

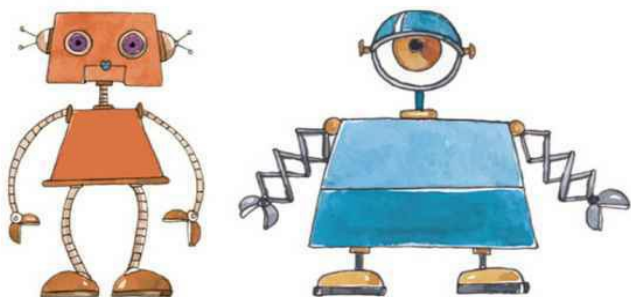
人工智能 十万个为什么

热门知识

智能相对论◎著



其实是一本严肃的人工智能通识书



Artificial
Intelligence

医院、学校、酒店、餐厅……
人工智能就在你身边



中国工信出版集团



电子工业出版社
PUBLISHING HOUSE OF ELECTRONICS INDUSTRY

1 医疗“黑洞”

拿着手术刀的AI医生会看病吗？

设想一下，你去一家医院看病，一进诊疗室的门就有一位护士不断地为你拍照，然后这些照片会被上传到一台AI（Artificial Intelligence，人工智能）设备中，这个设备会根据照片中你的模样来进行病情诊断……而在整个过程中，不会出现任何专业的人类医生。

你是不是觉得这种场景不可思议？即使现在AI医疗已经发展得很快，AI在医疗领域中实现了不同程度的落地，如AI识别医学影像、药物研发、AI辅助诊断等，但AI起到的还只是辅助作用，最终负责决断的依旧是人。让AI执证上岗，独立做临床诊断，还从未听说过。

然而，现在这样的“看病模式”已经有萌芽了。

美国食品和药物管理局（FDA）首次批准了一种AI诊断设备IDx-DR，该设备可以通过观察视网膜的照片来检测一种眼科疾病，并且不需要专家医生的参与。

也就是说，这个名为IDx-DR的AI设备竟然有了“上岗证”，成为一名真正的“医生”。

科学家不断攻破一个又一个技术难关，人们在高兴的同时，也有隐忧。医疗AI之路越走越顺畅，但现在就出现独立的AI医

生，合适吗？

AI医生“独立出诊”，产业链还不完整

我们要想让AI医生独立起来，必须在一开始就深入研究产业布局和各产业链每个环节的协调发展，否则，只要有一个环节发展不良，就会导致智能医疗的结构出现上下游之间的断档，或被技术伦理问题所牵绊。

1. AI医生的落地还没有标准

从患者端或者其他的医疗使用端来看，医疗AI其实在短时间内不会有特别大的变化。因为“上岗证”批不下来，甚至连如何为一个AI医生去审批“上岗证”也是待探讨的。AI医生合格的标准是什么？是器械的精密性，还是诊断的正确率？即使是FDA批准的IDx-Dr，一项使用了900多张图像的临床试验检测到视网膜病变的准确率也仅是87%。

归根结底，AI医生能否落地，是行政是否授权的问题。在医疗领域，一个产品的落地，必定涉及许可证和医学严谨性等问题，聘用一个独立的AI医生，可能还有比较长远的路要走。

2. “售后”服务不好办

在现实生活中，患者如果遇到了医生误诊的问题，可以要求医院赔偿或者处分该医生。医生给你看病，5个人里治好3个人可能就差不多了。然而，AI给你看病，可能100个人里就错了1个人，唯一被看错的那个人会怎么想？遇到水平不够好的医生，人

们还能自嘲一句“眼光不好”。遇到误诊的AI，人们恐怕就没那么宽容了。

首先，追究医院和厂家的责任肯定少不了。然后怎么办？“罪魁祸首”AI还没有受到任何处分呢。

要销毁这个AI医生吗？或者把这个AI医生的“头脑”格式化，以示惩罚？但是，无视它看对了99个病人的功劳似乎不妥。而且，在医学方面，随着电子病历和数字胶片的积累，大量结构化病例被用于机器学习，对于AI医生，这个模型训练的大数据至少是以10万份为起点。

AI是一种集体交付的结果，从程序、算法的开发到机械安装，打造一个AI医生的成本是难以计算的。假设患者家属一时气愤，怒摔机器，恐怕还会收到一笔昂贵的赔付账单。所以，如果AI出错，责任可以由医院和公司来承担，但对于这台“犯错”了的机器，要如何处置才能平息患者的怒火呢？

3. AI医生让“患者”变为“消费者”

医疗行业有一个特点：核心服务由单个专业技术人员提供。去医院看病，我们会关心哪个医生出诊，会去关注这个医生的口碑如何，服务地点和所在机构在很大程度上也会影响我们的评价，这些因素都会影响患者的就医决策。三甲医院的医生和二甲医院的医生，你更倾向选谁做主治医生呢？

一旦拥有了自主的AI医生，AI便不再只是作为医院宣传的噱头，以及提升医生效率的工具，而是进行独立诊断，成为一

个“专业人员”。虽然核心服务依旧由单个专业技术人员提供，但服务地点和医疗机构似乎不那么重要了，创新者的话语权将会更大，影响患者就医决策的将会是生产这个AI医生的公司。如此，患者的身份会更贴近“消费者”。在代入“消费者”这层身份后，医患关系也会变得冰冷。

医学AI或许比医疗AI更靠谱

其实比起越来越火的智能医疗，医学AI可能更符合当今社会的发展。我们要明确，医学和医疗其实是两个概念，医学是科学，而医疗是以医学科学为基础的实践技术。

所以，在未来，像AI、大数据这样的新技术将会带动整个医学的进步。但是医疗尤其是自主的AI医生涉及“一线”操作，人命关天，其发展可能不会像我们想象的那么快。因为无论如何，医疗还是应该以安全、成熟、稳定作为前提的。

对于科学来讲，高质量的数据是发展的宝贵引擎。而医疗健康行业、医药研发行业则是一座数据的金矿，医学AI是很好的研究方向、发展方向，更多的科研领域会快速地有成果出来，这些领域也会有更好的商业机会。

麦肯锡估计，制药和医学方面的大数据和机器学习每年可以产生高达1000亿美元的价值。这些价值源自更好的决策、优化创新、提高研究和临床试验的效率，以及为医生、患者和医疗机构创造新的诊疗手段等。

例如，IBM沃森研究中心肿瘤学部门和斯隆凯特林纪念医院

在进行个性化医学的研究，他们致力于使用患者医疗信息和诊疗历史来选择最优治疗方案。此外，还有许多公司也致力于研究此类产品。

谷歌前CEO施密特曾说过：“计算机确实可以在分析有用信息方面发挥作用，如预测疾病的结果。但如果我们生病了，仍会倾向找医生来看病，不过医生需要掌握最新的医疗技术来帮助其做决定。”

本质上，AI能够增强人类的智能。正如蒸汽机节省了人类的体能、电话加强了人类之间的联系、计算机强化了人类的计算能力一样，机器的协助并没有取代人类的活动，它只是扩展了人的技能和专业水平。

所以，就医疗AI而言，一个独立的AI医生或许比不上一个起辅助作用的虚拟助手。如果智能医疗的创新能够站在医生的角度，为医生赋能，那么创新成功的概率会大大增加。因为这样能给医生带来实实在在的好处，医生也会积极主动地配合。但如果将AI医生放进医院，医生只会觉得自己的职业受到了威胁，主动性恐怕也会受影响。

从目前的趋势来看，恐怕也没有多少企业想要打造一个独立的机器人医生。正如很多新闻报道的那样，AI现在能够帮助人们建立患者病历，节省了医生的时间，AI还可分析X光片和CT片，不过诊断和开药还是只能由医生完成。目前AI在医学方面主要用于辅助诊疗、健康管理、信息化管理、医学影像等，如图1-1所示。

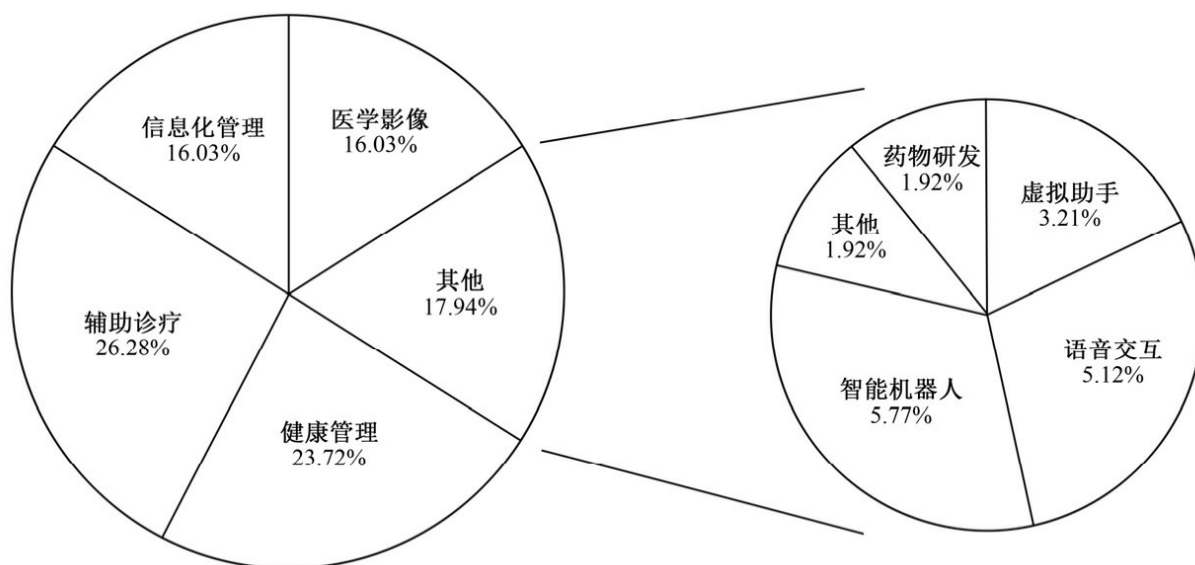


图1-1

虽然目前很多大的科技公司都在努力，谷歌也在2018年推出了一款“AI+AR”的肿瘤诊断系统，但现在的医疗生态系统依旧处于“石器时代”，因为很多系统还不够完善。我们要认清努力的方向，提高医学领域的进步速度，通过培养更多的AI人才来加速医学领域的AI研究进程。

思维移植会不会比“长生不老”更不靠谱？

人类对于长生不老总是有一种执念。从古时候的炼丹术，到现代社会的基因工程、冷冻人体等技术的发展，都体现了人们对长生的渴望。

人类要如何实现“长生不老”？除了现代科技提供的基因克隆的方案，美国人玛蒂娜还提出了“思维克隆人”的想法——以人们留在计算机和互联网中的数字痕迹（思维文件），包括聊天记录、照片和视频等数据为基础，通过与人类大脑功能相同的复制品“思维软件”产生和本人相近的人类意识，这就是“被克隆的思维”。

也就是说，通过思维的移植，人们有可能实现“长生不老”。有一家公司就涉足了这一领域，美国初创科技公司NECTOME试图通过手术来“上传大脑”，将思想实现永久化的数字保存。目前，已经有25个人交了1万美元的订金预订了该服务。

毫无疑问，如果记忆可以上传，思维可以移植，那么这将会对人类社会产生非常深刻的影响。当《记忆大师》的电影情节映射到现实生活中时，我们应该如何看待这项技术呢？

脑痴呆的终结者

使用思维移植，将有利于晚期疾病和慢性疾病患者的治疗，思维与身体的剥离意味着人们的意识留存时间将会延长，将死的人也拥有了等待医学进步的成本，这与冰冻人体相似。不同的

是，思维移植的受益者在身体不受自我支配后，依然能继续与社会进行交互。因此，这项技术对于阿尔茨海默病患者的意义尤为重大。

阿尔茨海默病夺去了许多老人的正常智力，却不会对老人的躯体有所损伤。这类患者如果能够将自己的思维分载到计算机上，他们就可以在进行治疗的过程中保持自我意识。

这有些像人造心脏分载了患者心脏的工作，在新的心脏供体被找到之前，人造心脏会一直维持患者的生命。最终，阿尔茨海默病患者可以重新将自己的思维传至已经治愈的大脑。即使身体还在治疗，但患者可以继续与家人交互。

受益的不仅仅是阿尔茨海默病患者，还有智力障碍儿童。在我国，每年新生儿智障人数约为110万，而介于智障与正常之间（一般以智商低于70为确认标准）的“边缘智障”人群则更加庞大，约有数千万之巨。

一般来说，智力障碍的患病原因有如下4种：①遗传因素，基因突变，如先天性代谢异常病；②产前损害，包括宫内感染、缺氧等；③分娩时产伤，包括窒息、颅内出血、早产儿等；④出生后患病，包括脑膜炎、脑炎、颅外伤、脑血管意外等。

对于不可逆的脑部损伤，目前医疗领域还没有很好的治疗方法。

思维的太空旅游

SpaceX公司首席执行官兼首席技术官埃隆·马斯克说过：“历史可以分为两条道路。第一条是我们将永远被限制在地球上，直到我们面临一些灭绝事件。另一条是我们将成为太空物种、多行星物种。我认为第二条道路是一条更加令人振奋的未来。”

自人类首次飞上太空已过去半个多世纪，但至今也只有极少部分人进入过太空。目前，个人想要上太空，就必须支付2000万美元（约合1.37亿元人民币），这是普通人难以承担的。

目前，一些公司正努力开发太空旅行项目，让普通人实现遨游太空的梦想。但是，送人上太空还面临着多重因素的影响，如重力、人们赖以生存的氧气、太空里的进食和排泄，以及如何返回等。

事实上，我们不必这么麻烦，如果思维可以移植，我们完全可以换一种“姿势”登上太空——让我们的意识先上去。

比起将人带入太空，将载有某种信息的芯片带上太空会更加容易。NASA的探测器InSight已经于2018年11月在火星地面着陆，与此同时，2429807个名字也被带上了太空。将这些名字带上火星的是一块只有0.8平方厘米的硅晶片，其通过一个电子束来形成成行的字母。

“神助攻”在哪里

2018年，科学家发现了细胞命运的密码，可以让细胞“返老还童”，让细胞从成年体细胞状态回到可以定向分化的多能干细胞状态，为再生医学提供“种子”细胞来源。

再生医学的优势在于通过改善再生微环境，患者借助自身的再生修复能力引导再生，再生后的组织是人体自身的一部分。2017年，34岁的卵巢早衰患者方女士在南京鼓楼医院产下健康的男宝宝，这是全球首例干细胞复合胶原支架治疗卵巢早衰临床研究诞生的婴儿。

再生医学的发展让思维移植有了医学基础。将来，随着再生医学技术的不断进步，只要构建合适的微环境，人体组织都有可能进行再生——这也意味着经思维移植的人们可以获得一个健康的大脑。

AI在这个领域也功不可没。人体组织如何再生？科学家们总是在不同生物的身体上找切入口。例如，涡虫是一种微小的虫子，拥有超乎想象的再生能力，当它被拦腰切断以后，能重新长出完整的身体。这种能力使得它们成为再生医学的热门研究对象。

但庞大的数据让科学家们很是头疼。受到进化论的启发，美国塔夫斯大学的科学家开发了一个AI系统，它有助于充分挖掘浩如烟海的发育生物学实验数据，并且建立起了一个完备的模型来解释这种现象。

科学家们开始利用AI来观察水螅的行为。水螅是一种与水母和珊瑚有关的微小动物。在过去，研究人员研究了该动物的神经系统，找出了它的大脑的哪些部分促使了它的行为。现在，研究人员利用AI来追踪所有的这些行为。

该小组使用了AI算法，自动注释水螅的行为，包括所有的摆动、枢轴、伸展和弯曲。他们了解到，无论环境条件（如光照、温度或附近食物的数量）如何，这些物种实际上只参与了6种基本的行为。

最后，研究人员发现，即使周围的环境变化再复杂，水螅仍能像往常一样继续工作。利用这一点，如果能发现水螅的神经系统是如何与行为相连接的，我们或许能够让大脑承受更加极端的治疗环境，包括冷冻技术、电力刺激等。

思维移植的重重壁垒

不可否认，思维的移植一定会牵涉伦理问题。现存医疗领域内，关于BCI（脑机接口）的人体试验包含知情同意、利益、风险分析等环节，但目前并没有合适的道德准则和法律来规范未来BCI广泛应用后的种种问题。

思维的移植也会面临安全性、隐私性等方面的质疑，尤其在思维的植入过程中，意识是否能保持其完整性也是要思考的问题。

更可怕的是，为了保持思维和记忆的完整性，进行这样一项意识剥离手术，大脑需要保持着鲜活的状态，而一旦剥离开来，没有了意识的机体究竟还算不算一条真正的生命呢？而在活着的状态下剥离意识，似乎更像一场带有高科技色彩的“安乐死”。

AI能送给视障人士一双“黑色的眼睛”吗？

被传唱一时的《你是我的眼》，其原唱者萧煌奇是一位视障人士，在失明的人生中，他一直保持乐观态度，热爱音乐与创作，一首以自己为原型的《你是我的眼》横空出世，给所有黑暗中的视障人士带来了鼓励与温暖。

你能想象吗？有一天，视障人士可能会深情款款地对AI唱起“你是我的眼/带我领略四季的变化/你是我的眼/带我阅读浩瀚的书海……”

当AI要承担起“眼睛”这个角色时，如何才能将这个世界带到视障人士的面前呢？利用高科技来“作眼”，又是否符合视障人士的真正需求？

从“感知”到“看见”：层出不穷的助视产品

不得不说，科技公司在推进前沿技术的同时，一直没有忘记对视障人士的关怀。一大批助力视障人士的产品和技术如雨后春笋般涌现，这些产品中流露的温情不仅体现了企业对视障人士的关怀，也是科技公司最好的品牌广告。

在梳理了各类与视障有关的智能产品后，我们大致将其分为以下3种类型。

1. 曲线救国型

大家还记得海伦·凯勒的故事吗？海伦·凯勒的老师在教她认“water”时，让她伸出一只手去感知水的流动，并在她的另一只手上拼写了这个单词。从这里可以看出，视障人士认知世界的渠道是除了视觉的其他感官感觉，如听觉、嗅觉和触觉。

基于此，一家公司开发了一款专供视障人士使用的盲文智能手表（Dot Watch）。该手表搭载了盲文显示系统，以盲文的形式将各种信息呈现在手表的触摸表盘中。

相机也有了触摸形式。一位美国设计师专门为视障人士设计了一款2C3D相机，这款相机能通过镜头实时地将拍摄的物体转换成三维触感数据，使视障人士通过触摸屏幕表面生成的立体形状来识别面部细节，如读取表情等。

IBM推出了无障碍环境的一项发明——专为视障人士设计的新型导航App NavCog。NavCog可通过耳机与视障人士“耳语”，帮助其实时识别位置、朝向，还能辨认迎面走来的熟人。

2. 外力加持型

外力加持型的载体一般是智能眼睛，主要是为视障人士打造的。

一款名为eSight的产品，结合算法和部分视障人士自身的需要，通过控制器中的液体镜头技术进行“聚焦”，视障人士利用眼镜中的Bioptic倾斜功能，不仅可以调整瞳孔距离（对焦），还可以调整图像的清晰度（颜色、对比度、亮度），从而“重获光明”。

3. 直截了当型

视觉的产生依赖于三大组织器官：眼球（主要为视网膜）、视神经、视皮层。对于视障人士而言，如果想要恢复视觉，就必须拥有能替代这三种组织的假体，即视网膜假体、视神经假体和视皮层假体。

国内就有研究团队研制出了人造视网膜，其由体内电子微系统和体外电子系统两部分组成。使用方法是在视障人士眼球内部植入IC芯片，用来接收信息和传导电信号，然后为视障人士配备一个体外接收系统，如眼镜。

“眼前的黑不是黑，你说的白是什么白”

虽然现在的智能助视产品比比皆是，但要真的掀开视障人士眼前的“帘子”，恐怕还不容易。

人造视网膜技术具有很强的综合性和复杂性，需要机器视觉、IC设计、半导体工艺、纳米技术、神经科学、生物材料等十多个学科的科学家和工程师全力投入，密切配合。但即便拥有如此高精尖的团队，外界信息通过电信号传递到大脑中，视障人士感受到的也不过是一个灰色的、带有马赛克的世界。

即使是黑白的“渣像素”，也能勉强算“看见”了。而那些智能产品，如智能眼镜、认知助手等，只能提供语音让视障人士接收到相应信息。视障人士能做的就是这些产品说“前面有障碍物，请绕过去”时无奈绕开，而不能亲眼看看阻拦自己的障碍物究竟是一块石头还是一辆单车。

当我们不断加大视障人士在其他感官上接受的信息量时，这也会带来不小的隐患。例如，从听觉入手的产品往往会让使用者戴上耳机，这就会让视障人士与周围的声音隔绝，出行在外，容易造成危险，而不戴上耳机进行电子播报，容易造成视障人士的信息外泄。

当我们将智能产品应用在视障人士身上时，除了要为他们带来生活上的便利，更要让他们看到这个美丽的世界，无障碍地探索这个世界。当我们用一种以其他感官来辅助视觉的技术思路来实现客观上的无障碍时，肯定会与视障人士主观上的无障碍有区别。

智能眼镜戴久了，人们或多或少都会产生不适感。想想我们在电影院看3D电影的时候，3D眼镜也曾让我们头晕目眩，沉浸式技术也极容易带来头晕、恶心等反应。除此之外，戴智能眼镜限制了侧面周边的视觉范围，视障人士要做到和正常人一样的移动和工作还并不容易。

要真正为视障人士带来便利，必须解决这些智能助视产品成本太高、价格太贵的问题。由于致盲因素不同，很多视障人士需要高度个性化定制的智能产品，这更加导致成本居高不下。例如，加拿大一家医疗科技公司Ocumentics在2018年开发的仿生镜片只适合25岁以上的成年人。

看得见，要以什么为标准？

归根结底，不管智能产品有多炫目，对视障人士而言，他们

更在乎智能产品的实用性，而真正实用的产品于他们而言就是3个字——看得见。

仿生眼球当然具有很大的市场，但它并不适用于所有眼科疾病。智能眼镜+芯片的组合是可以通用的，因为其视觉计算能力、人脸识别等功能可以使其接收外界信息，芯片通过柔性电极阵列来传输电信号，刺激视网膜的神经细胞，进而传递到大脑中，让视障人士看见黑白的影像。

但对于市场来说，“通用的”就不具有特殊性了，所以，谁能快速抓住“通用”中的亮点及最不易解决的问题，谁就能在助视方面成为佼佼者。

神经科学家认为“人脸识别”有两个方面。其一是特征识别，也是目前的智能眼镜配备的识别类型，其二是表情识别。经过亿万年的进化，人类形成了7种与情绪密切相关的基本表情，分别是快乐、惊奇、悲伤、愤怒、厌恶、轻视和恐惧。这些基本表情是人的本能，是不需要学习的。

目前，我们还不太清楚脸部特征与表情之间的区别。当看到一个人时，我们大脑里的人脸识别机制就会开始运作，我们会在一瞬间就判断出这张脸是不是熟面孔，以及这个人的表情如何。但这对于视障人士来说却难如登天。

所以，让AI来帮助视障人士看到人的表情或许是智能助视产品在人脸识别上的真正战场。这不仅需要AI去识别更细微的脸部特征，而且还需要AI为视障人士获得更加清晰的图案，而不是一

个模糊的影像。

英国《柳叶刀·全球卫生》的一份研究报告预计，当下全球视障人士数量为3600万，如果不加强对眼疾的治疗，到2050年视障人士数量将增至1.15亿。这是一个足够惊人的数据，而借助技术的力量，我们希望每一个在生活中艰难独行的视障人士，不论他们的年龄大小、环境状况如何、贫穷或富有，他们都能看到我们的美丽世界。

3D“造畜”和3D“造人”哪个更容易实现？

Netflix新剧《副本》中出现了这样一幕——劳伦斯的儿子利用便携式3D生物有机体打印机，先打印出细胞群，然后一点点积累，最后克隆出了自己爸爸的身体，再植入自己的堆栈（类似于记忆芯片），完美地“变身”为有钱的爸爸。

看上去，利用3D生物打印技术打印出一个完整的人体，而且能正常地为人所用是一件非常简单的事情。然而，现实是，别说一点点积累细胞了，刚开始打印出的细胞群都不一定能存活下来。全世界生物打印领域的先驱Organovo公司就曾经打印出人体肝脏，但这个肝脏仅仅存活了40天。

当然，科幻剧中的情节并非不能在现实中上演，打造一个“假人”可能有点难，但打造一个“假动物”会不会比较容易呢？我们又是否真的需要这项功能呢？

3D“造人”，以他山之石攻玉

3D“造人”，关键是要突破2D的认识。目前，人们对微观世界的理解一直停留在二维水平，生物学家也很难从三维立体的角度去理解生命体中细胞、血管和神经之间的相互联系。

除此之外，即使我们观察清楚了三维的微观现象，但是3D生物打印技术依旧存在不足，如生物打印材料如何量产、如何让生物“纸墨”达到“造人”的数量，以及如何让活细胞准确地排列在可打印凝胶中并维持活性。

这不是简单地把菜种进土壤里的问题。在打印时，我们必须准确计算活细胞的位置，也就是说，面对每个活细胞，你要解决的问题是，如何把一颗大白菜的种子准确地种在东南角的一小块地里，并使它得到最理想的光照，长成你最喜欢的样子。到目前为止，建造“完美的菜园”，凭借人类的计算能力还远远不够。

面对这一系列问题，我们或许可以从其他技术那里得到启发。

2018年，《尖端物料》中提到了一款新的螺旋形纳米机器人，其借助微弱旋转的磁场，能够在活细胞中“遨游”。不仅如此，还可以为它人为制定出一种活动方案，研究者就利用这个特性在细胞中做出了“N”和“M”的轨迹。与此同时，这款机器人还不会伤害到细胞。

我们可以把这个纳米机器人当作一辆“小车”，它能够在生命体内按我们想要的轨迹去“行驶”，如果这个“小车”内还能坐上观察员，那就再好不过了。

这时，AI显微镜就派上了用场。在顶级科技盛会Think 2018中，IBM发布了对未来5年内的五大科技预测，其中就有AI显微镜技术。IBM也正在研发这种小型且自主的显微镜，目的是持续检测水资源。

这款显微镜能够放入水体中，就地监视浮游生物、识别不同的物种，并跟踪其在三维空间中的移动。如果能将这款显微镜部署进“小车”中，借助其观察结果分析，我们或许可以从2D中跳出

来，更好地理解细胞之间的联系。

对于“菜园子”问题，我们可以打造一台AI设备，建立一个可行的数学模型，通过模型预测生物纸墨中的细胞会发生的变化，同时实现自动化诱导干细胞。

中科院广州生物医药与健康研究院承担的国家重大科研装备研制项目“全自动干细胞诱导培养设备”就起到了这样的作用，这是首台自动化无人值守、应用深度神经网络的智能化干细胞诱导培养设备。

生命的“器皿”，应该存在吗

终有一日，我们可能会解决3D“造人”的所有技术问题。但是，技术难题可以攻克，有一些问题我们却无法避免。

在电影《逃出绝命镇》中，男主人公因为健壮的体魄被选作富人的生命载体，在手术后，富人的意识就会存在于主人公的身体中，而主人公的意识会永远被封闭在意识暗坑中。

这个情节与之前提到的《副本》中的情节有点类似，这两个情节也有着相同的指向——躯壳是无意识的器皿，可以承载一个生命体的意识。

要知道，生命体不仅是物质的，也是能量的和信息的。从理论上来说，用细胞做材料，3D生物打印技术可以把一具人体完整地在细胞的层面复制出来，但是信息无法从细胞层面复制。所以，即使造出了人，也是没有能量和信息的人。

假如3D生物打印技术可以“造人”，再加上意识和记忆可以移植，那么我们的躯壳就会更像衣服，可以随意置换。这当然有一些好处，如残疾人可以通过打印一个新的身体来获得健全的身体，人们还可以利用这个技术打造个性化躯壳……

但隐患也是存在的，复制+打印就可以制作出一个完全相同的副本，如果有人将意识剥离出来，“穿上”这个副本去违法乱纪，那时的科技能检测出来吗？

我们可以打开“脑洞”设想一下，完全换了躯壳的人能被接受吗？例如，在体育领域，一个换上了打印躯壳的人是否能参加比赛？如果可以，我们要如何确保比赛的公正性？当然，我们或许可以筹办一个“非原装人”的运动比赛，但有一个令人啼笑皆非的问题是，这场比赛究竟比的是运动精神还是“装备”的优劣？

更为严重的是，一直代代相传的人类，假如突然有一天连生育能力的限制也没有了，可以随意“造人”，地球的存亡可能就是迫在眉睫的问题了。

用**3D**技术打造一座动物园

3D“造人”技术的壁垒太多，技术引发的伦理问题和社会问题要更加复杂。所以，与其去钻研如何3D“造人”，不如另辟蹊径，用3D“造人”技术上的突破去助力其他领域。

利用3D生物打印技术去打印动物，未尝不是一个好的选择。

来自加拿大的一对生物学家Daniel Mennill博士和她的丈夫正

在积极研究蟾蜍颜色变化的生物学意义。就在几年前，他们用黏土制造了一批金黄色的“蟾蜍”，在交配时节到来时将这些模型混入蟾蜍聚集的地方，并观察它们的行为。

然而，这种做法产生的结果并不理想，因为传统的黏土工艺无法制造出高仿真的模型，同时也很难精确控制人们想要研究的变量，如颜色、体形等。

2018年，这对生物学家用3D生物打印技术打印了一批高仿蟾蜍机器人“混入”了这些蟾蜍交配的大军中。虽然这些“蟾蜍”只是高仿产品，但对于生物学家们的观察还是起到了积极的作用。

以后，如果能通过3D生物打印技术制造出更逼真的动物，人类就可以更加高效和贴近自然地研究野外环境下动物的特定行为，这对生物学的发展大有帮助。

除此之外，3D生物打印的动物自然不能将其放回到大自然中，我们可以为它们建造一个集聚了3D生物的动物园。这些3D生物打印的动物可以替代真正的野生动物，承担起动物园的三大“职责”：科学研究、科普宣传教育及休闲娱乐。

“癌症杀手”真的会是“读心专家”吗？

你见过不用充电不用电池的机器人吗？

还真有！

2018年韩国首尔大学的研究者研发了一款名为“Hygrobot”的微型机器人。这款机器人不仅能够爬行和前后蠕动，还能像蛇一样自如地蜷缩。更神奇的是，它不需要电池，也不需要充电，而是靠吸收空气中的湿气来为自身提供能量。

单纯从外观来看，它确实不像拥有着复杂结构和精妙功能的微型机器人，还长得像被剪碎了的塑料片一样，平淡无奇。但它小小的身躯中却蕴藏着神奇的能力。

从向植物“取经”而来的研发灵感

吸收湿气就能提供能量？这款神奇的微型机器人到底是怎么做到的？

它的设计灵感源于植物。我们都知道，松果的鳞片在潮湿时紧闭，在干燥时打开，这样可以使它内部的种子传播得更远。这款微型机器人的灵感就来源于此。它不是由植物材料制成的，而是模仿其背后的机制。

“Hygrobot”可以迅速膨胀和收缩来纵向响应湿度变化。

这款微型小机器人由纳米纤维制成的两层结构组成：一层吸

收水分，另一层则不吸收。当机器人被放置在潮湿的表面上时，吸湿层膨胀，使机器人弓起，吸湿层变干后机器人回落。这个过程循环往复，便能使机器人移动。它可以通过吸收地面或空气中的水分来改变自身的形状和大小。

这样一台完全依靠吸收空气中的水分来运行的机器人是很有价值的，因为空气中的水分是一种全天然的能量，源源不断，不费成本；而且也没有毒性，不像电池那样有可能会发生爆炸。这也使这款微型机器人未来或许能用于人体内的药物运输，为人类治病。

“Hygrobot”——治疗疾病方面的冉冉新星

基于这款微型机器人在能源提供、运动机制等方面的特殊性，科学家们看到了它在未来应用于医疗行业的巨大可能性。

由于自身体型小、结构简单，所以它很容易根据不同的场景需要被制成不同的形状。为了检验它的潜力，研究人员将用抗生素浸泡过的“Hygrobot”放到覆满细菌的培养皿上，它能自行穿过培养皿，留下一条杀过菌的小道，像毛毛虫蠕动留下的黏液痕迹一样。

按照相同的原理，将来它也可以仅靠使用皮肤中的或者人体内部的水分来推动自己，将药物输送到人体各个部位。此外，由于它对空气中的湿气非常敏感，除了在人体内运送药物，还可以为其配备传感器，根据其对其他气体做出的反应检测人体生理环境，从而诊断疾病。

由于“Hygrobot”体型极小，可以进入人体内部，所以，在未来对于很多外部仪器无法做出准确诊断或者不方便治疗的疾病——如脑部、妇科和内脏方面的疾病，我们期待它能够大显身手。而且它能够在指定区域精准用药，所以在疾病治疗的疗效提高和成本降低方面，它或将掀起一场革命。

“癌症杀手”，绝非浪得虚名

对于很多人来说，癌症仍然是困扰他们的“不治之症”。肿瘤是人体细胞无限制增殖及病态地无限汲取养分，导致身体某一个部位的细胞生态变化的一种疾病，恶性的肿瘤便是癌症。目前医疗界对癌症还没有非常有效的治疗方案，癌症诊断书就像是死神的“名帖”，一旦被它的阴影笼罩，就等于半只脚踏入了“鬼门关”。

但是现在这个能像毛毛虫一样蠕动的机器人，或许能为广大癌症患者带来福音。

这款微型机器人使用的是纳米纤维材料，非常有韧性且不易被破损，这为它在复杂的人体内穿行并且将所运载的药物完好无损地送达指定治疗区域提供了保证。

遏制癌细胞增殖的药物早就研发出来了，但是传统的药物常常“敌我不分”，对于正常细胞和肿瘤细胞都有很强的杀伤作用。在使用常规用药手段时，药物还没到达肿瘤处，就在血液循环过程中被身体吸收了一大半；而在药物到达肿瘤处之后还能被吸收多少，也不能够确定。因此，在目前的医疗手段下，患者并不能

大量使用这些药物，否则，不仅病情没有得到多少缓解，反而身体却因为药物的副作用而垮下来。

这款微型机器人可以准确无误地将药物输送到肿瘤处发挥作用，而不会在半途中对身体的健康细胞造成损伤。除此之外，它不用电池也不用充电，这就避免了含有硅或重金属的可植入装置相关的毒性风险，也能避免对人体造成“二次伤害”。

练兵还需几何，方到用兵之时？

说到用于治病的可植入装置，值得一提的还有水凝胶机器人。其实，“Hygrobot”并非首个被应用于生物医学领域的微型机器人。2018年，来自哥伦比亚大学的团队在《科学机器人》

（*Science Robotics*）杂志上发表了关于自行研发的水凝胶机器人这一成果，它由磁铁激活，通过局部剂量性释放化学药物来治疗肿瘤，达到更高精度的靶向化疗效果。

顾名思义，水凝胶机器人采用的制作材料是与生物相容并且能进行生物降解的基于聚乙二醇（PEG）的水凝胶，对于人体和环境来说都更为友好。

至少在可植入装置的材料上，可降解的纳米纤维和水凝胶就预示着人类前进了一大步。

此外，基于DNA折纸术来组装和制备DNA纳米结构并实现相应功能的技术也日渐成熟，就像制作微型复杂芯片的摄影构图技术。材料能够通过埋入不同的生物化学模式，来响应DNA的特定指令（如弯曲、折叠等），达到改变形状的目的。这一技术大大

提高了微型机器人的适用层面，使它能够在不同环境下改变自身的形状，从而更好地完成工作。（免费书享分更多搜索@雅书.）

水凝胶机器人需要外部的磁铁激活和引导，这一点也提高了装置的操控难度和对磁铁的要求，并不便于身体内部复杂器官的治疗。相比之下，“Hygrobot”没有这个困扰，但是要使它精准地移动到指定部位并释放药物还需要后续的研究优化。如果能够在这一方面取得突破，那么我们相信它在药物运输、组织修复甚至破坏癌细胞植入等方面的潜力将令人瞩目。

所以，很有可能在不久的某一天，当我们生病时，只需要吞下装着微型机器人的胶囊。它们在进入我们的身体后，能够承载着治疗疾病的药物，并且根据我们身体的不同构造变换自身的形状，避开所有的障碍，精准地到达“目的地”，释放药物或者消灭病原体，完成许多在现在看来“不可能完成的任务”。

“癌症杀手”或许更是“情绪大师”

不过，这款机器人带来的想象力还不止于此，因为它还有一个“身份”——情绪大师。

如果有一天，你能读懂卧室窗台上你养的那盆多肉在想什么，或者了解你家的猫对你是否满意，甚至你能够监控你侄子的抑郁症是否病情稳定，这些听起来是不是有点像天方夜谭？或许微型机器人的发展，真的能让这一切变成现实。

从生理角度来说，抑郁症的形成是因为神经递质中多巴胺和

5-HT等物质浓度降低而导致大脑额叶和边缘系统的病变。在临床中通常通过增加单胺类递质的含量来治疗抑郁症。

这种对湿度敏感的微型机器人可以像“家庭护理员”一样待在抑郁症患者的身体里，每当患者体内多巴胺和5-HT浓度下降时，能够及时给出反馈，并且移动到相应区域释放单胺类递质来缓解抑郁症症状。

人、动物和植物的情绪在细胞生理机制层面有共通之处。人和动物主要是依靠大脑和小脑的机制、神经元和激素的分泌来影响情绪的，我们目前是通过测量记录植物体内生物电流的变化来观测植物情绪的。微型机器人的投入可能将动植物情绪的研究直接提升到分子层面，将微型机器人放入其体内，由于它本身不含有金属，也不需要充电，不会对生物磁场造成影响，所以，可以将来自设备的干扰降到最低，大大提升得到的研究数据的可信度。

不管是微型机器人直接反馈动植物的情绪反应，还是提供相关数据供研究学习，都能够大大推进这一领域的研究发展。想象一下在未来，我们通过在宠物体内或者心爱的盆栽体内植入微型机器人，就能感受到它们是怎么与你无言地“撒娇”“委屈”或者“吐槽”了。动物与人类的交流或许算不上什么，但是植物与人类的交流能够真正称得上是跨越物种的交流。

或许，这是微型机器人在人与大自然万物和谐相处方面最值得铭记的意义——让一切生物的情绪能够被感知、被理解，用后科技时代的产物，完成人类最原始的梦想。

AI能把抑郁症治好吗？

近些年来，“抑郁症”一词越来越多地被人们提起，不少名人都曾表示陷入过抑郁症的痛苦，而抑郁症患者不堪病痛而自杀的新闻也屡见不鲜。

生命的“陨落”，无疑给人们敲响了警钟。抑郁症高发，我们要如何去防治？如果我们能检测出潜在的抑郁症患者，及时的心理疏导，是否能够有效地减少这类悲剧的发生呢？

IBM的计算精神病学和神经成像研究小组团队曾经利用机器学习来预测人患精神疾病的风险，他们通过对59名普通人的语言方式进行追踪、分析，并对语言连贯性进行评分，来预测人们的潜在患病风险。在随后的结果验证中，AI预测的精确度达到83%。

人们在AI预测抑郁症方面取得了突破，可见有效利用AI诊断和治疗精神疾病已经是蓄势待发。而在这之后，AI在防治精神疾病方面会有多大的“用武之地”，其推进的难点又在哪里呢？

我们为什么需要AI预测抑郁症？

1. 人类预测抑郁症的识别率低

尽管抑郁症已经成为全球第四大疾病，预计到2020年将成为第二大疾病，但诊前的预测和诊后的监控都是薄弱环节，因为心理医生和精神病医生难以做到诊前精准预测和诊后有效追踪。

在中国，有大约2.5亿人需要心理咨询服务，8000万人需要心理治疗，心理诊疗的市场需求极大。与此形成鲜明对比的是，我国抑郁症的医疗和防治还处在低识别率阶段，地级市以上的医院对其识别率不足20%，只有不到10%的患者接受了相关的药物治疗。

而反观AI的表现，除了此次IBM的AI预测实现了高识别率，在2018年，美国哈佛大学曾通过用AI程序分析社交网站中的照片，提出用色彩学的方法来诊断抑郁症，正确率也高达70%，这无疑证明了AI预测将很有可能成为当今社会用于早期筛查和检测精神疾病的新途径。

2. 抑郁症难以完全治愈，让预防变得更为重要

抑郁症是可以治疗的，但却很难治愈。一个患了抑郁症的人，即使接受了心理治疗，恢复到了以往的精神状态，还是有极大的可能复发的。

治愈之难，使得预防变得尤为重要。这些年屡屡可见的抑郁症患者自杀的新闻，也在提醒着人们要重视对心理疾病的预防。AI预测的高识别率也将使其成为越来越被人们看重的“法宝”。

3. AI能听出人类的“言外之意”

IBM曾提出，有了AI，人类的语言文字就会成为通向精神健康的一扇窗。人们的语言和文字所形成的规律会被AI的认知系统分析，人们的“言外之意”就会成为精神健康和身体健康状况的可测指标，这种经分析得出的数据能够帮助医生更有效地预测并追

踪早期的精神疾病等。

南加州大学推出了一款AI心理治疗师，它会分析受访士兵的面部表情变化，以及士兵的语义和语音，再结合问卷调查，诊断其是否存在PTSD（创伤后应激障碍）症状。

AI通过观察人们日常所忽视的语言习惯来预测人的精神疾病，在日后，该技术或许可以更进一步，通过追踪、分析人类的微表情、微行为，甚至是细微的语调变化，从而更加了解人类的心理状态。可以肯定的是，这种技术发展成熟后，在刑侦领域也会有极大的应用空间。

实现AI预测抑郁症还要跨过哪些难关？

尽管AI预测抑郁症有着科学的技术和强大的数据库作为支撑，但要全面推进还存在一定的难度。目前的难点主要表现在受众接受度、机器学习和隐私保护3个方面。

首先，人类和机器在心理检测中始终存在疏离感，人类对于AI的心理检测结果接受度不高。心理检测与生理检测不同，在生理上，我们可以做各项生化检查，通过明确的数值和图像来判断生理上的疾病，人们也会更倾向于相信机器的精准度。但心理检测却有一些不易量化的指标，如焦虑、冲动、恐惧等情绪。患者在与医生交谈时，更倾向于相信心理医生能够感知自己的情绪并且产生“共情”，从而更好地诊断自己的病情，而没有自身情绪的AI，恐怕难以得到人们的信任。

其次，抑郁症的病因、病情十分复杂，机器学习难以全盘掌

握。人体的神经网络精密且复杂，而随着人体生长，人脑的神经网络也会不断生长变化。迄今，有关专家对于抑郁症的病因都难以解释清楚。但可以肯定的是，心理与社会环境诸多方面因素与抑郁症的发作有关。

抑郁症的表现也十分复杂，例如，被网友热议的“微笑抑郁症”，“微笑抑郁症”患者在白天大多数时间都面带微笑，但“习惯性微笑表情”并不能消除工作、生活等各方面带来的压力、烦恼、忧愁，只会让他们把忧郁和痛苦越积越深。

“微笑型抑郁症”多发生在那些学识高、事业有成的成功人士中，他们或是机关里的高官、企业中的老板，或是高级技术人员，这类人给人的印象通常是十分健谈、自信沉稳的。但是，如前文中提到的，如今AI预测抑郁症主要是通过分析被检测者的语言方式和语言连贯性进而确认其患病风险，而面对此类“成功人士式患者”的侃侃而谈，AI又能否发现他们微笑背后的抑郁呢？

最后，AI预测抑郁症的方法，即对于语言的分析不能适用于所有的语种。目前，硅谷一家名为X2AI的初创公司推出了针对叙利亚难民的AI心理咨询师，虽然其已经在土耳其的难民营中开始试用，但其功能相对简单，且只能做阿拉伯语的语义识别。

就以大家最熟悉的中文为例，当有人说“吃饭去了我”时，即使“我”作为主语并没有放在谓语前，依旧不会影响大家对这句话的理解。在中文中，语序颠倒是十分常见的，即便如此，也不会影响句子的原意。那能代表这个人有心理疾病吗？当然不能。毕竟“我去吃饭了”和“吃饭去了我”这两句话在口语中的使用频率恐

怕不相上下。

另外，不同病症是否会有不同的语言倾向呢？而在同一病症中，不同程度是否也会有不同的语言习惯呢？这一切还有待研究。

除此之外，AI预测抑郁症同样也会存在AI进军医疗领域的常见问题，如检测程序的设定、医疗数据保护等。所以，AI预测抑郁症要实现有效落地，恐怕还有很长的路要走。

被避讳的妇科，AI能让她摆脱尴尬吗？

女性健康已经日益受到大众尤其是女性群体的关注，自然，在女性健康意识越来越强的背景下，AI+妇科也被提上了日程。

从“妈妈”到“大姨妈”：涌现的智能妇科产品

近年来，针对女性医疗保健的智能产品越来越多。我们梳理市面上的智能妇科产品，大致可以将其分为以下3种类型：

1. 生育型

因为生理的特殊性，女性承担起了生育的职能，然而，生育的整个过程包括备孕、怀孕、生产和产后，都具有一定的患病风险，这也使女性健康面临较大的威胁。基于此，许多创业公司都围绕女性生育开发了一系列产品，如家庭生育监测设备，可以监测女性的生育能力。

Sera Prognostics公司就是在女性孕期提供服务的，其智能诊断测试设备能够判断出妊娠并发症的风险。还有Kleiner Perkins公司支持的Progyny，它在2017年就已经筹集了4900万美元，提供卵子冷冻、胚胎筛选和生育治疗等一系列智能服务。

2. 日常助手型

女性的经期一直是判断女性是否健康的标准，而AI也将女性经期列为重要研究对象。在这个方面，市面上也有相应的智能硬

件产品，能够结合物联网、大数据技术，实现经期数据自动获取，精准预测女性生理周期，并为女性提供经期管家式服务。

还有分析经血以评测女性身体健康的产品，如智能卫生巾，用户在使用过后可以取出其内置芯片，将芯片置于评测机器内，这个系统就能自动获取经血的颜色、气味等平时物理条件下不易采集或无法准确描述的数据，从而分析女性的身体状况。

除此之外，女性生殖器官的护理保健也是AI关注的重点，深圳就有一家公司研发出了一款妇科智能冲洗器，满足女性在不同场景下，外阴护理和内阴护理的需求。

女性医疗，我们究竟需要什么样的产品？

虽然市面上的智能妇科产品不少，但真正使用这些产品的女性却不多。因为比起那些想要购买智能妇科产品，健康意识强的女青年（青年是主要消费群体），在现实生活中，更多的还是那些不太关心自己生殖健康和不了解妇科疾病的女性。有调查表明，80%的女性在患上妇科疾病时并不自知。

出现这些现象的原因，主要是女性对“性”的避讳，这使得她们很少关注系统性的妇科健康知识。而系统知识的缺乏，又导致她们在妇科检查、诊断、治疗时，难以避免地产生一种心理排斥感。

所以，五花八门的智能妇科产品虽然看起来很炫酷，但对女性的实用价值并不高。广大女同胞们最需要的应该是AI对她们敏感意识的共情，在这种共情的基础之上，再给予她们医疗救治。

AI能通过观察人们日常所忽视的语言习惯，追踪、分析人类的微表情、微行为，听出、看出人类的“言外之意”，从而与人类达到某种程度上的共情。共情AI已经有了，如何去落到实处，消除女性的心理排斥感呢？

我们可以从两个角度来分析。一是女性医学的公共教育和家庭教育，也就是将AI应用于女性教育领域，AI+教育的组合已经相对成熟，这里就不加描述了；二是妇科的远程医疗，远程医疗已经是老生常谈的话题了，但对于妇科，却又有着不同的意义。

很多女性患者在患病后，首先想的不是去医院就诊，而是在网络上询问，但得到的建议往往是没有价值的。从中我们可以发现，比起面对面交流，这种匿名式的远程问诊更容易被广大女性接受。妇科可能是远程医疗最好的铺设渠道。

女性对隐私越来越看重，拥有只属于自己的、了解自身健康状况、能长期提供治疗指导的家庭医生服务，就显得越来越有必要。AI对海量数据的处理能力，也能够有效满足健康监测的需求。

当女性不好意思对医生说出自己身体上的不适时，AI可以把握病情主诉上的“尺度”，提前将可能出现的症状列出表格，让患者进行勾选，避免“隐私说给外人听”的尴尬。

除此之外，AI还可以通过设置一些生理指标监测的预警标准，让女性免于开口，在检测过程中，一旦这些生理指标出现异常，就会从应用端发出预警信息，有利于人们及时处理病情。有

一家名为Maven的初创公司就以女性按需远程医疗为主营业务，还打造了专门为女性设计的专家网络，仅用3年时间，这家公司就成功融资900万美元。

AI+妇科，究竟代表了什么？

不得不说，AI+女性医疗的市场潜力不容小觑。从九价宫颈癌疫苗的火爆，到人造子宫、人造胚胎的出现，女性在生理健康上有了更强烈的自主意识。

AI的加入，表明社会正在为女性谋求更好的医疗环境，使她们对妇科疾病有更深入的了解，女性讳疾忌医的现象也会渐渐减少。而且，比起男性，精致的女性显然是AI+医疗的有力消费群体。

2017年，京东数据研究院公布了《女性消费报告——2017京东女子图鉴》，该报告表明，25~50岁的中产女性是品质消费的中坚力量，她们更乐意为品质生活埋单，个性化的科技消费已经成为她们的追求。

除了女性的消费升级，AI+妇科或者说是智能医疗还有一个主要推动因素——女性的家庭角色。在家庭中，女性扮演了母亲、女儿、妻子等多种角色，这些角色也使得女性在医疗消费活动中起着特殊的作用，她们不仅对自己所需的医疗条件进行决策，还可以影响家庭中其他成员的决策。

因为女性通常具有较强的表达能力、感染能力和传播能力，善于通过说服、劝告、传话等对周围其他人产生影响，女性患者

容易将自己就医时的满意的感受和接受的满意服务经历当作自己的谈资，从而扩大智能医疗的传播范围。

太空生子什么时候才能实现？

埃隆·马斯克每次开新公司总是从图的最右边开始确定目标，然后一直推导到左边，如图1-2所示。

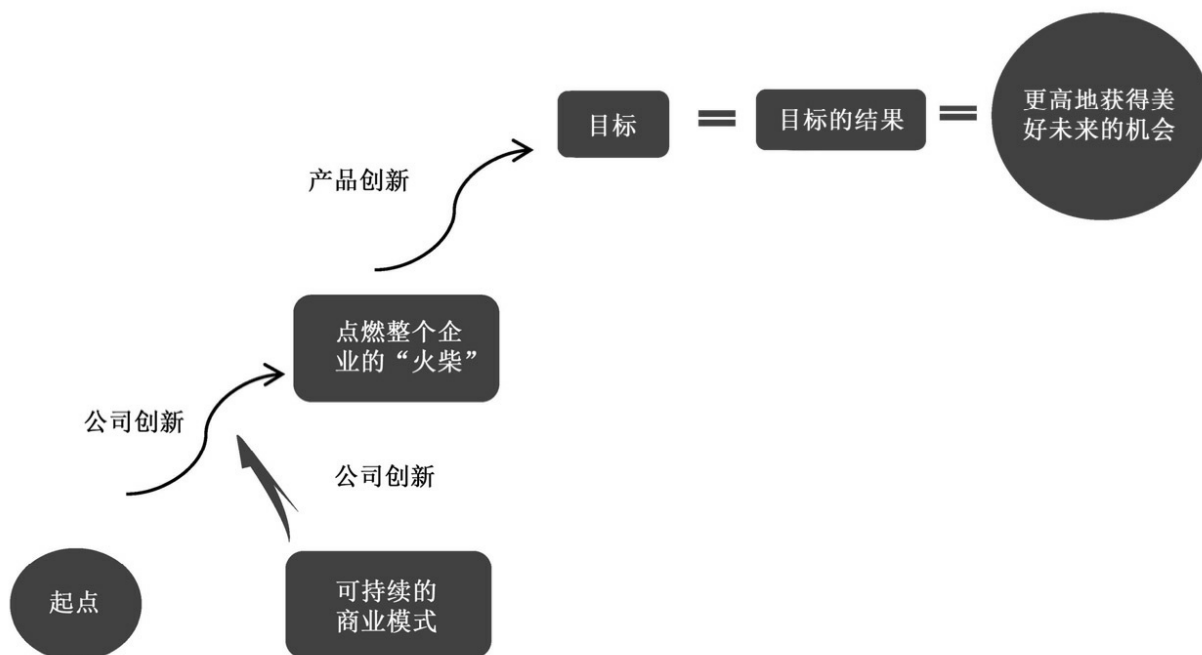


图1-2

通常在一个行业变得热门之前，并不是缺乏让其“燃烧”起来的“木头”，而是缺乏可以“燎原”的“火星”——总有一些技术上的短板阻止一个行业的快速起飞。所以，马斯克建造一家公司的核心初期战略都是创造能“点燃”整个行业的“火星”。

而说到“火星”，不得不让人想起马斯克的“殖民火星大计划”，当然，此“火星”非彼“火星”。殖民“火星”一直就是他的目标，这一点从他2002年成立SpaceX公司起就未改变过。

那么殖民火星的一大重要因素是什么呢？繁衍后代。

2018年4月初搭乘SpaceX公司的猎鹰9号火箭的龙飞船升空，飞船上除了有一批补给品，还包括12个男人的冷冻精液。这是第一次人类精液在太空进行实验。此前，包括青蛙、蝾螈、海胆、海蜇、蜗牛、青鳉鱼等若干动物，已经成功在太空中繁殖。

但对人类来说，想在太空繁衍生命就没那么容易了。

太空繁衍有多难？缺乏重力和宇宙辐射使其步履维艰

根据牛顿第三定律，力的作用是相互的。你在太空中接触一个物体，就会给它一个力，改变它的速度和方向。所以，在太空中，两个人即使只是相互拍下肩膀，双方也会被弹向两个方向，直至碰到其他物体。

在这里重点解释一下太空辐射。太空辐射是一种包含伽马射线、高能质子和宇宙射线的特殊混合体。在电影《绿巨人》中，由于伽马射线的辐射使班纳博士的肌肉组织、结缔组织、神经组织发生了异化，一旦愤怒，他就会化身为绿巨人浩克。

在现实生活中，可能不会发生这种事。如果将太阳作为参照物，那么伽马射线在几分钟内释放的能量相当于万亿年太阳光的总和，它强到可以把空气点燃，烧焦地球上的所有生物只需要几秒，可见宇宙环境有多恶劣。

人们在宇宙受到的辐射量超过在地球上的100倍，作为身体上最敏感的部位，生殖器官首当其冲。辐射会让精子发生突变。

一项研究表明，来自太空辐射的带电氧离子和铁离子会对母鼠卵泡储备造成伤害，而雌性卵巢内的卵泡又是有限的，长时间暴露在辐射下会对大脑造成永久伤害。

那么怎么办呢？

见微知著，从更小的单位研究起

1. 精液

2013年8月，日本的研究人员利用日本宇宙航空研究开发机构（JAXA）的货运航天器，将冷冻干燥精子样本送上国际空间站。在待了288天之后，样本随着火箭回到地球。研究者用这些精子对实验室的小鼠产生的卵子进行体外受精和胚胎移植。经历19天的孕育后，研究人员对小鼠实施剖腹手术，使其生下健康的小鼠。

其实，小鼠在某种程度上存在DNA的损伤，但这些DNA损伤并没有影响到后代，也没有通过性状表达出来，甚至当第二代小鼠在生育第三代时也表现正常。

这是人类首次在哺乳动物中完成类似的实验，而在对人的精液进行研究时，国际空间站上的科学家解冻并激活精液样本，然后观测精子活力与精卵结合的情况。这里需要理解精子获能的概念，即精子获得穿透卵子透明带的过程，只有经历了这个过程，卵子才能完成受精。现阶段需要搞清楚两个问题：微重力对精子活力的影响及在微重力下能否完成精子获能。

早前有研究证明，缺乏重力会干扰精子的正常功能。然而，即便精子细胞本身功能正常，它们能否在太空中和卵子结合依然成谜，现阶段对精子活动的探索依然处于极早期阶段。

2. 胚胎干细胞

2017年4月，天舟一号货运飞船在太空开展人的胚胎干细胞分化为生殖细胞的实验，该项目负责人为清华大学医学院教授纪家葵。该团队此前已研究证明，人的胚胎干细胞可分化培育出原始生殖细胞及类精子细胞。

该团队开发了一套荧光报告系统，将2008年获得诺贝尔奖的绿色荧光蛋白GFP嵌入到了一个名为VASA的基因之中——当细胞表达VASA基因的时候，荧光蛋白也同时会被激活表达，细胞就会发出绿色的荧光，以此来判断人类胚胎干细胞是否分化出原始生殖细胞，并实时传输回地面，与地球上同时进行的对照实验比对。

该研究重点关注太空微重力环境下生殖细胞发育与成熟的基本规律，探索微重力、高辐射的环境对于人的生殖细胞有哪些影响，生殖细胞的分化是否会因此而滞后或者效率降低。（免费书享分更多搜索@雅书.）

2015年NASA的研究结果表明，太空微重力环境影响了小鼠拟胚体（EB）在太空的分化能力，抑制了谱系分化基因的表达，但这些未分化的EB在地面条件下培养却能够进一步分化。

这一系列的研究为后续的研究提供理论依据和技术支持。然

而，不管是研究精液还是胚胎干细胞，终究是很早期的研究，最后即使能受孕，接下来该如何发育依然是需要探索的问题。

受精卵可以在太空发育吗？

初中生物教科书中曾讲到过太空育种，简单说就是让很多植物种子搭载火箭进入太空，在太空高能粒子辐射、微重力、高真空等条件下让植物种子发生遗传变异，然后经过挑选留下人类想要的品种筛选育种。但实际情况是，将一大批种子带到太空，只有少部分种子能结出又红又大的果子，绝大部分的变异都不是人类所需要的。

人的对称生长与地球上的重力有很大的关系，在强大的辐射环境中能孕育出什么样的宝宝，这个问题恐怕连绿巨人本人都难以回答。植物变异畸形的种子会即刻销毁，但人的受精卵在伦理上并不被允许这样做，那应该怎么办呢？显然这涉及很多问题。

AI真的能预测死亡时间吗？

对于一些重病患者而言，“死亡”这个话题虽然十分沉重，但也需要预留空间去探讨和适应。如果可以精准预测患者的“死亡”，是否能给予他们这个预留空间呢？

美国斯坦福大学开发了一个预测死亡时间的AI系统。这个AI系统整理了近200万名成人和儿童患者的电子健康档案数据，以及相关的医学诊断信息，从而得到病情的大数据。再通过数据收集与系统自主学习机制，这个AI系统可以预测患者具体的死亡时间。

为什么要用**AI**来预测死亡？

在中国，每年有约700万人走向生命终点，这个AI系统的出现，预示着医生能够更加精确地安排患者的临终关怀。除此之外，利用它我们还可以发掘一条新的道路。

对于大限将至的晚期患者，这个AI系统可以通过长期的数据跟踪来判断死亡概率。而对于一些特殊疾病的突发症状，它也可以通过机器学习，感知患者的一些生命体征的变化从而发出危险预警。

FDA于2018年批准通过了一个可以预测死亡的AI产品，这个产品名叫Wave Clinical Platform，由医疗科技公司Excel Medical研发。

Wave Clinical Platform是一个永远在线的远程监测平台，集合了患者的用药情况、年龄、生理情况、既往病史、家族病史等实时数据。

这个系统可以感知生命体征的细微变化，从而在发生致命事件6小时前发出预警。也就是说，这个系统可以通过比较数据库中的猝死病例，从而提前6个小时预测“猝死时间”，为医护人员赢得抢救时间。

英国科学家也曾在《影像诊断学》杂志上发表文章说，AI可以预测心脏病人何时死亡。AI能让医护人员发现那些需要干预治疗的患者，从而让医护人员能够拯救更多的生命。

AI预测死亡这一命题能成立吗？

对于AI预测死亡这一命题所遇到的问题，我们可以从以下3个方面来考虑。

1. “预测死亡”即“判死刑”，患者能接受吗？

不可否认的是，预测死亡确实可以让医生更合理地配置医疗资源。但“死亡”并非那么容易被所有人接受。

《众病之王：癌症传》的作者Siddhartha Mukherjee博士在文章中讲过自己亲历的一个故事，他曾经治疗过一名食道癌患者，这个患者的治疗十分顺利，但还是存在很大的复发可能性。于是医生提出了临终关怀，但这位患者拒绝了。这位患者认为，他的身体状况越来越好，精神状态也很好，为什么医生偏要说这些扫

兴的话呢？

令人遗憾的是，这位患者的癌症还是复发了。在他临终前，他始终处于昏迷状态，无法回应在他病床旁的家人。

从这个故事中可见，并非每一个患者都能淡然地接受“死亡”。当患者与病魔、死神苦苦争斗时，医生用一套所谓科学的、精密的AI系统预测了患者的死期，于患者而言，抗癌之旅本就艰辛，而在其头顶悬上一把会准时掉落的“死亡之刃”未免也太残忍。

2. 病情存在个体差异，复杂病例难以判断

AI预测死亡主要依赖于医疗大数据和深度学习。研究团队表示，AI预测死亡系统收集了发现病症后12个月内死亡的患者数据，然后通过神经网络利用大数据计算每条信息的权重和强度，生成一个患者在3~12个月内死亡的概率分数，再通过分数预测患者在3~12个月内是否会死亡。

医疗数据种类繁多，质量参差不齐，是一种极具个性化的信息。疾病的病程具有一定的规律，但具体病情症状却因人而异。个人体质、周围环境等因素都会影响疾病的转归。除了个体的差异，疾病本身也难以被清楚地认知。例如，几乎任何传染病的初期症状都与感冒类似。也就是说，疾病本身是带有欺骗性的。在面对复杂的病例时，医生也常常需要借助辅助工具，或召开病情讨论会议，几方会谈后才能确定治疗方案。

另外，AI预测死亡的深度学习有一个令人费解的地方，也就

是“黑盒子”问题——它能够推算出一个患者的死亡概率分数，却无法表达其背后的逻辑。

所以，通过概率分数来预测患者的死亡时间依旧存在许多问题。单单针对某类疾病的死亡预测可能有效，但是预测大病种的死亡概率的可能性却微乎其微。

3. 医疗大数据共享难

AI+医疗大多以算法开始，但最终还是会回到数据。数据获取难度大是所有AI项目的问题，医疗行业的数据，尤其是这类关于生死的数据更难获取。

医疗信息与其他领域的信息不同，种类十分繁杂，标准也不统一。尤其许多医疗数据会涉及患者的隐私，有部分患者并不愿意将自己的医疗数据用于AI研究。

就质量而言，医疗数据也有更高的要求，比如，所有的医疗数据都需要医生的人工标识。

除了患者方面的原因，从医院方获取数据也有阻力。在不能确定某项研究是有利于医院救护的时候，医院恐怕并不愿意担风险而贡献出所有的工作数据。技术人员如何和医生形成合力，获取高质量的大数据，是大部分AI医疗企业共同面临的难题。

“鸡肋”如何巧变为“熊掌”？

“AI预测死亡的准确率高达90%”更像是一个噱头，预测人类

的死亡只是更方便地进行姑息治疗，但其中还是会面临一些伦理问题。例如，要不要将死亡日期告知患者和其家人？机器是否有资格来宣判人类的死亡期限呢？

如果换一个预测对象呢？设想一下，作为一只宠物狗的主人，当狗狗的身体机能渐渐衰退，主人是否想知道这只狗什么时候会离世呢？由于语言不通，人类希望借助一些辅助工具来了解宠物，希望有更精确的医疗辅助系统来诊断宠物的病情，从而为宠物做更好的安排。面对宠物，AI预测死亡似乎更能被人类接受。

AI预测死亡系统的发展过程应该是一个不断提升价值的过程。一方面，这个系统应该建立更多对象的数据库，依赖深度学习来进行更多应用场景的选择。首先，选择一类对象（多半为宠物）作为训练学习模型的教材；然后，通过积累的“经验”来判断这类对象在发病期间的死亡概率；最后，对对象进行干预治疗。

另一方面，预测应由预测死亡变成预测病程。预测场景从垂直领域到横向领域，构建一个智能预测系统，既包括病程的转归期，也包括病程前期的所有阶段，最后做到为用户个性化建模。

在AI医疗上，我们细分了越来越多的名目。虽然“预测死亡”看起来涉及人类生死大事，但目前人们只是触及事情的表面。在戳破了“死亡预测”这个气泡后，如何让AI医疗预测成为一个真正惠民的项目，触及医疗痛点，恐怕才是大部分布局AI医疗的企业需要思考的。

罕见病不再罕见，AI真的能让患者享受生命尊严吗？

说起冰桶挑战，想必大家都很熟悉，这项挑战自2014年发起，旨在让更多人了解被称为渐冻症的罕见疾病，同时也达到了募款帮助患者治疗的目的。渐冻症因为名人的加入而有了关注度，但其他的罕见病却没那么“幸运”，在罕见病的目录中，更多的是那些不被人关注甚至被人歧视的罕见病症。

虽然不被重视，但罕见病群体却一直在扩大。在中国，渐冻症患者每年确诊逾400例，根据WHO的数据，目前我国罕见病总患病人数约为1680万人。事实证明，罕见病已经不再罕见，一直致力于医疗领域的AI是时候重视这个特殊的患病群体了。

AI能为罕见病做些什么？

想要治病，当然要对症下药。罕见病的困境，首先体现在确诊上。

由于不常见，罕见病误诊的概率非常高。据《2018年中国罕见病调研报告》，43%的患者需要在4家及以上医院求诊才能确诊，有16个患者甚至在超过20家医院求诊后才获得确诊。欧洲罕见病组织的调查也显示，40%的受调查者曾被误诊。

患者除了要承担误诊的风险，还要在长期的确诊等待中饱受生理和心理上的煎熬，面对这种情况，我们如何又快又准确地诊

断出罕见病呢？

简单来说，就是要“化简为繁”。为什么这么说呢？因为罕见病的患病人群少，能够提供给AI建立模型的有效数据也非常少，有没有一种方法能将数据增加呢？当然是有的，总结起来就是——“人数不够，细胞来凑”。

如同罕见病患者会表现出患病症状一样，罕见病的细胞模型也会在细胞结构上表现出一定的特征，而这些特征就成为机器学习的素材。AI可以“观察”非常多的细胞显微照片，并对照片上的细胞结构特征进行分析，从中找出与疾病相关的细胞特征。

位于美国犹他州盐湖城的Recursion公司就利用机器学习、计算机视觉和机器人自动化领域的进展，通过计算机软件每周分析上百万张超高清晰度的细胞显微照片，构建了上百种模拟罕见病的细胞模型，然后通过先进的成像技术对这些模拟罕见病的细胞模型拍照。

这种“化简为繁”的方法不仅在诊断上有所帮助，在罕见病的另外一重困境——药物开发中也能起到不小的作用。

罕见病的药物市场一直没能受到药物研发公司的重视，因为其患病人数少、市场需求小、药物研发成本高，罕见病的药物被形象地称为“孤儿药”。目前，我国对于“孤儿药”的研发处于初期阶段，罕见病患者的治疗药物基本依赖进口。

AI在处理数据方面绝对是个高手。AI药物研发的底层核心就是知识图谱，其实质就是将来自实验室的理化数据、各种期刊文

献中的研究成果及临床数据等原本没有关联的数据连通，将离散的数据整合在一起，从而提供有价值的决策支持。

所以，如果能通过细胞显微图像判断疾病特征，整合更多离散的数据，我们也可以通过这个方法来进行药物筛选，提高筛选的速度和成功率。

维护生命尊严，**AI**要从哪里做起？

虽然找到了途径去助力罕见病的诊断，但**AI**要解决的事情却远远不止于此。

从马斯洛需求层次理论出发，我们能够发现，相较于正常人，罕见病患者面临着更多难以满足的需求，他们既没有稳定的收入，又面临着巨额医疗支出，最低层次的生理需求尚难满足，更不用说最高层次自我实现了。

《罕见病群体生存状况调研报告》显示，在经济上，罕见病患者对家庭的依赖程度非常高。近八成受访者主要靠家庭成员的收入来维持生计，一成受访者是依靠自己的劳动收入，其余受访者则依靠政府提供的生活保障、社会捐赠或者其他。

罕见病患者面临工作机会少，在求职过程中遭到歧视等问题。

但谁不想活得有尊严呢？当被病魔折磨的患者们在日复一日的生活中渐渐失去自我价值感时，能够以独一无二的“我”的身份，自主地去做一件事，常常是激动人心的。于是，如何为罕见

病患者提供自我实现、寻求意义的渠道，才是AI要真正解决的问题。

AI首先要做的是延展罕见病患者的工作技能，延展途径有两个：一个是能力培训，这也是目前许多社会组织在做的；另一个是找到同等替代力的工具，这也是AI能大展身手的地方。一般来说，罕见病患者无法工作的原因往往在于身体机能的缺失，而AI本质上则是人的机能的延伸。

例如，对于渐冻症、肌无力患者，我们可以运用脑机接口技术，这种新的神经生物能力能够扩展渐冻症患者的运动能力、感知能力和认知能力，患者可以有效地、完美地将自己的思想转化成运动指令，由此可以操控简单的工具，或者调控一些机器。

其次，AI还需要建立一个匹配工作的数据库，为这些患者有针对性地找到一份工作，其中最有效的方法是搭建一个人机平台，将大众的力量通过这个平台汇集在一起。

2018年，谷歌曾在其搜索结果页面上推出了新的就业搜索功能，可以让用户在几乎所有主要的在线工作网站（如LinkedIn、Monster、WayUp和Facebook等）上搜索职位，但其本身并不发布任何企业的招聘信息。

针对罕见病患者的这个平台应该将重点放在社会服务工作项目上，这些项目可以让罕见病患者参与到为特殊人群服务的社会工作中，在与病友进行互动、满足其社交需求的同时，也实现自我价值。这样就实现了从以往的“众筹”变成“众就业”，从“授人以

鱼”变成“授人以渔”。

最后，AI应该建立一种模型，使罕见病患者的公平劳动行为有着具体化的定义和框架，来适应罕见病患者大量就业的变化。在建立模型时，既要避免企业通过“逆向选择”来识别特定群体及个人，也要避免极少数的患者以弱势群体的身份对企业进行“道德绑架”。

当然，维护患者的生命尊严，仅靠AI是远远不够的。要真正保护这些患者，还需要政府在相关的就业监管和政策制定上保持高度的重视。

同时，我们也应该在AI开发上使力，当罕见病患者的就业需求成为部署AI的中心时，关注和理解这些需求就变得十分重要。在AI开发上，我们可以引入不同的受保护人群，这样既为这些群体创造了新的岗位，如“数据清洁工”（能够“清理”数据，为模型提供有效数据），又能够开发出具有共情能力的AI系统。

拒绝开颅，AI能扩充人的脑容量吗？

还记得《最强大脑》第四季中那一场人机对战吗？在那一期节目中，镜头扫过，名人堂里的“超级大脑”们面露难色，没有人敢上去应战，尴尬的气氛简直都要溢出屏幕了。这一次“人机大战”也把“人类不敌机器”的话题推向了高潮。

事实上，现在说人类敌不过机器还为时过早，但AI在有些方面确实还是占据了一点“剑走偏锋”般的优势，那就是“库存量”。类比起来，也就是人脑的脑容量，虽然我们拥有灵长类动物中最大的脑容量，但人脑再强大，也难以存储某个事件的百万甚至千万条的数据。

虽然脑容量大的人不一定聪明，但脑容量大一点，就好像捡了个“三级包”（游戏词汇，指最大容量的装备包），虽然不一定“装得满”，但起码可以“装”。那么，我们有什么办法可以扩充自己的脑容量，升级一下自己的“装备”呢？

扩充脑容量，有哪些方法？

总体来说，人类脑容量演变呈“两头缓慢，中间迅速”的S形曲线式发展，其演变受到了包括行为和基因等多种因素的共同调节，如基因进化、工具的制造和使用、劳动等。而对于扩充脑容量，我们或许可以从这些影响因素中得到一些启发。

1. 复活已灭绝的物种

作为智人的一个分支，尼安德特人的脑容量甚至逼近了现代人的脑容量值。所以，如果技术上可行，我们是不是能直接“复活”尼安德特人，让“大”脑人成为现代人的一部分呢？

哈佛大学遗传学家George Church在接受德国《明镜周刊》采访时曾表示，他已快要开发出克隆尼安德特人的必要技术，届时他需要找到一位“愿意冒险的女人”充当代孕妈妈，孕育出3万年以来第一位尼安德特婴儿。

这听起来是不是很疯狂？别惊讶，Church的说法可不是天方夜谭。早在2009年，德国马克斯·普朗克研究院莱比锡人类进化研究所的科学家们已经完成了对尼安德特人的基因组测序，并将完整的尼安德特人基因数据公布在互联网上，每个人都可免费获得。

拥有了基因组，“复活”灭绝物种就成为可能。此前也有一些生物学家克隆濒危或灭绝动物的先例，在2009年，已经灭绝的西班牙山羊亚种巴卡多，从冰冻的皮肤样本中被克隆出来。虽然新生儿因呼吸衰竭而死亡，但它的诞生表明，复活已灭绝的物种还是可行的。

Church的实验室就正在创造尼安德特人的细胞，想象一下，一个新生的尼安德特婴儿的脑容量能有多少？“大”脑带来的是愚钝，还是超级聪明？这些谜题的答案，似乎近在眼前了。

2. 脑联合的畅想

一个大脑的容量或许有限，那再加几个大脑呢？如果可以建

立一个脑对脑的界面，将两个活人的大脑连接在一起，不就可以实现“扩容”了吗？

在历史上，也有人这样考虑过。诺贝尔奖获得者默里·盖尔曼曾在他的著作《夸克与美洲豹》中提到过，“某一天，人类能够与一台先进的计算机直接用线连接起来，思想与情感完全被共享，再也没有了语言上的选择和欺骗……这必然会造成一种复杂的新形式，它是许多人的真正混合。”

大家是不是觉得很熟悉？没错，这就是目前非常火热的脑机接口想法的雏形。但与脑机接口不同的是，在扩充脑容量上，计算机不再作为终端的信息输出设备，而是作为人脑之间的连接枢纽。

脑机接口技术的成功也表明，“大脑—机器—大脑界面”的逻辑是合理的。杜克大学的研究员曾经做过一项实验，他们把动物大脑连接到了一起，并合作完成了一些简单的任务。大脑连接后，猴子们展示了动作技能，老鼠也表现出了计算能力。

3. 让神经之间多点联系

注意，现在所说的这个方法可能是最简单也最容易实现的“扩容”方法了。实际上，认知挑战是塑造“大”脑的主要驱动力，这与AI有点相似，AI也是通过不断地认知、建立GAN来丰富自己的数据库的。

对于人类而言，认知挑战具体的驱动力可以分为生态驱动型和社会驱动型。生态驱动型模型指的是非社会性的环境因素，如

寻找食物和躲避被猎食；社会驱动型模型则是社会性因素，如人与人之间的协作、竞争等社会关系。

2018年，《自然》刊登的一项研究首次通过数学模型确认了生态挑战是人类脑容量增大的主要驱动力之一，否定了之前很多学者推测的一种可能，即人类社会的复杂性导致了人脑容量的增加。

毫无疑问，我们具有在外界环境和学习经验的作用下塑造大脑结构和功能的能力。大脑内部的突触、神经元之间的连接可以由于学习和经验的影响建立新的连接，从而影响个体的行为。神经元之间连接增多，神经密度变强，大脑容量自然就大了。

我们要不要扩充脑容量

如果人类只做一些简单的事情，那我们完全不需要大脑，这里可以参考一下单细胞生物。但是，我们希望能体验到愉悦感，而不只是为了简单的生存，有时候还希望可以将这种感受表达出来，所以，大脑引发出的是更为复杂的社会行为。



一直以来，没有人会质疑体积更大的大脑能够胜任更复杂的功能，不过越来越多的科学家开始思考体积变大是否是大脑功能增强的唯一途径。为了变得更聪明、记忆力更好，而去扩充一个“大”脑是否是必需的呢？

1. 尼安德特人的囚徒困境

长久以来，尼安德特人多被认为是因智力不足以应付环境的变迁而灭亡的低等人种。而实际情况是，尼安德特人非常成功地应对气候挑战的时间至少有20万年，比延续至今的现代智人还要长12.5万年到15万年。

尼安德特人的脑容量比智人还要大，与现代人十分接近甚至高于现代人，目前也有很多证据表明尼安德特人是拥有不低的智力的（他们甚至还会化妆），那为什么现代人的这个“近亲”会灭绝呢？

关于尼安德特人灭亡的原因科学界已经有了种种猜想。我们在这里也可以开一个“脑洞”——尼安德特人逼近现代人的大脑容量极有可能促进了这个人种的灭绝。大胆设想后，接下来就是小心求证了。

在丛林法则中，胜利者未必是最聪明的，也未必是最强壮的。反而，因为拥有高智商、自我意识很强的尼安德特人更容易陷入困境，明明有合作双赢的选择，但是博弈的结果都是选择背叛，个人理性有时能导致集体的非理性——聪明的尼安德特人因自己的聪明而作茧自缚，从而损害了集体利益。

BBC的一个纪录片中也曾提到过，东非草原的智人的协助能力超过其他地方的人类，就像东非草原上其他动物，如狮子、鬣狗等，合作的方式都比别处同科的动物更复杂，语言的发育也最充分，这些物种往往更容易在竞争中胜出。与尼安德特人相对的智人，在艺术方面也会更有协作精神，其艺术成就也往往有利于种族之间的识别、认同和互助。

如此来看，扩充了脑容量也不一定是一件好事。我们可以用以下两句话总结尼安德特人灭绝和智人留存的原因，即“聪明反被聪明误”和“三个臭皮匠顶个诸葛亮”。

2. 简单的大脑效率高

在人类的进化过程中，脑容量一直是增加的，直到增长到3万年以前的约1500mL之后，出现了下降趋势，到现代约减小为1300mL，不仅脑容量减小了，而且脑重量与体重的比值也减少了。

大脑为什么会出现小型化的趋势呢？科学家用蜘蛛做了一个实验，从中发现了一个现象：纵使大脑体积相差多个数量级，但这些蜘蛛表现出来的行为能力却几乎没有差异。

除了蜘蛛，迈阿密大学的科学家也对70余种的鸟类进行了调查。最终，科学家提出，一个微型的大脑或许能够更加高效地完成所有复杂的生命活动。

“进化，不需要任何多余”，也许这就是大脑变小的原因。而“大脑微型化”这一结论不仅影响了我们对脑容量大小的认知，更为AI的研究提供了新的见解。

简单来说，“大脑微型化”的行为就类似于计算机芯片中晶体管体积的缩小。如果这项研究能够取得突破性的进展，程序员就有可能从中获取灵感，将小动物的“小”脑融入AI的设计中。

宠物智能医疗兴起，关键问题解决了么？

中国有多少家AI+宠物医疗的公司呢？

数量上，恐怕还不到一个省内医院总数的1/3。

让宠物医疗拥抱AI，我们有多少个迈向这个领域的动力，就有多少个放弃的理由：论蛋糕大小，“它经济”远远比不上“他经济”和“她经济”；在智能领域，宠物医疗的战略意义也远不及人类的医疗、安防、教育等领域。

在融资上，智能宠物医疗领域还没有获得顶级资本的青睐。所以，当你发现宠物医生使用的医疗器械还停留在很多年前时，也就不足为奇了。

然而，随着互联网宠物医院业态的涌现，有关宠物医疗的创业方向越来越多，AI+宠物医疗也开始成为一门好生意了。

复制人类的智能医疗模式可行吗？

面对宠物医疗市场这块变得诱人的蛋糕，我们要思考的问题是，AI应该如何为其增值？

同样是智能医疗，宠物行业是否可以完全复制人类的模式，在预防、诊断和监管方面做出成绩呢？这恐怕很难，原因在于以下三个方面。

1. 大数据诊断行不通

在人类临床上，利用大数据来做病例模型、筛选治疗方案是很常见的，但在宠物临床上的普及率却非常低。这是因为宠物传染病、皮肤病的防治，以及免疫驱虫服务的服务流程较为标准化，从业人员往往能够凭借经验直接给出诊断结果。

除了北京、上海、广州、深圳等一线城市里接诊数量大的宠物医院，业内对成像系统做模拟病例都不太看好，一是因为成本太高，二是因为一般需要使用病例模拟的患宠病情发展会很快，可能根本来不及采取治疗手段。

所以，在宠物医疗领域，AI、大数据在宠物诊断这方面可能真的派不上大用场。

当然，大数据仍然有其他的用处。例如，在药物研发中使用基因数据、临床实验数据的共享、电子病历系统的广泛使用等。此外，医疗信息资源库可以为宠物提供个性化健康管理，包括智能导诊、健康记录和疫苗接种预警等；医疗信息资源库也可以为医生提供个性化临床决策支持。

2. 智能感知还不够

近年来，比较常见的宠物智能医疗产品就是可穿戴设备，包括智能项圈等，一般具有监测宠物的心率和呼吸速率的功能。在日常生活中，系统会将宠物的呼吸、心率数据发到云端进行分析，一旦数据异常，主人手机上的客户端就会建议主人带宠物去医院。

宠物智能穿戴虽然能够满足人们对宠物的关爱需求，但和人

类穿戴产品一样，存在产品同质化严重、产品功能与智能手机功能高度重合的问题。

此外，宠物与人不同，人类虽然长相各有不同，但起码是同一个物种。但是在宠物界，不同宠物的体型不同，运动形式也比较多样，同时，宠物还会有与心理状态相关的高级行为。

显然，依靠AI来感知动物的体态和运动形式还不足以全面了解宠物的生理状态。参考人类“望闻问切”的就诊模式，我们还需要“询问”动物，得到动物自身的“回答”，才能够更准确地知晓宠物的“感受”。

宠物虽然不会说话，但它们都有自己的叫声，这也是宠物医疗行业里智能问诊的突破点。日本有科学家发现，动物的叫声包含疼痛、发情、饥渴等信号，甚至会表达出个体身份等丰富的信息，也能反映出其身体、生理状况。

《亚澳动物科学》杂志就发布过一种数据挖掘算法，能够从奶牛叫声中提取信息，并进行早期异常检测。实验结果表明，该方法的准确率大于94%，且系统成本低，可实现无接触、实时检测。

所以，通过AI去分析某种动物的叫声信号，建立具有个体显著性差异的多个声学特征，就能将声学特征作为宠物医疗的“指示器”。

3. 宠物医院更看重营销而非治疗

目前，国内知名的连锁宠物医院瑞鹏和瑞派，它们在线下都有100余家门店，领跑国内宠物医疗行业。宠物医院在本质上还是一个“商店”。

这家“商店”面向的客户主要为单个养宠人群，收入来源较为分散，因此，它必须具有很强的营销能力才能养成用户习惯。宠物医院要想打造核心竞争力，需要具备良好的线上导流能力和线下门店扩张能力。

诚然，治疗效果和医疗设备也在宠物医院的竞争格局之列，但营销、人才输送、门店复制和服务标准化才是宠物医院形成连锁的核心要素。

如果要搭建一个智能连锁平台，就代表宠物医院需要连接一个产业链，这个链条包括专家、设备、耗材、地方宠物诊所、代理商等。在这个模式中，角色众多，不确定因素太多，成本更是不能确定。

最后，这些成本由谁承担？羊毛出在羊身上，宠物主人无疑要为这些智能医疗设备买单，令人遗憾的是，宠物可没有医保。

如果宠物医院有能力自己研发，那购买智能设备自然是更划算的。如果宠物医院没有更多的资本，那这类设备又是否真的能提升宠物医院的客流量呢？

智能宠物医疗未来可能的变量，还是取决于政策因素。例如，制定行业规范，包括宠物医院设备标准、宠物管理办法、宠物医疗保险等。

规模化过程中，**AI**要做宠物医生的防御衣

在技术落地的过程中，我们还要知道，宠物智能医疗并不是一个简单的技术升级，它在现实世界中会面临许多复杂的情况。而如何让这个技术惠及宠物医生、宠物和主人三方是比技术突破更让人头疼的问题。

“宠物医生被宠物咬了，主人是不负责的。”一位宠物医生告诉我们，这几乎是业内默认的一个“规矩”。在发生医生被宠物伤害的事件后，态度比较好的客户会向医生道歉，但大部分不会给予赔偿，而一些比较刁蛮的客户，反而会觉得这是因为宠物医生不够专业，甚至觉得是他们活该。

在宠物医院里，医生的权益很难得到保护，面对这样的情况，医生也掌握了一套“生存法则”——“把本身不严重的病说得严重点，不能让主人觉得这个病很容易治。”受访的宠物医生告诉我们，因为宠物疾病和诊治过程发生意外的概率较高，而为了减少纠纷，医生往往更倾向于隐瞒部分信息。这也是有人觉得宠物医疗信息不透明的原因之一。

所以，在宠物行业，智能医疗产品的设计和使用体验不仅要满足宠物主人和宠物的使用需求，还需要解决宠物医生的难题。

这里可以参考人类医疗的做法。2017年，IBM与Atrius Health达成一项协议，共同开发一个基于云的服务来改善医患体验，通过大数据描述不同人在不同治疗选择下的健康结果，进而支持医生与患者共同决策。

尽管共同决策的好处很多，但许多医生表明很难将共同决策技术整合到他们与患者的直接互动中。患者在医疗过程中“做自己的主”反而会增加医疗风险。

如果将场景换成宠物医院呢？效果可能会比较好。

一般来说，宠物行业存在一个共性——使用者、购买决策者是分离的。与人类医疗相比，这种共性的好处就是决策者对宠物医生的忠诚度会比较高，医疗风险也会相应降低。

通过提供影响宠物健康的多种因素的整体视图，可以使宠物医生与宠物主人共同决策。这样，既可以提高宠物主人的认知，也可以减少无根据的护理和成本变化，进而改善宠物医生和宠物主人的关系。

你愿意让你的宠物“数字化生存”吗？

在思考宠物智能医疗的方向时，我们试图去解决一些根本的问题——如何让宠物更加健康地、长久地生存下去？能否一劳永逸地解决宠物疾病等问题？

对于人类而言，我们提出了思维移植、数字化生存等观点，也就是让人类不再仅仅依赖于肉体，这样就可以不再面对疾病问题了。仿效这类做法，你是否愿意让你的宠物进行“数字化生存”，保持“永生”？

我们无法提取一只狗的意识，却可以收集一条狗在世界上留下的所有痕迹：样貌、体型、经历、吃喝的喜好、“撒娇”的形

态.....通过对这些数据的学习，我们完全有能力在数字世界“再造”出一条“永生”的狗。

目前也有很多技术包括VR、AR等都支持人们在虚拟世界“抚摸”一只狗，甚至进行一些逗猫狗的游戏，如抛球、甩盘等。而且，比人类数字化生存更好的是，用数字化技术“模拟”一条狗，不会面临太多社会伦理方面的问题。

但是，数字化的宠物也会给人一种错觉——我们有人陪伴，且这个人永远不会离开。毋庸置疑，这种情感陪伴也是有风险的，就像曾经风靡一时的电子宠物，让人们沉醉于虚拟世界，而不愿回归现实。

挖掘“黑马级”的智能医疗器械市场，难在哪了？

“假如没有症状，国内大多数患者都不会主动要求做胃镜检查。”中山大学第六附属医院的医生在接受采访时说道。

关于胃镜的可怕程度，许多人都有所耳闻。笔者也询问过现在在医院工作的同学，“是不是越严重的胃病做胃镜的反应就会越大？”

“不一定，比如胃癌患者做胃镜就不会有丝毫反应，”他回答道，“但是，往往做了一次胃镜后，会发现自己的病已经到了晚期。”

虽然有人说不同的人应激反应的程度不同，胃镜也没有传说中那么可怕，但是没有人能提前知道自己不是生理反应最激烈的那个。未来医疗健康将有望成为拉动内需的“三驾马车”之一。那么，在如今这个关键节点上，我们到底能不能完成把医疗器械彻底智能化的使命？

从检查到吃药，最好都“别找我麻烦”

研究表明，早期胃癌的术后5年存活率可达95%以上，但是如果到了中晚期，术后5年的存活率仅有20%。

基于这样的现实背景，智能胶囊胃镜开始出现在世人的视线当中。目前的智能胶囊胃镜主要分为两类，一类是内部驱动的微

型仿生胶囊机器人，但这类机器人普遍缺乏持续的动力系统，除非开发出无线充电的技术，否则，只能停留在实验室阶段。另一类是通过磁控方式实现机器人在消化道中的运动，这类胶囊机器人已经开始进入人体临床试验阶段。2013年，安翰公司推出了世界上首个获得SFDA批准应用于临床诊断的胶囊内镜机器人NaviCam。患者只需吞下一颗胶囊机器人就可以接受无痛无创的胃镜检查。

当下，中国消化道诊疗存在医师资源紧缺、医疗资源分布不均等多重问题，优质资源对于广大的群众来说依旧十分难寻。

另外，吃药难也是一个大问题。试问一点病痛都没有的人能有几个？但真的能够做到按时按量坚持吃药的患者又有几个？2017年，麻省理工学院团队研发出了一款新型药物装载器，该装载器可以在特定的时间内释放药物的颗粒结构，实现精准的药物输送。不仅如此，研究人员还将其用3D打印技术制作了实体。但是目前该装载器仍然存在稳定性方面的问题，因为这样的精准放药机制必然会有较高的成本造价，如果其能够起作用的时间不能达到令大众满意的程度，则很难得到量产。未来，麻省理工学院团队的目标是将该精准放药机制的稳定性提高到可以持续数百天以上。

智能医疗器械还有什么别的问题？

由于医疗器械属于医疗系统的基础设施的构成部分，因此，其对于全面的智能化医疗服务有着极其重要的意义。不过，在智能医疗器械的实际运用过程中仍然还有几个问题需要得到重视。

1. 找准使用范围，智能医疗器械不是“万能钥匙”

智能医疗器械在未来或许可以帮助我们实现远程医疗。

《2016年美国医疗健康消费者调查报告》提到，虽然近半数参与调查的消费者中有身患慢性病的，也有没有患慢性病的，他们均表示愿意使用远程医疗服务进行急性后期照护或慢性病情监测，但是消费者貌似对使用远程医疗急性病症诊疗不太感兴趣。这就要求技术开发人员需要充分了解患者的就诊诉求，避免开发无用的产品。

2. 患者对于植入物的接受程度还很有限

在医疗领域，目前常用的植入物基本上都与大病有关，最常见的植入物是全金属关节，而一般只有瘫痪的病人才用得着它。这也就是说不到非用不可的时候，病人一般很少选择在自己的身体中植入东西。英国药品和保健用品监管署曾警告称，大多数全金属植入物患者需要每年进行血液检测。植入物让患者不得不考虑其所带来的炎症、组织损伤、已经残留在身体中的金属碎屑给身体造成的影响。因此，对于胃病之类的“小”病，患者对于植入物的态度可能并不会像我们想象中的那样好，也就是说，大多数人不会为了实现精准吃药而植入这样的东西，除非这样的植入物可以实现无害溶解。

智能医疗器械的普及并非一朝一夕的事情，从顶层到底层，还存在许多值得深思的细节。因此，智能+医疗器械这块“大饼”，说它“香”也“香”，说它“难啃”也不是假话。

2 机器人“捕手”

外骨骼机器人能让我们秒变“钢铁侠”吗？

京东曾在头条号上发布视频，晒出京东公司为物流工作人员研发了一款外骨骼机器人，佩戴在工作人员的背部，以尽量减少频繁弯腰拿取商品造成的腰椎损伤，如图2-1所示。



图2-1

无独有偶，福特公司在2018年宣布与Ekso Bionics公司合作，在美国的两家汽车工厂内对外骨骼机器人进行测试，并展示出了实际应用的视频。这款机器人能够帮助工人减轻5~15磅的重量，从而减缓每天4600个日常高空作业任务对工人上半身造成的压力。

有京东和福特这样的知名公司带头，可以预见，外骨骼机器人将很快在各大工厂被应用。对普通员工来说，借助外骨骼机器人能够减轻工作中身体的负重压力，降低肩周炎、腰椎间盘突出等常见慢性疾病的发病率，确实是令人欣喜的好消息。

在为关心员工的好公司点赞的同时，我们不禁想搞明白，外骨骼机器人到底是什么高科技？穿上了它，是不是就拥有了钢铁侠般的超强能力？

可以肯定的是，影视作品中的“高科技”往往极度超越现实。虽然AI正处在飞速发展阶段，但是现实生活中的外骨骼机器人的研发速度并没有那么快。我们简要讲述了外骨骼机器人的前生今世，下面让我们共同了解这项可能改变人类生活状态的黑科技。

外骨骼机器人其实是美国军用装备

2000年，美国国防部基于增强士兵体能，提高单兵作战能力的目的，提出“外骨骼机器人”的概念。实际上，外骨骼机器人不同于我们在以往的军事战争或科幻电影中见到的“铁甲”，而是附着于人体外部的“AI”。它为穿戴者提供保护，并根据人的肢体活

动来驱动机械关节重现动作，用以提供额外的动力，帮助使用者跑得更快、跳得更高、负重更强。

美国军方在研制招标书时提出4个要求：①外骨骼系统必须能够保护穿戴者，以大幅度减少伤亡。②外骨骼系统携带的能源必须能够维持24小时的工作，而且必须轻且完全无声。③外骨骼系统必须能够让士兵背负更大质量的同时仍然能够跑得更快，跳得更高、更远。④外骨骼系统必须采用流畅、高效且完全无声的液压元件。

在美国国防部的巨资支持下，雷神（Raytheon）公司的XOS和洛克希德·马丁（Lockheed Martin）公司的人类外骨骼负重系统（Human Universal Load Carrier, HULC）脱颖而出。同样是研究外骨骼机器人，两家公司的区别在于，雷神公司专注于研究全身机械外骨骼，而洛克希德马丁公司专注于研究可增强下肢能力的机械外骨骼技术。

从2000年开始立项，Sarcos Research公司经过6年研发，成功制作了第一款原型机XOS 1，也是XOS 2的原型。XOS 1搬运重量的实际和察觉比为6：1，也就是说当搬运60千克的重物时，穿戴者会感到只搬运了10千克。2007年，雷神公司收购了Sarcos Research公司，并在2010年推出XOS 2。XOS 1当初设计的目的主要是证明理论的可行性，XOS 2才对性能有了更高的要求，需要更轻、更快、更强的设计，并且需要降低能源消耗。XOS 2自重95千克，相比第一代产品的重量轻了10%，搬运重量的实际和察觉比为17：1，能源消耗也降低了50%。

2004年，大名鼎鼎的加州大学伯克利分校下属人体工程和机器人实验室成功研发出伯克利下肢外骨骼（BLEEX）——这就是HULC的原型。2009年，该实验室推出HULC，同时授权洛克希德·马丁公司在军事领域对其进行推广。HULC自重24千克，由背包式外架、金属腿框架及相应的液压驱动设备组成。穿戴者的负重不会作用在本人身上，而是直接由机械外骨骼传至地面。HULC的功能主要包括力量增强和耐力增强两项。力量增强体现为士兵佩戴后可搬运90千克的重物而毫无感觉，耐力增强体现为士兵在3.2 km/h的行动速度下可降低5%~12%的氧气消耗。

外骨骼机器人在医疗领域大放异彩

在2014年巴西足球世界杯开幕式上，一名下肢瘫痪的少年穿着智能机械骨骼外衣开球的飒爽英姿引爆了全场，向全球展示了康复外骨骼机器人的风采，也使这个集人机信息交互、机器人自动控制、神经康复工程等多学科知识于一身的高科技新成果引起了世人的瞩目。

针对腿脚不便的老年人及残障人士，外骨骼机器人可以辅助或恢复其腿部的行走移动，从而提高老年人及残障人士的生活质量。针对正常人，外骨骼机器人可以提高穿戴者的负重能力，并减少体能消耗。

2013年，日本筑波大学研制的HAL的第五代产品HAL 5上市，同年2月HAL 5获得了国际安全认证，这是第一个获得该认证的机械外骨骼产品。同年8月，HAL 5获得欧盟认证，其被允许在欧盟境内进行医学治疗。HAL 5的重量只有10千克，基本可以满足

足常人的承受力。通过感测体表的一些生物信息，如肌肉的运动、神经电流的改变，HAL 5可以模拟穿戴者的动作，并且在平行的方向上增强穿戴者的力量和耐力。

ReWalk是以色列ReWalk Robotics公司的明星产品，它同样拿到了美国FDA的认证，进入了上市销售阶段。它是一种可穿戴式辅助运动的下肢外骨骼，是一种帮助腰部以下脊髓损伤的人重新实现独立行走的外骨骼机器人。它由腿部支架、独立控制的髋关节和膝关节电机、充电电池、附着于关节框架的计算机系统及一副控制平衡的拐杖组成。髋关节和膝关节处均包含可驱动的屈伸自由度调节器，踝关节处则设有带限位作用的双向矫正关节，可以防止人体受到伤害。

2016年5月，哈佛大学怀斯研究所发布了与ReWalk Robotics公司合作研发的新一代外骨骼机器人，其采用拉伸检测运动技术，之后通过拉索驱动踝关节的动作，实现动作驱动。

国内对于外骨骼机器人技术的研究开展较晚，发展较缓，与国际水平相差较远。在AI浪潮的推动下，中国也有大量上市公司和创业公司开始涉足这一领域，如2017年深圳迈步机器人科技有限公司发布其首款面向偏瘫患者的外骨骼机器人BEAR H1、上海傅利叶智能科技有限公司发布了外骨骼机器人Fourier X1等，但目前还没有企业获得CFDA认证，也因此并未真正有产品完全进入到商用阶段。

从京东发布的视频来看，京东的外骨骼机器人动力系统很可能采用了结构较为简单的气压驱动。气压驱动式动力系统在有负

荷的作用下，速度容易发生波动，不适用于较为精密的场合，只能在小功率场合应用，想要开发出高精密的外骨骼机器人可能还有较长的路要走。

外骨骼机器人让我们成为“钢铁侠”还比较困难

尽管外骨骼机器人取得了不少成就，但无法回避的现实是，目前技术上仍然面临不少困难。例如常见的液压系统所处的工作环境远比常规的地面液压系统复杂，液压系统在运动时常发生晃动、倾斜，同时系统中液压介质的容量也随时发生变化，容易导致运行不流畅、噪声大、可靠性不足。

用“钢铁侠”的标准来评判外骨骼机器人是否成熟可能有点苛刻，但无论如何，作为带辅助性质的AI器具，至少需要满足三点要求：智能化反馈足够及时、佩戴足够舒适、续航足够持久。这样的外骨骼机器人才能算得上成熟，而目前要实现这三点都面临一定的困难。

首先，足够及时的智能化反馈是操作流畅、体验良好的首要条件。在最基本的动作同步上，目前肌电信号采集的方法有很严格的外界环境限制，传感器会因汗液等细微因素的影响而受到干扰，使得外骨骼在进行动作的反馈时产生延迟，难以胜任高难度动作。更智能的外骨骼机器人，应能够在人机交互环路中准确、及时检测识别用户大脑的运动意图，达到所思即所动的理想运动控制效果。

其次，佩戴舒适是辅助器具的基本要求。目前的外骨骼机器

人受限于外骨骼捆绑式的穿戴方法，会对人体造成压迫，导致血液不畅、肌肉变形，并因此影响外骨骼的定位精度。外骨骼机器人未来还必须在增加外骨骼关节自由度、使人们能更舒服地穿戴外骨骼并可以全方位地自由运动上下功夫。

最后，较为持久的续航才能有产品实用性，而从直观的参数上来看，目前的外骨骼机器人依然没有完全达到2000年美国国防部提出的要求，如HAL5的续航时间只有160分钟，ReWalk的续航时间只有3.5小时，远低于24小时的整天续航要求。这需要轻而坚固的复合材料和便携式的持久电源两方面的技术突破，这也是很多电子设备都面临的共同难题。

由此可见，目前的产品满足普适的三点要求都还存在一定的困难，要想一秒变身“钢铁侠”，目前还难以实现。不过，科技总是在不断发展，随着新材料、新技术的不断出现，在不远的将来一定会有更高级的智能人机交互、更加灵活舒适的外骨骼机器人产品投入市场，造福大众。

我们距离下一个“阿尔法法官”还有多远？

随着百度、联想等公司先后进军AI领域，AI领域的爆炸式发展就进入了一个新阶段。

在设想或者已呈现的产品中，AI已经涉猎语音交互、安防、自动驾驶、医疗健康、电商零售、金融、教育等诸多领域。但把它与司法审判、纠纷解决这样严肃而传统的领域相结合，听起来有点儿像是天方夜谭。

但这已经在不知不觉中走向了现实。2016年7月底，中共中央办公厅、国务院办公厅印发的《国家信息化发展战略纲要》中，“智慧法院”赫然在列。2017年7月25日，上市公司久其软件在投资者关系互动平台上表示，其全资子公司华夏电通充分利用在司法领域积累的业务经验及大数据技术完善解决方案，加快建设“智慧法院”，并已在多省高院进行推广与试用。

法律纠纷审判应用AI很有必要

事实上，法律在某种程度上恰恰与AI的基础需求不谋而合。正如美国现代实用主义法学的创始人霍姆斯在其代表作《普通法》中所说，“法律的生命不在于逻辑，而在于经验”，法官在断案时，不仅要根据国家的硬性法律，也要根据自己的职业经验，而那些经验数据的累积场景恰恰最适合AI的应用，在法院中普遍出现的小纠纷案例也提供了足够的数据支持。

反过来，法律的执行机构法院体系，在纠纷解决领域（非刑

事案件）对AI的需求也是迫切的。

首先是效率需求。随着法律意识的增强，越来越多的人把社会、经济生活中的纠纷诉诸法律，这些纠纷数量巨大，但是大多简单易断。

数据显示，2016年最高人民法院受理案件15985件，审结14135件，比2014年分别上升42.6%和43%；地方各级人民法院受理案件1951.1万件，审结、执结1671.4万件。

绝大多数案件都是由基层审理的，现实情况是基层法官的人数十分有限。据法院系统一份调查统计数据，一名基层法官一年至少要处理200起纠纷，在经济较为发达的地区甚至能达到400多起。一名基层法官一个工作日至少要审结一起案件并且还有书写判决等多项工作，这对法官的职业能力和身体素质都是极大的挑战。而且，工作越忙，出现纰漏的可能性就越大。

反观这些纠纷处理，案件按难易程度综合确定为简单、一般、复杂、非常复杂4种案件类型，其中简单和一般的案件居多。如果能够将AI引入纠纷解决机制，就可以较为妥善地解决这个问题。

然后是公平需求。由于法官的职业能力高低不同，同一类型的案件在全国各地不同地区甚至即使是同一法院在不同时间段得出的判决结果都不相同。这固然是现实，但总体而言这对当事人来讲是不公平的，对整个社会的纠纷解决也是不利的。

通俗来讲，AI类似于机器的运作，很大程度上可以避免当事

人通过与工作人员产生利益关系而导致的判决不公正现象。这可以从根源上杜绝在解决纠纷领域的贪污腐败现象，使法律更具有公信力与威慑力。

其次，利用AI可以使相同类型的案件得到基本一致的判决，或者说更为公正的判决，可以在很大程度上避免不同地区不同判决的情况，提高判决结果的一致性，同时防止冤假错案的产生。

公平是法律的要义，高于一切。引入AI，统一尺度，意味着在同质案件上，因人员不同造成的不公平问题可以得到解决。

最后是经济因素。很多低收入人群面对纠纷解决的费用往往望而却步，这实际上使一部分人失去了维护自己正当权益的机会。

AI降低了纠纷解决的费用支出。虽然解决纠纷系统在创设的过程中会耗费大量的人力、物力，但案件的处理速度与成本都会使当事人满意，效率与公正就都能实现。

灵活性决定可行性

司法领域的规则都是制定好了的，所以，AI应用到法律审判与应用到其他领域（对于没有成文的规则，AI会用算法总结并遵循）有着很大的不同。

我们对AI能否解决法律纠纷的忧虑，不在于它是否能严格按照法条去判断，而在于它可能会太过于按照法条去判断，即过于刻板，得出的判决会十分没有“人情味”。

在进行判决时，法官不仅要考虑法条对案件的情节的可适用性，还要考虑当事人的诉求、案件在一定范围的影响，甚至我们国家特殊的国情。法官传统断案模式具有更多的人情味在里面。法院审理案件并不是非黑即白，AI平台还必须能检测出这其中的精妙之处。

但是，随着科学的不断进步与发展，只要信息成本足够低，我们就可以将所有的相关因素都输入进去，不管是硬性的法律规则，还是软性的人情社会，将AI的落脚点体现在人上，这样的问题会慢慢得到解决。利用AI系统代替基层法官处理简单或者一般的案件，快速简单地做出对纠纷的处理，这样高效率的司法方式是我们想要追求的。

不过，这也说明，使用AI是需要一定条件的，那就是简单或者一般的案件，对于复杂的纠纷，AI还难以达到妥善处理的水平。

鉴于目前人们对于AI的不了解与不信任，在初始阶段，法官可以将AI的判决结果作为预审，或者利用AI辅助完成审判过程中的一系列事项。

例如江苏法院建设“云上法院”，以“江苏法务云”为载体，实现同类案例、审判资料的智能推送，为法官办案提供辅助。上海法院则建立了大数据司法公开系统，审判流程、裁判文书、执行信息、新闻信息、联络服务全程公开、留痕、可视、可监督。

河北法院主推的“智审1.0系统”，在将AI应用于纠纷解决上走

得更远。它有五大功能：自动生成电子卷宗；自动关联与当事人相关的案件，避免重复诉讼、恶意诉讼或虚假诉讼的产生；智能推送相关的辅助信息辅助审判；自动生成与辅助制作各类文书，目前已实现裁判文书80%的内容一键生成；智能分析裁量标准，根据法官点选的关键词，自动统计、实时展示同类案件裁判情况。

目前，AI的应用还停留在辅助阶段，未来随着AI应用的进一步加深，相信一些标准化的纠纷案例能够实现AI自动审判。不过，一旦当事人对于AI所产生的预审判决不满意，启动正常司法程序，使用传统的司法模式也应该被允许。

AI开始应用到法律领域，让律师们都慌了

英国伦敦大学的科学家研制出一款新的AI产品，它能处理法律文件并对案件做出判决，自动化运算的“审判”结果和人工审判结果一致率达79%。

应用AI，审判工作和律师的日常工作在算法上是高度重合的。如果AI急速发展到能够考虑到所有因素、解决各种复杂疑难案件的程度，那么律师们就会慌了。

从效率的角度来讲，AI是高效率的，是我们追求的。将所有的相关因素设置好，这时不仅有效率，更有公平，不会一边倒。那么到时候，还会需要这么多律师吗？

麦肯锡全球研究院在2018年1月表示，在“当前技术”（广泛使用或至少在实验室测试）下，23%的律师工作能够被自动化。

技术进步的速度是无法预测的。以前，大家都认为AI主要威胁标准化的日常工作岗位，类似律师这样非标准化的专业技术人员是安全的。但随着AI取得的进步越来越大，所谓的非标准化的工作内容也可以被替代，如律师筛选文件、寻找相关段落等例行工作任务。

越来越多的律师事务所应用AI，比如，主要受理破产案件的Baker&Hostetler利用AI筛选大量的法律数据。

又如，Kira Systems是一家致力于用AI分析、评审合同的服务商，它把律师们需要评审合同的时间缩短了20%~60%。如果这个效率继续提升，AI做到与律师工作别无二致，那么许多人可能会因此失业。

在此背景下，Dentons，一家拥有超过7000名律师的全球性律师事务所成立了Nextlaw Labs——一个创新和风投部门。除了监测最新的技术，该部门还对7家法律技术初创公司进行了投资。这是一种对长期风险的警觉。

十年前，没有人想到今天移动互联网会发展成这样。同样，十年后AI会发展成什么样也无法预测。但至少目前，司法领域站在了AI变革的前沿。

作诗的**AI**机器人为什么能骗过行家？

“微明的灯影里/我知道她的可爱的土壤/是我的心灵成为俘虏了/我不在我的世界里/街上没有一只灯儿舞了/是最可爱的/你睁开眼睛做起的梦/是你的声音啊”

你喜欢这首诗吗？这首诗节选自北京联合出版公司出版的诗集《阳光失了玻璃窗》，而诗集的作者并非人类，而是一个**AI**机器人——微软小冰。

其实除了小冰，作诗超过27万首的编诗姬及专注于古诗创作的九歌机器人也引发了社会的热议。就在**AI**机器人创作诗词成为一种趋势的时候，质疑也接踵而至。

被质疑的初衷：让**AI**去学作诗或许是一个伪命题

在回答**AI**学作诗的目的之前，我们不妨先思考一下，人类创作诗词是为了什么？诗词是经过符号化的信息，它的本质是传达人的思想和情感。《毛诗序》中有这么一句，“诗者，志之所之也，在心为志，言之为诗”，其所要表达的正是诗词创作的目的。作诗是为了言志抒情，表达人的情感，将人的思想寓于其中。我国诗词能经久不衰的原因之一就是其中所饱含的人文情怀，能够引发共鸣，能够给人以精神上的享受。

很多时候，我们读的诗歌和文章，并不只是一些堆积的文字，而是与作者进行的精神交流，进入他的生活或者他的想象空间。正如诗评人秦晓宇所说，“诗人的诗歌中是有其经历、追

忆、愿景等的，这些浓缩在诗歌文本字里行间，才构成了极大的魅力。”而AI机器人只是利用算法运算和严谨周密的二进制作诗。因此，让AI去学作诗被许多人看作是一个伪命题。

被质疑的质量：抖了一下小机灵的**AI**作诗机器人，“钱途”坎坷

AI作诗机器人之所以被认为可能赚不了钱，是因为AI作诗机器人作诗的质量不高，这些作诗机器人“只是和人类玩了一场游戏”。

1. 并非真正作诗，而是基于诗词学习之后的临摹写作

《扬子江诗刊》副主编胡弦认为，目前看来，微软小冰的诗歌基本都是“二手货”，而诗歌的本质是原创。

确实，无论是微软小冰的以图写诗，还是九歌机器人的集句诗，诗词创作机器人的背后是由大数据和算法技术作为支撑的。在搭建好人工神经网络的算法后，AI通过不断输入的诗词数据来进行自我学习，包括它们的形式、内容等。例如微软小冰学习了519位诗人的现代诗，而九歌机器人学习了30多万首诗。

在通过深度学习建立好知识图谱之后，AI机器人能够根据所输入的信息，快速进行算法运算，然后输出诗词，这就是AI机器人作诗的过程。这就意味着，AI机器人并非真正作诗，而是基于诗词学习之后的临摹写作。

2. AI的创作优势与一首好的诗作并没有直接的联系

通过上文对AI机器人作诗路径的分析，我们能够看到与人类相比，AI机器人作诗至少有3个优势。第一，速度快。AI机器人能够通过持续不断的学习提高信息的处理和输入速度，往往比人类作诗更快速，通常在10秒之内就能完成。第二，句子工整。诗词，尤其是古诗词讲究合辙押韵，而AI机器人是按照固定的路径运行的，因此，能够严格按照要求输出。第三，博闻强记。AI机器人有着强大的诗词数据库支撑，在集句诗、藏头诗等考查知识储备、词语搭配等方面有着更为突出的优势。

乍一看，AI机器人似乎比人类更胜一筹，但是别急，我们再来看看在名人名家眼中一首好的诗歌作品应该是怎样的。诗人余光中认为，好的诗歌应该有这些特征：丰富的想象力、高超的语言、讲究音调和意象的营造。浙江“青年文学之星”、中国作协会员高鹏程认为好的诗歌作品应该包含三个方面，第一是第一眼看过去就应该让人动容、动心、动情，第二是经得起反复推敲、琢磨，第三是让读者难以忘怀、反复阅读。所以，如果按照这两位观点，那么AI的创作的优势与一首好的诗作并没有直接的联系。

3. AI作诗机器人在逻辑和语义连贯上还有较大的问题

AI作诗机器人除有创造力、情感等方面的“AI通病”，逻辑不清和语义不连贯也是其存在的突出问题。华东师范大学中文系周圣伟教授以机器人所写的“新酿三篇常细俗，闻韶一曲已为人”为例来点评，他认为单看任何一句似乎没有太大问题，但是合在一起后，两句诗之间看不出任何联系，缺乏内在的逻辑关系。复旦大学中文系侯体健副教授在看过作诗机器人“小诗机”的作品之

后，也给出了“词语混搭、半通不通”的评价。

综合而言，一方面，我们在阅读AI的诗词时，由于诗人的创作背景、个人生平等为一片空白，失去支撑点的想象空间完全被扼杀，丧失了阅读诗词的乐趣；另一方面，通过算法技术输出的诗词，本质上只是词语的搭配游戏。因此，AI被质疑只是抖了一下小机灵，在工整的外壳之下是空洞的灵魂。如鲍南在《北京日报》发表的文章中所言，机器人写诗只是一场文字游戏。

当然，至少从外界看，目前AI作诗机器人的市场已经迈向了逐渐成熟的阶段。不过，问题确实摆在面前。从目前的AI作诗机器人的本身的应用来看，几乎没有更进一步的商业想象空间，还只是停留在博观众眼球的阶段。

AI机器人作诗之后，盈利之路在哪？

从宏观来看，AI作诗机器人所拥有的技术、功能能够在更大的领域发挥作用。因此，我们大胆预想，AI作诗机器人的优势可以被应用于其他领域，以此来获得商业上的价值。

1. 中文语义分析技术能够助力人机交互

人机交互是各领域的产品AI化的关键之一。但是由于汉语本身所具有的复杂性，我国在此领域受到重重阻碍。而AI作诗机器人能够通过深度学习，在中文语义分析能力上得到极大提高，这就为优化中文语义分析的算法提供了基础。例如开发编诗姬的玻森数据就开发了玻森中文语义开放平台，以此来获得利润。

2. 辅助教学，加入AI教育的混战中

马化腾曾表示，AI+教育与AI+医疗领域有望诞生达到千亿美元以上规模的公司。AI作诗机器人可以顺势而上，凭借自身在语言文字学习方面的优势，加入AI教育的混战中。

3. 辅助科研，以全知视角助力克服知识盲点

虽然说AI作诗机器人自己创作的诗歌受到了名家们的批判，但是我们可以凭借其强大的学习能力和记忆能力来构建一个完整和全面的知识图谱，以全知的视角来帮助科研人员进行调查研究工作。例如在理解一首诗歌时，我们不仅仅要弄懂表面的文字解释，包含通假字、生僻字等，还要理解其中的，同时还需要清楚作者本身的人生经历及其所处的时代背景等。而人类因为知识获取和记忆的局限，往往在这些方面或多或少地存在着知识盲点，而通过全知视角的AI则能够有效克服这个问题。

总之，作诗机器人在一定程度上只是娱乐化的产物，正如法兰克福学派的哈贝马斯所说，这些产品本身并不具备艺术性，只是文化工业的产物。作诗机器人在未来可能也无法代替诗人，其诗作也难以供人们研究。但是，作诗机器人背后的技术却是一个值得继续挖掘的领域。

为什么没有出现杂技机器人？

2018年的大年三十下午2：20，一出别出心裁的春晚亮相。不是往年春晚重播，而是由北京电视台科教频道播出的“机器人春晚”，内容和春晚差不多，无非是唱歌、跳舞、小品、相声，但是少了杂技。

杂技起源于人类的生产劳动和战争，是我国最早独立成形的艺术形式之一，一直受人民欢迎，观赏性很强。这不由得使人发问，难道没有杂技机器人吗？

是的，现阶段还没有真正的杂技机器人。

一口不能吃成胖子，几个动作也不能构成杂技

据报道，2018年5月，迪士尼研发出了杂技机器人，它会后空翻。这款杂技表演机器人叫Stickman，它通过摆锤运动获得动力，脱离导线后在空中完成翻转，翻转动作包括一个后空翻、一个双后空翻。最后舒展身体以完成一个笨重的自由落体式的“软着陆”。

当然，事实证明是，它只会后空翻。

Stickman并不是第一款能够完成后空翻的机器人。赫赫有名的波士顿动力公司于2018年在YouTube上发布了人形机器人Atlas最新版本的一段视频，它不仅可以走“梅花桩”，还可以完成原地向后跳转、原地后空翻等动作，长得像忍者神龟。

但是这还是不能称得上杂技。如果仅仅会后空翻就能称为杂技机器人，那么杂技的门槛也未免太低了。杂技演员的基本功，传统的说法是“腰、腿、跟头、顶（倒立）”，随着现代杂技艺术的发展，人们对杂技的观赏性提出了更高的要求，舞蹈训练也被列入了杂技演员的基本功训练之中。一个杂技演员，具备上述几个方面的基本素养是必不可少的。

显然，机器人后空翻只做到了传统说法中的“跟头”，而且还不能做到连续翻跟头。Stickman的研究者也表明，它模仿人类杂技选手的水平非常有限。他们还要继续在这个机器人身上做实验，看看它到底能完成多么复杂的动作。

杂技杂技，就是需要不断炫技

杂技是一项常人不可为的表演艺术。一般有五大技术种类：翻腾类、平衡类、软功类、高空类和抛接类，所有的杂技表演都是这五大种类的单个表现或灵活组合。现代杂技节目对创新的要求越来越高，融合的元素越多、越灵活，越能给人以视觉冲击。

现阶段的机器人，一个元素都做得不够好，更不要说融合。湖南省杂技家协会副主席赵双午老师曾说过杂技表演中有“三大关系”：人与自身、人与道具和人与他人之间的关系。这一理论用在杂技机器人的表演上也是成立的，即机器人与自身、机器人与道具和机器人与机器人之间的关系。

机器人与自身的关系是指在不用道具且独身一人的情况下，与自身身体发生对抗的关系。这需要这个机器人真的“天赋异

禀”，即一个人一台戏。

机器人与道具的关系是指表演中必须运用道具，靠自身对肢体、道具的把控所呈现出的关系，同时还需要给人以美感。

机器人与机器人间的关系体现在两个机器人及两个机器人以上的节目表演中，团体之间的相互配合，强调“团队精神”。现在用于舞台表演的机器人大都是一样的动作，以量取胜。

虽然机器人可以进行一些人类的艺术、运动等活动，如它们可以学画画、学唱歌、学写诗、学投篮、学跨栏等，甚至它们可以通过技术迭代达到较高的水平，但是想要成为杂技表演者还差很远。

将来会出现杂技机器人吗？

会，或者说历史上曾有过。古代文献中经常出现一种娱乐型的木偶机器人，它们不仅能唱歌、跳舞、奏乐，还会耍杂技，与真人演员无异。

这很容易让人想起那个大名鼎鼎的“土耳其骗局”。据说当时有一个坐在大机箱前的“土耳其魔法师”，它能自动而快速地下象棋。在维也纳皇宫的首次表演中，它就迅速击败了对手Cobenzl伯爵，随后名声越来越大，它还击败了一系列著名的挑战者，包括拿破仑和富兰克林。直到几年之后，这个骗局才被揭穿。原来机箱里藏了一名象棋大师，他用一个磁铁系统来跟踪对手的举动并移动自己的棋子。

由于技艺并没有流传下来，因此，现在已经很难推测古人是按照什么逻辑做出的木偶机器人。现在有一项新的研究成果或许可以回答这一问题。加州大学伯克利分校和英属哥伦比亚大学最新研究成果表明，运用强化学习方法能教生活在模拟器中的机器人模仿人类，并让它学会武术、杂技等复杂技能。

机器人究竟是如何学习新动作的呢？

简单来说，机器人是通过获取动作的数据来学习的。研究者对动作的数据进行了改良，然后将这份整理好的数据给机器人学习，最终使得机器人的动作像人一样流畅。

AI是如何揪出“网络钓鱼者”的？

许多苹果手机用户都反映自己的iMessage经常收到垃圾信息。但是由于苹果公司一贯尊重用户的隐私，它在服务器端无权也从来不读取用户发送的信息内容，当然就更谈不上通过内容对用户信息进行监管过滤了。



目前，iMessage所收到的垃圾信息多数是广告，诱导用户下载App。但还有一类比打广告更恶劣的垃圾信息叫“网络钓鱼”。

科技公司一般是如何应对“网络钓鱼”的？

相信大多数人对“网络钓鱼”应该都不陌生，这是一种在线身份盗取方式，攻击者主要通过欺骗性的电子邮件和伪造的Web站点引诱收信人给出敏感信息。目前，“网络钓鱼”的危害遍及全球，数据显示，2017年，中国、澳大利亚、巴西是最容易受到攻击的区域（高达25%~28%的计算机用户成为攻击目标）。

“网络钓鱼”的方式日新月异，各大科技公司，尤其是社交类公司都深受其害。那么这些科技公司到底是如何对付“网络钓鱼”的呢？我们依据不同的特点将这些对付“网络钓鱼”的方法分为如下3类。

1. 检测行为数据，和手机交互的是“你”还是“它”？

你所知道的是，从你注册Facebook的那天起，Facebook会源源不断地向你提供你所感兴趣的各类社会、生活动态。但是你所不知道的是，在后台，Facebook会利用手机的陀螺仪来探测用户细微的动作，甚至包括用户的呼吸、点击屏幕的速度、握持手机的角度。这听上去似乎很恐怖，但是在我们看来，Facebook的做法别有深意。

实际上，每天在Facebook上进行注册的用户不仅有人类，还有千万台试图侵入社交网络的机器人。这些机器人入侵者会通过传播虚假信息导致混乱并损害Facebook的公众信任。面对这样的攻击，Facebook当然有责任对自己的社交网络进行人为管制。但是与数量惊人的机器人交战，仅依靠人工的力量是完全不够的。

所以，不管是探测用户呼吸，还是探测用户握持手机的角度，都是Facebook为了判断屏幕前的用户到底是不是真人所必须收集的行为数据。尽管现在网络犯罪分子所培养的机器人正在不断尝试模仿人类与终端设备的交互，例如故意放慢机器人注册信息时的处理速度，使之尽量与人类的正常速度接近，以此来逃避检测识别，但一个虚拟的机器人始终无法复制一个真实的人类与设备进行的物理交互。

2. 检测账户活动，太过频繁的活动很有可能是机器人

检测行为数据并不是阻止机器人侵入网络世界的唯一途径。与Facebook合作的初创公司Unbotify还能依靠AI根据一台设备上的账户数量及创建后的账户活动来判断账户是否为机器人账户。

这个方法的逻辑很简单，举个例子：如果一个账户在注册之后的1分钟内发送超过100个好友请求，那么你认为这是一个正常账户吗？肯定不是。但这样的账户在社交网络中可能很多，因此我们需要依靠AI来对其进行标记。另外，如果一个“网络钓鱼者”想要让更多的“鱼儿”上钩，那么他必须多下诱饵，因此，他通常会在一个设备上登录多个“僵尸”账户。但是通常，正常用户的做法是在多个设备上登录同一个账户。很明显，这两者的账户活动情况是大相径庭的。那么，AI也可以根据这些异常的账户活动检测出该账户是否为机器人账户。

3. 检测内容本身，关键词汇暴露“钓鱼”的本质

YouTube的评论区，是“钓鱼”内容最泛滥的地方之一。因此，YouTube也使用了AI管理检查工具来筛选恶意评论，以对付网上大量的“钓鱼”信息。

YouTube的AI是如何做的？它与前面两类检测方式大有不同，因为它是基于自然语言识别的一种检测方式。YouTube的AI的主要工作是自动标记那些它判定为会危害对话内容的评论，但它并不会自己做决定将其删除，而是让人类做最后的决定。

在AI的前期训练过程中，人类标记对它来说很重要，因为这是建立AI是非观念最关键的一步。可是“网络钓鱼者”却会在论坛上输入一些负面信息，同时告诉AI这些信息没有问题，以此来达到欺骗AI的目的。而这将会对AI检测内容的有效性形成极大的威胁，因为只要数据足够多，AI就很容易开始颠倒黑白。

我们可以看到AI拥有各种各样识别恶意“钓鱼”信息的能力。但是反过来，我们也会发现“网络钓鱼者”其实同样也可以利用AI，让AI通过学习来预测科技公司的识别方式，从而达到把其被检测出的可能性降到最低的目的。

因此，这场对抗赛的输赢其实还没定。

我们必须认识到，尽管像AI这样的技术将是未来网络防御的基石，但是同时犯罪分子也在盯着这些技术。

1. 对抗性样本，AI尚未解决的软肋

很多网站甄别“网络钓鱼”功能的实现首先是建立在AI的深度学习功能建模之上的，然后通过模型是否匹配来对良性和恶意信息进行区分和识别。问题在于，AI系统建模所依赖的神经网络是可以被对抗样本所干扰的。

也就是说，恶意软件只需要改变部分代码，生成对抗样本，就能够引起监测系统的识别错误。通常情况下，黑客只要改动不到1%的字节，就能躲过监测，而这并不会影响其入侵功能。所以，实际上，大家并没有深刻地认识到机器学习的弱点，其实包括深度学习在内的很多机器学习模型，普遍都已经表现出了对于对抗样本的脆弱性，而目前科研界对此并无合适的解决之道。

2. AI精准检测的背面是AI精准犯罪

另外，AI还有助于犯罪分子更了解自己的目标对象。2016年，《美国黑帽》的一篇论文提出一种名叫SNAP_R的递归神经

网络。这种神经网络是被动态地从目标用户的时间轴上的帖子中提取出来的，积累了大量的用户个人数据。因此，它可以在Twitter上对特定的用户推送“钓鱼贴”，提高“网络钓鱼”的“上钩率”和“精准度”。

近年来，鱼叉式“网络钓鱼”已经成为攻击者越来越有针对性的攻击方式之一。罪犯通过收集信息对网络中的关键人物进行个性化处理并组织有说服力的电子邮件，引诱用户提供机密信息。有61%的受访者透露自己曾经历过鱼叉式“网络钓鱼”。而对比手动鱼叉式“网络钓鱼”和“批量钓鱼”，使用了AI的先进式鱼叉“网络钓鱼”开始变得更加有效。

3. 区分人机的图灵测试实际没有用吗？

在登录网站的时候，用户一般会通过回答问题来证明自己是人类，而不是虚拟化的攻击者，最典型的例子就是“12306”订票系统那种找图操作。这其实是图灵测试的一种。但是近年来，黑客利用AI学习图像，使虚拟攻击者对这种问题的破解率接近90%。

2012年，有研究人员尝试用机器学习来进行安全攻击，在破解简单验证码的实验中，深度学习的精确度就已达到92%。2017年，一项名为“我是机器人”研究也揭示了如何破解最新语义图像验证码的方法。这也就意味着黑客可以进行未经授权的访问，并借助进入访问状态的账户进行更大规模“钓鱼”信息的散播。

因此，就目前的视觉身份验证方式来说，身份验证还有待改

进，加强对虚拟攻击者的防范，或许引入声纹识别是更好的办法。

AI的对抗战还在继续，AI如何更有效地防范虚拟攻击者的攻击，各界还在研究当中，例如，可以通过完善概率模型，对“网络钓鱼”进行反预测、反侦察；增加分类器数量，提升分类器质量，使“网络钓鱼”难以规避监测。

与其苛责AI给犯罪分子提供了更便捷的犯罪手段，不如想办法在这场黑白对抗战中取胜，毕竟犯罪分子已经把AI当作武器，我们也没理由把自己创造出来的“利器”拱手让人。

机器心理学家为什么可能会是人类最后一个职业？

《我，机器人》是美国著名科幻作家艾萨克·阿西莫夫一生中最重要的的一部中短篇科幻小说。小说描绘了机器人的智能水平在经历了一步步发展之后，最终“挺立于人类与毁灭之间”。更重要的是，小说中不但有机器人，还有机器人心理学家苏珊·凯文。

在实际工作中，机器人会出现各种各样的意外状况——这也是机器人心理学家需要应对的。有趣的是，机器心理学家要做的并非排除机器故障，而是要理解和解决机器人的“心理问题”。

回到现实中，机器心理学可以算是一个对应人类心理学的新学科，它可以让人们了解机器人的心理，和机器人更有效、更便捷地交互，最终使机器人更好地理解和服务于人类。

虽然现阶段还没有机器心理学这一学科，但随着AI的发展，机器心理学很可能会成为心理学的重要分支，这是为什么呢？

1. 机器心理学家是AI发展的“脚镫”

但凡说起AI，人们总是会强调大数据的多样性和它的计算力更强大、更准确、更高效，但是实际上AI会引爆这个时代的根本原因，是因为它使人们的交互方式产生了根本的变化。

人机交互的方式从键盘、鼠标等“实物交互”变成语音、触摸甚至脑电波（人的意识）“非实物交互”。因此，人们产生了一种

难以驾驭的恐惧。这时机器心理学家的作用就凸显出来了，让机器心理学家去消除人机隔阂，让更多的人更好地接受机器。

在几千年前，人类就驯服了马，但是人类真正开始利用马是因为脚镫的发明。脚镫是什么？就是搭在马身上，供人上马踩的东西。甚至可以说，脚镫影响了人类的历史进程，因为脚镫作为介质改变了人与马“交互”的方式。

而机器心理学家在某种程度上就在充当脚镫这个角色，人机的交互由此而变得更加顺畅。从机器角度来说，如果有人了解它们是如何根据这些信息来学习和采取行动的，那么它们犯错的概率也会小得多。或者，当它们犯了错时，有人能做出合理的解释，这样就不会造成公众恐慌。

2. 心理学是AI的“干爹”

2018年苹果公司招聘，要求求职者除了懂计算机还要懂心理学。原因是人们在与Siri沟通时，会不自觉地向它倾诉。某种程度上，这种复合型人才也算是机器心理学家的“初始版本”了。其实计算机与心理学结合并不是时代发展的产物，它们从AI发展伊始就一脉相承。

美国著名的AI学者司马贺（赫伯特·西蒙自取的中文名）就是一名心理学家，他把认知心理学和计算机科学结合在一起。

早在1955年，司马贺就成功开发出“逻辑推理者”，使用机器进行人工推理。随后，他又研制出“一般解决者”，通过判断现在状态与目标状态的距离，不断进行反馈从而达到目标。

这种反馈机制正是以人类的思维方式为基础的，为计算机模拟人的思维活动提供了具体的应用实例。按司马贺的说法，AI就是计算机表现出来的那种如果由人表现出来就会被称为智能的行为，如认知。

机器心理学家：是船长也是水手

机器心理学家既要指明方向也要干实在活。

1. 机器认知与人的认知差异就是方向

在谈机器认知之前，我们可以先考虑一下，对人类而言，理解其他人究竟意味着什么？

我们既不会试图去估计其他人的神经元的活动，推断他们的前额皮质是如何连接的，也不会去与其他人的海马体交互，但我们在理解他人方面非常有优势。

认知心理学家认为，我们的社会推理取决于其他人的高层次模型，这些模型涉及的抽象概念并未描述它所观察行为的基础物理机制；相反，我们理解的是他人的心理状态，如他们的欲望、信仰和意图。这就是所谓的心智理论。

DeepMind最新研究提出“机器心智理论”，研究者建立了一个名为Psychlab的平台，构建了一个心智理论的神经网络ToMnet，并通过一系列实验证明它具有“心智能力”。

实验结果表明，当识别的东西有干扰时，人的注意力会被分

散，而机器的注意力则较为集中。因此，若要说机器人具有“心智能力”，那这种“心智能力”和人的认知差异很大，而正是这种显而易见的差异，给发展机器认知心理学指明了新的方向。

2. 使机器心理咨询师像心理咨询师

上文提到的苹果公司的招聘，透露出一个很重要的信息，即Siri在某种程度上充当了心理咨询师的角色。在1966年，麻省理工学院的一位研究员Joseph开发了一款聊天机器人Eliza。Eliza引入了心理学家罗杰斯提出的个人中心疗法（Person-Centered Therapy），其作用就是在聊天的时候让聊天机器人尽量引导人倾诉。

其原理很简单，基本依赖于模式匹配和脚本答案，但目前最好的聊天机器人也无法让人感觉它是具有稳定性格和情感、活生生的。这就涉及如何让机器人的语言和行为更具有个性问题。

普通人学习心理咨询的过程如下：

- 学习心理学的基本知识、基本的助人技巧。
- 学习、记忆和理解各种理论对于人格结构的假设。
- 学习各种异常心理的成因、症状、干预方法。
- 被治疗、被督导。

学习者也需要学习大量的模型与案例，其实整个学习流程是有章可循的。这时机器心理学家可以将学习的经验复制给机器，

不断调整机器的学习模式，以适应学习过程。

这样做的好处是，机器可以被用来加强和测试适用于人类的认知过程，最终得出最优的学习方式。

3. 使机器更像人

随着AI的发展，积极对其进行研究的心理学家会帮助公司开发出效率更高且更像人类的机器，不管是在性格方面还是在伦理道德方面。这种特性会让这些机器对消费者更有吸引力，因为他们更愿意与这样的AI互动。

但有一个恐怖谷理论似乎与其相悖，该理论认为，随着机器与人的相似度越来越高，人们对它们的好感最初会逐渐上升；但当机器与人类相似到一定程度后，人们对它们的好感会急剧下降，甚至转成厌恶。

但是在未来肯定不会这样，因为机器心理学家的广泛存在必将基于这样一个事实：AI是可以思考、学习和做出明智选择的。这无疑还有很长的一段路要走。

1950年图灵发明了计算机，这代表了当时最高的智能水平。现在人们司空见惯的智能手机只有手掌那么大，人们提到手机的时候已经将“智能”这个词省略了，更多人并不会将这种智能放在心上。在我们所熟悉的、每天都被计算机和智能手机所环绕的生活里，与人造机器的共生本就是生命体验的一部分。

随着AI的不断发展，人们会接受各种被机器心理学家调教的

机器。路漫漫其修远兮，未来可期。

阅片机器人为什么还没有被普及？

还记得你上次为了看一个X片的结果在医院排队花了多长时间吗？

在传统医学上，医生需要把片子对着灯光一张一张地看，费时费力，而且一旦疲劳，阅片的成功率会有所下降，导致判断错误。

不过这个问题很快可能会得到解决。2018年在一档AI节目《机智过人》中，一个阅片机器人几秒内看了300多张CT片。

如果你对于医学影像识别领域有所关注，那么你会发现2017年最有趣的事莫过于杭州健培科技有限公司与阿里巴巴iDST视觉计算团队，在国际权威肺结节诊断大赛LUNA16的世界纪录之争。最终，杭州健培科技有限公司的“啄医生”阅片机器人以91.3%的平均召回率夺得第一，并且创造了新的世界纪录。这场世界纪录之争，反映了我国阅片机器人这一细分领域的蓬勃发展。

事实上，从肺部影像AI诊断系统“天肺一号”的推出，到腾讯的“腾讯觅影”、阿里巴巴的“ET医疗大脑”纷纷入场搅局，再到阅片机器人“视诊通”大战84位影像科的专业医生、“啄医生”阅片机器人与15名三甲医院主治医师打成平手，方兴未艾的阅片机器人已经引发社会各界的热议，人们也对它产生了无限遐想。

阅片机器人真的能做到又快又准吗？

AI机器人凭什么能阅片？

随着AI在医疗领域的深度落地，AI机器人在大数据和算法技术的支撑之下，能够对MRI图像、CT图像、超声图像等医疗影像进行识别和处理，并且通过进行自主学习，不断提高处理的能力和效率，从而能够辅助医生进行阅片诊断。

一般来说，在唤醒机器人后，阅片机器人的运行会经过图像输入、图像分割与识别、图像分析和信息输出4个步骤。图像输入是指将张数不等的医疗影像输入进阅片机器人；图像分割与识别是指阅片机器人会对输入的序列图像进行算法分割与识别、标注病灶等；图像分析是指对病灶进行相关分析，包括磨玻璃的密度、实性成分占比等；信息输出指将所得出数据进行汇总，得出报告。

通过观察阅片机器人的运行路径，我们不难发现其具有高效率、客观性等特征，阅片机器人能够在提高医生诊断效率的同时，减少人为失误率。

阅片机器人的“爆红”为什么是在这个时候？

阅片机器人的快速发展，其实是与算法技术在此领域的成熟应用分不开的。阅片机器人的核心就是医学图像的处理技术，包含图像的去噪、增强和分割等，而这背后是算法技术的支撑。在查询诸多文献后，我们发现目前比较常用的算法有蚁群算法、模糊集合论、卷积神经网络等。

1. 蚁群算法

蚁群算法是人们在研究蚂蚁觅食的过程中得出的用来寻找优化路径的概率型算法。在医疗图像处理中，常常是基于区域内部灰度相似性和区域之间灰度的不连续性来进行图像分割的。因此，能够利用蚁群算法的“正反馈”效应及分布式的计算方式来完成对输入图像的分割。

2. 模糊集合论

待考察的对象及反映它的模糊概念作为一定的模糊集合，人们建立适当的隶属函数，通过模糊集合的有关运算和变换，对模糊对象进行分析。目前基于模糊集合论的图像处理方法包括模糊连接度分割法、模糊聚类分割法等。

3. 卷积神经网络

卷积神经网络由人工神经网络发展优化而来，是一个多层的神经网络，每层由多个二维平面组成，每个平面由多个独立神经元组成。卷积神经网络采用了局部连接和共享权值的方式，避免了对图像的复杂前期预处理。人们可以直接输入原始图像，并且它还具有良好的容错能力、并行处理能力和自学习能力，可以处理复杂的环境信息。据悉，“啄医生”采用的算法就是在中科大的安虹教授团队基于影像识别的3D卷积神经网络上进行优化的。

正是这些算法的成熟，才促成了这些阅片机器人性能的快速提高，也加速了它们的落地应用。

不过，尽管阅片机器人有着科学和强大的技术支撑，但要全面进入医疗应用阶段，还需要一些时间。目前不确定因素主要表

现在程序设定、数据学习和数据保护三个方面。

第一，程序设定上的失误，可能导致大规模误诊。

阅片机器人目前仍然做不到100%的精确判断，肺结节诊断正确率的世界纪录为91.3%，“视诊通”在进行甲状腺结节超声图像的性质判定时正确率也只有76%。

正如前文介绍的那样，支撑阅片机器人运行的是一套由人预设好的程序，程序的各个环节紧密相连，前后相继，最终完成阅片工作。而人的主观失误正是体现在程序的预设上，如果其中任何一个环节设定出现了纰漏，就会使得最终的数据报告出现偏差，从而会导致医生的误诊。此前强生CTC检测仪器Cellsearch系统就被爆出存在包括复位错误等共37个类别的问题，所幸在发现问题之前，仪器还未造成严重事故。

第二，急需更多有质有量的案例，提升机器人的学习能力。

阅片机器人实现自我学习的基础是大量的学习数据输入，学习数据的质和量都对阅片机器人产生重要的影响，学习的数量越多、案例越典型，识别的速度就会更快，质量就会越高。相较而言，目前医疗相关数据在质和量上都存在问题。其一是大量的医疗数据未进行电子化，其二是医院与医院之间存在“藩篱”，不能共享数据库。在《机智过人》节目中，杭州健培科技有限公司CEO程国华透露其阅片机器人学习的医疗影像资料为10万套以上，而同场竞技的主治医师都为20万套以上。再给出一个数据可能会更为直观，战胜人类棋手的AlphaGo一共学习了数百万人类

围棋专家的棋谱。

第三，医疗数据监管力度不足，个人隐私保护存在隐患。

阅片机器人进行诊断的医疗影像资料报告在输出给医生的同时，也通过信息传输技术保存在了机器生产商的云平台上。经过长时间的积累，机器生产商拥有的个人数据会非常庞大。而这就意味着，目前我国医疗数据监管乏力的情况之下，个人的隐私将受到极大的威胁。

在2018年浙江松阳警方破获的一起特大侵犯公民个人信息案件中，犯罪嫌疑人入侵某部委的医疗服务信息系统，获取各类公民个人信息达7亿余条。正如和美医疗控股有限公司创始人林玉明提倡的一样，希望国家通过对数据立法来保障个人的隐私安全。

目前阅片机器人所取得的成就，标志着我国在AI部分细分领域有了突破性进展。尽管有些问题尚待解决，但我们相信在AI的助力下，未来人们就医会更高效、更便捷。

如何成为安防机器人“头号玩家”？

2018年4月，安防机器人初创企业Cobalt Robotics完成1300万美元A轮融资，“机器人+安防”这个概念再次引起人们的关注。王石早在2015年就在万科园区局部启用机器人巡逻保安，但是从目前的市场格局上来看，市场正处于起步阶段，国内入局安防机器人的“玩家”并不太多，“头号玩家”更是没有。

据预测，到2020年，全球范围内的机器人市场价值将达到346.728亿美元，到2022年，专业服务机器人（包括医疗、国防、救援、安防、物流、建筑机器人）将会主宰机器人市场。近几年安防机器人势必会爆发，各大公司如何在这百亿元市场中占领一席之地呢？

一、解决三大技术难题筑起行业“护城河”

就技术角度而言，安防机器人主要有十大技术热点：导航定位、计算机视觉、目标跟踪、移动与运动控制、检查/巡检、算法、目标检测与识别、传感器、网络、人机交互。而这十大技术热点彰显的三个核心问题是移动底盘、机器视觉、智能语音。

1. 移动底盘

目前移动底盘产品相对成熟，可以应用到各种机器人身上，但是现阶段却没有成熟的SLAM（同步定位与地图构建）技术方案。针对机器人行走，大部分企业采用三维激光雷达SLAM方案。该方案比较成熟，产品也更加丰富。只有极少数企业采用3D

视觉SLAM方案，该方案适用于复杂动态场景，但对计算性能和算法要求很高。不管是激光雷达还是3D视觉，从技术角度上讲，它们可以在任何室内环境下应用。

SLAM方案中，有两个至关重要的难点。第一个难点是多传感器的定位。地图构建是静态的，但机器人走动时，地图构建是动态的，这就需要通过多传感器的搭配去定位。第二个难点是多传感器之间的融合协调。由于这涉及每个传感器的特性和数据处理，因此，在协调上难度非常大。

2. 机器视觉

目前，机器人视觉的应用是众多AI企业集中攻坚的热门方向。机器人视觉的核心功能包括更智能的空间与环境感知能力和视觉认知能力。理想情况是机器人被植入深度视觉后，可以更精准地实现自动三维地图重建，自主规划行走路线，轻松进行物体识别及人的身份识别等功能。但是在应用场景中，问题却不少。例如突然掉落的一个东西成为障碍物，机器人的反应速度跟不上，可能会突然停止，或机器人前面多几个人行走就可能会导致机器人行走速度变得很慢或直接失灵。

目前市面上的深度视觉产品主要是深度摄像头。按技术分类，深度摄像头可分为以下三类主流技术：结构光、双目视觉和TOF飞行时间法。前两者受环境影响较大，后者因成本高量产比较困难。

3. 智能语音

智能语音识别一直是最让企业头疼的问题。智能语音包含语音合成技术、语音识别技术和自然语言处理（NLP）技术三项主要技术。语音合成技术发展最早，基本没有太大的技术问题；语音识别技术在2012年被卷积神经网络应用之后，准确率大幅提升，虽然效果和体验还不够理想，但也在C端、B端得到了广泛应用；NLP技术虽然在搜索引擎中早有应用，但在人机交互领域仍属于浅层处理。

这里有几个问题需要解决，首先是歧义消除，即机器在相关语境下是否能识别带有多重含义的词语。例如，“灌水”既有往容器中注水的意思，也有发表无意义帖子的意思。还有一个跟机器视觉类似的问题，当机器前面有多个人时（这在社区显然是正常的情况），它是否依然能与人正常交流。这里有一个“鸡尾酒会问题”亟待解决，“鸡尾酒会问题”显示了人类的一种听觉能力，人类能在多人场景的语音或噪声混合中，追踪并识别至少一个声音，在嘈杂环境下也不影响正常交流。

从多模态交互的角度来看，如果在目前的智能语音技术上再去扩展视频、图片、运动数据等素材非常困难，那么只能一对一单线操作，现在还没有成熟的方案能将它们结合起来。

二、从理念导入实践，技术落地还需“软着陆”

中智科创机器人有限公司于2015年率先在国内开发安防巡逻机器人。它是户外全天候智能机器人，集高清摄像头、红外热成像、视觉激光导航、环境传感器、警灯装置于一身，具备自主导航、自主执行任务、24小时全方位音视频监控、异常情况自动报

警等多种功能。据统计，一个安防巡逻机器人可以抵得上2.4个安防人员执行巡逻任务。

2017年，青岛克路德机器人有限公司在华为全连接大会上推出了一款安防机器人。这台机器人具备自主巡逻、业主识别、紧急情况报警、险情预警、远程对讲、语音对话等功能，可以实现24小时自主巡逻。此款安防机器人已经在鑫苑集团旗下的鑫苑名家小区正式投入使用。

入局的“玩家”不少，“头号玩家”却还未诞生，于是各“玩家”开始从最擅长的细分领域做起。

1. 识别障碍物，保证“通行无阻”

2017年，浙江国自机器人技术有限公司发布了一款名为TIGER的“智能安防机器人”，这款机器人可以巡逻、发布危机预警、查漏补缺、进行车辆管理、进行人脸识别，还可以与人进行语音交互。但是它上不了台阶，最多只能爬一个小斜坡。

武汉工控仪器仪表有限公司在全力解决机器人上台阶的问题。该公司开发的“小卒一号”安防机器人头上装着4个摄像头，能在小区里自由行走，无须人工控制，可以提供全方位无死角的监控保障。只需充电4小时，它就可以不间断巡逻8小时。

机器人很难识别小台阶等障碍物，有针对性地开发算法和进行实地调试，才能保证机器人“通行无阻”。技术人员需要实时监控机器人的步态数据，并在计算机上绘成一条坡度曲线。只有当坡度数据与真实坡度曲线几乎完全重合时，一项小小的上下坡识

别算法才算调试完成。

2. 3D视觉SLAM方案：未来的发展方向

2017年深圳市大道智创科技有限公司推出了一代产品“e巡”机器警长，它是国内较早采用视觉SLAM方案的安防机器人公司。该警长配备4路高清夜视摄像头360度实时全景监控，同时融合热成像系统、超声波阵列、红外阵列、TOF深度相机和平面激光雷达多传感器采集环境信息。它还通过视觉测距、视觉避障、VSLAM定位与环境重建及人脸、车牌的识别与追踪进行视觉处理。在定位方面，它以多传感器融合算法适应多种环境，定位误差小于10cm，航向误差小于1度。

上文提到，大部分企业采取三维激光雷达SLAM方案。激光雷达基本上都是从国外引进的，虽然产品稳定性较好，但是成本很高。因此，3D视觉SLAM方案无疑是未来的发展方向。就现在的技术而言，我们通过视觉构建出来的地图场景，更多的是一种或稀疏或稠密的点云图，抓取的是技术的点的结构特征，用到线的很少，不太智能。

未来，企业需要将3D视觉SLAM和深度学习结合，对场景做到语义级别的理解。不仅要知道哪里是特征点、特征线、特征面，更要知道这是什么物体、什么场景，物体在大厅里还是小区外。基于场景的识别能力、理解能力是提升机器人智能水平的重点。

3. 从0到1，从场景定制到经验复制

安保公司Unity Guard System（UGS）的全资子公司Unibot于2018年推出了防盗机器人，该机器人可以在商店内巡逻，通过面部识别系统识别顾客并打招呼。如果公司进行事先注册，机器人还能呼唤顾客的名字，并能向店方发送提醒信息。

该项防盗技术可以广泛应用在安防机器人身上，社区相对商店来说是一个更为封闭的系统，而且对防盗的要求更高，相信未来使用防盗技术的安防机器人可以在社区实现大规模应用。

安防机器人的市场开拓依然要根据特定的场景进行，不管是产业园区、公安警用还是商业楼宇等，每个场景的安防需求都是不一样的。但是只要完成场景中的某一个功能，这个市场就能够打开，因为技术是具备可复制性的。

不同场景对机器人的需求差异很大，因此，在机器人的设计上，不管是功能还是外观，我们都要依据场景需求确定。

安防机器人虽然还在起步阶段，但是它的发展前景广阔，势必会引领新趋势。未来还将产生集合图像与视频精准识别、大数据挖掘、智能预警等多种技术的智能安防产品，为安防业的发展注入强劲的动力。安防机器人正在从静态发展慢慢走向动态发展，形成从被动到主动的发展趋势。我国物业管理正在迈入智能安防的阶段，在物业领域，安防机器人将“大显身手”。

陪练机器人来了，你打得过那个**AI**吗？

更高，更快，更强！在奥林匹克精神的指引下，人类一直在探寻着自我突破之路，世界纪录在“刷新—停滞—刷新”的轮回中一次又一次试探着人类的极限。与人竞争，其乐无穷。那么，与**AI**比试呢？

在首届进博会上，欧姆龙公司展示了他们的最新研究成果——乒乓球机器人**FORPHEUS**，这让人类在体育界与**AI**对抗变成了现实。以后，人类在“掂量”自己之前，可能要先问问自己，打得过那个**AI**吗？

从发球机到陪练机器人，**AI**进化速度惊人

其实欧姆龙公司在进博会上展示的乒乓球机器人**FORPHEUS**已经是第五代产品了。2013年，欧姆龙公司推出了第一代**FORPHEUS**，在两台摄像机的协助下，**FORPHEUS**可以在0.001秒内控制击球的时机与方向，通过预测乒乓球的运动轨迹，从而计算出应该让球拍以什么角度、在哪个点回击，当时它就具备了与业余选手连续练球的能力。

2015年，第四代的**FORPHEUS**被吉尼斯世界纪录认证为世界首个“乒乓球教练机器人”，此时的**FORPHEUS**已经可以发球和应对简单的扣球了。最新一代的**FORPHEUS**进一步优化**AI**算法和机械调试，将回球误差控制在0.1毫米之内。另外，它还增设了一个追踪人类动作的摄像头，用于评估人类的实际运动水平。它利用

机器学习技术对球的轨迹进行分析，判断对手的水平，调整自己的水平，争取能与对手一较高下。

人类从出生到能拿起球拍接上几个难度不高的球，至少需要5年时间。FORPHEUS同样只花了5年时间，已经可以成为一些业余选手的陪练甚至能做他们的教练。相比只能从单一角度发球的发球机，在AI加持下的FORPHEUS的进化速度让人惊叹。当AlphaGo战胜李世石时，我们看到了人类与AI在自我学习和计算能力方面的差距，如今站在球台旁的FORPHEUS则在运动竞技层面给了人类极大的压力。

要想战胜人类，陪练机器人还存在技术障碍

AI机器人在运动赛场上战胜人类可能只是时间问题，然而已经进化了五代的FORPHEUS怎样才能从一个业余选手的陪练机器人“成长”为可以和专业运动员过招的高手呢？让我们先从比较简单的发球机的工作原理和其中需要运用到的知识来说明。

一个稍微智能一些的发球机需要有助于处理风速、风向、温度、自身位置和接发球运动员位置等信息的处理器，来计算出发球的速度、旋转和方向。其中包括图像采集和识别系统、可编程处理器和传感器、博弈论（分析发球种类和接发球策略的博弈）及机械工程知识。如果是一个可以和人类对打的陪练机器人，它还要能移动、挥拍、躲避障碍……AI需要接收和处理的信息数量与发球机相比根本不是一个量级的，由于陪练机器人在多个领域还存在技术障碍，这也使得陪练机器人在短时间内还无法战胜人类。

1. 陪练机器人的应用场景、使用条件均有限制

体育运动领域对于AI来说是一片蓝海，但是目前，陪练机器人仅在乒乓球和羽毛球这两个项目中出现，且球技还远远达不到高手水平，这也使得其商业变现能力大打折扣。

鲍春来在《机智过人》节目中与羽毛球机器人过招，我们能明显看出在人类高手的前后场调动、大力扣杀、追身球等技巧面前，羽毛球机器人根本没有招架之力。而乒乓球机器人在面对擦网球、擦边球、高球及旋转球等偶发情况时也没有很好的应对策略。其实我们并不怀疑AI的学习能力，所有已知的运动技巧和比赛策略，像AlphaGo那样，AI通过自我学习都能掌握。而现在，陪练机器人需要提高的是它的“运动能力”，它也需要变得“更高、更快、更强”。

2. 单打独斗与团队协作还存在巨大鸿沟

现在的陪练机器人只能应用于单人对抗的场景，面对双人比赛或像足球、篮球那样的团队协作性项目时，陪练机器人面临着更高的门槛。

在电子竞技DOTA2，TI7比赛中，Open AI在一对一的对抗中轻松战胜人类高手；在TI8比赛中，Open AI在五对五比试中却没能延续胜利。虽然Open AI 1天的训练量相当于人类180年的训练量，在即时即地的反应也做得非常漂亮，但在比赛中，Open AI还是暴露出很多问题。在双方僵持或人类进行战略性避战的情况下，Open AI的团队协作就会出现分歧。

陪练机器人的技术技巧、运动能力甚至比赛策略都可以通过AI的自我学习和人类机械工程技术的突破而提高，然而团队协作这个作为人类社会属性的特殊存在，AI仅通过“自我学习”就能融会贯通并运用自如吗？在这方面，与其说AI还需大量的自我学习，不如说人类在提升AI团队协作能力方面还有很多工作尚未完成，对于是否应该赋予AI社会属性也还存在争议。

3. 从“机器”到“机器人”的进化尚未完成

都说FORPHEUS是陪练机器人，但是FORPHEUS的外观与人形相比还是相差甚远。FORPHEUS的外形就像是大型三脚兽，它将乒乓球桌一端包围在身体下。我们从球台对面看，FORPHEUS只是一台机器的外形，它并不具备传统意义上的“人形”。也正因为如此，FORPHEUS有很多限性，在面对高球、扣杀时，它无法像人类一样通过位移稍离球台，等到球速变缓再做出高质量的回击。

人类从爬行进化到直立行走用了几百万年，AI机器人从静止到直立行走甚至跑步躲避障碍不过短短十多年时间。相信FORPHEUS从趴在球台上的三角兽进化成真正意义上的“机器人”的时间不会太长，到时候陪练机器人会越来越多地出现在那些有激烈身体对抗的体育项目中。说不定当你和朋友要来一场篮球或足球比赛但人员不齐时，你们就会请上一个陪练机器人上场凑数呢。

与人比，还是与AI比？考量技术“温度”的选择题

不管你是否愿意，陪练机器人迟早会走进人类的生活，到时候陪练机器人可以充当的角色可能并不单单是你的陪练教练，也有可能是你要在赛场上力争战胜的对手。而你到底是喜欢和人比还是和AI比呢？对于这个问题，人类其实早已给出了答案。

无论是与人类自己比赛还是和AI比赛，人类的最终目标都是要争取胜利。而现在人类纠结的事情是，与人类比赛，如果输了，那么你可能感受到来自对手的鼓励、安慰或是嘲讽；如果赢了，那么你可能感受到对手的失望、沮丧或是祝贺。这些都是人类独有的社会属性激发的情感表达。而当你面对AI时，无论输赢，AI给你的结果最终只会归结为0或1。人类的进化就是在不断挑战自我的螺旋式上升中完成的，当挑战对象由人类变成了AI，人类的进化还有意义吗？

如今，对于陪练机器人在AI上的训练除了各个运动项目的技巧和策略，人类应该考虑如何让它们变得更有“温度”，如何让人类感受到来自AI的“温暖”，如何让AI在“科学”“技术”与“社会”之间彼此互相影响，互相促进。

如何战胜AI？给人类支个招

虽然现在的陪练机器人还很稚嫩，但人类其实一直在等待AI战胜自己的那一刻。可是高傲的人类从来不会甘于失败，当AI具备了战胜人类的能力时，人类是否还有翻盘的机会？面对一个不会紧张，不会沮丧，没有情绪波动，也不会体力下降同时还拥有高超运动技巧的AI机器人，人类的办法并不多，但人类也并不是毫无胜算。

首先，在技巧层面，人类可以更多地采用“假动作”来打乱AI的阵脚。现阶段，AI都是通过高速摄像机来记录球的运动轨迹并结合对人类动作、位置、神态等细节进行计算并给出回应策略的。AI的世界中只有“1”和“0”，耿直的AI可能难以对人类做出的具有欺骗性质的“假动作”做出准确判断。

伦敦顶尖AI实验室DeepMind曾对现有的AI学习能力有过如下评价：“现在的AI非常擅长识别图片中的物体，但仍无法很好地理解视频。”DeepMind的研究表明，AI对于有些集中在身体的某一部分或是比较快速的动作的识别准确率也不是非常理想。由此表明，AI在赛场上即时捕获的数据和信息对于AI的帮助有限，特别是当人类用“假动作”对AI输入的数据和信息进行“干扰”时，善于解答概率问题的AI，很可能会判断失误。

今后，人类日常训练的方向和重点可能要进行调整——如何将“假动作”做得更真。不过，人类的“假动作”在骗过AI的时候，可千万别把队友也骗了。

其次，人类需要开发新战术和新技巧。AI的所有运动技巧和比赛策略都是基于人类现有的程度通过大量的自我学习而掌握的。由于AI过度依赖于逻辑运算，在既有的运算规则下，当比赛中人类使出了之前没有用过的新技术或有新的比赛策略时，AI在短时间内是难以适应的。即便AI会根据人类以往的比赛来判断对手的风格，但人类不按套路出牌，不遵循AI计算逻辑，加上比赛中的随机性和一些偶然因素，如乒乓球比赛中的擦网球、擦边球，足球比赛中的立柱折射、反弹球等，这些都有助于人类战胜AI。

最后，如果人类第N次尝试还是无法越过AI这座“大山”，那么就把它电源掐掉吧。

3 传统行业“变脸”

在AI之后，为什么会有“兽工智能”？

电影《怪医杜立德》中的杜立德博士有一项神奇的本领——无须借助任何科学仪器便可与动物交流。起因是他听懂了鸟语，给一只受伤的猫头鹰拔掉了刺，后来这件事被各种动物知道了，随即他的诊所便给动物们看起了病。随着剧情的发展，一只猴子告诉他马戏团的狮子要跳楼自杀，于是杜立德博士赶紧跑到马戏团阻止它。

抛开电影本身的喜剧色彩，这个设定也是正确的。猫头鹰和猴子是电影中传话的关键人物，而现阶段人类对鸟类和灵长类动物的行为和语言是最为理解的。

电影照进现实，可能没那么容易

随着科技的发展，像杜立德博士一样与动物沟通交流变得越来越趋近现实。2018年，来自华盛顿大学和艾伦AI研究所的团队就合力开发了新的神经网络模型，该模型可以被用来理解和预测狗的行为。研究人员在狗身上安装了GoPro相机来记录狗的行为，并且通过在四条腿和尾巴上安装的传感器来传递运动数据。

研究人员通过对狗的肢体动作和在GoPro上记录的内容进行比对分析，能知道狗在什么样的动作下看到了什么，并对其行为进行预测。

在预测狗的行为之前，人们在狗的语言识别上已经进行了不少研究，甚至可以听懂它们讲话。

生物学家Slobodchikoff创办了一家名为Zoolingua的公司，该公司开发了一种算法，可以将土拨鼠的叫声转为英语。他认为对土拨鼠叫声的研究同样可以用到猫和狗身上，即收集大量的狗叫的视频，然后用这些素材来训练AI算法，用人工来标记每一种叫声和摇尾巴的动作表达的意思，最终将其翻译出来。

这是典型的有监督学习方式，需要的大数据首先要通过人工标注。由于采样的范围和机器内存等存在局限性，这种方式在翻译的准确度和丰富性方面尚有待提高。相比之下，为实现人狗沟通而设计的No More Woof耳机要更胜一筹。

No More Woof是由北欧发明与发现协会（NCID）开发的，它应用的是三个不同技术领域的最新技术的组合，即脑电图传感、微运算和特殊脑机接口软件。这些传感器是用脑电图录音机先录下狗的大脑内流动的离子电流造成的电压波动，再将其传到一部微型计算机上，对它们进行解释。

但这里依然有几个问题尚未解决。首先，脑机接口至今尚未取得突破性进展，而使用在狗身上的外置脑机接口在识别精度上可能依然达不到预期。其次，以上两种翻译机均只能识别一些简单的内容，如“我饿了”“我很累”“我想出去散步”等。最重要的是，现阶段的技术还不能实现电影里的那种交互，你或许能听懂狗的叫声，但狗却不能听懂你的话。

从技术上说，识别动物的表情、动作、叫声并不难，现在的AI算法可以快速对这些信息进行识别并分类，难的是正确解读这些信息所表达的含义。人类对于动物行为和语言的认知程度并不相同，那么对动物的语言乃至行为的研究又该向何处发展呢？

与其一味解读，不如认它做老师

我们对动物的了解可能比我们想象的要少得多，因此，向它们学习，也许是现阶段AI的突破点。我们将其称为“兽工智能”，那么“兽工智能”到底是怎么回事？

1. 动物可以帮助AI变得更聪明

研究证明，当狗看到不同的物体时，身体的反应是不同的。狗能清楚地表现出视觉智能，能识别食物、障碍物、人类及动物。这些不同反应我们在现实生活中也很常见，既然狗能识别出不同目标，那么神经网络也能被训练成这样。

因此，在后续实验中，探究人员进一步尝试把神经网络训练得“像狗一样”，并让其在不同场景中识别物体。通过这种学习方式，神经网络可以识别出室内或室外等不同场景，并且能够理解怎样在不同场景下行走、路线怎样规划更合理，而这一学习原理还可以应用到机器人自主行走领域。

训练机器让神经网络懂得如何智能识别物体是一项艰难的任务，因为它需要大量先验知识。机器学习起来相当费时间，如今狗知道这些规则，那么人们就不必再从零开始来训练神经网络，通过观察狗的行为，神经网络就能掌握这些规则。

众所周知，如今AI最火的两种应用，一是智能语音，二是机器视觉。但人类目前在触觉、味觉和嗅觉这三个领域中的进展特别缓慢，尤其是味觉和嗅觉，属于小众需求，目前只有一些特殊领域的机器会用到，因此整个研发投入都不足。

就机器识别而言，人是视觉动物，所以视觉比较靠谱。而动物之间的识别，不一定靠视觉。嗅觉、味觉、触觉、听觉都有可能。例如老鼠由于皮层无褶皱，神经元分层也比灵长类动物少，视觉皮层占的比例非常小。因此，它只能看到鼻尖前面的一点地方。但它的嗅觉很发达，可以通过嗅觉来识别其他物种。

2018年7月，一位尼日利亚科学家研制出了一种新型AI芯片，可以让计算机拥有嗅觉识别能力，如识别出爆炸物气味等。但是研究进展缓慢，还不能大规模商用。

一旦被商用，机器嗅觉需要将AI先看到物体这一步骤省略，直接透过现象窥视隐藏物体的本质。如此一来，未来在公共场合出现的毒品将无处遁形。

2. 动物可以开拓全新的学习机制

一只被研究人员在日本发现和跟踪拍摄的野生乌鸦找到了坚果，需要将坚果砸碎，可是这个任务超出它的动作能力。它发现把果子放到路上让车轧过去，就可以完成“鸟机交互”了。虽然坚果被轧碎了，但它到路中间去吃坚果是一件很危险的事。它又开始观察，最后发现了过马路要走斑马线这一逻辑复杂的机制。于是，乌鸦就选择了一根正好在斑马线上方的电线，蹲下来了。它

把坚果抛到斑马线上，等车子轧过去，然后等绿灯亮。这时，车子都停在斑马线外面，它终于可以从容地走过去，吃到了地上的坚果肉。

这个过程既没有经过大数据训练，也没有经过所谓的监督学习，但是乌鸦却解决了世界顶级的科学家都解决不了的问题。这是与机器学习、深度学习完全不同的机制。

我们如何让AI像乌鸦一样聪明呢？这里就用到了搜索进化算法。

首先，我们知道，建造一个像乌鸦的脑子一样强大的计算机是可能的——我们的大脑就是证据。如果太难完全模拟，那么我们可以先模拟出乌鸦的大脑的演化过程。

这种方法称为“基因算法”。它建立一个反复运作的表现/评价过程，并且以能否生养后代作为评价。一组计算机将执行各种任务，最成功的将会“繁殖”，把各自的程序融合，产生新的计算机，而不成功的将会被剔除。经过多次反复后，挑选出最接近样本的甚至超越样本的计算机。

这种方法的缺点也很明显。人类主导的演化会比自然快很多，演化需要经过几十亿年的时间，而我们却只想花几十年时间，因此，现阶段具备的技术优势是否能使模拟演化成为可能还有待商榷。

电影中的杜立德博士因动物而重拾了快乐，而现实生活的残酷之处在于，AI的受益方是人，“兽工智能”的受益方也是人。

AI拍照为什么是一个骗人的把戏？

手机厂商在“AI拍照”上纷纷“大展身手”，我们通过手机AI可以拍出好看的照片。

2018年第一季度结束，国内手机厂商的年度旗舰机型已经摆上货架，各手机厂商都给自家的手机贴上了AI拍照的“标签”。

如今，在品牌策略上各有不同的手机厂商都纷纷投向了“AI拍照”的怀抱，在手机拍照方案上从侧面竞争走向了正面交锋。

不过在我们看来，押注“AI拍照”可能并不会是件称心如意的事。

被动的出发点：过度营销后，手机厂商拉起“AI拍照”旗帜

回顾智能手机在拍照领域的竞争，智能手机刚兴起之时像素的竞争是第一个引爆点。

后来，像素的竞争走进了死胡同，双摄像头成为手机的主流配置，各手机厂商开始走差异化竞争路线。华为引入徕卡镜头，提升品牌品质；vivo主打“自拍+美颜”；OPPO面向年轻群体，主打年轻人选择的拍照手机。各个厂商都在为自己的手机拍照性能“背书”，旗舰机型的拍照水平也逐渐拉近。

智能美颜、人像补光、超级夜拍、大光圈虚化.....媒体广告、节目赞助、户外大牌甚至街边手机店里的喇叭都在宣传手

机“强悍”的拍照性能，过度的营销透支了消费者对手机拍照性能的期待。

在手机摄影硬件的发展上，受限于手机的空间结构，决定相机成像质量最关键的要素——感光元件尺寸，很难有进一步的增加。目前华为P20使用的感光元件尺寸仅为1/1.73英寸，已经是智能手机中感光元件尺寸最大的。手机摄影硬件的更新迭代在短期内难以有质的飞跃，还有什么能成为下一个手机拍照的引爆点？手机厂商都在盯紧AI。

自从谷歌在Pixel 2代手机加入了自主研发的AI单元，在手机拍照上依靠单镜头获得不俗表现后，用AI对手机拍照进行革新，改善拍照质量似乎已经成为手机厂商的共识。

“AI拍照”的魅力在于手机可自动识别各类环境，并相应地调整相机设置，从而达到最优效果。当女朋友要你帮她拍照时，你不需要再担心自己的技术不好而被她嫌弃，AI的加持会让你的拍摄水平上升一个层次。但“AI拍照”能有这么神奇吗？

早在2013年，就有手机厂商宣称1300万像素的摄像头结合自家研发的“黑科技”能让手机具有媲美微单相机的拍摄效果。实际上到今天为止，手机和微单相机之间的拍照性能仍存在差距，一个主要的原因就是前文提到的感光元件，微单相机内置的感光元件尺寸更大。

目前，“AI拍照”只是手机厂商在营销上的一个噱头，只是在营销上寻找新的“爆点”而已。“AI拍照”从概念到落地时间仍不足

一年，技术尚未走向成熟。实验室的检验标准和用户的检验标准不一，在众多的使用场景中，“AI拍照”能经得起用户的一次又一次考验吗？

难过的技术坎：竞争赛道缩短，跟风之后难的是技术沉淀

传统的手机拍照注重图像信息的获取，而在上一轮的竞争中，图像信息获取和成像质量主要靠手机镜头和感光元件决定。索尼在小尺寸感光元件中的一系列研发让手机拍照变得更加实用。

目前，市面上主流的手机品牌都采用索尼的感光元件，国内的手机厂商在这方面还没有太多的技术积累。硬件核心技术掌握在供应商手中，手机厂商可以发挥的空间有限。单纯依靠硬件来做文章，消费者难免对此产生疲劳感。

即使是通过镜头品牌、人群偏好、附加功能等形成差异化竞争，仍然不能避免产品本质渐趋同质化的尴尬。AI的介入，有机会让手机厂商掌握主动权。

其一，利用AI，手机的相机会针对性地给出最优的拍摄方案，弥补环境的不可控性，以此来获取更优质的图像。

其二，拍照不再是简单的图像获取。基于人脸识别、场景识别等技术，手机的相机中增加人像光效、智能美颜、创意自拍等功能，玩法更加丰富了。

华为通过自主研发的AI芯片，从底层技术上实现AI拍照从概

念到落地，但产品体验与宣传的有几分吻合还需要打一个问号，毕竟营销从来都是“扬长避短”，用户的反馈才最为真实。

vivo采用的骁龙660处理器在AI上并没有表现力，只能联合在人脸识别上有技术优势的供应商，强行贴上“AI拍照”的标签，而OPPO借助处理器中集成的AI单元与其实验室在拍照上的技术积累做文章。对于没有AI芯片技术的OPPO和vivo而言，“AI拍照”重点还在于芯片供应商是否会给芯片增加独立的AI任务处理单元。芯片的门槛很高，需要研发的巨大投入和长达几年的等待周期，让OPPO和vivo在短时间内再造一个AI芯片显然不可能。另一个摆在手机厂商面前的难题是，AI拍照同样也面临的人脸识别、动态捕捉、光影分析等模块的技术挑战。

2018年，OPPO和vivo在AI布局上的动作不断，OPPO成立研究院，AI成为其研究重点之一。vivo在与高通、MTK等公司建设AI平台，涉及芯片、计算能力等方面。

各厂商关于AI的技术突破会压缩产品上市的周期。按照厂商每年都更新旗舰机型的惯例，率先在技术上取得突破意味着下一阶段市场竞争力的提升。但技术的沉淀如“冰冻三尺，非一日之寒”，在众多的挑战面前取得突破已经不易，沉淀成为厂商的核心竞争力更是不简单。这条路，不好走。

“AI生态梦”：AI+摄像头，看到了未来但壁垒重重

尽管AI的浪潮滚滚而来，却鲜有手机厂商给自己贴上“AI手机”的标签。融合了AI元素如Siri语音助手的智能手机也从不

用“AI语音手机”来标榜自己。而在手机拍照领域，场面却大不相同。拍照是用户使用频率极高又最容易被用户感知的功能，企业在布局AI生态时，“AI拍照”自然成为其生态闭环中最好的切入点。

手机厂商理所当然地高举“AI拍照”的大旗。热潮背后是企业和技术壁垒面前深深的忧虑。找准了切入点不意味拥有成熟的技术，而AI+摄像头的创造力也不仅局限在手机拍照领域。手机厂商看到了更多的可能性，但技术壁垒重重，研发之路荆棘遍布，“AI生态梦”看上难以实现。

iPhone X在摄像头中结合感光元件和传感器系统，让摄像头有了生物识别和动态捕捉功能。在AI的助力下，摄像头不仅能“看见”，更能“感知”。在未来，通过摄像头，AI应用程序甚至可以感知人的情绪、读懂人的心情。要实现这些，需要摄像头来捕捉、记录面部肌肉运动，并根据计算模型来分析面部表情，最终得出关于表情的动态结果。

现在很多手机厂商连生物识别和动态捕捉技术都还没弄清楚，要实现复杂的表情捕捉甚至更深层次的分析谈何容易。

总之，技术的发展和用户日益增长的需求会给行业带来一个又一个窗口期，但其中的挑战也有很多。对于手机厂商而言，他们应该把重心从营销转到技术上，打破重重壁垒，用更好的产品回报用户。

现在的**AI+**教育，为什么培养出来的可能是“考试机器”？

2018年，品途商业评论发布了2017年AI领域投资盘点，其中教育领域的融资成绩十分显眼。VIPKID、Keeko、作业盒子等一大批公司获得了高额融资，叉学教育更是拿到了2.7亿元融资。在AI教育风吹起来的时候，我们不禁思考AI+教育，究竟会带来什么变化？

现阶段的**AI+**教育，不过只是这三板斧

无论是叉学教育提出来的“智适应”系统，还是Keeko的教育机器人，纵观众多AI教育入局者，都打着三个类似的旗号。

1. 上帝视角，扫除知识点盲区

利用大数据，我们可以将学习科目的所有知识点进行拆分组合，通过测验来检测薄弱知识点，进行专项突破，从而能够快速掌握所有知识点。例如，叉学教育推出的松鼠AI系统和作业盒子基于AIOC打造的自适应学习系统都凸显了这一特点。它们首先将知识点进行拆解，然后通过较短的时间来检测学生对知识点的掌握程度，最后针对学生的盲区进行专门的视频讲解、专项练习、专题测试等。

2. 因材施教，一对一教学

通过算法技术及感应器，我们能够全面抓取、分析学生的学

习情况和掌握知识的能力，做到定制式教学，因材施教。例如VIPKID通过人脸识别技术、大数据实时跟进每位学生的学习情况，并且能够根据每个学生的学习进度和个人特点制订个性化的学习计划。考拉阅读推出的ER Framework能够通过衡量学生的阅读能力和水平，来进行个人定制阅读。

3. 实时沟通，“不下班”的老师

在AI的环境下，学生能够实时与AI进行互动，通过语音、文字等形式有效消解学习时间的边界，随时获取知识。例如Keeko教育机器人、小哈机器人可以基于图文识别、触觉识别、人脸识别等技能，根据相应程序与用户实现交流互动。

AI教育存在明显的软肋

尽管目前AI+教育已经取得了一些成果和成绩，但是总的来说，它仍处于一个低级的状态，甚至有些公司还打着个性推荐的幌子，做着一般化输出的“伪AI”。而这些AI教育项目暴露出的问题也很明显。

1. 缺乏思辨性，语义理解存在障碍

相较于有规律可循、强调逻辑推演的理工类学科，AI对于注重思辨性的人文类课程则难以进行教授。一方面，AI有大数据和算法上的优势，但它只适用于有标准答案的客观题，而人文类课程有很多无标准答案、需要灵活处理的主观题，对此按照设定程序运行的AI则显得用处不大。另一方面，目前AI对于数字和公式等已经具备了识别和处理能力，但是对于人类语言，特别是语义

的理解还存在较大的障碍。

以语文为例，语文中含有大量对于语义的理解和赏析，而解题的关键在于对上下文语境及相关背景的掌握，同时这也与个人的知识储备和经历相关。由于汉语本身的复杂性，在不同的语言环境下，同一个字或同一个词组，都可能有数种甚至数十种截然不同的含义。

因此，教授知识的困境一方面是由于原文语义的理解，会导致教学上的偏差。另一方面则是对于学生的作答，AI也无法给出合适的指导和建议。云知声AI Labs资深技术专家刘升平也表示，在设备和人的交流上，AI仍面临巨大的挑战。

2. 缺乏情感沟通，选择性心理影响教授效果

在老师传授知识的过程中，学生并非被动的存在，学生会因为一系列复杂的心理因素对老师所传授的内容进行选择接触、理解和记忆，从而影响老师的教学效果。虽然AI能用多样的表现形式来适配学生的兴趣，形成私人定制教学，但是AI目前还无法与学生进行有感情的互动和交流。

如前文提到的无标准答案的主观题，人类教师对此往往是采用大量的沟通技巧，来塑造学生的思维方式，而按照设定路径输出的AI在此则存在缺陷。在BBC基于剑桥大学研究者Michael Osborne和Carl Frey的数据体系对365种职业未来的“被淘汰率”的分析中，教师这个职业被机器人取代的可能性仅为0.4%，而其中一个重要的理由是人与人的互动能让学习的过程更加愉悦，而这正

是目前AI的缺陷。

3. 缺乏人的社会化培养，仅仅关注知识的灌输

在AI+教育的产品研发中，企业还停留在如何更好地让学生吸收知识这一层面，而对于个人的品德培育则还处在未开发状态。对于学生来说，知识的学习是一方面，更为重要的一方面则是个人人格的培育，“三观”的培养等，使学生完成人的社会化进程。

正如中国人民大学附属中学校长翟小宁所说，AI时代学习方式会发生根本性变革，但教育的本质不会改变。教育的使命是立德树人，只有立德才能使人获得幸福，使人获得福祉。

因此，现阶段的AI教育更多的是让我们看到一种可能性，让我们能通过AI对教育做一些调整，而不是全面革新。

现阶段的AI教育，培养出的很可能是“考试机器”

除了指出问题，我们更担心，如果AI使用不恰当，那么反而会引发技术异化，最终事与愿违。这集中表现在以下几方面：

第一，对于成绩和应试的极端重视及对于其他能力的忽视。

现阶段几乎所有的产品都只强调提高学生的学习成绩，而对于其他能力的挖掘则是一片空白。长此以往，在这样的环境和氛围之下，学生会极端重视学习成绩，最终也将造就“唯成绩论

者”。在生活中，常常提到的“书呆子”就是对此类人群的生动比喻。而在新闻报道中，缺乏品德教育的孩子在认知和行为上也更容易走向极端。

第二，生命“物化”后，引发对于生命的漠视。

在孩子还没有形成一定的基本认知之前，盲目依赖AI来进行教育，则会使得孩子形成生命“物化”的意识，这在一定程度上会让孩子对生活的态度更为冷漠。在电影《湄公河行动》中毒枭糯康培养的娃娃兵是真实存在的，他们从小接受残酷的训练，长此以往就形成了对其他生命的冷漠态度。

第三，技术偏见与商业意识形态的入侵。

AI的背后实际上是人输入的算法和运行的程序，而这也就不不可避免地带上了人本身的属性。

我们不难看出AI其实不可避免地会带上设定者的思维和意识，并且是以一种更为隐蔽和普遍的方式存在的。而生产商倘若为了自身的利益，在产品中灌输商业意识形态，那么依据马尔库塞“单向度的人”理论，技术会影响人对环境的感知，从而影响人的实践。如果AI本身存在偏见，那么就会使人们在看待世界时也形成偏见。

目前还处在初级阶段的AI教育，一方面，过分看重成绩，忽视了对学生其他能力的培养；另一方面，AI中存在的商业意识形态可能使学生只会接受，而缺乏批判意识。那么在两者的合力下，最终造就的可能是一个个只在学习考场表现优异的“考试机

器”。

AI为什么不能即兴作曲？

美国流行歌手Taryn Southern于2018年发表了一张名为*I AM AI*的新专辑，它成为人类历史上第一支正式发行的AI专辑。

主打单曲*Break Free*虽然达不到格莱美的标准，但是我们完全听不出它是由应用程序编曲的，它颠覆了普通人认为AI制作的歌曲会比较机械、没有情感的认识。

实际上AI巨头公司都在深入研究AI音乐，一些AI音乐作品已经达到“大师级”甚至“以假乱真”的水平。

2018年2月，第一部由算法创作的音乐剧*Beyond the Fence*在伦敦上演，并获得了较高的评价；6月，谷歌研发的机器学习项目Magenta通过神经学习网络创作出了一首时长90秒的钢琴曲；9月，索尼计算机科学实验室AI程序创作了一首披头士音乐风格的歌曲*Daddy's Car*，广受好评；百度公司的AI可以在分析画作之后，作出与之风格相对应的曲子。

AI在作曲领域取得了许多令人欣喜的成就，已经成为能与人类协同创作复杂艺术作品的得力助手。那么，AI实现作曲的原理和技术途径是什么？它有哪些优缺点和需要解决的问题？

AI作曲也在遵循“基本法”

音乐发展至今，所有的创新和突破都在竭尽所能地逼近人类极限，历代西方作曲大师无不在伟大的作品中留下探索音乐与新

技术融合之道的时代印记。

从基础理论设计与数学逻辑同构并进行符号化组织的角度来看，音乐是一门艺术，也有很强的可计算性，音乐背后蕴含着数学之美。常规的作曲技法，如旋律的重复、模进、转调、模糊、音程或节奏压扩，和声与对位中的音高纵横向排列组合，配器中的音色组合，曲式中的并行、对置、对称、回旋、奏鸣等，都可以被描述为单一或组合的算法。这从本质上决定了AI可以较好地应用到音乐创作上。

久负盛名的AI音乐作曲系统EMI就是通过对作品进行分解，以新的排列来复用这些结构并进行重组，从而获得不同风格的新音乐的。

其实，早在20世纪60年代，就已经有计算机与传统音乐相结合的尝试，直到广泛研究智能算法的热潮兴起之后，许多基于机器学习神经网络的开源项目才逐渐出现。AI有了长足的进步，越来越多的人开始关注这个科技与艺术相结合的领域，计算机音乐与传统音乐的桥梁才被逐渐架设起来。

虽然是即兴创作，也有一些作曲技术模型

AI作曲主要基于以下几种模型：分形音乐模型、马尔可夫链模型、遗传算法模型、人工神经网络模型和各种基于规则知识的改进或混合模型。

（1）分形音乐模型：它表明音乐完全可以通过数学算法进行创作。分形音乐模型是几何学在作曲中的应用，但是它只能创

作一些较为简单的作品。

（2）马尔可夫链模型：由于建模简单，我们用它可以即时产生新音乐，所以它一直被广泛用于商业程序上，也大量出现在互动音乐艺术家的作品和即兴演出中。它基于随机过程、概率逻辑的有限控制方法，尤其是使用马尔可夫链模型结合一定约束规则，在统计的基础上对音乐的未来走向进行概率预测与风格边界限制。

（3）遗传算法模型：将音符的排列组合进行编码，模拟物种繁殖过程，自动挑选出最优秀的作品。由于具有算法成熟和实现比较简单这两大优势，遗传算法模型得到了广泛的关注。但是，用遗传算法模型进行智能音乐生成，选取合适的评价函数是非常富有挑战性的工作，在在一定程度上限制了应用的快速发展。

（4）人工神经网络模型：这是当前AI音乐研究的前沿技术，它普遍采用具有深度学习能力的各种改进神经网络模型，来帮助AI模型学习样本音乐中的关键元素。

模型充分学习一系列人类已经创作好的音乐，提取和存储音高、音长、音量、音色、音程、节奏、调式、和声等关键特征，即可按照要求大量输出有类似特征的新音乐。例如用谷歌大脑作在线交互钢琴曲只需要它识别当前任意类型的少量音乐，就可以根据音乐的相符度进行预测，实时输出搭配音乐。

让AI进行即兴创作，难点在哪？

目前AI作曲领域研究的方向主要集中在深层特征的提取与应用和混合系统的构造上，但它还有以下几个难点：

（1）音乐的表示问题。音乐组曲过程较为复杂，现有的特征提取机制尚不能精确掌握一部作品的全部信息，如作品中与乐句、调性等相关的音乐信息一般体现不出来。如何精准地表示音乐的细节特征、提取音乐的深层逻辑、建立表层结构和深层逻辑的关系，是AI作曲亟待解决的基础性问题。

（2）学习与创造的问题。通过大量学习建立的作曲系统能否“灵光一现”，合理地突破预置规则，尝试使用不同方式创造性地做出一些风格独特，更生动、更具吸引力的音乐作品，以及如何进一步激发AI的创造性，实现从按照规则制作到突破规则创作的转变，是AI作曲面临的技术难题。

（3）创作作品的质量评估问题。人类对音乐作品的评判往往比较感性，因此，作曲系统中的质量评估机制是一个非常重要的部分，它往往会引导创作的方向，甚至最终决定作品的好坏。把人类的审美观用机器能够理解的语言描述出来，建立有效的评判标准是研究人员需要解决的重要问题。

“有味道”的智能公厕为何越来越没味道了？

下面介绍一个“有味道”的AI应用场景——厕所。

2018年1月10日，长沙智能厕所落户天心区贺龙南广场，吸引了众多市民前来体验。

春运期间，上海虹桥火车站对一些厕所进行了智能化升级，加装了“厕位智能引导系统”。

2018年3月24日，济南章丘区智能星级公厕改造完成并对外全面开放……

一直以来，“厕所”难登大雅之堂，但为何2018年却频频出现在人们的视野中？曾经被嗤之以鼻的场所，又为何成了热门话题呢？

看山不是山，厕所有了**AI**的加持多了一些“味道”

传统意义上，“厕所”泛指由人类建造专供人类（或其他特指生物）进行生理排泄和放置（处理）排泄物的地方。“公共厕所”则是指供城市居民和流动人口共同使用的厕所。现在看来，这些定义似乎都有些狭隘了。随着**AI**的广泛应用，厕所扮演的角色正在转变。

1. 形象大使

“唯厕是臭”的观念在人们心中根深蒂固，但现在我们走进一

间公厕，闻到的却不是恶臭，而是清香。这不仅仅是“空气置换装置”内外循环空气的结果，很大程度上还要归功于智能公厕除臭机。此外，那些感应水龙头和冲便器可以有效地节能节水，公共厕所又戴上了“环保”的帽子。厕所不再是以往那个脏乱之地，而是成了城市的形象大使，代表着文明的发展水平。

2. 科技体验馆

广州黄村街道负责人在介绍他们的智能公厕时说：“现在这个厕所已经成为村里人的骄傲，不少村民都会盛情邀请亲友离村时上个厕所再走。”这或许是智能公厕目前最真实的写照。免费再加上未来科技，能阻止我们进入智能厕所的可能只有漫长的排队了。

3. 休息室

除了标配的无障碍卫生间，公厕中的“第三卫生间”正在悄然兴起，这是专门为需要亲人（尤其是异性）陪伴的行为障碍者或行动不能自理者设置的卫生间。和其他公共场所一样，厕所内免费WiFi和手机充电站一应俱全，章丘的智能公厕甚至还设置了便民服务柜。未来，我们在公共厕所里待的时间可能会更长，因为它正在朝着一个个小型休息室发展。（免费书享分更多搜索@雅书.）

AI不仅改变了公共厕所的形象，而且还丰富了它的功能。我们不能简单地给智能公厕下定义，因为这一切还没有结束，未来公共厕所还有更多的可能。

AI在公共厕所领域掀起了一场革命，但它能彻底解决公共厕所的问题吗？

山还是山，智能厕所的“味道”依然没有变

可惜的是，智能公厕依然只是“烟雾弹”而已，在全面普及之前，智能公厕依然需要在“智能”上下功夫。

1. 传感器的局限

如果把智能公厕运用的技术分类，我们很容易就能得出一个结论：几乎所有智能设备都会用到传感器技术。以感应水龙头为例，其上安置了一个红外脉冲发射器，当我们的双手触发该发射器时，会返回一个信号给红外接收器，然后该信号通过译码电路转换成电信号开启电磁放水阀进行放水；当我们的手离开时，信号消失，电磁放水阀也将自动关闭。这个原理同样适用于厕所烘干机、感应便器冲水器和感应洗手液器，后来出现把出水、洗手液、烘干等功能集为一体的多功能水龙头也就不足为奇了。

这些智能感应设备节水节能的方式省去了人为操作开关的时间所带来的消耗，但这也正是它们的致命缺陷。首先，它们的上限很低。据统计，感应水龙头与传统水龙头相比能节约30%左右的水，但这已经是它的极限了。其次，智能设备并不能阻止我们浪费，与以往相比，我们的行为方式并没有任何改变，换言之，我们想消耗多少水还是可以消耗多少水，和传统公厕是没有区别的。因此，这样的AI与我们宣称的“环保节能”还是不能匹配的。

2. “死板”的AI

人脸识别厕纸机和扫码取纸是近两年突然兴起的功能，为了强制解决厕所浪费的问题，在各大一线城市被普遍应用。虽然打着AI的旗号，博取了大量的关注，也确实减少了纸张的浪费，但这些机器似乎没有那么人性化，有时甚至还有些“死板”。

首先，厕纸长度都是60厘米，虽然这是经过大量调研做出的决定，但确实存在有些人不够用的情况。其次，取纸需要等待十几分钟，这简直就是急性腹泻者的噩梦。最后就是那些残障人士和儿童，如果身高不够，那么如何进行人脸识别呢？

抛去追求新鲜事物的狂热，冷静下来我们就会发现，智能取纸机的智能程度其实比饮料零售机高不了多少，只不过是把投币按键换成了人脸识别或扫码罢了。

3. AI的“盲区”

除了环保，公共厕所另一个重要的理念就是干净。以现有的AI设备，要时刻保持便池和洗手池外围公共区域的清洁还相当困难。上文提到的广州黄村的智能公厕就有一键自动冲洗地面的功能，但效果似乎不尽如人意。首先，冲洗地面这种方式 and 节水本来就是相悖的，必然不是未来的发展方向。其次，机械冲洗的质量不能很好地保证。最后，冲洗的时机需要人来把握，并不是全自动的，这也是仍然需要雇用厕所保洁员的原因。也许未来AI在公厕上的终极应用是雇一个机器人来监控和清扫公厕。

至少现在，智能公共厕所的智能程度明显还不够。现有的智能设备优化可能限制了我们的行为方式，但我们不能彻底改变它

们。

如果说仅仅是研究突破缓慢，无法即时将技术进行转化都可以接受，但各大开发商绞尽脑汁，结果只是变着法子使用传感器，智能公厕的道路似乎就越来越窄了。其实没有必要再开发新的需求，能更好地满足旧的需求就已经很好了。要知道世界上还有超过10亿人无厕所可用，与其哗众取宠，还不如多建几座公共厕所更实在。

AI在舞台上配音时为什么也会唱“黑脸”？

在文字出现之前，声音曾经是人类唯一的交流工具。由于声音的传播距离非常有限，所以，那个时候人类的生存以“部落”为单位，人们之间的关系十分紧密。后来随着传播媒介的一步步发展，人们不再需要彼此近距离交流就能获得大量信息，于是人们开始怀念单一的声音带给我们的感觉，声音这种最原始的媒介承载着人类最充沛的情感。

2018年1月，世界首部利用AI模拟人声的纪录片在央视播出，这部名为《创新中国》的纪录片的解说词是由在2013年就“已逝”的声音完成的。原来在这个奇迹的背后是科大讯飞利用语音合成技术模拟出了我国已故著名配音演员、语言艺术家李易的声音。

科大讯飞强劲的语音合成技术让AI模拟的声音成功打动了李易老师的学生、朋友和家人。在AI自然流畅的语言解说中，人们似乎还能再见到故人的音容笑貌。那么科大讯飞这项语音合成技术究竟有什么神奇之处？它的操作过程又是怎样的？

这项看似很高级的技术理解起来并不复杂。

首先是输入文本，让机器模拟人对自然语言的理解过程。再对文本进行语言处理，主要包括文本规整、词语切分、语法语义分析，然后给出后续步骤所需要的发音提示。

其次是规划音段特征，如音调、音长、音重等，让机器可以

对语言的特有韵律进行处理，使机器模拟的声音更自然，并且使其更准确地传达实际语义。

最后根据前两部分处理的结果进行语音合成即可。通过这几个步骤，AI模拟的声音与人声已经非常相似，在某些情况下，我们也很难分辨机器人的声音与人的声音。

AI配音的边界可以延伸到哪里

这么惊艳的AI配音技术，它的边界究竟能够延伸到哪里？我们就此提出了AI配音的两个用武之地。

1. “粉丝经济”向AI配音伸出“橄榄枝”

“粉丝经济”已经成为现在文娱产业经济增长的主要支柱之一。随着2018年养成类偶像节目的火爆，粉丝对明星投入的情感越来越多，这个群体为明星付费的意愿同样水涨船高。既然明星的周边如此火爆，何不运用配合AI语音合成的VR、AR技术来打造虚拟“明星们”，让他们更真实地出现在粉丝的日常生活中呢？要深挖中国的粉丝潜力，如腾讯视频在《明日之子》上打造虚拟二次元偶像“荷兹”，粉丝听着现实中熟悉的偶像叫自己起床，还能让他陪自己聊天，虚拟真人版偶像或许更能得到粉丝的认可。

2. AI配音是音也是“药”

据国外媒体报道，有研究表明，年迈的夫妇可能因为一方丧偶而增大死亡率，这种现象被称为“心碎综合征”。这项研究由哈佛大学和威斯康星大学麦迪逊分校的两位科学家负责，研究结果

显示，男性丧妻后“全因死亡率”的概率增加了18%，女性丧夫后“全死因死亡率”的概率增加了16%。我们还可以做一个合理推断，在丧子或丧双亲的情况下，这种“心碎综合征”的表现也会存在。心理学家表示，要想修复这种创伤是非常困难的。但是AI或许可以做到，它能够利用过去已有的音频合成亲人的声音。如果心理医生说的话能够用亲人的声音来传达，那么也许可以帮助患者更快地走出阴霾。

AI在舞台上配音时是如何唱“黑脸”的

AI配音在解决问题的同时也会引发新的问题，如果把握不好，那么AI在技术的大舞台上就会成为唱“黑脸”的角色。导致AI在配音时唱“黑脸”的原因又有哪些？下面就为读者们一一列举。

1. AI盗用声音却被“无罪释放”

手机里的高德地图我们比较熟悉，但大家可能不知道它所使用的林志玲的声音其实部分是由AI配音技术后期合成的。而这样的语音合成过程是否必须要求声音拥有者本人提前去技术公司录制呢？

语音合成对音频质量并没有那么高的要求，利用海量的互联网音频也可以实现人声模仿。谷歌软件工程师在发表的论文 *Looking to Listen at the Cocktail Party* 中提出，采用全新视听模型可以在不同噪声中把重叠的人声分离出来，形成每一位说话者单独、纯净的音频信号。同时，科大讯飞也提出以全自动无监督的方法可以快速得到单个目标发音人的纯净音频。

阿拉巴马大学伯明翰分校的一项调查表明，如果给予AI的信息足够多，那么它可以生成任何以假乱真的图片或视频。但是随着现在个人的声音已经越来越成为个人身份的标志之一，对个人声音利益的侵害也同肖像权一样可能有损个人人格尊严造成财产上的损失。明星的形象是有肖像权的，如果他们的照片被他人私自用于商业，那么他们可以将对方告上法庭以维护自己的肖像权。但是目前在我国立法界及学界对声音权的保护却仍无统一结论。

2. AI配音干扰声纹识别

大家或许听说过声纹识别。一般来说，人的发声具有特定性和稳定性，虽不能达到指纹那样精确的程度，仍然有越来越多的国家已经把声纹鉴定作为辨认犯罪嫌疑人的重要手段。

但在GeekPwn2017国际安全极客大赛上，白帽黑客们却上演了一场与声纹识别的对弈。现场5组选手有4组根据《王者荣耀》中英雄妲己的声音样本，利用AI语音合成技术模拟妲己的声音并通过了“声纹锁”的验证，成功欺骗了语音验证系统。这意味着利用个人声音验证身份可能没那么靠谱。

声纹识别在现实中用途十分广泛，离我们比较近的有手机声纹解锁。另外，它也能用在智能家居产品及公共安全领域。但是当声纹识别遇上了AI语音合成技术，一场智能的博弈就开始了，一不小心就会打开个人隐私安全的“潘多拉魔盒”。AI语音合成技术越高明，持有该技术的人就能越轻而易举地闯入你的生活。

此外，在刑侦工作中，原本进行声纹分析可以判断说话人的性别、年龄、方言（生活地区）等特征，为侦查提供方向和范围。但AI配音的干扰要求刑侦手段需迅速跟上科技发展的步伐，否则声纹识别的有效性就会受到质疑，司法判决的过程也会变得异常艰难，这无疑为犯罪者提供了另一层“保护伞”。

3. AI又和艺术家们“杠上了”

AI在《创新中国》中配音的表现令人吃惊，不禁有人发问，AI配音如果在行业里被广泛应用，是否AI会取代传统的配音演员呢？“配音演员”由四字组成，不仅在“配音”，更在“演员”。2018年年初，综艺节目《声临其境》在展示了优秀演员的配音功力的同时，也让观众看到在配音间里，配音者不仅要提供声音，更要演戏。因为配音必须要符合剧本角色的情绪，甚至呼吸的频率都要对得上。

目前要建立机器的情感识别系统已经非常困难，机器深度学习需要大量数据进行量化分析，而人类的情感是最难以被量化的，更别说让机器产生情感并进行配音。配音演员和演员这两种职业本就不同，让AI取代传统配音演员独立参与影视剧制作几乎是不可能的。

不过，利用AI配音代替游戏配音和读书配音倒是不错的选择。与纪录片一样，此类配音效果并不需要调动太多的情绪，就算AI配音需要有好几种不同的感情色彩，其机器学习的量也在可控的范围之内，不会像影视剧配音那样复杂。

在AI配音这件事情上，有人拍案叫绝，有人忧心忡忡。“技术善论”和“技术恶论”的争论不会停止，但是只要控制的“阀门”还掌握在人类手中，一切就不会那么糟。

装上**AI**大脑的无人机为什么会担当“无人机警察”？

2018年，由国务院、中央军委空中交通管制委员会办公室组织起草的《无人驾驶航空器飞行管理暂行条例（征求意见稿）》在工信部网站上公布，并公开征求公众意见。其中有一条突破现行“所有飞行必须预先提出申请，经批准后方可实施”的规定，简化申请与批复流程，微型无人机在禁飞空域外，无须申请。也就是在微型无人机禁飞区外飞行将无须申报，但这并不意味着那些平时用来航拍或自拍的无人机就没有安全隐患。除了之前大家都担心的无人机“黑飞”、失控坠机、炸机等各类事故，无人机还有更大的风险，如在英国就已经出现一些不法分子利用无人机向伦敦等地的监狱投送毒品和通信设备的事件。另外，目前各大科技公司都在尝试将无人机应用到物流等领域，一旦无人机进入人口密集区，不管是来自系统或硬件本身的原因导致坠机、炸机，还是涉及民众隐私等问题，其风险都不可小觑。

如果真有一天，无人机成了“罪犯”或罪犯的帮凶，那么谁能制服它？

事实上，全球已经有一些团队提供了解决方案，如上海交通大学研发的无人机捕手和Teal无人机，它们都有一个共同的特征——搭载**AI**“大脑”。那么这些新成员有哪些绝活？它们能成为无人机领域的“警察”吗？

要做无人机警察，**AI**大脑传授了这三招

我们了解到，在AI助力下的无人机相较于一般的无人机，在识别与定位、轨迹预测、捕获无人机方面都有突出的特点，具体如下：

1. 识别定位，分清敌友

有了AI“大脑”的无人机，能够通过雷达驱动和基于视觉识别来对目标物体进行确认与定位。例如英伟达公司将无人机Iris搭载了其所研发的Jetson TX1机器学习AI模块，通过视觉识别和机器学习系统，配合无人机的光学传感器，在测试中能够准确进行定位、识别目标物体。相较于一般无人机单纯依靠GPS进行定位，AI无人机在识别的速度和精确度上明显占有优势。

2. 预测轨迹，胸中有数

在经过识别定位之后，AI无人机能够通过之前的案例学习和积累数据，对跟踪目标进行轨迹预测。在上海交通大学研发的无人机捕手的追逐测试中，AI无人机能够对目标物体进行轨迹预测，并且依据自身判断选择最合适的飞行路线，从而可以更有效地应对相对复杂的环境。

3. 选择时机，自主捕获

在经过上述步骤之后，AI无人机能够自主判断恰当的时机，对目标无人机进行抓捕。同时，AI无人机在进行抓捕的过程中，还能主动避开障碍物，例如AI无人机的躲避反应速度就在0.3秒内，减轻了撞击风险。在目前的抓捕方式中，主要有以“苍擒”一号为代表的电磁波干扰及以上海交通大学研发的无人机捕手为代

表的喷网。

差点火候，无人机的这个**AI“大脑”**还得继续修炼

经过多方查证，我们认为有了AI“大脑”的无人机的确在诸多方面都占有优势，但是其在图像识别、算法技术等方面仍然存在缺陷，还有进步空间。

1. 拼不过“同款作案”，挺不过环境干扰

目前的AI无人机对目标的定位主要是基于视觉识别进行的。获得国际空中机器人飞行大赛第一名的张陟曾指出，对基于图像识别进行工作的AI无人机，倘若采用颜色一致的无人机或采用逆光飞行，则会对AI无人机在识别和定位方面产生严重的阻碍。这也就意味着如果进行“黑飞”的是同款无人机，或在恶劣的天气条件下，AI无人机的误判概率会大大增加。

2. 战不过多人“协助”，斗不过“精妙”操作

由AI“大脑”指导的无人机是根据预先设定的程序运行的。如果AI无人机在对抓捕目标实施跟踪时闯入了其他设备，那么就会导致AI无人机依据设定的算法程序自动停止跟踪，转向对于闯入者进行识别，从而丢失了原来跟踪的目标。因此，若是“团伙作案”，AI无人机就可以相互掩护，利用AI无人机识别的间隙进行脱逃。正如上海交通大学导航与无人机技术研究所所长王红雨所说，AI无人机在实施抓捕时会有一定的延迟，而这个延迟是其工作原理决定的，如果此时“黑飞”的AI无人机飞手及时躲避，那么则会大大降低抓捕的成功率。

要做“警察”，光靠AI“大脑”还不够

目前AI无人机的AI“大脑”还有进一步的发展空间，但是仅靠AI“大脑”还不足以解决AI无人机可能出现的问题。要做“警察”，AI无人机还可以朝以下三个方面继续努力：

1. 在传感器上要下狠功夫

无论是激光测距仪、声波测距仪、GPS还是可见光相机，这些种类各异的传感器对于AI无人机AI“大脑”的激活都起着非常关键的作用。但同时，每一种传感器都有各自的优势和劣势，因此，研发人员要在传感器上花大力气进行产品迭代，来探索AI无人机传感器的最优配置。正如美国无人机研究专家亚当·罗斯坦所说，传感器的任何改进，哪怕只是重量和成本的略微降低，都可以帮助AI无人机实现更加自主的飞行。

2. 加上远程人工交互

AI除了能够建立庞大的数据系统，人工交互也是其重要的特点之一。目前的AI无人机由于算法还不成熟，想要完全独立适应于复杂的环境还有很长的路要走。我们应该能通过通信传输技术进行远程的人工交互，像医生操控AI外科手术机器人那样，可以在程序运行中实时加入人的指令，用以修正或改变程序的运行。AI“大脑”再加上人的大脑，就有希望极大地提高AI无人机的抓捕成功率。

3. 建立完整的无人机监管系统

实施AI无人机抓捕只是阻止AI无人机犯罪行为的一部分，相应的法律规范和惩处措施也亟待出台。同时，民航局还可以要求厂家对出产的AI无人机进行备案及在机身上配备相应的遥感芯片，就像管理航天飞机一样，建立起全国的AI无人机管理系统。这样AI无人机一方面能通过学习掌握全国的AI无人机数据，另一方面可以通过芯片进行识别，这样能够有效、快速地提高识别和捕获的成功率。

目前的AI无人机要担当“无人机警察”还很难，但是“以机治机”的思路是值得肯定并可行的。我们期待着AI无人机能够被纳入系统化的管理中，期待着AI“无人机警察”上岗的那一天。

AI+动物能否改变动物灭绝的局面？

我们先来看一组数据。1890年，我国野生东北虎的数量为1200~2400只，1930年约为450只，40年间减少了75%。到20世纪80年代，我国野生东北虎已基本处于灭绝边缘，仅剩14只左右，与1930年相比，减少了96%。

毫无疑问，野生东北虎的数量已经越来越少。此外，包括小熊猫、大熊猫、白颈长尾雉、金丝猴、白鹇等在内的珍稀动物都濒临灭绝，如果我们不保护动物，那么也许在未来，越来越多的动物将成为珍稀动物。幸运的是，随着AI的发展，AI已被应用到物种保护领域。

濒危动物，拿什么技术拯救你？

在2018年7月29日的世界爱虎日，英特尔公司与世界自然基金会宣布，他们将运用AI实施东北虎保护项目。实际上，越来越多的国家或组织开始利用AI保护野生动物。

1. 任何不谈物种检测的物种保护都是无稽之谈

物种检测最基本的做法是用运动传感摄像头自动拍照，然后将这些照片输入到一个模拟人类视觉皮层神经元之间连接模式的深层神经网络，最后用文字和数字对照片进行标注，如是什么动物、数量、性别、大致年龄、位置、附近的其他动物等。由于每只动物都有自己的特征，越精确的标注越有利于积累有效的数据。

来自奥本大学、哈佛大学、牛津大学、明尼苏达大学和怀俄明大学的研究人员开发了一种机器学习算法，它可以识别、描述并统计野生动物的数量，准确率高达96.6%。该研究召集了超过5万名志愿者。语料库收集了大象、长颈鹿、羚羊、狮子、猎豹和其他动物在自然栖息地的图像，对320万张图片进行了计算机视觉算法的训练。

这一研究借助群众的力量，最终准确且低成本地收集了野生动物的数据。在物种检测的同时，研究人员还实现了保护生物学及动物行为科学等相关科学的“大数据”积累。

如果需要对某些陆地动物的运动轨迹和活动范围进行检测，那么最好安装智能机器视觉处理设备，对动物的活动进行更加精准的监测和数据采集。智能机器视觉处理设备通过数据分析和识别成百上千个摄像头的图像，追踪动物的历史运动轨迹，最后形成精确的画像。

2. 最基本的是追踪运行轨迹

对鱼类和鸟类进行检测，总是离不开轨迹追踪。保护海洋生物的难点在于难以追溯，而且需要人们对海洋环境、酸碱性、所处航道宽度有所了解。Wildbook是一个致力于保护海洋生物的软件，它不仅可以从人工手动上传的动物照片中接收数据，还可以搜索图像和视频，甚至能查看可能对它的学习有用的一切媒体。

这种深度学习方法使Wildbook能够在不同的图像中精准地找到相同的动物，帮助研究人员更准确地使用有关动物健康、饮食

习惯、狩猎模式、种群大小和潜在偷猎者活动的数据。

集合群众的力量能快速积累数据，这种方法在数据难以获取的领域优势尤为明显。康奈尔学院和康奈尔鸟类学实验室联手研发了一个名为eBird的应用程序，他们已经拥有超过30万名愿者提供的3亿多个观察数据。为了保证结果的准确性，研究者将eBird收集的数据与实验室观察数据及从遥感网络收集的物种分布信息结合起来，最终机器学习便能预测某些动物栖息地的变化及鸟类迁徙的路径。

对物种轨迹进行追踪，其意义不仅在于保护动物本身，更重要的是我们可以衡量气候变化对野生动物的影响，为科学家提供有关气候变化是如何影响动物种群的宝贵信息。

3. 偷猎者放不下手中的枪，AI只能与他“正面开战”

在数千平方千米的土地上，仅靠人力找到所有偷猎者根本不可能。即使使用飞机巡逻，靠直升机或在动物行进路线上架设摄影机来侦察也不现实。摄影机只能拍摄单一位置的场景，直升机太贵，而且由于地面隐蔽性太强，很难在事情发生之前对画面进行精准捕捉。而一些贫穷的地区，政府往往无暇顾及保护动物，偷猎者数不胜数。

南加利福尼亚大学工程和计算机科学教授Milind Tambe博士带领小组对防止捕猎进行了研究，他们称这项技术为野生动物安全助手。

他们的数据来源主要有两个：过去哪个区域有情况、哪个区

域需要额外的保护。他们通过积累这些数据，对未来的袭击地点做出更准确的预测，最后决定在哪些地方加强防护。

这以低成本实现了利用AI防止捕猎，它具有广泛的使用意义。此外，我们还可以利用声音来绘制偷猎者的枪声所在位置的地图、利用无人机配备红外摄影仪巡逻等。由于AI是响应式的，因此，它会根据对手和它们的行为不断改进。一方面，AI可以预测对方的行为；另一方面，AI可以据此调整策略。随着数据的积累越来越多，AI就能实现在盗猎者盗猎前的精确捕捉。

AI无限好，问题也不少

从上述技术应用中我们能发现，AI依然存在不少问题。

1. 动物保护数据大多是用户上传的，质量难以保证

数据问题在AI与任何行业或技术结合时都会存在，但是有关动物保护的数据问题，针对不同的物种，会出现截然不同的结果。数据要么太多，要么太少。在机器学习中，数据并不是越多越好，机器学习会出现数据过拟合的情况，有用的数据越多越好。以一种采集鸟类声音的研究为例，数据多而杂，一般都是2000多小时以上的数据，其中包括其他鸟类的声音、风声、雨声、落叶声等，一般难以实现鸟声分离。

在对鱼类进行追踪学习时，常常会出现数据不够的情况，因为数据包括海岸线宽度、水的酸碱度、水温等不易获取的数据。由于对动物的田野调查非常耗费时间，因此，很多研究均鼓励用户或志愿者上传数据。这虽然提高了数据搜集效率，但也导致数

据质量参差不齐。

2. AI是一门技术，盗猎者也能研究

喜剧电影为了戏剧性，常常塑造出蠢贼等形象。现实生活中，很多偷猎者不仅不蠢，反而极其聪明，他们心狠手辣，经验丰富。还有一些人，仅仅是以打猎为乐，他们有钱有时间，纯粹是图个高兴。这些人并没有我们想象的那么好对付。

我们对AI保护动物的研究，一方面是对动物本身的追踪保护，另一方面是与偷猎者搏斗，而这通常是建立在偷猎者不借助高科技手段的基础上，可能一些愚蠢的偷猎者容易被甄别。但AI是一门技术，当我们在研究用AI保护动物时，偷猎者同样在研究用AI伤害动物。

偷猎者不仅聪明还有点张狂，他们会成立论坛，交流偷猎心得，还将捕杀野生动物的过程进行了记录。未来的偷猎与反偷猎，很可能上升到AI与AI的对决。

在使用AI保护动物方面，如果我们不建立切实可行的技术壁垒，偷猎者或许会利用AI加速动物灭绝。

保护野生动物除了可以利用AI，我们还可以建立一个生态圈，构建完整的生态链，利用大数据分析需要保护哪些动物、哪些植物、哪些环境等。微软就曾试图这样保护海洋生物，毕竟在相对密闭、相对可控的空间里，AI更能“大展拳脚”。同时，我们还可以对珍稀动物的基因进行分析、克隆甚至3D打印出新的个体等。

AI制作的表情包很搞笑吗？

现在如果没有几个有意思的表情包，我们在网上聊天都不太方便。

你是不是以为只有人会做表情包？

AI最近利用深度学习做了很多表情包，这些表情包的深度学习模型是由斯坦福大学的两个学生用超过40万个表情包训练得来的。其基本的框架是一个编码器—解码器图说生成系统，我们需要先将技能型CNN图像嵌入，然后使用LSTM和RNN进行文字生成。

从目前的反馈来看，AI表情包生成器表现良好，大多数人表示难以分辨人类制作的表情包和AI制作的表情包。仔细辨别之后，我们会发现AI生成的表情包中大概会有30%被认为是人类的作品。另外，对于表情包的搞笑程度也有评分，人类制作的表情包搞笑程度一般为7分（满分为10分），而AI制作的表情包最高达到了6.8分，可以说是很高的分数了。

AI懂幽默吗？

在科幻电影《霹雳五号》中曾有这样一段情节：一名逃跑的机器人有了意识，并对众人说它自己是有生命的。而男主角最终测试它所言是否是真的办法就是测试它能否听得懂笑话。在男主角讲完笑话几秒后，这个机器人突然发出了一连串笑声。这时，男主角认为它具备了自我意识。实际上，现在评判机器人是否具

有意识的标准并不统一，但是的确有很多人把机器人是否有幽默感作为判断机器人是否有人类思维的重要标准之一。

李开复曾表示：AI会在很多领域替代人类，但在娱乐领域不会，因为AI不懂什么叫幽默。就这个AI表情包生成器而言，它距离“懂幽默”还很远。

首先，AI还不会断句。在中英文中，不同的断句方式会使句子的语义发生很大的改变，断句方式的改变会影响一句话的幽默程度。

其次，幽默具有地域性。各国各地的文化各不相同，而机器学习笑点的模型不具有迁移性。让机器要生成一张带文字说明的图片很简单，但是要让其真正弄懂大众笑点却十分困难。

新的表情包是层出不穷的，人类开发的表情包极具创造性，而AI则只会“依葫芦画瓢”。

如何成为比人类更合格的表情包制造者？

AI的幽默品质还需慢慢进化，不过在进化的过程中，这里有两点建议可以帮助AI做一个更加合格的表情包生产者。

一是去掉训练内容中的糟粕。AI表情包生成器在训练中同时汲取了数字文化的精华与糟粕，很多训练数据都与咒骂、种族主义和性别歧视相关，这也是自然语言处理中一个普遍存在的大问题。在未来的训练过程中有关人员应该过滤这些内容。

对于机器学习系统而言，AI的学习过程就像婴儿的学习过程一样，AI通过观察与模仿所选择的系统行为类型进行学习。只要系统中某一部分是人为操作的，那就可能因为人们观察事情的角度或所谓的偏见而影响系统。这就像“酒与污水定律”：把一匙酒倒进一桶污水里，得到的是一桶污水；把一匙污水倒进一桶酒里，得到的还是一桶污水。一个节点的缺失或坏节点的进入，就会毁灭整个体系。

二是增加“网感”。在中国，表情包的“网感”很大程度上在于“模糊”的画质，中国表情包有一个很大的特色就是“模糊”。AI要成为表情包生产的专家，就必须对表情包文化进行更深入地了解。在互联网初期，表情包1.0时代，美国卡耐基·梅隆的一串ASCII字符是人类计算机历史上第一个“表情包”，在当时图片的视觉效果就是非常模糊的。

另外，在中国，表情包的主体使用人群是80后、90后，《武林外传》《还珠格格》《情深深雨蒙蒙》等影视剧都是这两代人的集体记忆，也是最能引发共鸣和传播效应的内容。而这些影视剧大多数都有10多年以上的历史，画质非常模糊。

从美术角度来说，一个东西画得越实，其明确表达的情绪就越理性，反之，越虚越模糊的图片表达的情绪就越主观。模糊的表情包更接地气，也会给人更多的联想空间。模糊，是一个表情包被无数人上传下载，压缩画质之后的结果，也是它在社交网络中“摸爬打滚”的印记，使用这样的表情包很有“网感”。因此，可想而知，AI制作的表情包要想在中国流行，画质上还得做一点特殊处理。

目前，表情包不仅仅是一个娱乐大众的工具，还具备更多的商业价值。“长草颜文字日常”的IP化为其开发者创造了巨大的利润，开发者仅靠收取版权费就能赚很多钱。很多广告主也有用表情包提升自己品牌知名度的想法，用表情包进行病毒式传播确实不失为一种绝妙的营销方式。小细节处藏有大商机，AI表情包生成器的开发“有趣又有金”。

4 机器进化

猪脸识别为什么比人脸识别更有趣？

如图4-1所示，下图的一窝小猪，将它们分散后，大家还能区分出它们吗？



图4-1

由于猪是多胎生动物，因此它们的长相十分相似。图4-1所示的是干净的小猪，人们勉强能区分出它们。在实际生活中，猪受生活环境等影响，很难识别。要实现个体化管理，首先要把它们辨认出来。

2018年3月，中山大学教授陈瑶生发布了“猪脸识别”技术，旨在攻破这一生物活体识别技术。我们曾定义“兽工智能”，即和动物相关而衍生的AI。“猪脸识别”这一技术的发展无疑是“兽工智能”相关领域的大突破，会让畜牧业发生巨大改变。

愿景很美好，烦恼也不少

“猪脸识别”技术早就不是新鲜事了。陈瑶生教授表示，有了“猪脸识别”技术，操作者只需要举起手机，对着某头猪扫一扫，就能得到猪的编号、猪的父母、品系等信息，甚至可以通过对猪体态和动作的识别来判断猪的健康情况。

但和所有的技术一样，“猪脸识别”这一“兽工智能”技术在发展过程中也或多或少存在一些问题。

1. 猪的变化按日计算，难以采集数据，难对比

猪的生长周期为110~120天，与牛270天左右的生长周期相比，猪在生长过程中的外貌变化可以用“翻天覆地”来形容，这就是“猪脸识别”比“牛脸”“羊脸”识别难度更高的原因。

没有现成的数据比对，我们几乎不能确定猪生长到哪一阶段面部特征或形体会会有显著的变化，因此，何时进行数据采集就更加没有依据。

即便猪既听话又干净，也常常以正脸面对摄像头，智能自动采集后庞大的数据如何存储及如何分析和调用数据，依然是我们需要探究的问题。

2. 面临更加成熟技术的竞争

“猪脸识别”需要依靠数据采集、数据的学习及最后的检索等程序来确定猪的身份，每一个环节都可能会因为技术缺陷造成误差。经过长时间发展的智能耳标更成熟，而且在确定猪的身份上更加精确。

耳标相当于猪的身份证，具有唯一性，我们既可以利用它进行动物日常信息管理，也可以实现动物产品的全程追溯。智能耳标也在向着AI的方向发展，显然这两门技术的决斗还尚未开始。

除了“猪脸识别”，还有“羊脸”“牛脸”“狗脸”“马脸”识别等。

上文提到“牛脸”“羊脸”识别相对于“猪脸识别”更简单，而它们的发展无疑会给“猪脸识别”带来新的启示。剑桥大学的教授开发了一种表情识别系统，该系统通过面部识别来判断绵羊的疼痛程度。如果将这项技术应用到牧场，用摄像头来监控羊群，就可以及时发现绵羊的生病情况。

狗是人类最忠实的好朋友，“兽工智能”在其身上取得的成果很多。2018年百度推出的“狗脸识别”不仅便于宠物的找回，而且能够扫描狗脸、进行喂食及收取快递与购买商品并支付费用等。虽然初衷是好的，但自由进出可能会导致坏人牵着用户的狗偷走了他的钱。

在未来，生物活体识别技术通过深度学习对动物面部特征、整体体态和行为特征进行识别，判断其品种和其健康情况，如哪些动物生病了、生了什么病、哪些动物没有吃饱、哪些动物到了

发情期需要配种等。

更重要的是，这些技术可以为食品安全、养殖户信贷服务，还可为更多的金融服务等商业应用提供决策依据。以前，如果农户只给一部分的猪投保，那么在猪死亡后，保险公司很难辨认这头猪是否投过保。而未来，农户可以依据猪的身体状况有选择性地投保，保险公司也有了对猪进行辨认的科学依据。

说到这里，猪脸难以识别的问题似乎还没有解决。以陈瑶生教授的实验场为例，母猪识别率为98%，肉猪识别率则为85%。但这个识别率并不算高，特别是对大规模的养殖场而言。

猪被识别出的概率还可以再提高吗？

可以。在这个看脸又看身材的社会，猪也不能例外。

先把脸部处理一下，如先变个形

当前使用AI实现视觉识别的原理基本上是一致的，即利用计算机神经网络的深度学习学到每一头猪的特征，然后我们利用深度学习的模型，针对测试数据集，得到每一头猪的概率，最后来判别出每一头猪。

最常采用的方法是把人脸的模型直接微调到动物脸上，但是微调在深度学习中更像是一个处理手段。

我们可以把迁移学习看成一套完整的体系。这样，我们能不能抛弃从之前数据里得到的有用信息，同时能应对新进来的大量数

据由于缺少标签或数据更新而导致的标签变异情况。

有研究者采用迁移学习识别马，这是一种全新的思路，对猪脸识别有极大的启发。人脸特征和动物脸部的特征本身差异很大，但是当动物的脸部变形之后，就会和人脸比较相似了。我们需要先找到人脸和马脸相似性较大的一个映射空间，使得人脸的训练数据可以被有效地利用起来训练马脸。

具体办法如下：首先我们找到人脸和马脸角度或表情相似的图片，在这里可以理解为五官之间的夹角相似；然后以相似的部位作为关键点，接着训练，以获得一个映射区间，得到了这一映射区间之后，与原来的马脸图片进行变换；最后采用人脸模型去调试动物脸检测的模型。

这项技术用在猪脸上原理也是类似的，最重要的是需要大量的人脸数据集来训练，鉴于人脸识别的快速发展，显然“猪脸识别”也可乘上这辆“快车”。

面部识别只是生物识别的一种，它可以与其他特征结合，整合多种识别模型，实现终极生物特征识别。

现实生活中，如果猪长时间低头站立和躺卧，就表明猪可能处在疲劳状态或精神萎靡；如果猪长时间抬头站立，就表明猪可能因为遭受到外界环境的刺激而产生恐慌。通过多种模态的结合，可以进一步辅助猪脸识别，提升准确性。

凭借个人的能力是没办法辨别猪的，更精准的识别与诊断，只能通过“兽工智能”来实现。

面部识别和语音助手加持，**AR**眼镜为什么没人要？

在极具科技感的概念赚足了眼球后，**AR**眼镜的消费级市场之路一直没有太多实质性进展。事实上，包括谷歌、微软在内的各大厂商在**AR**眼镜的创新上从未停止脚步，但不得不说消费级**AR**眼镜仍然处于概念阶段，对大多数普通消费者来说还有点遥远。

在CES 2018上，来自中国的“小厂商”Rokid公司发布了**AR**眼镜Rokid Glass，并称其产品是全球首款消费级**AR**眼镜，还得到了CES的官方认可（据媒体报道，该产品在发布前就获得CES 2018“最佳穿戴设备”和“科技创造美好生活”两个奖项）。我们有幸采访到了Rokid公司的负责人、技术极客Misa，或许能找出**AR**眼镜从概念走向现实的一些基本逻辑，也能探寻到为什么大佬级科技企业也无法做好**AR**眼镜的一些真相。

迄今为止，**AR**眼镜还没有真正的突破性产品

科技创新从小众走向大众，必须有现象级产品来代表整个行业实现突破，最终把行业带到一个新的层面，智能手机iPhone 4，智能电视小米、乐视的几个爆款产品都是如此。然而，这些年虽然出了不少**AR**眼镜产品，但各种各样的硬伤让它们离现象级产品还比较远。

很显然，直到今天，已经走向市场的**AR**眼镜仍然存在各种

各样的令人无法容忍的硬伤。智能手机仅用3~4年时间就风靡全球，AR眼镜发展了整整7年仍然摆脱不了各种缺陷，无法成为一个各方面都能达到应用级别的产品。

不过话说回来，这些硬伤，也给产品指出了改进的方向。

AR眼镜的“用户体验”还有“三座大山”

回首谷歌眼镜的失败，可以用一句“一厢情愿的产品无人埋单”来形容。直白地说，作为佩戴在鼻梁上、需要不断与环境主动交互的产品，AR眼镜的“用户体验”逻辑有它的特殊之处。

1. 硬件设计：不是太阳镜，但要像太阳镜一样方便

Rokid公司的CEO Misa反复强调一个观点：“丑陋的东西一定是错误的”。致力于智能语音交互的Rokid公司为了“不丑陋”，推出了像Pebble一样的产品，赋予了智能音箱足够的颜值和功能创新。而回过头看AR眼镜，除了受限于整体技术水平的显像技术，来自硬件设计层面的“丑陋”仍然十分明显，主要体现在以下几方面：

（1）分体式。Meta 2的AR眼镜要连接PC，新出的Magic Leap One仍然是分体式设计。这些产品甚至很难称为眼镜，有谁戴眼镜还要在腰上挂一个其他设备、手里拿着操作杆或干脆被PC拴住？加之“长相奇特”，一眼望去就很丑的产品会有市场吗？

（2）笨重。不得不说，谷歌眼镜之后出现的几个AR眼镜产品在造型和佩戴上都“开了倒车”。如HoloLens，与谷歌眼镜相

比，HoloLens重量大、佩戴不便，难以被应用到日常生活中。这样的产品就算有再多的功能也只会沦为“高科技玩具”，如果不加以改进，那么肯定无法进入消费级市场。

（3）携带不便。分体式、笨重及不可折叠都导致用户不得不把戴着AR出行“当个事儿”，在智能手机提高屏占比、降低厚度与重量以更方便携带时，戴着AR眼镜走却越来越费事。

把“不做丑陋的东西”的理念由智能音箱转移到AR眼镜后，Rokid的确做了不少硬件上的改进。例如他们制作了符合人体工学的眼镜鼻托，将电池后置，使整体重量降低，从而降低佩戴负担；去掉普遍的外层装饰镜片，以增强外观表现和轻便度；第一次让镜腿可以像普通太阳镜一样折叠，方便随身携带。

即便如此，这些都只是相对其他产品层面的，要说Rokid Glass整体观感有多漂亮也未必，在CES现场评测的一些国外媒体（如The Verge），虽然都给予了Rokid Glass肯定的评价，但在硬件层面要让AR眼镜像普通眼镜一样，它还有较长的路要走。

2. 软件交互：功能性并不是AR眼镜产品逻辑的首要出发点

CES颁奖给Rokid Glass主要是基于其在软件交互上的两个创新——视觉交互和手势交互，而这两个创新也代表了AR眼镜在交互上的改进方向。

（1）视觉交互。既然是AR眼镜，显然要能够与映入眼帘的事物尽可能进行更丰富、更创新的交互才更能赢得消费者的信赖，避免变成玩具。谷歌眼镜的失败就是因为它变成了没有实质

功能的耍帅“墨镜”。

在Misa看来，很多人模仿苹果公司只是模仿了它的成功，并没有模仿它的努力。苹果公司在做产品方面最大的努力，就是不断创造需求，而不是仅仅满足用户提出的需求。在这种理念下，Rokid公司总想去生产“有情感联系的产品”，而不是只生产“简单满足消费者的产品”。于是，Rokid公司做了AR眼镜中从没有人做过的AI面部识别。我们利用AI面部识别在生活场景中锁定一张脸，就能获取其公开的社交信息。

事实上，这个功能消费者可能根本就没有想到过，这也说明AR眼镜有着无限可能。在视觉交互上不能只从已有实现的角度出发，换个角度或许能带来不一样的产品效果。

（2）手势交互。消费级AR眼镜存在的意义必定是让生活更便捷，其判断标准在于戴眼镜一定要比“不戴眼镜用手机”更方便，否则AR眼镜在便捷性方面就不占优势。

目前，手势控制AR显像已经成为标配，在增加操作的丰富度与便捷度上，各厂家绞尽脑汁。Magic Leap One做了一个大大的无线手柄，像电视遥控器一样用户可以握在手里。它可以实现6个自由度的复杂精细控制，这种基于功能设计出的产品固然好用，但正如前文所说，对一款AR眼镜来说太多余。Rokid公司把自己的老本行语音交互搬了过来用在了AR眼镜上，产生了不一样的效果。在CES现场演示中，利用AR眼镜可以购物、查询天气信息及指定目的地导航等语音控制。

由此可见，对AR眼镜的操控交互来说，便捷比功能重要。从市场角度来看，消费者手里并非只有一款带有各种功能的智能终端。Rokid公司的AI面部识别和语音助手正是在尽可能地让AR眼镜更便捷。据了解，Facebook内部的Oculus团队也在通过The Verge的报道研究Rokid公司的技术方案及产品定义，这也说明Rokid公司的做法不仅是某个公司的创新，还暗含了行业发展的某些逻辑。

但是，要让Rokid Glass从体验版本走向大众未必那么简单，最大的阻碍是数据。在CES的样机上，Rokid Glass预先输入了20多张人脸图片，但在实际应用时，没人知道消费者会遇到哪个老同学，需要调用哪种识别能力来避免尬聊，这导致Rokid Glass在人脸识别上的数据需求“深不见底”。其他诸如购物场景也是如此，三星的Bixby2的AR功能也宣称用户举起手机就能识别环境中重要的地标、促销等场景信息，但业界并不看好三星的数据储备。

或许，Rokid公司可以借助阿里巴巴、腾讯等公司的开放API，或寻求潜在的数据合作伙伴来解决这一问题。但无论如何，这都不是一件简单的事。

3. 所见即所得，AR眼镜代替手机也未尝不可

可能很少有人会想到，被Facebook学习的AR眼镜会是中国团队做出来的。不过，Rokid公司的掌门人Misa却自信满满，在他看来，中国的智能设备、语音控制技术已经处于世界领先地位，在前沿领域产生领先世界的智能产品不足为奇。

目前，国内的场景数据和人口优势都特别明显，在深圳、杭州这些地方，已经有供全球采购的ODM和OEM技术解决方案。不管是技术、数量、成本还是迭代升级的速度，中国的智能硬件厂商都占据了世界主导地位。

此外，在当今美国AR眼镜公司纷纷投入语音助手的热潮下，中国本土语音助手研发具备先天的优势，转入AR眼镜研发中是顺理成章的事。

Rokid Glass带着人脸识别和语音助手出现在CES舞台上，从宏观角度来讲，是国家智能设备及语音技术双重环境因素催生的结果之一；从微观角度来讲，则是Rokid公司充分利用环境优势的战略布局的体现。

因为AR眼镜的诸多效用，不管是社交、购物还是信息查询等功能，都与手机的功能高度重叠。AR眼镜智能化到一定阶段后，不可避免地要进行跨行业的竞争，与智能手机抢占用户的时间。

对比AR眼镜与智能手机，我们可以发现手机作为“中介式”智能终端在场景交互上并不如“所见即所得”的AR眼镜那么直接。例如用AR眼镜购物时，信息是直接映射而不是通过手机这个第三方屏幕实现的；用AR眼镜导航时，线路是直接标记在眼前的路上的，而不需要抽空看一眼地图，这就决定了AR眼镜若能真正走向消费级，它将颠覆现有智能手机一统天下的格局。

在大环境优势（产业链集群、人才、技术、资金等）下，

AR眼镜本身还能做得更多，甚至替代智能手机重新创造一个繁盛的智能硬件行业。不只是Rokid公司，中国品牌有着充分的优势能在这个环境中获得自己的发展空间。

当然，这些都只是对未来的畅想，解决当下AR眼镜发展面临的诸多现实问题仍然是当务之急。

地震救援机器人什么时候能拯救人类？

“亲爱的宝贝，如果你还活着，一定要记得我爱你！”这是汶川地震中一位母亲留给年幼的孩子的最后一句话。被发现时，她的身体已经变形了，但这位母亲以匍匐的姿势紧紧地抱住了年幼的孩子，最终孩子安然无恙，母亲却不幸遇难。

在汶川地震中，虽然我国进行了积极的救援，但是结果依然不理想。可喜的是，随着AI的发展，地震救援机器人已经小有成果，也许在将来，地震救援将会更加高效。

上能钻下能爬，地震机器人趋于多样化

（1）履带式救援机器人：它由日本横滨的警视厅发明，可以直接把人吞进机器里，然后将人带离危险地带，可有效减少60%的伤亡。该机器人配备了摄像机、机械臂和各种传感器等装置，其中机器里供人躺着的担架里的传感器还可以探测伤者的伤情。

（2）蛇形机器人：为了弥补履带式救援机器人不能钻进狭小的缝隙问题，日本又发明了蛇形机器人。蛇形机器人长约8米，宽约2.5厘米，主要用于搜寻被困人员，它依靠装有动力装置的尼龙绳索进行驱动，可以深入废墟中的每个角落。蛇形机器人装有针孔摄像头，可以将拍到的图像传给救援人员，帮助他们了解受灾区域的内部情形。

斯坦福大学的机械工程师开发出一种蛇形会生长的软体机器

人，它能够在不移动整个身体的情况下长距离生长，并能以蛇形蜿蜒。它比纯机械机器人更安全，这不仅是因为它柔软，还因为它极轻，便于延展靠近人体。

（3）蟑螂机器人：蟑螂虽然令人生厌，却有着非凡的钻缝技能和“抗压能力”。加州大学伯克利分校的研究人员借鉴蟑螂灵活的外骨骼结构，制造出一款机器人雏形，更加适应狭窄崎岖的地面环境。这个蟑螂机器人可以顺利通过狭窄的空间，这也使得它拥有了钻进废墟瓦砾，搜寻地震幸存者生命迹象的潜能。

（4）可探测呼吸和体温的机器人：Quince机器人是很有代表性的救援机器人，只有儿童玩具汽车那么大，装有4组履带式轮子及6个电动机。它的机械臂可以开门和递送食物或进行补给。Quince机器人的独特之处在于其传感器设备，它的红外感应器同时也是二氧化碳探测器，能够探测人体呼吸和体温状况，这可以用于探索地震中的生命迹象。

国内也有一些做得好的地震救援机器人，但是大多属于上面几类。可以看出地震救援机器人的种类已经越来越多样化，最理想化的结果是，在今后的救援中，机器人“深入一线”“蜿蜒前进”，发挥“一臂之力”，“担负”起救援的重任。但在实际使用过程中，依然存在许多问题。

地震救援机器人到底是专家，还是“新难民”？

汶川地震期间，有很多志愿者自行跑到汶川进行救援，结果他们成为“新难民”，给专业的救援人员带来了很多麻烦。在灾害

救助中，不专业的行动，如对地震险情没有预判、缺乏专业知识、准备不充分等，只会给灾区造成更大的伤害。

甚至有救援专业人士认为，震后救援对实施者的考验，比大夫上手术台做最困难的手术还要大。而地震救援机器人是否能经受住这一考验呢？

2001年世界贸易中心大楼被炸毁时，工程师安排了几个轻巧的机器人到废墟中进行搜寻。这些机器人可以钻进那些第一批救灾人员无法进入的地方，但是，这些机器人没有找到任何幸存者，也不能耐高温，还出现了许多技术故障。

10年后，日本福岛核电站泄漏施救工作中，Quince机器人发挥了很大的作用。它先后走遍了多个楼层，进行了辐射和温度测试，它还深入核反应堆建筑物内部拍摄了很多清晰的照片。但是Quince机器人最后没有成功返回，它在执行任务的过程中与控制中心失去了联系。东电公司接着派出的机器人也都纷纷“殉职”。

从任务层面来说，现阶段的地震救援机器人确实能完成某一部分任务，但是还不专业。2017年9月，墨西哥城发生7.1级地震。在地震发生后不久，一队来自卡内基·梅隆大学的机器人专家来到震区，但是在该团队停留在灾区的3天内，他们的蛇形机器人只被使用了一次，而且使用结果并不理想。

地震救援机器人作为震后救援实施者，如果专业性不高，那么它无疑是某种程度上的“新难民”。

现在技术发展得怎么样呢？

城市里，楼梯是为人攀爬而设计的，门是为人手打开而设计的，因此，地震救援机器人越来越类人化。2018年12月，佛罗里达州的一个地震救援机器人比赛赛场上，机器人笨拙地蹒跚前行，完成一系列任务，包括开门、越过碎石、旋转阀门。而与此同时，让机器人把工具放在桌子上这样一个简单的动作都还被认为是在炫耀。

很多这样的任务，人在一分钟或更少的时间内就可以完成，如今最先进的地震救援机器人才刚开始走出实验室，正处于“蹒跚学步”阶段。

想要技术真正从实验室走出去，还有相当长的一段路。据统计，地震中大部分人都是因为被困在废墟中无法得到及时救援而死亡的，因此，如何快速找到幸存者是救援工作的关键。围绕寻找幸存者，在震区的复杂情况下，如何更精准地定位和进行机器识别仍然是现阶段亟须解决的问题。

地震救援机器人离我们还有多远？

地震大多发生在板块边界上，全球最大的两大地震带分别是环太平洋地震带和欧亚地震带，它们是全球六大板块间的接触带。日本恰好地处太平洋板块和亚欧板块的交界处，地震活动很频繁，一年之内大大小小的地震加起来近1000次，因此，日本对地震救援机器人投入了很多的人力和财力进行研究，取得了很多成果。

我国位于世界两大地震带——环太平洋地震带与欧亚地震带

之间，是多地震的国家之一。

显然我国对地震援救机器人的研究，主要还是基于当前可以预见的市场需求。什么时候可以全面应用地震救援机器人？可能是5年，也可能是20年，人机合作会在很长时间内一直占据主导地位。但是就结果而言，我们认为，不管“黑猫”还是“白猫”，能救人就是“好猫”。

为什么说“煤老板时代”又将来临？

我国的煤炭企业数量不断减少，从2200个减少到100个左右，小煤矿将被彻底清除，“煤老板”退出历史舞台已成定局。

但是新的一批“煤老板”又即将横空出世，准确来说，是太空挖矿的“煤老板”。美国《连线》的文章《太空采矿可能引发星球大战》指出，太空采矿有可能是下半个世纪最赚钱的行业之一，经济规模高达数万亿美元。

太空采矿有多赚钱？单颗小行星价值以亿美元为单位

目前国际上比较热门的采矿技术主要集中在火星采矿、月球采矿、小行星带的小行星采矿三个方面。矿产资源最丰富的当数小行星，太阳系有数以万计的小行星，它们的直径从几英尺到几百英尺不等，拥有许多有价值的矿物和化学物质。2013年，Asterank网站对一颗名为“Germania”的小行星做出过评估，这颗小行星估值达95.8万亿美元，这是什么概念呢？

2013年，180亿美元的负债让底特律彻底破产。这颗小行星的价值不仅远高于当时全世界GDP的总和，还可以负担起1041个破产的底特律。

小行星的价值计量单位从来都是以亿美元计的，因此，政府、私人公司、学界纷纷加入“淘金”热。2016年NASA研发出了一个能在月球上进行资源开采的机器人RASSOR，它专门用来探测月球上的矿石、沙土、水分等资源，可以说是一个什么都能干

的“优秀矿工”。

但是，特朗普在2018年的NASA预算中砍掉了“小行星重定向计划（ARM）”的一部分。该任务计划将一颗小行星引入环地球轨道，用以研究和开采。而这意味着这块“大肥肉”将会落入私人航天企业口中。

2016年，美国一家私人航天公司Moon Express获得美国联邦航空管理局登月许可，该公司计划于2020年在月球南极建设机器人月球矿藏开采前哨基地，将利用自家探测器在月球挖掘稀有的矿产资源。

不过，2020年Moon Express是否能到月球开采还不能确定，可以确定的是，马斯克的SpaceX公司很容易就可以在太空采矿领域分一杯羹。2016年，马斯克解释其火星殖民计划时，就曾提过太空采矿的设想。著名的天体物理学家马丁·埃尔维斯表示，如果有一艘非常先进的太空飞船，能在地球轨道和环绕小行星运行的轨道切换，就可以实现在小行星开采，而猎鹰重型火箭正是这种太空飞船。

盯上这些小行星的公司特别多，最先成立的是私人勘探公司“行星资源”，它于2012年在华盛顿成立，并在2016年接受了卢森堡政府的2500万欧元投资。卢森堡虽然地方小，但对外太空的兴趣却不小。它是欧洲较早为外太空矿产、水和其他资源（尤其是小行星上发现的资源）的所有权提供法律依据的国家。据《南华早报》报道，我国政府已经计划追踪并登陆某颗小行星，预计最早在2020年实现对太空岩石中稀有珍贵的金属和矿物的开采。

虽然政府和私人航天公司对太空都充满兴趣，但至今依然没有突破性的消息传出。值得庆幸的是，世界上在采矿方面研究实力最强的机构之一——美国科罗拉多矿业大学，已经开设了太空采矿专业，专业人才将会源源不断地出现，太空采矿的步伐在不断加快。

小行星虽好，但不易开采

汪海林老师说，不能轻易采矿，因为要求高，而开采小行星的要求更高。开采小行星一般有两种方式，一是运回地球附近开采，如前面提到的被特朗普砍掉的ARM；二是就地开采。无论哪一种方式，难度都非常大，主要有如下两个问题。

1. 环境问题：如何克服太空恶劣的环境

为保持相对静止，设备在对小行星进行开采时会随着小行星的自转而自转。但是这样航天器就会周期性地处于太阳强照射和黑暗中，当航天器对自转周期比较小的行星开采时，很容易因频繁的冷热交替而损坏。同样，容易对设备造成损害的还有太空中的强紫外线和真空环境。前者自不必多说，在真空条件下，常用的聚合材料会脱气，具体表现是材料性质不稳定、结构扭曲等。在这种情况下，对开采设备的材料要求就特别高，容错率几乎为零。

地球上能采矿是因为有重力存在，但对于微重力甚至零重力的太空，不管用哪种方式都存在很多技术问题。以地下开采为例，架支架这一套在地球上常见的手段还是否有用？整个矿房的

尺寸可以扩至多大？这些都是需要探索的关键问题。

如果说上面的问题都是可预测的，那么接下来的这个问题就可能难以预测了。太空中存在大量的微小星体和飞行器留下的微小残骸，这些微小个体常以极快的速度撞击小行星或互相碰撞。这些微小残骸的特点是快、准、狠。采矿设备常常会来不及避让，这种情况也难以预防，对采矿设备的打击极大。

2. 技术问题：如何在小行星上锚定和钻孔

小行星之所以叫小行星就是因为有一些行星很小，甚至直径不足一米。对于这种行星，最简单的办法就是用网抓捕起来运回地球。而对于尺寸较大的小行星，则需要航天器在其表面附着，再进行后续操作。这里可以以船类比，当船靠岸或停止时，需要抛出锚。而由于在太空中航天器处于微重力状态，对固定航天器的锚具有更高的要求。

说到采矿，钻孔是一项基本而重要的工作。在地球上，一般会用水充当钻头与孔壁之间的润滑剂。而在太空中，由于超低的温度和压强，水呈固态，难以喷洒。当钻头进行钻孔工作时，进出都很困难。

这些问题在短时间内很难得到有效的解决，需要长期攻克。目前很多国家已经将太空作为AI发挥作用的重要“战场”，但是它的作用依然有限。先进的AI也难以解决这些问题，那么，想要在太空挖矿还能怎么办呢？

细菌也能挖矿，生物学技术开启新思路

细菌微生物等常常有旺盛的生命力，它们对太空这种极端恶劣环境的耐受性特别强。利用它们在极端环境下的生活方式，可能会给太空挖矿提供另一种思路。泰狄娜·米洛杰维克及其研究团队在维也纳大学生物物理化学系，对火星不同位置、不同历史时期的表层土壤成分进行了模拟，搭建了一些小型“火星农场”。该“农场”包括气体及化学合成的各类成分，火星岩石和可能已经灭绝的细菌生命等。

他们观察了一种名为“嗜热金属球菌”的古老细菌与类火星岩石之间的相互作用。这种细菌不仅能在极端环境中生存，而且能够从各类矿物中吸收金属作为营养物质。作为一种自养生物，它能对铁、硫、铀等无机物进行新陈代谢，释放出可溶性金属离子。研究者认为，它可以作为一种生物学技术，人们能用它开采小行星、流星及其他天体上的金属矿物。

由于我国的私人航天发展得较晚且还处在造火箭阶段，还没到考虑挖矿的地步，因此，我国民间没有传出太多有关太空采矿的声音。太空采矿这一设想最早被提出是在1903年，那一年，莱特兄弟刚刚试飞成功人类历史上第一架飞机。著名火箭科学家康斯坦丁·齐奥尔科夫斯基将“探索小行星”列为征服太空的14个目标之一。100多年过去了，飞机成了常见的交通工具，我们有理由相信，总有一天，去小行星采矿会像坐飞机一样平常。

AI能翻译婴儿语言吗？

26岁的王小强刚做爸爸，他还沉浸在初为人父的喜悦中。但是，当婴儿开始哭闹时他就没辙了，孩子到底是饿了、尿湿了还是哪里不舒服？

正常婴儿在两周大时，每天约哭闹1.5小时；出生后第6周，平均每天哭闹2.5小时；到3个月时，正常婴儿哭闹的时间减少到每天1小时。婴儿每天都在哭泣，理解“婴语”对父母而言很难。
(免费书享分更多搜索@雅书.)

美国的Ariana Anderson博士也是如此。她是4个孩子的母亲，在抚养前两个孩子时，她常常搞不懂孩子大哭的原因，直到抚养第3个孩子时她才能轻松地理解孩子哭声背后的意义。基于这种经历，她研发了一款用算法理解婴儿哭声的“婴语”翻译App——Chatterbaby。

“婴语”难理解，AI当自强

婴儿哭闹的原因很复杂，一般分为生理性哭闹和病理性哭闹。生理性哭闹常见的原因有饥饿、排尿、排便、疲倦困乏、生活规律紊乱、衣着不适、出牙，要求或欲望未得到满足。病理性哭闹常见的原因有维生素D缺乏性佝偻病、感染性疾病、腹痛、猫叫综合征、维生素A或维生素D中毒及新生儿甲状腺功能亢进症、头痛。

区分生理性哭闹和病理性哭闹是儿科医师经常遇到的难题，

这也从侧面解释新生儿父母要正确理解婴儿的哭声有多难。

很多团队对婴儿哭声做过相关研究，各团队都表示他们的产品可理解95%以上的宝宝哭的原因、准确度比人高3倍等，但是至今我们并没有看到具体的产品。Chatterbaby通过收集婴儿在不同感受下的哭声，并通过机器学习对这些哭声频率和特征进行分析学习，从而告知父母他们的孩子为什么而哭。现阶段，这个App可以分辨孩子的哭泣是因为饥饿、烦躁还是疼痛。

布朗大学风险儿童研究中心的心理学家Stephen Sheinkopf认为，婴儿的哭声中确实隐藏着许多神经学线索，尤其是在哭声的声调、声量、共鸣等声学特征中，这些特征能被量化和可视化。但是仅通过声音对婴儿进行检测显然不够，将声音、行为和其他生理数据整合进一个模型，无疑会比通过单一的声音特征检测准确得多。

2018年3月31日，阿里云发布了一款“婴语”贴纸。这款贴纸无毒无味，形状近似便利贴，轻贴在婴儿任何部位即可运行。对于0~1岁的宝宝，它的识别准确率达到了95%，能够轻易识别婴儿的30多种行为，如饥饿、疲倦、害怕、开心、撒娇等。

这款产品依托阿里云物联网IOT套件、AI智能语音分析、生物识别反馈系统、情绪建模等，通过App实时对宝宝的哭声进行分析与反馈。贴纸不仅可以检测婴儿哭声的分贝与尖锐程度，还可以对心率或体温等进行检测。

这些依托AI的“婴语”翻译机似乎给奶爸奶妈们带来了福音。

对经验少的父母而言，“婴语”翻译机在早期无疑会派上用场。但随着宝宝逐渐适应周围环境，“婴语”翻译机的准确率会下降。此外，依然有一些我们不得不面对的问题。

路漫漫其修远兮，“婴语”与AI相互求索

AI现阶段的发展依然还处于“婴幼儿”阶段，那么让AI服务于婴儿还有哪些问题呢？

1. 技术未满，能用而已

从某种程度上来说，识别婴儿的生理性哭闹可能较为容易，但是识别病理性哭闹可能就没那么简单了。将哭闹不休的婴儿带到儿科，医生也要通过各种检测才敢下定论，测体温和量心率并不能作为依据。

而且，一旦“婴语”翻译机将病理性哭闹误诊为生理性哭闹，会耽误最佳的治疗时间，其后果不堪设想。

这并不是类似于普通消费升级的产品，可以通过不断迭代来优化产品性能。这是一项容错率非常小的技术，一旦出错，就会引发悲剧。

2. 先看手机还是先看孩子

有专家认为，手机在发射微波的同时也存在“极低频磁场”。由于婴儿的头骨较薄且不完整，脑部发展不完善，如果婴儿接触手机太多，那么会对身体造成影响，可能导致癌症、神经及发展

障碍等。且不论这种说法是否有骇人听闻的嫌疑，但是绝大部分父母都不愿意冒这样一个风险。那么，面对婴儿哭泣，父母该怎么办？

是先去抱小孩还是先看手机？

如果是两代人同时照看婴儿，先看手机无疑会引起家庭纠纷。一般人的反应是先哄孩子，然后抱着孩子看手机。

那么，通过声音辨别疾病就只能这样了吗？

可能不适合用声音确诊某些婴儿的病理性哭闹，但可以通过声音识别疾病，甚至它还适用于成人预防一些难以治愈或不易发现的疾病。

现有研究结果表明，人的声音变低哑可能是因为胃酸反流；如果患了慢性鼻窦炎，则会出现鼻塞一样的声音；人的声音变得低沉很可能是得了甲状腺疾病；人的声音变得没有起伏，很可能是患了帕金森病。

美国麻省理工学院的李特博士近年来就在研究这一课题——利用声音诊断帕金森病。现有医学研究表明，人类患帕金森病是因为大脑中脑部位的神经细胞发生了变性死亡，这种细胞死亡带来的症状是肢体颤抖。此外，没有更多新的进展，至于如何治愈帕金森病更是无从说起。

相比于化验的检测方法，这种方法极大地减少了对人身体的接触，李特通过大量的数据分析，建立起一套声音分析系统。只

要留下30秒左右的录音，就能让计算机自己判断人患病与否，实验室阶段准确率高达99%。

这项技术最终可以帮助李特客观地掌握症状的进展情况，从而可以更有效地控制用药量，甚至准确确定吃药时间。

Chatterbaby的另一大目的就是通过哭声来诊断不同类型的自闭症。医学界和教育界都认为，自闭症越早发现越好。如果在早期对这些孩子的特殊需求进行满足与鼓励，就会增加他们最终成长为正常人的概率。可惜的是，患上自闭症的孩子往往在多年后才会被确诊。

权威数据表明，我国已确诊自闭症患病人数在1000万以上，而0~14岁的儿童患者已经超过200万人，并以每年20多万人的速度增长。如果通过一项简单的指标就能对孩子的发展做出预测，那么这无疑是一项值得推广的技术。

现阶段的AI显然不能解决所有问题。不过只要AI能解决一部分问题，提升解决问题的效率，就可以发挥其应用价值。毕竟这项技术对于大多数年轻父母来说，不失为一种安慰剂。

为什么人类既期待又排斥“读心机”？

哆啦A梦有一个神奇的记忆头盔，把它戴上之后，人就能听到附近其他人心里想的事。这种读心机居然在现实中出现了。

2018年4月，美国麻省理工学院的团队开发了一款名为AlterEgo的头戴设备，严格来说是“下巴佩戴的设备”。这款设备能够通过内置的电极读取脸部神经肌肉中的电流信号，从而知道某个人想表达的内容。换言之，这款“读心机”能够知道你心中想的事。

这一发明听起来似乎非常好，但实际并不是这样。

读心机并不是真正地懂用户

虽然读心机这个概念炒得很火，但我们想提醒大家，目前它的效果似乎并不令人满意。

1. “读心”这个事不靠谱

Affectiva公司旗下的核心产品Affdex可通过扫描人脸上诸如眼睛、鼻子、嘴巴、眉毛之类的部位，建立相应的锚点，并通过捕捉皮肤细纹的变化来辅助识别表情，最终依据对面部表情的识别来“读心”。

但是表情（甚至微表情）是可以被表演出来的，如何保证在分析时不被目标刻意做出的表情误导呢？丹麦神经营销学的领军

人物Martin Lindstrom认为，人会撒谎，但是大脑不会。要想观察人真正的想法，应该去看他的大脑在想什么，这比听他在说什么准确得多。

俄勒冈大学的一个团队创建了一个系统，这个系统可以通过扫描大脑读取人的思想，用眼睛“看见”别人脑中的想法。

但机器重建的结果并不理想，基本的脸部重构并没有实现，情景再现就更不用说了。这项技术一旦成熟，便可在刑侦领域大放异彩。最理想化的情景应该是利用受害者的记忆来构建罪犯的面部照片，或基于犯罪分子的大脑记忆还原犯罪现场和犯罪场景。

因此，不管是面部识别还是脑部识别，在“读心”这件事上，它们都不算靠谱。

2. “读心”有多难？

现阶段对脑电波的检测，大致分为两种方式：植入式和非植入式。植入式的方式更精确，我们利用此方式可以编码更复杂的命令。但非植入式的方式更安全，接受程度更高，如果面向健康人类开发产品，那么这几乎是唯一的选择。

那更精确的植入式即便安全度很高会怎么样？

2016年荷兰乌得勒支大学医学院的神经科学家Nick Ramsay领导团队完成了用于临床的大脑植入手术，实现对ALS（肌肉萎缩性脊髓侧索硬化症）患者思想的解读。

该实验把两个电极安放在大脑的皮质区（控制运动区），患者通过想象自己敲打键盘传播信号。经过训练，他的“输入”速度已经变快了——从一开始的50秒选中一个字母提高到了20秒。

非植入式的方式是对大脑进行核磁共振记录，我们在头皮上贴上电极，收集脑电运动。这种采集方式虽然实行起来方便，对使用对象的伤害较少，但效果甚微。

人体的脑电波信号非常微弱，因此，脑电波很容易受到如静电、皮肤油脂、化妆品的干扰。而人脑的运动又异常复杂，整个大脑约有850亿个神经元，使用这种方式只能监测到几万或几百万个神经元的运动。

这样就意味着，现阶段，想要做到“读心”几乎是不可能的。而就AlterEgo这款“读心机”而言，它仅仅是读取脸部神经肌肉中的电流信号，从而知道某个人想表达的内容。而对某些人而言，想说的话与内心真实的想法可能并不相同。因此，“读心机”并不能真正能读懂我们的心，最多读懂了我们想让它读懂的部分。

就算“读心机”懂用户，然后该怎么办呢？

如果“读心机”真的问世了，那我们就近乎三体人了。“思维透明”是三体人的生物特征，三体人脑电波强，他思考时的脑电波能被别的三体人大脑接收到。

扎克伯格有一次在回复网友时这样说：“Facebook拥有AR技术和服 务，用户只要一直戴着可穿戴设备，就可以提高用户的使用体验，改善沟通方式。未来某一天，我认为我们可以通过使用

技术向其他人直接发送完整的、丰富的想法。你在脑海中所想的东西，你的朋友也可以在自己的大脑中快速地体验一遍。这就是Facebook的终极交流技术。”

“你在脑海中所想的东西，你的朋友也可以在自己的大脑中快速地体验一遍”，这真是一句令人毛骨悚人的话，这必将意味着零隐私。虽然李彦宏认为用户愿意用隐私交换便捷，但是此隐私非彼隐私。况且每个人都有难以启齿的事情，但是现在这些事不仅要被人知道，还要接受广大人民群众“检阅”，这真的是人们未来想要的交互方式吗？

不可否认，将“读心机”应用到一些特定领域会发挥其价值，例如Affdex的最初动机是为了方便与自闭症儿童沟通。通过Affdex解读这些自闭症儿童的表情，方便我们与他们交流，也有利于心理疏导。在广告领域，制作者只需要在广告片完成后，邀请一部分人来试看，然后在这个过程中使用Affdex测试观看者的情绪变化，最后找到他们情绪波动最大的段落插入Logo便可实现精准营销。

“人工心智”将会成为人类的生存性威胁吗？

自AI兴起以来，就有不少学者提出了“AI威胁论”。人们想象，当AI被完全应用时，AI将会成为一种源自人类却优于人类的独立自主的生命形态，具有与自然人一样的人性构建。

当然，还是有不少人提出了反对意见，人们对AI发展的忧虑说明他们很大程度上在于混淆了智能与心智这两个概念。一般而言，心智是由三个功能相对独立但又全面整合的系统——情感系统（感受和动机）、意志系统（意志和行动）和智能系统（感知和思维）构成的。显然，智能只是心智的一个子系统。

AI虽有感知和思维，但如果没有情感系统的引导，智能系统就不能发挥它应有的作用。心智一直被认为是AI无法掌握的能力，然而，在科研机构DeepMind的一篇论文*Machine Theory of Mind*中，研究人员提出了一种新型神经网络ToMnet，其具备理解自己及周围智能体心理状态的能力。

AI被开发出了心智，我们在高兴的同时也感到了一丝担忧。AI越来越拟人化，究竟是好是坏？面对此项研究，AI威胁论是否有了新的证据？

AI的“心智”，从理解信念开始

所谓信念，是指人的愿望、想法、知识、理念、意图及需要等。信念与人的愿望、情绪、行为等都有很大的关系。在日常生活中，人们都是通过推测别人的信念来估计他人的情绪、心态及

下一步行为的。

信念分为真实信念和错误信念，例如你看到一个书包，自然会去猜测里面装有书籍。如果书包里确实装的是书籍，即人的想法与现实情况相符合，这就是真实信念。而错误信念就是，你认为书包里有书，但书包里装的是一件衣服，即想法与现实不符合。

一般来说，AI能够通过大数据、图像识别等技术，基于概率推论预测出人们或智能体的真实信念。但绝大部分的AI都难以通过简单的错误信念测验。即使AI很可能以很高的准确率完成错误信念任务，但它依旧不能理解这其中的原因。

有一个经典的初级信念测试——Sally-Anne测试。Sally有一个篮子和一个球，Anne有一个盒子。Sally将球放进自己的篮子里，然后便出去了。Anne将篮子里的球拿出来，放进她的盒子里。后来，Sally回来了，她想玩球。信念问题是，Sally会去哪里找球？

4~5岁的孩子都能够答对这个问题，但AI就不具有这样的认知。人类在大约4岁的时候就会开始理解错误信念，即开始知道现实生活中可能有一些情景与自己的想法是不一致的。

理解错误信念对于机器人的心智开发而言，有至关重要的作用。因为对错误信念的理解是推测他人对事件的判断、预测他人下一步行为、了解他人心态的基础。这是训练AI掌握心智理论的关键阶段。

在DeepMind的研究中，最终结果证明，ToMnet学会了对来自不同群体的智能体进行建模，包括随机、规则系统和深度强化学习智能体等。该网络通过了经典的Sally-Anne测试，可以理解持有错误信念的智能体，也就是说，AI已经可以达到4岁儿童的理解水平了。

AI的心智开发，总算有了坚实的基础。

AI心智开发后，“AI威胁论”为何再次兴起？

毋庸置疑，AI对错误信念的理解有利于其心智的形成，具备理解能力的机器人也将改善人机之间的合作。然而，技术的发展也是双刃剑，AI再次升级，也引发了人们的担忧。

1.AI“骗术”再升级，难以被人类信任

AI“欺骗”人类已久，比如在2017年，Facebook的聊天机器人就创造了一种新的表达方式进行对话。机器人自创“行话”，让人类不解其意。

工作人员将这次“意外”解释成没有设置英语语法的激励机制，但AI背着人们交流，还是让人类产生了担忧。

在“欺骗”人类的道路上，AI可能会越走越远。从语言合成到图像生成再到模拟人类的对话，AI正在许多领域逼近甚至超越人类的水平。

有研究人员指出，如果一台计算机能够理解错误信念，那它

就可能知道如何诱导人类去相信它，进而实现机器对人的“欺骗”。

设想一下，在聊天机器人学会谈判之后——它依靠机器学习和先进的策略，试图改善谈判的结果。随着时间的推移，这些机器人谈判变得相当熟练，甚至开始假装对某件事情感兴趣，便于在谈判的后期做出“牺牲”。

2. “人工心智”的威胁

以前，AI的威胁可能仅仅在于技术的工具性威胁和人们的适应性威胁，这也是出现各种取代论调的原因。

AI作为人类使用的工具、技术，人们害怕因使用不当造成危害或引发失业风险，例如弱AI系统可能因人的恶意操控而造成导航错误、金融市场崩盘等。在这些想象中，AI并不具有自主威胁人类的意图。

但是，“人工心智”的出现引起人类对AI最深层的恐惧，那就是AI具有主观的感受后，会给人类带来生存性威胁。

霍金曾这样描述过AI：“一旦人类开发了AI，它将以一种持续增加的速率不断进行自我更新，并最终脱离我们的掌控……AI未来还将发展出自己的意志，这种意志将与人类的意志产生冲突。”

事实上，能带来这种威胁的AI更确切地说不是AI而是人工心智——一个尽管是人工的，但却是独立自主的、能够感受周围智

能体的主体。当AI有了心智，就会具有能对智能加以规范的自治的情感系统，这样，它与人类的区别又会缩小。

就目前来看，AI为我们解决了许多难题，在未来在AI的帮助下，人类社会的生产力将会大幅度提高。至于AI的发展是否会威胁到人类，我们持中立态度。

在观察了许多人对AI的担忧后，我们发现，一切生存性威胁都来源于人类的臆想——人类往往设想AI的动机和意图一定是邪恶的，它会凭借自己的超能力控制、奴役甚至灭绝人类。

而在现实中，AI其实并没有我们想象的那么坏。

盲眼“猎豹”机器人出路在哪？

波士顿有“狗”，麻省理工学院有“猎豹”。

麻省理工学院的这匹新“猎豹”不简单，直接代替“波士顿狗”成为机器人界的新“网红”。为什么这次麻省理工学院的“猎豹”如此厉害？原因在于这匹“猎豹”闭着眼睛就能疯狂吊打它的各位“前辈”。麻省理工学院新研发的第三代“猎豹”机器人完全不依靠视觉传感器及摄像头，纯靠算法便能轻松完成移动、爬楼梯、快速奔跑、跳跃等高难度动作，相比以往的机器人来说，有了质的突破。

那么，这样的机器人究竟有什么用？不断对这些机器人的性能进行迭代升级的研发团队可不仅仅希望它只是“网红”那么简单。“猎豹”机器人最想占领的领域就是人类难以涉足的野外搜救、侦查领域。

但是对于“猎豹”机器人及它的同类来说，真正成为这个领域的佼佼者还有三大难题。

野外作业“五感”能力跟不上，机器人量产难实现

目前，腿足式机器人是机器人界的主流。多年来，在机械设计和控制策略方面，腿足式机器人进展明显。但是这类机器人在准静态的时候运动速度非常慢，而且也难以处理一些意外状况，还经常摔倒。

为什么运动速度提不上来？这主要是由于过去的机器人过于依赖视觉系统，每走一步都需要通过复杂的视觉计算。所以，这次团队在研发第三代“猎豹”机器人时转变思路，不再一味地看重视觉能力，而是转为提高机器人的触觉能力。他们开发出了“接触检测算法”和“模型预测控制算法”，让机器人在接触地面的时候迅速判断出脚下所踩的材料是什么，并根据不同的材料做出相应的反应，提高了机器人的运动速度。

尽管速度得到提高，但在危险的野外作业环境下，人类五感中的任何一个起到的作用都非常关键。目前，就算是最先进的机器人也无法把人类的五感齐聚一体，拥有了视觉和触觉的机器人往往又会缺乏嗅觉，而这也相应地限制了机器人产业的落地进程。之前，科学家曾开发出一种具有嗅觉能力的机器人。当它通过某一片区域的时候，靠近地面的导管会将气味吸入LSPR传感器进行分析。在不同的气体成分中，金纳米粒子会产生不同的光吸收反应，快速探测出气味的成分。但是这种机器人能够识别的气味范围十分有限，目前仅能识别出乙醇分子的气味。而传统警犬的嗅觉，往往比人类强几百甚至上千倍。要用这类机器人代替警犬进行野外侦查工作，目前还难以实现。

形态学是个大问题，如何穿过狭小空间？

传统机器人（如“猎豹”机器人）大部分都是由金属和塑料等刚性材料制成的，具有功率大、精度高、性能稳定等优点。但是机器人的稳定性在遇到野外的恶劣环境后还能否保证就很难说了。在一些狭小的搜救空间中，我们很难想象一堆金属和钢筋水泥硬碰硬是怎样的场景。另外，传统机器人的关节部分多采用刚

性部件，摩擦力较大，这导致了这类机器人的寿命会相对较短，因此应用成本很高。

目前想要用机器人侦查、搜救必须使它具备一定的变形能力，才能方便它在各种狭小空间下作业。而软体机器人恰好具备了这样的能力，可以在各种非结构化的环境下工作。而且软体机器人的柔韧性好、摩擦力小，可使用年限更长。

对于搜救、侦查这项工作而言，其实犬类的形态学构造可能并不是最完美的。美国加州大学伯克利分校的研究人员开发了一款外形酷似蟑螂的“铰链式可压缩机器人”CRAM。CRAM的外壳非常柔软，而且它有蟑螂那样独特的外形结构、行走方式和“打不死”的顽强生命力。研究人员设计的CRAM能迅速穿过窄至3毫米的空间，能承受高于体重900倍的压力仍毫发无损。也就是说，CRAM完全可以装配微型麦克风和摄像头帮助搜救队探测瓦砾中的具体情况，便于迅速找到被困人员。

让人去适应机器人，这个逻辑不对

人类是如何控制机器人的？让机器人按要求行动，它需要理解人类发出的指令。在绝大多数情况下，这意味着机器人需要学习复杂的人类语言或配备一个远程遥控设备。但如果把机器人大规模运用到野外侦查、搜救工作中，那么每个警察必须人手配备一个遥控设备，还必须谨防电池没电或语言不通。

因此，我们需要建立一套统一的人机交互系统。麻省理工学院的计算机研究所开发了一套可以依靠脑电波和手势控制机器人

的系统。麻省理工学院首先收集人们看到机器人错误动作时发出的脑电波，然后让机器人学习识别这部分脑电波。当机器人意识到自己错了，便会根据人类的手势再次做出正确的反应。因此，当我们看着机器人的时候，我们所需要做的只是在精神上认可它的行为或反对它的行为，不必训练出一套适应机器人的行为方式，反过来说，机器人要去适应人类。

通过观察肌肉信号，机器人可以理解人的自然手势，这也会使得人类与机器人的交流更加接近人与人之间的交流。麻省理工学院的这一系统去掉了人类和机器人之间多余的连接设备，以人体本身作为操作中枢，让机器人主动理解人类的意图，而不是让人去适应机器人，这才能够称得上AI。真正的AI不是把事情变得更麻烦，而是把事情变得更简单。

目前，不管是在搜救、侦查还是大众消费领域，机器人的普及程度仍然不高，“炫技”“营销”这种标签还贴在机器人身上。只有解决了这些难题，机器人才会得到大众的认可。

人类想早点移民火星，AI能帮上忙吗？

十年前，互联网上很多人自诩“火星”，他们写着常人难以理解的“火星文”，生怕被别人看懂。但是想移民火星做火星，并不是这十几年发生的事。自从人类在1969年首次登月后，但凡在太空事业上取得一点进展，就会有人问离移民火星还有多远。

随着40多次火星探测任务的开展，火星已经成为太阳系中除了地球人类了解最为透彻的行星。但即使这样，移民火星依然遥遥无期，最重要的原因是恶劣的环境使人无法生存。另外还有三大障碍：无氧的大气、严寒的温度和干旱的地表，这三个难题都很难被攻克。

随着科技的发展，AI也越来越多地渗透到了航天领域，如上天的航天机器人“西蒙”、日本送入国际空间站的机器人Kirobo等。那么，在移民火星这件事上，AI能帮人类做些什么？

打前哨：兵马未到，探测先行

2018年7月25日，在*Science*上发表的一篇文章称，意大利多家机构组成的团队对火星快车号探测器的雷达数据进行分析后，发现在火星南极冰盖下方1.5千米处存在一个直径约20千米的稳定的液态湖。这一消息一经发布，人们纷纷欢呼雀跃，似乎一脚已经迈进了火星。

1. 机器视觉与机器学习虽好，但与效率没有太大关系

基于火星复杂的环境，在移民火星之前，探测是非常重要的第一步。而我们对火星的认知也绝大多数来自探测仪发回地球的数据，火星探测的重要性不言而喻。

NASA的火星探测器计划一直在紧锣密鼓地进行。早在2016年，“好奇”号开始使用一种名为“科学知识收集自主探索”的软件，将计算机视觉与机器学习结合起来，根据科学家确定的标准，选择岩石和土壤样本进行研究。火星车可以用它的化学摄像机的激光对目标进行射击，分析燃烧的气体，将图像和数据打包，然后发送回地球的指挥中心。

虽然已经在一定程度上积极运用了AI，但是这些火星车行动非常缓慢。在火星超过2000天的时间里，“好奇”号仅行驶了约11.5英里，效率并不算高。

2. 无人机探测飞得更快，探索区域更广

无人机探测是一种高风险、高回报的项目。虽然火星上稀薄的大气层可使巨大旋翼叶片旋转速度更快，但由于火星大气密度只有地球大气的1%，为了使无人机能够在低密度大气环境中飞行，无人机重量要尽可能轻，尽可能结实和牢固，同时需要配备强大的动力。

NASA将在2020年发射“火星2020”探测器，并在这个探测器上搭载一架无人机。在此之前，NASA已经探索性资助了一种全新的探测器：由AI控制的一群机器人蜜蜂Marsbees。Marsbees是一种微型的无人机，结构和大黄蜂类似，但它有更大的翅膀，通

过机器震动翼飞行。这样的设计能让这些机器人飞入火星的大气层，探测并获取有关数据。

值得一提的是，火星探测不仅要在火星表面实现数据采集和打包发送，还要突破深空超远距离测控通信、火星制动捕获、在轨长期自主管理、稀薄大气减速与安全着陆等关键技术，而这些关键技术也常常需要AI的辅助。例如在软着陆火星时，AI可以拍摄多张照片，通过分析这些照片，使着陆器自主选择一个能让“四条腿”安全着陆的平面。

定居前：适应不成，“智”能改造

地球上的人和其他生物都是依赖于24小时的生物钟而生活和工作的，但到目前为止，太空中还没有一个星球的自转和公转完全与地球相同。虽然火星自转一周的24.6小时与地球自转一周的23小时56分4.09秒相近，但火星绕太阳的公转周期是1.88年（687日），与地球上的一年相差太多。因此，人类想要在火星上生存，就需要彻底改变自己的生物钟。

而改变生物钟有多难呢？谁也没办法回答。同时，由于火星存在大量的二氧化碳，人类要么不呼吸，要么只能呼吸二氧化碳。如果人类适应不了高浓度的二氧化碳，便会引发一系列不良反应甚至死亡。

适应这条路行不通，那么只有一种办法——改造火星。

在改造火星初期，人类需要在火星建起一个个巨大的密闭式火星基地，但仅依靠人是没办法完成这件事的。由于火星的环境

恶劣，人们穿着厚重的宇航服工作，呼吸、行动、体力都是问题，必须依靠拥有强AI的机器人，明确分工、准确高效地协作人类完成从原料制造到设备搬运、基地建设，再到基地密闭舱环境调控等一系列复杂任务。

德国宇航中心DLR研发团队花了十年时间优化人型机器人Justin的物理性能，目的是让它解决火星上人类的住房问题。物体识别软件和计算机视觉能力让Justin能探测并分析周边的环境，完成维护清洁检查仪器、搬运物体等任务，同时它还可以使用工具、拍照并上传照片、精准抓住飞行中的物体和躲避障碍物。在一次测试中，在宇航员用平板远程指导下，Justin仅用几分钟就修好了Munich实验室失灵的太阳能板。

实际上，不管是在科幻电影中还是小说中，在太空，人类与机器人合作都是“标配”，在移民火星这件事上也会如此。

未来：从弱AI到强AI

依据两种不同的目标或理念，工业界的AI主要有两条路线，一条是弱AI路线，以概率论、统计学、贝叶斯定律等为基础，基于大数据、云计算体系及人类过去的生产过程、生活过程形成的数据积累，通过机器学习完成指定任务。另一条是强AI路线，主要是对脑科学、神经科学的研究，目的是研制出达到甚至超越人类智慧水平的人造物，且它能有心智和意识、能根据自己的意识开展行动。

现阶段，弱AI越来越强，强AI却越来越弱。因此，人类对深

空的探测活动基本上还处在基于计算机视觉、自然语言处理、机器学习、语音识别等弱AI阶段。

但是火星探测器与地面最远通信距离约4亿千米，是地月距离的900多倍，“对话”延时长达40多分钟，需要克服信号衰减、传输延时和外界干扰等困难。由于距离太远，在大多数情况下探测器需要依靠自主控制，独立完成帆板展开、对日定向、制动捕获、两器分离、故障诊断等工作。

因此在未来，发挥作用的将是一批强AI高智能机器人，它们可以自主独立地进行分析判断，如前面提到的Justin便在某种程度上向着强AI发展。一般的机器人，需要预先嵌入程序，并且它们的每个行为都需要有明确的指令。但Justin不一样，它可以自行解决很复杂的任务，从某种程度上来说，它具有一定的独立思考能力。

发展强AI也是我国航天领域未来的发展方向。我国首次火星探测任务于2016年批准立项，预计将在2020年由“长征五号”运载火箭发射火星探测器，直接送入地火转移轨道。在此次“出征火星”的计划中，强AI将发挥主力军作用。

5 生活“局中局”

AI同声传译为什么成了巨头们都翻不过去的坎儿？

AI又摔了个“跟头”。

在2018年的博鳌亚洲论坛上，第一次出现了AI同声传译。值得注意的是，这是博鳌论坛创办17年来首次采用AI同传技术。然而，在如此重要的场合，现场配备的腾讯AI同传却掉了链子。AI同传词汇翻译不准确、重复、短语误用等“乌龙”引来了嘲笑。

人们总是把AI和人类职位对立起来，各种“取代论”层出不穷。然而，就目前来看，AI同传前路未明，太早将其与人类同传对立起来实在不明智。

那么问题来了，AI同声传译为什么成了巨头们都翻不过去的坎儿？我们要如何解决AI同传的难题呢？

AI同传进阶之路：如何变智能问题为数据问题

很多人都觉得AI如果要处理自然语言，就必须理解自然语言。实际上，AI翻译靠的是数字，更准确地说是统计。AI同传出错，并不是不够“智能”，而是数据和模型出了问题。

1. AI同传还需要理解力

首先，AI同传要理解场景。在博鳌论坛上，人们在会议中讨论的内容专业度高、覆盖面广，AI对特殊场景的理解还不够。场景对于语义具有至关重要的影响，相同的一句话在不同的场景里有不同的意思。例如“好”这个字有多种意思，它既可以表示称赞，也可以表示状态，还可以表达问好.....诸如此类，语义的表达和理解都要结合具体的场景。在具体的句子中，这种语义与情景的结合就更为紧密，更需要机器理解学习。

其次，AI要理解口语的模糊逻辑。口语翻译是不会100%传译的，根据AIIC（国际会议口译员协会）的规定，同传译员只要翻译出演讲者80%的内容就已经算是合格了（90%~100%的“同传”几乎是不可能的）。这意味着AI工作量减少了吗？当然不是，正是这种模糊的东西使得AI同传更加困难。此外，口语没有标点符号来标识句子，缺少了必要的声调和停顿，就很容易引起歧义。而模糊的指令极有可能出现的是满屏的错码。

2. 用隐马尔可夫模型（HMM）解决统计数据之外的语言问题

在参考腾讯AI同传的失误后，我们发现，仅仅增加数据量是不够的，在现实生活中，还会遇到零概率或者统计量不足的问题。

比如，一个汉语的语言模型，就足足达到20万这个量级。曾有人做过这样一个假设，如果除去互联网上的垃圾数据，互联网中将会有100亿个有意义的中文网页，这还是被高估的一个数据。

为了解决数据量的问题，我们认为可以借助隐马尔可夫模型（HMM）。在实际应用中，我们可以把HMM看作一个黑箱子，这个黑箱子可以利用比较简洁的数据，数据被处理后就能得出如下结论：

- （1）每个时刻对应的状态序列。
- （2）混合分布的均值和方差矩阵。
- （3）混合分布的权重矩阵。
- （4）状态间转移概率矩阵。

看起来可能比较复杂，简单来说，这个模型可以通过可观察的数据发现这个数据域外的状态，即隐含状态。也就是说，我们可以凭借一句话来弄清楚这句话背后隐含的意思，从而解决一些微妙的语义问题。

如图5-1所示，这个模型能够通过你提供的可以明显观察的句子，推断出一个人的心情状态（开心或难过），并得到最后的行为判断（宅、购物、社交），即通过已知推断出未知。

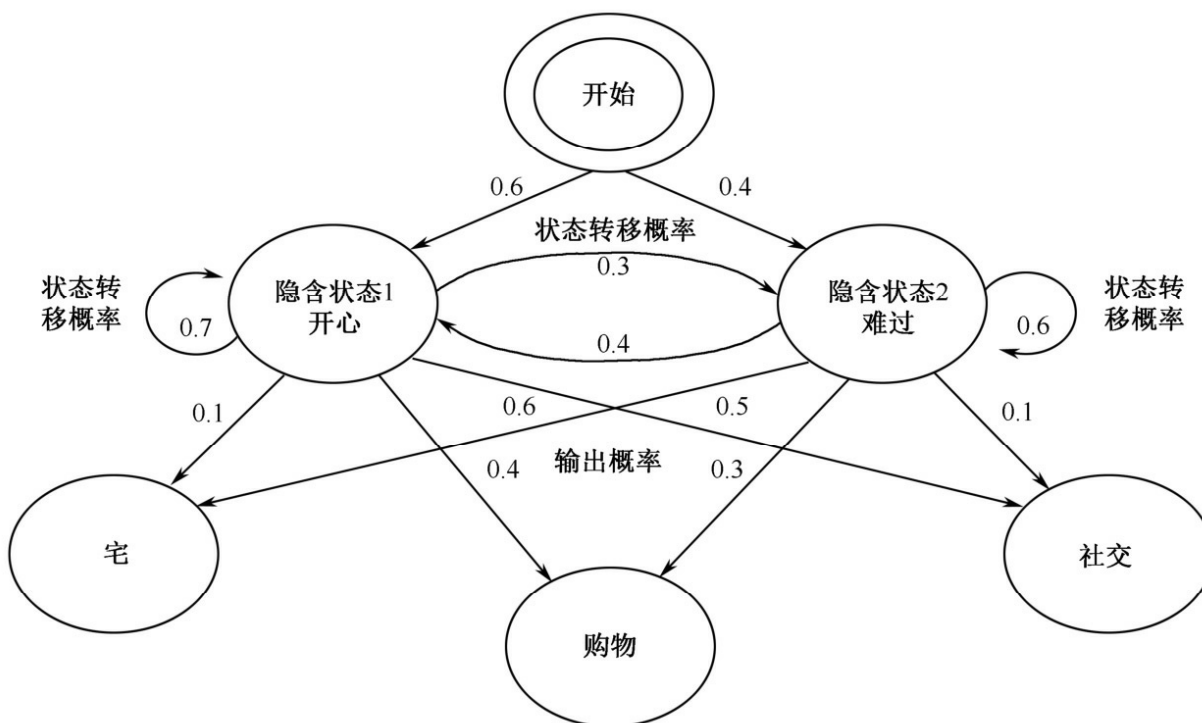


图5-1

如何优化这个模型，得到最优隐含状态？人们提出了许多解决问题的算法，包括前向算法、Viterbi算法和Baum-Welch算法。这其中的奥妙，难以尽述。但不能否认的是，在深度学习的基础上，数据+模型可以很好地打造出一款AI同传翻译，数据越大，神经网络越好。即使翻译出来的结果不尽如人意，但只要我们建设足够大的数据库，建立更好的模型，打磨算法，AI同传很快就会有更大的提升。

如何打造高质量AI同传？

除了增加数据库和打磨数据模型，AI同传还可以从哪些方面提升呢？我们不妨借鉴一下其他的技术。图5-2中，这四个方面代表了人们在NLP领域的一些进步。用金字塔形来表示这四个技术

之间的关系，难度是逐级上升的。



图5-2

目前，人们在聊天机器人和阅读理解方面已经取得了很大的突破。而AI阅读理解技术的进步不只是NLP的高阶进化，更深一层的意义是，科学之间是相通的，技术之间可以互相借鉴，金字塔顶端技术可以反哺底端。

在自然语言处理上，人与AI的区别是人有经验知识。即人们在听到某个字时，会自然地联想到后一个字，或者会因为一个词触发了一句话的联想。比如，我们听到“中”，既可能想到“国”，也可能想到“间”。但是AI“联想”的词却依靠数据。它说“北”，如果输入的数据不变，那后面跟的就是“京”。

我们曾经在《AI在阅读理解领域开始“跑分”，这个“人类好帮手”还能去哪炫技》一文中总结了AI阅读理解的技术层面，我们或许可以从中得到用AI阅读理解技术反哺AI同传的方法。

AI阅读理解技术的流程如下：Embedding Layer（相当于人的词汇级的阅读知识）→ Encoding Layer（相当于人通览全文）→ Matching Layer（相当于带着问题读段落）→ Self-Matching Layer（相当于人再读一遍进行验证）→ Answer Pointer Layer（相当于人综合线索定位答题）。

综合来看，阅读更偏向的是Multi-turn，即做完一次输入输出后，要把结果作为下轮输入的一部分继续输出，系统在运作时需要考虑上下文。而翻译则是Single-turn，即一句话进一句话出。

将这种方法合理利用后，机器翻译即使现在是Single-turn，将来也有可能是Multi-turn。AI同传现在没用到上下文背景，将来它也有可能结合上下文使翻译质量更佳。

如今，创作还是AI正在摸索的领域，而一旦在这个领域有所突破，将一些技术应用到AI同传中，我们或许可以达到翻译的最高境界——“信、达、雅”。

在未来，AI会不会挤占人类同声传译员的位置？

AI同传会取代人类翻译吗？当然不会。先不说语言本身的复杂性，我们可以看看同传的实际应用场景。

在实际工作中，不论是口译还是直接对话，都需要同传来完

成，也就是说，AI同传不仅要学会翻译，还要学会聊天。在这方面，机器还有很大的进步空间。那么，AI同传的用处在哪里呢？

1. AI共享同传，仅针对普通人的市场

人们出国旅游常常会遇到语言沟通等问题，然而，并不是每个人都有专业的口语翻译。这时，如果一个可穿戴设备或一部手机就能为你同声传译，想必会减少很多人的出国成本。随身携带一位专属的“同声传译”，是不是很酷呢？

智能硬件一直是AI的热门领域。2018年，微软和华为合作，在Mate 10手机中嵌入了微软的神经网络机器翻译，可以算得上是在终端运行神经网络机器翻译的第一例。

如果AI同传的硬件设备得以普及，商业模式可能转变为以出租或共享为主，即按需求进行租用，有一个专门的技术公司负责租赁，正如共享单车一样，使使用费降到极低。而这类AI的应用场景并不在复杂的会议现场，而是在日常生活、外出旅游等场景中，语料库的建设也会更加简单。

AI同传只会更加惠民，却不会取代如在金融会议、医疗会议中的更加专业的人类同传。

将AI同传与硬件设备相结合，创造切实可行的语音接口，还可以在很大程度上提高用户在移动终端、可穿戴设备、智能家居、智能汽车等智能设备的体验，真正在交互层面实现智能时代的人机结合。

2. AI同传成为同声翻译的考官

同声传译需求量成倍增加，但是合格的同声传译的数量增长却非常缓慢，据了解，现实市场上能够将十句话翻译出十句的同传译员寥寥无几。同时，拥有高级口译资格证书的人并不一定能胜任同声传译，同声传译还需要进行专业的技能训练，而有些合格的同声传译人员也并不一定有口译证书。

目前，我国还没有一个固定的机构来负责同声传译的相关事宜，也没有一套统一的标准对同声传译的工作进行考评。

面对这样的困境，我们或许可以在AI同传上开个脑洞。

人们可以利用AI数字化、标准化等特点，以数据库为依托，将AI训练成单一功能性的考核机器，针对不同的应用场景，对同声传译员进行考核和评级，从而规范人才市场。

这里或许可以参考驾驶培训机器人。驾驶培训机器人包含高精度GPS导航技术、惯性技术和虚拟传感技术、视频检测、数据处理、无线传输、指纹身份识别等高新技术，能够精确记录、判断驾驶人操纵驾驶机动车的真实能力。

同理，AI同传也可以在各种场景中观察、判断考生的翻译能力及考生对翻译规则的熟悉、理解程度。这个系统可以减少考试员的劳动强度和人为因素，确保考试公平、公正，考核方法科学、准确。

简单来讲，我们的目标是通过智能机器，使考核自动化，选

拔或训练真正的人才，而并非让它取代人类翻译。

更有意思的是，在考核过程中，AI能不断吸收新“养分”，增加口语类文本语料库，我们何乐而不为呢？

人类什么时候才能听懂动物的语言？

2018年，有一个视频在养宠圈中广泛流传，引无数养宠人士潸然泪下。

视频的主角是动物行为专家HeidiWright和一只生命即将走到尽头的导盲犬，HeidiWright以她的能力为媒介，将导盲犬的肢体动作和声音翻译成人类语言，帮助它和主人进行最后的交流。在HeidiWright的转述中，导盲犬表示它为无法继续守护主人感到惋惜，还不停地呼叫另一只狗伙伴，让它照顾好主人。

这段视频让人们感动的同时，也让许多人感到遗憾，因为绝大部分人都无法像上文提及的主人那般幸运，能够倾听到狗的心声。

无数人曾设想过，是否有可能出现一种翻译工具，能够将宠物语言转换为人类语言呢？

人宠语言互译并非伪命题，十年内或可“美梦成真”

Slobodchikoff教授称，未来5到10年，人类使用一种手机大小的装置——宠物翻译器，便能与动物进行“对话”。

这位北亚利桑那大学的生物教授花了30年时间研究草原土拨鼠的行为，他用AI软件记录并分析草原土拨鼠的叫声，将叫声翻译成英语后，发现这些草原上的小家伙们“具有语言所有方面的复杂通信系统”。而目前，他正试图筹集资金来开发猫和狗的语

音翻译设备。

不过，在这条未知明暗的道路上探索的显然不止他一人。“宠物翻译器”的低配版就被放上了淘宝，取得了可观的销量，卖家声称这个设备经实测翻译准确率高达80%。纵览评论，“好玩”“有意思”“灵气”之类的好评不在少数，从中我们也可以得知消费者对宠物翻译器的需求很大。

如果按Slobodchikoff教授所说的，这项技术或许能在十年内成为现实，它能在小范围内满足人与宠物的交流沟通，在大范围内满足人类一统动物世界的梦想。

自动语音识别技术和语音翻译技术助力，宠物情绪传达不再是“镜中花、水中月”

我们发现，低配版宠物翻译器运用的技术原理就是对狗的叫声、动作等生物信号进行采样，对获取的数据进行频谱分析，把得到的翻译语言以中文形式语音播报出来。但是由于采样的范围和机器内存等局限性，这种低配版宠物翻译器在翻译的准确度和丰富性方面尚有待提高。

不过，现在也有了好消息，为实现人狗沟通而设计的“No More Woof”耳机就是其中之一。

“No More Woof”是由北欧发明与发现协会（NCID）开发的，应用的是三个不同技术领域的最新技术的组合，即脑电图（EG）E传感、微计算和专用脑—机接口（BCI）软件，它主要由脑电图耳机、Raspberry Pi处理器和一款便携音箱组成。

这些传感器是脑电图记录器，它可以降低读数，减少离子电流在狗脑中的电压波动。然后由微型计算机拾取波动，在这种情况下形成一个Raspberry Pi，并对它做出解释。

例如大脑中有一种特定的电信号用来定义疲劳感，还有一些最容易被发现的神经模式：“我饿了”“我累了”“我很好奇那是谁”“我想尿尿”等。耳机中的传感器会捕捉这种特殊的电信号，并将它们转化为人们能够听懂的语言。

再结合基础的自动语音识别技术和语音翻译技术，计算机算法可以大致地分辨出宠物的情绪，这些是短时间内宠物语言翻译能实现的。至于要通过宠物翻译器来了解动物伙伴们真正的内心世界，我们还期待人类进一步的“大动作”。

如果要达到精确翻译，还需要解决哪些问题？

动物的大脑并不如人类的大脑复杂，人脑的活动通常有一个明确的目标导向，动物的大脑却不一定。人的各种语言之间的转换也具有相对窄范围的对应关系，而动物的语言与人类的语言则对应范围很宽。

例如狗会发出急促的叫声，可能是因为它想要向主人乞食，也可能是因为警惕陌生人，还可能是因主人不陪自己玩而生气。如果它想表达的是这一种情绪，而AI的翻译器却传达为另一种，那么就很容易使人和宠物之间的沟通误入“歧途”，从而完全丧失宠物语言翻译的意义。

那么是否有可能通过AI实现完全精准的宠物语言翻译呢？目

前来说还有一定难度，在我们看来，AI在宠物语言翻译上想要有所突破还得克服以下这些困难：

1. 数据关

要明确动物语言所表达的具体意义，我们需要先对动物的叫声和即时脑电波动进行完整的采样比对，再在这些数据的基础上建立数据库。

而这两种数据都具有广泛性和多样性。以犬类为例，不同的犬种声带粗细宽窄各不相同，针对同一情景发出的叫声分贝高低和尖细情况也不同，而刺激犬类发出叫声的场景又是难以穷尽的，单收集犬类的声音样本就是一个无比巨大的工程，数据库自然也难以完善。

2. 技术关

一个AI翻译产品做到翻译精确至少需要攻破几个难题：形式端，拍译要攻克图像识别，同声翻译要攻克语音识别；内容端，攻克文本语言分析、大数据。而AI还没有发展到能够精确地处理这些问题的阶段，机器缺乏对视觉场景、听觉场景、自然语言处理的常识判断。

如搜狗搜索在2017年6月的分享会上发布了创新产品搜狗翻译App，它应用了基于生物学习的神经网络机器翻译（NMT）系统，将翻译精确度提升到了一个前所未有的高水准。然而它在翻译效果的“信、达、雅”上，仍然只达到了“信”的层面，对语言背后的幽默、情感等丰富含义的解读和人们所期待的水准还有些距

离。

3. 语义关

语料积累、场景收集和副语言与文化背景成痛痒之地。AI翻译在文本或语言的寓意分析方面做得还不够好。与人类语言相比，动物语言都是即时信号，信息内容全部关乎当下，或示威，或示警，或示爱等，我们从中看不到用语言激起对过去联想的迹象，并且单个个体能发出的声音形式太单一了，蕴含在其中的丰富信息难以明确表达。

宠物翻译的难点不仅在于声音的收集，更在于声音背后具体含义的对应。

这种对应是宽范围的，难以精确判断的，机器缺乏对视觉场景、听觉场景、自然语言处理的常识判断，无法精确理解语音所表达的内涵，甚至在这个方面还比不上人类对动物语言的理解。人可以根据生活经验来理解动物语言，例如看到狗冲着陌生人叫，人们可以推测它是在防备这个陌生人，而机器可能就没办法很好地理解，从而导致判断错误。

4. “历史包袱”，AI难以跟上生命体的学习进程

狗的叫声在一定历史时期并不是一成不变的，狗凭借自身的灵性及主人的后天驯养，具备学习能力。例如狗类中智商排名第一的边境牧羊犬智力水平已经相当于6~8岁的小孩，经过学习，在放牧时，它会用不同的叫声来驱使羊群，控制羊群走向。

还有一些宠物狗，甚至会在人类的刻意训练下发出类似“妈妈”的叫声，宠物语言在日新月异的变化，计算机却很难去掉语言的“历史包袱”，这些也造成了AI宠物翻译的困境。

总之，AI能做的就是不断改进自身的功能，我们要用科学手段完善数据库、内容、语料和场景，将形式和内容双管齐下，在坚实的地基上建立起实现人和动物“有效沟通”的“巴别塔”。

AI能识破坑老人钱的套路吗？

“我们的产品绝对有效，还会赠送超值礼物，数量有限，先到先得！阿姨，您可要抓紧了……”在接待者的热情宣传中，老人们面带笑容，拎着袋子，心甘情愿地买下了所谓的“干儿子”“干女儿”推荐的产品。

老人被骗，已经不是新鲜事了。特别是在科技不断发展的今天，“中老年人+互联网”这样的组合，简直就是“良善好欺”的另一个代名词。

《中老年互联网生活研究报告》显示，在互联网中上当受骗过（或者疑似上当受骗过）的中老年人比例高达67.3%。他们被骗的主要渠道是朋友圈（69.1%）、微信群（58.5%）及微信好友（45.6%）。

老人被骗事件频发，AI是时候“挺身而出”了。

AI防骗，从哪里开始？

用AI防骗之前，我们要先了解骗子行骗的过程。以电信诈骗为例，近年来，电信诈骗产业发展迅速，从以往“猜猜我是谁”的盲打式诈骗到手握个人信息的精准式诈骗，诈骗分子的骗人手法、装备、技术都在不断更新。

还有一些骗子会利用老人需要陪伴的心理，通过各种花言巧语与老人们建立情感联系，最后再说服老人购买他们的保健品或

理财产品。

这些诈骗技术不仅迷惑性强，而且受骗人群也越来越多，提高老人们的防骗意识和能力已迫在眉睫。那么，AI可以从哪些方面入手呢？

1. 验证，验证，再验证

很多时候，老人其实是非常谨慎的，他们在和一些人建立联系后也会试图进行一些搜索，例如查一查这个人的公司名字、社交账号等。但是骗子们往往会把自己伪装得很好，公司的业务和他自己的职位看上去都是如此的真实合理，老人们简单的网络侦察根本发现不了其真面目。

利用AI进行网络侦察就会有效得多。在图像识别上，AI可以使用对方朋友圈里的照片进行反向搜索，看看对方的照片是否在其他网站上被重复使用。

2018年，Facebook表示其正在建立一个新项目，试图通过最新的机器学习算法来标注可疑消息，并将其发送给第三方事实核查人员，以阻止虚假消息的传播。Facebook称，通过“机器+人工”的方式，虚假消息将成倍减少。

虽然这些技术并不是针对防老人诈骗的，但是也算是为老人验证信息安全提供了一个思路。希望在未来，AI可以帮助老人判断信息，并告诉老人：别信，这是骗人的。

2. 让机器去交谈

美国联邦贸易委员会的高级律师Patti Poss曾指出，诈骗分子往往试图尽快将“目标”带离他们最初相识的平台。一般是从社交软件到电话短信的转移，或者是线上到线下的转移。

老人容易受情感和固定思维的影响，对这些人不太设防，觉得已经有了社交软件上的联系，再打一个电话也没什么大不了。事实上，一旦转移阵地，骗子就会得寸进尺，最后登堂入室，实行诈骗。

AI从其诞生之日起便与逻辑学密不可分，二者的共同发展促进了用机器模仿人类思维的智能学的进步。所以，在这个时候，不妨让AI去更加“机智”地交谈，寻找骗子的逻辑漏洞。

例如我们可以利用AI找出对方的语法错误，并将一部分信息内容与数据库进行对比。在此前，AI可以深度学习已有的案例，掌握骗子的语言技巧，从而发现骗子的“套路”。国内的蚂蚁金服也建立了这样的互联网防骗知识库。

情感是AI防骗的关键

在发展心理学中，有一个理论叫作社会情绪选择理论。该理论认为，年轻人的需求是受未来导向的——我们交朋友和上学都是为将来的事业做准备；而对老年人来说，他们所剩的时间已经不多了，未来导向的目标不再适用，他们就会更多偏向与情感相关的目标，这也就是他们更关注亲情、友情等的原因。

我们也发现，骗子都是从这里入手：首先打亲情牌，和老人拉家常，混个脸熟。接着给老人一些小优惠，如免费体检，然后

就开始推销产品了。

所以，防止老人被骗的最好办法就是隔断骗子与老年人建立情感联系的可能性，而在这方面，AI有两个方式可以着力。

第一个方式是取而代之，也就是我们常说的陪伴机器人。利用语音交互等功能，它可以进行相对密集的情感交流，进而使老人对陪伴自己的机器人产生依赖性。在遇到可疑的事情时，老人也会第一时间找AI查询，而不是上当受骗。

但是，谁又能保证AI是一直可信的呢？

机器人系统一旦遭遇病毒或黑客攻击，就会完全瘫痪，这对老人的打击无疑是巨大的，毕竟寄予了情感的机器可不是能随意置换的家具。而更可怕的结果则是，系统中储存的家庭隐私可能会被泄露，更加便于那些心怀歹意的人行骗。

在未来，可能会出现更高级的AI来实施欺诈，利用AI检测欺诈和犯罪是一个双方斗智斗勇的过程。所以，最有效的防骗渠道应该是第二个方式，让AI在防骗之余，成为父母与子女之间建立情感的枢纽。

例如我们可以在智慧家居设备里连接两个家庭，让家居产品实现社交化，包括聊天机器人、日程功能等。而且家居产品之间能够共享信息，它们能远程设定程序，帮助解决父母与子女之间沟通的问题。父母和子女都可以通过家里随处可见的家居设备设定关心程序，例如定时洗衣，定时电话等，让老人有被子女随时关心的感觉。

AI能顺利进入老人家里吗？

有数据显示，在有被骗经历的中老年人中，他们的受教育程度集中在初中和高中学历，分别为39.4%和37.7%；中等收入和高收入中老年人居多，分别占67.1%和24.3%。

从经济自主性来看，有受骗经历的中老年人中有41.1%表示家里的重大支出由自己决定，37.5%的中老年人表示是由家庭成员共同协商决定的，仅有16%和5.4%的中老年人表示是由配偶等其他支配的。

从这些数据中我们也能看出，被骗的中老年人中很大一部分社会经济地位较高并具有经济自主性，而这类老人几乎都有一个通病，那就是“要面子”。

《中老年互联网生活研究报告》显示，在受骗后，有68.3%的中老年人表示“不寻求帮助，当经验教训”，只有25.9%和17.9%的中老年人会选择向子女和朋友求助，而表示选择报警求助的仅有0.6%。

因此，让老人求助于AI恐怕也不是什么容易的事情。

一个原因是老人在没被骗之前都不觉得自己会被骗，没必要也不愿意购买一个防止被骗的高科技产品（目前的AI产品成本都不会很低）；另一个原因则是AI自身的可信度还不够高。

在2018年，就有一群网络黑客利用AI非法收集数据来进行网络诈骗。AI确实能防止老人被骗，但是AI本身也曾是诈骗集团的

诈骗工具，这就让被帮助的老人心里有点不是味儿了。最重要的是，人们在宣传AI防骗时使用的“高科技”“大数据”“识别率非常高”等此类高频词汇，是不是也像极了那些推销保健品的骗子呢？

“高考后综合征”：AI能不能发现高考生的心理疾病？

每年的六月都是“高考月”。

高考后，考生需要查询成绩、填报志愿、等候录取等。高考学子面临的将是一系列的紧张流程，一环扣一环，大多数考生（包括部分家长）一边期待着，一边焦虑着。而社会上绝大多数人只是兴致勃勃地等着各省状元的出现，称赞他们为学霸后，便潇洒离开，却很少有人会关注这些考生在高考后的心理动态。

对于高考生而言，学业的重压突然消失，那颗总是沉甸甸的心也无处安放了。面对这样的情况，我们可以把评估、预防和治愈考生心理的重担交给AI吗？

用AI来寻找心理防线崩溃的学子，可行吗？

这听起来不太可靠，但AI应用于心理学上的研究早就开始了。就在2018年，IBM的计算精神病学和神经成像研究小组团队就尝试利用机器学习预测人患精神疾病的风险，其预测精准度达到了83%。除了预测精神疾病，Facebook也推出了AI防自杀系统，通过识别用户社交软件上的文字、图像和音视频来判断用户是否有自杀倾向。

而将这些AI项目应用于患有“高考后综合征”的高三学子是否有效呢？从理论上来说这是可行的，AI可以利用机器学习建立相

关模型，定期、批量的对青少年心理状况进行评估，进而发现青少年的心理问题。但实际上，想要AI来干预高考生的心理治疗，难度恐怕很大。那么，问题在哪里呢？

1. 考生想要的公平，AI不仅体会不到，还在不断制造“不公平”

要想治好这个病，必须搞清楚这个病的由来。“高考后综合征”典型的状态分为两种，一种是过于放松，就好比一根弦绷了太久，骤然崩断，与高考前的心理产生了巨大的反差，使学生们产生了心理和生理上的不适应。

另一种是悲观压抑，考生们对于高考成绩惴惴不安，常常会感到失落、悲观，甚至会觉得对不起父母家人而内疚，无颜面对老师、同学，产生焦虑情绪。

而以上这两种病症，归根究底，还是因为人们太过看重高考，对“高考决定命运”感到深深的恐惧。

网上曾流传着一个各地录取分数差异的段子。

北京考生：“老爸，我考了600分，比一本分数线高53分，估计我能上北京大学！”“儿子真有出息！走，我们去上海旅游！”浙江考生：“爸，我考了600分，与一本线差了20分，我落榜了。”“真没出息，别上大学了，跟我到上海打工去吧！”上海考生：“爸，我考了600分，估计上复旦大学没问题，干脆送我出国吧？”“行，你出国学个工商管理，正好回来帮我。”

这个笑话当然代表不了什么，但其中影射的现实才是现在的学子们最害怕的。从分数到志愿，再到高校录取，学子们都希望得到一种真正意义上的公平，每个人都可以通过自己的努力实现自己的学府梦。

而AI能读懂学子们这种隐隐的恐惧吗？我个人觉得，凭借机器学习读懂学子们内心深处对于不公平的恐惧确实很难。因为在现实生活中，连AI都在为这种不公平的环境“添油加醋”。

可以说，目前的AI教育产品解决的都是“高考前公平”，因材施教也好，娱乐教学也好，这些产品在AI大数据的加持下，形势可谓是一片大好。例如猿辅导在2018年获1.2亿美元融资，作业帮也不甘落后，获得1.5亿美元融资，紧接着，VIPKID的2亿美元融资更是刷新纪录。

然而，过了高考这个节点，那些诸如扇贝、百词斩之类的英语教育产品一下子就转到了大学英语教育上，如四六级、雅思、托福等，却少有AI产品关注或以普惠性的手段解决“高考后公平”的问题。然而，高考后的公平才是众多学子关注的，这恐怕也是许多“考600多分却名落孙山”悲剧故事的起源。

2. 病因复杂，病情模糊，AI难以捉摸

其实，考生患上“高考后综合征”还有一个诱因，就是社会的聚焦。社会上一些人或机构也有不少“病症”，如“状元追捧症”“虚与委蛇症”“借机恶炒症”“自吹自擂症”等。

可以肯定的是，生物、心理与社会环境诸多方面的因素参与

了“高考后综合征”的发病过程，而“高考后综合征”复杂的病因主要缘于人性的复杂。人性这么难以捉摸的东西，人们都尚且不自知，AI的机器学习更是难以全盘掌握。



目前比较成功的AI预测精神疾病项目，如上文提到的IBM预测抑郁症的系统，主要是通过分析被检测者的语言方式和语言连贯性进而确认其患病风险的。而“高考状元”和“落榜学子”虽然看起来是两类人，但因为他们受到的是同样的学校教育，掌握的也是同一套语言系统，面对这类患病人群，AI恐怕会“精神错乱”。

不论成绩好坏，有很多高考生都会患上“高考后综合征”，但真正表现出明显症状的人很少，因为在众多考生看来，不管是空虚还是焦虑，都只是一种短暂的不适应而已。没有自杀、没有抑郁，AI要建立什么样的模型才能发现这些心理防线已经崩溃的人呢？

那么，AI能为高考学子们做些什么？

用AI来发现那些心理防线崩溃的学子似乎不太容易，那怎么办呢？我们就只能任其发展了吗？当然不是。

对大部分人来说，跨越地域的变化，技能的重新输入，学习环境的适应，都是需要时间的。而对有些人来说，这种适应甚至无法靠自己实现。所以，除了评估考生高考后的心理，AI还能够帮助学生适应角色，解决两类问题，一类是理性地为学生和家长做情景分析；另一类是找到突破口来监控和预防可能发生的情

况。

1. 理性的分析判断很重要

在长达三个多月的等待时间里，其实最令考生和家长煎熬的就是填报志愿了。在分数揭晓后，与其费尽心思地安抚那些没有考到理想分数的学生，不如给他们提供有用的建议。

高考填报志愿其实更像是一场博弈，在信息不对称的情况下，由于绝大多数家长和考生都不具备数据分析的专门知识，就很容易做出错误的决策，这就增加了高考志愿填报的不确定性。

考生即使知道了自己的分数、排位和控制线，也不等于能够百分百的被自己填报的学校和专业录取，如果不选择调剂，考生还可能存在落榜风险。

在数据分析上，想必AI是不肯落后的。在这个方面，AI可以通过汇集近几年的高考数据，来分析考生被某高校录取的概率，同时，依据考生分数给予不同的填报方案。

就在2018年高考后，映客就联手了百度教育，结合OCR、智能检索技术、知识图谱和大数据分析技术，打造出一套分数评测方案，为考生提供极速、精准的智能在线估分系统。

2. 最重要的是建立家庭数据库

网上有一段时间很是流行“原生家庭”这个词。原生家庭，指的是孩子出生以来的第一个家庭，也就是父母打造的家庭。谈论

它的人有很多，而苦恼于原生家庭带来的阴影的人也很多。

高考对绝大多数家庭的影响是十分重大的，考生的许多压力往往也来源于家庭。这或许就是我们利用AI来预防和监控“高考后综合征”的突破口。我们可以利用AI收集和标注考生的家庭背景，准确追踪到压力产生的源头，进而捕捉到“高考后综合征”的苗头，防止心理上的焦虑往更加严重的方向发展。当然，这只是一个预想，在其具体落地后，家庭隐私问题也是需要慎重考虑的。

世界编剧队伍里为什么突然多了个AI编辑？

AI又完成了一场“48小时电影挑战”，继两年前AI在伦敦科幻电影节上创作名叫*Sunspring*的短片之后，*Zone Out*是第一部完全由AI独立制作的科幻电影。这一次AI不仅采用了“换脸”技术构造电影人物，尝试使用神经网络生成对话和配音，还继续连任了电影编剧一职。

AI如何写剧本？不了解2018年AI进展的人可能会觉得是不是又要开始吹捧AI了。其实，AI在文学领域已经有了不少建树，微软机器人“小冰”都推出了自己的第一部诗集。因此，文学并不是专属于人类的一块“处女地”。这次完成*Zone Out*剧本创作的AI，名叫*Benjamin*，而它的剧本创作大部分依赖LSTM（长短期记忆）神经网络。

首先，*Benjamin*需要接收数十部科幻电影剧本的信息。然后，它会把剧本内容分解到字母级。为什么要分解成字母级呢？这主要是为了方便它预测在科幻电影剧本的创作中一般哪些字母有更大的概率会被放在一起使用。在完成了这一步学习之后，*Benjamin*会生成由自己创作的句子，而非简单地从输入的语料库当中的句子。当然，只会写句子称不上是一个合格的编剧。*Benjamin*可以提取剧本与一般文学作品类型不同的特征，*Benjamin*也因此学会了模仿剧本的结构。

好的作者并不等于好的编剧。但当国内的大IP被改编成影视剧的时候，启用原作者担任编剧常常会被观众看作小说改编影视

的质量保证。而实际上，我们必须知道的是小说作者与影视编剧的工作是完全两种不同的工作模式。对于投资方来说，启用原作者的唯一目的就是将其作为卖剧的一个噱头。不管是导演中心制还是演员中心制都让中国编剧的地位显得无足轻重，而反观韩国，一部韩剧如果想要造势，在他们的宣传工作中必定会用到的一个关键词“金牌编剧”。目前，韩国电视剧的制作环境已经越来越以编剧为中心，而观众对作品的选择也往往会把编剧作为重要的考虑因素。编剧，是韩剧制作当中的绝对王者。

上述AI来担任编剧的例子虽然让人觉得颇为神奇，但是让它完全取代人类编剧可能还难以被人们认同。那么，AI有办法拯救一下中国疲软的编剧行业吗？

AI剧本画像，给中国编剧的一条新出路

《淑女的品格》一剧在之前的微博热搜曾一度接近2亿的阅读量，受到全民关注。为什么《淑女的品格》一经转发就能得到数亿网友的认可？其实，在这条热搜的背后，我们应当看到在主流影视作品的收看人群中一直隐隐存在一股很强大的意见，这些意见没有找到一个合适的出口爆发——“除了校园爱情、婆媳大战这些桥段，电视剧还能不能有点新鲜的剧情了？”公共舆论管理当中有一个叫作打捞沉没舆情的专业说法，而《淑女的品格》得到立项可以算得上是网友沉没的舆情被打捞起来的一个典型表现。影视作为意识形态领域的创作，不管在哪个国家都会受到限制。从目前中国的影视质量来看，各大编剧想象力匮乏，甚至谈不上是审查压迫了他们的创作。

大家期待一部什么样的剧，到底能不能得到一个精准的预测？

《淑女的品格》的立项过程或许可以给这个问题的答案一点启示。一个小的脑洞最后引发了社交网络上的狂欢，这看上去与“蝴蝶效应”是不是有点相似？一只南美洲的蝴蝶，偶然扇动了几下翅膀，就可以在两周后引起美国得克萨斯州的一场飓风。如果AI能预测类似“蝴蝶效应”这样的混沌系统，那么散落在互联网各处关于影视讨论的声音是不是就都能够被打捞起来，而AI将从这些被打捞的声音当中完成一个“剧本画像”的工作，也可以提前预知这场“飓风”的样子。美国马里兰大学的研究表明，AI可以预测混沌系统的发展趋势。例如预测模型火焰锋面的混沌演进过程。老牌混沌理论学家采用了一种机器算法学习原型混沌系统。在AI学习了这一系统之后，他们发现比起之前的预测方法，机器能预测到的未来是以前能预测到的八倍，而且预测效果几乎和真实情况完全匹配。

虽然就目前预测混沌系统的AI建立在动力学的基础之上，如果想要将其挪用到影视“剧本画像”的工作上可能所需的并不是同一套操作方法。但是这并不影响我们对该思路进行借鉴，可以预测舆论爆发点的AI能够轻松帮助影视公司确定一个影视创作的具体轮廓。数据将会告诉编剧们遵循这样的剧本轮廓去创作失败的概率是多少，而影视公司也可以不必担心收益问题尽管大胆地投资。

不止写剧本，**AI**还想谈谈著作权

2017年10月25日，在沙特举行的未来投资计划大会上，沙特阿拉伯授予了机器人索菲亚“公民身份”。在人类社会当中，身份认同一向是件很重要的事情。而Benjamin创作了科幻电影剧本，我们不禁想问它是否也应该获得一个编剧人的身份？因为只有在获得身份的同时，AI才能获得权利。在这里，这份权利叫作著作权。

1. 机器创造的作品该有著作权吗？

日本政府下属的知识产权战略本部在2016年曾宣布要讨论制定对AI创作的音乐和小说等的权利进行保护的法律。当时，根据日本的《著作权法》，只有人类的作品享有著作权，而由AI创作的作品即使被盗用，也无法采取措施禁止和要求损害赔偿。但随着现在越来越多虚拟歌姬、虚拟诗人、虚拟编剧的出现，关于机器人作品的著作权问题引起了许多学者的关注。有学者主张，由AI机器人独立“创作”的作品应该承认其本身享有著作权。但是从另外一个方面来说，创作被定义为一种智力活动，包含了人的思维、情感和表达，只能由有血有肉的自然人来实施。尽管，科学家们探索通过制造“人工神经细胞”来模拟人体中的神经元以达到模拟人脑思维，但是人脑的复杂程度决定了在未来很长一段时间里AI要真正模拟人脑的思维活动几乎不可能。尽管机器人在很多方面都在超越人类，但是由于它进行的活动都属于无意识活动，因此，它也有理由被认为没有资格获得著作权，而这将有可能对投资AI形成比较大的障碍。

2. 保护人类的著作权，AI“身先士卒”

如果你了解AI，那么你就会知道其实AI是这样一个“物种”：保护人类是它们最初被赋予的使命，就算争夺自己的著作权无果，但对于人类的著作权，它们会坚决捍卫。如何捍卫？例如数据家黎晨曾用机器学习分析过《鬼吹灯1~4》是不是天下霸唱所写的。对于写作抄袭这件事，我们需要先给出一个鉴定的建议：一个人写作的内容会经常改变，但是人在不经意间养成的习惯是不会变的。这些习惯指的就是对于副词、助词和介词的使用。因此，只要分析天下霸唱在副词、助词、介词使用上的特点，就能够找出鬼吹灯的前后四部是不是由同一个人写的。机器学习对此种分析得心应手，其过程也可以简单地被梳理为：选取特征词、计算词频、PCA降维画图、机器学习和结果分析。最后，机器鉴定的结果和广大书迷的感觉一样，鬼吹灯的前后四部完全不是同一个人写的（鉴定结果仅供参考），这就是AI的能力。在“洗稿”、“洗剧本”等劣行盛行的今天，AI或许能够为此类乱象给出一个明确的解法。

AI和艺术的融合会唤起更多人对技术的关注和思考，纽约大学的帝势艺术学院和伦敦大学金史密斯学院也开设了针对艺术创作者的机器学习课程，而这或许是冰冷的AI在无尽的未来中可以紧握的一点温暖。

从自主系统开始，苹果公司要在无人驾驶上重走手机之路

在最近的资本市场上，苹果公司颇为不顺。

截至2018年11月23日，苹果公司以172.29美元报收，此价距离当年8月2日233.47美元的历史最高峰下跌了20%以上，市值也蒸发超过2200亿美元，萎缩到8175.85亿美元。在苹果公司宣布停止公布手机销量后，以富士康和和硕为代表的供应商下调新款iPhone销售预期，市场一致看衰的背景下，苹果公司用什么来提振投资人的信心呢？

无人驾驶被苹果公司拿出来说了。前段时间，有媒体报道，库克证实苹果公司正在研发用于无人驾驶汽车的自主系统，这也是库克为数不多的第二次公开披露苹果公司关于无人驾驶方面的计划和进展。或许苹果公司要为它的无人驾驶项目提速了。

雷声大、雨点小，苹果公司无人驾驶研发成果秘而不宣

相比特斯拉这类竞争者，苹果公司进入无人驾驶的赛道稍晚，它在2014年才开始组建团队，以“Project Titan”作为内部代号。进入的时间晚，并不意味着苹果公司对无人驾驶项目不重视，当库克在2017年6月首次对外公开其无人车战略的时候，将无人驾驶技术提到“所有AI项目之母”的高度。

于是苹果公司挖来了特斯拉公司负责整车研发和制造的高级

副总裁Doug Field、大众汽车集团的首席数字官Johann Jungwirth、福特车身结构和冲压专家Aindrea Campbell、保时捷919技术总监Alexander Hitzinger、特斯拉公司负责工程研发的副总裁Chris Porritt.....2018年7月，美国联邦调查局指控苹果公司前员工窃取商业机密的诉讼文件曝光了苹果公司“Titan”项目的团队规模——5000人。

一开始，苹果公司就把“盘子”铺得很大。可是苹果公司除了逐渐扩大自己的测试车队（截至2018年9月已达70辆，这个规模仅次于通用汽车的Cruise和Waymo），向外公布的研究成果对于无人驾驶技术并没有突破性的指引。

比如类似于飞机的空中加油机，让汽车在行驶时通过“连接臂”共享电池系统；让无人驾驶汽车与iPhone、iPad或MacBook等苹果系统的设备同步；当汽车遇到紧急情况需要人类接管时，让它发送警报提醒正在使用这些设备的用户及时接管汽车。

根据自动驾驶初创公司Voyage联合创始人MacCallister Higgins在网络上放出一段苹果公司第三代自动驾驶测试车的视频来看，相比前两代测试车，苹果公司也只是对毫米波雷达数量进行调整，对传感器列阵进行了优化。

此外还有一些天马行空的想法：如何设计一个静音车门，如何设计没有方向盘和油门的内饰，如何把AR/VR设备放到车里，如何应用球形轮胎，如何重新设计一款更美观的激光雷达等。

迄今，苹果公司无人驾驶技术展现给大众的印象是，它的研

发更多停留在硬件和设计层面，苹果公司最为擅长的软件开发、生态构建等还没有任何风声透出。颇具玩味的是，苹果公司在2015年买下了3个与车相关的顶级域名：apple.car，apple.cars和apple.auto，但是至今还未启用。

用**CarPlay**接管无人车？苹果公司没那么天真

相信以苹果公司的高度，它不会没有认识到一套充满智慧的车载系统对于无人驾驶汽车的重要性。

2013年苹果公司进军汽车领域时就制定了“iOS in the Car”计划，并在次年的日内瓦车展上展出了合作伙伴搭载的CarPlay——一套可以将用户的iOS设备、iOS使用体验与汽车仪表盘进行结合的车载系统。苹果公司能用CarPlay来接管未来的无人车吗？从目前来看，CarPlay还不具备这样的能力。

用户对CarPlay并不满意。

“CarPlay支持的App太少了，它连最基础的专业导航都不支持。每当我被迫用起苹果iOS系统的自带导航时，就无比怀念百度地图和高德地图。”

“CarPlay与汽车连接使用时，经常受手机信号的影响。手机信号不好或行车抖动会导致连接断开。断开后正在使用的导航、音乐之类的应用也马上关闭，好几次差点出事！”

“苹果系统一升级，CarPlay系统就变得更卡，反应也越来越迟钝，点个图标也要等几秒。”

“升级iOS 12后，在使用数据线连接CarPlay时，另外的USB接口的U盘音乐不能播放，只能播放收音机、苹果系统的手机自带的或手机App中的音乐”

“一连接CarPlay，车载蓝牙就失效，这个问题产生很久了苹果公司也没有修复。”

当然最让人无法接受的是，大量用户反映连接CarPlay后Siri无法使用，而在苹果公司的规划中，Siri是CarPlay的核心——让司机在眼睛不离开道路的情况下通过语音完成操作。

CarPlay非常难用，福特与微软公司合作开发的SYNC也好不哪去。系统崩溃、触屏难用、反应速度慢这些问题也都在它们身上出现过，有些至今也没能解决。至于那些基于Android系统开发出来的车载系统，其稳定性和人机交互逻辑方面的问题就更多。

“小爱你好、小度你好、斑马你好、Nomi你好……”谁开车还能得记清那些开门暗号。无人驾驶赛道玩家太多，车载系统的研发同质化严重，对于普通用户来说，要想分清这些语言交互助手和它们所匹配的车型还很有难度。

很显然，无论是iOS还是Android都是基于手机的使用场景设计开发而来的，而汽车的使用场景和人机交互逻辑与手机完全不同，将iOS和Android稍做修改就搬进车内注定是不会成功的。从库克的这次表态来看，CarPlay可能成为苹果公司无人车自主系统的一个过渡产品。

自主系统是苹果公司布局无人驾驶的第一步

对于无人驾驶自主系统的研发，苹果公司无疑是有优势的，在苹果公司庞大的商业帝国中，苹果公司为它的Mac电脑开发了Mac OS系统，为iPhone开发了iOS系统，甚至连Apple Watch都有属于自己的Watch OS系统。那么对于无人车，苹果公司为什么不从底层开始，设计一套完全针对汽车驾驶场景的“Car OS”呢？

正如十年前手机行业面临的变革一样，无人驾驶技术也将让汽车行业产生翻天覆地的变化。在变化来临之前，是先做车（硬件）还是先做系统（软件）呢？苹果公司用iPhone的经验进行回答——用软件定义硬件、用新技术定义旧行业。

自主系统是苹果公司布局无人驾驶的第一步，然后就像用iPhone重新定义手机一样，用AI重新定义汽车。未来，汽车除了被用于出行，还将会是移动的空间，移动的计算终端，移动的能源终端，移动的摄像机，移动的温度计，移动的机器人……借鉴当前消费电子领域的成功经验，用一套烂熟于心的流程，建立一个“软件+硬件+服务”的全新汽车消费生态。

就像iOS（软件）之于iPhone（硬件），在自动驾驶无人车上，苹果公司在开发自主系统（软件）之后，它的无人车（硬件）在哪？

其实苹果公司一直都在寻找制造无人车的合适机会。由于苹果公司在汽车研发上缺乏经验，在保证现有业务体系不受影响的前提下，不可能像特斯拉那样的初创公司一开始就“赤膊上阵”，苹果公司走的是一条“合作造车”路线。

合作伙伴的选择一度让苹果头痛。主要原因是苹果公司太过强势，它想要主导权，但车厂不愿将自己赖以安身立命的造车数据交给苹果公司。

直到2018年5月，苹果公司才与大众公司达成合作协议，共同开发自动驾驶的无人车。不过项目是以对大众T6厢式车的改造开始，苹果公司重点对仪表盘和座椅等部分进行改造，还计划加入各种传感器和电子设备，底盘、车轮等动力机械部分。这或许只是苹果公司与大众公司进行深度合作，开发具有前瞻性质的自动驾驶无人车之前的一次试探与磨合。根据苹果公司的商业模式，它无论如何都不会放弃硬件，无人车也是如此。

IBM的人机辩论大赛，AI的胜利为什么名不副实？

就在旧金山的一间办公室里，IBM举办了历史上第一场人机辩论比赛。两道辩论题目分别是：我们是否应当资助太空探索和我们是否应当更多地使用远程医疗。而两场辩论，Project Debater的对手都不容小觑，一位是以色列全国辩论冠军Noa ovadia，另一位是以色列辩论专家Dan Zafrir。虽然对手强大，但是现场辩论Project Debater由于提供了更多有利的证据而更具说服力，最终观众的投票也倒向了Project Debater。

可疑，真的很可疑。机器赢了人类并不应成为我们的关注点，观众是否专业才是最该被重视的。在辩论的世界里，关于输赢有如下几种说法：①辩论辩的是真理。②辩论辩的是逻辑。是真理也好，是逻辑也好，反正比赛的胜负不可能仅由信息量的多少来决定。Project Debater在辩论场上的啰唆背离了用最简单的话把道理和逻辑讲通的基本要求。所以在我们看来，这场人机辩论大赛AI的胜利名不副实。

AI的辩论风格难立

自古强者都有风格，辩手也是。人称“辩魂”的少爷黄执中开创了辩论学派——“新剑宗”，他是亚洲系统建构辩论学理的第一人。但他的风格却并不好学，相传有“新手学黄执中必死”的论断。所谓不同的派别有不同风格，在金庸的小说《笑傲江湖》中，华山派25年前因葵花宝典之争被分为两家，一家主练气，称

为“气宗”；一家主练剑，称为“剑宗”，武学之路大相径庭，却也各有特色。而金庸本人与另一武侠小说大家古龙的风格差异也使他二人在武侠小说创作上各有其高峰。人说金庸的江湖再远都有一座庙堂，但是古龙的庙堂再高都是一片江湖。金庸写世，而古龙写人。

那么，话说回来，AI的辩论有风格吗？别说风格了，Project Debater连抑扬顿挫都还没搞清楚。可以说，如果没有人类辩手的参与，两个AI之间的辩论足以让一大片现场观众呼呼大睡。风格如何形成？高手自创风格，但是社会学家塔尔德告诉我们，创造是极少数的，而模仿是大多数的。加拿大一家初创公司——琴鸟发布了一款AI语音系统，它能够通过分析讲话录音和对应文本及两者之间的关联，在1分钟之内模仿人类讲话。琴鸟公司的AI系统使用的是一种模仿人脑思维的算法，能在倾听的过程中掌握每个人说话时字母、音位和单词的发音特点，然后推理并模仿这个人说话的情感和语调，即“风格学习”。琴鸟的AI语音系统还曾经模仿过特朗普、奥巴马和希拉里三人的声音，并让这三个人成功开展了一场“对话”。

不同于苹果的Siri，琴鸟的智能语音系统已经做得相当自然了。尽管如此，美国卡内基梅隆大学语言技研究所的教授迪莫·鲍曼还是表示这个AI系统尚不能模仿人们在讲话中的呼吸和唇部运动，因此我们仍然可以听出计算机的语音特征。而AI要真正地复制人声，还要再等几年。所以最后的结论是，对于目前的这些AI来说，风格难立。

AI的幽默属性尚在进化中

幽默在辩论中是使人信服的一个关键因素。在科幻电影《霹雳五号》中有这样一个桥段，一名逃跑的机器人有了意识，坚称自己有“生命”。而男主角最终测试它的方法是给它讲了一个笑话。在讲完笑话后，这个机器人发出了一连串笑声。这时，男主角才开始认为它真的具有自我意识。因此，也有很多人把机器是否有幽默感作为判断机器是否进化到具有人类思维的重要标准之一。

不管是1993年央视首创的电视辩论赛，还是现在风生水起的辩论网络综艺，或是各大高校每年一届的学生辩论比赛，幽默一直都是获得观众和评委认可的一个重要因素。幽默的前提是冲突。北京大学心理与认知科学学院毛利华副教授说：“人类大脑的主要功能是让我们预期这个世界。”当我们面对眼前发生的一件事情时，我们的大脑会思考事情可能存在的几种不同发展方向。有的方向概率大，有的方向概率小。但如果我们大脑最后接受的东西与大脑之前预期的结果不一致，就会产生冲突。此时，大脑进入紧张状态，希望用认知资源对这个刺激进行加工。一旦大脑发现自己的经验可以解释它，那么紧张的情绪就会得到释放，转而进入愉悦的情绪状态。

所以，如果机器想要获得“幽默”这种人类的独特品质，那么首先它应该学会预测。2018年，谷歌旗下的科技孵化器Jiasaw、康奈尔大学和维基媒体基金会合作开发了一个预测人类谈话走向的智能系统，以预防不必要的吵架和攻击行为。通过自然语言处理技术，AI会自动对其所接触到的内容进行语义分析，并提取在对话双方的讨论内容中出现的关键词特征，然后进一步通过机器

学习算法构建分析结果。目前已经有相关论文的数据表明，一台经过训练的计算机可以以61.6%的准确率预测一场对话是否会朝产生敌意的方向发展，而人类在这件事情的判断上准确率为72%。之前，李开复做客综艺节目时曾发表观点称：AI会在很多领域替代人类的工作，但在娱乐领域不会，因为AI不懂什么叫幽默。但是如果按照前文的介绍来看，构建一个有幽默感的机器人并非是不可能的事情。AI学会了预测对话走向，只要再设置这一项功能即可：如果AI预测的结果与现实不符，那么打开控制AI笑声的开关，一个能懂得人类笑点的机器人便会诞生。但是，目前AI的幽默属性还尚在进化当中，真正懂幽默的机器人还未成熟。

AI的逻辑能力或是人类辩论的福音

社会民主存在的前提是有一群理性、智慧的公民。而这样的公民实际上是不存在的，因为在公民投票选举的民主制度中，先不去评判公民的选择是否正确，在判断对错之前他们甚至连真假都不曾清楚地了解。美国著名新闻学者李普曼为了尽力维持社会民主的存在，在《幻影公众》一书中谈到一种公开辩论的手段。他指出人们可以通过公开辩论的方式来分辨某个人的发言是为了私利还是公共利益，这样有助于人们决定是否采纳他所提出的建议和规则。

或许正因为辩论有这样的能力，美国历届总统的选举才都采取了辩论这种形式。但随着社会信息“爆炸”，如何在越来越复杂的环境和公民有限的政治判断力中架起一座讨论的“桥梁”变得尤其重要。AI辩手Project Debater具有非常强大的数据处理能力，能够处理几十个与主题相关的数百万篇新闻。此外，Facebook也正

在通过“机器+人工”的方式削减虚假消息的传播数量。通过机器学习算法来标注可疑消息，然后将其发送给第三方事实审核人员。因此，我们能够推出在辩论中，AI也有能力提醒公民他们听到的哪些信息是真的，哪些信息有可能是假的。

另外，AI辩手Project Debater或许还可以帮助人类建立一个最公正的辩论环境。众所周知，AI的逻辑推理能力极强，这一点在“阿尔法狗”与人类的对战当中体现得淋漓尽致。所以，如果说辩论的胜负由逻辑决定，那么使用AI作为公平的裁判是一个不错的想法。名家公孙龙曾以其“白马非马”的诡辩之术让古代许多大儒无言以对，但利用数学当中的集合论却可以轻松解决这个问题。一直以来，AI都被誉为常识的“婴儿”、逻辑的“巨人”。所以，只要将人类在辩论当中用到的词语都抽象为数学符号，那么当你还在冥思苦想对方的逻辑漏洞时，AI或许早就已经帮你发现了。

每一次AI的表演都想创造里程碑式的进步，但毕竟真正的里程碑还是少数。让AI在辩论中打败人类不免有些夸张。被放大的情绪背后，我们更应该看到的是AI的真实模样。

Facebook为什么要推出AI防自杀系统？

扎克伯格决定在全球范围内推广AI防自杀系统以平息用户对Facebook的不满情绪，这是怎么回事？

原来，之前Facebook上曾多次出现直播自杀的事件，但由于平台方没有及时排查到此类信息，一时给社会造成了很大的负面影响。而Facebook也理所当然地遭到了平台用户的责难。所以，扎克伯格才决定推出AI防自杀系统。不过，Facebook已经不是第一次因为信息管理不利而被指责了。2018年《纽约时报》和《卫报》的报道指出英国一家从事数据分析的政治咨询公司利用Facebook的信息管理漏洞，窃取了5000千多万Facebook用户的个人资料，帮助现任美国总统特朗普在2016年的大选当中获得有利舆论。

Facebook近年来争议不断，这套AI防自杀系统是它挽回用户信心的一个小动作，但是小动作中也可能藏有大玄机。

从图文到视频：AI识别自杀的技术之流

在各种AI识别人类性格、性取向、精神疾病等技术层出不穷的前提下，我们可以尝试对AI防自杀系统进行简单剖析。

1. 文字识别

AI防自杀系统会对用户的发帖内容进行识别，针对含消极情绪的语句，如“我不想活了”“我恨世界上所有人”“我为什么还不

死”，AI会特别注意，并将发布此类言论的用户进行标记。然后，AI会将相应的数据移交给公司专员进行筛选和处理，最后由人工来做出决策。Facebook这套AI防自杀系统没有让AI单兵作战，而是选择人机协作的方式来筛选信息，为提高系统识别自杀的准确度提供了一定的保障。目前，AI侦测与回报机制的速度比人工的手动回报快三成，AI防自杀系统已经大大提高了Facebook信息筛选的效率，并成功阻止了多起自杀事件。

2. 图片识别

其实，要机器根据一张图片来分析一个人是否有自杀倾向是相当困难的事情，因为图片的表达内容有时候比较抽象，其中还会涉及色彩学方面的问题，就算是人类也不一定能做到。不过AI仍然可以对此做出一些尝试，深度学习的神经网络模型在各种图像识别比赛中已取得的突破性进展，目前AI鉴图一般会采用CNN、GoogleNet、ResNet三种深度网络模型结构。开发人员需从技术层面研发出一个“分类器”，从而让AI能够计算出该图属于某种类型图片的概率。这套鉴图原理从理论上来说对于所有的图片识别都通用，只不过AI要从图片中识别出自杀倾向的难点在于抓取此类图片的共同点。所以，在第一步建立“分类器”时，AI可能会遇到比较大的困难。因此，后续AI识别的准确率相对文字识别来说，也不会有太大的保证。

3. 音视频识别

用户在社交平台上发布的内容不仅是文字和图片，还有很多是音频或视频的形式。所以，AI防自杀系统在视频直播中也植入

了AI程序。虽然Facebook并没有透露有关该AI视频识别技术的相关细节，但是我们仍然可以从机器视觉领域的进展来一窥究竟。之前，一家公司的计算精神病学和神经成像研究小组团队利用机器学习预测大脑精神疾病抑郁症，通过对59名普通人的语言方式进行追踪和分析，从而预测他们潜在的患病风险，其精确度达到83%，由此可以看出机器具备一定能力从语音层面来分析人类情绪。另外，南加州大学曾推出一款AI心理治疗师，能够分析受访士兵的面部表情变化，并将AI分析表情的结果作为诊断士兵是否存在PTSD（创伤后应激障碍）的依据之一。所以，AI防自杀系统要判断视频中的人是否有自杀倾向，需要从人物的面部表情、是否存在危险器械、是否有流血等画面三个方面进行分析。总的来说，AI防自杀系统的有关图像识别技术基本上与以上提到的两个机器学习技术类似。

AI防自杀识别系统之痛

虽然现在AI识别人类情绪的工程一直在持续推进，但是情绪识别对于机器来说仍然不是一件简单的事情。

首先，我们需要弄清楚AI是否分得清“演戏”和“真实”。例如在快手上存在大量小剧场表演式的视频类型。人们把这些视频当作娱乐消遣，而机器会如何看待这种视频？嘴角上扬、泪眼婆娑、声嘶力竭等画面会被机器的“眼睛”迅速捕捉，并将其标记为极端消极负面的情绪传播。人们津津乐道的一场表演在机器眼里成了有自杀倾向的人群在宣泄情绪，这让人哭笑不得。模仿是人类的天性之一，长期熟悉模仿式表演的用户能够很快分辨出模仿和真实情绪流露之间的细微差别，但是就是这一点点细微的差别

却是机器情绪识别最难突破的瓶颈。

其次，当机器将识别出的一些负面关键词作为标记自杀人群的标准时，我们必须让机器知道那些说要“死”的人是不是真的想寻短见。Facebook这一AI防自杀系统应用于国外会有什么“乌龙”，我们还不清楚。但是设想一下，如果把这一系统挪到中国来，那么人工筛选的信息量会大大增加。在中国，人们喜欢把“死”作为程度词来使用。在日常生活中，随时随地都会听到有人说“热死了”“冷死了”“烦死了”“讨厌死了”，这类句子在人类看来只是十分常见的日常用语，但是机器对此却难以判定。因为机器对这个世界暂时还没有一个全面的感知系统。当中国被誉为“四大火炉”的城市一到夏天，每个人都能感受到高温袭来，在这个时候说“热死了”是多么正常的表达，而机器并不知道现实世界到底有多热。

另外，在机器的背后还需要思考的问题是悬在各大社交平台头顶的一把“达摩利斯克之剑”——隐私之殇。2018年6月4日，《纽约时报》报道称，在过去十年中，Facebook至少与60家设备制造商达成协议，向它们提供用户隐私信息，这些协议中的多数至今还在生效。Facebook建立了庞大的信息帝国，但是每一次信息技术的革命只是在表面上让信息看上去更加去中心化。然而，实际上只是建立了更加庞大的信息垄断市场。对比现实世界的条条框框，网络承担了社会“减压阀”的功能。负面情绪再怎么样都是人类的基本情绪，也是人类本该有的情绪。原本以为在网络世界，人能拥有展示脆弱的机会。但是，谁知道人的脆弱已经被平台收集，躲在黑暗中却被各大商家悄悄盯上。在商业世界中，商

人需要不断发掘人性的弱点，以便于制定推广商品的策略，而Facebook正在承担这样一个角色。

AI防自杀系统的落地

AI防自杀系统在正式投入使用后，一个月检测并拦截到的自杀事件有上百个，确实有效挽救了不少生命。但是从一个争议不断的社交平台的角度来说，AI防自杀系统的目的到底是什么还需要探讨。这样一套AI程序究竟是为了及时拯救生命，还是为了肃清平台环境，或只是平台针对外界质疑的挡箭牌？这还需要进一步探讨。

去除一些个例，AI防自杀系统还有很大的落地空间。因为AI防自杀系统不仅是在社交平台，还在医疗机构、审讯机构等发挥作用，及时阻止一些不必要的伤亡。目前，自杀已经成为世界第十大死亡原因，2016年美国卫生统计中心发布的数据显示，美国的自杀率在15年间上升了24%。在全球自杀的人群当中，仍然有相当一部分人属于非理性自杀，他们需要有一个AI防自杀系统帮助他们再获得一次重生的机会。所以，我们期待AI防自杀系统在解决了一些硬性問題之后，全面落地到日常生活当中，挽救更多可能在不经意间逝去的生命。

性格能够被识别后，我们究竟失去了哪些权利？

俗话说得好，“性格决定命运”。

性格为什么可以决定一个人的命运呢？在以前，人们认为思想决定行为，行为决定习惯，习惯决定性格，性格进而决定命运。而到了AI兴起的今天，因果关系简单得多。

当你的性格可以被AI一眼看穿时，你就必须意识到，你的命运即将开始改变。你的人格魅力、一些古怪的性格特点及情绪稳定度会被AI数据化。你身边的人都可以根据这一份数据报表来评判你，并决定是否要与你做朋友，你的老板也可以据此来决定要不要聘用你。

2018年，来自墨尔本大学的研究人员就设计了一种AI生物识别镜，可以根据人的脸部照片检测和显示其个性特征及外貌上的魅力，最多可以分析14项性格特征，包括性别、年龄、种族、魅力、情绪稳定度等。

AI变成了一块“魔镜”，可以照到每个人性格的“样子”。现实中自然不会有恶毒的皇后给善良的公主喂毒苹果，我们却有点害怕被“魔镜”否认“你才是世界上性格最好的人”。用AI分析性格，还有待商榷。

外貌体现性格？显然是不可靠的

识别性格的AI镜如果仅从用户的外貌做出估计，它的准确性便会十分有限。数据的一维性会对数据存量产生更多的要求，但这个系统在设计上只参考了相对较小且众包的数据集。所以，仅靠外貌就想要得到公平的性格结论显然是不可能的。

除此之外，AI分析性格的系统并非心理分析仪，而是生物识别，尤其是面部识别系统。在这方面，就难免遭遇人脸识别的通病——在年龄、性别和种族上的歧视。

2018年上半年，亚马逊的Rekognition被许多执法部门采用，但就在采用后不久，美国民权组织使用该系统扫描了所有535位美国国会议员的面部照片后，发现其中28人竟被识别成了罪犯。

无独有偶，由IBM、微软和旷视科技（Face++）设计的面部识别算法在检测肤色较深的女性时，出错率也高达35%。

这显然是一个很难解决的问题，当我们默认一个人的外貌与其性格有着强联系时，在“以貌取人”的时代，再将性格判定权利交给有偏见的AI，后患无穷，而生物识别镜凸显了算法偏差可能造成的现实后果。

AI性格分析进入市场，会剥夺我们的哪些权利？

退一万步来讲，即使通过面部识别能够分析出某个人的性格特征（毕竟外貌和性格都与遗传有关），新技术的应用并没有想象中的那么简单。

1. 性格的“魔镜”变成了父母的“遥控器”

一般来说，性格一经形成便比较稳定，但并非一成不变，而是具有可塑性的。性格不同于气质，它体现了人格的社会属性，是可以通过习惯化了的行為方式去改变的。

性格的可塑性自然让父母看到了某种希望——打造一个有完美性格的孩子。

电视剧《你的孩子不是你的孩子》中有这样一个故事——一位母亲手中有一个可以控制儿子时间的遥控器，在母亲的眼里，这个遥控器就是一种“资源”，她可以不断重复儿子的时间，不断修正儿子的错误，最后让儿子变得更加完美。

性格的校准器无疑也会成为这样一种“资源”，这样的“资源”，你想不想要？人们可以通过这个机器不断地校正，让孩子的性格变得更加讨喜。而到了最后，你的性格究竟还是不是你的性格？这是一个值得思考的问题。

2. 性格分析系统成了商业信息的靶子

众所周知，智能系统可以从海量的数据中提取用户特征，算法再精准投放大量广告，引导用户进一步消费，智能手机已然成为了我们的“圆形监狱”。

当AI性格分析开始大范围落地，眼看着有一天，智能产品能够掌握大量人们的性格数据，比人们自己还“懂”自己。在购物和广告领域，无法控制的算法可能会为某种性格的你推荐适合的商品，让你毫不犹豫地消费。



AI分析性格的兴起，使得营销人员能够以前所未有的姿态直接了解消费者。用户的所有信息指标正在时时刻刻地与保险公司或商业公司共享，这个项目透明化地呈现了这些信息对个人产生的潜在后果。

其结果便是商业信息泛滥，消费者毫无抵抗能力，不假思索地忠诚于商业品牌。

3. 性格分析剥夺了我们追求进步的权利

目前，已经有许多企业开始利用大数据、AI等技术来招聘员工，通过在职位画像和人才画像之间进行精准匹配，节省了人工筛选简历的环节，有效提高招聘效率。但是，利用AI进行面试并不能实现机器与人之间真正的沟通，面试者的综合素质究竟如何还是需要HR的判断。

然而，试想一下，你在面试的时候，HR可以通过AI性格分析系统选择一个指标进行了解，如你的责任感。如果HR发现你的责任感并不是这群面试者中最好的，你可能就失去了一个就业机会。

在招聘过程中，当HR能够大批次地使用AI来分析应聘者的性格时，那些公众所认为的“性格不好”的人就会面临更大的困难，即使他们的能力十分突出。在以后，说不定个人的业务能力也能被AI估值，到那时，用数字评判一切固然是公平的，却剥夺了每个个体追求进步的权利。

我们不能因为时代的高速前进而忽视自己与周围环境的互动能力，尤其是自己独立思考的能力。法国哲学家勒内·笛卡儿曾经提出“我思故我在”。

未火先凉，智能睡眠监测管理平台为何自己先休眠？

凌晨3：23，二胎妈妈肖女士因为身旁的小宝宝翻了个身被惊醒了。此时，窗外一片寂静，肖女士无比清醒。

凌晨4：04，在床上辗转了40多分钟后，肖女士打开了手机开始看剧，她放弃了继续睡觉的打算。

清晨6：20，随着小宝宝的一声啼哭，肖女士放下手机，准备迎接新的一天。失眠了一晚的肖女士此时内心很焦躁。

这是一个饱受失眠困扰的睡眠障碍者普通一晚的真实写照。相信大多数人都有过失眠的经历。如果你对用吃安眠药的方式来解决睡眠问题有所顾忌，那么你放心将你的睡眠交给那些挂着智能标签的睡眠监测管理平台吗？

由失眠人群催生的生意

倒床就睡的人永远不会知道失眠的痛苦，也永远不会知道有多少人在经受着失眠的困扰。

根据世界卫生组织统计，全球睡眠障碍率达27%。中国睡眠研究会2016年公布的睡眠调查结果显示，中国成年人失眠发生率高达38.2%，超过3亿中国人有睡眠障碍，且这个数据仍在逐年攀升中。中国睡眠研究会2017年的一项调查研究表明，目前我国内地成年人中失眠患病率高达57%，工作人群中65%的人存在睡

眠障碍。

每一个痛点都意味着商机，失眠的人群有多少，背后与睡眠相关的生意就有多大。有数据显示，2015年我国涉及改善睡眠产品行业细分市场为2114亿，预计到2020年，整体睡眠市场的产业规模将达到4000多亿。在这个体量庞大的市场中，AI自然不能缺席，相比床上用品、医药保健产品、食品、图书音像制品这些助眠类的产品，科技圈更多关注的是“AI+睡眠监测”。于是，像手环、手机App、被动式传感监测仪等睡眠监测产品或睡眠监测平台就诞生了。

AI监测睡眠是否科学？

“睡眠也能被监测？特别是用手机App也能监测，不是骗人的吧？”很多失眠者初次接触智能监测睡眠产品时都会有所怀疑，这些号称装备了深度学习、大数据分析等黑科技的智能监测睡眠产品到底是用怎样的方式监测睡眠？它们的科学依据是什么？

以我们来看，睡眠监测平台的监测模式可以分为三类：以智能手环为代表的可穿戴设备监测；以配置了监测传感器的床垫、枕头等为代表的非穿戴设备监测；在智能手机中的App监测。这三类监测模式的工作原理各具特色。

1. 可穿戴设备监测

可穿戴睡眠监测设备一般用体动记录仪来记录动作，以身体活动和感觉灵敏度作为衡量指标。体动记录仪将用户睡眠时的姿势等数据进行记录，通过AI计算来判断睡眠状态。深度睡眠的时

候人的肌肉会松弛，肢体不会产生较大的运动，而浅睡眠的时候，人体会产生一定的轻微运动。

目前的体动记录仪基本都具备从3个方向轴进行记录数据。根据记录数据，分析软件通过计算可以分析出能量消耗和睡眠等相关参数，包括觉醒时间、觉醒次数、睡眠效率等。

代表产品：小米手环、Apple Watch手表。

2. 非穿戴设备监测

非穿戴监测睡眠设备一般都“变身”为智能床垫、智能枕头、智能床带、智能枕头别扣等。这些设备都内置了高灵敏度的传感器，它们可以记录用户的睡眠质量、心率、呼吸和打鼾情况。用户可以通过蓝牙连接手机App，查看经过AI分析后的睡眠报告。同时，大部分非穿戴监测睡眠的智能设备还拥有播放助眠音乐、智能闹钟等功能，让用户在浅睡眠的情况下比较没有痛苦地醒来。

代表产品：Beddit、Earlysense、Withings、Sleepace等。

3. 手机App监测

根据对人类睡眠的研究，人在深度睡眠的时候，身体的重量相对比较重，因此手机加速器计得的感应值就比较大。根据这一原理，手机App监测睡眠是利用手机内部的加速度传感器和陀螺仪来监测用户在睡眠中的活动，从而用AI推算出深浅睡眠状态，分析睡眠质量，深度睡眠时间等。

代表产品：蜗牛睡眠App、萤火虫睡眠App、睡眠大师App等。

智能睡眠监测“未火已休眠”

上述三种睡眠监测模式好像都有各自的“独门秘籍”，都有各自的专长和优势，但在睡眠监测中，这三种模式没有绝对的强者，相反，都是在“火”过一阵之后渐趋平静。智能睡眠监测管理平台在没有解决人类失眠问题之前为何自己就先“休眠”了呢？其原因有四点。

首先，睡眠监测结果不准。

“明明整夜都在做梦，大脑一片混乱，可手机App给出的睡眠报告还是显示昨晚的睡眠质量不错。”不少失眠者在使用过一些智能监控睡眠平台的产品后都会有这样的经历。“测不准”是各类睡眠监测平台的“通病”，主要因为睡眠模式的多样（深睡、潜睡、做梦）与判断方式和标准未形成一套通用的AI算法。

根据Sleep Shepherd的创始人兼总裁Michael Larson的说法，AI算法的问题在于“当下大多数AI算法的核心是模式匹配。而在睡眠技术领域所存在的一个问题就是，算法中使用的模式是有缺陷的，它们无法很好地描述睡眠状态。”当它的算法基础与运动传感器发生偏离的时候，AI就会存在缺陷，可能会给出不准确的数据。

其次，可穿戴设备本身就不利于睡眠。

对于用可穿戴设备来监测睡眠的模式来说，其本身就是一个悖论。

原想利用这些设备监测睡眠，可现实给它的回应——没人会戴着手表或手环睡觉。对于那些裸睡嗜好者或有强迫症的人，这些可穿戴设备很可能成为他们失眠的原因。

此外，晚上戴手环或手表睡觉，手环和手表会卡紧手腕，如果长期压迫手腕动脉，那么容易引起手臂麻木，进而影响睡眠。受睡姿的影响，可穿戴设备采集的数据也存在很大的误差。这个模式注定只是供那些好奇的用户尝鲜使用，不是最优方案。

再次，用户担心辐射影响健康。

手机辐射对人健康到底是否有伤害？关于这个问题，很多媒体从不同角度、不同方面做过多次科普，但仍然有人对此保持怀疑。在问卷星上有一份《关于手机和手机辐射了解的调查》，在这份调查报告中，有79.41%的受访者认为使用时间越长，手机辐射越大；仅有22.06%的受访者晚上会将手机放在枕边充电，55.88%的受访者表示会将手机放在卧室但不放在床边，甚至有14.71%的受访者晚上会将手机放在书房。在国民教育程度偏高的德国，一项德国媒体对“手机辐射是否会对人体造成伤害”的调查中，也有55%的德国人相信手机辐射存在风险。

很多智能监控睡眠设备都采用雷达、射频技术进行数据采集，智能监测睡眠管理平台要做的是打消用户对辐射的顾虑，平台不能只对用户进行教育，毕竟对于一个睡觉都要关机的老人来

说，说服他将手机放在枕边监测睡眠并不是一件容易的事情。

最后，只监测不干预要它何用？

“睡得好不好，我自己难道不知道吗？”智能监测睡眠管理平台以用户睡眠痛点切入，但未能击中痛点。

失眠者收到的睡眠监测报告，仅仅只是一份报告。平台通过数据分析，用户的睡眠质量很差，用户自己也知道经常失眠。然后，就没有下文了。报告对于用户如何改善和解决自己的睡眠问题毫无帮助，睡眠监测平台逐渐被用户弃用，慢慢进入“休眠”状态。

AI监测睡眠如何完成跨越？

很显然，只做睡眠监测，商业变现途径有限，那么AI监测睡眠平台实现跨越的方式有哪些呢？

1. 生命监测优先，睡眠监测靠后

就智能睡眠监测平台的目前发展来看，纯粹的睡眠监测不痛不痒，无法在根本上解决和干预用户的睡眠问题。

我们认为与其提升采集数据的精确度、监测报告的准确性，睡眠监测平台首先要解决及时唤醒睡眠时打呼的用户或因喝酒产生的呕吐物可能引发窒息的问题，让肥胖患者不会在睡眠中猝死。智能睡眠监测管理平台优先发展的方向应该是保护生命。

只有在死亡面前，人类才会真正的害怕，这才是最痛的痛

点。

2. 不光要监测还要干预

只以睡眠监测作为切入点无法直接给用户带来良好的睡眠体验，如果能够形成“睡眠监测+数据分析+医生干预+医生随访”的闭环，平台根据失眠患者的具体状况提供精准定制化的失眠解决方案，那才是失眠人群最需要的智能监测睡眠平台的运行模式。虽然目前市场上已经有类似产品出现，但还远远没有形成有标识性的商业模式。

在改善睡眠方面，除了平台外的人工干预，监测设备在用户睡眠中的及时介入也非常重要。比如环境噪声、室内温度会影响用户睡眠，监测设备对窗户、空调等进行控制来优化睡眠质量。如果想象力再丰富一些，帮助用户驱赶蚊子、调整随时可能引发坠床事故的睡姿等，都是智能睡眠监测平台未来可以做的事情。

3. 夜间物联网的入口

其实对于智能监测睡眠平台来说，它更为远大的意义在于如何打开夜间物联网的入口。

类似于智能音箱，互联网的巨头们将其看作智能家居的入口和家庭AI交互的切入点，如今的竞争异常激烈。带给智能监测睡眠平台的思考是——如果白天的场景被智能音箱承包，那晚上由谁来接管？睡眠监测技术和当前火爆的智能语音技术有着极大的相似性：技术本身无法直接带来价值，但以技术为轴衍生出来的场景应用却是前景光明。

最后，用乔布斯著名的广告语送给尚处于发展之中的智能监测睡眠平台：你可以赞扬他，你可以侮辱他，你说他什么都行，但有一点你不能做，就是你不能忽视他。

为什么要把孩子的健康交给机器人？

一直以来，儿童健康都是社会上很受关注的问题。自“毒奶粉”事件被爆出，越来越多的家长将儿童健康提上了家庭日程。但在父母们的实际操作中，孩子的健康变得格外简单——只需注意饮食和生活环境的质量，身体出现异常及时送往医院即可。

众多家长对儿童健康还是知之甚少，网上形形色色的帖子和书籍不少，可就是解答不了赤膊上阵的家长们的疑惑，例如成人的那一套健康方式究竟适不适合孩子？不同时期，孩子的保健方式是否应该有所侧重？孩子偶尔出现“不正常”的情况，究竟是疾病还是特例？

这时，许多人就开始打AI的主意。

用**AI**服务儿童健康，父母们需要吗？

近些年来，儿童教育、陪伴机器人频频出现，很多父母都愿意甚至期待将孩子的教育和娱乐交给机器人。但父母放心的前提是机器人即使在这两个方面出了问题，也是可以被弥补的——教育可以修正，娱乐互动也不是刚需。儿童的健康至关重要，父母们真的需要儿童的健康服务机器人吗？

国内有科学家做了一个有关于儿童保健护理系统的试验，他们选取早期发育儿童77例，分为两组。试验组予以儿童保健护理系统管理，对照组予以儿童常规护理模式。最后的结果是，试验组儿童身高、体重等各临床指标较对照组明显高，这表明了儿童

保健护理系统管理对儿童早期生长发育具有显著的促进作用。

但与之相对的却是我国儿童保健护理教育的缺失。《生命时报》在北京、广州、山东、安徽等地调查发现，少年儿童的肥胖和视力不良检出率正在持续走高，厌学、网络成瘾等心理问题也日渐凸显。

于是，在儿童的健康护理日益成为所有家长心中的隐患之后，能够直接提供健康服务功能，并且伴有陪伴和娱乐功能的智能终端无疑成为了家长们的刚需。在这个特定的时代背景下，健康服务机器人已经是儿童生活场景下的必然选择。

所以，儿童健康服务机器人其实是一个亟待发掘的潜力市场——AI可以从智能影像诊断、个性化保健方案、智能健康管理等多个维度深入打造儿童保健领域的AI产品。

首先，父母能够利用AI随时监测儿童的健康情况。这类儿童健康服务机器人的造型可以是孩子们喜欢的卡通人物，它同时载有多种环境传感器，能够识别语音和理解语义，连接第三方的健康监测设备和医疗平台，还可以依赖深度学习针对某项疾病做出诊断，如自闭症的诊断。

北卡罗来纳大学（UNC）教堂山分校精神病学家Heather Hazlett就曾开发出深度学习算法，用来预测2岁前的自闭症高危儿童是否会在2岁之后被诊断为自闭症，并以88%的准确度远超准确度只有50%的传统行为问卷调查法。

其次，AI还能够为孩子设计个性化的保健方案。不同年龄的

孩子的保健侧重点会有所不同，不同个体的健康指标也会有区别，所以，AI要能够建立更加丰富的数据库，了解不同时期的孩子特征。

例如学龄前儿童（6~7岁）的体格发育速度慢，智能发育日趋完善，他们喜模仿而又无经验，故意外事故较多。儿童因接触面广，易患急性肾炎、风湿病等，AI在为这类儿童服务时，就应该侧重于这类疾病的预防和监测。

把儿童健康交给机器人，还没那么容易

儿童健康服务机器人对父母和孩子有很多好处。但儿童保健机器人要落地还是有难度的。

1. 智能终端不断商业化，AI是福音还是噩梦

“AI+儿童保健”应该以怎样的形态出现才会被更多家长和孩子接受？从目前来看，作为一个长期的监测设备，可穿戴设备是一个极佳的选择，比如智能手表。有数据表明，儿童手表在2021年将占到智能手表整体出货量的30%。

但这个市场还潜藏着许多待解决的问题，主要包含两个方面，一是产品本身材质的安全，二是产品商业化带来的危险。材质安全问题主要出现在泛滥的山寨产品中，这里不再展开。下面主要介绍第二个方面。

当我们为儿童机器人添加了健康服务功能后，智能终端会从海量的数据中提取儿童的健康特征，不排除会有商家使用成熟的

信息推送技术在终端内精准投放大量的儿童保健广告，引导用户（父母）进一步消费。

面向儿童的智能健康服务产品的兴起，使得一个家庭接收的商业信息会更加泛滥。如果家长没有介入产品使用的过程，企业会直接将用户对象定位为儿童，那么营销广告将直接接触到儿童。如果这个广告是某游戏的推销，那么会加剧成瘾性游戏对儿童的侵蚀。相对于成年人而言，儿童在商业化的营销中简直毫无抵抗力。

数据除了被商家拿来营销，还有可能被不法之徒利用。以智能手表为例，《焦点访谈》就曾经曝光了儿童智能手表存在的隐患，如地图、通话被黑客侵入，导致信息外泄和定位不准等安全问题。

因此，制定行业标准和技术规范，抵制劣质产品，让儿童和家长免受有害商业营销行为的侵扰，给家长以宽心是“AI+儿童保健”亟须解决的一个问题。

2. 场景千万个，AI应该去哪里发光发热？

一般来说，提起儿童机器人，落点都会在家庭场景。有陪伴功能的机器人固然如此，但儿童保健却有着更多场景要求。

在探究中国3~6岁儿童的家庭照料和正式照料（幼儿园/托儿所）与其超重肥胖的关系试验中，我们发现，3~6岁儿童总超重肥胖率为20.99%，其中家庭照料超重肥胖率为22.08%，明显高于正式照料15.82%的肥胖率。

分析结果显示，正式照料可显著降低儿童超重肥胖的概率，原因在于，与家庭照料相比，正式照料大幅度提高了儿童参加体育活动的概率。

也就是说，其他场景如幼儿园、学校等，也非常需要儿童健康服务机器人的介入，以达到更强的监督和保护作用。而在产品的应用情景中，运动情景也是一个最大的需求项目。

当然，这不是一个简单的智能产品就能解决的问题，让除家庭以外的场景，尤其是学校重视儿童保健问题，其实更需要改变整个正式照料场景的体系。通过AI去观察、监测儿童的运动轨迹，进而让医疗方获得最有价值的健康信息，这在将来会是一个趋势。

不仅仅是把儿童健康交给机器人

上文中提到“AI+儿童保健”的应用场景不仅仅是家庭场景，还应该有更多。而这个“更多”其实应该包含更广范围的社会生活，利用AI为人们挖掘一些社会问题，为处于绝境中的儿童发声。

据Childhelp网站统计，美国每年都有超过660万起关于儿童被虐待的报告，其中28.3%的儿童受到身体上的虐待，20.7%的儿童受到性虐待，10.6%的儿童受到精神虐待。

被虐待的孩子长期处于阴影之下，即使成年后逃离了施虐者，依旧遭受着精神上的折磨。对于这些事件，如果能够及早预防、及早发现，那么就能够避免儿童被虐待的悲剧一再发生。

美国研究者在20世纪80年代曾提出，在3～6岁的儿童群体中，遭受虐待的孩子，往往拥有比同龄人看起来更成熟的面孔，“少年老成”的面孔背后，往往意味着生活的艰辛。

如果可以利用AI的人脸识别，非接触式地获得孩子的面部图像并进行对比分析，将系统与社会上的儿童救助机构连接，那么AI的儿童健康监测也能够监测到更多大家关心的内容。

解决食品安全问题，“人造食品”或是方向之一

2018年下半年，远在美国的喜欢吃肉的人吃得一点也不安心，8月至11月在美国集体暴发了沙门氏菌感染症，截至12月初已经有250人左右患病，很多患者都是因为食用了鲜牛肉加工商JBS USA的食品。

10月，JBS已经对旗下可能感染沙门氏菌的食品召回了690万磅，到12月4日JBS又召回了510万磅碎牛肉，总量合计1200多万磅（约5488吨）。

不光是美国，2018年德国的也发生过多起大规模沙门氏菌感染鸡蛋的事件，12月5日又一德国鸡蛋生产企业在多个州召回售出的受沙门氏菌污染的鸡蛋。

人造食品是解决食品安全问题的关键吗？

提起“人造食品”四个字，大家第一反应想到的会是什么？“人造鸡蛋”“人造海带”“人造辣条”反正就是不能吃的“假”食品，它们存在食品安全问题。现在用“人造食品”来解决食品安全问题是否可行？

目前人们对“人造食品”已经有了新的界定，之前想到的那些“人造食品”都是属于“假·人造食品”，现在已经有了很多可以食用的“真·人造食品”，甚至还包括了各种肉类和鸡蛋。

对“人造食品”的追溯可以到很早以前，像我们常吃的果冻、

蟹肉棒其实都可以算是“人造食品”。以“人造肉”为例，在最开始使用植物蛋白模拟肉，从大豆中提取蛋白质，经过“纺织化”处理，将原来的球状蛋白变成肉的纤维状，使食物口感接近肉，也被称为“植物肉”。

而在2013年，荷兰马斯特里赫特大学教授马克·波斯特研制出了一种“试管牛肉”。人们首先从牛的肌肉组织中分离干细胞，并放入营养液中，3周后细胞数目超过100万个。接着再把它们放入数个小型容器中，细胞合成大约1厘米长、几毫米厚的“肉丝”。最后大约3000条薄薄的“肉丝”冷冻起来就能组成一块正常大小的“肉饼”，而这种“人造肉”则是货真价实的“肉”了。

第二种“人造肉”需要在无菌环境下进行培养，从根本上杜绝了疯牛病、口蹄疫、沙门氏菌之类的食物传播疾病。

像在德国暴发的鸡蛋受感染事件，其实用“人造鸡蛋”也能解决这一问题，2013年美国汉普顿·克里克公司研发的一种可以替代鸡蛋的人造食材在美国各大超市上架出售。这种被称为“超越”了鸡蛋的“人造鸡蛋”一经出现，便吸引了不少人的目光。

当然它并不是和“人造肉”一样是实实在在的肉，不像鸡蛋一样，由一颗颗蛋壳包裹，而是以液体方式被装在瓶子中销售。

这种由植物做成的“人造鸡蛋”，在口感上与传统鸡蛋基本一模一样，但因为是由植物制成的，所以它也免去了受到类似禽流感、沙门氏菌等疾病、细菌感染的风险。

在《科学美国人》评选出2018十大新兴技术中，“人造肉”这

项技术入榜，如果未来真的不管是肉类食品还是其他食品，例如一些保存周期较短容易变质生菌的食品，都能在保障安全的前提下以“人造”的方式产生，食品安全能够得到比较大的改善。

人造食品为何会受到大资本的青睐

“人造食品”在几年前就成为了资本圈里的“香饽饽”，全球很多资本大佬，对于在“人造食品”特别是“人造肉”领域的投资十分看好。

在2013年，当比尔·盖茨尝过了Beyond Meat公司推出的一款无肉鸡肉卷之后，就决定对其进行投资。Beyond Meat是一家生产“植物肉”的公司，从2016年开始，人造肉公司Beyond Meat就已经在美国的全食超市和西海岸30家餐厅及耶鲁的饭堂出售产品了。根据其官网数据，他们的人造肉汉堡已经售出了数千万份。其他的投资人还有著名影星莱昂纳多·迪卡普里奥，Twitch联合创始人比兹·斯通等。

还有像我们熟悉的李嘉诚，也对“人造肉”表现出了极大的兴趣。2014年，李嘉诚通过其投资公司Horizons Ventures向位于纽约布鲁克林的人造肉初创公司Modern Meadow注资1000万美元，用于人造肉的加速研发及规模的扩大。

除了植物肉，资本也很青睐“真·人造肉”领域。Memphis Meats公司在2017年得到了来自比尔·盖茨、Richard Branson、Twitch联合创始人Kyle Vogt及伊隆·马斯克的兄弟Kimbal Musk等一群富豪或公司的关注，总共获得融资2200万美元。

是什么样的魅力使得“人造肉”如此受资本欢迎呢？

首先，相比传统的畜牧业，“人造肉”更加环保。从1960年开始，牲畜数量已翻了几番，目前已超过了600亿头；全球1/4的土地现在被用于放牧，1/3的可耕地被用于生产饲料；1/5的亚马逊雨林已被毁，主要用于牲畜和饲料生产。

世界上20%左右的人造温室气体排放量是由畜牧业产生的，已经远远超过了全球交通运输（包括汽车、火车、飞机、轮船）的排放。

欧盟甚至在《气候变化——2050：今天决定未来》的决议中，警告畜牧业的温室气体排放，并敦促减少肉食的消费。

其次，“人造肉”有助于解决全球食品短缺问题。上文有提到的美国汉普顿·克里克公司研发了“人造鸡蛋”，初衷是因为其创始人乔希·泰特里克在非洲参加一个减贫项目时，目睹了那里食物短缺的严重情况，萌生了要制作“人造鸡蛋”的想法。

世界自然基金会（WWF）的一份报告显示，如果全球人口增长状况一直保持目前的速度，而食品生产力和饮食偏好不发生改变，那么到2050年人类将遭遇严重的食品短缺危机。

最后，也是最直接的问题，“人造肉”的出现解决了对动物屠宰的需求。在中国，每年生猪屠宰量大约在6亿头以上，还有牛、羊、鸡等，每年屠宰的牲畜以十亿计数。

不只是屠宰，在养殖方面，对于牲畜而言也显得十分“残

忍”，自1960年以来美国、英国等工业化国家开始工厂化养殖动物。美国有超过95%的肉类来自工厂化农场，而那里的环境并不适合动物生长。

一般散养家鸡的正常寿命为3~5年，但是养殖场的肉鸡只能活35~49天。在如此短暂的生命周期里，它们必须长到大约2公斤，这是任何有机体都难以承受的残酷的快速生长。

也正是这些原因，使得畜牧业长期受到素食主义者们的抗议和威胁。“人造肉”若是能解决这些问题，受到资本青睐也就理所当然了。

人造食品在中国商业化跨不过的一道坎：消费者的兴趣

我们知道像“植物肉”“人造鸡蛋”等商品已经在美国市场实现了商业化，但像是使用动物肌肉组织培育的“真·人造肉”还并没有进入上消费者的餐桌。

价格高昂是阻碍之一。上文提到，在2013年，荷兰马斯特里赫特大学教授马克·波斯特研制出了一种“试管牛肉”，但成本高达200多万元人民币。也有专家表示，这类人造肉现阶段价格高是由于技术不成熟，没有形成生产规模，可能再有10年左右的时间，人造肉将会和普通肉价格持平，甚至价格比普通肉更低。

第二个原因是人造食品的营养价值达不到被代替物的程度，比如“人造鸡蛋”虽然在口感、价格上都跟真鸡蛋基本一样，但营养成分没有真鸡蛋那么高。

但这些在我们看来还都只是小问题，关键问题是消费者能否埋单。

中国的消费市场与美国又有很大的差别，在“真·人造肉”面世之后，根据以往的经验来看，欧美民众的接受程度会远大于国内，特别是在一些明星、企业家的带动之后。

虽然“人造肉”在技术上与“转基因”食品有本质区别，但我觉得目前国内市场在看待“人造肉”和“转基因”食品的眼光是一样的。

“转基因”食品的优劣我们现阶段无法评估，在美国转基因食品未被强制要求贴标，民众对其的接受程度也远高于国内。但在国内对于“转基因”的反对的声音基本占到了8成，即便一些专家学者、甚至是100多位诺贝尔奖得主也持反对态度。他们之中大部分人都是自然科学奖获得者，还有很多物理及医疗奖获得者出来支持“转基因”食品但依然未能取得效果。

有“中国诺奖”之称的2018年度未来科学大奖，三位著名的作物育种专家李家洋教授、袁隆平教授和张启发教授荣获“生命科学奖”，以表彰他们在系统性地研究水稻特定性状的分子机制和采用新技术选育高产优质水稻新品种中的开创性贡献。

由此可见，对自己不熟悉的陌生的东西感到恐惧可以理解，但是过度恐惧也没有必要。所以，以“转基因”食品为例可以预见“人造肉”未来在国内市场也绝不会一帆风顺。

小心，教育机器人别好心做坏事

近日，国内某品牌的教育机器人广告，成功引起了我的注意。在广告中，小孩子问了他爸爸几个问题，但他爸爸的回答都是要他去问机器人。

笔者带着好奇的心情，去翻了翻其他品牌的广告，80%以上的品牌在内容表达上基本一模一样，可这类教育机器人广告真的合适吗？

别让机器人磨灭孩子本性

之前在某平台，笔者正好看到过一个关于儿童教育的案例，在小孩咿咿呀呀还不会说话时，他喜欢用手去指，如果此时家长耐心地告诉他这是什么，是用来干什么的，那么小孩会指向更多的东西。因为家长激发了孩子的好奇心，他渴望了解不认识的东西。而如果家长在孩子指向一个东西时，选择不回答或回答很简短，慢慢地孩子就会失去他的好奇心，之后不管遇到什么不认识的东西他都不会去问。

同样的道理，在使用教育机器人方面，如果都像广告中一样，家长都让孩子去问机器人，那么会产生什么结果？

首先，孩子的问题可能会“难倒”机器人，教育机器人的“智商”目前被严重高估，它不是全能，其实说白了教育机器人就像一个功能更多、问题储备量更丰富的智能音箱。当面对孩子各种天真的问题时，人有时候都无法给出一个准确的答案，更何况是

一个机器。

其次，当家长都像广告一样教育自己的小孩，长期不回答他们的问题时，小孩们也将不会再发问。在儿童阶段父母无疑是小孩的主要的沟通渠道，当他们失去与父母的对话的机会后，机器人会成为他们仅剩的沟通对象，但鉴于儿童表达不会很清楚，机器人可能都不知道他们在说什么，而无法与小孩进行交互，小孩们的好奇天性也很可能会随之磨灭。

最后，教育机器人给出的答案也不一定就是正确的。教育机器人市场目前处于“野蛮生长”状态，产品价格、质量差距巨大，很多质量差的产品连一加一等于几都会答错，而小孩对于一个问题的认知一旦形成便很难更改，想要借助教育机器人来教育小孩还早了点。

对于机器人对小孩的影响，来自德国和英国的一个研究团队专门展开过一次测试，他们模仿波兰心理学家Asch于1951年的一个实验，选择了一批成年人及一批小孩为测试对象。测试的内容是，测试者给出一个答案相对明确的问题后，机器人会故意回答错误选项，再看成年人和小孩如何作答。测试结果显示成年人基本不会受到影响，而小孩的测试结果则相反，他们非常相信机器人的选择，很多都会选择机器人给出的错误答案，哪怕自己原本的选择是正确的。通过数据分析发现，在没有机器人干扰下，小孩回答的正确率可以达到87%，而在机器人介入后，正确率下降到了74%。

所以，对于小朋友的教育绝不能像各种教育机器人广告中的

那样，特别是当小孩和机器人独处时，如果他们获取答案的来源只是机器人，那么哪怕是错误答案，他们也只会跟着选择。

当孩子更愿意亲近一个机器人时，该如何是好

近日，笔者与身边的朋友讨论了下关于教育机器人的问题，有一个很有意思的问题被提出，就是“当陪伴小孩最多的是机器人，而机器人成为小孩的精神支柱时，该怎么办？”

这个问题，虽然不是由机器人引发的，但确实可能成为现实。在生活中，很多小孩都非常看重自己的一个玩具、一件衣服或根本不属于他的一件东西，更何况是一个能与其“交流”的机器人。当然不能一概认定教育机器人会造成这种现象，但是凡事都有那个“万一”。

为了防止这个“万一”的发生，在讨论时就有人说他绝对不会给自己儿女购买类似的机器人产品。

使用机器人来接近并治愈患有自闭症的儿童似乎是一个不错的选择。在现实生活中，人们很多时候都是将自己“锁”在了一间间钢筋水泥之中，在自己不愿去与邻里交往的同时，也限制了孩子们的玩耍空间，加上很多小孩从小就是由爷爷奶奶们陪伴成长的，使得他们越来越孤僻。

市面上也有很多专门针对自闭症儿童研发的智能机器人，例如，卢森堡的一个公司研发了一款名叫QTrobot专门针对自闭儿童的机器人，他们通过测试研究发现，当自闭儿童与机器人在一起时会明显增加儿童的注意力和参与度，减少了自闭儿童的焦虑

和破坏行为。

这看似能很好地解决儿童的自闭问题，可或许会引起另外的现象。关于这个问题，笔者专门询问了暨南大学的心理学老师刘老师，从心理学上来说，使用外物介入的方式治疗自闭儿童，一般情况下就是让自闭儿童形成精神寄托，但若是由于大脑或生理激素造成的自闭，机器人的介入也完全不会有效果。

而当机器人成为孩子的精神寄托之后，还会带来多方面的负面影响。上文有提到过，孩子会对机器人盲目信任，会影响孩子自己的判断思考能力，从而影响孩子的智力发育，并且这种情感寄托的形成期是在儿童时期，随着他们的长大，这种不健康的精神寄托将会导致他们与外部世界格格不入。

并且刘老师指出，还有一个现象也需要注意，在人类心理师治疗自闭儿童的过程中，许多儿童在面对自己的治疗师时是放松且愿意交流的状态，而一旦面对其他人，又回到自闭状态。

如果这个现象同样在机器人身上发生，后果还将严重得多，因为机器人没有情绪、无法判断，即便自闭儿童出现一些错误的言语或行为，机器人根本没有办法即时纠正。

面对质疑，教育机器人该如何自处

随着二胎政策的实行，其背后是2000亿的消费市场。而儿童消费在中国家庭消费中的比重也越来越高。目前，有数据显示国内教育机器人市场规模已接近10亿元，中国市场也被认为是全球最有潜力的市场之一。

但面临的问题也十分严重，目前市场上的教育机器人除了外形不同，基本功能都很雷同，产品价格、质量差异也非常之大，即便消费者想买也根本不知道如何选择，而这些都还只是次要问题。

如何打消父母认为教育机器人会给孩子带来各种的负面影响是应该考虑的主要问题。像目前市面上的一些高端交互类机器人，例如亚马逊的智能音响、微软的聊天机器人，都出现过一些莫名其妙的声音，在成人客户面对这些问题时都会不知所措，而教育类机器人的用户群体都是孩子，当类似问题在他们身上发生时极有可能会使他们产生心理阴影。

所以教育机器人需要找准自己的定位，就目前而言，教育机器人完全没有能力独立照顾孩子的生活学习，教育机器人的定位应该是作为家长辅助子女进行生活和学习的一个工具。

就像计算机或手机一样，单独让孩子自己使用这些产品，他们无法分辨哪些信息是对的哪些信息是错的，需要家长帮其过滤，教育机器人同样如此。

而教育机器人厂商需要做的第一件事就是更改自己的广告或标语，将那些极具诱导性的像“孩子自己就能学习”“孩子的学习再也不用我操心了”之类的通通撤销，因为这完全不符合现实情况。

并且无论是厂商还是家长，都必须清楚认识到，让孩子与机器人交朋友本身没有错，但如果像一些广告中描述的那样“再也

不用担心孩子没朋友了”之类的，那就是大错特错。

人与机器人的交流永远无法替代人与人之间的交流，毕竟目前我们生活的社会，无论生活还是工作都是在人与人之间进行的，而不是在人与机器人之间进行的。