

Примеры вопросов из теста “Суперкомпьютеры” (2016)

Давалось 20 вопросов на 30 минут

Пользоваться можно чем угодно

В системе sigma пропускать вопросы, чтобы потом вернуться - нельзя (просто пропустить вопрос тоже нельзя)

Итоговая оценка складывается из 2-х оценок за праки и оценки за тест, но если тест на 2, то и итоговая 2, округление в пользу студента, но с учётом статистики посещения лома и джинни.

**Господа, может обсудим ответы на эти вопросы?
старайтесь выбрать свой фоновый цвет, чтобы не сливаться
друг с другом**

при возможности оставляйте обоснование своему ответу

Ломоносов и BG, архитектура компьютеров:

- 1) Какой будет архитектура большинства вновь создаваемых суперкомпьютеров?
(гибридной, это модно и интересно/как блюджин/еще какой-то вариант)
гибридной
- 2) сколько соседей у узлов BG
6
- 3) Где в блюджин используется топология/информационная сеть “дерево”
вроде как при коллективных операциях MPI
- 4) у Blue Gene максимально 4 потока на узле? да, см. сл вопрос
- 5) Сколько потоков в процессоре, который стоит в блюджине (*видимо речь об узле?*)
4
- 6) Blue Gene: при каком методе запуска доступно больше всего памяти:
 - SMP
 - Dual
 - VN
 - всем одинаково

SMP? если имеется в виду “доступной одному процессу”, то SMP

Ну да, в этом главный вопрос

<http://hpc.cmc.msu.ru/bgp/obs/modes> - цитата:

В каждом из режимов MPI-процессам доступен приблизительно следующий объем памяти:

VN — 472 МБ

DUAL — 978 МБ

SMP — 1992 МБ

это на процесс. Если запускать с числом узлов N при разном кол-ве процессов, то у тебя на каждом узле 2гб памяти и будет одинаково всегда. Так что, грубо говоря, вопрос фиксируешь ты кол-во процессов или узлов при запуске в разных режимах. Либо вопрос вообще про память на 1 процесс и тогда то, что выделено зеленым верно.

7) был вопрос про архитектуру Ломоносова

ВОПРОС №17
Максимальный объем доступной памяти достигается при запуске параллельных программ на Blue Gene/P в режиме:
☐ не зависит от режима запуска
☐ DUAL
☐ SMP
☐ VN
Ответить

ВОПРОС №11
Коммуникационная сеть «дерево» Blue Gene/P используется для:
☐ для доступа к файловой системе
☐ выполнения 2-ух точечных передач большого объема
☐ выполнения коллективных операций MPI при соблюдении определенных условий
☐ выполнения любых коллективных операций MPI
Ответить

3? или 4? "Используется для коллективных операций и коммуникатора WORLD", так что скорее 3

ВОПРОС №16
С каким количеством процессоров непосредственно связан каждый вычислительный узел Blue Gene/ коммуникационной сетью «решетка» :
☐ 6
☐ 16
☐ 8
☐ 128
☐ 4

6. Ну тор, все дела

ВОПРОС №19

Почему явные итерационные методы успешны для решения задачи на собственные значения на суперкомпьютере BlueGene/P?

- ☐ Неявные итерационные методы показывают более низкую точность вычислений.
- ☐ Существуют эффективные методы распараллеливания перемножения матрицы на вектор с учетом параллельного быстрого преобразования Фурье.
- ☐ Для прямых методов нельзя эффективно использовать параллельное быстрое преобразование Фурье.

2?

ВОПРОС №2

Выберите верные (-ое) утверждения (-е):

- ☐ С/к "Ломоносов" имеет гибридную архитектуру
- ☐ Вычислительные узлы IBM Blue Gene/P имеют двусторонние связи с шестью соседями
- ☐ IBM Blue Gene/P позволяет запускать программы с использованием GPU
- ☐ Домашний каталог каждого пользователя с/к "Ломоносов" виден на любом вычислительном узле

Ответить

1,2. 4? я думаю без 4, т.к. я бы считал домашним в этом вопросе домашний на access без 4 вроде, потому что на parallel.ru есть такое:

1. Быстрое хранилище (tier 1) – предназначено для проведения расчетов.
2. Основное хранилище (tier 2) – предназначено для хранения рабочих данных пользователя (например данные проекта над которым пользователь работает в данный момент)
3. Хранилище архивных данных (tier 3) – предназначено для хранения данных, которые в данный момент пользователю не нужны, но понадобятся в будущем, и хранения архивов данных.

Домашняя директория пользователя (/home/users/\$user) расположена на быстром хранилище (tier 1).

Важно: доступ с вычислительных узлов на основное хранилище (tier 2) или хранилище архивных данных (tier 3) невозможен. Тоже думаю, что без 4.

Ок: 1,2

ВОПРОС №15

Максимальное число потоков, которые можно запустить на одном узле Blue Gene/P:

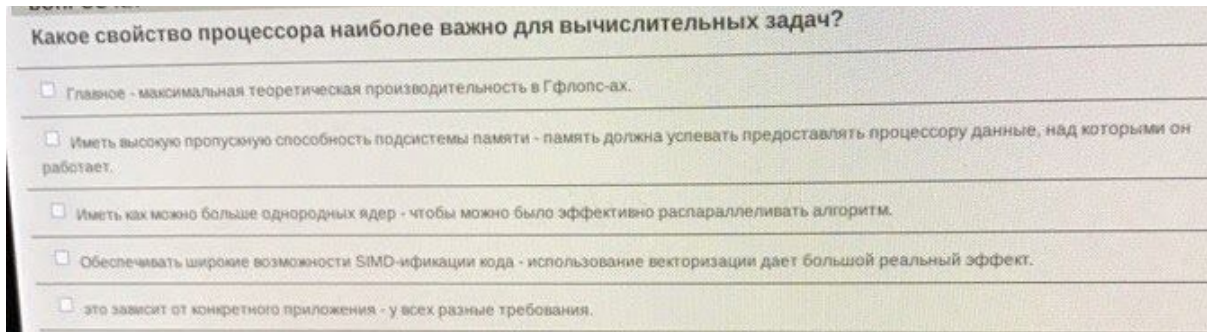
- ☐ 16
- ☐ 8
- ☐ 4
- ☐ Число потоков не ограничено
- ☐ 2
- ☐ 32

6 4

Да-да, ошибочка, 4 То есть число потоков не ограничено?

Имеется в виду 4 потока, а вариант ответа 3

Спасибо



5

Ресурс параллелизма, сложность алгоритмов:

1) Сложность алгоритма перемножения плотных прямоугольных матриц? - ответ $O(N^3)$

2) Вычислительная сложность перемножения квадратных плотных матриц: $O(n)$, $O(NN)$, $O(nnn)$, нет правильного ответа (у меня была вычислительная Мощность)

мощность - $O(N)$?

Сложность плотных - $O(N^3)$

Это сложность такая, мощность = сложность/объем входных-выходных данных

А тыртышников говорил, что матрицы за $N^{\log_2(7)}$ перемножаются...

Ну в теории они вообще за N^2 перемножаются, но на практике получается

$N^{(2.38)^{}}$:)

Ф Л У Д И Л К А

че вы хотите от поповой. она может еще мощность со сложностью сама перепутала

автограф в зачетке и ведомости хотим. тогда обязательно посмотри в конце что жирным шрифтом выделено в параметрах sbatch Affinity/Multi-core options? В конце этого гугл дока, внимательней!!! Подумой!!! ох петросянчик ЖЖЖ((

интересно, каковы шансы, что вопросы будут те же самые?

50/50, раз не знаем, то энтропия максимальна, значит распределение равномерное. Это не равномерное распределение, равномерное - непрерывное, а это дискретное

ну эти наверное не удалят, но добавят новые

всем удачи

спасибо

3) вид параллелизма в двойном цикле

- конечный
- координатный
- скошенный
- нет верного ответа

ВОПРОС №7

Каким параллелизмом обладает фрагмент программы:

```
for(i=1; i<=n; ++i)
for(j=1; j<=m; ++j)
A[i][j] = (A[i-1][j] * A[i][j-1])/2;
```

- ☐ не обладает
- ☐ другим
- ☐ скошенным
- ☐ конечным
- ☐ координатным

скошенный?

+++

ВОПРОС №16

Каков порядок вычислительной мощности алгоритма перемножения плотных квадратных матриц?

- ☐ $O(n^2)$
- ☐ правильного ответа нет
- ☐ $O(1)$
- ☐ $O(n^3)$
- ☐ $O(n)$

$O(n)^+$

ВОПРОС №14

Какое свойство процессора наиболее важно для вычислительных задач?

- ☐ Обеспечивать широкие возможности SIMD-ификации кода - использование векторизации дает большой реальный эффект.
- ☐ Главное - максимальная теоретическая производительность в Гфлопс-ах.
- ☐ Иметь высокую пропускную способность подсистемы памяти - память должна успевать предоставлять процессору данные, над которыми он работает.
- ☐ Иметь как можно больше однородных ядер - чтобы можно было эффективно распараллеливать алгоритм.
- ☐ это зависит от конкретного приложения - у всех разные требования.

5+

Сопряженные градиенты и скорейший спуск, разбиение сетки:

- 1) название метода и типа что это за метод, например метод скорейшего спуска
- 2) является ли метод сопряженных градиент обобщением метода скор спуска?

ВОПРОС №4

Краевая задача:

$$\begin{cases} -\frac{\partial u}{\partial x^2} - \frac{\partial u}{\partial y^2} = f(x, y), & (x, y) \in D, \\ u(x, y) = \varphi(x, y), & (x, y) \in \partial D \end{cases}$$

- ☐ является смешанной краевой задачей для уравнения параболического типа
- ☐ является задачей Дирихле для уравнения Пуассона
- ☐ является задачей Трикоми для уравнения Бицадзе-Самарского
- ☐ является задачей Неймана для уравнения Лапласа

Ответить

По идее 2, но тут немного не такой оператор лапласа.

Это оператор лапласа со знаком минус, и 2 верно. Я про то, что там сверху не du , d^2 и должно быть, не? Хм, да, действительно, но больше похоже на их опечатку (составителей теста/книги), чем на то, что так было задумано. Ну просто это скрин же откуда-то (из книги/статьи). Варианта “не является задачей”, увы, не предусмотрено. Ну ок, остановимся на 2.

Указать **неверное** утверждение. Метод сопряженных градиентов

- ☐ является итерационным методом проекционного типа, предназначенным для решения системы линейных алгебраических уравнений (СЛАУ) с симметричной положительно определенной матрицей
- ☐ является обобщением алгоритма метода скорейшего спуска
- ☐ сходится к точному решению любой системы линейных алгебраических уравнений с квадратной невырожденной матрицей
- ☐ позволяет получить точное решение симметричной положительно определенной СЛАУ за конечное количество итераций при условии, что все действия совершаются без ошибок округления

3? Вроде бы матрица должны быть симметричной и положительно определенной, а не только невырожденной. Может 1?

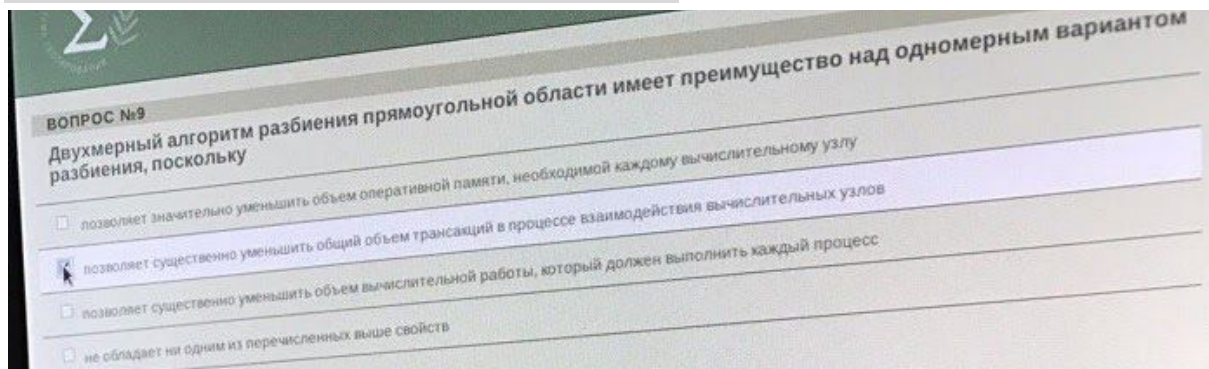
Судя по всему на последующий 2-х изображениях один и тот же вопрос, то есть “перечисленных выше” значит “перечисленных и выше и ниже”

ВОПРОС №6

Двухмерный алгоритм разбиения прямоугольной области имеет преимущество над одномерным вариантом разбиения, поскольку

- ☐ позволяет существенно уменьшить общий объем транзакций в процессе взаимодействия вычислительных узлов
- ☐ позволяет значительно уменьшить объем оперативной памяти, необходимой каждому вычислительному узлу
- ☐ не обладает ни одним из перечисленных выше свойств
- ☐ позволяет существенно уменьшить объем вычислительной работы, который должен выполнить каждый процесс

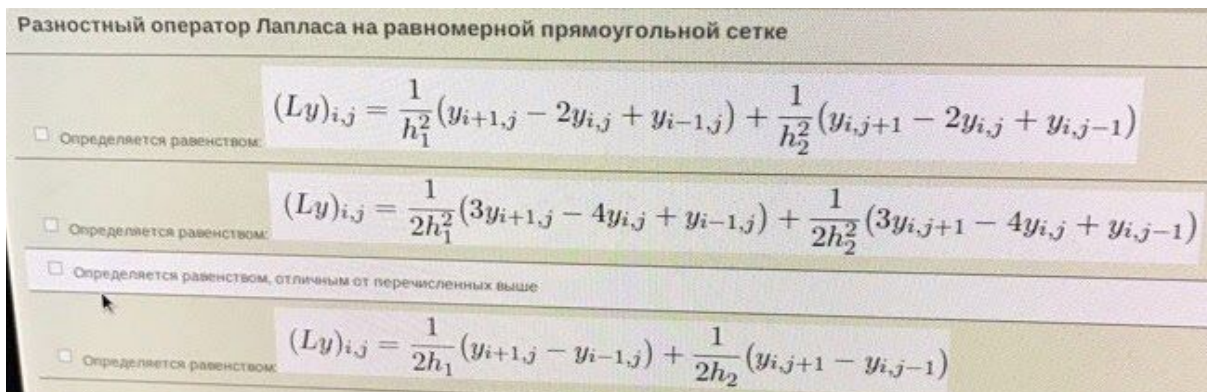
^ 1 (так как обмениваться нужно границами, а тут граница совпадает со всей подобластью [линия], в итоге вся сетка гоняется)



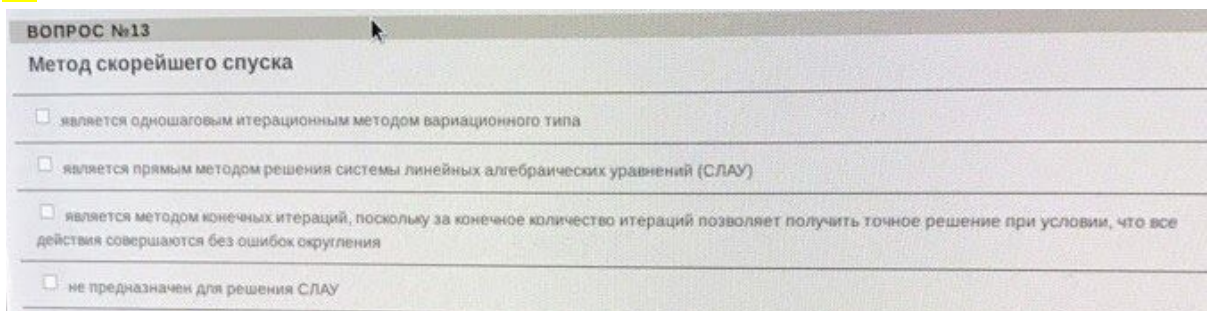
1? скорее 2

А, перепутал. Я тоже за 2. Вот тут

<https://я/drive.google.com/drive/folders/0B0X-oQW4pjUUd2htak5UUWFsemc> на второй картинке об этом.



1? вроде да да



1? вроде да

CUDA:

- 1) Про `cuda kernel<<< >>>` - найти верные сигнатуры запуска, какие будут параметры запуска
 - 2) Для каких задач используются графические ускорители:
 - обработка видео
 - обработка изображений
 - общие вычисления
 - обращения к файловой системе
 - программирование рекурсивных функций
 - 3) Укажите верные утверждения:
 - cuda - расширение c/c++ +
 - cuda специально для nvidia sad but true, vendor lock
 - cuda это расширение Fortran

cuda теперь вообще поверх LLVM идёт, какой хочешь фронтенд теперь пиши, хоть NodeJS
 - 4) CUDA только для nvidia ГПУ AMD вроде всё держится за OpenCL
- Upd: не все GPU, proof: <https://developer.nvidia.com/cuda-gpus> , тут говорят только про терпу и некоторые другие - чисто может кому интересно
- 5) Какие опции команды `sbatch` позволяют ограничить количество выделенных GPU-карт при выбранном определенном количестве узлов?
 - `gpu 1`
 - `prn 1`

- -s gpu 0
- используются все GPU
- все GPU? ++
- вроде да

6) Фрагмент кода:

```
строка1- cudaMemcpyAsync(arr1, arr2, count, cudaMemcpyHostToDevice, st1);
строка2- kernel<<count / 256, 256, 0, st2 >>(arr1, arr3, count);
строка3- cudaMemcpyAsync(arr2, arr1, count, cudaMemcpyDeviceToHost, st1);
```

- могут выполняться параллельно строки 1, 3 и строка 2?
- строка 3 выполняется после строки 1

7) Про cudaMemcpyAsync: откуда куда пересылка данных; указать правильные утверждения про вызов cudaMemcpy(arr1, arr2, count, cudaMemcpyHostToDevice)

На последующих 2-х изображениях есть разница в коде. Вообще у Колганова много похожих, но немного отличающихся вопросов.

Выберите все верные утверждения относительно следующего кода, при условии, что st1, st2 отличны от потока по умолчанию, а ядро меняет массив arr1:

```
строка1- cudaMemcpy (arr1, arr2, count, cudaMemcpyHostToDevice);
строка2- kernel<<count / 256, 256, 0, st2 >>>(arr1, arr3, count);
строка3- cudaMemcpyAsync(arr2, arr1, count, cudaMemcpyDeviceToHost, st1);
```

- ☐ ядро выполнится параллельно с копированием в строке 3
- ☐ копирование в строке 3 выполнится только после копирования в строке 1
- ☐ ядро выполнится только после завершения копирования в строке 1
- ☐ Значения элементов массива arr2 после завершения функции в строке 3 будут совпадать со значениями элементов массива arr1, полученные после завершения функции в строке 1
- ☐ Значения элементов массива arr2 после завершения функции в строке 3 будут совпадать со значениями элементов массива arr1, полученные после завершения функции в строке 2

9

1,4? Или без 4?

вроде без 4, потому что одновременно выполняется ядро, где меняется arr1, и копируется arr1 в arr2.

без 4 5. Так как 2 и 3 строки запущены в разных потоках, а kernel мутирует arr1. Поэтому в третьей строке может быть скопирована рандомная белиберда.

2,3 почему не подходят? первая строка же вроде бы синхронная операция, которая блокирует все потоки?

Да, точно, чёт думал там Async

В итоге 1,2,3

Выберете все верные утверждения относительно следующего кода, при условии, что st1, st2 отличны от потока по умолчанию, а ядро меняет массив arr1:

```

строка1- cudaMemcpyAsync(arr1, arr2, count, cudaMemcpyHostToDevice, st1);
строка2- kernel<<<count / 256, 256, 0, st2 >>>(arr1, arr3, count);
строка3- cudaMemcpyAsync(arr2, arr1, count, cudaMemcpyDeviceToHost, st1);

```

- ☐ Значения элементов массива arr2 после завершения функции в строке3 будут совпадать со значениями элементов массива arr1 , полученные после завершения функции в строке2
- ☐ ядро выполнится параллельно с копированиями в строке1 и строке3
- ☐ ядро выполнится только после завершения копирования в строке1
- ☐ копирование в строке3 выполнится только после копирования в строке1
- ☐ Значения элементов массива arr2 после завершения функции в строке3 будут совпадать со значениями элементов массива arr1 , полученные после завершения функции в строке1

2,4? + верно второе, там поток st1 != st2 **Понял, спасибо**

Отметьте все верные утверждения про данный запуск ядра:
 Kernel<<<512, dim3(32, 4), 0, 0>>>();

- ☒ ядро будет использовать 128 нитей
- ☐ Запуск ядра выполнится асинхронно
- ☐ ядро будет использовать 256 * 256 нитей
- ☒ ядро будет запущено в потоке по умолчанию
- ☐ Запуск ядра выполнится синхронно
- ☐ ядро будет использовать 512 нитей

2,3,4?+

В основном графические ускорители применяются для:

- ☒ обработки видео
- ☒ обработки изображений
- ☐ написания рекурсивных алгоритмов
- ☐ файлового ввода вывода
- ☐ работы с базами данных
- ☐ вычислений общего назначения

-для игр гонять нейронки, зарабатывать на кеггле резюме и призовые фонды
 лол,много уже на кэгле заработал(призовых фондов)?

а для общего назначения? матрицы ж на нём тоже неплохо множатся

просто тут не сказано типа для того, что хорошо параллелится. Матрицы параллелятся,но есть типа много чего, что на проце гораздо быстрее из-за кешей (и не параллелится).

Короче в любом случае под вопросом это

Отметьте все неверные конфигурации запуска ядра:

- ☐ kernel<<<75, 0, 0, 0>>>()
- ☐ kernel<<< dim3(5, 55, 1, 1), dim3(32, 4), 0, 0>>>()
- ☒ kernel<<< dim3(-11), dim3(1), 0>>>()
- ☐ kernel<<<1, 1024>>>()
- ☐ kernel<<< dim3(13, 55, 1), dim3(32, 5, 4), 0>>>()

2,3 Почему 2? Конструктор dim3(5,55,1,1) скорее всего упадет, потому что 4 параметра. Да, по этой логике отметил

1 - Можно ли запускать с нулем нитей? Просто не запуснуться тогда?

Хрен его знает, cuda host api толком нигде не прописан, nvidia должна была закрыться со стыда от такой хуйни

Какие опции команды sbatch позволяют ограничить количество выделенных GPU-карт при выбранном определенном количестве узлов?

- ☐ никакими, нам будут доступны все GPU-карты, выделенных задаче узлов.
- ☐ -gpus 1
- ☐ -s gpus 0
- ☐ -p gpus 1

1?

Отметьте все верные конфигурации запуска ядра:

- ☐ kernel<<< dim3(5, 1, 1), dim3(32, 4), 0, 0>>>()
- ☐ kernel<<< dim3(5, 78), dim3(2,5,11), 0>>>()
- ☐ kernel<<<16, dim3(1, 50, 80), 0>>>()
- ☐ kernel<<<1, 1, 5, -1>>>()
- ☐ kernel<<<0, 45, 0, 0>>>()

1,2,3 +

Опять же, можно ли как с 0 блоками запустить (как в 5 варианте)?

ВОПРОС №1

Отметьте все верные факты про технологию CUDA:

- ☐ CUDA является расширением стандартного языка Fortran
- ☐ CUDA является новым языком программирования на базе C/C++
- ☒ Это программно-аппаратная архитектура параллельных вычислений, позволяющая значительно ускорить код с использованием любых GPU.
- ☐ CUDA является расширением стандартных языков Java/C#/Python
- ☐ CUDA является расширением стандартных языков C/C++
- ☒ Это программно-аппаратная архитектура параллельных вычислений, позволяющая значительно ускорить код с использованием GPU NVidia

3,5,6? CUDA есть на Fortran, C/C++ и является технологией NVidia То есть 1,5,6 или как? 1,5,6

В лекциях не было про фортран, кто как думает надо ответить? CUDA вообще на LLVM есть, пиши фронтенд и будет тебе счастье <https://developer.nvidia.com/cuda-llvm-compiler>

Можно по человечески?

ВОПРОС №8

Отметьте все верные факты про вызов данной функции `cudaMemcpyAsync(array2, array1, count, cudaMemcpyHostToDevice)`:

- ☐ Копируется count байт
- ☐ Операция асинхронна, выполняется в потоке stream
- ☐ Происходит копирование данных с GPU на CPU
- ☐ Происходит копирование данных с GPU на GPU
- ☒ Происходит копирование данных из массива array1 в массив array2
- ☐ Происходит копирование данных с CPU на GPU
- ☐ Операция асинхронна, выполняется в потоке по умолчанию
- ☐ Происходит копирование данных из массива array2 в массив array1

1,5,6,7+

Турбулентность:

- 1) Что такое число Рейнольдса (отношение вязкости к инерции как-то так) --- или инерции к вязкости
- 2) Причина возникновения турбулентности
- 3) коэффициент рейнольдса? - вязкость и инерция
- 4) Какие уравнения используются для описания осредненных характеристик турбулентных течений? навье-стокс? Рейнолдса?

Укажите правильный (правильные) ответ(ответы):

- ☐ Метод прямого численного моделирования (DNS) турбулентности разрешает только небольшие пространственные масштабы
- ☐ Метод прямого численного моделирования (DNS) разрешает все возможные пространственные масштабы
- ☐ LES моделирует все возможные пространственные масштабы
- ☐ LES моделирует только достаточно крупные вихревые структуры

2,4?

Вроде да

Что является причиной возникновения турбулентности?

- ☐ Гидродинамические неустойчивости
- ☐ Политическая нестабильность на ближнем востоке
- ☐ Случайные внешние силы

3?1

И то , и то мб? скорее всего и то, и то. На вики написано про 3, вот тут (<http://old.icad.org.ru/docs/1.pdf>) нашла про 1.

вопрос №12

Какие уравнения используются для описания осредненных характеристик турбулентных течений?

- ☐ Уравнения Рейнольдса
- ☐ Уравнения Навье-Стокса
- ☐ Уравнения Эйлера

1? По вики, Уравнения Рейнольдса (англ. *RANS (Reynolds-averaged Navier–Stokes)*) — уравнения Навье — Стокса (уравнения движения вязкой жидкости), осреднённые по Рейнольдсу.

Квантовая (молекулярная) динамика:

- 1) Атомы и электроны, квантовые коды молекулярной динамики, указать соотношение между величинами M (число электронов) и N (число коэффициентов быстрого преобразования Фурье)
 - $\sim N$
 - $\sim N^4$
 - $\sim N^3$
- 2) в каком методе в квантовой динамике скрыто главное ограничение масштабируемости? матрица перекрытия, фурье или перемножение матрицы на вектор?

В каком численном методе, используемом в коде квантовой молекулярной динамики, скрыто главное ограничение масштабируемости?

- ☐ Быстрое преобразование Фурье
- ☐ Перемножение матрицы на вектор
- ☐ Вычисление матрицы перекрытия

1?точно 1, поскольку в пересдаче был вопрос “почему бфп - главное ограничение масштабируемости на bluegene”

Еще мнения?

Какова масштабируемость кодов квантовой молекулярной динамики на основе теории функционала плотности в зависимости от количества атомов?

Если считать, что N - количество электронов, а M - количество базисных векторов в разложении Фурье, и M пропорционально N .

☐ N

☐ N^3

☐ N^4

Говорят, что N , но без обоснований

Это что вообще за хуйня?

Странно, что наверное масштабируемость должна расти при увеличении кол-ва электронов и при большем кол-ве базисных векторов. Как-то логически кажется, что должно быть от N^2 и больше

Был еще какой-то вопрос на локальность данных, которого здесь нет.

О чем примерно?

У меня он описан как "локальность по времени? по данным?" Конкретнее, увы, не могу сказать.

Это типа help по sbatch

Parallel run options:

- A, --account=name charge job to specified account
- begin=time defer job until HH:MM MM/DD/YY
- c, --cpus-per-task=ncpus number of cpus required per task
- comment=name arbitrary comment
- d, --dependency=type:jobid defer job until condition on jobid is satisfied
- D, --workdir=directory set working directory for batch script
- e, --error=err file for batch script's standard error
- export[=names] specify environment variables to export
- export-file=file|fd specify environment variables file or file descriptor to export

- get-user-env load environment from local cluster
- gid=group_id group ID to run job as (user root only)
- gres=list required generic resources
- H, --hold submit job in held state
- i, --input=in file for batch script's standard input
- l, --immediate exit if resources are not immediately available
- jobid=id run under already allocated job
- J, --job-name=jobname name of job
- k, --no-kill do not kill job on node failure
- L, --licenses=names required license, comma separated
- P, --popova start nedikvatniy course and add gorelye gopi to students**
- m, --distribution=type distribution method for processes to nodes
(type = block|cyclic|arbitrary)
- M, --clusters=names Comma separated list of clusters to issue commands to. Default is current cluster.

Name of 'all' will submit to run on all clusters.

- mail-type=type notify on state change: BEGIN, END, FAIL or ALL
 - mail-user=user who to send email notification for job state changes
 - n, --ntasks=ntasks number of tasks to run
 - nice[=value] decrease scheduling priority by value
 - no-requeue if set, do not permit the job to be requeued
 - ntasks-per-node=n number of tasks to invoke on each node
 - N, --nodes=N number of nodes on which to run (N = min[-max])
 - o, --output=out file for batch script's standard output
 - O, --overcommit overcommit resources
 - p, --partition=partition partition requested
 - propagate[=rlimits] propagate all [or specific list of] rlimits
 - qos=qos quality of service
 - Q, --quiet quiet mode (suppress informational messages)
 - requeue if set, permit the job to be requeued
 - t, --time=minutes time limit
 - time-min=minutes minimum time limit (if distinct)
 - s, --share share nodes with other jobs
 - uid=user_id user ID to run job as (user root only)
 - v, --verbose verbose mode (multiple -v's increase verbosity)
 - wrap[=command string] wrap command string in a sh script and submit
 - switches=max-switches{@max-time-to-wait}
- Optimum switches and max time to wait for optimum

Constraint options:

- contiguous demand a contiguous range of nodes
- C, --constraint=list specify a list of constraints
- F, --nodefile=filename request a specific list of hosts
- mem=MB minimum amount of real memory
- mincpus=n minimum number of logical processors (threads) per node
- reservation=name allocate resources from named reservation
- tmp=MB minimum amount of temporary disk
- w, --odelist=hosts... request a specific list of hosts
- x, --exclude=hosts... exclude a specific list of hosts

Consumable resources related options:

- exclusive allocate nodes in exclusive mode when
 cpu consumable resource is enabled
- mem-per-cpu=MB maximum amount of real memory per allocated
 cpu required by the job.
- mem >= --mem-per-cpu if --mem is specified.

Affinity/Multi-core options: (when the task/affinity plugin is enabled)

- B --extra-node-info=S[:C[:T]] Expands to:

--sockets-per-node=S number of sockets per node to allocate

--cores-per-socket=C number of cores per socket to allocate

--threads-per-core=T number of threads per core to allocate

each field can be 'min' or wildcard '*'

total cpus requested = (N x S x C x T)

--ntasks-per-core=n number of tasks to invoke on each core

--ntasks-per-socket=n number of tasks to invoke on each socket

Help options:

222 -h, --help show this help message

-u, --usage display brief usage message

Other options:

-V, --version output version information and exit

-P, --popova start neadikvatniy course and add gorelye gopi to students,+1