

國立交通大學
National Chiao Tung University

出國報告（出國類別：出國短期研究）

Software and Hardware Techniques
for Big Data Analysis
適用於巨量資料分析之軟硬體技術

服務機關：交通大學電子研究所
姓名職稱：許書餘 博士後研究員
派赴國家：美國 史丹佛 史丹佛大學統計系
出國期間：102/07/31~11/04
報告日期：102/12/10

摘要

隨著數位資料產生的速度與量愈來愈高，巨量資料之相關應用在近年逐漸受到重視，為了有效的利用巨量資料，最重要的一個應用即為巨量資料分析，藉由重要訊息之萃取，以告知此資料所表示之意義。

然而，巨量資料之分析遭遇到兩個難題，一是演算法在大規模計算時的複雜度，以及大規模資料分析的準確度。另一點是系統愈大規模資料分析時的速度。現有之統計演算法複雜度高，或是準確率有限。而在運算平臺方面，由於一般之中央處理器並非為資料分析做設計，故效能較差，或是需要串聯多臺伺服器，才可達到可接受之分析速度。

有鑑於此，本團隊與史丹佛王永雄老師共同合作，於軟硬體協同設計，在降低演算法複雜度的前提下，仍能保有具競爭力之分析密度估測準確率，並可進一步延伸至各式分類、分群、壓縮等應用，並導入硬體加速的概念，設計彈性化之硬體加速電路與平臺，相較於中央處理器與通用繪圖處理器，可達到 297 倍與 11 倍之運算加速，並且僅消耗小於一瓦之功率，提供省電而迅速之解決方案予巨量資料分析之應用。

目次

一、目的	4
二、過程	6
三、心得及建議	10
心得	10
建議	11
四、附錄	12

本文

一、目的

數位化時代與無線、網際網路的興起，加上蓬勃發展的消費性電子與個人化行動裝置，資料的產生量與產生速度正持續成長中，由國際數據資訊研究機構 (IDC) 統計，於 2010 年的數位資料總量為 1,200 艾位元組 (Exa Byte)，而 2020 年資料更將暴增至 35,000 艾位元組，有鑑於此，如何有效率的減少資料量，進行資料的結構化、儲存機制、管理機制，並即時從巨量資料 (Big Data) 中獲得有意義的資訊將會是極具挑戰的課題。

巨量資料容量可達兆位元組 (Tera Byte) 甚至是超過拍位元組 (Peta Byte)。資料來源包含結構化與非結構化之資料，例如供應鏈管理、金融分析、網路串流、社群網路、多媒體、感測網路、醫學保健、科學計算…等。此外，資料增長的速度亦遠高於傳統資料的產生速率。舉例來說，沃爾瑪 (Walmart) 超市每天在其全世界的分店共有 267 百萬件購物交易紀錄。報價與交易數據之選擇權報價組織 (Options Price Reporting Authority) 指出，金融交易市場的資料產生尖峰速度可達 1.57 吉位元率 (Gigabit/s)，單日產生資料可達 5.5 吉位元。網路巨擎科高 (Google) 每小時則需要處理超過拍位元組的資料串流。在低於 5,000 美元的價位下，每週可產生數百吉位元組的脫氧核糖核酸 (DNA) 與核糖核酸 (RNA) 定序資料，更進一步的基因分析則需動用到數百兆位元組至數拍位元組以上的定序資料。位於瑞士日內瓦，世界最快的大型強子對撞機，撞擊所產生的科學資料每秒最多可達 40 兆位元組，適當處理後可將儲存資料降低至每秒一吉位元組，但每年須更進一步分析的科學資料量仍遠大於 10 拍位元組。

這樣大量的資料，唯有在分析之後才能真的得到有意義之訊息，傳統的數位訊號處理演算法，無論在分析的速度上，或是準確率上，都有其限制，並需要許多人為決定之參數，使得可應用的範圍較被侷限。有鑑於此，機器學習之技術是個較為適合的資料分析方案，藉由各式不同的數據統計方法，使自動化資料分析成為可能，並提高分析資料的能力。然而，現有之機器學習演算法，例如類神經

網路 (Neural Network)、支援向量機 (Support Vector Machine)、K 最鄰近點 (K Nearest Neighbor)、核密度估測 (Kernel Density Estimation)、適應性增強分類器 (Adaboost)，在高資料維度之分析建立之模型準確率有限，並且並未對於大量的資料存取最佳化，此將造成過長的分析時間，而無法即時的反應分析之結果。

除了準確分析的困難之外，大量的資料要如何有效的儲存與管理，也成為了難以處理的課題。在此新興的巨量資料領域，目前巨量資料分析多數仰賴超級電腦或雲端設備來執行，各家網際網路相關廠商與科學研究單位，例如科高 (Google)、臉書(Facebook)、亞馬遜(Amazon)、歐洲核子研究組織(CERN)，皆自行設立大型與高速之伺服器資料中心，或是倚靠甲骨文(Oracle)的資料庫伺服器廠商來協助平臺之架設，然而目前仍成本高昂，若非使用超級電腦等高速運算中心，仍需要數日至數周以上的分析時間，難以提供即時與準確的決策。為了有效解決目前遭遇到的困境，分析方法與硬體設施的突破，便是我們所需要探討的議題。

此次希望結合史丹佛大學王永雄院士長期於統計與機器學習演算法之開發經驗，以及交通大學電子工程學系所於硬體方面之長處，藉由兩校研究團隊的互補性，提出軟體硬體共同設計的機制，包含演算法開發與相關硬體設施的規劃，在巨量資料分析領域中，提出更有競爭力的解決方案。

二、過程

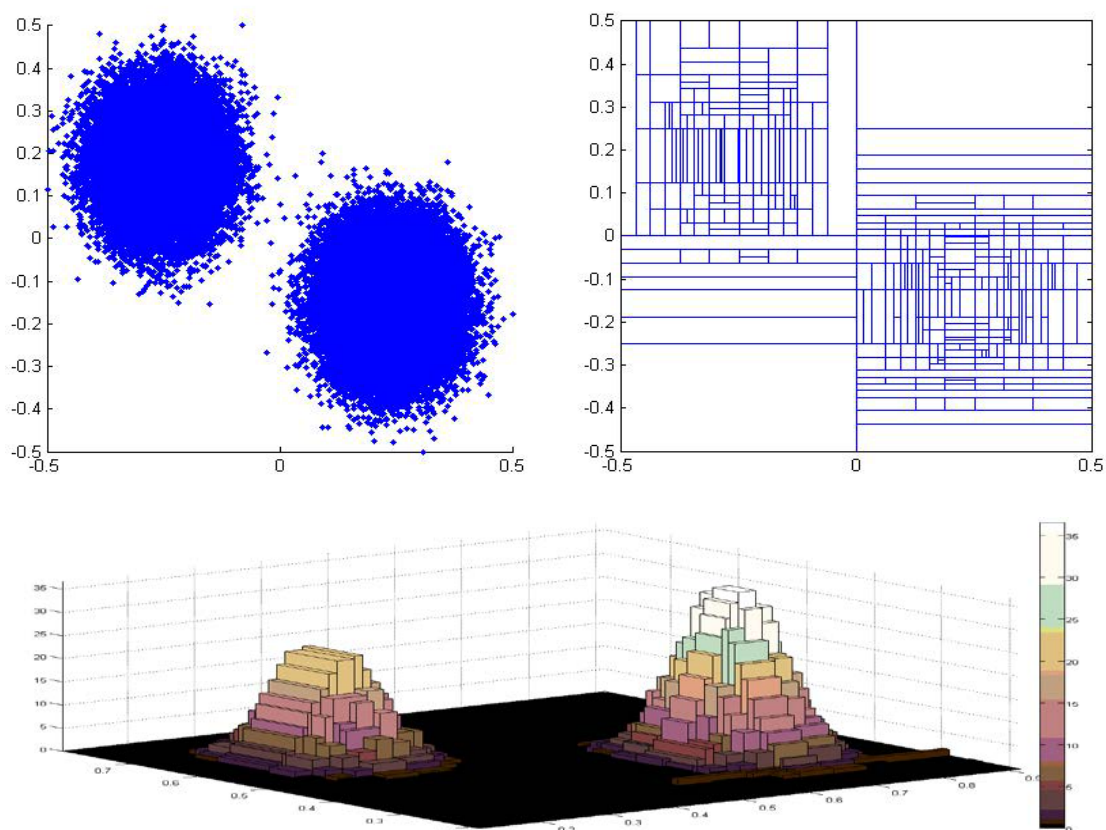
參訪者於七月底到達史丹佛大學，並於八月一日起固定待在王院士實驗室，座位位於史丹佛大學之克拉克中心 (James H. Clark Center)，研究內容包含適用於巨量資料分析之軟體與硬體協同設計，與相關研究的閱讀及比較，並與其他實驗室成員合作。

前述巨量資料分析的瓶頸包含軟體演算法的設計，以及硬體的加速，因此軟硬體兩者同時兼顧將會是巨量資料分析中非常重要的一環。在演算法開發方面，史丹佛的團隊可說是相當領先，有別於前述的傳統學習演算法，王院士團隊開發了一貝式循序切割演算法 (Bayesian Sequential Partition)，用以估測資料之密度分布，藉由循序二元切割的方式，演算法之複雜度大幅降低為線性，並且減少資料的相依性，適合平行處理。由於資料密度與分佈估計可說是統計學習的基礎，因此，經由貝式循序切割演算法估測資料密度後，可再進一步延伸至資料分類、資料分群、甚至是資料壓縮的動作，進一步萃取巨量資料的重要訊息，並且減少需要儲存的資料量。除了二元切割的方式之外，有向性切割也是一個被討論的做法，藉以使用更有效率的切割數，來達到相同甚至更佳의效能。

然而，雖然貝式循序切割演算法已將密度估計的複雜度降至線性，但在巨量資料的前提下，資料動輒數兆位元組，甚至更高，單一個人電腦已難以迅速的執行相關運算，這樣的困境可由效能平行化的硬體加速電路解決，藉由硬體加速來提供較中央處理器更快之巨量資料計算效率，一個不需改變太多程式的作法是使用通用目的繪圖處理器，可使用一般的 C 程式語言來撰寫演算法，然而其加速倍率仍有限制，並且需要較大的功率消耗，因此使用特定目的之硬體加速器將有機會得到更高之效能。

為了達到更迅速之效能，現有硬體公司如甲骨文、國際商業機器、富士通，大多提出更高速之中央處理器，更多層次更大容量的快取空間，更大容量之動態記憶體，以及更快速的網路及介面設計，然而其目的是用於通用目的之系統，導致在資料分析上之效率較差，需要大幅提高系統價格來獲取較佳的效能，或是需要大幅串聯多臺伺服器，以加速分析之時間。

現今雖有文獻提出機器學習的相關硬體加速電路，但大多僅探討分類器或是分群器之設計，如使用支援向量機、K 平均算法、或模糊神經分類器，其演算法複雜度並非線性，此外，其設計之資料吞吐量 (throughput) 亦不夠高，並且未考慮兆位元組至拍位元組等級之資料量，難以直接應用於巨量資料之分析。有鑑於此，我們嘗試設計一加速電路，用以支援貝式循序切割之運算。

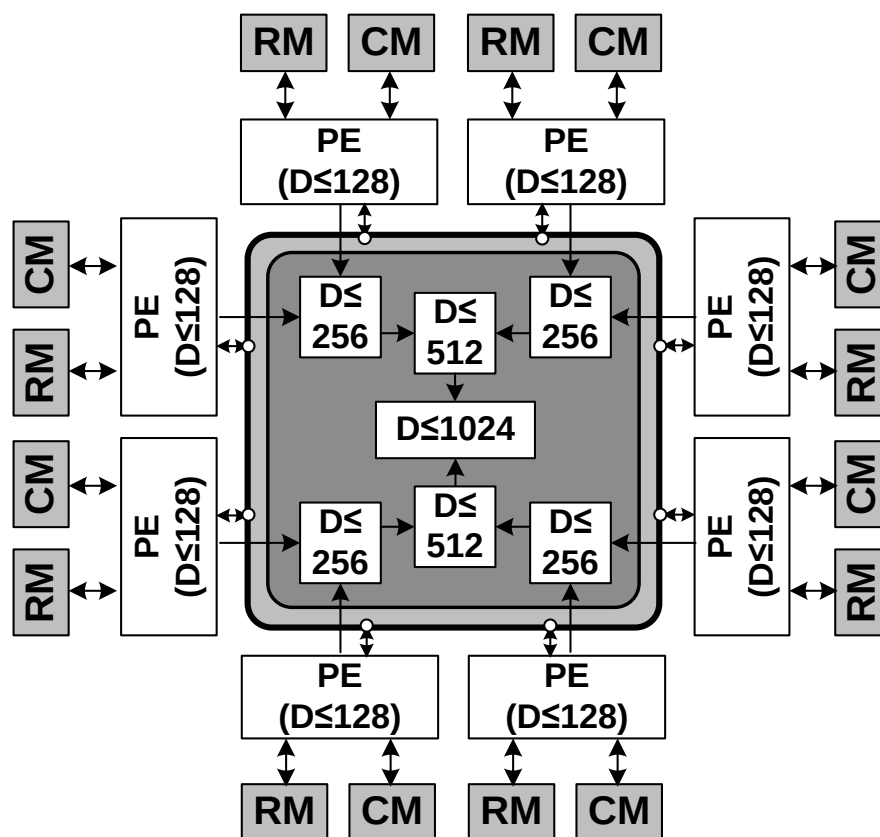


圖片一 資料分布與密度估測

使用二維資料為例，圖片一顯示原始資料之分佈，以及估測之結果，為了確認資料之分布位置，此演算法必須執行大規模的資料篩選，確認資料位置，再藉由區域的面積進一步的計算出資料分佈之密度，由於資料點之分佈並不會重複，因此，增加運算速度最快的途徑，即為在軟體支援的前提下，增加硬體之運算平行度。

我們的作法是設計特定的處理引擎，來完成所需的資料位置確認，以及計數的動作，考慮到資料點數與資料維度的不確定性，適當的硬體彈性是必要的，然

而，彈性與單位晶片面積之平行度難以同時共存，因此我們以時間多工的模式以及多層次平行度的做法，來同時滿足硬體彈性與平行度的需求，架構如圖片二所示，每一個處理引擎 (PE, Processing Engine) 與兩個本地記憶體庫連結，處理 128 維度以下之資料，待資料維度上升時，再以階層式的連結方式來傳遞資料，藉以計算高維度資料所需之分析。這樣的架構有助於我們提升硬體利用率，低資料維度時可以提高系統平行度，並可以高硬體利用率完成高維度資料分析。



圖片二 處理引擎架構

我們將設計之硬體以電路首先以九零奈米製程實現，最大資料維度 1024 維，不限資料量，使用八組平行之處理引擎，平行度最高可達六十四，佔據晶片面積達 11 毫米平方(含界面電路)。使用 128 維度資料與一百萬筆三十二位元資料做比較，此加速電路僅需 4.51 毫秒以完成一次之資料篩選與計數，相較於中央處理器，可達到 297 倍之加速，相較於繪圖處理器，可達十一倍之加速。我們並使用場效可程式化閘陣列搭配兩倍資料速率動態記憶體來完成初步的成果驗

證，證明此架構之可行性，小於一瓦的功率消耗，更是遠小於中央處理器與繪圖處理器所需之功率。有鑑於此，使用硬體加速的機制更有機會於巨量資料分析中，提供更即時的服務。

三、心得及建議

心得

此三個月的短期研究，令我獲益良多。王永雄院士本身橫跨多個領域，從計算機工程、統計學、生物計算皆有涉獵，同時擔任美國科學院院士及中研院院士，並獲得統計學界最高榮譽考普斯會長獎，王院士曾於芝加哥大學、加州大學洛杉磯分校、哈佛大學任教，最後被史丹佛大學聘請至統計系，多數王院士的學生皆投入學術界，並於生物及統計研究方面有優異的成果，投入產業界的學生亦有不錯的表現。

從各個大小事中，皆可令我感受到王院士對研究的積極與熱誠。印象最深刻的是王院士每天皆會到實驗室詢問實驗室成員之進度，包含博士生或是研究員，或是詢問成員是否有任何問題，並給予適當的指引，王院士清楚每一位實驗室成員的研究內容細節，所以若有任何問題，大多能快速解決，並且會回去繼續思索這些結果或是疑問，看看是否有未周全之處或是有可以再探討的部分，所以每一次的討論也都令我有些新的收穫，這是我非常感激的一點。

我的部分研究領域是數位電路相關，與王院士的研究領域不盡相同，但王院士不排斥不熟悉的領域，而是會去做更多的學習，讓自己能夠對相關內容皆有相關之理解。從我的觀察，王院士不僅記得我與他的討論內容，並且會在討論之後試圖自行去了解不懂的部分，甚至是做延伸閱讀，所以對於先前沒討論過的東西，也會有一定的了解，甚至是在下一次的討論時，提出新的看法。這種自我學習的精神很值得參考，令我印象深刻不論是學生或是教授，都應該要保有願意自我學習的態度，才能保持研究的動力。

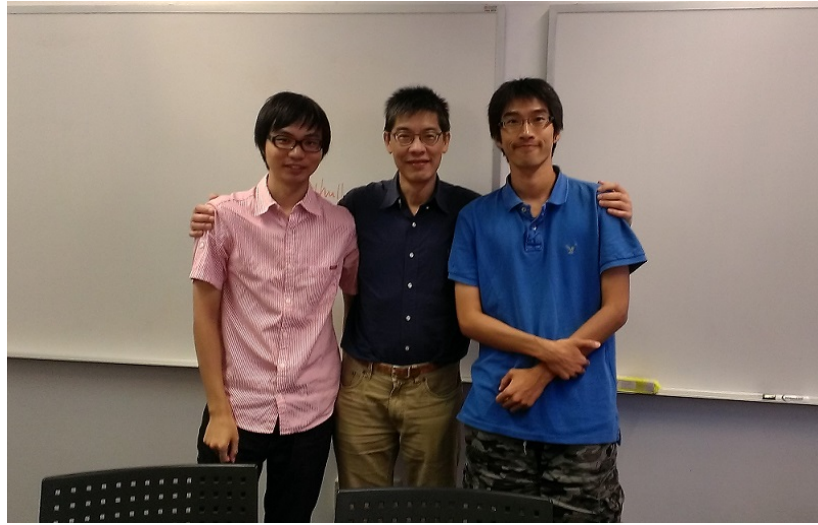
加州的好天氣，讓史丹佛大學成為更適合研究的地點，總是晴朗的天氣讓人心情愉悅，並可以用積極的態度更投入研究，不論是在室內或是室外都可以看到許多醉心於研究的人們。此外，鄰近矽谷也使得學界與產業界之間有強大的連結，但是如何在與產業連結之餘，仍然可保持其研究領域的前瞻性，這是交大可借鏡的一點。

建議

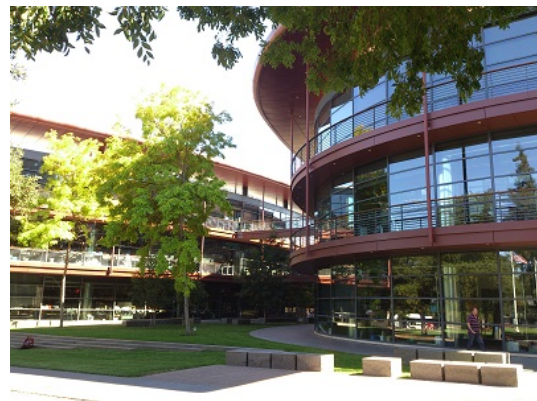
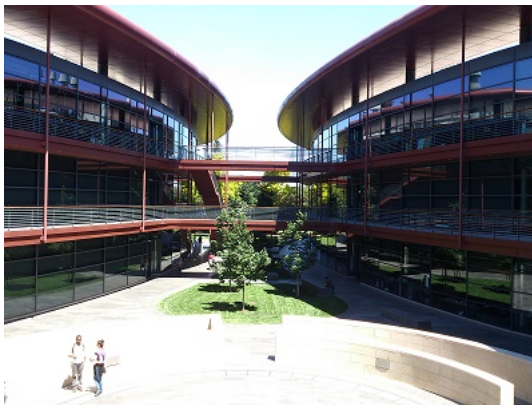
1. 定期與國際團隊合作，並非形式上之交流，而應有實質的研究成果，藉此提升研究題目之前瞻性，並藉由較佳的研究成果，提升交通大學之國際知名度。
2. 目前開設的課程與國外見到的相比，教學之技術內容似乎是有些落後，應考慮開設新課程並以新技術為主題，現有跨界整合之課程是不錯的，但一般來說不夠深入。

四、附錄

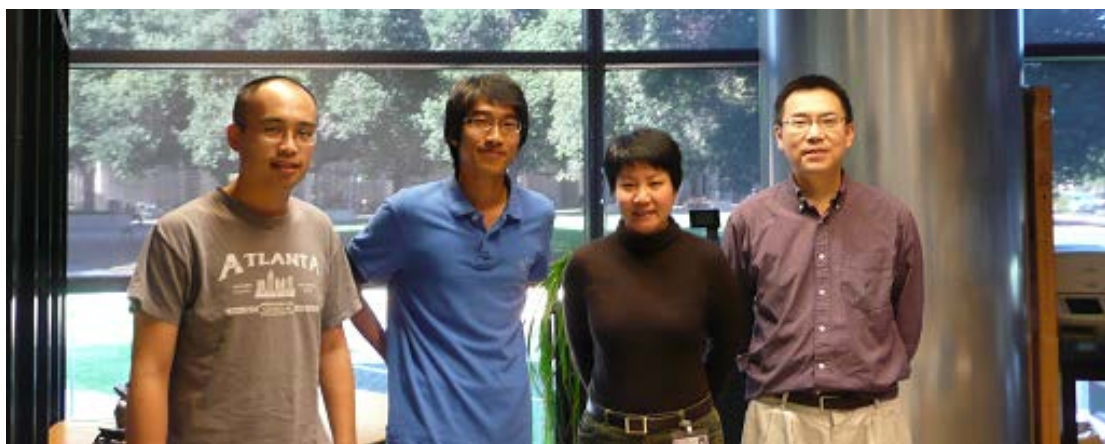
史丹佛大學與王永雄教授實驗室相關照片



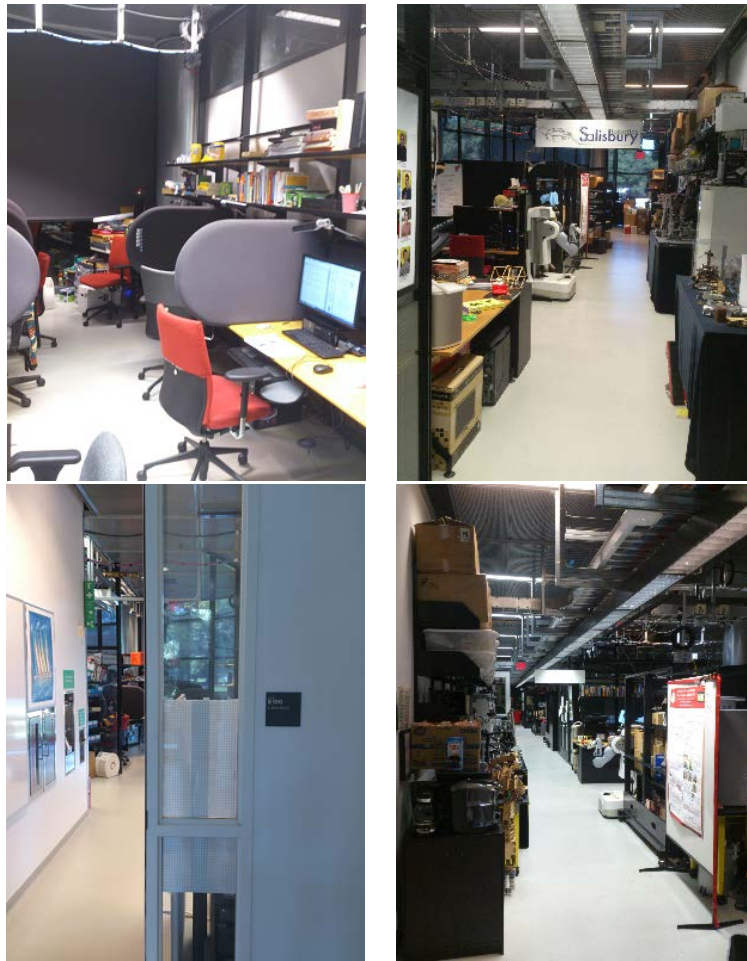
圖片三 訪問者許書餘(右)與王永雄教授(中)及交大邱奏翰同學(左)報告後合影



圖片四 王永雄老師所在實驗室外觀照片(克拉克中心)



圖片五 與實驗室成員合照 (來自生物資訊組)



圖片六 實驗室內部照片



圖片七 與實驗室成員合照 (統計計算組)



圖片八 訪問者於實驗室作業情形