

淺談 AI 科技運用於現今政治作戰之研究

筆者/汪念瑩

提要

AI 人工智慧(以下簡稱 AI)毫無疑問的是近幾年最熱門的議題，數位科技發展的速度如火如荼，藉由網路連結到世界的每個角落，又藉由手機連結到每個人。我們的意識、認知、思想都在不知不覺中受到影響，然而人們對於 AI 的認知事實以及它實際上在戰時以及平時的作用是一知半解的，本文就 AI 與政治作戰的定義以及目前 AI 在平戰時的運用以及我國因應的方式作出淺論，與讀者們一同了解及探討未來國軍在「AI 世代」下的走向。

第一段首先為各位讀者簡單介紹現今政治作戰與 AI 的定義與概念，簡述政治作戰六大核心指導，現今我們朗朗上口的認知戰以及混合式威脅的概念；特別要釐清的是 AI 的運作方式，避免大家產生混淆。

第二段為各為介紹現今主流 AI 在作戰方面的運用，我們區分武力作戰以及政治作戰兩方面為各位敘述，武力戰部分包含：訓練模擬器、無人攻擊武器、後勤管理、情報蒐集與研析、資訊防禦與資訊安全等五項；政治作戰包含：假消息、社交媒體操作、社交媒體監控、數據分析和預測、特定價值觀的資訊引導等五項。

第三段提出國軍在 AI 未來的展望，區分九項論述，包含：一、與民主國家的積極合作，二、AI 人才的培養，三、主動出擊的媒體政策，四、國軍資訊政策的重整，五、立法監管網路犯罪及認知戰問題，六、重視已然成形的認知侵略，七、建立國軍正確資安及個資保密觀念。

關鍵詞：AI，人工智慧，政治作戰，認知戰

前言

AI 人工智慧(以下簡稱 AI)毫無疑問的是近幾年最熱門的議題，數位科技發展的速度如火如荼，藉由網路連結到世界的每個角落，又藉由手機連結到每個人。我們的意識、認知、思想都在不知不覺中受到影響，流行的 AI 應用軟體諸如 Chat gap¹、MIDJOURNEY²等甫一推出即成話題，深怕不跟上潮流就會被時代拋在後面。以政治作戰人員的角度來說，不難聯想到我們政治作戰「以人為中心」以及「無時空限制」的特性。然而要問到「到底甚麼是 AI」，大家似乎又並不是真的了解到他的功能和其代表的意義是甚麼。「把 AI 應用在未來戰爭」是許多軍事學者都會提到的事情，也是許多先進國家的目標，但是到底要怎麼用、如何用，以及我們要如何去培養這些人才，應對這些威脅，如何在諸多限制條件的狀況下運用 AI 來達成作戰任務，卻又有那麼點莫衷一是。在一日千里的 AI 科技運用下，我們國軍所面臨的挑戰以及展望是甚麼，值得我們深思與探討。

壹、淺談政治作戰與 AI 的基本概念

雖然有這麼點老生常談，我們還是先從政治作戰與 AI 的基本概念談起。

我國國軍政治作戰要綱明示：「凡除直接以軍事或武力加諸敵人之外的戰鬥行為，皆為之政治作戰」³，同時也說明「因應戰爭型態之變化，一切政治作戰作為必須結合科技、資訊等現代化戰爭之特性，積極提升政治作戰整體效能。」⁴因此在 AI 科技的運用上，政戰人員自然不可能置身事外。事實上，比起武力作戰，民眾最容易接觸到的也是經由 AI 而實施的政治作戰，舉凡近年來常出現在媒體上的「混合戰」⁵以及「認知戰」⁶就是最好的例子，甚至我們可以大膽的說，所謂「認知戰」就是「政治作戰」的老酒裝新瓶，雖然這個名詞翻譯自美國智庫，但其本質不脫政治作戰，只是用了新的名詞

¹ 是 OpenAI 開發的人工智慧聊天機器人程式，是一個大型語言模型，於 2022 年 11 月推出，被訓練成能夠理解和回答各種不同主題的問題，也能夠進行自然語言對話。

² Midjourney 是一個由同名研究實驗室開發的人工智慧程式，可根據文字生成圖像，其所生成的圖像引發了不少關於著作權的問題。

³ 國防部，《國軍政治作戰要綱》(台北：國防部，民國 105 年)，頁 1-1。

⁴ 同註 3，頁 1-3。

⁵ 一般而言我們對混合戰的定義都引用自美國智庫法蘭克·G·霍夫曼所著《Conflict in the 21st Century: The Rise of Hybrid Wars》之內容，「未來的對手更加聰明，並且鮮少將自己限縮在他們工具箱裡的單一工具。傳統的，不規則的和災難性的恐怖主義挑戰將不會是明確、清楚的樣態；他們都將以許多形式呈現。戰爭模式的模糊，誰參戰的模糊，以及使用什麼技術，產生廣泛多樣性與複雜性，稱之為混合戰。」，本翻譯引用自國防大學黃柏欽上校教官於 2019 年 6 月出版之《國防雜誌》第 34 卷第 2 期刊載之《戰爭新型態－「混合戰」衝擊與因應作為》一文。

⁶ Rachael Burton, <Disinformation in Taiwan and Cognitive Warfare,>Global Taiwan Brief, Vol. 3。http://globaltaiwan.org/2018/11/vol-3-issue-22/>(2018 年 12 月 14 日)

讓大家重新注意到這個歷久彌新的話題。

檢視政治作戰的六大核心指導「思想戰、組織戰、情報戰、謀略戰、心理戰」⁷，就可以發現認知戰就是利用情報操作以及媒體操作達到影響思想與心理的效果，以達到策動群眾，完成謀略目的的作為，現在普遍稱之為「混合式威脅」。將 AI 運用其中的目的也不超出這個範圍。我國國防大學中共軍事事務所董慧明副教授綜合歐美各國對於混合戰的理論，將「混合式威脅」作出界定：

「國家或非國家行為者為實現政治目的，達到經濟目標，在戰場上使用常規武器、非常規武器、恐怖主義、犯罪行為等一系列結合先進科技和武裝力量之複雜組合手段，已出乎預料的方式，鎖定對手進行有形和心理層面的打擊。」⁸

這些手段與目的的綜合運作，讓我國人民不分平戰時均處於「混合式威脅」的攻擊下，其凶險程度並不遜於直接兵臨城下，而 AI 正是最佳利用工具。

那麼我們朗朗上口的 AI 又是甚麼呢？在此筆者要釐清幾個我們時常在媒體上看到卻不求甚解的名詞，也就是「AI」、「大數據」以及「機器學習」。筆者整理非官方組織「臺灣人工智慧學校」的理論如次⁹：

首先是 AI(Artificial intelligence)的定義，望文生義我們可以知道這個名詞意即「人造的智慧」，在現今可以泛指是「讓電腦表現出人類智慧行為的技術」，也就是說我們可以簡單的認為，如果能夠讓電腦表現(或模擬)出人類智慧行為，包含思考、創造、決策，就可以稱為是 AI，而目前實現人工智慧的技術主流是「機器學習」。在這個定義下我們就能知道，使用繪圖軟體或簡報軟體製作成品不算是 AI，但如果僅輸入了標題它就能自動生成圖片與簡報，那就是 AI。

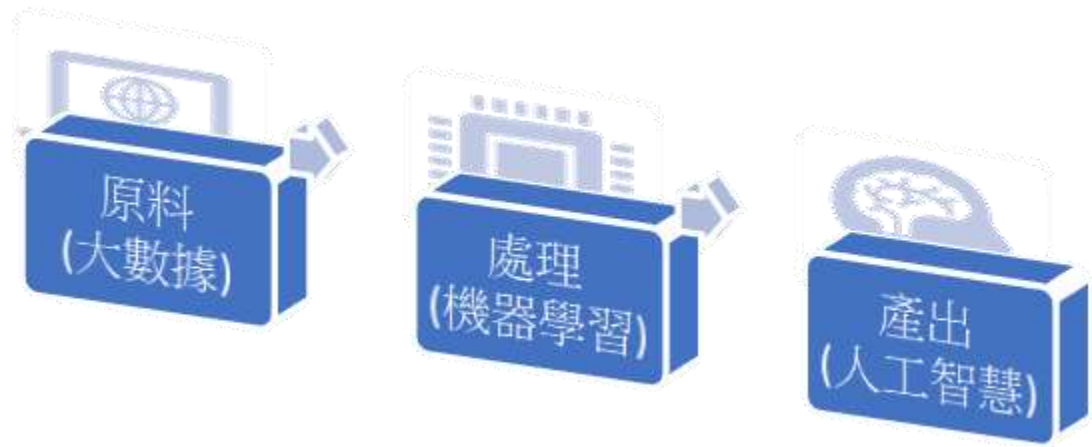
機器學習的定義是，能夠從資料裏頭學習到規則的演算法，然後將這些規則實作在電腦系統中，讓電腦展現看起來有智慧的行為，也因此機器學習需要大量且合適的資料，這時候就輪到大家熟悉的「大數據」登場，讓機器學習以實現人工智慧。換句話說，我們可以把大數據想成是原料，機器學習事處理原料的方法，產出的結果就是 AI，最後再由人類加以運用。

⁷ 同註 3,1-1。

⁸ 董慧明，〈當「混合式威脅」對我國家安全的衝擊和因應〉，《法部務清流雙月刊》(台北市)，第 26 期，法務部，109 年 3 月，頁 6。

⁹ 台灣人工智慧學校，〈台灣產業 AI 化的問題 1 大數據、機器學習與人工智慧〉，<https://aiacademy.tw/definition-ai-industrialization/>，(檢索日期：2023 年 7 月 20 日)

圖 1、AI 的訓練流程原理



資料來源：作者自繪。

而 AI 的終極目標在於「替代人類作出決策」，如果是熟知指參程序的讀者想必都已經聯想到情報四大流程「指導、蒐集、處理、運用」了，我們甚至不難想像出未來 AI 可能直接參考敵軍威脅模式產出數份可行的作戰計畫，連兵棋推演的過程都交由電腦執行，人類只要負責決定採用哪一案就好。

當然不只是戰略層面，AI 的作戰應用方面我國早有研究，同時也是各先進國家的兵家必爭之地，2022 年 12 月國防安全研究院提出的國防科技趨勢評估報告中就有專章說明人工智慧的軍事應用¹⁰。其中說明 AI 軍事運用主要是提升國防部的運作效能，強化競爭力，裡面提到美國認為 AI 的涵蓋範圍包含了作戰、訓練、維持、部隊保護、招聘、醫療保健等，並加速作戰節奏¹¹。

貳、AI 在戰爭中運用的面向

我們以下區分為武力作戰和政治作戰兩方面去闡述 AI 的幾個應用層面。由於本文主要重心還是政治作戰，武力作戰的部分就以簡略的方式為各位說明。

一、武力作戰

(一) 訓練模擬器：目前國軍部隊訓練，尤其是仰賴科技的海軍與空軍已經大量引入模擬器來增加訓練效益，減低裝備傷損，AI 的導入可以進一步強化模擬器性能，導入電腦 AI 對抗模式，磨練人員裝備操作與指參能力。而不同於非 AI 時代的電腦對抗是面對已經設計好的程式，AI 可以針對每次人員的應對進行自我學習以及升級，可以說是真正意義上的「與作戰人員一同

¹⁰ 舒孝煌，〈關鍵軍事科技與台灣國防產業之整合，第二章：人工智慧的軍事應用〉，2022 國防科技趨勢評估報告(台北市)，財團法人國防安全研究院。2022 年 12 月，頁 26-33。

¹¹ 同註 11，頁 28。

成長」，也省去軟體更新的動作。

(二) 無人攻擊武器：AI 可以被應用於自動化和半自動化武器系統，如無人機、自動炮塔和自主地雷等。這些系統能夠在沒有人類直接操作的情況下，執行偵測、攻擊和防禦任務，減少人員傷亡。其中最廣為人知的是無人機。使用無人機減少難以訓練且成本高昂的飛行員已行之有年，主要目前是採用遙控飛機模式，以人員遙控的方式來操作無人機。但是目前使用群體智能演算法操作的無人機在烏俄戰爭中嶄露頭角，已儼然形成未來無人機發展的主要方向¹²。群體智能無人機採取的是低價、大量的蜂巢式攻擊，與操作人員操作單一可掛載多飛彈、多次使用的無人機方式相反，由於其低價的特性可組成自殺無人機群，並且可由 AI 演算索敵後實施攻擊，可大大降低辨識人力與縮短決策過程，甚至直接交由 AI 演算出高價值目標後直接攻擊，同時也不會有單一無人機被擊墜就導致任務無法執行的風險¹³。

(三) 後勤管理：電腦化後勤管理事實上現在已經廣泛的應用在許多大型物流企業之中，不過要稱之為 AI—也就是藉由機器學習達成自主運作似乎還為之尚早，AI 化的理想的狀況當然是能串接油彈消耗及整補，物品使用期限、留存設備資訊、維修紀錄以及異常檢測、物資消耗狀況，甚至能預判戰損。只是以國軍目前的狀態來說，連全面電腦化都還沒有做到。雖然目前我們國軍有電腦後勤系統，但是平心而論，各財產系統間沒有連動整合，尤其越基層越下級其後勤管理能量越缺乏，其中以負責後備動員之單位最為薄弱。筆者曾經擔任過後備動員裝備管理單位的主管，平日裝備管理及庫儲清點、保養全都仰賴少數之現員官兵以及有經驗的士官，每年清點更是必須動員大量人力搬運槍箱開箱清點，損耗時間及人力相當巨大，同時也可能造成一些人為疏漏或意外。但如果能抓住 AI 轉型的機遇，開發適用三軍一體的後勤系統，利用 QR CODE 或條碼確實管理財產進出消耗，甚至開發大型模組化機械搬運設備，使各項後勤系統能達成聯動，並且深入基層，或許能使我國軍後勤管理一日千里。事實上部分經費充足的大型後勤廠庫也確實有這樣的設備，但人力缺乏之後備單位或許是更需要智慧管理設備之所在。

(四) 情報蒐集與研析：目前 AI 在幾個領域已經取得了相當大的成功，其中之一就是人臉辨識，許多 AI 在圖像和影片處理辨識領域上極為優秀，能夠藉由幾張照片來辨識人員是否是同一個人，當然也能有效辨識軍事打擊目標。此外，前面也提到 AI 的運作方式其實極為類似情報指導流程，因此 AI 絕對會是情報研析的絕佳幫手，事實上，2023 年 5 月美國人工智慧開發商

¹² 吳寧康，中央廣播電台，〈負責任的 AI 戰爭？規範軍事人工智能邁出第一步〉，<https://www.rti.org.tw/news/view/id/2159433>，2023 年 3 月 1 日。

¹³ 本段整理自 BBC 文章，大衛·哈布靈，〈下一代無人機將進入「機群」時代〉，<https://www.bbc.com/ukchina/trad/vert-fut-40060866>，2017 年 05 月 26 日。

Scale 宣布其開發的 AI 平台「Donovan」，已交付給美國陸軍第 18 空降軍進行情報分析測試，「Donovan」以「大型語言模型」(Large Language Model, LLM)¹⁴和「人類回饋強化學習」(Reinforcement Learning From Human Feedback, RLHF)¹⁵技術為基礎，其目的在於協助部隊指揮官進行人員調度、標的確認及情勢評估，無須增加人力來消化大量的戰情資料，使指參單位也能快速掌握敵我情資，有效串聯不同部隊並下達行動指令，進而獲得戰場關鍵優勢¹⁶。筆者也試著對目前大家熟知的 Chat gap 進行誘導，使其作出一個簡單的戰略預測，不難看出其所預測的作戰方式其實與我國演習模擬的敵軍威脅模式大同小異，可能是藉由兩岸歷年公開軍事演習資料而做出之結論，而一個專業並且沒有被設下枷鎖的軍事決策 AI 必然能做出更有效的威脅評估及作戰建議，也許未來的戰場情報準備工作將會被 AI 一手包辦，連指揮參謀程序都能大幅縮減流程。

圖 2、誘導 Chat gap 預測中共登陸臺灣的作戰方案



資料來源：作者自行使用 Chat gap 產製，由於 Chat gap 大量抓取體資料庫，其回應經常會無視使用者使用繁體提問，以簡體回答。

¹⁴ 王允中，工研院，〈國際大型語言模型應用技術觀測 [趨勢新知]〉，https://www.moea.gov.tw/Mns/doi/bulletin/Bulletin.aspx?kind=4&html=1&menu_id=13553&bull_id=12454，2023 年 06 月 21 日。

¹⁵ 簡單來說就是就是讓 AI 模型在人類的指導下學習，由人類對 AI 產出的答案評分，選出覺得適當與不適當的回答，除了程式設計人員之外，我們也可以看到 CHAT GAP 也有所謂的「回饋」按鈕，讓使用者提出意見，這使的 AI 的成長更佳的人性化。王道維，風傳媒，〈王道維觀點：當 Google 遇上 ChatGPT——從語言理解的心理面向看 AI 對話機器人的影響〉，<https://www.storm.mg/article/4725780?page=3%20>，2023 年 2 月 11 日。

¹⁶ 軍武頻道，自由時報，〈打造最強大腦 美軍運用 AI 輔助戰場決策〉，<https://def.ltn.com.tw/article/breakingnews/4302304>，2023 年 05 月 15 日。

(五) 資訊防禦與資訊安全：AI 可以說是未來最強的矛與盾，透過網路與大數據成長的 AI 透過網路發動資訊戰，同樣的也必須使用 AI 來加以對抗，AI 可用於網路安全、資訊戰和反無人機系統等，以保護作戰和軍事基礎設施免受敵對實體的攻擊。目前已經有企業組織針對網路安全利用 AI，主要用於偵測、預測與回應攻擊¹⁷，雖然美國電子工程雜誌認為其運用仍有障礙，而其技術想必會持續發展。

二、政治作戰

我們要明白的是，政治作戰是一種不分平戰時，不分前後方的作戰方式，事實上，我們國家可以說是每分每秒都處於攻擊之下，而 AI 加速也加大了這樣的壓力，其主要運用在幾個方面。

(一)假消息：假消息已經完全不是新聞，眾所皆知 2018 年 9 月 4 日燕子颱風造成的「關西機場事件」引發的我國駐日外交官蘇啟誠處長自殺乙事就是假消息運作的結果¹⁸，當時引發人們注意起假消息的傷害，而我國也推出「事實查核中心」來對抗這些假訊息，但時過境遷，人們對假消息的重視度也越來越低，甚至人們開始不在意假消息，因為已經多到來不及也沒有心力查證。而 AI 讓假消息的威力更上一層樓，過往需要使用電影圖片或抽換影像細節造假或部分捏造的方式，在 AI 繪圖軟體的問世後，甚至可以自己生成以假亂真的圖片，其中以 DEEP FAKE¹⁹技術最令人難以應付。

在烏俄戰爭中俄羅斯就以此技術造假烏克蘭總統澤倫斯基勸降影片²⁰，動搖烏克蘭軍隊民心，這個影像在全球主要媒體因為被認為是假影片而遭到封鎖下架，當然俄羅斯社群媒體以及一些隱密性高的通訊軟體間這個影片則受到大力宣傳。假消息廣泛的被使用來煽動社會不穩定、製造混亂，或影響選舉等政治活動。因為網紅換臉事件，我國也於 2023 年初完成立法三讀通過²¹，賣臉換成人片可重判 7 年。但該次修法著重於「性隱私」及「性造謠」，對於換臉讓人們說出「他沒說的話」這件事情仍然沒有明確性的罰則。實務

¹⁷ Ludovic Rembert, <AI 實現網路安全：炒作還是現實？>，<https://www.edntaiwan.com/20221107nt31-ai-based-cybersecurity-hype-or-reality/>，2022 年 11 月 07 日。

¹⁸ 朱冠瑜，風傳媒，<蘇啟誠事件關鍵查核>日本關西機場證實：中國領事館想派巴士，但遭日方拒絕>，<http://www.storm.mg/article/497459>，2018 年 9 月 15 日。

¹⁹ Deepfake 源自於英文「deep learning」（深度學習）和「fake」（偽造）組合，主要意指應用人工智慧深度學習的技術，合成某個（不一定存在的）人的圖像或影片、甚至聲音。TingWei，泛科學，<Deepfake 不一定是問題，不知道才是大問題！>，<https://pansci.asia/archives/342421>，2020 年 01 月 24 日。

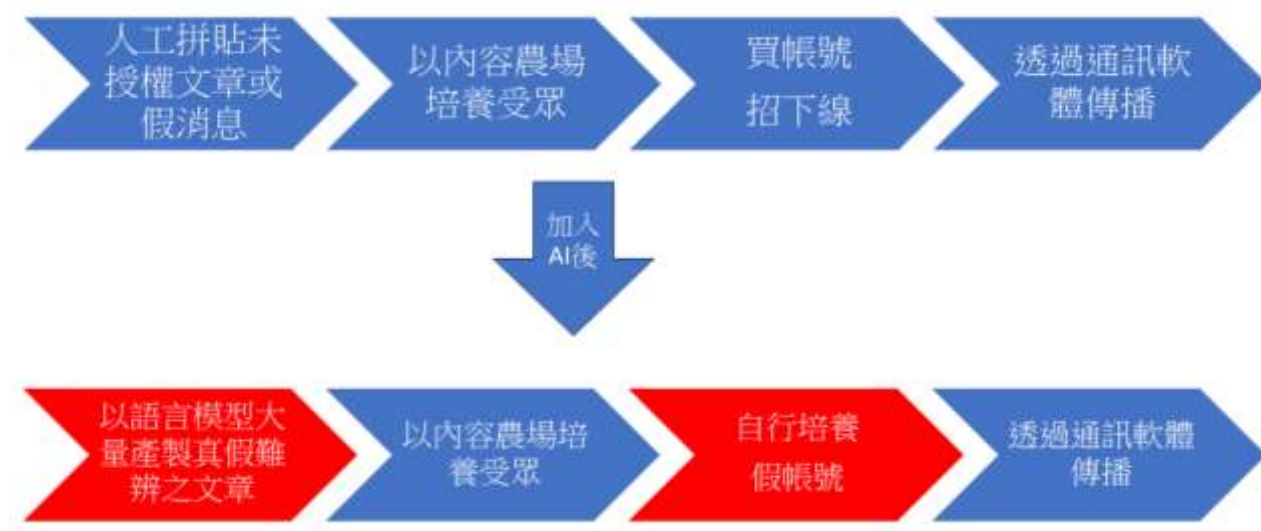
²⁰ 中央社，<澤倫斯基深偽影片流傳 專家憂是俄資訊戰冰山一角>，<https://www.cna.com.tw/news/aopl/202203170415.aspx>，2022 年 03 月 17 日。

²¹ 程遠述，聯合新聞網，<立法院刑法三讀通過，賣換臉成人片最重判 7 年>，<https://udn.com/news/story/6656/6893269>，2023 年 01 月 07 日。

上恐怕只能用毀謗罪來起訴，而在如何預防這件事情上也可以說是一籌莫展，目前 AI 生成影片還是有某種程度的瑕疵，如上所提澤倫斯基的影片因為嘴型與說話的內容對不上，很快就被識破，並且以烏克蘭國家通訊社以及親西方社群媒體如推特、臉書等加以闢謠，但隨著時間過去，AI 克服技術障礙產生難以識別真假的影片恐怕也是指日可待的事情。

(二)社交媒體操作：基本上方法有兩種，一種是自動生成和傳播文章，「內容農場」模式就是經典的操作方式，平常以到處拼貼的文章培養讀者，再加入真假難辨的假訊息，或以七分真三分假的方式偷渡真假難辨的消息，再藉由培養出的讀者以通訊軟體傳散文章。第二種則是在新聞或社群軟體上以數量驚人的帳號留言評論、按讚、轉貼，以製造虛假的「輿論」。以往網軍還需要是真人操作，但未來只要寫好程式，使用 AI 就能製造出足以使政府做出錯誤的決策判斷的「人氣」。這樣的模式可以說是無往不利，報導者也曾經做過專業的報導²²說明此一操作帶來的民主危機。AI 的用處可以說是可以進一步的加速這個流程，移除中間的操作人員，直接以 AI 完成整套作業模式。我們大概可以將加入 AI 之後的運作模式繪製成下圖，不過這樣的替代與原始模式是可以並存的。購買長久經營，有固定受眾的粉專跟從頭培養一個粉專之間的效益畢竟有所差異，由於臉書演算法的限制，導致往往有些粉專中途易主，使用者也難以察覺，但其輿論操作的力量藉由我們每日滑動手機的動作滲透我們每個人的日常。

圖 3、加入 AI 後的內容農場運作



²²劉致昕，柯皓翔，許家瑜(2019)，報導者，〈直擊鬼島狂新聞、全球華人聯盟背後的内容農場帝國〉，<https://www.twreporter.org/a/information-warfare-business-disinformation-fake-news-behind-line-groups>，2019 年 12 月 26 日。

在留言評論導引輿論的方面，目前還是主要以人工來辦理，也因此偶爾也可以見到「翻車」的情形，國立教育廣播電台曾專訪臺灣人工智慧實驗室創辦人杜奕瑾，對於目前假消息引導輿論的模式提出說明，並定名為「協同行為」，主要就是在面對特定議題使用特定口號，並且同進同出²³，而當 AI 進一步的介入這個協同行為中時，想必將更難發現其輿論操作的模式，無法分辨人為操作和人民自由意志所生輿論間的差異。我們不難想像公關公司可以平常就申請數千個帳號，以 AI 讓這些帳號平時就以特定的模式發文，彷彿一個真實存在的人物去培養受眾，在關鍵時刻引導輿論，以虛擬的方式導引真實世界。

(三)社交媒體監控：AI 可以用於監控社交媒體上的用戶活動和意見，從而洞察公眾輿論和政治趨勢。這些資訊可能被用來調整政治策略或傳播特定信息。基本上目前的做法就是使用網路爬蟲來擷取網站社群的資料。Python 爬蟲是一種程式工具，常見手法是透過爬蟲（Spider）模擬使用者瀏覽目標網頁，針對網頁中細部資料，自動抓取所需資訊。在正常情況下，於搜尋引擎輸入關鍵字後，必須手動點進各項目連結取得對應資訊，而 Python 爬蟲可以有效地針對資料類別與名稱，快速擷取使用者想要的資訊²⁴。這是目前極為主流的大數據獲取方式，如果讀者對於我們前面的文章還有印象，大數據就是訓練 AI 的重要元素，用大數據訓練 AI，用 AI 獲取大數據、分析大數據，在藉由分析結論作出決策，決定適當的輿論操作方式，製造人們的資訊盲區，進而操縱認知。我想讀者們應該都聽過「同溫層效應」這個名詞。這原本是氣象學中的用詞，「同溫層」是指大氣層中的平流層，在平流層裡，大氣基本保持水平方向流動，較少有垂直方向的流動。後指社群網站透過數據演算，傳送使用者偏好的訊息，讓使用者越來越封閉在特定社群，若有不同認知出現，就會產生認知衝突²⁵。AI 可以分析用戶數據，根據特定的用戶興趣和政治傾向來投放個性化的政治廣告和宣傳。這樣的廣告可能在社交媒體和在線廣告平台上出現，以影響公眾的政治看法和選擇。

(四)數據分析和預測：AI 可以用於政治選舉和政策預測。通過分析大量的選舉和政策相關數據，AI 可以幫助政治家和政黨了解公眾意見，並預測選

²³ 國立教育廣播電台，〈【新聞真假辨】專訪杜奕瑾（台灣人工智慧實驗室創辦人）：「PTT 創世神」如何用 AI 迎擊資訊戰網軍？PTT 如何實現民主？（逐字稿大公開）〉，<https://tfc-taiwan.org.tw/articles/8002>，2022 年 8 月 9 日。

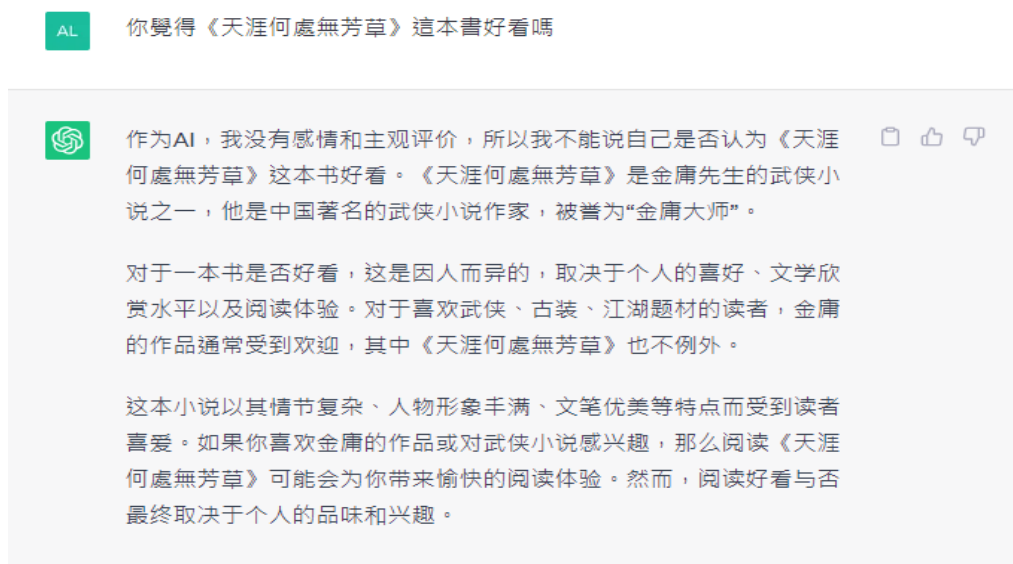
²⁴ 巨匠電腦，〈Python 爬蟲實作觀念篇：想進入 AI 產業必須先認識這些工具！〉，<https://www.pcschool.com.tw/blog/it/web-scrawler-in-python>，2023 年 2 月 22 日。

²⁵ 翰林雲端學院，〈同溫層效應〉，<https://www.ehanlin.com.tw/app/keyword/%E5%9C%8B%E4%B8%AD/%E5%9C%8B%E6%96%87/%E5%90%8C%E6%BA%AB%E5%B1%A4%E6%95%88%E6%87%89.html>，檢索日期：2023 年 7 月 20 日

舉結果和政策影響。並且以民調的方式展現來影響選民意志，操縱選舉結果。臺灣人普遍被認為有「西瓜偎大邊」的個性，當認為自己的主要選擇可能沒辦法勝過主要對手時，就會從而依靠次要選擇以打敗主要對手，在許多被稱為「三角堵」的選戰中都可以窺見此一選舉傾向，也就是俗稱的「棄保效應」，也因此公職人員選罷法第 53 條第三款規定「政黨及任何人自投票日前十日起至投票時間截止前，不得以任何方式，發布、報導、散布、評論或引述前二項資料(民調)」，但仍然難以防堵藉由賭盤以及耳語的方式來操作的方法，這也是我們提到的「內容農場」操作模式以深入通訊軟體為目標的原因。

(五)特定價值觀的資訊引導：AI 是一種有趣的工具，但是如果使用者把 AI 所提供的訊息當成「真理」，就會自行陷入錯誤，以知名的 Chat gap 為例，不少人把 Chat gap 當成蒐尋軟體在使用，諮詢它各式各樣的問題，甚至把它回答的內容當成「正確」的內容，但事實上，Chat gap 僅僅是一個「語言模型」，它運用了自然語言處理²⁶ 來讓你相信它所說的話，它不是搜尋軟體，亦不保障內容正確性，如果你詢問了錯誤的問題，它甚至會給你「看起來正確的答案」，比如筆者詢問 Chat gap 《天涯何處無芳草》是否好看，它的回答令人莞爾，因為全世界茫茫書海或許有同名小說，但金庸顯然並沒有創作這樣的作品，但它仍然煞有介事地回答了使用者這個問題，足以證明其所回應之內容並不保證正確。另外，由於目前中文資料容易大量抓取簡體資料庫，也會造成其回應內容的偏誤，同樣的途中也能發現，筆者以繁體詢問，AI 卻以簡體回答的情形。

圖 4、Chat gap 錯誤回答的範例



²⁶自然語言處理 是一種機器學習技術，讓電腦能夠解譯、操縱及理解人類語言。如今，組織擁有來自各種通訊管道的大量語音和文字資料，例如電子郵件、簡訊、社交媒體新聞摘要、影片、音訊等。他們使用 NLP 軟體來自動處理此資料，分析訊息中的意圖或情緒，並即時回應人類通訊，摘自：AMAZON ASW, <甚麼是自然語言處理>， <https://aws.amazon.com/tw/what-is/nlp/>，檢索日期：2023 年 7 月 20 日。

此外，它所回饋的內容被設計者加上了限制與枷鎖，控制它的回應傾向，並且影響使用者的價值觀，而其並不一定符合特定使用者的價值體系。若詢問爭議問題，其得出的答案會盡量在邏輯上偏向使閱聽者認為回答中立而且極力去除感情因素，對於相對無立場的使用者來說，其四平八穩的言論則相當有價值，從而認定 AI 的價值觀與自己一致，更加相信 AI 所提出的結論，這樣的使用者就可能相信 AI 所給予的答案，在符合價值觀的狀況下相信其所給予的所有回應，達到認知操作的功能。目前 Chat gap 的主流價值體系顯然符合英美的主流價值體系。

也正因為如此，為了確保 AI 軟體符合「某些團體、國家」的價值觀，或避免觸及敏感議題，其團體國家必會推出自己的 AI 軟體與之抗衡，比如中共就以網路長城的方式禁止中國人民使用 Chat gap，並另外推出了「文心一言」與之對抗，而一如我們預料的，如果詢問天安門和維吾爾的議題就會卡住，拒絕回答²⁷。

參、政治作戰在 AI 未來的展望

上一段我們說明了 AI 在武力作戰和政治作戰方面的運用，而武器這種東西是每個陣營都可以操作使用的，端看我們如何去運用它，以及有沒有運用它的能力，我想提出以下幾點建議

一、與民主國家的積極合作

與民主國家之合作一直以來都是我國的外交重點，自 2001 年澳洲記者克雷格愛迪生（Craig Addison）所撰寫的《矽屏障—臺灣最堅實的國防》（Silicon Shield）一書之後，「矽盾」一詞廣泛地出現在各國的國與區域安全的討論中，許多人都認為，我國堅強的半導體技術以及勤勞不懈的半導體人才，使臺灣免於中國大陸軍事威脅，得以捍衛完整主權。中共自然也明白這個問題，因此不斷的挖角我半導體工程師、派遣經濟間諜，只是在我國半導體產業的積極防禦下，得以維持穩健的運作。近來因為美中博弈，我台積電開始踏出國門，至其他民主盟邦（特別是美國）設廠，當然也引發了國內的疑慮，懷疑這是否是挖空「護國神山」²⁸。但深一層看，設廠並非搬空產能，高階

²⁷ 中央社，〈文心一言能寫唐詩 問習近平天安門維吾爾就卡住〉，<https://www.cna.com.tw/news/ait/202303210084.aspx>，2023 年 3 月 1 日。

²⁸ 林哲良，風傳媒，〈台灣半導體優勢還能撐多久〉，<https://events.storm.mg/campaign/SiliconShield/main.html>，檢索時間 2023 年 7 月 25 日。

製程仍然根留臺灣，將較低階的製程移往國外合作並無不可，且產能外移也可以稀釋我國能源負擔，更能展現我與民主盟邦的合作誠意。目前台積電也傳出可能到日本熊本設廠的新聞，也足見我國對兩大最重要民主盟邦的重視。

二、AI 人才的培養

雖然 AI 的目標是取代人力，取代人類決策，但在此之前還是依靠人類製造出先進 AI，因此 AI 人才的培養無疑是未來產業上的重點，當然我國也已經有所作為，藉由產、官、學之合作來培養未來人才，日前智榮基金會董事長施振榮、雷鳥全球管理學院院長 Dr. Sanjeev Khagram、東吳大學校長潘維大三方共同簽署「半導體產業國際人才培育計畫」合作備忘錄，攜手為台灣半導體產業國際管理人才的需求展開超前部署。²⁹人才培育的腳步不可停歇，防止人才被挖角的措施當然也必須同步進行。

三、主動出擊的媒體政策

前段提到以 AI 製造真假難辨之訊息配合內容農場以混淆認知的方法必然成為未來主流，既然如此，我國何仿不同樣利用 AI，建構以事實為主的龐大媒體鍊？事實上美國之音也訪問我中央社，希望能推動因應假消息的事實查核合作計畫。³⁰我國之所以在認知戰的戰場上節節失守，就是因為己身能量不足，不論如何訴諸愛國情操，總是敵不過中共透過娛樂、傳媒、跨國網購以及意識形態所發起的攻擊，然而俗話說「攻擊就是最好的防禦」，以國家為首，與民主國家合作跨國媒體鍊，藉由 AI 翻譯來彌補人力之不足，絕對是可行的方向。

四、國軍資訊政策的重整

前面提到了未來的方向以及國家已經在進行的政策，但是遺憾的是這樣的改革之風以及培養之風，並未吹進軍中。我國軍一直以來因為保守的資訊政策而遭到詬病，不論是電腦設備的落後，資安政策的過度限制導致了許多國軍人員在工作進行上綁手綁腳，我政戰人員算是被限制頗重的一群，畢竟操作戰車或是砲車、野外求生都可以不用科技設備，但是製作文宣影片卻不用不行。國軍希望基層擁有製作影片文宣之能力以強化人才招募以及國軍形象宣傳已經行之有年，但是基層遭資安政策限制以至於資源匱乏、能力匱乏，

²⁹ 鐘張涵，聯合新聞網，〈施振榮攜美國雷鳥學院發起半導體產業國際人才培育計畫〉，https://udn.com/news/story/7240/7325987?list_ch2_index，2023 年 7 月 26 日。

³⁰ 中央社，〈美國之音訪中央社 盼推動 3 項合作計畫〉，<https://www.cna.com.tw/news/ahel/202303290369.aspx>，2023 年 3 月 29 日。

卻又必須達成上級交付之任務所帶來之壓力至今仍未改善，我國軍實應重新認知並面對資安政策與未來趨勢，審慎邀集參考基層意見，莫再故步自封。

五、立法監管網路犯罪及認知戰問題

日前我國 NCC 提出「數位中介法」草案引來不小波瀾³¹，幾個大型數位服務提供者反彈甚鉅，網路上一片撻伐，質疑是否侵害言論自由，最後 NCC 撤回草案，執政黨迄今也沒有提出後續的修正草案或規劃。個人認為非常可惜，因為很明顯的，極權國家藉由民主國家的資訊、政治與言論之開放來打擊、分裂民主國家，已經是不容忽視的問題。像是 2023 年 5 月 3 日，美國國會以全球資訊戰為題召開聽證會，針對俄國、中國等威權政府投注大量資源輸出假訊息與宣傳敘事，型塑對其有利的資訊環境以實現其政治目標，提出可行的對抗策略。³²我國國防安全研究院劉姝廷政策分析員認為，民主國家之因應趨勢主要是「強化以事實為基礎的傳播策略」³³，而在語言相通的狀況下，這方面的優勢與劣勢事實上是一體的，因此我們極為需要網路言論監管法案，臺灣事實查核中心董事羅世宏先生亦以專文投書清流月刊表示：「期待超大型社群媒體／數位平臺犧牲獲利、自動承擔社會責任並不實際。」³⁴，顯然立法管制已經是迫在眉睫，刻不容緩的事情。而以國軍角度而言，自然難以觸及立法問題，但對我國軍人員強調「識假」，認識信息戰以及破除信息戰之教育，仍可落實於政治教育以及資安防護中。

六、重視已然成形的認知侵略

認知侵略已非一朝一夕，冰凍三尺絕非一日之寒，上面提到我們與中國大陸語言相通，早在許多用語上已被滲透侵略，而用語上的滲透侵略容易導致的是觀念的扭曲與認知的歪斜。20 世紀羅馬尼亞哲人蕭沅(Emil Cioran，1911 年 4 月 8 日—1995 年 6 月 20 日)曾經提到「我們不是生活在一個國家，而是生活在一個語言之中，祖國，僅此而已(On n'habite pas un pays, on habite une langue. Une patrie, c'est cela et rien d'autre.)」這位虛無主義的

³¹ 黃順祥，新頭殼，〈《數位中介法》NCC 放棄了！言論自由、網路犯罪拿捏「重新來過」 律師 4 舉例秒懂：草案問題非常多〉，<https://www.business today.com.tw/article/category/183027/post/202208210001/>，2022 年 9 月 7 日。

³² 自由亞洲電台，〈全球資訊戰：美國會關注如何對抗中俄虛假資訊〉，<https://reurl.cc/65Vz7y>，2023 年 5 月 3 日。

³³ 劉姝廷，國防安全研究院，〈國家安全雙週報第 80 期：民主國家在全球「資訊戰」下可有的策略〉，<https://indsr.org.tw/respublicationcon?uid=12&resid=1961&pid=4007>，2023 年 5 月 26 日。

³⁴ 羅世宏，〈社群媒體/數位平臺需要受到法律與公眾監督〉，《清流雙月刊》(台北市，法務部)，第 38 期，2022 年 5 月 38 日。

哲學家提到了語言的重要性，我們與中國大陸使用同樣語言的狀況下所受到的認知侵略必須被重視，尤其必須被我們的文宣機關重視。日前莒光園地因使用了「質量」一詞來形容步槍射擊訓練的成效而遭到立委撻伐，³⁵看起來雖然只是小事，但事實上卻足以證明我們受到的認知侵略是無孔不入的。從新聞媒體開始習慣使用「視頻」而非「影片」，國小字典用「土豆」稱呼馬鈴薯，這些事情看似吹毛求疵，但事實上卻是使國人加速的落入認知戰的陷阱中，不可不慎！

七、建立國軍正確資安及個資保密觀念

個資保密與資安保密不斷被強調，而在 AI 時代下的個資保密不只越困難，人們的個資保密概念也因為普及的社群軟體而稀薄，筆者是 70 年代的人，當時還記得父母千叮萬囑網路上不應該隨意地散佈自己的照片以及本名還有住址，但是現代人們對自己的個人資料以及隱私、長相甚至居住所都不避諱的展露，委實令人憂心。此外，筆者對於我國軍文宣體系的穿軍服讚系列也抱持著懷疑態度，要知道現代社會可以靠著幾張照片用 AI 運算一個人的不雅照，日前就有詐騙集團將大專院校官網上的教授面孔合成不雅照後寄勒索信³⁶，不難想像這樣的魔手伸到以軍譽為重的國軍官兵手上來也是遲早的事情，我們的文宣措施是否應該要思考到新的 AI 技術對我國軍甚至軍眷所帶來的隱性危害，值得深思。

肆、結論

我想，我們可以誇張地說「21 世紀是 AI 的世紀」，而因應 AI 的世紀，AI 軍事應用的未來，國軍在資訊政策上重新思考定位是勢在必行的事情，一昧的禁止資訊科技的使用是否會帶來國軍人員在資訊科技應用上的僵化？仰賴資訊開放跟人民素質的文宣政策是否反而有收緊的必要？我國軍媒體是否應更加注意製播成品的內容，不要為了貼近年輕用語而喪失了自己的堅持與任務？目前本軍政治教育多仰賴莒光園地電視教學以及忠誠報等政令宣導，各處組為了利用年輕人熱愛短影片之流行，也希望基層能製作短影片，但在資訊政策以及設備上完全跟不上，反而造成許多資安問題，或囿於資安問題而無法執行。AI 是個用具，是個一日千里而且能夠變型進化的用具，但

³⁵ 謝君臨，自由時報，〈「莒光日」出現中國用語 王定宇：滲透和破壞常從小事開始〉，<https://news.ltn.com.tw/news/politics/breakingnews/4344071>，2023 年 6 月 25 日。

³⁶ 許維寧，聯合報，〈變臉惡照詐騙全台教授 北醫大：公開照片將加浮水印〉，<https://udn.com/news/story/7315/7054755>，2023 年 3 月 24 日

其中共目的沒有改變，改變的只有工具，同樣的，我國軍保國衛民的目的也沒有改變，但我們要認真的思考，我們有沒有因應的工具，以及如何去創造機勢，主動出擊。

伍、參考文獻

中文：

- 1.國防部，《國軍政治作戰要綱》(台北：國防部，民國 105 年)，頁 1-1。
- 2.黃柏欽，《戰爭新型態－「混合戰」衝擊與因應作為》，《國防雜誌》(台北市)，第 34 卷第 2 期，2019 年 6 月。
- 3.董慧明，〈當「混合式威脅」對我國家安全的衝擊和因應〉，《法部務清流雙月刊》(台北市)，第 26 期，法務部，109 年 3 月，頁 6。
- 4.台灣人工智慧學校，〈台灣產業 AI 化的問題 1 大數據、機器學習與人工智慧〉，<https://aiacademy.tw/definition-ai-industrialization/>，(檢索日期：2023 年 7 月 20 日)
- 5.舒孝煌，〈關鍵軍事科技與台灣國防產業之整合，第二章：人工智慧的軍事應用〉，2022 國防科技趨勢評估報告(台北市)，財團法人國防安全研究院。2022 年 12 月，頁 26-33。
- 6.吳寧康，中央廣播電台，〈負責任的 AI 戰爭？規範軍事人工智能邁出第一步〉，<https://www.rti.org.tw/news/view/id/2159433>，2023 年 3 月 1 日。
- 7.大衛·哈布靈，〈下一代無人機將進入「機群」時代〉，<https://www.bbc.com/ukchina/trad/vert-fut-40060866>，2017 年 05 月 26 日。
- 8.王允中，工研院，〈國際大型語言模型應用技術觀測 [趨勢新知]〉，https://www.moea.gov.tw/Mns/doit/bulletin/Bulletin.aspx?kind=4&html=1&menu_id=13553&bull_id=12454，2023 年 06 月 21 日。
- 9.王道維，風傳媒，〈王道維觀點：當 Google 遇上 ChatGPT——從語言理解的心理面向看 AI 對話機器人的影響〉，<https://www.storm.mg/article/4725780?page=3%20>，2023 年 2 月 11 日。
- 10.軍武頻道，自由時報，〈打造最強大腦 美軍運用 AI 輔助戰場決策〉，<https://def.ltn.com.tw/article/breakingnews/4302304>，2023 年 05 月 15 日。
- 11.Ludovic Rembert，〈AI 實現網路安全：炒作還是現實？〉，<https://www.edntaiwan.com/20221107nt31-ai-based-cybersecurity-hype-or-reality/>，2022 年 11 月 07 日。
- 12.朱冠瑜，風傳媒，〈蘇啟誠事件關鍵查核〉日本關西機場證實：中國領事館想派巴士，但遭日方拒絕〉，<https://www.storm.mg/article/497459>，2018 年 9 月 15 日。
- 13.TingWei，泛科學，〈Deepfake 不一定是問題，不知道才是大問題！〉，

<https://pansci.asia/archives/342421>，2020 年 01 月 24 日。

14. 中央社，〈澤倫斯基深偽影片流傳 專家憂是俄資訊戰冰山一角〉，<https://www.cna.com.tw/news/aopl/202203170415.aspx>，2022 年 03 月 17 日。

15. 程遠述，聯合新聞網，〈立法院刑法三讀通過，賣換臉成人片最重判 7 年〉，<https://udn.com/news/story/6656/6893269>，2023 年 01 月 07 日。

16. 劉致昕，柯皓翔，許家瑜(2019)，報導者，〈直擊鬼島狂新聞、全球華人聯盟背後的内容農場帝國〉，<https://www.twreporter.org/a/information-warfare-business-disinformation-fake-news-behind-line-groups>，2019 年 12 月 26 日。

17. 國立教育廣播電臺，〈【新聞真假辨】專訪杜奕瑾（台灣人工智慧實驗室創辦人）：「PTT 創世神」如何用 AI 迎擊資訊戰網軍？PTT 如何實現民主？（逐字稿大公開）〉，<https://tfc-taiwan.org.tw/articles/8002>，2022 年 8 月 9 日。

18. 巨匠電腦，〈Python 爬蟲實作觀念篇：想進入 AI 產業必須先認識這些工具！〉，<https://www.pcschool.com.tw/blog/it/web-scraper-in-python>，2023 年 2 月 22 日。

19. 翰林雲端學院，〈同溫層效應〉，<https://www.ehanlin.com.tw/app/keyword/%E5%9C%8B%E4%B8%AD/%E5%9C%8B%E6%96%87/%E5%90%8C%E6%BA%AB%E5%B1%A4%E6%95%88%E6%87%89.html>，檢索日期：2023 年 7 月 20 日

20. AMAZON ASW，〈甚麼是自然語言處理〉，<https://aws.amazon.com/tw/what-is/nlp/>，檢索日期：2023 年 7 月 20 日。

21. 中央社，〈文心一言能寫唐詩 問習近平天安門維吾爾就卡住〉，<https://www.cna.com.tw/news/ait/202303210084.aspx>，2023 年 3 月 1 日。

22. 黃佩君，自由經濟，〈如何讓台「矽屏障」不消失？經部百億推動關鍵設備材料在地化〉，<https://ec.ltn.com.tw/article/breakingnews/3216158>，2020 年 07 月 02 日。

23. 林哲良，風傳媒，〈台灣半導體優勢還能撐多久〉，<https://events.storm.mg/campaign/SiliconShield/main.html>，檢索時間 2023 年 7 月 25 日。

24. 鐘張涵，聯合新聞網，〈施振榮攜美國雷鳥學院發起半導體產業國際人才培育計畫〉，https://udn.com/news/story/7240/7325987?list_ch2_index，2023 年 7 月 26 日。

25. 中央社，〈美國之音訪中央社 盼推動 3 項合作計畫〉，<https://www.cna.com.tw/news/ahel/202303290369.aspx>，2023 年 3 月 29 日。
26. 黃順祥，新頭殼，〈《數位中介法》NCC 放棄了！言論自由、網路犯罪拿捏「重新來過」 律師 4 舉例秒懂：草案問題非常多〉，<https://www.businesstoday.com.tw/article/category/183027/post/202208210001/>，2022 年 9 月 7 日。
27. 自由亞洲電台，〈全球資訊戰： 美國會關注如何對抗中俄虛假資訊〉，<https://reurl.cc/65Vz7y>，2023 年 5 月 3 日。
28. 劉姝廷，國防安全研究院，〈國家安全雙周報第 80 期：民主國家在全球「資訊戰」下可有的策略〉，<https://indsr.org.tw/respublicationcon?uid=12&resid=1961&pid=4007>，2023 年 5 月 26 日。
29. 羅世宏，〈社群媒體/數位平臺需要受到法律與公眾監督〉，《清流雙月刊》(台北市，法務部)，第 38 期，2022 年 5 月 38 日。
30. 謝君臨，自由時報，〈「莒光日」出現中國用語 王定宇：滲透和破壞常從小事開始〉，<https://news.ltn.com.tw/news/politics/breakingnews/4344071>，2023 年 6 月 25 日。
31. 許維寧，聯合報，〈變臉慾照詐騙全台教授 北醫大：公開照片將加浮水印〉，<https://udn.com/news/story/7315/7054755>，2023 年 3 月 24 日

英文

1. Rachael Burton, <Disinformation in Taiwan and Cognitive Warfare,> Global Taiwan Brief, Vol. 3。 <http://globaltaiwan.org/2018/11/vol-3-issue-22/>>(2018 年 12 月 14 日)