.
- [11]Charnes A, Cooper W W, Rhodes E. Measuring the efficiency of decision making units[J]. *European Journal of Operational Research*, 1978, 2(6): 429-444.
- [12]安庆贤,陈晓红,余亚飞,等. 基于 DEA 的两阶段系统中间产品公平设定研究[J]. *管理科学学报*, 2017, 20(1): 32-40.
- An Qingxian, Chen Xiaohong, Yu Yafei, et al. Fair setting for intermediate products in two-stage system based on DEA[J]. *Journal of Management Sciences in China*, 2017, 20(1): 32-40. (in Chinese)
- [13]谢建辉,李勇军,梁 樑,等. 随机环境下的多投入多产出生产前沿面估计[J]. *管理科学学报*, 2018, 21(11): 50-60.
- Xie Jianhui, Li Yongjun, Liang Liang, et al. Estimation of multiple inputs and multiple outputs frontiers in stochastic environment[J]. *Journal of Management Sciences in China*, 2018, 21(11): 50-60. (in Chinese)
- [14]李光金. 评价相对效率的投入-产出型 DEA[J]. *管理科学学报*, 2001, 4(2): 58-62.
- Li Guangjin. Input-and-output-oriented DEA for assessing relative efficiency[J]. *Journal of Management Sciences in China*, 2001, 4(2): 58-62. (in Chinese)
- [15]周卓儒,王 谦,李锦红. 基于标杆管理的 DEA 算法对公共部门的绩效评价[J]. *中国管理科学*, 2012, (3): 72-75.
- Zhou Zhuoru, Wang Qian, Li Jinhong. Performance evaluation by means of improved DEA method[J]. *Chinese Journal of Management Science*, 2012, (3): 72-75. (in Chinese)
- [16]李 峰,朱 平,梁 樑,等. 基于最近距离投影的 DEA 两阶段效率评价方法研究[J]. *中国管理科学*, 2022, 30(10): 198-209.

- Li Feng, Zhu Ping, Liang Liang, et al. A two-stage DEA efficiency evaluation approach based on closest targets[J]. Chinese Journal of Management Science, 2022, 30(10): 198–209. (in Chinese)
- [17] 宋马林, 吴杰, 杨力, 等. 非期望产出, 影子价格与无效决策单元的改进[J]. 管理科学学报, 2012, 15(10): 1–10.
- Song Malin, Wu Jie, Yang Li, et al. Undesirable outputs, shadow prices, and the improvement of inefficient DMUs[J]. Journal of Management Sciences in China, 2012, 15(10): 1–10. (in Chinese)
- [18] Khezrimotlagh D, Zhu J, Cook W D, et al. Data envelopment analysis and big data[J]. European Journal of Operational Research, 2019, 274(3): 1047–1054.
- [19] Chen W C, Cho W J. A procedure for large-scale DEA computations[J]. Computers & Operations Research, 2009, 36(6): 1813–1824.
- [20] Ali A I. Streamlined computation for data envelopment analysis[J]. European Journal of Operational Research, 1993, 64(1): 61–67.
- [21] Barr R S, Durchholz M L. Parallel and hierarchical decomposition approaches for solving large-scale data envelopment analysis models[J]. Annals of Operations Research, 1997, (73): 339–372.
- [22] Zhu Q, Wu J, Song M. Efficiency evaluation based on data envelopment analysis in the big data context[J]. Computers & Operations Research, 2018, (98): 291–300.
- [23] Jie T. Parallel processing of the build hull algorithm to address the large-scale DEA problem[J]. Annals of Operations Research, 2020, 295(1): 453–481.
- [24] Dulá J H. An algorithm for data envelopment analysis[J]. Informds Journal on Computing, 2011, 23(2): 284–296.
- [25] Chen W C, Lai S Y. Determining radial efficiency with a large data set by solving small-size linear programs[J]. Annals of Operations Research, 2017, 250(1): 147–166.
- [26] Dellnitz A. Big data efficiency analysis: Improved algorithms for data envelopment analysis involving large datasets[J]. Computers & Operations Research, 2022, (137): 105553.
- [27] Liu J S, Lu W M, Yang C, et al. A network-based approach for increasing discrimination in data envelopment analysis[J]. Journal of the Operational Research Society, 2009, 60(11): 1502–1510.
- [28] Liu J S, Lu W M. DEA and ranking with the network-based approach: A case of R&D performance[J]. Omega, 2010, 38(6): 453–464.
- [29] de Blas C S, Martin J S, Gonzalez D G. Combined social networks and data envelopment analysis for ranking[J]. European Journal of Operational Research, 2018, 266(3): 990–999.
- [30] Aydın U, Karadayi M A, Ülengin F. How efficient airways act as role models and in what dimensions? A superefficiency DEA model enhanced by social network analysis[J]. Journal of Air Transport Management, 2020, (82): 101725.
- [31] 钟韬, 彭勤科. 基于社会网络分析的投资组合优选方法[J]. 系统工程理论与实践, 2015, 35(12): 3017–3024.
- Zhong Tao, Peng Qinke. Portfolio selection method based on social network analysis[J]. Systems Engineering: Theory & Practice, 2015, 35(12): 3017–3024. (in Chinese)
- [32] Ang S, Zheng R, Wei F, et al. A modified DEA-based approach for selecting preferred benchmarks in social networks[J]. Journal of the Operational Research Society, 2021, 72(2): 342–353.
- [33] Kang H J, Kim C, Choi K. Combining bootstrap data envelopment analysis with social networks for rank discrimination and suitable potential benchmarks[J]. European Journal of Operational Research, 2024, 312(1): 283–297.
- [34] 魏权龄. 数据包络分析方法[M]. 北京: 科学出版社, 2004.
- Wei Quanling. Data Envelopment Analysis[M]. Beijing: China Science Publishing & Media Ltd. (CSPM), 2004. (in

Chinese)

[35] 范小云, 荣宇浩, 王 博. 我国系统重要性银行评估: 网络层次结构视角[J]. 管理科学学报, 2021, 24(2): 48 – 74.

Fan Xiaoyun, Rong Yuhao, Wang Bo. Identifying systemically important banks in China: A network hierarchy structure perspective[J]. Journal of Management Sciences in China, 2021, 24(2): 48 – 74. (in Chinese)

Identify benchmarks from a large-scale sample: A social network DEA

AN Qing-xian^{1, 2}, WANG Ping³, WEN Yao^{4*}

1. School of Business, Central South University, Changsha 410083, China;

2. School of Economics, Hefei University of Technology, Hefei 230601, China;

3. School of Management, Hefei University of Technology, Hefei 230009, China;

4. School of Economics and Management, Fuzhou University, Fuzhou 350108, China

Abstract: To promote positive values and foster a good social atmosphere, all walks of life launch some benchmark selection activities, such as “the most beautiful doctor” and “the most beautiful teacher”. These role models (or benchmarks) tend to be a minority among numerous candidates. However, traditional methods have difficulty evaluating tens of thousands of candidates and do not consider the learning relationships and specific advantages among candidates. Thus, it is difficult to realize benchmark identification by traditional methods. To reduce the workload of benchmark identification, this study proposes a social network data envelopment analysis method for identifying benchmarks from a large-scale sample, by combining the data envelopment analysis (DEA) with social network analysis (SNA). First, a pairwise evaluation process is constructed using the DEA method to explore the efficiency state of each decision-making unit (DMU) relative to another DMU, and a pairwise evaluation matrix that reflects the reference relationship between any two DMUs is obtained. Then, a social network is built based on this matrix, and role-model DMUs in a large sample are identified by comparing the in-degree centrality values of all DMUs. Next, the specific advantages of the selected benchmarks are analyzed based on the interpretation of weights in the DEA model. Finally, an experiment is carried out using online diagnosis data from 10 418 doctors across 55 departments on The Chunyu Doctor Platform to verify the social network data envelopment analysis method.

Key words: data envelopment analysis; large-scale sample; benchmarks; social network analysis; pairwise evaluation