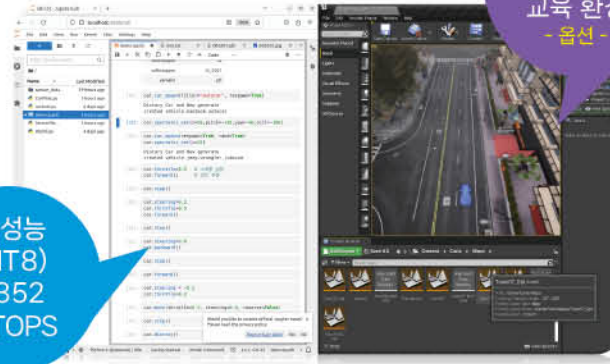


생성형 인공지능 실습장비 GenWorks

국방과학연구소(ADD)
AI 프로젝트 지원 설계로
CUDA를 비롯한 인공지능 프레임워크
사전 설치

AI 성능
(INT8)
3,352
AI TOPS

한백전자
CARLA
교육 환경
- 옵션 -



- 최대 24코어, 32스레드 CPU와 3,352 AI TOPS GPU 내장
- 대용량 데이터 모델 파라미터 처리를 위해 64GB 메인 메모리와 2TB NVMe SSD 기본 탑재
- 장시간 풀로드 생성형 AI 워크로드를 안정적으로 수행할 수 있도록 GPU-CPU 발열을 고려한 고효율 쿨링 및 전원 설계 적용
- 교육 현장에 최적화된 최소 시스템 패키지 우분투 LTS 배포판 제공으로 불필요한 서비스는 제거하고 생성형 AI 개발 환경 중심으로 재구성
- 실습 후 언제든지 초기 상태로 되돌릴 수 있도록 시스템 이미지 기반 파일 시스템 복구 모드 및 원격 원격 초기화 기능 지원
- 4K 대화면 디스플레이를 통해 프롬프트-응답 과정, 이미지-영상 생성 결과, 시뮬레이터 화면 등 생성형 AI 결과물의 시각적 확인 지원
- CUDA 기반 PyTorch-TensorFlow는 물론 Hugging Face Transformers/Diffusers 등 생성형 AI에 특화된 라이브러리 및 실행 환경 사전 구성
- 텍스트 생성, 코드 생성, 이미지 생성 등 대표적인 생성형 AI 실습을 위해 오픈 소스 기반 대형 언어 모델 및 확산 모델과 예제 프로젝트를 사전 탑재
- 개방형 디지털 자산이 포함된 오픈 소스 기반 CARLA 시뮬레이터를 옵션으로 제공해 자율주행·로보틱스와 생성형 AI를 연계한 시나리오 생성·평가 실습 지원
- 전문적인 응용 개발을 위해 Visual Studio Code 기반 통합 개발환경과 Jupyter, Git, Docker 등 생성형 AI 개발에 필요한 확장 기능 지원
- 강의에 바로 적용 가능한 생성형 인공지능 교육 콘텐츠를 제공하여 이론-실습-프로젝트로 이어지는 일관된 교육 과정 구성 지원

교육 콘텐츠

- 생성형 인공지능 개요와 Transformer, 실습 환경 구축
- 딥러닝 모델의 기초 (Forward/Backward, LeNet, Autograd)
- 컴퓨터가 자연어를 이해하는 첫 걸음, Tokenization과 Attention의 모든 것
- GPT를 중심으로 한 생성형 언어모델의 구조와 활용
- 대규모 언어 모델을 더 빠르고 가볍게
- GPT 학습하기
- RAG 이해
- 능동적인 에이전트로서의 LLM, MCP 에이전트
- 멀티모달 LLM의 이해
- Diffusion의 원리 및 실습
- Stable diffusion에 대한 이해

HANBACK ELECTRONICS

HANBACK
ELECTRONICS
(주) 한 백 전 자
Since 1984

생성형 인공지능 실습장비 GenWorks



HANBACK ELECTRONICS

대전광역시 유성구 유성대로 518

TEL. 042. 610. 1111 (1114) FAX. 042. 610. 1199

E mail. edu@hanback.co.kr

본 카탈로그의 제품사양 및 외형은 품질개선을 위해 예고 없이 변경될 수 있습니다. V1.0.0



홈페이지 바로가기

운영프로그램 사양

생성형 인공지능 실습장비

List		Specifications
OS	Kernel	Linux Kernel 6.x
	Desktop	Window Manager: OpenBox Display Manager: LightDM Panel: Tint2 System Monitor: conky Bluetooth Manager: blueman Network Manager: network-manager Graphics: NVIDIA Driver, CUDA V13.x, Runtime [cudart]
	CLI	Shell: Zsh, Oh-My-Zsh with powerlevel9k thema and nerd fonts Terminal Multiplexer: Tmux Base Tools: fzf, bat, lsd
	Tool Chain	Python3, NodeJS, Java, Clang, GCC, LLVM
	IDE	Visual Studio Code for LLM
	Library	Data analysis: numpy, matplotlib, Pandas, SciPy, seaborn, plotly, Folium Image Processing: OpenCV, Mahotas, SimpleCV, Pillow, Scikit-Image, SimpleITK, Pgmagick Audio Processing: Pydub, Librosa, Soundfile, Pedalboard, Supriya, Libsox, Pyo, Pippi
AI	CUDA Toolkit, cuDNN, TensorRT	Pre-installed CUDA Toolkit, cuDNN, and TensorRT stack optimized TensorRT-LLM components for high-throughput FP8/INT8/FP16 large language model inference and KV-cache optimization
	Framework (PyTorch, TensorFlow, Keras, Scikit-Learn)	PyTorch, TensorFlow, Keras, and Scikit-Learn with full GPU support enabled Hugging Face Transformers and Diffusers libraries for large language models and diffusion-based image generation LangChain and related libraries for retrieval-augmented generation (RAG) and agent orchestration NVIDIA NeMo framework and NVIDIA TAO Toolkit containers for NVIDIA-optimized model training, adaptation, and deployment
	Algorithm	Machine Learning: Linear Regression, Logistic Regression, Decision Tree, SVM, Naive Bayes, KNN, K-means, Random Forest, Dimensionality Reduction (e.g., PCA), Gradient Boosting (XGBoost/LightGBM) and AdaBoost Object Detection: HOG, R-CNN and Fast/Faster R-CNN, SSD, YOLO family, RetinaNet Generative AI: Transformer encoder-decoder and decoder-only architectures, BERT, GPT-style language models, Llama 2/3 family, text embedding models, RAG pipelines, LoRA/QLoRA configuration utilities
CARLA - 옵션 -	NVIDIA Model	NVIDIA TAO pre-trained computer vision models: TrafficCamNet, DashCamNet, PeopleNet, PeopleSegNet, VehicleMakeNet, VehicleTypeNet, LPDNet and related variants NVIDIA face and gaze models: FaceDetect, FaceDetectIR, FaceNet, FPENet, GazeNet, EmotionNet, HeartRateNet NVIDIA speech and audio models: WaveGlow, Flowtron, and fine-tuning configurations for Flowtron and other NeMo-based TTS models NVIDIA large language and generative models optimized for RTX 5090, including the Nemotron family and NVIDIA-optimized open LLM checkpoints Language identification models including LangID PearlNet and related NVIDIA NeMo language ID models
	Unreal Engine 4.x	
	Scalability via s server multi-client architecture	
CARLA - 옵션 -	Actors and blueprints	Managing the blueprint library Actor life cycle: Spawning, Handling, Destruction Types: Sensors, Spectator, Traffic signs and traffic lights, Vehicles, Walkers Add a new vehicle
	Maps and navigation	Maps: Changing the map, Landmarks, Waypoints, Lanes, Junctions, Environment Objects Navigation: Navigating through waypoints, Generating map navigation CARLA maps: Non-layered maps, Layered maps, Custom maps HBE maps: SimpleTrack Generate maps with OpenStreetMap Generating Maps in RoadRunner
	Traffic Manager	Localization State Collision State Traffic Light State Motion Planner State Vehicle Lights State In-Memory Map Path Buffer and Vehicle Tracking PID controller
	Sensors	Cameras: Depth, RGB, Optical Flow, Semantic segmentation, Instance segmentation, DVS Detectors: Collision, Lane invasion, Obstacle Other: GNSS, IMU, LIDAR, Radar, RSS, Semantic LIDAR
	Custom assets	Adding props, Adding vehicles, Packaging assets, Custom materials

HANBACK ELECTRONICS

구성품

• GenWorks



• 교재



CARLA 시뮬레이터를 지원하는 인공지능 워크스테이션

하드웨어 사양

List		Specifications
Body	Size	505 x 230 x 480 mm
	Weight	20kg
	Material	SGCC Steel, Ultra-clear Tempered Glass
	Color	White
Technical	CPU	Performance cores 8 + Efficient cores 16 Total Threads 32 MAX Turbo Frequency 6.00GHz Performance-core MAX Turbo Frequency 5.60GHz Efficient-core MAX Turbo Frequency 4.40GHz Performance-core MAX Base Frequency 3.20GHz Efficient-core MAX Base Frequency 2.40GHz Cache 36MB, Total L2 Cache 32MB
	GPU	CUDA Cores: 21,760 Boost Clock: 2.41GHz Base Clock: 2.01GHz Memory Size: 32GB (GDDR7) Memory Interface Width: 512bit Memory Bandwidth: 1,792 GB/s Ray Tracing Cores: 4th Generation (RTX Blackwell) Tensor Cores: 5th Generation, up to 3,352 AI TOPS NVIDIA Architecture: Blackwell NVIDIA Encoder (NVENC): 3x 9th Generation NVIDIA Decoder (NVDEC): 2x 6th Generation PCI Express Interface: PCIe 5.0 x16 Display Outputs (FE): 3x DisplayPort 2.1b, 1x HDMI 2.1b Maximum Resolution & Refresh Rate: 4K at 480Hz or 8K at 120Hz [with DSC, HDR] Maximum GPU Temperature: 90°C Total Graphics Power (TGP): 575W
	Base Board	Memory: 64GB DDR 5x, Tested Speed 6400MT/s Storage: PCIe 4.0 NVMe M.2 SSD 2TB. Reading MAX 7,000MB/s, Writing: MAX 6,500MB/s Audio: ALC1220 Codec, 32-bit/92kHz DAC, 120dB SNR DAC with Differential Amplifier Connectivity: RTL8125BG 2.5GbE LAN, Wi-Fi 6E AX211, Bluetooth v5.2 VRM: 16[VCORE] + 1[VCCGT] + 2[VCCIN_AUX] Interface: PCIe x16 x 3ea, PCIe x1 M.2Socket[Key E] x 1ea USB 2.0 x 2ea, USB 3.2 [20Gbps] Type-C x 1ea, USB 3.2 [10Gbps] x 1ea, USB 3.2 [5Gbps] x 1ea,
	Cooling System	Water Block: Copper Housing, 75 x 56 mm Pump: Motor Speed 800~2800RPM Radiator: Aluminum, 121 x 394 x 27 mm Tube: Ultra-low Evaporation Rubber with Nylon Braided Sleeve, 400mm FAN: Top 120 x 120 x 26mm x3ea Front 120 x 120 x 26mm x 3ea Rear 140 x 140 x 26mm x 1ea
	Power	Cable Type : Full Modular DC Output : 1000W[+3.3V 20A, +5V 20A, +12V 83A, -12V 0.3A, +5Vsb 3A]
	Display	Resolution 3840x2160 [4K UHD] Refresh Rate 60Hz HDR HDR10 support Response speed: 4ms [GTG] Brightness 300cd Contrast ratio 5000:1 I/O Port HDMI x 2ea, USB 2.0 x 2ea, USB type-C x 1ea

HANBACK ELECTRONICS

- 모든 실습기자재 및 각종 디바이스는 정품으로 납품합니다.
- 실습실 내 모든 구성품들은 즉시 사용이 가능한 상태로 설치를 완료합니다.
- 설치 및 시운전에 필요한 모든 경비는 납품 업체에서 지불하며 전기, 네트워크 사용에 문제가 없도록 설치합니다.
- 설치 완료 후 성능테스트를 실시합니다.
- 제품 보증 기한은 검수일로부터 1년으로 하며 기간 내 48시간 내 A/S가 가능합니다.
- 보증 기한 내 기자재 이동에 문제가 없어야 하며 이동 완료 후 장비 가동에 문제가 없어야 합니다.

기타사항