

目 錄

壹、計畫辦理情形

一、計畫基本資料

二、計畫摘要

三、執行成效綜合說明

貳、實際指標達成情形

一、實際達成指標一覽表

二、實際達成情形與預期達成情形不符原因說明

參、經費執行成效

一、經費執行總表

二、整體經費支用成效說明

肆、計畫辦理情形整體檢討說明

伍、附件

一、指標項目成果佐證資料

二、其他

壹、計畫辦理情形

一、計畫基本資料

計畫名稱	以區塊鏈為基礎建立安全與公平之代幣交易制度精進師生實務職能方案					
申請學校	國立臺北商業大學		領域/學門別		04商業、管理及法律領域 /041商業及管理學門	
執行期程	自113年8月1日起至114年7月31日止					
計畫主持人姓名	吳威震		單位/職級		副教授	
聯絡電話	0223226477		行動電話		0926813835	
E-MAIL	weichen@ntub.edu.tw					
學生參與名單 (在籍學生至少三名；不得為在職專班學生)	姓名	性別	年級	國籍	學制	系/科/所
	沈容	女	一	本國籍	大學(二技)	財務金融系
	吳若綺	女	一	本國籍	大學(四技)	資訊管理系
	丁俞豪	男	一	本國籍	大學(二技)	財務金融系

二、計畫摘要

本計畫提出以區塊鏈和智能合約架構下，建立一個以安全及可信任的區域市場可流通的代幣交易制度，其特性為安全及可信任、區域市場交易與可跨區域流通的代幣交易制度，以解決目前各區域市場各自發行的代幣無法流通，透過本計畫的架構下，未來可以建跨區域的聯盟鏈，以活絡代幣的流通性，促進與振興國內市場消費。

產業

學院

三、執行成效綜合說明

(一)計畫主持人服務研究之產出成果

1. 期刊論文(SCI):MYu-Chi Wei, Yu-Chun Chang, Wei-Chen Wu. Multi-language IoT Information Security Standard Item Matching based on Deep Learning, Computer Science and Information Systems, Computer Science and Information Systems 21(2):663 – 683, DOI:10.2298/CSIS230822012W
2. 研討會論文:Zhu-Xuan Huang, Yu-Chih Wei, Wei-Chen Wu (MLFSP 2025), Reversible Data Hiding in DICOM Images with Regional Sensitivity Analysis and Layered Protection, The 2025 International Conference on Machine Learning on FinTech, Security and Privacy (MLFSP 2024), Osaka, Japan.
3. 發明專利:金融機構之跨平台操作系統及其方法
4. 開授課程:保險科技與資安防護
<https://aps.ntut.edu.tw/AAOffice/syllabus/113>

(二)學生參與服務研究之產出成果

- 吳若綺:以區塊鏈架構建立區域市場可流通的代幣交易制度
- 丁俞豪:金融機構之跨平台操作系統及其方法
- 沈容:有效運用區塊鏈技術於健康醫療照護之研究

(三)企業需求達成情形及預估提升之效益

1. 需求達成情形

金融系統需求：完成代幣交易制度導入黑快馬消費金融與供應鏈系統，並加入 加密簽章與風險控管，自動化處理貸款、分期、還款及兌換。

通路與支付需求：已成功與 APP、POS、Portal 及多元支付平台串接，並通過 跨平台 API 資安測試，確保支付資料不洩漏。

資訊安全需求：建置 異常交易偵測、智能合約安全審查 與 身份驗證機制，有效防止重放攻擊、雙重支付與未授權存取，合作機構對成果非常滿意。

2. 提升效益

技術效益：強化金融與支付系統的 安全性、透明性與合規性，建立可落地的「代幣＋智能合約＋資安防護」支付模組。

營運效益：透過 安全可信的支付環境 提升用戶信任度，增加黑快馬產品差異化與市場競爭力，並吸引更多金融機構與商家合作。

產業效益：提供 具資安保障的區塊鏈代幣制度示範案例，推動跨區域「代幣聯盟鏈」，並促進學研成果技術移轉與產業升級。

(四)其他事項檢討說明(無則免填)

產業

學院

貳、實際指標達成情形

一、實際達成指標一覽表

指標項目	預期達成指標		實際達成指標	
	量化指標說明	量化指標	量化指標說明	量化指標
1. 期刊論文				
SCI	發表相關國外學術期刊	國內 0 篇 國外 1 篇	Yu-Chi Wei, Yu-Chun Chang, Wei-Chen Wu. Multi-language IoT Information Security Standard Item Matching based on Deep Learning, Computer Science and Information Systems, Computer Science and Information Systems 21(2):663 – 683, DOI:10.2298/CSIS230822012W	國內 0 篇 國外 1 篇
SSCI		國內 0 篇 國外 0 篇		國內 0 篇 國外 0 篇
TSSCI		國內 0 篇 國外 0 篇		國內 0 篇 國外 0 篇

2. 研討會論文	發表於國際學術研討會	國內 0 篇 國外 1 篇	Zhu-Xuan Huang, Yu-Chih Wei, Wei-Chen Wu (MLFSP 2025), Reversible Data Hiding in DICOM Images with Regional Sensitivity Analysis and Layered Protection, The 2025 International Conference on Machine Learning on FinTech, Security and Privacy (MLFSP 2024), Osaka, Japan.	國內 0 篇 國外 1 篇
3. 研究著作		國內 0 本 國外 0 本		國內 0 本 國外 0 本
4. 專書論文		國內 0 章 國外 0 章		國內 0 章 國外 0 章
5. 技術報告		國內 0 篇 0 篇		國內 0 篇 0 篇
6. 產學合作成果報告		0 篇		0 篇
7. 學生實務專題報告(必填，每位學生至少1篇)	3篇區塊鏈專題報告	3 篇	吳若綺：以區塊鏈架構建立區域市場可流通的代幣交易制度 丁俞豪：發明專利 https://tiponet.tipo.gov.tw/S092_OUT/2022 輸入I856852 沈容：有效運用區塊鏈技術於健康醫療照護之研究	3 篇
8. 編製實務教材		0 件		0 件
9. 開授相關實務課程	開授區塊鏈課程	1 門	1. 保險科技與資安防護 https://aps.ntut.edu.tw/AA0office/syllabus/113金融資安產碩專秋季班第三學期/1225 1.pdf	1 門

10. 產學合作 簽約		0 件		0 件
11. 產學合作 金額		0 萬元		0 萬元
12. 技術移轉		0 件		0 件
13. 專利佈局	區塊鏈專利一項	1 項	金融機構之跨平台操作系統及其方法：證書號I856852（查詢 https://tiponet.tipo.gov.tw/S092_OUT/2022 輸入I856852）	1 項
14. 商品化		0 項		0 項
15. 協助員工 教育訓練		0 人次		0 人次
16. 學生短期 實務實習		0 人時		0 人時
17. 學生實作 作品		0 件		0 件
18. 畢業後獲 合作機構聘 用				
副學士		0 人		0 人
學士		0 人		0 人
碩士		0 人		0 人

19. 學校認列本方案符合本校內教師進行產業研習或研究之規定	本方案符合本校教師進行產業研習或研究作業要點之產業實地服務，如教師依該要點完備文件及程序，即可認列。	是	依「技專校院教師進行產業研習或研究實施辦法」規定，教師可被認列下列情形。結合本計畫與黑快馬的合作情境，認列情況如下： 1. 實地服務或研究：教師與師生團隊至黑快馬進行研究工作，包含系統設計、原型開發與測試、智能合約與代幣系統的實務導入與驗證，這些工作在黑快馬或其提供之辦公／研發環境中進行。 2. 產學合作計畫：本計畫若與黑快馬簽訂合作契約，明定技術移轉、系統商品化目標或成果、業界應用場域效益，例如將設計出的代幣交易制度或支付整合系統導入黑快馬的產品或對黑快馬客戶供應鏈／通路生態系中應用，則可認列為產學合作計畫，具有具體貢獻與對產業發展之成效。 3. 深度實務研習：教師及學生若在黑快馬舉行工作坊、在黑快馬的團隊中實作、參訪其研發流程並進行 hands-on 分析與驗證（例如雲端架構、APP／POS 系統與支付系統安全性測試等），並與黑快馬共同規劃與執行這些研習活動。	是
20. 赴合作機構實地服務或研究				

計畫主持人	進行實地服務或研究 至少6個月	是	1. 合作場域與角色定位 本計畫與黑快馬股份有限公司合作，由黑快馬提供研發／雲端系統與軟體 OEM 貼牌服務之技術支援與場域。黑快馬作為業界合作機構，提供其研發中心、雲端平台、開發團隊與 OEM 貼牌軟體系統作為測試與落地場景。 2. 服務形式與任務內容 在黑快馬的協助之下，教師與師生團隊將於其技術研發部門或專案環境中全程參與： *系統需求與功能規格分析：與黑快馬技術團隊共同評估「區塊鏈代幣交易制度」如何整合到黑快馬現有的 SaaS/E-commerce/OEM 軟體架構中。 *原型建置與測試：在黑快馬的雲端、APP 或網頁系統環境部署私有鏈 + 智能合約架構，並執行代幣發行、折抵、商家分級與現金兌換等流程的測試。 *網路安全與風險管理： #建立交易加密與簽章機制（如公私鑰管理、Hash 驗證）以確保代幣交易不可竄改。 #測試黑快馬系統在面對駭客攻擊、重放攻擊、雙重支付等威脅下的防護能力。 #評估 API 串接過程的安全性，確保代幣交易與支付模組在跨系統資料交換中不會洩漏敏感資訊。 #設計異常交易偵測（Fraud Detection）機制，利用黑快馬的大數	是
-------	--------------------	---	--	---

據與風控模型，監測不正常的代幣使用行為，降低金融犯罪風險。

*實務驗證與流程整合：利用黑快馬現有客戶或案例（例如其通路商、電商平台等客戶群）進行代幣使用／折抵的真實場景模擬，以檢驗交易速度、安全性、API 接口、支付流程與用戶體驗。

*安全、風控與監管設計：在黑快馬環境中設計與測試風險控管機制，如智能合約漏洞防範、使用者身份驗證／授權、資料隱私與交易記錄可追蹤性。

3. 實施時間與地點

*時間安排：例如計畫第一階段（3個月）進行需求分析與原型設計，“中期階段”進行系統開發與初步測試（2個月），“後期階段”於黑快馬場域或客戶場景中驗證與調整（2個月）。

*地點：以黑快馬公司本部／研發中心、其雲端平台／伺服器環境為主；若黑快馬設有實體工作空間或客戶端應用場景，也可於該處進行現場觀察與測試。

4. 成果與效益

*技術移轉與商品化可能：若原型在黑快馬平台／OEM 軟體產品中整合成功，有機會成為黑快馬正式提供的服務或模組。

*業界應用與改善：黑快馬能在其產品或客戶服務中加強代幣支付、會員獎勵／折抵之功能，提升其 SaaS／OEM 產品線的競爭性。

*學術／教學貢獻

			：師生可從中取得實作經驗，撰寫論文、申請專利；亦可將案例與實務經驗納入課程或教學範本。 *對產業發展之具體貢獻：示範校企合作如何推動區域代幣制度與智能合約應用，促成技術創新與系統化支付／激勵制度之落地。	
沈容	進行實地服務或研究至少6個月	是	工作分配：消費金融系統研究、風控模型設計，偏重在金融交易流程的安全測試、風控與異常偵測。 1. 學習並分析黑快馬的消費金融系統（線上貸款、分期、還款流程），研究如何導入代幣交易。 2. 建構代幣使用與風控模型的整合，找出降低交易風險的策略。 3. 資安任務： *測試代幣折抵與還款流程是否有異常交易風險（如雙重支付）。 *與黑快馬的大數據模型結合，模擬不正常交易行為，設計異常偵測機制。 *分析代幣交易記錄的透明性與追蹤性，提出如何兼顧隱私與合規。	是

吳若綺	進行實地服務或研究 至少6個月	是	工作分配：系統與平台整合、通路及支付系統整合、API 串接、平台安全測試。 1. 在黑快馬的研發環境中，負責部署代幣交易制度於雲端平台、APP、POS 系統與電商 Portal。 2. 協助 API 串接與資料流程設計，確保代幣與支付模組間的正确交互。 3. 資安任務： #測試 API 串接與跨平台交易過程的安全性，確保資料不會洩漏。 #驗證行動支付模組是否能防止重放攻擊與未授權操作。 #設計基本的使用者認證與授權流程（如多因子驗證），提升交易安全。	是
丁俞豪	進行實地服務或研究 至少6個月	是	工作分配：供應鏈金融與資產管理研究，偏重在供應鏈金融與資產管理場景的智能合約安全測試與合規檢驗。 1. 學習黑快馬的供應鏈融資與應收帳款買賣系統，研究如何利用智能合約強化資產透明與分配管理。 2. 協助設計代幣與現金兌換（Q 值比率）的應用場景，測試跨系統資金流通的可行性。 3. 資安任務： *驗證智能合約在兌換過程中是否存在漏洞（如合約被惡意利用）。 *測試自動扣款與還款流程的安全性，避免資料竄改或未經授權存取。 *設計合規性測試，確保代幣兌換機制符合金融監管需求。	是

二、成果報告摘要

(一)計畫主持人成果摘要

成果類別：期刊論文	頁數：20
題目：Multi-language IoT information security standard item matching based on deep learning	
作者：Yu-Chi Wei, Yu-Chun Chang, and Wei-Chen Wu	
關鍵字：Information Security, Information Security Standards, IoT Security, Text mining, Deep Learning	
摘要：In the realm of IoT information security and other domains, various information security standards exist, such as the IEC 62443 series standards published by the International Electrotechnical Commission and ISO/IEC 27001 by the International Organization for Standardization. Business organizations are striving to improve and protect their operations through the implementation and study of these information security standards. However, comparing or pinpointing applicable control measures is becoming increasingly labor-intensive and prone to errors or deviations, especially given the plethora of information standards available. Identifying specific control measures scattered across different information security standards is gradually becoming an important issue. In this research, we utilise a range of domestic and international information security standards as the foundation, employing text mining and deep learning methods to map the similar parts of control measures between standards, thereby enhancing the efficiency of comparison tasks and allowing human resources to be allocated to more pertinent issues.	

(二)參與學生成果摘要

學生實務專題報告-1	頁數：7
題目：以區塊鏈架構建立區域市場可流通的代幣交易制度	
作者：吳若綺	指導教授：吳威震
關鍵字：區塊鏈，智能合約，可流通代幣，金融科技	

摘要：為了解決跨區域市場可流通代幣的交易方式，亦不受相關金融監理機構法規所限制，本論文將提出一個以區塊鏈和智能合約架構下，建立一個安全公平及可信任的跨區域市場可流通的代幣交易制度，其特性為安全、公平及可信任、跨區域市場的代幣交易制度。

學生實務專題報告-2		頁數：5
題目：金融機構之跨平台操作系統及其方法		
作者：丁俞豪		指導教授：吳威震
關鍵字：區塊鏈、身份驗證、金融科技		
摘要：本專利針對現今存在的身分偽造、盜領問題，設計出一項能確保個資安全，同時得以節省開戶時間的系統，欲提供不失保障又便利的開戶體驗。藉由此系統將顧客資料上傳到區塊鏈，便能在往後的開戶流程中，減少相同的徵信手續，以及仰賴他行帳戶做為驗證的流程。以解決民眾對於現今開戶的等待時間過長的問題，並能防範有心人士利用客戶個資以挪取財產之惡行。最終達成資訊安全的防護縱深發展，也加強金融數位發展環境之韌性。此外，我們將整個系統套用至三個框架進行運用，分別是一般銀行開戶、申辦信用卡以及證券戶之開戶。希望藉由各銀行間的資訊共享和區塊鏈去中心化、透明性、不可竄改性的三大特色，將客戶資料在安全性良好的前提下，去除繁瑣步驟以活用至各情境。省下雙方時間同時保障資安，共構雙贏局面。		

學生實務專題報告-3		頁數：17
題目：有效運用區塊鏈技術於健康醫療照護之研究		
作者：沈容		指導教授：吳威震
關鍵字：電子病歷、區塊鏈		

摘要：隨著資訊科技的進步以及資料量的急速增加，將大數據運用於健康照護系統的研究，是近年來各國積極發展的策略，醫院將原本記錄在紙本病歷上的病歷紀錄，透過電腦系統與網路技術，存放於醫院中的病歷資料庫中，透過電子病歷系統與臨床資料庫系統，醫師可以更快的取得病患的病歷資料，而醫院也可以經由系統電子化減少許多調閱病歷所需要的時間和人力。本研究將提出一個HealthCoin的交換點數，提供醫院、醫生與病人彼此交換的點數，以利後續做交換紀錄的憑證。本研究提出的健康區塊鏈架構只要有三種成員，每個成員在健康區塊鏈架構中都是一個節點，然而病歷涉及病人個人隱私，醫療院所不僅要具備電子病歷交換機制，還必須更加重視資通安全與個人資料保護。推動電子病歷的權責單位事出多頭，在缺乏資源配置的管理及協調機制下，欠缺橫向協調整合，不但導致計畫推動容易出現本位主義而各行其是，難以呈現分工合作之綜效。也由於缺乏執行統合機制，即使是目前較有成效的醫院部分，也出現執行進度差異頗巨的問題。推動電子病歷交換，攸關智慧醫療應用的推動，因此本計畫將以健康區塊鏈技術解決第三方公正的問題，藉由HealthCoin的交易點數彼此對健康存摺作交易。

三、預期達成指標與實際達成情形差異說明(無不符者免填)

學院

產業

參、經費執行成效

一、經費執行總表

單位：新臺幣元

項目	計畫核定		學校執行情形	
	核定計畫金額 (補助款+自籌款)	核定補助金額 (補助款)	實際支用金額 (補助款+自籌款)	結餘款 (補助款+自籌款)
人事費	0	0	0	0
業務費	210000	210000	210000	0
設備及投資	21000	0	21000	0
合計	231000	210000	231000	0
整體執行率 (%)	100%	應繳回結餘款	0	

二、整體經費支用成效說明

(一)經費重點項目支用成效

1. 報名費123572:為配合合作機構須做相關產品的資安檢測，安排三位學生與計劃主持人參與資安相關教育訓練，並考取相關證照，成果如附件。
2. 印刷費34104:用於相關活動海報印刷，與網路資源授權電子書之列印。
3. 雜支33030:與本計劃相關之耗材。

(二)經費執行差異說明(無則免填)

無差異

肆、計畫辦理情形整體檢討說明

1. 行政作業

本計畫在行政作業推動上，整體進展順利。計畫主持人與學校行政單位能及時完成相關流程，包括產學合作契約簽署、經費核銷、會議召開及成果送審。行政支援效率高，能有效協助師生與合作機構進行溝通，減少了跨單位協調上的困難。建議未來可建立更標準化的行政作業指引，縮短文件往返與核可時間。

2. 課程開設

課程設計以「區塊鏈代幣制度與金融科技應用」為核心，融入產業研習的實務內容，讓學生在修課過程中即可同步參與合作機構的專案。課程中安排了區塊鏈基礎、智能合約實作、金融風控應用與資訊安全測試等模組，學生回饋學習成效佳，且能將課堂所學直接應用於實地研究。未來可增加更多跨領域模組（如 AI + FinTech），進一步擴大應用層面。

3. 合作機構配合

合作機構黑快馬股份有限公司全程高度配合，除提供研發環境與實務案例外，亦安排技術人員與業務主管定期與師生進行討論。特別是在系統安全測試與智能合約應用場景設計方面，企業專業意見大幅提升了研究的可行性與實用性。企業表示本計畫成果符合其需求，且具有技術移轉與商業化價值。未來可進一步建立長期合作關係，持續深化代幣制度於不同業務的應用。

4. 學生輔導

參與的三位學生能依其專業背景（資管、財金）承擔不同的研究工作，並在教師指導下完成系統分析、金融模型研究及資安測試等任務。學生不僅提升了專業知識，更培養了跨領域合作與解決實務問題的能力。輔導過程中，教師每週安排固定會議，檢討進度並提供技術與研究指導，學生最終皆能產出專題報告與實作成果，達成研習目標。

結論

整體而言，本計畫在行政支援、課程開設、產學合作與學生培育等方面均達到預期目標。計畫不僅成功完成代幣制度與資訊安全應用的實作與測試，也促成與黑快馬的產學合作成果，學生學習與研究表現亦達標。未來若能延伸計畫成果至更多應用領域，將可進一步擴大產學合作效益與教育價值。

學院

產業

伍、附件

學院

產業

計劃主持人部分

類別	名稱	名稱
期刊論文	SCI 期刊	MYu-Chi Wei, Yu-Chun Chang, Wei-Chen Wu. Multi-language IoT Information Security Standard Item Matching based on Deep Learning, Computer Science and Information Systems, Computer Science and Information Systems 21(2):663–683, DOI:10.2298/CSIS230822012W
研討會論文	研討會論文	Zhu-Xuan Huang, Yu-Chih Wei, Wei-Chen Wu (MLFSP 2025), Reversible Data Hiding in DICOM Images with Regional Sensitivity Analysis and Layered Protection, The 2025 International Conference on Machine Learning on FinTech, Security and Privacy (MLFSP 2024), Osaka, Japan.
專利	發明專利	金融機構之跨平台操作系統及其方法
開授課程	保險科技與資安防護	https://aps.ntut.edu.tw/AAOffice/syllabus/113

學生部分

學生姓名	專題報告	備註
吳若綺	以區塊鏈架構建立區域市場可流通的代幣交易制度	
丁俞豪	金融機構之跨平台操作系統及其方法	此報告已申請專利不公開
沈容	有效運用區塊鏈技術於健康醫療照護之研究	

Multi-language IoT Information Security Standard Item Matching based on Deep Learning *

Yu-Chi Wei¹, Yu-Chun Chang², and Wei-Chen Wu³

¹ National Taipei University of Technology
Taipei, Taiwan
vickrey@mail.ntut.edu.tw

² National Taipei University of Technology
Taipei, Taiwan
t109ab8013@ntut.org.tw

³ Department of Finance, National Taipei University of Business
Taipei, Taiwan
weichen@ntub.edu.tw

Abstract. In the realm of IoT information security and other domains, various information security standards exist, such as the IEC 62443 series standards published by the International Electrotechnical Commission and ISO/IEC 27001 by the International Organization for Standardization. Business organizations are striving to improve and protect their operations through the implementation and study of these information security standards. However, comparing or pinpointing applicable control measures is becoming increasingly labor-intensive and prone to errors or deviations, especially given the plethora of information standards available. Identifying specific control measures scattered across different information security standards is gradually becoming an important issue. In this research, we utilise a range of domestic and international information security standards as the foundation, employing text mining and deep learning methods to map the similar parts of control measures between standards, thereby enhancing the efficiency of comparison tasks and allowing human resources to be allocated to more pertinent issues.

Keywords: Information Security, Information Security Standards, IoT Security, Text mining, Deep Learning.

1. Introduction

With the proliferation of Internet of Things (IoT) technologies, everyday life has become increasingly digitized. IoT devices have a wide range of practical applications, whether in office environments, transportation, financial transactions, healthcare, or even in standard household smart appliances [22]. Broadly speaking, any device that can connect to the internet falls into this category, from those with basic network functionality to those combining various sensor devices, specialized software, or even capable of receiving and transmitting data from other complex IoT devices. The advent of IoT and the digital economy is a double-edged sword, on one hand making our lives more convenient to some extent, but on the other hand, escalating the information security threats associated with IoT devices and applications.

* An extended version of The 12th Frontier Computing Conference/FC2022 paper

In the current era where the Internet of Things (IoT) is burgeoning and the inter-connection of all things is becoming a trend, the potential risks behind its applications warrant our deep reflection and assessment. Revising new standards is a time-consuming and labour-intensive project, requiring information security professionals to reference, organise, and summarise the contents of various different standards. In this research, based on textual data exploration, existing international IoT standards are automatically pre-processed into numerous features, and then trained using deep learning models. This enables the automatic analysis of existing standards' information security requirements and their alignment with those of other international IoT information security standards. Finally, members of the standards drafting unit can directly refer to and assess whether the automatically generated corresponding results are suitable for use, thus saving a substantial amount of labour and time costs.

This research aims to utilise text mining to automatically translate and reference various existing international IoT standards. After textual preprocessing, these standards are trained using machine learning and deep learning models. The objective is to segment and automatically analyse the information security requirements of the existing standards, matching them with the requirements listed in other international IoT security standards. This not only assists the Mobile Application Security Alliance in continually updating IoT security verification standards, but also allows for the practical examination of whether domestic information security standard-setting processes comply with international IoT security standards. This study uses both domestic and international information security standard content as its dataset, with the capability to swiftly identify similar content. Furthermore, the content is not limited to being in the same language, and the overall output process can be finely tuned based on the input dataset to achieve the best matching results. This application is not limited to comparing and analysing the content of information security standards alone. It can also be based on other existing data and literature to explore and analyse their similarities, providing a reference for researchers looking to implement text processing, text analysis, machine learning, deep learning, and information security standards in their workflow.

2. Related Works

2.1. IEC 62443 Standards

IEC 62443 is a series of international standards for Industrial communication networks - IT security for networks and systems, which contains a series of technical procedures for the security of control systems, and the standard classified the user roles into operator, integrator and manufacturer, designs risks and potential problems for each role to help users of the standard to design and evaluate their own industrial automation systems and improve network security. The IEC 62443 series of standards is divided into four parts. The first part includes terminology and explanations of concepts related to automated industrial control systems, as well as examples of their use; the second part describes the security planning, operation, and management of the structure of industrial automation and control systems; the third part details technologies related to information security, information security risk assessments, and other definitions concerning information security; the fourth part focuses on the description of various security requirements, including

the product security development lifecycle, components, and technologies. In the process of developing the Mobile Application Security Alliance IoT security certification series, IEC 62443 Part 4-2 [8] (IEC 62443-4-2) is also one of the key reference standard and will be introduced separately in subsequent sections. The table below, Table 1, shows the structure and orientation of each content of the "IEC 62443" series of standards.

Table 1. The list of IEC 62443 series of standards

Standard structure	Parts	Standard content
General: Defines the standard concepts, models, terminology interpretation and examples, etc.	1-1 TS	Concepts and models
	1-2 TR	Master glossary of terms and abbreviations
	1-3	System security compliance metrics
	1-4	IACS security life cycle and use-cases
Policies and Procedures: Provide defined system management requirements for IACS asset owners and services.	2-1 TS	Secure program requirements for IACS asset owners
	2-2	Security Protection Rating
	2-3 TR	Patch management in the IACS environment
	2-4 IS	Requirements for IACS service providers
	2-5 TR	Implementation guidance for IACS asset owners
System: Security risk assessment and security requirements defined for industrial control systems.	3-1	Security technologies for IACS
	3-2	Security risk assessment and system design
	3-3	System security requirements and security levels
Component: The safety product development process and component safety requirements as defined by the product supplier.	4-1 IS	Secure product development lifecycle requirements
	4-2	Technical security requirements for IACS components

2.2. OWASP Top 10

OWASP, known as the Open Web Application Security Project, is an open, non-profit organization dedicated to helping governments and businesses improve web software security, tools, and technical documentation, as well as gain practical insight into the vulnerabilities and security of the information assets they use. Every few years, OWASP produces a list of the top 10 web application security vulnerabilities and provides some easy ways and directions to educate users on how to avoid these vulnerabilities. Table 2. below shows the ten web application security vulnerabilities pro-posed in "OWASP Top 10:2021 [18]".

Despite all the vulnerabilities presented in the OWASP Top 10 are carefully organized and filtered to the top ten most common web application security vulnerabilities of our time, there is still a ranking hierarchy among the vulnerabilities, and the higher the ranking, the more important the web application security vulnerability is in the current information environment.

Among the existing web application security vulnerabilities, there are several items that have appeared in the previous version of the "OWASP Top 10", but their ranking has been changed in response to the changing times and environment. For example, A01:

Table 2. The list of OWASP Top 10:2021

Vulnerabilities ID	Vulnerabilities Top 10 of application security ver.2021
A01:2021	Broken Access Control
A02:2021	Cryptographic Failures
A03:2021	Injection
A04:2021	Insecure Design
A05:2021	Security Misconfiguration
A06:2021	Vulnerable and Outdated Components
A07:2021	Identification and Authentication Failures
A08:2021	Software and Data Integrity Failures
A09:2021	Security Logging and Monitoring Failures
A10:2021	Server-Side Request Forgery

Access Control Failure in "OWASP Top 10:2021", which was ranked fifth in the previous version of OWASP Top 10:2017, was moved from fifth to first in the latest version. According to the officials, more than 90% of the applications they tested had a category access failure problem, and the number of occurrences was much higher than other vulnerability categories.

In addition to the ten most common security weaknesses of web applications, OWASP also has responded to the increasing use of APIs and Internet of Things devices in the industry, they presented the "OWASP API Security Top 10" and "OWASP IoT Top 10", which includes ten most common security vulnerabilities of network applications. Despite there is a newer version of "OWASP IoT Top 10", which is the version 2018, but overall and detailed information of "OWASP IoT Top 10:2014 [16]" is relatively more abundant than the 2018 version on the official OWASP website, more information is definitely more helpful for deep learning model to classify information security controls into similar categories, that was the main reason we chose to use "OWASP IoT Top 10:2014 [16]" instead of "OWASP IoT Top 10:2018 [17]". Table 3 below shows the list of top 10 security vulnerabilities of "OWASP IoT Top 10:2014".

Table 3. Introduction to the OWASP IoT Top 10:2014

Vulnerabilities ID	Vulnerabilities Top 10 of Internet of Things ver.2014
I01:2014	Insecure Web Interface
I02:2014	Insufficient Authentication/Authorization
I03:2014	Insecure Network Services
I04:2014	Lack of Transport Encryption
I05:2014	Privacy Concerns
I06:2014	Insecure Cloud Interface
I07:2014	Insecure Mobile Interface
I08:2014	Insufficient Security Configurability
I09:2014	Insecure Software/Firmware
I10:2014	Poor Physical Security

For the showcase of this study, we attempted to make IEC 62443-4-2 controls automatically classified into the closest of the ten specified "OWASP IoT Top 10" categories through text mining and deep learning methods, thus saving the time and cost required for manual comparison of information security standards.

2.3. Text Similarity Matching

Text similarity matching methods are becoming increasingly important in many applications. Existing methods often compute similarity based on shallow syntax or POS tagging or by comparing basic syntax similarity, generating vectors, and then inferring similarity from this set of vectors. However, due to the variability in natural language expression, these methods often struggle to predict actual semantic content and implications. To address these issues, researchers have attempted various approaches. At the beginning, using a lexicon to note the positions of words within sentences, forming a one-hot encoding vector representation. However, this method couldn't link related words. Mikolov et al. [23] tried combining neural networks in their research. Mikolov et al. [23] tried using pre-trained word representations in conjunction with supervised learning methods as extra features, which showed significant improvements over traditional word embedding methods. These methods evolved into larger frameworks, like sentence embedding [10] or paragraph embedding [11]. Matthew et al. [19] extracted context-sensitive features from language models, integrated these features into training for specific tasks, and gradually began to understand the variability and actual semantic content in language expressions.

Researchers also considered the context and situation of sentences. The Skip-gram model [5] is a renowned method that trains and identifies using the context of target words. WordNet [4] is a network primarily focused on the "word-semantics" in English, storing the structures and potential relationships between words, quantifying the semantic relationship between two different words. ConceptNet [12] uses a dictionary-based embedding model, aligning with the hierarchical structure of predefined words in WordNet, defining various relationships between words. Emrah [7] proposed a method focused on calculating sentence similarity without using machine learning, relying on dependency parsers and lexical embedding models, achieving results better than most traditional methods.

In research on machine learning for text similarity analysis and comparison, Ji and Eisenstein [9] introduced a supervised machine learning method that measures semantic similarity between sentences using a discriminative term or proper noun, in conjunction with a set weighting index, giving higher importance to certain features, then computing sentence similarity. The authors claimed their new method outperforms the widely-used TF-IDF weighting method. Mohamed and Oussalah [15] proposed a similarity calculation method that uses WordNet to obtain dependency relationships for words which based on instances extracted from Wikipedia and normalized Google distance. The normalized Google distance calculates the hit count returned for a set of keywords using the Google search engine. Hassan [6] proposed a method based on Wikipedia's content for context determination, called Salient Semantic Analysis (SSA). Mihalcea et al. [13] combined corpus-based semantic similarity with knowledge-based semantic similarity, using data from WordNet and the British National Corpus, which reduced the error rate compared to traditional methods. Wang et al. [24] focused on the similar and dissimilar parts of

sentences. They constructed a similarity matrix and corresponding vectors for each word meaning, decomposing the resulting matching vectors to identify similar and dissimilar parts, eventually using matrix decomposition to extract sentence vectors to compute sentence similarity.

BERT [2], introduced by Google's AI team in 2018, used BooksCorpus and over 800 million entries and data from Wikipedia for pre-training. The operation is divided into two stages: pre-training and fine-tuning. In the pre-training phase, there are two training methods: Masked LM and Next Sentence Prediction. Then in the fine-tuning phase, the model is adjusted based on specific tasks. BERT performs well in sentence classification, tagging, and text classification. However, Reimers and Gurevych [20] found that while BERT and RoBERTa achieve effects in many sentence regression tasks, such as text semantic similarity, they need to input two sentences to be compared into the model repeatedly until the closest two sentences are found. The excessive computational cost makes BERT unsuitable for semantic similarity searches. Hence, they introduced SBERT (Sentence-BERT). SBERT, unlike BERT, which repeatedly attempts to combine two sentences, calculates the similarity distance between two sentences directly by matching their word embedding representations, significantly reducing computation. This model also achieves good results in some STS and transfer learning tasks.

3. Research Methodology

3.1. Data Pre-processing

In this study, we use python and jupyter notebook as the test environment. In the data pre-processing progress, we first need to retrieve the contents of the information security standard, and split the contents into each column, including its control number, control name and control description as a spreadsheet. After this, the contents of the information security standard form are stored in memory and ready to go.

These manually retrieved control contents in the spreadsheet does not require data pre-processing, they can simply import into the deep learning model in their original format for training. The pre-trained models provided by SBERT [21] are already trained from various types of datasets, familiar with the original word patterns, so there is no need to perform steps such as words and sentences segmentation, word lemmatization, stemming or other data pre-process methods you can find in other NLTK tasks to filter the features.

As shown in the figure above, the following figure is a screenshot of the jupyter notebook after importing and reading the information security standard content into the spreadsheet. This experiment uses IEC 62443-4-2 [8] content as the training set, and tries to classify the content of the controls in each of the ten categories of "OWASP IoT Top 10:2014 [16]" as the test set.

It is also possible to match similar contents between different language information security standards. In the data pre-processing state, while keeping the unique control id number field legible, translation modules can be used to translate control descriptions into the specified language and then perform a similarity comparison exercise with other information security standards. Usually, it is better to translate other languages into English and perform similarity matching between the two standards using English as the common language, because most of the pre-training data for deep learning models are trained from English data as shown in the Fig. 2 below.

```

In [17]: train_blist[0]
Out[17]: 'Components shall provide the capability to identify and authenticate all human
users according to IEC 62443-3-3 SR 1.1 on all interfaces capable of human user
access. This capability shall enforce such identification and authentication on
all interfaces that provide human user access to the component to support segre
gation of duties and least privilege in accordance with applicable security pol
icies and procedures. This capability may be provided locally by the component
or by integration into a system level identification and authentication syste
m.'

In [18]: test_blist[0]
Out[18]: 'Default passwords and ideally default usernames to be changed during initial s
etup'

```

Fig. 1. A schematic diagram of train/test data content

	影像監控系統標準_zh	影像監控系統標準_translated
0	產品預設不應透過實體介面存取產品作業系統之除錯模式。若需經實體介面存取，則應透過身分鑑別作...	Product presets should not be removed by the p...
1	產品應具有實體埠插拔操作記錄功能。	The product should have the function of the ph...
2	產品應具備相關警示功能於實體操作發生斷訊時。	The product should have the relevant warning f...
3	產品外部不應有徒手即可還原預設通行碼的功能。	The outside of the product should not have the...
4	產品應支援安全啟動(Secure Boot)功能，不應以未經授權的韌體、驅動程式及作業系統執...	The product should support the Security Boot f...
5	產品之作業系統與網路服務，不應存在美國國家弱點資料庫所公開的常見弱點與漏洞資料，且漏洞評鑑系...	The operating system and network services of t...
6	產品之作業系統與網路服務，不應存在美國國家弱點資料庫所公開的常見弱點與漏洞資料，且漏洞評鑑系...	The operating system and network services of t...
7	產品開啟之網路服務應為廠商提供必要服務之所需，防止產品因啟用網路介面而被侵入的可能性，且廠商...	The online service of the product opening shou...
8	產品所收集之遙測資料應告知使用者，且未告知之遙測資料不應被收集。	The remote measurement data collected by the p...
9	韌體應具備更新機制。	The ligament should have a update mechanism.
10	產品若支援離線手動更新，則更新檔案應加密保護以確保機密性，且應採用NIST SP 800-1...	If the product supports offline manual update,...
11	產品若支援線上更新，其更新路徑應通過安全通道，且安全通道版本應符合「附錄 A」的要求，同時金...	If the product supports online update, its upd...
12	產品應具備驗證韌體之完整性及真實性的功能。	The product should have the function of verify...
13	產品應具備備援更新功能，即發生更新失敗時，系統能回復至更新前之狀態。	The product should have the renewal function, ...
14	產品所儲存的敏感性資料，應僅由獲授權個體存取。	The sensitive data stored in the product shoul...
15	產品所儲存之身分鑑別因子、加解密用之金鑰(不含非對稱加密用之公鑰)不應明文儲存，而保護資料的...	The identity of the product's identity identit...
16	產品應提出金鑰管理程序，以確保金鑰管理的品質。	The product should propose the golden key mana...

Fig. 2. A schematic diagram illustrating the successful prediction of control items between two different standards

In the Fig. 2 above, we used a local IoT security standard for this showcase, which is “IoT-1001-1 v2.0 Image Monitor System Information Security Standard - Part 1: Information Security Requirements [14]” from Mobile Application Security Alliance, which is an IoT product certification alliance dedicated to the promotion of domestic IoT information security in Taiwan. According to the figure, any standards in different languages can be translated into English by the translation module and then start the comparison process of information security standards directly, this allows the process to be able to compile information security standards in different languages without any limitation due to language.

3.2. Model Training

BERT [1], the abbreviation of Pre-training of Deep Bidirectional Transformers for Language Understanding, which is already highly characterized by the endless training data based on the Google search engine, so that BERT only needs to specify the form of its output data, and then fine-tune it according to the task, finally, it can be used for various common natural language processing tasks.

But Reimers and Gurevych found that although both BERT and RoBERTa achieve some good results in many sentences regression tasks, such as textual semantic similarity, they both need to pass both sentences to be compared into the model and repeat this process until the two most similar sentences are found. This is a very costly process, especially when the data is large. According to this, they considered BERT is not suitable for the task of semantic similarity search because of the limitation of the algorithm, so they proposed SBERT [21] (Sentence-BERT), which does not need to try to combine two sentences repeatedly like BERT, but by directly matching and calculating the words similarity distance of two sentences using word's embedding representations, which greatly reduces the computational effort and achieves very good results in some STS and migration learning tasks.

The deep learning method SBERT provides a number of pre-training models, which allow users to train their own research data directly to make further predictions. In the official guidance document of SBERT, 13 pre-training models are provided. The 13 pre-training models are listed with their performance of sentence embeddings, performance of semantic search, average overall performance, running speed and model size, so that users can select them according to their task requirements. The five models with the best performance based on the above five indicators were shown as Table 4.

In this study, the best average overall performance one: *all-mpnet-base-v2*, were selected, which was an all-round model tuned for many use-cases, trained on a large and diverse dataset of over 1 billion training pairs.

As shown in the figure above, IEC 62443-4-2 controls were successfully classified by deep learning models into the ten categories of OWASP IoT Top 10:2014. Matching similar contents including controls or descriptions between several information security standards, which often requires a lot of labor and time, but this study showed that it is totally possible to quickly generate similarity comparison results between certain information security standards by using text mining and deep learning methods. It can also be said that this study, corresponding contents between information security standards and standards is also one of the typical NLP tasks, i.e., the application of semantic textual similarity tasks.

Table 4. Comparison of SBERT best performance pre-trained models

Model name	Performance of sentence embeddings	Performance of semantic search	Average overall performance	Encoding speed	Model size
all-mpnet-base-v2	69.57	57.02	63.30	2800	420 MB
multi-qa-mpnet-base-dot-v1	66.76	57.60	62.18	2800	420 MB
distiluse-base-multilingual-cased-v2	60.18	27.35	43.77	4000	480 MB
paraphrase-MiniLM-L3-v2	62.29	39.19	50.74	19000	61 MB
paraphrase-multilingual-mpnet-base-v2	65.83	41.68	53.75	2500	970 MB

	IEC 62443 Controls	OWASP Top 10:2014 Category	Distance
0	5.3 CR 1.1	I5	0.515
1	5.4 CR 1.2	I5	0.519
2	5.5 CR 1.3	I5	0.385
3	5.6 CR 1.4	I5	0.416
4	5.7 CR 1.5	I2	0.469
...	...	...	...
83	15.9 NDR 3.12	I5	0.575
84	15.10 NDR 3.13	I5	0.580
85	15.11 NDR 3.14	I10	0.528
86	15.12 NDR 5.2	I3	0.455
87	15.13 NDR 5.3	I3	0.530

Fig. 3. A comparative schematic illustrating the distance between control measures across different standards

3.3. Evaluation Methodology

In section 3.2, we have demonstrated that it is possible to perform similarity comparisons between information security standards using deep learning methods. But, how was the predictive accuracy? To find out the predictive accuracy of the model, first, a reference answer that cross-validates the model prediction results is necessary. For example, a table which providing an official mapping of the controls of a standard itself to the controls of another standard, such as a table which maps IEC 62443-4-2 [8] controls to EN 303-645 [3] controls. But unfortunately, no such mapping table is provided in the official documents of these two parties.

For this reason, we use the official mapping of Appendix D of IoT-1001-1 v2.0 Image Monitor System Information Security Standard - Part 1: Information Security Requirements [14], which is a Taiwanese IoT information security standard focused on image monitoring systems, includes a mapping table to the OWASP IoT Top 10:2014 [16], these two information security standard have built a explicit relations between their controls, which allows this study to use the information in this table as a reference for the accuracy of automated comparisons with deep learning models. Fig. 4 below shows a screenshot of the controls in Appendix D of the standard "IoT-1001-1 v2.0 Image Monitor System Information Security Standard - Part 1: Information Security Requirements" against each standard specification.

附錄 D

(參考)

技術要求事項與各標準規範對照表

Control's serial number of the standard

表 D.1 技術要求事項與各標準規範對照表

技術要求	OWASP 對應項目(4)	ANSI/UL 2900-1 對應項目	ONVIF 對應項目(19- 20)
5.1.1.1	I10 : Poor Physical Security Ensuring only required external ports such as USB are required for the product to function. Ensuring the product has the ability to limit administrative capabilities.		-

Fig. 4. A screenshot of the standard "IoT-1001-1 v2.0 Image Monitor System Information Security Standard - Part 1: Information Security Requirements, Appendix D" against OWASP IoT Top 10:2014

The reference answer mapping of the standard comparison is based on the "IoT-1001-1 v2.0 Image Monitor System Information Security Standard - Part 1: Information Security Requirements" standard, and the official mapping table of the standard to OWASP IoT Top 10:2014 in Appendix D of the standard as the reference answer. In other words, a total of 38 screened security items in the Image Monitor System Information Security Standard will actually be classified into the ten corresponding categories of "OWASP IoT Top 10:2014". Although each category of "OWASP IoT Top 10:2014" has from 4 to 14 information security controls, it was found that it is difficult to match the information security control specified in the reference answer for information security standards from different sources. In addition to the difference in terminology between different standards, it is assumed that the accuracy of the wording of the original Chinese standard will be affected after translation. Therefore, in this section, we choose to convert the accuracy of the base standard information security control into reference values by whether they are correctly classified or not. Figure 5 below shows the schematic diagram of the two experimental approaches.

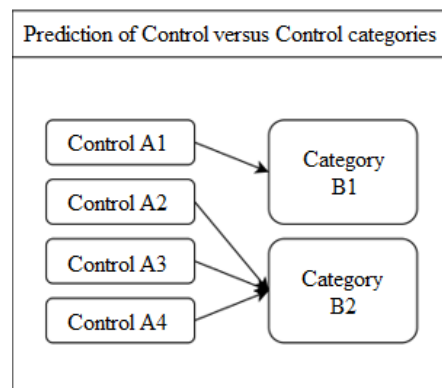


Fig. 5. A schematic diagram of controls versus control categories comparison

As shown in the Fig. 4, which shows that the control numbered 5.1.1.1 of "IoT-1001-1 v2.0 Image Monitor System Information Security Standard - Part 1: Information Security Requirements [14]" can actually corresponded to category I10 of "OWASP IoT top 10:2014 [16]".

However, since the standard itself is written in Chinese, it needs to be translated into English and then fed into a deep learning model for comparison, so we have used the translation module mentioned in section 3.1 to automatically complete this task for us. After checking the table, a total of 38 filtered security controls in the "IoT-1001-1 v2.0 Image Monitor System Information Security Standard - Part 1: Information Security Requirements" will actually be classified into the ten corresponding categories of "OWASP IoT Top 10:2014".

4. Evaluation Results

4.1. Initial Evaluation Results

We used the five models with the best performance in Table 4. and *distiluse-base-multilingual-cased-v2*, which is a multilingual model that supports more than 50 different languages, and more balanced in the scores of the indicators, were selected and compared with the "OWASP IoT Top 10: 2014", and the following Table 5 shows the experimental results.

Table 5. Comparison of experimental results with different SBERT models

Exp. No.	Model	k=1	k=2
S1	all-mpnet-base-v2	61 %	68 %
S2	multi-qa-mpnet-base-dot-v1	50 %	66 %
S3	paraphrase-MiniLM-L3-v2	39 %	50 %
S4	distiluse-base-multilingual-cased-v2	68 %	68 %
S5	paraphrase-multilingual-mpnet-base-v2	50 %	74 %

In the above table, the k represents the prediction of the k most similar outcomes at the end of each prediction. In other words, when the model can output the least number of predictions, the more accurate it can hit the same category of predictions, which means that the model has a better performance on the task of matching information security standards. The number of successful hits is one of the important indicators of the effectiveness of the reference model for this task.

Under this condition, experiment number S1 and S4 have the best performance, which are *all-mpnet-base-v2* and multilingual model *distiluse-base-multilingual-cased-v2*, achieving 61% and 68% hit rate respectively under the restriction of $k=1$, and 68%, 74% hit rate respectively under the restriction of $k=2$, which means at least three quarters of controls in the standard were successfully predicted to the correct categories by the deep learning models.

In addition to the difference in terminology between different standards, the accuracy of the wording of the original standard will also be affected if it is translated, not to mention the fact that there are also controls or requirements that meet several OWASP IoT Top 10 categories after review and analysis, but the reference answer only has a given category and thus cannot be included. However, when it comes to the actual use for the information standards, even though they are for the same domain-oriented information security standards, there are some parts that are not similar. In practice, when an information security consultant is looking for controls or requirements that are suitable for a particular case, the items that are suitable for the case may be scattered in different information security standards, or different categories inside the same standard. Among those that are not successful, there must be some items that are not in the same category but have similar practical applications and application methods.

4.2. Discussion of Evaluation Results

In Section 3.2 of this paper, the five SBERT models that performed better on average were compared with the results of the comparison experiments between their scores provided in

the official guidance documents and the English standards. This means that nearly a quarter of the information security items are difficult to classify correctly by the model. The actual list of information security item numbers that were not predicted by each model shows that these unpredictable information security item numbers are specific numbers, as shown in Fig. 6 below shows the prediction status of each model for the specified corresponding security item at $k=3$. The dark squares indicate that the number was successfully predicted by the specified model, while the light squares indicate that the number was not successfully predicted by the specified model.

Models \ Controls which can not be predicted	1	2	7	8	9	15	16	17	18	19	20	21	23	24	26	27	29	33	34	37	38
all-mpnet-base-v2																					
multi-qa-mpnet-base-dot-v1																					
paraphrase-MiniLM-L3-v2																					
distiluse-base-multilingual-cased-v2																					
paraphrase-multilingual-mpnet-base-v2																					
unsupervised k-nearest neighbor																					

Fig. 6. Standard controls for which none of the plural models can be predicted

The distribution of the light-colored squares in the above figure shows that the information security items that cannot be successfully predicted by the specified models are very similar for the above five deep learning models, especially for Experiment Numbers. 2, 7, 8, 9, 20, 21, and 23. Experiment number 2, 7, 8, 9, 20, 21, 23, these experiment numbers correspond to the following items in the original standard: "IoT-1001-1 v2.0 Image Monitor System Information Security Standard - Part 1: Information Security Requirements [14]".

Based on the above table, it can be inferred that it is more difficult for the deep learning models to classify the contents of information security controls into two categories, category 2 and category 8. When encountering the above information security controls in practice, both category 2 and category 8 will not be the first choice of the deep learning models, but other categories. In "OWASP IoT Top 10:2014" [16], category 2 is Insufficient Authentication/Authorization, which translates to unreliable authentication mechanism, and category 8 is Insufficient Security Configurability, which translates to unreliable security configuration.

From the evaluation results, the two deep learning models with the best prediction results, *all-mpnet-base-v2* and *distiluse-base-multilingual-cased-v2*, are the best predicted models for the translated "IoT-1001-1 v2.0 Image Monitor System Information Security Standard - Part 1: Information Security Requirements [14]", which corresponds to the prediction of "OWASP IoT Top 10:2014", achieved 61% and 68% hit rate at $k=1$ respectively. The prediction results are shown in Table 7 below.

According to the above table, the prediction results can be easily classified into two categories, one is the category where Model 1, denoted as $M1$, and Model 2, denoted as $M2$, have the same prediction results for the information security sub-category, but both predict failure. Table 8 explores the potential causes of prediction errors in the classification results. In Experiment Number 2, the correct category should in Category 2, the

Table 6. List of standard controls that cannot be predicted by the majority of models.

Exp. No.	Control Number	Correct Category	Control Description
2	5.1.3.1	2	The product should not have the ability to restore the default pass code with your bare hands.
7	5.2.4.1	8	Sensitive data stored in the product shall be accessible only by authorized individuals.
8	5.2.4.2	8	The identity authentication factor and key for encryption and decryption (excluding the public key for asymmetric encryption) stored in the product should not be stored in clear text, and the data should be protected by the security functions approved by NIST SP 800-140C, CMVP Approved Security Functions.
9	5.2.4.4	8	Sensitive data should be stored in the security domain of the product, isolated from the normal operating environment.
20	5.3.3.1	8	The product should provide the user to turn on/off the WPS PIN function of "Wi-Fi Protected Setup (WPS)" and its default value should be off.
21	5.3.3.2	8	By default, the Wi-Fi security mechanism should be "Wi-Fi Protected Access (WPA)" and the version of Wi-Fi Protected Access should meet the requirements of Appendix C.
23	5.4.1.1	2	Before accessing the product resources, the identity identification mechanism with protection against retransmission attacks should be adopted.

Table 7. Prediction results of the two models for the specified standard controls

Exp.No.	Control	M1 prediction result	M2 prediction result
2	5.1.3.1	10	10
7	5.2.4.1	5	5
8	5.2.4.2	4	4
9	5.2.4.4	5	10
20	5.3.3.1	7	5
21	5.3.3.2	7	7
23	5.4.1.1	10	7

corresponding information security subdivision is that the product should not have the ability to restore the default passcode externally with bare hands. It should be the word "external" that causes the deep learning model to predict this information security item as Category 10: Poor physical security. In Experiment Number 7, the corresponding information security breakdown is that sensitive information stored in the product should only be accessed by authorized individuals. In terms of this information security control, it is reasonable to predict to Category 5: privacy concerns because it also describes user privacy. In Experiment Number 8, The corresponding information security itemized content is: the identity authentication factor and key for encryption and decryption (excluding the public key for asymmetric encryption) stored in the product should not be stored in clear text, and the data protection method should be used with the security functions approved by NIST SP 800-140C, CMVP Approved Security Functions. In terms of this information security category, the prediction to Classification 4: Lack of Transport Encryption is reasonable because it contains key words in the field of encryption such as encryption and decryption, key and plaintext. In Experiment Number 21, for the information security control, the default security mechanism for Wi-Fi is "Wi-Fi Protected Access (WPA)" and the version of Wi-Fi Protected Access should meet the requirements of Appendix C. In terms of this information security control, the predicted classification is Category 7: Insecure Mobile Interface, which is not accurate. It is guessed that in the pre-training data of the two models, Wi-Fi usually appears together with key words such as cell phone and mobile, so the models classified it as Category 7.

Table 8. The two models jointly classify the error causes of the false security item

No.	Control	Result ($M1 \& M2$)	Correct category
2	5.1.3.1	10	2
7	5.2.4.1	5	8
8	5.2.4.2	4	8
21	5.3.3.2	7	8

The reason behind the inaccurate predicts are that different deep learning models have different prediction judgments for the same information security item, but the classification is basically similar to the former one: it is influenced by specific wording, or the information security item may apply to both plural "OWASP IoT Top 10:2014" classification, resulting in its misclassification. The actual results of the respective predictions are listed for analysis, and the reasons for the wrong classification results are speculated in Table 9. In Experiment Number 9, the corresponding information security sub-section is: sensitive data should be stored in the security domain of the product, isolated from the normal operating environment. Model 1 predicts that Category 5: Privacy Concerns are reasonable, and sensitive data are indeed related to user privacy; Model 2 predicts that Category 10: Poor Physical Security is not reasonable, and the model presumes that the information security item is not related to physical security because of the terms "operating environment", "isolation", and "security domain". The model predicts that the information security item is related to the description of physical security because of the terms "operating environment," "isolation," and "secure area. In Experiment Number 20,

the information security control is: the product should provide users to turn on/off the WPS PIN function of "Wi-Fi Protected Setup (WPS)", and the default value should be off. Model 1 predicts Category 7: Insecure Mobile Interface, which is a relatively inaccurate classification. It is guessed that in the pre-training data of both models, Wi-Fi usually appears together with key words such as cell phone and mobile, so the model classifies it as Category 7. After all, if Wi-Fi is automatically connected to public networks, it may cause user privacy leakage, which is a user privacy concern.

In Experiment Number 23, the corresponding information security control is: Before accessing product resources, identity authentication mechanism with protection against retransmission attacks should be used. Model 1 predicts a classification of 10: Poor Physical Security, which is inaccurate. Model 2 predicts a classification of 7: Insecure mobile interface, which is more reasonable than the prediction of Model 1, but not correct. In Experiment Number 21, for the information security control, the default security mechanism for Wi-Fi is "Wi-Fi Protected Access (WPA)" and the version of Wi-Fi Protected Access should meet the requirements of Appendix C. In terms of this information security control, the predicted classification is Category 7: Insecure Mobile Interface, which is not accurate. It is guessed that in the pre-training data of the two models, Wi-Fi usually appears together with key words such as cell phone and mobile, so the models classified it as Category 7.

產業

學院

Table 9. The two models each classify the wrong security category of the error cause speculation

No	Control	Result(M1)	Result(M2)	Correct category
9	5.2.4.4	5	10	8
20	5.3.3.1	7	5	8
23	5.4.1.1	10	7	2
21	5.3.3.2	7	7	8

From Table 9 and the speculation on the failure of the prediction of the security category for which none of the plural models could be predicted, it is clear that at least half of the security categories that failed to be predicted may also apply to the plural "OWASP IoT Top 10:2014[16]" classification, plus the fact that in the standard "IoT-1001-1 v2.0 Image Monitor System Information Security Standard - Part 1: Information Security Requirements [14]", the corresponding OWASP In Appendix D of the original Top 10:2014 mapping table, the security subcategory does not specify a mapping to another subcategory, even though the subcategory is similar for that security subcategory, resulting in model prediction failure. By actually viewing the table and the information security controls that failed for the seven security controls that could not be predicted by the plural model, if the predictions that were judged to be reasonable were categorized as correct predictions, with model 1 representing *all-mpnet-base-v2* and model 2 representing *distiluse-base-multilingual-cased-v2*, the two models The final revised prediction results for these seven information security controls are shown in Table 10 below.

Table 10. Results of the error analysis of the two models for the unpredictable security controls breakdown

Exp. No. / Ctrl. No.	2 / 5.1.3.1	7 / 5.2.4.1	8 / 5.2.4.2	9 / 5.2.4.4	20 / 5.3.3.1	21 / 5.3.3.2	23 / 5.4.1.1
Model 1	X	V	V	V	X	X	X
Model 2	X	V	V	X	V	X	X

According to the above table, model 1: *all-mpnet-base-v2* and model 2: *distiluse-base-multilingual-cased-v2* achieve 61% and 68% hit rate respectively for $k=1$. If the predictions with reasonable classification are classified as correct and recalculated, the hit rate will increase to 69% and 76%.

Finally, the experimental results proved that the use of deep learning models for fast and automated comparison of information security standard content has good accuracy and retains considerable room for improvement.

4.3. Final Evaluation Results

SBERT [21], as an enhanced version of BERT [1] for text similarity search task, provides a pre-training model with higher accuracy than the native pre-training model provided by BERT. Table below shows the best two models in SBERT, *all-mpnet-base-v2*

and *distiluse-base-multilingual-cased-v2*, with the same $k=1$, i.e., each information security sub-prediction only outputs one closest information security sub-prediction, and this output value is the only consideration for accuracy. Under the condition that the translated "IoT-1001-1 v2.0 Image Monitor System Information Security Standard - Part 1: Information Security Requirements [14]" corresponds to the prediction of "OWASP IoT Top 10:2014"[16].

Table 11. Deep Learning Approach to Information Security Standard Prediction Implementation Results

Exp. No.	Model name	Predict accuracy	Predict / All
SS1	all-mpnet-base-v2	69%	26 / 38
SS2	distiluse-base-multilingual-cased-v2	76%	29 / 38

According to the above table, the better model, *distiluse-base-multilingual-cased-v2*, successfully predicted 26 of the 38 information security items with $k=1$, while the remaining items failed to be predicted by the plural model in sub section 4.2 of this study. In this study, we examined the seven information security items for which both models failed to predict, and confirmed that three of the items were also applicable to the plural "OWASP IoT Top 10:2014" classification, although they did not match the answers.

In spite of the information security standards targeting the same aspect, there will still be parts where they differ significantly from each other. In practice, when information security consultants are looking for suitable sub-items or control measures for a specific case, the relevant items may be spread across different categories. Among the items that don't align perfectly, there are bound to be some that, while not in the same category, are very similar in practical application and usage. This suggests that, in practical terms, using deep learning models for comparing information security standards has shown, from experimental results, to be not only faster but also fairly accurate. Its performance surpasses the implementation using traditional machine learning. In the future, this research will experiment with generative AI, attempting to produce more general terms related to the control items of different standard, and then apply the SBERT method for further experimentation to enhance the readiness of successful classification.

5. Conclusion

This study utilises the contents of multiple international information security standards and translated domestic standards as its dataset, possessing the ability to rapidly identify similar control items. The content is not restricted to a single language and demonstrates good predictive accuracy. The study also proposes an automated process, streamlining a workflow that would otherwise require significant labour to review and compare. Ultimately, this can serve as a reference for scholars wishing to conduct future research in text processing, text mining, deep learning, and information security standards.

Although this research has achieved commendable results in comparing similarities among different information security standards, there are still many areas that warrant

further exploration in the future. For instance, automated data processing procedures or the application of machine learning methods such as Few-Shot Learning for data with lower volume, greater diversity, and insufficient annotations. Additionally, the use of generative AI represents another avenue to explore. Some standards may feature different customary terminologies across various standards organisations or publishers. Generating more general terms related to control and then utilising the SBERT method for further experiments might enhance the accuracy of successful classifications.

Acknowledgments. This research was partially funded by National Science and Technology Council (NSTC 112-2221-E-027-067-).

References

1. Devlin, J., Chang, M.W., Lee, K., Toutanova, K.: Bert: Pre-training of deep bidirectional transformers for language understanding. arXiv preprint arXiv:1810.04805 (2018)
2. Devlin, J., Chang, M.W., Lee, K., Toutanova, K.N.: Bert: Pre-training of deep bidirectional transformers for language understanding (2018), <https://arxiv.org/abs/1810.04805>
3. European Telecommunications Standards Institute: EN 303 645:CYBER; Cyber Security for Consumer Internet of Things: Baseline Requirements, v2.1.1 edn. (2020)
4. Fellbaum, C.: WordNet, pp. 231–243. Springer Netherlands, Dordrecht (2010), https://doi.org/10.1007/978-90-481-8847-5_10
5. Guthrie, D., Allison, B., Liu, W., Guthrie, L., Wilks, Y.: A closer look at skip-gram modelling. In: Proceedings of the Fifth International Conference on Language Resources and Evaluation (LREC'06). European Language Resources Association (ELRA), Genoa, Italy (May 2006), http://www.lrec-conf.org/proceedings/lrec2006/pdf/357_pdf.pdf
6. Hassan, S.: Measuring semantic relatedness using salient encyclopedic concepts. Ph.D. thesis (2011), <https://www.proquest.com/dissertations-theses/measuring-semantic-relatedness-using-salient/docview/1011651248/se-2>
7. Inan, E.: Simit: a text similarity method using lexicon and dependency representations. *New Generation Computing* 38(3), 509–530 (2020)
8. International Electrotechnical Commission: Security for industrial automation and control systems - Part 4-2: Technical security requirements for IACS components. (2019)
9. Ji, Y., Eisenstein, J.: Discriminative improvements to distributional sentence similarity. In: Proceedings of the 2013 conference on empirical methods in natural language processing. pp. 891–896 (2013)
10. Kiros, R., Zhu, Y., Salakhutdinov, R.R., Zemel, R., Urtasun, R., Torralba, A., Fidler, S.: Skip-thought vectors. In: Cortes, C., Lawrence, N., Lee, D., Sugiyama, M., Garnett, R. (eds.) *Advances in Neural Information Processing Systems*. vol. 28. Curran Associates, Inc. (2015), https://proceedings.neurips.cc/paper_files/paper/2015/file/f442d33fa06832082290ad8544a8da27-Paper.pdf
11. Le, Q., Mikolov, T.: Distributed representations of sentences and documents. In: Xing, E.P., Jebara, T. (eds.) *Proceedings of the 31st International Conference on Machine Learning. Proceedings of Machine Learning Research*, vol. 32, pp. 1188–1196. PMLR, Beijing, China (22–24 Jun 2014), <https://proceedings.mlr.press/v32/le14.html>
12. Liu, H., Singh, P.: Conceptnet—a practical commonsense reasoning tool-kit. *BT technology journal* 22(4), 211–226 (2004)

13. Mihalcea, R., Corley, C., Strapparava, C.: Corpus-based and knowledge-based measures of text semantic similarity. In: Proceedings of the 21st National Conference on Artificial Intelligence - Volume 1. p. 775–780. AAAI’06, AAAI Press (2006)
14. Mobile Application Security Alliance: IoT-1001-1 v2.0 Image Monitor System Information Security Standard - Part 1: Information Security Requirements (2021)
15. Mohamed, M., Oussalah, M.: A hybrid approach for paraphrase identification based on knowledge-enriched semantic heuristics. *Language Resources and Evaluation* 54, 457–485 (2020)
16. OWASP IoT Security Team: OWASP Internet of Things (IoT) Top 10 2014. (2014)
17. OWASP IoT Security Team: OWASP Internet of Things (IoT) Top 10 2018. (2018)
18. OWASP IoT Security Team: OWASP Top 10 vulnerability 2021. (2021)
19. Peters, M.E., Neumann, M., Iyyer, M., Gardner, M., Clark, C., Lee, K., Zettlemoyer, L.: Deep contextualized word representations. *CoRR* abs/1802.05365 (2018), <http://arxiv.org/abs/1802.05365>
20. Reimers, N., Gurevych, I.: Sentence-BERT: Sentence embeddings using Siamese BERT-networks. In: Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP). pp. 3982–3992. Association for Computational Linguistics, Hong Kong, China (Nov 2019), <https://aclanthology.org/D19-1410>
21. Reimers, N., Gurevych, I.: Sentence-bert: Sentence embeddings using siamese bert-networks. *arXiv preprint arXiv:1908.10084* (2019)
22. Swamy, S.N., Kota, S.R.: An empirical study on system level aspects of internet of things (iot). *IEEE Access* 8, 188082–188134 (2020)
23. Turian, J., Ratnoff, L., Bengio, Y.: Word representations: a simple and general method for semi-supervised learning. In: Proceedings of the 48th annual meeting of the association for computational linguistics. pp. 384–394 (2010)
24. Wang, Z., Mi, H., Ittycheriah, A.: Sentence similarity learning by lexical decomposition and composition. In: Proceedings of COLING 2016, the 26th International Conference on Computational Linguistics: Technical Papers. pp. 1340–1349. The COLING 2016 Organizing Committee, Osaka, Japan (Dec 2016), <https://aclanthology.org/C16-1127>

Yu-Chih Wei is an Associate Professor in the Department of Information and Finance Management at the National Taipei University of Technology. He holds a Ph.D. in Information Management from National Central University, and a B.S. and a M.S. in Information Management from YuanZe University. His research interests include FinTech security, health informatics security, ISRA, SupTech, VANET security, information security management, and business continuity management. Before pursuing an academic career, Dr. Wei was a researcher at the Information & Communication Security Laboratory of Chunghwa Telecom Co., Ltd.

Yu-Chun Chang received his M.S. degree in Department of Information and Finance Management, National Taipei University of Technology in 2023. His research interests include information security and text mining.

Wei-Chen Wu is Assistant Professor in the Department of Finance at the National Taipei University of Business. He received his Ph.D. degree in Information Management from National Central University in 2016. From 2020-2021, He was Assistant Professor in the Department of Finance at the Feng Chia University. From 2008-2016, he was also

Assistant Professor and Director of the Computer Center at Hsin Sheng College of Medical Care and Management. His teaching interests lie in the area of programming languages, ranging from theory to design to implementation and his current research interests include blockchain technology, fintech cybersecurity, network security, and deep learning. Wei-Chen Wu has collaborated actively with researchers in several other disciplines of computer science. He has served on many conference and workshop program committees and served as the workshop chair for Frontier Computing Conference (FC2017 FC2021) and Machine Learning on FinTech, Security and Privacy Conference (MLFSP2019 MLFSP2023).

Received: August 22, 2023; Accepted: November 21, 2023.

學院

產業



中華民國專利證書

發明第 I 856852 號

發明名稱：金融機構之跨平台操作系統及其方法

專利權人：丁俞豪、周其綦、陳楷閔、陳永瑩、吳威震

發明人：丁俞豪、周其綦、陳楷閔、陳永瑩、吳威震

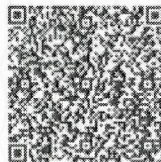
專利權期間：自 2024 年 9 月 21 日至 2043 年 10 月 26 日止

上開發明業經專利權人依專利法之規定取得專利權

經濟部智慧財產局 局長

廖承威

中華民國 113 年 9 月 21 日



注意：專利權人未依法繳納年費者，其專利權自原繳費期限屆滿後消滅。

·課程標準

·上課時間表

·註冊編班名單

·選課表

·退出(關閉瀏覽器)

上課時間表

113金融資安產碩專秋季班 第三學期 【114/9-115/1】

課號	課程名稱	學分	時數	修	教師	上課時間、地點	人	撤	教學大綱與進度表	備註
1222	論文	3.0	3	▲			16	0		
1223	金融科技與資訊安全實習(三)	3.0	14	★			9	0		
1225	保險科技與資安防護	3.0	3	★	吳威震 林敬皇	114年09月09日星期二 第 5 節 【科研大樓243e】 115年01月06日星期二 第 7 節 【科研大樓243e】	11	0	查詢 查詢	
1226	人工智慧與金融大數據	3.0	3	★			1	0		跨114金融(1209)
小計		12.0	23							

114-1 學期教學大綱及進度表

課號 Course Number	課程名稱 Course Name	授課教師 Instructor	開課班級 Class	學分數 Credits	時數 Hours	授課語言 Language
1225	保險科技與資安防護	吳威震	113 金融資安產碩專秋季班	3	3	中文
課程大綱 /Syllabus	本課程探討保險科技在承保前調查與事故調查中的應用，並以數位鑑識為核心，強調如何透過科技提升保險業務的安全性與準確性。課程包含保險科技（InsurTech）的發展與資安防護，並奠定數位鑑識的基礎，包括法規與倫理考量、探討承保前調查與風險評估技術，如背景調查、風險建模及 AI 欺詐偵測、聚焦於數位鑑識工具，如 EnCase、Autopsy、FTK，並分析電子設備及物聯網的鑑識挑戰、透過案例研究與實務演練，探討保險理賠爭議的解決方案與企業資安機制的建立，並進行綜合討論與學員專題報告。結合理論與實務，適合希望理解保險科技與資安防護的學習者。					
課程進度 /Schedule 請註明 a.週次 /Week b.單元主題 /Topic	1. 保險科技概論與數位鑑識基礎 2. 承保前調查與風險評估技術 3. 數位鑑識技術與工具介紹 4. 事故調查的數位鑑識應用 5. 案例研究與實務演練 6. 數位鑑識未來趨勢與挑戰 1. W12 保險科技概論與區塊鏈 2. W13 Solidity 與智能合約 1 3. W14 Solidity 與智能合約 2 4. W15 保險智能合約邏輯設計 5. W16 Chainlink Oracle 6. W17 保險區塊鏈資安風險 7. W18 期末報告					
評量方式與標準 Evaluation and grading policy	1. 課堂課堂參與課堂實作 30% 2. 指定作業 30% 3. 期末專題報告 40%					
使用教材、參考	*使用外文原文書： ☐ 是 ☒ 否					

書目或其他 Materials	*使用自編教材/Use of self-made materials : ☐ 是 ☒ 否 1. 自編教材
課程諮詢管道 The access to curricular consultation (上限 200 字)	週二 234、週三 567
*課程對應 SDGs 指標 (可複選) The course corresponds to the SDGs (mark all that apply)	參閱 SDGs 項目：網頁連結(https://osausr.ntut.edu.tw/p/426-1001-76.php?Lang=zh-tw) SDG8：尊嚴就業與經濟發展 (Decent Work and Economic Growth) SDG9：產業創新與基礎設施 (Industry, Innovation and Infrastructure)

以區塊鏈與智能合約建立安全與公平之跨區域市場代幣交易制度

吳若綺

國立臺北商業大學資訊管理系

摘 要

為了解決跨區域市場可流通代幣的交易方式，亦不受相關金融監理機構法規所限制，本論文將提出一個以區塊鏈和智能合約架構下，建立一個安全公平及可信任的跨區域市場可流通的代幣交易制度，其特性為安全、公平及可信任、跨區域市場的代幣交易制度。

關鍵詞：區塊鏈，智能合約，可流通代幣，金融科技。

1. 緒論

近幾年金融科技(Fintech)相關應用與議題不斷地被提出，加上虛擬貨幣慢慢也被大眾所接受，甚至拿來做交易[3,4,6,7,9]，也被許多國家認可能與實體貨幣作兌換或當作金融商品做交易，尤其以目前兩大虛擬貨幣比特幣(Bitcoin)與以太坊(Ethereum)[2,11]為最大市場，這兩套虛擬貨幣系統都架構在公有區塊鏈的信任機制上，比特幣是區塊鏈(Blockchain)技術最本質的應用，它也因為在沒有政府的干預下，達成了上百萬美元的匿名交易市場，而成為最具爭議性的區塊鏈應用。然而虛擬貨幣錢包匿名的特性，也衍生出需多金融犯罪等問題，且錢包私鑰遺失將造成虛擬貨幣永久無法使用，再者虛擬貨幣須接受當地國家相關金融監理機構所規範，若相關應用只限在各區域市場所交易或流通，限制將非常大，此外目前有許多區域市場的應用，例如光幣[12]、MCoin[13]、東海酷幣[14]與南臺幣[15]，但各區域貨幣皆是建立在的私有鏈架構上，且並無智能合約來確保雙方的公平性，因次本論文將提出以區塊鏈和智能合約架構下，建立一個以安全及可信任的區域市場可流通的代幣交易制度，其特性為安全及可信任、區域市場交易與可跨區域流通的代幣交易制度。

2. 區塊鏈與智能合約

區塊鏈是一種分散式儲存的資料結構，並以密碼學方式保證交易資料是不可篡改和不可偽造的分散式公開帳本。在區塊鏈上的區塊記錄了一定時間內被審核通過的交易記錄。一旦被記載到區塊鏈上，相關的交易資訊就不能被修改、刪除。區塊鏈技術本身可以被應用在金融和非金融領域，傳統的金融交易是建立在可靠的、可信任的公正第三方的基礎上，使用者在網路交易時必須依賴於公正第三方，並確認交易雙方身份的安全。但依靠公正第三方維繫安全與個人網路資產隱私的背景中，並不是絕對的安全，當公正第三方一旦被侵入，將無法再保證其所維護相關資料是安全的。因此區塊鏈的技術解決了不需過度依賴公正第三方，通過分散式的共識的機制和匿名特性，在複雜的網路環境中解決彼此互信任的問題，有以下的特徵：

(1)去中心化

區塊鏈在沒有公正第三方的情況下，達成了所有資料的分散式記錄與儲存，並且能夠保證資料記錄的真實性。區塊鏈是由點對點(P2P, Peer To Peer)的網路組成，在這種去中心化的網路環境中，所有網路節點都是相同對等的，所有節點享有相同的權利和義務。由於所有節點並沒有公正第三方或者信任機構背

書，所以在去中心化的區塊鏈中，並無單一節點的攻擊，因此不會因為單一節點失效而對整個區塊鏈網路產生影響。

(2) 資料及操作公開性

區塊鏈作為分散式帳本技術，系統內所有的資料記錄及操作對於所有在網路節點都是公開的。在典型的區塊鏈中，每一個節點都能夠儲存所有發生的歷史交易記錄。區塊鏈使用非對稱加密演算法等密碼學技術的組合應用，保證區塊鏈資訊在網路的高度透明性，並且區塊鏈網路運行的程式、規則、節點的接入方式都是公開的，這是區塊鏈信任的基 <http://web.fintech.mcu.edu.tw/zh-hant/content/銘傳幣礎>。這些機制的運用，確保區塊鏈中的記錄可以被所有節點審查和追蹤。

(3) 資訊不可篡改性

區塊鏈中的資訊是不可篡改的，一旦資料資訊被驗證通過寫入區塊中並加入區塊鏈，就無法被篡改。區塊鏈的資料資訊必須經過網路大部分節點的審核以後，才能允許被記錄。除非能夠控制系統中 51% 以上網路節點，否則對單一節點的區塊記錄篡改是沒有意義的，這種不可篡改特性保證了區塊鏈資料的穩定性與可靠性。

(4) 匿名性

區塊鏈在複雜的網路環境中解決了在節點間的信任問題，因此區塊鏈中的交易節點可以在無需瞭解對方身份的情況下進行交易。區塊鏈中的交易是基於加密位址，而不會對交易雙方身份進行認證。交易雙方僅需要公佈自己的位址就可以與對方進行交易，區塊鏈的節點使用非對稱加密技術，建立節點間在匿名環境下的信任。所有節點維持自身的公鑰與私鑰對，節點公開發佈自己的公鑰，保留自己的私鑰。進行資訊傳遞，資訊接收方公佈自己的公鑰給發送方，發送方使用其公鑰加密。資訊接收方在接收到傳遞的加密資訊後，使用自己的私鑰對加密過的資訊進行解密。通過這樣的方式，節點間可以在不需要身份認證的情況下，完成匿名環境下的信任交易[1,5]。

智能合約(Smart Contract)由密碼學家 Nick Szabo 提出[10]，智能合約可以將交易的條款透過資訊化的方式來落實；並在沒有第三方的情況下，透過自動的方式與互不信任的人用區塊鏈信任的機制完成交易。而目前以太坊(Ethereum)與 HyperLedger 都是運用區塊鏈的特性來落實智能合約的概念，同時也是目前兩大的智能合約區塊鏈運作平台，以太坊使用 Solidity 開發工具撰寫智能合約[2]，而 HyperLedger 使用 Golang 撰寫[8]。

智能合約部署方式是事先使用開發工具寫好合約內容，並編譯產生 bytecode 與 ABI (Application Binary Interface) 智能合約的 bytecode 存入至區塊鏈中，合約部署到區塊鏈後所產生的帳戶位址即合約位址，當使用者要與智能合約互動時，就要透過 ABI 以及合約位址來發起交易，以執行智能合約中的程式。

3. 代幣交易制度

本研究建立的制度因為屬區域市場，因此所建置之區塊鏈架構為私有鏈，如圖 1 所示。代幣發行者負責將多筆交易紀錄儲存在新的區塊中，一旦儲存在區塊內的交易紀錄即成永久儲存，任何人皆無法修改。

表 1 符號說明

符號	說明
$Token_{value}$	代幣金額，例如 10 單位代幣 $Token_{10}$
U_i	代幣使用者 i

S_k	代幣發行者 k
M_j	合作商家 j
(e_i, d_i)	代幣使用者 i 的公鑰 e_i 與私鑰 d_i
(e_k, d_k)	代幣發行者 k 的公鑰 e_k 與私鑰 d_k
(e_j, d_j)	合作商家 j 的公鑰 e_j 與私鑰 d_j
$E_{key}\{msg\}$	使用 key 金鑰對訊息 msg 做加密
$Sig_{key}\{msg\}$	使用 key 金鑰對訊息 msg 做簽章
P_{value}	商品價值
Pay_{value}	商品支付費用
Q	代幣與現今的兌換率
Hash	雜湊函數
$SC(x, y)$	x 與 y 協議的智能合約
T	當下的時間戳記

區塊的產生由代幣發行者執行，產生時機可以選擇在固定的時間(T)內或達到最大交易記錄(N)時，將所完成的交易紀錄全部直接儲存在新的區塊中。

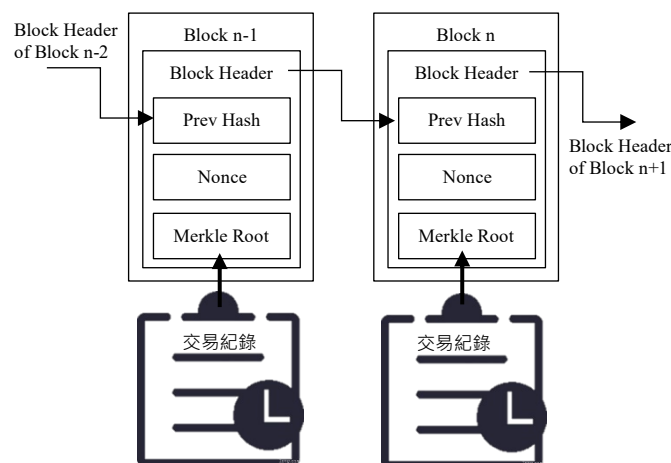


圖 1 區塊鏈架構

在代幣使用與現金兌換過程中，必須先建立各種情境的智能合約佈署與相關的參數的定義。本架構考量了兩種形式的智能合約：

- 代幣使用的智能合約

$SC(M_j, U_i)$ 並須定義不同店家等級的折抵上限 U 值與使用比例 W 值。

- 現金兌換的智能合約

$SC(U_{i+1}, U_i)$ 、 $SC(S_k, U_i)$ 、 $SC(M_j, S_k)$ 可分別針對不同兌換雙方定義不同的智能合約與兌換比率 Q 值。

3.1 代幣分配

當有新的使用者完成註冊後，代幣發行者將分配基本的 $\text{Token}_{\text{value}}$ ，此分配也是一筆交易紀錄，也會記錄在區塊鏈中，當使用者傳送註冊訊息與自己的公鑰 e_i 當做交易位址，代幣發行者也將回復自己的交易位址 e_k 與時間戳記 T 並透過使用者公鑰 e_i 對 $\{T, \text{Sig}_{d_k}(\text{Token}_{\text{value}})\}$ 做加密，當在 T 時間內或達到一定使用者註冊量時即最大交易記錄 N ，將所有註冊申請的新使用者所分配到的 $\text{Token}_{\text{value}}$ 紀錄在新的區塊中，此時也是建立此私有鏈之創世區塊的時機，如圖 2 所示。

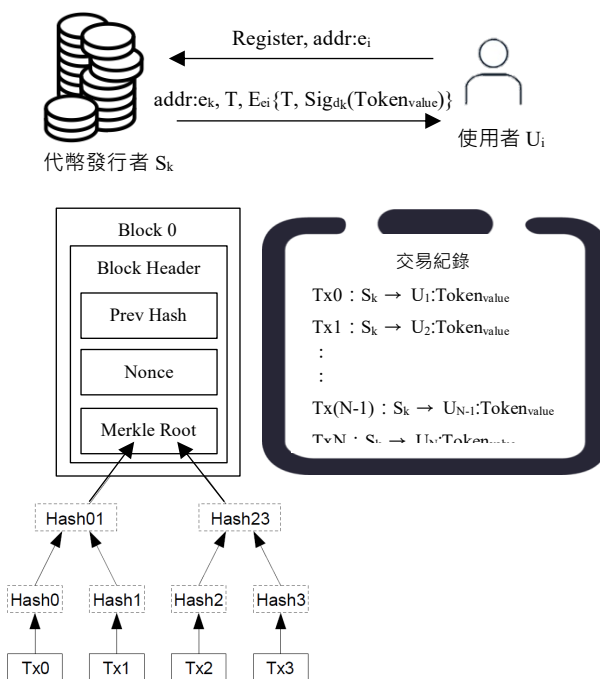


圖 2 代幣分配

3.2 代幣使用

代幣發行者可以透過地理位置尋找合作商家，並透過合作商家等級劃分來決定代幣使用的比例，如表 2 範例，透過代幣的折抵，提供使用者折抵商品的部分金額，假設商品價值 P_{value} ，使用者可以使用 $\text{Token}_{\text{value}}$ 來折抵部分金額，因此使用者只需支付部分實際金額 $\text{Pay}_{\text{value}}$ ，如公式(1)。

$$\text{Pay}_{\text{value}} = \begin{cases} P_{\text{value}} - \text{Token}_{\text{value}}, & \text{if } \text{Token}_{\text{value}} < W \times P_{\text{value}} \\ P_{\text{value}} - \text{Token}_u, & \text{if } \text{Token}_{\text{value}} > W \times P_{\text{value}} \end{cases} \quad (1)$$

當使用者扣掉代幣 $\text{Token}_{\text{value}}$ 折抵後支付給合作商家即 $\text{Pay}_{\text{value}} = \text{Cash} + \text{Token}_{\text{value}}$ ，當商家想出售自己的商品或服務時，將傳送自己的公鑰 e_j 當做交易位址，使用者將回復自己的交易位址 e_i 與時間戳記 T 並透過商家公鑰 e_j 對 $\{T, \text{Sig}_{d_i}(\text{Pay}_{\text{value}})\}$ 做加密，當交易完成，合作商家將提供商品或服務給使用者使用，如圖 3 所示。

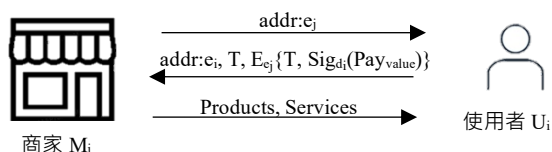


圖 3 代幣使用

根據公式(1)與表 2 範例，合作商家等級決定 $\text{Token}_{\text{value}}$ 可以折抵現金的多寡。

表 2 合作商家等級劃分範例

	折抵上限 Token_U	使用比例 W
一般商家	Token_{100}	10%
合作商家	Token_{200}	20%
重要商家	Token_{300}	30%
:	:	:
自營商家	Unlimited	100%

在代幣使用的智能合約中 $\text{SC}(M_j, U_i)$ ，代幣發行者有權利和義務與合作商家協議出對使用者最有利的折抵上限 U 值和使用比例 W 值，並將協議好的內容放至智能合約中，並事先編譯好 bytecode 放置區塊鏈中，如圖 4 所示。

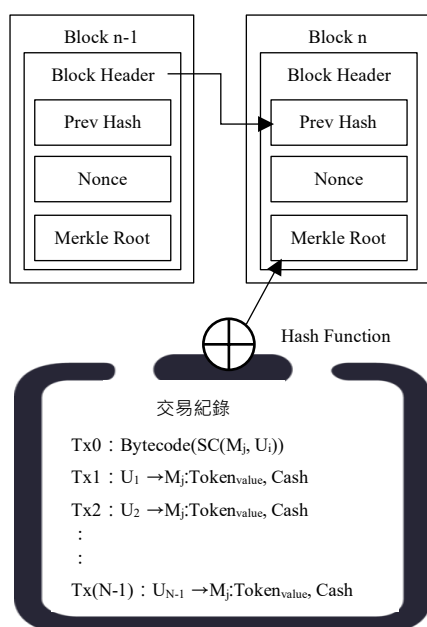


圖 4 代幣交易與智能合約

3.3 現金兌換

如果代幣發行者合作商家數越多，代表能使用代幣的機會就更多，更能創造出代幣的流通， Q 為兌換現金比率，可以由市場機制來決定，或由代幣發行者來制定，因此 $\text{Cash} = Q \cdot \text{Token}_{\text{value}}$ 。本架構考慮了以下幾種模式進行現金兌換：

- 使用者 U_i 對使用者 U_{i+1}

當使用者 U_i 想出售自己的 $\text{Token}_{\text{value}}$ 給使用者 U_{i+1} ，假如 U_{i+1} 同意此筆交易， U_{i+1} 將傳送自己的公鑰 e_{i+1} 給 U_i 當做交易位址， U_i 將回復自己的交易位址 e_i 與時間戳記 T 並透過 U_{i+1} 的公鑰 e_{i+1} 對 $\{T, \text{Sig}_{d_i}(\text{Token}_{\text{value}})\}$ 做加密傳送給 U_{i+1} ，當 U_{i+1} 收到 $E_{e_{i+1}}\{T, \text{Sig}_{d_i}(\text{Token}_{\text{value}})\}$ 後， U_{i+1} 將使用自己的私鑰 d_{i+1} 解開 $\{T, \text{Sig}_{d_i}(\text{Token}_{\text{value}})\}$ ，並使用 U_i 的公鑰驗證 $\text{Token}_{\text{value}}$ 是否來自 U_i 本人，當確定無誤後， U_{i+1} 將回傳 T ，

$E_{e_i}\{T, \text{Sig}_{d_{i+1}}(\text{Cash})\}$ 給 U_i ， U_i 使用自己的私鑰 d_i 解開 $\{T, \text{Sig}_{d_{i+1}}(\text{Cash})\}$ ，並使用 U_{i+1} 的公鑰 e_{i+1} 驗證 Cash 是否來自 U_{i+1} 本人，即完成現金兌換的交易。

- 使用者 U_i 對代幣發行者 S_k

當使用者 U_i 想兌換自己的 $\text{Token}_{\text{value}}$ 給代幣發行者 S_k ， S_k 將傳送自己的公鑰 e_k 給 U_i 當做交易位址， U_i 將回復自己的交易位址 e_i 與時間戳記 T 並透過 S_k 的公鑰 e_k 對 $\{T, \text{Sig}_{d_i}(\text{Token}_{\text{value}})\}$ 做加密傳送給 S_k ，當 S_k 收到 $E_{e_k}\{T, \text{Sig}_{d_i}(\text{Token}_{\text{value}})\}$ 後， S_k 將使用自己的私鑰 d_k 解開 $\{T, \text{Sig}_{d_i}(\text{Token}_{\text{value}})\}$ ，並使用 U_i 的公鑰驗證 $\text{Token}_{\text{value}}$ 是否來自 U_i 本人，當確定無誤後， S_k 將回傳 $T, E_{e_i}\{T, \text{Sig}_{d_k}(\text{Cash})\}$ 給 U_i ， U_i 使用自己的私鑰 d_i 解開 $\{T, \text{Sig}_{d_k}(\text{Cash})\}$ ，並使用 S_k 的公鑰 e_k 驗證 Cash 是否來自 S_k 本人，即完成現金兌換的交易。

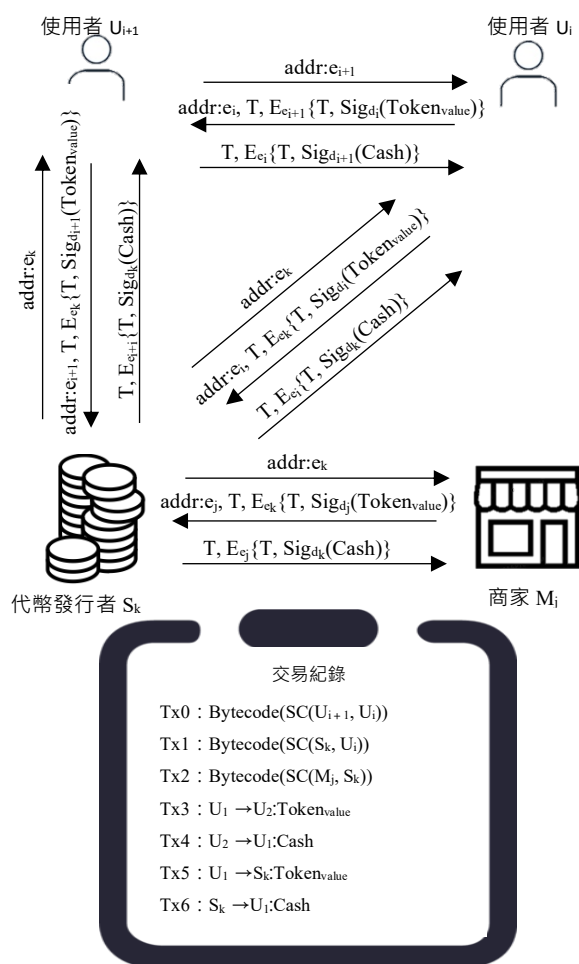


圖 5 現金兌換與智能合約

- 商家 M_j 對代幣發行者 S_k

當商家 M_j 想兌換自己的 $\text{Token}_{\text{value}}$ 給代幣發行者 S_k ， S_k 將傳送自己的公鑰 e_k 給 M_j 當做交易位址， M_j 將回復自己的交易位址 e_j 與時間戳記 T 並透過 S_k 的公鑰 e_k 對 $\{T, \text{Sig}_{d_j}(\text{Token}_{\text{value}})\}$ 做加密傳送給 S_k ，當 S_k 收到 $E_{e_k}\{T, \text{Sig}_{d_j}(\text{Token}_{\text{value}})\}$ 後， S_k 將使用自己的私鑰 d_k 解開 $\{T, \text{Sig}_{d_j}(\text{Token}_{\text{value}})\}$ ，並使用 M_j 的公鑰驗證 $\text{Token}_{\text{value}}$ 是否來自 M_j ，當確定無誤後， S_k 將回傳 $T, E_{e_j}\{T, \text{Sig}_{d_k}(\text{Cash})\}$ 給 M_j ， M_j 使用自己的私鑰 d_j 解開 $\{T, \text{Sig}_{d_k}(\text{Cash})\}$ ，並使用 S_k 的公鑰 e_k 驗證 Cash 是否來自 S_k ，即完成現金兌換的交易。

在現金兌換的智能合約中 $SC(U_{i+1}, U_i)$ 、 $SC(S_k, U_i)$ 、 $SC(M_j, S_k)$ 可分別針對不同兌換雙方定義不同的智能合約與兌換比率 Q 值，並將協議好的內容放至智能合約中，並事先編譯好 `bytecode` 放置區塊鏈中，如圖 5 所示。

3.4 建立聯盟鏈

各區域所管理的代幣發行者，可以彼此建立合作關係，將各自的私有鏈串聯產生聯盟鏈，聯盟鏈中的代幣發行者彼此也可以交換自己的 $Token_{value}$ ，各代幣發行者將可以分享合作商家的各項優惠與代幣折抵使用。

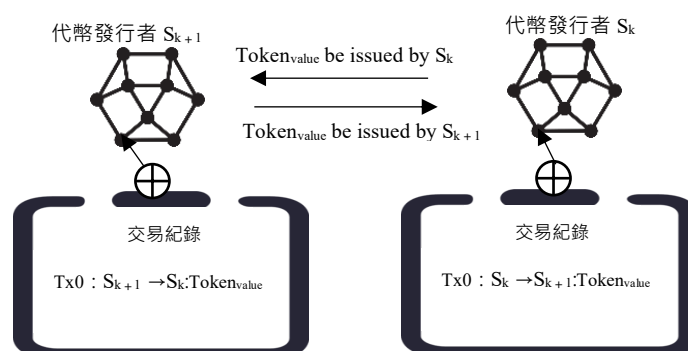


圖 6 聯盟鏈交易

當兩家代幣發行者建立聯盟關係，將各自交換自己的 $Token_{value}$ ，一但交易，並各自儲存在自己的私有鏈中，達成雙方不可否認的交易特性，如圖 6 所示。

4. 使用情境：以校園為例

代幣發行者可以配合相關活動承辦單位規劃免費獲取代幣的機制，鼓勵活動參與者的參與，創造更多的代幣流通，本論文將以一個校園生活為範例，規劃如何活絡代幣的使用與流通。此時校方即為代幣發行者，教師與主辦單位為分配者，學生為使用者。

4.1 賺取代幣

- 鼓勵學生註冊使用

學生一旦註冊使用，帳戶上就會有一個基本代幣金額，例如 $Token_{10}$

- 融入課程參與

學生每次上課簽到即可獲的 $Token_1$ ，老師可以根據學生課堂表現，給予額外的 $Token_{value}$ ，但只限本課堂的學生。

- 活動或研習參與

參與者每次活動簽到即可獲的 $Token_1$ ，演講者可以根據參與者課堂表現，給予額外的 $Token_{value}$ ，但只限本活動的參與者。

- 政府政策配合

由於學校不定期會有教育部或其他政府單位要求配合相關政策，且需要學生的參與和配合，主辦單位即可使用 $\text{Token}_{\text{value}}$ 來鼓勵學生參與，例如近日配合政府防疫措施，入校必須做自主健康管理，主辦單位可以針對每日有完成自主健康管理者給予 $\text{Token}_{\text{value}}$ ，相對也可減少每日需派遣相關人力留守校門口的人力成本。

- 校內措施配合

學校各單位每學期應有相關措施需要學生參與和配合，主辦單位即可使用 $\text{Token}_{\text{value}}$ 來鼓勵學生參與，例如每學期每堂課程於期末皆須完成總結性評量，以了解學生的學習成效，但每次所花的人力成本非常高，每一堂課程皆會使用一節課來施作問卷，不但浪費人力也影響老師授課時間，如果推行自行完成線上評量，學校可以給予 $\text{Token}_{\text{value}}$ 來鼓勵學生自發性參與。

- 獎勵方式

過去學校給予學生的獎勵方式多為獎狀或獎金，未來可以使用部分獎金，部分 $\text{Token}_{\text{value}}$ 混合方式給予獎勵。

4.2 使用代幣

- 合作商家商品使用

學生可以使用代幣至學校附近商家進行商品費用折抵，可以依據合作商家等級劃分表進行商品費用折抵參考，如表 2。

- 校內行政費用使用

學生可以使用代幣至學校相關行政單位進行使用，例如影印費用或圖書逾期罰款等。

- 校內設施租賃使用

學生可以使用代幣租賃學硬體設施，例如體育場所的環境使用、汽機車停車位的使用等。

- 折抵部分學雜費

學生可以使用代幣繳交部分學雜費，此部分可以配合銀行繳費端進行拆帳。

4.3 兌換現金

- 學生對學生兌換

學生可以透過智能合約，將想要的兌換的 $\text{Token}_{\text{value}}$ 與想要的現金放至智能合約上，給其他想要兌換的學生兌換。

- 學生對學校兌換

學校定期發布兌換訊息於智能合約上，提供給有想兌換的學生兌換。

- 商家對學校兌換

學校與合作商家於當初簽訂合作時，就將相關合約發布於智能合約上，當滿足智能合約上的所有條件時，商家就必須給予 $\text{Token}_{\text{value}}$ ，學校也必須給予現金至商家指定帳戶。

5. 安全性與成本分析

5.1 安全性分析

本論文一個以私有鏈的區塊鏈架構建立安全及可信任的區域市場可流通的代幣交易制度，以確保金鑰是安全的，不會被竊取。在區塊鏈架構下排出了公正的第三方，各自保管自己的公鑰與私鑰對(e, d)，所有的交易與簽章是被(e, d)加密保護的。

此架構也保證任何之惡意竊聽者無法在網路上取得之相同資訊後，重新傳送後並獲得驗證通過。在所有的交易訊息傳送過程都使用時間戳記 T，其中時間戳記 T 即是為避免重送攻擊，只有接收者驗證來自發送者的 T 在一定時間範圍內才是合法發送者。

在區塊鏈中的資訊是不可篡改的，一旦資料資訊被驗證通過寫入區塊中並加入區塊鏈，就無法被篡改。區塊鏈的資料資訊必須經過網路大部分節點的審核以後，才能允許被記錄。除非能夠控制系統中 51% 以上網路節點，否則對單一節點的區塊記錄篡改是沒有意義的，這種不可篡改特性保證了區塊鏈資料的穩定性與可靠性。

5.2 成本分析

假設每位學生註冊後，學生可得到 Token₁₀，如果以一間 1 萬名學生的學校，每學期 200 門實體課程或活動，每門課程給與授課教師 Token₁₀₀ 做使用，假設商家與學校協議的 Q 值介於為 0.5 至 1 之間，在此課程規模環境下，學校最多只需投入 6 萬至 12 萬成本作為此代幣交易的資金(Token₁₀ × 1 萬 + Token₁₀₀ × 200 = Token₁₂ 萬)。

使用不到 12 萬的成本，可以促進周邊商家的市場消費，增加學生出席狀況、已可促進學生參與學校政策措施的與提高對學校的配合程度。

6. 結論

本論文提出以區塊鏈和智能合約架構下，建立一個以安全及可信任的區域市場可流通的代幣交易制度，其特性為安全及可信任、區域市場交易與可跨區域流通的代幣交易制度，以解決目前各區域市場各自發行的代幣無法流通，透過本論文的架構下，未來可以建跨區域的聯盟鏈，以活絡代幣的流通性，促進與振興國內市場消費。

參考文獻

- [1] Andreas M. Antonopoulos. 2014. *Mastering Bitcoin: unlocking digital cryptocurrencies*. O'Reilly Media, Inc.
- [2] Chris Dannen. 2017. *Introducing Ethereum and solidity*. Springer.
- [3] Keke Gai, Meikang Qiu, and Xiaotong Sun. 2018. A survey on FinTech. *Journal of Network and Computer Applications* 103, (2018), 262–273.
- [4] In Lee and Yong Jae Shin. 2018. Fintech: Ecosystem, business models, investment decisions, and challenges. *Business Horizons* 61, 1 (2018), 35–46.
- [5] Satoshi Nakamoto. 2019. *Bitcoin: A peer-to-peer electronic cash system*. Manubot.
- [6] Bernardo Nicoletti, Nicoletti, and Weis. 2017. *Future of FinTech*. Springer.
- [7] Thomas Philippon. 2016. *The fintech opportunity*. National Bureau of Economic Research.
- [8] Rob Pike. 2012. Go at google: Language design in the service of software engineering. *SPLASH* (2012).
- [9] Thomas Puschmann. 2017. Fintech. *Business & Information Systems Engineering* 59, 1 (2017), 69–76.
- [10] Nick Szabo. 1997. Formalizing and securing relationships on public networks. *First Monday* (1997).
- [11] Gavin Wood. 2014. Ethereum: A secure decentralised generalised transaction ledger. *Ethereum project yellow paper* 151, 2014 (2014), 1–32.

- [12] 竹山數位鎮民計畫. Retrieved from <http://www.findglocal.com/TW/Taichung/146451755369062/天空的院子>
- [13] MCoin. Retrieved from <http://web.fintech.mcu.edu.tw/zh-hant/content/銘傳幣>
- [14] 東海酷幣. Retrieved from <http://csld.thu.edu.tw/web/page/page.php?scid=54&sid=59>
- [15] 南臺幣. Retrieved from <https://news.stust.edu.tw/pid/3517>

學院

產業

有效運用區塊鏈技術於健康醫療照護之研究

沈容

國立臺北商業大學財務金融系

摘要

隨著資訊科技的進步以及資料量的急速增加，將大數據運用於健康照護系統的研究，是近年來各國積極發展的策略，醫院將原本記錄在紙本病歷上的病歷紀錄，透過電腦系統與網路技術，存放於醫院中的病歷資料庫中，透過電子病歷系統與臨床資料庫系統，醫師可以更快速的取得病患的病歷資料，而醫院也可以經由系統電子化減少許多調閱病歷所需要的時間和人力。本研究將提出一個 HealthCoin 的交易點數，提供醫院、醫生與病人彼此交換的點數，以利後續做交換紀錄的憑證。本研究所提的健康區塊鏈架構只要有三種成員，每個成員在健康區塊鏈架構中都是一個節點，然而病歷涉及病人個人隱私，醫療院所不僅要具備電子病歷交換機制，還必須更加重視資通安全與個人資料保護。推動電子病歷的權責單位事出多頭，在缺乏資源配置的管理及協調機制下，欠缺橫向協調整合，不但導致計畫推動容易出現本位主義而各行其是，難以呈現分工合作之綜效。也由於缺乏執行統合機制，即使是目前較有成效的醫院部分，也出現執行進度差異頗巨的問題。推動電子病歷交換，攸關智慧醫療應用的推動，因此本計畫將以健康區塊鏈技術解決第三方公正的問題，藉由 HealthCoin 的交易點數彼此對健康存摺作交易。

關鍵詞：電子病歷；區塊鏈；HealthCoin

1.緒論

隨著資訊科技的進步以及資料量的急速增加，將大數據運用於健康照護系統的研究，是近年來各國積極發展的策略，醫院將原本記錄在紙本病歷上的病歷紀錄，透過電腦系統與網路技術，存放於醫院中的病歷資料庫中，透過電子病歷系統與臨床資料庫系統，醫師可以更快速的取得病患的病歷資料，而醫院也可以經由系統電子化減少許多調閱病歷所需要的時間和人力[1 - 4]。台灣自 2002 年已開始推動「醫療院所病歷電子化試辦計畫」，衛福部於 2003 年開始推行醫療院所病歷電子化試辦計畫。電子病歷可靠且持久，利於病歷隨時取用，因此可提高醫師的診斷決策，從而改善護理品質，並提高病患的安全。目前世界各國都將電子病歷視為醫療服務的重要政策，如美國配合全民健保政策推動 Meaningful Use[5]，中國大陸推動的電子病歷系統功能應用水準分級評價方法及標準[6]。由於電子病歷有助於提升醫療品質及降低醫療成本。因此，世界各國均極力推行電子病歷(Healthcare Information Management Systems Society, HIMSS)。目前實施電子病歷的國家，包括：北歐各國、荷蘭、澳洲、紐西蘭、日本、韓國、部份東南亞國家等。

2007 年推動「國民健康資訊基礎建設」，主要在促進病歷電子化及建立醫事憑證管理中心，2009 年提出「智慧醫療服務計畫」，包括遠距健康照護、推動電子病歷及醫療影像傳輸計畫、健保 IC 卡改善計畫、醫院安全關懷 RFID 計畫與健康資料庫加值應用服務等項目[7], [8]。台灣自 2009 年起，便已著手建立電子病歷交換中心，發展至今已有相當成效。根據衛生福利部電子病歷交換中心提供的資料，目前已完成電子病歷交換建置的醫院共計 276 家。至 2012 年底，全台 500 多所醫院，其實只有 142 家連結上電子病歷交換中心，卻能在短短一年之內，迅速增加近一倍的數量，健康資訊交換第七層協定(Health Level Seven；HL7)的普及，扮演相當重要的角色。至 2015 年全台約有八成醫院可提供電子病歷交換，而各醫院之間可交換互通之電子病歷項目有：門診病歷、門診用藥、血液檢查、醫療影像及報告、出院病摘等內容[9]。而在全民健保的資料運用上，「健保醫療資訊雲端查詢系統」運用於臨床的診間服務，診間醫師可查詢病患最近三個月內至各門診、住院及藥局所彙集整理成的即時就醫紀錄。

由於病歷涉及病人個人隱私，甚至在全球化發展導致傳染病管理的複雜度日益增加的情況下，病歷資料甚至可能會涉及國安問題，醫療院所不僅要具備電子病歷交換機制，還必須更加重視資通安全與個人資料保護。衛福部針對病歷資料交換規劃的資安機制，首先是在健保 VPN 封閉式網路下才能使用，減少與外界接觸的可能，而在使用者認證方面，將採「雙卡認證」，必須同時通過病患健保 IC 卡與醫事人員卡，方能交換病歷；同時還要保存病患簽署的紙本同意書，以及留存 Log 紀錄，任何查看、更改病歷的行為，都會被記錄並永久保存。至於影像交換中心，基本上只有索引資料庫，不會保留任何病歷資料，自然沒有洩露病患隱私的疑慮。然而要讓電子病歷交換機制正式上線，還需要全國多家衛生所及診所全部加入，而全國已連結上電子病歷交換中心的醫院，實際交換的病歷數量也相當有限，不足以呈現電子病歷交換應有的效果。

電子病歷交換上路困難，主要問題其實與技術關係不大。推動電子病歷的權責單位事出多頭，在缺乏資源配置的管理及協調機制下，欠缺橫向協調整合，不但導致計畫推動容易出現本位主義而各行其是，難以呈現分工合作之綜效。也由於缺乏執行統合機制，即使是目前較有成效的醫院部分，也出現執行進度差異頗巨的問題，如署立醫院及大型醫學中心，幾已全數實施電子病歷，但其他中小型醫院的實施成效就相當有限，也間接影響診所的電子病歷建置工作。推動電子病歷交換，攸關智慧醫療應用的推動，因此本計畫將以健康區塊鏈技術解決第三方公正的問題，藉由 HealthCoin 的交易點數彼此對健康存摺作交易。

2.文獻探討

2.1 健康存摺

健康存摺系統為一線上健康資料查詢系統，提供健保保險對象可隨時隨地便利地查詢個人的健康資料，掌握健康大小事、做好自我健康管理！也可以在就醫時，提供醫師參考，幫助醫師快速掌握個人健康狀況，提升醫療照護安全與品質。健康存摺系統於 103 年 9 月 25 日在健保署官網上線，只要持自然人憑證或已註冊密碼之健保卡，通過身分驗證，即可於「健康存摺」查詢或下載個人的健康資料。自 105 年 7 月起，為提升服務效能，健康存摺系統進行全面改版，包括：資料視覺化呈現、增加外部衛教連結、擴大資料提供範圍，以及提升「健康存摺」取得可近性及使用便利性等功能，希望能讓「健康存摺」為民眾發揮最佳功能[10]。

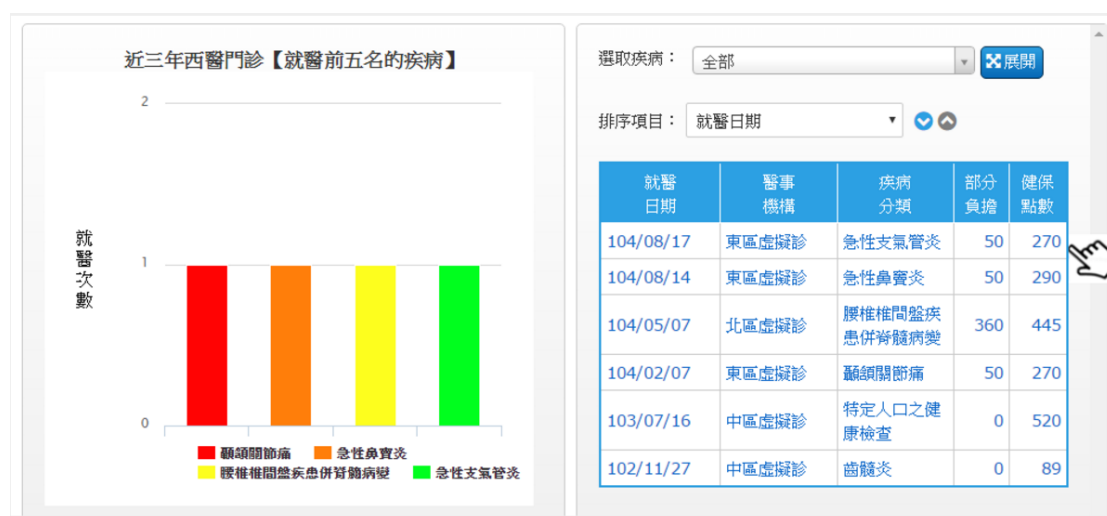


圖 1 健康存摺資料視覺化呈現

2.2 區塊鏈技術

區塊鏈技術是一種對一段時間內所有交易或者網路行為進行記錄的分散式資料記錄技術，即一種按照時間順序將資料區塊以順序相連的方式組合成的 chain 資料結構，並以密碼學方式保證的不可篡改和不可偽造的分散式公開帳本技術。在區塊鏈上，每筆交易都可以被系統的參與者通過多數節點共識的機制進行審核。一旦被記載到區塊鏈上，相關的交易資訊就不能被修改、刪除。區塊鏈上的區塊記錄了一定時間內被審核通過的每筆交易記錄。比特幣是區塊鏈技術最本質的應用，它也因為在沒有政府的干預下，達成了上百萬美元的匿名交易市場，而成為最具爭議性的區塊鏈應用。然而，區塊鏈技術本身可以被金融和非金融領域應用。傳統的網路金融交易是建立在可靠的、可信任的公正第三方的基礎上。使用者網路交易必須依賴於公正第三方對交易雙方進行身份和交易資格的確認。人們生活的網路世界在依靠公正第三方維繫安全與個人網路資產隱私的背景中，並不是絕對的安全。因為公正第三方這種背景下扮演了強有力的中心角色。而中心機構一旦被侵入，將無法保證以公正第三方所維護安全風險。在這種背景下，區塊鏈的誕生通過對分散式共識的應用，對依靠公正第三方信任的網路世界規則提出了巨大改進。分散式共識和匿名性是其最重要的兩個關鍵特徵。通過這些特徵的綜合

應用，在複雜的網路環境中，區塊鏈在沒有公正第三方信任機構的參與下，對網路交易記錄、網路資產記載提出了分散式帳本的方式，有以下的特徵。

(1)去中心化

去中心化是區塊鏈技術最基本的特性。區塊鏈技術的即在沒有中央處理節點的情況下，達成了所有資料的分散式記錄、儲存並且能夠保證資料記錄的真實性。區塊鏈技術是由點對點(P2P)的網路組成。不同於中心化網路模式，P2P 網路中各節點的地位平等，每個節點有相同的網路權力。在這種去中心化的網路環境中，所有網路節點都是相同對等的，所有節點享有相同的權利和義務。區塊鏈網路中的在網路節點必須遵守同樣的密碼學規則，共同維護系統中的資料記錄。對資料的記錄、儲存過程，必須得到區塊鏈網路內其他節點的批准後才能執行。由於所有在網節點並沒有公正第三方或者信任機構背書，所以在去中心化的區塊鏈網路中，並無單一節點的攻擊而對整個區塊鏈網路產生影響。

(2)資料及操作的透明性

區塊鏈技術作為分散式帳本技術，系統內所有的資料記錄及操作對於所有在網路節點都是透明的。在典型的區塊鏈網路中，每一個節點都能夠儲存所有發生的歷史交易記錄的完整與一致帳本。區塊鏈通過對非對稱加密演算法、散列加密等密碼學技術的組合應用，保證區塊鏈資訊在網路的高度透明性。並且區塊鏈網路運行的程式、規則、節點的接入方式都是公開的，這是區塊鏈網路信任的基礎。這些機制的運用，保證了區塊鏈中記錄的資料可以被全網所有節點審查、追蹤。

(3)資訊不可篡改性

區塊鏈區塊中的資訊是不可篡改的，一旦資料資訊被驗證通過寫入區塊並加入區塊鏈中，就無法被篡改。區塊鏈的資料資訊必須經過網路大部分節點的審核以後，才能允許被記錄。除非能夠控制系統中 51% 以上網路節點，否則對單一節點的區塊記錄篡改是沒有意義的。即對個別節點的帳本資料的篡改、攻擊不會影響網路總帳的安全性。這種資訊的不可篡改性保證了區塊鏈資料的穩定性與可靠性。

(4)匿名性

區塊鏈技術在複雜的網路環境中解決了在節點間的信任問題，因而區塊鏈網路中的交易節點可以在無需瞭解對方身份的情況下進行交易。區塊鏈網路中的交易是基於加密位址，而不會對交易雙方身份進行認證。交易雙方僅需要公佈自己的位址就可以與對方進行交易通信。這種匿名性的技術基礎就是非對稱加密演算法。區塊鏈網路的節點使用非對稱加密技術，構建節點間在匿名環境下的信任。所有節點維持自身的公私密金鑰對，對區塊鏈網路節點間的通信資訊進行加密和解密。節點公開發佈自己的公開金鑰，保留自己的私密金鑰。進行資訊傳遞的發送方，使用資訊接收方公佈的公開金鑰對將要傳遞的資訊進行加密。資訊接收方在接收到傳遞的加密資訊後，使用自己的私密金鑰對加密過的資訊進行解密。通過這樣的方式，節點間可以在不需要身份認證的情況下，完成匿名環境下的信任交易[11], [12]。

3. 研究架構

本研究將提出一個 HealthCoin 的交易點數，提供醫院、醫生與病人彼此交換的點數，以利後續做交換紀錄的憑證。本研究所提的健康區塊鏈架構只要有三種成員，每個成員在健康區塊鏈架構中都是一個節點：

□醫院：主要儲存病人健康存摺的地方，醫院之間彼此可以對病人健康存摺作交易，在此我們假設每筆的 HealthCoin 為 1，只有醫院可以做批次交易。

◇醫生：醫生為了了解病人的相關醫療或生理資訊，可以對醫院提出健康存摺交易，醫生之間彼此可以對病人的健康存摺作交易，在此我們假設每筆的 HealthCoin 為 1

○病人：病人為了知道自己目前的醫療或生理資訊，可以提出對醫院或醫生提出健康存摺交易，因為病人有權了解自己的健康資訊，在此我們假設所需的 HealthCoin 為 0。

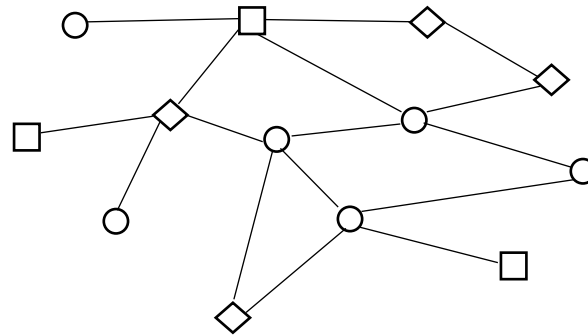


圖 2 交易網路

3.1 交易方式

1. 醫院 A 對醫院 B 取得 n 筆健康存摺

Step 1: 醫院 B 產生一組公開金鑰對，公鑰為 $P_{\text{醫院 B}}$ ，私鑰為 $S_{\text{醫院 B}}$ ，醫院 B 以 $P_{\text{醫院 B}}$ 為位址給醫院 A。

Step 2: 醫院 A 產生一組公開金鑰對，公鑰為 $P_{\text{醫院 A}}$ ，私鑰為 $S_{\text{醫院 A}}$ ，醫院 A 用 $S_{\text{醫院 A}}$ 簽章 $\text{HealthCoin}=n$ 傳給位址 $P_{\text{醫院 B}}$ 。

Step 3: 醫院 B 使用醫院 A 的 $P_{\text{醫院 A}}$ 加密 n 筆健康存摺並把這次的交易傳送到 Health P2P 網路，讓所有人都知道此筆交易紀錄。

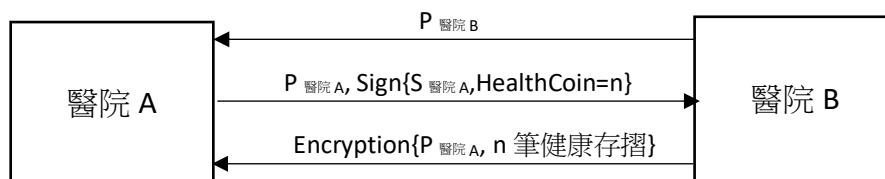


圖 3 醫院之間的交易

2. 醫生 A 對醫生 B 取得 1 筆健康存摺

Step 1: 醫生 B 產生一組公開金鑰對，公鑰為 $P_{\text{醫生 B}}$ ，私鑰為 $S_{\text{醫生 B}}$ ，醫生 B 以 $P_{\text{醫生 B}}$ 為位址給醫生 A。

Step 2: 醫生 A 產生一組公開金鑰對，公鑰為 $P_{\text{醫生 A}}$ ，私鑰為 $S_{\text{醫生 A}}$ ，醫生 A 用 $S_{\text{醫生 A}}$ 簽章 HealthCoin=1 傳給位址 $P_{\text{醫生 B}}$ 。

Step 3: 醫生 B 使用醫生 A 的 $P_{\text{醫生 A}}$ 加密 1 筆健康存摺並把這次的交易傳送到 Health P2P 網路，讓所有人都知道此筆交易紀錄。

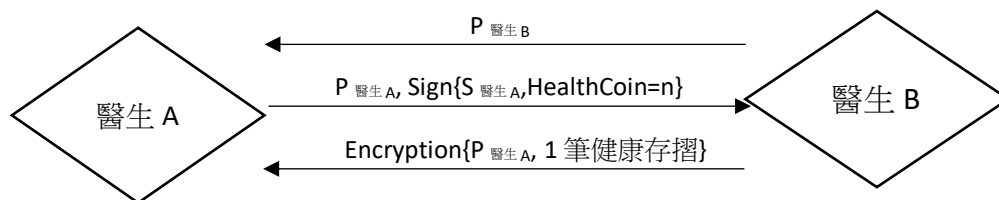


圖 4 醫生之間的交易

3. 病人 A 對醫院 B 取得健康存摺

Step 1: 醫院 B 產生一組公開金鑰對，公鑰為 $P_{\text{醫院 B}}$ ，私鑰為 $S_{\text{醫院 B}}$ ，醫院 B 以 $P_{\text{醫院 B}}$ 為位址給醫院 A。

Step 2: 病人 A 產生一組公開金鑰對，公鑰為 $P_{\text{病人 A}}$ ，私鑰為 $S_{\text{病人 A}}$ ，病人 A 用 $S_{\text{病人 A}}$ 簽章 HealthCoin=0 傳給位址 $P_{\text{醫院 B}}$ 。

Step 3: 醫院 B 使用病人 A 的 $P_{\text{病人 A}}$ 加密 1 筆健康存摺並把這次的交易傳送到 Health P2P 網路，讓所有人都知道此筆交易紀錄。

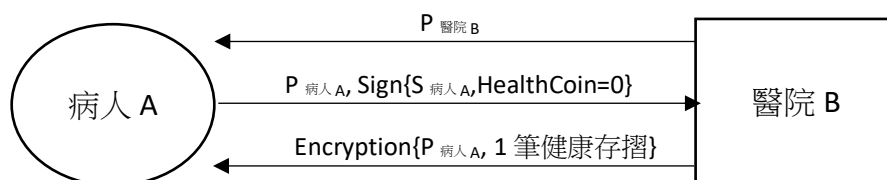


圖 5 病人與醫院之間的交易

4. 病人 A 對醫生 B 取得健康存摺

Step 1: 醫生 B 產生一組公開金鑰對，公鑰為 $P_{\text{醫生 B}}$ ，私鑰為 $S_{\text{醫生 B}}$ ，醫生 B 以 $P_{\text{醫生 B}}$ 為位址給病人 A。

Step 2: 病人 A 產生一組公開金鑰對，公鑰為 $P_{\text{病人 A}}$ ，私鑰為 $S_{\text{病人 A}}$ ，病人 A 用 $S_{\text{病人 A}}$ 簽章 HealthCoin=0 傳給位址 $P_{\text{醫生 B}}$ 。

Step 3: 醫生 B 使用病人 A 的 $P_{\text{病人 A}}$ 加密 1 筆健康存摺並把這次的交易傳送到 Health P2P 網路，讓所有人都知道此筆交易紀錄。

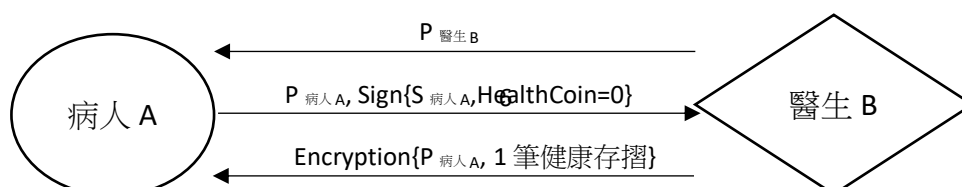


圖 6 病人與醫生之間的交易

3.2 健康區塊鏈

如上的交易方式可以組成一個區塊，當完成一個區塊後，就會增加到健康區塊鏈，每個區塊中的交易會產生一個 HASH 雜湊值，每個區塊中的 Prev. HASH 是前一個區塊的雜湊值，透過 Prev. HASH 連結的方式，構成本研究的健康區塊鏈。

3.3 驗證追蹤

本研究假設目前健康區塊鏈中流通的 HealthCoin 為 n ， n 值由衛生署根據目前加入健康區塊鏈中的醫療院數，醫療人員數來決定，整個健康區塊鏈即為紀錄 HealthCoin 的公開交易帳簿，使用 HealthCoin 進行的所有交易，都會記錄在這個全世界唯一的公開交易帳簿內，因此只要透過這個公開交易帳簿，就可以將所有可使用的 HealthCoin 的位址、所有的交易記錄下來，就可以決定哪一個位址現在可以交易多少個 HealthCoin，為了達到這個驗證目的，公開交易帳簿就必須形成健康區塊鏈。

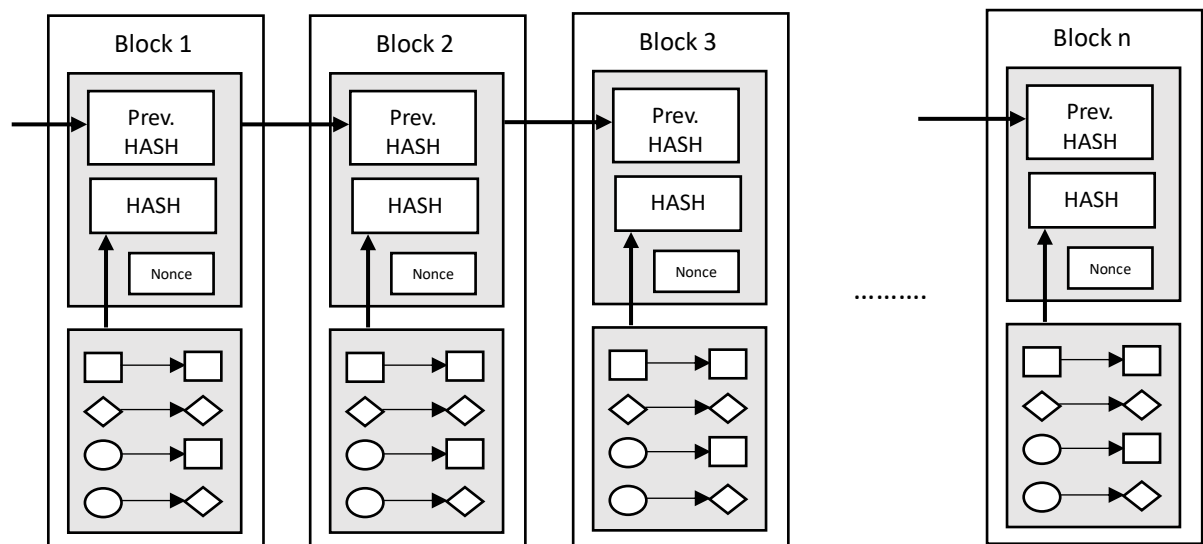


圖 7 健康區塊鏈

如圖 7 所示，一個區塊是由交易的集合與表投所形成，表投中儲存了前一個區塊表頭的 HASH 雜湊值，假設記錄在 Block 2 內的交易有其中 1 為源遭竄改，這樣

就得換成正確的 HASH 雜湊值，可是 Block 2 的內容改變，Block 3 儲存的 HASH 雜湊值也必須更動，因此只要稍微更改健康區塊鏈，後面的區塊全部都得改變才行，區塊表頭的兩個雜湊值，增加了竄改區塊鏈的難度，能避免遭到竄改。

4. 預期成果

本計畫預期解決病人個人隱私與醫療院所電子病歷交換機制與解決電子病歷的橫向協調整合問題，資源配置有效管理。推動電子病歷交換，攸關智，慧醫療應用的推動，因此本計畫將以健康區塊鏈技術解決第三方公正的問題，藉由 HealthCoin 的交易點數彼此對健康存摺作交易。並重視資通安全與個人資料保護。並符合衛福部針對病歷資料交換規劃的資安機制，在健保 VPN 封閉式網路下才能使用，減少與外界接觸的可能，而在使用者認證方面，將採「雙卡認證」，必須同時通過病患健保 IC 卡與醫事人員卡，方能交換病歷；同時還要保存病患簽署的紙本同意書，以及留存 Log 紀錄，任何查看、更改病歷的行為，都會被記錄並永久保存。至於影像交換中心，基本上只有索引資料庫，不會保留任何病歷資料，自然沒有洩露病患隱私的疑慮。然而要讓電子病歷交換機制正式上線，還需要全國多家衛生所及診所全部加入，而全國已連結上電子病歷交換中心的醫院，實際交換的病歷數量也相當有限，不足以呈現電子病歷交換應有的效果。

參考文獻

- [1] I. Iakovidis, “Towards personal health record: current situation, obstacles and trends in implementation of electronic healthcare record in Europe,” *Int. J. Med. Inf.*, vol. 52, no. 1, pp. 105 – 115, 1998.
- [2] D. Delamarre et al., “CARDIOMEDIA: a communicable multimedia medical record on Intranet and digital optical memory card,” *Int. J. Med. Inf.*, vol. 55, no. 3, pp. 211 – 222, 1999.
- [3] A. Rafiq, X. Zhao, S. Cone, and R. Merrell, “Electronic multimedia data management for remote population in Ecuador,” in *International Congress Series*, 2004, vol. 1268, pp. 301 – 306.
- [4] R. A. Greenes, M. Collen, and R. H. Shannon, “Functional requirements as an integral part of the design and development process: summary and recommendations,” *Int. J. Biomed. Comput.*, vol. 34, no. 1 – 4, pp. 59 – 76, 1994.
- [5] D. Blumenthal and M. Tavenner, “The ‘meaningful use’ regulation for electronic health records,” *N Engl J Med*, vol. 2010, no. 363, pp. 501 – 504, 2010.
- [6] 李华才, “推进电子病历系统深入应用的重要举措,” *中国数字医学*, no. 2012 年 05, pp. 1 – 1, 2012.
- [7] 廖淑君, “從英國 NHS 國家 IT 計畫看電子病歷之推動: 以病患個人資訊隱私保護為中心,” *科技法律透析*, vol. 25, no. 5, pp. 25 – 49, 2013.

- [8] 衛生福利部, 台灣健康雲-105 年度綱要計畫書 (A1006), 審議編號: 105-0324-07-12-01, 計畫全程. 103 年 1 月至 105 年 12 月.
- [9] 衛生署: 電子病歷推動專區. <http://emr.mohw.gov.tw/>.
- [10] 健康存摺. <https://myhealthbank.nhi.gov.tw/IHKE0002/IHKE0002S07.aspx>.
- [11] S. Nakamoto, Bitcoin: A peer-to-peer electronic cash system. 2008.
- [12] L. I. Dong, W. E. I. Jinwu, and others, “區塊鏈技術原理, 應用領域及挑戰,” 電信科學, vol. 32, no. 12, pp. 20 – 26, 2017.

學院

產業

1. 計劃主持人取得 CEH 認證
2. 計劃主持人取得 CEI 認證
3. 計劃主持人與丁俞豪(Yu-Hao Ting)同學共同發表一篇研討會論文，Session2-7，編號 IC-09，<https://www.fronticomp.com/ic2025>
4. 吳若綺同學完成道德駭客實務入門及 CEH 認證班 35 小時
5. 丁俞豪同學完成成為白帽駭客 2 滲透測試實務課程 18 小時





吳若綺

成為白帽駭客1 | 道德駭客實務入門及CEH認證班

培訓日期: 2024/11/24 — 2024/12/22

訓練時數共計 35 小時

wriedu

證書編號: WW24TBAZ15865

發證日期: 2024/12/22

完訓證書

Certificate of Completion

緯育 TibaMe

丁俞豪

成為白帽駭客2 | 跟著方丈學滲透測試實務

培訓日期: 2025/02/09 — 2025/02/22

訓練時數共計 18 小時

wirededu

證書編號: WW25TBAZ01143

發證日期: 2025/02/22

學院