

Największy na świecie system przechowywania danych

2 września 2011

IBM pracuje nad najbardziej pojemnym magazynem danych w historii. W laboratoriach w Almaden powstaje 120-petabajtowy system przechowywania danych. Będzie się on składał z 200 000 wspólnie pracujących dysków twardych. W sumie ma on pomieścić biliard plików, a jego twórcy mają nadzieję, że wspomogą symulowanie tak złożonych systemów jak klimat i pogoda. Same informacje dotyczące położenia poszczególnych plików, ich nazw oraz atrybutów zajmą około 2 petabajtów.

Magazyn danych powstaje na potrzeby jednego z klientów IBM-a, który zamówił też nowy superkomputer do symulowania procesów zachodzących w przyrodzie. Bruce Hillsberg, dyrektor ds. badań nad systemami przechowywania danych, który odpowiada za powyższy projekt mówi, że doświadczenia zdobyte podczas tworzenia takiego systemu przydadzą się do opracowania podobnych komercyjnych magazynów danych. Jego zdaniem w ciągu najbliższych kilku lat firmy oferujące chmury obliczeniowe zaczną składać zamówienia na podobne systemy przechowywania danych.

Inżynierowie IBM-a mają do wykonania bardzo ambitne zadanie. Obecnie największe systemy przechowywania danych liczą sobie około 15 petabajtów.

Na potrzeby obecnego zamówienia IBM opracował nowy sprzęt i oprogramowanie. Wiadomo, że całość będzie chłodzona wodą, a inżynierowie zastanawiają się, w jaki sposób umieścić dyski tak, by zajmowały jak najmniej miejsca. Kolejnym poważnym wyzwaniem jest radzenie sobie z nieuniknionymi awariami poszczególnych dysków. Wykorzystano standardową technikę przechowywania licznych kopii danych na różnych urządzeniach, ale jednocześnie udoskonalono ją tak, by mimo awarii

poszczególnych dysków całość pracowała z maksymalną wydajnością.

Gdy jakiś dysk ulegnie awarii, to po jego wymianie system pobierze dane z innych dysków tak, by stworzyć dokładną kopię zepsutego nośnika. Wgrywanie danych ma odbywać się na tyle wolno, by nie wpływało na wydajność systemu. Jeśli natomiast jednocześnie zepsuje się kilka sąsiednich dysków, tworzenie ich kopii ma przebiegać bardzo szybko, by uniknąć niebezpieczeństwa, że dojdzie do kolejnej awarii, która spowoduje całkowitą utratę danych. Hillsberg ocenia, że dzięki takim rozwiązaniom system nie utraci żadnych danych przez około milion lat, a jednocześnie nie wpłynie negatywnie na wydajność superkomputera.

Magazyn będzie wykorzystał system plików GPFS, który powstał w Almaden na potrzeby superkomputerów. System ten zapisuje wiele kopii plików na różnych nośnikach, co pozwala na błyskawiczny ich odczyt, ponieważ różne fragmenty pliku mogą być później odczytywane jednocześnie z różnych dysków. Ponadto umożliwia on informacji o dokładnym położeniu każdego pliku, dzięki czemu uniknięto konieczności skanowania dysków w poszukiwaniu potrzebnych plików. W ubiegłym miesiącu, korzystając z systemu GPFS inżynierowie IBM-a zaindeksowali 10 miliardów plików w ciągu zaledwie 43 minut, znacznie poprawiając poprzedni rekord wynoszący miliard plików w trzy godziny.

Opracowanie: Mariusz Błoński

Na podstawie: Technology Review

Źródło: [Kopalnia Wiedzy](#)