

The White Horse Paradox and Inheritance Logic

Miguel López-Astorga

Institute of Humanistic Studies
Research Center on Cognitive Sciences
University of Talca
Av. Lircay s/n
Talca 3460000, Chile
milopez@utalca.cl

Abstract: Several studies have addressed the White Horse Paradox from different perspectives. Although the essential elements of the sentence in the paradox can already be found in Mozi's work, it is thought that Gongsun Long presented the paradoxical version of that very sentence. Here, I comment on some possible accounts related to logic, cognitive science and philosophy of science. However, my main goal is to show how Inheritance Logic deals with the paradox. It has been thought that the acceptance of the paradox depends on its interpretation: whether it relates the animal 'horse' to the color 'white' or the term 'horse' to the term 'white'. I claim that, while it is easy to give reasons for rejecting the first interpretation, a logic such as Inheritance Logic is necessary to support the second interpretation. The reason is that the latter is a term logic, not a propositional or predicate logic.

Keywords: *extension, Inheritance Logic, intension, term logic, White Horse Paradox*

Introduction

The *White Horse Paradox*, which is related to Gongsun Long (see, e.g., Sun, 2013), has been analyzed from different perspectives. It refers to the circumstance in which 'a white horse is not a horse' (from now on, I will refer to the latter sentence between quotes as 'P'). The basics of the problem are earlier than Gongsun Long's work. The first author that proposed a sentence such as 'A white horse is a horse' was Mozi (*Mohist Canons*, Book 45:4).

Mozi lived between the 5th and 4th centuries BCE. His intention was to show that we can arrive at general concepts such as ‘horse’ from expressions such as ‘white horse’ (e.g., Schumann, 2024). In particular, the sentence appeared in *Dialectical Chapters*. However, there is a discussion about Mozi’s texts. It appears that Liu Xiang, a librarian whose life probably spanned between 77 and 6 BCE, provided the form and order of the texts as we know them today (see, e.g., Theobald, 2010a; for further information on the school of thought following Mozi’s philosophy, see also, e.g., Graham, 2004; for a more general account of the different sources of ancient logic, see also, e.g., Schumann, 2023).

On the other hand, it is thought that Gongsun Long lived between 325 and 250 BCE. It seems that he wrote 14 chapters, but we have only 6 of them (Theobald, 2010b). What he proposed with P has been addressed in many ways. For example, the concepts of ‘use’ and ‘mention’ have been considered (Sun, 2013, indicates in this regard works such as that of Liu, 2007), cross-cultural philosophy has been taken as an interpretation tool (Havelange, 2020), and the issue has been dealt with even from a contemporary psychology approach (López-Astorga, 2021, where P is revised resorting to the theory of mental models; for a description of the latter theory, see e.g., Khemlani, Byrne and Johnson-Laird, 2018). However, I will focus on a logical consideration of P and the two senses P has, following its original words in ancient Chinese.

The problem seems to be the negation (*fēi*) as it enables two interpretations of P (e.g., Indraco, 2017):

- (A) A being that is a white horse cannot be a horse.
- (B) The expression ‘white horse’ does not have the same meaning as the word ‘horse’.

It is not difficult to offer arguments to undermine (A). I will comment on some examples below. However, the main point I wish to make is linked to Inheritance Logic (e.g., Wang, 2013). Beyond the rejection of (A), Inheritance Logic (IL) allows showing why (B) is the suitable interpretation.

First, I describe some examples opposing (A). I will start with the explanation of the acceptance of P that can be derived from the method of extension and intension proposed by Carnap (1947). As I will sustain, it is not difficult to challenge that explanation. Second, I will consider the account of the rejection of (A) that can be offered from the theory of mental models. Third, I will argue that even an alternative account based on Carnap’s (1936; 1937) theses about reduction and rejecting (A) is possible. However, my chief aim is to show that IL, in addition, is able to reveal that (B) is the correct option.

P and state-descriptions

There is an analysis of P based on a part of Carnap's philosophy (López-Astorga, 2021). The part is the method of extension and intension (Carnap, 1947). That is not the only account using the concepts of extension and intension to address P. For example, Wen (1983; explanation in English in Sun, 2013) pointed out that Gongsun Long was mostly interested in the difference between the intension of 'white horse' (WH^I) and the intension of 'horse' (H^I). In Wen Gongyi's view, Gongsun Long tended to forget the intersection between the extension of 'white horse' (WH^E) and the extension of 'horse' (H^E), that is, Gongsun Long tended to forget the set $\{WH^E \cap H^E\}$ (as shown below, Wen's account has obvious links to my account based on IL).

But Carnap's (1947) method works with state-descriptions, which

... are complete possible worlds. In them, all the atomic logical formulae (i.e., all the simple logical formulae without operators) that can be built are to be found, either affirmed or negated. State descriptions are consistent: they cannot include a formula and its negation. (López-Astorga, 2021, p. 170)

Within Carnap's (1947) framework, P is correct (the explanation below is based on that in López-Astorga, 2021). A definition provides a logical equivalence relation. If the definiens is true, the definiendum has to be true, and vice versa. Both can be false in a particular scenario or situation. What cannot happen is that the definiens is true and the definiendum is false, or that the definiens is false and the definiendum is true. Given these predicates,

$T =_{df}$ to be a horse

$U =_{df}$ to be white

we can build a well-formed formula in first-order predicate calculus capturing Mozi's idea, such as the following:

$$\forall x [(Tx \wedge Ux) \Rightarrow Tx]$$

Symbols ' $\forall$ ', ' $\wedge$ ', and ' $\Rightarrow$ ' in the previous formula express, respectively, universal quantification, conjunction, and conditional relation.

However, following Carnap's proposal of definition, we can build (1), which is also a well-formed formula in first-order predicate calculus, for P.

$$(1) \quad \forall x [(Tx \wedge Ux) \Leftrightarrow Tx]$$

The symbol ' $\Leftrightarrow$ ' in (1) refers to a biconditional relation.

Formula (1) cannot be true in all the state-descriptions (following the terminology in Carnap, 1947, it is not L-true, i.e., it is not true in all the state-descriptions). There are state-descriptions in which an element 'a' is a horse (Ta) but it is not white ($\neg Ua$; symbol ' $\neg$ ' indicates negation). In those state descriptions, (1) is false. Accordingly, P can be accepted (as mentioned, a similar explanation is to be found in, e.g., López-Astorga, 2021).

This appears to be an account supporting (A). However, it is easy to contest it. P does not actually provide that 'a white horse is not a horse' because the relation between 'white horse' and 'horse' is not biconditional. 'A white horse is not a horse' can mean that 'if an animal is a white horse, that very animal is not a horse'. This last interpretation would transform the biconditional relation in (1) into a conditional relation; that is, in the formula above, capturing Mozi's idea. If the relation in (1) is conditional, the formula becomes true in all the state-descriptions (i.e., L-true).

P and the theory of mental models

There are a lot more accounts about P in the literature (see also, e.g., Graham, 1990; Thompson, 1995). In this section, I will consider the account that the theory of mental models can propose. Based on the latter account, (A) does not hold.

Following the theory of mental models, the human mind works by reviewing possibilities or models (e.g., Espino, Byrne and Johnson-Laird, 2020). In the particular case of the conditional, individuals often think about the possibilities or models corresponding to the material conditional, that is, that of classical logic (see also Johnson-Laird and Ragni, 2019). Nevertheless, the theory of mental models does not assume the classical propositional logic. There are many differences between the latter theory and classical logic (see also, e.g., Johnson-Laird, 2023).

What is interesting for us is that, for instance, for a conditional such as (2),

- (2) If this item is a horse, then this item is white

As indicated, its possibilities or models can match those in classical logic. In other words, its possibilities or models can correspond to the possible situations in which (2) can be taken up, that is, to the possible situations in (3).

- (3) Possible (horse & white) &
Possible (not-horse & white) &
Possible (not-horse & not-white)

In the theory of mental models, sentential connectives lead to ‘conjunctions of possibilities’, such as (3) (see also, e.g., Khemlani, Hinterecker and Johnson-Laird, 2017).

Another example is that of a disjunction such as (4) (for disjunction in the theory of mental models, see, in addition, e.g., Johnson-Laird, Quelhas and Rasga, 2021).

- (4) Either this item is a horse or a donkey

Sentence (4) is an exclusive disjunction. Its possibilities in the theory are also, in principle, the same as in classical logic:

- (5) Possible (horse & not-donkey) &
Possible (not-horse & donkey)

The theory of mental models helps understand why people tend to consider P to be incorrect (the explanation below is based on that in López-Astorga, 2021). According to this theory, a fact refers to a circumstance in which the linguistic message only enables one possibility (e.g., Khemlani, Hinterecker and Johnson-Laird, 2017). This is the circumstance in which, for example, several individuals state that an item is a white horse. Only one possibility exists:

- (6) Possible (horse & white)

If (6) is true, (7) is also true.

- (7) Possible (horse)

In other words, if (6), which matches the first possibility in (3), holds, P (or Gongsun Long’s proposal) does not hold.

This explains why people tend to think that a white horse is a horse. If it is possible that a particular item is a white horse, it is also possible for that very

item, as Mozi claimed, to be a horse (as indicated for a similar account, see, e.g., López-Astorga, 2021).

Beyond the real scope of these proposals, the actual point they make, and their possible connections to other works cited in the present paper, Carnap's (1936; 1937) framework leads to rejecting interpretation (A) as well. One might think that the application of Carnap's (1936) reduction sentences to P makes little sense. Still, the latter framework also seems to show that (A) is not the case.

P and reduction sentences

Carnap's (1936) reduction sentences allow us to test whether certain properties can be related. If we suppose that 'Q' is a property, and we want to know whether properties 'R' and 'S' can be related to 'Q', following Carnap (1936), this reduction sentence can be built:

$$(8) \quad \forall x [Rx \Rightarrow (Sx \Rightarrow Qx)]$$

According to Carnap (1936), (8) can be a reduction sentence for 'Q' if and only if (IFF) (9) holds.

$$(9) \quad \exists x (Rx \wedge Sx)$$

In (9), '∃' represents the existential quantifier.

Formula (9) ensures that (8) is not true just by virtue of the fact that the antecedents in its conditionals are false. (8) can be deemed as a reduction sentence IFF there are cases of 'R' and 'S'.

As shown, if only Carnap's (1947) framework is considered, one might argue in favor of P. Nevertheless, if Carnap's (1947) framework is considered along with that in Carnap (1936; 1937), the situation changes. It is not hard to build a reduction sentence related to P. It would be sentence (10).

$$(10) \quad \forall x [Tx \Rightarrow (Ux \Rightarrow Vx)]$$

Formula (10) would be a reduction sentence for 'V'. To link (10) to P, it is only necessary to assume (11).

$$(11) \quad \forall x (Vx \Leftrightarrow \neg Tx)$$

By means of (11), 'V' is defined as the property of not being a horse. One more formula is needed. Restriction (9) implies, in this case, (12).

$$(12) \quad \exists x (Tx \wedge Ux)$$

Without (12), (10) is not a reduction sentence for 'V'. This is not a problem. It is a fact that there are horses, as well as it is a fact that some of those horses are white.

After this point, there are, at a minimum, two ways to show that (10) cannot be endorsed. One resorts to the concept of state-description. The other is based on a derivation in first-order predicate calculus.

In the first case, given the material interpretation of the conditional, (10) would be true in a state-description in which (13) were true.

$$(13) \quad \forall x (\neg Tx)$$

However, that state-description should not be considered. (10) is a reduction sentence for 'V' IFF (12) is the case. Hence, only the state descriptions in which (12) is true should be considered. In the latter state-descriptions, there is at least one element 'a' with the properties in (14).

$$(14) \quad Ta \wedge Ua$$

Formulae (10) and (14) allow deriving (15).

$$(15) \quad Va$$

Formula (15), in turn, leads to (16) by virtue of (11).

$$(16) \quad \neg Ta$$

Thus, (14) and (16) provide Ta and $\neg Ta$ in the same state-description. This is not possible: state-descriptions must be consistent.

Likewise, a derivation in first-order predicate calculus can show this. To check whether or not (10) is admissible as a reduction sentence, it is necessary to start the derivation taking (17) as an assumption.

$$(17) \quad \forall x [Tx \Rightarrow (Ux \Rightarrow Vx)] \wedge \forall x (Vx \Leftrightarrow \neg Tx) \wedge \exists x (Tx \wedge Ux)$$

By virtue of a process akin to the previous one, it is possible to come to (14), (15) and (16). Therefore, the contradiction between Ta and $\neg Ta$ is deduced again. This reveals that the initial assumption, that is (17), is false. The conclusion is (18).

$$(18) \neg\{\forall x [Tx \Rightarrow (Ux \Rightarrow Vx)] \wedge \forall x (Vx \Leftrightarrow \neg Tx) \wedge \exists x (Tx \wedge Ux)\}$$

Formula (18) implies that, at a minimum, one of the three formulae of the set consisting of (10), (11) and (12) is false. But the three formulae need to be true at once. If not, (10) is not a reduction sentence for ‘V’ linked to P. So, this account, based on first-order predicate calculus and reduction sentences, supports Mozi’s idea and rejects P.

P and IL

Both the previous explanation based on reduction sentences and the account below following IL can also have connections to the literature about P (especially, but not exclusively, for the case of IL, see, e.g., Wen, 1983, and Pang, 1979; English explanations in Sun, 2013). Nonetheless, as I will indicate, IL is linked to computing and artificial intelligence (e.g., Wang, 2023). At present, that is not trivial.

Within IL, (B) explicitly makes sense. This is relevant because, as shown in the previous sections, it is not difficult to offer accounts that reject (A). What seems to be necessary is to present accounts revealing why (B) requires to be assumed. IL is able to do that because there are only terms in IL (all my descriptions of IL in this section will follow that in Wang, 2013). The main copula in IL is the ‘inheritance copula’: ‘ $\rightarrow$ ’. This copula allows relating, for example, concepts such as ‘horse’ and ‘mammal’ (for an explanation of the inheritance copula in IL, see, e.g., Wang, 2013, Definitions 2.2 and 2.8):

$$(19) \text{Horse} \rightarrow \text{Mammal}$$

Let M^E be the extension of ‘Mammal’. What (19) provides is that

$$\text{‘Horse’} \in M^E \text{ and ‘Mammal’} \in H^I$$

Hence, IL seems compatible with Mozi’s intuitions.

From this logic, if we want to deny (B) and, hence, to state that ‘white horse’ has the same meaning as ‘horse’, we must assume (20):

$$(20) \text{Horse} \rightarrow \text{White-horse}$$

which means:

‘Horse’ $\in WH^E$ and ‘White-horse’ $\in H^I$

And (21):

(21) White-horse $\rightarrow$ Horse

which means:

‘White-horse’ $\in H^E$ and ‘Horse’ $\in WH^I$

IL includes a simpler form of expression. It also has a ‘similarity copula’: ‘ $\leftrightarrow$ ’ (for an explanation of the similarity copula in IL, see, e.g., Wang, 2013, Definition 6.1). The similarity copula expresses the inheritance relation from both the left term to the right term and from the right term to the left term. So, in IL, to say that ‘white horse’ has the same meaning as ‘horse’ can be expressed as (22).

(22) Horse $\leftrightarrow$ White-horse

This is because, in the latter logic, if (22) holds, (23) holds.

(23) (Horse $\leftrightarrow$ White-horse) $\leftrightarrow$ [(Horse $\rightarrow$ White-horse) $\wedge$ (White-horse $\rightarrow$ Horse)]

But IL has four theorems, Theorems 6.1, 6.2, 6.3 and 6.4, in Wang (2013), related to sentences such as (22) or (23). Except for 6.1, the other theorems are problematic if (B) is not the case.

Theorem 6.1 indicates that “Similarity is a reflexive, symmetric and transitive relation between two terms” (Wang, 2013, p. 78). Regarding the terms in (22) and (23), this does not necessarily imply difficulties. Even statements such as (24) and (25) can be endorsed.

(24) Horse $\rightarrow$ Horse

(25) White-horse $\rightarrow$ White-horse

The case of 6.2 is different. Following the latter theorem, given a statement such as (22), a statement such as (26) must be true regardless of the content of (22) and (26).

(26) (Horse $\leftrightarrow$ White-horse) $\Rightarrow$ (Horse $\rightarrow$ White-horse)

Statement (26) is not acceptable. H^E can be defined as follows:

$H^E =_{df} \{\text{White-horse, Color}_1\text{-horse, ..., Color}_n\text{-horse}\}$

where ‘Color₁-horse’ to ‘Color_n-horse’ refer to horses with colors different from white.

H^{E2} is a subset of H^E .

$$H^{E2} =_{df} \{\text{Color}_1\text{-horse}, \dots, \text{Color}_n\text{-horse}\}$$

That is,

$$H^{E2} \in H^E$$

the problem is (27).

$$(27) (H^E \in WH^E) \Rightarrow (H^{E2} \in WH^E)$$

This is not the case. The consequent in (26) requires the antecedent in (27). However, the consequent in (27) is not true. Horses with colors other than white are not elements in WH^E .

Theorem 6.3 has problems as well: it provides that, given a sentence such as (22), (28) should be true too, again, regardless of their contents.

$$(28) (\text{Horse} \leftrightarrow \text{White-horse}) \Leftrightarrow [\text{Horse} \in (WH^E \cap WH^I)] \Leftrightarrow [\text{White-horse} \in (H^E \cap H^I)]$$

Beyond what can be said about the third clause in (28), the problems in its second clause, that is, $[\text{Horse} \in (WH^E \cap WH^I)]$, are obvious. If, as indicated, the consequent in (27) does not hold, (29) holds.

$$(29) H^{E2} \notin WH^E$$

And (30) also holds.

$$(30) H^{E2} \notin (WH^E \cap WH^I)$$

Finally, Theorem 6.4 indicates that if (22) is the case, (31) has to be the case, once again, regardless of their contents.

$$(31) (\text{Horse} \leftrightarrow \text{White-horse}) \Leftrightarrow (H^E = WH^E) \Leftrightarrow (H^I = WH^I)$$

Beyond the problems of the third clause in (31), the second one, that is, $H^E = WH^E$, already reveals that (31) is false, since (29) is true.

Hence, in addition to the fact that Mozi seems to be right and Gongsun Long seems to be wrong, IL shows that the denial of (B) cannot be supported within

its framework. This is important because IL is an ideal logic. Non-Axiomatic Logic (NAL) is a limited-resource logic that can be derived from IL. The Non-Axiomatic Reasoning System (NARS) is a computer program trying to reason like human beings that can be derived from NAL (e.g., Wang, 2013). One might think that, as described above, the account of P based on IL depends on previous studies in literature. Nevertheless, the NARS tries to make inferences like us. So, it is not trivial whether it captures why sentences such as P seem non-sense sentences to people in daily life.

Furthermore, this paper has considered only a minimal part of IL. IL (and, therefore, NAL and NARS) have many more components. It even has resources to deem terms such as 'White horse' as compound terms such as 'White $\cap$ horse' (Wang, 2013). However, although the potential of IL is much higher, the characteristics of it described in this paper already suffices for an account of P.

Conclusion

Much has been written about P. A relevant point is the difference between (A) and (B). As shown, from a logic based on state-descriptions, it is possible to argue that (A) is correct. However, that account is easy to challenge. From a theory such as the theory of mental models, (A) is not valid. Something similar occurs from an explanation following the concept of reduction and philosophy of science. Reduction sentences appear to show us that (A) is unacceptable. The consequence seems to be the need for accounts in favor of (B).

I did not wish to discuss the scope or strength of previous proposals in detail. My aim has been to claim that from IL, a strong account sustaining (B) can be given. The latter logic is a term system. In it, concepts such as those of extension and intension are the key. Its inheritance relations make it clear that the extensions and intensions of terms such as 'horse' and 'white horse' allow noting that to negate (B) is not possible. That is what the theorems associated with the similarity relations in IL show.

Likewise, all of this opens up interesting perspectives. NAL depends on IL and NARS, in turn, depends on NAL. If NARS emulates human reasoning, it may reveal the cognitive reasons why people deny P.

Acknowledgments

Writing this paper was supported by PIA Ciencias Cognitivas, Centro de Investigación en Ciencias Cognitivas, Instituto de Estudios Humanísticos, Universidad de Talca and by project ‘Dialéctica virtuosa de la educación: De la fundamentación filosófica y antropológica a implementaciones concretas’, 15/I122 SAL127/19, UNMdP.

References

- Carnap, R. (1936) ‘Testability and meaning’, *Philosophy of Science*, 3(4), pp. 419–471. Available at: <https://doi.org/10.1086/286432>
- Carnap, R. (1937) ‘Testability and meaning—continued’, *Philosophy of Science*, 4(1), pp. 1–40. Available at: <https://doi.org/10.1086/286443>
- Carnap, R. (1947) *Meaning and necessity: a study in semantics and modal logic*. Chicago: The University of Chicago Press.
- Espino, O., Byrne, R.M.J. and Johnson-Laird, P.N. (2020) ‘Possibilities and the parallel meanings of factual and counterfactual conditionals’, *Memory & Cognition*, 48, pp. 1263–1280. Available at: <https://doi.org/10.3758/s13421-020-01040-6>
- Graham, A.C. (1990) *Studies in Chinese philosophy and philosophical literature*. Albany: State University of New York Press.
- Graham, A.C. (2004) *Later Mohist logic, ethics, and science*. Hong Kong: The Chinese University of Hong Kong Press.
- Havelange, M. (2020) ‘White horse not horse’: an analysis of modern interpretative approaches to Kung-sun Lung’s White Horse Discourse. MSc thesis. University of Amsterdam. Available at: <https://eprints.illc.uva.nl/id/eprint/1749/1/MoL-2020-06.text.pdf> (Accessed: 30 March 2025).
- Indraccolo, L. (2017) ‘The “white horse is not a horse” debate’, *Philosophy Compass*, 12(10), article e12434. Available at: <https://doi.org/10.1111/phc3.12434>
- Johnson-Laird, P.N. (2023) ‘Possibilities and human reasoning’, *Possibilities Studies & Society*, 1(1–2). Available at: <https://doi.org/10.1177/27538699231152731>
- Johnson-Laird, P.N., Quelhas, A.C. and Rasga, C. (2021) ‘The mental model theory of free choice permissions and paradoxical disjunctive inferences’, *Journal of Cognitive Psychology*, 33(8), pp. 951–973. Available at: <https://doi.org/10.1080/20445911.2021.1967963>
- Johnson-Laird, P.N. and Ragni, M. (2019) ‘Possibilities as the foundation of reasoning’, *Cognition*, 193, article 103950. Available at: <https://doi.org/10.1016/j.cognition.2019.04.019>

- Khemlani, S., Byrne, R.M.J. and Johnson-Laird, P.N. (2018) 'Facts and possibilities: a model-based theory of sentential reasoning', *Cognitive Science*, 42(6), pp. 1887–1924. Available at: <https://doi.org/10.1111/cogs.12634>
- Khemlani, S., Hinterecker, T. and Johnson-Laird, P.N. (2017) 'The provenance of modal inference', in G. Gunzelmann, A. Howes, T. Tenbrink and E.J. Davelaar (eds) *Computational foundations of cognition*. Austin: Cognitive Science Society, pp. 663–668.
- Liu, L. (2007) *Zai Yuyan Zhang Panxun*. Chengdu: Sichuan Daxue Chuban She.
- López-Astorga, M. (2021) 'School of Names and mental possibilities: is a white horse a horse?' *Studii de Lingvistică*, 11, pp. 167–175.
- Pang, P. (1979) *Gongsun Longzi Yanjiu*. Beijing: Zhonghua Shuju.
- Schumann, A. (2023) *Archaeology of logic*. Boca Raton: CRC Press and Taylor & Francis Group. Available at: <https://doi.org/10.1201/9781003316466>
- Schumann, A. (2024) 'Mozi's square of opposition and logemes as new logical approach', in J. Lemanski, M.W. Johansen, E. Manalo, P. Viana, R. Bhattacharjee and R. Burns (eds) *Diagrammatic representation and inference. Diagrams 2024*. Lecture Notes in Computer Science, 14981. Cham: Springer. Available at: https://doi.org/10.1007/978-3-031-71291-3_20
- Sun, Z. (2013) 'Cracking the white horse puzzle', *Journal of East-West Thought*, 3(3), pp. 97–106.
- Theobald, U. (2010a) 'Mozi 墨子 and the Mohist Canon (*Mojing* 墨經)', *ChinaKnowledge.de—An Encyclopaedia on Chinese History and Literature*. Available at: <http://www.chinaknowledge.de/Literature/Diverse/mozi.html> (Accessed: 30 March 2025).
- Theobald, U. (2010b) 'Gongsun Longzi 公孫龍子', *ChinaKnowledge.de—An Encyclopaedia on Chinese History and Literature*. Available at: <http://www.chinaknowledge.de/Literature/Diverse/gongsunlongzi.html> (Accessed: 30 March 2025).
- Thompson, K.O. (1995) 'When a "white horse" is not a "horse"', *Philosophy East and West*, 45(4), pp. 481–499. Available at: <https://doi.org/10.2307/1399790>
- Wang, P. (2013) *Non-Axiomatic Logic: a model of intelligent reasoning*. Singapore: World Scientific. Available at: <https://doi.org/10.1142/8665>
- Wang, P. (2023) 'The role of copulas in reasoning'. Technical Report #17. Temple AGI Team. Temple University. Available at: <https://cis.temple.edu/tagit/publications/TAGIT-TR-17.pdf> (Accessed: 30 March 2025).
- Wen, G. (1983) *Xian Qin Luo Ji Shi*. Shanghai: Renmin Chuban She.

Miguel López-Astorga, PhD in logic and philosophy of sciences from the University of Cádiz, Spain, is currently a professor at the University of Talca, Chile. He has a bachelor's degree in philosophy and education sciences (philosophy section) from the University of Seville and a teaching degree

for the high school level in philosophy from the University of Chile. López-Astorga's recent research focuses on human inferential activity. He is especially interested in the relations that can be found between syntax, semantics and pragmatics in that activity. He also addresses the possible cognitive potential of some ancient logics, particularly, Stoic logic.