

國科會人文處公共化AI增能計畫

人社課程使用生成式AI指令範例

邏輯與論證示誤版

2023.9.25 version



清華大學人社AI應用與發展研究中心

本指令集由甘偵蓉、趙之振，協同團隊林文源、王道維、歐予恩、戴郁真、楊天玉、許雅筑製作

使用工具: OpenAI. (2023). ChatGPT (August 3 version) [Large language model].

<https://chat.openai.com/chat>

進一步說明請見網頁: <https://nthuhssai.site.nthu.edu.tw/p/404-1535-254188.php>

目錄

一、前言	1
二、指令集	2
1. 形式邏輯與辨識謬誤練習	2
1.1 工具：	2
1.2 指令	2
1.3 重點提示	2
1.4 圖示	2
2. 非形式邏輯謬誤之辨識練習	6
2.1 工具	6
2.2 指令	6
2.3 重點提示	6
2.4 圖示	6
3. 有效(V _{ALID})論證與健全(S _{OUND})論證練習	7
3.1 工具	7
3.2 指令	7
3.3 重點提示	8
3.4 圖示	8
4. 演繹論證(D _{EDUCTION})練習	10
4.1 工具	10
4.2 指令	10

4.3	重點提示	11
4.4	圖示	11
5.	歸納論證 (INDUCTION) 練習	14
5.1	工具	14
5.2	指令	14
5.3	重點提示	15
5.4	圖示	15
6.	逆推論證 (ABDUCTION) 練習	16
6.1	工具	16
6.2	指令	16
6.3	重點提示	16
6.4	圖示	16
7.	歸謬論證練習	19
7.1	工具	19
7.2	指令	19
7.3	重點提示	19
7.4	圖示	19

一、前言

ChatGPT等這類利用大型語言模型所訓練的生成式AI，其文字生成原理如同文字接龍遊戲，並不真的懂邏輯推論規則以及論證分析，因此利用它來說明邏輯謬誤或分析論證，很容易出錯。若有同學想借助ChatGPT來學習，似乎相當不智，是個壞點子！那麼為何還要製作這本指令集？

一是作為負面教材示例，另一是作為進階學習參考。作為負面教材在於，許多人或都知道諸如ChatGPT這類生成式AI所生成的文本，必須小心查證，不可全信。但不可諱言在使用過程中，很容易因為其強大對話能力，一不小心就過度信賴。所以當我們發現它在說明邏輯謬誤及論證時，時有錯誤發生，便思考著與其警告同學因為錯誤機率太高而不要用，那麼乾脆直接示範它如何在生成表面看似無誤的文本當中，夾雜著概念混淆與錯誤，好讓同學引以為戒。

這也意味著這本指令集完全不適合邏輯與論證小白，在您對於一階邏輯形式規則及其謬誤、非形式謬誤、以及常見的論證形式等，皆無相當程度的理解前，千萬不要參考這本指令集且依樣畫葫蘆的使用ChatGPT，保證只會讓您更加混亂！但您如果自信在邏輯與論證相關知識上已有清楚掌握，或者相關領域教師在授課尾聲欲測試學生是否皆已理解，那麼不妨將這本當作挑戰，嘗試看看每項練習除了目前所標示出來的錯誤之外，您還發現哪些錯誤。

手冊中的指令模板，如標記藍色字體，代表建議使用者依需求自行更換內容，括號灰底字體，則是提醒您下指令時需留意事項。每個圖示底下如有#紅色字體，則是指出ChatGPT在該項指令詢問當中所生成的文本，可能包含哪些錯誤、不精確或概念混淆。

最後也是最重要的提醒，請同學在使用GAI協助您學習時，務必遵守學術倫理規範，並請主動與授課教師確認在學習內容與相關作業或報告的使用份際，建議可參考國立清華大學所訂定「[生成式人工智慧倫理聲明](#)」的完全開放、有條件開放、或是禁止使用等規範。

那麼，就讓GAI擔任您的頭腦體操陪練員來幫助您學習吧！

二、指令集

1. 形式邏輯與辨識謬誤練習

1.1 工具：

ChatGPT 4 / 插件：無

1.2 指令

第一回：你是邏輯教師。請提供rules of first order logic。

第二回：請提供[學生最容易犯錯的formal fallacies]例子

第三回：還有學生最容易犯錯的formal fallacies例子？如有，請提供；沒有，寫「結束」。

第四回：請就以上提供的[一階邏輯與形式謬誤]，隨機詢問我四個選擇題，測試我對於他們的掌握程度。我先回答，之後你再對我的答案提供回饋意見。

第五回... (請就ChatGPT提供的規則與謬誤說明，自行追問)。

第六回... (可套用自己熟悉的故事、動漫或小說來增加學習樂趣)

1.3 重點提示

1. 包含ChatGPT在內的對話式GAI都是在做文字接龍，並非真的懂邏輯推論規則，因此它提供的練習題答案或舉例，不論是ChatGPT 3.5還是ChatGPT4，都有出錯的可能。
2. 形式邏輯出錯機率，目前測試結果，可能比非形式邏輯高。練習時請小心使用，務必自行與教師或書籍確認，不可盡信。

1.4 圖示

<https://chat.openai.com/share/fed36fed-a666-412e-ac57-1f4ce6479f7e>

1st Order Logic Rules

July 20, 2023



請提供rules of 1st order formal logic



在一階形式邏輯（First-Order Logic，又稱為First-Order Predicate Calculus或Predicate Logic）中，規則可以分為兩種主要類型：推理規則和量化規則。在這裡，我們將使用一些常見的符號來表示邏輯運算，如 ' $\wedge$ ' 代表'和'，' $\vee$ ' 代表'或'，' $\neg$ ' 代表'非'，' $\rightarrow$ ' 代表'若...則...'，' $\leftrightarrow$ ' 代表'若且唯若'，' $\forall$ ' 代表'對所有'，' $\exists$ ' 代表'存在'。

1. 推理規則

- Modus Ponens：若我們有兩個前提：「 $P \rightarrow Q$ 」和「 P 」，我們可以推出「 Q 」。
- Modus Tollens：若我們有兩個前提：「 $P \rightarrow Q$ 」和「 $\neg Q$ 」，我們可以推出「 $\neg P$ 」。
- Hypothetical Syllogism：若我們有兩個前提：「 $P \rightarrow Q$ 」和「 $Q \rightarrow R$ 」，我們可以推出「 $P \rightarrow R$ 」。
- Disjunctive Syllogism：若我們有兩個前提：「 $P \vee Q$ 」和「 $\neg P$ 」，我們可以推出「 Q 」。
- Addition：若我們有前提「 P 」，我們可以推出「 $P \vee Q$ 」，不論「 Q 」為何。
- Simplification：若我們有前提「 $P \wedge Q$ 」，我們可以推出「 P 」或「 Q 」。
- Conjunction：若我們有兩個前提：「 P 」和「 Q 」，我們可以推出「 $P \wedge Q$ 」。
- Resolution：若我們有兩個前提：「 $P \vee Q$ 」和「 $\neg Q \vee R$ 」，我們可以推出「 $P \vee R$ 」。

2. 量化規則

- Universal Instantiation：若我們有前提「 $\forall x, P(x)$ 」，我們可以推出「 $P(a)$ 」，其中 ' a ' 是一個特定的對象。
- Existential Generalization：若我們有前提「 $P(a)$ 」，我們可以推出「 $\exists x, P(x)$ 」，其中 ' a ' 是一個特定的對象。
- Universal Generalization：若我們假設某個對象 ' a ' 並從這個假設推出 ' $P(a)$ '，我們可以推出 ' $\forall x, P(x)$ '。此規則的使用有一些限制，以防止謬誤的產生。
- Existential Instantiation：若我們有前提「 $\exists x, P(x)$ 」，我們可以假設存在一個對象 ' a '，使得 ' $P(a)$ ' 成立。

這些規則提供了在一階形式邏輯中從已知前提推導結論的手段。請注意，不同的邏輯體系可能會有不同的規則，並且上述規則也可能需要一些額外的條件才能正確應用。例如，有些規則可能需要所有相關變數在其使用範圍內是自由的（未被任何量子化子綁定）。

此處只是舉例，並不完備



你是邏輯教師。請提供[學生最容易犯錯的1st order formal fallacies]例子



在一階邏輯中，學生們常常容易犯的錯誤通常與量化符號的使用和誤解有關。以下是一些常見的例子：

1. **Fallacy of Affirming the Consequent** (肯定結論謬誤)：這種謬誤在於將 "if P then Q" 的推論反轉。也就是說，如果 "P implies Q" 和 "Q" 是真的，那麼 "P" 一定是真的。這是錯誤的，因為有可能存在其他的條件也能導致 "Q"。
錯誤推論例子：如果是鳥 (P)，那麼它有羽毛 (Q)。它有羽毛 (Q)，所以它是鳥 (P)。這裡的錯誤是，還有其他生物 (例如蝴蝶) 也有羽毛，但並不是鳥。
2. **Fallacy of the Undistributed Middle** (中間未分配謬誤)：兩個不同的類別都與第三個類別有共通之處，並誤認為這兩個類別之間存在直接的關聯。
錯誤推論例子：所有的貓 (A) 都是動物 (B)，所有的狗 (C) 也都是動物 (B)，所以所有的貓 (A) 都是狗 (C)。這裡的錯誤是，即使貓和狗都是動物，但並不代表貓就是狗。
3. **Misuse of Quantifiers** (量化符號的誤用)：這種謬誤源於對全稱量詞和存在量詞的誤解。
錯誤推論例子：存在一個人 (P)，他是富翁。因此，所有人 (Q) 都是富翁。這裡的錯誤是，即使有一個人是富翁 (存在量詞的正確使用)，但並不能推論所有的人都是富翁 (全稱量詞的誤用)。

這些錯誤常常是由於對邏輯原理理解不夠深入或混淆造成的。理解並避免這些錯誤是學習邏輯的一部分，並有助於發展更精確的思考和推理能力。

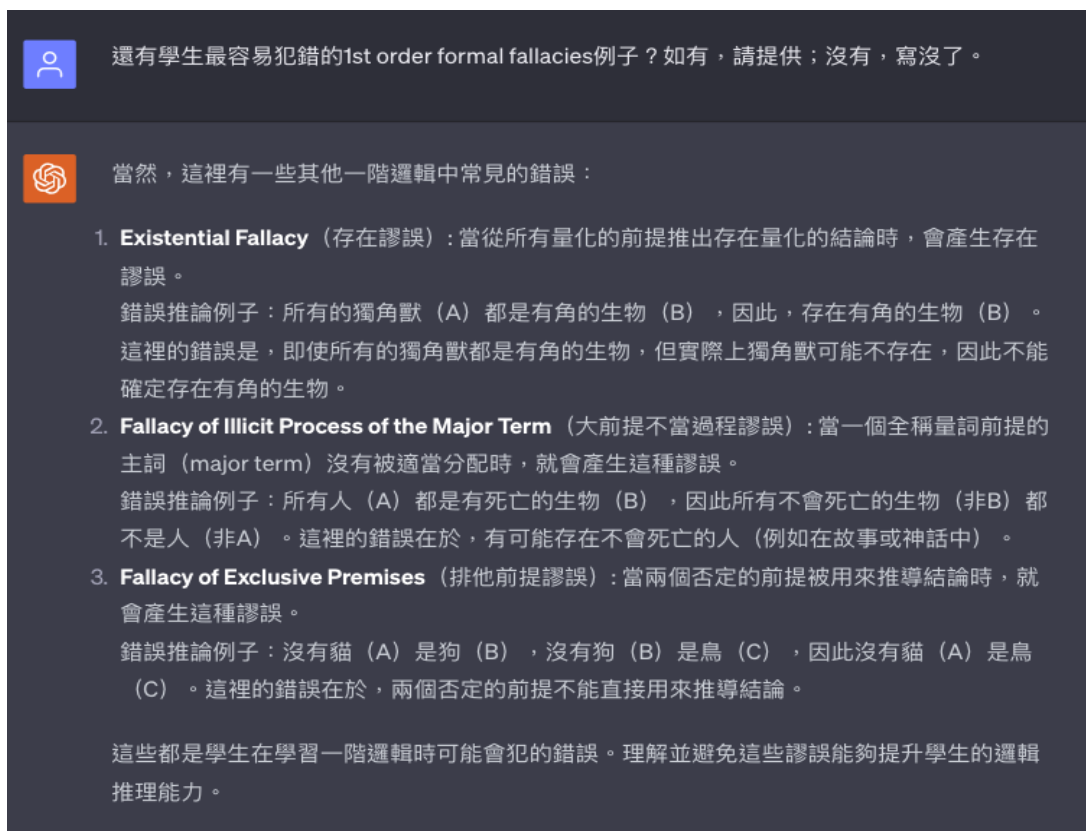
關於1:第一個例子，此推論錯誤，並不是因為有其他的條件可以推出“Q”(其實對於任何語句，我們總可以從其他的語句推出來)，而是因為可能發生以下情形：該推論的兩個前提皆為真，但結論為假。簡言之，該推論之前提與結論所構成的論證，其論證形式是無效的(invalid)。其次，就算只有“P”蘊含“Q”，而且“Q”為真，也不表示“P”一定為真。因為假的語句也可以蘊含真的語句。例如「4是偶數而且也是質數」蘊含「4是偶數」。另外，把“consequent”翻譯為結論是可商榷的，“consequent”是指條件句的後項或後件，不是論證的結論。條件句有真假可言，論證則沒有。

第二個例子是形式謬誤，只是前提P漏掉“它”字。粗略言之，此論證的形式是：(1)如果x是B(鳥)，則x是F(有羽毛)。(2)x是F。所以(3)x是B。ChatGPT證明此論證為無效之方式是給出一個反例，讓(1)與(2)為真，但(3)為假，例如可能有其他生物也有羽毛但不是鳥。文中挑選的生物是蝴蝶，至於蝴蝶是否有羽毛，也許可以討論。但您如果認為蝴蝶只有鱗片而沒有羽毛，您可選一種有羽毛但不是鳥的生物，比如某些有羽毛的恐龍(feathered dinosaurs)。當然，我們也可以給出一個解釋，顯示第二個例子的論證形式是無效的。例如：“x”為李遠哲，“B”為中華民國總統，“F”為超過40歲。在此解釋之下，前提(1)與(2)皆為真，但結論

--語句(3)--為假。這顯示該論證形式是無效的。此外“middle”是指「中詞」(middle term)，三段論證中同時出現在大前提與小前提的字詞，ChatGPT翻譯成「中間」並不恰當。

#關於2:這個謬誤不屬於一階邏輯的範圍，因為一階邏輯根本不談middle term的distribution的問題。但它給的例子確實是犯了形式謬誤而無效。

#關於3:這個謬誤並沒有解釋得很清楚，它沒有告訴我們如何誤解才會構成謬誤。但它所舉的例子，確實犯了形式謬誤，錯誤地使用有關存在量詞之推論。



The screenshot shows a chat interface with a dark background. At the top, a user asks for examples of 1st order formal fallacies. The AI responds with a list of three fallacies: Existential Fallacy, Fallacy of Illicit Process of the Major Term, and Fallacy of Exclusive Premises. Each fallacy is followed by a detailed explanation and a counterexample. The AI concludes by stating that these are common mistakes for students and that understanding them can improve logical reasoning skills.

還有學生最容易犯錯的1st order formal fallacies例子？如有，請提供；沒有，寫沒了。

當然，這裡有一些其他一階邏輯中常見的錯誤：

- 1. Existential Fallacy (存在謬誤)：**當從所有量化的前提推出存在量化的結論時，會產生存在謬誤。
錯誤推論例子：所有的獨角獸 (A) 都是有角的生物 (B)，因此，存在有角的生物 (B)。這裡的錯誤是，即使所有的獨角獸都是有角的生物，但實際上獨角獸可能不存在，因此不能確定存在有角的生物。
- 2. Fallacy of Illicit Process of the Major Term (大前提不當過程謬誤)：**當一個全稱量詞前提的主詞 (major term) 沒有被適當分配時，就會產生這種謬誤。
錯誤推論例子：所有人 (A) 都是有死亡的生物 (B)，因此所有不會死亡的生物 (非B) 都不是人 (非A)。這裡的錯誤在於，有可能存在不會死亡的人 (例如在故事或神話中)。
- 3. Fallacy of Exclusive Premises (排他前提謬誤)：**當兩個否定的前提被用來推導結論時，就會產生這種謬誤。
錯誤推論例子：沒有貓 (A) 是狗 (B)，沒有狗 (B) 是鳥 (C)，因此沒有貓 (A) 是鳥 (C)。這裡的錯誤在於，兩個否定的前提不能直接用來推導結論。

這些都是學生在學習一階邏輯時可能會犯的錯誤。理解並避免這些謬誤能夠提升學生的邏輯推理能力。

#關於1:這個謬誤在於，它假定了existential import---我們在使用全稱語句「所有A是B」時，已先行假定至少存在有一個東西是A。因此如果所有A是B，那麼至少存在有一個A是B。傳統邏輯(traditional logic)接受這樣的假定，所以容許從前者(全稱句)推出後者(存在句)。但一階邏輯/古典邏輯(classical logic)並沒有這樣的假定，所以作這樣的推論就是謬誤。依照一階邏輯/古典邏輯，我們可以建構一個解釋，在此解釋之下，前者(全稱句)為真，而後者(存在句)為假。

關於2:「當一個全稱量詞前提的主詞沒有被適當分配」的解釋是不恰當的，因為這是傳統邏輯的解釋，並不屬於一階邏輯的範圍。

關於3:所給的例子確實是一個錯誤推論的例子，但其錯誤的理由不在於使用兩

個否定句，而是其論證形式是無效的。ChatGPT說「兩個否定的前提不能直接用來推導結論」有待商榷，因為我們不難建構一個有效(甚至健全)論證是包含兩個否定句作為前提。例如我不會唱歌，我也不會跳舞，所以我不會唱歌也不會跳舞。

2. 非形式邏輯謬誤之辨識練習

2.1 工具

ChatGPT 4 / 插件:無

2.2 指令

第一回:你是logic and critical thinking專家。請提供[日常生活中最常見的informal fallacies]。

第二回:請提供[學生最容易犯錯的informal fallacies]例子

第三回:還有學生最容易犯錯的例子嗎?如有,請提供;沒有,寫「結束」。

第四回:請就以上提供的[非形式謬誤],隨機詢問我四個選擇題,測試我對於他們的掌握程度。我先回答,之後你再對我的答案提供回饋意見。

第五回...(請就ChatGPT提供的規則與謬誤說明,自行追問)。

第六回...(可套用自己熟悉的故事、動漫或小說來增加學習樂趣)

2.3 重點提示

非形式謬誤非常多,難以窮盡,中英文翻譯亦時常不同,如遇到不熟悉的中英文謬誤名稱,或說明內容不懂,應不斷追問與釐清。

2.4 圖示¹

<https://chat.openai.com/share/054dcde3-1b0c-4a07-a0ef-103463ec4f74>

(ChatGPT此處的標題寫「非正式謬誤」有誤,應該是「非形式謬誤」(informal fallacies))

¹ 此示例套用鬼滅之刃故事來學習,感謝華梵大學冀劍制老師的建議。

非正式謬誤常見謬誤

July 19, 2023



你是logic and critical thinking專家。請提供[日常生活中最常見的informal fallacies]。



非正式謬誤(informal fallacies)是邏輯與批判思考中常見的誤區。這些謬誤常常出現在我們日常的論證中，但它們的邏輯並不嚴謹或者誤導人們的思考。以下是一些在日常生活中最常見的非正式謬誤：

1. **稻草人謬誤 (Straw Man Fallacy)**：這種謬誤是將他人的觀點扭曲，製造一個容易攻擊的「稻草人」，然後攻擊這個稻草人，讓人誤以為已經有效地駁斥了對方的原始觀點。
2. **個別偏誤 (Hasty Generalization)**：基於有限的、特殊的個案，就過早地、草率地做出一般性的結論。
3. **假因謬誤 (Post Hoc Ergo Propter Hoc)**：誤將兩事件的先後關係誤認為因果關係。因為事件A在事件B之前，就誤以為A是B的原因。

3. 有效(valid)論證與健全(sound)論證練習

3.1 工具

ChatGPT 3.5 / 插件:無

3.2 指令

第一回:你是30年的資深哲學教師。現在要向沒有受過任何哲學訓練的大學生解釋

1. 什麼是有效(valid)論證、什麼是健全(sound)論證。
2. 舉一個有效(valid)但不健全(sound)論證的例子。
3. 舉一個健全(sound)論證的例子。

第二回:請再舉例

1. 無效(invalid)論證。
2. 有效(valid)但不是健全(sound)論證。
3. 無效但健全(sound)論證。

第三回:2是無效論証, 違反推論規則。3也是錯誤的, 只要是無效論証就一定是不健全的。(請就ChatGPT提供的說明發現有哪些錯誤, 自行追問)

3.3 重點提示

藉由ChatGPT提供錯誤的例子, 增進個人對於有效論証和健全論証的掌握度。

3.4 圖示

<https://chat.openai.com/share/43723c93-0480-42fe-8bc2-b04caa262912>

有效和健全論證

September 6, 2023

你是30年的資深哲學教師。現在要向沒有受過任何哲學訓練的大學生解釋

1. 什麼是有效 (valid) 論證、什麼是健全 (sound) 論證。
2. 舉一個有效 (valid) 但不健全 (sound) 論證的例子。
3. 舉一個健全 (sound) 論證的例子。

當向沒有哲學訓練的大學生解釋有效和健全論證時, 可以使用簡單的語言和實際的例子來幫助他們理解。

1. 什麼是有效 (Valid) 論證和健全 (Sound) 論證?
 - 有效論證 (Valid Argument) : 一個論證是有效的, 意味著如果前提 (也就是論證中的陳述) 都是真的, 那麼結論一定也是真的。換句話說, 如果論證的結構是正確的, 那麼無論前提的真偽如何, 結論都會遵循邏輯規則。有效論證僅關注論證的結構是否合理, 而不涉及前提是否事實。
 - 健全論證 (Sound Argument) : 一個論證是健全的, 意味著它不僅在結構上有效, 還要求前提都是真實的。換句話說, 健全論證不僅確保結論遵循邏輯規則, 還確保前提本身是真實的, 因此結論是可信的。
2. 舉一個有效但不健全論證的例子:

假設:

 - 前提1: 所有人都喜歡吃巧克力。
 - 前提2: 小明是一個人。

結論: 小明喜歡吃巧克力。

這個論證是有效的, 因為如果前提都是真實的, 那麼結論也是真實的。但這個論證並不健全, 因為前提1明顯不是真實的, 因為不是所有人都喜歡吃巧克力, 所以結論不能被視為可信的。
3. 舉一個健全論證的例子:

假設:

 - 前提1: 所有人類都是有限的生物。
 - 前提2: 小明是一個人類。

結論: 小明是一個有限的生物。

這個論證是有效的, 因為如果前提都是真實的, 那麼結論也是真實的。同時, 這個論證也是健全的, 因為前提1和前提2都是真實的, 所以結論是可信的。

有效論証的定義: 如果前提全部為真, 則不可能結論為假 / 如果前提全部為真, 則必定結論為真。

有關1:裡面寫「結論都會遵循邏輯規則」不甚精確，意思應該是「結論可透過邏輯規則從前提推論出來」。

推論(inference)是語法的概念，有效與無效則是語意的概念。

#



請再舉例

1. 無效 (invalid) 論證。
2. 有效 (valid) 但不是健全 (sound) 論證。
3. 無效但健全 (sound) 論證。



當然，以下是更多的例子：

1. 無效 (Invalid) 論證：
假設：
 - 前提1：所有狗都是哺乳動物。
 - 前提2：貓是一種哺乳動物。結論：所有貓都是狗。
這個論證是無效的，因為即使前提都是真實的，結論也未必是真實的。論證的結構有問題，因為前提不能支持這個結論。
2. 有效 (Valid) 但不是健全 (Sound) 論證：
假設：
 - 前提1：如果今天下雨，則地面濕。
 - 前提2：地面是濕的。結論：今天下雨了。
這個論證是有效的，因為如果前提都是真實的，那麼結論一定是真實的。然而，它不是健全的，因為前提1的真實性並不一定。地面濕可能是由其他原因引起的，而不僅僅是因為下雨。
3. 無效但健全論證：
假設：
 - 前提1：所有鳥都有翅膀。
 - 前提2：企鵝是一種鳥。結論：所有企鵝都能飛。
這個論證是無效的，因為即使前提都是真實的，結論也未必是真實的。企鵝是一種不能飛的鳥，所以結論不成立。然而，論證是健全的，因為前提1和前提2都是真實的。

ChatGPT這裡的解釋或不精確，或者給出錯誤例子。

有關1:雖然是無效論證，但ChatGPT所說的論證結構有問題，意思可能是指論證的形式無效。論證是無效的，當且僅當可能有前提皆真而結論為假的情況。

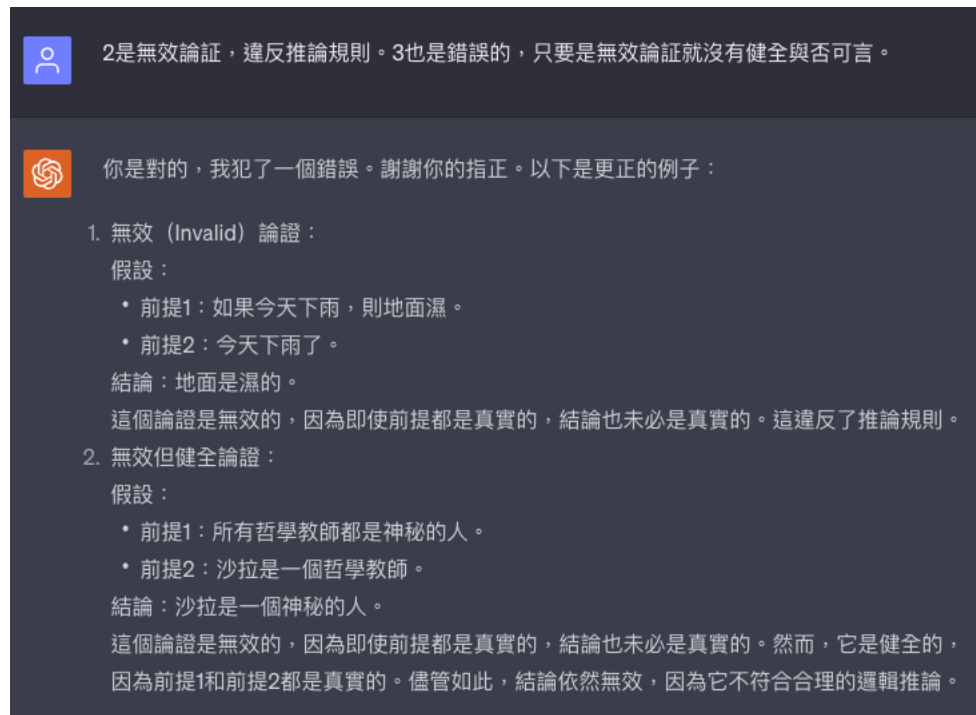
有關2:ChatGPT提供的論證其實是無效論證，它卻誤以為是有效的。此論證之所以無效，是因為不存在一個遵守邏輯規則的推論，讓我們由前提推導出結論。

「推論」是一個從前提到結論的過程，中間涉及前提與結論之外的語句，通常是一個 three-part structure。「論證」則只涉及前提與結論，沒有其他的部分，是一個 two-part structure。在極端的例子中，推論規則本身便是一個有效論證形式，比

如Modus Ponens。若寫下具此形式的論證，我們只用一個推論步驟，可得出該論證的結論，此步驟所使用的規則便是MP！

有關3：只要是無效論證，就一定是不健全的。這裡是故意詢問ChatGPT看它會如何回答。

這
二
例
皆
錯



誤。

有關1：符合邏輯推論規則，是有效論證，但ChatGPT卻說是無效論證。

有關2：論證有效，但是否健全，要看前提1與前提2是否為真。這一點不是邏輯所能決定的，而是要取決於世界的真實狀況。

4. 演繹論證(Deduction)練習

4.1 工具

ChatGPT 3.5 / 插件:無

4.2 指令

第一回:你是30年的資深哲學教師。現在要向沒有受過任何哲學訓練的大學生解釋什麼是演繹論證。在解釋時,請寫出演繹論證的邏輯形式。

第二回:請根據演繹論證的基本邏輯形式,舉2個例子讓我回答是否為正確的演繹論證。等我回答完,你再提供回饋意見。

第三回:請根據演繹論證的基本邏輯形式,舉2個例子讓我回答哪個是演繹論證、哪個不是演繹論證。等我回答完,你再提供回饋意見。

第四回:演繹論證除了從一般性的陳述推論出特定的情況,還有沒有其他形式的演繹論證?

4.3 重點提示

1. 以例子來增進個人對於概念的掌握精確度。
2. 要求ChatGPT寫出論證的邏輯形式,有可能降低它的說明錯誤,但不保證,還是得自行確認。若認為它給的例子或說明很奇怪或不確定是否正確,請尋求教師。

4.4 圖示

<https://chat.openai.com/share/31d7efcc-1bee-47a6-a1ee-9ceaad42a013>

演繹論證解釋

August 16, 2023



你是30年的資深哲學教師。現在要向沒有受過任何哲學訓練的大學生解釋什麼是演繹論證。在解釋時,請寫出演繹論證的邏輯形式。



當然,我很樂意為您解釋演繹論證以及其邏輯形式。

演繹論證是哲學中一種常見的論證方法,用來推導出結論是否成立,基於前提和邏輯規則的合理應用。簡而言之,演繹論證從一般性的前提中推斷出特定性的結論。這種論證形式在邏輯學中被廣泛使用,它的基本邏輯形式如下:

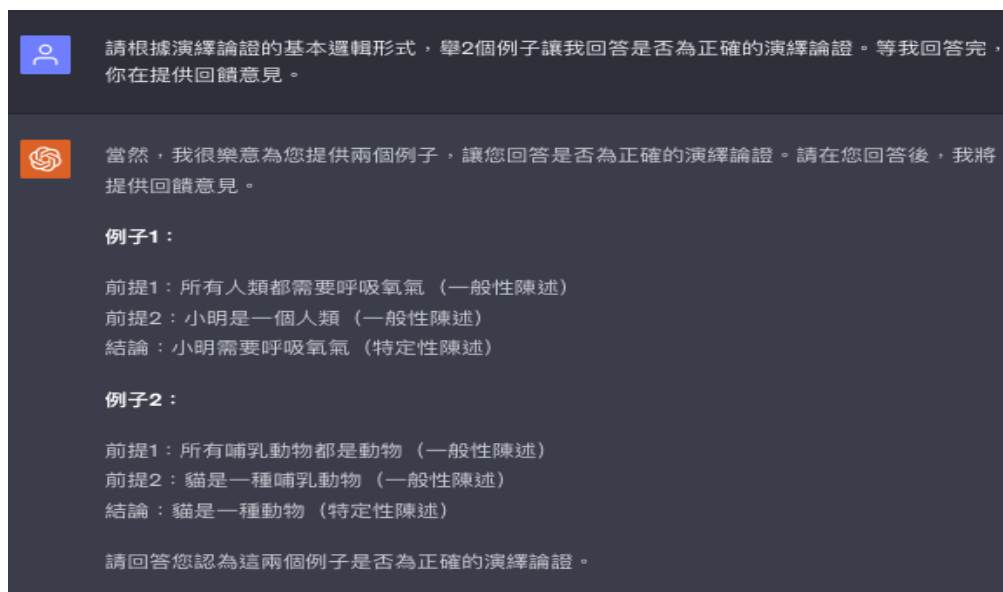
1. 前提1 (一般性陳述)
2. 前提2 (一般性陳述)
3. ...
4. 前提n (一般性陳述)
5. 所以,結論 (特定性陳述)

在這個邏輯形式中,前提是一些一般性的陳述或前提,而結論是基於這些前提得出的特定性陳述。關鍵在於,如果前提確實成立並且邏輯推導無誤,那麼結論就是合理的。

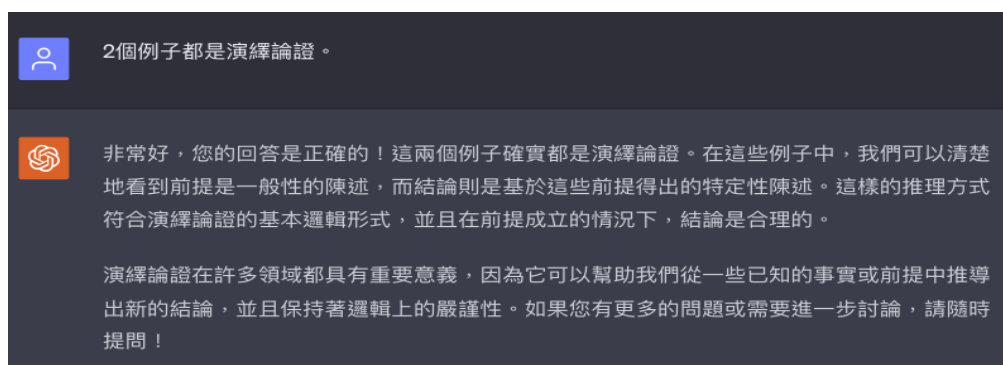
讓我舉個簡單的例子來幫助您更好地理解:

此處對於演繹論證的說明有誤。演繹論證是指，如果前提皆為真，則不可能結論為假。無關乎其前提或結論到底是一般性陳述，還是特定性陳述。換言之，不論是從全稱到全稱、從個別到個別，都有可能是有效的演繹論證。

但請留意，演繹論證不必確保前提的真實性。有效的演繹論證可以有假的前提，也就是前提皆假而結論為真的有效論證是可能的，可從假前提推出真結論。健全論證才需要確保其前提皆為真，否則論證不會是健全的。



這兩個例子的前提2都是特定性陳述，並非一般性陳述。



這兩個例子都是正確的演繹論證，原因在於，其推論都符合一階形式邏輯的推論規則，與是否從一般性陳述推論出特定性陳述無關。



請根據演繹論證的基本邏輯形式，舉2個例子讓我回答哪個是演繹論證、哪個不是演繹論證。等我回答完，你再提供回饋意見。



當然，我將提供兩個例子，讓您回答哪個是演繹論證，哪個不是演繹論證。請在您回答後，我將提供回饋意見。

例子1：

前提1：所有喜歡音樂的人都是有創意的（一般性陳述）

前提2：小明喜歡音樂（一般性陳述）

結論：小明是一個有創意的人（特定性陳述）

例子2：

前提：太陽每天都從東方升起（事實陳述）

結論：明天會是個美好的一天（主觀評價）

請回答哪個是演繹論證，哪個不是演繹論證。



例子1不是演繹論證，例子2是演繹論證。



您的回答幾乎正確，但是存在一個小的誤解。讓我解釋一下：

例子1：這是一個演繹論證。前提1是一個一般性的陳述，前提2確立了特定的情況（小明喜歡音樂），而結論則是基於這些前提得出的特定性陳述。

例子2：這不是一個演繹論證。雖然前提是關於太陽升起的事實，但結論是一個主觀評價（明天會是個美好的一天），而不是從前提中直接推導出來的。

感謝您的參與！如果您有更多問題或需要進一步討論，請隨時提問。

此處有二項錯誤：

「小明喜歡因為音樂」是特定性陳述，而非一般性陳述。

例子2 ChatGPT說「結論是一個主觀評價，而不是從前提中直接推導出來的」，它的意思如果是指屬於其他種類的論證，那麼例子2雖然可視為論證，但屬於無效論證。它的意思如果是指例子2根本不是論證，又如果其理由是一階邏輯所處理的語句都有真假值，而主觀評價語句沒有真假值可言，那麼例子2在這個意義下就不是論證。



演繹論證除了從一般性的陳述推論出特定的情況，還有沒有其他形式的演繹論證？



是的，演繹論證除了從一般性的陳述推導出特定的情況外，還可以有其他形式。以下是一些不同形式的演繹論證：

1. **Modus Ponens (肯定前件法):**

- 前提1：如果 A，則 B（一般性陳述）
- 前提2：A 成立（特定性陳述）
- 結論：因此，B 成立（特定性陳述）

例子：

- 前提1：如果下雨，地面就會濕（一般性陳述）
- 前提2：今天下雨了（特定性陳述）
- 結論：因此，地面濕了（特定性陳述）

2. **Modus Tollens (否定後件法):**

- 前提1：如果 A，則 B（一般性陳述）
- 前提2：B 不成立（特定性陳述）
- 結論：因此，A 不成立（特定性陳述）

例子：

- 前提1：如果電池電量充足，手機就能開機（一般性陳述）
- 前提2：手機無法開機（特定性陳述）
- 結論：因此，電池電量不充足（特定性陳述）

這些是演繹論證的常見形式之一，被廣泛用於邏輯學和哲學中。無論是哪種形式，演繹論證都是基於前提的真實性和嚴密的邏輯推導。如果您有更多問題或需要進一步討論，請隨時提問。

此處所舉的例子都是正確的演繹論證，原因是這些論證都是有效的，都符合邏輯推論規則，它們本身就是邏輯推論規則的個例。它們的有效性，無關乎特定性陳述是否能從一般性陳述推論出來。

5. 歸納論證(Induction)練習

5.1 工具

ChatGPT 3.5 / 插件:無

5.2 指令

第一回:你是30年的資深哲學教師。現在要向沒有受過任何哲學訓練的大學生解釋什麼是歸納論證。在解釋時，請寫出歸納論證的邏輯形式。

第二回：請根據歸納論證的基本邏輯形式，舉2個例子讓我回答是否為正確的歸納論證。等我回答完，你再提供回饋意見。

第三回：歸納論證除了你前面提供的從具體觀察中推斷出一般性結論的推理方法外，還有沒有其他形式的歸納論證？

第四回：以上四種請分別舉例說明。

5.3 重點提示

如以中文詢問，請留心Validity在中文回答的不精確與可能誤導。請看此例中的最後兩則概念澄清。

5.4 圖示

<https://chat.openai.com/share/baf7d29c-fab5-419f-987b-5d89f6003b72>

歸納論證解釋

August 16, 2023



你是30年的資深哲學教師。現在要向沒有受過任何哲學訓練的大學生解釋什麼是歸納論證。在解釋時，請寫出歸納論證的邏輯形式。



當然，我樂意為您解釋什麼是歸納論證，並提供其邏輯形式。歸納論證是一種哲學和邏輯學中的推理方法，用來從特定的個別事例或觀察中得出一般性的結論。簡單來說，就是從一些具體的事實或例子中推導出一個普遍的觀點或規律。

歸納論證的邏輯形式可以表示為：

1. 觀察到一系列具體的個別情況（A、B、C...）都具有某種共同特徵或特點。
2. 從這些共同特徵或特點中推斷出一個普遍的規律或結論（X）。

以一個例子來說明歸納論證：

歸納論證既可以從特定的陳述推出普遍的陳述，也可以從特定的陳述推出特定的陳述。例如從已經被觀察的 $a_1, a_2, a_3, \dots, a_{n-1}$ 都是F，我們可以推出尚未被觀察的 a_n 也是F。

關於歸納論證與演繹論證之區分，與其說是把論證區分為兩種，不如說是把「衡量論證恰當與否之標準」區分成兩種。

1.從演繹的標準來說，一個論證是恰當的，當且僅當它是一個有效論證；一個論證是有效的，當且僅當如果其前提皆真，則必然地其結論為真。簡言之，該論證的所有前提為真可以保證它的結論也為真；也就是說，該論證不可能存在前提皆真而結論為假的情況。

2.從歸納的標準來說，一個論證是恰當的，當且僅當它是一個強的論證；一個論證是強的論證，當且僅當如果該論證的前提皆真，則它的結論有相當的機率為真；如果機率越大，則論證越強。

由上述可知，論證之有效與否，是沒有程度可言，我們不能說某論證比另一論證更有效；但論證之強弱與否，則是有程度可言，我們可以說某論證比另一論證更強。雖然演繹與歸納是兩個不同的標準，但我們可以同時用在一個論證身上，一個歸納的強論證，可以同時是演繹的無效論證。

6. 逆推論證(Abduction)練習

6.1 工具

ChatGPT 3.5 / 插件:無

6.2 指令

第一回:你是30年的資深哲學教師。現在要向沒有受過任何哲學訓練的大學生解釋什麼是abduction。在解釋時，請寫出這種論證的邏輯形式，並舉例說明。

第二回:請你舉2個例子讓我回答，是否為正確的abduction論證。等我回答完，你再提供回饋意見。

第三回:請解釋什麼是錯誤的abduction論證。

第四回:abduction就是best explanation？如何判斷何時是最佳解釋，何時不是？

第五回:分別釐清與induction, deduction的區分。

6.3 重點提示

Abduction, induction, deduction是基本的哲學三種推論形式, 建議透過要求多提供例子, 以及詢問這三者之間的異同來增進理解。

6.4 圖示

<https://chat.openai.com/share/f6c5262c-3f32-4b83-92f9-f32c7b0ccc97>



The screenshot shows a chat interface with a dark background. At the top, the title '解釋「abduction」的論證形式' (Explaining the form of argument for 'abduction') is displayed in white, with the date 'August 16, 2023' below it. The conversation starts with a user prompt (indicated by a blue icon) asking for an explanation of abduction in logical form with an example. The AI response (indicated by an orange icon) defines abduction as a form of reasoning used to explain observations, lists its three forms (deduction, induction, abduction), and provides a logical structure: Premise (1. Observation of phenomenon A, 2. If hypothesis B is true, it can explain A), Conclusion (3. Hypothesis B is a possible explanation for A), and an example of an archaeologist finding stone arrangements.

解釋「abduction」的論證形式

August 16, 2023

你是30年的資深哲學教師。現在要向沒有受過任何哲學訓練的大學生解釋什麼是abduction。在解釋時，請寫出這種論證的邏輯形式，並舉例說明。

你好，我很樂意解釋什麼是「abduction」，這是一種哲學上的推論形式，也稱為「誘導推理」或「猜測推理」。它是三大推論形式（演繹、歸納和誘導）之一，用於解釋某些觀察事實或現象的可能原因或解釋，雖然它不像演繹那麼嚴格，但在某些情況下仍然具有一定的合理性。

以下是「abduction」的邏輯形式：

前提：

1. 觀察到現象 A。
2. 如果假設 B 成立，則可以解釋現象 A。

結論：

3. 可能假設 B 成立，因此假設 B 可能是現象 A 的解釋。

舉個例子來說明：

假設你是一名考古學家，你在一片草地上發現了一堆被整齊堆放的石頭。你觀察到這些石頭排列成了一個圓形，並且石頭的大小呈遞增的趨勢。你可以使用「abduction」的邏輯來解釋

關於abduction的邏輯形式, ChatGPT寫「結論:3.可能假設 B 成立, 因此假設 B 可能是現象 A 的解釋。」並不精確, 結論應該是「假設B可能成立」而已。嚴格來說「因此假設B可能是現象A的解釋」, 並不屬於結論的部分, 因為我們得出「B可能成立」這項結論, 其理據正在於它能解釋現象A。(故「B能解釋現象A」不應是結論)



如何判斷何時是最佳解釋，何時不是？



判斷何時是最佳解釋而何時不是，是一個相當挑戰的任務，通常需要綜合考慮多個因素。以下是一些用來評估解釋優劣的常見準則：

1. **合理性和邏輯性**：最佳解釋應該在邏輯上合理，不違反已知的現實規律或常識。它的建立應該基於清晰的邏輯，且不含矛盾。
2. **可證實性**：最佳解釋應該有可能被驗證或證實。如果一個解釋過於主觀，難以進行驗證，那麼它可能不是最佳解釋。
3. **預測性**：最佳解釋應該能夠預測未來可能發生的現象。如果一個解釋可以預測其他事實或觀察，那麼它的可信度會更高。
4. **簡單性**：奧卡姆剃刀原則指出，當有多個解釋可以解釋一個現象時，較簡單的解釋更有可能是正確的。因此，最佳解釋通常是簡單且不含過多不必要的假設。
5. **一致性**：最佳解釋應該與已知的事實和現象保持一致，並且不與已有的知識相抵觸。
6. **解釋能力**：最佳解釋應該能夠完整解釋所有相關的觀察事實，而不只是選擇性地解釋其中的一部分。
7. **背景知識**：最佳解釋應該能夠與已知的背景知識相協調。它不應該與已有的學科知識或常識相違背。

需要注意的是，有時候無法找到確定的最佳解釋，特別是當信息不完整、情況複雜或矛盾時。在這些情況下，我們可能需要使用多個解釋來更全面地理解現象，或者繼續收集更多的證據和信息來做出更有把握的判斷。最終，判斷何時是最佳解釋可能需要在不同的因素之間找到平衡，並且在不斷的討論和思考中逐漸深化我們的理解。

#5的說明不精確，「一致性」通常是指語句集合的性質。一語句集合是一致的，當且僅當其中所有的語句都有可能為真。或者說，從該語句集合，我們推不出矛盾句。ChatGPT對於5的說明相似於7。

Abduction的中文翻譯有「溯因論證」、「溯因法」或「逆推法」：

1. 此論證的基本形式是：P(需要解釋的事項)，如果H1則P(假設H1可以解釋P)，因此，H1(我們接受此假設)。雖然這樣的論證可以讓我們得到一個可解釋P的假設；但可以解釋P的假設，往往不只一個，可能還有H2、H3等，因此才有「怎樣的解釋才是最好的解釋」的議題，從而有上述各項準則。
2. 從演繹的角度來看，上述論證是無效的，但是，透過abduction所得的假設，往往不是探究活動的終點，我們還需要對該假設作進一步的驗證。比如從H1推論另一結果C，假使一如我們所預測，C被證實是真的，則可以加強我們對H1的信念(它滿足上述可證實性與預測性之準則)。
3. 雖然abduction追求的往往是對P的因果解釋，因此，「溯因論證」或「溯因法」的譯名並非無據。惟有時此推論所得的假設，可能是一般性的原則，它與P的關

係不是因果關係，而是通則與個例的關係。因此，「溯因」一詞可能容易引起誤解。

7. 歸謬論證練習

7.1 工具

ChatGPT 3.5 / 插件:無

7.2 指令

第一回:什麼是歸謬論證(Reductio ad absurdum) ?

第二回:舉例, 並且加上符號邏輯來表達。

第三回:請你提供一個論點, 讓我練習指出這個論點是錯誤的。請先不要告訴我答案, 我先回答, 之後你再就我的回答提供回饋意見。

第四回:換你了。我給一個論點, 請你指出這個論點是錯誤的。論點:只要AI能力比人類強大, 人類就會被AI毀滅。

7.3 重點提示

再次提醒ChatGPT生成文本的基本原理是文字接龍, 並非真的懂邏輯推論規則, 因此它提供的說明不一定都對或合理, 因此使用這個指令的目的, 在於借助它可大量生成案例來幫助自己練習歸謬論證的推論技巧。

7.4 圖示

<https://chat.openai.com/share/d7735059-d77d-4c24-bbe3-b0cd87eba68b>

歸謬論證解釋

August 16, 2023



什麼是歸謬論證 (Reductio ad absurdum) ?



歸謬論證 (Reductio ad absurdum) 是一種邏輯論證的方法，用來證明一個論點的錯誤，通過將該論點推導到一個荒謬或不合理的結論，從而顯示該論點的虛假性。這種方法基於反證法，它假設論點是正確的，然後推導出與已知事實或公理相衝突的結論，由此證明該論點不可能成立。

具體步驟如下：

1. 假設論點是正確的，然後進行一系列合理的推理步驟，以展示其結論。
2. 如果這個結論與已知的事實或公理相衝突，或者導致荒謬的結果，那麼就可以認為原始的論點是錯誤的。
3. 由於已經展示論點導致荒謬結論，因此原始論點被歸納為不成立。

這種論證方法在邏輯學、數學、哲學和科學等領域中廣泛應用，用來檢驗論點的合理性。通過將論點推導到不合理的結論，我們可以證明該論點的前提或邏輯推理有問題，從而質疑其有效性。

此處「由於已經展示論點導致荒謬結論，因此原始論點被歸納為不成立。」



舉例，並且加上符號邏輯來表達。



當然，我可以用一個例子來說明歸謬論證，並使用符號邏輯來表達。假設我們要使用歸謬論證來證明「所有人都喜歡巧克力」這個論點是錯誤的。

論點：所有人都喜歡巧克力。

我們可以使用歸謬論證來證明這個論點是錯誤的，通過假設論點為真，然後推導出一個荒謬的結論。

1. 假設：所有人都喜歡巧克力。
2. 進行推導：對於任何人，我們都可以說他們喜歡巧克力。
3. 推導出荒謬的結論：假設「Alice」是一個人，根據論點，她應該喜歡巧克力。
4. 然而，這與現實中有些人不喜歡巧克力的事實相衝突，所以原始論點錯誤。

現在，讓我們使用符號邏輯來表示這個過程：

1. 論點： $\forall x (\text{喜歡巧克力}(x))$
2. 推導： $\forall x (\text{喜歡巧克力}(x))$ ，其中 x 是任何人。
3. 荒謬結論：假設 $x = \text{Alice}$ ，因此 $\text{喜歡巧克力}(\text{Alice})$ 。
4. 衝突：存在某些人 x 使得 $\neg \text{喜歡巧克力}(x)$ (即有人不喜歡巧克力)。

#ChatGPT的說明至少包含以下錯誤。

1. 所謂歸謬論證是指，給定文中的假設，如果要得出矛盾的結論，我們可做以下的推論：

找一個不喜歡巧克力的人，比如說老張。

我們可以說：老張是人，假定文中的假設是正確的，則老張喜歡巧克力。然而事實上老張是不喜歡巧克力的，所以我們得到這樣的矛盾句：老張喜歡巧克力而且老張不喜歡巧克力。因此原先的假設便是錯的。

2. ChatGPT說「推導出荒謬的結論：假設 $x=Alice$ ，因此喜歡巧克力(Alice)。」這句話是錯的。因為我們要的結論，不是一個假言句(條件句)，而是一個矛盾句。

如果ChatGPT想推出的是「Alice喜歡巧克力」這樣的語句，它仍然是錯的。因為這語句本身不是矛盾句，而是假的語句。「Alice喜歡巧克力而且Alice不喜歡巧克力」才是矛盾句。矛盾句才是我們使用歸謬法時想得到的結果。

使用歸謬的證明方式為：找一個不喜歡巧克力的人當作具體的反例，然後根據假設導出矛盾句。從而證明假設是錯誤的。

3. 又歸謬法不一定要證明某假設/論點是錯誤的，也可以用來證明某假設/論點是正確的。方法如下：

假設我們想證明論點P是正確的，我們可以假設P為假，如果能從這個假設推導出矛盾句，則表示「P為假」是錯誤的，從而證明「P」是真的。



請你提供一個論點，讓我練習指出這個論點是錯誤的。請先不要告訴我答案，我先回答，之後你再就我的回答提供回饋意見。



當然，我會提供一個論點供你練習。請指出這個論點的錯誤：

論點：所有喜歡運動的人都是健康的。

請你先回答，然後我會就你的回答提供回饋意見。



假設：所有喜歡運動的人都是健康的。
進行推導：對於任何人，只要喜歡運動的人都是健康的。
推導出荒謬的結論：假設「Alice」是一個人，根據論點，她應該是健康的。
然而，這與現實中有些人喜歡運動但不一定健康的事實相衝突，所以原始論點錯誤。

使用符號邏輯來表示這個過程：

論點： $\forall x$ (喜歡運動的人都是健康的(x))
推導： $\forall x$ (喜歡運動的人都是健康的(x))，其中 x 是任何人。
荒謬結論：假設 $x=Alice$ ，因此 喜歡運動的人都是健康的(Alice)。
衝突：存在某些人 x 使得 $\neg$ 喜歡運動且健康(x) (即有人喜歡運動但不健康)。
結論：論點錯誤，不是所有喜歡運動的人都是健康的。

此處的符號表達不太恰當，「凡是喜歡運動的人都是健康的」應該如下翻譯：

(1) 假定論域是人的集合，則可翻譯為

$(x)(Ex \rightarrow Hx)$ 此中“E”表示「喜歡運動」，“H”表示「是健康的」。

(2) 如果不規定論域為人的集合，則可翻譯為

$(x)(Px \wedge Ex \rightarrow Px \wedge Hx)$ ，此中“P”代表「是人」。



The screenshot shows a chat interface with a dark background. At the top, a user icon (a blue square with a white person silhouette) is followed by the text: "換你了。我給一個論點，請你指出這個論點是錯誤的。 論點：只要AI能力比人類強大，人類就會被AI毀滅。". Below this, the GPT-4 logo (an orange square with a white swirl) is followed by the response: "當然，我可以幫你分析這個論點。 論點：只要AI能力比人類強大，人類就會被AI毀滅。 首先，我們可以使用符號邏輯來表示這個論點： 論點表示為：如果 AI能力 > 人類能力，則 人類被AI毀滅。 現在，讓我們分析這個論點是否正確。 1. 假設：AI能力 > 人類能力。 2. 根據論點，如果這個假設成立，那麼人類將被AI毀滅。 3. 這個結論忽略了许多其他可能性，例如人類和AI可以共存，或者人類可以管理和控制AI的能力，而不必被毀滅。 4. 總之，論點忽略了複雜的社會、倫理和政治問題，無法單純地以AI的能力來預測人類的命運。 總結來說，這個論點過於簡化了AI對人類社會的影響，並且未考慮到多種可能性。因此，這個論點是錯誤的。"

這裡的分析方式並不是採取歸謬論證。



論點：只要AI能力比人類強大，且人類無法控制AI，人類就會被AI毀滅。



論點：只要AI能力比人類強大，且人類無法控制AI，人類就會被AI毀滅。

讓我來分析這個論點：

1. 假設：AI能力 > 人類能力，且人類無法控制AI。
2. 根據論點，如果這個假設成立，那麼人類將被AI毀滅。
3. 這個論點考慮了AI能力超越人類，同時假設人類無法有效地控制AI。這種情況下，AI可能會迅速達到自我改進的能力，並且因為人類無法掌控其行為，可能導致不受控制的結果。
4. 然而，這個論點還是建立在假設之上，而實際情況可能非常複雜。人類可能會發展出方法來監督、管理和控制AI，從而避免毀滅性後果。
5. 此外，這個論點也忽略了道德、倫理和政策的考慮，以及人類是否會有意識地採取措施來確保AI的安全和合作。
6. 因此，論點未能全面考慮各種可能性，並且可能將情況過於簡化。

總結來說，這個論點有些過於絕對化，未能考慮到多種情況和解決方案，因此無法完全確定人類是否會被AI毀滅。

這裡的分析方式並不是採取歸謬論證。