

ChatGPT 模擬 OSCE 護病溝通中學習感知與溝通表現之性別差異與相關性分析

Gender Differences and Correlations Between Learning Perceptions and Communication

Performance in ChatGPT-Simulated OSCE Nurse–Patient Communication

楊馥宇^{1*}, 吳美玲²

¹長庚科技大學臨床技能中心

²長庚科技大學護理系

* fyyang@mail.cgust.edu.tw

【摘要】本研究探討護生性別在 ChatGPT 模擬 OSCE 護病溝通之學習感知與溝通表現之差異與關聯。28 位護生（男 12，女 16）參與 ChatGPT 模擬問診與衛教，完成學習感知問卷，依 OSCE 評量指標評分溝通。採用 Mann-Whitney U 檢定性別差異，Spearman 檢定變項相關，並以 Z 分數比較學生學習感知與溝通表現落差。結果顯示男護生學習感知顯著高於女護生，但學習感知與溝通表現無顯著相關，男護生傾向高估自身表現，女護生則相對一致，顯示性別可能影響 AI 模擬學習感知與評估一致性；對話數與溝通得分具正相關，無性別差異。未來生成式 AI 輔助學習應納入性別與對話參與指標，確保教學回饋與評估準確性。

【關鍵字】護病溝通；學習感知；客觀結構式臨床評量；性別差異；ChatGPT

Abstract: This exploratory study examined gender differences in nursing students' learning perceptions and communication performance during ChatGPT-simulated OSCE communication. Twenty-eight students (12 male, 16 female) participated, completing OSCE evaluations and perception questionnaires. Mann–Whitney U tests showed that male students had significantly higher perception scores than female students ($p < 0.05$), but Spearman's rho found no correlation between perceptions and performance. Z-scores revealed that male students tended to overestimate their performance, whereas female students' learning perceptions were more aligned with their actual scores. Additionally, the number of dialogue turns was positively correlated with communication scores, though neither variable showed significant gender differences. Despite the small sample size, these findings suggest that gender may influence perceptions in AI-assisted learning environments, supporting the inclusion of gender and engagement metrics in AI educational design.

Keywords: Nurse–Patient Communication, Learning Perception, OSCE, Gender Differences, ChatGPT

1.前言

護病溝通能力被廣泛認為是影響醫療照護品質的關鍵要素。具備熟練溝通技巧的護理人員，能更有效地提供醫療服務，減輕患者的焦慮、疼痛與疾病症狀，並於護理過程中建立患者對醫療體系的信任(Siokal et al., 2023)。因此，護理人員需持續精進其溝通能力，根據患者的年齡、情緒狀態、生理條件、語言與文化背景、社經地位及溝通環境等因素，調整溝通策略，採取同理、專業與個別化的訊息傳遞方式，進而提升患者對醫療資訊的理解與接受度，

增進其滿意度(Naz et al., 2024 ; Siokal et al., 2023)、治療依從性(Kerr et al., 2022), 並改善生理與功能狀況(Nummer et al., 2024)。相對地, 若溝通技巧不足, 將可能增加醫療錯誤風險, 對病人造成負面後果。

然而, 實務中護病溝通常面臨多重障礙, 其原因包括護理人員溝通訓練不足、專業知識不全、同理心與自信心缺乏、對患者態度消極, 以及語言與文化隔閡等因素, 皆可能影響其傾聽與回應能力, 進而削弱溝通成效(Alshalawi et al., 2025)。因此, 護理教育機構有責任在教育階段強化學生之溝通技巧與態度, 協助其建立積極正向的護病關係, 以因應臨床現場之實際需求。

近年研究亦指出, 護病溝通技巧可能存在性別差異。護理人員的性別會影響其與患者之間的互動方式, 這可能與性別刻板印象、溝通風格差異或文化期望有關(Alharazi et al., 2025)。Graf 等人(2017)指出, 女性護理系學生在溝通技巧表現上顯著優於男性護理系學生, 女性對溝通訓練態度更積極, 且溝通方式較具體; 此外, 女性在解釋與澄清、提問技巧, 以及正向回饋等面向的得分也顯著較高(Kalati et al., 2025)。然而, 性別是否必然影響溝通表現, 尤其是在人工智慧模擬 OSCE (客觀結構式臨床評量) 情境中, 仍有待進一步驗證。

一般而言, 護病溝通可區分為語言、非語言與副語言三種形式。本研究聚焦於語言溝通的表現, 亦即護理學生在 OSCE 評量中與 ChatGPT 模擬病人互動時, 與任務相關的語言對話次數與得分。需要注意的是, 若對話內容未涵蓋評分指標中所規範之任務項目, 則即使互動頻繁, 亦不代表其溝通表現良好。反之, 若對話內容過度集中於特定任務, 忽略其他評分項目, 也可能導致整體分數偏低。換言之, 對話次數與溝通得分並不一定呈現正向關係, 其間是否存在性別差異, 亦值得探討。

目前, 多數台灣護理教育機構採用 Harden 與 Gleeson (1979) 所提出的 OSCE 模式, 作為評估學生臨床技能與畢業門檻的依據。OSCE 強調標準化病人、情境模擬與結構化評量, 旨在提升護學生的臨床信心與專業能力(Onwudiegwu, 2018)。惟傳統 OSCE 仰賴人工病人與大量人力資源, 成本與訓練時間高昂。隨著生成式人工智慧(Generative AI)技術快速發展, 包含 ChatGPT 在內的大型語言模型(LLMs)逐漸被導入醫護臨床與教育場域中, 作為教學與學習的輔助工具(Pan & Ni, 2024)。

ChatGPT 具備自然語言互動、即時回饋、低成本與高度可調整等優勢, 能有效支援護理學生進行臨床情境模擬訓練, 包括病史採集、病歷記錄、問診、衛教、交班與臨床決策等(Amin & Alanzi, 2024 ; Saad et al., 2025 ; Huang et al., 2024 ; Tsang, 2023)。於 OSCE 應用層面, 亦可協助教師快速設計案例、標準化情境與評分機制, 有效降低評估成本(Misra & Suresh, 2024)。透過 ChatGPT 取代標準化病人, 可讓學生進行反覆練習, 緩解考試壓力, 並可即時提供溝通表現的評估與回饋, 進而提升其自我反思與溝通能力。然而, 目前針對 AI 模擬 OSCE 在實際溝通訓練中之成效與性別影響, 仍缺乏實證研究。

另有研究指出, 性別可能影響學習者的溝通風格、科技接受度與自我效能評估。性別不僅可能影響對科技工具的接受態度, 也會形塑其學習經驗與知覺學習成效。Pan 與 Ni (2024) 調查來自五個醫護相關專業、不同學年的 1,718 名學生, 結果發現男性對大型語言模型的理解與信任度顯著高於女性, 並較可能認為其益處大於風險。此一差異亦可能影響護理學生於 ChatGPT 模擬 OSCE 中的態度與自我評估, 導致學習感知與實際表現之間出現落差。然而, 目前針對此類學習模式中性別差異所造成之主觀與客觀表現落差之探討, 仍付之闕如。

綜上所述, 儘管現有文獻已初步揭示 ChatGPT 於模擬 OSCE 護病溝通訓練的應用潛力, 惟針對性別差異對於學習感知與實際溝通表現間關係之影響, 實證研究仍屬有限。有鑑於此,

本研究擬進一步探究性別對 AI 模擬 OSCE 護病溝通學習歷程之影響，並聚焦於學習感知、互動頻率與溝通表現之間的關聯性，本研究欲探討以下三項研究問題：

一、不同性別的護理學生在 AI 模擬 OSCE 護病溝通中，於 OSCE 溝通得分、對話次數與學習感知是否存在顯著差異？

二、OSCE 溝通得分、對話次數與學習感知三者之間，是否具有顯著相關性？

三、護理學生的學習感知是否能真實反映其實際的 OSCE 溝通表現？

2.方法

2.1. 受試者與測試環境

本研究受試者為臺灣某大學護理系三、四年級學生共 28 名（男性 12 名，女性 16 名），平均年齡為 20.96 歲，採課堂宣導方式招募，並以自願參與為原則。所有受試者於研究開始前，均已簽署研究參與同意書。其中 21 名學生曾於研究前使用過 ChatGPT，另有 7 名從未使用。所有受試者皆接受過基本護病溝通訓練，並具備 OSCE 評量與臨床實習經驗。

測試環境設於一間安靜研究室內，學生透過筆記型電腦以語音轉文字方式與 ChatGPT 互動，螢幕為 30 吋之外接顯示器，作為模擬對話視覺介面。測試進行時僅有一位受試者與一位研究人員在場，以避免他人干擾或造成訊息洩漏，確保資料的客觀性與真實性。

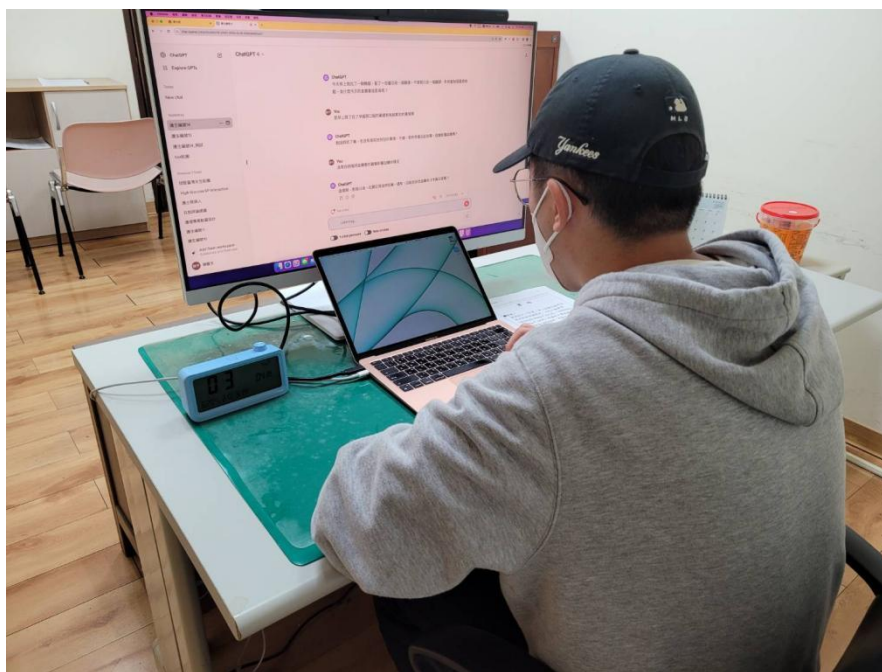


圖 1 ChatGPT 模擬 OSCE 護病溝通測試環境

2.2. 研究流程

本研究為個案研究，使用 OpenAI 開發的 ChatGPT-4 模型，模擬 OSCE 護病溝通與評分，藉由觀察、記錄和分析個別護生實際測試 ChatGPT 模擬 OSCE 問診評估與衛教之內容和行為，瞭解不同性別的護生使用 ChatGPT 進行高齡高血糖患者的護病溝通表現和學習感知。

研究分三階段進行：(1)設計提示語階段：利用具有專家信、效度的 OSCE 教案開發護病溝通的提示語，經過數十次測試，確認 ChatGPT 可以正確模擬 OSCE 中的標準病人，與護生溝通對話，確保護生練習問診和衛教之回應內容，皆符合正確病例情境訓練資料的範圍。(2)專家評估階段：請護理專家實際與 ChatGPT 進行問診與衛教，修正不符合實體 OSCE 過

程的互動說法與行為。(3)護生實測階段：護生實際進行 OSCE 護病溝通，瞭解 ChatGPT 模擬成效，如評分、回饋準確性和護生自評測試過程的學習感知。

正式評估流程為：(1)研究者說明本實驗目的和作法。(2)護生填寫受試者同意書。(3)系統操作說明：研究者教導護生如何運用語音轉文字與 ChatGPT 互動，並讓護生實際練習。(4)護生用一分鐘時間詳讀情境與考生任務。(5) ChatGPT 模擬 OSCE 護病溝通實作：護生在八分鐘時間內與 ChatGPT 模擬老年高血糖患者進行血糖、用藥、飲食與運動問診評估與衛教。(6)護生填寫自編「ChatGPT 學習感知問卷」。

2.3. 操作工具及提示語設定

由於護理教育相當重視專業術語和特定疾病判斷標準的正確性，因此，在護生與 ChatGPT 模擬 OSCE 護病溝通時，必須確保 ChatGPT 能根據 OSCE 臨床情境劇本回應護生，並能依照事先提示語設定的評分標準，正確評估護生的 OSCE 溝通能力。本研究在 ChatGPT 的指令中，輸入評分標準和回應模式，例如當護生說「請評分」時，ChatGPT 就要根據預設評分標準，將前面問診與衛教的對話逐一檢查，提供護生各項評分結果與原因，以及當護生問到飲食或運動的問題，ChatGPT 模擬的病人必須主動追問相關問題。

研究者會先告知 ChatGPT 接下來對話，請它扮演老年糖尿病患者，接著把病患的病史貼在對話中，請 ChatGPT 根據這些資料回答，最後再放一張由 ChatGPT 繪製的老年女性糖尿病患者圖片，讓護生在一開始和 ChatGPT 對話時，會先看到那張病患圖片，幫助護生投入 OSCE 溝通情境中。

2.4. 資料蒐集與分析

本研究蒐集之資料包括：ChatGPT 所記錄的護理學生與模擬 OSCE 病人之完整對話歷程，以及學生填寫之學習感知問卷。所有統計分析皆使用 IBM SPSS Statistics 28 版進行處理。在溝通表現方面，本研究採用兩項指標進行評估：其一為護生與 ChatGPT 模擬病人互動的總對話次數；其二為語言溝通能力評分，依據黃湘萍等人(2017)所建構之 OSCE 案例「老年高血糖病患血糖控制評估與指導」之標準化溝通評量準則進行評分。該量表原始總分為 28 分，考量本研究採用之 ChatGPT 無法進行非語言表達，本研究刪除其中一項非語言評分指標，調整後滿分為 26 分。原評量工具具良好信效度，整體評分者一致性指標 ICC 為.98，評分者信度達 100%，原始 Cronbach's α 為.54。

本研究邀請兩位具臨床與教學經驗之護理專家，分別針對護生的語言表達、同理展現與內容正確性進行獨立評分，針對評分不一致之項目進行討論修正，最終互評者信度 Cronbach's α 達.96，具高度一致性。

學習感知部分則以自編問卷進行評估，共包含 20 題，採李克特氏五點量表作答，分數愈高代表自我學習成效感知愈佳。其中 8 題為對 ChatGPT 學習工具之態度評價，12 題為「學習成效感知」構面，後者之內部一致性信度為 Cronbach's α = .94，整體問卷信度為.95。研究分析中使用之學習感知分數，係取上述 12 題之平均分數。

在統計分析方面，本研究採用 Mann-Whitney U 檢定比較不同性別護生在對話次數、OSCE 溝通分數與學習感知三項變項之差異，並計算 Cohen's r 作為效果量，評估性別影響之實質程度。為進一步分析變項間之關聯性，採用 Spearman 等級相關分析進行檢驗，顯著水準設為 α = .05。此外，為探討學生的主觀學習感知是否與其實際表現相符，將學習感知與 OSCE 得分進行 Z 標準化處理 (Z-score)，並計算其差距值 (Z 感知-Z 表現)。Z 值為正，代表學生自我感知高於實際表現，反之，則表示自評偏低。最後，再以 Mann-Whitney U 檢定

分析不同性別學生在 Z 落差分數上的顯著性差異，以評估性別對學習感知與實際表現一致性的潛在影響。

3.結果

本研究探討 ChatGPT 模擬 OSCE 護病溝通時，護生的性別對於對話數、OSCE 溝通分數與學習感知之差異與關聯。

3.1. 性別在學習感知、對話數與溝通分數比較

表 1 的 Mann-Whitney U 檢定結果顯示，男護生在 OSCE 溝通分數 ($M = 12.00$, $SD = 4.41$) 上略低於女護生 ($M = 12.44$, $SD = 2.83$)，男護生在對話數 ($M = 13.17$, $SD = 6.04$) 亦略低於女護生 ($M = 14.50$, $SD = 5.14$)，但 OSCE 溝通表現 ($U = 91$, $p > .05$) 與對話數 ($U = 84.5$, $p > .05$)，皆無顯著性別差異。

問卷中，與學習感知相關題分數 ($M = 4.35$, $SD = 0.54$)。Mann-Whitney U 檢定結果顯示，男護生在學習感知分數上 ($M = 4.59$, $SD = 0.55$) 顯著高於女護生 ($M = 4.16$, $SD = 0.46$)， $U = 41.5$, $p = .011 < .05$, Cohen's $r = 0.482$ ，具中等效果量，結果顯示性別可能影響學生對 AI 模擬學習經驗的主觀評價學習感知。

表1 Mann-Whitney U 檢定性別在學習感知、對話數與OSCE溝通分數比較

項目	性別	數字	平均等級	等級總和	U	Z	顯著性
OSCE 溝通分數	男	12	14.08	169	91.0	-0.234	0.815
	女	16	14.81	237			
對話數	男	12	13.54	162.5	84.5	-0.535	0.593
	女	16	15.22	243.5			
問卷分數	男	12	19.04	228.5	41.5	-2.552	0.011*
	女	16	11.09	177.5			

3.2. 學習感知、對話數與溝通分數之關聯性分析

本研究採用 Spearman 等級分析學習感知、對話數與 OSCE 溝通分數之相關性，結果顯示對話數與 OSCE 溝通分數呈現顯著正相關 ($r = 0.4$, $p < .05$)，溝通時對話數較多者，可能 OSCE 溝通分數亦較佳。相對地，護生在 AI 模護護病溝通的學習感知、對話數與 OSCE 溝通分數之間則無顯著相關 ($p > .05$)。

3.3. 感知與表現 Z 分數落差之性別差異

為進一步探討學生在學習感知和實際 OSCE 溝通分數，是否會因不同性別而有所差異，本研究將學習感知與 OSCE 得分進行 Z 標準化後，計算其 Z 差距值 (Z 感知-Z 表現)。結果顯示男護生整體在「主觀學習感知」的 Z 分數高於實際 OSCE 溝通分數 ($Z = -2.365$, $p = .018 < .05$)，而女護生之主觀學習感知與實際表現則相對一致，統計檢定顯示兩者落差具有顯著性別差異 $p = 0.043$ ，代表男護生相較於女護生，傾向高估自身 OSCE 溝通表現，顯示性別對於感知與實際表現落差具有統計顯著性。

4.討論

本研究運用 ChatGPT 模擬病人情境，進行 OSCE 護病溝通測試，並由具備臨床經驗之護理專家依據標準化 OSCE 評分準則進行評量，目的在於探討生成式 AI 導入後，不同性別護生在學習感知、對話次數及實際溝通能力之異同與關聯性。

研究結果顯示，男、女護生在 OSCE 溝通分數與對話數方面，雖然平均表現上女性略優於男性，但未達統計顯著差異，且效果量屬極小或小幅度。此結果意涵，在生成式 AI 模擬護病溝通任務中，性別並非影響實際溝通表現的主要因素。此一發現與過去部分研究所指出之性別影響溝通技巧的觀點(如 Kalati et al., 2025)有所不同，或可說明在標準化 AI 對話環境中，溝通表現受到性別以外因素，如任務理解、專業知識的影響可能更為顯著。

相較於實際表現，學習感知得分則呈現明顯的性別差異。本研究發現男性護生在學習感知上的自我評估顯著高於女性，且具有中等效果量，顯示性別可能影響學生對 AI 輔助學習的主觀感受。進一步以 Z 標準化進行「學習感知與實際表現一致性」分析時，也發現男護生整體傾向高估自身表現，其主觀感知與實際 OSCE 分數存在顯著落差，反觀女性護生則在主觀評估與客觀表現間呈現相對一致的趨勢。此發現呼應了先前關於性別在科技使用與學習自我效能上差異的相關研究(Pan & Ni, 2024)，亦與 Kirkpatrick & Kirkpatrick (2007)對於學習成效評量層次的觀點一致，指出僅以學習者主觀反應評估學習成效可能有所偏差。

上述結果亦突顯教育評量過程中，性別公平性所應關注之問題。若 AI 教學系統設計者、教師或研究者僅依據問卷回饋判斷學生學習效果，可能高估男性學生的學習成果而低估女性學生的真實表現，進而產生不公平的教學回饋或評估結論。因此，建議未來在設計生成式 AI 作為護理教育溝通訓練工具時，應納入性別作為學習分析與教學調整的參考變項，並結合多元資料來源，如問卷、實作成績、反思紀錄等，以提升教學回饋的準確性與評量的信效度。

整體而言，本研究結果指出生成式 AI 具備模擬護病溝通的潛力，不同性別學生在實際表現上差異不大，但其在學習感知上的認知偏差，值得教育者與系統設計者關注。此發現對於未來 AI 應用於護理教育溝通能力培訓與學習回饋系統的設計，具有理論與實務的參考價值。

5. 限制和建議

本研究具備初步探索價值，惟仍存在數量限制。首先，研究樣本僅來自臺灣一所科技大學之護理系學生，且樣本規模有限 ($n=28$)，因此研究結果在統計推論與外部效度上仍具侷限性。其次，儘管本研究揭示性別在學習感知層面上具有顯著差異，但此結果是否能擴及不同背景、年齡層或跨文化樣本，仍需進一步驗證。

因此，建議未來研究應擴大樣本數，以提升研究之代表性與結果的普遍性，並納入更多控制變項，如學生先備知識或溝通風格，進一步釐清影響學習感知與表現差異的潛在因子，及結合質性訪談資料，深入瞭解學生對生成式 AI 輔助學習的實際經驗，與性別角色預期對評估的影響。本研究提供有關 AI 模擬 OSCE 護病溝通應用之初步實證基礎，期盼能作為未來護理教育整合人工智慧科技與強化性別敏感教學設計的重要參考依據。

6. 結論

本研究旨在探討護理學生於 ChatGPT 模擬 OSCE 護病溝通情境中的性別差異，聚焦於學習感知、對話次數與溝通能力分數三項指標之差異與關聯性。透過分析護生與生成式 AI

進行問診與衛教任務過程中之對話紀錄、OSCE 溝通評分與學習感知問卷結果，發現若干具啟發性的現象與趨勢。

研究結果顯示，男護生在學習感知之自我評估得分顯著高於女護生，然此感知與其實際 OSCE 溝通表現間並無顯著相關。進一步以 Z 分數分析學習感知與實際表現的一致性，亦發現男護生傾向高估自身能力，女性護生則呈現較高的一致性表現，顯示性別可能影響主觀學習感知與客觀學習表現之落差。相對而言，對話次數與 OSCE 溝通分數之間存在顯著正相關，代表對話次數可能有助於提升溝通任務達成表現，但此關係未因性別而產生顯著差異。

綜合而言，本研究初步指出在 AI 模擬 OSCE 溝通情境中，性別可能影響學生的主觀感知與評估一致性，惟對實際溝通能力與互動參與之影響有限。此一結果對未來 AI 輔助學習系統之設計與回饋機制具有實務啟示，建議教學者在進行生成式 AI 應用教學評估時，應將性別變項與對話次數納入參考指標，以提升評估準確性與教學公平性。

致謝

本研究承蒙臺灣國家科學及技術委員會（NSTC）專題研究計畫「NSTC-113-2410-H-255-002-」和長庚科技大學臨床技能中心(Grant No. ZRRPF3N0081)之補助，謹此致謝。

參考文獻

- 黃湘萍、趙莉芬、王瑜欣、劉英妹、倪麗芬、簡淑慧（2017）。建構及檢測護理客觀結構式臨床技能測驗考題之信效度、鑑別度及難易度。*護理雜誌*, 64(6), 67-76。
<https://doi.org/10.6224/JN.000084>
- Alharazi, R. M., Abdulrahim, R. J., Mazuzah, A. H., Almutairi, R. M., Almutary, H., & Alhofaian, A. (2025). Barriers and factors affecting nursing communication when providing patient care in jeddah. *Clinics and Practice*, 15(1), 19. doi:<https://doi.org/10.3390/clinpract15010019>
- Alshalawi, S., Aljedaani, S., Fintyanh, D., Salim, S., Fairaq, W., & Refaei, S. (2025). The Effect of Nurse-Patient Communication on Patient Satisfaction in the Emergency Department. *Journal of Nursing & Midwifery Sciences*, 12(1). <https://doi.org/10.5812/jnms-158740>.
- Al-Kalaldeh, M., Amro, N., Qtait, M., & Alwawi, A. (2023). Barriers to effective nurse-patient communication in the emergency department. *Emergency Nurse*, 31(5).
- Amin, H. A. & Alanzi, T. M. (2024) Utilization of Artificial Intelligence (AI) in Healthcare Decision-Making Processes: Perceptions of Caregivers in Saudi Arabia. *Cureus*, 16(8), e67584. doi:10.7759/cureus.67584
- Chang, C. Y., Yang, C. L., Jen, H. J., Ogata, H., & Hwang, G. H. (2024). Facilitating nursing and health education by incorporating ChatGPT into learning designs. *Educational Technology & Society*, 27(1), 215-230.
- Harden, R. M., & Gleeson, F. A. (1979). Assessment of clinical competence using an objective structured clinical examination (OSCE). *Medical education*, 13(1), 39-54.
- Huang T, Hsieh P, Chang Y. (2024). Performance Comparison of Junior Residents and ChatGPT in the Objective Structured Clinical Examination (OSCE) for Medical History Taking and Documentation of Medical Records: Development and Usability Study. *JMIR Med Educ*, DOI: 10.2196/59902

- Kalati, Seyed Hossein, Ahmadzadeh Tori, Neda¹, Kalati, Javaneh, Norouzi Parashkouh, Nastaran, Karimi, Hengameh (2025). The effect of blended educational methods in communication skills in nursing with students' interactions with patients: A quasi-experimental study. *Journal of Education and Health Promotion* 14(1), 1-8. DOI: 10.4103/jehp.jehp_1479_23
- Kirkpatrick, D. L., & Kirkpatrick, J. D. (2007). Implementing the four levels: A practical guide for effective evaluation of training programs. Berrett-Koehler Publishers.
- Lin, H. L., Liao, L. L., Wang, Y. N., & Chang, L. C. (2024). Attitude and utilization of ChatGPT among registered nurses: A cross-sectional study. *International Nursing Review*. <https://doi.org/10.1111/inr.13012>
- Misra, S. M., & Suresh, S. (2024). Artificial Intelligence and Objective Structured Clinical Examinations: Using ChatGPT to Revolutionize Clinical Skills Assessment in Medical Education. *Journal of Medical Education and Curricular Development*, 11. <https://doi.org/10.1177/23821205241263475>
- Naz, R., Khaleel, S., & Naz, S. (2024). Evaluation of nurse-patient communication and its impact on satisfaction. *Biological and Clinical Sciences Research Journal*, 2024(1), 936. <https://doi.org/10.54112/bcsrj.v2024i1.936>
- Onwudiegwu, U. (2018). OSCE: Design, development and deployment. *Journal of the West African College of Surgeons*, 8(1), 1–22.
- Pan, G., & Ni, J. (2024). A cross sectional investigation of ChatGPT-like large language models application among medical students in China. *BMC Medical Education*, 24(1), 908. <https://doi.org/10.1186/s12909-024-05871-8>
- Saad, O., Saban, M., Kerner, E., & Levin, C. (2025). Augmenting Community Nursing Practice With Generative AI: A Formative Study of Diagnostic Synergies Using Simulation-Based Clinical Cases. *Journal of Primary Care & Community Health*, 16, 21501319251326663.
- Sibiya, M. N. (2018). Effective Communication in Nursing. *Nursing*, 19, 20-34. doi: 10.5772/intechopen.74995
- Siokal, B., Amiruddin, R., Abdullah, T., Thamrin, Y., Palutturi, S., Ibrahim, E., ... & Mallongi, A. (2023). The influence of effective nurse communication application on patient satisfaction: a Literature Review. *Pharmacognosy Journal*, 15(3), 479-483. doi:10.5530/pj.2023.15.105
- Tsang, R. (2023). Practical applications of ChatGPT in undergraduate medical education. *Journal of Medical Education and Curricular Development*, 10, 23821205231178449.