

漫談專利檢索(下)

張元銘

執行階段

前述的原則、方法和策略或許較為抽象，但也突顯出檢索本身是一項冗長、需要細心的工作，有如偵探在做調查一般。真正執行時其實祇要掌握當初設定的目標，善用線索、反覆修正檢索條件，使結果符合要求的標準即可。就一般的執行情況而言，大致可以分成幾個階段：

- 確定有興趣的主題，尤其是遇到關於科技方面的主題。有時在尚未檢索之前，對於自己的目標或預期的結果祇有很模糊的概念，此時除了仿照前述概念解析的方式，盡量使要搜尋的標的變得明確，另外也可從簡單、初步的檢索結果中再調整、研擬出真正想要找的目標。
- 確認檢索的範圍，包括檢索的年份、國家／地區，這可依據檢索需求而定成先決條件，以限制最後得到的專利數量。
- 選擇適當的專利資料庫，包括免費、付費或自建的資料庫，有時還須參考非專利的資料來源以提供線索。
- 反覆進行檢索、篩選、確認。
- 視需要將結果電子化成為自有的資料庫，方便日後查詢、轉用。

舉例而言，剛開始檢索時，按照初步已有的線索，可以

先選擇以下比較簡單的一個條件或組合做為出發點：「案號或日期」、「發明人或申請人」、「申請人加分類號」、「摘要」、「摘要加來源國別」、「申請人加分類號或摘要」…等等。

從初步結果所提供的線索再進一步檢索時，不僅原有條件會做修改，常常還要再納入其他條件(尤其是查詢摘要所用的關鍵字)，然後把各檢索條件做進一步組合、調整，內容時而擴充時而縮小，如此反覆以達最佳化的條件。最後的條件有可能是一長串複雜的檢索指令。

假設以美國專利資料庫為例，透過網際網路進入美國專利商標局查閱近十年公告的專利中符合設定之技術主題者。經由反覆不斷修正各欄位的條件，而由最初數千筆專利逐步縮小至二百筆專利。人工判斷標題或摘要確為所需的專利約有一百筆左右，然後下載專利說明書全文(文字檔或影像檔)。若是利用自動輔助檢索、分析的軟體工具，則下載速度和判別的專利數目都可再增加許多。

視需要可將這些電子檔建構成資料庫以方便日後查詢和利用。而當收集的專利數量眾多時，就應有系統的整理和分類，達到知識管理的目的，而非做圖書館式的累積。例如利用相關的分析工具將重要的結果繪製成各種「專利地圖」，從人、事、時、地、物的角度做策略分析，以明白該技術的發展趨勢、競爭對手的專利佈局、技術的密集區／空乏區…等等深入的情報。

由於資訊變化的速度往往快得超出想像，所以定期重新檢索以更新既有的資料庫也是不可忽略。

檢索指令的注意事項

明白了基本的原則和方法，但有時仍難免會受限於資料庫檢索介面所提供的有限功能。所以宜盡可能靈活運用，否則一樣事倍功半。故查詢前必須明白資料庫的特性，尤其要遵循其檢索指令語法中有關文字 / 數字的格式、字數限制、國籍代碼…等規定。

不同的資料庫檢索介面常採用不同的布林邏輯運算符號來構成檢索指令，常見者如下：

- 與、及、和：「and」、「+」、「*」…

例如「A+B」表示資料中必須同時具有 A 和 B 兩者。

- 或：「or」、「,」、「+」…

例如「A,B」表示資料中具有 A、B 兩者其中之一即可。

- 非、無：「without」、「not」、「-」、「#」…

例如「A-B」表示資料中具有 A 但不含 B。

- 全無、全有：「none」、「all」分別表示該欄資料是空的、不是空的。

- 各種括號、引號：（）、[]、{}、“”、‘’…

用以調整運算優先順序，例如(A,B)+C表示資料是A和B的聯集再與C做交集。或者用以表示多字詞要在一起查詢，例如“heating means”表示heating後面必須接means而視為一個整體，不能分開出現。

- 大小範圍：「>」、「<」、「>=」、「<=」、「<>」、「to」、「…」…

表示數值資料的上限、下限和不等於，有的也可應用於英文字母的順序。

針對例如英語的拼字語文，同一字因不同的詞性甚至誤繕而有多種寫法的變化，也衍生出許多相關字。為了避免漏掉這類的相關字，就以文字字首、字尾、字中的切截方式(以萬用字元代替)進行專利檢索，表示法如下：

- 開放式切截功能：「*」、「\$」、「!」...

用以代替任意數目(包括零)的字母/文字，例如「cataly*」可能找到 catalyze、catalyst、catalysts、catalytic...，「電動*」可能找到「電動車」、「電動玩具」、「電動刮鬍刀」...，「h*ane」可能找到 humane、hexane、heptane...，「*nese」可能找到 Chinese、Japanese...，「*flam*」可能找到 inflammable、flammable...。

- 固定的切截功能：「?」、「#」...

用以代替特定數目的字母/文字，例如「ca?」可能找到 cab、cam、cat...，「電動??」可能找到「電動馬達」、「電動機具」...，「l?ce」可能找到 lace、lice...，「??ine」可能找到 chine、whine...。

另外有的還提供相近運算符號，以找出位置相隔一定字數或以內的兩個字，表示法如下：

- 有特定次序的相近：「adj@」、「@w」...(@為自然數)。

例如「heating adj3 means」可能找到「heating is a means」、「heating it carefully means」...。

- 無特定次序的相近：「near@」、「@n」…(@為自然數)。

類似前者，但是不分前後順序，因此上例改用「heating near3 means」可能還會再找到「means used for heating」、「means the furnace heating」…。

在非拼字的語文(例如中文)的檢索中，可能因為基本文字表示方式的差異，前述切截、次序的功能造成語意選取 / 切割的效果經常迥異於拼字語文，所以就比較少見。不過此種功能在擬定關鍵字時倒是可以做為輔助參考。

此外還有所謂禁止查詢的字(stop words)，例如頻繁出現的「the」、「of」、「is」、「and」、「claim」…等等，一方面比較不具檢索意義，另一方面可能造成系統處理上的過重負擔，故常被資料庫系統排除在可檢索的字以外。這類字也不少，必須事先留意。

結論

循序漸進、善用線索、反覆驗證是專利檢索的法則。而正確的邏輯思考與長期累積的經驗，則是增加檢全率和檢準率的不二法門。也別忘了專利檢索的真正目的在於挖掘和應用當中蘊藏的「策略性情報」。因此後續應針對檢索的結果做適當的整理、分析和解讀，以了解競爭環境的技術層次、法律保護現況、預測技術發展趨勢、刺激創新，進而規劃公司投資方向、促進研究開發、降低侵權風險，才能為企業創造更多的價值和利潤。