

生成式人工智能对零工工资的影响

张 艺 姜 珊

目录

附录 I 《中华人民共和国职业分类大典》(2022 年版) 职业分类规则	1
附录 II 零工招聘数据职业分类过程	2
附录 III 生成式人工智能暴露度指数构建说明	4
附录 IV 对加入企业固定效应的探讨	5
附录 V 安慰剂检验	7
附录 VI 稳健性检验	8
附录 VII 异质性检验	12
附录 VIII 附图与附表	14
参考文献	16

附录 I 《中华人民共和国职业分类大典》(2022 年版) 职业分类规

则

《中华人民共和国职业分类大典》(以下简称《职业分类大典》)由原劳动和社会保障部、国家市场监督管理总局、国家统计局联合编制,于 1999 年首次颁布,后分别于 2015 年和 2022 年进行修订。其中,2022 年版《职业分类大典》进一步更新与扩展了对新兴职业的分类,特别是首次增加“数字职业”标识“S”,共标识数字职业 97 个,更好地反映了人工智能等新技术对职业结构和分类体系的影响。因此,本文选用 2022 年版《职业分类大典》作为零工岗位的分类标准,以确保分类结果更加符合当前职业发展趋势。

具体而言,2022 年版《职业分类大典》按照各类职业的工作性质,将职业归为 8 个大类、79 个中类、449 个小类、1639 个细类,并详细描述了 1639 个细类职业的工作内容及职责,能够全面、客观、准确地反映中国社会职业的分工与功能。因此,本文选择使用 2022 年版《职业分类大典》的细类职业标准作为零工岗位的分类依据,这有助于提高分类的准确性和一致性。以“翻译”职业为例,其相关分类如下:

大类:2 专业技术人员:从事科学研究和专业技术工作的人员。

中类:2-10 新闻出版、文化专业人员:从事新闻采访报道、文图编辑校对、翻译、播音与节目主持、图书资料与档案管理、考古及文物保护等工作的专业人员。

小类:2-10-05 翻译人员:从事外国与中国语言和文字互译、中国各民族语言和文字互译以及听力残疾人士与非听力残疾人士之间互译的专业人员。

细类:2-10-05-01 翻译:以口头或书面形式,从事两种及以上语言文字间信息转换的专业人员。从事的工作内容主要包括:①运用双语或多语技能、相关专业领域知识和辅助工具,把一种语言所表达的信息通过口头形式用另一种语言同步或接续表达出来,或通过书面形式用另一种语言文字表达出来;②对照原文对笔译译文进行修改和审定,或对机器翻译的结果进行译后编辑审核;③创建、审核和管理双语或多语对照术语表、词汇库和翻译记忆库;④进行翻译理论与实践研究;⑤进行翻译教学。

在上述分类中,“2”代表“翻译”所属的大类代码,“2-10”是中类代码,“2-10-05”是小类代码,而“2-10-05-01”是细类代码。

附录 II 零工招聘数据职业分类过程

为系统、准确地识别零工岗位对应的标准职业类别,本文构建了一套利用机器学习预测标准职业的识别方法。零工招聘数据职业分类过程可以参考图 III 的流程图,具体过程如下:

(一) 第一阶段: 人工标注标准职业

在第一阶段,本文首先利用零工岗位中含有工作内容描述的岗位样本,以《职业分类大典》中的标准职业名称及职责为参照基准,结合岗位名称与其工作内容描述进行人工对比与精确标注。对于部分难以直接判断的岗位,本文在人工标注过程中还借助生成式人工智能辅助确定其对应的标准职业名称。最后,本文标注了 87 万条零工招聘广告的标准职业名称,作为后续机器学习预测岗位名称的训练集。

(二) 第二阶段: 机器学习模型训练与选择

基于零工岗位的信息特征与研究目标,本文选取“岗位名称”和“公司名称”作为机器学习模型训练与后续标准职业名称预测的依据,并通过十折交叉验证评估不同模型的预测分类性能,进而选出用于后续预测标准职业名称的最佳模型。该阶段包括以下四个步骤:

1. 文本预处理

第一步:文本分词。由于中文文本在自然状态下缺乏空格等明确的词语边界,直接使用原始文本将影响后续的特征提取与模型训练效果(胡涟漪等,2024)。因此,本文使用 Python 软件包“jieba”对合并后的岗位名称和公司名称进行分词处理。第二步:删除停用词。为提高职业识别的准确性和模型效率,本文进一步删除了常见停用词,以及文本中的无效符号(如标点符号和特殊字符)。

2. TF-IDF 向量化

在完成文本预处理后,本文采用 TF-IDF 向量化方法,将分词后的文本转换为机器学习模型可识别的数值特征。这些特征随后将被输入到贝叶斯、逻辑回归、随机森林和支持向量机等模型中,通过学习零工岗位与标准职业名称之间的映射关系,进行标准职业名称预测。具体来说,TF-IDF 是文本挖掘领域的特征加权算法,结合单文档词频(TF)与全局逆文档频率(IDF)计算词汇综合权重,通过词语的文档覆盖范围量化其分类区分能力。该算法可有效降低缺乏区分度的常见词汇(如“岗位”、“工作”)等通用词汇的特征权重,相对强化各类职业专属差异化词汇的表达权重,进而提升零工岗位标准化职业分类的预测精度。

3. 职业权重调整

由于招聘数据中不同职业类别的样本分布不均,直接进行模型训练可能导致样本量较少的职业的预测精度较低。因此,本文在模型训练过程中引入类权重机制,通过对样本量较大的职业类别赋予较低权重,对样本量较小的职业类别赋予较高权重,从而平衡类别分布不均的问题,提升模型对样本量较小职业的分类能力。

4. 模型选择

为了选择最佳的机器学习模型,本文采用十折交叉验证法评估了包括贝叶斯、逻辑回归、随机森林和支持向量机在内的多种机器学习算法的职业分类效果。该方法将训练集随机划分为 10 个子集,每次用 9 个子集训练模型,剩余 1 个子集用于验证,循环执行 10 次,以确保模型的稳定性与泛化能力。最终选择随机森林模型作为最佳机器学习模型。

(三) 第三阶段: 基于随机森林模型的职业名称预测

与第二阶段相同,本文首先对零工岗位的岗位名称和公司名称进行文本预处理,并采用TF-IDF向量化方法将其转化为机器学习模型可识别的数值特征。然后,利用随机森林模型预测其标准职业名称。最后,本文将全部零工岗位标准化为152个细类职业名称。

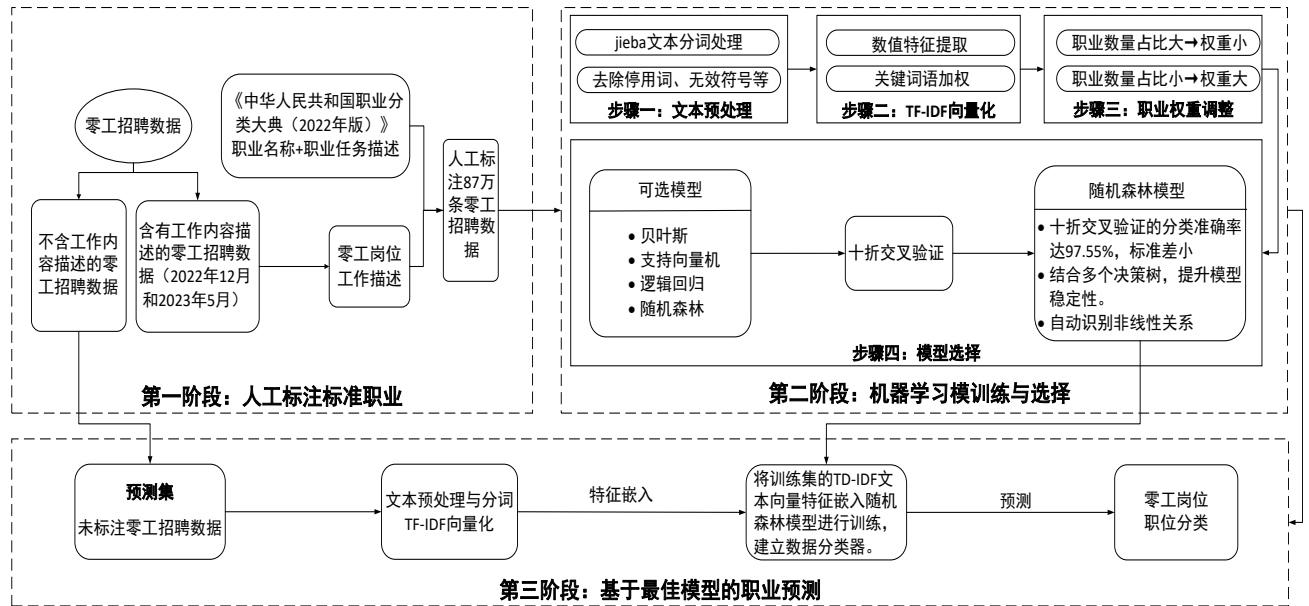


图 111 零工招聘数据职业分类过程

附录 III 生成式人工智能暴露度指数构建说明

本文借鉴已有研究 (Eloundou et al., 2024), 基于任务暴露等级构建职业层面的生成式人工智能暴露度指数。具体而言, 本文利用 ChatGPT 从任务层面评估各项工作任务受生成式人工智能影响的程度, 并根据影响大小将所有任务划分为四个暴露等级: ①E0(无暴露): 使用 ChatGPT 等类似大语言模型不会显著减少完成任务所需时间或会降低完成质量; ②E1(直接暴露): 仅使用 ChatGPT 等类似大语言模型可使任务完成时间至少减少一半; ③E2(间接暴露): 仅凭 ChatGPT 等类似大语言模型无法显著提高效率, 但若结合其他软件(如检索、自动化或图像生成等系统)可将任务完成时间减少一半; ④E3(图像暴露): 当 ChatGPT 等类似大语言模型结合具备图像识别或生成能力的系统时, 任务完成时间可减少一半。

为了将任务暴露等级转化为可量化指标以便计算职业层面暴露度, 本文参考已有研究 (Eloundou et al., 2024) 的赋分方式, 对 E0 类任务赋值 0 分, E1 类任务赋值 1 分, E2 与 E3 类任务因需结合其他工具或系统才能部分替代人工任务, 赋值为 0.5 分。接着, 对职业中所有任务的暴露等级得分取加权平均值, 得到该职业的 GAI 暴露度, 其计算公式如下:

$$GAI\ Exposure_i = \frac{\sum_1^{N_i} Score_k}{N_i}$$

其中, i 表示第 i 个职业, N_i 表示职业 i 的任务总数, $Score_k$ 是职业 i 的第 k 个任务的得分值。由于《职业分类大典》并未给出每项任务占职业的权重, 因此假定每项任务具有相同权重。该暴露度指数的取值区间为 $[0, 1]$: 取值 0 代表该职业下每一项任务的完成时间不受生成式人工智能的影响, 取值 1 代表在生成式人工智能的帮助下, 该职业的每一项任务的完成时间在确保相同完成质量的前提下至少缩短一半。

附录 IV 对加入企业固定效应的探讨

在基准回归中本文未加入企业固定效应，主要基于以下两方面考虑。第一，若加入企业固定效应，估计量将反映在同一企业内部，高暴露岗位相对于低暴露岗位在 ChatGPT 发布后的工资变化。然而在现实情形中，企业对工资的调整通常并非局限于企业内部，而更可能体现为行业层面的调整。鉴于工资具有较强的名义刚性，企业在短期内难以对相同岗位频繁或大幅度调薪，因此企业内部的相同岗位工资差异并不能完全反映技术冲击下的真实工资变化。第二，在零工招聘数据中，同一企业往往只发布单一或少量职业类型（例如仅出现高暴露岗位或仅出现低暴露岗位）。同时，受工资刚性影响，同一企业在短期内公布的岗位工资变动空间有限。在此情形下，企业固定效应将与处理变量高度共线，导致估计效率显著下降、标准误增大。基于以上原因，本文在基准回归中不加入企业固定效应。

为了检验上述分析的可靠性，本文在表 IV1 中展示了在不同样本选择下加入企业固定效应前后交互项 $Treat \times Post$ 的估计结果，以进行稳健性检验。下面依次比较表 IV1 中列（1）至列（6），并说明各列所对应的样本构成及其对估计的影响。

首先，列（1）与列（2）的回归基于全样本数据，即包含所有企业发布的招聘岗位。结果显示，在不加入企业固定效应的列（1）中，高暴露职业在 ChatGPT 发布后出现显著工资下降；在加入企业固定效应后的列（2）中，估计系数绝对值变小且不再显著。这是由于企业内部通常仅发布少量职业类型，可用于组内比较的有效变异有限，因此引入企业固定效应使得估计效率大幅下降、标准误上升，从而导致主效应无法被显著识别。

其次，列（3）与列（4）将样本限制为仅包含同时发布高暴露与低暴露岗位的企业。在此更严格的样本中，企业内部的横截面差异更充分，可用于识别企业内工资变化。结果表明，在不含企业固定效应的列（3）中，主效应显著为负；在加入企业固定效应后的列（4）中，系数依然显著为负，但估计值变小。这一组结果说明，当企业内部存在足够的岗位结构变异时，企业固定效应的加入依然能得到显著的结果，但该结果仅反映生成式人工智能冲击下的相同企业内部工资调整。这个工资下降幅度显然比不固定企业固定效应的工资下降幅度要小，原因是短期内企业普遍存在名义工资刚性，相同岗位在企业内部难以大幅度降低工资，因此企业内部相同岗位的工资下降幅度并不大。

表 IV1 不加入企业固定效应与加入企业固定效应的基准回归结果对比

	全样本		高低暴露度并存样本		仅高/低暴露度样本	
	(1)	(2)	(3)	(4)	(5)	(6)
	工资	工资	工资	工资	工资	工资
Treat×Post	-39.545*** (12.456)	-0.592 (6.257)	-15.511*** (2.571)	-5.160** (2.259)	-39.980*** (9.877)	-0.811 (6.206)
控制变量	是	是	是	是	是	是
月份固定效应	是	是	是	是	是	是
职业固定效应	是	是	是	是	是	是
城市固定效应	是	是	是	是	是	是
企业固定效应	否	是	否	是	否	是
样本量	5147530	5130152	1051439	1051439	4096059	4078577
调整后 R ²	0.569	0.811	0.657	0.791	0.574	0.828

注：括号里的数字为聚类到职业层面的稳健标准误；*、** 和 *** 分别表示在 10%、5%、1%的水平上显著，下表同。

最后，列（5）与列（6）针对的样本为仅发布单一职业类型（仅高暴露或仅低暴露岗位）

的企业。在这类企业中,处理变量在企业内部不具备横截面差异,因此,在加入企业固定效应后的列(6)中,交互项 $Treat \times Post$ 与企业固定效应高度共线,导致系数无法被有效识别并呈现不显著的结果;而不包含企业固定效应的列(5)则仍呈现显著负向影响。该对比进一步验证了加入企业固定效应的识别基础:只有当企业在暴露度上具有组内变异时,企业固定效应设定才是可行且有意义的。

综合三组结果可见,企业固定效应会吸收大量跨企业工资差异,在样本缺乏组内变异时显著降低估计效率,但不会改变主效应的方向。当识别基于具有岗位暴露度差异的企业子样本时,加入企业固定效应后的估计依然显著为负,这说明本文的基准回归结果是稳健的。

附录 V 安慰剂检验

关于冲击时点的选择,本文对其进行了一系列安慰剂的检验,使用提前一个月和两个月的虚假发布时间,以及剔除发布冲击当月的数据,以避免由于不可观测的固有差异或时间变化导致的估计误差。

(一) 虚假发布时间的安慰剂检验

为了避免处理组职业和对照组职业的零工工资差异是由于不可观测的固有差异或时间变化导致的,本文参考王锋和葛星(2022)的处理方法,将 ChatGPT-3.5 发布冲击时间分别提前 1 个月和 2 个月(即 2022 年 11 月和 10 月)来构建虚假发布时间,并将检验窗口期设定为还未受冲击的 2022 年 7 月至 2022 年 12 月。若上述研究结果是由于处理组和对照组之间的不可观测因素或时间变化所导致的,那么使用虚假的 ChatGPT-3.5 发布冲击时间也可能得到显著的结果。表 V1 的第(1)和(2)列分别将 2022 年 10 月和 11 月设定为 ChatGPT-3.5 发布冲击时间点,分别用 $Post_2210$ 和 $Post_2211$ 表示。结果显示,交互项 $Treat \times Post_2210$ 和 $Treat \times Post_2211$ 的系数均不显著。这表明处理组职业和对照组职业的工资在干预前的时间趋势不存在系统性差异,也验证了使用 2022 年 12 月作为 ChatGPT-3.5 发布冲击时间点的合理性。

(二) 剔除 ChatGPT 发布冲击当月

为了避免可能的测量误差,本文参考曹春方和张超(2020)的处理方法,将 ChatGPT 发布冲击当月 2022 年 12 月的样本剔除。具体而言,使用 $Post_2212$ 来表示剔除 2022 年 12 月的冲击时点,当月份为 2022 年 7 月至 2022 年 11 月时, $Post_2212$ 取 0; 当月份为 2023 年 1 月至 2023 年 6 月时, $Post_2212$ 取 1,其他模型设定与正文模型(1)相同。表 V1 的第(3)列为剔除 ChatGPT 发布冲击当月的回归结果,可以看出交互项 $Treat \times Post_2212$ 的系数为 -43.218,且在 1%水平上显著,验证了前文基准回归结果的稳健性。

表 V1 安慰剂检验结果

	(1)	(2)	(3)
	2022 年 10 月	2022 年 11 月	剔除冲击当月
$Treat \times Post_2210$	-6.169 (7.461)		
$Treat \times Post_2211$		-2.400 (6.041)	
$Treat \times Post_2212$			-43.218*** (13.968)
控制变量	是	是	是
月份固定效应	是	是	是
职业固定效应	是	是	是
城市固定效应	是	是	是
样本量	2101956	2101956	4646353
调整后 R^2	0.557	0.557	0.517

附录 VI 稳健性检验

为验证基准回归结果的可靠性，本文进行了多角度的稳健性检验。（1）更改职业 GAI 暴露度的计算方式，通过调整不同暴露等级任务的赋分方式调整职业 GAI 暴露度；（2）更改处理组构造方式，将处理组由二元处理组改为不同形式的连续型处理组；（3）替换被解释变量；（4）更换职业预测机器学习算法，采用多种机器学习算法进行多数投票法重新预测职业；（5）将零工工资前置一期，以考察技术冲击对工资的滞后效应；（6）仅以“岗位名称”为依据进行机器学习职业预测，以检验机器学习算法的稳健性；（7）更换标准误聚类层级至企业、城市和省份；（8）加入行业-月份和行业-职业高维固定效应；（9）更换企业层面的控制变量。

（一）更改 GAI 暴露度衡量方式

为验证基准回归结果的稳健性，本文进一步基于不同的暴露度计算方式进行稳健性检验。本文参考已有研究（Eloundou et al., 2024）的赋分调整方式，将原指标分别替换为两种不同的计算口径：一种为保守暴露度指数，仅考虑直接受生成式人工智能影响的任务（E1），忽略间接及图像暴露，即将间接暴露（E2）与图像暴露（E3）类任务的得分由 0.5 更改为 0；另一种为宽口径暴露度指数，将间接暴露（E2）与图像暴露（E3）类任务的得分由 0.5 提升为 1，以反映更广泛的技术替代潜力。回归结果参见表 VI1 的列（1）和列（2），核心解释变量 $Treat \times Post$ 的系数依然显著为负，与基准回归一致。

（二）更改处理组构造方式

基准回归中处理组的二元化处理不可避免地丢失了职业间暴露强度的细微差异，因此本文进一步采用两种不同的处理组构造方式：一种为阈值连续暴露度处理组，将暴露度大于等于 0.5 的职业保留为连续值，小于 0.5 的职业记为 0；另一种为连续暴露度处理组，在全部职业中保留连续数值，以反映完整的梯度信息。回归结果参见表 VI1 的列（3）和列（4），核心解释变量 $Treat \times Post$ 的系数依然显著为负，结果保持稳健。

表 VI1 更改职业 GAI 暴露度计算方式和处理组构造方式

	(1)	(2)	(3)	(4)
	工资	工资	工资	工资
Treat×Post	-39.990*** (13.192)	-39.547*** (12.449)	-53.203*** (16.952)	-90.467*** (27.419)
GAI 暴露类型	保守暴露度	宽口径暴露度	基准暴露度	基准暴露度
处理组类型	二元处理组	二元处理组	阈值连续处理组	连续处理组
控制变量	是	是	是	是
月份固定效应	是	是	是	是
职业固定效应	是	是	是	是
城市固定效应	是	是	是	是
样本量	5147530	5147530	5147530	5147530
调整后R ²	0.569	0.569	0.569	0.573

（三）替换被解释变量

为缩小零工工资的绝对差异，并减少异常值对估计结果的影响，本文对零工工资取自然对数处理。这有助于提高回归模型估计结果的可靠性，同时使回归系数具备更清晰的经济解释。本文在表 VI2 的第（1）列采用零工招聘工资的对数值作为被解释变量重新进行回归，结果显示， $Treat \times Post$ 的回归系数依然显著为负，表明与低 GAI 暴露度的零工岗位相比，高 GAI 暴露度的零工岗位工资下降了 18%。这一结果与基准回归的结果（19.9%）相近，验证了前文基准回归结果的可靠性。

（四）更换机器学习算法

在基准回归中，本文使用了职业分类效果最佳的随机森林算法进行职业预测。然而，使用单一算法预测职业可能存在预测偏差和局限性，因此本文结合逻辑回归、支持向量机、贝叶斯、随机森林四种机器学习算法，采用多数投票法进行职业预测。具体来说，当四种机器学习算法中存在某种在次数上占优的预测职业时，例如某预测职业出现 3 次，或出现 2 次且其他两种算法的预测职业不统一时，本文将该岗位标注为该职业；当四种算法的预测结果不存在某一种占优预测职业时，本文则将该岗位标注为空值，并在回归分析中剔除此类样本。表 VI2 的第（2）列展示了相应的回归结果，交互项 $Treat \times Post$ 的回归系数依然显著为负，验证了基准回归结果的稳健性。

（五）零工工资前置一期

生成式人工智能带来的冲击，可能会对工资产生滞后效应。为检验工资下降效应的动态合理性，本文将被解释变量零工工资前置一期，即下一个月的零工工资，以考察本期冲击对下一期工资的影响。表 VI2 的第（3）列回归结果显示，ChatGPT-3.5 发布当月的冲击对下一期零工工资的下降效应依然显著为负，且数值与基准回归的结果接近。

表 VI2 稳健性检验

	(1)	(2)	(3)	(4)
	工资对数	工资	工资 t+1 期	工资
$Treat \times Post$	-0.180*** (0.062)	-28.218*** (9.372)	-44.654*** (14.945)	-27.974** (12.702)
机器学习模型	随机森林	四种算法投票结果	随机森林	随机森林
职业名称预测依据	岗位和公司名称	岗位和公司名称	岗位和公司名称	岗位名称
控制变量	是	是	是	是
月份固定效应	是	是	是	是
职业固定效应	是	是	是	是
城市固定效应	是	是	是	是
样本量	5147530	4811822	4887747	5147436
调整后 R^2	0.549	0.600	0.739	0.573

（六）岗位名称预测

基准回归将岗位名称与公司名称作为预测依据，以提升职业预测的准确率。但为了排除公司名称对结果的潜在影响，本文仅使用岗位名称进行职业预测，以检验机器学习算法的稳健性。表 VI2 的第（4）列的回归结果中，核心解释变量 $Treat \times Post$ 的系数依然显著为负，

表明 ChatGPT-3.5 发布冲击对零工工资的负向影响是稳健的。

(七) 改变标准误聚类方式

本文在基准模型设定中将标准误聚类到职业层面,但由于零工工资也会受到地区和企业层面不可观测异质性因素的影响,可能存在除职业层面以外的误差相关性。因此,本文在表 VI3 的列(1)至列(3)中依次将模型的标准误聚类到企业、城市、省级层面,容许同一层级内的企业、城市或省份之间存在相关性。从实证结果来看, $Treat \times Post$ 系数的标准误有所变化,但其系数均在 1%水平上显著为负,表明标准误聚类方式的调整对估计结果的影响较小,回归结果保持稳健。

(八) 增加高维固定效应

为了进一步缓解行业层面随时间变化的不可观测因素,以及行业与职业交叉层面的不可观测因素对回归结果的影响,并减少基准回归模型中的内生性问题,本文进一步在表 VI3 的第(4)和(5)列控制了行业-月份交互固定效应和行业-职业交互固定效应。回归结果显示,核心解释变量 $Treat \times Post$ 的回归系数均在 1%水平上显著为负,说明本文的基本结论是稳健的。

表 VI3 改变聚类方式和固定效应稳健性检验

	(1)	(2)	(3)	(4)	(5)
	工资	工资	工资	工资	工资
$Treat \times Post$	-39.545*** (6.908)	-39.545*** (1.105)	-39.545*** (1.400)	-30.473*** (9.557)	-36.878*** (10.046)
控制变量	是	是	是	是	是
月份固定效应	是	是	是	是	是
职业固定效应	是	是	是	是	是
城市固定效应	是	是	是	是	是
行业-月份固定效应	否	否	否	是	否
行业-职业固定效应	否	否	否	否	是
聚类层面	企业	城市	省份	职业	职业
样本量	5147530	5147530	5147530	5147528	5147387
调整后R ²	0.569	0.569	0.569	0.625	0.597

(九) 更换控制变量

在基准回归模型中,考虑到数据的可得性和控制变量的完整性,本文引入了企业成立年限作为企业层面的控制变量。然而,由于企业成立年限变量的引入会导致回归样本数量的减少,因此本文选择剔除该变量以扩大样本。表 VI4 的第(1)列的回归结果显示,核心解释变量 $Treat \times Post$ 的系数绝对值有所增大,但依然在 1%的水平上保持显著,验证了基准回归结果的可靠性。

此外,为了进一步提高回归模型的准确性和全面性,本文在控制企业成立年限的基础上,分别引入注册资本和实缴资本作为企业层面的控制变量。表 VI4 的第(2)列的结果显示,在控制注册资本后,核心解释变量 $Treat \times Post$ 的回归系数大小和显著性水平与基准回归结果基本一致。考虑到注册资本并不能完全反映企业的实际资本投入,因此,本文在表 VI4 的第(3)列中改为控制企业实缴资本,由于该变量缺失值较多,最终只剩下约 12%的样本。结

果显示,核心解释变量 $Treat \times Post$ 的回归系数绝对值减小,但依然显著。这表明即使在加入了更为严格的控制变量后,基准回归结果依然保持稳健。

表 VI4 更换控制变量稳健性检验

	(1)	(2)	(3)
	工资	工资	工资
Treat $\times$ Post	-40.725*** (12.219)	-40.946*** (12.905)	-16.693* (8.656)
注册资本	否	是	否
实缴资本	否	否	是
企业成立年限	否	是	是
其他控制变量	是	是	是
月份固定效应	是	是	是
职业固定效应	是	是	是
城市固定效应	是	是	是
样本量	5254605	4823268	623497
调整后R ²	0.564	0.600	0.588

附录 VII 异质性检验

本节从劳动者特征与岗位属性的角度,系统的讨论了生成式人工智能的影响是否在不同群体与不同类型岗位之间存在异质性。

（一）收入水平

传统技术变革通常认为低收入岗位更容易被技术替代。然而,生成式人工智能对高收入岗位的任务替代能力更强,导致其工资下降幅度更大;而低收入岗位的替代空间有限、工资降幅较小。为验证上述推论,本文根据零工招聘工资中位数将样本划分为高收入和低收入零工岗位进行分组回归检验。表 VIII 中第(1)和(2)列结果显示,生成式人工智能对高收入岗位的零工工资的负向影响显著,而对低收入岗位则不显著,且两者的系数差异在 1%水平上显著。

（二）工作场景

生成式人工智能导致线上岗位工资显著下降,而对线下工作内容的替代有限,因此线下岗位的工资下降幅度较小。本文以岗位名称是否包含“线上”关键词作为样本分类依据,将样本划分为线上、一般工作两个子样本。具体来说,岗位名称含“线上”关键词的样本,定义为“线上工作”;而岗位名称不含“线上”关键词的样本,因未明确限定工作场景,不存在清晰的场景指向,故定义为“一般工作”。表 VIII 中第(3)和(4)列的回归结果显示,生成式人工智能对线上工作工资的负向影响要远高于一般工作,且组间系数差异在 1%水平上显著。

表 VIII 收入水平与工作场景异质性检验

	(1)	(2)	(3)	(4)
	高收入	低收入	线上工作	一般工作
<i>Treat × Post</i>	-30.548*** (5.851)	-7.023 (10.352)	-54.715*** (6.437)	-36.149*** (13.144)
控制变量	是	是	是	是
月份固定效应	是	是	是	是
职业固定效应	是	是	是	是
城市固定效应	是	是	是	是
样本量	3152394	1995117	29287	5118235
调整后 R^2	0.354	0.429	0.893	0.568
组间系数差异检验	P值=0.000		P值=0.000	

（三）行业类型

(1) 制造业与服务业。生成式人工智能对制造业岗位工资的影响有限,而服务业岗位更易受到生成式人工智能的影响,导致工资下降更明显。按《国民经济行业(GB/T 4754—2017)》与《三次产业划分规定(2018)》将样本划分成制造业与服务业两个子样本。表 VII 2 中第(1)和(2)列结果显示,生成式人工智能在服务业样本中的影响显著为负,而在制造业样本中的影响不显著,且两者组间系数存在显著差异。

(2) 高技术行业与非高技术行业。高技术行业零工工资下降幅度小于非高技术行业。根据国家统计局印发的《高技术产业(制造业)分类(2017)》和《高技术产业(服务业)分类(2018)》中定义的高科技行业目录,并结合《国民经济行业分类(GB/T 4754—2017)》

中企业所属行业小类信息,将样本划分为高技术行业与非高技术行业进行分组回归。表 VII 2 中第(3)和(4)列结果显示非高技术行业工资下降幅度更大,且两者组间系数差异在 1% 水平上显著。

表 VII 2 行业层面异质性分析

	(1)	(2)	(3)	(4)
	服务业	制造业	非高技术行业	高技术行业
<i>Treat</i> × <i>Post</i>	-41.870*** (12.511)	-12.147 (7.487)	-43.167*** (14.930)	-17.569*** (5.993)
控制变量	是	是	是	是
月份固定效应	是	是	是	是
职业固定效应	是	是	是	是
城市固定效应	是	是	是	是
样本量	4779900	365170	3124846	357397
调整后 R^2	0.562	0.742	0.557	0.606
组间系数差异检验	P值=0.000		P值=0.000	

附录 VIII 附图与附表

如图 A1 所示，在 2022 年 12 月初，“ChatGPT”关键词的百度搜索词指数出现快速上涨，这说明在 ChatGPT-3.5 上线后，公众对这一突破性技术产生了极大的兴趣。相比之下，其他通用 AI 关键词（如“AI”、“人工智能”、“GPT”）的搜索量并未同步上升，这一现象说明公众关注度聚焦于 ChatGPT，而不是整个人工智能技术，或者是其他一般 GPT 技术。因此可以将 ChatGPT-3.5 的发布时间作为生成式人工智能在中国产生技术冲击的始点。

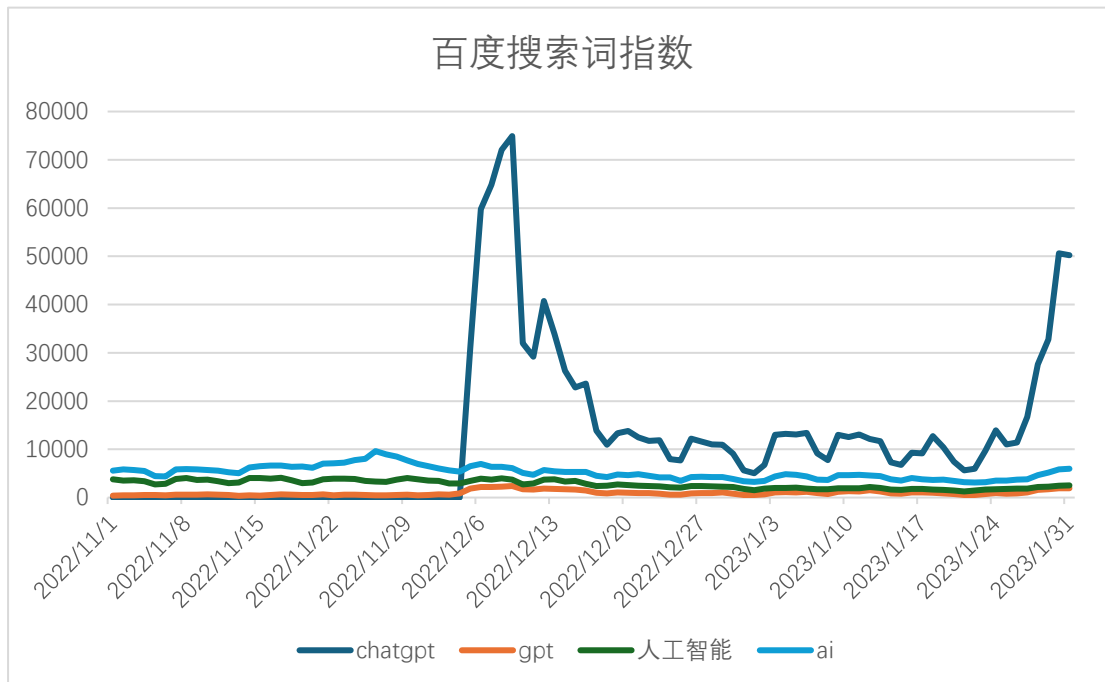


图 A1 ChatGPT 在 2022 年 11 月至 2023 年 1 月的百度搜索词指数变化

注：数据来源于百度指数：<https://index.baidu.com/v2/index.html#/>。

由于样本中包含 152 个职业，若将所有职业的出现频次逐一列示显得过于冗长。因此，本文在表 A1 与表 A2 分别统计并汇报干预组与对照组出现次数最多的前 10 个职业的频数和 GAI 暴露度，以展示受到生成式人工智能冲击的干预组和对照组的职业集中度。

表 A1 干预组出现频次最多的 10 个职业

排名	职业名称	出现频数	GAI 暴露度
1	网络主播	74342	0.825
2	数字媒体艺术专业人员	55329	0.714
3	互联网营销师	31157	0.714
4	营销员	22314	0.606
5	剪辑师	19846	0.517
6	市场营销专业人员	8433	0.510
7	招聘师	3877	0.700
8	会计专业人员	3576	0.675
9	人工智能训练师	770	0.690
10	打字员	175	0.550

表 A2 对照组出现频次最多的 10 个职业

排名	职业名称	出现频数	GAI 暴露度
1	模特	2302070	0.100
2	网约配送员	1207762	0.240
3	包装工	802058	0.156
4	理货员	306086	0.033
5	保安员	129715	0.086
6	餐厅服务员	43750	0.143
7	其他教学人员	22580	0.253
8	快件处理员	20085	0.238
9	汽车代驾员	19406	0.040
10	客运车辆驾驶员	18761	0.013

为了展示工资总体分布形态,图 A2 汇报了零工岗位日工资的核密度分布。可以观察到,该分布呈现明显的三峰特征,对应约 100 元、200 元和 270 元的低、中、高日工资,反映出零工市场存在清晰的工资分层。

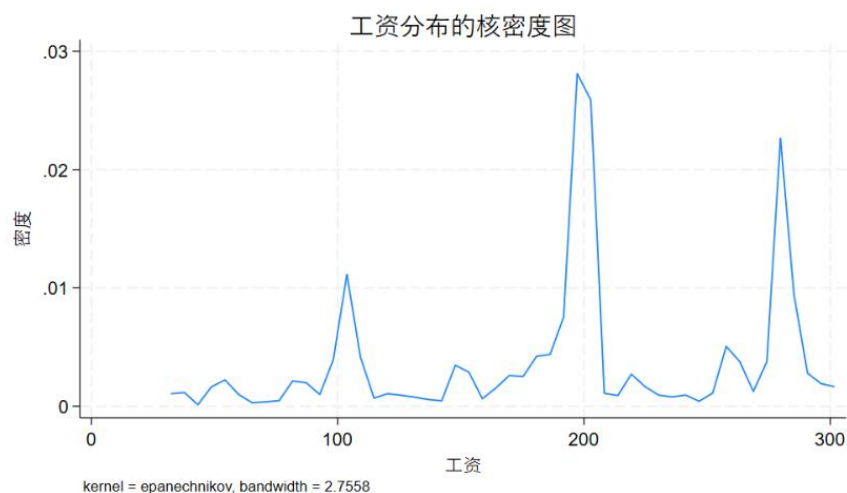


图 A2 工资核密度图

参考文献

- [1] 曹春方、张超,“产权权利束分割与国企创新——基于中央企业分红权激励改革的证据”,《管理世界》,2020 年第 9 期,第 155—168 页。
- [2] Eloundou, T., S. Manning, P. Mishkin, and D. Rock, “GPTs Are GPTs: Labor Market Impact Potential of LLMs”, *Science*, 2024, 384(6702), 1306-1308.
- [3] 胡涟漪、盖庆恩、朱喜、郭士祺,“中国职业技能结构转型:任务内容的视角”,《经济研究》,2024 年第 1 期,第 188—207 页。
- [4] 王锋、葛星,“低碳转型冲击就业吗——来自低碳城市试点的经验证据”,《中国工业经济》,2022 年第 5 期,第 81—99 页。

注:该附录是期刊所发表论文的组成部分,同样视为作者公开发表的内容。如研究中使用该附录中的内容,请务必在研究成果上注明附录下载出处。