

HUN-REN Cloud új GPU erőforrás menedzsment módszertana az AI4Science programban

Farkas Attila

attila.farkas@sztaki.hun-ren.hu

HUN-REN SZTAKI

INTÉZMÉNYI AI NAGYKÖVETI HÁLÓZAT



Milyen témákkal lehet a kutatóhelyi AI Nagykövethez fordulni?

Minden kutatóhelyen dedikált AI Nagykövet segíti a kutatókat, hogy a mesterséges intelligencia használatát beépítsék a kutatásaikba. A nagykövetek inspirációt adnak az AI használatához, elősegítik az új módszerek megismerését és támogatják, hogy minden kutató megtalálja a számára legnagyobb hozzáadott értéket, amelyet az AI hozhat.

AI OKTATÁS ÉS INSPIRÁCIÓ



Az AI milyen területein mélyíthetik el a kutatók a tudásukat?

Alapozó képzések

A cél az általános alapok elsajátítása, amelynek nyomán a kutatók tovább tudnak lépni személyes témák felé.

Specialista képzések

A hálózaton belül fejlesztett saját képzések, amelyek egy-egy tudományterület egyedi kérdéseit járják körül

Adatgazdász képzés

SZEMÉLYES AI TÁMOGATÁS



Hová lehet fordulni az AI-t igénylő kutatási ötlettel?

HUN-REN Központ AI szakértői csapat

1. kutatási ötletek validálási és kutatástervezési támogatása
2. infrastruktúra-támogatás
3. alkalmazástámogatás

Külső szakértői csapat

Neuron Solutions külső szerződéses partnerként

AI kutatási partnerkereső

AI kompetenciákkal rendelkező, más kutatási területek iránt nyitott kutatók + AI alkalmazási ötletekkel rendelkező kutatók.

ELÉRHETŐ AI TECHNOLÓGIÁK



Milyen biztonságos AI technológiák használhatók érzékeny kutatási területekhez?

Számítási infrastruktúra

1. Biztonságos GenAI futtatása, akár saját környezetben
2. Gépi tanulás elősegítése akár elosztott, több GPU kártya támogatásával

Saját, hazai AI keretrendszerek

1. HUN-REN Cloud
2. Komondor

Miért kellett új projekt és GPU igénylési lehetőségeket bevezetni?

- Az **AI előretörésével meredeken nőtt az igény** a GPU-t használó projektek indítására, így a kiosztható GPU kapacitás elfogyott
 - főleg a hosszabb, így gyakran hektikus/részleges kihasználtságú projektek miatt
- Egyre nagyobb igény mutatkozik még **újabb és kényelmesebb AI szolgáltatásokra**, ld. GenAI4Science
- Köszönhetően (többek között) az AI4Science programnak növekednek az AI kompetenciák, egyre **haladóbb és nagyobb felhasználási esetek** jelentek meg, ami más (kötegetelt, job-alapú) megközelítést igényelnek
- A felhő **következő bővítésének ütemezése elhúzódik: 2026**

A GPU használati módok kibővítése új lehetőségekkel

1. IaaS mód
 - 1 éves igényű projekt
 - **Szolgáltatást nyújtó projekt**
 - **Ernyőprojekt**
 - 3 hónapos igényű projekt
 - Előkészítő projekt
 - Haladó projekt
2. **PaaS mód SLURM job manager alapon**
 - Interaktív üzemmód
 - Batch üzemmód
3. **GenAI4Science platform használata**

Új GPU elosztási rendszer

Az új lehetőségek mellett új GPU elosztási rendszer kialakítása is szükségessé vált. Olyan új rendszer bevezetése volt a cél, amely

- Elegendő erőforrás esetén úgy működik, mint az eddigi rendszer
- Ha viszont nincs elegendő erőforrás, akkor automatikusan átmegy egy a versenyhelyzetet fair módon lekezelő mechanizmusba

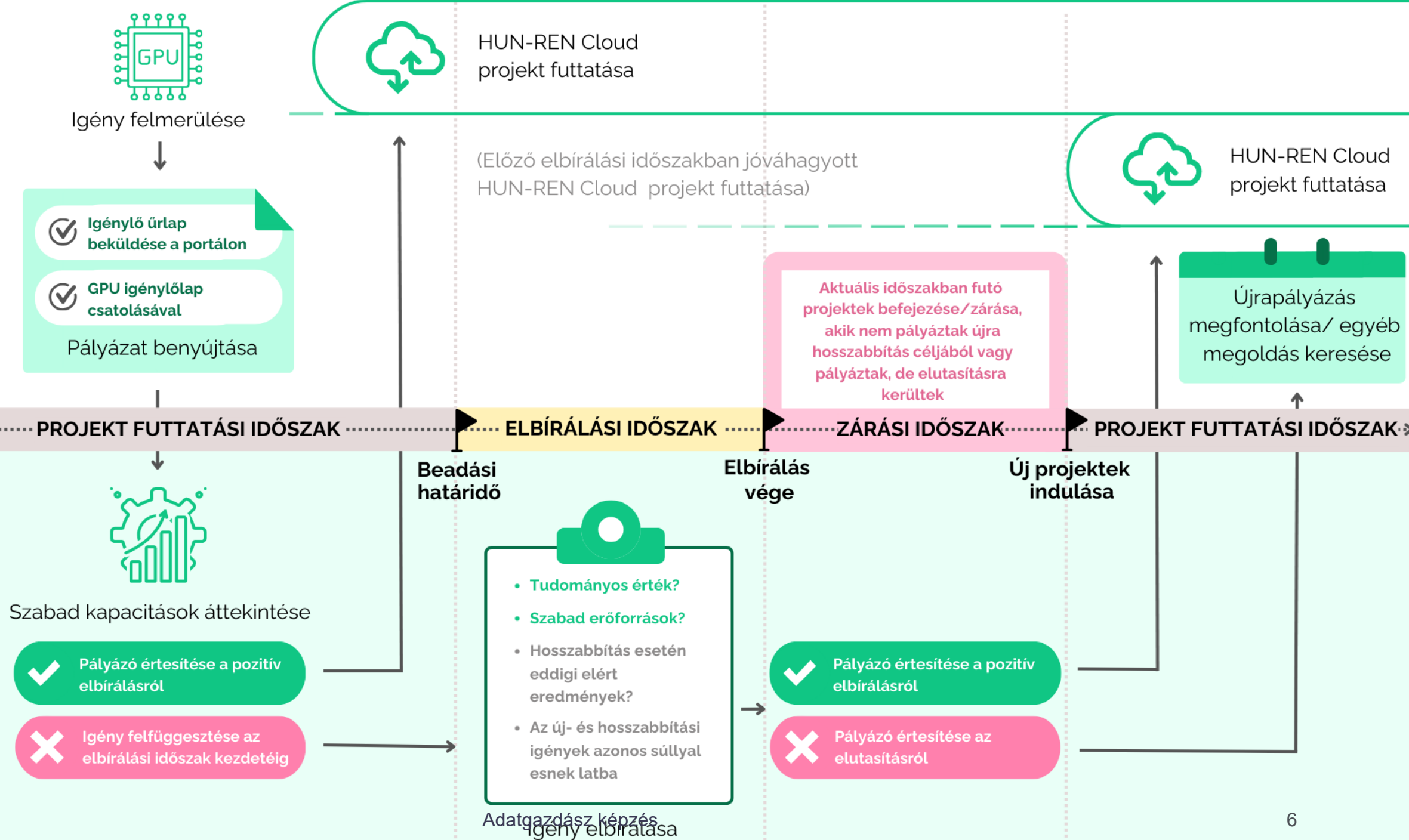
A versenyhelyzet kezelése 2 szempontból is fair:

1. Nem kell adott esetben egy évet várni a projekt indítására, legkésőbb 3 havonta elbírálásra kerül minden projektigény
2. A versengő projektigények egységes szempontok (tudományos teljesítmény, impakt, erőforrásigény, stb.) alapján kerülnek elbírálásra

GPU Erőforrás-gazdálkodási Szabályzat

<https://science-cloud.hu/gpu-eroforras-gazdalkodasi-szabalyzat>





GPU igénylőlap

- A projekt igénylés során megadandó új mezők
- A GPU igénylőlap feltöltésének helye

GPU nélküli erőforrások használatának tervezett vége *

mm / dd / yyyy



Maximum 6 hónap + két hét az igénylés napjától számítva előkészítő és haladó projektek esetében, valamint maximum 1 év ernity és szolgáltató projektek esetében. Ha nincs szüksége a projekt típusához maximálisan igényelhető teljes időtartamra, kérjük, válasszon rövidebbet.

GPU-s erőforrások használatának tervezett vége *

mm / dd / yyyy



Előkészítő és haladó projektek esetén maximum a következő projekt futtatási időszak vége adható meg tervezett zárási dátumnak, ennél hosszabb (maximum 1 éves) időtartam csak szolgáltató és ernity projektek esetében állítható be. Kérjük a GPU igénylőlapra kitöltötteknek megfelelő dátumot adjon meg. Ha nincs szüksége a projekt típusához maximálisan igényelhető teljes időtartamra, kérjük, válasszon rövidebbet.

▼ GPU igénylőlap

Kérjük, töltsen le a [GPU igénylőlap nyomtatványt](#) és kitöltve csatolja az online kitöltött igénylő űrlap beküldésekor. A GPU-erőforrásokra vonatkozó kapcsolódó szabályok részleteit a [GPU Erőforrás-gazdálkodási Szabályzatunkban](#) találják.

Új fájl hozzáadása

Browse...

No files selected.

Ezzel a mezővel korlátlan számú fájl tölthető fel.
4 MB korlát.
Engedélyezett típusok: docx.



Speciális projektek

- Szolgáltatást nyújtó projekt:
 - Nyilatkozat tétele: Tudomásul veszem, hogy a HUN-REN Cloud best effort jellegű szolgáltatást nyújt SLA nélkül és ennek megfelelően az igényelt projekt által nyújtott szolgáltatás is best effort jellegű lesz SLA nélkül. Kijelentem, hogy a HUN-REN Cloud felé semmilyen SLA követelésem nincs és nem is lesz.
- Ernyőprojekt:
 - Meg kell adni, hogy min. hány projekt fog működni az ernyőprojekt keretében és ebből hány haladó projekt van
 - Amely projektekről már lehet tudni, ott megadni a projekt címét, résztvevőit és hogy hány és milyen publikációt, ill. disszeminációs tevékenységet vállalnak.
- Haladó projekt:
 - Meg kell adni, hogy miért nevezheti magát haladónak a projekt:
 - Előkészítő projekt sikeres befejezése
 - A projekt tagoknak legalább egyike már futtatott haladó projektet a HUN-REN Cloudon

Fontos időpontok

- **Pályázási határidők**

- február 28.
- **május 31.**
- aug. 31.
- nov. 30.

- **3 hónapos Projekt időszakok**

- jan. 1. – márc. 31,
- ápr. 1. – jún. 30.
- júl. 1. – szept. 30.
- okt. 1. – dec. 31.

- **Elbírálási időszakok**

- márc. 1-14,
- **jún. 1-14**
- szept. 1-14.
- dec. 1-10.

- **Felkészülési/zárási időszakok**

- márc. 15-31.
- **jún. 15-30.**
- szept. 15-30.
- dec. 11-23.

<https://science-cloud.hu/hirek/2024/felhasznalok-tajekoztatasa-az-uj-gpu-eroforras-gazdalkodassal-kapcsolatban>



Slurm



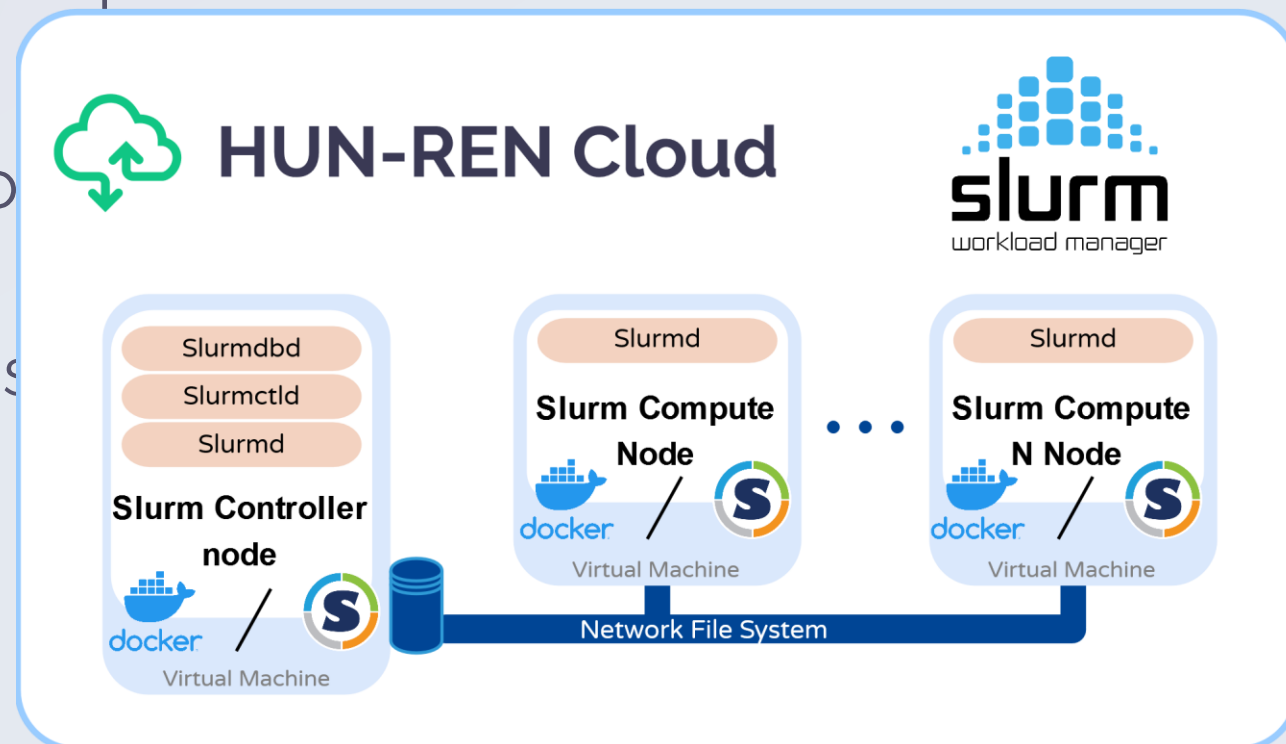
- Nyílt forráskódú erőforrás-kezelő és ütemező rendszer
- Feladatok (Job) / Kötegek (Batch) ütemezése
- HPC rendszerekben régóta használt megoldás (pl.: Komondor)
- Erőforrások (CPU, GPU, memória stb.) dinamikus elosztása
- Feladatok sorba állítása és párhuzamos futtatása
- MP/MPI/MPI4PY futtatási lehetőség



Slurm szolgáltatás



- Slurm referencia architektúrára épül
- Felhasználó kezelés
- NFS fájlmeosztás a csomópontok között
- FIFO logika – Fair erőforrás használat
- Singularity
- Különböző erőforrás partíciók



<https://docs.slurm.science-cloud.hu/>

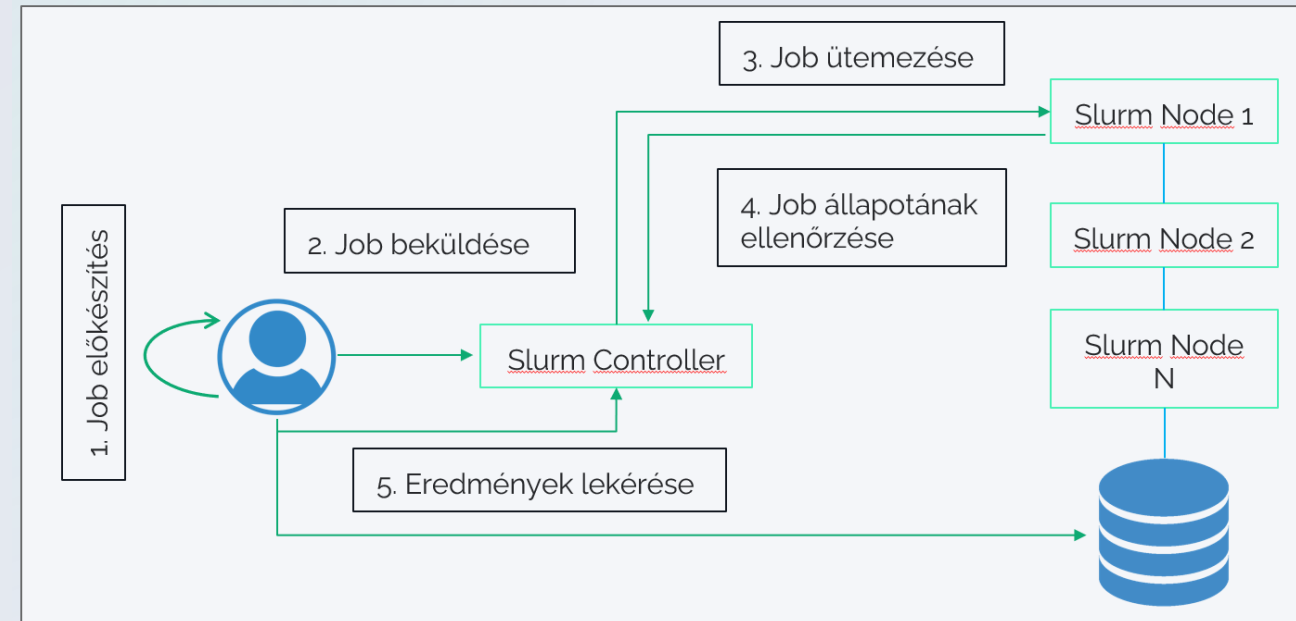
Kvóták és partíciók

- Felhasználó szintű limitációk
 - 1 futtatható job
 - 3 job kerülhet ütemezésre
 - 100 GB tárhely
- GPU specifikus partíciók
 - GPU vRAM által meghatározott partíciók
 - Interaktív
 - 7 nap után a job leállításra kerül
 - Preferált partíció a fejlesztői környezetek kialakítására
 - Batch
 - 1 nap után a job leállításra kerül
 - Gyorsan lefutó joboknak létrehozott partíció



Slurm jobok futtatásának menete

1. A job előkészítése (felh.)
2. A job beküldés (felh.)
3. A job ütemezése a klaszteren (slurm)
4. A job állapotának nyomon követése (slurm)
5. A sikeres futást követően a felhasználó eléri a számítás eredményét (felh.)



Milyen feladatokra javasoljuk?

- Rövid ideig szükséges erőforrás
- Nagy számítási igényű feladatok
- Egymásra épülő feladatok
- „Parameter sweep” típusú feladatok
- Párhuzamos feldolgozást igénylő feladatok
- Hosszabb felügyelet nélküli számításokat



Regisztráció menete

- Új projekt igénylés során
 - Slurm erőforrások igénylése (új checkbox)
- Meglévő projekt igénylés esetén
 - Projekt módosítás -> Slurm erőforrások igénylése
- Projekt szintű hozzáférés
 - A projekt felhasználói saját profiljukon megadhatnak egy SSH kulcsot
- Igényelhető intervallum 3 hónap
 - Hosszabbítás módosítással igényelhető

2 Igényelt erőforrás

☒ Igényelek SLURM erőforrásokat

A SLURM egy skálázható és moduláris erőforrás-kezelő rendszer. Feladatütemezést és monitorozást biztosít a számítási feladatokat végző csomópontok között, valamint képes meghatározott időszakokra dinamikusan erőforrásokat allokálni a felhasználók között. Lehetőséget kínál nagy teljesítményű számítási igényű feladatok futtatásához a GPU-ütemezés támogatásával. [Olvassa el részletes SLURM dokumentációnkat.](#)

SLURM erőforrások használatának tervezett vége *

dd/mm/yyyy



Maximum 3 hónap az igénylés napjától számítva. Ha nincs szüksége a maximálisan igényelhető teljes időtartamra, kérjük, válasszon rövidebbet.

Projekt azonosító *

A projektazonosító lesz az SSH felhasználóneve. Maximum 24 karakter hosszú lehet és csak angol kisbetűket, számokat, valamint alulvonást tartalmazhat.



Regisztráció menete

- Jóváhagyást követően hozzáférés biztosítása a rendszerhez
- SLURM erőforrás-allokációval rendelkező futó projektek adatlapján a tagok számára megjelenik egy új szekció:
 - Aktuális terheltség
 - SSH kulcs beállításához segítség
 - Hozzáférés
- Lejárat előtt a **felhasználó felelősége** a számítási adatok kimentése

<https://science-cloud.hu/eloadasok>

GenAI4Science szolgáltatás

- Nagy nyelvi modell (LLM) alapú csevegő szolgáltatást nyújt (mint a ChatGPT)
 - Teljesen a HUN-REN Cloud erőforrásokon biztosítva
- Nyílt forráskódú szoftverkomponensekre épül
 - Nem függ harmadik fél szolgáltatásaitól
 - De képes kommunikálni OpenAI API-n keresztül más szolgáltatásokkal
- Adatvédelem biztosított
 - Az adatok nem hagyják el a HUN-REN Cloudot
- Egyszerre több LLM modellt is támogat

<https://doc.genai.science-cloud.hu/>



GenAI4Science szolgáltatás felépítése

■ Ollama

- Támogat számos LLM modellt (llama, deepseek, gemma, mistral, Qwen, stb.)
- Modell katalógust biztosít
- Támogatja a testreszabott modelleket
- REST API-t biztosít a modellek eléréséhez
- GPU-alapú modell futtatás (korlátozott a HUN-REN Cloudon)
 - Terv: erőforrás bővítés -> elosztott futtatás több GPU felhasználásával vLLM alapon



■ Open-WebUI

- Felhasználóbarát
- Felhasználói jogosultság kezelés
- Többfelhasználós támogatás
- OpenAI API alapú elérés támogatása (integrációs lehetőségek)
- RAG támogatása
- Helyi, távoli, webes keresés (Google, DuckDuckGo stb.)



DEMO – Bejelentkezés, modell választás

Sign in to GenAI4Science (Open WebUI)

Email
Enter Your Email

Password
Enter Your Password

Sign in

Search a model

All Code Reasoning Vision

- codellama:13b 13B CODE
- codestral:22b 22.2B CODE
- deepseek-r1:70b 70.6B REASONING
- deepseek-r1:32b 32.8B REASONING
- gemma2:27b 27.2B
- gemma3:27b 27.4B VISION

Temporary Chat

New Chat

Workspace

Search

Chats

llama3.3:70b
Set as default

llama3.3:70b
New state-of-the-art 70B model from Meta that offers similar performance compared to Llama 3.1 405B model....

How can I help you today?

+ Web Search Code Interpreter

Suggested

- Show me a code snippet of a website's sticky header
- Explain options trading if I'm familiar with buying and selling stocks
- Give me ideas for what to do with my kids' art

Attila Farkas

DEMO – chat felület, RAG

Generate a code to finetune a Bert model



llama3:8b

Finetuning a BERT model involves modifying the pre-trained model's weights to better fit your specific task or dataset. Here's an example code in PyTorch to finetune a BERT model:

```
python
Run Copy Code

import torch
from transformers import BertTokenizer, BertForSequenceClassification
from sklearn.metrics import accuracy_score

# Set the device (GPU or CPU)
device = torch.device("cuda" if torch.cuda.is_available() else "cpu")

# Load the pre-trained BERT model and tokenizer
tokenizer = BertTokenizer.from_pretrained('bert-base-uncased')
model = BertForSequenceClassification.from_pretrained('bert-base-uncased', num_labels=8)

# Set the hyperparameters for finetuning
batch_size = 16
epochs = 3
learning_rate = 1e-5

# Load your dataset (e.g. IMDB, Amazon reviews, etc.)
train_data = ...
test_data = ...
```

+ Send a Message

LLMs can make mistakes. Verify important information.

According to the provided information, the HUN-REN Cloud has:

- Number of vCPU: 7,344
- Memory (RAM): 35 TB
- HDD storage: 1,248 TB
- SSD storage: 338 TB
- Internal bandwidth: 100 Gbit/s
- Number of GPU cards: 72
- GPU memory (RAM): 2,560 GB
- GPU double precision TFLOPS: 622
- GPU single precision TFLOPS: 1,252
- GPU FP16 tensor TFLOPS: 14,984

1 <https://science-cloud.hu/en/hun-ren-cloud> 2 <https://science-cloud.hu/en/faq>

What type of GPU cards has the HUN-REN Cloud?



llama3:8b

Searched 3 sites

According to the provided information, the HUN-REN Cloud has NVIDIA Tesla V100 GPUs with a maximum of 32GB of GPU memory. These GPUs can be divided into smaller vGPUs, so it is possible to request them with 8 or 16 GB of memory. Additionally, the Wigner data center also has NVIDIA Tesla A100 GPUs with a maximum of 40GB of GPU memory, which can be divided into vGPUS with 5, 10 and 20 GB of memory.

1 <https://science-cloud.hu/en/faq> 2 <https://science-cloud.hu/en/hun-ren-cloud> 3 <https://www.cgldirector.com/types-of-gpus/>

Web Search



Upload Files

+ Send a Message

LLMs can make mistakes. Verify important information.



HUN-REN Cloud

2025.05.14.

Adatgazdász képzés

Hozzáférés igénylés

Hozzáférés igénylés

Jelenleg korlátozott erőforrás áll rendelkezésre a szolgáltatás mögött, ezért csak korlátozott számban tudunk hozzáférést biztosítani.

Az erőforrás korlátok miatt jelenleg csak a HUN-REN Kutatóhálózat kutatói számára tudunk hozzáférést biztosítani, intézmény méret arányos módon.

Belépés a [GenAI4Science](#) portálon keresztül lehetséges a jóváhagyott igénylést követően.

Hozzáférés igénylése

GenAI4Science dokumentáció



GenAI4Science szolgáltatás igénylése

A GenAI4Science szolgáltatáshoz az alábbi adatok megadásával igényelhető hozzáférés. [További információ a GenAI4Science-ről.](#)

Name *

Farkas Attila

E-mail cím *

farkas.attila@sztaki.hu

Intézmény *

HUN-REN intézmények ▾

Számítástechnikai és Automatizálási Kutatóintézet ▾

Megjegyzés

☐ Az [Adatkezelési Nyilatkozatot](#) elolvastam, és elfogadom

Igény leadása

Használati esetek

- AI4Science program támogatása
- Mesterséges Intelligencia Nemzeti Labor támogatása



MESTERSÉGES INTELLIGENCIA
Nemzeti Laboratórium



HUN-REN Cloud

2025.05.14.

Adatgazdász képzés

Segítség a GPU igénylőlapok kitöltésében, új szolgáltatások használatában

 HUN-REN Cloud

Rólunk ▾

Aktualitások ▾

Projektek

Dokumentáció ▾

GenAI4Science

← Címlap | Támogatás kérése

Támogatás kérése

Vezetéknév *

Keresztnév *

Összefoglalás

- **Új GPU allokálási rendszer** került bevezetésre az erőforrások dinamikusabb kiosztása érdekében
 - Különböző projekt kategóriák pályázhatók periodikusan
- Új **PaaS** szolgáltatások a IaaS hozzáférés mellett
 - **SLURM szolgáltatás** – kötegetelt vagy interaktív feladat végrehajtás ütemező segítségével
 - **GenAI4Science** – nagy nyelvi modellek elérése chat üzemmódban vagy API hozzáférés



Köszönöm a figyelmet!

Farkas Attila

attila.farkas@sztaki.hun-ren.hu

HUN-REN SZTAKI

