

「機器學習是人工智慧的一個分支，關注可從資料中學習的系統的建構與研究...機器學習的核心涉及「表示」與「泛化」。資料例項的「表示」是所有機器學習系統的一部分。「泛化」是系統在未見過的資料例項上表現良好的屬性；可得到此結果的條件是運算學習理論中的一項研究關鍵對象。」

— 法國哲學家、
數學家暨科學家勒內·笛卡爾

有更好的方法嗎？

Cylance® 如何使用數學擊敗惡意軟體

雖然很少人願意承認這一點，但問題是企業安全人員所捍衛的城堡千瘡百孔，秘密通道到處都是，用來防衛的屏障也沒有效果。這些弱點都是使用品質不佳安全軟體、低劣硬體的後果，某些狀況下也可能是有惡意內部人士植入後門。最終結果就是只能眼睜睜看著攻擊者贏得這場戰爭。

攻擊者的動機多種多樣，他們從世界各地發起攻擊，並且隨著技術的進步，攻擊的複雜性也不斷演進。在這波演進當中，現代威脅通常會採用規避技巧繞過現有的安全措施。光在事件發生後偵測出這些進階威脅就已經夠難了，更不用提如何在事前保護整家企業抵禦這些威脅。

人類因素

為了跟上現代攻擊者的步伐，安全性技術必須能在無人類介入的狀況下跟著演進。這就是數學和機器學習的優勢所在。若我們能根據數學風險因素客觀地分類「好」檔案與「壞」檔案，我們也能教導機器對這些檔案即時作出適當決策。

將數學與機器學習方法運用於電腦安全上，能徹底改變我們瞭解、分類與控制每個檔案執行的方式。Cylance® 產品採用上述方式，有別於市場上其他所有安全產品。多年來，醫療、保險和高頻率交易等產業已經運用了機器學習的原則，分析大量商業資料並推動自主決策制定。Cylance 每項實施作業的核心，都是可大規模擴充的「大腦」，能將高度調校的數學模型近乎即時地運用於龐大的資料量。

將機器學習運用於檔案分類

過去幾十年間已建立了數十億個檔案—包括惡意與非惡意檔案。在檔案建立的進化過程中，出現了規定特定檔案類型建構方式的模式。這些模式有變異性與不規則性，但整體來說，電腦科學過程是相當一致的。

若檢視 Microsoft®、Adobe® 和其他大型軟體廠商等不同開發業者，會發現模式更加一致。著眼於特定開發商與相似攻擊者所用的開發過程，所觀察到的模式一致性還會增加。挑戰就在於識別模式、瞭解它們如何在數百萬種屬性與檔案之間顯現出來，以及辨識哪些一致的模式可以告訴我們這些檔案的性質。

由於所涉及的資料大小、偏見的傾向以及所需的運算量，讓人類無法運用此資料判斷檔案是否屬於惡意。遺憾的是，大部分安全性公司仍仰賴人類作出這些判斷。這些公司聘用大量人員審視數百萬個檔案，判定哪些是「好」檔案，哪些是「壞」檔案。

現代威脅數量龐大、複雜程度高，無論是人類的腦力還是體耐力都無法負荷。雖然人類在行為與漏洞分析以及識別入侵指標方面已取得進展，但這些「進展」仍有相同的嚴重缺陷。也就是全都以人類觀點與分析為基準—而人類往往會犯下過度簡化的錯誤。

然而機器並沒有相同的偏見，機器學習與資料探勘能攜手並行。機器學習的重點在於根據從先前資料習得的屬性進行判斷。Cylance 就透過這種方式將惡意檔案與安全或合法檔案區分開來。資料探勘的重點在於探索先前未知的資料屬性，讓這些屬性能被用在日後的機器學習決策上。

機器學習運用四階段流程：收集、擷取、學習和分類。

「機器學習是人工智慧的一個分支，關注可從資料中學習的系統的建構與研究...機器學習的核心涉及「表示」與「泛化」。資料例項的「表示」是所有機器學習系統的一部分。「泛化」是系統在未見過的資料例項上表現良好的屬性；可得到此結果的條件是運算學習理論中的一項研究關鍵對象。」 —維基百科

收集

檔案分析與 DNA 分析或精算評估非常相似，一開始都要收集大量資料 — 在此領域中要收集的是特定類型的檔案 (可執行檔、PDF、Microsoft Word® 文件、Java、Flash 等)。我們從產業來源、企業專屬儲存庫和裝有 Cylance 代理程式的使用中電腦上即時輸入的內容，透過「饋送」方式收集到數十億個檔案。

收集的目的是要確保：

- 具備統計上顯著的樣本數量
- 具備涵蓋最廣泛檔案類型與檔案作者 (或 Microsoft、Adobe 等作者群) 的樣本檔案
- 不過度收集特定檔案類型而讓收集內容有任何偏誤

收集到這些檔案之後，會檢閱檔案並將其分成三類：已知且確認為正當、已知且確認為惡意，以及未知。務必要準確地將檔案分成這三類 — 正當類別中包含惡意檔案或惡意類別中包含正當檔案，都會產生錯誤偏差。

擷取

機器學習程序的下一個階段是擷取屬性。這個程序和現今威脅研究人員執行的行為識別或惡意軟體分析程序有著實質上的不同。

Cylance 不尋找人們認為或暗示是惡意的特徵，而是運用機器的運算功能和資料探勘技巧，識別某檔案中最廣泛的一組特徵。這些特徵可以基本至 PE 檔案大小或所用的編譯工具，也可複雜至以二進位格式檢閱最開始的邏輯跳躍。我們根據檔案類型 (.exe、.dll、.com、.pdf、.java、.doc、.xls、.ppt 等) 擷取出獨特的細微特徵。

Cylance 藉由識別最廣泛的屬性組合，去除人力檔案分類所帶來的偏差。使用上萬種屬性也讓攻擊者難以創造出 Cylance 尚未偵測出的惡意軟體，從而大幅提高其作業成本。

這種屬性識別與擷取程序的結果，就是建立起和生物學家所用人類基因組極為相似的「檔案基因組」。接著將此基因組當作基準，建立數學模型以判斷預期檔案特徵；就像運用人類 DNA 分析判斷細胞特徵與行為一樣。

學習

收集到屬性之後，會將輸出結果正規化，轉換為可用於統計模型中的數值。這正是運用向量化與機器學習消除人造偏差並加速分析處理速度的時機。接著 Cylance 的數學家就運用擷取階段識別出的數百萬種檔案屬性，開發能精確預測檔案屬於正當或惡意的統計模型。

我們以關鍵量測標準建立出數十種模型，確保 Cylance 產品所用的最終模型能做出準確預測。無效的模型會被丟棄，有效的模型則會歷經多種等級的測試。第一級會先從幾百萬個已知檔案開始，之後的階段會涉及整個檔案語料庫 (數千萬個檔案)。最後將最終模型從測試語料庫中擷取出來，載入 Cylance 的生產環境以用於檔案分類。

請務必記住，每種/每個檔案都歷經數千種屬性的分析，以區分出合法檔案與惡意軟體。這就是 Cylance 引擎識別惡意軟體 (無論封裝或未封裝、已知或未知) 並達到空前準確性的方式。此方式將單一檔案區分出無數特徵，並對照其他數十億個檔案分析每項特徵，取得關於各特徵常態的判斷。

RosAsm Base3.exe PE File Structure	
Dos MZ Header	
DOS Stub	
PE File Header	
PE Signature	
Image_Option_Header	
Section Table	
Array of Image_Section_Headers	
Data Directories	
Sections	
.idata	
.rsrc	
.data	
.text	
.src	

分類

建立出統計模型之後，就能使用 Cylance 引擎分類未知檔案 (例如未曾見過的檔案，或未由其他白名單或黑名單分析過的檔案)。這樣的分析只需花費幾毫秒，但由於分析過龐大的檔案特徵，結果卻極為精準。

因為是用統計模型進行分析，因此分類並不是在黑箱作業狀況下完成的。Cylance 在分類過程中，會提供「可信度分數」給使用者。這個分數賦予使用者更多洞察力，使他們能用來權衡該對特定檔案執行什麼動作——封鎖、隔離、監控或進一步分析。

機器學習方法與傳統威脅研究方法之間有個很重要的區別。Cylance 透過數學方法建立了專門確定檔案是正當還是惡意的模型。若我們對其惡意意圖的信心程度不到 20%，且沒有其他惡意意圖的指標，也會傳回「可疑」的回應。這麼做可讓企業對於自己環境上所執行的檔案具備全面性觀點。目前業界中，威脅研究人員判斷某檔案是否為惡意，白名單廠商只判斷某檔案是否正當；上述方法也能去除其中的偏差。

Cylance 與真實世界

Cylance 曾於外界完全未觀察到威脅的情況下阻止了 Microsoft Word RTF (CVE-2014-1761) 零時差惡意軟體威脅的執行，並且是在先前毫不知情的情況下。

Cylance 於 2014 年 3 月就已發現並隔離了此威脅，而 malwr.com 直到 4 月才予以列出，當時 51 款防毒引擎中只有 4 款能夠偵測到此威脅。

而 Cylance 引擎在無須任何更新的情況下立即偵測到相同惡意軟體 (a2fe8f03adae711e1d3352ed97f616c7)。Cylance 阻止了此漏洞攻擊的執行，如以下螢幕擷取畫面所示。

不會過時的安全性

將數學模型運用於端點上，讓 Cylance 引擎在惡意軟體偵測與防禦方面輕鬆超越所有傳統方法。我們的任務是要在壞檔案造成任何損害之前予以阻止。透過這種方式，即使壞檔案位於磁碟上，端點仍可維持安全而不受侵害。

Microsoft Security Advisory 2953095

4 out of 5 rated this helpful - Rate this topic

Vulnerability in Microsoft Word Could Allow Remote Code Execution

Published: March 24, 2014 | Updated: April 8, 2014

Version: 2.0

General Information

Executive Summary

Microsoft has completed the investigation into a public report of this vulnerability. We have issued MS14-017 to address this issue. For more information about this issue, including download links for an available security update, please review MS14-017. The vulnerability addressed is the Word RTF Memory Corruption Vulnerability - CVE-2014-1761.

Acknowledgments

Microsoft thanks the following for working with us to help protect customers:

- Drew Hintz, Shane Huntley, and Matty Pellegrino of the Google Security Team for reporting the Word RTF Memory Corruption Vulnerability (CVE-2014-1761)

Other Information

Microsoft Active Protections Program (MAPP)

To improve security protections for customers, Microsoft provides vulnerability information to major security software providers in advance of each monthly security update release. Security software providers can then use this vulnerability information to provide updated protections to customers via their security software or devices, such as antivirus, network-based intrusion detection systems, or host-based intrusion prevention systems. To determine whether active protections are available from security software providers, please visit the active protections websites provided by program partners, listed in Microsoft Active Protections Program (MAPP) Partners.

Feedback

- You can provide feedback by completing the Microsoft Help and Support form, Customer Service Contact Us.

Support

- Customers in the United States and Canada can receive technical support from Security Support. For more information, see Microsoft Help and Support.
- International customers can receive support from their local Microsoft subsidiaries. For more information, see International Support.
- Microsoft Technical Security provides additional information about security to Microsoft products.

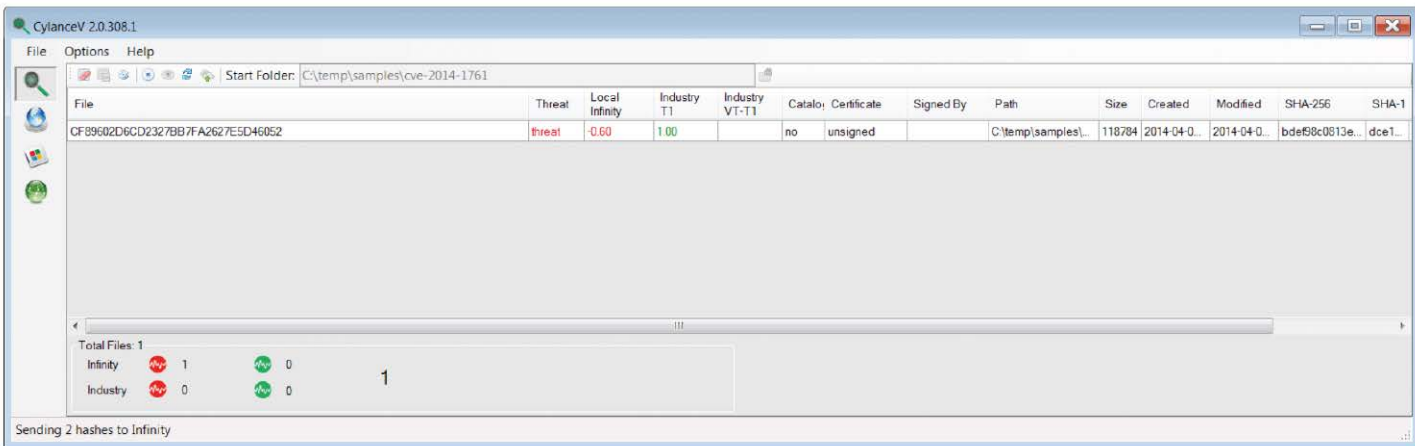
Disclaimer

The information provided in this advisory is provided "as is" without warranty of any kind. Microsoft disclaims all warranties, either express or implied, including the warranties of merchantability and fitness for a particular purpose. In no event shall Microsoft Corporation or its suppliers be liable for any damages whatsoever including direct, indirect, incidental, consequential, lost of business profits or special damages, even if Microsoft Corporation or its suppliers have been advised of the possibility of such damages. Some states do not allow the exclusion or limitation of liability for consequential or incidental damages so the foregoing limitation may not apply.

Revisions

- V1.0 (March 24, 2014): Advisory published.
- V1.1 (March 27, 2014): Updated Advisory FAQ to clarify that Microsoft WordPad is not affected by the issue and to help explain how the issue is specific to Microsoft Word.
- V2.0 (April 8, 2014): Advisory updated to reflect publication of security bulletin.

Page generated 2014-05-14 17:53Z-07:00.

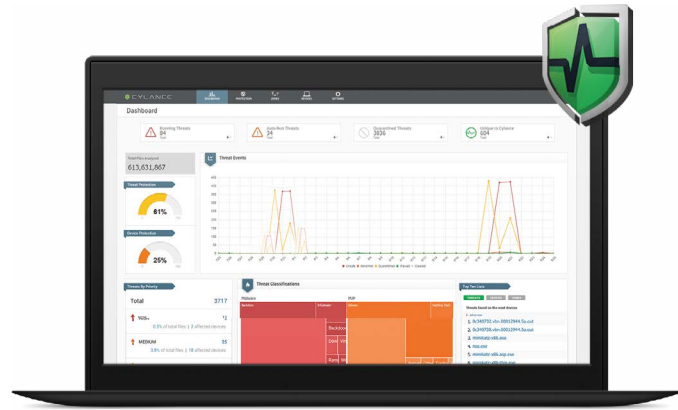


CylancePROTECT® 重要功能

- 防護與偵測先前無法偵測到的進階威脅
- 具備雲端功能,但在敏感環境中無須依賴雲端
- 無須每日 .DAT 更新,因此不需一直處於連線狀態
- 對效能影響極低;運作階段執行大幅減少開銷
- 能透過專門打造的網頁介面輕鬆部署與管理

CylancePROTECT®

CylancePROTECT 是我們的旗艦企業產品,配備強大的 Cylance 引擎,能在企業的各端點上即時預防進階威脅的執行。



CylancePROTECT 提供惡意軟體的即時偵測與防禦。其運作方式是分析作業系統 (OS) 和記憶體層是否有潛在惡意軟體執行,預防惡意裝載得逞。記憶體防護專門設計為接觸量極低,不會產生過大效能負荷。另外,記憶體防護還會提供額外的保護層,偵測和拒絕漏洞通常利用的特定行為,強化 DEP、ASLR 和 EMET 基本作業系統防護功能。

我們以多種企業職能必備附屬功能支援上述兩種核心功能,包括:

- 白名單與黑名單支援管理精細程度
- 單純偵測模式 (稽核模式)
- 自我防護 (防止使用者竄改)
- 從管理主控台進行完整控制、更新與設定

CylanceV™

CylanceV 是事件回應與取證工具,能讓企業將 Cylance 引擎的功能流暢整合至 SOC、CSIRT、服務台和之間的任何位置。CylanceV 是輕量、易於整合的端點檔案分析公用程式,能加以擴充以滿足任何部署規模。CylanceV 提供兩種形式。第一種是能直接存取 Cylance 引擎的 REST API。使用者可撰寫程式以運用 API (可向 Cylance 取得 Python 樣本)。使用這種選項的企業須將檔案傳到 Cylance 的安全性雲端進行分析。

第二種形式是本地公用程式,其中包含的屬性擷取引擎與統計模型執行於企業界線內 Windows 或 Linux 機器上 (無須將檔案傳到 Cylance 的安全性雲端)。

CylanceV 專為迅速部署與完整企業評估等目的所打造,讓機器學習的強大功能充分發揮在追捕惡意軟體上。

關於 Cylance

Cylance 是唯一一間提供預防性網路安全解決方案, 在最脆弱的位置 — 端點處阻止先進威脅與惡意軟體的公司。Cylance 的端點安全性解決方案 CylancePROTECT 運用革命性的人工智慧方法, 於程式碼在端點上執行之前分析程式碼的 DNA, 發現和阻止其他解決方案識別不出的威脅, 同時, 僅使用與當今企業部署的端點防毒及偵測和回應解決方案關聯的一小部分系統資源。如需更多資訊, 請造訪 www.cylance.com

Cylance 諮詢

專家組成的 Cylance 諮詢團隊能以多種方式協助企業。我們的團隊能從現有基礎架構中發掘弱點、識別先前無法偵測出的威脅、協助於攻擊後復原、實施「最佳實務」政策, 以及部署產品以在業務受到影響之前偵測並預防攻擊。

所提供的主要服務有:

- 入侵評估
- 穿透測試
- 取證調查
- 專門與客製服務 (CIKR、嵌入式、ICS)
- 安全性軟體開發

Cylance 諮詢團隊將 Cylance 引擎的強大功能運用於各種服務, 提供其他任何廠商都無法做到的深層洞察力與分析。

總結

Cylance 確信, 數學建模、人工智慧和機器學習是通往安全未來的關鍵。我們提供的每一項產品與服務都與 Cylance 引擎密切整合, 讓您對於現代威脅態勢具備無與倫比的精準洞察力。最棒的是 Cylance 引擎能從新資料不斷學習與訓練, 真正做到「不會過時」, 即使攻擊者改變自己的策略, 也不會隨著時間喪失效率。

+1 (844) CYLANCE
apac@cylance.com
www.cylance.com
18201 Von Karman Ave., Suite 700, Irvine, CA 92612

