

知识领域	核心技术/概念	典型应用场景	代表工具/框架	关键挑战	2023年技术指标
机器学习	- 监督学习 (SVM、决策树)	- 金融反欺诈 (F1值 >0.92)	Scikit-learn、XGBoost、LightGBM	- 特征工程耗时 (占项目周期 60%)	自动特征工程工具 (Feature Tools) 减少 80% 人工参与
	- 无监督学习 (K-means、DBSCAN)	- 零售商品推荐 (CTR提升 15-30%)		- 冷启动问题	
	- 半监督学习	- 工业设备预测性维护		- 概念漂移	
	- 集成学习 (Bagging/Boosting)				
深度学习	- CNN (ResNet、EfficientNet)	- 医学影像病灶检测 (灵敏度 >97%)	PyTorch、TensorFlow、Keras	- 模型参数量达千亿级	混合精度训练加速 30%，显存占用降低 40%
	- RNN (LSTM、GRU)	- 视频内容审核 (准确率 99.5%)		- 训练能耗高 (GPT-3 耗电 1,287MWh)	
	- Transformer (ViT、Swin)	- 语音合成 (自然度 MOS >4.2)		- 黑箱决策风险	

计算机视觉	- 神经网络搜索 (NAS)			
	- 2D/3D 目标检测	- 无人机巡检 (缺陷识别率 92%)	- 小样本学习 (<100 标注样本)	
	- 实例分割 (Mask R-CNN)	- 智能仓储 (拣选效率提升 3倍)	- 跨域泛化能力	实时目标检测 FPS>120
	- 光流估计	- 虚拟试妆 (AR 贴合误差 <2mm)	OpenCV、MMDetection、Detectron2	(NVIDIA Jetson AGX平台)
	- 超分辨率重建 (SRGAN)			
自然语言处理	- 预训练语言模型 (BERT、GPT-4)	- 法律文书审查 (效率提升50倍)	- 长文本连贯性 (>1,000 字)	
	- 文本生成 (T5、PEGASUS)	- 舆情分析 (情感分类准确率 89%)	HuggingFace、spaCy、NLTK	大模型上下文窗口扩展至32k tokens (如 Claude 2)
	- 知识蒸馏	- 多语言客服 (支持 83种语言)	- 事实一致性校验	

强化学习	- 多语言对齐			
	- 深度Q网络 (DQN)	- 能源网络优化 (成本降低12-18%)	- 稀疏奖励场景	
	- 策略优化 (PPO、SAC)	- 芯片布局设计 (效率提升1000倍)	Stable Baseline s3、Ray RLLib	MuZero 算法在 Atari 游戏超越人类水平30%
	- 多智能体系统	- 金融交易策略	- 模拟到现实的差距	
	- 逆向强化学习			
知识图谱	- 本体建模 (OWL)	- 金融反洗钱 (关系网络分析)	- 多源异构数据融合	
	- 图神经网络 (GAT、GraphSAGE)	- 制药知识发现 (研发周期缩短20%)	Neo4j、Apache Jena、DGL	- 时效性维护 (日更新量>百万条)
	- 知识推理 (规则引擎)	- 智慧城市事件推理	- 隐含关系挖掘	千亿级三元组知识库构建 (如 Wikidata)
	- 动态知识更新			
	- 文本到图像 (ControlNet、LoRA)	- 影视可视化 (成本降低70%)	- 版权归属界定	

生成式AI	- 视频生成 (Pika、Sora)	- 个性化教育内容生成	- 深度伪造检测 (AUC<0.7)	- 文本到3D生成质量 (PSNR>28dB)
	- 3D生成 (NeRF、Gaussian Splatting)	- 新材料发现 (周期从年缩短至天)	Midjourney、ComfyUI、Diffusers	(NVIDIA GET3D)
	- 分子生成 (AlphaFold 3)		- 物理规律符合性	
	- 模型量化 (INT8训练)	- 智能驾驶 (端到端延迟<50ms)	- 异构硬件适配	
边缘智能	- 知识蒸馏 (教师-学生网络)	- 工厂设备异常检测 (准确率>95%)	TensorRT、ONNX Runtime、TF Lite	1W功耗设备运行 (ResNet-18 (30FPS))
	- 联邦学习 (纵向/横向)	- 隐私医疗数据分析	- 安全攻击防护	
	- 边缘推理优化			
	- 差分隐私 ($\epsilon < 2$)	- 信贷审批偏差消除 (公平性提升40%)	- 隐私-效用权衡	

AI伦理与治理

- 模型可解释（LIME、SHAP）
 - 公平性指标（统计均等、机会均等）
 - 红队测试
 - 招聘算法去偏见
 - 内容审核透明度提升
- IBM AIF360、Fairlearn、Captum
- 跨文化伦理标准差异
 - 对抗样本攻击成功率 >30%
- 欧盟AI法案分类监管体系（4级风险分类）