

資安轉捩時刻

為何漸進式修補措施不再有效

2024 年 10 月



來自 Google 的问候

隨著工作型態改變，企業環境持續重塑，全球各地機構面臨的安全威脅也急遽增長。Mandiant 近期發布的《M-Trends 報告》顯示，光是去年，國家級間諜活動和商業導向的勒索軟體攻擊就層出不窮，大型政府機關和商業機構深受其害；各產業的資料侵害事件也屢見不鮮，從電信業到娛樂業無一倖免。

為協助機構進一步瞭解如何因應新的威脅情勢，Google 與 [Hypothesis Group](#) 合作，近日做了一項全球資安研究。我們訪問超過 2,000 名中型和企業級機構的決策者，探討他們採取的安全防護措施、觀察到的變化，以及未來展望。

我們深知未來的工作模式已來臨，新時代的安全風險也如影隨形。但許多機構並沒有順應趨勢，反而採用更多傳統策略與工具。結果呢？安全事件與日俱增且代價不菲，不僅干擾業務營運、影響客戶資料，也使機構信譽受損。現行的資安策略已不敷使用，是時候採取更完善的做法脫胎換骨了。

企業需要的不是更多資安產品，而是更安全的產品。我們相信安全事件和資料侵害應為特例，不該是常態。藉由在解決方案中融入安全考量的設計，我們希望機構更加胸有成竹，積極防範攻擊。

期待這份研究可以成為契機，讓我們持續溝通，瞭解貴機構的安全防護需求，同時探討 Google Workspace 如何負起策略合作夥伴的責任，協助您打造更安全的工作環境。



Andy

Andy Wen

Sr. Director of Product Management
Google Workspace Security

目錄

簡介	4
重要洞察資訊	5
現行做法並非長久之計	6
機構明白提升安全性刻不容緩，但固守之前的策略並非解決之道	11
未來必須改變	16
中型機構適合推動轉型	20
全球各地採取不同做法因應資安情勢變化	23
重點摘要	27
詳細研究目標與方法	29



簡介

現今工作環境瞬息萬變，混合辦公模式已成為數百萬名員工的日常，加上 AI 技術日新月異，隨之而來的是安全威脅快速增加、更趨複雜。各機構逐漸意識到這些日益嚴重的新風險，但多數仍猶疑不決，並未大膽做出必要決策確保環境安全。

這份白皮書將直探現今企業環境的矛盾核心：有些主管對自家策略信心滿滿，但安全事件的發生頻率與嚴重程度，卻顯露他們的認知與現實差之千里。只有做出必要的改變，而非墨守成規，機構才能弭平這道落差，更妥善應對資安情勢變化。

為進一步瞭解現況和機構提升安全防護機制的方法，Google Workspace 與 Hypothesis Group 這家深入分析、設計及策略公司合作，執行了一項多方位研究。此研究計畫包含以下項目：

- 量化問卷調查：在 2024 年 7 月 22 日至 8 月 8 日之間，訪問超過 2,000 名 IT、業務和資安決策者，受訪者遍布美國、英國、印度和巴西的中型機構 (300 到 999 名員工) 和企業級機構 (超過 1,000 名員工)。
- 質性訪談：在 2024 年 8 月 22 日至 26 日之間，與美國的六位決策者對談。
- 與下列業界專家討論目前的資安趨勢和未來狀況：

John McCumber

前 NIST 共同主席與 Gartner 總監
資安專欄作家

Joshua Scarpino 博士

TrustEngine 資安副總監
Assessed Intelligence 執行長

Rachel Tobac

SocialProof Security 執行長、Women in Security and Privacy 董事會主席、CISA 技術諮詢委員會 (Technical Advisory Council) 志工

重要洞察資訊

資安問題迫在眉睫：現行做法並非長久之計且成本高昂，機構急需改變策略，才能預先防範日新月異的威脅。

若要因應快速發展的威脅情勢，必須採取新的安全防護措施，但只有少數機構投入推動革新。

1 現行做法並非長久之計

雖然機構表示對自家安全防護機制有信心，但許多行為皆有風險，資安主管也日漸憂心。

2 機構明白提升安全性刻不容緩，但固守之前的策略並非解決之道

許多主管為了提升安全性，均採用「加法」策略，例如增加工具、投入更多心力、提高保險費率。但是，安全事件卻仍屢見不鮮，且代價不菲。

3 改變才能順應未來趨勢

許多主管明確表示想簡化資安措施，並使用 AI 保護機構。比起只在事後補救，解決根本的不足之處才是王道。

機構規模和所處地區各異，因應資安情勢變化的做法也不同，但唯一共識是改變勢在必行。

4 中型機構 (300 至 999 名員工) 適合推動轉型

考量到機構複雜程度和老舊技術，許多主管對未來都充滿焦慮，但多半仍會積極反思安全防護措施，願意做出改變。

5 全球各地面臨的挑戰互異，因此必須採取不同做法，才能因應資安情勢變化

雖然世界各地的主管都明白工作環境將持續變遷，但只有部分國家/地區有長遠有效的資安策略，因此需要全面重新評估。

1

現行做法並非長久之計

雖然機構表示對自家安全防護措施和工作模式有信心，實際行為卻漏洞百出，讓資安決策者夜不成寐。

信心滿滿，但工作模式轉變導致安全性下降

大多數機構認為自家安全防護措施足夠，至少理論上如此。高達 **96%** 的決策者（包含 IT、業務和資安主管）皆表示對自家機構的資安深具信心，例如能有效管理資料、應用程式、裝置和生成式 AI

等新興技術。然而，這種信心掩蓋了實際工作模式的重大轉變。據我們瞭解，這類轉變無所不包：從混合辦公模式日益普及，到導入其他新工具（尤其是生成式 AI）都是如此。

74%

的人表示工作模式大幅改變

「我們活在變革的時代。雖然可用的工具不可勝數，但我非常憂心，因為似乎沒有相應的政策，也沒有任何計畫可以妥善運用這些新工具。」

— 製藥業研究總監



具有風險的工具和行為屢見不鮮

決策者多半表示許多做法會對機構帶來風險：

- 只有 **56%** 表示會依循 IT 政策
- 機構內部使用的工具有 **19%** 並未獲得授權 (也就是未經過機構審查及核准)
- 將近半數 (**44%**) 認為未授權的工具「安全無虞」
- 三分之二 (**63%**) 表示每週均會使用未經授權的生成式 AI 工具
- 一半 (**48%**) 的決策者信任未授權的生成式 AI 工具

機構說法與實際行為有落差，會放大現行資安機制的固有風險。此外，如果又使用影子 IT 服務 (這種行為既常見又難保安全)，不僅會導致安全事件的風險增加，也可能違反法規。

「機構不一定理解風險會如何隨著新技術演變，這可說是資安主管的挑戰，因為他們經常只能事後補救。」

— **Joshua Scarpino 博士**

Assessed.Intelligence 執行長



資安主管備感焦慮

資安決策者 (SDM) 已察覺這些風險。

63%

認為自家機構的技術環境與過去相比
較不安全

4 分之 1 的工具經認定為**具有風險**



「大家之所以認為現在環境較不安全，我猜是因為覺得缺乏掌控權和相關知識。無論是誰，都很難跟上並理解手上的全部工具、使用情境、應用方式，以及有心人士如何以各種手段利用這些工具對付我們。」

— John McCumber

前 NIST 共同主席與 Gartner 總監、資安專欄作家



這類具有風險的行為令人焦慮不安，**93%** 的 SDM 都擔心會發生安全事件。導致焦慮的原因很明確，雖然 SDM 最擔心安全漏洞和外部攻擊，像是生成式 AI 攻擊和後續的資料侵害事件，但他們也怕使用者不夠謹慎，無論對方有意或無意。

「決策者告訴我，他們憂心忡忡，夜裡輾轉難眠。技術工具並不完美，需要大量人力監督，因此決策者備感壓力。」

— **Rachel Tobac**
SocialProof Security 執行長

 **93%**

SDM 擔心會發生安全事件

攻擊/侵害事件主要疑慮

生成式 AI 攻擊	31%
資料侵害	30%
網路勒索	26%

使用者相關的主要疑慮

未經授權的存取	24%
員工過失	20%
憑證遭駭	17%



2

機構明白提升安全性刻不容緩，但固守之前的策略並非解決之道

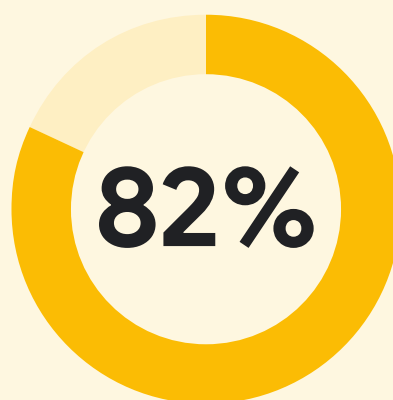
對進階安全防護措施的需求從未如此迫切。SDM 已經察覺現行策略有所不足，也明白必須大範圍改進，但許多 SDM 的做法顯然不夠周全，往往還適得其反。

明確呼籲要提升安全性

82% 的 SDM 指出必須提升安全防護措施。過半數表示現今工作環境複雜，難以確保機構安全無虞。還有 59% 承認他們仰賴過時技術，未準備好因應日後的安全防護需求。

「過去幾年來累積了許多資安解決方案，有些舊工具確實需要替換，或至少需要重新評估。」

— IT、媒體與娛樂業副總監



資安決策者表示機構需要提升安全防護措施



逐步改變並不夠

雖然半數 SDM 表示自家機構的資安策略在過去兩年經歷重大變化，但這些改變通常是漸進的，並非徹底革新。事實上，大部分機構似乎只是繼續採取相同策略。



過去兩年內的 重大資安變革

強化資安保險 37%

增加資安支出 35%

升級授權 34%

新的資安供應商加入 31%



雖然這些調整或許能帶來一定助益，但較偏向是強化現有做法，而非採取截然不同的策略。

「為了建立多重防護，很多人會對同樣的產品投入資金。我聽同事和客戶說過，他們已嘗試投入一大筆錢，卻未實際解決問題。他們花大錢投資高價平台，卻沒為團隊設置密碼管理工具，這就像花錢安裝防彈窗戶，卻大開家門。」

— Rachel Tobac
SocialProof Security 執行長

62%

的 SDM 表示會在需要新功能時擴充工具 (而非替換舊工具)

10

項資安工具 (企業級機構使用的平均數量)

61%

機構使用的資安工具比兩年前多

27%

的資安工具功能重複



過猶不及，必須採取更統合的策略

為確保環境安全，**超過三分之二**的機構比以往投入更多時間金錢，但仍遇上許多付出慘痛代價的安全事件。

相比之下，採用較少工具的機構卻未如此。報告顯示，機構若使用超過 10 項資安工具，遇到安全事件的頻率和代價就越高。

81%

的機構一年發生至少一次安全事件

8

每年回報的安全事件平均數量

\$488 萬

美元

2024 年處理侵害事件的全球平均成本¹

至少採用 10 項資安工具的機構

採用少於 10 項資安工具的機構



一年內發生的安全事件數量

14

對比

6

一年中安全事件相關支出超過 \$25 萬美元

34%

對比

19%

「如果每項資安問題都需要專門的廠商處理，絕對會難以管理，且不符成本效益，因此務必要統合精簡。現有的廠商過多，並非長久之計。」

— IT、媒體與娛樂業副總監

1. IBM,《Cost of a Data Breach Report》
(資料侵害事件處理成本報告), 2024 年



3

未來必須改變

我們的研究顯示，SDM 願意採取不同的安全防護措施，而積極尋求新策略與技術的 SDM 皆獲得卓越成效。

渴望改變

受目前形勢所趨，機構也改以不同方式考量資安策略。

決策者越來越傾向簡化的資安解決方案，**65%** 表示有興趣使用統合的資安系統。他們認為若平台能將多項資料來源整合至單一介面，確實會帶來實質效益，包括簡化程序、增進營運效率、減少管理多項工具的負擔，以及節省寶貴時間。

78%

的機構正在重新考量安全防護措施，因應目前的資安形勢

73%

願意接受新的安全防護措施，無論成本高低





此外，生成式 AI 越來越受歡迎，**59%** 的資安決策者認為這是重要工具，能對抗瞬息萬變的威脅。

「AI 在資安方面最大的影響，是讓使用者有能力因應安全事件、使用現有工具，以及實現自動化調度管理，因而有辦法應變、回報及善後。」

— **John McCumber**

前 NIST 共同主席與 Gartner 總監、資安專欄作家

 **59%**

的 SDM 認為生成式 AI 是保障資料安全的重要工具

採用生成式 AI 的主要原因

促進資料隱私權	43%
強化威脅情報功能	37%
提升資安營運效率	36%
威脅防護	33%

投入改變的優點

儘管有採用新做法的強烈意願，卻只有 **38%** 的 SDM 積極替換過時工具，**62%** 仍選擇擴充現有配置。

替換舊工具的 SDM (對比擴充工具的 SDM)



每年回報的
安全事件少 **20%**



使用的資安工具較可能
小於或等於 **2 年前** 的數量
(44% 對比 36%)



表示比過去**花費較少時間**
確保環境安全
(33% 對比 28%)

「團隊必須盤點購入的工具，才能知道已具備哪些功能，避免一再購買相同的防護措施。瞭解哪些功能可用，以及哪些地方需要安全防護，有助機構進一步掌握自家工具，進而減少資安事件和問題。」

— Rachel Tobac

SocialProof Security 執行長



4

中型機構適合推動轉型

中型機構 (300 至 999 名員工) 的安全機制並未完全發揮潛力，而他們也瞭然於心。面對複雜的現代工作環境時，中型機構的疑慮遠超過企業級機構 (超過 1,000 名員工)。

此外，中型機構決策者表示對目前的技術環境越發不安。隨著工作環境中快速採用生成式 AI，這類疑慮也有增無減。

73%

的機構表示內部越來越常使用生成式 AI，而這導致安全事件增加

(相較之下，這類企業級機構僅 58%)



目前難題

與企業級
機構的差距

71% 表示他們依賴舊技術，
對未來準備不足 +23 pp

63% 表示複雜的工作方式使資
安措施窒礙難行 +11 pp

73% 表示使用生成式 AI 已導致
安全事件增加 +15 pp



自身感受

與企業級
機構的差距

63% 因潛在安全事件備感焦慮 +13 pp

63% 認為技術環境比起以往較
不安全 +12 pp

36% 表示資安團隊因安全威脅
而不堪負荷 +16 pp

願意採用新做法

82% 的中型機構資安決策者指出，自家機構正積極反思安全防護措施。越來越多中型機構決策者相信雲端原生解決方案具有優勢，其中 **81%** 認為比起傳統電腦版應用程式，雲端原生應用程式更為安全。他們對雲端技術的信心顯示，中型公司不僅意識到改變勢不可擋，也已準備好採用更創新的解決方案。

「中型機構情況特殊，因為他們正在擴大規模，準備拓展至更大的市場。在此階段必須奠定根基，否則就會遠遠落後，擔心的事情也會成真。」

— **Joshua Scarpino 博士**

Assessed.Intelligence 執行長

82%

的中型機構資安決策者正積極
反思安全防護措施



5

全球各地採取不同做法 因應資安情勢變化

雖然決策者一致同意工作模式在過去一年大幅轉變，但應變做法卻因國家/地區而大相逕庭，尤以資安策略為甚。

所有區域的決策者均指出，工作環境與一年前相比變化很大，各市場的資安決策者也在重新思考安全防護措施。



「我們的工作方式與 12 個月前相比大幅轉變/有些改變」

	印度	85%
	巴西	74%
	英國	73%
	美國	62%

「我們的機構正在反思安全防護措施」

	印度	83%
	巴西	76%
	英國	82%
	美國	71%



印度

因情況複雜而不堪負荷，但渴望改變

印度的資安情勢極為複雜，相關工具繁多，因此格外突出。不過，當地機構十分願意提升安全性和發掘新方案，強烈渴望改變。印度顯然處於轉捩時刻，決策者明白必須採用更簡化高效的資安策略。

資安數據匯報

與全球平均的
差距

平均使用 **35 項工具**
(例如產品、服務、應用程式)

+11

平均使用 **19 項資安工具**

+9

67% 表示複雜的工作方式使資安措施
窒礙難行

+10 pp

渴望改變

82% 認為自家機構需要改進安全防
護措施

+8 pp

81% 願意發掘新的安全防護措施，無
論成本高低

+8 pp



巴西

極度焦慮，但較少關注具體威脅

巴西境內的決策者表達出對安全事件和相關後果的深層焦慮，卻較少關注導致安全問題的具體原因。這相當矛盾：他們害怕後果，而不是害怕導致事件發生的安全漏洞。當地機構反映出對改變的渴望，也是最有意願採用生成式 AI 的區域之一 (71%，而全球平均為 62%)。

擔心安全問題

與全球平均的
差距

41% 表示會不停擔心發生安全事件

+19 pp

42% 擔心失去客戶或使用者信任
(若發生安全事件)

+8 pp

37% 擔心失去客戶 (若發生安全事件)

+7 pp

39% 擔心損害品牌形象 (若發生安全
事件)

+5 pp

較少關注具體問題

40% 表示他們依賴舊技術，對未來準
備不足

-19 pp

48% 表示使用生成式 AI 已導致安全
事件增加

-17 pp

44% 表示複雜的工作方式使資安措施
窒礙難行

-13 pp



英國

在乎新興技術，但反應慢半拍

英國境內的決策者相當留意新興技術對資安的影響，尤其是在舊版系統和生成式 AI 威脅方面。然而，英國在推行降低風險的政策上向來步伐較慢，決策者也有嚴重過勞的問題。儘管面臨這些挑戰，英國當地機構卻較無意願採取新的安全防護措施，且傾向認為安全性漏洞在所難免。

擔心技術

與全球平均的
差距

75% 表示他們依賴舊技術，對未來準備不足

+16 pp

77% 表示使用生成式 AI 已導致安全事件增加

+12 pp

儘管面臨挑戰，仍不太願意改變

43% 表示資安團隊因安全威脅而不堪負荷/過勞

+15 pp

44% 認為安全性漏洞在所難免

+14 pp

27% 已導入生成式 AI 專屬的資安政策

-14 pp

60% 願意發掘新的安全防護措施，無論成本高低

-13 pp



美國

過度自信可能忽視風險

美國境內的決策者對自家安全防護機制有深厚信心，較不擔心侵害事件，尤其是生成式 AI 相關威脅。樂觀的態度雖然值得嘉許，但也引起疑慮：美國的機構是否低估急遽增長的威脅情勢？畢竟美國境內處理資料侵害事件的成本居全球之冠，如未明確關注新興威脅，過度自信的機構日後就很可能受到攻擊。

對資安情勢充滿自信

與全球平均的
差距

10% 表示會不停擔心發生安全事件

-12 pp

55% 表示使用生成式 AI 已導致安全事件增加

-10 pp

52% 認為技術環境比起以往較不安全

-4 pp

處理資料侵害事件的成本最高

2024 年處理資料侵害事件的平均成本為 \$936 萬美元¹

+\$500
萬美元

1. IBM,《Cost of a Data Breach Report》
(資料侵害事件處理成本報告), 2024 年

重點摘要

「務必讓全機構上下瞭解安全風險，高階主管常常不理解資安的價值主張，資安專家必須能用淺顯的方式說明，指出風險所在和相應的解決方案。」

Joshua Scarpino 博士

Assessed.Intelligence 執行長

① 資安主管與業務主管需要合作

資安決策者與資深高階主管需要通力合作，透過公開溝通交流與健全的風險管理程序，制定出色的資安策略。就 SDM 而言，若取得相關人士的支持，並建立更深入的合作關係，資安團隊應該就能有效運作，而非坐立難安。

② 不該等到發生侵害事件才推動革新

資料侵害事件等危機也可能是有力的催化劑，資安決策者可藉此徵求高階主管核准必要的安全防護措施。然而，勒索軟體等網路攻擊一夜之間幾乎就能實際干擾任何企業營運，因此積極做出必要改變在如今越發重要。對大部分企業來說，重點不在於資料是否遭到侵害，而是何時會發生。

③ 關注攻擊的起因

根據 Mandiant 最新發布的《M-Trends 報告》，在 2023 年成功入侵的案例中，72% 源自於身分遭駭（如網路釣魚或竊取憑證），或是軟體漏洞攻擊。換言之，只要簡單改進這些方面的措施，就對機構大有助益。舉例來說，Google Workspace 等產品能夠防範網路釣魚，且具備雙重驗證功能，可自動阻擋超過 99.9% 的垃圾郵件、網路釣魚攻擊和惡意軟體。Google Workspace 更無電腦版用戶端應用程式或地端部署系統，不需費力修補或更新。



「很多人認為需購買昂貴的大型程式或工具，必須要價數十萬美元，才能發揮作用。實際上，機構可推行的最大改變之一，就是制定機構政策和調整公司文化。」

Rachel Tobac

SocialProof Security 執行長

「AI 將整合至資安工具中，加強我們應對事件和威脅的能力。隨 AI 而生的眾多功能必將帶來深遠影響。在我看來，AI 的潛力尚未完全發揮，未來還有無限可能。」

John McCumber

前 NIST 共同主席與 Gartner 總監
資安專欄作家

④ 踏出一小步也能改變

全面革新舊版軟體可能令人卻步，不妨逐步推動重要的安全防護措施，盡可能減少對使用者的影響。舉例來說，您可以部署 [Chrome Enterprise](#)，為使用者提供安全的瀏覽體驗，並在他們慣用的網路瀏覽器中，提供內建的網路釣魚與惡意軟體防範功能、不含 VPN 的零信任存取機制，以及資料保護控制選項。

⑤ 生成式 AI 可能增加風險，但也能提供防護

為了降低資料遺失風險，請務必使用安全且遵守業界標準的生成式 AI 工具，例如 [Gemini](#)。此外，生成式 AI 也有助機構對抗新興威脅。舉例來說，Gmail 中的 AI 防護機制已採用大型語言模型 (LLM)，可更妥善防範垃圾郵件和網路釣魚攻擊。事實上，多虧有 LLM，Gmail 擋下的垃圾郵件數量增加 20%，且能更高效評估使用者回報的垃圾郵件，單日處理量是先前的 1,000 倍。採用 AI 後，數十億名 Gmail 使用者的生活品質皆大幅提升。



研究目標

- 01

瞭解目前的資安情勢 (特別要務、心態、疑慮、挑戰)
- 02

評估資安與技術對機構的影響，尤其是在生成式 AI 崛起方面
- 03

探討資安的未來、新興趨勢、機構日後的投資計畫

研究方法

量化問卷調查

20 分鐘的線上問卷調查，於 2024 年 7 月 22 日至 8 月 8 日執行，作答者有 2,025 名，分別位於美國、英國、巴西和印度。

作答者需具備以下條件，才符合篩選標準：

- 任職於員工超過 300 人的機構 (規模不一)
 - 擔任業務決策者或 IT 決策者，可掌管機構或團隊層級的技術/效率提升產品；或者
 - 擔任資安決策者，可掌管機構的資安產品
- 涵蓋各行各業，無論是否受監管

區域	(n=)
美國	501
英國	523
巴西	501
印度	500
機構規模	
中型 (300 至 999 名員工)	834
企業級 (超過 1,000 名員工)	1,191
角色	
IT 決策者	673
業務決策者	762
資安決策者	590
總計	2,025

深度訪談

在 2024 年 8 月 22 日至 26 日，於美國境內執行六 (6) 場一小時的深度線上訪談。受訪者均在超過 300 人的機構中擔任長字輩或下一階的高階主管，包含 IT、業務和資安決策者。

專家訪談

達成重要的專案里程碑時，我們會諮詢 3 位資安專家 (John McCumber、Joshua Scarpino 博士、Rachel Tobac)，取得深入分析、專業知識，並交流想法。