

科技部人文社會科學研究中心

學術研究群成果報告

人工智慧與法律規範學術研究群

學術研究群編號：MOST 107-2420-H-002-007-MY3-SG10708

學術研究群執行期間：107 年 7 月 1 日至 108 年 6 月 30 日

學術研究群召集人：李建良

執行機構及系所：中央研究院法律學研究所

中 華 民 國 108 年 7 月 30 日

中文摘要

人類運用智慧製造出機器，更進一步讓機器學習並運用人類的智慧，正是人類智慧的展現。經由人造機器表現的智慧（人工智慧）改善人類生活品質，成為人類政治、經濟、社會與文化模式的重要元素之一，由來已久。隨著人工智慧研發進化、運作方式與應用領域的不斷擴增，可以預見將對人類社會各種面向與層次帶來不容小覷的衝擊與變革，無疑也會帶來或顯或潛的社會問題、倫理議題與法律爭議。科技發展與法律規範相互碰撞，牽動的不會只有特定的科學技術領域，更大程度地可能是對規範板塊與法秩序的衝撞與推移。法律人的基本思維是否受到波及？法制運行如何知所因應？法律學門是否在臺灣乃至世界的人工智慧研究圈與學術鏈中扮演關鍵性的角色？如何進行法律學門內部與跨學門之間的整合研究？哪些因人工智慧而起的法律與風險等議題需要研析探索防範，以降低人工智慧應用的負面衝擊與社會成本？面對紛亂複雜的多方各種法律爭議，乃有以「人工智慧與法律規範」為主題建立法律學術研究群構想的緣起。

本學術研究群的建立與運作，旨在嘗試集臺灣法學界群體之力，以憲法、行政法、刑事法、民事法、商事法、金融法、智財法為經，以法哲學、國際法、區域法、法律史為緯，結合國內關注此一領域且學有專精的法律學者，共同投入人工智慧法律問題的研究，進行廣度與深度兼備的理論探索與實務分析。本學術研究群首先將分別耙梳各法領域與人工智慧的交互關係，依各領域的規範特性，分析人工智慧對各該法領域產生的影響與法律問題，並以法哲學、法律史與國際法/區域法作為橫向的連結，進行貫穿各該法領域的基礎研究。其次，延續基礎研究的成果，擇定若干主題，例如自駕車的法律責任問題，進行深入研究。接著是法律學門自身的融貫研究與跨學門領域的科際整合，前者是憲法、行政法、刑事法、民事法、商事法、金融法、智財法、法哲學/法律史、國際法/區域法之間的問題整合與關聯研究；後者是法律學門與其他領域如資訊科技、哲學、倫理學、社會學等學門之間的問題比較與觀點砥礪。

本學術研究群採定期聚會與不定期座談的方式進行，以一年為期，每月固定聚會一次，共計 12 次，並視邀請參與學者的時間而定，舉行兩次中型座談會，與法律學群及其他學科領域交流與對話。定期聚會除依前述規劃的子題進行文獻研讀討論外，將由成員輪流提出報告或邀請相關學者專家演講，同時陸續將活動中演講或報告的相關資料撰寫成論文，發表於期刊或研討會，或集結成專書，作為本學術研究群的計畫成果。

第一次會議紀錄	
時 間	107 年 7 月 30 日（星期一）
主 題	人工智慧在空品物聯網上的應用
內容摘要	
講者：陳伶志研究員（中央研究院資訊所） 本次會議的主題是物聯網（物品連上網路後，並產生價值，主要是透過統計分析、人工智慧）如何應用於空氣品質（空品）的監測與改善，特別是室外的空氣污染—PM2.5（小到可以進入身體且人體無法排出，累積後造成心肺負擔）。陳研究員指出：過去監控機制（環保署測站）是放在距地面至少 10 公尺、空曠通風良好的地方，且需避開污染源（不希望量測到極端值），因此。與人們實際上呼吸到的空氣不同。基於精確度、準確度的考量，2012 年開始發展、試驗每個人都可以擁有（便宜）、攜帶（小）的監測器，希望可以做到確知實際呼吸到的空氣品質，並且實驗室負責分析、製圖，開放的硬體把取得資料公開，提供民眾即時的空氣監測，把它變成淺顯易懂的圖，比如說可以抓到霾害過境的路徑，或了解人們集體拜拜對空氣品質造成的影響。問題層次有三： 一、什麼事發生了？為什麼會發生？要找出哪些機器是可信的（民眾未必會放在戶外），找出空間性的異常（方圓 3 公里以內，空氣品質應該是差不多的）、時間性的異常（短時間內衝高），長時間統計後，判斷該感測器是否可靠，亦可進一步判斷是何種異常：室內、污染源。 二、未來會發生什麼？關注未來 4.5 個小時的空氣預報（在此範圍內，無須考慮大氣模型）兩難：Accuracy 準確性 v. Scalability 運算速度。模型有：1. Hybrid(NNAR+ARIMA)：非常準，但速度慢；2. Matrix Optimization：非常快（30 秒內做完全台灣 2000 台測站），但非常不準；3. Cluster-based Hybrid：用特徵演算法，找出特徵類似者，把學習到的套用在相同特徵的人上面用地理來切割，把學習到的套用在附近的人上面。 三、影響未來會發生什麼：如果可以成功預測，就會影響到你今天傍晚要不要去跑步，甚至推薦你最佳路徑。 陳研究員同時提到採取的研究取徑是一種公民科學，各地區推廣空氣盒子，學校老師、學生、一般民眾對資料的利用，建立生態圈：中研院—FB 社群—台北市智慧城市展（產業：空氣盒子專案）—廠商捐贈給各地方政府—供各中小學環境教育，發展-資料分析-跨領域合作，學術研究-公民科學-公私協力-共創共好，成為物聯網＋大數據＋人工智慧的練兵場，產官學研＋公民合作新典範：開放、合作、共學、共創、共好。 Q&A： 1. 機器想要偵測不同的污染時，是否都需要調整？ 加裝偵測揮發性有機氣體的感測器，需要客製化。 2. 用於環境訴訟？ 針對一般民眾，而非針對特定污染，用於訴訟。 3. 工廠是否會用這些資料來反證污染沒有那麼嚴重？ 我們的機器不準，只能提醒個人，要用於執法，必須有更準確的儀器。環保單位最初不認	

同，但目前已經知道要怎麼使用物聯網了（抓偷排）。

4. 目前只是參考值，是否可能提升到客觀準確的程度？

民間先試著做，成熟後政府才會有意願。

5. 是否有調查過空氣盒子的分佈，分布不均是否會影響資料的準確？

參與式的，沒辦法先過濾；分布不均可能是人口。技術上面有誤差、規範上面無法反映目前人類所遭遇的危害，要分開來看，後者的問題是沒辦法透過前者來解決。

6. 現在的物聯網是否有辦法解決規範上的問題？

資料不準確，所訓練出來的 AI 就未必正確，在未來的發展會受限，只能建議，不具有強制力。

7. 這樣不準確資料的預測，是否會造成危害？

環保署的是小時值：累積重量來量測污染重量；物聯網做的是五分鐘的值，兩者不能相提並論。且我們應該試著去使用這樣不準確的資料。大數據本來就容許在不那麼精準的資料中，發展出有價值的東西。

第二次會議紀錄

時 間	107 年 8 月 27 日（星期一）
主 題	Artificial intelligence in the Court: Black Box, Due Process, and the Need for Greater Algorithmic Accountability（人工智慧技術、演算法於司法實務之應用）

內容摘要

講者：林勤富副教授（清華大學科技法律研究所）

本次會議的主題是人工智慧技術、演算法於司法實務之應用，焦點集中在 State v. Loomis 一案引發的法律爭議及其關連課題。本案被告在槍擊案中被起訴，承認輕罪（開贓車、拒絕盤查），不承認重罪，進入認罪協商程序。本案法院運用了 COMPAS 統計模型（其為較傳統的演算法，只是一個比較複雜的統計系統，而非先進的人工智慧），判斷了被告的再犯率。本案 COMPAS 判斷結果結果有為高程度的再犯率，從重量刑。

一、問題提出：許多政府職能被演算法化

在決策過程中納入演算法（包含了各種高/低階的演算法），讓決策程序半自動化，在私部門信貸、保險費率等領域中已被廣泛使用，而在金融、交通、社福領域中，尚有爭議。而在司法程序中納入演算法，協助、作成決策，亦引發了許多爭議。

二、案例介紹：法院中的演算法—State v. Loomis

- COMPAS 之背景：美國希望透過自動化工具，準確、有效率、一致的做成司法決定。

- Loomis 不服，上訴至 Wisconsin 最高法院，理由如下：

- ◆ 違反正當法律程序：憲法規定量刑必須基於 accurate information，但其沒有辦法確認 COMPAS 的正確性，無法知道其怎麼運作、無法為交互詰問。

- ◆ 違反「量刑需基於個人化」之原則：COMPAS 是根據全國性資料來做成決定，並非根據其個人資料作成之決定。換言之，類似的人可能有類似的再犯率，但不代表我就會這樣。

- ◆ 違反 due process right：COMPAS 有專為女性設計的量表，考量性別並不公平（為什麼不主張平等權？或許是訴訟策略上的考量）。

■ Wisconsin 最高法院駁回上訴：

- ◆ 資訊均由被告提供，或公開資料，COMPAS 正確性可以驗證。

- ◆ COMPAS 只是法院考量的因素之一，法院仍會針對個案為判斷。

- ◆ 必須證明「參考性別是錯誤的」，且「此錯誤被法院採用」。

- ◆ 法院引用某個專家的意見，考量性別是為了讓整個司法體系、個別判斷更準確。

三、案例分析與問題討論

■ Wisconsin 最高法院採用了寬鬆的審查基準，保障了法官的裁量權。

- ◆ 承審法官不僅考量了 COMPAS，亦考量其他因素即可。

- ◆ 法院只考量了資料輸入及正確性，不考慮 COMPAS 的運作過程。法院根本不知道 COMPAS 如何運作。

- ◆ COMPAS 本質上不是拿來量刑用的，是給矯正機關用來判斷再犯率的，再犯率只是量刑的其中一個因素。

- ◆ COMPAS 必須在報告前放警示標語，指出系統的風險。

- ◆ COMPAS 運用的相關問題：

- 判斷過程是不公開的。

- 不是基於個人，而是基於群體做成的判斷。

- 有學者認為 COMPAS 有種族歧視之情形，可能會高估黑人男性的再犯率。

- 根據全國資料作成的，不是基於各州的資料。

- COMPAS 系統不是設計給量刑用的，只是用來估計再犯率。

- 寬鬆審查：法院不想正面面對演算法化的問題。

- 亦有文獻指出，警示標語是沒有用的，因為人對於科學或電腦做出的結果都很順從，法官仍可能被限制。

- Data Input, Accuracy, and Process, Input 是否正確？Output 是否正確？結論的 accuracy 如何判斷？

■ 兩種 Black Box Problems：

- ◆ Legal Black Box：trade secret 原則造成的黑箱（法律所保障的黑箱）

- ◆ Technique black box：演算法本身的黑箱（技術上的黑箱）

■ Safeguards：

- ◆ 完全不接受自動化決策

- ◆ 只能用公部門開發的、或是開放原始碼的演算法

- 只能解決上述第一種黑箱（法律所保障的黑箱）

- ◆ 增加透明度
 - 在此法院應要求其公開。
 - 只能解決上述第一種黑箱。
- ◆ 透過科技解決：
 - 在程式裡面加入 affirmative actions（只能解決上述第一種黑箱）。
 - 發展 explainable AI（訓練一個 AI 來解釋另一個 AI 做的事）。
- ◆ 透過權利解決：
 - right to explanation
 - GDPR techniques：提供「拒絕自動化決策，讓人類介入」的權利

■ 反思：

- ◆ 大部分的解決方法都是針對 legal black box，但其實 technical black box 的問題比較難解決，且未來的應用會越來越頻繁。
- ◆ 在批評演算法的 accountability 時，也要注意人類法官亦有不透明的缺點。
- ◆ 演算法 accountability 要如何更好？
 - 定義自動化決定。
 - 確認哪些決定是不能完全自動化的。
- ◆ 不同的 safeguard 可以混用來確保 accountability。
- ◆ gatekeepers 的專業、資源是重要的。
- ◆ regulatory impact assessment：當政府的決策引入自動化系統時，即需做 regulatory impact assessment，accountability 不足、帶有歧視時即不得引入。

Q&A

■ Q：量刑為什麼要考量再犯率？

- ◆ 法院首先必須論證，在量刑時為什麼要考量再犯率。台灣再參考此案時，也應注意這一點。
- ◆ counterfactual contest: 建立於我們傳統對於因果關係的想像。
 - 缺少某個因素，是否會改變結果？
 - 非 p 則非 q。若現在發生了 p，我們想知道他與 q 的關係，就看沒有 p 的情況下，q 是否會出現？
 - 但 AI 未必是線性因果關係，可能是多重因素底下，綜合考量的結果，在此未必容易操作。

■ A：這是折衷的解決方法。counterfactual contest 找出的並非真正的因果關係，因為連程式設計師都不知道中間的因果關係。只能用測試的，告訴被告如果改變某個因素，結果即會不同。

■ Q：為什麼在不受監督的情況下，可以用 COMPAS 系統？

- ◆ 以私部門為例，機器人理財是受監督的，其必須透明，也一定有問責機制。那為何法院可以使用不受監督的 COMPAS 系統？

■ A：Wisconsin 法院：紐約州、加州證實了此系統正確性蠻高的，也引用了專家的說

法，且下級法院僅能將其作為參考。

- Q：在 deep learning 之前，有一個絕對、可驗證的資料，那 COMPAS 是否有前後對照的機會？換言之，人類在判斷再犯率時，是怎麼判斷的？是否有「正確答案」？

- ◆ 可能一：有正確答案，期待演算法達到正確。

- ◆ 可能二：沒有，人類有偏見，期待用演算法來取代之。

- A：部分結果可以驗證（若有再犯），非百分之百可驗證。在可能二的情形下，其實 programmer 在此成為法律解釋適用的角色。

- Q：若連 reverse engineering 也無法使用，似乎沒有人能夠 verify、explain？

- ◆ 要釐清 verification 的意思，在 training 階段，有標準可以驗證這套系統是否達到預定標準（ex 拿一半的資料來訓練，拿另一半的資料來驗證並修正，以達成在現有資料中，系統能完全正確）；然而在出廠以後（實際使用），各個 case 的確沒有辦法驗證，要等到其再犯時才能驗證，COMPAS 蒐集了全國的資料，不斷的修正。但我們永遠沒有辦法說，對下一個案子的判斷是正確的。

- A：在這個討論下，好像真的正確就可以了，但另一種價值選擇會認為，就算結果是正確的，也會希望知道過程，希望能夠理解，尤其是在司法程序中，會希望知道論理過程。驗證的問題會牽涉救濟的可能性，如果沒辦法驗證，沒辦法挑戰他可能是錯的。

- Q：美國這樣的做法是否普遍？是否可以單純針對刑期上訴，不爭執事實、法律適用？COMPAS 是否有法律依據？

- A：美國普遍使用，但聯邦最高法院拒絕審理。可能沒準備好或案件不夠多或沒那麼嚴重。可以單純針對刑期上訴。沒有法律依據，法官可以選擇要不要用。犯罪資料在美國是公開的，且美國近幾年都在推動 evidence base sentence，想自動化、減少成本，政府在此扮演了很大的角色，有強烈的公私協力色彩。

第三次會議紀錄

時 間	107 年 9 月 28 日（星期五）
主 題	初探自駕車的電車難題：刑法學觀點、AI 個資在英國與歐盟的經驗
內容摘要	

講者一：薛智仁副教授（台大法律學系）

本次會議的主題之一是從刑法學觀點初探自駕車的電車難題。內容摘要如下：

一、什麼是自駕車的電車難題？

- 情境 1（3 人 v. 5 人）：直行會撞死 5 人，轉彎會撞死 3 人，轉彎可不可以？
- 情境 2（1 人 v. 1 人）：直行會撞死駕駛，轉彎會撞死路人？

二、刑法如何回應人為駕駛的電車難題：

■ 緊急避難？

- ◆ 為保全「重大優越」的利益，所為之不得已之行為。換言之，保全之利益必須顯然重大於其所犧牲的利益。
- ◆ 生命對生命的衝突？
 - 1 v. 1 生命質量不可衡量：生命質量不可衡量，每個人的生命是一樣的，無法符合「保全重大優越的利益」的標準。
 - 3 v. 5 社會連帶義務：生命質量不可衡量，但承認社會連帶義務，惟嚴格限縮，不可把風險轉嫁給第三人。

三、刑法如何回應自駕車的電車難題：問題：在第 5 級全自動自駕車，若是程式設計者指示自動巡航系統遵守「最小損害原則」和「優先保護乘客原則」，自駕車因此撞死他人，程式設計者或使用者應該負殺人罪或過失致死罪的刑責嗎？

■ 德國刑法學概況：

◆ 緊急避難：不可。

- 是否為避難過當，能否阻卻/寬恕罪責？
- 非自然人可寬恕罪責，係因其在該狀況下是特別緊張的。但在自動車的困境下，車子的行為（係因演算法的設計、駕駛啟動車輛）並無緊張困境。

◆ 適用義務衝突或容許風險。

- 理由 1：降低肇事率
- 理由 2：程式設計者的決策困境和人為駕駛不同。
- 只要程式設計者在設計時的決策是對社會上所有人都有利即可。

◆ 本文見解：

- 降低肇事率與電車難題無關。
- 決策困境不同並不會改變評價標準。ex 自動射擊設備，預先的正當防衛。判斷射擊行為是否合法時，仍使用正當防衛的標準

四、總結：

- 自駕車的轉彎不得造成他人死亡。
- 自駕車的技術只能因應相對單純的兩難困境。
- 自駕車使用者的生命並非優於其他人，設計者應告知使用者。
- 阻卻罪責的緊急避難。
- 與阻卻違法的差異：無辜者可否主張正當防衛。

Q&A

- Q：這個問題值得討論，是否因自駕車出現後，這樣的情境會增加？
- A：不確定，猜測不會。但之所以會拿來討論，是因為程式設計時，即須考量要怎麼設計才符合法律規定。
- Q：如果自駕車被恐怖份子當成武器，正當防衛、義務衝突，在刑法上的討論？在設計上可否讓它毀滅，來保全他人？
- A：
 - ◆ 主流見解：仍不能主張緊急避難，因生命無法衡量。
 - ◆ 少數見解：例外可以被容許。
- Q：只有兩種結果：不作為，雙方都會死；如果射擊/摧毀該車輛，僅有車上的人會死。換句話說，車上的人無論如何都會死。我們犧牲那些人不是因為他們不重要，而是他們無論如何都無法救助的？
- A：反對：滑坡效應，可能會被推論成：未來無救助可能的人的生命都會被犧牲。
- Q：如果可以合理預測車上的人會死的話，那有個道德上的戒命是不存在的（要救比較多/比較少的人），開放選擇，在道德上是開放的。法律上可以阻擋滑坡效應？
- A：在刑法上無法是開放的，必須要表態。可能的做法：1.還是去區分哪個行為是對的，區分哪個生命值得被保護；2.自然人在遇到這樣的情境下，是否可能用理智來做決定，不能的話就原諒他。
- Q：為何不能是開放的？
- A：必須給程式設計者方向；且在利益衝突時，仍要給雙方一個理由。
- Q：在生命價值不可衡量的前提下，如何給理由？會不會導出「無辜者的命比較重要」的結論？
- A：不會，只是要求他自承風險。
- Q：迫使自駕車需要選擇的情境，若是他人所招致？
- A：無論如何，在刑法上不會改變它是否為緊急避難。
- Q：若避難者對於危難的發生是與有過失的？
- A：故意：不可主張；過失：可以主張，但空間變小（標準提高）。
- Q：若生命不可衡量，法律評價即為「不可為了救 20 萬人，犧牲 500 人」，背後還有人性尊嚴的問題？
- A：
 - ◆ 以鄧如雯殺夫案為例，不可主張正當防衛，緊急避難若為最後手段，是否符合「保全顯然優越利益」的要件？鄧如雯所面對這樣的危險狀況，完全是先生所引起的，被犧牲的人正好是危難來源，所以他所犧牲的範圍會比較大。犧牲的範圍跟危難的來源有關。
 - ◆ 生命利益不可衡量，生命利益非永遠不能犧牲，只是無法抽象的恆量，在個案中仍可衡量。

講者二：何之行助研究員（中研院歐美所）

本次會議的主題之二是 AI 個資在英國與歐盟的經驗。

一、從 DeepMind 案談起

- 事實經過：DeepMind 想要做一個 app 與輔具，偵測急性腎衰竭，需要大量的個人資料，與 NHS 達成協議，可以取得未加密的，未得同意的病患資料。ICO 介入調查。
- 爭議焦點：
 - ◆ 未得同意：在直接照護的情境下，可以免除同意。
 - ◆ 未與 ICO、HRA 進行 consult。
- ICO 建議：
 - ◆ 須得病人明確同意。
 - ◆ DeepMind 未作隱私風險評估。
 - ◆ Data transfer 的範圍已超過必要範圍（由於資料庫未分類，故傳輸了全部的資料）。
- DeepMind 的 defense：
 - ◆ 直接照顧。
 - ◆ 非用於 AI 訓練。
- concerns from the society：
 - ◆ Care data program (2014)：想要把資料從家庭醫生手中拿回來。
 - ◆ 在併購時，無需再得到同意

二、AI and Privacy：

- AI 需要大量的訓練資料，與隱私保護衝突。
- 如果沒辦法解釋的話，如何滿足 GDPR 可解釋的要求？

三、GDPR：

- 特色：
 - ◆ 規範範圍廣
 - ◆ 處罰重
 - ◆ 擴大個人資料定義。
- 設立背景：
 - ◆ 受到經濟貿易影響：希望透過這個法律讓其他國家的企業負擔較高的法遵成本。
- 原則：
 - ◆ 保護資料主體。
 - ◆ 目的特定、資料最小化：如何兼顧個資保護與 AI 發展（科學研究例外規定）。
 - ◆ 透明性
 - ◆ 可歸責性
 - ◆ 告知後同意，AI：難事前預測資料用途。
 - ◆ 被遺忘權
- 若 AI 想再 GDPR 下發展，該怎麼做？（關於科學研究的例外）

- ◆ 33 允許概括同意：
- ◆ 放寬對於研究的定義：
 - 不限公/私部門，允許商業公司
 - 公共利益、醫療、公共衛生
- 目的外利用：
 - ◆ compatible：二次目的與原始目的是否 compatible。
 - ◆ 科學研究（包含商業）
 - ◆ 6(4)
 - 新目的與原始目的有關聯
 - 區分敏感/非敏感
 - 須符合安全防護措施
 - ◆ 放得很寬，是否能達到保護個資的目的值得討論
- 無需同意的敏感性個資處理，GDPR 9(2)
- 假名化：GDPR：基本要求，無法因此免除同意。

● Q&A：

- Q：金融領域中，個人可否主張匿名化來接受金融服務？
- A：不可，有洗錢防制的考量。
- Q：那是否會使金融業以此為由規避 GDPR 的規定？
- A：要將資料區分為洗錢防制/商業目的，才比較方便討論。許多資料是洗錢防制法要求金融機構才搜集的。
- Q：是否需在每個領域個別立法？ex：需要特別訂立洗錢防制法，而非引用個資法中的例外規定？
- A：GDPR 不會使原本的規定放寬。
- Q：個資法/GDPR 是否可以當作搜集資料的法律依據？ex：衛福部要把資料交給國衛院，其法律依據可否為個資法？
- A：法院實務：符合個資法例外規定，可以，其實不可。
- Q：如果侵害行為都需要個別法律授權，這樣的話個資法例外規定要做什麼用？
- A：法律移植下的錯誤。
- Q：假名化後，既然非個資，為何還需要經過同意？
- A：
 - ◆ 只有在完全去連結/去識別化，才能排除個資法
 - ◆ 我國法務部的態度：只要經過假名化，就是新的資料
 - ◆ 但若真正去連結，會失去資料的意義
- Q：資訊隱私 v. 資訊自主？
- A：
 - ◆ 隱私：去識別化後，即不違反資訊隱私
 - ◆ 自主：去識別化後，仍可能違反資訊自主

- ◆ 人性尊嚴 ->資訊自主 ->資訊隱私

第四次會議紀錄

時 間 107 年 10 月 25 日(星期四)

主 題 人工智慧、理財機器人與金融法規、AI 醫療民事法責任之研究

內容摘要

講者：楊岳平助理教授（台大法律學系）

本次會議的主題之一是人工智慧、理財機器人與我國金融法規之關係。內容摘要如下：

一、理財機器人的意義：

- 應用人工智慧來提供金融消費者來做理財規劃，未必是「機器人」，而是理財程式。
- 過去就有，只是現在機能更強大：
 - ◆ 基金股票篩選、預測未來
 - ◆ 具有情感溝通的功能

二、理財機器人的分類：

- 提供者：
 - ◆ 金融業者：投顧公司、銀行。
 - ◆ 非金融業者：電子、資訊公司。
- 提供之服務：（為金融業者）建置整體系統 v. 具體投資建議
 - ◆ 兩者的監管方式不同：前者較類似金融業之技術提供者；後者直接提供消費者理財程式。
 - ◆ 是否涉及自然人：自然人介入 VS. 機器人提供
 - 前者是指理財專員，使用理財程式提供服務；後者則直接由程式來提供服務。
 - 區分實益：監管方式（責任歸屬）：有人介入時，較容易找到自然人來認定責任；完全沒有自然人時，責任歸屬會是個問題。
- 理財服務內容的強度：建議性質 VS. 代操性質
 - ◆ 前者是指程式運算的結果並不會自動執行，仍由投資人自行決定是否要接受程式的建議；
 - ◆ 後者程式直接接管帳戶，代替客戶快速下單（優點：快(投資講求時間差)、但需

要更高程度的信任)

三、理財機器人的優點 (VS.一般理財專員)

■ 優點：

- ◆ 優勢的數據消化與計算能力：更正確、能處理更複雜的問題
- ◆ 較低機率的不當行為（例如：詐欺、利害衝突、歧視）
- ◆ 不受情感影響（自然人之工作狀況受情緒影響）
- ◆ 全天的服務提供，無時差問題
- ◆ 可以服務投資規模較小的投資人
- ◆ 理財機器人在大量使用下，成本較小（可以用固定的成本服務大量的投資人），就算個人投資規模較小，機器人仍可提供服務

■ 缺點：

- ◆ 顧問內容較單一缺少變化
 - 可能會提供所有人一模一樣的建議，此取決於程式的複雜程度
- ◆ 缺少情感交流：
 - 自然人可以透過觀察情感，來提供更準確的建議。
 - 目前為 5 秒搓合一次 (5 秒內下單的都是一樣的)，可能會變成逐筆交易制：
 - ◇ 設備較好者可以搶快來套利。
 - ◇ 散戶投資人可能會退出，專注於長期投資。
 - ◇ 理財機器人可以為散戶提供更快的交易。
- ◆ 理財機器人的消費者保護風險
 - 廣告不實，沒有真的提供 AI
 - 程式中植入詐欺及利害衝突
 - 程式未即時學習更新市場或法規變動
 - 呈現方式可能誤導消費者
- ◆ 理財機器人的系統性風險
 - 大家會在類似的時候買&賣
 - 造成閃跌閃崩，造成股票市場的波動
 - 為了監管，主管機關需要知道程式的設計邏輯，會需要透明的系統
 - 資本市場探索與嘗試錯誤功能，有礙金融創新。如果都能得知真實價格，金融市場要如何獲利？

四、理財機器人的監管問題

■ 證券投資顧問業之管制

◆ 類型：

- 投資顧問：提供分析意見
- 全權委託：直接交易

◆ 管制：

- 許可制：投信投顧法§63：證券投資顧問事業，應經主管機關許可，並核發營業執照後，始得營業。

◇ 需要許可的根本理由：

- 過去美國經濟大蕭條，是投顧事業詐欺、資訊不對稱造成的
- 個人沒有能力知道投顧業者的決定怎麼做成的，需要靠主管機關來彌平這樣的資訊不對稱

● 其他監管

◇ 最低實收資本額新臺幣 2,000 萬（\$5I）

- 如果只是提供建議，不代管資產為何需要這麼多錢？值得探討。

◇ 設置「投資研究」、「財務會計」部門，配置適足適任之經理人、部門主管及業務人員

◇ 應作成投資分析報告

◇ 投顧人員之資格審查（通過同業公會的審查）

◆ 程式設計公司是否應受此管制？

● 在此是否構成「投顧事業」

◇ 大盤趨勢還不算投顧事業

◇ 提供個別有價證券之買賣建議，則涉及投顧

● 金管會態度：

◇ 為投顧業者，尚未開放全權委託

◇ 未來可能開放（全權委託）

◆ 程式設計公司是否應取得此許可？

● 在此是否構成「投顧事業」

◇ 大盤趨勢還不算投顧事業

◇ 提供個別有價證券之買賣建議，則涉及投顧

● 考慮開放有限的投顧執照

◇ 資本額不用那麼高、為差異性管理、進入金融監理沙盒做測試

● 監管的可能問題

◇ 應作成投資分析報告

- 但業者也不知道 AI 如何做成決定

◇ AI 是否應該通過測試？（傳統對於投顧人員有資格審查）

- 金管會不知道怎麼測試

- 同業公會測試可能有營業秘密的問題

■ 金融消費者保護

◆ 提供投資建議，為何要受這麼大的管制？

● 建議買球票、演唱會票，不需要許可制，那為什麼建議買股票要？

- 原因：美國經濟大蕭條後檢討發現，當時股票崩跌，有部分是因為投顧業者的詐欺，造成投資人為錯誤投資。故希望解決此種詐欺、利害衝突、資訊不對稱問題，於是設計了一整套投顧事業的監管規範。

◆ 內涵：

- 核心：善良管理人注意義務、忠實義務，並搭配了刑事責任，讓監管機構

有介入空間

✧ 介入原因：消除資訊不對稱

✧ 大眾找投顧業者來協助投資，但不知道業者在想什麼，也沒有能力知道投顧判斷如何做成，必須靠法規、行政管制來彌平此資訊不對稱

● 具體方式（台灣）：

✧ KYC，了解客戶之財務資產狀況

✧ 廣告促銷時不得誇大偏頗

✧ 不得有利益衝突

✧ 不得有虛偽欺罔謾罵

✧ 不得意圖利用投資分析來謀求自己的利益

◆ 理財機器人的金融消費者保護

● 要告誰？幾個可能：

✧ 業者僅就程式的設計與操作負過失責任

▪ 過失認定時點提前到設計階段（而非服務提供階段），難認定有過失

✧ 比照消保法，負無過失責任

▪ 造成沒人敢提供理財機器人，扼殺創新

✧ 賦予機器人法人格，其就提供服務時的過失負責，要求業者就個別機器人提撥責任保證金

▪ 功能導向，回復原本的投顧責任判斷模式，以機器人提供服務之時點來認定，是否符合善良管理人之水準

▪ 業者負有限責任/無限責任？

✧ 責任保證金用完後是否應繼續補？

● 投信投顧公會「證券投資顧問事業以自動化工具提供證券投資顧問服務作業要點」：

✧ 內部定期審核演算法、KYC 作業、利益衝突、內部組成專責委員會來監督程式、資訊揭露

✧ 欠缺檢查機制與執行機制

Q&A

■ Q：把「機器人」的人拿掉，可以討論的範圍更大？

■ A：

◆ 英文正確的翻譯應該是「機器顧問」，確實不見得帶有人

◆ 部分研究指出，真的實體的人形機器人是扣分的，金融消費者害怕機器人

◆ 科技部的發展重點：

● 科技面：想發展最先進的 AI

● 人文面：背後的法律，以及防弊機制

■ Q：理財機器是商品、服務、行業？（與洗衣機、洗衣業比較）金融消費者保護法

or 消費者保護法？

- A：適用金融消費者保護法，消費者保護法不包括投資人（基於投資目的而使用商品或服務的人），使用理財機器人的目的大多是投資，大體上是適用金融消費者保護法。

- Q：理財機器與金融監理沙盒的關係？

- A：

- ◆ 針對具有創新性的金融科技發明，在有限範圍實驗，豁免相關管制，金管會定期觀察實驗結果，再評估是否符合現行法、是否需修改法規

- ◆ 應用於理財機器人：不樂觀

- 是否具備創新性？

- ◇ 科技上創新，但有法規上的破壞性創新嗎？

- ◇ 亦有業者依照既有規定運作，看似沒有問題

- 實驗完如何接軌？

- ◇ 若實驗完成發現有法規問題，金管會管的規定，金管會可以自己改；有些涉及立法院修法，他不修怎麼辦

- 業者考量：

- ◇ 一但進去了可能等同承認自己違法，也可能就被關在裡面出不來了，不如在外面偷偷地做。

- Q：

- ◆ 最早提出理財機器人概念的是王道銀行，實際上是提供一個線上服務，只是程式，而非實體的人

- ◆ 全自動、半自動駕駛的差別？

- ◆ 自動駕駛出現很久了，但一直是輔助，若出現了一個全自動駕駛，會需要另一套管制模式？

- A：

- ◆ 有沒有自然人介入會是關鍵，目前台灣允許的理財機器人還是需要搭配投顧專員，監管上的破壞相對較小。下一個階段（全權委託）會直接接管帳戶，就會不太一樣。

- ◆ 人有沒有介入會是管制方式的關鍵。

- Q：全權委託下，法律關係為何？是否有辦法成立契約關係？

- A：

- ◆ 投資人與程式設計公司之間的，事前的概括授權，一開始下載程式時，就同意了未來的交易。在不承認程式有獨立法人格時，只能抓程式設計公司。

- ◆ 視商業模式而有不同：程式設計公司可以直接賣程式給消費者回家用；也可以透過投顧公司來取得這套系統（還有一個更上層的程式設計公司）。

- Q：法人格化的困難？法人格化會使民事上的咎責變得容易，刑事責任方面會變得困難（一種 trade off，放棄刑事法的管制）

- A：

- ◆ 美國習慣用民事責任來監管（在金融業裡，某種程度上民間較政府更強大），比

較不擔心 trade off 的問題。

- Q：(刑事責任) 證券詐欺：怎麼認定詐欺故意？
- A：台灣採善良管理人注意義務標準(客觀標準)，不需知道 AI 的主觀意識，在此反而容易認定。
- Q：忠誠義務、受託人義務中，告知的程度？傳統自然人與自然人交易時，我是否有把關鍵資訊告訴你。AI 對 AI 交易中，告知義務的內容可能要調整。
- A：
 - ◆ 過去的判斷標準是，投資者提供重大不實(=影響投資人判斷)的錯誤資訊給投資人，構成詐欺。此標準非固定，反而可以。
 - ◆ 可以思考到底為什麼要有法人制度(絕對不是為了處罰)
 - ◆ 其中一個角度是資產隔離(而非有限責任)。透過創造法人，把資產分成兩半，把責任分離，各自發展。
- Q：價格與管制的關係？
- A：
 - ◆ 在其他領域談 AI，是要破壞或取代原本價格帶來的資訊，有 AI 以後，不需要比價，直接購買最喜歡的那個。在金融領域，價格就是本質，價格不會被取代，價格內的取代則會受到影響。
 - ◆ 在金融領域過去強調程序上的管制(因為無法確切知道這個投資判斷對不對，只能透過程序來確保)，但 AI 是無法理解的，在程序上與過去有很大不同，要如何管制，也許要另外找規範的依歸。
- Q：如何管制？會比投顧業更難管
 - ◆ 投入資源面：
 - 目前對投顧業的管制好像不是程序面，而是資源面，要有人才、要有錢，但這其實是間接的，input 面。但 output 到底怎麼樣也不一定，但這種管法是最容易的，因為 input 是可以衡量出來的，對營業自由影響也不會太大，也不用仔細看他的程序面。
 - 但機器可以這樣管嗎？要管他吃的大數據是好不好的？也涉及是服務還是產品。產品通常會直接對他做檢測，有沒有錯誤(defect)、結果(performance)好不好。AI 可能兩個面都要有，看他有沒有 bug，也看他 performance 好不好。看起來可以做，但比對資源面的管理會難一些，例如 performance 到底要多好(目前投顧的 performance 也沒有多好)
 - ◆ 程序面：AI 會不會有營業秘密的問題，還是真的是 Big Data 的問題？
- A：
 - ◆ 也許沒有必要對 AI 非常嚴格，畢竟我們對投顧業的期待也沒有那麼高，只要兩者相同就好了，不要讓 AI 業者有不當利益。
 - ◆ 對投顧從業人員目前也只是考試，考完之後的成長也不清楚，AI 也是類似，一開始做測驗，後來學什麼也不知道。
 - ◆ 檢測內容：

- defect：要測
- performance：對目前的從業人員也沒有做這樣的測試
- disclosure：揭露過去的成績
- ◆ 可以依照 AI 的程式複查度（ex 有無 Deep Learning）來分類，但不完全等同他的 performance。
- ◆ 關於 performance 的管制：目前理專人員有依照他的實績，讓客戶可以選擇。那未來也可以透過管制，要求揭露 AI 過去的實績。換句話說，performance 是透過事後揭露來達成，而不是事前。
- ◆ 關於做成建議的理由（reasoning）：客戶需要的不是演算法，而是需要知道做此建議的理由（counterfactual explain），在此提供一個資訊給消費者，讓他知道這個建議是基於這樣的 correlation。此理由 AI 未必不能提供，或因涉及營業秘密而不需提供：
 - counterfactual explain：如果不這樣的話，那會怎樣。你做出此決定的依據，有跟無的差別是什麼。
 - 可能的困難點：data driven 的 AI 比較常使用 correlation inference，而非因果關係。它可能真的講不出來有此因素跟無此因素的差別，也許只能跟客戶坦承。
- Q：在行銷領域，無法說明沒有關係。但在做預測時，我們到底允許多少、哪一種決策，可以純粹使用 correlation，而且講不出來兩者的關聯。
- A：
 - ◆ 目前的理專人員說他因為..所以..，好像也不認為他有因果關係的認知，可能是基於過去的經驗。或許 correlation 可能足以作為投資報告中的理由。
 - ◆ 現行投顧法也沒有要檢驗投資報告分析是對或錯，只要提出即可，主管機關不會檢驗報告寫的有沒有道理，對 AI 或許也不用有那麼高程度的要求。

第五次會議紀錄	
時 間	107 年 11 月 29 日(星期四)
主 題	Can Nonhuman Author? Rethinking Copyright Through Algorithmic Authorship
內容摘要	

講者：陳舜伶助研究員（中研院法律所）

本次會議的主題是 AI 可否做著作權創作？及反思著作權本身的問題。內容摘要如下：

一、著作權基本原則（美國法）：

- 保護的不是想法，而是想法的表現形式（固著在實體的媒介上 ex 寫出來錄起來）
- 保護科學發展
- 提供著作權作為經濟誘因
- 有限時間內有專屬權（非常長，ex 死後 70 年）
- 門檻低，只要有最小的原創性及固著即可
- 反思：這麼容易達到，那真的只有人能為之嗎？

二、非人可以是作者嗎？

- 人畫的畫其實沒有那麼特別或有創意？
- 如果大家都看不出來哪些畫是動物/AI 畫的，那他們應受著作權保護嗎？
- 動物可以是作者嗎？Monkey selfie 案
 - ◆ 案情簡介：
 - 報紙跟作者買了某張照片：照片說明為「這張照片是猴子自己拍的」
 - 作者 slater：選鏡頭、固定攝影機→對這張照片有某些控制，想主張著作權
 - Wikimedia commons（開放授權的網站）收錄這張照片，被大眾拿去用，作者寄信告知 wikimedia，wikimedia 反駁作者是猩猩
 - 某動物權組織告該攝影師 Naruto PETA v. Slater
 - ◆ 爭點：該保留區的靈長類學者主張，那些猩猩有能力做值得被著作權保護的行為，猩猩是有自我意識的，它或許沒有意識要創作著作權，但此點與著作權是否成立無關。這張照片的利益應歸於 Naruto。
 - ◆ 美國著作權辦公室的解釋：
 - authorship 應該要是人
 - ◆ 法院見解：
 - humanity：主體能不能繼承經濟上的利益
 - 動物無法在法院提告，除非立法
- AI 可以是作者嗎？
 - ◆ OBVIOUS 案：法國的某個計畫，創作者創造了一個假想的族譜：
 - 神經網絡 neuwork network1，餵 15k（14-20 世紀）肖像畫來讓 AI 學會創造
 - 神經網絡 neuwork network2
 - 高價賣出
 - ◆ AI 社群的批評：
 - OBVIOUS 對創作的貢獻是有限的
 - 人類介入的可能：
 - ✧ 人類設計 network
 - ✧ 選擇資料庫

- ✧ 訓練 network
- ✧ curated the resulting outputs

- 誰提供貢獻？

- ✧ Goodfellow- GANs developer
- ✧ OBVIOUS?

- ◆ 著作權學者的見解：

- 見解 1：可以有著作權，但權利享有者必須是人類（programmer 或雇主）
- 見解 2：AI 有某種心智能力，可以是作者，但不能享有所有權（找使用 AI 的人或開發程式者）
- 見解 3：可以保護他但不要放在著作權裡（立特別法）
- 見解 4：AI 會征服人類，所以現在就要阻止他，不要給他任何權利
- 見解 5：不需要給 programmer 或雇主任何經濟誘因他們也會用 AI
 - ✧ 給很有限的著作權保障
- 見解 6 現階段都還有人類的介入，暫時不需討論這個，目前只要微調即可

- ◆ 反思：

- 談到動物時，好像比較會明白提到人類中心，反正動物就是不會創作
- 談到 AI 時，被視為工具，比較願意承認 AI 有創作能力，給著作權
- 兩種非人的對待上有點不同

三、非人法律地位的討論

- 我們預設了人為中心，在討論要不要讓其他事物進入人類世界
- 跟人與非人之間的互動有關，是否要把他放入人的道德 or 法律社群，沒有一定的標準，會不斷變動
- 寵物、虐待動物、某條河取得法人格的議題發展中，大多是保護其不受傷害，而不是讓他取得某種權利
 - ◆ ex 紐約有兩隻猩猩被關著，以人身保護令提起訴訟，法院認為動物沒有 moral agency 判敗訴，但我們不會說被關著的嬰兒、植物人沒有 moral agency。（對非人的要求比較多）
 - ◆ 馬被虐待後，主人被判刑、罰金，但不足醫藥費，動物保護團體代替馬提起訴訟，要求其支付醫療費
 - ◆ 法律學者的滑坡憂心：若允許動物提起訴訟，會不會像過去奴隸被解放般一發不可收拾？
- 線要畫在哪？
 - ◆ 自我意識？知覺？
 - ◆ Copyright 一直都有畫線的問題
 - ◆ 不要問 AI 是否有 legal personhood，一直變動而且無解，而是要問那個法律的目標到底是什麼？
 - 如果 AI 有能力生產有創造力的東西，給他著作權也有促進 AI 使用的效果，那就給（功利主義）

- 要給誰，就由給誰最能促進發展來決定
- 沒有思考到的：
 - ✧ 該給多高程度的保障
 - ✧ 在創造過程中可能會影響別人的著作財產權，可能是反動機
 - ✧ 大量生產可能會越來越限縮未來的創新（越來越可能侵害著作權）
 - ✧ 著作權保障跟公共使用的衝突
 - ✧ authorship 畫在 intelegent
 - ✧ 所有權人需要能夠分享該作品，所以需要 moral agency

■ 非人作者對著作權的挑戰？

- ◆ Copyright 有特定的政策目的，所以可以接受只有人有著作權
- ◆ 但 creativity 可能可以被動物、AI 滿足
- ◆ 究竟什麼時候要把非人放盡著作權保護，可能還要一段時間
- ◆ 因為目前人都有介入，所以著作權可能還不用調整太多
- ◆ 最小原創性會不會太低？固著的意義是什麼？要不要有創作意圖？

Q&A

■ Q：AI 的討論每次都會讓我們去反省原本就存在的議題有哪些？

■ A：

- ◆ 著作權原本要求的原創性就很低
 - 攝影：拍的都是現存的東西，創作的成分可能很少
 - 很像 AI，所有素材都是人類提供的，但最後是由 AI 判斷的，可能還是創作
- ◆ 誰是作者&所有權？
 - 著作財產權：可以轉讓
 - 著作人格權：綁定作者
 - 分開處理似乎是比較合適
- ◆ 什麼是 Personhood：
 - 財產權、人格權的前提都是人格主體
 - 著作的要件是人類的精神創作

■ Q：值得討論的是，分開處理要考慮的是什麼？

■ A：

- ◆ 理論基礎：
 - 人格權說：人精神分離的一部分
 - 功利主義：給動機（給 AI 一點用也沒有），動機要給後面那個人（提供更好的 AI or 設備）
- ◆ 要留意動機的反面就是封鎖效應
- ◆ 作者留給 AI？
 - Introbuton 誰做的，就要正確地去承認他的貢獻

✧ 不問 AI 是否有人格，就要誠實的寫出來，人類不需要剽竊 AI

- 細部問題：存續時間（AI 不會死 or 馬上就會死）

◆ 著作權法是要創造誘因給大家創作，AI 的確無需給予動機，但若從其他角度：

- 我們要鼓勵使用 AI 創作，所以鼓勵背後使用的那個人
- 但這樣會在討論時執著於人的貢獻會不會太少
- 給 AI 法人格會不會讓這個問題簡單一點

■ AI 與處罰，目的是什麼？

◆ 如果目前工程師已經盡力設計了，那為什麼還要處罰他？

◆ 在 AI 有法人格以後，是否會切斷工程師應負的刑事責任？

■ A：

◆ 權利義務主體性的抽象思考

- 很多問題在建構法律制度時就存在的問題
- ex 動物權是為了動物還是為了人類的永續發展
- 我們賦予某種東西人格，支持目的論
- 法律是為了規制人類的活動，要因應人類的社會需求
- 除了要討論著作權，刑事責任、甚至基本權都有討論價值

◆ 為什麼要與動物做區隔？

- 差異在於動物也是生物體（有物理上的獨立性），AI 則是人類創造出來的

◆ 延伸討論：

- Alphago zero 已不用人類的棋譜
- 法人化的問題：會讓刑事責任成立的可能低很多（故意很困難）
- 可以用別的立法模式來解決問題，例如基金補償

■ Q：賦予 AI 法人格會有哪些問題？

■ A：

◆ 法人格到底要承載什麼？

◆ 部分權利能力主體的觀念與應用可能

◆ 是否涉及故意的問題：電車難題就是故意的問題

◆ 是否有創作理念是否會有影響？

◆ AI 的大量創作怎麼辦？

- 音符就那幾個，若 AI 快速的窮盡所有排列組合怎麼辦？
- 一開始著作權會強調 incentive，是因為他的假設是取之不盡的
- 但若是有限解（ex 西洋棋、圍棋、五言絕句），這個假設會不會是完全相反
- 回應：著作權如果你沒有先看到，就算創造一樣的東西，也不會違反著作權
- 著作權的目的到底是什麼？挖盡所有知識金礦？

✧ 五言絕句的目的不是要把所有文字組合玩殘，而是找到那些感動人的組合

✧ AI 可以窮盡所有組合，但不見得真的會有人這麼做，這麼做未必會有

市場價值（五言絕句全破解未必有人會買），如果沒有的話，這樣的討論就只有理論價值。

✧ 玩殘了以後可能會有另一個產業的興起，幫你找出有意義的五言絕句，例如「一定要讀的 10 首五言絕句」

第六次會議紀錄

時 間	107 年 12 月 26 日(星期三)
主 題	Casual Explanation as a Partical Solution to Mitigating Algorithmic Harms

內容摘要

講者一：邱文聰副研究員（中研院法律所）

本次會議的主題之一是探討如何對演算法可能造成的傷害進行調適。內容摘要如下：

一、AI 的演進、分類與問題意識：

■ 演進：

◆ 第一波與第二波的最大差異：

- 第二波 AI 需仰賴大量資料，故會有「資料使用」的相關問題

◆ 第三波：類似多拉 A 夢，全能型

- 現在還沒出現，故先處理第二波 AI 的挑戰

■ 第二波 AI 的分類

◆ 第一類 AI：模仿人類行為

- 一般人對 AI 的想像，用機器、演算法來模仿人的行為，ex 用機器來辨識動物。
- 科技部 AI 醫療影像資料庫，希望能透過醫療影像學習，讓電腦也能像專業醫生，判讀醫療影像（病患是否正常，有無癌細胞）ex：機器寫作、創作。
- 此類 AI 產生的法律問題會有重大不同

◆ 第二類 AI—Association Discovery：

- 過去做研究的方式是：先有理論，根據理論來說明哪些資料之間，有特定的關聯，再透過統計方法來證明此理論為真。
- 資料的使用（尤其是 BIG DATA），可以讓機器發現人類還未知的，資料之間的邏輯關係。這種演算法不需要仰賴特定的理論，就可以去尋找資料與資料之間的關係。只要有足夠大量的資料，就可以發現以前的理論沒有發現的關聯性。

✧ 此種 AI 的法律挑戰：其所發現的是 correlation（關聯性），而非因果關

係。Correlation：例如「冰淇淋銷售量」與「被鯊魚攻擊」之間在時間軸上有 correlation，我們清楚知道這兩個事件背後的成因，知道這兩者間沒有因果關係(可能有共同的原因—天氣很熱)。其他例子：「自閉症個案數」與「有機食品銷售量」有 correlation；「漫畫書銷售量」與「computer science 博士學位授與」有 correlation。

- ✧ 如果我們對各個事件的成因有基本認識，我們可以辨別出他們只是關聯性而非因果關係；但當我們欠缺足夠的知識去了解事件的成因，correlation 就會引導我們去做一些不當推論。例如 2000 年有一本書，主要在談人類智商跟行為、社經地位有關聯性，用智商來預測某個人未來的收入、工作表現、是否會犯罪。此書引發美國社會的爭論，這是合理推論嗎？可以用 IQ 來推論社會現象嗎？

■ 實際舉例：

- ◆ 例 1：以「是否購買沒加香料的乳液」來預測「購買者（女性）是否懷孕」，此是否為理性的預測？
- ◆ 例 2：以「是否購買傢俱保護貼」來預測「購買者信用狀況」，挑戰：若認為這種類型 AI 的應用應更透明(例如 GDPR 內的 transparency)(例如貸款被拒絕者，會希望知道為什麼)，但揭露此訊息，會造成緊張關係，例如債信不好的人，會去買保護貼，也許就能獲得貸款。
- ◆ 例 3：美國應用某個演算法（其功能為預測某個人是否會再犯）來量刑（而非假釋）。換言之，其將 AI 運用於預測人類的未來行為，這個預測會影響到他現在的身份地位關係，例如說他會受到多少刑罰、是否允許他假釋、是否羈押或交保。

■ 問題意識：

- ◆ 我們使用第二類 AI，什麼樣的使用是適當的？
- ◆ 會有哪些可能的傷害？

二、Association Discovery 可能帶來的傷害及解決辦法

■ 可能帶來的傷害

- ◆ 傷害 1：隱藏歧視
 - 歧視係指使用不當的特徵（種族、性別）作為判斷標準、操作變項，其違反了平等保障。而在此種 AI 底下，歧視可能被隱藏。
 - 例如種族，可以用居住地區（ex 郵遞區號）來作為替代指標，尤其是美國這種黑白分區居住的情形下。
- ◆ 傷害 2：將現有社會的歧視作為、現象，轉換為演算法
 - 例如在前述例 3，有人質疑用過去美國的資料庫來預測哪種人（黑人、白人）比較容易犯罪，會有偏誤，似乎黑人都比較容易犯罪。而發展該 AI 的公司，不過是忠實呈現社會，並沒有特別歧視黑人。

■ 傳統的解決辦法

- ◆ 解決辦法 1（在下游解決產生的歧視問題）：

- 以現有的反歧視法、平等權保障，處理這種「在表象上、統計上是正確的，而實際上有可能會產生歧視」的推論。
- 其內涵為工具理性，判斷歧視時，以手段與目的間是否符合工具理性（即手段與目的間的關聯性）作為判斷標準。

◆ 解決辦法 2（在上游檢驗演算法的知識生產過程是否正確）：

- 透過資訊揭露來處理，例如揭露在發展演算法的過程中，到底輸入了什麼樣的資料。惟此為商業機密，多未對外公開；而演算法過於複雜，亦可能使專家無法預估，產生黑箱結果。
- 此方法的挑戰是，到底要揭露多少資訊、揭露什麼樣的資訊，才夠我們檢驗演算法的發展是否合理。

■ 傳統解決辦法的問題：

◆ 解決辦法 1（反歧視法、平等權保障）無法處理 AI 的問題

- 反歧視法、市場機制所能處理的歧視：工具非理性及毫無緣由的嫌惡

◇ 工具非理性：

- 當手段沒有辦法達成目的時，這個手段就是不合理的，也就是歧視。例如目的為找到最佳員工，以膚色、性別為判斷依據，是歧視；而若以學校畢業成績、律師考試分數或名次，則不是歧視，因為我們能以過去學校的表現來推論未來是否為好律師。此種歧視通常可以用市場機制來淘汰（例如對手可以找到真正優秀的人，歧視者就會逐漸被淘汰）

◇ 毫無緣由的嫌惡：

- 若為普遍存在於社會的嫌惡，通常難以發現，且若大家皆如此，此種歧視是不會被市場淘汰的，最後的結果可能是 separate but equal，會區隔出不同的市場（例如白人的市場及黑人的市場）。要解決這種歧視，過去的方法是去特定某些特徵，這個特徵可能是社會上典型存在的歧視，我們不應該以此標準來判斷。
- 例如，種族、性別之所以會被稱為 sensitive category、suspect class，就是因為這種類型在歷史上長久存在，被用來歧視時難以發現，故除非有很好的理由，否則以此為標準即是不對的，現在的 Anti-discrimination law 就是如此。

- 用這樣的方法來處理 AI 難以成功：

- ◇ AI 宣稱的都是 correlation 上的合理性，當我們在判斷誰是好的貸款者、債務人，是否會還錢時，不知道背後的原因為何、欠缺足夠資訊，只知道買了保護貼的人比較會還錢，好像很難說這樣的推論是不合理（irrational）的。所以，傳統上以 irrationality 當作是標準的 Anti-discrimination law，會遇到挑戰。

- ◇ 通常會迴避掉 suspect class，很少會直接在 AI 的操作上，直接挑選種族、性別，會轉換為其他的特徵，讓人看不出事實上是以 suspect class

為標準。故這種下游的處理策略是難以成功的。

- ✧ correlation 有可能是 rational 的，因為其往往是人類欠缺必要知識時，提供一個可以推論的依據。這樣的推論到底可不可以，還是要看目的是什麼，如果目的是為了防止鯊魚的攻擊，而禁止販賣冰淇淋是不行的，因為目的跟手段無法吻合。若只是要預測冰淇淋的銷售量，也許鯊魚攻擊數量，是一個好的參考。

◆ 解決辦法 2 透明化 transparency

- 讓這個演算法的生產能夠更透明，讓大家知道結果是怎麼來的。
- 但有一個基本上的困難在於，由於演算法太複雜了，就算他不是商業機密，連資料科學家都很難解釋為什麼結果是如此。
- 拆解演算法：當 transparency 的呼聲一出現，第一個出來反對的人就是資料科學家，連專家都看不懂了，給一般人沒有意義。
- 提出相當解釋，而就算不把演算法拆解給大家看，起碼要告訴我結果是怎麼來的，提出相當解釋：

✧ 「解釋」的意涵

- 關聯性 correlation 解釋：例如不把貸款核發給我的原因是我沒有買保護貼，是依據保護貼來做成判斷的，起碼我可以做一些挑戰，質疑決策的合理性。
- 反於事實的解釋 counterfactual 解釋：告訴當事人，什麼條件的轉變之後，結果會有所不同。例如，如果你買了保護貼，你就有可能拿到貸款；因為你沒有買保護貼，所以你沒有拿到貸款。

✧ 上述解釋可能面臨的挑戰

- 把買保護貼揭露出來的結果就是，大家都去買保護貼，但沒辦法真的確保這個人有債信，會造成 accountability 跟 transparency 的緊張關係。
- 多以營業秘密為由拒絕透露
- Self-fulfilling prophecy 因為演算法算出保護貼跟債信有關，用這個結果來解釋時，是一個自我證成，而不代表沒有買保護貼的人債信比較差。

三、本文見解—應提供因果的解釋 Causal explanation

- 強調的並不是「改變了哪些條件，就會得出不同的結果」這種結果上的解釋；而是在於，必須告知所使用的兩個變項之間，在自然世界中，存在著因果關係。例如，必須告訴我，防曬乳跟懷孕之間，有沒有因果關係，如果有，是怎麼樣的因果關係。如此才能真正去讓演算法所提出的變項之間的關係，不再只是停留在 correlation 的層次，而必須到因果的層次。
- 可能的問題：
是不是所有的 Association Discovery 的演算法都無法存在？因為大多數 AI 都提不出因果的說明。所以必須做類型的限制，哪些應用必須提供因果的解釋。

■ 類型化：

- ◆ 「實質影響當事人法律或相似權利」的 AI 應用才需要提出因果說明；若非則不須提供。
- ◆ 另外，某些法律特別允許的無知推論（在欠缺充分證據時必須做出的判斷），例如環境法上常需要在欠缺充分資訊、不確定因果關係的情形下做成決定（precautionary principle，當可能的危害很大，縱然目前欠缺足夠知識來做出行動，仍不能因此而不行動，通常用於物，例如食品安全、環境保護）。
- ◆ 可能有以下類型：
 - 實質影響當事人法律或相似權利，且要求較嚴格的佐證資料：
 - ◇ 用演算法來推測是否再犯，決定是否假釋、或決定刑度，對其權利影響較大、且需要充分證據。不得單純使用 Association Discovery 的演算法，若當事人要求提出因果說明，機關即須提出。
 - 不影響當事人法律或相似權利，不要求較嚴格的佐證資料：
 - ◇ 對當事人權益不會造成太大的影響，也允許做無知推論。例如預測消費行為。一般的 correlation 就可以接受了。
 - 實質影響當事人法律或相似權利，但不要求較嚴格的佐證資料：
 - ◇ 提出因果說明，或提出可以訴諸無知的正當化理由（causal explanation or justification for appeal to ignorance）

四、GDPR 中的 right to explanation：

■ 本文的 Causal explanation 與 right to explanation 的關聯性：

- ◆ 相同：演算法的結果會影響其法律權利時適用
- ◆ 不同：GDPR 沒有寫明其所要求的解釋是不是因果的解釋，相關規定講的蠻模糊的（Article 13: meaningful information about the logic involved），常被解讀為拆解式的，把演算法的訓練資料、model、如何調教的講出來；或只是 counterfactual 的解釋。跟本文所提出的 Causal explanation 是有差異的。但不代表 GDPR 的 explanation 不能是 Causal explanation，是有可能的，因為也是限於法律權利侵害較大時。

Q&A：

- Q：Causal explanation 的內涵是什麼？
- A：不是指演算法「考慮過的因素」，必須以科學證據來證明其因果關係。而其內涵可能包含：因果機制、機率。讓 AI 的研發者承擔更多的義務，的確可能會遭遇反彈。
- Q：似乎與機器學習的價值有衝突，機器學習的價值在於沒有科學證據時，仍能找出關聯。
- A：的確，所以也沒有完全封殺這種 AI 的應用，但限於特定種類的應用。
- Q：看起來可以解決隱藏性歧視，但在「實質影響當事人法律或相似權利，且要求較嚴格的佐證資料」的類型中，還有正當法律程序的要求，Causal explanation 是否同時解決了正當法律程序的問題？

- A：解決了大部分，正當法律程序很重要的一點是 reasoning，要說明決定怎麼來的，Causal explanation 一定能提供更充分的理由來正當化決定。也提供了當事人提出挑戰的可能。
- Q：「實質影響當事人法律或相似權利，且要求較嚴格的佐證資料」舉的都是刑事法的例子，而在發動搜索時的 probable cause，其實比較像 correlation 的概念而非 Causal explanation。例如某人看到警察就跑，警察未必是基於有因果關係而逮捕他，只是依據過去的經驗所做的決定。
- A：這就是 profiling 的問題，警察可不可以因為很多黑人在販毒，所以是有問題的？profiling 在刑事法上可不可以被接受，在刑事法上有很多討論。另外，刑事法上依照強制處分侵害程度高低，相對應的證據程度也不同；在本文討論的問題上，國家行為強度的高低，與是否要求嚴格的佐證資料亦有關聯。
- Q：AI 的優點在於幫助我們釐清人類價值選擇的偏誤及盲區，但要求因果關係似乎是反向而行，由人類來糾正 AI。
- A：這沒辦法，我們本來期待 AI 有辦法糾正人類的各種偏見，看是否能更客觀地做成決定，但 AI 本身是有價值取舍的（非如資料科學家所稱的價值中立），宣稱具有 correlation，本身就是價值判斷，要不要基於這樣的 correlation 做成行動，更是一個價值判斷，這都不是 AI 能夠回答的。
- Q：適用於 Supervise AI 還是 Unsupervise AI？
- A：Supervise AI 是指先標註這是癌細胞，讓 AI 用 1000 張去學什麼是、什麼不是癌細胞，模仿醫生來判讀有沒有癌細胞；Unsupervise AI 是指把所有影像都丟給 AI，讓他自己去找哪些因素之間有關聯，例如金牛座的人比較容易有什麼疾病，可能會找出超越人類想像的東西。兩者皆可使用。
- Q：如果是涉及到生命、健康，新型的醫療疾病，過去沒有資料，要讓 AI 自己去找關聯性，那如果利用 AI 找出某個治療方式，在沒有其他選擇之下（既沒有過去的資料、也不可能不醫治，採用了 AI 的療法），但病人仍然死亡，要求嚴格證明、因果解釋，但在新的疾病是沒辦法提出因果解釋的。
- A：
 - ◆ 這種情況未必是「實質影響當事人法律或相似權利，且要求較嚴格的佐證資料」，有可能是「實質影響當事人法律或相似權利，但不要求較嚴格的佐證資料」。另外這可能要由病人自己決定，醫師也不知道為什麼 AI 會做這個判斷。對主管機關來說，問題則是要不要讓這樣的 AI 上市。
 - ◆ 還是要看怎麼看待這個 AI 的判斷，怎麼使用這個判斷，對當事人的影響是什麼。例如，要因此否決當事人擔任部長職位的機會，或是強制開刀。換句話說，就算是同一個 AI，說明義務會因為使用方法而有不同。
- Q：依照使用方式來決定開發者的證明責任，若使用者跟開發者不同怎麼辦？
- A：AI 的使用者跟開發者的確有可能不同，也有可能是同一人。在不同的情況下，使用者必須要求開發者提供更多資料，使用者才能使用這個 AI。換句話說，會課與使用者義務，但這個義務還是會轉嫁到開發者。

- Q：有可能去限制開發者的責任嗎？例如開發新的藥，但還沒通過完全的實驗階段，可能人體實驗只走到第一階段結束，還無法使用。某個疾病只有這個藥可以治療，醫生、藥商有可能說實驗還沒做完，但問病患要不要試試看嗎？
- A：
 - ◆ 這就是 intermedium rule，有一個專業的中介者來判斷，這時候的責任會在中介者身上，如果中介者搜集了各種資料，在沒有更好的資訊協助他判斷時，跟病人說明他也不知道怎麼辦，現在 AI 提供了這個資訊給我，你要不要試試看？在這個脈絡之下，可以拿那個藥來用。前提是現行法的管制架構是允許的，未上市的藥，但對病患有可能有效時，或是唯一可能的希望時，有可能用專案方式來允許使用這個藥物。
 - ◆ 最後的判斷還是在專家身上，比如在 Alternative Dispute Resolution 替代糾紛解決程序中，法官跟律師或被告討論，他不知道怎麼判，參考了 AI 的判斷，看律師或被告是否同意，不同意的話會回到原來的救濟方式。
 - ◆ 在這裡也可以看到，是否應提出因果說明，會因 AI 使用方式不同而有差異。例如一般的法官模式可能要求較嚴格的佐證資料，在替代糾紛解決程序中可能不同。
- Q：在法官模式中，以後的解釋義務會不會高於現在的法律義務？AI 的使用者及研發者，會不會反過來問，當在做接近道德或法律判斷時，還是必須由人來做決定，但這跟發展 AI 的初衷相違背？
- A：可以不要用 AI，就自己寫判決，當然現在的法官模式並不完美。回到前面所提的兩個模式：模仿人類行為及 Association Discovery，如果是前者，相關的責任義務都是相同的。要模仿法官，法官本來要有的 reasoning 也要具備。當然可以說現在法官的 reasoning 很爛，但起碼有一個形式上可以課責的對象。
- Q：如果只是要處理課責，是不是最後階段有人為介入就可以了？需要要求演算法給出解釋嗎？例如，correlation 再加上人的決定即可。
- A：
 - ◆ 這是 GDPR 的模式（22 條），這裡設想的情形是，人為介入在 AI 的時代會變成 superficial 的，沒什麼太大意義。假設 AI 告訴醫生應該這樣診斷，他會面臨到的困境是，當 AI 已經如此強大，有多大的信心做出相反的決策。所以要真正解決問題，不是只給他人為介入的可能性。
 - ◆ 在討論 GDPR 時，也有提到人為介入可能是假的，沒什麼太大意義，就算人類在法律上有實質改變的可能，但實際上也很難做出不同的判斷。例如，某個人就是沒買保護貼，分行經理要不要推翻 AI 貸款給他，可能是很大的挑戰。
- Q：本文與 GDPR 的關係？
- A：發想時與 GDPR 有關，但 GDPR 對此問題的還在討論階段，解釋到什麼程度、操作上要怎麼做都還不清楚，也有很多挑戰——演算法太複雜、商業機密、只告訴你考量了什麼因素。如果解釋的內容很表面，那其實等於沒有，沒有辦法真的解決我們面臨的困境。GDPR 自己也不知道他們 explanation 的內涵是什麼，技術還在發展，

甚至有人否定 right to explanation 的存在。總之，因為不滿意 GDPR 的處理，所以提出本文。

■ Q：Causal explanation 要怎麼具體實現？

■ A：13 條有可能成為其法律基礎，然而立法者在當時確實沒有想到 Causal explanation，只想到透明。本文做的是下一個層次，如果真的要透明，應該要提供因果說明。把因果說明植入現在的法律架構（22 條提供了基本框架，當重大決策使用了自動化演算法，必須提供當事人權益保障）就可以了。

■ Q：如果要立法來實現這個觀念，是通案性還是事務性的？

■ A：通案性的，連結著決策，才會有因果說明的義務。Judgement（怎麼開車）並非決策（會影響權利義務關係）。縱然放寬，把開車也當成決策，也不會是「實質影響當事人法律或相似權利，且要求較嚴格的佐證資料」，不會要求有 Causal explanation 才能向右轉，如果是好的 prediction，就可以接受該決策。

■ Q：假設某位專門治頭痛的醫生，有非常豐富的經驗，再累積許多病例後，發現早上起床會頭痛的人，是因為本身有憂鬱症，但病人沒有意識到他有憂鬱症，該醫生建議病人朝憂鬱症的方向去治療。像這樣的決策，要求醫生說明因果關係，但目前僅存在關聯性，必須進一步的研究才知道中間的因果關係，此時怎麼課與醫生義務有說明義務？

■ A：

◆ 醫生在診療時，會進行鑑別診斷（differential diagnosis），用排除法，慢慢的試，排除各種疾病，但不會一下子就做不可回復的處置。如果不是影響重大的處置，只是在嘗試各種治療選項（例如以憂鬱症的藥物治療看看），在不會有太多副作用的情形下，這樣的成本效益是大家接受的；如果是影響重大的處置，則會採用更多的診斷手段來確認是何種疾病，才會採取行動。採取何種處置，也會影響 AI 在此的說明、解釋義務。

◆ 這種 Association Discovery AI 的優點是，提供了我們採取行動的提示，但不應該是我們做成最後判斷的依據。所以在做嘗試時，她是一個好的工具，例如找不出其他頭痛原因時，AI 告訴我們可能是憂鬱症。他會是科學研究的第一步，但不是最後一步。

■ Q：網路仲介借錢與還錢的 P2P 公司，自己有一套特殊的機制，與聯徵中心、銀行徵信均不同，例如有幾個手機號碼、多久換一次手機號碼（這些看似普通，但從統計數據上來看，與是否還款有關），但不可透露。原因是，一旦公布，就可能有人去製造信用，影響他的評估。

■ A：

◆ 貸款會是「實質影響當事人法律或相似權利，但不要求較嚴格的佐證資料」的類型，需提出因果說明，或提出可以訴諸無知的正當化理由。當然涉及公益，例如偏鄉、特殊扶助的貸款時，會更複雜。

◆ 這個模式沒辦法解決所有的 AI 應用的態樣，但是他可以至少在理想型的單純情況（例如單純的貸款），可以去思考，講不出因果關係時，訴諸無知能不能夠被

正當化。

- ◆ 在貸款的情形，通常是可以被正當化的，畢竟要不要貸款是銀行決定的，銀行就是做一個預測。若更複雜的因素出現時，特定群體（例如黑人）無法取得貸款時，這時候訴諸無知可能會有問題，可能有偏誤存在。

講者二：莊錦秀助理教授（東華法律系）

本次會議的主題之二是探討 AI 的醫療民事責任。內容摘要如下：

- 傳統醫療態樣：
 - ◆ 醫師親自診療患者（理學檢查、檢驗）->臆診（可能是疾病 A、B、C）->針對 ABC 為相應處置->持續觀察病患（持續觀察注意義務）->患者病況變化（轉好、穩定、轉壞）注意義務及防果義務->再次診療
- 傳統醫療發生不利結果時之責任：
 - ◆ 歸責主體：
 - 有可歸責或違約：
 - ◇ 傳統醫療：行為人或機構
 - ◇ 遠距醫療？
 - ◇ AI 醫療？
 - 純粹是自然病程或不幸：患者自行承擔
 - ◆ 民事責任：
 - 侵權行為：行為人行為、損害發生、因果關係
 - 契約：契約目的違反、損害發生、因果關係
 - 醫療機構之責任：契約責任主體、或因違反監督義務而負擔侵權責任
- AI 醫療發生不利結果時之責任：
 - ◆ AI 定義廣泛，可能包含機器學習、神經網絡、深度學習
 - ◆ 目前許多醫院均有機器學習層級之 AI（北醫、長庚），甚至要求醫生寫結構式病例，以利輸入、接軌
 - ◆ 目前國內引進的 AI，無法解釋 AI 做成某建議之原因
 - ◆ 目前北醫、長庚均在訓練機器，將 AI 原本內建的歐美資料庫轉換為國內的資料庫
 - ◆ 若未來出現不利結果時，是哪一個環節出錯（AI 原廠、各醫院餵的資料...），將無法舉證
 - ◆ 美國「學習型的中間人」（類似沙盒理論），醫生若使用 AI 發生不利結果，由於尚在研發階段，故不課與責任
 - ◆ 報告人無法接受此理論，既然雙方訂定醫療契約，亦未告知將使用 AI，不應由病患承擔 AI 尚未成熟之不利利益
- 機器學習演算法之基本模式：取決於演算法的設計、類型和可用的數據量
 - ◆ 輸入：輸入 醫生之紀錄、臨床結果、時間序、圖像

- ◆ 演算法：透過線性迴歸、神經網絡、決策術等模式，處理上述輸入之資料，做成預測
- ◆ 評估：衡量演算法之優劣（例如分類準確性，預測誤差，誤報率）
- ◆ 輸出：評估治療結果、判斷腫瘤位置、判讀心電圖

■ 機器學習演算法之責任歸屬

- ◆ 若演算法僅提供資訊（輔助），由醫師做最終之醫療決定，產生不利結果時較好處理，類似美日採取之見解。
- ◆ 若由演算法獨立做成決定，產生不利結果時應如何處理？消保法產品之問題、（民事）法人化、（刑事）嚴格責任

■ 遠距醫療：

- ◆ 特色：透過患者端之儀器或人員（患者自己、家屬、專技人員）搜集醫療資訊，由遠端的醫生根據此資訊做臆診及下處方，指示相關人員做處置，並將處置後的結果傳給醫生，醫生會再次為相應的處置
- ◆ 由於責任較難釐清，故遠距醫療僅適用於緊急治療中之國際會診，或偏遠地區病患之簡易治療（例如調整藥物），以及在家之安寧病人
- ◆ 可能產生的問題：
 - 在哪個環節出問題，舉證更困難
 - 病人、家屬與有過失（未按照醫師指示）
 - 傳統侵權責任以行為人行為為判斷歸責依據是否妥適
 - 儀器維護之責任

■ 醫療之過失責任：

- ◆ 侵權責任：醫生可能會因未能獲得知情同意，或違反適當的照護標準的義務，致患者受損害
- ◆ 照顧標準：醫療常規
 - AI 輸入的可能是美國的醫療常規，與台灣的不同
 - 醫療常規可能會因為各地區資源之不同而有不同，但機器學習會採用統一的標準

■ 封閉循環：

- ◆ 機器學習仰賴的是人類醫師過去一百年努力的成果，若未來 AI 取代人類，人類醫師消失後將導致資料庫變成封閉式的循環（均為醫療 AI 所做的決策），將不再有創新，人類醫師的培訓也將產生問題。
 - 可能可以立法強制臨床上必須加入人類醫師，避免陷入封閉循環
 - 照顧義務標準之形成，亦需有人類醫師之參與

■ 機器學習訓練與隱私保障

- ◆ 違反醫療法 72 條：醫療機構及其人員因業務而知悉或持有病人病情或健康資訊，不得無故洩漏。
- ◆ 違反醫師法 23 條：醫師除依前條規定外，對於因業務知悉或持有他人病情或健康資訊，不得無故洩露。

- ◆ 病患之資料限於醫療使用，若要用於研究需經病患同意

Q&A

■ Q：

- ◆ 醫療過去的補償機制（疫苗、藥害救濟）是否可以運用在發展 AI 上？
- ◆ 既有的民刑事責任認定上，是以人類行為為標準，而現在變成物（AI）的行為時，應如何處理？
- ◆ AI 是由醫療機構所研發，難道我們要為醫療機構的疏失負責？
- ◆ 若要參考疫苗、藥害救濟之補償機制，要如何提撥？
- ◆ Q：在官僚體系追不上科技進步的情況下，除了仰賴主管機關，是否有替代機制（例如民刑事責任、自律組織）？

■ A：

- ◆ 課稅，再提撥到補償機制中（類似保險）
- ◆ 醫療組織非常重倫理難以自律

■ Q：嚴格責任之適用？

■ A：

- ◆ 適用於刑事責任，以處理法人格化後之刑事責任
- ◆ 追問：有趣的是，在法律史上，隨著醫藥發展，過去刑法的過失理論不斷放寬（對被告有利），鼓勵科技發展。而若如報告人所說，反而會有遏止創新的結果。
- ◆ 在無法管控時，先往嚴格責任走，未來再調整

■ Q：自然人的醫療過失，目前的走向有除罪化的傾向，AI 與自然人在此似乎是朝著兩個完全不同的方向走？

■ A：

- ◆ 醫療法修正並非完全除罪，報告人亦不贊成此方向，應該從增加人員設備、安全措施，處理醫療機構的問題，來保護醫療人員，而非除罪化。
- ◆ 若醫療人員可以除罪化，所有專技人員均應除罪化（例如消防人員亦同樣緊急）。

■ Q：遠距醫療現況？

■ A：目前仍由醫師在遠端接收資訊後開立處方、給藥，未來可能由機器為之

■ Q：使用病患資料訓練 AI，或是由 AI 為醫療行為，是否須告知病患？

■ A：前者目前沒有告知病患；後者需告知病患

■ Q：為什麼前者不告知？也許病患很樂意提供資料？

■ A：

- ◆ 成本很高、自找麻煩、他們也自認有權利這樣做。更大的問題是健保署要求醫院把影像資料都上傳健保署，健保署手上有一堆影像資料，未來可能會拿出來用（爭議在於是否符合「健保行政目的」）。
- ◆ 做得比較好的是台北榮總，在將醫療影像用於 AI 學習之前會先告知（opt-out 的方式），以掛號信告知既有的病患，真的 opt-out 的病患比例很低。

第七次會議紀錄	
時 間	108 年 1 月 25 日(星期五)
主 題	作為思想試驗的電車難題—自駕車碰撞倫理系統初探
內容摘要	
講者：陳弘儒助理教授（清大通識中心） 本次會議的主題是探討作為思想試驗的電車難題。內容摘要如下： 一、為何要處理此議題？ ■ 有用論： ◆ 觀點 1：電車難題是自駕車必然會面對的問題 ◆ 觀點 2：電車難題對於思索自駕車之法正倫理管制有助益 ● 電車難題一開始出現是為了處理墮胎問題 ■ 無用論： ■ 本文論點： ◆ 無用論忽視了「AI 替人們做決定」的意義，判斷權限移轉至 AI ◆ 無用論是將「技術與工程」與「價值或目的」分開，但技術中立是不可能的 二、本文基本論點 ■ 電車難題「不僅」是一實例。 ◆ 「思想試驗」是電車難題的重要定位，其關鍵貢獻在於針對「道德直覺」或「判斷」提出反省。 ◆ 為何法律需要考量思想試驗的討論？ ● 法律判斷對於社會行動的道德屬性具有定錨功能。 ● 「決定的道德屬性」會影響到法律如何針對「決定」的法律判斷進行論證。如果這個行為有、沒有道德屬性（=直行或左轉都符合道德），會影響法律在此的角色，影響法律決定（直行或左轉）的理由。 ■ 作為思想試驗的電車難題可以對自駕車之碰撞倫理系統提出貢獻 ◆ 協助自駕車的碰撞系統的倫理考量與人類判斷相趨近。 ● 讓我們知道我們有哪些道德直覺 ◆ 碰撞系統的價值排序進行釐清。 ● 效益主義、有沒有道德義務要選擇損害最小 ◆ 直覺反應 (human-like) v. 理性判斷 (rational agency) ● 在駕駛行動中，我們常常是靠直覺在反應。Ex：人們會直覺地把方向盤往左打，保護自己。 ● 未來技術成熟以後，AI 可以理性計算，我們要決定哪些要哪些不要納入考量。 ◆ 設計者負有一個可被證成的義務，需要強制揭露其如何設計碰撞系統之倫理考	

量因素。(ex 極度保護駕駛的設計)

三、法律管制面的思考

- 4.1. Top Down v. Bottom Up
- mandatory v. non-mandatory
 - ◆ 要/不要有統一的倫理判斷系統
- 價格 (price) 與倫理判斷 (moral judgment) 的關係
- 歸責問題
- 因果關係

四、何謂思想試驗？

- TE (Thought Experiment) 的三個要素：
 - ◆ (1). 對於場景或某個問題的介紹。
 - ◆ (2). 實驗者在被給訂條件下，進行判斷或是得到某個結果。
 - ◆ (3). 實驗者針對結果得出某個結論。
- 對於 TE 的質疑
 - ◆ (1). 想像與真實之間的關係？
 - ◆ (2). 檢驗方式的不可複製性？

五、思想試驗的功能

- 認識論上的功能：
 - ◆ (1). 建構性：指出某個命題是可能成立的。
 - ◆ (2). 摧毀性：某個命題（或主張）不一致之處。
- 闡述性或修辭性：
 - ◆ 讓某個觀念更容易理解。
- 啟發式功能（最重要的功能）
 - ◆ (1). 獨立道德判斷因素，思考特定因素是否有重要性！
 - ◆ (2). 對於不同理論相互比較，思考為何相異理論判斷會有差異！
 - ◆ 單純，不允許加入其他條件，在經過比較以後，我們可以知道哪些因素影響我們的判斷

六、思想試驗的基本特徵

- 例子本身不具有重要性！
 - ◆ 關鍵不在於他是電車、或是他會撞死人，而是他揭示了...
 - ◆ 倫理學上的 TE 重點在於思考，道德判斷的依據 (ground or reason) 何在！
 - ◆ 思想試驗與實證科學沒有互斥！
 - 很多思想試驗後來可以被實證科學驗證
 - ex 近期出現透過實驗的方式，來驗證人的道德直覺
 - ◆ 敏感於「道德事實」(moral facts) 或因素。

七、Foot 如何討論電車難題- 反對雙重效果原則

- 脈絡：墮胎的 intention (好的效果) 是為了救媽媽，就算胎兒會死 (不好的效果) 是可預見的

■ 反對「雙重效果原則」(the doctrine of double effect)

◆ Foot 認為「可預見某個結果發生」跟「有意圖讓某個結果發生」的區分是有問題的！

- 可以找到某個例子，就算都有好的/不好的效果，但我們的判斷是不一樣的

◆ Foot 認為 DDE 無法解釋為何，某個行動 A，促成某個效果 B 是可以預見 (foresee) 時，但不是我們所意圖 (intend) 時，A 有時是道德上許可的，有時候是被禁止的。這意味著，DDE 不是一個完整的原則。

- 論證：

✧ 提出法官難題作為比較：出現了一個兇手，某暴民說如果法官沒有做成對的決定，就要雪洗某村莊，我們不會認為法官在此應該陷害一個無辜者來救村莊

■ Foot 提出的主張- 積極義務 v. 消極義務

◆ 義務的類型：

- 積極義務：每個人都負有積極照顧或是提供救助的義務。
- 消極義務：每個人不得造成他人傷害的義務。

◆ 義務的優先性：消極義務 > 積極義務

◆ 同類義務相衝突時，則可以透過效益主義式的方式比較哪個義務應該優先被遵守。(哪個選擇死的人比較少)

◆ 因此，「效益主義」沒有最優先的位置！

八、Thomson 對於電車難題的論證

■ 認為 Foot 的說明有問題，採用積極與消極義務的論證將導出「司機必須 (must) 轉向電車」。

■ Thomson 認為，原則上司機是允許 (permissible) 轉向電車，但是若他不轉向也是可以的！換言之，在轉向與否上，司機具有一個道德自由。

■ Foot 的可能錯誤

◆ 按照 Foot 的論證，「司機」跟「乘客」轉動方向盤有何不一樣？

- 司機「必須」轉動方向盤，因為兩個消極義務衝突。
- 乘客「必須」啥都不做，因為是消極與積極義務之衝突。

◆ 行動者之角色差異似乎不足以造成我們道德判斷上的差異！Something is wrong here!

■ Thomson 如何論證這個道德自由- 健康卵石例子

- 有一個神奇的健康石頭 (摸到他的人會痊癒)，在河面上飄著，漂向某個人 A，你 (分配者) 有機會把這個石頭轉向另外五個人 B，這些人沒有石頭都會死。
- 一個人可以改變「善的分配狀態」，若且為若，預期得到善的享有者相較於其他人沒有更多主張 (例如契約) 可以提出之時。
- 在電車難題的例子中，Thomson 上述的主張可轉變成：「電車司機在道德上可以將電車轉向 (道德自由=道德上不可非難的行為)，若且為若以下條件

滿足：鐵軌上的那一個人相較於其他五個人沒有更多主張可以提出之時！」

九、電車難題對於碰撞倫理系統的啟發

- 1. 哪些因素是碰撞倫理系統「應該」要考量的，不是一個純粹「技術上的課題」。因素的考量本身就是一個「道德上」的決定，決定了擁有哪些特徵的人們具有（算計）理性決策下的優先性。
- 2. 相較於人類駕駛的碰撞行動的反應，自駕車的碰撞反應更具有工具理性的運用，也因此，更要在碰撞倫理系統設計上敏感於道德因素。
- 3. 自駕車設計者「必須」揭露碰撞系統中採用了哪些倫理考量！
- 4. 在管制面上，法規必須思索歸責與因果關係的問題。
- 5. 懸而未決的議題是，是否設計者真是事先決定了碰撞可能時的「算計理性」的判斷標準？（如果用 machine learning 的方式，是不是「設計者事先決定的」）->會影響設計者的責任

Q&A：

- Q：為什麼不是設計者決定的？
- A：
 - ◆ 要釐清哪些是「事先被決定的」，data input 是，但 machine learning 後的結果要再思考。
 - ◆ ex：uber 死人案：uber 辨識出這是行人，但他不煞車的原因是，他判斷這是偽陽性的人（原本認知他是人，後來判斷是誤判）。
 - ◆ 基礎事實錯誤 v. 倫理考量
 - uber 案裡設計者決定了倫理判斷，但事實認定錯誤。
 - 不可能完整考量所有的倫理，當涉及沒有事先考量到的問題，還是有成本分配的問題
- Q：為什麼不是人就可以撞上去？違反我們開車的常理
- A：在碰撞系統中，加入哪些決定因素是道德判斷，必須受到檢視。
- Q：
 - ◆ 當設計者沒有設定明確的規則，AI 自己形成規則，在此怎麼辦？
 - ◆ 除非設計者刻意餵歧視性資料，形成錯誤規則，否則要怎麼課責？
 - ◆ 行為與結果間因果關係判斷困難時，民法過去的作法：
 - 從寬認定因果關係
 - 不要考慮行為，直接認定物跟結果的因果關係，有管領權限者負責
- A：
 - ◆ 贊成把行為拿開
 - ◆ 規則：
 - 程式設計者給的規則
 - pattern recognition，但 pattern 未必是行為準則
- Q：這種尖銳問題是不是交給機器比較好？

- A：AI 的確更理性，但我們要考慮要不要把這部分讓出去。
- Q：電車難題（前提是煞車失靈）跟自駕車好像不太一樣？自駕車也要假設煞車失靈時，要怎麼做決定
- A：自駕車是由許多系統組合而成
 - ◆ 的確有可能
 - ◆ 電車難題中，結果是固定的，自駕車是機率問題
- Q：
 - ◆ 有沒有新的理論？
 - ◆ 當車商將倫理系統外包給其他公司，刑事法中的討論是可否主張專業人士？
 - ◆ 如果我們將醫療交給機器，是否要設定前提條件
- A：40 年的討論都沒有結論，只是現在有 AI，想要把原本直覺決定理性化，會不會緣木求魚？道德直覺的顯明化。

第八次會議紀錄

時 間 108 年 2 月 27 日(星期三)

主 題 Medical Robots: A New Challenge to Bioethics?

內容摘要

講者：吳全鋒副研究員（中研院法律所）

本次會議的主題是 Medical Robots: A New Challenge to Bioethics？內容摘要如下：

一、問題意識：AI 運用在醫療上，會不會對 Bioethics 造成根本價值的改變

- 所討論的 Medical Robots (AI) 不限於巨量資料的分析（仍掌握在手中），已超過廠商設計的反應範疇，可以自主做成決定
- 衝突原因：當 AI 出現時，會對傳統生命價值產生挑戰。Ex 健保署希望在分析後，判斷醫療資源該怎麼分配；以機器動手術時，產生問題，應該由醫生或機器做決定？
 - ◆ 傳統生命倫理的基礎在於「醫生-病人」「研究者-被研究者」的衝突，但 AI 是由非醫學領域者發展出來的，未必熟悉生命倫理。
 - ◆ 挑戰 1：帶入其他領域的專業、商業模式，與生命倫理衝突時該怎麼解決？
 - ◆ 挑戰 2：醫病關係是不對等的，所以會用比較高的強度去管，AI 也是嗎？會不會限縮了 AI 的發展？

二、Autonomy 自主原則：

- 每個人應該得到充分資訊，尊重各種價值，由病人自主選擇
- informed consent: well informed, free consent
- 挑戰：
 - ◆ 有沒有把病人關心的價值，都放進系統了？
 - ◆ 醫師只做醫療判斷，但病人會有其他考量，系統有辦法全部納進去嗎？

- ◆ 假設可以把所有 consideration 都納入系統...
- ◆ 在「搜集資料-產生 model」模式中，是否考量了少數群體
- ◆ 若把 AI 當成工具，最後由人做成決定。給所有資訊：給過多的資訊，也會造成病人沒辦法決定。給 material information：怎麼決定哪些是 material information？
- ◆ Black box：在 machine learning 下，醫師也不知道是怎麼做成的
- ◆ 誰負 informed consent 的義務？醫生？廠商？

三、Beneficence and Nonmaleficence

- 醫生必須以病人的利益為最高考量
- 發展醫療機器人時，把科學事實和價值分開，只要在科學事實下，能做出正確判斷的機器人就是好的。很難在考量利益時，納入每個人重視的價值。
- 縱然嘗試量化，每個文化重視的價值也不同，不會是一個統一的標準
- 某些人主張必須把主觀價值、情感因素排除，AI 正可以幫助人做成客觀判斷。回應：...我們是不是真的想要客觀的決定，是不是最有利的
- 健康也強調 social functioning，人的情緒會影響診斷。若某人因情緒失落拒絕吃藥，是否要強迫她吃藥。此為醫療 AI 的特殊之處，涉及人的情感

四、Justice

- 藥品的智慧財產權，會影響不同國家國民的健康
- 有能力發展 AI 者，直接決定了 SOP，決定哪些資料應該被標準化...。會對目前沒有能力發展者，造成深遠的結果
- 在一國之中，也會惡化城鄉差距，
- 若在沒有排除 stereotype，沒有納入價值，就直接做成決定...
- 傳統上，資料的匿名化、去識別化足以保障病人隱私，但在 AI 之下會有隱私問題。影像資料要怎麼去識別化。

Q&A

- Q：Bioethics 跟規範一樣嗎？可以直接轉換嗎？規範可能促進/阻礙科技發展，倫理也是嗎？
- A：部分核心項目，已經變成法律的要求。有些沒有那麼直接，會用其他形式存在於法律中。Bioethics 比較像專業自律，比較軟性的規範。
- Q：這是不是所有 AI 都會涉及的問題—AI 有沒有辦法處理價值？如果最後有人為介入，是不是可以解決這個問題？
- A：歐盟的立論基本上是，AI 是不涉及價值的，只是個工具，所以 Bioethics 可以直接套入。
- Q：的確所有 AI 都會受到挑戰，但醫療涉及病人，且已經有既定的原則了，直接套用可以嗎？
- A：用 informed consent 轉嫁風險（不管基本的安全...）
 - ◆ 研究倫理：知識生產過程中，風險是內含的無法排除的風險
 - ◆ 臨床倫理：起碼把風險、安全掌握到一定程度，經過 balance，決定讓他上市
 - ◆ AI 的情況下是沒有人知道，永遠也沒辦法知道（不像研究倫理）

- Q：把它當成醫療工具，會不會使監管手段更多？
 - ◆ 把它當成手術刀，審查時會漏掉很多東西
 - ◆ 從醫學史的角度來看，似乎是醫師責任越來越低，醫療器具製造商責任越來越重
 - ◆ 把 AI 當成工具，是不是醫師對自身專業的捍衛
- A:
 - ◆ 醫師的重要性可能越來越低
 - ◆ 回應：在民事或刑事責任上，也許會出現相反的主張
- Q：有沒有可能用 informed consent 解決一切？
- A：
 - ◆ 例外：沒有 informed 也可以做
 - ◆ 例外：同意也不免責
- Q：informed consent 要保護什麼？提升正確性？Autonomy？
- A：
 - ◆ 提升做自主決定的 capability
 - ◆ 在遇到真的無知的時候，求籤是理所當然
- Q：
 - ◆ 科學知識有自己的邏輯、內涵，當 AI 提出的建議，沒辦法被科學知識解釋，那要怎麼辦？
 - ◆ 醫療問題好像很快就可以判斷出哪個 AI 比較精準？
- A：非不可能分類人類偏好，提供各個病患選擇，但目前都沒有討論

第九次會議紀錄

時 間 108 年 3 月 28 日(星期四)

主 題 人工智慧之成果與著作權保護之探討

內容摘要

講者：王怡蘋教授（台北大學法律系）

本次會議的主題是人工智慧之成果與著作權保護之探討，內容摘要如下：

一、背景：AI 可以寫文章、寫詩集、譜樂曲、畫畫。慢慢地看到 AI 的成果開始進入市場，在著作權角度上就會出現問題

二、問題意識：

- 成果是否/要不要受到著作權保護
- 權利怎麼歸屬
- 輸入 AI 學習用資訊之性質（重製權）

三、AI 成果侵害他人之著作權

■ 英國：承認

- ◆ 由電腦生成之文學、戲劇、音樂、藝術著作，已採取創造該著作所必要安排者為其著作人。
- ◆ 以程式設計者為著作權人，受到著作權保護

■ 歐盟：有爭議

- ◆ 2017 關於機器人之民事規範建議
- ◆ 關於機器人無需特別規定，僅需解釋既有法律
- ◆ 但最困擾的地方是「自己精神創作」
- ◆ AI 的學習模式與人類相似，但並無創作的想法

■ 德國：很難受到著作權保護

- ◆ 著作權法：受本法保護之著作僅只人類精神創作受保護
- ◆ 精神在於，著作完成時，即享有之天賦權利，本質即為對人之保護
- ◆ 可能以著作鄰接權保護

■ 我國：

- ◆ 法律未特別說明要保護的是什麼，從立法精神（第一條）也看不出來傾向歐陸（保障人格）/英美（促進文化、產業發展）
- ◆ 判決常引用之文句：「凡具有原創性...等人類精神力參與的創作，均受著作權法所保護」
 - 反面是否即不受保護
- ◆ AI 之成果要怎麼保護
 - 無著作鄰接權
 - 但把著作鄰接權塞進不同的規定裡

四、到底要不要保護？

■ 保護理由：

- ◆ 誘因理論：以提供創造的誘因，確保為發展 AI 技術之投資得以獲得回報
- ◆ 市場失靈理論
- ◆ 但是賦予專屬權，也意味著限制競爭。促進跟限制之間要如何拿捏？

五、如果要保護他，要怎麼保護？

■ 著作鄰接權—以電子資料庫為例（歐盟資料庫保護指令）

- ◆ 保障目的：投資，包含時間工作經歷
- ◆ 內容：重置、散布、公開傳達權
- ◆ 對象：建置者
- ◆ 保護期間：15 年

■ 著作鄰接權—以表演人為例

- ◆ 表演人的原創性好像沒那麼高，他只是在表演已存在的作品
- ◆ 權利內容：
 - 著作財產權：錄音錄影其表演、重製與散步錄音物語錄影霧
 - 著作人格權：姓名表示權、禁止醜化權

◆ 保護期間：

- 著作財產權：50-70 年
- 著作人格權：50-70 年。表演人死亡即消滅（維護的是人與作品之間的連結）

六、如果以著作鄰接權保護 AI 的成果

■ 目的：保護投資

■ 權利歸屬：

◆ 程式設計者？

- 電腦程式已受著作權保護，產出成果還要給他嗎？

◆ AI 擁有者？

◆ 利用人？

- 允許約定權利歸屬於利用人？

■ 權利內容：

◆ 著作財產權：等同於現行之著作財產權

◆ 著作人格權？

- AI 不具有人格，其本身不該享有
- 著作人格權原則上不應歸屬於非創作之人

◆ 台灣的奇怪規定：可以約定雇用人或出資人取得著作人格權

- 美國：只取得著作財產權，而非人格權
- 姓名表示權：權利 or 義務

◆ 要不要強制標示他是 AI 的作品，或是否可以拿著 AI 的成果，聲稱是自己的來投稿。

◆ 這樣的話還是權利嗎？或者是義務？

- 禁止醜化權：

◆ 可否任意變更 AI 之成果，甚至醜化

◆ 該給誰？跟著作財產權一起賦予給同一人？

■ 保護期間？

◆ AI 生產速度？AI 生產速度快

◆ AI 成果品質？品質差是否僅給予低度保護？

七、輸入 AI 學習用資訊之性質：

■ AI 需要大量資料來學習

◆ 透過分析這些資料，來發展規則

◆ 重製權？

■ 歐盟議會關於資料探勘之規定：

◆ 基於學術研究目的，關於重製與擷取，得有例外規定

■ 德國著作權法第 60d 條

- ◆ 為學術研究之目的，可以自動化、系統性重製原始資料，建立資料系統。
- ◆ 不得以營利為目的

■ 直接適用資料探勘的規定？

- ◆ 學術研究（意義、界線為何？）、非營利（只拿工本費算是盈利嗎？）的限制
- ◆ 後續的商業利用與獲利？
- ◆ 是否需要針對 AI 設計之授權模式？
 - 需要極大量資料
 - 目的為供 AI 學習，以建立規則
 - AI 成果未必會出現他人著作之內容

八、AI 成果侵害他人之著作權

- AI 成果整體近似於他人著作
 - ◆ 容易處理，可透過比對篩出
- 擷取他人著作之部分（ex 擷取他人作品有特色的部分）
 - ◆ 比例小，較無法用比對判斷
- 目前之判斷標準
 - ◆ 實質近似
 - ◆ 接觸：後作品的作者有無接觸過前作者的作品
- AI 可否適用目前之標準
 - ◆ 在資料庫內的東西，必然接觸過。資料庫越大越容易成立抄襲？
 - ◆ 人類學習也是從模仿開始，但我們期望人類最終得以創造、轉化。但 AI 會嗎？
- 侵權責任之歸屬（應不應該有人負責、誰該負責）
 - ◆ 「擷取他人著作之部分」是否有控制、防免的可能性？
 - ◆ 更根本的，誰負責，侵權責任的基礎為何？防免 or 損益同歸
 - 現在的侵權責任多建構在防免（是否可歸責），惟亦有損益同歸的想法
- 防免：程式設計者
- 損益同歸：AI 擁有者、利用人
 - ◆ 過失/無過失責任

Q&A：

- Q：人類接觸的知識越多會不會使個人越容易？
- A：「接觸」在人類方面就很困擾了，人類更難去知道創作人知道了什麼（AI 則可以明確的知道資料庫裡面有什麼），必須要去證明創作人有無連續或不中斷的接觸過前作品。在人類的情形下，我們好像不在乎「接觸」與「侵權」之間的因果關係。另外，著作權重要的不是成果，而是內在反省後的啟發。
- Q：禁止醜化（有）與著作人格權（無）的差異？
- A：沒有禁止醜化的話，任何人都可以去改 AI 的成果嗎？還是要禁止大家任意的變更他。這確實是對著作人格權的挑戰
- Q：AI 的創作意圖？
- A：學習過程與人類的確相似，但無想要創作的意圖或想法，也少了內在反思，所以不會想用著作權來保護他，用其他方式（著作鄰接權）來保障投資。
- Q：抄襲在著作權法的意涵？
- A：

◆ 抄襲非著作權法中的名詞，實務上通常是用在整體性的抄襲（而非部分的借用）。但著作權中還有合理使用的規定，範圍還蠻寬的

◆ 學術倫理的關懷重點：

- 避免學術評價錯誤，錯誤的成果歸屬
- 著作權的關懷重點則有市場的考量

■ Q：著作可能都受言論自由的保障，在此為什麼要有原創性？著作權的門檻可能會限制到言論自由

■ A：

◆ 著作權=言論自由+...。著作權也可能使取得前人資料門檻太高，影響到言論自由

◆ 法律讓著作權人享有排他權利，影響市場，故只保護具有原創性的內容

■ Q：為什麼人類接觸大量訊息時，我們不會認為他有侵害著作權的疑慮，而 AI 有？

■ A：

◆ 人類會用合法的管道來取得，背後即有著作財產權再處理。AI 的問題在於建立學習資料庫時，必須先轉換為數位碼，即有重製的問題。

◆ 著作權法第 3 條第 1 項第 5 款對「重製」的定義：「重製：指以印刷、複印、錄音、錄影、攝影、筆錄或其他方法直接、間接、永久或暫時之重複製作。於劇本、音樂著作或其他類似著作演出或播送時予以錄音或錄影；或依建築設計圖或建築模型建造建築物者，亦屬之。」

■ Q：為什麼著作權的保護要有原創性？如果他是人類創作的必要要素，那 AI 的作品是否先天就不具備原創性，所以著作權拿來規範 AI 作品就是不可能的，那是不是用另一套制度來規範他比較好。而重置在 AI 學習過程中的意義也不同，是否可以？

■ A：

◆ 著作權是一種專屬權，所以需要門檻來區分受著作權保護/不保護的言論，我們區分的方法就是其對於人類文化是否具有貢獻。

◆ 如果從嚴來解釋原創性，AI 作品的確不可能受著作權保護；若從寬解釋，則可能受保護。

◆ 在創作上的確跟人類不一樣，因為當我們要保護其作品時，必然要畫線，也許借用著作權的規定（原創性）會比較簡單。

◆ 目前沒有太確定的想法，但基礎都是保護投資，若市場機制足以回收成本，那也未必需要法律管制。

■ Q：若著作權跟言論自由是交錯的，那他必然有一塊是不受言論自由保障的，那這部分是法律上 or 憲法上權利？這個「原創性」的條件是立法者 or 憲法給的？

第十次會議紀錄

時 間	108 年 4 月 26 日(星期五)
-----	---------------------

主 題	人工智慧治理的法學研究路徑
內容摘要	
講者：劉靜怡教授（台大國發所）	
本次會議的主題是人工智慧治理的法學研究路徑，內容摘要如下：	
一、背景：在 AI 的趨勢之下，雖然不敢說未來一定會如何如何，但就像過去網際網路初期我們想像的美好發展未必出現，反而出現了許多問題。	
■ 前言：人工智慧如何走到這裡	
◆ 從科技發展史來看，人工智慧已經持續發展了很久，在產業上有非常多元的應用。	
◆ 各國都在談，是否會因國家而有不同？歐美在此領域是否落後於中國？	
二、人工智慧型塑的世界：朝向某種「簡化」的方向發展？	
■ 自動化、規格化，越來越習慣接受「演算法」的建議？	
◆ 疾管署醫生做了一個毒蛇系統，利用 AI、演算法、影像辨識，把蛇照片上傳到系統上，就可以判讀出蛇的種類、毒性、如何處理。（開發者在此背後的責任？）	
◆ 人本來就有惰性、缺乏專業知識，這種 AI 系統的吸引力是非常大的。	
◆ 光現在的購物、電影建議，大家都會被吸引了，更何況是不熟悉的專業領域，會更依賴 AI。	
■ 應該如何管制此模式？	
◆ 機器未經明確人工或程式指示而「自動學習」的進展將愈加明顯，對於規範體制而言，必須如何因應？	
三、人工智慧型塑的世界：法律人（至少應該）看到（了）什麼（？）	
■ 誰的演算法？誰的智慧？	
■ 誰的「人格」？誰是權利義務主體？誰的決定？誰的責任？誰有「能力」負責？「負責」的意義是什麼？	
■ 責任分配：誰來承擔風險？AI 會不會改寫法律上權利義務關係？	
■ 「演算法社會」中的誰來「監督」與「管制」的課題？	
四、人工智慧時代的法理學爭辯路徑	
■ 道德在決定法律內涵時所扮演的角色	
■ 法律在指引我們的行動時，應該扮演的角色為何？如果，我們有遵守法律的道德義務，那麼，我們在決定法律的內涵時，是否必須參照道德的要求來做決定？	
■ 假設有朝一日，在預測法律評價方面，AI 科技比人類更為進步準確（AI 更能預測法院判決結果）。	
■ 除了法律問題之外，法律人往往也必須納入道德爭議的考量——例如遵守法律的道德義務，或者是在決定如何行動時將道德列入考量——，AI 在這些問題上，可以取代人類嗎？	
◆ 統計上顯示自駕車比人類親自駕車還要安全的說法？	
● 目前自駕車還沒上路，如何得到這樣的判斷？	

- ◆ Google 與 Facebook 等科技公司認為依照目前 AI 科技的迅速發展，未來人工智慧在道德判斷上甚至法律判斷上的表現，可能會比人類更好的主張
- ◆ 可能出現 AI 在分析過程中及分析結果上無法排除某些可能會導致「偏見」(bias) 和「差別待遇」(discrimination) 的因素，因此就可能產生道德、倫理甚至法律層面的爭議
- ◆ 這樣的偏見、歧視、差別待遇可能是無法察覺的，畢竟 AI 的決策過程不透明
 - AI 透明化的困難：
 - ◇ 時間不夠充裕
 - ◇ 演算法不斷演化
- ◆ 無法理解 AI 如何「思考」(think)：一方面可能是人類「不夠聰明」，所以無法理解其思考模式，但是，也可能是因為 AI 的運算模式，對於人類而言，並無意義
 - GDPR 中被認為是自動決策解釋權(right to explanation)的規定，在實踐上，自然就會產生困難

■ AI 道德判斷的模式

- ◆ Top-down approach：提供抽象原則，指引 AI 做判斷
 - 我們目前對很多事情沒有一致的道德標準，那要提供 AI 什麼原則？
- ◆ Bottom-up approach：給他各式各樣的情境&結果，用案例累積(類似 common law)，來學習如何做道德、法律判斷
 - 是否有辦法有效排除干擾因素，把道德/法律判斷區分出來。例如中文的法律論述可能很艱澀，人都不一定看得懂了
- ◆ 缺點：
 - 無一致的道德標準
 - 可能會複製法院或法學傳統既有的偏見與錯誤
 - AI 的法律詮釋可能可以很符合法律，但沒辦法提供正當性（詮釋結果能夠提供給法律怎樣的道德程度）

- 更嚴肅地看待道德在法律中所扮演的角色，進一步地推進我們對道德在法律詮釋中應該扮演什麼角色的理解

五、人工智慧對法律制度的影響：以行政管制為例

- 過去常仰賴電腦來分析資訊、效益評估
- 利用 AI 決定判決時的刑期長短或者假釋與否等決策
 - ◆ AI 需不至於正當化甚或放大了某些長期存在的偏見
 - ◆ 須維持決策的監督與審查
 - 人類心靈活動本身就有複雜性和不透明性，AI 模仿了人類心靈活動，會加劇不透明
 - 不能以科技為中立的為由來拒絕監督，且縱然沒有更好的方法，提出此方法可能的問題就是很大的貢獻了。
- 行政管制：

- ◆ 機器學習至少到目前為止應該依然無法直接處理行政管制規範中複雜的面向，理由在於要理解行政管制規範背後的規範意旨，往往並不是透過機器學習單純的預測功能，就可以達成任務
- ◆ 禁止授權原則：
 - 法院還是可能會斷然否決行政機關採取機器學習行使行政管制權力所得的管制結果
 - 法院還是可能會斷然否決行政機關採取機器學習行使行政管制權力所得的管制結果
- ◆ 正當法律原則：
 - AI 決策的過程應如何規範？人在此過程中的角色？如何論證及說明理由？
 - Right to Exit：是否應該容許受到 AI 決策影響的個人，在符合一定條件時，可以選擇「退出」由 AI 做成的執法決策？甚至是 affirmative 的 opt-in 才能對我使用 AI 決策

六、Jack Balkin 的「資訊受託人」(information fiduciaries) 主張：

- 客戶、消費者及最終使用者的關係：
- 「資訊受託人」(information fiduciaries) 的角色
- 非客戶、消費者及最終使用者：
 - ◆ 演算法使用者則應該負擔公共責任 (public duties)，應該避免將演算法運用所帶來的成本和傷害直接予以外部化
 - ◆ 私部門使用 AI 涉及或影響公共利益時即應負擔公共責任
- FDA 的管治模式：
 - ◆ 根據演算法的效能、可預測性和可解釋性建立分類：上市前許可、設計測試和執行指引

Q&A：

- Q：關於 AI 對法律產生的衝擊，我們不怕人類的不透明，因為我們有各種機制可以控制人類，而且人類是多元的，可以平衡各自的缺陷。但 AI 的錯誤可能是系統性的，導致我們特別害怕。兩種可能的管制模式：控制 AI 讓它不要太大，讓他不要產生系統性風險 / 鼓勵發展，盡可能地多元？
- A：
 - ◆ 因為不了解，也不知道怎麼了解 AI (甚至是研發者) (在此宣稱 AI 是中立的也很奇怪)，所以導致害怕。
 - ◆ 蘇：我們害怕的是人無法制衡。例如，我們可以舉反證推翻測速照相。但對於那些難以舉證推翻的，人類會比較害怕，例如酒測、DNA。
 - ◆ 林：之所以會害怕是因為不夠瞭解 (想我們就不會害怕測速照相)；法律學者也不該把 AI 當成一個概念來討論，應該要在各領域具體特定應用下，來討論各種風險，所以 FDA 模式從風險角度來管制，也許是正確的。
- 甲、
- Q：

- ◆ AI 是主體/客體？測速照相是單獨作成決定嗎？還是受到人類的某種控制？
- ◆ AI 的問題真的那麼大嗎？我們在怕什麼？AlphaGo 打敗人類代表什麼意義？AI 最恐怖的會不會還是與公權力的結合？
- A：AlphaGo 代表有一種智慧出現了，高於人類且不可控制。
- Q：警方配戴隨身攝影機，拍攝執法過程，交給媒體播放，法律問題在哪裡？
- A：
 - ◆ 佩戴所取得的影像在刑事訴訟法上有證據能力
 - ◆ 但與「配戴隨身攝影機，拍攝執法過程」是否合法、正當無關。
- Q：為什麼相信 AI 會比人做出更好的價值/法律判斷？人對 AI 的想像到底是什麼，如果只是複製人類，看不出有什麼吸引力，似乎是要創造出神？
- A：這些公司都有科技樂觀論，過去也沒什麼失敗經驗。而他們對價值/法律的想像太簡化。
- Q：機器期待的是更快的處理資料，對 AI 的期待是更好的學習能力？
- A：AI 應用以服務人類。

第十一次會議紀錄

時 間 108 年 5 月 30 日(星期四)

主 題 命運的轉軸點？AI 時代政治資金管理之展望與挑戰

內容摘要

講者：陳柏良博士（中研院法律所博士後研究）

本次會議的主題是 AI 時代政治資金管理之展望與挑戰，重點包含政治資金 保障言論自由、預防貪腐等，內容摘要如下：

一、序論

- 言論自由的預設
 - ◆ 資訊稀缺：供給極少
 - ◆ 閱聽者具備充分的資訊審議時間：需求極大
 - ◆ 政府是意見自由競爭市場的主要威脅
 - ◆ 極重要政府意義定義為何
 - ◆ 取得言論自由和貪腐的平衡點
- 廉價自由時代
 - ◆ Eugene Volokh ：主張廉價言論(cheap speech)
 - ◆ Tim Wu ：言論為審查者的武器
 - ◆ Gary King：
 - 言論自由戰略

- 轉移話題(distraction)
- 西瓜效應(cheerleading) 製造大量人工帳號轉貼文
- 預防使用者在實體世界的集體行動(collective action)

■ 言論審查戰術

- ◆ 線上騷擾與攻擊發言者
- ◆ 海量資訊的洗版
- ◆ 控制主要言論平臺

■ 廉價言論的大量出現

- ◆ 資訊氾濫(information flood)
- ◆ 閱聽者關注產業(attention industry)的興起 閱聽者的關注是稀有財，具備極高經濟價值。
- ◆ 過濾氣泡（同溫層效應） 閱聽者關注行銷或仲介業者，過濾資訊以提升使用者黏著性

二、AI 是否可以作為溝通言論的主體？有無主體身份？

■ AI 是否需要區分？

- ◆ 強 AI：可以模仿人類並具有創造性、社交能力
- ◆ 弱 AI：大量複製或模仿人類的行為

■ AI 產出之內容（演算輸出）是否要歸類為政治言論處理？是否為言論？國家是否可以進行管制？

- ◆ 內容創造：針對不同議題自動轉推內容自動調整
- ◆ 內容分享：單純轉貼

三、AI 的演算輸出是否是「表意」？

■ Stuart Benjamin：

- ◆ 應設立寬泛的標準
 - Sendable
 - receivable
 - actually choose

◇ → 停留在弱 AI 的時代，本質仍言論，AI 則為工具

■ Tim Wu：

- ◆ 憲法僅保障具有表意成分(expressive aspect)的言論
- ◆ 搬運/ 傳遞訊息：僅為行為，並非言論
- ◆ 高度工具性資訊：非屬言論（google map）

四、AI「表意」類比

■ 動物言論

■ 法人言論

- ◆ 1976 Virginia State Pharmacy Board v. Virginia Citizens consumer Council
 - ⇒ 言論自由的保障轉從 speaker 到 listener 閱聽人資訊近用權

五、政治資金與言論

- 是否設立獨立監管機關（台灣體系不同）
- 捐獻是否有門檻上限
- 獨立政治開支（民眾）
 - ◆ 1976 Buckley v. Valeo 案（言論越多越好）
- 捐錢的行為本身就是政治言論
- 政治資金管理不能用平等權衡量（排除平等權）
- 政治言論自由 v. 政府是否有反貪腐的重大政府意涵
- 僅有門檻限制，避免買通
- 台灣政治獻金法：收受政治獻金
 - ◆ 政治獻金之定義
 - ◆ 其他經濟利益之定義
 - ◆ 獨立開支去管制，無須揭露
 - ⇒目前獨立開支規避政治獻金法，獨立開支應管制及揭露
 - ⇒機器大量貼文時，公投法廣告是否應須揭露
- 美國聯邦法
 - ◆ 其他利益(anything of value)：是否如此限縮
 - ◆ 獨立支出強制揭露
 - ◆ 是否設下審議性之標準

Q&A

- Q：
 - ◆ 面對 Cheap speech 問題，是否應提倡「免於受到隱私干擾的自由」？Google map 是否屬於言論？
 - ◆ 共和主義的對比是政治主義和共和主義式，閱聽人權利有無限制，權利基礎何在？保障利益為何？listener 及 speaker 中間的中介者，是否可以規範不得藉由資訊篩選創造出不同的社會實相？會影響行動者很具體的行動，同一套「管制」？
- A：資訊近用權 閱聽人權利是否是重大的公眾利益，正確義務 (accuracy) 不得有欺詐、猥褻，設有限制，不能 micro targeting。
- Q：是否要求大家都看同一個東西？
- A：除非行銷手法會造成損害，不然如何限制媒體，另外一個問題是資訊多元、混亂，獨家新聞時間差不能認證是真假新聞，將會落入相同的論證，單獨來看每個面向都是真的，但同一事情有不同八個面向，問題難管制。
- Q：世界能否出現複數的真實？若允許，即可以允許複數的權利？是否可以聚合一些價值？
- A：應該是仍允許可以有複數的真實，對於中介媒體自己不能有複數。
- Q：政治獻金來源大宗？美國不透明，台灣對自然人的上限很低，法人上限額度較高，傾向法人做政治獻金，AI 誰來決定可以做政治獻金？(AI 也是一種政治獻金的手

段)？

- A：弱 AI 保障背後的那個人，強 AI 是否要創設新的概念規範，應揭露。
- Q：Sunstein republic.com 網路會極化，社會就會極化，Sunstein 如何回應關於特定政治人物粉絲的引戰？如何回應想看垃圾言論的人？想看假話的人？→人民有說謊的自由，但人民有無基於自我實現接收謊言的自由？
- A：什麼是自由？審議和審視之後才是自由？做自己喜歡做的事就是自由？市場競爭式的民主：投票就是民主，審議式民主：政治審議過程，關鍵非投票，而是經過層層審議過程（是否公開、論證過程）。
- Q：閱聽人若有權利，是否有責任？為何？基礎假設「言論越多越好」，在 AI 的時代之下，是否仍是如此？假設仍存在嗎？
- A：資訊守門員或濾鏡的角色，通常是視聽媒體，因此有針對媒體的公平原則 diversity，強調 second order diversity？要求電視台平衡報導，網路新聞是否也可以如此？到了政治場域，即為政黨，之後政黨及視聽（中介團體）弱化之後（因網路崛起），當有更多機會接觸到權利後，是否有義務存在？義務來源？如何課予？資訊是否越多越好？反過來，若資訊超過人類時間乘載量時，自己如何篩選？社群網路篩選？我們還剩下什麼權利及義務？政府是否可以介入管制？
- Q：回到政治獻金，如何解釋柯文哲、連勝文關於政治獻金的現象？核心在於在無法管制的狀況下，就放任？實行自我實現？
- A：美國猥褻定義，脫衣舞仍被歸納在言論自由，免於受到隱私干擾的自由，應該把道德排除。

第十二次會議紀錄

時 間 108 年 6 月 14 日(星期五)

主 題 「合理的人」判斷標準之省思—以 AI 醫療觀察為中心
AI as Medical Device （AI 醫療器材的管制）

內容摘要

講者：莊錦秀助理教授（東華法律系）

本次會議的主題之一是探討以 AI 作為醫療器材的相關問題。內容摘要如下：

一、前言：

- 醫療用品之自主性計算功能：
 - ◆ AED：觀測、給予人類指示
 - ◆ 心電圖：過去 12 張圖/R3 才能判讀，目前由系統協助判斷
 - ◆ 功能性核子共振 Fmri：觀測組織內部血流是否正常
 - ◆ 心律調整器：偵測病人心跳並回應之

- 目前均用醫療過失/產品責任處理
- 醫療專家系統 (ex MYCIN、Watson)
 - ◆ 問題輸入->建議 (使用者平台<->推理機<-知識庫 (人類專家提供知識)
 - ◆ 醫療過程：病徵->輸入專家系統->可能的各種疾病 ABC，以及各疾病可能性的百分比

二、AI 系統造成損害的民事責任

- 醫療過失：醫療服務提供者
 - ◆ 契約
 - ◆ 侵權行為法：
 - 注意標準
 - 過失：可預見風險下採行為違反合理人標準
 - ◆ 替代責任 187、188
 - ◆ 產品責任 (消保法)：警告買方原則、學習中介原則 (爭議：AI 非產品，而是內部軟體)
 - ◆ 企業責任：AI 所有人負被害人損賠，內部再依實際具體分擔責任
- 目前第一代 watson：採企業責任/未來感知系統：準法人

二、 計算機所致侵權

- 從嚴格責任 (無過失責任) 轉為過失責任
 - ◆ ex Fmri 技師錯誤識別操作 64 切導致患者受傷害 (技師過失)
 - ◆ ex Fmri 系統缺陷，採嚴格責任
 - ◆ ex 由計算機操作 Fmri，錯誤識別操作 64 切導致患者受傷害？
 - 目前多數見解傾向 2.3 相同處理
 - 本文認為 1.3 相似，宜相同處理
- 過失責任：若 Watson 未判讀出一般醫師可判讀之資訊，應負責
- 合理的計算機標準
 - ◆ 目前的過失標準：可預見情況下，所為行為違反合理人標準。
 - ◆ 行業慣用、平均或最安全技術

Q&A

- Q：目前 Watson 需要執照嗎？
- A：台灣：只要符合醫療器材管理辦法就好 (只要經過國外批准，台灣只是形式審查)
- Q：合理的計算機標準是否會反過來影響傳統的醫療慣習？
- A：AI 必定比人類強 (沒有比人類強就沒有使用必要)，其輸入的資料是人類專家提供，兩者差異應該不會太大
如果認為 AI 必定比人類強，會導致人類不敢去挑戰 watson 的決定。
- Q：當醫師遇到一個全新的案件，注意義務會放寬。若 Watson 有學習功能，遇到相同情況，似乎不會比較小心而直接產出結果
- A：若有學習系統，則為準法人
- Q：「當未來 AI 系統比人類行為更安全更精準」之命題，是有爭議的，縱然上述為真，

是否代表人類要讓渡人類活動權能給 AI？在此不只是風險計算，亦包涵人類行為的自由

- A：也許有些領域 Ai 就是比人類強，人類通通不准做。
- Q：「當未來 AI 系統比人類行為更安全更精準」是規範還是事實？

講者：邱文聰副研究員（中研院法律所）

本次會議的主題之二是探討以 AI 作為醫療器材的相關問題。內容摘要如下：

一、問題意識：管制模式，以及可能遇到的問題

二、AI 作為醫療器材

- 藥事法：軟體為醫療器材的一種
- 其他國家 Software as a medical device：沒有被包含在其他硬體醫療器材的 Software

三、管制方式：

- 過去的軟體不會改變，上市時是什麼就是什麼
- 有自主學習能力的軟體會不斷變動，上市許可要怎麼核發？
 - ◆ FDA 未來可能的做法：
 - 對未來的可能修正提出計畫，包含確認安全性有效性的方法
 - ◇ SaMD Pre-Specification：預計會修正的地方
 - ◇ Algorithm Change Protocol：遇到風險應如何處理
 - ◆ 從飛航軟體的角度來看，也許每次修正都要重新申請許可
- 可能的管治模式：
 - ◆ 依照生理狀態的嚴重程度(嚴重 or 輕微)/應用的方式 (ex 診斷)，區分不同的管治方式
 - ◆ 有效性評估方式

Q&A

- Q：如果把演算法分成：傳統的/視當時情形再做成決定/會自我學習的，自我學習目前存在嗎？
- A：

存在

 - ◆ 飛航軟體的創新誘因高、創造力超強，但風險非常高，所以還是會要求每次更新都要經過批准，也不會允許沒有人的民用航空器
 - ◆ FAA 飛航系統考量因素：創新、安全性、自我學習能力/速度
 - ◆ 回應莊老師上一場：也許過失責任+舉證責任可以兼顧「提供廠商創新誘因&維持安全誘因」

計畫執行收穫

經由 1 年 12 次密集不間斷的會議與討論，本研究群計畫的執行有以下的收穫：

- 一、AI 法學的橫向認識與聯繫：結合不同領域的法學者，探討 AI 與基本權利、行政管制、智慧財產、刑事責任、民事責任、金融管理等各種法律議題。
- 二、AI 法學的縱向思考與延伸：在各自不同法領域中，開掘並拓展新興研究議題與思考資源，包括言論自由的、權利主體的反思、金融治理的新模式、醫療人權的新視野、司法制度與實踐的新發展等。
- 三、AI 法學與人文科學的整合：將哲學、倫理學、社會學的元素與理論導入法律問題的探討，形成法學與人文科學的科際整合。
- 四、AI 法學與自然科學的交流：透過自然科學領域學者的演講與討論，認識並吸收自然科學知識，一方面補充法學思維的不足，另一方面強化與其他領域的對話。
- 五、AI 研究群的建立與運作：結合不同領域的法學者，透過定期會議以及年度工作坊的舉辦，業已形成國內 AI 的法學研究群，對於促進各年齡層法學者投入 AI 的法律研究，頗有助益。
- 六、AI 基礎研究的奠基與開展：本年度透過持續的 presentations 與 discussions，逐漸彙整出 AI 與法律的基礎研究，為下一年度 AI 專題的開展與深化奠下基礎。
- 七、AI 法學集體論著能量的籌備：本年度的成果即將透過集體論著的方式呈現，形塑法學專題合著的典範。

參考文獻

中文部分

期刊論文

- 王伯藹 (2018)。〈美國聯邦審計署表人工智慧技術評估報告〉，《科技法律透析》，30:8 期，頁 4-5。
- 王明禮 (2014)。〈論資訊隱私：科技與商業發展脈絡下的觀察 (Information Privacy: A Contextualized Analysis)〉，《中原財經法學》，32 期，頁 59-105。
- 甘琳 (2018)。〈歐盟執委會公布歐洲人工智慧通告〉，《科技法律透析》，30:6 期，頁 6-7。
- 李建良 (2018)，〈New Brave World?—AI 與法律、哲學、社會學的跨界遇合〉，《人文與社會科學簡訊》，第 20 卷第 1 期，頁 75-80。
- 何亦婕 (2015)。〈日本推動智慧醫療照護與巨量資料應用之趨勢觀察〉，《科技法律透析》，27 卷 12 期，頁 51-69。
- 吳旻純 (2011)。〈人工智慧新革命——超級電腦「華生」〉，《生活科技教育月刊》，44 卷 5 期，頁 1-9。
- 吳柏凭 (2016)。〈人工智慧對於著作權概念的衝擊——日本著作權的新政策發展方向〉，《科技法律透析》，28 卷 12 期，頁 26-31。
- 宋皇志 (2017)。〈人工智能在專利檢索之應用初探〉，《全國律師》，21 卷 10 期，頁 27-37。

- 林利芝 (2018)。〈初探人工智慧的著作權爭議—以「著作人身分」為中心〉，《智慧財產權》，237 期，頁 61-78。
- 林芝余、陳婷 (2018)。〈人工智慧與雲端運算法律相關議題〉，《理律法律雜誌雙月刊》，107 年 1 月號，頁 4-5。
- 林勤富 (2018)。〈人工智慧法律議題初探〉，《月旦法學雜誌》，274 期，頁 195-215。
- 張保生 (2001)。〈人工智能法律系統的法理學思考〉，《法學評論》，19 卷 5 期，頁 11-21。
- 張麗卿 (2019)。〈人工智慧時代的刑法挑戰與對應——以自動駕駛車為例〉，《月旦法學雜誌》，286 期，頁 87-103。
- 曹源 (2016)。〈人工智能創作物獲得版權保護的合理性〉，《科技與法律》，2016 卷 3 期，頁 488-508。
- 郭雨嵐、汪家倩、侯春岑 (2017)。〈法律科技與人工智慧時代，科技法律人才的養成與挑戰〉，《萬國法律》，214 期，頁 51-59。
- 陳良基 (2017)。〈打造人工智慧創新環境機制〉，《國土及公共治理》，5 卷 4 期，頁 60-71。
- 陳譽文 (2017)。〈人工智慧規範性議題綜觀〉，《科技法律透析》，29 卷 4 期，頁 43-51 頁。
- 黃崧洺 (2018)。〈從人工智慧與法律科技展望法律服務之價值——以專利服務為例〉，《專利師》，35 期，頁 70-80。
- 黃詩淳、邵軒磊 (2017)。〈運用機器學習預測法院裁判——法資訊學之實踐〉，《月旦法學雜誌》，270 期，頁 86-96。
- 楊秋敏 (2016)。〈AI 人工智慧與老人照護〉，《開南法學》，8 特刊，頁 219-242。
- 廖淑君 (2011)。〈智慧聯網之發展與個人資訊隱私保護課題：以歐盟之因應為例 (Internet of Things vs. information privacy protection: take European Commission's actions as an example)〉，《科技法律透析》，23 卷 11 期，頁 18-42。
- 潘俊良 (2017)。〈美國白宮公布「為人工智慧的未來準備」報告〉，《科技法律透析》，29 卷 1 期，頁 5-7。
- 潘俊良 (2017)。〈簡析德國自動駕駛與車聯網發展策略〉，《科技法律透析》，29 卷 4 期，頁 25-33。
- 潘俊良 (2018)。〈自駕車之發展與挑戰—以德國法制為借鑑 (Development and Challenges of Autonomous Driving—Taking the German legal Issues as a reference)〉，《科技法律透析》，30 卷 12 期，頁 48-72。
- 蔡碩庭 (2018)。〈網路服務提供者侵害著作權之民事責任 (Civil Liability for Copyright Infringement of Internet Service Providers)〉，《智慧財產評論》，15 卷 1 期，頁 163-216。
- 蕭仁豪 (2018)。〈日本人工智慧 (AI) 發展與著作權法制互動課題之探討〉，《科技法律透析》，30:1 期，頁 46-72。

- 賴俊傑 (2009)。〈迎接一個人與機器人共存的社會—日本次世代機器人安全性確保準則之淺介〉，《科技法律透析》，21 卷 6 期，頁 2-7。
- 關光威 (2008)。〈E 世代知識管理管理平臺中隱私權與智慧財產權法的爭議 (A New Mode of Knowledge Transfer and the Relevant Intellectual Property and Privacy Issues in Taiwan)〉，《智慧財產評論》，6 卷 2 期，頁 79-92。
- 羅文妙 (2018)。〈淺談機器人及其專業佈局〉，《理律法律雜誌雙月刊》，107 年 1 月號，頁 9-10。

專書論文

- 劉靜怡 (2018)。〈人工智慧潛在倫理與法律議題鳥瞰與初步分析—從責任分配到市場競爭〉，收於：劉靜怡(主編)，《人工智慧相關法律議題芻議》，頁 3-45。台北:元照。
- 顏厥安 (2018)。〈人之苦難，機器恩典必看顧安慰—人工智慧、心靈與演算法社會〉，收於：劉靜怡(主編)，《人工智慧相關法律議題芻議》，頁 50-85。台北:元照。
- 吳從周 (2018)。〈初探 AI 民事責任—聚焦反思臺灣之實務見解〉，收於：劉靜怡(主編)，《人工智慧相關法律議題芻議》，頁 89-116。台北:元照。
- 李榮耕 (2018)。〈初探刑事程序法的人工智慧應用—以犯罪熱區為例〉，收於：劉靜怡(主編)，《人工智慧相關法律議題芻議》，頁 120-148。台北:元照。
- 邱文聰 (2018)。〈初探人工智慧中的個資保護發展趨勢與潛在的反歧視難題〉，收於：劉靜怡(主編)，《人工智慧相關法律議題芻議》，頁 153-175。台北:元照。
- 沈宗倫 (2018)。〈人工智慧科技與智慧財產權法制的交會與調和—以著作權法與專利法之權力歸屬為中心〉，收於：劉靜怡(主編)，《人工智慧相關法律議題芻議》，頁 181-214。台北:元照。
- 黃居正 (2018)。〈與人工智慧相關的國際法議題—從國際人道法到生命體法〉，收於：劉靜怡(主編)，《人工智慧相關法律議題芻議》，頁 219-241。台北:元照。

英文部分

專書

- Ajay Agrawal, Joshua Gans & Avi Goldfarb (2018). Prediction Machines: The Simple Economics of Artificial Intelligence. Cambridge, MA: Harvard Business Review Press.
- Ethem Alpaydin (2016). Machine Learning: The New AI. Cambridge, MA: The MIT Press.
- Joseph E. Aoun (2017). Robot-Proof: Higher Education in the Age of Artificial Intelligence. Cambridge, MA: The MIT Press.
- Kevin D. Ashley (2017). Artificial Intelligence and Legal Analytics: New Tools for Law Practice in the Digital Age. New York: Cambridge University Press.

- Nick Bostrom (2014). *Superintelligence: Paths, Dangers, Strategies*. New York: Oxford University Press.
- Meredith Broussard (2018). *Artificial Unintelligence: How Computers Misunderstand the World*. Cambridge, MA: The MIT Press.
- Erik Brynjolfsson & Andrew McAfee (2014). *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies*. New York: W.W. Norton & Company.
- Taina Bucher (2018). *If...Then : Algorithmic Power and Politics*. New York: Oxford University Press.
- Ryan Calo, A. Michael Froomkin & Ian Kerr eds. (2016). *Robot Law*. Cheltenham, UK/Northampton, MA: Edward Elgar Publishing.
- Fred H. Cate & James X. Dempsey eds. (2017). *Bulk Collection: Systematic Government Access to Private-Sector Data*. New York: Oxford University Press.
- John Cheney-Lippold (2017). *We Are Data: Algorithms and the Making of Our Digital Selves*. New York: NYU Press.
- Samir Chopra & Lawrence White (2011). *A Legal Theory for Autonomous Artificial Agents*. Ann Arbor, Michigan: University of Michigan Press.
- Brian Christian & Tom Griffiths (2017). *Algorithms to Live By: The Computer Science of Human Decisions*. New York: Picador.
- Glenn Cohen, Holly Fernandes Lynch, Effy Vayena & Urs Gasser, *Big Data, Health Law, and Biethics*, New York: Cambridge University Press.
- Ronald K. L. Collins & David M. Skover (2018). *Robotica: Speech Rights and Artificial Intelligence*. New York: Cambridge University Press.
- Paul R. Daugherty & H. James Wilson (2018). *Human + Machine: Reimagining Work in the Age of AI*. Cambridge, MA: Harvard Business Review Press.
- Pedro Domingos (2015). *The Master Algorithm: How the Quest for the Ultimate Learning Machine Will Remake Our World*. New York: Basic Books.
- Paul Dumouchel, Luisa Damiano & Malcolm DeBevoise (2017). *Living with Robots*. Cambridge, MA: Harvard University Press.
- Virginia Eubanks (2018). *Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor*. New York: St. Martin's Press.
- Ariel Ezrachi & Maurice E. Stucke (2016). *Virtual Competition: The Promise and Peril of the Algorithm-Driven Economy*. Cambridge, MA: Harvard University Press.
- Joshua A. T. Fairfield (2017). *Owned: Property, Privacy, and the New Digital Serfdom*. New York: Cambridge University Press.
- Cyrus Farivar (2018). *Habeas Data: Privacy vs. the Rise of Surveillance Tech*. New York: Melville House.

- Primavera De Filippi & Aaron Wright (2018). *Blockchain and the Law: The Rule of Code*. Cambridge, MA: Harvard University Press.
- Ed Finn (2017). *What Algorithms Want: Imagination in the Age of Computing*. Cambridge, MA: The MIT Press.
- Franklin Foer (2017). *World Without Mind: The Existential Threat of Big Tech*. New York: Penguin Press.
- Martin Ford (2018). *Architects of Intelligence: The Truth about AI from the People Buidling it*. Packet Publishing.
- Brett Frischmann & Evan Selinger (2018). *Re-engineering Humanity*. New York: Cambridge University Press.
- Sean Gerrish & Kevin Scott (2018). *How Smart Machines Think*. Cambridge, MA: The MIT Press.
- George Gilder (2018). *Life After Google: The Fall of Big Data and the Rise of the Blockchain Economy*. New York: Gateway.
- David Gray & Stephen Henderson eds. (2018). *The Cambridge Handbook of Surveillance Law*. New York: Cambridge University Press.
- Ian Goldin & Chris Kutarna (2016). *Age of Discovery: Navigating the Risks and Rewards of Our New Renaissance*. New York: St. Martin Press.
- Stephen Goldsmith & Susan Crawford (2014). *The Responsive City: Engaging Communities Through Data-Smart Governance*. New York: Jossey-Bass.
- David Gray (2017). *The Fourth Amendment in an Age of Surveillance*. New York: Cambridge University Press.
- Jeffrey Greene (2014). *Moral Tribes: Emotion, Reason, And the Gap Between Us and Them*. New York: Atlantic Books.
- David J. Gunkel (2017). *The Machines Question: Critical Perspectives on AI, Robots, and Ethics*. Cambridge, MA: The MIT Press.
- David J. Gunkel (2018). *Robot Rights*. Cambridge, MA: The MIT Press.
- Woodrow Hartzog (2018). *Privacy's Blueprint: The Battle to Control the Design of New Technologies*. Cambridge, MA: Harvard University Press.
- Mireille Hildebrandt (2015). *Smart Technologies and the End(s) of Law: Novel Entanglements of Law and Technology*. Cheltenham, UK/Northampton, MA: Edward Elgar Publishing.
- Sharona Hoffman (2016). *Electronic Health records and Medical Big Data: Law and Policy*. New York: Cambridge University Press.
- Amir Husain (2017). *The Sentient Machine: The Coming Age of Artificial Intelligence*. New York: Scriber.
- Sarah E. Igo (2018). *The Known Citizen: A History of Privacy in Modern America*. Cambridge, MA: Harvard University Press.

- Engin Islin & Evelyn Ruppert (2015). *Being Digital Citizens*. London: Rowman & Littlefield International, Ltd.
- Meg Leta Jones (2016). *Ctrl + Z: The Right to Be Forgotten*. New York: NYU Press.
- Jerry Kaplan (2016). *Artificial Intelligence: What Everyone Needs to Know*. New York: Oxford University Press.
- Garry Kasparov & Mig Greengard (2017). *Deep Thinking: Where Machine Intelligence Ends and Human Creativity Begins*. New York: PublicAffairs.
- John D. Kelleher (2018). *Data Science*. Cambridge, MA: The MIT Press.
- Lucas Kello (2017). *The Virtual Weapon and International Order*. New Haven, Conn.: Yale University Press.
- Ronald Krotoszynski (2016). *Privacy Revisited: A Global Perspective on the Right to be Left Alone*. New York: Cambridge University Press.
- Susan Landau (2017). *Listening In: Cybersecurity in an Insecure Age*. New Haven, Conn.: Yale University Press.
- Josh Lauer (2017). *Creditworthy: A History of Consumer Surveillance and Financial Identity in America*. Columbia University Press.
- Patrick Lin, Ruan Jenkins & Keith Abney (2017). *Robot Ethics 2.0: From Autonomous Cars to Artificial Intelligence*. New York: Oxford University Press.
- John Markoff (2015). *Machines of Loving Grace: The Quest for Common Ground Between Humans and Robots*. New York: HarperCollins Publishers.
- Viktor Mayer-Schonberger & Thomas Ramge (2018). *Reinventing Capitalism in the Age of Big Data*. New York: Basic Books.
- Andrew McAfee (2017). *Machine, Platform, Crowd: Harnessing Our Digital Future*. New York: W.W. Norton & Company.
- Gina Neff (2018). *Self-Tracking*. Cambridge, MA: The MIT Press.
- Martha C. Nussbaum (2003). *Upheavals of Thought: The Intelligence of Emotions*. Cambridge University Press,
- Arlindo Oliveira (2017). *The Digital Mind: How Science is Redefining Humanity*. Cambridge, MA: The MIT Press.
- Cathy, O'Neil (2016). *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. New York: Crown.
- Ugo Pagallo (2013). *The Laws of Robots: Crimes, Contracts, and Torts*. New York, USA: Springer-Verlag.
- Frank Pasquale (2014). *The Black Box Society: The Secret Algorithms That Control Money and Information*. Cambridge, MA: Harvard University Press.
- Mark R. Patterson (2017). *Antitrust Law in the New Economy: Google, Yelp, LIBOR, and the Control of Information*. Cambridge, MA: Harvard University Press.

- Kenneth Payne (2018). *Strategy, Evolution, and War: From Apes to Artificial Intelligence*. Washington, D.C.: Georgetown University Press.
- Nick Polson & James Scott (2018). *AIQ : How People and Machines Are Smarter Together*. New York: St. Martin's Press.
- Jeremy Rabkin & John Yoo (2017). *Striking Power: How Cyber, Robots, and Space Weapons Change the Rules for War*. New York: Encounter Books.
- Jennifer Rothman (2018). *The Right of Publicity: Privacy Reimagined for a Public World*. Cambridge, MA: Harvard University Press.
- Ann Rudinow Sætnan, Ingrid Schneider & Nicola Green eds. (2018). *The Politics and Policies of Big Data: Big Data, Big Brother?* London: Routledge.
- Paul Scharre (2018). *Army of None: Autonomous Weapons and the Future of War*. New York: W.W. Norton & Company.
- Bruce Schneider (2018). *Click Here to Kill Everybody: Security and Survival in a Hyper-Connected World*. New York: W.W. Norton & Company.
- Klaus Schwab (2016). *The Fourth Industrial Revolution*. World Economic Forum.
- Maurice E. Stucke & Allen P. Grunes (2016). *Big Data and Competition Policy*. New York: Oxford University Press.
- Terrence Sejnowski (2018). *The Deep Learning Revolution*. Cambridge, MA: The MIT Press.
- Richard Susskind & Daniel Susskind (2015). *The Future of the Professions: How Technology Will Transform the Work of Human Experts*. New York: Oxford University Press
- Max Tegmark (2017). *Life 3.0: Being Human in the Age of Artificial Intelligence*. New York: Alfred A. Knopf.
- Edward Tenner (2018). *The efficiency Paradox: What Big Data Can't Do*. New York: Knopf.
- Joseph Turow (2017). *The Aisles Have Eyes: How Retailers Track Your Shopping, Strip Your Privacy, and Define Your Power*. New Haven, Conn.: Yale University Press.
- Jeffrey L. Vagle (2017). *Being Watched: Legal Challenges to Government Surveillance*, New York: NYU Press.
- Siva Vaidhyanathan (2018). *Antisocial Media: How Facebook Disconnects Us and Undermines Democracy*. New York: Oxford University Press.
- Paul Vigna & Michael J. Casey (2018). *The Truth Machine: The Blockchain and the Future of Everything*. New York: St. Martin Press.
- Paul Voigt & Axel von dem Bussche (2017). *The EU General Data Protection Regulation (GDPR): A Practical Guide*, Switzerland: Springer

- Wendell Wallach and Colin Allen (2009). *Moral Machines: Teaching Robots Right from Wrong*. New York: Oxford University Press.
- Wendell Wallach (2015). *A Dangerous Master: How to Keep Technology from Slipping Beyond Our Control*. New York: Basic Books.
- Ari Ezra Waldman (2018). *Privacy as Trust: Information Privacy for an Information Age*. New York: Cambridge University Press.
- Kevin Werbach (2018). *The Blockchain and the New Architecture of Trust*, Cambridge, MA: The MIT Press.
- John Frank Weaver (2013). *Robots Are People Too: How Siri, Google Car, and Artificial Intelligence Will Force Us to Change Our Laws*. New York: Praeger.
- Darrell M. West (2018). *The Future of Work: Robots, AI, and Automation*. Washington D.C.: Brookings Institution Press.
- Tim Wu (2018). *The Curse of Bigness: Antitrust in the New Gilded Age*, New York: Columbia Global Reports.
- Aleš Zavrašnik ed. (2018). *Big Data, Crime and Social Control*, London: Routledge.

期刊論文

- Benjamin Alarie et al., *Law in the Future*, 66 *University of Toronto Law Journal* 423 (2016).
- Jack M. Balkin, *The Three Laws of Robotics in the Age of Big Data*, 78 *Ohio State Law Journal* 1217 (2017).
- Jack M. Balkin, *The Path of Robotics Law*, 6 *California Law Review Circuit* 45 (2015).
- Solon Barocas & Andrew D. Selbst, *Big Data Disparate Impact*, 104 *California Law Review* 671 (2016).
- David J. Barron & Elena Kagan, *Chevron's Nondelegation Doctrine*, 2001 *Supreme Court Review* 201 (2001).
- Oren Bracha & Frank Pasquale, *Federal Search Commission - Access, Fairness, and Accountability in the Law of Search*, 93 *Cornell Law Review* 1149 (2008).
- Ryan M. Calo, *Peeping HALs: Making Sense of Artificial Intelligence and Privacy*, 3 *European Journal of Legal Studies* 2 (2010).
- Ryan M. Calo, *Open Robotics*, 70 *University of Maryland Law Review* 571 (2011).
- Ryan M. Calo, *Robotics and the Lessons of Cyberlaw*, 103 *California Law Review* 513 (2015).
- Ryan M. Calo, *Robots as Legal Metaphors*, 29 *Harvard Journal of Law & Technology* 209 (2016).
- Anthony J. Casey & Anthony Niblett, *Self-Driving Laws*, 66 *University of Toronto Law Journal* 429 (2016).

- Bryan Casey, Amoral Machines, or: How Robotists Can learn to Stop Worrying and Love the Law, 111 Northwestern University Law Review 1347 (2017).
- Ronald A. Cass, Vive la Deference?: Rethinking the Balance Between Administrative and Judicial Discretion, 83 George Washington Law Review 1294 (2015).
- Anupam Chander, *The Racist Algorithm?*, 115 Michigan Law Review 1023 (2017).
- Danielle Keats Citron, Technological Due Process, Washington University Law Review 85 (2008).
- Michael Chertoff (2018). Exploding Data: Reclaiming Our Cyber Security in the Digital Age, New York: Atlantic Monthly Press.
- Danielle Keats Citron & Frank Pasquale, The Scored Society: Due Process for Automated Predictions, 89 Washington Law Review 1 (2014).
- Kate Crawford & Jason Schultz, Big Data and Due Process: Toward a Framework to Redress Predictive Privacy Harms, 55 Boston College Law Review 93(2014).
- Cary Coglianese & David Lehr, Regulating by Robot: Administrative Decision Making in the Machine-Learning Era, 105 Georgetown Law Journal 1147 (2017).
- Cary Coglianese, The Emptiness of Decisional Limits: Reconceiving Presidential Control of the Administrative State, 69 Administrative Law Review 43 (2017).
- Cary Coglianese, Presidential Control of Administrative Agencies: A Debate over Law or Politics?, 12 University of Pennsylvania Journal of Constitutional Law 637 (2010).
- Andrea Scripa Els, Artificial Intelligence as a Digital Privacy Protector, 31 Harvard Journal of Law & Technology 217 (2017).
- Amitai Etzioni & Oren Etzioni, Keeping AI Legal. Vanderbilt Journal of 19 Entertainment and Technology Law 133 (2016).
- Ariel Ezrachi & Maurice E. Stucke, Artificial Intelligence & Collusion: When Computers Inhibit Competition, 2017 University of Illinois. Law Review 1775.
- Andrew Guthrie Ferguson, Big Data and Predictive Reasonable Suspicion, 163 University of Pennsylvania Law Review 327 (2015).
- Roger Allan Ford & W. Nicholson Price II, Privacy and Accountability in Black-Box Medicine, 23 Michigan Telecommunications & Technology Law Review 1 (2016).
- Erica Fraser, Computers as Inventors - Legal and Policy Implications of Artificial Intelligence on Patent Law. 13 SCRIPTed: A Journal of Law, Technology and Society 305 (2016).
- Ben Hattenbach & Joshua Glucoft, Patents in An Era of Infinite Monkeys and

Artificial Intelligence, 19 Stanford Technology Law Review 32 (2015).

- Mireille Hildebrandt, Law as Information in the Era of Data-Driven Agency, 79 Modern Law Review 1 (2016).
- Ignatius Michael Ingles, Regulating Religious Robots: Free Exercise and RFRA in the Time of Superintelligent Artificial Intelligence, 105 Georgetown Law Journal 507 (2017).
- Elizabeth E. Joh, Policing by Numbers: Big Data and the Fourth Amendment, 89 Washington Law Review 35 (2014).
- Elizabeth E. Joh, Policing Police Robots, 64 UCLA Law Review Discourse 516 (2016).
- Meg Leta Jones, The Ironies of Automation Law: Tying Policy Knots with Fair Automation Practices Principles, 18 Vanderbilt Journal of Entertainment and Technology Law 77 (2015).
- Ian Kerr, The Devils is in the Defaults, 1 Critical Analysis of Law 4 (2017).
- Joshua A. Kroll et al., Accountable Algorithms, 165 University of Pennsylvania Law Review 633 (2017).
- Sabine Gless, Emily Silverman & Thomas Weigend, If Robots Cause Harm, Who Is to Blame? Self-Driving Cars and Criminal Liability, 19 New Criminal Law Review 412 (2016).
- Pamela S. Katz, Expert Robot: Using Artificial Intelligence to Assist Judges in Admitting Scientific Expert Testimony, 24 Albany Law Journal of Science & Technology 1 (2014).
- Fazal Khan, The “Uberization” of Healthcare: The Forthcoming Legal Storm over Mobile Health Technology’s Impact on the Medical Profession, 26 *Health Matrix* 123 (2016).
- Pauline Kim, Data-Driven Discrimination at Work, 48 William & Mary Law Review 857 (2017).
- Bruce H. Kobayashi & Larry E. Ribstein, Law’s Information Revolution, 53 Arizona Law Review 1169 (2011).
- Joshua A. Kroll, Joanna Huey, Solon Barocas, Edward W. Felten, Joel R. Reidenberg, David G. Robinson & Harlan Yu, Accountable Algorithms, 165 University of Pennsylvania Law Review 633 (2017).
- Peter Margulies, Surveillance by Algorithm: The NSA, Computerized Intelligence Collection, and Human Rights, 68 Florida Law Review 1045 (2016).
- Toni M. Massaro & Helen Norton, Siri-Ously? Free Speech Rights and Artificial Intelligence. 110 Northwestern University Law Review 1169 (2016).

- John O. McGinnis & Russell G. Pearce, The Great Disruption: How Machine Intelligence Will Transform the Role of Lawyers in the Delivery of Legal Services, 82 Fordham Law Review 3041 (2014).
- Laura Myers, Allen Parrish, & Alexis Williams, Big Data and the Fourth Amendment: Reducing Overreliance on the Objectivity of Predictive Policing, 8 Federal Courts Law Review 231(2015).
- Cristian-Vlad Oancea, Artificial Intelligence Role in Cybersecurity Infrastructures. 4 International Journal of Information Security & Cybercrime 59 (2015).
- Paul Ohm, Broken Promises of Privacy: Responding to the Surprising Failure of Anonymization, 57 UCLA Law Review 1701 (2010).
- Paul Ohm, Sensitive Information, 88 Southern California Law Review 1125 (2015).
- Frank Pasquale & Glyn Cashwell, Four Futures of Legal Automation, 63 UCLA Law Review Discourse 26 (2015).
- Frank A. Pasquale, A Rule of Persons, Not Machines: The Limits of Legal Automation, _George Washington Law Review_ (2018).
- Frank A. Pasquale, Toward a Fourth Law of Robotics: Preserving Attribution, Responsibility, and Explainability in an Algorithmic Society, 78 Ohio State Law Journal 1243 (2017).
- Richard Primus, The Future of Disparate Impact, 108 Michigan Law Review 1341 (2010).
- Andrew Proia, Drew Simshaw & Kris Hauser, Consumer Cloud Robotics and the Fair Information Practice Principles: Recognizing the Challenges and Opportunities Ahead, 16 Minnesota Journal of Law, Science & Technology 145 (2015).
- Anjanette H. Raymond & Scott J. Shackelford, Jury Glasses: Wearable Technology and Its Role in Crowdsourcing Justice, 17 Cardozo Journal of Conflict Resolution 115 (2015).
- Neil M. Richards & Jonathan H. King, Big Data Ethics, 49 Wake Forest Law Review 393 (2014).
- Neil M. Richards & Jonathan H. King, Three Paradoxes of Big Data, 41 Stanford Law Review Online 41 (2013).
- David Marc Rothenberg, Can Siri 10.0 Buy Your Home? The Legal and Policy Based Implications of Artificial Intelligent Robots Owning Real Property, 11 Washington Journal of Law, Technology & Arts 439 (2016).
- Burkhard Schafer, Editorial: The Future of IP Law in an Age of Artificial Intelligence. SCRIPTed: 13 A Journal of Law, Technology & Society 283 (2016).

- Matthew U. Scherer, Regulating Artificial Intelligence Systems: Risks, Challenges, Competencies, and Strategies. 29 *Harvard Journal of Law & Technology* 353 (2016).
- Gregory Scopino, Preparing Financial Regulation for the Second Machine Age: The Need for Oversight of Digital Intermediaries in the Futures Markets, 2015 *Columbia Business Law Review* 439.
- Andrew D. Selbst & Julia Powles, Meaningful Information and the Right to Explanation, 7 *International Data Privacy Law* 233 (2017).
- Brian Sheppard, Incomplete Innovation and the Premature Disruption of Legal Services, 2015 *Michigan State Law Review* 1797.
- Reva B. Siegel, Equality Talk: Antisubordination and Anticlassification Values in Constitutional Struggles Over Brown, 117 *Harvard Law Review* 1470 (2004).
- Drew Simshaw, Nicholas Terry, Dr. Kris Hauser, & M.L. Cummings, Regulating Healthcare Robots: Maximizing Opportunities While Minimizing Risks. 22 *Richmond Journal of Law & Technology* 3 (2016).
- Omer Tene & Jules Polonetsky, Privacy in the Age of Big Data: A Time for Big Decisions, 64 *Stanford Law Review Online* 63 (2012).
- Darin Thompson, Creating New Pathways to Justice Using Simple Artificial Intelligence and Online Dispute Resolution, 2 *International Journal of Online Dispute Resolution* 4 (2015).
- Andrew Tutt, An FDA for Algorithms, 69 *Administrative Law Review* 83 (2017).
- Sandeep Vaheesan & Frank Pasquale, The Politics of Professionalism: Reappraising Occupational Licensure and Competition Policy, 14 *Annual Review of Law & Social Science* 309 (2018).
- Suzanne Van Arsdale, User Protections in Online Dispute Resolution, 21 *Harvard Negotiation Law Review* 107 (2015).
- Harry Surden & Mary-Anne Williams, Technological Opacity, Predictability, and Self-Driving Cars, 38 *Cardozo Law Review* 121 (2016).
- Cass R. Sunstein, Nondelegation Canons, 67 *University of Chicago Law Review* 315 (2000).
- Latanya Sweeney, Discrimination in Online Ad Delivery, 3 *ACMQueue* 11(2013).
- Nancy B. Talley, Imagining the Use of Intelligent Agents and Artificial Intelligence in Academic Law Libraries, 108 *Law Library Journal* 383 (2016).
- David C. Vladeck, Machines without Principals: Liability Rules and Artificial Intelligence, 89 *Washington Law Review* 117 (2014).
- John Villasenor, Technology and the Role of Intent in Constitutionally Protected Expression, 39 *Harvard Journal of Law & Public Policy* 631(2016).

- Albert H. Yoon, The Post-Modern Lawyer: Technology and the Democratization of Legal Representation, 66 University of Toronto Law Journal 456 (2016).
- Kenji Yoshino, The New Equal Protection, 124 Harvard Law Review 747 (2010).
- Larry N. Zimmerman, Artificial Intelligence in the Judiciary, 85 Journal of the Kansas Bar Association 20 (2016).