

人工智能,是幸福还是威胁?

如果说2017年互联网行业最热的词是什么,也许非“人工智能”莫属。

从1956年前后人工智能技术开始系统发展直至20世纪80年代,在这段时间里,受制于当时较低水平的计算机技术,人工智能只能执行来自人类的指令和完成一些简单的计算。它将人类从繁琐的重复性劳动中解脱出来,因而被当做人类解决问题、处理工作的好工具和好帮手。可以说,那时候的人工智能只是自动化程度较高的工具,为了增进“绝大多数人的最大幸福”而缓慢发展着。

转折性的事件发生在1997年5月,IBM公司研制的深蓝(Deep Blue)计算机战胜了国际象棋大师卡斯帕洛夫(Kasparov)。从那时起至今,人工智能在挑战人类的道路上高歌猛进:无人驾驶汽车、自动撰写新闻程序、视觉识别系统……而最具颠覆性的事件发生在2016年3月,谷歌Deep Mind研发的AlphaGo在围棋对弈中击败韩国职业九段棋手李世石。人类突然惊慌失措地发现,一直以来被人类视为体现着人类最高智慧的围棋运动已被人工智能超越。但人工智能并未就此停下脚步。2017年10月19日,Alpha Zero(阿尔法元)的机器系统仅训练3天就战胜了曾经大败李世石的人工智能系统AlphaGo,比分为达到惊人的100:0。更令人震惊的是,Alpha Zero是谷歌全新设计的人工智能程序,其对围棋的学习完全从零开始,不需要任何历史棋谱和人类的先验知识,并且仅仅通过3天的训练就完胜屡战屡胜、炙手可热的前代人工智能AlphaGo……

在意识到人工智能日新月异的自我进化速度之后,人类社会哀声震天。棋手古力在微博上无不忧伤地表示“20年不抵3天,我们的伤感,人类的进步”,柯洁则感叹“对于AlphaGo的自我进步来讲,人类太多余了”。来势汹汹的人工智能技术令“人工智能威胁论”一时间甚嚣尘上。Deep Mind公司创始人之一谢恩·莱格警告称,人工智能是“21世纪第一大危险”,他相信人工智能将是人类灭绝的主因之一。物理学家斯蒂芬·霍金则表示,几乎可以确定的是,在未来1000年到10000年内,一场严重的技术灾难将对人类带来威胁,将可能把人类的生存带向“错误的方向”……

但是,人工智能真的会威胁人类吗?

BBC 基于剑桥大学研究者 Michael Osborne 和 Carl Frey 的数据体系分析了 365 种职业在未来的“被淘汰概率”,最后的结论认为:具有重复性劳动、无须天赋的工作被取代的可能性非常大,而那些要求社交、协商、创意和审美的工作是最难被取代的。此外,人工智能的运行成本非常高昂。加州大学洛杉矶分校(UCLA)“视觉、认知、学习与自主机器人中心”主任朱松纯曾提到:以现在的技术,要让一个机器人长时间像人一样处理问题,可能要自带两个微型的核电站,一个发电驱动机械和计算设备,另一个发电驱动冷却系统。而相比较下,人脑的功耗仅仅是 10—25 瓦。

《人类简史:从动物到上帝》一书认为,人类花费了 250 万年的时间才爬上了食物链的顶端。即便是从智人算起,人类成为万物主宰也花费了 10 万年的时间。而人工智能从 1956 年出现至今不过短短数十年,如何跨越这 10 万年的进化长河从而比肩人类?况且人工智能进化的同时人类也在进化着。“魔高一尺、道高一丈”。或许当人工智能在社交、协商、创意和审美的工作上追上人类的时候,人类已经想出了制服人工智能的诸多办法。尤其需要特别指出的是,要想创造出具有创意和审美等自主意识的人造物体或个体,就必须破解意识的形成机制,即破解“意识是如何从物质中转化产生的”这一终极谜题。在这一终极谜题被破解之前,人工智能恐怕无法像人类一样拥有自主意识,也不可能拥有颠覆人类社会的邪恶意念,因而必将只能是唯人类马首是瞻的好伙伴、好帮手、好管家而已。

因此,在“意识是如何从物质中转化产生的”终极谜题被破解之前,我们法律人大可不必惊慌失措于“人工智能威胁论”,而是应当秉持“仰天大笑出门去,我辈岂是机器人”的自信心态,为人工智能技术的发展创造良好的制度环境,让人工智能技术更多、更好地服务于“绝大多数人的最大幸福”!

朱艺浩

丁酉鸡年十二月二十七日