

УДК 004.91

Я. Зьола, О. Назаревич

(Тернопільський державний технічний університет імені Івана Пулюя)

ДЕЯКІ АСПЕКТИ ВИБОРУ ПРОГРАМНИХ ЗАСОБІВ ПРИ ПОБУДОВІ ОБЧИСЛЮВАЛЬНОГО КЛАСТЕРУ ДЛЯ НАУКОВИХ ДОСЛІДЖЕНЬ НА БАЗІ ТИПОВОГО КОМП'ЮТЕРНОГО КЛАСУ

Наукові дослідження часто ставлять перед колективом дослідників потребу виконання трудомістких обчислень, які на "звичайному" комп'ютері можуть виконуватись багато днів, місяців, а то й років, що не сприяє ефективному використанню часу, зусиль та наявних обчислювальних ресурсів. Виходом з такої ситуації є застосування механізму паралельних обчислень. В загальному значенні термін паралельні обчислення - це процес обробки даних, при якому одночасно можуть виконуватися більше, ніж одна машинна операція.

Майже від самого моменту виникнення комп'ютерів цей напрям їх розвитку став надзвичайно важливим, а в багатьох випадках - визначальним. Впровадження паралелізму часто визначало особливості архітектури як процесорів (наприклад, впровадження математичного сопроцесора, SIMD-розширень тощо), так і комп'ютера в цілому (прикладом паралелізму можуть бути навіть звичні всім плати розширення). Це яскраво демонструє сам факт появи та високий рівень популярності багатоядерних процесорів. Окремим випадком паралельних обчислень є розподілені обчислення - спосіб вирішення трудомістких обчислювальних завдань з використанням мережі комп'ютерів. У свою чергу, такою спеціалізованою мережею є кластер.

Мета дослідження - вибір і обґрунтування технології кластеризації, що дозволить організувати обчислювальний HPC (високої продуктивності) кластер на базі типового комп'ютерного навчального класу (7-12 машин) із дотриманням таких додаткових умов: використання спільної пам'яті, спільне використання наявних дискових ресурсів та прозорість для користувача завдяки створенню ілюзії єдиної системи SSI ("суперкомп'ютера") із багатьма терміналами для роботи, роль яких виконують робочі станції.

Зазвичай, паралельна обчислювальна система функціонує за рахунок чіткої координації дій (реалізація математичної моделі на рівні протоколу та вищих рівнях: low-, middle- та high levels) усіх її складових - розподілення самих обчислень, узгодження та перевірка на надійність і достовірність отриманих результатів, менеджмент доступних ресурсів тощо. Такий підхід у свою чергу ставить перед науковим співтовариством проблему масштабованості паралельних систем та робить її однією з найважливіших вимог поруч із такими характеристиками, як надійність, співвідношення ціна/якість, сумісність та мобільність програмного забезпечення тощо. Масштабованість передбачає можливість нарощування обчислювальних потужностей за рахунок збільшення кількості процесорів, об'єму пам'яті, також важливою є рівень незалежності пропускну здатності комунікаційної підсистеми від кількості процесорних вузлів, задіяних у обчисленнях.

Отже, звичайний комп'ютерний клас ми можемо об'єднати найпростішим способом і організувати спільні математичні обчислення це Octave+MPI на рівні транспортного протоколу TCP/IP (наприклад PelicanHPC liveCD). Недоліком є великі накладні витрати та неможливість міграції паралельних процесів, а також складність розподілення навантаження. Більш прогресивний підхід - використання grid-клієнта, наприклад, boinc (проект@home) та використання разом із CPU потужностей GPU nVidia (починаючи з версії 6.4) як математичного сопроцесора (на основі API CUDA). Більш професійним підходом є рух у напрямку систем з SSI, Single Root, shared memory, кластерних файлових систем та організації на базі комп'ютерного класу "ілюзії супер комп'ютера" з багатьма терміналами. Робочі станції уподібнюються до "тонкого клієнта", а всі їхні локальні потужності CPU&GPU, RAM, HDD на 90% задіяні для побудови такого обчислювального кластера (супер комп'ютера) для наукових обчислень із можливістю балансування обчислювального навантаження в межах комп'ютерного класу.