



第五章

考虑用户心理状态转换的职位推荐方法

当深入审视数智化时代的运行逻辑时，可以发现，以大数据和智能方法为核心的营销决策范式，其应用领域远不止于电子商务情境，而是凭借其强大的普适性，正在重塑各行各业的价值创造方式。广义而言，任何存在需求匹配的在线平台，其运营机制本质上都蕴含着营销的内核，即将最适合的信息、产品或服务高效地传递给最需要的用户。网络招聘平台便是一个典型例证。在此情境下，电商中的“商家—产品—消费者”关系被映射为“企业—职位—求职者”的新型互动结构，企业通过发布职位广告来展示其人才需求，实则是在营销职位这一特殊产品，而求职者则通过投递简历来表达其个人偏好与能力。在此过程中，海量的职位与求职者同样构成了一个复杂且动态的双边市场。与电商平台面临的信息过载困境类似，网络招聘平台上的求职者极易在纷繁的职位中迷失方向，企业也苦于无法精准筛选出其目标人才。因此，该情境下的职位推荐亦可被理解作为一种至关重要且极具价值的数智化营销决策。本章研究即从广度上拓展前文视角，旨在设计一种考虑用户心理状态转换的职位推荐智能方法。通过解析求职者的历史行为及其隐含的心理动态，该方法能够为每位求职者推荐其最可能感兴趣且符合条件的高潜力职位，从而实现精准的人岗匹配，提升求职者体验，降低企业招聘成本，并增强平台用户黏性和商业价值。

5.1 引言

在日趋激烈的“人才大战”背景下，招聘已成为人力资源管理的关键环节（Williamson et al., 2010）。招聘方借助多种渠道，力求吸引高素质求职者加入合适岗位，而求职者则致力于找到自身满意的工作。随着信息技术的快速发展，网络招聘渠道对招聘方和求职者的重要性日益凸显。市场调研显示，2025年全球网络招聘市场的规模约为557亿美元，预计未来5

年的复合年增长率为12.2%，至2030年年底，该市场规模将达到1027亿美元^①。一般而言，网络招聘平台允许招聘方发布职位招聘广告并查看感兴趣的求职者信息，求职者则可在平台上发布自己的简历，并将简历定向投递至心仪的职位。国内外知名的大型网络招聘平台包括LinkedIn、Monster、51Job、猎聘网等。网络招聘凭借其便利性和广泛覆盖，显著改善了市场双方的决策过程，具体体现在降低成本、提高效率、增强信息流动以及扩大职位或求职者的选择范围等方面。然而，这种新型招聘渠道也面临着传统线下招聘中不常见的问题和挑战，包括信息过载、隐私保护，以及因缺乏互动而导致的判断偏差等（Brandao et al., 2019）。

从现实问题来看，信息过载通常表现为用户个体层面的认知过载。对于求职者而言，求职是一项对其产生重大且长期影响的重要决策，因此求职者希望掌握充分的职位信息，以便慎重做出最优决策。然而，由于信息处理能力有限，网络招聘平台上海量的职位信息会给求职者带来沉重的处理负担，导致求职决策效率降低，甚至干扰决策质量（Roetzel, 2019）。事实上，求职者希望在决策过程中获取更多个性化信息，从而只需对那些高度相关的职位信息进行处理和反馈。作为新兴的在线招聘与求职渠道，网络招聘平台通过提供优质的针对性服务和高级信息管理工具，提升用户满意度并创造收益。为应对信息过载问题，平台设计了标签筛选、搜索引擎和推荐系统等嵌入式的解决方案。标签筛选允许求职者依据企业类型、行业、城市等基本属性对职位进行初筛，但仅限于有限的通用职位属性。搜索引擎则根据求职者输入的关键词在后台数据库中进行文本匹配，返回符合条件的职位信息。相比之下，推荐系统能够根据平台掌握的求职者和职位信息，主动为求职者优先呈现其可能感兴趣的职位，个性化程度更高。在标签筛选和搜索引擎难以充分解决信息过载的情况下，推荐系统通常能实现更优的信息匹配效果，因此成为学术界和工业界关注的重点，推动了网络招聘平台技术开发和方法的持续改进。

总体而言，求职者的在线求职过程具有两个不同于其他决策情境（如购物决策）的主要特点。

（1）求职决策存在明确的目标导向（Creed et al., 2009）。求职者以成功找到适合自己的满意职位为目标，在决策过程中表现出较为专注和审慎的行为，且通常不会出现多个决策过程混杂的情况（Da Motta Veiga and Turban, 2014）。在目标导向的驱动下，求职者的心理状态会持续发生

① <https://www.researchandmarkets.com/reports/5337633>.

转换,使得其感兴趣的、认为适合自己的职位也随之动态变化(Osborn, 1990)。人力资源领域的研究表明,求职者的心理动态由连续且有系统性的内在序贯状态构成,即准备状态和活跃状态(Blau, 1993),并由此驱动其呈现出外在多变的求职行为模式。具体而言,在准备状态下,求职者进行职位信息的广泛查找,意图从大量分散信息中快速掌握职位类型、机会质量和潜在适合度等,从而获得可能投递的候选职位集合。活跃状态则意味着集中查找,求职者经过筛选后仅聚焦于少量相近的感兴趣职位的详细信息,期望深入了解职位的具体要求、待遇和申请过程等(Albert, 1966)。因此,求职者在准备状态下兴趣较为分散,关注的职位之间可能差异较大;而在活跃状态下则兴趣逐渐收敛,可能只与最具吸引力且最适合自己的相似职位产生交互。

(2) 求职者决策本质上是一个双向选择的过程。鉴于职位和求职者均为稀缺资源且具有明显的排他性(Hong et al., 2013b),求职者在决策过程中既要考虑选择的职位相较其他职位是否更具吸引力,也要考虑自身是否相较其他竞争者更能满足职位要求,从而保证较高的求职成功率和整体效率(Malinowski et al., 2006)。此外,双向选择对求职过程的影响不仅体现在求职者在某个时点过滤掉不感兴趣或无法胜任的职位,更会驱动其持续进行自适应和自我调节(Kanfer and Bufton, 2018; Wanberg et al., 2020)。具体而言,如图5.1所示,求职者在整个求职过程中会不断收到来自招聘方的信息反馈,这种反馈既可能是明确的接受或拒绝,也可能是申请职位后缺乏明确反馈所隐含的负面信号。因此,求职者的心理状态和求职行为会持续受到当前已接收反馈的影响,令其自尊和自我效能产生变化,也容易因短期内反馈不佳而发生序贯状态的回退(Barber et al., 1994)。已有实证

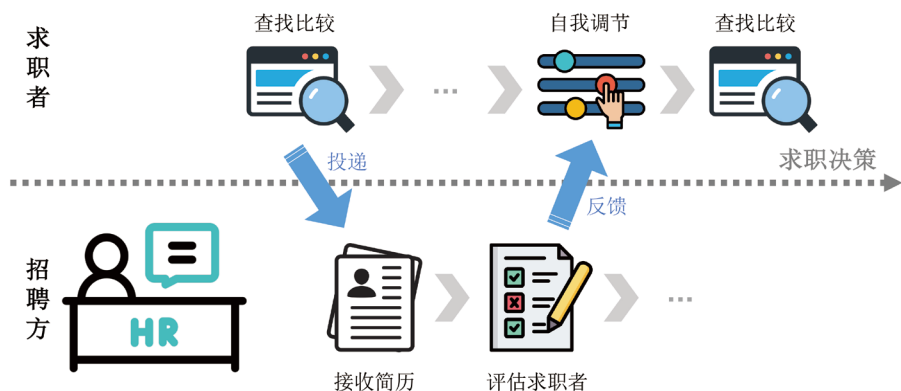


图5.1 在线求职决策过程示意图

研究表明,若求职者已处于活跃状态末期但仍未被聘用,则为了推进求职过程,其很有可能重启整个求职决策,退回至准备状态,重新查找比较一些不同类型或更能胜任的职位,以尽快实现成功求职(Saks and Ashforth, 2000)。例如,某求职者向咨询公司投递简历后,若在一段时间内并未收到满意反馈,则很可能重新开始考虑金融机构、互联网企业等其他同样符合其专业、兴趣和能力的相关职位。

然而,现有的职位推荐方法大多是将电子商务等情境下较为成熟的方法迁移至网络招聘情境中应用。尽管这些方法能够在一定程度上生成有效的职位推荐,但尚未能充分将网络招聘和求职决策的情境特点融入推荐方法的设计,使得方法在技术创新及实践应用上仍有提升空间,其推荐表现也有待改进。具体体现在以下两方面。

(1) 现有方法对求职过程的动态性和阶段性刻画不足。绝大多数职位推荐方法属于静态推荐,即根据求职者的基本信息进行静态画像(Drigas et al., 2004),或综合求职者在网络招聘平台上产生的所有历史行为来推断其偏好(Lee and Brusilovsky, 2007)。然而,在求职者心理状态持续转换的情况下,静态推荐方法很难准确捕捉其真实兴趣,导致职位推荐效果受限。事实上,尽管求职者的可观测行为特征会随心理状态发生变化,但在谨慎决策的过程中,其浏览职位广告、投递简历等各次行为之间并非相互独立,而是具有较强的时序关联性。并且,由于求职者在一时间段内通常只专注于寻找一份工作,其在线行为记录中的“噪声”相对较少,这为动态方法从求职者可观测行为序列中学习清晰的动态模式和规律,并据此预测其后续感兴趣的职位提供了便利。近年来,已有研究将循环神经网络(Quadrana et al., 2017)、长短时记忆模型(Liu et al., 2016; Benabderrahmane et al., 2017)、图神经网络(Wang et al., 2020)等深度学习方法引入动态职位推荐任务,但由于未融入求职者心理的内在特征和机理,这些方法仅能获得浅层、稀疏的行为规律,难以甄别出不同求职状态并据此提供针对性的职位推荐。此外,深度学习方法复杂隐层结构和大量参数令模型的可解释性较弱,决策者无法直观理解职位推荐的过程和结果,难以从中获得明确的管理启示,从而限制了方法对现实管理决策的支持。

(2) 现有方法依赖简历信息来捕捉求职过程中的双向选择。从人职双向选择或匹配的视角出发,已有研究关注互惠推荐方法(reciprocal recommendation),即在职位推荐中同时考虑职位条件满足求职者兴趣的程度,以及求职者能力满足职位要求的程度。这些方法通过互惠值排序来

优先推荐最适合求职者的职位 (White et al., 2012), 综合面向求职者的职位推荐模型和面向招聘方的求职者推荐模型的结果 (Yu et al., 2011a), 或约束每个职位被推荐给求职者的次数以缓解竞争强度 (Keim, 2007)。然而, 现有的互惠职位推荐方法大多基于求职者简历中的基本信息和职位广告中的要求来计算求职者在多大程度上适合相应职位。实际上, 求职者的许多潜在能力和性格特征等并不会完全体现在简历信息中, 这使得方法难以全面、准确地捕捉其内在自我评估的结果及其与职位间的真正契合度。同时, 简历信息通常不会即时或频繁更新, 因此无法体现双向选择所驱动的求职者持续自适应和自我调节过程。特别是在发生序贯状态回退时, 推荐方法不仅需刻画求职者的动态职位兴趣, 还应考虑其与各职位之间适配度的变化, 而这些难以从简历信息中反映出来。尤其值得关注的是, 近期关于求职者简历信息隐私保护的问题日益受到学界和业界的重视, 来源于网络招聘平台的简历数据的公开性和可用性均逐步受限, 使得依赖求职者简历信息的职位推荐方法在现实管理应用中面临多方面的挑战。

为应对上述问题, 本研究考虑网络招聘情境特点, 从求职者序贯心理状态的视角出发, 提出融入混合相似度的动态贝叶斯网络 (dynamic Bayesian network with hybrid similarity, DBN-HS) 方法用于职位推荐。具体而言, 本方法在动态贝叶斯网络的基础结构上构建模型, 刻画求职者决策过程中在准备与活跃两个心理状态间的转换以及其相应的兴趣变化, 通过求职者在网络招聘平台上可观测的职位广告浏览行为序列, 学习其中潜在的动态模式。根据序贯状态的相关理论发现, 方法设计可量化的隐层识别策略, 通过约束各状态与职位相似度和兴趣变化有关的条件概率, 为模型中的隐状态赋予显式的概念标定。特别地, DBN-HS由预先设定的识别策略监督动态贝叶斯网络的训练过程, 对求职过程的阶段性进行先验的可解释建模, 并据此合理缩小模型参数估计的搜索空间, 从而提高学习效率。在此基础上, 方法捕捉有关求职者心理动态的普适性内在规律, 及时追踪探测其所处的状态, 再结合推断和预测结果生成有针对性的职位推荐。同时, 模型从求职者实际发生的行为 (即在线职位交互) 中学习其对自身潜在特征和能力的评估, 避免采用敏感的简历信息。考虑求职者依据招聘方即时信息反馈而进行自适应的过程, 模型允许求职者从活跃状态退回至准备状态, 并重新生成职位兴趣, 从而捕捉求职者在两个状态下的不同需求, 也将双向选择的情境特点更灵活地融入求职过程的刻画, 以提升职位推荐的准确性和现实应用价值。

此外, 因求职者在两种序贯状态下对职位兴趣的收敛程度不同, DBN-

HS方法采用其连续两次浏览职位之间的相似度来衡量其兴趣收敛的程度。由于单纯根据职位属性计算相似度难以确定各属性的权重，而基于求职者交互行为衡量的职位相似度又太过稀疏，DBN-HS将两种思路进行融合。首先采用普适性较高的随机森林模型，根据共同交互视角的相似度学习各属性相似度在职位整体相似度中的权重，再用属性相似度对稀疏的行为视角相似度进行补充校正。该策略既能保证最终得到的职位相似度足够准确且稠密，在求职者行为记录之外引入职位属性信息，有效构建整体的混合推荐方法，也能通过兼顾职位的客观描述及其对求职者的要求属性，进一步融合求职过程的双向选择特征。基于此，职位相似度的计算可根据随机森林模型学到的各属性相似度的重要性分值进行事后解释，而动态贝叶斯网络模型所刻画的机制也与经典理论研究发现的求职者序贯心理状态相对应，令提出的DBN-HS方法具有较强的内在可解释性，从而为网络招聘市场中的各方决策者提供有价值的管理启示。

5.2 相关工作

5.2.1 职位推荐方法

一般而言，根据不同的建模假设和输入信息，通用推荐方法整体上可分为三类：基于内容的推荐（content-based）、协同过滤推荐（collaborative filtering）和混合推荐（hybrid）（Adomavicius and Tuzhilin, 2005）。与此分类体系相适应，职位推荐领域也涵盖三类方法。在基于内容的职位推荐方法中，已有研究较为关注求职者简历的文本信息提取，通过对简历内容进行自动分区（Yi et al., 2007）、标签提取（Yu et al., 2005）等方式来刻画求职者的兴趣或属性。也有研究基于求职者简历和网络招聘广告间的文本语义匹配进行职位推荐，通过计算能力（Kumaran and Sankar, 2013）、性格（Faliagka et al., 2014）等非结构化维度的相似性为求职者匹配最适合的职位。在协同过滤职位推荐方法中，Rafter et al.（2000）根据平台记录的历史求职行为计算求职者之间的相似度，Wu et al.（2014）则据此计算职位间的相似度，进而为求职者推荐其可能感兴趣的职位，即采用基于记忆的协同过滤方法生成职位推荐结果。此外，也有基于贝叶斯模型的协同过滤方法，根据模型预测到的求职者对职位感兴趣的概率提供推荐（Färber, 2003）。目前，职位推荐中采用的混合推荐方法以简单结合的方式为主，例如，Heap et al.（2014）提出的级联推荐先通过协同过滤方法对候选职位进行初步排序，再用基于内容的方法选出强推荐的职位；Hong et

al. (2013a) 则将各推荐方法得到的职位排序加权求和, 进而输出最终的推荐结果。

除上述三类方法外, 因职位推荐的问题场景具有一定特殊性, 互惠推荐也成为较有特色的一类职位推荐方法。例如, 通过整合求职者兴趣与职位条件、求职者能力与职位要求的两类相似度来计算互惠值, 并将互惠值更高的职位优先推荐给求职者 (Yu et al., 2011a)。由于在网络招聘情境下, 求职者和职位均为稀缺资源, 平台上的招聘方和求职者持续进行双向选择, 因此职位推荐的目标不仅局限于为求职者提供感兴趣且满足其需要的职位建议, 还应考虑求职者的能力是否符合职位的要求 (Malinowski et al., 2006)。从求职者的角度来看, 推荐其无法胜任的职位容易错误引导其对求职结果的期望, 也可能令其错失真正适合的工作机会。而对招聘方而言, 不考虑双向选择的职位推荐会使部分招聘方接收到大量不满足要求的简历, 因处理无效信息而降低人才招聘效率, 同时另一些招聘方则完全接收不到简历, 即便可能存在非常适合其职位的求职者。长期看来, 若人职匹配的成功率下降, 平台两方用户均会降低对网络招聘平台的满意度和使用黏性。实际上, 每次职位推荐的潜在影响在于, 它会降低竞争求职者获得该职位的成功率, 以及竞争职位获得合适人才的概率 (Hong et al., 2013b), 因此职位推荐应尽可能以成功的人职匹配为目标, 帮助实现全局最优的资源分配。此外, 因人力资源领域的研究长期积累了大量有关招聘或求职的定性发现, 一些基于知识的职位推荐方法 (knowledge-based) 也较为常用。例如, Ali and Rajamani (2012) 根据招聘方专家所提供的关联规则进行推荐; Paparrizos et al. (2011) 采用决策表和朴素贝叶斯模型构造的分类器, 预测求职者下一份工作的特征。

从方法演化的角度来看, 最早的自动化职位推荐方法产生于2000年左右, 即Rafter et al. (2000) 提出的考虑求职者相似度的协同过滤方法。2004年起开始出现一些基于内容的职位推荐方法 (Drigas et al., 2004), 关注求职者简历和职位招聘广告的内容提取, 并据此进行人职分析和匹配。2007年, Meo et al. (2007) 提出基于主动调试的职位推荐方法, 通过求职者对历史职位推荐的采纳率来动态调整推荐数目和内容, 发现采用动态策略调整的职位推荐系统能够获得更优的求职者反馈。同时, 基于知识的方法也将领域内已经检验的影响因素纳入分类器等机器学习模型, 从而预测并推荐求职者后续可能感兴趣的职位 (Lee and Brusilovsky, 2007)。2010年, 互惠职位推荐方法被提出 (Pizzato et al., 2010), 求职者与职位之间的双向选择得到广泛重视。2011年, Lops et al. (2011) 将社交信息引入

职位推荐任务，使求职者本人及其所在社团的特征共同构成求职者画像，实现更准确的职位推荐。2013年出现了较为成熟的混合推荐方法（Hong et al., 2013a）以及基于关系图的职位推荐方法（Lu et al., 2013），后者通过有向、加权、多连接的图结构来刻画职位之间、求职者之间基于内容和基于互动的关系，并采用图模型相应算法预测并推荐求职者感兴趣的职位。深度神经网络在2016年左右被引入网络招聘情境下的职位推荐问题（Liu et al., 2016）。职位推荐中的深度学习通常用于对求职者简历或职位招聘广告等文本信息进行表示学习（Qin et al., 2020; Mishra and Rathi, 2021），或刻画求职者点击流等在线行为记录中的复杂时序关系（Benabderrahmane et al., 2017; Wang et al., 2020）。此外，近期的职位推荐方法研究也致力于采用主题模型提取求职者或职位所需的技能（Bansal et al., 2017），通过极端梯度提升等树模型预测人职交互（Sato et al., 2017），以及借助图模型挖掘求职者间、职位间或人职间的潜在关系（Shalaby et al., 2017; Yang et al., 2017）等，从而改进职位推荐的效果。

从技术演化的过程可以发现，尽管互惠推荐等方法在一定程度上融入了网络招聘的情境特点，大多已有的职位推荐方法仍主要由其他领域（如电子商务等）较为成熟的推荐方法迁移而来。迁移过程中，甚至会针对具体可获取的数据对方法进行情境化适应，但未能充分基于与职位推荐问题场景相关的经典理论和独有特点来设计有针对性的新方法。一方面，若缺乏足够的先验知识融入和情境特点捕捉，则职位推荐方法刻画并预测求职者兴趣、能力和求职行为的准确性会较弱。尤其现有职位推荐仍以静态方法为主，而鉴于求职者在网络招聘平台上持续发生的行为已能够被充分捕捉和记录，且求职决策本身也是具有较强连续性和阶段性的动态过程，有效的动态职位推荐方法在准确性上呈现相对优势。另一方面，若无法保证对问题场景的针对性，则职位推荐方法的作用过程及输出结果所能为现实决策者提供的解释性和管理启示也会受到局限。此外，很多现有职位推荐方法高度依赖于求职者的简历信息，在数据隐私问题日益敏感的情况下，这些方法的可实现性和可扩展性均容易受到制约。

本章研究考虑用户求职过程中的序贯心理状态及其相应的信息查找行为模式，根据求职者历史的职位交互记录，融合属性信息以校正职位相似度，进而采用动态贝叶斯网络对求职者的职位兴趣进行准确刻画和预测。方法通过序贯状态的相关理论发现来设计隐层识别策略，为模型中的隐变量标定显式含义，以增强职位推荐的可理解性。同时，方法仅基于求职者行为和职位属性信息生成推荐，通过描述心理状态转换来捕捉求职者决策

过程中的目标导向和自我调节,且兼顾职位的客观属性及其对求职者的要求属性,从而在规避使用求职者简历中隐私信息的前提下,保证足够有效的职位匹配。

5.2.2 求职者序贯心理状态

完整的求职决策过程通常由求职者浏览职位招聘广告、投递简历、线上线下参加面试等一系列行为构成。通过分析这些行为及其对应的求职者心理状态,可以更好地辅助求职者决策,也为企业招聘提供现实管理启示。由于传统招聘易受到时间和地理因素的制约,且招聘方和求职者所能接触到的信息均十分有限,随着互联网技术的高速发展,网络招聘平台逐渐成为求职者找工作的主要渠道(Williamson et al., 2010)。网络招聘平台能够显著提升求职与招聘的效率和便利性,但因平台实现的主要功能是招聘方发布职位广告和求职者投递简历等,求职者从参加面试到最终入职等线下行为难以被平台捕捉。即便如此,仅记录求职者的在线求职行为,对于理解其决策及其潜在的求职需求、职位兴趣等也大有帮助。原因在于,与在线购物等情境不同,职位相较于产品或服务所带来的效用影响通常更大且更深远,这使得求职者的决策过程更为严肃,具有更明确的目标导向(Creed et al., 2009; Da Motta Veiga and Turban, 2014),也较少出现同时寻找多份工作的情况,从而在线行为中的“噪声”相对较少,有助于数据分析模型通过可观测的变化行为模式来捕捉求职者潜在的心理动态,并生成对其后续行为的有效预测。

尽管求职决策是由目标驱动连续过程,但求职者的心理状态并非一成不变,相应地,其行为特征也会在求职过程中呈现出明显的动态变化(Osborn, 1990)。关于求职者决策过程中所经历的心理状态及其对行为的影响,目前尚未有非常明确的定论。人力资源领域的已有研究发现,求职者决策存在学习效应,即求职者会在求职过程中逐渐掌握更具效率和效果的方式,从而改变其自身行为,以期获得更好的求职结果(Barber et al., 1994)。同时,求职决策过程也可能受到情绪反馈效应的影响,即求职者经历较大压力、沉浸于沮丧情绪时,容易因焦虑而改变其求职行为的特征(Eden and Aviram, 1993)。而应用最广的求职者决策过程分析角度则是考虑其中的序贯效应(sequential effect),即认为求职过程是由连续且有系统性的某种状态序列构成,如准备状态和活跃状态(Blau, 1993)。具体而言,求职者在准备状态下意图查找大量信息以获得候选职位集合,在活跃状态下则进一步在脑海中进行筛选,获取感兴趣的少量职位的详细信息,

并投递简历申请。这两种心理状态驱动了求职者从广泛查找（*extensive search*）至集中查找（*intensive search*）的行为转变（Albert, 1966）：广泛查找涉及大量分散的职位信息，以便求职者快速掌握职位类型、机会质量和潜在适合度等；集中查找则聚焦于少量相近职位的信息，求职者通过深入了解职位具体要求、待遇和申请过程等，对比选择出有意愿申请的职位。因此，处于准备状态的求职者对职位的潜在兴趣可能非常分散，使得其关注和交互的职位之间通常差异较大；而进入活跃状态后，求职者兴趣逐渐收敛，只明确考虑最具吸引力、最适合自己的一些较为相似的职位，并产生交互行为。

然而，序贯效应并不意味着现实中求职者只会按照顺序依次经历两个状态，进而结束其求职过程。与一般性决策过程不同，求职决策是一种自适应过程，其间求职者会持续进行自我调节（Kanfer and Bufton, 2018; Wanberg et al., 2020）。原因在于，招聘情境下的“产品”（即职位）是稀缺资源，同一职位的不同求职者之间存在很强的竞争性和排他性（Hong et al., 2013b），使求职过程本质上成为双向选择。换言之，最终的求职结果不仅取决于求职者对职位的兴趣，也依赖于求职者能够胜任职位的程度，即由招聘方反馈所体现的求职者竞争力。在这种情境下，求职者会持续受到当前信息反馈的影响，容易因短期内反馈不佳等发生心理状态的回退。已有实证研究表明，处于活跃状态末期的求职者若仍未被聘用，则可能重启整个求职决策过程，即退回至准备状态，重新对差异较大的职位信息展开广泛查找（Barber et al., 1994; Bertrand et al., 2010）。由此，考虑到求职者在两种心理状态下的不同信息需求以及在两种状态间转换的可能性，网络招聘平台在为其推荐或匹配职位广告时，应通过灵活建模其行为特征来及时追踪求职者所处状态的变化，进而准确预测求职者在当前心理状态下最有可能感兴趣的职位，并提供可理解的洞见。

5.3 职位推荐方法

本章提出一种新的职位推荐方法——融入混合相似度的动态贝叶斯网络（DBN-HS）方法，整体框架如图5.2所示。结合网络招聘的情境的特点，本研究从求职者在序贯状态下的心理与行为机制出发进行方法设计。通过基于随机森林模型的职位相似度校正和基于动态贝叶斯网络的求职过程刻画，推断求职者潜在的心理状态和兴趣，并为其推荐可能适合的职位。具体而言，首先利用训练集中职位在各属性上的相似度及求职者历

史行为中的共同交互标签训练随机森林模型，并采用无权重有放回抽样（bagging）和分类回归树算法（classification and regression tree, CART）学习随机森林的模型参数。在此基础上，DBN-HS基于已知的职位间属性相似度，通过加权平均方式对稀疏的共同交互相似度进行校正。校正后得到的混合相似度与求职者的历史行为序列共同输入动态贝叶斯网络，采用最大期望（expectation maximization, EM）和最大后验概率（maximum-a-posteriori, MAP）估计算法来学习求职者决策过程中的序贯状态转换、职位兴趣变化、交互行为生成等参数。此外，方法设计了隐层识别策略以监督模型训练过程，从而对序贯状态进行显式标定。最终，基于训练得到的动态贝叶斯网络，DBN-HS采用维特比式推断策略解码测试集中求职者后续最可能处于的序贯状态及对应的职位兴趣，并通过概率预测生成个性化的职位推荐结果。

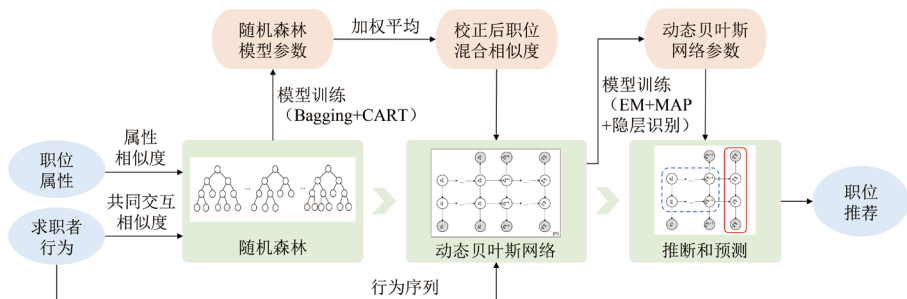


图5.2 DBN-HS方法框架

5.3.1 问题定义

假设某网络招聘平台记录了一段时间内求职者与职位之间的可观测交互行为。给定职位集合 J 和职位属性集合 R ，以及职位 $\forall j \in J$ 在属性 $\forall r \in R$ 上的取值 j_r 。例如，职位属性 r 为企业性质，职位 j 在该属性上的取值 j_r 为“中外合营”。对于某位求职者 u ，其求职过程可表示为序列 $(j_u^1, (j_u^2, s_u^2), \dots, (j_u^{N_u}, s_u^{N_u}))$ ，其中 j_u^n 表示求职者 u 在第 n 次行为所交互的职位ID， s_u^n 表示该职位与其上一次交互职位之间的混合相似度，即 $s_u^n = \text{hsim}(j_u^{n-1}, j_u^n)$ ， N_u 则为求职者 u 的行为记录总数。对于 $\forall n = 1, 2, \dots, N_u$ ，有 $j_u^n \in J$ ；对于 $\forall n = 2, 3, \dots, N_u$ ，有 $s_u^n \in [0, 1]$ 。将求职者 u 在第 n 次行为时的心理状态记为 z_u^n ，对应兴趣记为 i_u^n ，本研究的目标是：根据求职者历史发生的第1~ N_u-1 次职位交互行为，推断其在各次行为中的心理状态和对应兴趣，进而为其在第 N_u 次交互时推荐可能感兴趣的职位。

5.3.2 职位相似度校正

为刻画求职者在不同心理状态下对职位兴趣的收敛性，DBN-HS方法预先计算求职者连续两次交互职位之间的相似度，将其作为反映其兴趣收敛性的可观测变量，并与职位序列共同输入动态贝叶斯网络，以切分和甄别求职行为中体现的不同序贯状态。与基于内容和协同过滤两类推荐方法的核心假设一致，职位间相似度也可从基于内容的相似度和基于共同交互的相似度两个视角进行度量。

在职位推荐的相关研究中，根据具体属性计算人职间、职位间、求职者间相似度的方法较为常见，其中最常用的度量指标为余弦相似度（Hong et al., 2013b; Heap et al., 2014）。然而，余弦相似度难以捕捉不同属性对职位相似度计算的重要性差异。例如，求职者对于薪资水平、工作地点等职位属性可能比对企业投资阶段、招聘形式等更为敏感，因此不加区分的加权方式往往无法反映求职者对职位相似性的真实感知，导致推荐结果不够准确。鉴于此，已有研究在等权重的职位相似度计算基础上进行改进，如令用户或专家根据经验对不同职位属性区分加权（Lu et al., 2013），采用层次分析法或网络分析法进行属性重要性两两比较（Faliagka et al., 2012），或通过设置理想选项来求解各属性相似度在整体相似度中的权重（Kelemenis and Askounis, 2010）。但这些方法依赖人工主观判断，其科学性和代表性难以验证，限制了在大规模应用中的有效性。

基于共同交互的职位相似度度量较好地解决了加权困难问题。该方法假设同一求职者交互过的不同职位之间具有较高的相似性（De Pressemier et al., 2016; Pacuk et al., 2016），从实际行为中挖掘求职者对职位相似度的综合认知，以更全局的角度判断求职者评估的职位间相似关系。然而，网络招聘平台上的职位和求职者数量庞大，导致求职者与职位的交互记录整体稀疏（Heap et al., 2014; Bansal et al., 2017），使得基于共同交互的相似度估计存在偏差——无法区分“无共同交互”是因不相似还是因未被同一求职者看到。

综合考虑上述两种职位相似度计算方法的优势和局限，本研究提出基于随机森林模型的职位相似度校正策略。先分别计算训练集中两个职位在各属性上的相似度，再将其输入随机森林模型进行学习，训练标签为这两个职位基于求职者共同交互行为计算的相似度值。由此得到各属性相似度与职位整体相似度之间的关系，对于任意两个职位，无论是否存在共同交互记录，均可利用训练后的随机森林模型预测其相似度。最后，将该预测

值与已知的共同交互相似度加权平均,得到校正后的两职位间相似度。该策略既避免了测试集中大量相似度为0的情况,又保证了相似度度量源自全局求职者的认知视角。

基于随机森林模型的职位相似度校正策略受启发于相关研究中人职相似度的计算方法,即将属性相似度与交互相似度相结合。已有方法多采用线性回归(Malherbe et al., 2014)或逻辑回归(Wu et al., 2014)进行拟合,但职位属性分布不确定且可能存在较强依赖(如职称等级和薪资水平),线性模型的全部假设条件难以满足,拟合准确性有限。本方法采用灵活性更强的随机森林模型,能够较好地处理属性相似度与职位整体相似度之间可能存在的非线性关系,并有效平衡标签不均衡的数据集,从而获得更准确的拟合结果。为避免过度稀释共同交互中包含的信息,DBN-HS仍保留基于共同交互的相似度信息,仅以一定的权重用随机森林模型的预测值进行校正,兼顾职位相似度的稠密性和认知一致性。特别地,为突出网络招聘情境的双向选择特征,输入至随机森林模型的属性既包含职位的客观描述,也包含其对于求职者的要求,从而构建职位之间的混合相似度,为求职者推荐感兴趣又潜在适合的职位。

具体而言,给定所有职位的 $|R|$ 个属性,DBN-HS首先根据数据类型计算任意职位 j 和职位 k 在各属性 $r(r=1, \dots, |R|)$ 上的相似度 $\text{fsim}_r(j, k)$ 。令 j_r 和 k_r 分别为两个职位在属性 r 上的取值,所有属性上的相似度取值范围均为 $[0, 1]$ 。

(1) 连续属性或有序分类属性。对于职位薪资水平、工作年限要求等连续属性,或企业规模、职称等级等有序分类属性,基于属性取值之差的绝对值,采用式(5.1)来计算两个职位在该属性上的相似度。

$$\text{fsim}_r(j, k) = 1 - \frac{|j_r - k_r|}{\max(r) - \min(r)} \quad (5.1)$$

其中, $\max(r)$ 和 $\min(r)$ 分别为所有职位在属性 r 上取值的最大值和最小值。

(2) 无序分类属性。对于职位类型、企业性质等无序分类属性,可直接根据属性取值是否相同,采用式(5.2)计算两个职位在该属性上的相似度。

$$\text{fsim}_r(j, k) = \begin{cases} 1, & j_r = k_r \\ 0, & j_r \neq k_r \end{cases} \quad (5.2)$$

(3) 层级分类属性。除上述属性外,职位也有职能分类、所属行业、所在地区等具有多层级概念的分类属性,且不同招聘方在平台上填写的属性取值可能处于不同层级。例如,所在地区分为省、区、市等层级。对于

此类属性，本方法借鉴已有研究提出的基于概念层次结构的相似性度量方式来计算两个职位在该属性上的相似度（Bizer et al., 2005）。令取值 j_r 和 k_r 在该属性概念结构中所处的层级分别为 $l(j_r)$ 和 $l(k_r)$ ，二者的最临近公共概念取值为 $c_{j,k}$ （如朝阳区和海淀区的最临近公共概念为北京市），二者的取值距离为 $d_r(j_r, k_r)$ ，则职位 j 和职位 k 在属性 r 上的相似度可由式(5.3)计算得到。

$$\begin{aligned} \text{fsim}_r(j, k) &= 1 - d_r(j_r, k_r) = 1 - (d_r(j_r, c_{j,k}) + d_r(k_r, c_{j,k})) \\ &= 1 - \left(\frac{1/2}{a^{l(c_{j,k})}} - \frac{1/2}{a^{l(j_r)}} + \frac{1/2}{a^{l(c_{j,k})}} - \frac{1/2}{a^{l(k_r)}} \right) \\ &= 1 - \left(\frac{1}{a^{l(c_{j,k})}} - \frac{1/2}{a^{l(j_r)}} - \frac{1/2}{a^{l(k_r)}} \right) \end{aligned} \quad (5.3)$$

其中，超参数 $a > 1$ （常用取值为2）。

经上述计算得到两个职位在各属性上的相似度后，DBN-HS基于训练集中求职者的历史职位交互行为记录，采用如式(5.4)所示的Jaccard系数计算职位 j 和职位 k 的共同交互相似度 $\text{isim}(j, k)$ 。

$$\text{isim}(j, k) = \frac{|U_j \cap U_k|}{|U_j \cup U_k|} \quad (5.4)$$

其中， U_j 和 U_k 分别为与职位 j 和 k 发生过历史交互的求职者集合。

进而，给定训练集中的职位集合 J ，对于 $\forall j, k \in J, r=1, \dots, |R|$ ，以 $\text{fsim}_r(j, k)$ 为特征、 $\text{isim}(j, k)$ 为标签，训练随机森林模型。随机森林（random forest）是一种集群分类器，由多个树型分类器构成（Breiman, 2001），能够克服单一决策树过拟合、易收敛到局部最优解等问题。其工作机制为：从训练集中有回放地抽取 M 个子集（大小约为原集的2/3），并建立决策树；节点分裂时采用分类回归树算法（Breiman et al., 1984），以基尼系数最小为原则优先选择节点特征，并按照节点特征划分以实现决策树的分支生长。为使不同决策树之间相关性较小并提升每棵树的分类精度，随机森林模型在进行节点分裂时并不考虑所有的输入特征，只随机选择几个特征参与比较，由此确定节点特征。在不剪枝的前提下，到达每个节点时均重复随机特征的选择和比较，直至所有数据均成为叶子节点。对于连续变量输出（如职位相似度），随机森林模型通常取 M 棵决策树的平均预测值作为最终结果。DBN-HS方法在训练随机森林的过程中，通过预实验来确定模型最优的超参数 M 。

根据随机森林模型的训练结果，可得到任意两个职位间相似度的预测值 $\text{psim}(j, k)$ ，并通过各输入特征在树中的平均精度下降和平均误差上升指标

来评估其在数据拟合中的重要性 (Breiman et al., 1984), 即确定各属性相似度对于职位整体相似度度量的相对贡献大小, 从而在保证获得较高预测精度的同时增强职位相似度计算的可解释性。最后, 采用加权平均的形式, 用随机森林输出的职位相似度预测值校正共同交互相似度, 即由式(5.5)计算职位 j 和职位 k 之间校正后的混合相似度 $hsim(j, k)$ 。

$$hsim(j, k) = w \times psim(j, k) + (1 - w) \times isim(j, k) \quad (5.5)$$

其中, w 和 $hsim(j, k)$ 的取值范围均为 $[0, 1]$ 。权重 w 代表预测值对共同交互相似度的校正比例, w 越大则校正比例越高。极端情况下, 若 $w=0$ 则意味着仅采用求职者共同交互记录来衡量职位间的相似度, 若 $w=1$ 则仅依据随机森林模型的预测结果来评估职位间相似度。 w 值可根据动态贝叶斯网络的训练和测试效果调优确定。

5.3.3 求职过程建模

研究提出的DBN-HS方法通过构建一个动态贝叶斯网络, 对求职者决策过程中的序贯心理状态、职位兴趣和交互行为的生成过程进行建模, 模型结构如图5.3所示。假设求职者集合为 U , 其中包含的求职者总数为 $|U|$ 。对于任一求职者 u , 第 n 次职位交互行为所对应的心理状态 z_u^n 和兴趣 i_u^n 均为不可观测的隐变量, 可随时间变化并由数据推断得到。可观测变量 j_u^n 和 s_u^n 分别代表求职者 u 第 n 次行为所交互的职位及其与第 $n-1$ 次交互职位间的混合相似度。

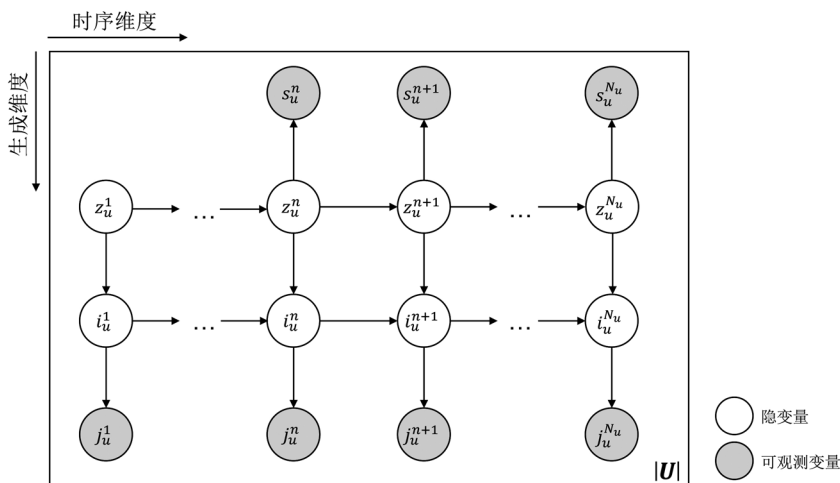


图5.3 刻画求职过程的动态贝叶斯网络

如图5.3所示, 模型描述的求职过程可分解为时序和生成两个维度。

(1) 时序维度。在横向的时序维度上,模型主要刻画求职者在决策过程中动态的序贯心理状态转换,即在“准备”和“活跃”两个状态之间转换(Albert, 1966; Blau, 1993)。由于求职者同时受到目标导向和自适应调节的影响(Barber et al., 1994; Creed et al., 2009; Kanfer et al., 2018),其前后两次职位交互行为可能从准备状态前进至活跃状态、从活跃状态退回至准备状态,或者停留在同一状态。模型捕捉心理状态间的时序依赖关系,不约束转换的方向或范围。同时,配合心理状态转换,动态贝叶斯网络也对求职者随时间的兴趣变化进行建模(Osborn, 1990),赋予其灵活的变化模式。在序贯状态和兴趣的心理动态驱动下,求职者在决策过程中也会产生一系列可观测的职位交互行为(Osborn, 1990)。由于求职者的心理状态转换灵活且因人而异,DBN-HS方法保留求职者行为序列中完整的时序关系,以便自动学习并准确捕捉其求职过程中可变且不等长的状态转换。

(2) 生成维度。在纵向的生成维度上,模型深入刻画求职者决策过程中不可观测的心理动态如何生成其每次可观测行为。一方面,求职者在准备状态下会广泛查找大量分散的职位信息,在活跃状态下则集中查找并筛选少量最适合的详细信息,导致准备状态下交互的职位更加多样,职位间相似度低于活跃状态。因此,模型捕捉求职者当前职位与上次交互职位间的混合相似度对心理状态的依赖性。另一方面,心理状态通过兴趣间接作用于交互职位(Osborn, 1990)。考虑到求职者对职位的需求和偏好具有天然的异质性,并非所有求职者在同一心理状态均有相同的职位兴趣分布,故构建的动态贝叶斯网络将兴趣描述为状态和产品之间的中介:准备状态下求职者兴趣较为分散且易变,活跃状态下兴趣相对聚焦且变化少(Albert, 1966; Blau, 1993)。兴趣决定求职者与不同职位进行交互的概率,即求职者更可能与其感兴趣的职位进行交互(Osborn, 1990)。由此,DBN-HS方法能够建模求职者心理状态对其兴趣变化的影响,以及动态兴趣对其交互职位的影响。

为能先验地实现对模型中心理状态和兴趣的合理概念化,DBN-HS方法设计了隐层识别策略,基于求职者序贯心理状态的相关理论,推衍出量化的约束条件,通过限定动态贝叶斯网络中部分条件概率来增强职位推荐的可理解性。具体包括:①令活跃状态下求职者相邻两次交互职位间的平均相似度始终高于准备状态,利用可观测的平均混合相似度差异来识别隐式的心理状态;②若同一求职者由准备状态转换至活跃状态,则令其职位兴趣更加集中;反之,若由活跃状态转换至准备状态,则令其职位兴趣更加分散。据此,可通过序贯心理状态识别隐式的兴趣分布,赋予其显式的管

理含义。

基于动态贝叶斯网络所描述的机制,生成求职者决策过程的形式化表达如下。求职者 u 初始处于心理状态 z_u^1 ,可由所有求职者初始状态的多项概率分布得到,即 $z_u^1 \sim P(z_u^1)$ 。接着,求职者的初始状态依据所有兴趣的多项分布生成其初始职位兴趣,即 $i_u^1 \sim P(i_u^1 | z_u^1)$ 。而后,基于所有职位的多项分布,由求职者 u 的初始兴趣生成其求职过程中第1次行为交互的职位,即 $j_u^1 \sim P(j_u^1 | i_u^1)$ 。对于求职者 u 从第 n 次推进至第 $n+1$ 次的行为:①鉴于求职者 u 在产生下次行为时可能转换至任一心理状态,故由第 n 次行为所处的状态依据多项分布生成第 $n+1$ 次行为的状态,即 $z_u^{n+1} \sim P(z_u^{n+1} | z_u^n)$ 。②求职者 u 决策过程中第 $n+1$ 次行为所处的状态和第 n 次行为对应的兴趣依据多项分布共同生成其第 $n+1$ 次行为的兴趣,即 $i_u^{n+1} \sim P(i_u^{n+1} | i_u^n, z_u^{n+1})$ 。③与初始职位交互的生成方式相同,由求职者 u 第 $n+1$ 次行为对应的兴趣生成该次行为交互的职位,即 $j_u^{n+1} \sim P(j_u^{n+1} | i_u^{n+1})$ 。④根据高斯分布,由其第 $n+1$ 次行为所处的心理状态生成第 $n+1$ 和第 n 次交互的两个职位之间的混合相似度,即 $s_u^{n+1} \sim P(s_u^{n+1} | z_u^{n+1})$ 。

5.3.4 参数估计

给定用于训练的所有求职者决策过程中的历史行为记录,动态贝叶斯网络模型的对数似然函数如式(5.6)所示。其中, θ 表示模型中的所有参数, j_u 和 s_u 则表示求职者 u 所有可观测的交互职位和职位间相似度。

$$l(\theta) = \sum_u \ln P(j_u, s_u; \theta) = \sum_u \ln \sum_{z_u} \sum_{i_u} P(j_u, s_u, z_u, i_u; \theta) \quad (5.6)$$

本研究采用EM算法(Dempster et al., 1977)迭代优化模型参数。期望步根据可观测记录和当前的模型参数估计来推断隐变量,记为 $P(z_u, i_u | j_u, s_u; \theta)$;极大步最大化对数似然度期望的下限,如式(5.7)所示。

$$Q(\theta, \theta^{\text{old}}) = \sum_u \sum_{z_u} \sum_{i_u} P(z_u, i_u | j_u, s_u; \theta^{\text{old}}) \ln P(j_u, s_u, z_u, i_u; \theta) \leq l(\theta) \quad (5.7)$$

当两个步骤交替进行至对数似然度收敛,即可获得模型参数的一组较优估计。根据动态贝叶斯网络中变量间的条件依赖和独立关系,联合概率 $P(j_u, s_u, z_u, i_u; \theta)$ 可分解为式(5.8)的形式,进而将对数似然期望拆分为6个独立部分,对应6类参数,如式(5.9)所示。

$$P(j_u, s_u, z_u, i_u | \theta) = P(z_u^1) P(i_u^1 | z_u^1) \prod_{n=1}^{N_u-1} P(z_u^{n+1} | z_u^n) \prod_{n=1}^{N_u-1} P(i_u^{n+1} | i_u^n, z_u^{n+1}) \times \prod_{n=1}^{N_u} P(j_u^n | i_u^n) \prod_{n=1}^{N_u-1} P(s_u^{n+1} | z_u^{n+1}) \quad (5.8)$$

$$\begin{aligned} Q(\theta, \theta^{\text{old}}) &= \sum_u \sum_{z_u^1} P(z_u^1 | j_u, s_u; \theta^{\text{old}}) \ln P(z_u^1) + \sum_u \sum_{z_u^1} \sum_{i_u^1} P(z_u^1, i_u^1 | j_u, s_u; \theta^{\text{old}}) \ln P(i_u^1 | z_u^1) + \\ &\sum_u \sum_{n=1}^{N_u-1} \sum_{z_u^n} \sum_{z_u^{n+1}} P(z_u^n, z_u^{n+1} | j_u, s_u; \theta^{\text{old}}) \ln P(z_u^{n+1} | z_u^n) + \\ &\sum_u \sum_{n=1}^{N_u-1} \sum_{i_u^n} \sum_{z_u^{n+1}} \sum_{i_u^{n+1}} P(i_u^n, z_u^{n+1}, i_u^{n+1} | j_u, s_u; \theta^{\text{old}}) \ln P(i_u^{n+1} | i_u^n, z_u^{n+1}) + \\ &\sum_u \sum_{n=1}^{N_u} \sum_{i_u^n} P(i_u^n | j_u, s_u; \theta^{\text{old}}) \ln P(j_u^n | i_u^n) + \\ &\sum_u \sum_{n=1}^{N_u-1} \sum_{z_u^{n+1}} P(z_u^{n+1} | j_u, s_u; \theta^{\text{old}}) \ln P(s_u^{n+1} | z_u^{n+1}) \end{aligned} \quad (5.9)$$

定义联合概率 $\alpha(z_u^n, i_u^n)$ 和条件概率 $\beta(z_u^n, i_u^n)$ ，如式 (5.10) 和式 (5.11) 所示，分别从前向后和从后向前递归计算。

$$\begin{aligned} \alpha(z_u^n, i_u^n) &= P(j_u^1, \dots, j_u^n, s_u^2, \dots, s_u^n, z_u^n, i_u^n; \theta) \\ &= \sum_{z_u^{n-1}} \sum_{i_u^{n-1}} \alpha(z_u^{n-1}, i_u^{n-1}) P(z_u^n | z_u^{n-1}) P(i_u^n | i_u^{n-1}, z_u^n) P(j_u^n | i_u^n) P(s_u^n | z_u^n) \end{aligned} \quad (5.10)$$

其中，初始的 $\alpha(z_u^1, i_u^1) = P(j_u^1, z_u^1, i_u^1; \theta) = P(z_u^1) P(i_u^1 | z_u^1) P(j_u^1 | i_u^1)$ 。

$$\begin{aligned} \beta(z_u^n, i_u^n) &= P(j_u^{n+1}, \dots, j_u^{N_u}, s_u^{n+1}, \dots, s_u^{N_u} | z_u^n, i_u^n; \theta) \\ &= \sum_{z_u^{n+1}} \sum_{i_u^{n+1}} \beta(z_u^{n+1}, i_u^{n+1}) P(z_u^{n+1} | z_u^n) P(i_u^{n+1} | i_u^n, z_u^{n+1}) P(j_u^{n+1} | i_u^{n+1}) P(s_u^{n+1} | z_u^{n+1}) \end{aligned} \quad (5.11)$$

其中，初始的 $\beta(z_u^{N_u}, i_u^{N_u}) = 1$ 。

而后定义辅助变量 $\gamma(z_u^n, i_u^n)$ 和 $\xi(z_u^n, i_u^n, z_u^{n+1}, i_u^{n+1})$ ，如式(5.12)和式(5.13)所示，用于计算隐变量的条件概率。

$$\begin{aligned} \gamma(z_u^n, i_u^n) &= P(z_u^n, i_u^n | j_u, s_u; \theta) = \frac{P(j_u, s_u | z_u^n, i_u^n) P(z_u^n, i_u^n)}{P(j_u, s_u)} \\ &= \frac{\alpha(z_u^n, i_u^n) \beta(z_u^n, i_u^n)}{\sum_{z_u^n} \sum_{i_u^n} \alpha(z_u^n, i_u^n) \beta(z_u^n, i_u^n)} \end{aligned} \quad (5.12)$$

这里 $P(j_u, s_u | z_u^n, i_u^n) = P(j_u^{1:n}, s_u^{1:n} | z_u^n, i_u^n) P(j_u^{n+1:N_u}, s_u^{n+1:N_u} | z_u^n, i_u^n)$ ，因存在条件

独立关系 $(j_u^{1:n}) \perp (j_u^{n+1:N_u}) | z_u^n, i_u^n$ 和 $(s_u^{2:n}) \perp (s_u^{n+1:N_u}) | z_u^n, i_u^n$ 。

$$\begin{aligned}
 & \xi(z_u^n, i_u^n, z_u^{n+1}, i_u^{n+1}) \\
 &= P(z_u^n, i_u^n, z_u^{n+1}, i_u^{n+1} | j_u, s_u; \theta) \\
 &= \frac{\alpha(z_u^n, i_u^n) \beta(z_u^{n+1}, i_u^{n+1}) P(z_u^{n+1} | z_u^n) P(i_u^{n+1} | i_u^n, z_u^{n+1}) P(j_u^{n+1} | i_u^{n+1}) P(s_u^{n+1} | z_u^{n+1})}{\sum_{z_u^n} \sum_{i_u^n} \sum_{z_u^{n+1}} \sum_{i_u^{n+1}} \alpha(z_u^n, i_u^n) \beta(z_u^{n+1}, i_u^{n+1}) P(z_u^{n+1} | z_u^n) P(i_u^{n+1} | i_u^n, z_u^{n+1}) P(j_u^{n+1} | i_u^{n+1}) P(s_u^{n+1} | z_u^{n+1})} \quad (5.13)
 \end{aligned}$$

令式(5.12)、式(5.13)与式(5.9)合并, 则可将对数似然度期望改写为6个独立的模型部分, 即将式(5.9)进一步转换为式(5.14)。

$$\begin{aligned}
 Q(\theta, \theta^{\text{old}}) &= \sum_u \sum_k \sum_{i_u^1} \gamma(z_u^1 = k, i_u^1; \theta^{\text{old}}) \ln \pi_k + \\
 & \quad \sum_u \sum_k \sum_b \gamma(z_u^1 = k, i_u^1 = b; \theta^{\text{old}}) \ln F_{kb} + \\
 & \quad \sum_u \sum_{n=1}^{N_u-1} \sum_k \sum_m \sum_{i_u^n} \sum_{i_u^{n+1}} \xi(z_u^n = k, i_u^n, z_u^{n+1} = m, i_u^{n+1}; \theta^{\text{old}}) \ln A_{km} + \\
 & \quad \sum_u \sum_{n=1}^{N_u-1} \sum_b \sum_m \sum_c \sum_{z_u^n} \xi(z_u^n, i_u^n = b, z_u^{n+1} = m, i_u^{n+1} = c; \theta^{\text{old}}) \ln H_{bmc} + \quad (5.14) \\
 & \quad \sum_u \sum_{n=1}^{N_u} \sum_b \sum_{z_u^n} \gamma(z_u^n, i_u^n = b; \theta^{\text{old}}) \ln \prod_{j=1}^{|J|} E_{bj}^{1_j(j_u^n)} + \\
 & \quad \sum_u \sum_{n=1}^{N_u-1} \sum_m \sum_{i_u^n} \gamma(z_u^{n+1} = m, i_u^n; \theta^{\text{old}}) \ln G_{ms_u^{n+1}}
 \end{aligned}$$

其中, π_k 、 F_{kb} 、 A_{km} 、 H_{bmc} 、 E_{bj} 和 $G_{ms_u^{n+1}}$ 分别表示 $P(z_u^1 = k)$ 、 $P(i_u^1 = b | z_u^1 = k)$ 、 $P(z_u^{n+1} = m | z_u^n = k)$ 、 $P(i_u^{n+1} = c | i_u^n = b, z_u^{n+1} = m)$ 、 $P(j_u^n = j | i_u^n = b)$ 和 $P(s_u^{n+1} | z_u^{n+1} = m)$ 。此外, $1_j(j_u^n)$ 为指示函数, 当 $j_u^n = j$ 时取值为1, 否则为0。

总结而言, 在期望步中, 算法将参数用于计算 $\gamma(z_u^n, i_u^n)$ 和 $\xi(z_u^n, i_u^n, z_u^{n+1}, i_u^{n+1})$, 其中 $u=1, 2, \dots, |U|$; $n=1, 2, \dots, N_u$ 。在极大步中, 算法则相互独立地最大化6个模型部分的参数。

为降低过拟合风险, 本方法另用MAP估计 (Murphy, 2012) 对模型中涉及的所有多项分布施加狄利克雷先验。令 $|Z|$ 、 $|I|$ 和 $|J|$ 分别表示心理状态集合 Z 、兴趣分布集合 I 和职位集合 J 的大小, μ_m 和 σ_m 另分别表示求职者在心理状态 m 下相邻两次行为所交互职位间的相似度均值和标准差, 则有概率 $\pi, A_{k,:} \sim \text{Dirichlet}(x; \eta_1, \eta_2, \dots, \eta_{|Z|})$ 、 $F_{k,:}, H_{b,m,:} \sim \text{Dirichlet}(x; \tau_1, \tau_2, \dots, \tau_{|I|})$ 、 $E_{b,:} \sim \text{Dirichlet}(x; \phi_1, \phi_2, \dots, \phi_{|J|})$ 、 $G_{m,:} \sim N(\mu_m, \sigma_m^2)$, 其中 $\eta_z = \eta / |Z|$ 、 $\tau_i = \tau / |I|$ 、 $\phi_j = \phi / |J|$ 。由此, 模型所有参数的最大后验概率如式(5.15)~式(5.22)所示。

$$\pi_k = P(z_u^1 = k) = \frac{\sum_u \sum_{i_u^1} \gamma(z_u^1 = k, i_u^1) + \eta_k - 1}{\sum_u \sum_{i_u^1} \sum_{i_u^1} \gamma(z_u^1 = k, i_u^1) + \eta - |I|} \quad (5.15)$$

$$F_{kb} = P(i_u^1 = b | z_u^1 = k) = \frac{\sum_u \gamma(z_u^1 = k, i_u^1 = b) + \tau_b - 1}{\sum_u \sum_{i_u^1} \gamma(z_u^1 = k, i_u^1) + \tau - |I|} \quad (5.16)$$

$$\begin{aligned} A_{km} &= P(z_u^{n+1} = m | z_u^n = k) \\ &= \frac{\sum_u \sum_{n=1}^{N_u-1} \sum_{i_u^n} \sum_{i_u^{n+1}} \xi(z_u^n = k, i_u^n, z_u^{n+1} = m, i_u^{n+1}) + \eta_m - 1}{\sum_u \sum_{n=1}^{N_u-1} \sum_{i_u^n} \sum_{i_u^{n+1}} \sum_{i_u^{n+1}} \xi(z_u^n = k, i_u^n, z_u^{n+1} = m, i_u^{n+1}) + \eta - |S|} \end{aligned} \quad (5.17)$$

$$\begin{aligned} H_{bmc} &= P(i_u^{n+1} = c | i_u^n = b, z_u^{n+1} = m) \\ &= \frac{\sum_u \sum_{n=1}^{N_u-1} \sum_{z_u^n} \xi(z_u^n = b, z_u^{n+1} = m, i_u^{n+1} = c) + \tau_c - 1}{\sum_u \sum_{n=1}^{N_u-1} \sum_{z_u^n} \sum_{i_u^{n+1}} \xi(z_u^n = b, z_u^{n+1} = m, i_u^{n+1}) + \tau - |I|} \end{aligned} \quad (5.18)$$

$$E_{bj} = P(j_u^n = j | i_u^n = b) = \frac{\sum_u \sum_{n=1}^{N_u} 1_j(j_u^n) \sum_{z_u^n} \gamma(z_u^n, i_u^n = b) + \phi_j - 1}{\sum_u \sum_{n=1}^{N_u} \sum_{z_u^n} \gamma(z_u^n, i_u^n = b) + \phi - |I|} \quad (5.19)$$

$$G_{ms_u^{n+1}} = P(s_u^{n+1} | z_u^{n+1} = m) = \frac{1}{\sqrt{2\pi\sigma_m^2}} \exp \left[-\frac{(s_u^{n+1} - \mu_m)^2}{2\sigma_m^2} \right] \quad (5.20)$$

$$\mu_m = \frac{\sum_u \sum_{n=1}^{N_u-1} \sum_{i_u^{n+1}} \gamma(z_u^{n+1} = m, i_u^{n+1}) s_u^{n+1}}{\sum_u \sum_{n=1}^{N_u-1} \sum_{i_u^{n+1}} \gamma(z_u^{n+1} = m, i_u^{n+1})} \quad (5.21)$$

$$\sigma_m = \sqrt{\frac{\sum_u \sum_{n=1}^{N_u-1} \sum_{i_u^{n+1}} \gamma(z_u^{n+1} = m, i_u^{n+1}) (s_u^{n+1} - \mu_m)^2}{\sum_u \sum_{n=1}^{N_u-1} \sum_{i_u^{n+1}} \gamma(z_u^{n+1} = m, i_u^{n+1})}} \quad (5.22)$$

在模型训练过程中, 需要使用DBN-HS方法中的隐层识别策略来设计相关约束, 以保证参数估计与理论发现一致。

(1) 状态约束。通过概率交换来约束不同心理状态下求职者相邻两次交互职位间平均相似度的相对大小关系。具体地, 预先标定 $z_u^n=0$ 为准备状态, $z_u^n=1$ 为活跃状态, 其中 $n=1, \dots, N_u$, 并估计 μ_0 和 μ_1 。若在迭代过程中出现 $\mu_0 > \mu_1$, 则交换所有与 z_u^n 相关的概率参数 (即 π 、 F 、 A 、 H 、 μ 和 σ) 中用于表示状态的角标次序, 如将 μ_0 和 μ_1 互换, 以此类推。由此, 保证 $z_u^n=0$ 始终为准备状态, $z_u^n=1$ 始终为活跃状态, 且活跃状态下求职者相邻两次交互职位间的平均混合相似度始终高于准备状态。

(2) 兴趣收敛性约束。对于同一求职者 u , 根据其相邻两次职位交互行为的所处状态 z_u^n 和 z_u^{n+1} ($n=1, \dots, N_u-1$), 通过峰度计算来约束其兴趣变化模式。在DBN-HS方法的动态贝叶斯网络中, 兴趣节点作为职位集合上的概率分布, 可由峰度来计算其收敛性, 峰度越大则求职者在各职位上的兴趣分布越集中或收敛, 峰度越小则兴趣越分散或均匀。将兴趣 $i_u^n=b$ 的峰度记为 $K(b)$, 则有式(5.23)。

$$K(b) = \frac{\sum_{j=1}^{|J|} E_{bj} (R_{bj} - \bar{R}_{bj})^4}{(\sum_{j=1}^{|J|} E_{bj} (R_{bj} - \bar{R}_{bj})^2)^2} \quad (5.23)$$

其中, E_{bj} 为兴趣 b 下求职者与职位 j 交互的概率。 $R_{bj} \in [1, |J|]$ 为职位 j 在兴趣 b 峰度计算中的次序, 则排序时令 E_{bj} 最大的职位位于最中间次序, 即 $R_{bj} = |J|/2$ 或 $(|J|+1)/2$, 而令 E_{bj} 最小的职位的次序位于两端, 即 $R_{bj} = 1$ 或 $|J|$ 。 $\bar{R}_{bj} = \sum_{j=1}^{|J|} E_{bj} R_{bj}$ 为兴趣 b 下求职者所有交互职位的平均次序。类似附录A可以证明, 若简化假设求职者对其感兴趣的职位并无兴趣程度差异, 则其感兴趣的职位数量越多 (如准备状态), 相应兴趣分布的峰度越小, 说明峰度这一量化指标及其约束能够作为求职者序贯心理状态所描述理论规律的合理推衍。由此, 对于同一求职者 u , 若 $z_u^{n+1} > z_u^n$, 则约束概率 $P(i_u^{n+1} | i_u^n, z_u^{n+1}) > 0$ 必有 $K(i_u^{n+1}) \geq K(i_u^n)$; 若 $z_u^{n+1} < z_u^n$, 则令 $P(i_u^{n+1} | i_u^n, z_u^{n+1}) > 0$ 必有 $K(i_u^{n+1}) \leq K(i_u^n)$ 。训练中, 可采用指示函数 $1((z_u^{n+1} - z_u^n)(K(i_u^{n+1}) - K(i_u^n)) \geq 0)$ 来限制 α 、 β 和 ξ 的计算, 从而保证同一求职者由准备状态转换至活跃状态时兴趣更集中, 由活跃状态转换至准备状态时兴趣更分散。

针对上述DBN-HS的参数估计过程, 可进行时间复杂度分析。可以发现, 在采用的EM估计算法的各次迭代中, 时间消耗主要在于计算辅助变量 $\alpha(z_u^n, i_u^n)$ 、 $\beta(z_u^n, i_u^n)$ 、 $\gamma(z_u^n, i_u^n)$ 和 $\xi(z_u^n, i_u^n, z_u^{n+1}, z_u^{n+1})$ 。假设训练数据集中, 求职者、心理状态和兴趣分布的总数分别为 $|U|$ 、 $|Z|$ 和 $|I|$, 且求职者决策过程中平均包含的行为数量为 $\bar{N}_u$ 。则根据式(5.10)~式(5.12), 计算 $\alpha(z_u^n, i_u^n)$ 、 $\beta(z_u^n, i_u^n)$ 和 $\gamma(z_u^n, i_u^n)$ 的复杂度均为 $O(\bar{N}_u |U| |Z| |I|)$ 。而据式(5.13)可得, 计算 $\xi(z_u^n, i_u^n, z_u^{n+1}, z_u^{n+1})$ 的时间复杂度为 $O(\bar{N}_u |U| |Z|^2 |I|^2)$ 。因此, DBN-HS方法参数估计中各次迭代的复杂度为 $O(\bar{N}_u |U| |Z| |I| (3 + |Z| |I|))$, 或可等价记为 $O(\bar{N}_u |U| |Z|^2 |I|^2)$ 。鉴于 $|Z|$ 和 $|I|$ 均为预先设置的有限大小的超参数, 则DBN-HS的训练时间消耗随数据集中的求职者总数或其求职过程中产生的平均行为数呈线性增长。因此, 在现实的大规模职位推荐应用场景中, DBN-HS方法也应能展示良好的可拓展性。

5.3.5 推断和职位推荐

学习完动态贝叶斯网络模型的所有参数后, 对于测试集中的新求职者 u , 本方法首先采用维特比式推断策略确定其第 $1 \sim N_u - 1$ 次职位交互行为对应的最可能心理状态和兴趣。然后基于概率预测策略预测其第 N_u 次行为最可能感兴趣的职位, 并生成职位推荐列表。

(1) 维特比式推断。推断的目标为求职者 u 在其决策过程中第 $1 \sim N_u - 1$ 次职位交互行为所对应的概率最大的心理状态和潜在兴趣, 如式 (5.24)

所示。

$$\widehat{z}_u^{1:N_u-1}, \widehat{i}_u^{1:N_u-1} = \arg \max_{z_u^{1:N_u-1}, i_u^{1:N_u-1}} P(j_u^{1:N_u-1}, s_u^{1:N_u-1}, z_u^{1:N_u-1}, i_u^{1:N_u-1} | \theta) \quad (5.24)$$

针对求职者 u 的第 n 次行为, 令 δ_{ukb}^n 表示其处于心理状态 k 、对应兴趣分布 b , 并呈现可观测的交互职位 j_u^n 及其与第 $n-1$ 次交互职位间相似度 s_u^n 的最大概率。

对于 $n=1$, 有

$$\delta_{ukb}^1 = P(z_u^1 = k, i_u^1 = b, j_u^1 | \theta) = P(z_u^1 = k)P(i_u^1 = b | z_u^1 = k)P(j_u^1 | i_u^1 = b) \quad (5.25)$$

对于 $n=2$, 有

$$\delta_{ukb}^2 = P(j_u^2 | i_u^2 = b)P(s_u^2 | z_u^2 = k) \max_{z_u^1, i_u^1} \delta_{uz_u^1 i_u^1}^1 P(z_u^2 = k | z_u^1)P(i_u^2 = b | i_u^1, z_u^2 = k) \quad (5.26)$$

以此类推, 通过迭代方式可以得到,

对于 $n=N_u-1$, 有

$$\delta_{ukb}^{N_u-1} = P(j_u^{N_u-1} | i_u^{N_u-1} = b)P(s_u^{N_u-1} | z_u^{N_u-1} = k) \times \max_{z_u^{N_u-2}, i_u^{N_u-2}} \delta_{uz_u^{N_u-2} i_u^{N_u-2}}^{N_u-2} P(z_u^{N_u-1} = k | z_u^{N_u-2})P(i_u^{N_u-1} = b | i_u^{N_u-2}, z_u^{N_u-1} = k) \quad (5.27)$$

因此, 针对求职者 u 第 N_u-1 次职位交互行为, 可以推断如下:

$$\widehat{z}_u^{N_u-1}, \widehat{i}_u^{N_u-1} = \arg \max_{k,b} \delta_{ukb}^{N_u-1} \quad (5.28)$$

进而, 再回溯推断求职者 u 在其决策过程中经历的所有序贯心理状态 z_u^n 和兴趣分布 $i_u^n (n=1, \dots, N_u-2)$, 即

$$\widehat{z}_u^n, \widehat{i}_u^n = \arg \max_{z_u^n, i_u^n} \delta_{uz_u^n i_u^n}^n P(\widehat{z}_u^{n+1} | z_u^n)P(\widehat{i}_u^{n+1} | i_u^n, \widehat{z}_u^{n+1}) \quad (5.29)$$

(2) 概率预测。为提供准确且鲁棒的职位推荐, DBN-HS方法提出一种概率预测策略, 即除维特比式解码得到的最大概率状态和兴趣外, 也考虑其他所有 $z_u^{N_u-1}$ 和 $i_u^{N_u-1}$ 的可能性, 同时利用式(5.27)实现剪枝。具体地, 对于 $z_u^{N_u-1}=k$ 和 $i_u^{N_u-1}=b$, 采用式(5.27)计算求职者 u 第 N_u-1 次行为处于心理状态 k 且对应兴趣 b 的最大概率, 剪去其他以较低概率到达状态 k 和兴趣 b 的路径, 但考虑求职者第 N_u-1 次行为所有可能的状态 k 和兴趣 b , 由此有效应对数据中的不确定性和噪声。进而, 即可由式(5.30)直接计算其第 N_u 次行为与职位 j 交互的概率。

$$\begin{aligned} & P(j_u^{N_u} = j) \\ &= \sum_{z_u^{N_u-1}} \sum_{i_u^{N_u-1}} \sum_{z_u^{N_u}} \sum_{i_u^{N_u}} P(z_u^{N_u-1}, i_u^{N_u-1}, z_u^{N_u}, i_u^{N_u})P(j_u^{N_u} = j | z_u^{N_u-1}, i_u^{N_u-1}, z_u^{N_u}, i_u^{N_u}) \\ &= \sum_{z_u^{N_u-1}} \sum_{i_u^{N_u-1}} \sum_{z_u^{N_u}} \sum_{i_u^{N_u}} \frac{\delta_{uz_u^{N_u-1} i_u^{N_u-1}}^{N_u-1}}{\sum_{z_u^{N_u-1}} \sum_{i_u^{N_u-1}} \delta_{uz_u^{N_u-1} i_u^{N_u-1}}^{N_u-1}} \times \end{aligned} \quad (5.30)$$

$$P(z_u^{N_u} | z_u^{N_u-1}) P(i_u^{N_u} | i_u^{N_u-1}, z_u^{N_u}) P(j_u^{N_u} = j | i_u^{N_u})$$

其中, 由于可以假设求职者 u 在决策过程中, 经过从第 $1 \sim N_u-2$ 次职位交互行为对应的概率最大的心理状态和兴趣序列到达第 N_u-1 次行为的状态 k 和兴趣 b , 故用

$\frac{\delta_{uz_u^{N_u-1}, i_u^{N_u-1}}^{N_u-1}}{\sum_{z_u^{N_u-1}} \sum_{i_u^{N_u-1}} \delta_{uz_u^{N_u-1}, i_u^{N_u-1}}^{N_u-1}}$ 来近似 $P(z_u^{N_u-1}, i_u^{N_u-1})$ 。而因本方法在

动态贝叶斯网络的构建中融入了职位相似度, 且能够根据求职者历史共同交互、职位属性以及随机森林模型训练后的参数等信息来计算获得数据集中任意两个职位间的混合相似度, 则如式(5.31)所示, DBN-HS方法在概率预测策略中另纳入各候选职位与求职者 u 第 N_u-1 次行为交互的职位之间的混合相似度, 从而进一步增强对于其第 N_u 次行为预测的准确性以及职位推荐的鲁棒性。

$$\begin{aligned} & P(j_u^{N_u} = j, s_u^{N_u} = s) \\ &= \sum_{z_u^{N_u-1}} \sum_{i_u^{N_u-1}} \sum_{z_u^{N_u}} \sum_{i_u^{N_u}} P(z_u^{N_u-1}, i_u^{N_u-1}, z_u^{N_u}, i_u^{N_u}) \times P(j_u^{N_u} = j, s_u^{N_u} = s | z_u^{N_u-1}, i_u^{N_u-1}, z_u^{N_u}, i_u^{N_u}) \\ &= \sum_{z_u^{N_u-1}} \sum_{i_u^{N_u-1}} \sum_{z_u^{N_u}} \sum_{i_u^{N_u}} \frac{\delta_{uz_u^{N_u-1}, i_u^{N_u-1}}^{N_u-1}}{\sum_{z_u^{N_u-1}} \sum_{i_u^{N_u-1}} \delta_{uz_u^{N_u-1}, i_u^{N_u-1}}^{N_u-1}} \times \\ & P(z_u^{N_u} | z_u^{N_u-1}) P(i_u^{N_u} | i_u^{N_u-1}, z_u^{N_u}) P(j_u^{N_u} = j | i_u^{N_u}) P(s_u^{N_u} = s | z_u^{N_u}) \end{aligned} \quad (5.31)$$

其中, s 表示职位 j 与求职者 u 第 N_u-1 次行为所交互职位之间的混合相似度, 即 $s = \text{hsim}(j_u^{N_u-1}, j)$ 。最终, 根据式(5.31)的计算结果, 按照概率降序将其求职过程中第 N_u 次行为最有可能交互的职位推荐给求职者。

在推断和预测过程中, DBN-HS方法的时间消耗主要在于计算 δ_{ukb}^n 。给定序贯心理状态数和兴趣分布数分别为 $|\mathbf{Z}|$ 和 $|\mathbf{I}|$, 则对于共产生 N_u 次行为的求职者 u , 由式(5.25)~式(5.27)可知, 计算 δ_{ukb}^n 的复杂度为 $O(N_u |\mathbf{Z}| |\mathbf{I}|)$ 。鉴于 $|\mathbf{Z}|$ 和 $|\mathbf{I}|$ 均为有限大小的超参数, 即 δ_{ukb}^n 的计算时长仅随求职者行为数呈线性增长。若设置 N_u 的最大取值并做序列截断, 则推断并预测各求职者所经历心理状态和兴趣序列的时间消耗仅在常数级, 较为可控。因此, DBN-HS方法在基于历史训练数据集学习到随机森林以及动态贝叶斯网络的模型参数后, 在现实的职位推荐应用中能够面向每位新到达的求职者实现效率足够高的即时动态推荐。

5.4 真实数据实验

实验从国内最大的网络招聘平台之一——猎聘网采集真实数据集, 以评估所提出的DBN-HS方法及各类基线方法的职位推荐表现。该数据集记录

了2018年4月至6月期间大量的求职者职位浏览行为，每条记录包含求职者ID、职位ID、浏览时间等交互信息，以及职位的12个结构化属性：企业性质、企业规模、职能分类、所在地区、所属行业、年薪下限、年薪上限、工作年限、性别要求、最低年龄要求、最高年龄要求和学历要求。参考已有推荐方法的相关研究，实验对原始数据集进行预处理，仅保留足够频繁的职位和足够活跃的求职者的相关记录（Paparrizos et al., 2011; Sahoo et al., 2012）。具体地，过滤掉历史交互总数低于5次的职位和历史行为总数低于5次的求职者后，得到最终用于真实数据实验的数据集。该数据集共包含24 683条职位浏览记录，覆盖1479位求职者和2750个职位，即每个职位在数据集中平均出现8.98次，每位求职者平均产生16.69次职位浏览行为。

本实验采用十折交叉验证对各推荐方法进行训练和测试，即随机将数据集中的求职者分成10等份，每次取其中一份作为测试集，剩余9份作为训练集，计算10次测试的指标平均值以评估方法的推荐表现。训练阶段学习相应求职者的全部行为记录，测试阶段则根据每位求职者已产生的所有历史行为，预测其最后一次行为将会交互的职位。参考求职者序贯心理状态的相关理论研究（Albert, 1966; Blau, 1993），实验将DBN-HS方法中动态贝叶斯网络的隐状态数设置为2。根据预先在相同数据集上进行的超参数调节实验结果（详见5.4.4节），将模型中的兴趣数设置为80，以平衡职位推荐的效率和效果，且将职位相似度的校正比例调节为0.8以获得最优表现。给定各方法为每个测试集序列生成的前 K 个得分最高的职位列表，实验采用命中率（HR@ K ）和归一化折损累计增益（NDCG@ K ）两个指标，分别衡量推荐方法的精度和排序效果，如式(5.32)和式(5.33)所示。

$$\text{HR@}K = \frac{1}{|U|} \sum_{u \in U} 1(r_{u,j_u^{N_u}} \leq K) \quad (5.32)$$

$$\text{NDCG@}K = \frac{1}{|U|} \sum_{u \in U} \frac{1(r_{u,j_u^{N_u}} \leq K)}{\log_2(r_{u,j_u^{N_u}} + 1)} \quad (5.33)$$

其中， $r_{u,j_u^{N_u}}$ 代表在方法为求职者 u 所提供的推荐列表中，其第 N_u 次行为实际交互的职位 $j_u^{N_u}$ 对应的排列次序。 $1(r_{u,j_u^{N_u}} \leq K)$ 为指示函数，当 $r_{u,j_u^{N_u}} \leq K$ 时取值为1，否则取值为0。

5.4.1 职位混合相似度

在学习求职者决策过程中潜在的心理状态和兴趣分布之前，DBN-HS方法先借助随机森林模型对职位相似度进行校正，获得输入动态贝叶斯网络的职位间混合相似度。随机森林模型的训练标签为训练集中基于求职者历

史共同交互记录的职位间相似度 $isim$ ，由Jaccard系数计算。在实验采用的猎聘网真实数据集中，所有职位间的两两共同交互相似度平均为0.0133，其中相似度为0的职位对占比高达97.65%，相似度在0.5以内的职位对占比为99.41%，而相似度为1的职位对占0.50%。这表明，现实中求职者与职位间的交互确实非常稀疏，仅基于历史共同交互的职位相似度衡量难以完全反映求职者实际的兴趣相似性。尤其对于历史交互极少甚至尚无交互的新发布职位，即使其符合求职者潜在的职位兴趣，也很难据此得到推荐。

作为随机森林模型的输入特征，实验计算训练集中所有职位在各属性上的两两相似度 $fsim_r(r=1, \dots, |R|)$ 。其中， $|R|$ 为职位的属性个数，在实验数据集中即为12。具体地，实验采用式(5.1)计算五个连续属性（年薪下限、年薪上限、工作年限要求、最低年龄要求 and 最高年龄要求）以及两个有序分类属性（企业规模和学历要求）上的职位相似度；采用式(5.2)计算两个无序分类属性（企业性质和性别要求）上的职位相似度；采用式(5.3)计算三个层级分类属性（职能分类、所在地区和所属行业）上的职位相似度。进而，调节不同的决策树数量 M 来训练随机森林模型，得到各 M 取值下模型预测的职位间共同交互相似度的均方根误差（root mean squared error, RMSE），如图5.4所示。RMSE值越小，表示预测越准确。

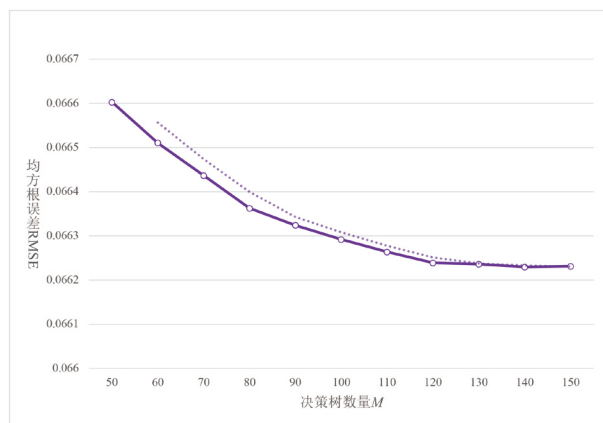


图5.4 随机森林模型在不同决策树数量下的预测效果

由图5.4可知，随机森林模型对职位共同交互相似度的预测效果整体上随决策树数量的增大而提升，但提升幅度逐渐缩小。当决策树数量达到约130后，模型预测值的RMSE基本趋于稳定，为0.0662。鉴于在给定规模的训练数据下，随机森林模型的训练时长近乎随决策树数量线性增长，本实验选用超参数 $M=130$ 下的职位相似度预测值 $psim$ ，以兼顾模型的计算效率和拟合效果。

此外，随机森林模型可以通过平均精度下降和平均误差上升的量化计算来实现事后可解释性，即给出训练中学到各职位属性相似度对预测整体的共同交互相似度的重要性，从而说明在求职者的历史交互行为记录中所体现的各职位属性对其决策的影响程度。图5.5展示了真实数据实验训练后的随机森林模型所学习到的重要性结果，其中将各职位属性按照重要性从大到小的次序从左到右排列，所有属性的重要性总和为1。

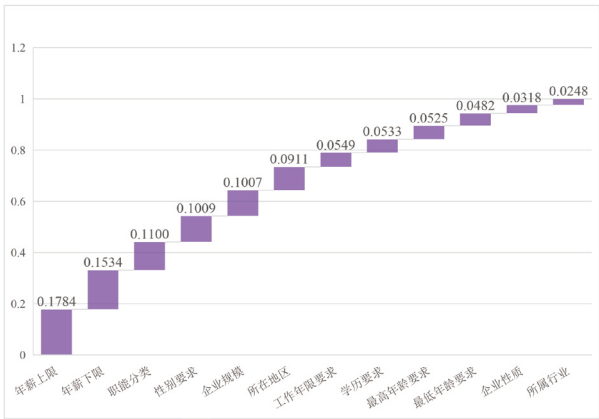


图5.5 随机森林模型习得的职位各属性相似度的重要性

如图5.5所示，在所有职位属性中，对于其共同交互相似度预测最重要的是年薪上限和年薪下限，反映出薪资水平是求职者决策过程中最关键的影响因素，这与已有的相关研究和经验相一致。此外，重要性较高的属性还包括职能分类、性别要求、企业规模和所在地区。性别为职位对求职者的硬性要求，而职能与求职者长期积累的知识和工作经验高度相关，且地区为较难切换的求职考虑因素，因此求职者倾向于在这些属性上选择较为相似的职位。而企业性质和所属行业的重要性较低，说明求职者对这些属性的依赖性整体不高，可能对较多样的职位均有一定兴趣。例如，会计师、数据分析师、法律顾问等职位可能较容易适应于不同行业，使得所属行业不会很大程度上影响求职者决策。由此，随机森林模型能在保证较优训练效果的同时，提供关于各属性相似度重要性的有效量化说明，从而使得基于其的职位相似度校正以及DBN-HS方法具有良好的管理可理解性。

5.4.2 基线方法对比

为说明DBN-HS方法在职位推荐任务中的效果优势，实验基于猎聘网的真实数据集，将其与现有的基线推荐方法进行对比。具体而言，实验采用的基线方法可分为一般性静态推荐方法、一般性动态推荐方法和针对性职

位推荐方法三类，共十种基线方法。各方法的简要描述和面向研究场景的实现方案如下。

(1) 推荐热门产品 (POP)。作为一种常见的非个性化推荐方法，POP将全部历史记录中被交互最多的热门职位无差异地推荐给各求职者 (Prawesh and Padmanabhan, 2014)。

(2) 基于用户相似度的协同过滤 (User-CF)。在本实验的数据场景下，User-CF方法将求职者的职位交互行为视为隐式反馈，以历史交互次数作为评分计算求职者间的相似度，据此为各求职者推荐与其相似的求职者交互过的职位 (Breese et al., 1998)。

(3) 基于内容的推荐 (CBR)。CBR根据实验数据集中的12个结构化职位属性，基于求职者的职位交互历史计算其在各属性上的平均兴趣，并通过计算候选职位的属性向量与求职者兴趣向量间的余弦相似度推荐最可能感兴趣的职位 (Basu et al., 1998)。

(4) 极端梯度提升 (XGBoost)。在拓展为动态方法的实验场景下，XGBoost基于求职者历史交互过的职位ID及其所有属性信息进行树学习，预测求职者下次可能浏览或产生交互的职位 (Chen and Guestrin, 2016)。

(5) 深度兴趣网络 (DIN)。由于数据集中缺少求职者的属性信息，DIN仅生成职位的属性向量，结合求职者ID和职位ID的独热编码，根据历史交互行为来学习各次行为的权重，据此推荐后续可能感兴趣的职位 (Zhou et al., 2019a)。

(6) 隐马尔可夫模型 (HMM)。作为一种经典的动态贝叶斯网络，HMM仅用一类隐层 (隐状态) 描述求职者决策过程中的职位兴趣变化 (Sahoo et al., 2012)。参考5.3.5节，实验另设计针对HMM模型结构的维特比式解码和概率预测策略。

(7) 融入职位混合相似度的隐马尔可夫模型 (HMM-HS)。由于DBN-HS方法在建模时引入了随机森林模型校正后的职位间混合相似度，其中同时蕴含共同交互视角与属性视角的相似度信息，基线对比实验在HMM基础上构造HMM-HS混合推荐方法。具体地，HMM-HS描述求职过程中的隐性兴趣变化，由当前兴趣生成其交互的职位以及该职位与上次交互职位间的相似度，并实现定制化推断和预测策略。

(8) 综合网络招聘系统 (iHR)。针对职位推荐任务，iHR系统提出了一种经典的混合推荐方法，即分别实现基于内容的推荐、协同过滤推荐和互惠推荐，再通过对三种推荐排序的加权求和获得最终的职位推荐列表 (Hong et al., 2013a)。由于数据集中不包含求职者的简历或属性信息，实

验仅根据其历史交互记录和职位属性信息实现前两种推荐，并调节基于内容的推荐和协同过滤的排序权重，以获得效果最优的混合推荐结果。

（9）基于人职关系建模的推荐（RM）。在本实验的数据场景下，RM采用有向、加权、多连接图刻画职位间基于内容的关系以及求职者与职位间基于互动的关系（Lu et al., 2013）。进而通过3A排序算法学习图中求职者和职位两类节点间的连接关系，预测未来发生的人职交互行为。

（10）基于图模型的推荐（GBR）。根据求职者对职位的共同浏览记录以及职位本身的属性信息，GBR基于图模型从多角度刻画职位间的相似关系，为各求职者推荐与其交互过的职位相似的其他职位（Shalaby et al., 2017）。适应于实验中的可获数据，不采用GBR原方法中与求职者简历分析相关的模块。

针对测试集中每位求职者的最后一次行为，各方法在训练后分别生成排序前5和前10的职位推荐列表，并进行十折交叉验证。以求职者实际交互的职位为正确推荐，表5.1总结了真实数据实验中基线方法对比的结果。

表5.1 网络招聘真实数据集上各方法的推荐效果

名 称		HR@5	NDCG@5	HR@10	NDCG@10
静态方法	POP	0.0262	0.0185	0.0448	0.0290
	User-CF	0.2653	0.1829	0.2857	0.2002
	CBR	0.2467	0.1647	0.2603	0.1762
动态方法	XGBoost	0.4025	0.3455	0.4630	0.3651
	DIN	0.4666	0.3667	0.5362	0.3911
	HMM	0.2449	0.1946	0.3061	0.2251
	HMM-HS	0.3961	0.3055	0.4539	0.3332
职位推荐	iHR	0.3220	0.2251	0.3422	0.2440
	RM	0.3873	0.2854	0.4229	0.3043
	GBR	0.4300	0.3196	0.4778	0.3530
DBN-HS		0.4886	0.3678	0.5627	0.3848

如表5.1所示，DBN-HS的职位推荐表现优于所有基线方法，说明其效果优势及方法设计的合理性。具体分析如下：

第一，所有个性化方法在推荐精度和排序效果上均明显优于非个性化的POP方法，说明考虑求职者个体兴趣差异来进行职位推荐的必要性。

第二，动态推荐方法整体呈现出优于静态方法的趋势，表明求职者在其决策过程中的职位兴趣确实存在动态性，刻画其动态兴趣有助于提升职位推荐效果。

第三，相较于能够灵活融合职位属性信息及求职者时序行为信息的

XGBoost和DIN方法，DBN-HS呈现出可比乃至更高的职位推荐准确性，且建模的有关求职者心理状态、兴趣和行为的依赖关系和隐层识别策略令其较深度学习方法具有更强的可理解性，说明DBN-HS在职位推荐任务的管理实践中具有足够的应用潜力和价值。

第四，无论在是否融入相似度的数据场景下，研究构建的动态贝叶斯网络均显著优于HMM的推荐表现。这说明，相较于HMM中的单层隐状态，DBN-HS所刻画的由序贯心理状态与兴趣分布共同驱动的求职者职位交互行为更符合其实际求职过程中的心理动态和行为机制，从而能有效捕捉HMM难以处理的、由潜在心理状态转换导致的动态模式复杂性和多变性。

第五，在职位推荐问题下，DBN-HS的效果显著优于三种现有的针对性职位推荐方法，说明在同样融合求职者历史职位交互行为记录和职位属性的数据场景下，本研究提出的方法能够更精确地捕捉求职者与职位间的潜在关联关系，并展现出对于求职者动态的职位兴趣明显更高的预测效力。

5.4.3 消融实验

除基线方法对比外，另构造一组消融实验，用以探究DBN-HS中各方法要素对其职位推荐表现的贡献程度。实验中，消融方法的构建主要从职位间相似度计算和动态贝叶斯网络结构两个角度出发。

(1) 职位间相似度计算方面。分别将DBN-HS方法中采用的混合相似度 h_{sim} 替换为仅共同交互相似度 i_{sim} 、仅随机森林模型预测相似度 p_{sim} ，以及各职位属性相似度的均值 $\overline{f_{sim}} = \sum_{r=1}^{|R|} f_{sim_r} / |R|$ 。

(2) 动态贝叶斯网络结构方面。分别移除DBN-HS中可观测的职位相似度节点和不可观测的兴趣节点，前者仅从求职者交互的职位ID序列中学习其潜在的心理状态和兴趣，后者忽略兴趣在求职者心理状态与职位间的中介作用，仅由状态同时生成每次行为所交互的职位及其与上次交互职位间的相似度。消融实验同样令各方法为测试集中的求职者分别生成长度为5和10的职位推荐列表，并由两个指标评估其推荐表现，结果如表5.2所示。

表5.2 DBN-HS方法的消融实验结果

方 法	HR@5	NDCG@5	HR@10	NDCG@10
i_{sim}	0.4145	0.2813	0.4798	0.3128
p_{sim}	0.4784	0.3476	0.5304	0.3719
$\overline{f_{sim}}$	0.3673	0.2512	0.4122	0.2758
无相似度	0.3188	0.2135	0.3849	0.2451
无兴趣	0.2531	0.1717	0.3469	0.2114
DBN-HS	0.4886	0.3678	0.5627	0.3848

结果表明：

一方面，相较于DBN-HS中采用的职位间混合相似度，换用其他相似度计算方式均会对推荐效果带来不同程度的削弱。仅采用共同交互视角的职位相似度（ $isim$ ）很容易受到历史交互记录稀疏性的影响，对于测试集中历史交互较少甚至尚未被交互的职位，难以准确预测求职者的潜在兴趣。而 $\overline{isim}$ 等价于在职位相似度计算中赋予所有属性相同的权重，其显著的推荐效果劣势说明各职位属性对求职者决策的影响程度存在明显差异；即使属性信息能缓解稀疏历史交互的负面作用，忽略权重差异也会更大程度地扭曲求职者视角下的职位间相似关系，从而降低职位推荐的准确性。

另一方面，结合HMM和HMM-HS的职位推荐表现可以发现，无论对于DBN-HS还是HMM，融入职位相似度均能有效提升推荐效果，说明更丰富的可观测信息有助于挖掘求职者在决策过程中的潜在心理动态。此外，兴趣隐层对于DBN-HS方法的推荐表现也具有显著贡献，表明兴趣的确是求职过程中不可忽视的重要中介，序贯心理状态需经由兴趣共同作用于求职者的职位交互行为。综上，DBN-HS方法能面向网络招聘情境下的求职过程生成机制，有机整合各方法要素，并有效融合求职者认知和职位属性信息，从而形成推荐效果上的整体优势。

5.4.4 参数分析

基于猎聘网的大规模真实数据集，本节针对DBN-HS方法的学习过程和结果进行参数分析，具体包括参数可视化和超参数讨论两部分。

1. 参数可视化

实验首先对研究提出的DBN-HS方法中动态贝叶斯网络学习到的关键参数进行可视化，以反映求职者在决策过程中由潜在心理动态生成其可观测职位交互行为的综合机制。根据隐节点识别策略，DBN-HS方法可以先验地区分和标记求职者不可观测的心理状态和兴趣，从而使其在管理实践中具有良好的可理解性。

作为直接甄别两个序贯心理状态的重要依据，实验关注各状态下求职者相邻两次行为所交互职位间的混合相似度均值和标准差，即 μ_z 和 $\sigma_z(z \in Z)$ 。由动态贝叶斯网络训练后得到：隐状态0（准备状态）的 $\mu_0=0.3150$ 、 $\sigma_0=0.1297$ ；隐状态1（活跃状态）的 $\mu_1=0.8297$ 、 $\sigma_1=0.0984$ 。可以发现，当求职者处于活跃状态时，其当前行为所交互的职位与上次交互职位间的平均混合相似度显著高于准备状态，且相似度的平均波动更小。这与DBN-HS的建模假设和隐层识别策略一致：求职者在准备状态下以广泛

查找多样的职位信息为主,在活跃状态下则集中查找相似职位的信息。两个状态下求职者相邻交互职位间混合相似度分布的明显差异,使得DBN-HS方法能够对其进行合理推断。

图5.6展示了求职者相邻行为在两个序贯心理状态间的转换概率,即 $P(z_u^{n+1}|z_u^n)$ 。其中,各箭头的起点代表求职者决策过程中第 n 次行为所处的心理状态,终点代表其第 $n+1$ 次行为的状态,箭头上的数值则为相邻两次行为对应状态的转换概率,从各状态出发的转换概率总和均为1。如图5.6所示,处于准备状态的求职者在下次职位交互行为发生时,有40.55%的概率转换至活跃状态,说明其在求职目标的驱动下会发生心理状态的前进。而当求职者处于活跃状态后,其职位兴趣以及对自身适合职位的判断已相对清晰,再次发生心理状态转换的概率明显减小。但由于在自我调节的求职过程中受到招聘方即时反馈等因素的影响,求职者仍有约23.19%的概率退回至准备状态,即重新探索其他可能感兴趣的职位。以上结果与理论研究表明的网络招聘情境下求职者的决策过程特征相一致。

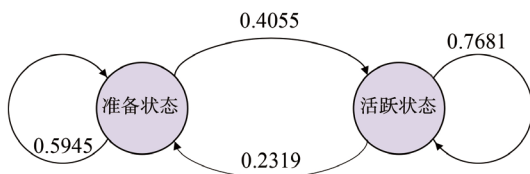


图5.6 序贯心理状态间的转换概率

图5.7采用两张热力图说明各心理状态下求职者相邻两次行为间的兴趣变化概率,即 $P(i_u^{n+1}|i_u^n, z_u^{n+1})$ 。其中,两张子图分别对应求职者第 $n+1$ 次职位交互行为可能所处的准备状态或活跃状态;每张子图中,行代表其第 n 次行为对应的兴趣,列代表第 $n+1$ 次行为的兴趣,每行概率之和为1。可视化结果表明,求职者在不同心理状态下的兴趣分布变化模式明显不同。具体地,在准备状态下,求职者的职位兴趣仍非常分散且容易改变;而在活跃状态,其相邻两次行为所对应兴趣的变化概率几乎收敛为一条对角线,即职位兴趣已较为集中且基本不再发生变化,这也与序贯心理状态的相关理论分析相一致。

综上所述,动态贝叶斯网络习得参数的可视化结果说明,DBN-HS方法能够有效学习到驱动求职者决策过程中可观测职位交互行为的潜在心理状态和兴趣的作用机制,从而实现准确的职位推荐。此外,融合情境特点和经典理论的建模及识别策略也有助于合理概念化模型中的隐节点,使其在可理解性上呈现优势。

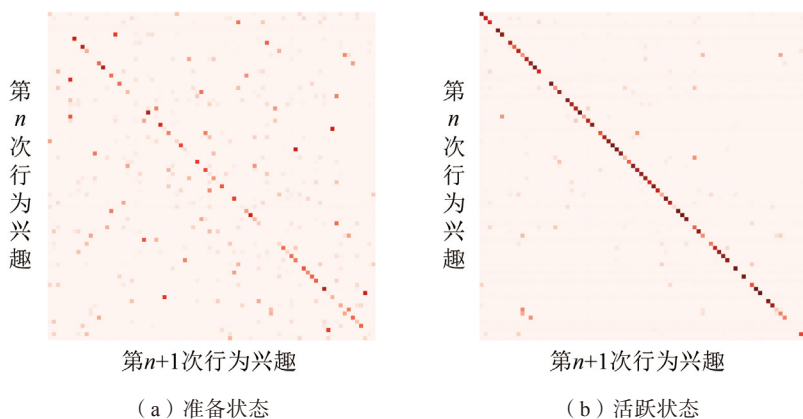


图5.7 各心理状态下的兴趣变化概率

2. 超参数讨论

实验另探究DBN-HS方法中的若干重要超参数对其推荐效果的影响，即采用HR@5、NDCG@5、HR@10和NDCG@10指标比较和讨论不同超参数设置下的职位推荐表现，为其现实应用提供启示。

DBN-HS方法借助校正后的职位间混合相似度学习求职者在线行为中的动态模式。由于混合相似度是动态贝叶斯网络中甄别求职者潜在心理状态的重要可观测变量，式(5.5)中的校正比例 w 可能成为影响方法推荐效果的关键因素。图5.8展示了不同 w 取值下，DBN-HS方法的职位推荐表现指标。结果表明，随着校正比例的增大，DBN-HS方法的职位推荐准确性呈现先升后降的整体趋势；当 $w=0.8$ （即按0.8和0.2的比例混合随机森林模型预测的职位相似度和共同交互视角下的职位相似度）时，其推荐效果达到最优。由此，通过随机森林模型的训练和预测校正，可借助职位的属性信息较大程度缓解共同交互稀疏性带来的相似度偏差，但也需调节校正比例以控制其对共同交互相似度所反映的有效信息的稀释作用，将合适的职位间混合相似度输入动态贝叶斯网络后，方法得以实现相对准确的求职过程刻画和职位推荐。

另一个需要调节的超参数是DBN-HS动态贝叶斯网络中的兴趣数 $|I|$ ，图5.9展示了在训练和测试中采用不同兴趣数时方法的推荐效果。结果表明，尽管整体上DBN-HS方法的职位推荐表现会随兴趣数的增大而提升，但当兴趣数足够大后（如80左右），继续扩大兴趣维度对推荐准确性的增益不大，甚至可能出现下降。因此，在保证动态贝叶斯网络已能充分学习求职者决策过程中的心理状态转换、兴趣变化和行为生成等动态模式的前提下，考虑到参数学习时间复杂度随兴趣数呈指数增长，同时为避免过细粒

度潜变量带来的过拟合现象，需要调节合适的参数以兼顾DBN-HS方法的推荐效率和效果。

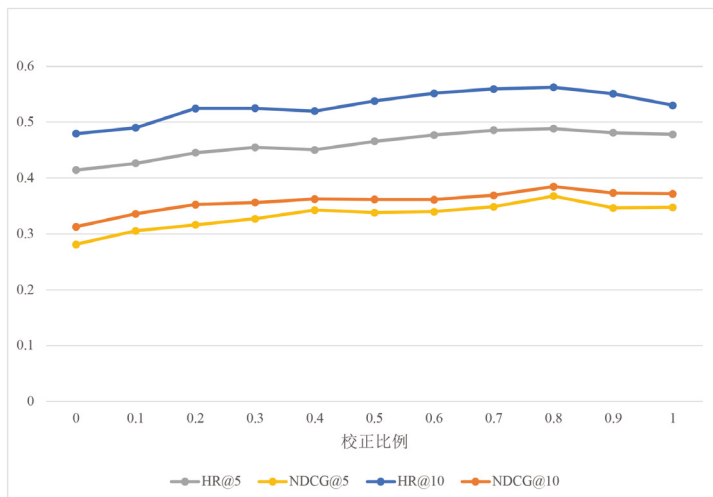


图5.8 DBN-HS在不同校正比例下的推荐表现

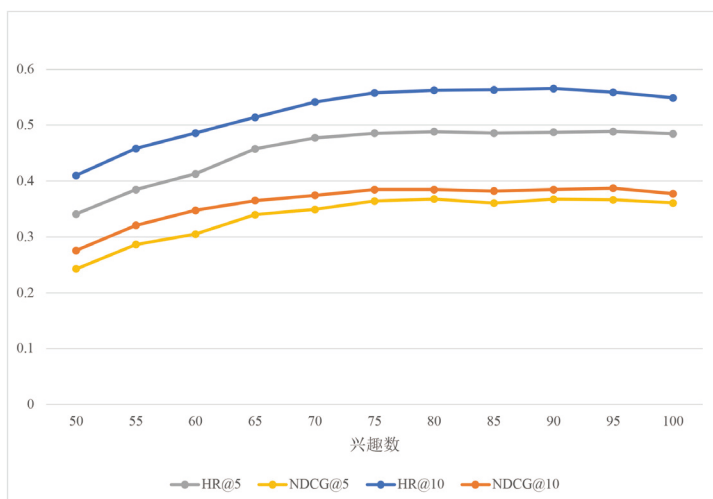


图5.9 DBN-HS在不同兴趣数下的推荐表现

5.4.5 实例讨论

结合真实数据实验中的代表性实例，可辅助说明DBN-HS方法如何实现准确且易于理解的职位推荐。

实例一：求职者A

求职者A所在地区为上海市。在数据采集期内，A共在猎聘网上产生24次职位浏览行为，历时约7天。表5.3总结了该求职者5次代表性浏览行为对

应的职位属性，以及由DBN-HS方法通过维特比式推断解码出的最大概率心理状态和兴趣（省略了部分说明性不高的职位属性）。

表5.3 求职者A职位交互行为记录

职位ID	浏览时间	企业性质	企业规模 / 人	职能分类	所属行业	心理状态	兴趣ID	兴趣峰度
2154	5/31 18:28	中外合营	≥ 10000	商务专员 / 助理	医疗 / 保健 / 美容	准备	7	11.43
306	6/5 9:13	私营 / 民营企业	500 ~ 999	销售专员 / 助理	房地产开发 / 建筑	准备	82	13.09
2100	6/5 9:15	外商独资 / 外企办事处	2000 ~ 5000	客服经理 / 主管	食品 / 饮料 / 日化	准备	90	12.11
2090	6/5 15:08	外商独资 / 外企办事处	1000 ~ 2000	助理 / 秘书 / 文员	制药 / 生物工程	活跃	3	37.70
2394	6/7 14:13	外商独资 / 外企办事处	500 ~ 999	助理 / 秘书 / 文员	制药 / 生物工程	活跃	3	37.70

如表5.3所示，A在求职前期的职位兴趣较为分散，所浏览的职位在企业性质、企业规模、职能分类、所属行业等属性上均有较大差异，因此，DBN-HS方法推断A前期更有可能处于准备状态，且最大概率的兴趣分布峰度较低，即对很多职位产生兴趣的概率相近。随着求职过程推进，A逐渐进入活跃状态，职位兴趣相应地收敛至ID为3的兴趣分布（具有37.70的较高峰度）。由此心理动态产生的职位交互行为中，各职位的属性明显更相似。根据模型推断出的A产生最后一次行为时的最大概率心理状态（活跃状态）和兴趣分布（ID为3），DBN-HS方法为A生成职位推荐列表。A的ID为2394，真实交互的职位位于列表中的第2位，其他推荐职位也与之具有高度相似的属性值。

实例二：求职者B

求职者B所在地区为北京市。在数据采集期内，B在猎聘网共产生18次职位浏览行为，历时约9天。表5.4列出了该求职者5次代表性浏览行为对应的职位属性及DBN-HS方法解码出的最大概率心理状态和兴趣（省略部分说明性不高的职位属性）。

表5.4的记录表明，B与A显著不同，B在记录前期呈现出的职位兴趣较为集中，浏览的职位在企业规模、职能分类、所属行业、年薪下限等属性上非常相似。因此，DBN-HS方法推断B此时更可能已处于活跃状态，且最大概率的兴趣分布（ID为61）峰度较高，即仅对少量的职位具有潜在兴趣。几天之后，由于求职过程中可能遇到招聘方反馈未达预期等原因，B退回至准备状态，且最大概率的兴趣发生变化，转变为峰度较低的兴趣

表5.4 求职者B职位交互行为记录

职位ID	浏览时间	企业规模 / 人	职能分类	所属行业	年薪下限 / 万	心理状态	兴趣ID	兴趣峰度
2106	4/10 13:13	≥ 10000	总裁 / 总经理	基金 / 证券 / 投资	128	活跃	61	47.33
2179	4/10 14:21	5000 ~ 9999	总裁 / 总经理	基金 / 证券 / 投资	120	活跃	61	47.33
1144	4/19 17:20	5000 ~ 9999	总裁 / 总经理	房地产开发 / 建筑	96	准备	44	18.05
2395	4/19 17:23	100 ~ 499	副总裁 / 副总经理	基金 / 证券 / 投资	96	准备	33	9.67
2189	4/19 17:27	100 ~ 499	副总裁 / 副总经理	房地产开发 / 建筑	80	准备	78	24.93

分布，即在各职位上的兴趣更为分散。由此心理动态产生的职位交互行为中，职位之间的属性差异相对更大，尤其是在企业规模、职能分类和年薪下限等属性上，B对于职位的要求明显降低，可能借此提升求职成功率。根据模型推断的B产生最后一次行为时对应的最大概率心理状态（准备状态）和兴趣分布（ID为78），DBN-HS方法为B生成职位推荐列表，其中ID为2189的真实交互职位位于列表中的第5位，其他推荐职位的属性值与之具有一定差异。相对而言，由于准备状态下求职者的职位兴趣更为分散，准确预测处于该状态的求职者所交互职位的难度大于活跃状态，但通过学习大规模训练数据中普遍存在的求职者心理机制和行为模式，DBN-HS方法仍然取得了良好的推荐效果。

综上所述，根据求职者已产生的职位交互行为记录及其中的动态模式，DBN-HS方法能够准确捕捉求职者经历的序贯心理状态转换（如求职者A由准备状态前进至活跃状态、求职者B由活跃状态后退至准备状态）以及相应的兴趣分布变化，从而对其后续最有可能感兴趣或交互的职位进行有效预测。同时，通过隐层识别策略和模型参数的可视化，可以获得各心理状态下求职者相邻两次交互职位间的平均混合相似度和各兴趣在职位集合上的概率分布峰度等量化结果，有助于将求职者潜在经历的心理状态和兴趣对应至较为明确的显式含义，使DBN-HS职位推荐方法在现实管理应用中保证足够的可理解性。

5.5 结论

在网络招聘日趋流行的背景下，为求职者提供精准且可解释的个性化职位推荐成为领域内关注的焦点。由于求职过程天然具有目标导向、双向

选择等情境化特点，直接从其他问题场景迁移而来的推荐方法难以实现最优效果。基于用户在线求职过程中持续接收信息反馈而呈现的动态行为模式，本章研究创新性地提出一种融入混合相似度的动态贝叶斯网络（DBN-HS）方法，用于刻画求职者心理和行为的动态特征，并据此生成职位推荐。通过将人力资源领域序贯心理状态相关的理论发现与方法设计进行深度耦合，提出的方法能够兼顾推荐的准确性和可解释性。

具体而言，针对稀疏的职位间共同交互相似度，DBN-HS方法采用随机森林模型对其进行一定比例的校正，自动学习各职位属性相似度在整体相似度衡量中的相对重要性。基于求职者的可观测职位交互行为序列及职位间的混合相似度信息，DBN-HS对其求职过程中潜在的心理状态转换和兴趣分布变化建模，并通过维特比式推断和纳入职位相似度的概率预测策略来获得职位推荐结果。由此，在协同过滤推荐思路的基础之上，纳入职位的属性或内容信息，形成效果更为稳健的混合推荐方法。真实数据实验的结果表明，提出的DBN-HS方法相较于现有基线方法具有更优的职位推荐表现。此外，DBN-HS方法由先验知识推衍的隐层识别策略可以对求职者不可观测的心理状态和兴趣进行概念化标定，结合参数可视化结果和实例讨论，能够说明本方法也具备良好的管理可解释性。

由此，研究提出的DBN-HS方法在一定程度上完善了职位推荐领域的智能分析，同时也体现出足够的实践价值。

第一，经随机森林模型校正后的职位间混合相似度能量化反映求职者视角下的职位相似关系，缓解共同交互相似度的稀疏性问题，不仅有助于职位推荐等相关求职过程分析和预测方法提升精度，也支持企业决策者识别其竞争企业的竞争职位（包括新进入人才市场的职位）。

第二，结合网络招聘情境特点的建模能有效地刻画求职过程中求职者的心理动态和行为机制，通过针对心理状态和潜在兴趣的准确职位推荐来提升其求职效率和成功率，并帮助平台实现高效优质的人职匹配。由于方法中未使用求职者简历数据，仅由状态转换和职位要求属性捕捉动态的双向选择过程，因此能在较好地保护求职者隐私信息的前提下提供推荐服务。

第三，借助维特比式推断和隐层识别策略，DBN-HS方法可以从求职者实时产生的职位交互行为中探测其当前所处的求职心理状态和兴趣，在为职位推荐提供充分的理由之余，也能发现有关求职者心理和行为动态模式的显式规律，从而辅助平台和招聘方进行求职过程分析和决策制定。

延续该研究问题，未来工作可进一步探索下列方向。

第一，补充结合招聘方的即时反馈信息，更深入地刻画求职过程中的双向选择，从而更精准地推荐求职者能胜任的职位。

第二，探究职位间相似度可能存在的异质性和动态性，采用更多角度的数据和分析技术来完善职位间相似关系的度量，并支持职位推荐等方法的优化设计。

第三，更充分地利用职位属性信息，通过捕捉求职者在各具体属性上的潜在兴趣和偏好，进一步提升对其职位交互行为预测的准确性，同时也为推荐结果提供更丰富的解释。