

圧倒的なコストパフォーマンスを実現したGPUクラウドサービス



# GPUSOROBAN

こんな課題はありませんか？

- ✓ GPUが手に入らない・価格が高い
- ✓ データ転送料や隠れコストが重い
- ✓ ドル建て決済でコストが急増(為替変動リスク)
- ✓ 国内法規制に対応したデータセンターが必要

すべてGPUSOROBANが解決します

## 豊富なNVIDIA® GPUラインナップ

A4000

GPUメモリ

**16GB**

A100

GPUメモリ

**40GB / 80GB**

H200

GPUメモリ

**141GB**

B200

GPUメモリ

**180GB**

## GPUSOROBANが選ばれる3つの理由



### 業界最安級&円建て固定料金

高価な外資系クラウドとは異なり、為替の変動に左右されない「日本円」での料金体系。データ転送料は無料。予算超過の心配なくご利用いただけます。



### 安心の国内データセンター

GPUSOROBANは国内の自社データセンターで運用しています。日本の法律に準拠したデータ管理で、機密情報の取り扱いも安心です。



### 日本語対応の技術サポート

「環境構築が不安」という方も安心。導入後の技術サポートは無料ですべて日本語対応しており、マニュアルも日本語となります。

## 業界最安級で提供できる理由

### 建設コストを抑制

研究所(香川県)や未活用廃校(佐賀県)を活用することで建設コストを抑えています。



### 設備コスト・SLAの工夫

AIの学習や研究所のシミュレーションに特化することで設備コストを抑えています。SLA保証を行わないことで低価格を実現しています。



各プランの料金・お問い合わせはこちら



Email: [ln-gpu-sales@highreso.co.jp](mailto:ln-gpu-sales@highreso.co.jp)  
TEL: 03-6775-7888

株式会社ハイレゾ  
〒162-0843  
東京都新宿区市谷田町3丁目24-1



# LLMデータ開発の 成功方程式

AIの成否を分けるのは、  
モデルではなく「データ」だった。

AI開発成果の80%は、データの「量と質」で決まります。  
「集まらない」「品質がばらつく」「専門性が足りない」  
—APTO様はData-Centric(データ中心)なアプローチで解決。  
導き出した方程式が【高品質データ × 適切なチューニング  
× GPU環境】です。

## LLM 開発を成功に導く方法



**高品質データ**  
AIの土台となる  
データの量と質



**適切なチューニング**  
ドメインに合わせた  
最適な「教え方」



**GPU環境**  
試行錯誤の回数を  
最大化する環境

## — 高品質データ

❗ 「良いモデルに、大量のデータを与えれば  
精度は上がる」は間違いです。

### ■ スタイル・ミスマッチの罠

大型モデルの立派な回答を小型モデルに無理に学ばせた結果、  
精度が -4.6pt 低下する失敗が起きました。

### ■ 約20%が削除すべき問題データという現実

データセット全体を精査した結果、5件に1件(約20%)が  
問題データであることが判明しました。

| 問題の種類  | 件数      | 内容                      |
|--------|---------|-------------------------|
| 文字化け   | 1,730 件 | エンコーディング不整合などによる読取不能データ |
| データの重複 | 1,434 件 | 学習時にバイアスを生み出す重複データ      |
| データの矛盾 | 46 件    | 論理的に整合性が取れていないデータ       |

APTO様ではこうした問題データを徹底排除し、モデルに最適化されたデータのみを抽出しています。

## — 適切なチューニング

データだけでなく「考えるプロセス」まで  
学習させることで、OSSモデルでもトップクラスの  
商用モデルに迫る性能を引き出せます。

### 数学理論：答えより解き方

「正解」ではなく「途中の思考プロセス」を教える  
データ設計で、AIME 正答率 66.7% を達成  
(GPT-4o 参考値 12.0% を大幅に超越)。



### 医療 AI：専門性の獲得

医師の思考回路を模したチューニングにより、  
医療応答品質のスコアを倍増させました。



チューニング前

16%

チューニング後

32%

2倍

## — GPU環境 GPUクラウド「GPUSOROBAN」利用

試行錯誤と圧倒的に回せる環境こそが、  
LLM開発の成功確率を最大化します。

LLM 開発では、GPU 環境が検証スピードと成果に  
直結します。同じ3週間でも、選ぶ環境によって  
試せるアプローチの数が大きく変わります。

A100 × 4

3 週間の  
1 回の学習時間 試行錯誤回数  
約 1.5 日 ▶ 2~3 回

H200 × 8

3 週間の  
1 回の学習時間 試行錯誤回数  
約 12 時間 ▶ 5 回

### △ 失敗を翌日にリカバリ

H200×8などの最新環境なら3週間で5回のイテレーション  
(試行錯誤)が可能で、失敗しても翌日に方針転換し、次の手を  
打てます。

### 🔄 速い GPU の真の価値

速い GPU の真価は、単に学習が早く終わることではありません。  
圧倒的なスピードで実験サイクルを回し、仮説検証を積み重ねながら  
「最適解にいち早くたどり着くこと」にこそあります。

