

環境適応サービス開始後の再構成処理

山登庸次[†]

[†] NTT ネットワークサービスシステム研究所, 東京都武蔵野市緑町 3-9-11

E-mail: tyoji.yamato.wa@hco.ntt.co.jp

あらまし 著者は, 通常コードを, 環境に応じて自動変換し, 高性能で運用可能とする環境適応ソフトウェアを提案し, GPU 等への自動変換や適正配置での性能向上を検証してきた. 本稿では, これまで個々のデバイス等で検討してきた運用中再構成について, 行う全体処理を検討し, リソース量再構成を新たに検討し, さらに, 運用中再構成の処理時間を測定する.

キーワード 環境適応ソフトウェア, 自動オフロード, 最適化, 運用中再構成

Reconfiguration processing after service launch on environment adaptation

Yoji YAMATO[†]

[†] Network Service Systems Laboratories, NTT Corporation, 3-9-11, Midori-cho, Musashino-shi, Tokyo

E-mail: tyoji.yamato.wa@hco.ntt.co.jp

Abstract We have proposed environment-adaptive software that automatically converts normal code according to the environment and enables high-performance operation. In this paper, we study the overall processing performed for reconfiguration during operation, and newly propose resource amount reconfiguration, measure the processing time for reconfiguration during operation, and discuss reconfiguration timings.

Key words Environment Adaptive Software, Automatic Offloading, Optimization, Reconfiguration during Operation

1. はじめに

近年, ムアの法則の鈍化が予想されている. そのような状況から, CPU だけでなく, FPGA (Field Programmable Gate Array) や GPU (Graphics Processing Unit) 等のデバイスの活用が増えている. 例えば, Microsoft 社は FPGA を使って検索効率を高め [1], Amazon 社 [2] は, FPGA, GPU をクラウド技術を用いて (例えば, [3]-[7]) 提供している. また, IoT 機器利用 (例えば, [8]-[19]) も増えている.

しかし, CPU 以外のデバイスを適切に活用するためには, デバイス特性を意識したプログラム作成が必要で, OpenMP (Open Multi-Processing) [20], OpenCL (Open Computing Language) [21], CUDA (Compute Unified Device Architecture) [22] や組み込み技術といった知識が必要になり, スキルの壁が高い. そこで, 著者は, 一度記述したコードを, 配置先環境の FPGA や GPU 等を利用できるよう, 変換, 配置等を自動で行い, アプリを高性能に動作させる, 環境適応ソフトウェアを提案した. 合わせて, その要素として, アプリのループ文等を, FPGA, GPU に自動オフロードする方式, アプリのリソース量や配置を適切化する方式を提案評価している [23]-[27]. さらに,

に, 運用開始後に, 利用状況に応じて構成を変更する FPGA, GPU, 配置再構成方式も提案評価している.

本稿では, これまで個々に検討してきた運用中再構成について, 行う全体処理を検討し, その中で, 未検討だったリソース量再構成を新たに検討し, さらに, 運用中再構成の処理時間を測定し, 実施タイミングを検討する.

2. 運用中再構成

2.1 既存提案

ソフトウェア環境適応実現のため, 著者は図 1 の処理フローを提案している. 環境適応ソフトウェアは, 環境適応機能を中心に, 検証環境, 商用環境, テストケース DB, コードパターン DB, 設備リソース DB が連携し動作する.

ここで, Step 1-6 は, 運用開始前に必要となる, コードの変換, リソース量の設定, 配置場所の設定, 動作確認を行うが, Step 7 は, 運用開始後に, 運用中利用状況に応じて, 構成を変更した方がよい際に再構成する.

2.2 運用中再構成全体処理

運用中再構成を Step 7で行うが, 運用中再構成の目的を明確にする. 環境適応ソフトウェアでは, 利用開始前に, ユーザの

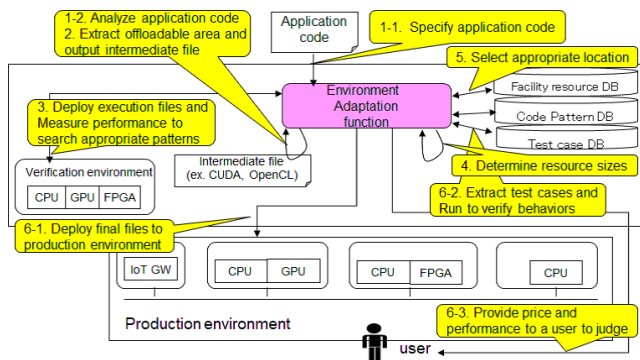


図 1 Processing flow of environment-adaptive software

想定する試験パターンで最適化を行い、コードの変換、リソース量の設定、配置場所の設定を行う。しかし、実際に運用開始後は、事前に想定したリクエスト傾向ではなく、想定外の利用傾向になる可能性もある。例えば、SQL 処理をアクセラレートする FPGA ロジックで運用開始したが、半年後は、NoSQL のクエリが主流になっていた場合等である。このような場合でもより適切な構成に再構成することで、高速な性能を維持するため、実際の利用データを分析した運用中再構成を自動で行う。運用中再構成処理では、運用開始前に最適化を行う環境適応処理を、事前の想定データでなく、実際の利用データに適切に行う。行う環境適応処理は、FPGA ロジック再構成、GPU ロジック再構成、リソース量再構成、配置再構成である。これらは、全て一括して行っても良いし、各処理だけ独立して行っても良く、ビジネス形態に応じて行ってもよい。

2.3 各運用中再構成処理

2.3.1 FPGA, GPU ロジック再構成

FPGA ロジック再構成、GPU ロジック再構成については以前論文で詳細を提案評価している。以下の 6 段階で、ロジックの再構成は行われるが、FPGA も GPU も再構成の大きな考え方としては同じである。

1. 一定期間の、商用リクエストデータ履歴を分析し、処理時間負荷上位の複数アプリを特定し、そのアプリ利用時のデータサイズの最頻データを取得する。

2. 複数の負荷上位アプリで、最頻データのテストケースを高速化する、オフロードパターンを検証環境測定試行を通じて抽出。

3. 現オフロードパターンと抽出した複数の新オフロードパターンの処理時間を測定し、商用利用頻度に基づく性能改善効果を求める。

4. 新オフロードパターン性能改善度効果が現オフロードパターンのその閾値以上であるかで再構成提案を判断。

5. ユーザに FPGA/GPU 再構成実施を提案し、OK/NG の返答を得る。

6. 商用環境で別 OpenCL/OpenACC (Open ACCelerators) [28] を起動することで静的再構成を実施。

2.3.2 リソース量再構成

運用開始前リソース量設定については、以前論文で詳細を提

案評価しているが、その再構成は、本稿で新たに検討する。

運用開始前に、CPU とオフロード先の適切なリソース比を決めるため、[29] も参考に、何れかのデバイス処理がボトルネックとなる事を避けるため、想定テストケースの処理時間を元に、CPU とオフロードデバイスの処理時間が同等オーダーになるように、リソース比を定める。次に商用環境へのアプリの配置を行うが、商用環境への配置の際は、ユーザが指定したコスト要求を満たすように、リソース比を可能な限りキープして、リソース量を決定する。

運用開始後は、事前想定テストケースでなく、そのアプリ利用時最頻データを取得し、その最頻データでのテストケースを用いて、適切なリソース比を計算する。運用開始後の大きな価格上昇はユーザも許容しがたいと考えるため、リソース量決定時は、当初の価格と同程度か低価格になる場合だけ、再構成を提案する形とする。

2.3.3 配置再構成

配置再構成については、以前論文で詳細を提案評価している。オフロードアプリの配置については、運用開始前は [30] 等も参考に他者アプリの配置状況に応じて適切に配置できるが、運用開始後は、他者アプリ配置も増えているため、それらの配置も考慮した全体最適の配置が必要となる。

再構成配置は (1)-(3) 式に応じ、GLPK [31] 等の線形計画ソルバで計算され決定されるが、ユーザ満足度に応じる値 S の計算式 (1) が、運用開始前から新たに加わる。個別ユーザの満足度は、再構成前応答時間 R_k^{before} が再構成後 R_k^{after} で X 倍になったら X が、再構成前価格 P_k^{before} が再構成後 P_k^{after} で Y 倍になったら Y が満足度関連値となる。再配置計算の目的関数は、再構成対象の全ユーザ群満足度に関連した値であり、複数アプリに対して $(X+Y)$ の総和を最小化する配置を計算して、全体最適の配置を求める。

$$S = \sum_{k \in App} \left(\frac{R_k^{after}}{R_k^{before}} + \frac{P_k^{after}}{P_k^{before}} \right) \quad (1)$$

$$\sum_{i \in Device} (A_{i,k}^d \cdot B_{i,k}^p) + \sum_{j \in Link} (A_{j,k}^l \cdot \frac{C_k}{B_k^l}) = R_k^{after} \leq R_k^{upper} \quad (2)$$

$$\sum_{i \in Device} a_i \left(\frac{A_{i,k}^d \cdot B_k^d}{C_i^d} \right) + \sum_{j \in Link} b_j \left(\frac{A_{j,k}^l \cdot B_k^l}{C_j^l} \right) = P_k^{after} \leq P_k^{upper} \quad (3)$$

3. 処理時間測定

今までは、個々の再構成の評価のみしていたが、本稿では環境適応ソフトウェアの行う再構成処理全てを実装し、各処理の処理時間を測定し、実施タイミング等を考察する。

3.1 測定条件

FPGA 再構成では、信号処理 tdFIR [32] を運用開始前に既存手法 [24] で FPGA にオフロードしておき、CPU で動作する画像処理 MRI-Q [33] 含めて、運用開始する。一定期間リクエスト負荷をかけ、性能改善効果が高いオフロードパターンへの再構成提案を確認する。

FPGA オフロード対象ループ文数：tdFIR 6, MRI-Q 16

算術強度絞り込み：算術強度分析の上位 4 つのループ文に絞り込み

リソース効率絞り込み：リソース効率分析の上位 3 つのループ文に絞り込み

実測オフロードパターン数：4（1 回目は上位 3 つのループ文オフロードを測定し，2 回目は 1 回目で高性能だった 2 つのループ文オフロードの組合せパターンを測定．）

リクエスト負荷：tdFIR 200 req/h, MRI-Q 10 req/h の負荷を 3 つのデータサイズで行う．tdFIR では，162KB, 2.06MB, 33.0MB のサンプルデータを 75:120:5 の比でリクエストする．MRI-Q では，32*32*32, 64*64*64, 64*64*64 のダブルサイズ（64*64*64 をコピーして追加）のサンプルデータを 3:5:2 の比でリクエストする．

負荷分析時間：1 時間

負荷上位アプリケーションの数：2

性能改善効果閾値：2

GPU 再構成では，フーリエ変換 NAS.FT [34] を運用開始前に既存手法 [27] で GPU にオフロードし，CPU で動作する流体計算姫野ベンチマーク [35] を含めて，運用開始する．一定期間リクエスト負荷をかけ，性能改善効果が高いオフロードパターンへの再構成提案を確認する．

GPU オフロード対象ループ文数：NAS.FT 81，姫野ベンチマーク 13

個体数 M：NAS.FT 30，姫野ベンチマーク 10

世代数 T：NAS.FT 30，姫野ベンチマーク 10

適合度：(処理時間)^{-1/2} 処理時間が短い程高適合度になる．(-1/2) 乗とすることで，処理時間が短い特定個体の適合度が高くなり過ぎ，探索範囲が狭くなるのを防ぐ．

選択：ルーレット選択．ただし，世代での最高適合度遺伝子は交叉も突然変異もせず次世代に保存するエリート保存も合わせて行う．

交叉率 P_c ：0.9

突然変異率 P_m ：0.05

リクエスト負荷：NAS.FT 20 req/h, 姫野ベンチマーク 30 req/h の負荷を 3 つのデータサイズで行う．NAS.FT では，クラス W, A, B のサンプルデータサイズで，3:5:2 の比でリクエストする．姫野ベンチマークでは，M, L, XL のサンプルデータサイズで，2:5:3 の比でリクエストする．

負荷分析時間：1 時間

負荷上位アプリケーションの数：2

性能改善効果閾値：2

リソース量再構成では，NAS.FT と姫野ベンチマークを，当初想定 of データサイズでリソース量を決めて運用開始後，実際の利用データサイズが想定と大きく異なる際に，リソース量が再構成されることを確認する．

当初想定 of データサイズは，NAS.FT はサイズ W，姫野ベンチマークはサイズ XL とし，実際の負荷をかけるデータサイズは，NAS.FT はサイズ B，姫野ベンチマークは M とする．GPU リソースの分割，再割当には NVIDIA vGPU [36] を用い

る．リソースは，CPU は 1 Core 単位，GPU は 4GB RAM 単位で割当てでき，CPU 1 Core：GPU 4GB RAM のコスト比は 1:2.5 とする．

配置再構成では，NAS.FT を GPU に，MRI-Q を FPGA にオフロードした際の，複数ユーザの平均満足度を向上する，全体配置最適化をシミュレーションする．300 アプリを順に配置し，新たに 100 アプリを配置した際に，100, 200, 400 アプリを再配置計算対象とした際の最適配置を計算する．シミュレーション条件は以下の通りである．

配置トポロジは 4 層で構成され，クラウドレイヤー拠点数は 5，キャリアエッジレイヤーは 20，ユーザエッジレイヤーは 60，インプットノードは 300 とする．

クラウドでは，サーバは CPU 8 台，GPU 16GB RAM 4 台，FPGA 2 台，キャリアエッジでは，CPU 4 台，GPU 8GB RAM 2 台，FPGA 1 台，ユーザエッジでは，CPU 2 台，GPU 4GB RAM 1 台とする．サーバ 1 台の全リソース使用月額額はクラウドでは 5 万，10 万，12 万とした．キャリアエッジ，ユーザエッジは割高になり，クラウドの 1.25 倍，1.5 倍とした．リンクについては，クラウド-キャリアエッジ間は 100Mbps，キャリアエッジ-ユーザエッジ間は 10Mbps の帯域とする．リンクコストは，100Mbps のリンクは月額 8,000 円，10Mbps のリンクは月額 3,000 円とした．

アプリが利用するリソースとして，NAS.FT は，利用リソース量は GPU 1GB RAM，利用帯域 2Mbps，転送データ量 0.2MB，処理時間 5.8 秒である．MRI-Q は，利用リソース量は FPGA サーバの 10%，利用帯域 1Mbps，転送データ量 0.15MB，処理時間 2.0 秒である．

アプリの配置依頼は，300 のインプットノードから上位ノードにランダムに生じさせる．配置依頼数として，NAS.FT：MRI-Q=3:1 の割合で最初に 300 回アプリの配置依頼をする．ユーザ要求条件として，配置依頼する際に価格か応答時間かその両方が選ばれる．NAS.FT の場合，価格に関しては月 7,500 円 (a) か 8,500 円 (b) か 10,000 円 (c) 上限か，応答時間に関しては 6 秒 (A) か 7 秒 (B) か 10 秒 (C) 上限かが選択される．MRI-Q の場合，価格に関しては月 12,500 円 (x) か 20,000 円 (y) 上限か，応答時間に関しては，4 秒 (X) か 8 秒 (Y) 上限が選択される．ユーザ要求として，NAS.FT では a, b, c, A, B, C, aC, bB, bC, cA, cB, cC をそれぞれ 1/12 ずつ，MRI-Q では x, y, X, Y, xY, yX, yY をそれぞれ 1/7 ずつの確率で選択する．

3.2 測定環境

測定環境を図 2 に示す．FPGA 再構成には，Intel FPGA PAC D5,005 を用いて，Intel Acceleration Stack 2.0 [37] で制御する．GPU 再構成には，NVIDIA GeForce RTX 2,080 Ti を用いて，PGI Compiler 19.10 [38] で制御する．リソース量再構成では，NVIDIA Tesla T4 の GPU リソースを，NVIDIA vGPU 12.2 [36] で分割利用する．配置再構成では，GLPK 5.0 [31] でシミュレーションを行う．

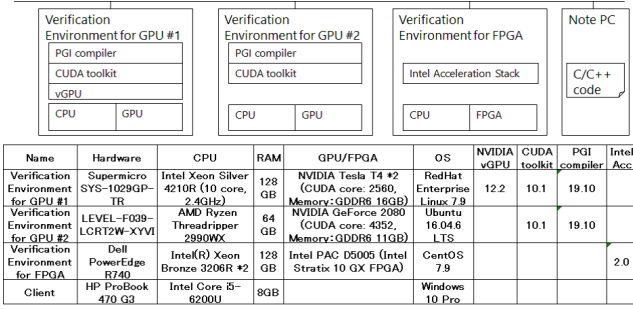


図 2 Verification environments

3.3 結果

FPGA 再構成, GPU 再構成, リソース量再構成, 配置再構成の再構成による改善結果と処理時間を示す。

図 3(a) は, 再構成提案前後の FPGA オフロードの処理時間改善度とそれに関連する一定期間の処理時間合計を示している。再構成前は, tdFIR がオフロードされており, リクエストの処理時間合計の 79.7 秒が 1 時間の負荷である。分析した結果, tdFIR と MRI-Q の 2 つが負荷上位アプリケーションとなる。運用開始後の最頻データを使ってオフロードパターンを探索し, 商用利用回数をかけることで, 処理時間削減改善度は, 再構成前 tdFIR で 41.1 秒/時間, 再構成後 MRI-Q で 252 秒/時間となる。図 3(a) から, 改善度閾値 2.0 をチェックし, tdFIR から MRI-Q にオフロードアプリを変更した再構成が提案される。アプリのコンパイル時間が支配的であるため, 処理時間はそれに依存するが, 両アプリとも 1 回のコンパイルが 6 時間程度で, 4 パターン測定で 1 日程度かかっている。

図 3(b) は, 再構成提案前後の GPU オフロードの処理時間改善度とそれに関連する一定期間の処理時間合計を示している。再構成前は, NAS.FT がオフロードされており, リクエストの処理時間合計の 1,210 秒が 1 時間の負荷である。分析した結果, 姫野ベンチマークと NAS.FT の 2 つが負荷上位アプリケーションとなる。運用開始後の最頻データを使ってオフロードパターンを探索し, 商用利用回数をかけることで, 処理時間削減改善度は, 再構成前 NAS.FT で 308 秒/時間, 再構成後姫野ベンチマークで 1,180 秒/時間となる。図 3(b) から, 改善度閾値 2.0 をチェックし, NAS.FT から姫野ベンチマークにオフロードアプリを変更した再構成が提案される。for 文数等のアプリサイズに処理時間は依存するが, NAS.FT の場合で, 6 時間程度で再構成後のオフロードパターン探索を行っている。

図 4(a) は, NAS.FT をデータサイズ W, 姫野ベンチマークをデータサイズ XL で最適化して運用開始した後, 運用開始後の最頻データがそれぞれサイズ B, M になった際の, 再構成前後の設定リソース量, コスト, コスト対効果を示している。NAS.FT では, データサイズが W から B に大きくなり, 計算量が増えたため, CPU に対して GPU のリソースを増やした方が良いと計算し, GPU を 16GB RAM に増やした方がコスト対効果が上がると考えられたが, 価格が大きく上昇するため, 再構成は提案されない。一方, 姫野ベンチマークでは, データサイズが XL から M に小さくなり, 計算量が減ったため, CPU

(a)	Offload application	Improvement of processing time	Summation of processing time
Before reconfiguration	tdFIR	41.1 sec/h	79.7 sec
After reconfiguration	MRI-Q	252 sec/h	274 sec

(b)	Offload application	Improvement of processing time	Summation of processing time
Before reconfiguration	NAS.FT	308 sec/h	1,210 sec
After reconfiguration	Himeno benchmark	1,180 sec/h	1,400 sec

図 3 (a) FPGA logic reconfiguration results. (b) GPU logic re-configuration results

(a)	Offload application	Set resource amount	Cost	Cost performance
Before reconfiguration	NAS.FT	CPU : GPU = 2 core : 8GB	7,000	1
After reconfiguration	NAS.FT	CPU : GPU = 2 core : 8GB	7,000	
Before reconfiguration	Himeno benchmark	CPU : GPU = 1 core : 16GB	11,000	1
After reconfiguration	Himeno benchmark	CPU : GPU = 1 core : 8GB	6,000	1.5

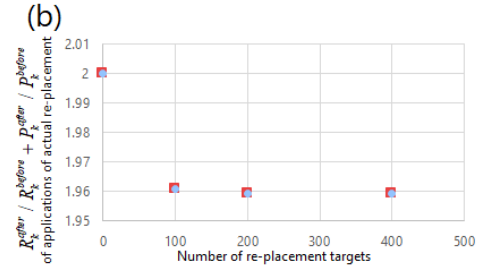


図 4 (a) Resource amount reconfiguration results. (b) Applications re-placement results

に対して GPU のリソースを減らして良いと計算し, GPU を 8GB RAM に減らす提案をし, コスト対効果を 1.5 倍にしている。リソース量再構成の提案は数秒以内にされている。

図 4(b) は, 配置再構成シミュレーションで, 再配置したアプリの $R_k^{after} / R_k^{before} + P_k^{after} / P_k^{before}$ の平均値を縦軸に取ったグラフである。若干ばらつきはあるが, 再配置対象アプリ数の約 1 割弱が, 実際に再構成されている。図 4(b) より, 再配置を行ったアプリは $R_k^{after} / R_k^{before} + P_k^{after} / P_k^{before}$ の平均が 1.96 程度に改善されていることが分かる。この値は 2 から大きく改善される値ではないが, 例えば, NAS.FT をキャリアエッジからクラウドに配置を変えた際は応答時間が 6.6 秒から 7.4 秒になるが, 価格が約 8,400 円から約 7,000 円となるため, 値は 2 から 1.954 になる。再配置計算時間は, 再配置対象のアプリが増えるにつれ, 線形計画の条件式が増えることになるが, 400 アプリでも 1 分以内で終わっている。

3.4 考察

FPGA, GPU, 配置の再構成は別論文で評価しているが, 今回新たに提案したリソース量再構成も, データサイズ傾向を分析して, 価格が同程度から低くなる場合に, コスト対効果を上げる再構成を提案しており, ユーザ利便性は高い。

運用中再構成の各処理時間については, アプリによるが, 今回の測定では, FPGA 再構成は 1 日, GPU 再構成は 6 時間,

リソース量再構成は数秒、配置再構成は 400 アプリ再構成計算の場合に 1 分となっている。FPGA の場合は、アプリのコンパイルに 6 時間程度かかるため、検証環境測定数が 4 でも 1 アプリで 1 日程度かかっている。GPU の場合は、再構成パターン探索に遺伝的アルゴリズム [39] を用いるが、多世代多個体で多数の測定を行うため、1 アプリで 6 時間程度かかっている。リソース量再構成は、最頻データ選出とその際の処理時間に応じたりソース比、量計算であり数秒である。配置再構成は、ソルバの線形計画計算時間に依存するが、100 アプリの場合 10 秒以内だが、400 アプリの場合 1 分と、アプリ数増加に伴い、長大化していく。

リソース量や配置の変更は、価格に影響するため、頻繁な変更は望ましくない。そのため、再構成試行は 1-数か月に一回程度を想定している。ただし、運用開始時は、全ての適応処理を行ってから開始するが、運用中再構成は、個々の処理で独立に行ってよい。例えば、FPGA、GPU ロジックやリソース量の再構成は 3 か月に 1 回と一定期間毎にして、配置再構成は、100 アプリ増えるごとに等一定増加数毎にする等、ビジネス形態に応じて定めればよい。

4. ま と め

本稿では、著者が提案している環境適応ソフトウェアの重要要素の運用中再構成について、行う全体処理を検討し、その中で、リソース量再構成を新たに検討追加し、運用中再構成の各処理時間を測定した。

FPGA、GPU ロジックについては、商用最頻データを用いて、検証環境で定期的に最適化試行することで、適切なパターンを見つける。リソース量は、商用最頻データでの CPU とオフロード先処理時間から、適切なリソース比、量を見つける。配置は、線形計画手法で、応答時間と価格条件で全体最適化を計算することで、適切な配置を見つける。全再構成処理を実装し、処理時間を測定した。FPGA 再構成は 1 日、GPU 再構成は 6 時間、リソース量再構成は数秒、配置再構成は 400 アプリで 1 分であった。

運用中再構成の確認した処理時間に応じて、今後は、各再構成処理の商用実施タイミングについて議論する。

文 献

- [1] A. Putnam, et al., "A reconfigurable fabric for accelerating large-scale datacenter services," Proceedings of the 41th Annual International Symposium on Computer Architecture (ISCA'14), pp.13-24, June 2014.
- [2] AWS EC2 web site, <https://aws.amazon.com/ec2/instance-types/>
- [3] O. Sefraoui, et al., "OpenStack: toward an open-source solution for cloud computing," International Journal of Computer Applications, Vol.55, No.3, 2012.
- [4] Y. Yamato, "Automatic Verification for Plural Virtual Machines Patches," The 7th International Conference on Ubiquitous and Future Networks (ICUFN 2015), pp.837-838, 2015.
- [5] Y. Yamato, "Optimum Application Deployment Technology for Heterogeneous IaaS Cloud," Journal of Information Processing, Vol.25, No.1, pp.56-58, Jan. 2017.
- [6] Y. Yamato, "Use Case Study of HDD-SSD Hybrid Storage,

- Distributed Storage and HDD Storage on OpenStack," 19th International Database Engineering & Applications Symposium (IDEAS '15), pp.228-229, July 2015.
- [7] Y. Yamato, et al., "Fast Restoration Method of Virtual Resources on OpenStack," IEEE Consumer Communications and Networking Conference (CCNC 2015), pp.607-608, Jan. 2015.
- [8] M. Hermann, et al., "Design Principles for Industrie 4.0 Scenarios," Technische Universitat Dortmund. 2015.
- [9] H. Noguchi, et al., "Autonomous Device Identification Architecture for Internet of Things," 2018 IEEE 4th World Forum on Internet of Things (WF-IoT 2018), pp.407-411, Feb. 2018.
- [10] Y. Yamato, et al., "Proposal of Real Time Predictive Maintenance Platform with 3D Printer for Business Vehicles," International Journal of Information and Electronics Engineering, Vol.6, No.5, pp.289-293, 2016.
- [11] Y. Yamato, et al., "Ubiquitous Service Composition Technology for Ubiquitous Network Environments," IPSJ Journal, Vol.48, No.2, pp.562-577, 2007.
- [12] Y. Yamato, et al., "Development of Service Control Server for Web-Telecom Coordination Service," IEEE International Conference on Web Services (ICWS 2008), pp.600-607, Sep. 2008.
- [13] M. Takemoto, et al., "Service Elements and Service Templates for Adaptive Service Composition in a Ubiquitous Computing Environment," The 9th Asia-Pacific Conference on Communications (APCC 2003), Vol.1, pp. 335-338, 2003.
- [14] Y. Yamato, et al., "Study and Evaluation of Context-Aware Service Composition and Change-over Using BPEL Engine and Semantic Web Techniques," IEEE Consumer Communications and Networking Conference (CCNC 2008), pp.863-867, 2008.
- [15] H. Sunaga, et al., "Ubiquitous Life Creation through Service Composition Technologies," World Telecommunications Congress 2006 (WTC 2006), 2006.
- [16] Y. Yamato, et al., "Method of Service Template Generation on a Service Coordination Framework," 2nd International Symposium on Ubiquitous Computing Systems (UCS 2004), Nov. 2004.
- [17] H. Noguchi, et al., "Device Identification Based on Communication Analysis for the Internet of Things," IEEE Access, DOI: 10.1109/ACCESS.2019.2910848, Apr. 2019.
- [18] H. Noguchi, et al., "Distributed Search Architecture for Object Tracking in the Internet of Things," IEEE Access, DOI: 10.1109/ACCESS.2018.2875734, Oct. 2018.
- [19] P. C. Evans and M. Annunziata, "Industrial Internet: Pushing the Boundaries of Minds and Machines," Technical report of General Electric (GE), Nov. 2012.
- [20] T. Sterling, et al., "High performance computing : modern systems and practices," Cambridge, MA : Morgan Kaufmann, ISBN 9780124202153, 2018.
- [21] J. E. Stone, et al., "OpenCL: A parallel programming standard for heterogeneous computing systems," Computing in science and engineering, Vol.12, No.3, pp.66-73, 2010.
- [22] J. Sanders and E. Kandrot, "CUDA by example : an introduction to general-purpose GPU programming," Addison-Wesley, 2011.
- [23] Y. Yamato, "Study and Evaluation of Automatic GPU Offloading Method from Various Language Applications," International Journal of Parallel, Emergent and Distributed Systems, Taylor and Francis, DOI: 10.1080/17445760.2021.1971666, Sep. 2021.
- [24] Y. Yamato, "Automatic Offloading Method of Loop Statements of Software to FPGA," International Journal of Parallel, Emergent and Distributed Systems, Taylor and Francis, DOI: 10.1080/17445760.2021.1916020, Apr. 2021.

- [25] Y. Yamato, "Improvement Proposal of Automatic GPU Offloading Technology," The 8th International Conference on Information and Education Technology (ICIET 2020), pp.242-246, Mar. 2020.
- [26] Y. Yamato, "Proposal of Automatic Offloading for Function Blocks of Applications," The 8th IIAE International Conference on Industrial Application Engineering 2020 (ICIAE 2020), pp.4-11, Mar. 2020.
- [27] Y. Yamato, "Study and Evaluation of Improved Automatic GPU Offloading Method," International Journal of Parallel, Emergent and Distributed Systems, Taylor and Francis, DOI: 10.1080/17445760.2021.1941010, June 2021.
- [28] S. Wienke, et al., "OpenACC-first experiences with real-world applications," Euro-Par 2012 Parallel Processing, pp.859-870, 2012.
- [29] K. Shirahata, et al., "Hybrid Map Task Scheduling for GPU-Based Heterogeneous Clusters," IEEE Second International Conference on Cloud Computing Technology and Science (CloudCom), pp.733-740, Dec. 2010.
- [30] C. C. Wang, et al., "Toward optimal resource allocation of virtualized network functions for hierarchical datacenters," IEEE Transactions on Network and Service Management, Vol.15, No.4, pp.1532-1544, 2018.
- [31] GLPK website, <https://sites.google.com/site/nssvdabb/glpk>
- [32] Time domain finite impulse response filter web site, <http://www.omgwiki.org/hpec/files/hpec-challenge/tdfir.html>
- [33] MRI-Q website, <http://impact.crhc.illinois.edu/parboil/>
- [34] NAS.FT website, <https://www.nas.nasa.gov/publications/npb.html>
- [35] Himeno benchmark web site, <http://accr.riken.jp/en/supercom/>
- [36] NVIDIA vGPU software web site, <https://docs.nvidia.com/grid/index.html>
- [37] Intel Acceleration Stack website, <https://www.intel.com/content/www/us/en/software-kit/665814/intel-fpga-pac-d5005-acceleration-stacks-v2-0-1.html>
- [38] M. Wolfe, "Implementing the PGI accelerator model," ACM the 3rd Workshop on General-Purpose Computation on Graphics Processing Units, pp.43-50, Mar. 2010.
- [39] J. H. Holland, "Genetic algorithms," Scientific american, Vol.267, No.1, pp.66-73, 1992.