

Microsoft Azure上でのタンパク質間相互作用 予測システムの並列計算と性能評価

大上 雅史¹

山本 悠生^{1,2}

秋山 泰^{1,2}

1 東京工業大学 情報理工学院 情報工学系

2 東京工業大学 情報生命博士教育院



東京工業大学

パブリッククラウドの普及

・パブリッククラウド



Google Cloud Platform



－ 利点

- ・ 利用料金を払えば誰でも利用可能、運用はアウトソース
- ・ 環境や資源を共有するしくみが充実
- ・ 計算資源の増減が簡単

パブリッククラウド上での並列計算の例



Mrozek D, *et al.* Cloud4Psi: cloud computing for 3D protein structure similarity searching. *Bioinformatics* 2014, **30**: 2822-2825.

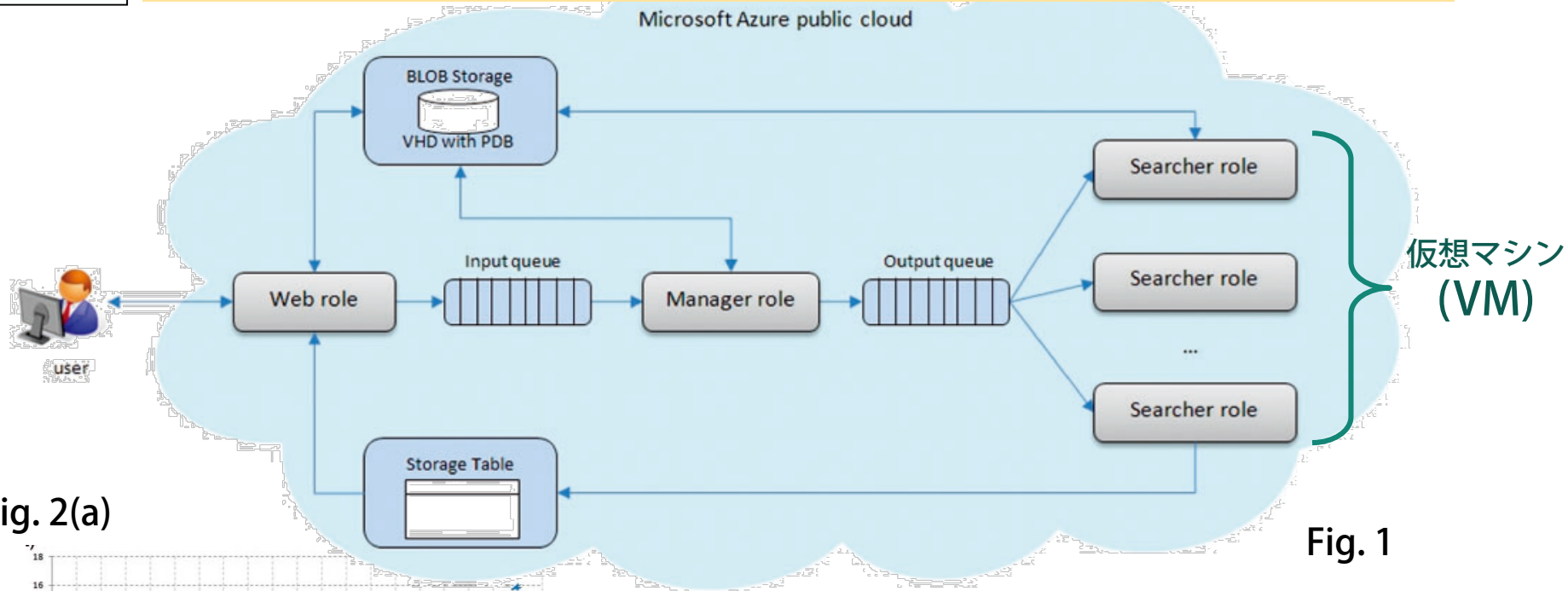


Fig. 2(a)

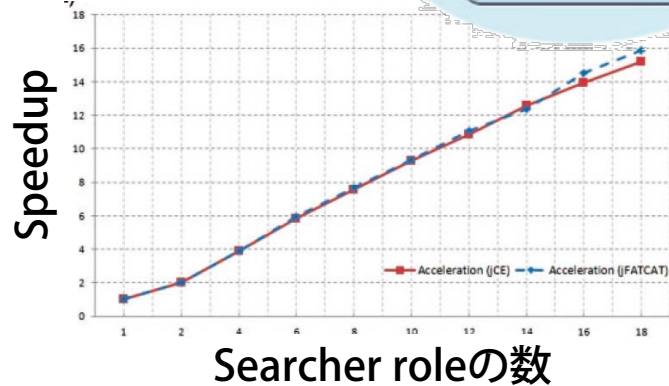
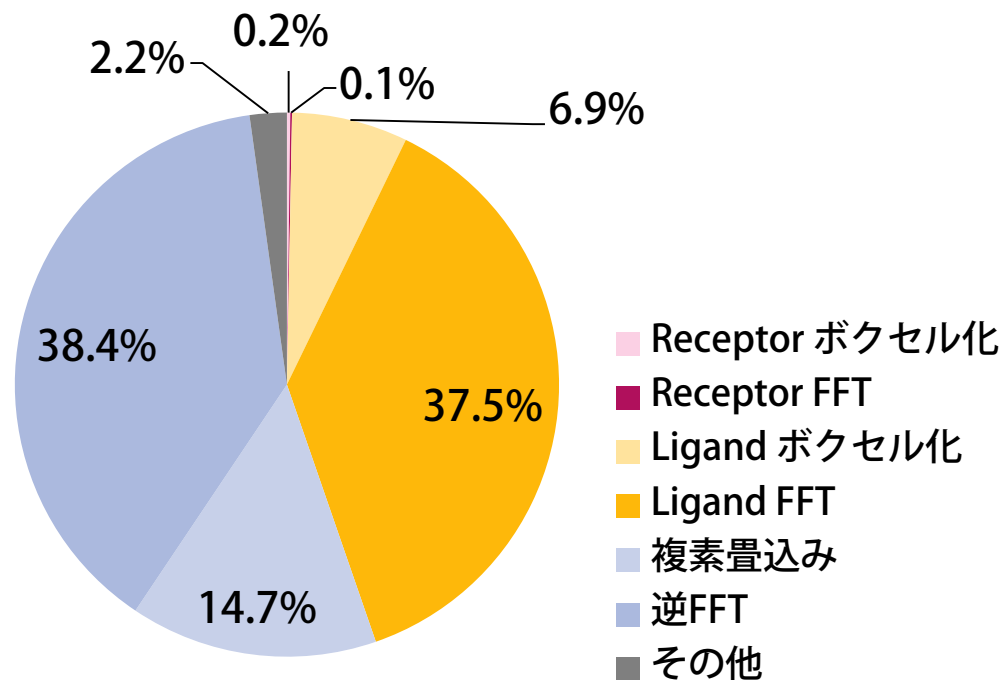
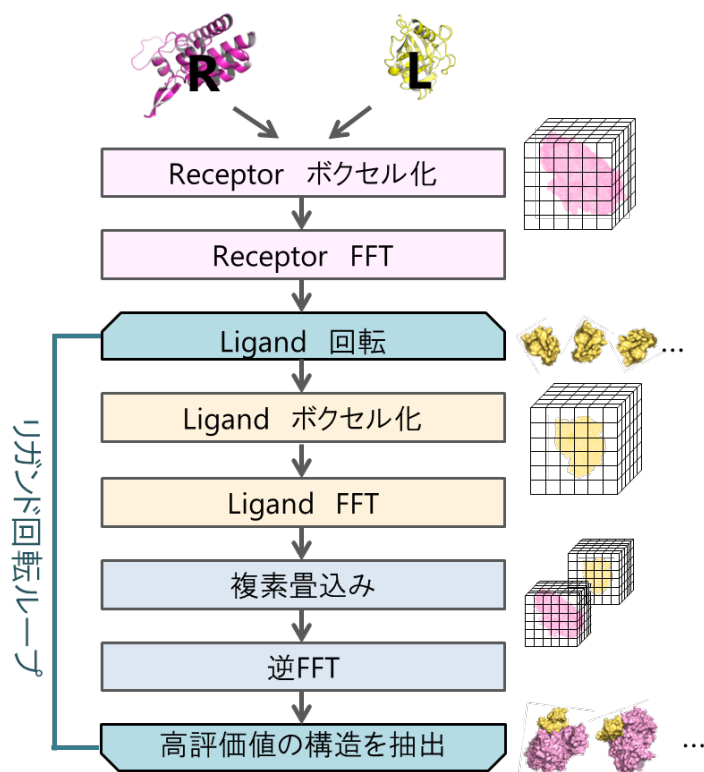


Fig. 1

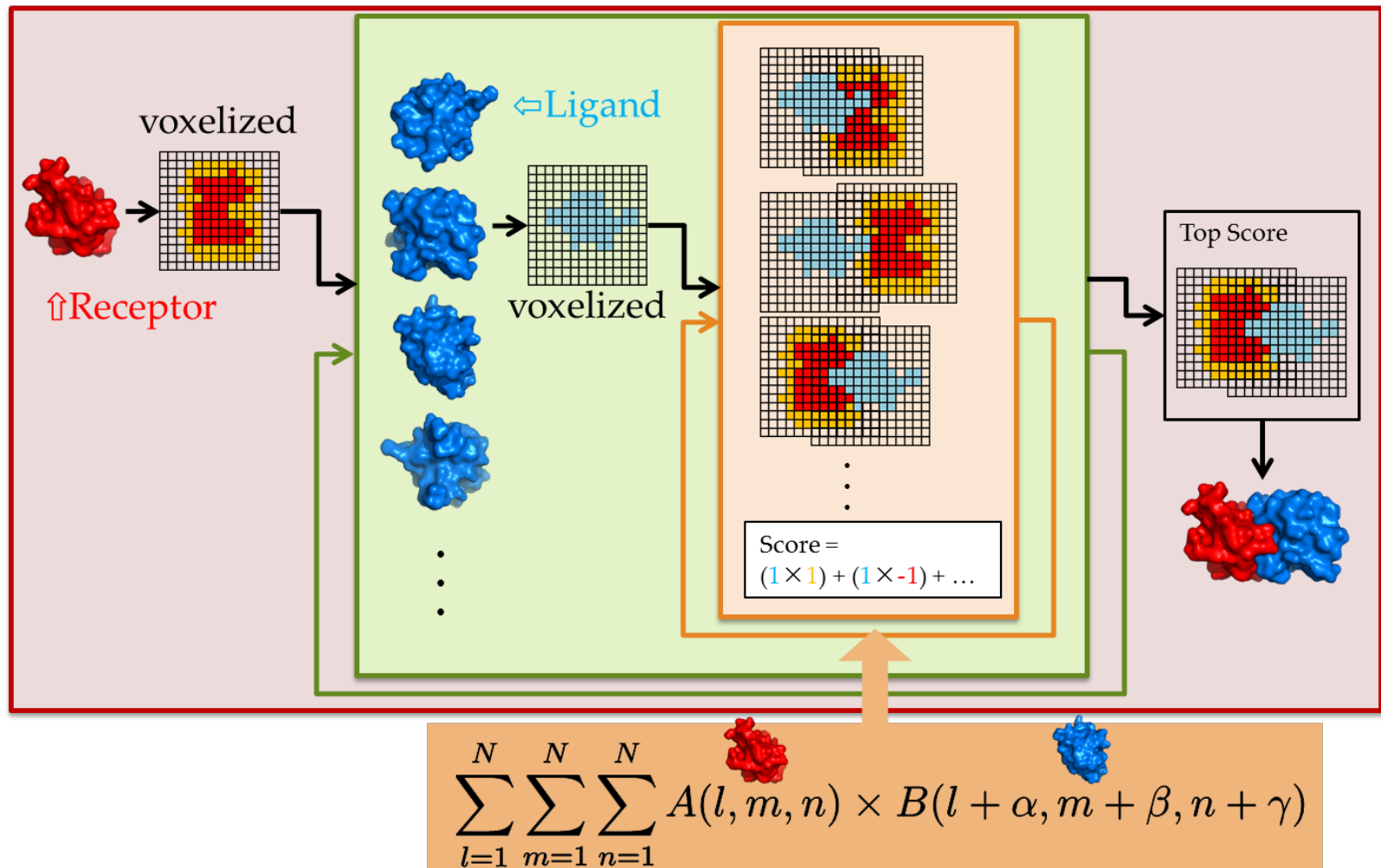
- タンパク質の類似構造検索タスク。
- Searcher role 6台まで、ほぼリニアに高速化。
- Searcher role 8台超で性能低下するも、18台でも約89%の効率。

• MEGADOCK Ohue M, et al. *Bioinformatics*, **30** (2014), Ohue M, et al. *Prot Pept Lett*, **21** (2014)

- 入力された2つのタンパク質立体構造に対し、相互作用の有無を予測
- 剛体モデルを組合せて評価値の良い複合体構造をサンプリングし、それらの評価値に基づいて予測
- 計算の大部分(約80%)はFFT (高速フーリエ変換)



タンパク質が結合したときの評価値を高速に評価



Katzhalski-Katzir E, et al. *Proc Natl Acad Sci* (1992), **89**(6): 2195–2199.

本研究の目的

- **Microsoft Azure上でMEGADOCKの並列計算を実施**
 - Azure上の複数のVMを用い、MPIによって計算を並列実行
 - 数十VM規模での性能を計測
 - 特に、GPU付きVMに注目

Azureの代表的なインスタンス(VMの種類)

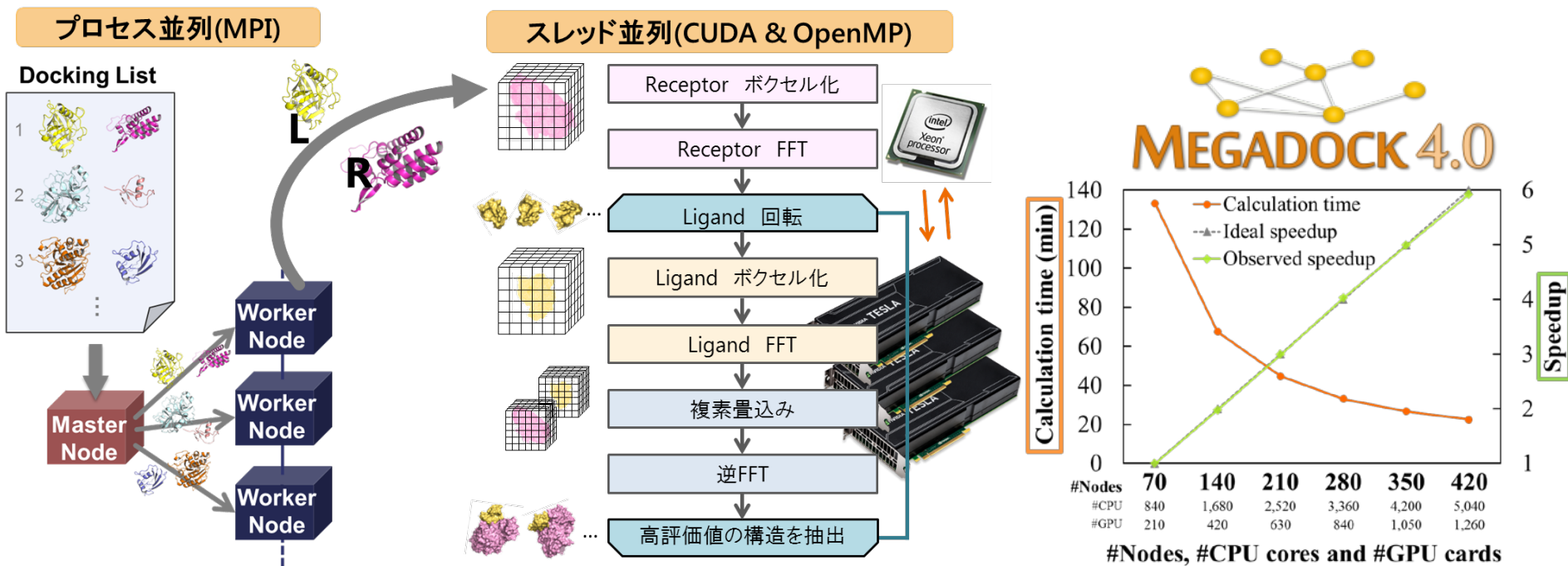
Instance	CPU	GPU(*)	RAM	Network	価格
NC24	24コア	Tesla K80 x 4	1,440 GB		440.65円/h
NC24r	24コア	Tesla K80 x 4	1,440 GB	RDMA (InfiniBand)	484.71円/h
DS14	16コア	なし	112 GB		141.48円/h
A9	16コア	なし	112 GB	RDMA (InfiniBand)	197.06円/h
H16	16コア	なし	112 GB		178.50円/h
H16r	16コア	なし	112 GB	RDMA (InfiniBand)	195.84円/h

(*) ただし、物理的にはK80が2枚。
(K80は1枚の中に2つのGPUチップを含む)

マルチGPUノード並列

• MEGADOCKは既にGPUクラスタ向け並列化が実装済

- プロセス単位でMasterとWorkerに資源を割り当て
- Masterがタンパク質ペアをWorkerに渡し、Workerが計算
- GPUはWorkerが使用

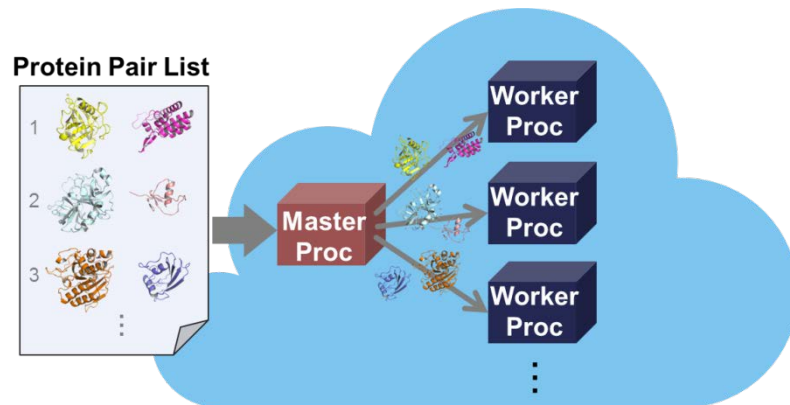


Ohue M, Shimoda T, Suzuki S, Ishida T, Akiyama Y. *Bioinformatics*, 30(22): 2014.

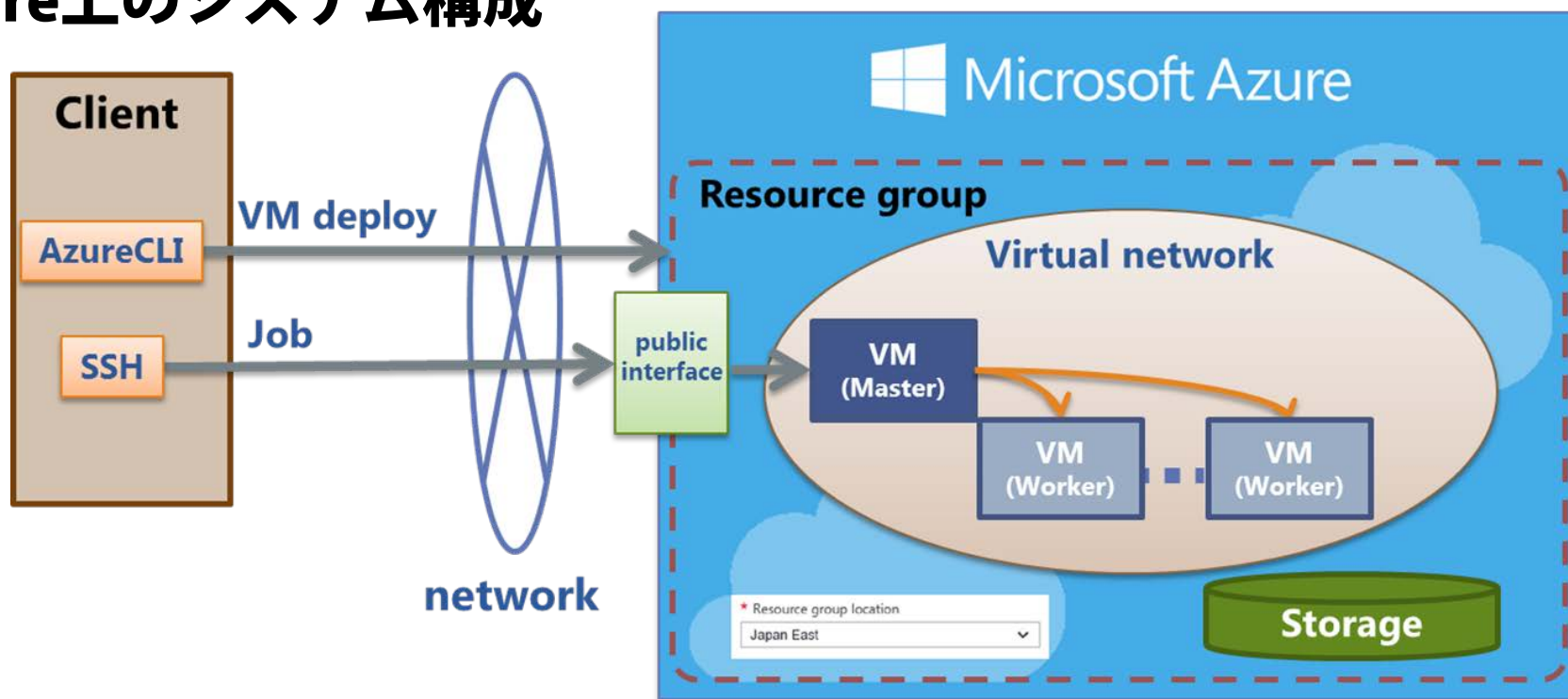
Azureへの実装

概要

- Masterプロセス(1つ)がタンパク質ペアをWorkerに割り当て。
- 各VMのWorkerプロセスが割り当てられたペアの予測計算を実施。



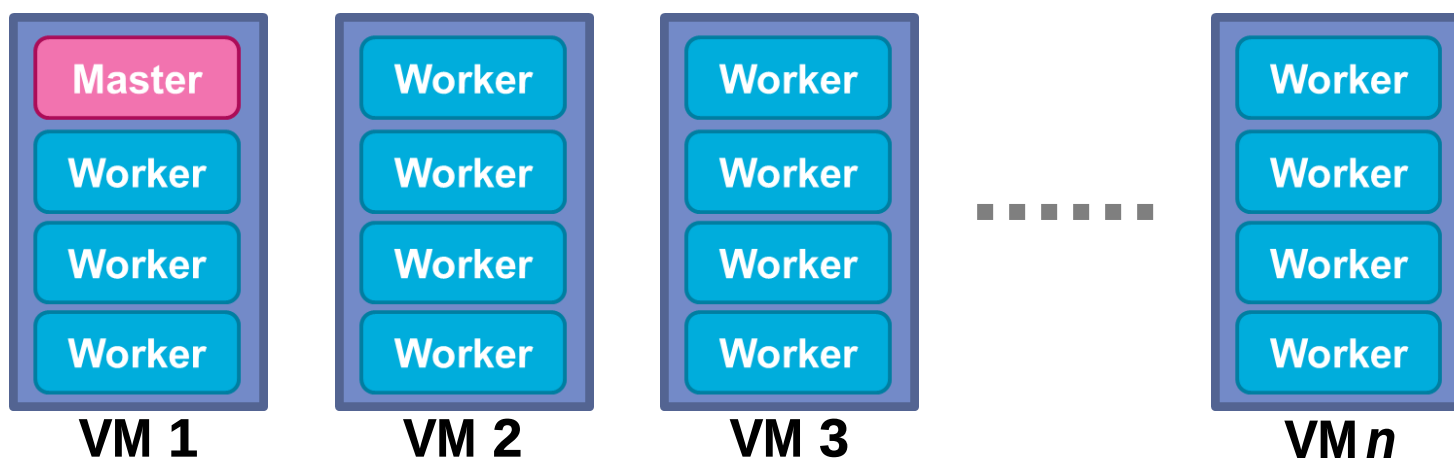
Azure上のシステム構成



• プロセス並列とスレッド並列の粒度を決める

- これまでのクラスタ上での実装の場合、主にメモリ制限のため、各ノード・1プロセスとしていた。
- 最近のクラウドのVMはメモリが豊富。
- (今回対象の)GPUインスタンスはGPUチップが4枚/VM

→ 本研究では、各VM 4プロセスずつ割り当てた。



• データセット

- ZDOCK Docking Benchmark 1.0 (Chen R, *Proteins* 47, 2002)の59タンパク質ペアの相手を入れ替えた全組み合わせ
→ $59 \times 59 = 3,481$ 件の予測計算

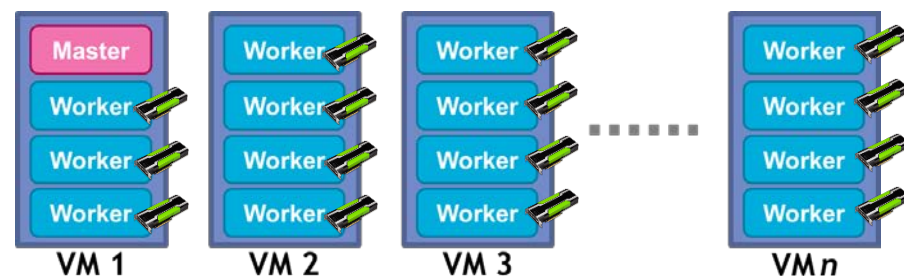
• 使用インスタンス

Instance	CPU	GPU(*)	RAM	Network	価格
NC24	24コア	Tesla K80 x 4	1,440 GB		440.65円/h
NC24r	24コア	Tesla K80 x 4	1,440 GB	RDMA (InfiniBand)	484.71円/h

(*) ただし、物理的にはK80が2枚。
(K80は1枚の中に2つのGPUチップを含む)

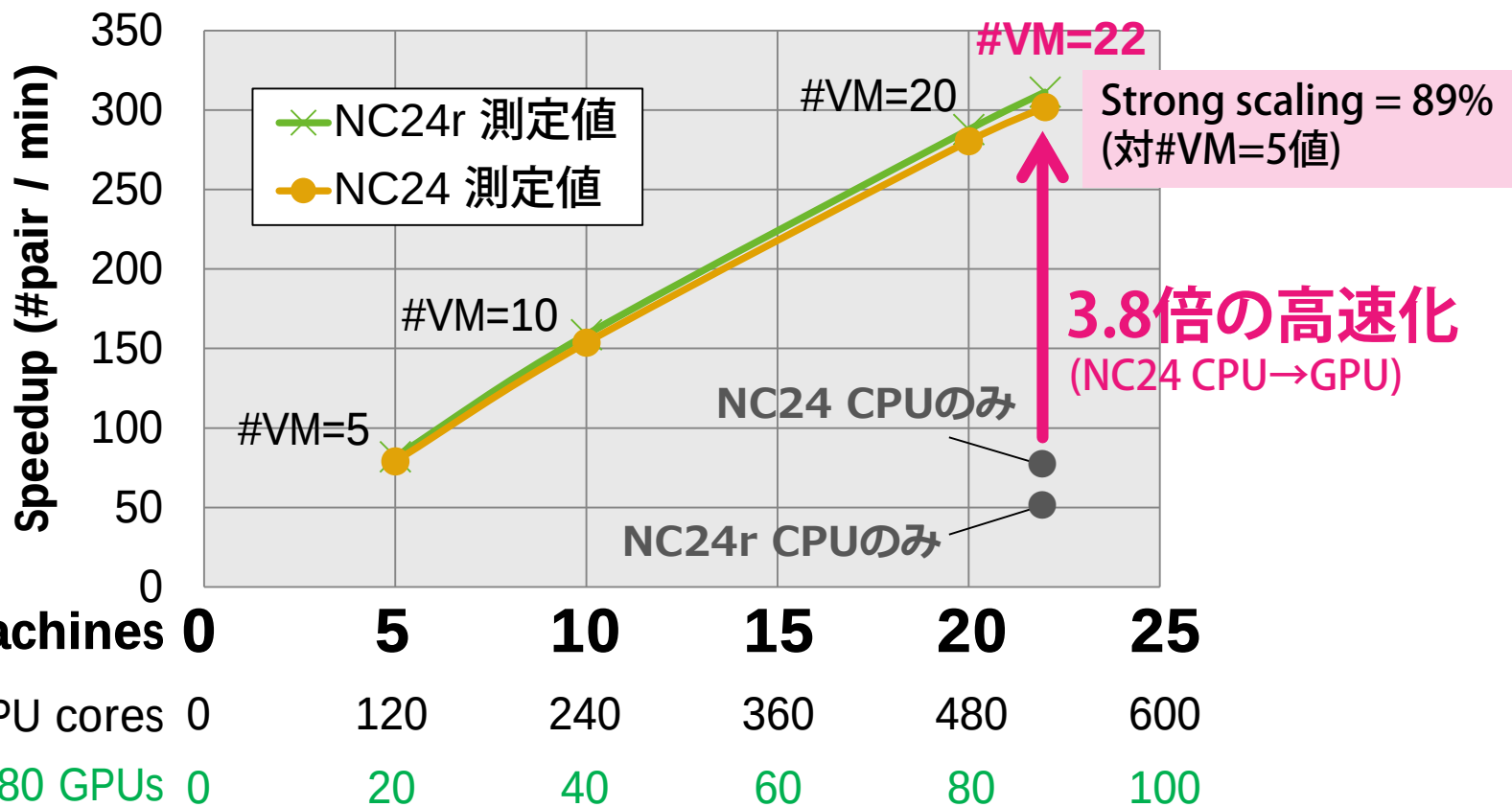
• VM内の実行プロセス

- 各VMに4プロセス
- Masterを1つ立てる
- Worker: 6 CPUコア + 1 GPU



GPUインスタンスでの測定結果

※マイクロソフトのクォータコア制限のため、最大22VMまで測定。



- 89%の並列化効率(強スケーリング)を達成
- GPUを使わずに計算した場合に比べて3.8倍高速

• GPUインスタンスの費用対効果

Instance	CPU	GPU	Network	価格
NC24	E5-2690v3 24コア	Tesla K80 x 4		440.65円/h
NC24r	E5-2690v3 24コア	Tesla K80 x 4	RDMA (InfiniBand)	484.71円/h
A9	E5-2670 16コア	なし	RDMA (InfiniBand)	197.06円/h

→ NC24rとA9を比較

- A9を1.5倍して、24コア→295.59円/hと見積もり。
- NC24r 484.71円/h ÷ 295.59円/h $\div$ 1.64

→GPUによる高速化率が1.64倍以上ならば、GPUインスタンスが得。

今回の場合、GPU利用時で3.8倍高速だったので、MEGADOCKにおいてはGPUインスタンスの利用が推奨される。

- **MEGADOCKをAzure上に移植、AzureのGPUインスタンスを用いた並列実行を実施し、その性能を評価した。**
 - 89%の並列化効率 (NC24, 5 VM → 22 VM)
 - GPUによる3.8倍の高速化 (NC24, CPU → CPU+GPU)
 - InfiniBandはMEGADOCKの計算にはほとんど寄与しない
 - クォータコア制限が41 VMまで緩和されたので、今後測定を実施
 - コンテナ型仮想化技術を用いた場合の検証も今後の課題

(3) 15:50-16:15

コンテナ型仮想化による分散計算環境におけるタンパク質間相互作用予測システムの性能評価
青山 健人、山本 悠生、大上 雅史、秋山 泰

- **謝辞**
 - JSPS科研費 (24240044, 15K16081)
 - JST CREST (Extreme Big Dataプロジェクト)
 - JSTリサーチコンプレックス推進プログラム 川崎市殿町拠点
 - Microsoft Business Investment Funding
 - (株)リバネス 第30回リバネス研究費 (日本マイクロソフト賞)