

gustav / December 03, 2016 04:26PM

[發展資安AI機制 阻絕駭客攻擊威脅 - 青年日報社論 2016/12/01](#)

發展資安AI機制 阻絕駭客攻擊威脅 - 青年日報社論 2016/12/01

媒體報導，日本防衛省及陸上自衛隊的資訊系統，今年9月疑似遭駭客侵入。日本當局雖否認情資遭竊，但專家認為駭客手法高明；且由於防衛學術單位的「開放系統」遭駭，只要透過可在「開放系統」與「封閉系統」間切換的電腦，即能進一步入侵連接「封閉系統」的其他防衛省和自衛隊資訊設施，以及駐地與基地的網路，情資恐有外洩之虞。

解讀這段訊息，可從目前常見駭客攻擊著手。一般而言，攻擊前階段，駭客會先鎖定目標作業系統尋找漏洞；其後即正式發動攻擊。透過「足跡蒐集」(Footprinting)、「列舉」或「掃描」工具，針對鎖定的區域或系統作全面偵查、分析，找到弱點或漏洞，設計攻擊方案，發動攻擊。

其間關鍵是，除非做到網路實體隔離，若單憑切換機制區分內外網路管制，百密總有一疏；倘若管制機制存有未被發現的漏洞或弱點，駭客只要以自動化工具「踩點」或掃描，要找到一個網路系統的入侵路徑，並非難事。

以惡意程式入侵被鎖定的弱點來說，方式大略有兩類。一是透過人為的系統操作，將惡意程式帶入系統。例如，當我們被誘騙打開不明郵件的夾檔，或開啟惡意郵件時，便會啟動內嵌的檔案，無意間促成惡意程式入侵；又如，瀏覽惡意網站時，被誘騙啟動夾帶惡意程式的超連結圖樣或按鈕，也可能透過瀏覽器，將惡意程式帶入作業系統。另外一類，則是以網頁或網路服務系統為攻擊對象，藉由網頁或服務系統中的程式互動，將惡意程式寫入系統；或竊取暫存於瀏覽器上的個人隱私資訊或系統資訊。後者透過健全的系統資安策略，即可進行防堵；唯人為無意間的惡意程式散布，則是資安控管中最具風險者。

尤值吾人警惕的是，各式各樣的「蠕蟲」攻擊，並不需要人為主動的特定動作來協助傳播，而是由蠕蟲程式，透過作業系統的安全漏洞，以類似癌細胞增生擴散方式，經由電子郵件與內部網路(LAN)，鎖定未被蠕蟲感染的網路節點擴散，並感染整個互聯網路中的作業系統。2003年8月間大爆發的「衝擊波蠕蟲」感染，便是利用微軟系統的安全漏洞，透過電子郵件與區域網路，於微軟 Windows 2000與Windows XP系統之間增生擴散，難以收拾。

一般而言，惡意程式侵入，經由跳板擴散到被鎖定的網路區域後，其程式將直接竊取資訊、監控系統、操控系統，抑或是將大量的系統串聯形成「殭屍網路」，進行「阻斷服務攻擊」(DDoS, distributed denial-of-service)，令被操控的所有作業系統同時針對一個網頁或網路服務，進行大量的瀏覽，導致該網頁或網路服務有限頻寬被占盡而癱瘓。

最可怕的是，以上動作完全不需要人為操作，即能讓惡意程式如蠕蟲、癌細胞一樣，隨意且大規模的擴散，防不勝防。

在資安要求特別高的軍事場域，完全的硬體分離最為有效。但，為了更徹底防堵駭客大規模攻擊，將網路資訊安全納入國防自主項目，確有其必要。具體的做法，包括發展加密的網路基礎建設，並制定獨立的通訊協定，避免使用開放通訊協定。最理想者，是發展自有的作業系統與資訊規格，這些目前在歐盟、日本與中國大陸，都是備受重視的國防發展目標。

隨著駭客技術日新月異，加上大數據應用與機械學習等人工智慧(AI)技術的發展，駭客攻擊手法也愈來愈難防範。然應用人工智慧技術，抵抗駭客攻擊的科技也已開始發展。一組由美國麻省理工學院電腦科學與人工智慧研究室(MIT CSAIL)偕民間機構專家組成的研究團隊，今年4月即公開發表一項運用機械學習抵抗駭客攻擊的人工智慧系統，有效預判駭客可能攻擊方式，並進行防禦的成功率已達85%。

析言之，面對現今資訊時代，「國家級」大規模駭客集體攻擊的格局，已經不是最大規模的資安威脅；基於人工智慧技術發展的攻擊技術，才是真正難以應付的危機。因此，援引最新的人工智慧科技，並發展防範資安威脅的人工智慧機制，實是刻不容緩的任務。

出處：<http://news.gpwb.gov.tw/News/172856>