

규 격 / 사 양 서

1. 개요	<table><tr><td>품명</td><td colspan="3">서버 (심층 학습 전용 자체 제작 서버)</td></tr><tr><td>수량</td><td>1EA</td><td>제조사</td><td>-</td></tr></table>			품명	서버 (심층 학습 전용 자체 제작 서버)			수량	1EA	제조사	-
	품명	서버 (심층 학습 전용 자체 제작 서버)									
수량	1EA	제조사	-								
2. 용도 및 특징	고성능 CPU와 GPU를 탑재한 서버로 딥러닝 연산에 최적화 되어있음										
3. 기기구성 (세부내용, 수량 등)	아래 사양 참조 (서버 1EA)										
4. 규격 및 사양	* 본 기본 규격서는 최소한의 규격을 규정하며 반드시 본 규격서에 명시된 동등 및 동급 이상의 제품을 납품하여야 한다. 장비: SuperMicro GPU Workstation 741GE-TNRT - Dual Socket Xeon 4th Gen Scalable processors - 16 DIMMs; up to 4TB ECC Reg DDR5-4800MHz - 4 PCI-E 5.0 x16 (double-width), 3 PCI-E 5.0 x16 (single-width) - 2 M.2 NVMe - 8x 3.5" hot-swap NVMe/SATA/SAS drive bays (8x 2.5" NVMe hybrid) - 2x 10GbE BaseT with Intel® X550-AT2 - 2000W (1+1) Redundant Power Supplies CPU: Intel 5th Gen Xeon Scalable Processor (8-core)** 6534 (8/16) 3.9GHz 22.5M * 2EA RAM: Samsung 32GB DDR5 Registered ECC PC5 5600 * 10EA Storage: Samsung Enterprise NVMe PM9A3 3.84TB (2.5in NVMe U.2) * 1EA GPU: NVIDIA Blackwell RTX PRO 5000 D7 48GB, 300W Blower * 4EA S/W 설치지원 (Ubuntu, CUDA Toolkit, cuDNN, Nvidia-Docker, Slurm 포함) MLOps: UYUNI MLOps (또는 이에 준하는 S/W) - Docker/Kubernetes 기반 딥러닝 플랫폼 솔루션 - GPU 클러스터 모니터링 및 사용자별 사용 이력 확인 - Job Scheduling 관리 - Multi-instance GPU 기능 지원을 통한 GPU 인스턴스 분할 사용										
5. 설치, 교육 및 사후관리	- 추가로 설치가 필요한 OS/소프트웨어: UYUNI MLOps, Ubuntu, CUDA, nvidia-driver, cudnn, Docker, Slurm 등 (본교 사설망 환경에 맞는 보안 환경 설정 포함) - 보증기간 동안 추가 설치 소프트웨어에 대한 사용자 방문 교육 연 2회씩 진행 - 보증기간 동안 HW 및 추가 설치 SW에 이상 발생시 익일 방문 필										

6. 품질보증 기간 (Warranty)	- 2년 무상 부품 및 현장 지원, HW/SW 문제 발생시 익일 방문 필
7. 납기일자	계약일로부터 60일 이내
8. 계약업체 제출서류	
9. 참가자격	- 전담 HW/SW 엔지니어가 업체 내에 상주하여 HW/SW 문제 발생시 당일/익일 방문 지원이 반드시 가능해야함 - 부품만 납품하는 업체는 참가 불가 - 위에 포함된 제품보증, CUDA version 등 모든 내용이 납품될 장비에 반드시 반영되어야 하며, 추가로 기존 보유 장비 또한 CUDA version, MLOps 등 동일한 개발 환경으로 셋업이 반드시 가능해야 함
10. 참가자격 제출서류	
11. 기타사항	- 입찰시 사전에 구매자와 설치/교육/사후관리 및 납품 일정에 대해 사전 설명 필 - 문의처: 이화여자대학교 통계학과 주원영 (010-7497-4447)