

Copilot のよる生成 Ai

(ローカルでの使用や画像生成についての比較検討)

Core i7-12650H、32GB 1+1TB SSD NVIDIA GeForce RTX 4060 Laptop GPU(8GB)
140W の Acer Nitro 5 を既に持っています。これらの三機種で比較をしてください。

了解しました、あなた。 比較対象は **3 機種**ですね：

1. MacBook Pro M5 (32GB)
2. Windows+大容量 VRAM GPU (例：RTX 4090/5090 24～32GB)
3. あなた所有の Acer Nitro 5 (i7-12650H / 32GB / RTX 4060 Laptop 8GB 140W)

ここでは ローカル生成 AI (LLM・画像生成・LoRA 学習) の観点で、「何ができて、どこが限界で、どれが速いか」を構造的に比較します。

結論 (最初に要点だけ)

- M5 (32GB) → 30B～70B 級の巨大 LLM が動く。速度は中堅。静音・省電力。
- Windows+大型 GPU (4090/5090) → 画像生成・7B～14B LLM は圧倒的に高速。ただし VRAM 上限で巨大モデルは不可。
- Nitro 5 (4060 Laptop 8GB) → 7B～14B はそこそこ速いが、30B 以上は VRAM 不足で厳しい。→ ただし CPU が強く、ハイブリッド推論で意外と健闘。

🧠 1. LLM（ローカル推論）性能比較

■ モデルサイズの上限（最重要）

機種	実行可能なモデルサイズ	理由
M5 32GB	30B～70B 級が動く （量子化前提）	Unified Memory で VRAM 制限がない
Windows + 4090/5090 （24～32GB VRAM）	最大 34B 前後	VRAM が壁。70B は基本不可
Nitro 5 （4060 Laptop 8GB）	最大 13B～30B （ハイブリッド）	VRAM 8GB ではフル GPU 不可。CPU オフロード必須

👉 巨大モデル（30B～70B）を動かしたいなら M5 が圧勝。

■ 推論速度（tok/s）

● 7B～14B モデル

- 4090/5090 : 150～300 tok/s
- Nitro 5 (4060) : 40～70 tok/s
- M5 32GB : 30～60 tok/s

👉 この領域は Windows + 大型 GPU が圧倒的に速い。

● 30B モデル

- M5 32GB : 25～40 tok/s
- Nitro 5 (4060) : 10～15 tok/s（ハイブリッド）
- 4090/5090 : VRAM 不足で遅い（5～15 tok/s）

👉 30B は M5 が最も快適。

● 70B モデル

- M5 32GB : 量子化で動く（10～20 tok/s）
- 4090/5090 : 基本不可
- Nitro 5 : 不可

👉 70B をローカルで動かせるのは M5 のみ。

🎨 2. 画像生成（Stable Diffusion / SDXL）

機種	SD1.5	SDXL	高解像度
M5 32GB	○	△（遅い）	△
4090/5090	◎	◎（最速）	◎
Nitro 5 （4060）	○	○～△	△

👉 画像生成は Windows+大型 GPU が圧倒的。M5 は SDXL が苦手。
Nitro 5 は中堅。

⚙️ 3. LoRA 学習・ファインチューニング

機種	LoRA 学習	大規模学習
M5 32GB	△（遅い）	×
4090/5090	◎（最速）	○（VRAM 次第）
Nitro 5 （4060）	△	×

👉 学習用途は Windows+大型 GPU が圧倒的。

🍷 4. 運用性・エコシステム

項目	M5	Windows+大型 GPU	Nitro 5
静音・省電力	◎	△	△
環境構築の楽しさ ◎（Ollama 最強）		△	△
CUDA エコシステム	×	◎	○
持ち運び	○	×	○

💰 5. コスト

機種	コスト感
M5 32GB	本体のみで完結
4090/5090	PC 本体+GPU で 30~60 万円
Nitro 5	既に所有 (追加コストなし)

📊 総合比較表 (あなたの用途向け)

用途	M5 32GB	Windows + 大型 GPU	Nitro 5
巨大 LLM (30B~70B)	◎ (最適)	×	△ (30B はギリ)
高速 LLM (7B~14B)	○	◎ (最速)	○
SDXL 画像生成	△	◎	○
LoRA 学習	△	◎	△
静音・省電力	◎	△	△
環境構築の楽さ	◎	△	△

🎯 最終結論 (あなた向けに最適化)

- 1. 巨大モデル (30B~70B) をローカルで動かしたいなら → M5 (32GB)
 - VRAM 制限がないため Windows では載らないモデルが動く
 - 速度は中堅だが、モデルサイズの自由度が圧倒的
- 2. 画像生成・高速推論・LoRA 学習をしたいなら → Windows + 大型 GPU (4090/5090)
 - SDXL が圧倒的に速い
 - 7B~14B の LLM は M5 の 2~3 倍以上の速度
- 3. Nitro 5 (4060 Laptop) はどう位置づける？
 - 7B~14B は十分実用的
 - 30B はハイブリッドでギリ動く
 - すでに所有しているので、→ 軽いローカル LLM 用途ならまだ戦える