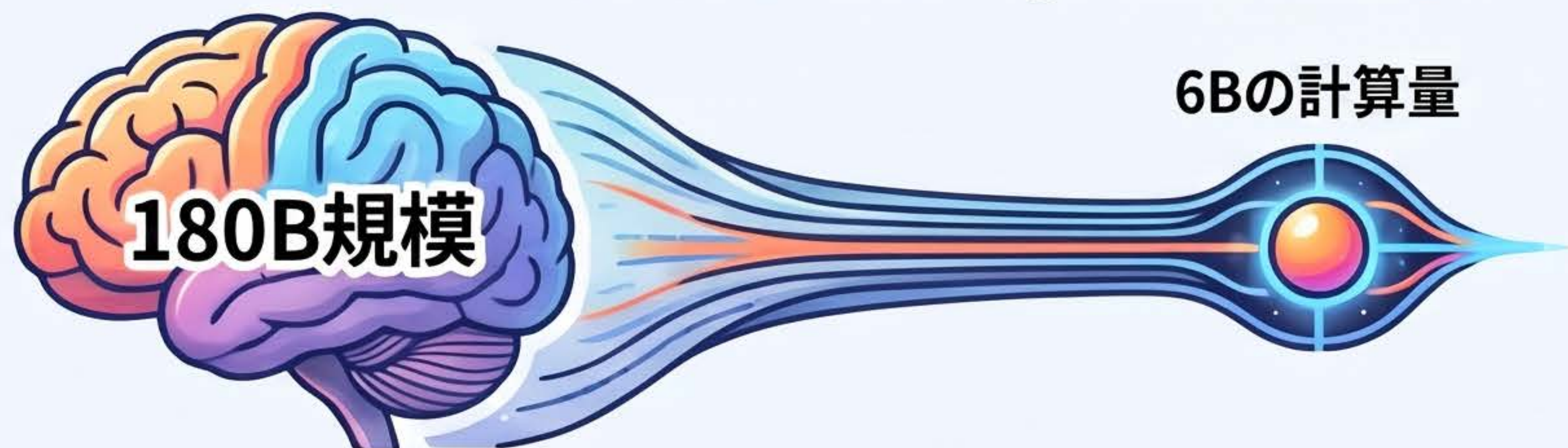


Qwen3.8-Flash-Next：次世代AIアーキテクチャの全貌と実務ガイド

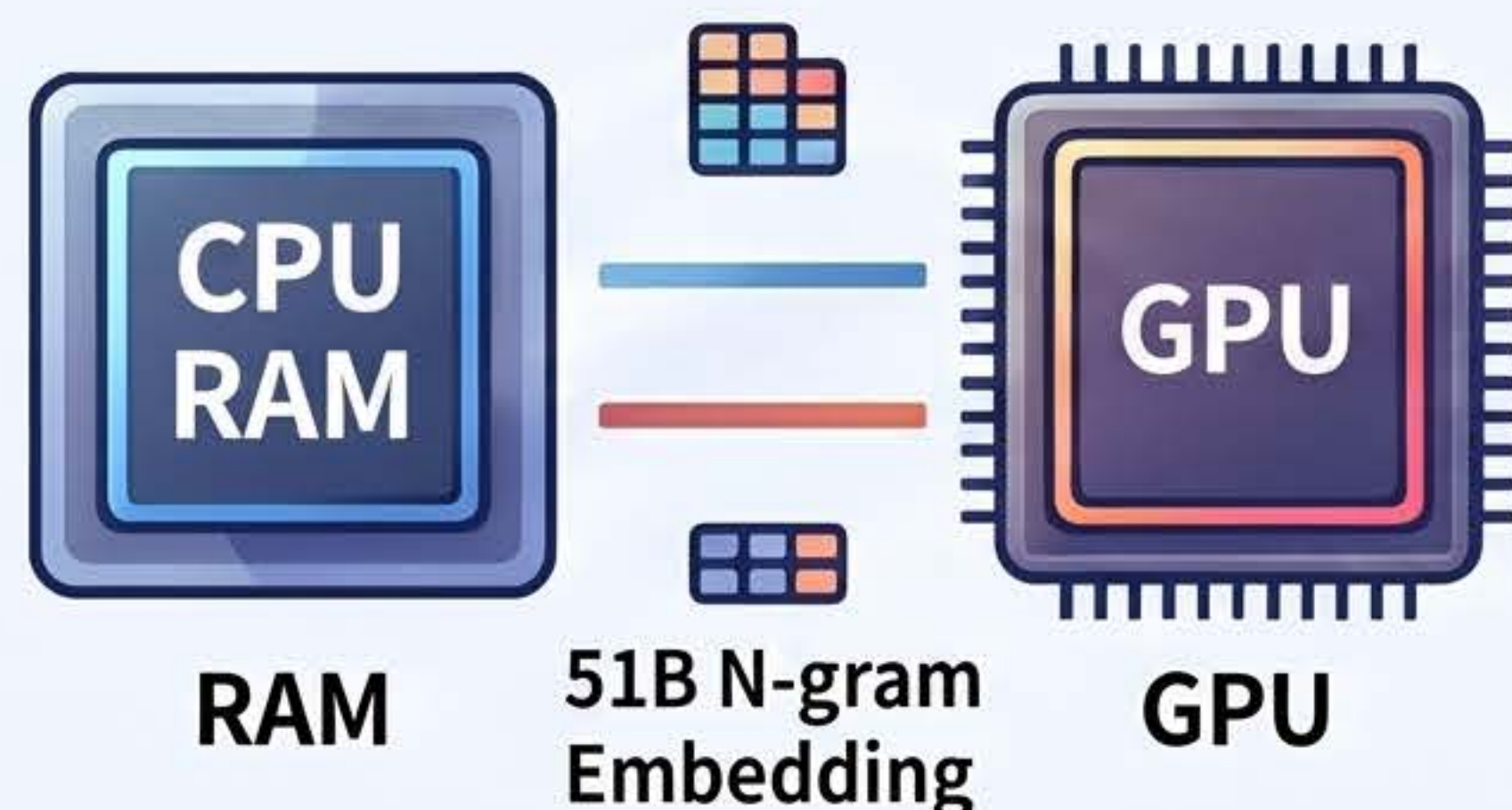
技術的革新：Qwen4への道



125BのMoEモデルのうち1トークン当たり約6Bのみを活性化し、推論を高速化。

長文処理を加速する GDN+QSA

micro-block単位で重要領域を選択し、
1Mトークンの長文処理を最大8.6倍高速化。

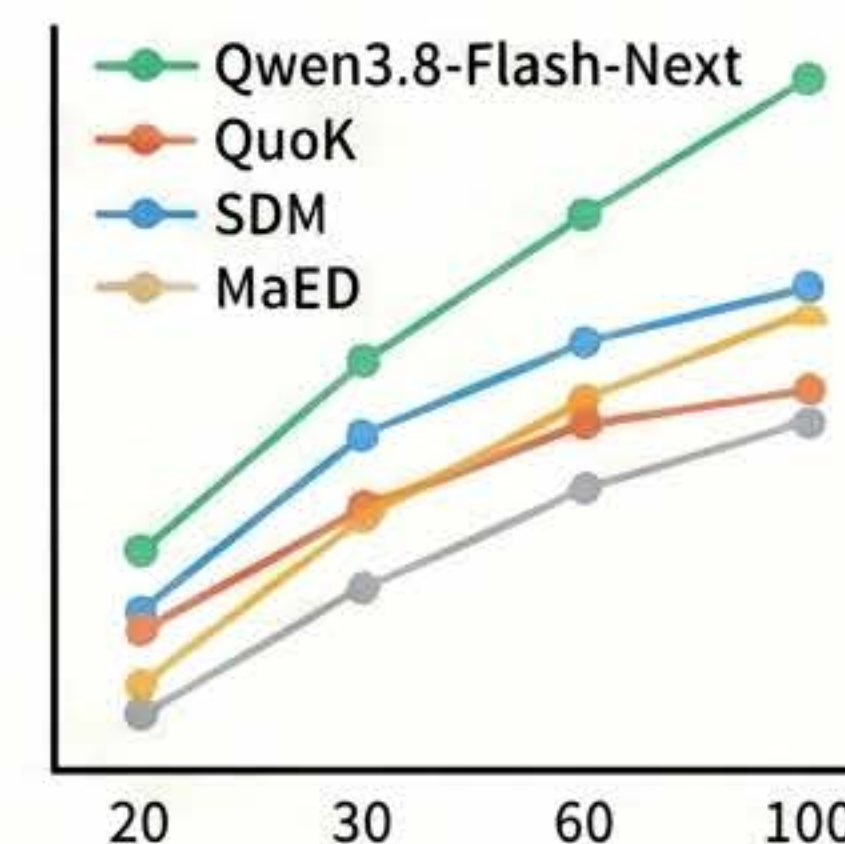


CPU RAMを活用する N-gram Embedding

51BのパラメータをホストRAMへ
オフロードし、GPU負荷を抑えつつ
知識容量を拡大。



実務導入の検討事項とリスク



コーディング・推論で 高い性能を発揮

SWE-bench等で既存モデルを
凌駕するが、一部のベンチマー
クはベンダ報告値に留まる。



Community License 1.0の制約

Apache 2.0ではなく、MaaSや
AIアシスタント事業には別途
ライセンスが必要。



運用リソースの現実

計算量は少ないが、重み一式の
保持に約180GB~360GBの
メモリ空間を必要とする。