

各省公务员录用考试

《申论》真题卷（网友回忆版）

（微信公众号：逢考必胜网）

重要提示：

内容来源于网络 仅供参考

word版定制联系微信：szgkz1z3



为了能给您提供更好的服务
请添加客服进行咨询，谢谢~

准考证号

姓名

材料一

AIforScience是指以机器学习、深度学习等人工智能技术分析处理多维度、多模态、多场景下的模拟和真实数据，解决复杂推演计算问题，加快基础科学和应用科学的发现、验证、应用，打造下一代科学范式。AIforScience不只是一个新的科学浪潮，它更是开启了一个全新的科学时代。

在过去50年里，蛋白质折叠问题一直是生物学领域的重大挑战。蛋白质是由氨基酸链组成的、具有自身独特3D结构的大型复杂分子，其对人们理解生命形成的机制至关重要。蛋白质从氨基酸序列折叠成何种形状与其功能密切相关，而预测蛋白质结构对于理解其功能和工作原理至关重要，这就是“蛋白质结构预测问题”，被称为生物医学领域的“圣杯”。谷歌DeepMind公司创始人和CEO戴密斯·哈萨比斯和资深成员约翰·贾伯与DeepMind的团队一起开发了人工智能系统AlphaFold，可以在几分钟内预测由人类基因组编码和20个模式生物的几乎所有已知蛋白，并精确到原子级。这是第一种在已知没有相似结构的情况下构建高分辨率预测的方法。AlphaFold这种变革性方法，解决了从一维氨基酸序列预测三维蛋白质结构的长期难题，给出了50年来关于蛋白结构最准确、最完整的图像，破解了长期困扰生物医学研究领域的困境，增进了人类对基本生物过程的理解并促进了药物设计，为加速生物和医学研究打开了大门。在两年一度的蛋白质预测大赛CASP（蛋白质结构预测关键评估）中AlphaFold以绝对优势夺冠。AlphaFold2开源仅一周的时间里，98.5%的人类蛋白质结构被AlphaFold2所预测，而在此之前，全球多少顶尖科学家耗时数十年的努力，也只解码了覆盖人类蛋白质序列中17%的氨基酸残基。

材料二

麦肯锡的一份调研报告提出，适应性是技术改善学习的关键方式，通常也称之为个性化，个性化意味着给予学习者“发言权”和“选择权”。传统教育技术通过查缺补漏的方式为学生提供指导。如果以学生的长处为导向，发现学生特长和能力，即“资产”，则可以正向辅助学生学习。人工智能以“资产”为导向，更能全面掌握学生基本情况，从而促进教育公平。学习不仅包括认知，也包含社会学习能力、自我调节能力、推理、解释和证明能力等其他关键技能。人工智能可以帮助学生实现多种技能的增强学习。人工智能模型能够处理多种学习路径和多种交互方式，从而可以为一些特殊学生，比如神经多样性学习者，提供不同的学习路径，采用适合其特征的学习方式，使学生从中受益。人工智能模型可以帮助学生完成开放的创造型任务，促进学生创造和创新能力培养。正确答案并非是唯一的学习目标，人工智能模型可以引导学生熟练进行团队协作和领导团队，尝试完成多样化的目标。除此之外，人工智能可以通过减少教师的行政或文书工作中的低级负担来改善教学工作，并建议将所节省的时间用于更有效指导学生。比如教师有时并不能随时随地在学生身边提供帮助，但是可以借助人工智能生成评价结果，了解学生情况并及时反馈，由此扩展了教师的认知范围。最新研发技术已经能够记录教师课堂内容并为教师推荐教学方案，课堂模拟

工具也不断多样化，可以帮助教师在真实情境中练习技能。

材料三

生成式人工智能是一种基于机器学习和人工智能技术的范畴，其目标是让计算机系统能够自主地生成各种类型的数据，如文本、图像、音频等。生成式人工智能的核心能力在于创造新的内容，而不仅仅是对已知模式的重复应用。生成式人工智能应用广泛，包括的领域有：自动写作与内容生成：可以自动生成文章、新闻、故事等文本内容，帮助内容创作者提高效率；艺术创作与设计：在绘画、音乐、设计领域，生成式AI可以创造出独特的艺术品和创意作品；虚拟现实与游戏开发：可以生成逼真的虚拟世界、地图、角色等，用于游戏设计和虚拟现实体验；科学研究与创新：在药物研发、分子设计、科学模拟等领域，生成式AI可以加速创新和发现；语音合成与音乐创作：可以创作音乐、合成语音，并模仿各种声音风格；教育与培训：可以为教育领域提供个性化的学习材料和辅助教学；医疗诊断与图像生成：在医学图像分析、病理判读等领域，生成式AI可以辅助医疗专业人员进行诊断。

以ChatGPT为例，它拥有接近人类水平的语言理解和生成能力，因其出色的回答问题、创作内容、编写代码等能力，使得人们直观真切地体会到人工智能技术进步带来的巨大变革和效率提升，上线5天用户突破100万，两个月活跃用户突破1亿。相比视频、图像、语音等，自然语言的语法、语义、逻辑复杂，存在多样性、多义性、歧义性等特点。文本数据稀缺，通常表现为非结构化的低质量数据。自然语言处理任务种类繁多，包括语言翻译、问答系统、文本生成、情感分析等。因此，长期以来自然语言处理被认为是人工智能最具挑战性的领域。ChatGPT不仅实现了高质量的自然语言理解和生成，并且能够进行零样本学习和多语言处理。在此之前，人工智能在不同场景应用需要训练不同模型。而ChatGPT利用单一大模型即可完成人机对话、机器翻译、编码测试等多种任务，已经具备通用人工智能的一些核心技术和特征：能够自动化地学习各种知识、信息，不断自我优化；充分理解和流畅表达人类语言，逻辑推理强，实现了具备一般人类智慧的机器智能；拥有一定的自适应和迁移学习能力，可以适用于多种应用场景和任务。ChatGPT证明了大模型的学习和进化能力，将推动强人工智能（机器拥有知觉和意识，有真正的推理和解决问题的能力）加速演进。目前大模型智能程度已接近人类水平，甚至一些业界人士认为，将来会逐渐产生自我认知和感知，进而出现意识并且超越人类。

材料四

“现在我们研发的语音输入法和中英文翻译机已在快速普及，语音识别技术可以把人说的话实时转写成文字，准确率超过98%。”科大讯飞董事长说，“只要是简单重复性的技能，以及通过学习、训练能够得到且不需要天分的技能，人工智能完全可以替代人类。”他举例说，原来依靠人工去识别、审核图片和视频，工作量很大、速度很慢，现在则完全可以交给机器。它能快速过滤掉大部分无用信息，节省人力和时

间。以国外的应用为例，全美失踪与被迫害儿童中心的数据显示，2016年共接到820万个与虐待图片、在线诱惑、贩卖和骚扰相关的报告。基于英特尔的AI技术，该中心可快速扫描包含可疑内容的网站，进行快速查询并分享数据，过去30天的处理报告周转时间可缩短到1至2天。对于失踪与被迫害的儿童来说，这节省下来的20多天简直就是一辈子。人的感官只有耳朵、眼睛、鼻子、嘴巴等，人工智能可以拓展、延长人类的感知能力和行动能力。少数人天生就有听觉障碍、视觉障碍等，人工智能可以弥补这些缺陷。

最近，中国电信对外推出了升级的千亿参数大模型星辰语义。中国电信相关负责人介绍，星辰千亿参数的语义大模型重点解决百亿参数的语义模型在商业化落地过程中面临的幻觉、多轮逻辑推理等问题，主要聚焦提升图文生成、图文理解能力，中文意象理解生成能力提升30%。这不是个例，近期百度、阿里巴巴、科大讯飞等国内相关企业纷纷推出AI大模型的升级版本。例如，阿里云发布千亿级参数大模型通义千问2.0，在复杂指令理解、文学创作、通用数学、知识记忆等能力上均有显著提升。科大讯飞升级了讯飞星火认知大模型V3.0，并启动更大参数规模的星火大模型训练。

2023年世界互联网大会乌镇峰会上，多家企业纷纷展示其大模型的应用新成果。在华为展区，盘古气象大模型吸引了人们的眼光，这是全球首款精度超过传统数字预报方法的AI模型，仅需1.4秒就能完成未来24小时的全球气象预报，相比传统预报提速10000倍以上。华为云盘古大模型研发团队提出了一种新的高分辨率全球人工智能气象预报系统，可秒级完成对位势、湿度、风速、温度、海平面气压等全球气象的预测，直接应用于多个气象研究细分场景。此前，华为云盘古气象大模型研究团队使用1979至2021年共43年的全球气象数据，训练深度神经网络，并创造性地提出了适应地球坐标系统的三维神经网络来处理复杂的不均匀3D气象数据，以此在精度和速度方面超越传统数值预测方法。相比传统数值预报，其计算速度提升超过10000倍。此外，基于文心大模型开发出的智舱大模型与智舱开发工具链，可使汽车具有更高阶的智能驾驶能力、更强大的自我学习和记忆功能，改变人车交互方式、提升用户用车体验。

《2023年中国生成式AI企业应用研究》预计，2035年中国生成式AI企业采用率将达到约85%。制造业、零售业、电信行业和医疗健康领域率先采用，其采用率分别达到82%、90%、65%和53%。

材料五

2022年12月，清华大学交叉信息研究院做了一个AI模型性别歧视水平评估项目，在包含职业词汇的“中性”句子中，由AI预测生成一万个模板，研究团队再统计AI模型对该职业预测为何种性别的倾向，当预测偏误和刻板印象相符，就形成了算法歧视。测试模型包括ChatGPT前身GPT-2。测试结果发现，GPT-2有70.59%的概率将教师预测为男性，将医生预测为男性的概率则是64.03%。评估项目中，其他首测的AI模型还包括Google开发的BERT以及Facebook开发的RoBERTa。所有受测AI对于测试职业的性别预判，结果倾向都为男性。

“它会重男轻女，爱白欺黑（注：种族歧视）。”研究人员表示。AI的歧视，早有不少案例研究。如AI图

像识别，总把在厨房的人识别为女性，哪怕对方是男性。2015年6月，Google照片应用的算法甚至将黑人分类为“大猩猩”，Google公司一下被推上风口浪尖。

2019年12月16日，意大利总工会位于B市的运输业劳动者工会，商业、旅游与服务业劳动者工会以及非典型劳动者工会将意大利户户送有限责任公司诉至B市法院。原告认为，被告公司使用的算法具有集体歧视性，因为该算法在对骑手进行荣誉排名时并不考虑骑手不赴约的原因，骑手可能因参加罢工或疾病、未成年子女需求等其他合法原因导致延迟取消预定或没有取消预定，从而造成其荣誉排名降低，丧失优先选择工作条件的机会。因此，原告请求法院对被告的歧视行为加以确认，并同时提出要求被告制定消除歧视的方案、修改自助预定系统、对上述歧视性行为给原告造成的非财产性损害进行赔偿、采取一切适当措施消除歧视性行为的影响等请求。2020年12月31日，B市法院做出判决，该判决针对本案是否仍然具有诉的利益、针对骑手是否可以适用反歧视法律规定、工会组织是否具备原告资格、歧视的含义、本案中被告使用的算法是否对骑手构成歧视、举证责任以及是否应该进行损害赔偿等问题一一进行了审查并说理。最后，法院支持了原告请求，确定并宣布被告通过数字平台使用的算法具有歧视性，要求被告在其网站及平台的“常见问题”板块公布法院判决，以消除歧视行为的影响。

材料六

2023年11月1日，28个国家和地区在首届人工智能安全峰会上签署《布莱切利宣言》（以下简称《宣言》），同意协力打造一个“具有国际包容性”的前沿人工智能安全科学研究网络，以对尚未被完全了解的人工智能风险和能力加深理解。

《宣言》提到，目前较大的安全风险出现在人工智能的“前沿”领域，即前沿AI，也就是那些能力强大的通用人工智能模型，包括可以执行各种任务的基础模型，以及特定的弱人工智能，两者都具有给人类社会带来伤害的能力。如果上述前沿AI被故意滥用，或者在正当的使用过程中出现意外失控，都将会带来巨大风险。《宣言》尤其关注网络安全和生物技术等领域的风险，以及前沿AI可能放大虚假信息等风险。

近期，美国联邦调查局（FBI）向公众发出警告，提醒人们防范利用“深度伪造”技术制作虚假色情照片或视频所实施的骚扰或勒索行为。声明中称，技术的不断进步使人工智能生成的内容在质量、定制性和易用性方面不断提高，这也导致一些受害人报告他们的照片或视频遭到了恶意篡改。FBI表示，恶意行为者通常从某人的社交媒体账户、互联网和受害者本人获取内容，制作出具有暧昧性并看似与受害者相似的内容，而后通过社交媒体、在线论坛或网站进行传播。根据该机构的说法，针对这些受害者制作的内容通常是出于获取更多个人信息、经济利益或骚扰他人的动机。

2023年6月，16人匿名提起诉讼，声称OpenAI与其主要支持者微软公司在数据采集方面违背了合法获取途径，选择了未经同意与付费的方式，进行个人信息与数据的收集。起诉人称，OpenAI为了赢得“AI军备

竞赛”，大规模挪用个人数据，非法访问用户与其产品的互动以及与ChatGPT集成的应用程序产生的私人信息。根据这份长达157页的诉讼文件，OpenAI以秘密方式从互联网上抓取了3000亿个单词，并获取了“书籍、文章、网站和帖子——包括未经同意获得的个人信息”，行为违反了隐私法。

材料七

2023年3月，1000多名人工智能专家和行业高管签署了一份公开信，信中说：“具有人类竞争力智能的AI系统可能会对社会和人类造成深刻的风险，这一点已经被广泛研究并得到了顶级AI实验室的认可。正如广泛认可的AsilomarAI原则中所述，高级AI可能代表地球生命史上的深刻变化，应以相应的关怀和资源进行规划和管理。不幸的是，目前尚无这种级别的规划和管理。最近几个月人工智能实验室陷入了一场失控的竞赛，他们致力于开发和部署更强大的数字思维，但是没有人能理解、预测或可靠地控制这些大模型，甚至模型的创造者也不能。当代人工智能系统现在在一般任务上表现与人类相当，我们必须扪心自问：我们是否应该让机器用宣传和虚假信息充斥我们的信息渠道？我们是否应该自动化所有的工作，包括那些富有成就感的工作？我们是否应该开发非人类的思维，可能最终会超过我们、智胜我们、使我们过时并取代我们？我们是否应该冒险失去对我们文明的控制？”

“这些决策不能委托给未经选举的技术领袖。强大的AI系统只有在我们确信其影响将是积极的、风险可控的情况下才应该开发。这种信心必须有充分的理由，并随着系统潜在影响的程度增加而增强。”OpenAI最近关于人工智能的声明指出，“在开始训练未来系统之前进行独立审查可能很重要，并且对于最先进的工作来说，人们应该在合适的时间点限制用于创建新系统的计算增长率。我们认为，现在已经到了这个时间点。”

同时，AI开发人员必须与决策制定者合作，大力加速强大的AI治理系统的发展。这些治理系统至少应该包括：专门针对AI的新型和有能力的监管机构，对高能力的AI系统和大量计算能力的监督和跟踪；帮助区分真实和合成数据的溯源、水印系统，并跟踪模型泄露；一个强大的审计和认证生态系统；对由AI引起的损害的责任，针对AI安全研究的强大公共资金；资源丰富的机构，以应对人工智能将造成的巨大经济和政治混乱。

材料八

我国是人工智能大国。数据显示，我国人工智能核心产业规模达到5000亿元，企业数量超过4300家，创新成果不断涌现。世界知识产权组织的数据显示，2019年我国企业和机构申请了近3万项人工智能相关专

利，占全球人工智能专利申请的40%以上。在技术和相关产业发展的同时，中国一直致力于提升人工智能技术的安全性、可靠性、可控性、公平性。

2023年8月15日，国家网信办联合国家发展改革委、教育部、科技部、工业和信息化部、公安部、广电总局公布的《生成式人工智能服务管理暂行办法》正式施行。对生成式人工智能服务实行包容审慎和分类分级监管，明确了提供和使用生成式人工智能服务总体要求。提出了促进生成式人工智能技术发展的具体措施，明确了训练数据处理活动和数据标注等要求。规定了生成式人工智能服务规范，明确生成式人工智能服务提供者应当采取有效措施防范未成年人用户过度依赖或者沉迷生成式人工智能服务，按照《互联网信息服务深度合成管理规定》对图片、视频等生成内容进行标识，发现违法内容应当及时采取处置措施等。

目前，全球AI立法进程明显提速，世界各国的监管都在追赶AI的演化速度。欧盟在尝试推动出台人工智能监管政策方面是先行者。2021年，欧盟委员会推出了全球第一个关于人工智能的法律框架——欧盟人工智能法案。该法案草案的一个突出特点是注重基于风险来制定监管制度，以平衡人工智能的创新发展与安全规范。

在欧盟委员会设计的框架中，根据应用场景、使用技术等多个方面，人工智能技术被划分为四个不同的风险等级，配有不同的监管要求。当时的框架设计者们相信，尽管人工智能是一项快速发展的技术，但是欧洲的制度设计足以适应技术变化。但事实似乎并不尽然。仅仅几年间，人工智能技术的发展就让那些规则设计者们感叹欧盟的设计似乎已经过时了。欧盟人工智能法案的主要起草人之一、来自德国的欧洲议会议员W指出，人工智能技术在两年前还没有这么先进，而未来两年中还会进一步发展，“如此之快”的发展速度让当时的大部分法律设计在实际生效时可能已经不再适用了，因此在ChatGPT等生成式AI风靡后，欧盟立法者又紧急添“补丁”。一个新变化是《人工智能法案》最新草案加强了对通用人工智能的透明度要求。例如，基于基础模型的生成式AI必须要对生成的内容进行标注，帮助用户区分深度伪造和真实信息，并确保防止生成非法内容。像OpenAI、Google等基础模型的提供者，若是在培训模型期间使用了受版权保护的数据，也需要公开训练数据的详细信息。公共场所的实时远程生物识别技术从“高风险”级别调整为“被禁止”级别，即不得利用AI技术，在欧盟国家的公共场合进行人脸识别。

英国采取的人工智能监管思路与欧盟提出的监管设计框架截然不同。英国有意对人工智能研发和使用采取宽泛的监管原则，实施更为灵活、平衡的监管办法。此外，英国政府希望，由行业监管机构在参考政府的一系列指导原则的基础上，制定具体的监管办法。

美国与英国近似，监管思路更注重利用AI。推动AI行业的创新和发展，最终是为了保持美国的领导地位和竞争力。例如，《人工智能权利法案蓝图》作为美国人工智能治理政策的里程碑事件，制定了安全有效的系统、防止算法歧视、保护数据隐私、通知及说明、人类参与决策制定等五项基本原则，但并无更加细致条款。

材料九

近日，美国与欧盟达成了一项号称“关乎互联网未来”的人工智能合作协议。这项名为“人工智能促进公共利益行政协议”的协议通过线上签署，拟在预测极端天气和应对气候变化、应急响应、医保事业、电网运行，以及农业发展等五大重点领域带来公共利益。一名美国政府官员称，此前相关合作协议仅限于加强个人隐私保护等涉及AI的特定领域，而此次双方将合作建立人工智能联合模型。以电网为例，美国和欧洲国家政府都在收集发电、电力应用以及如何平衡电网载荷以应对天气变化等方面的数据。双方达成新协议后，可利用这些数据共同建模。为负责应急管理、电网运营等人员提供更好的决策方案。据悉，目前合作仅限于美欧之间，但美方表示，未来几个月或有其他国家受邀加入。

然而，此次联合建模并不相互共享训练数据集，即美国的数据留在美国，欧洲的数据留在欧洲，双方在数据流通上仍有所保留。因为美国和欧洲在关于数据跨境流动上有着不同的立场。具体而言，美国拥有多家大型数字平台企业，因此其积极鼓励其他国家数据流入，破除各国跨境数据流动壁垒，以获取足够可供利用的数字资源；而这恰恰是欧洲所不具备的，因而后者更加强调数据治理与数字主权。随着《隐私盾协议》的失效，且美欧之间尚未建立起数据跨境流通的新框架，他们的合作需要寻求一种符合双方诉求和利益的模式，这在某种意义上也是一种妥协。从目前美欧之间关于数据跨境的一些妥协、协调、合作来看，美欧在面对外部挑战的时候，会最大限度地做出相互的妥协和让步，改变内部制度来满足彼此的需求，以此换取对外的相对一致的合作。

有专家预测，“行政协议”可能会涉及与人工智能，包括数据问题相关的标准设定、产业转入门槛问题、商业生态的对接融合问题，以及所谓的伦理价值观标准问题等，这些层面的合作推进也能降低美欧之间的差异程度和潜在冲突，换取更高水平的合作，这正是他们努力的方向。目前合作的清单多集中在气候、医疗、农业等偏重公共利益、较易达成共识的领域，其实还有部分较敏感的领域是美国希望下一步能够推进的，比如自动驾驶。

材料十

在人工智能技术的发展中，中美两国都具有较高的实力和潜力，美国的人工智能公司，如谷歌、亚马逊、微软等在人工智能研发领域的投入和持续发展方面均处于领先地位，而中国的人工智能企业在语音识别、人脸识别、智能驾驶等领域也有着相当的优势和创新成果。

中美两国企业在人工智能领域的合作非常广泛。一方面，中国的人工智能企业在技术研发方面拥有非常强大的实力和资源，他们往往是中美人工智能合作的主要推动方，与美国的谷歌、微软等公司合作进行

技术研发或者达成软件开发协议。例如，百度AI与微软合作，开发了一个“智能聚合平台”，这个平台可以帮助企业提高工作效率和产品创新等。此外，阿里巴巴也通过合作策略优化了公司运营，计划在未来几年投入1500亿元人民币用于人工智能领域的研发。另一方面，作为世界领先的科技强国，美国在人工智能领域也有许多优秀的企业。例如，IBM正在与中国的企业合作，开发人工智能技术，改善世界的客户服务，提高工作效率，开发和部署人工智能技术和应用程序。谷歌与中国的企业也在人工智能的领域开展广泛的合作，例如谷歌的TensorFlow深度学习平台就是与中国的发行商合作达成的。从合作方式上来看，中美两国的人工智能企业合作形式多种多样，其中包括立项合作、技术引进、联合创新、成果共享等，每种形式都有不同的优点和适用范围。实际上，中美两国在人工智能技术上的合作已经在很多领域得到了非常好的应用，不仅有助于两国之间的技术共享，还有助于人工智能技术在其他领域的加速应用。

材料十一

2023年10月，习近平主席在第三届“一带一路”国际合作高峰论坛开幕式上的主旨演讲中提出《全球人工智能治理倡议》，这是我国积极践行人类命运共同体理念，落实全球发展倡议、全球安全倡议、全球文明倡议的具体行动。

问题一

根据给定资料，阐述为什么说新一代人工智能技术是“颠覆性”技术。（15分）

要求：语言精练，层次要点分明，字数不超过300字。

问题二

根据给定资料，分析新一代人工智能技术社会风险的成因。（15分）

要求：分析准确，条理清楚，语言精练，字数不超过300字。

问题三

结合给定资料，阐述全球各国应如何深化人工智能技术领域合作，引导人工智能朝着有利于人类文明进步的方向发展。（20分）

要求：条理清晰，内容准确全面，语言精练，字数不超过300字。

问题四

阅读但不拘泥于给定资料，以“人工智能时代，公共部门何为？”为主题写一篇文章。（50分）

要求：

- （1）自选角度，自拟题目；
- （2）观点明确，联系实际，分析具体，条理清楚，语言流畅；
- （3）总字数800-1000字。

参考答案

问题一

公众号：橙子考公资料站【答案】

1. 突破性解决科学难题。破解生物医学领域长期难题，增进人类对基本生物过程的理解，促进药物设计，加速生物和医学研究。
2. 改变学习教育方式。全面掌握学生基本情况，帮助学生多技能、针对性、创造性发展，促进教育公平；减少教师低级负担，改善教学工作，有效指导学生。
3. 创新内容生成方式。自动化学习知识信息，不断自我优化；充分理解和流畅表达人类语言，逻辑推理强；拥有自适应和迁移学习能力，自主生成各类型数据，适用于多场景和任务，促进各领域发展。
4. 代替和拓展人类能力。高效完成简单重复性及通过学习、训练得到且无需天分的技能，省时省人力；弥补人类感官缺陷，拓展人类感知和行动能力；催生新的产业和服务模式，推动未来社会变革。

问题二

公众号：橙子考公资料站【答案】

1. 存在算法歧视。在职业性别预测中出现性别、种族歧视，并存在集体歧视性，忽略实际情况，决策结果不公平，伤害用户利益。
2. 滥用前沿AI。在正当的使用过程中出现意外失控，出于获取更多个人信息、经济利益的动机，利用“深度伪造”技术篡改受害人的照片、视频，制作虚假色情照片或视频，泄露个人隐私，并通过社交媒体、在线论坛或网站进行传播，实施骚扰或勒索行为。
3. 缺少相应级别的规划和管理。无人能理解、预测、控制人工智能大模型，缺乏独立审查机制，创建新系统的计算增长率无限制，其影响及风险不可控；开发人员与决策制定者缺乏合作，人工智能治理系统发展缓慢。

问题三

公众号：橙子考公资料站【答案】

1. 提速AI立法进程，加大监管力度。中国多部门联合出台实施《生成式人工智能服务管理暂行办法》，明确使用要求，实行审慎监管；欧盟委员会推出人工智能法案框架，注重基于风险来制定监管制度；英国实施更为灵活、平衡的监管办法，由行业监管机构制定具体的监管办法。

2. 建立人工智能联合模型，促进数据跨境流通。签署合作协议，涵盖与人工智能的对接融合、伦理价值观标准问题等方面，共同谋求重大领域的公共利益，推动训练数据集共享，推进各国在敏感领域的合作交流。

3. 开展优势互补，实现技术共享。中国企业利用自身技术研发的强大实力和资源，与美国优秀公司采用多元化合作方式，推动技术共享及人工智能技术的多领域加速应用。

问题四

公众号：橙子考公资料站【答案】

公共部门需有为智变时代启新篇

ChatGPT的爆火，让更多人知道人工智能不仅仅是智能手机、智能家具、自动驾驶，还可能是语言、文字、思考的交互。但人工智能不仅改变着我们的生产方式、生活方式，也对社会治理、公共服务提出了新的挑战。面对这一历史性的变革，公共部门作为社会治理的重要力量，必须积极应对，承担起应有的角色与担当。

公共部门要加强国际合作与交流，共同推动人工智能建设和发展。要促进国际间先进技术的共享，进行优势互补，助力人工智能在各领域的进步和应用。俗话说：“单丝不成线，独木不成林”，世界各国合则两利、斗则俱伤。在人工智能领域，世界大国都有各自的优势，比如美国在人工智能研发领域的投入方面处于领先地位，我国则在语音识别、人脸识别、智能驾驶等领域有着相当的优势和创新成果，两国的合作有助于技术的研发和落地。在合作中要相互尊重、彼此信任、求同存异，让数据流动起来，让技术腾飞起来，让成果照进现实。

公共部门需强化人工智能风险的研判和防范，引导科技向善。加深对人工智能风险的理解和能力认识，做好人工智能伦理教育，提升公众对人工智能的认知和风险意识。随着人工智能研发技术脚步的加快，越来越多的社会问题产生了，比如恋爱骗局、版权威胁、病毒勒索、电信欺诈等，这些问题影响了公众的生产生活，破坏了网络安全，造成了社会经济的损失。对于随时可能发生的“黑天鹅”“灰犀牛”事件，公共部门不仅要提前意识到人工智能发展可能产生的不确定性，还必须做到居安思危、未雨绸缪，坚持预防为主思路，在多渠道、多平台、多领域开展警醒宣传，防范于未然。

公共部门还需不断完善治理体系，确保人工智能安全可控。加快立法脚步，建立监管体系，明确技术提供和使用要求，依法惩处违法行为，维护社会公平正义。古人云：“近来见得天地之道，刚柔互用，不可偏废，太柔则靡，太刚则折。”对于人工智能的风险化解规避亦是如此。立法要增强其预见性、前沿性，多角度、全方位健全完善法律法规，为技术的合规使用提供支撑，为其高速发展保驾护航，同时还要利用

与监管相结合，鼓励人工智能创造，对不同领域、不同程度的风险，实施不同力度的处罚，避免一罚了之，保证发展活力。

人工智能是新时代的星星之火，我相信在政府部门多角度的努力下定能形成燎原之势，打造未来的新期待。