

گروه بندی دو مرحله ای کارها با استفاده از الگوریتم های ژنتیکی و اِعمال سیاست های مختلف زمانبندی در گریدهای محاسباتی

صادق وهابزاده زرگری^{*}، عادل ترکمان رحمانی[†]

چکیده

گریدهای محاسباتی با توان عملیاتی زیاد، به طور گسترده برای اجرای کاربردهای محاسباتی تشکیل شده از کارهای مرتبط به هم مورد استفاده قرار می گیرند. پردازش این کاربردها به صورت رویهم رفته در گریدها، هزینه و زمان زیادی را در بر دارد که ناشی شده از سربار به وجود آمده از (۱) ارسال کارها به منابع، (۲) پردازش کارها و (۳) انتقال نتیجه پردازش شده به سمت کاربر می باشد. در این مقاله استراتژی زمانبندی مبتنی بر گروه بندی دو مرحله ای، با اعمال الگوریتم های ژنتیکی توسعه یافته است. الگوریتم های ژنتیکی ابتدا در مرحله اول، یک کاربرد متشکل از کارهای وابسته به هم را به مجموعه ای از کاربردهای وابسته به هم تبدیل می کند و سپس در مرحله دوم با استفاده از اعمال روش گروه بندی ژنتیکی بر روی تک تک کاربردهای حاصل، آنها را زمانبندی می کند. در ضمن سعی شده است امکان اعمال سیاست های مختلف زمانبندی بر اساس کمینه نمودن هزینه و زمان اجرا وجود داشته باشد. نتایج و جزئیات ارزیابی با استفاده از اجرای یک شبیه ساز مبتنی بر شبیه ساز *GidSim* نشان داده شده است.

کلمات کلیدی

وابستگی کارها، گروه بندی دو مرحله ای کارها، زمانبندی در سطح کاربرد، گریدهای محاسباتی با توان عملیاتی زیاد، کاربردهای متشکل از وظیفه های مجزا، الگوریتم های ژنتیکی

Two-stage Grouping Algorithm using Genetic Algorithms Based on Different Scheduling Policies on Computational Grids

Sadegh Vahabzadeh Zargari[‡], Adel Torkaman Rahmani[§]

Abstract

Although high throughput computational Grids have been used extensively for executing applications with compute-intensive jobs, there exist several applications with a large number of lightweight jobs. The overall processing undertaking of these applications involves high overhead time and cost in terms of (i) job transmission to and from Grid resources and, (ii) job processing at the Grid resources.

In this paper, we present two-stages scheduling strategy and use genetic algorithms to optimizing both of those stages. In the first stage a collaborative-tasks application transforms to subset of bag-of-tasks applications. In the second stage, we apply job grouping algorithm to scheduling these applications. Our scheduling strategy uses both of cost-based and time-based policies that use genetic algorithms for performing dynamic job grouping activity at runtime. The detailed analysis is conveyed by running simulations.

Keywords

Job Dependencies, 2-stage Job Grouping, Application Level Scheduling, High Throughput Computational Grids, Bag-of-Tasks Applications, Genetic Algorithms

^{*} دانشگاه علم و صنعت ایران، دانشکده مهندسی کامپیوتر، Vahabzadeh@comp.iust.ac.ir

[†] عضو هیأت علمی دانشگاه، دانشگاه علم و صنعت ایران، دانشکده مهندسی کامپیوتر، Rahmani@iust.ac.ir

[‡] Iran University of Science and Technology, Computer Engineering Department, Vahabzadeh@comp.iust.ac.ir

[§] Academic Staff, Iran University of Science and Technology, Computer Engineering Department, Rahmani@iust.ac.ir

۱- مقدمه

با پدید آمدن ایده‌های محاسباتی، که توسط فوستر و کیلمن [۱] در سال ۱۹۹۹ مطرح گردید، یک سکوی نرم افزاری جدید برای اجرای کاربردها بر روی تعداد زیادی از منابع محاسباتی نامتجانس و توزیع شده، و با عبور از حوزه‌های مدیریتی، و اعمال سیاست‌های مختلف به وجود آمد. با توجه به تعریف، یک کاربرد مجموعه‌ای از وظیفه‌ها و کارهای متعدد است و هر کار^۱ به عنوان یک واحد اتمی از محاسبات کاربرد تعریف می‌گردد.

تمام کاربردها از پتانسیل لازم برای مورد اجرا واقع شدن در گرید محاسباتی برخوردار نیستند و امکان دارد هزینه و زمان اجرای برخی از کاربردها در گرید بسیار بالاتر از اجرای آن کاربرد در یک ماشین خاص باشد. با وجود این، کلاس ایده‌آلی از کاربردها وجود دارد که از پتانسیل بسیار بالایی برای اجرا شدن در گرید برخوردار هستند. این کاربردها از یک سری وظیفه‌های مستقل و مجزا^۲ که نیازی به ارتباط درون پردازشی ندارند، تشکیل شده‌اند [۲]. هر یک از کارها و وظیفه‌های موجود در این کاربردها را می‌توان بدون توجه به سایر کارها زمانبندی و اجرا نمود. اگر این نوع کاربردها را در رده اول قرار دهیم، کاربردهای رده پایین‌تری هم وجود دارند. به عنوان مثال می‌توان به کاربردهایی اشاره نمود که از مجموعه‌ای از وظیفه‌های مرتبط و کارهای وابسته به هم تشکیل شده‌اند [۶]. در این کاربردها به دلیل وابستگی‌های موجود، پتانسیل بالای رده اول برای مورد اجرا واقع شدن در گرید وجود ندارد.

یک چالش مهم تبدیل کاربردهای رده پایین‌تر به کاربردهای رده اول که از قابلیت موازی سازی بالاتری برای مورد اجرا واقع شدن در گرید برخوردار هستند، می‌باشد.

چالش دیگر در این زمینه می‌تواند اجرای کاربردهای رده اول در یک محیط گرید محاسباتی باشد. با توجه به اینکه اجرای هر کاربرد عبارت است از زمانبندی کارهای تشکیل دهنده آن کاربرد، چالش مهم دیگر و قابل بحث در این زمینه، زمانبندی این کارها می‌باشد.

با توجه به تعریفی که توسط برمن، فاکس و هی [۴] ارائه شده است، یک زمانبند^۳ مسئول انتخاب مناسبترین ماشین‌ها و منابع محاسباتی برای پردازش کارها در گرید می‌باشد و این انتخاب با هدف رسیدن به توان عملیاتی بالاتر، و با توجه به مشخصه‌های هر کاربرد، دسترس پذیری، نیازمندیها و قابلیت پردازشی منابع انجام می‌پذیرد. مثلاً در مورد کاربردی که از تعداد زیادی کار با مشخصه پردازشی کم وزن^۴ تشکیل شده است، کل زمان ارتباط میان هر کار و منابع بسیار بالاتر از کل زمان محاسبه و پردازش آن کار می‌باشد. راه حل مناسب برای

Job-^۱

Bag-of-tasks applications-^۲

Scheduler-^۳

Light weight jobs-^۴

۲- کارهای مرتبط انجام شده

پروژه GrADSolve^۵ [۹] یک سیستم RPC^۶ را به منظور انجام پردازش‌ها و محاسبات موازی در گرید، ارائه می‌دهد. زمانبندی GrADSolve، بر اساس زمانبندی در سطح کاربرد و پروژه AppLeS^۷ می‌باشد و از ترکیب دو پروژه عظیم مرتبط با گرید به نام های GrADS^۸ و NetSolve^۹ به وجود آمده است. GrADS فن آوریها و ابزارهای لازم را برای توسعه و اجرای کاربردها در محیط گرید فراهم می‌آورد و یک روش زمانبندی در سطح کاربرد، ارائه می‌دهد. NetSolve یک سیستم گرید محاسباتی بر اساس RPC است که با

Fine-grained jobs-^۵

Coarse-grained jobs-^۶

Meta-jobs-^۷

Grid Application Development Solver-^۸

Remote Procedure Call-^۹

Application Level Scheduling-^{۱۰}

Grid Application Development Software-^{۱۱}

$$n = \sum_{i=1}^m n_i \quad (3)$$

در روش زمانبندی Sub optimal [۸] با ارائه یک زمانبند فوقانی^{۲۰}، یک زمانبندی دو سطحی ارائه شده است. در سطح اول گره‌های محاسباتی قرار دارند که هر کدام زمانبند محلی و سیاست زمانبندی خاص خودشان را دارند و در سطح دوم زمانبند فوقانی قرار دارد که هدف آن پیدا کردن دنباله‌ای مناسب از کارها برای تخصیص به هر گره با استفاده از اعمال الگوریتم‌های ژنتیکی می‌باشد.

در روش زمانبندی با استفاده از گروه بندی کارها [۵]، یک روش زمانبندی کارها برای کاربردهایی که از وظیفه‌های ریزدانه و مجزا تشکیل شده‌اند، ارائه شده است. گروه بندی کارها با هدف تبدیل چندین کار ریزدانه به یک کار درشت دانه و به منظور کاهش نرخ ارسال کارها انجام می‌گیرد. گروه بندی بر اساس پارامتر ویژه‌ای به نام Granularity Size انجام می‌شود. Granularity Size مدت زمانی است که انتظار می‌رود یک کار در یک منبع محاسباتی مورد پردازش قرار گیرد. این پارامتر در قابلیت پردازشی منابع ضرب می‌شود تا تعداد مورد انتظار کارهای قابل پردازش در یک زمان معین، برای هر منبع محاسباتی به دست آید. در روش گروه‌بندی با استفاده از الگوریتم‌های ژنتیکی [۱۳] ضمن بهینه‌سازی روش گروه‌بندی، انجام زمانبندی با توجه به سیاست‌های مختلف میسر شده است. روش زمانبندی مبتنی بر گروه‌بندی دو مرحله‌ای کارها [۱۴] الگوریتم زمانبندی را برای کاربردهای متشکل از وظیفه‌های وابسته به هم قابل اسنفاده نموده است.

۳- روش گروه بندی دو مرحله‌ای

با وجود اینکه روش زمانبندی تطبیقی با ارائه مدل‌های پیش بینی در هر دو سطح سیستم و کاربرد، زمانبندی مناسبی ارائه نموده است، اما به ماهیت واقعی بسیاری از موجودیت‌ها توجه نداشته است. به عنوان مثال فقط به تعداد کارها توجه شده است و بار پردازشی کارها مورد بحث قرار نگرفته است. در ضمن به هزینه اجرای روش توجه نشده است.

روش گروه بندی کارها با آنکه توانسته است به میزان قابل اعتنائی هزینه و زمان پردازش کارها را کاهش دهد، اما فقط برای زمانبندی کاربردهای رده اول که از کارهای مجزا تشکیل شده‌اند، ارائه شده است. این روش کارآیی لازم را در مورد کاربردهایی که از وظیفه‌های مرتبط و کارهای وابسته به هم تشکیل شده‌اند، ندارد.

روش پیشنهادی با اعمال الگوریتم‌های ژنتیکی، یک استراتژی گروه بندی دو مرحله‌ای را برای زمانبندی کارها ارائه می‌دهد. این روش در مرحله اول یک کاربرد متشکل از کارهای وابسته به هم را به مجموعه‌ای از کاربردهای وابسته به هم که از رده اول هستند، تبدیل می‌کند و سپس در مرحله دوم با استفاده از اعمال روش گروه بندی [۱۳] بر

هدف اجرای کاربردهای موازی از راه دور به وجود آمده است. هسته اصلی پروژه GrADSolve بر اساس بانک اطلاعات XML^{۱۲} است.

در روش زمانبندی اکتشافی^{۱۳} [۳] با ارائه الگوریتم Qsufferage، سعی شده است یک روش زمانبندی با احترام به پارامترهای کیفیت سرویس و برای رده خاصی از کاربردها که متشکل از وظیفه‌های مجزا هستند، ارائه گردد. هدف این روش کمینه کردن زمان پاسخ می‌باشد. نسبت پاسخ^{۱۴} (IT) به عنوان پارامتر کیفیت سرویس (QoS) به صورت رابطه ۱ تعریف شده است.

$$rr = \frac{\text{the} \therefore \text{duration} \therefore \text{of} \therefore \text{running}}{\text{the} \therefore \text{completed} \therefore \text{time} - \text{the} \therefore \text{submitted} \therefore \text{time}} \quad (1)$$

در روش زمانبندی مجزا^{۱۵} [۱۰] زمانبند و قطعات ویژه کاربرد که در انجام پروسس زمانبندی مورد نیاز هستند، از هم تفکیک شده‌اند. زمانبند دارای یک پروسس‌چر جستجو است که اطلاعات ورودی آن سیستم‌های اطلاعاتی‌گرید، مدل کارآیی^{۱۶}، داده‌های تعیین احتمالات، لیست تطبیق^{۱۷}، و اجزای دیگر هستند و خروجی آن زمانبندی نهایی می‌باشد.

مدل کارآیی هر کاربرد با استفاده از بررسی آن کاربرد، و با هدف پیش بینی زمان اجرا و میزان حافظه مورد نیاز به دست می‌آید. در این روش سعی شده است، با توجه به ماهیت پویا بودن گرید، درصد خطای پیش بینی^{۱۸} کاهش یابد. درصد خطای پیش بینی با فرمول رابطه ۲ محاسبه می‌شود.

$$\text{predictionError} = 100 * \frac{\text{predictionTime} - \text{actualTime}}{\text{actualTime}} \quad (2)$$

در روش زمانبندی تطبیقی^{۱۹} [۷] با ارائه مدل‌های پیش بینی، زمانبندی کارها در هر دو سطح، یعنی هم در سطح سیستم و هم در سطح کاربرد انجام می‌شود. در سطح کاربرد، با اعمال الگوریتم‌های ژنتیکی، ضمن تخصیص بهینه کارها به منابع، زمان متوسط تکمیل کارها کاهش می‌یابد. در این روش، گونه‌های مختلف جمعیت ژنتیکی با کنار هم قرار دادن تعداد کارهای احتمالی تخصیص یافته به هر منبع به دست می‌آید. طول هر عضو جمعیت برابر تعداد منابع و هر ژن نشان دهنده تعداد کارهای تخصیص یافته به یک منبع است. به عنوان مثال با داشتن n کار و m منبع اعضای جمعیت ژنتیکی به صورت n₁ n₂ n₃ . . . n_m خواهد بود و شرط صحت، برابر بودن مجموع تعداد کارهای منابع با تعداد کل کارها طبق رابطه ۳ می‌باشد.

¹²-eXtensible Markup Language

¹³-Heuristic Scheduling

¹⁴-Response Ratio

¹⁵-Decoupled Scheduling

¹⁶-Performance Model

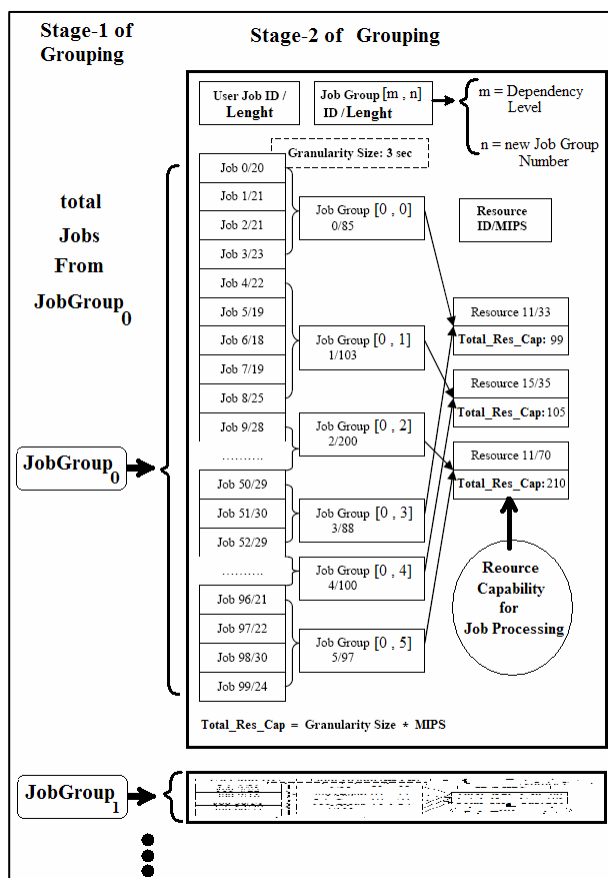
¹⁷-Match List

¹⁸-Prediction Error

¹⁹-Adaptive Scheduling

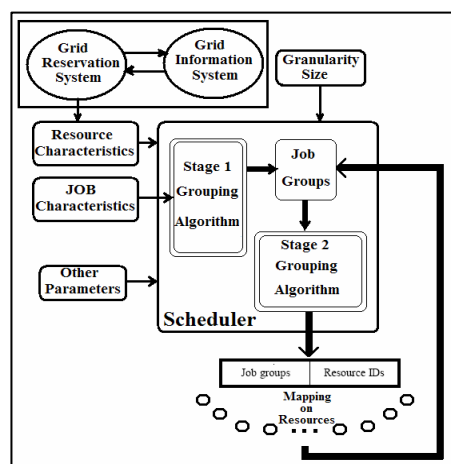
²⁰-Superscheduler

Granularity Size مدت زمانی است که انتظار می رود یک کار در یک منبع محاسباتی مورد پردازش قرار گیرد. این پارامتر در قابلیت پردازشی منابع ضرب می شود تا تعداد مورد انتظار کارهای قابل پردازش در یک زمان معین، برای هر منبع محاسباتی به دست آید.



شکل (۲): یک مثال از استراتژی مرحله دوم گروه بندی

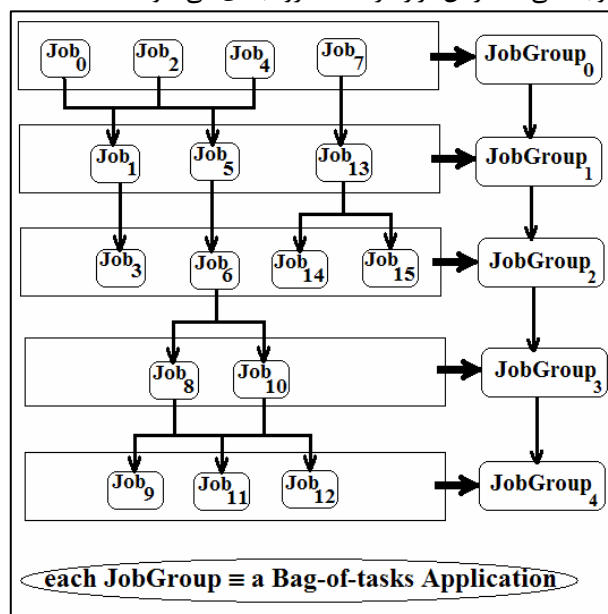
شمای کلی روش گروه بندی دو مرحله ای در شکل ۳ نشان داده شده است. به علت پویا بودن گرید، برای اطمینان از در دسترس بودن منابع در زمان اجرای کارها، وجود یک سیستم رزرو منابع که با یک سیستم اطلاعاتی در ارتباط باشد، ضروری است.



شکل (۳): شمای کلی زمانبند ارائه شده

روی تک تک کاربردهای حاصل، آنها را به ترتیب وابستگی، زمانبندی می کند.

بنابراین لازم است که ابتدا کاربردهای متشکل از کارهای وابسته به هم، به کاربردهای رده اول تبدیل شوند. برای این منظور باید کارها را به صورتی گروه بندی نمود که به جای کارهای وابسته به هم، گروه های کاری وابسته به هم داشته باشیم و کارهای دسته بندی شده در هر گروه کاری نیز مستقل از هم باشند. در این صورت می توان با هر گروه کاری حاصل شده، مثل یک کاربرد رده اولی برخورد نمود. استراتژی انجام این کار در شکل ۱ با ذکر مثالی بیان شده است. همانطور که در شکل ۱ قابل رؤیت است، کارها با توجه به سطح وابستگی که در آن قرار گرفته اند گروه بندی می شوند.



شکل (۱): یک مثال از استراتژی مرحله اول گروه بندی

هر یک از گروه های کاری به دست آمده از مرحله اول گروه بندی را می توان هم ارز یک کاربرد رده اول در نظر گرفت. بنابراین در مرحله دوم می توان بر روی هر یک از گروه های کاری به دست آمده از مرحله اول، الگوریتم گروه بندی را که برای کاربردهای رده اول ارائه شده است اعمال نمود.

آنچه که مسلم است تعداد دفعات اعمال الگوریتم گروه بندی در مرحله دوم برابر تعداد گروه های کاری ایجاد شده است و از آنجا که تعداد گروه های کاری ایجاد شده با تعداد سطوح وابستگی نسبت مستقیم دارد، پارامتر Dependency Level اهمیت بالاتری پیدا می کند

در شکل ۲ استراتژی مرحله دوم گروه بندی کارها با استفاده از ذکر یک مثال نشان داده شده است. همانطور که از شکل ۲ قابل فهم است، الگوریتم گروه بندی مرحله دوم برای هر یک از گروه های کاری به دست آمده از مرحله اول یک بار اجرا می شود.

گروه بندی کارها با هدف تبدیل چندین کار ریزدانه به یک کار درشت دانه و به منظور کاهش نرخ ارسال کارها انجام می گیرد. گروه بندی بر اساس پارامتر ویژه ای به نام Granularity Size انجام می شود.

// Stage 2 Grouping Algorithm

```

24 NextResource := 0;
// Dependency Loop
25 for deLevel := 0 to DependencyLevel-1 do
// get Jobs List from current Dependency Level
26 JobList := JobGroupsDependencyLevel;
// Grouping Algorithm
27 m := 0;
28 for i := 0 to JobList.Size-1 do
29 for j := NextResource to ResourceList.Size-1 do
30 if NextResource = ResourceList.Size-1 then
31 NextResource := 0;
32 endif
33 Total_JobLoad := 0;
34 Total_ResourceCapability :=
ResourceListj.MIPS * Granularity_Size;
35 while (Total_JobLoad ≤ Total_ResourceCapability
and i ≤ JobList.Size-1) do
36 Total_JobLoad :=
Total_JobLoad + JobListi.JobLenght;
37 i++;
38 endwhile
39 i--;
40 if Total_JobLoad > Total_ResourceCapability then
41 Total_JobLoad :=
Total_JobLoad - JobListi.JobLenght;
42 i--;
43 endif
44 GroupJobListm :=
Add_new_Job (Total_JobLoad, assign unique ID);
45 GroupJobListm.TargetResource := ResourceListj.ID;
46 m++;
47 endfor
48 endfor
// Sending jobs for computation
49 for i := 0 to GroupJobList.Size-1 do
50 Send GroupJobListi to TargetResource
for job computation;
51 endfor
// Receiving computed results
52 for i := 0 to GroupJobList.Size-1 do
53 Receive computed GroupJobListi
from TargetResource;
54 endfor
55 endfor

```

شکل (۵): شبه کد مرحله دوم گروه بندی

زمان کل پردازش با استفاده از رابطه (۴) محاسبه می شود. همانطور که قابل مشاهده است، زمان سربار حاصل از عمل گروه بندی هم به کل زمان افزوده شده است.

$Total Processing Time =$

$Grouping Algorithm Time +$

$Genetic Algorithm Time +$

$Total Processing Overhead Time +$

$Total JobGroups Computing Time +$

$Time taken to receive all computed jobs$

(۴)

۴- روش پیشنهادی و محیط شبیه سازی

در روش پیشنهادی، الگوریتم های ژنتیکی یک بار در مرحله اول گروه بندی و بار دیگر در مرحله دوم آن اعمال می گردد.

در مرحله اول هر عضو از جمعیت^{۱۱} ژنتیکی، از جایگشت های تصادفی گروه های کاری حاصل از عمل گروه بندی به وجود می آید.

به عنوان مثال اولین عضو در جمعیت ژنتیکی حاصل، به صورت رابطه ۵ خواهد بود.

$Population_1 =$

$JobGroup_1 JobGroup_2 \dots JobGroup_{DependencyLevel}$

(۵)

$Population_{-21}$

اطلاعات منابع، کارها، و پارامتر Granularity Size، و سایر پارامترهای لازم به زمانبندی اعمال می شوند و زمانبندی، در مرحله اول با گروه بندی کارها بر اساس وابستگی آنها، کارهای وابسته به هم را به گروه های کاری وابسته به هم تبدیل می کند و در مرحله دوم با استفاده از روش گروه بندی، کارهای موجود در هر گروه کاری را زمانبندی می کند و ماتریس نگاشت کارها را به منابع مناسب تولید می کند. هر گروه کاری منتظر دریافت نتیجه پردازش شده گروه کاری قبلی می ماند و با استفاده از روش گروه بندی زمانبندی می شود.

شبه کد الگوریتم مرحله اول گروه بندی در خطوط ۱ الی ۲۳ شکل ۴ ارائه شده است و با توجه به مثال ذکر شده در شکل ۱ نیاز به توضیحات بیشتر وجود ندارد. در خط ۱۱ گروه کاری جدید مربوط به یک سطح وابستگی ایجاد می شود و در خط ۱۶ وابستگی ها بررسی می شود. از نتایج مهم این الگوریتم معرفی پارامتر سطح وابستگی یا Dependency Level است و هسته مرکزی آن، افزودن کارها به گروه کاری سطح وابستگی مربوطه است که در خط ۱۸ بیان شده است.

// Stage 1 Grouping Algorithm

```

1 ResourceList := Resource_characteristics(GRS, GIS);
2 JobList := Job_characteristics_from_application();
3 DependencyLevel := 0;
4 for i := 0 to JobList.Size-1 do
5 JobListi.used = false;
6 endfor
7 for i := 0 to JobList.Size-1 do
8 if JobListi.used = true then
9 continue;
10 endif
11 JobGroup.AddnewGroup ( DependencyLevel );
12 for j := 0 to JobList.Size-1 do
13 if JobListj.used = true then
14 continue;
15 endif
16 if ( JobListj.beforJob = Null ) or
( JobListj.beforJob.used = true ) then
17 JobListj.used = true;
18 JobGroupDependencyLevel.add( JobListj )
19 endif
20 endfor
21 DependencyLevel ++;
22 endfor
23 DependencyLevel --;

```

شکل (۴): شبه کد مرحله اول گروه بندی

همچنین شبه کد الگوریتم مرحله دوم گروه بندی در خطوط ۲۴ الی ۵۵ شکل ۵ ارائه شده است. در این مورد نیز با توجه به مثال ذکر شده در شکل ۲ نیاز به توضیحات اضافی احساس نمی شود. در خط ۲۵، حلقه تکرار الگوریتم مرحله دوم به تعداد سطوح وابستگی قرار دارد. در خط ۲۶ کارهای موجود در گروه کاری ایجاد شده از مرحله اول گروه بندی که مربوط به سطح وابستگی جاری هستند، در لیست کارهای الگوریتم گروه بندی مرحله دوم قرار می گیرند.

در خطوط ۲۷ الی ۴۸ گروه بندی کارها با توجه به پارامتر Granularity Size انجام می گیرد. در خطوط ۴۹ الی ۵۱ گروه های کاری ایجاد شده به منابع ارسال می شوند و در خطوط ۵۲ الی ۵۴، الگوریتم در انتظار دریافت نتایج پردازش شده در منابع می ماند.

در خط ۵۵ شاهد ادامه تکرار حلقه for و تکرار دوباره خطوط ۲۶ الی ۵۴ برای کارهای موجود در سطوح وابستگی پایین تر که منتظر دریافت نتیجه پردازش کارها در سطوح بالاتر بودند، خواهیم بود.

$Fitness_function =$

$$\frac{|Total_JobLoad - Total_ResourceCapacity|}{Total_JobLoad} \quad (9)$$

در صورتیکه سیاست زمانبندی، کمینه نمودن هزینه^{۲۵} محاسبه باشد، تابع برازندگی از رابطه ۱۰ محاسبه می‌گردد.

$Fitness_function =$

$$ResourceList_i.Cost * \frac{Total_JobLoad}{Total_ResourceCapacity} \quad (10)$$

اگر سیاست زمانبندی این باشد که یک مُوازنه در کمینه نمودن هزینه محاسبه و زمان انجام آن برقرار شود، تابع برازندگی از رابطه ۱۱ محاسبه می‌گردد.

$Fitness_function =$

$$ResourceList_i.Cost * \frac{|Total_JobLoad - Total_ResourceCapacity|}{Total_ResourceCapacity} \quad (11)$$

در هر تکرار با اِعمال توابع جهش^{۲۶} ژنتیکی و پیوند^{۲۷} ژنتیکی، اعضای جمعیت به جواب بهینه همگرا می‌شوند. نحوه تنظیم پارامترهای ژنتیکی، تاثیر قابل توجهی در سرعت همگرایی و نتیجه نهایی دارد. بنابراین پیدا کردن یک مُوازنه^{۲۸} مناسب در تنظیم پارامترهای ژنتیکی از چالش‌های مهم می‌باشد.

جهت مقایسه الگوریتم پیشنهادی با روش‌های مشابه، یک محیط شبیه‌سازی مبتنی بر شبیه‌سازی GridSim [۱۱] که به زبان جاوا نوشته شده، پیاده‌سازی شده است.

GridSim یک جعبه ابزار است که برای مدل‌کردن و شبیه‌سازی مدیریت منابع گرید و زمانبندی اجرای کارها طراحی شده است. این جعبه ابزار متشکل از دو بسته نرم‌افزاری به نامهای gridsim و simjava می‌باشد.

کلاسهای موجود در بسته gridsim شبیه‌سازی موجودیت‌های گرید، هستند. به عنوان مثال، کاربر گرید با کلاس GridUser، کارها با کلاس Gridlet، و منابع با کلاس GridResource نمونه‌سازی می‌شوند. GridSim از امکانات لازم جهت ایجاد کارهای با طول‌های متفاوت و استاندارد، تعیین توپولوژی، ارسال کارها به منابع، دریافت نتیجه محاسبه، و شبیه‌سازی زمان محاسبه برخوردار است.

با توجه به شکل ۷ و محیط شبیه‌سازی، اندازه متوسط کارها برابر ۵۰۰ در نظر گرفته شده است و برای تولید کارها با طولهای متفاوت، حداکثر میزان انحراف ۹۰٪ ± در نظر گرفته شده است. با این توضیح که نتیجه آزمون شبیه‌سازی با تعیین تعداد کل کارها برابر ۱۰۰، ۲۰۰، ۵۰۰، ۱۰۰۰ و ۵۰۰۰ مشابه هم بوده است، و با توجه به اینکه پارامتر Dependency Level نقش تعیین‌کننده تری دارد، تعداد کارها در

تابع برازندگی^{۲۲} هم به گونه‌ای انتخاب می‌شود که حداکثر تعادل بار^{۲۳} را در بین گروههای کاری مختلف داشته باشیم. بین جمعیت‌های ایجاد شده، جمعیت‌هایی مورد قبول خواهند بود که وابستگی کارها را همانند وابستگی موجود بین آن کارها در جمعیت اولیه حفظ کنند. برای این کار ابتدا تابع Δ_{ij} را به صورت رابطه ۶ تعریف می‌کنیم.

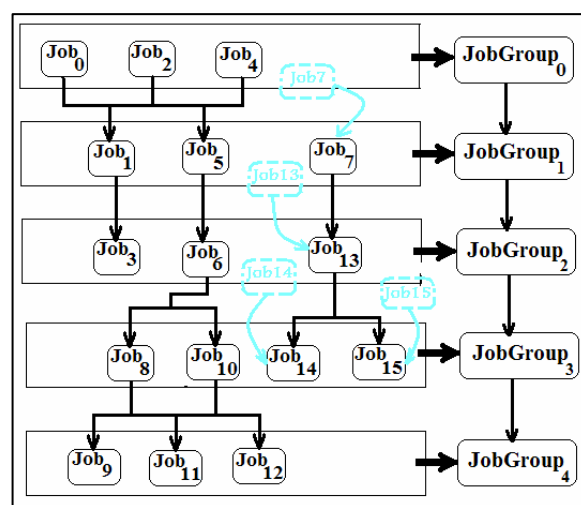
$\Delta_{ij} =$

$$JobGroup_i.NumOfJobs - JobGroup_j.NumOfJobs$$

Δ_{ij} اختلاف تعداد کارهای موجود در گروه کاری i و گروه کاری j را نشان می‌دهد و تابع برازندگی به صورت رابطه ۷ خواهد بود.

$$Fitness\ Function = \sqrt{\sum_{(For_Each_i,j)} (\Delta_{ij})^2} \quad (7)$$

در شکل ۶ می‌توانید نتیجه اجرای الگوریتم ژنتیک را بر روی مثال شکل ۱ و مرحله اول گروه بندی ببینید. همانطور که قابل مشاهده است اجرای الگوریتم ژنتیک منجر به ایجاد یک تعادل باری میان تعداد کارهای موجود در گروههای کاری شده است و در هر گروه کاری به طور متوسط ۳ کار قرار گرفته است و فقط $JobGroup_3$ دارای ۴ کار است. یعنی ۱۶ کار موجود به طور مساوی در بین گروههای کاری تقسیم شده‌اند.



شکل (۶): نتیجه اجرای الگوریتم ژنتیک در مرحله اول

در مرحله دوم هر عضو از جمعیت ژنتیکی، از جایگشت‌های تصادفی کارها به وجود می‌آید و با استفاده از یک تابع برازندگی مناسب، جمعیت‌نخبگان انتخاب می‌شود. به عنوان مثال اگر n کار داشته باشیم اولین عضو از جمعیت ژنتیکی به صورت رابطه ۸ خواهد بود.

$$Population_1, JobList = Job_1 Job_2 Job_3 \dots Job_n \quad (8)$$

اگر سیاست زمانبندی، کمینه نمودن زمان^{۲۴} محاسبه باشد، تابع برازندگی از رابطه ۹ محاسبه می‌گردد.

Cost-based-²⁵

Mutation-²⁶

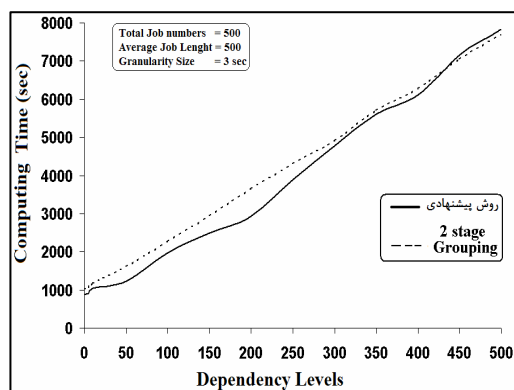
Cross over-²⁷

Trade-off-²⁸

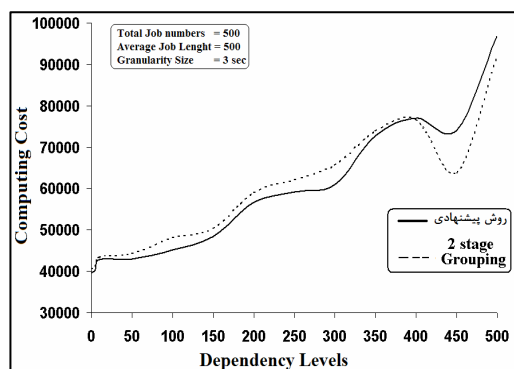
Fitness function-²²

Load Balancing-²³

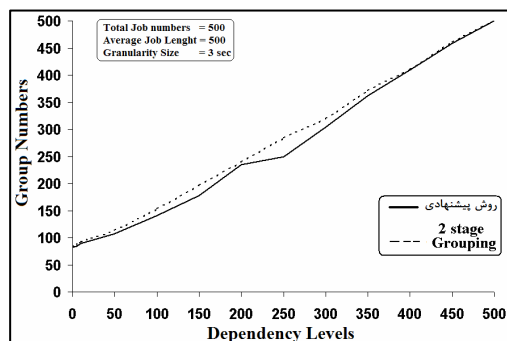
Time-based-²⁴



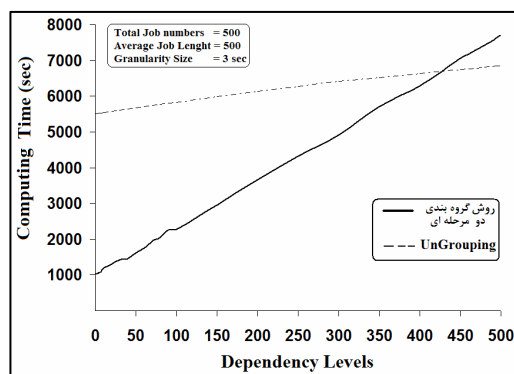
شکل (۹): مقایسه کل زمان محاسبه دو روش



شکل (۱۰): مقایسه هزینه محاسبه دو روش



شکل (۱۱): نمودار تعداد گروه‌ها



شکل (۱۲): مقایسه روش گروه‌بندی با روش Ungrouping

محیط شبیه سازی ثابت و برابر ۵۰۰ در نظر گرفته شده است. پارامتر Granularity Size برابر عدد ۳ و زمان سر بار ارسال و دریافت کارها برابر ۱۰ ثانیه فرض شده است.

Number of Jobs	300	Average Job length	500.0	Genetic Parameters	
Deviate % (random Job length = Average + . D %)	90	Overhead time for Submit (s)	10	Population Size	50
Session cost with Resources	0.0	Granularity time (s)	3	Individual #	20
Scheduling Policy	☒ Time-based ☐ Cost-based ☐ Cost&Time-based		Scheduling Strategy	☐ Grouping ☒ GA + Grouping ☐ Ungrouping ☐ GA + Ungrouping	
☒ "Resource2", 400 MIPS, MAC, LINUX, cost : 110.0 ☒ "Resource1", 1600 MIPS, SOLARIS, UNIX, cost : 220.0 ☒ "Resource3", 2400 MIPS, 48C, RED HAT, cost : 300.0 ☒ "Resource4", 220 MIPS, IBM, WINDOWS, cost : 50.0 ☒ "Resource5", 1000 MIPS, SOLARIS, MANDRAKE, cost : 30.0 ☒ "Resource6", 800 MIPS, MAC, WINDOWS, cost : 200.0 ☒ "Resource7", 1200 MIPS, SOLARIS, LINUX, cost : 170.0				Iteration #	50
				Mutation #	2
				Cross Over #	1
				Cross Over %	50
				Mutatin %	50
				Exploration %	25
				Best individual %	25

Name	Platform (Architecture , OS)	MIPS	Cost
R1	SOLARIS , Unix	1600	220
R2	MAC , Linux	400	110
R3	48C , Redhat	2400	300
R4	IBM , Windows	220	50
R5	SOLARIS , Mandrake	1000	30
R6	MAC , Windows	800	200
R7	SOLARIS , Linux	1200	170

شکل (۷): محیط شبیه سازی و مشخصات منابع

۵- نتایج شبیه سازی

نتایج به دست آمده حاکی از این است که روش پیشنهادی، هزینه و زمان اجرای خیلی بهتری نسبت به روش گروه‌بندی دو مرحله ای دارد. در جدول شکل ۸ جزئی از نتایج حاصل شده از شبیه سازی نشان داده شده است.

Dependency Levels	2-stage Grouping Algorithm			2-stage Grouping+Genetic Algorithm		
	Computing Cost	Computing Time	Groups#	Computing Cost	Computing Time	Groups#
0	40587	1015	84	39513	884.2	82
5	41243	1080	87	40368	912.3	84
10	43388	1182	93	42876	1042	89
50	44223	1616	114	42997	1223	107
100	48169	2280	154	45224	1972	141
150	50390	2951	197	48356	2491	178
200	59068	3656	240	56724	2936	235
250	62147	4315	284	59122	3892	249
300	65700	4917	320	61003	4781	304
350	73884	5704	372	72724	5605	362
400	76686	6286	411	76953	6118	409
450	63887	7057	461	74122	7156	459
499	91709	7702	500	96791	7824	500

شکل (۸): جزئی از نتایج حاصل از شبیه سازی

در شکل ۹، نتایج کل زمان پردازش و محاسبه کارها به زبان نمودار بیان شده است. با توجه به نمودار روش پیشنهادی تا قبل از رسیدن به سطح وابستگی حدود ۹۰٪ بهتر از روش دو مرحله ای می باشد.

مراجع

- [1]. Foster, I. and Kesselman, C. (1999): *The Grid: Blueprint for a New Computing Infrastructure*. San Francisco, Morgan Kaufmann Publisher, Inc.
- [2]. A.L. Rosenberg, *Optimal schedules for cycle-stealing in a network of workstations with a bag-of-tasks workload*, IEEE Trans. Parallel Distribut. Syst. 13 (2) (2002) 179-191.
- [3]. Chuliang Weng, Xinda Lu (2003): *Heuristic scheduling for bag-of-tasks applications in combination with QoS in the computational grid*, Elsevier B.V.
- [4]. Berman, F., Fox, G. and Hey, A. (2003): *Grid Computing- Making the Global Infrastructure a Reality*. London, Wiley.
- [5]. Nithiapidary, M., Junyang, L., Nay, L.S., Srikumar, V., Anthony, S. and Rajkumar Buyya (2005): *A Dynamic Job Grouping-Based Scheduling for Deploying Applications with Fine-Grained Tasks on Global Grids*, GRIDS Laboratory-The University of Melbourne, Australia.
- [6]. Bart, J., Luis, F., Norbert, B., Candice, G., Jean-Yves, G., Roman, S., Seong Yu (2003): *Enabling Applications for Grid Computing with Globus*, ibm.com/redbooks.
- [7]. Yang Gao, Hongqiang Rong, Joshua Zhexue Huang (2004): *Adaptive grid job scheduling with genetic algorithms*, Elsevier B.V.
- [8]. V. Di Martino *, M. Mililotti (2004): *Sub optimal scheduling in a grid using genetic algorithms*, Elsevier B.V.
- [9]. Sathish S. Vadhiyara, and Jack J. Dongarra (2003): *GrADSolve-a grid-based RPC system for parallel computing with Application-Level Scheduling*, Elsevier Inc.
- [10]. Holly Dail, Fran Berman, and Henri Casanova (2002): *A decoupled scheduling approach for Grid application development environments*, Elsevier Science (USA).
- [11]. Buyya, R. and Murshed, M. (2002): *GridSim: A Toolkit for the Modeling, and Simulation of Distributed Resource Management, and Scheduling for Grid Computing*, Journal of Concurrency and Computation: Practice and Experience (CCPE), 14(13-15):1175-1220.
- [12]. Arun, T., Michael, B., Fabiano, L., Huang, R., Linda, L., Paul, M., Jeff, M., Nasser, M., Karthik, S., Olegario, H., Luis, F. (2004): *Grid Services Programming and Application Enablement*, ibm.com/redbooks.

[۱۳]. صادق وهابزاده زرگری، عادل ترکمان رحمانی (۲۰۰۶): یک

روش زمانبندی پویا مبتنی بر گروه بندی کارها با استفاده از الگوریتم

های ژنتیکی برای کاربردهای متشکل از کارهای ریزدانه در گریدهای

محاسباتی، یازدهمین کنفرانس انجمن کامپیوتر ایران

(CSICC)، تهران.

[۱۴]. صادق وهابزاده زرگری، عادل ترکمان رحمانی (۲۰۰۶): یک

روش زمانبندی پویا مبتنی بر گروه بندی دو مرحله ای کارها برای

کاربردهای متشکل از کارهای وابسته به هم در گریدهای محاسباتی،

یازدهمین کنفرانس انجمن کامپیوتر ایران (CSICC)، تهران^۱.

در نمودار شکل ۱۰، نتیجه ارزیابی هزینه محاسبه کارها به منعکس شده است. با توجه به نمودار، نمودار هزینه روش پیشنهادی در بیشتر موارد زیر نمودار هزینه روش گروه بندی دو مرحله ای و بهتر از آن می باشد. در ضمن با افزایش تعداد سطوح وابستگی، نمودار هزینه دو روش به سمت یکدیگر همگرا شده است که آن را به نحوی می توان تایید نتیجه حاصل از نمودار شکل ۹ دانست.

با توجه به نمودار شکل ۱۱، تعداد گروه ها در صورت اعمال الگوریتم ژنتیک کاهش می یابد. در ضمن تعداد گروه های حاصل، متناسب با افزایش سطوح وابستگی و نزدیک شدن آن به حد نهایی که برابر تعداد کل کارها می باشد، افزایش می یابد و در نهایت برابر تعداد کل کارها می شود.

در شکل ۱۲، لزوم گروه بندی به زبان نمودار بیان شده است. با توجه به کارایی بهتر اعمال الگوریتم ژنتیک و با توجه به نمودار شکل ۱۲ می توان نتیجه گرفت که روش پیشنهادی نسبت به روش Ungrouping نیز از کارایی به مراتب بهتری برخوردار می باشد.

۶- نتیجه گیری و کارهای آتی

با گسترش روزافزون گریدها، مخصوصاً گریدهای محاسباتی، مدیریت بهینه منابع و بهره وری از توان محاسباتی آنها از چالش های مهم و قابل بحث می باشد. با افزایش کاربردهای حجیم تشکیل شده از وظیفه های وابسته به هم و چالش اجرای آنها در گریدهای محاسباتی، زمانبندی صحیح و پویای این کاربردها، تاثیر قابل توجهی در مدیریت بهینه منابع دارد. در روش ارائه شده در زمانبندی تطبیقی، به ماهیت واقعی بسیاری از موجودیت ها توجه نشده است. روش گروه بندی کارها با آنکه توانسته است به میزان قابل اعتنايي هزینه و زمان پردازش کارها را کاهش دهد.

روش پیشنهادی، با اعمال الگوریتم ژنتیک بر روی روش گروه بندی دو مرحله ای آن را تا حد قابل قبولی بهبود داده است. نتایج شبیه سازی حاکی از کارایی بهتر روش پیشنهادی نسبت به روش گروه بندی دو مرحله ای دارد. با توجه به شکل ۹ و نتایج به دست آمده، ۱۰/۵ درصد بهبود در زمان محاسبه کارها حاصل شده است که این مقدار در مورد کاربردهایی که سطح وابستگی آنها زیر ۹۰٪ است به ۳۶/۸ درصد می رسد و از آنجاییکه معمولاً سطح وابستگی کاربردها به ندرت به ۹۰٪ می رسد، بهبود به دست آمده قابل اعتنا می باشد. در ضمن با توجه به شکل ۱۰ و نتایج حاصل از آن ۱/۴ درصد بهبود در هزینه محاسبه کارها حاصل شده است که این مقدار در مورد کاربردهایی که سطح وابستگی آنها زیر ۹۰٪ است به ۱۳/۳ درصد می رسد.

پیشنهاد می شود که در آینده بر روی زمانبندهای نامتمرکز و زمانبندهای مبتنی بر مدل های پیش بینی کارایی، کارهای بیشتری صورت پذیرد.