

[illegible]

FALL 2026

EDITORIAL

HUGO RENNINGS | EDITOR-IN-CHIEF

Third time's a charm, right? Well, here is the third issue of the Illogician, the second of 2026 and the first under (mostly) new management!

When we started in the programme last year, we were the first group to be greeted by the Illogician: a magazine that lies somewhere between the bulletin of symbolic logic and a DIY meme up on the wall of the MoL room. Now it has become a fixture of the logic programme.

Since students come and go so often, the baton of this magazine must be passed on quickly. We have tried our best to learn all the tricks (and quirks) of the trade from the previous editors and despite the occasional illogical step, the magazine has made its way to you.

So take your time to solve the puzzles, read the articles written by your peers, and even glimpse your future in the Illogician. We hope you will be inspired to keep this project going by contributing, and then one year from now it will be you writing, editing and designing.

Issue 3: Fall 2026

The Illogician is a semiannual, semiserious periodical specialised in the areas of logic, language, and computation. It is published by Master of Logic students at the Institute for Logic, Language, and Computation of the University of Amsterdam.

EDITOR-IN-CHIEF

Hugo Rennings

CREATIVE DIRECTOR

Idske Roest

TECHNICAL SUPERVISOR

Giuliano Gorgone

LAYOUT

Josje van der Laan, Giuseppe Manes, Hugo Rennings, Idske Roest

ILLUSTRATIONS

Hugo Rennings, Idske Roest, Ziqi Wang

COVER

Ziqi Wang

EDITORS

Matteo Celli, Thiago Cocco Roque, Emke de Groot, David Kühnemann, Josje van der Laan, Bardo Maienborn, Giuseppe Manes, Alexandre Mazuir, Hugo Rennings, Mariana Rio Costa, Idske Roest, Isabel Trindade, Louise Wilk

CONTACT

Website: resources.illc.uva.nl/TheIllogician/

E-mail: hugo.rennings@student.uva.nl

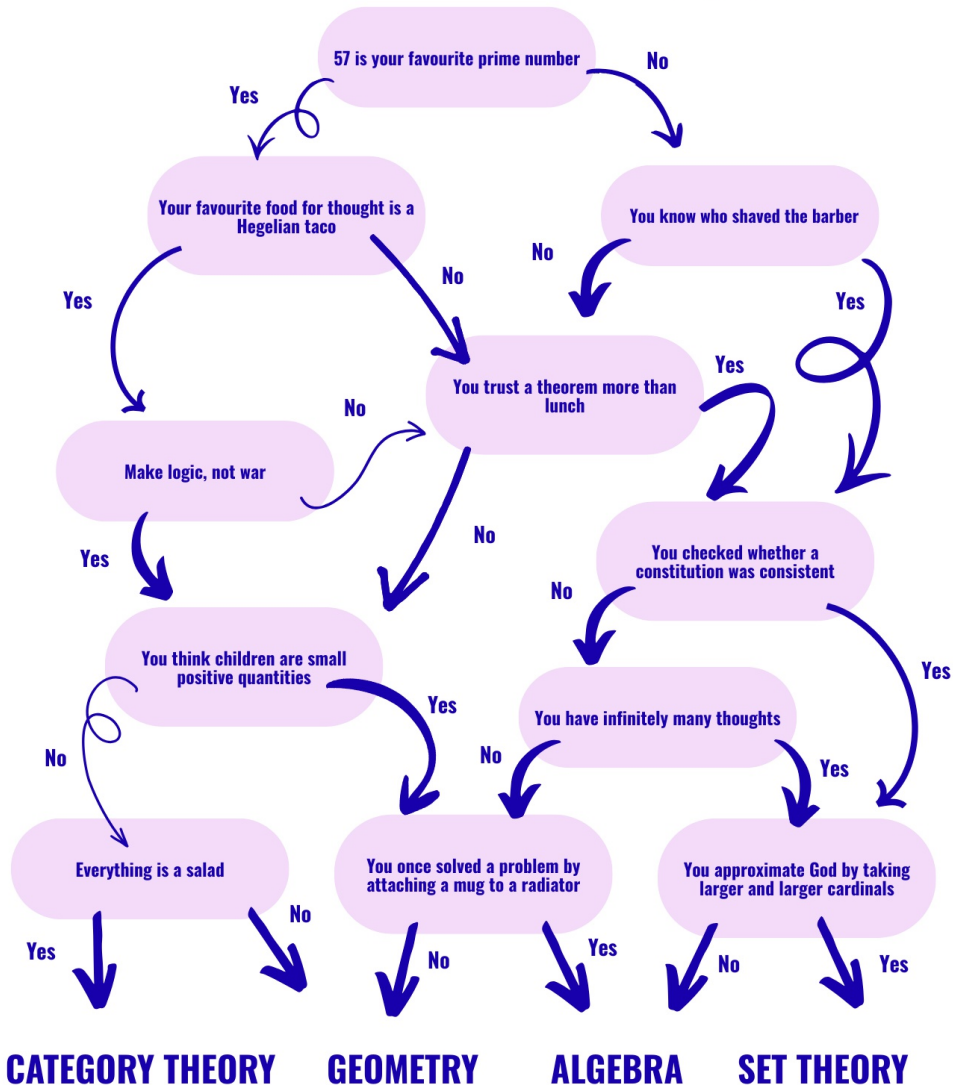
TABLE OF CONTENTS

LOGIC LORE AND REBUS	4
TOWARDS A TRUTHMAKER MODEL FOR MEINONGIAN LOGIC JIXIAN YU	6
THE HAT PUZZLE, OR HOW LOGIC SAVES LIVES TENYO TAKAHASHI	10
THE INFINITE HAT PUZZLE YIPU LI	14
CRYPTIC CLUES AND ILL-SING	17
DO LLMS DESERVE THIS MUCH ATTENTION IN ACADEMIA? FILIP REHBURG	19
NICK QUOTES AND WORLD CUP	21
ROTATING BREAD: UNDERSTANDING HIGH-DIMENSIONAL HOLES USING PASTRY MAX WEHMEIER	24
WHAT THEY DON'T TELL YOU ABOUT THE GMT-THEOREM JUSTUS BECKER	26
FUNDAMENTALS OF COMPUTATIONAL PERPLEXITY THEORY MATTEO CELLI	30
STRUCTURALIST FOUNDATIONS FOR MATHS MAGNUS KJAERGAARD	33
REAL NEWS: ABSTRACTOPHOBIA HUGO RENNINGS	38
(IL)LOGICAL HOROSCOPE MATTEO CELLI, JOSJE VAN DER LAAN, BARDO MAIENBORN, GIANNIS RACHMANIS	40
DIVERSITY COMMITTEE	43
ILLUNCHTIME IDSKE ROEST	44

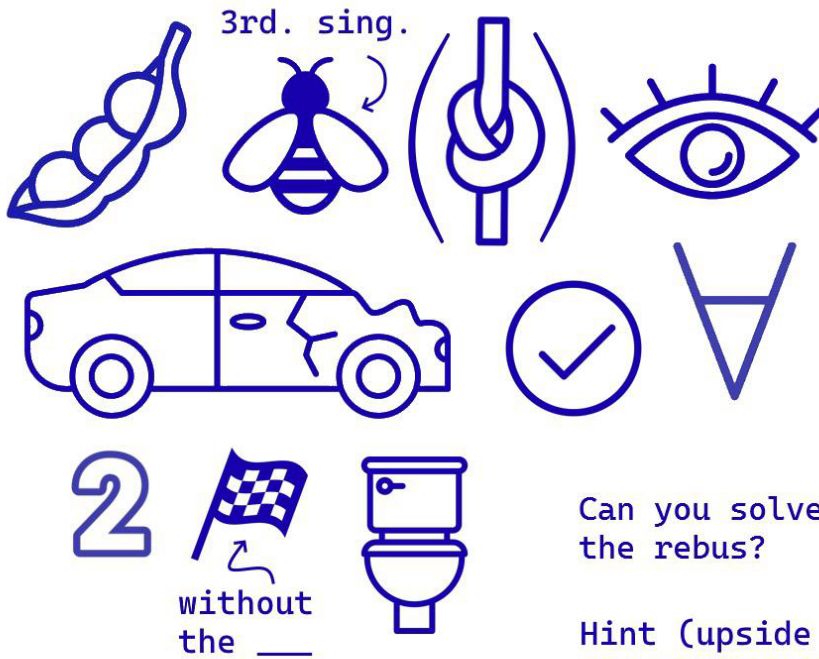
which LOGIC LORE defines you?

Follow the flow chart to find out where you belong in the landscape of logic! To maximize illogicality, all context has been omitted. However, if you ever happen to run into one of us, we will be able to restore the missing lore behind our beloved logicians, and we may be persuaded to speculate on why certain areas of logic attract the stories they do.

*Rodrigo Almeida
Josje van der Laan*



REBUS



Can you solve
the rebus?

Hint (upside down):
THIS IS RELATED

SCRIBBLE SPACE

TOWARDS A TRUTHMAKER MODEL FOR MEINONGIAN LOGIC

JIXIAN YU | PHILOSOPHY

When we say that Sherlock Holmes is a detective, or that the golden mountain is golden, or that the round square is impossible, we seem to be saying something. We can distinguish Holmes from Pegasus, ask whether the golden mountain has any properties beyond its name, and disagree about what makes the round square impossible in the first place. None of these objects exists. Yet they have properties, they get talked about, and they get reasoned about, and ordinary logic does not give us a clean way to handle them. **Meinongian logic** is the family of logics that does. This piece is a short introduction to it, and to a recent way of building one of these logics out of truthmakers rather than possible worlds.

WHAT MAKES A LOGIC MEINONGIAN

A Meinongian logic is one that allows reference to, predication of, and quantification over objects that do not exist.¹ The position is named after Alexius Meinong (1853–1920), and its defining thought is what Meinong called the *independence of Sosein from Sein*: an object's having of properties does not depend on its existing. Holmes is a detective whether or not he exists, the golden mountain is golden and a mountain whether or not there is any such thing, and our talk about them is about *those very objects*, not about something else.

1. A Meinongian logic should not be confused with a free logic. Both make room for terms that seem to name nothing, such as the golden mountain, but they do so in opposite ways. A free logic lets a singular term fail to denote, and it keeps the quantifiers committed to existence, so that they range over existing objects alone (Zalta 1983). On this approach there are no non-existent objects, and the golden mountain denotes nothing. A Meinongian logic takes every such term to denote an object, and it treats existence as a property that an object may have or lack (Jacquette 1996; Berto 2013). The domain over which its basic quantifier ranges then contains non-existent objects alongside existent ones. In the system sketched here every characterizing condition has a referent, and existence is given by a separate predicate. The existentially loaded quantifiers, defined from that predicate, range over the existent objects alone, as the quantifiers of a free logic do.

Formally, this has two consequences for the language. First, the domain D of a Meinongian model is a domain of *objects*, not of existents. Holmes and the golden mountain are in D alongside Napoleon and the Eiffel Tower. Second, existence is not built into the quantifier. Most Meinongian logics use two quantifiers, an *unloaded* one Σ ranging over all of D , and a *loaded* one $\exists$ ranging only over the existing objects. The two are connected through a primitive **existence predicate** E :

$$\exists x \varphi := \Sigma x (E(x) \wedge \varphi).$$

So existence becomes a real first-order property that some objects have and others do not, rather than something built into the meaning of the quantifier. To say there are non-existent objects is then to say $\Sigma x \neg E(x)$, which is no longer the contradiction it would be in classical logic.

That gives us the basic shape. What still needs settling is how non-existent objects get into the domain in the first place. They get there by being *characterized*. A characterizing condition, written $\alpha[x]$, is a formula that specifies a combination of properties, and a Meinongian logic associates with each such condition an object of D that satisfies it.

The golden mountain is the object characterized by “x is golden and x is a mountain.” Holmes is the object characterized by the relevant conjunction of Conan Doyle’s predications. The round square is the object characterized by “x is round and x is square.”

THE CHARACTERIZATION PRINCIPLE

Naively, the way to say this is: For every $\alpha[x]$, there is an object d such that $\alpha[d]$.

This is the **Characterization Principle (CP)**. It looks innocent, but Russell long ago observed that an unrestricted version trivializes existence. Take the condition “x is a golden mountain and x exists.” The naive CP introduces an object satisfying this condition, hence an existing golden mountain, hence a golden mountain exists. We have read existence out of thin air, by writing it into a definition.

The standard contemporary response is **Modal Meinongianism**, developed by Graham Priest (Priest 2016) and Francesco Berto (Berto 2013). The key move is to weaken CP so that a characterization is satisfied not necessarily at the actual world, but at *some* world, possibly a non-actual or impossible one. So we have a set of worlds $W = W^p \cup W^i$, partitioned into possible and impossible worlds, and: (QCP) For every $\alpha[x]$, there is an object d and a world w such that $\alpha[d]$ holds at w .

The golden mountain, then, is golden and a mountain at some non-actual possible world. The round square is round and square at some impossible world (you cannot have it at a possible one, since being round and square is impossible). And “the existing golden mountain” has its characterization satisfied at some non-actual world, which is no commitment to anything actual: the object exists *there*, not here. Russell’s worry is defused.

This is a workable Meinongian logic. It has a

domain of objects, including non-existents and inconsistent ones; it interprets predicates relative to worlds; it has unloaded and loaded quantifiers; and the QCP populates the domain.

STATES, PARTS, AND EXACT VERIFICATION

There is, however, a more discriminating way to build a semantics for a Meinongian logic, and that is what the second half of this piece is about. The shift is not particular to Meinongianism. It is part of the larger *hyperintensional turn* in philosophy: an account of content is hyperintensional when it distinguishes contents that are necessarily, or even logically, equivalent. Possible-worlds semantics is not hyperintensional, since two formulas true at exactly the same worlds get the same content. For many purposes this is fine. For some it is not. Theories of essence, of ground, of subject matter, and of belief have all needed a finer notion of content than worlds provide.

One leading framework here is **Fine’s exact truthmaker semantics** (Fine 2017). Instead of asking at which worlds ϕ is true, it asks what *exactly verifies* ϕ . A state s exactly verifies ϕ when it is wholly relevant to ϕ , with no irrelevant part. Rain exactly verifies “it is raining.” The state of rain together with the state of wind does not, because the wind is irrelevant.

Three features distinguish the framework from worlds semantics. First, content is **bilateral**: each formula has both a set of *verifiers* $\llbracket \phi \rrbracket$ and a set of *falsifiers* $\llbracket \phi \rrbracket^*$, with falsifiers carrying positive content of their own rather than being the absence of verification. Second, states have **part-whole structure**: a state can be a part of another, and two states can be fused into a third, $s \sqcup t$. The state of raining is a part of the state of raining-and-windy; their fusion is what makes the conjunction true.

Third, content is **compositional in this structural sense**:

$$\llbracket \varphi \wedge \psi \rrbracket^+ = \{ s \sqcup t : s \in \llbracket \varphi \rrbracket^+, t \in \llbracket \psi \rrbracket^+ \}.$$

A verifier of a conjunction is built by fusing verifiers of the conjuncts. Disjunction works analogously on falsifiers, and negation simply swaps verifiers and falsifiers. The clauses are uniform: they make no reference to worlds, because there are no worlds; there are only states.

The setting is hyperintensional because two formulas true at exactly the same worlds can have different exact verifiers, simply because they are made true by different things. So “ $2 + 2 = 4$ ” and “ $P \vee \neg P$ ” are true at the same worlds (all of them) but have different verifiers.

The connection between truthmaking and Meinongian objects is not new. Sendlak, for one, reads talk of possible worlds in terms of Meinongian states of affairs (Sendlak 2022). What follows takes a further step, and builds the Meinongian logic itself out of states.

A MEINONGIAN TRUTHMAKER MODEL

What I want to suggest in the rest of this piece is that this is the right setting in which to do Meinongian logic. The result is a **Truthmaker Meinongian Model (TMM)**. The components are a space of states S ordered by parthood, a *possible region* $P \subseteq S$ (everything outside it is impossible), a *concrete region* $C \subseteq P$ in which existence lives, a Meinongian domain D of objects, an actual state $s_0 \in C$, and a valuation that gives each predicate a verifier-falsifier pair and gives each characterizing condition a definite object as referent. There are no worlds.

The truthmaker setting is Fine’s and the Meinongian ideas are the tradition’s; the way of

putting the two together, the Truthmaker Meinongian Model, is my own proposal, and I work it out in full in my master’s thesis. It is still in progress, so I offer it here as a sketch rather than a settled result: please do not cite this piece, though any comment is most welcome.

Three things change relative to the world-based picture. First, the Characterization Principle now requires that the object assigned to a condition α actually *carries* verifying content for α . Some state in the model verifies α of the assigned object. This gives a *characterization term*, written $\iota x \alpha[x]$ and read “the object characterized by α ,” a definite referent, rather than only the existential statement that some object satisfies α .

Second, **impossibility is structural**. Take any predication $F(d)$. The single rule called Exclusivity says that if r verifies $F(d)$ and r' falsifies $F(d)$, then their fusion $r \sqcup r'$ lies outside P . Impossibility, in other words, is generated: it falls out of the predicate valuation and the fusion operation, rather than being stipulated by sticking sentences into arbitrary sets. The round non-round, characterized by $R(x) \wedge \neg R(x)$, has verifiers built from R -verifiers and R -falsifiers fused together, and by Exclusivity those fusions are impossible. We did not have to tell the model that they were. This is where the truthmaker model earns its keep against the modal one. Modal Meinongianism has to place the round square at some impossible world, but which impossible world, and which impossibilities exist at all, is left to stipulation, since its impossible worlds are just arbitrary sets of sentences. Here there is nothing to stipulate. The round non-round is impossible because Exclusivity makes the fusion of a verifier and a falsifier impossible, and that follows from the content of its condition alone.

Third, **objects are individuated by their full content**. Each object d has a *verification assignment*,

VA_d : a function that records, for every predication, what verifies and what falsifies it of d . Two objects are identical exactly when they have the same verification assignment. This is the formal counterpart of Meinong's *Sosein*: the totality of how an object is characterized, independent of whether it exists. Two distinct inconsistent objects can now differ structurally, not just by stipulation. The round square has verifiers built from R-verifiers and Sq-verifiers; the round non-round has verifiers built from R-verifiers and R-falsifiers. Different parts, different content. Different objects, by construction.

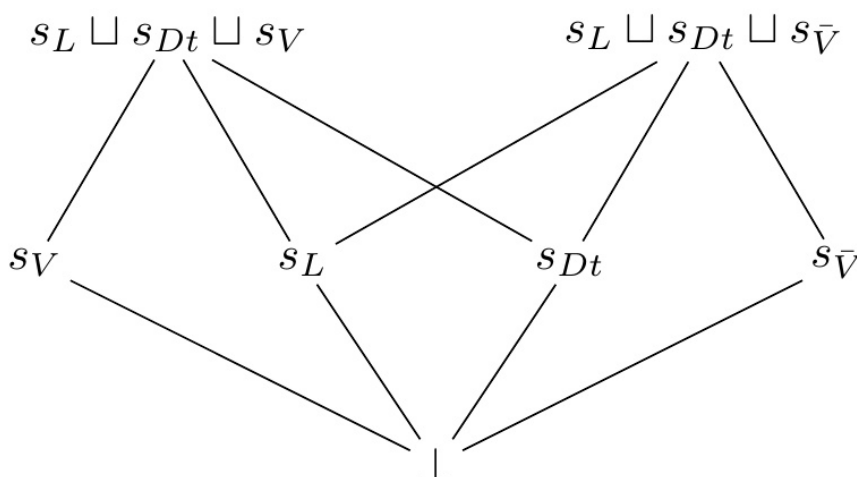
Take the following example for illustration (see the Figure at the end of the article). Consider Holmes and a stipulated near-twin, Schmolmes, characterized exactly as he is but for one change: where Holmes plays the violin, Schmolmes never took it up. Everything else, the Baker Street address, the detective work, the rest, they share. So the model gives them the same verifiers and the same falsifiers everywhere but on one predicate:

"plays the violin" has a verifier at Holmes and a falsifier at Schmolmes. Their verification assignments differ by exactly that one pair, and the identity criterion asks for no more than that. They are two objects, not one, even though a list of all their other properties would not tell them apart.

That is the kind of distinction a Meinongian logic needs, and it is one that the part-whole structure of states yields on its own. The Meinongian who once seemed to be populating their logic with creatures stitched together by stipulation now has them put together out of parts. These objects were always more sharply defined than the tradition allowed; we simply lacked a semantics with parts fine-grained enough to bring the differences out.

BIBLIOGRAPHY

For the bibliography, visit the Illogician website.



THE HAT PUZZLE, OR HOW LOGIC SAVES LIVES

TENYO TAKAHASHI | MATHEMATICS

PROLOGUE

In a dim room, 100 prisoners sat in various poses, wracking their brains. Everyone looked desperate, except for one prisoner who was strangely excited by the situation: the *logician*.¹

“We should use logic—” someone began.

“No, that’s useless!” another interrupted.

The logician considered this. “Well, not always.”

They were scheduled to be executed the next morning. The method was well-known. That was why they were gathered here, trying to save themselves or, in the worst case, as many lives as possible.

THE RULES

The execution proceeds as follows.

All 100 prisoners stand in a line on a staircase, each wearing either a black or a white hat. The distribution of hats is arbitrary: it could be all black, all white or any combination in between. The prisoners do not know it in advance.

Each prisoner can see the hats of all prisoners in front of them, but not their own hat or the hats behind them. Let us count the prisoners from the top of the staircase, as mathematicians often do, starting from 0. Thus, prisoner 0 can see 99 hats,

prisoner 1 can see 98 hats, and so on, until prisoner 99, who can see no hats at all.

Starting with prisoner 0, they must each guess the color of their own hat aloud, one by one. A prisoner survives if their guess is correct; a wrong answer costs their life. Everyone can hear the guesses that have already been made, but no one hears the judge’s verdicts. In other words, prisoner 23 knows what prisoners 0 through 22 said, but not whether any of them survived.

The prisoners are not allowed to communicate any extra information once the execution begins. They may say only “black” or “white” on their own turn. However, before the hats are placed on their heads, they may agree on a strategy that gives special meaning to those words. The judge does not know or care about such a convention; the judge only checks whether the spoken word matches the speaker’s own hat.

The prisoners are seeking the best possible strategy. A strategy is better if it guarantees more survivors in the worst case.

At first sight, the situation looks hopeless. Without any strategy, each prisoner can only guess randomly. In the worst case, no one may survive. Can we do better?

Try to think about the puzzle before reading on. You

1. I originally learned this puzzle from Ruiting Jiang, a friend of mine who also graduated from the Master of Logic program.

are in the same position as the logician, although perhaps with a more comfortable chair.

A FIRST TRY

A simple strategy improves the situation: every other prisoner sacrifices themselves by announcing the color of the next prisoner's hat.

Here is how it works. Prisoner 0 can see prisoner 1's hat, so prisoner 0 says that color aloud. Prisoner 1 hears this and repeats it, thereby saving themselves. Similarly, prisoner 2 saves prisoner 3, and so on. With this strategy, prisoners 1, 3, 5, ..., 99 all survive. The others may or may not. In the worst case, at least 50 prisoners survive.

"Not bad," the logician thought, "but can we do better than this?"

The weakness of this strategy is that it keeps starting over. Prisoner 0 gives useful information to prisoner 1, but prisoner 1's answer adds nothing new for prisoner 2, even though prisoner 2 can hear it. So prisoner 2 has to begin again by sacrificing themselves to save prisoner 3, and the same pattern repeats.

To improve the strategy, we should ask for more from each answer. Ideally, prisoner 1's answer should both save prisoner 1 and help prisoner 2. This suggests the following question:

What information can prisoner 0 give so that prisoner 1 can determine their own hat color, and prisoner 2 can also determine their own hat color by combining prisoner 0's message with prisoner 1's answer?

Again, it is worth pausing here before reading the next section.

A BETTER IDEA

The key is that prisoner 0 does not need to tell prisoner 1's hat color directly. Instead, prisoner 0 can describe the relationship between prisoner 1's hat and prisoner 2's hat.

Suppose prisoner 0 says "black" to mean "prisoners 1 and 2 have the same hat color", and says "white" to mean "prisoners 1 and 2 have different hat colors". This is allowed: prisoner 0 is still only saying one of the two permitted words.

Now prisoner 1 can survive. Prisoner 1 sees prisoner 2's hat. If prisoner 0 said that the two hats are the same, prisoner 1 says the color they see on prisoner 2. If prisoner 0 said that the two hats are different, prisoner 1 says the opposite color.

Prisoner 2 can also survive. Prisoner 2 cannot see prisoner 1's hat, but prisoner 2 can hear prisoner 1's correct answer. Together with prisoner 0's message, this tells prisoner 2 their own hat color.

This gives a strategy for groups of three. Prisoner 0 encodes whether prisoners 1 and 2 are the same or different; prisoner 3 does the same for prisoners 4 and 5; and so on. The prisoners can guarantee two survivors in each block of three. Since 99 is left over at the end, this guarantees at least 66 survivors.

This is a significant improvement over 50. However, the logician was still not satisfied.

"We saved 16 more lives," they thought, "but who wants to be prisoner 0, 3, 6, ..., or 96?"

The issues of the previous strategy are reduced but still remain. Prisoner 0 can see the hats of prisoners 1 through 99, but the strategy uses only the relation between prisoners 1 and 2. All the information about prisoners 3 through 99 is thrown away. Moreover, the information flow stops after

prisoner 2, so prisoner 3 must start a new little scheme.

What we want is a single stream of information running all the way down the staircase. Prisoner 0 should give one piece of information about the whole line. Prisoner 1 should use it to survive, and then prisoner 1's answer should help prisoner 2. Prisoner 2 should do the same for prisoner 3, and so on.

At this moment, an excellent and terrifying idea struck the logician. It took them a while to be brave enough to share it with the other prisoners.

This is the final warning: the next section gives the best strategy.

THE PERFECT SOLUTION, EXCEPT FOR THE SELFLESS HERO

The logician's idea is to use the *parity* of the number of black hats. Instead of describing many hats separately, prisoner 0 says whether the total number of black hats in front of them is even or odd.

Here is the agreed strategy:

- If prisoner 0 sees an even number of black hats among prisoners 1 through 99, prisoner 0 says: "black".
- If prisoner 0 sees an odd number of black hats among prisoners 1 through 99, prisoner 0 says: "white".

The words "black" and "white" could be swapped. The important point is that the first answer encodes one bit of information: even or odd. Prisoner 0's own survival is not guaranteed, but everyone else can use this initial parity information.

Let us see how prisoner 1 uses it. Suppose prisoner 0 said "black", meaning that among prisoners 1 through 99 there is an even number of black hats. Prisoner 1 counts the black hats they see among prisoners 2 through 99. If prisoner 1 sees an even number, their own hat must be white; if prisoner 1 sees an odd number, their own hat must be black. Thus prisoner 1 can determine their own hat color exactly.

The same reasoning works if prisoner 0 said "white", except that the target parity is odd rather than even.

Now consider prisoner 2. Prisoner 2 heard prisoner 0's parity announcement and prisoner 1's correct answer. Combining these, prisoner 2 can calculate what the parity must be among prisoners 2 through 99. Then prisoner 2 counts the black hats they can see among prisoners 3 through 99. The difference tells prisoner 2 whether their own hat is black or white.

This procedure continues down the entire staircase. Each prisoner knows the original parity, hears all previous answers, and sees all later hats. Therefore, on their turn, each prisoner from 1 through 99 can determine whether their own hat is the missing contribution needed to make the parity come out right.

Even prisoner 99, who sees no hats, is saved by the accumulated information. By the final turn, all other hats from 1 through 98 are known, and the original parity tells prisoner 99 whether their own hat must be black or white.

Thus, the prisoners can guarantee that at least 99 of them survive, no matter how the hats are arranged. Only prisoner 0 remains at risk.

This is optimal. Prisoner 0 has no information

about their own hat color. If prisoner 0 would say “black” in a given situation, the judge could have placed a white hat on prisoner 0 without changing anything prisoner 0 can see or hear before speaking; similarly for the other color. So no strategy can guarantee prisoner 0’s survival. At most 99 prisoners can be guaranteed, and the parity strategy achieves this bound.

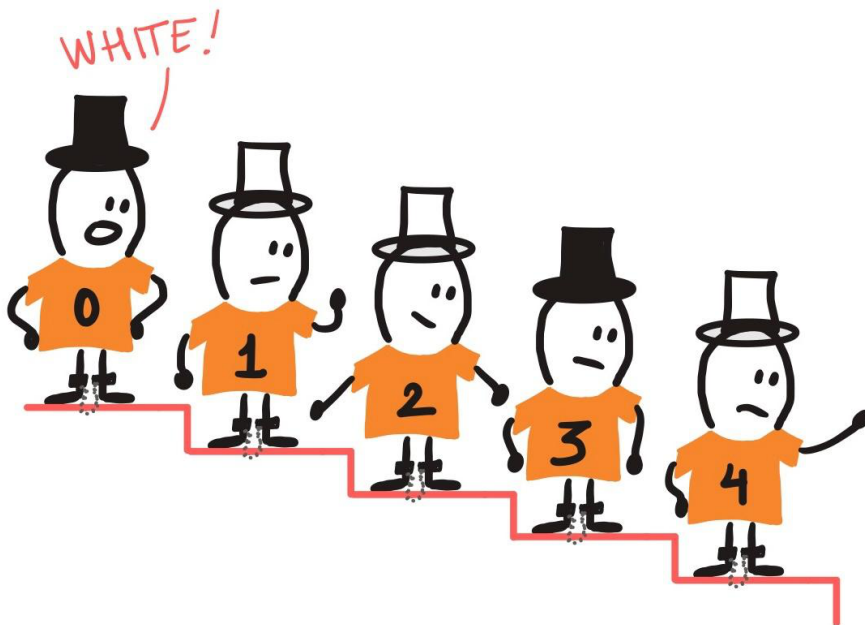
One obvious problem remained, and the logician had an obvious solution.

“Not many logicians have contributed to the real world,” said the logician. “I will stand at the top of the staircase.”

EPILOGUE, AND MORE PRISONERS

This closes our story. The strategy generalizes to any finite number of prisoners. If there are m prisoners for $m \geq 1$, the same parity argument guarantees that at least $m - 1$ survive in the worst case, which is optimal.

“But wait,” says a *mathematician*. “What if there are infinitely many prisoners?”



THE INFINITE HAT PUZZLE

YIPU LI | MATHEMATICS

PROLOGUE¹

“Not bad” the judge murmured reluctantly. He was annoyed by the logician’s triumph, but a sinister grin suddenly spread across his face. “Maybe, poor logician, there are more lives for you to save yet”

The judge opened the dungeon and revealed a silent line of infinitely many prisoners with panic written on their faces

“Let’s see if luck is on your side this time.”

THE RULES

The execution proceeds as follows.

Prisoners stand in a line on the staircase as in the last time, but this time, there are **infinitely** many of them, and they are ordered as the natural numbers, each wearing either a black or a white hat. The distribution of hats is arbitrary.

Each prisoner can see the hats of all prisoners in front of them, but not their own hat or the hats behind them. For instance, prisoner 0 can see all the hats but their own; prisoner 1 can see all hats but her hat and the hat on prisoner 0’s head, etc. As before, starting with prisoner 0, they must each guess the color of their own hat aloud, one by one. They are allowed to agree on a strategy before the execution begins.

At first sight, the situation looks dire. The strategy which the prisoners developed for the finite case ceases to work, for this time it may be the case that infinitely many prisoners are wearing black and infinitely many prisoners are wearing white. Announcing the parity of the number of black hats simply does not make sense.

But our logician has secret weapons in his tool box. He decides to open his toolbox and take out one of his most powerful secret weapon:

Axiom of Choice

For a collection $(X_i)_{i \in I}$ of non-empty sets, there is a function $F : \{X_i \mid i \in I\} \rightarrow \bigcup_{i \in I} X_i$ s.t. $F(X_i) \in X_i$ for all $i \in I$. This function is known as the *choice* function.

“With the help of the Axiom of Choice, I might work out something.”

We invite you to think about the puzzle before reading on.

HELP FROM THE AXIOM OF CHOICE

Let us be a little bit more formal here. We may model a distribution of the black and white hats as a function from the natural numbers to $\{0, 1\}$. Let $f : \mathbb{N} \rightarrow \{0, 1\}$, $f(n) = 0$ means that the n -th prisoner is wearing a white hat, and $f(n) = 1$ if his hat is black.

1. This is a follow up of the last article. We strongly encourage the reader to read the last article before venturing into the bizarre journey in the realm of infinity. 2. The informed guess is that no strategy that guarantees finitely many errors exists without using Axiom of Choice. But I do not know a proof (a proof would usually involve working in the Solovay model, readers may refer to my article in the first issue of Illogician for a introductory reference). However, if you do manage to develop such a strategy without Choice, I would be really surprised!

The idea is to define an equivalence relation E_0 on the collection of functions: we say for $f, g : \mathbb{N} \rightarrow \{0, 1\}$, let fE_0g iff they are different on finitely many natural numbers. i.e. $\{n \in \mathbb{N} \mid f(n) \neq g(n)\}$ is finite. Check for yourself that this is indeed an equivalence relation! We are viewing two distributions of hats as E_0 equivalent if all but finitely many prisoners agree on the colour of their hats in the two distributions.

Let $[f] = \{g : \mathbb{N} \rightarrow \{0, 1\} \mid fE_0g\}$ be the equivalence class of f . Thus, we may use the Axiom of Choice to obtain a choice function F that chooses an element from each equivalence class. i.e. $F([f]) \in [f]$, in particular, this means that $F([f])E_0f$.

Now we are ready to describe a strategy such that when carried out, only finitely many prisoners would guess wrong. First, let's observe a nice fact: based on what they see, each prisoner is able to decide which equivalence class the entire distribution of hats is in. Say the distribution is f , then our prisoner n knows all but the first n values of f . Thus he is able to determine the unique equivalence class, namely $[f]$, such that for any $g \in [f]$, g and f are different only on finitely many natural numbers beyond n .

Let the prisoners agree on a choice function F that chooses an element from each equivalence class in advance. Having decided on their own which equivalence class the distribution of hats is in, say it's $[f]$, prisoner n makes the guess $F([f])(n)$.

This is a strategy that guarantees finitely many mistakes! Why? Let's say the distribution of hats is f . For prisoner n , he does not know that the distribution is f , for he lacks the information of the first n many hats, but what he does know is which equivalence class the distribution is in. Thus, the collective guesses of the prisoners is modeled by

$F([f])$, i.e. the n -th prisoner guesses $F([f])(n)$. Now as F is a choice function, $F([f]) \in [f]$ and thus $F([f])$ and f are different only on finitely many natural numbers. meaning that all but finitely many prisoners got their guess right. This is what we wanted.

In case this strategy looks bizzare to you, you are not alone. It involves a blatant application of the Axiom of Choice, which is known for producing weird objects. We will discuss the role of Choice in the last section. But before that, the logician did manage to find a strategy that guarantees finitely many errors, this is of course an achievement. But can we do better? A hint: combine the idea developed here with the optimal strategy from the finite case.

This is the final warning: the next section gives the best strategy.

THE PERFECT SOLUTION, EXCEPT FOR THE SELFLESS HERO

The idea is to let the first prisoner announce the parity of the number of difference between what he sees and what $F([f])$ is. i.e. if $\{n \mid n > 0 \text{ and } f(n) \neq F([f])(n)\}$ is odd, he announces white and otherwise he announces black.

This is a brilliant combination of the ideas. With this information, all prisoners except for the first one are able to reason inductively based on what they heard to make the correct guess of the hat. Look back on Tenyo's article and convince yourself that this is truly the case.

WHAT IF WE HAVE NO CHOICE?

Surprisingly, the logician managed to find a strategy that guarantees at most one error, even in the infinite case. However, the logician's solution

used the Axiom of Choice in a significant way.

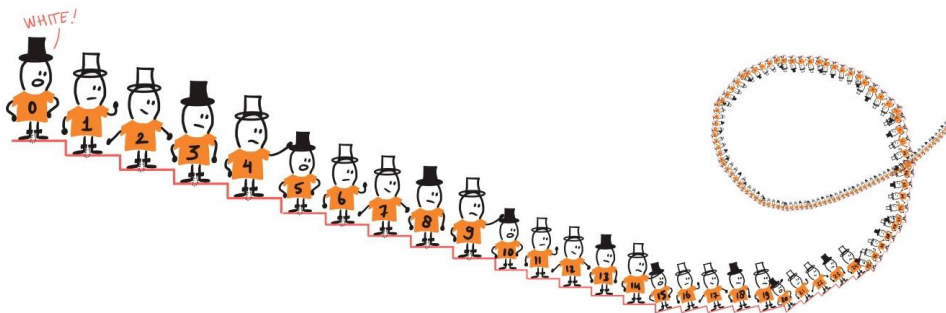
Is there a strategy **without** using the Axiom of Choice that guarantees at most one error? Unfortunately, Geschke, Lubarsky and Rahn showed that this is not the case (Lubarsky 2015). They found a set theoretic universe where the Axiom of Choice fails, and there is no strategy for the prisoners that guarantees at most one error. In other words, in a Choiceless world, it may well be the case that no strategy guarantees at most one mistake.

Still, interesting questions remain. Is there a strategy without using the Axiom of Choice that guarantees at most k ($k > 1$) errors? To further loosen the requirement, is there a strategy without using the Axiom of Choice that guarantees finitely

many errors, i.e. a strategy such that when carried out, only finitely many prisoners get it wrong? It is clear that our optimal strategy satisfies the requirement, but the interesting question is whether such strategies exist without Choice. If the executioner deprives us of the right of using Axiom of Choice, what would be the optimal strategy? Personally, I do not know an answer to these two questions. These are for you, the illogicans, to figure out!²

BIBLIOGRAPHY

Lubarsky, Robert S. 2015. "Choice and the Hat Game." *Mathematical Logic Quarterly* 61 (1-2): 68–80.



CRYPTIC CLUES

**MARIANA RIO COSTA, BARDÓ MAIENBORN, DAVID KÜHNEMANN,
LOUISE WILK, HALLIE WAGNER, IDSKE ROEST**

Cryptic clues originated in the early 20th century Britain as an evolution of regular crossword puzzles. These crosswords are called cryptic for the fact that the clues include some kind of wordplay or hidden meaning, and that the answer usually does not match the plain reading of the clue. This (often more difficult) variety of crosswords gained popularity in the UK and beyond, with cryptic crosswords appearing regularly in major newspapers like The Times and The Guardian.

In their modern format, cryptic clues generally include the following:

1. A definition: Found at the beginning or end of the clue, the plain reading of this part by itself will describe the answer to the puzzle, just as a clue in a regular crossword. However, this description will often only match the answer in a broader, humorous, or more roundabout way and hence the definition itself is often not enough to solve the puzzle.
2. Wordplay: The other part of the clue will form some sort of wordplay that, when solved, also yields the answer to the puzzle. The wordplay itself can usually itself be divided into wordplay indicator, words that hint at what kind of wordplay needs to be applied, and fodder, i.e. words on which the wordplay acts.

EXAMPLE

A colouring that is sound in all possible worlds. (4)

Here, the wordplay is "that is sound in all possible worlds." The first indicator, sound, is a homophone indicator, and tells us to find a homophone of something in the clue. The second indicator is "in", which is a hidden word indicator, and tells us that we need to find a hidden word in "all possible worlds". This hidden word is *blew*, which has as a homophone *blue*, which is a colouring.

CLUES

- | | |
|---|--|
| 1. First-order, second-order, intermediate modality, what kind of exam is this? (4) | 4. Teacher stresses core proof is all messed up. (9) |
| 2. Lowkey, Kripke sounds essential for a good life? (4) | 5. Here we think it illicit to start without wiping your eyes, you hear? (4) |
| 3. Headless hamster returns furious to the city. (9) | |

ILL-SING

THE TYPE THEORY ANTHEM

Auto. 1 2 3 4

Sing. 1 2 3 4

From sequent calculus To Gentzen proofs we span "Gentzen" for the likes of ours

Reach for the stars Reach for the boxes And you will make the cut

From sequent calculus
 To Gentzen proofs we span
 "Gentzen" for the likes of ours
 Reach for the stars
 Reach for the boxes
 And you will make the cut
 In all possible worlds
 We decorate our homes
 Let housing be decidable
 Reach for the stars
 Reach for the boxes
 And you will make the cut
 We strongly normalise
 The right to one's own type
 Let love be what we want alone
 Reach for the stars
 Reach for the boxes
 And you will make the cut
 And through hardships and luck
 Soon enough we will rise
 To levels of a constructor
 Reach for the stars
 Reach for the boxes
 And you will make the cut
 Reach for the stars
 Reach for the boxes
 And you will make the cut

Brought to you (unknowingly) by
 the involuntary ILLC choir:

Giuliano Gorgone
 Navid Kianfar
 Magnus Kjærgaard
 Lois Lagerweij
 Frédérique Laliou
 Dionysis Yusofy

Orchestrated and rewritten by
 Alexandre Mazuir

DO LLMs DESERVE THIS MUCH ATTENTION IN ACADEMIA?

FILIP REHBURG | COMPUTATION

While the first public release of ChatGPT, which kicked off the AI¹-boom, happened less than four years ago, Large Language models (LLMs) now seem to have permeated every part of our digital existence. While this is indubitably harmful for the environment and worrisome for a host of other reasons, this article focuses on the impacts of LLMs on academia. Investigations into the capabilities of LLMs are becoming more numerous, leading to a general waste of time and effort, publications of low impact, and may even hold back more rigorous research.

But first, what is a LLM? The increase in performance on language tasks exhibited by LLMs can mainly be attributed to two factors: the Transformer architecture and the size as measured in number of parameters. Transformers implement attention, which allows the model to process all words in the input in conjunction, instead of one by one. This permits more accurate modeling of the interplay between words in a sentence, paragraph, or text. Moreover, LLMs are trained on corpora of previously unheard of sizes and have billions of weights. It cannot be discounted that there is a direct correlation between model size and performance, which AI companies heavily indulge when spending millions of dollars on compute resources to train a single model. I want to alert you to the fact that this was highly uncommon

before the AI-boom.

As is commonly known, LLMs have a number of problems, including their ability to confidently present hallucinated opinions as fact, the immense power and water consumption of data centers, copyright issues when companies nonconsensually scrape protected materials from the internet, and I could keep going on. Of particular note are the ties of AI companies, datacenter providers and hardware manufacturers with each other, as well as in the financial sector and government. While chat bots are generating losses, the companies who provide them gain trillion dollar valuations and make billion dollar deals with each other. This poses a significant risk to the stability of the economy.

Given the drawbacks of LLMs, why are they becoming the subjects of investigation across academic fields? There are many articles asking “Can LLMs solve X?”² or reviews comparing the abilities of different LLMs. I think that such research categorically misunderstands the capabilities of LLMs and language modeling in general. Further, such publications do not seem a worthwhile investment of grant money and research time.

It is important to mention that LLMs present a generational performance improvement in

1. Please note that I am deliberately misusing the term here. AI is the well-established field of research into intelligent behavior of machines. LLMs fall under the umbrella term of AI (Machine Learning, Neural Networks, NLP) but are not representative of the field. In a case of semantic broadening it has become common to refer to LLMs as AI.

2. I encourage you to search “Can LLMs” on Google Scholar yourself.

language modeling tasks including text generation, machine translation, sentiment analysis and named entity recognition. Further, since LLMs are trained on vast datasets including repositories of knowledge (e.g., Wikipedia), it is to be expected that some of this knowledge is absorbed by the model and reproduced when generating text. LLMs are not trained to perform reasoning, however often Sam Altman and Dario Amodei tweet that superintelligent AGI is around the corner. It is, however, more reasonable to view the release of the Transformer architecture as a generational improvement, while interpreting the work of AI-companies as iterative gains since. It is to be expected that real AGI will take coordinated, principled efforts by the academic community around AI for years or decades to come.

I advise researchers to keep realistic expectations with respect to the capabilities of LLMs in mind when formulating hypotheses. As an example, while LLMs are capable of generating output which reads as legal text to the layperson, LLMs are not LawyerAI™ and should not be used to generate legal advice. There is now a plethora of investigations asking naively whether LLMs can solve complex tasks in various domains. Frustratingly, this often even involves data from other modalities such as image or video data. Such research commonly disregards established results and methods, while glancing over relevant nuances of the field. The answer is too often “yes” for toy examples, but “no” for cases which experts regard as interesting and complex.

I support the publishing of negative results. However this does not apply, when the result is a forgone conclusion. When such research realizes that LLMs cannot solve a highly complex problem, for which specialized methods already exist or are

in active development, they often do not offer meaningful improvements. Finding out that LLMs cannot perform a complex task, such as solving a challenging reasoning problem, is not research, but to be expected due to their nature as language models and does thus not present a result in itself.

A respectable publication should not just note that LLMs can't solve the issue, but provide a rigorous, principled, and well-researched method addressing the shortcoming. Otherwise we risk peer-reviewed research becoming a review outlet for private, closed-source, closed-weights models; that is to say: product-testing. This will lead to reproducibility problems and publications being outdated quickly, since LLMs are updated frequently and the guardrails which are placed around chat bots are not made public.

In general, the rate of submissions to journals have multiplied since the start of the AI-boom. Amongst these submissions are many written partially or entirely by LLMs, studies investigating the capabilities of AI, and methodologies relying on code written almost entirely by AI. Such submissions do not represent good faith contributions to academia, but aim to play the system to increase measurables. This overloads the unpaid volunteer peer reviewers. This is bad, since it slows down the review process and inhibits the publication of valuable contributions. Unfortunately, it also increases the already high pressure to publish and review faster and faster, creating destructive incentives for good-faith researchers.

It seems that the marketing of large AI companies has convinced not just users, but grant givers and researchers that LLMs are the future of intelligence. This is why we are suddenly seeing LLMs in every aspect of the academic process. This

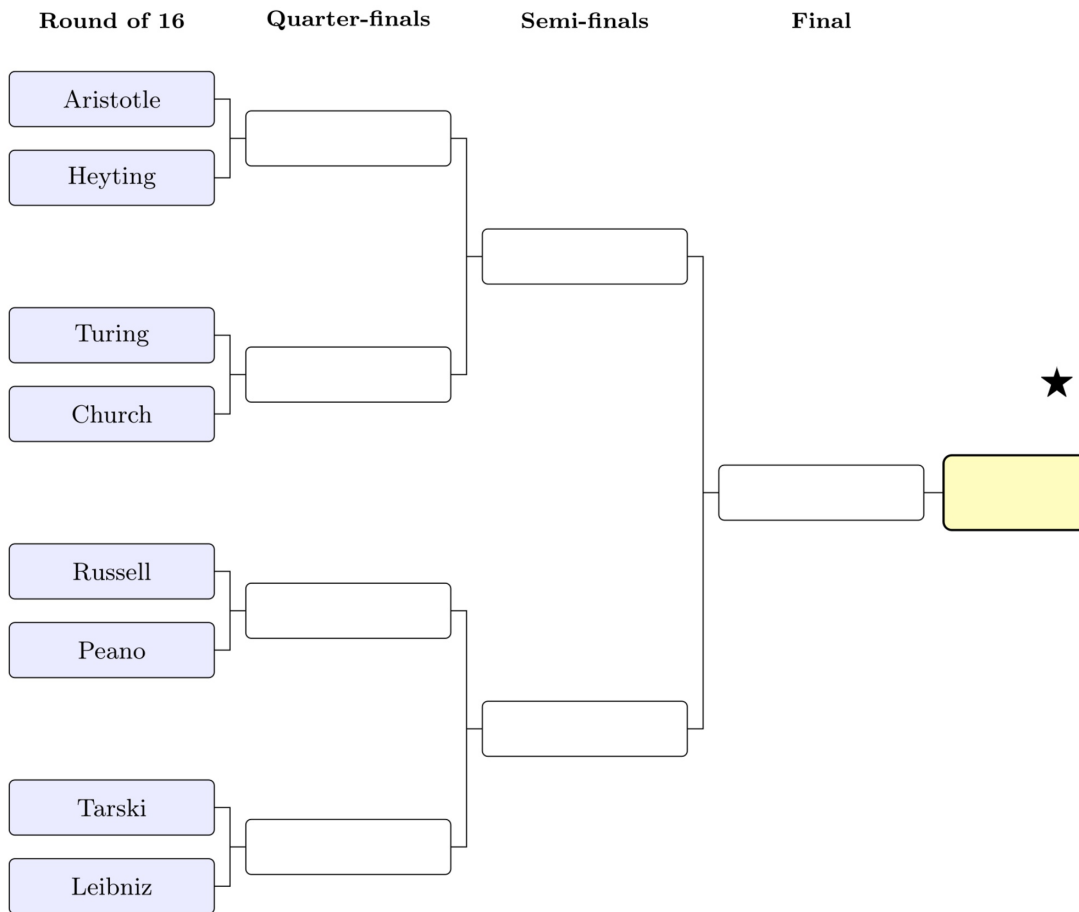
article is meant to point out the problems of using LLMs in academia. I urge you the reader to stay true to your field of expertise and see LLMs as what they are: tools to process language. In this light, I recommend you to continue exploring (hopefully LLM-free), well-founded hypotheses on the basis of

prior, principled research. Lastly, I hope to convince you that when the AI-hype settles and realistic views of LLMs become more commonplace in the academic community, it will be the quality of your research that matters, not the quantity of papers you published with "AI" in the title.

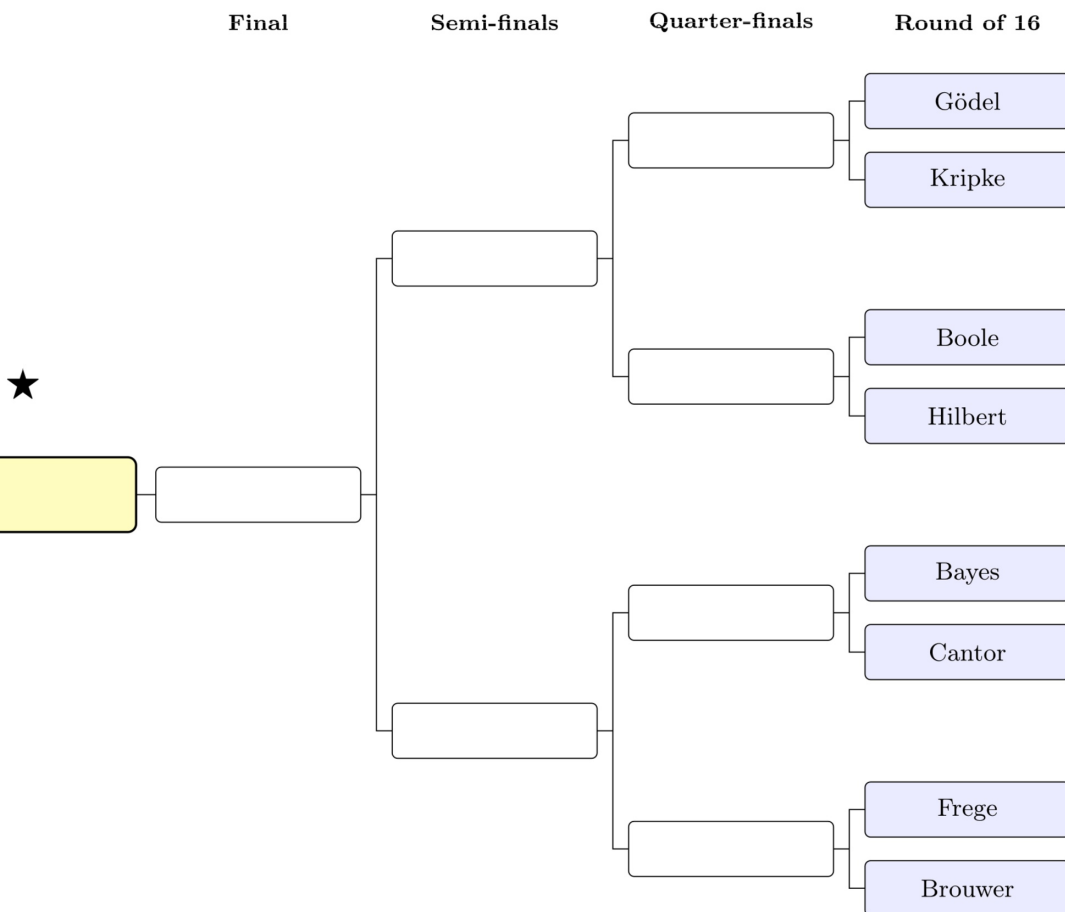
NICK QUOTES BINGO

In our proof, this is good news. In life, I don't know if this is good news...	STONE IS NOT A BRICK  BUT A LAST NAME	 <u>Jazz</u> . You might not even know what it is. It's a music genre for old people 
[about point-free topology] I have a friend that doesn't like it, calls it pointLESS topology	IF YOU SINNED ONCE YOU ARE A SINNER	I'm giving you a DRAMA and then I'll try to solve it
Forget about modal logic COMPLETELY	I promise I won't draw all the finite posets	His proof was shot, but at least he shot two other famous logicians

Who is The Greatest L



t Logician of All Time?



ROTATING BREAD: UNDERSTANDING HIGH- DIMENSIONAL HOLES USING PASTRY

MAX WEHMEIER | MATHEMATICS

There is a joke that topologists cannot distinguish between a bagel¹ and a coffee mug. The reason for this is that topologically speaking, they are equivalent (or homeomorphic to use the precise term). If we had a coffee mug that was made out of modelling clay, we could collapse the “inside” of the mug and be left with just the handle, i.e. a bagel. In a sense, they both just have one hole, the one in the middle/handle. Therefore, they are topologically equivalent. We call this shape a torus of genus 1.



ONE-DIMENSIONAL HOLES

Topologists consider this hole to be one-dimensional. Intuitively, this may seem confusing: “Inside” the hole, there is a disk missing, i.e. a two-dimensional manifold. By trying to give the hole a dimension from this approach, we are

characterizing it by what is missing. For this approach to work, the topological space has to be contained in another bigger space. For the bagel, this is $\mathbb{R}^3$. Furthermore, the containing space should not contain any holes itself, because if something is already missing in the bigger space, we cannot detect it in the smaller one. To avoid the issue of needing this containing space, we use a different idea to find holes: We cast out nets and try to catch them. (This analogy is due to (Kreck 2010).) For one-dimensional holes, our net is the circle S^1 , which is a one-dimensional manifold. Then we name the hole after the net needed to catch it instead of what is missing. We cast the net by mapping it into the space we are investigating. If we can collapse it to a single point with a continuous map, then our net came back empty. However, if we cannot do that, then something got stuck in our net and we have caught a hole. If we cast this net around the hole in the middle of the bagel or the handle of the mug, then we cannot collapse it, meaning we have found a hole. A little bit more checking shows that this is the only way we can do this². So they both have exactly one hole of dimension one. If we look at a filled ball, we see that all nets can easily be collapsed to a point by pulling them to the center. However, the same can

1. The usual formulation of this joke involves donuts, but they are unhealthy and this gave me a reason to make bagels.

2. Technically, a torus of genus 1 is hollow, so we can also cast the net around the vertical crosssection, finding another hole. However, the bagel and the mug are not hollow, so we can collapse this net.

be done on the hollow sphere by imagining the circle like a rubber band on the ball, which pulls itself tight. So they both do not have any one-dimensional holes, although there is a clear hole in the hollow sphere. Similarly, after cutting open a bagel, we see that there are more holes inside. Since we just concluded that a bagel contains only one one-dimensional hole, this must be a different kind of hole.

TWO-DIMENSIONAL HOLES

For one-dimensional holes, it seemed like there was a disk missing. Now the hole in the hollow sphere and the ones in the bagel seem spherical, so of dimension three. Earlier our intuition was one dimension off, so let's use a two-dimensional manifold to try and catch this hole. Then our net is the hollow sphere S^2 . If we map it to the filled ball, we see that we can still just pull everything to the center, collapsing it. Thus there is no two-dimensional hole here. If we map it to the hollow sphere though, we cannot collapse it anymore. If we imagined again it was made out of rubber, it would still not be able to pull itself tight without ripping it somewhere. This is similar to a filled balloon: It wants to collapse, but this is not possible, until you puncture it somewhere. Thus the hollow sphere has indeed a two-dimensional hole. We can map S^2 in the same way around all the little holes in the bagel and see that they are also two-dimensional holes.

HOLES IN HIGHER DIMENSIONS

We can also use the hyperspheres S^n to characterize holes of even higher dimensions. For example, we can use the three-dimensional hypersphere S^3 to detect three-dimensional holes. But since these can only appear in manifolds of dimension four or higher, they are hard to imagine

and are not useful for baking or life in general. One fun example actually comes from the other direction: What if we consider S^0 ? By everything we have seen, this should show us zero-dimensional holes. Since $S^0 = \{-1, +1\}$, we can collapse it in every path-connected topological space, because then we have a path between every two points by definition and can move -1 to $+1$. A space consisting of two disconnected parts, like two separate balls, has a zero-dimensional hole, since we can place $+1$ on one ball and -1 on the other. Then there is no way to bring them together. Thus, zero-dimensional holes characterize the path-connected components in a topological space.

TOPOLOGICAL BAKING

As we have seen, a bagel contains many small two-dimensional holes. The dough just after mixing however is solid and contains none or just a few holes. After fermenting for a while, there are more. Since a continuous deformation can never create new holes, the map that takes the mixed dough to its risen state is not continuous. So in case someone tells you that in reality everything is nice and smooth, you can answer that the rising deformation is not even continuous. Personally I am happy about this, because who wants to eat non-risen bread all the time?



BIBLIOGRAPHY

Kreck, Matthias. 2010. *Differential Algebraic Topology. From Stratifolds to Exotic Spheres*. Vol. 110. Grad. Stud. Math. Providence, RI: American Mathematical Society (AMS).

WHAT THEY DON'T TELL YOU ABOUT THE GMT THEOREM

JUSTUS BECKER | PHILOSOPHY, MATHEMATICS

SOME BETTER-KNOWN FACTS

Let me begin this article like a good book on mathematical logic: with Kurt Gödel and a definition. In a short paper from 1933 titled “Eine Interpretation des intuitionistischen Aussagenkalküls” (An interpretation of the intuitionistic propositional calculus), Gödel came up with the following translation from intuitionistic propositional logic IPL into modal logic (Gödel 1933a).¹

$$\begin{aligned} tr^G(p) &= \Box p \\ tr^G(\neg A) &= \Box \neg tr^G(A) \\ tr^G(A \rightarrow B) &= \Box(tr^G(A) \rightarrow tr^G(B)) \\ tr^G(A \star B) &= tr^G(A) \star tr^G(B), \\ &\text{with } \star \in \{\wedge, \vee\} \end{aligned}$$

The motivation behind this interpretation was to capture the constructive meaning of intuitionistic logic, which Heyting had introduced as a complete formal system only three years earlier (Heyting 1930). Gödel added the $\Box$ modality to classical logic as a rather informal operator, for which he came up with a somewhat reasonable axiomatisation:

$$\begin{array}{l} \Box A \rightarrow A \quad (\Box(A \rightarrow B) \wedge \Box A) \rightarrow \Box B \\ \Box A \rightarrow \Box \Box A \quad \frac{A}{\Box A} \end{array}$$

As readers familiar with modal logic will have recognised, this logic is nowadays known as the classical modal logic S4. In the same paper, Gödel shows that this interpretation of intuitionistic logic inside classical modal logic is sound, i.e.: if $\text{IPL} \vdash A$ then $\text{S4} \vdash tr^G(A)$. This was rather easy and could be

done using only the axiomatic Hilbert systems of intuitionistic propositional logic and S4. However, the question of completeness – while conjectured to be true – was left open, until proved in 1948 by McKinsey and Tarski (McKinsey and Tarski 1948). This is also why the following is known as the Gödel-McKinsey-Tarski Theorem (GMTT):

Theorem 1. $\text{IPL} \vdash A$ iff $\text{S4} \vdash tr^G(A)$.

This still left the question about *first-order* intuitionistic logic open, which was proved a good few years later in 1963 by Rasiowa and Sikorski (Rasiowa and Sikorski 1963) who proved full faithfulness of an embedding from first-order intuitionistic logic into first-order modal logic.

MODALITIES AS PROVABILITY PREDICATES

It is not clear from where exactly Gödel had the idea to use a modal operator to express an informal notion of provability of a sentence. But just as many other ideas throughout scientific history, it seems to me as if using a modality to denote provability was somewhat of a folklore idea among modal logicians the early 20th century.

One obvious candidate for furthering this idea and making it formal is Gödel himself and his incompleteness theorems, the first of which he had published in 1931 (Gödel 1931), two years prior to his interpretation. For proving incompleteness of set theory² he constructed a formal provability predicate enabling set theory to talk about its own

provability. This is also what later led to the development of modal provability logic, where $\Box$ is interpreted as precisely such a formal provability predicate. However, the basic modal provability logic, known as Gödel-Löb logic, turned out to be inconsistent with S4.³ Still, one can directly connect the two notions of informal and formal provability in two ways. One approach is to interpret $\Box$ as “provable and true” leading to an extension of S4 known as Grzegorczyk logic (Grz).⁴ The completeness of Grz with respect to this provability interpretation was proved independently by Rob Goldblatt (Goldblatt 1978), as well as Alexander Kuznetsov and Alexei Muravitsky (Kuznetsov and Muravitsky 1976). Another approach is to have indexed modalities by terms that carry explicit information about the proofs which leads to an S4-like justification logic known as Logic of Proof (Artemov and Fitting 2019). Notably, it was proved by Andrzej Grzegorczyk (Grzegorczyk 1967) in 1967 that the faithfulness of the Gödel embedding also holds for his logic:

Theorem 2. $\text{IPL} \vdash A \text{ iff } \text{Grz} \vdash \text{tr}^G(A).$

Remarkably, this gives us two ways in which we can explicitly connect formal mathematical provability to intuitionistic logic. Either via the Logic of Proof and Theorem 1, or via the provability interpretation of Grz and Theorem 2.

WHAT YOU MIGHT NOT KNOW ABOUT THE GMTT

These results are all well and nice, and this is usually where the main story ends. But it is notable that all commonly known proofs of the previous two theorems rely on topological and algebraic ideas. This is in contrast to my own research interests. Thus, enter Proof Theory.

It was while I was doing research for my Master of

Logic thesis (Becker 2024) that I was led down a rabbit hole. This rabbit hole began with a proof system of a certain *Maehara-style* found in (Kuznets and Straßburger 2019), a central system in my master’s thesis. The style of the calculus is attributed to Shôji Maehara who introduced it in his paper titled “Eine Darstellung der Intuitionistischen Logik in der Klassischen” (A representation of intuitionistic logic inside classical logic) from 1954 (Maehara 1954) for which no English translation exists.

As the keen reader might guess by the title of this paper, I noticed that there could be a more groundbreaking result that Maehara had presented there. Clearly, the title suggests that one might find a result similar to the GMTT. So I started reading that paper out of curiosity. Let me highlight the main points that I have found.

The paper begins by remarking on the well known double negation translation. More specifically, Maehara mentions the interpretation of classical logic inside intuitionistic logic by Sigekatu Kuroda using a double negation translation (Kuroda 1951). Similar results have also been independently discovered by Kolmogorov (Kolmogorov 1925), Gödel (Gödel 1933b) and Gentzen (Gentzen 1974). The paper then motivates itself by considering a reverse interpretation of intuitionistic logic in classical logic. The idea is very similar to the one of Gödel, however it seems as if Maehara was not aware of the previous developments in modal logic. Note also, that Maehara considers not a translation from intuitionistic propositional logic to its classical relative, but instead from intuitionistic *first-order* logic to classical first-order logic.

Maehara introduces a new connective *Bew* which stands for *Beweis* or *beweisbar* (proof or provability), essentially a unary modality expressing informal

provability. Thus, let us identify $\text{Bew} = \Box$ to make the connection to the previous sections explicit. After recalling Gerhard Gentzen's sequent calculi for classical and intuitionistic logic (Gentzen 1935), he introduces his own *Hilfskalkül* (helping calculus) which is the previously mentioned Maehara calculus. He shows that his calculus is equivalent to intuitionistic logic. The motivation for this calculus is to use it as a tool to show for his main theorem.

He then goes about introducing two sequent rules for his provability operator. As is usual in sequent calculus, there is a left rule and a right rule:

$$\Box_R \frac{\Box B_1, \dots, \Box B_n \Rightarrow A}{\Box B_1, \dots, \Box B_n \Rightarrow \Box A}$$

$$\Box_L \frac{A, \Gamma \Rightarrow \Theta}{\Box A, \Gamma \Rightarrow \Theta}$$

Γ and Θ are lists of formulas (or sequences of formulas, which is why it's called the *sequent* calculus).⁵ Maehara remarks, almost coincidentally, that these sequent rules are equivalent to the following:

$$\Box A \rightarrow A$$

$$\Box A \rightarrow \Box \Box A$$

$$\frac{B_1, \dots, B_n \Rightarrow A}{\Box B_1, \dots, \Box B_n \Rightarrow \Box A}$$

It is now a very easy exercise to show that this logic is a first-order version of S4 which we call QS4. Moreover, we can actually derive the converse Barcan formula $\Box \forall x A \rightarrow \forall x \Box A$ which can be a desired feature of a first-order modal logic (see (Cresswell and Hughes 1996) p. 245). As is natural in sequent calculus, Maehara shows consistency and closure under modus ponens for QS4 by proving cut-elimination.

Formulas of the form $\Box A$ should now be read as “A is provable”. Thus, he makes the connection between the classical reading of a formula A as “A is true” and its intuitionistic reading as “it is provable that A is true” explicit via the following translation:

$$\begin{aligned} tr^M(P(\bar{x})) &= \Box P(\bar{x}) \\ tr^M(\neg A) &= \Box \neg tr^M(A) \\ tr^M(A \star B) &= \Box (tr^M(A) \star tr^M(B)), \\ &\quad \text{with } \star \in \{\rightarrow, \wedge, \vee\} \\ tr^M(QxA) &= \Box Qx tr^M(A), \\ &\quad \text{with } Q \in \{\forall, \exists\} \end{aligned}$$

Simply put, this translations just puts a $\Box$ in front of every subformula. It is a nice exercise to show that, in the propositional case, this translation is actually equivalent to $tr^G(\cdot)$ from before. Thus, it almost seems unsurprising to find a theorem similar to the GMTT. However, here we connect first-order intuitionistic logic (IFOL) with the first-order modal logic QS4.

Theorem 3. $\text{IFOL} \vdash A$ iff $\text{QS4} \vdash tr^M(A)$.

Showing the left-to-right direction is rather straightforward, indeed Gödel could prove it axiomatically for the propositional case. The converse is quite involved and makes use of the Maehara calculus from before. In essence, a classical proof of $tr^M(A)$ gets non-locally transformed into an intuitionistic Maehara-style proof such that all occurrences of $\Box$ are removed. We might therefore rightfully claim that this is the first GMTT for first-order intuitionistic logic preceding the result of Rasiowa and Sikorski by about a decade.

But wait, there is more! This is merely the *Hauptsatz* (main theorem) of the paper. There is one additional theorem in the paper which is not there to derive the main theorem, connecting to *intuitionistic* first-order modal logic IQS4:⁶

Theorem 4. $\text{IFOL} \vdash A$ iff $\text{IQS4} \vdash \text{tr}^M(A)$.

Thus, we have even another variant of a GMT-like theorem. Initially, it seemed that this would have been an unknown result, as I did not know anyone who knew about this. But then I realised that I only know proof theorists, and after some digging found that algebraicists of course also found this out (Wolter and Zakharyashev 2014). But the story does not even end there: it was rather recently found by Jan von Plato (Plato 2025), who translated handwritten notes by Gödel, that he had also obtained a completeness result in 1941, although he never published it. The proof by Gödel also uses sequents, however his method differs from the one Maehara used. I do not claim that Maehara's result has been a "forgotten gem" that no one knew about. To the contrary, there are multiple sources I could find from the recent decades which explicitly cited this result (see (Inoué 2023; Yamasaki and Sano 2017) for the most recent ones). It still seems remarkable though that a result of such a scale has received so little attention; and that the vast majority of citations of Maehara's *Darstellung* is only due to his proof system. Thus, the story about "Maehara's Gödel-McKinsey-Tarski Theorem" can also be seen as a story about the sociology of mathematics. Especially deep results have a tendency to be discovered more than once, almost as if they want to be discovered. At the same time, some discoveries may lie dormant for decades, either because they sit unpublished in some handwritten notebook, or they must be read by the right person to be appreciated. Thus, one may only wonder how many more such theorems are quietly sitting in the literature, only waiting for someone curious enough to follow a rabbit hole.

FOOTNOTES

1. As a historical note, Gödel actually considered different connectives for intuitionistic and classical logic. For example, he used $\cdot$ to denote classical conjunction while $\&$ denotes intuitionistic conjunction. 2. While Gödel's incompleteness theorems are nowadays often phrased in terms of Peano Arithmetic, he himself talks explicitly about the incompleteness of Principia Mathematica (by Russel and Whitehead) and Zermelo-Faenkel set theory. 3. This could have already been noticed by Gödel himself, as we can read the second incompleteness theorem as $\Box \neg \Box \perp \rightarrow \Box \perp$ (read $\Box \perp$ as "the theory is inconsistent" and $\neg \Box \perp$ as "the theory is consistent") and that $\neg \Box \perp$ (and by necessitation also $\Box \neg \Box \perp$) is provable in S4. 4. As a historical note, the logic Grz is usually defined over Kripke logic K by adding the axiom schema $\Box(\Box(A \rightarrow \Box A) \rightarrow A) \rightarrow A$ whereas Grzegorzczuk himself defined it by adding the axiom schema $((\Box(B \rightarrow \Box A) \rightarrow \Box A) \wedge ((\neg B \rightarrow \Box A) \rightarrow \Box A)) \rightarrow \Box A$ to S4. 5. For readers not familiar with sequent calculus, read a sequent $\Gamma \Rightarrow \Theta$ as an implication $\bigwedge \Gamma \rightarrow \bigvee \Theta$ or informally as "assuming all of Γ we can derive some of Θ ". 6. In the intuitionistic modal logic tradition, logics denoted with a capital I are usually reserved for logics with explicit $\Diamond$ s, which are intuitionistically not interdefinable with $\Box$, as well as logics validating $\neg \neg \Box A \rightarrow \Box \neg \neg A$. Both these, do not apply to what we call IQS4, thus one might also call this logic iQS4. However, one can easily convince oneself that this theorem also holds for a logic which does fulfill these additional requirements.

BIBLIOGRAPHY

For the bibliography, see the Illogician website.

FUNDAMENTALS OF COMPUTATIONAL PERPLEXITY THEORY

MATTEO CELLI | COMPUTATION

INTRODUCTION

This article lays the groundwork for a new field of theoretical Computer Science, which we call *Theory of Computational Perplexity*. Computational Perplexity is the study of how much puzzlement is prompted by a computational task. It is strongly related to, albeit distinct from, computational complexity theory insofar as the existence of an efficient algorithm for a given problem does not imply that the working computer scientist will be any less perplexed by the aforementioned problem. In the present article we will introduce the main concepts and define the main perplexity classes, Perp and NPerp. The first major open problem of the field, $\text{Perp} \stackrel{?}{=} \text{NPerp}$, will be formulated, hoping to pave the way to a new and fruitful avenue of research.

SOME PERPLEXITY CLASSES

Let us begin with a familiar example. Consider the problem which the snobs call *Boolean satisfaction problem*, and to which the *connoisseurs* refer as SAT. In plain prose, for a formula of propositional logic, we ask whether there exists an assignment that satisfies it. From a computational perspective, this is perhaps the most well-known NP-complete problem, i.e., an NP problem that is as hard as any other NP problem (i.e., NP-hard *and* in NP). Now, recall that a problem is in NP iff there exists a non-deterministic turing machine that solves it in

polynomial time, or, alternatively, if a deterministic machine can check in polynomial time whether a candidate solution is correct. As far as *Perplexity* goes, on the other hand, we shall not focus on Turing Machines anymore. Perplexity is about humans, and deals with intimately human themes such as hope and despair, and it does so in a systematic and systematisable way. From now onwards, Turing machines will be replaced by something they will never be able to replace: the working computer scientist.

But then, how will we understand SAT from a perplexity-theoretic perspective? We begin by framing what computational perplexity is about:

Definition 1 (Computational Perplexity). *Computational Perplexity is the study of how much puzzlement or confusion is caused by a computational problem to the working computer scientist.*

Clearly, boolean satisfaction will not be as daunting for a human than for a machine. Trying out a bunch of truth tables cannot be so bad, can it. Therefore, intuitively, SAT should be classified as one of the least perplexing problems. Finding an independent set in an undirected graph (INDSET), on the other hand, while being as computationally hard as SAT (they are both NP-complete), is clearly more baffling than the latter. Not to mention *Private-neighbour independent set*, the version of INDSET in which no nodes of the independent set

are adjacent to more than one node in the complement. Here, on top of the puzzlement for even understanding what this means, we add the confusion of why on earth someone called it PNIS without questioning its disastrous giggling potential among students, and, even if they are not willing to admit it, professors. This is to say, complexity and perplexity must have different classes, since provably equally hard problems are obviously not similarly puzzling. We now begin by introducing some core perplexity classes in a way that it adequately captures the present intuitions. To do so, we understand “unpacking a question” algorithmically: we read it, we understand each word, then each sentence, then sentence compounds linked by connectives, and then conjecture what possible solutions may look like.

Definition 2 (The class Perp). A problem is in Perp if and only if there exists a working computer scientist who is puzzled by the question, but manages to understand it in polynomial time by deterministically unpacking it in its elementary notions.

From the definition, SAT is clearly in Perp: understanding what “finding an assignment that makes ϕ true” once we are given ϕ is a task that is clearly understood in polynomial time (the obvious proof is left to the reader). INDSET, on the other hand, is a beast of a different breed. We read the question, and, not without feeling a bit of shame, we recall by drawing on a piece of paper the difference between a directed and an undirected graph. After some (hopefully short) trial-and-error, we try to make sense what an independent set is. We read the definition, and colour some nodes to see if we understood it. No, we have not. Reread, retry, erase, make a more general drawing, recolour some nodes, yes, these form an

independent set, the definition makes sense. This process exhibits an exquisitely non-deterministic behaviour. Thus, we define the following class:

Definition 3 (The class NPerp). A problem is in NPerp if and only if there exists a working computer scientist who is puzzled by the question, but manages to understand it in polynomial time by nondeterministically unpacking it in its elementary notions.

Alternatively, we can rephrase it as follows:

Definition 4 (The class NPerp, alternative definition). A problem is in NPerp if and only if there exists a working computer scientist who is puzzled by the question, but manages to check in polynomial time if a candidate explanation is sensibly correct, by deterministically unpacking it in its elementary notions.

The existence of NPerp problems is clearly witnessed by INDSET.

AN OPEN PROBLEM

But can we be sure that Perp and NPerp are not identical? And what would their equivalence imply? As in the familiar case of P and NP, we can easily see that we can obtain an inclusion $\text{Perp} \subseteq \text{NPerp}$. This is particularly evident if we consider the alternative definition, where the candidate interpretation of the question of what satisfiability is is the only possible one. But how about the other inclusion? I claim that understanding it may be as puzzling as it is hard to prove $\text{NP} \subseteq \text{P}$. Suppose however that we understand it, and we manage to show that $\text{NPerp} \subseteq \text{Perp}$, thus obtaining $\text{Perp} = \text{NPerp}$. The consequences for the practice of computer science, I contend, would be disastrous. If anything, the working computer scientist will not be anymore entitled to whine about how weird or

confusing computational task are, since they will now that there is an efficient way to understand what they mean. The comforting conclusion “understanding this issue is left for future research” will be followed by sceptical glares of colleagues

who will ask “but did you really think hard enough?”. As of now, however, I indulge in such privilege, maybe for the last time. Understanding what $\text{NPerp} \subseteq \text{Perp}$ means is a task that must be taken up in the years to come.



Inspired by Scott Aaronson's Complexity Zoo. Informed by textbooks and other sources in complexity theory. Many thanks to the researchers who made these ideas accessible! ☺

— Ziqi Wang

"THE COMPLEXITY ZOO" BY ZIQI WANG

STRUCTURALIST FOUNDATIONS FOR MATHEMATICS

MAGNUS KJAERGAARD | PHILOSOPHY, LINGUISTICS

Structuralism in philosophy of mathematics is the view that mathematical statements rather than describing objects, describe only structures on objects, and that mathematical entities as such are or positions in a structure, rather than objects/entities in their own right. This view has been around in some form for most of the history of contemporary philosophy of mathematics, and has been defended by figures such as Dedekind, Hilbert, Poincaré, Quine, Benacerraf, and others (Colyvan 2012, 40). However, following (Benacerraf 1965), interest in structuralism was renewed due to its potential as a solution to Benacerraf's underdeterminacy problem. This problem, simply put, consists in the fact that in the standard foundation of mathematics, the Zermelo/Fraenkel axioms (ZFC), all mathematical objects are reducible to sets, and yet there are often multiple equally valid choices for *which* sets a collection of objects should be reducible to. Famously, Benacerraf used the example of the Zermelo and Von Neumann constructions for the natural numbers. Both systems identify the natural number 0 with the empty set, but the constructions differ in their choice of successor function. In the Zermelo construction $\text{succ}(x) = \{x\}$ and in the Von Neumann construction $\text{succ}(x) = x \cup \{x\}$. This yields the interesting consequence that although these constructions are elementarily equivariant in the language of Peano Arithmetic - both systems prove

exactly the same theorems - one system, the Von Neumann construction, would have it that $1 \in 42$, the other not. This underdeterminacy problem presents a substantial barrier for any prospect of an *ontological* reduction of numbers into sets, and has largely caused the view that numbers are sets to fall into disrepute (Colyvan 2012). Structuralism promises an exceedingly elegant and straightforward solution to the problem. Instead of identifying numbers with any particular object - abstract or not - they should instead be thought of as referring to positions in the structure of the natural numbers, which in turn is instantiated in *both* the Zermelo and Von Neumann constructions. Being the number 2 is to stand in a certain relation to a surrounding structure, and both $\{\{\emptyset\}\}$ and $\{\emptyset, \{\emptyset\}\}$ stand as such to their respective constructions. Not only does this avoid underdeterminacy in the choice of which structure should be the one, but it seems at the face of it that only some sort of structuralism is able to avoid having to justify one construction or another as the canonically correct choice, which constitutes an inherent advantage over other accounts of the ontology of mathematics. Furthermore, structuralism is also uniquely harmonious with mathematical practice. As noted by (Colyvan 2012)[p. 40]; "Structuralism is able to explain why mathematicians are typically only interested in describing the objects they study up to isomorphism - for that is all there is to

1. Emphasis on the fact that this is a verbatim quote.

2. An interesting observation, that I won't go further into here, is that the Zermelo construction for natural numbers actually doesn't fall under this definition of the natural numbers, as indeed $\{\{\emptyset\}\}$ is not a 2-element set.

describe." Structuralism is, of course, not without issues of its own, but exactly what these are can be hard to pin down given the sheer variety of different versions of structuralism that have been proposed. Broadly speaking, most fall into one of two camps; realist or *ante rem* structuralism *a la* (Shapiro 1997) or eliminative/*in re* structuralism *a la* Hellmann (G. Hellman 1989) (Reck and Schiemer 2025). Both positions of course have to deal with the usual problems that realist (respectively anti-realist) philosophies of mathematics have to contend with, e.g. the dilemma posed in (Benacerraf 1973). And yet, in both cases the arguably bigger issue lies in spelling out what exactly counts as a structure and when two structures should be considered the same. It should therefore come as no surprise that alternative foundations for mathematics such as category theory and homotopy type theory, that are inherently structuralist in nature have recently been of great interest in the contemporary discourse on the philosophy of mathematics.

THE DISCUSSION SO FAR

The question was discussed in (Geoffrey Hellman 2003). Just one year later an answer was given in (Awodey 2004) - that answer being: ¹. The reason for the terseness of this reply being that Hellman's doubts stem largely from concerns about whether category theory makes for an autonomous foundation for mathematics. A question that -(Awodey 2004) points out - is separate from the question of whether category theory can give a precise notion of the term on top of which a formal structuralist system can be built. This, (Awodey 2004) concludes, should be relatively obvious, and indeed a structuralist need not argue for an outright replacement of the ZFC with a structuralist alternative in order for their argument to succeed. However, a successful argument for the

replacement of ZFC with one of the proposed structuralist alternatives would still be a strong point in structuralism's favor, and as such I claim the questions raised by Hellman are still relevant for structuralists to consider. The issue of whether category theory specifically makes for such an autonomous foundation has been discussed at length ever since the publication of Lawvere's Elementary Theory of the Category of Sets (ETCS) (Lawvere 1964) and has more recently been given a comprehensive treatment by (Linnebo and Pettigrew 2011). They argue that a (worthwhile) alternative foundation to set theory must be logically, conceptually, and justificatorily autonomous from ZFC. They conclude that the ETCS - as well as a few other proposed categorical foundations - indeed enjoy both logical and conceptual autonomy, but that their justificatory autonomy hinges on the extent to which the practice of working mathematicians may be appealed to as justification for its adoption. More recently (Tsementzis 2017) has made an intriguing case that categorical foundations such as the ETCS are not actually structuralist in a strict sense and has argued convincingly that the **Univalent Foundations** (UF) based on homotopy type theory provide a better alternative foundation that is truly structuralist. He subsequently outlines a formal method for translating standard results from mathematics into the structural language by means of UF. This method does not allow us to translate all the mathematics, as indeed UF is a constructive system unlike ZFC, but it does translate quite a lot, and according to (Tsementzis 2017) it is open ended enough that it can always be expanded in the future. Although there will inevitably be some compared to ZFC, this is an intended consequence of using a constructive system, and many mathematicians will probably consider this a point in UF's favor. In the essay this

article is based on, I try to evaluate where the ETCS and UF respectively stand in light of the ongoing discussion, but due to this format I will not be able to do that adequately here. Instead, I will spend the remainder of the article making it a bit more explicit what it means for a mathematical framework to be structuralist.

WHAT DOES IT MEAN TO BE STRUCTURALIST?

As noted already, there is a panoply of different structuralist positions that all conceptualize structures slightly differently, making it impossible to consider them all here. However, both (Shapiro 1997) and (Tsementzis 2017) have adopted a view inspired by the Quine slogan: according to which the relevant issue is not what structures *are*, but when structures should be considered identical. And in this regard, we are in the fortunate position that mathematicians already have a rigorous notion of structural identity: isomorphism.

Building on this idea, (Tsementzis 2017) argues that whether or not a system is structuralist is not dependent on what our ontology of structures are, but whether or not our formal system considers isomorphic objects to be identical (SFOM). By this there can be little doubt that UF is a structuralist system. Indeed, this is baked essentially into the core of the program in the form of the axiom of univalence. This axiom states;

$$(A =_{\mathcal{U}} B) \cong (A \cong B)$$

Or in words; for any two types A,B the type of ‘proofs of identity between A and B’ is isomorphic to the type of ‘proofs that A and B are isomorphic’, or as (Tsementzis 2017) succinctly puts it: . In UF, equality and identity are simply equivalent. Tsementzis further notes that this notion of identity

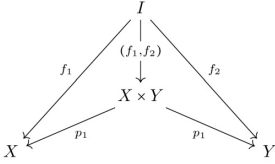
is so strongly structuralist that at first it appears ungrammatical. Indeed, how can what appears to be not a structure, but a statement - that $A = B$ - be isomorphic to anything? The answer being simply that, in UF, $(A =_{\mathcal{U}} B)$ *just is* a structure. The statement is simply the same as . In UF, structure is quite literally all there is to describe. The Elementary Theory of the Category of Sets of (Lawvere 1964) does not contain anything quite like the axiom of univalence, but it is still often considered a structuralist foundation. Briefly put, the ETCS is a recreation of standard set theory within a category theoretical framework. As such, sets are not determined by their members in the ETCS, but rather characterized in terms of *morphisms* and *universal properties*. Starting with the latter, a universal property is something that specifies an object/construction up to isomorphism. For instance; the object **1** is identified not as some particular set, but rather defined in terms of the properties shared only by sets with exactly 1 element. Such sets are called *terminal objects* and are defined thus

Definition 1 (from (Leinster 2012)). *A set T is called terminal if for every set X , there is a unique function $X \rightarrow T$.*

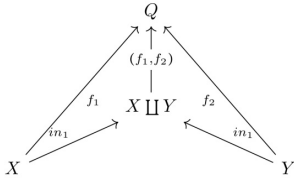
Here, it should be quickly apparent that, though there are many sets that obey the definition of terminal objects, all of these will be mutually isomorphic to each other. The significance of this is that the object **1** is precisely the category of things isomorphic to **1** (or $\{\emptyset\}$ or $\{*\}$ and so on). Further, the notions of sum and product are similarly specified in category theory by means of universal properties.

Definition 2 (Categorical product). *Let X, Y be sets. A **product** of X and Y is a set $X \times Y$ together with a pair of projection functions $X \xleftarrow{p_1} X \times Y \xrightarrow{p_2} Y$ such that for any set and pair of functions $X \xleftarrow{f_1} I \xrightarrow{f_2} Y$ the

following diagram commutes



Definition 3 (Categorical sum (or co-product)). *Let X, Y be sets. A **sum** of X and Y is a set $X \amalg Y$ together with a pair of injections $X \xrightarrow{\text{in}_1} X \amalg Y \xrightarrow{\text{in}_2} Y$ such that for any set and pair of functions $X \xrightarrow{f_1} Q \xleftarrow{f_2} Y$ the following diagram commutes



It is worth mentioning that the cartesian product and disjoint union operations are categorical products and sums, respectively. Once again, notice that this specifies the respective notions up to isomorphism, so for instance the set 2 is specified up to isomorphism as $1 \amalg 1$, which means that the number 2 is similarly the class of objects isomorphic to 2 (or $\{\emptyset, \{\emptyset\}\}^2$, and subsequently all the natural numbers may be similarly defined. This certainly seems like a thoroughly structuralist system, so are we safe to conclude that the ETCS is a structuralist foundation for mathematics?

As alluded to earlier, the main source of doubt on this matter again comes from (Tsementzis 2017), who argues that not only does the ETCS not in fact satisfy the SFOM, it in principle *cannot*. This is for the simple reason that the ETCS is, by design, a reconstruction of standard ZFC within a category theoretic framework, which means that violations of SFOM in ZFC carry over into the ETCS. For an easy example of such a violation, he gives the example that while $\mathbb{Z}$ and $2\mathbb{Z}$ are isomorphic, we of

course have it that $\text{ZFC} \models 1 \in \mathbb{Z}$ but $\text{ZFC} \not\models 1 \in 2\mathbb{Z}$. Therefore, set theory as a foundation is indeed not structuralist in the sense of (Tsementzis 2017) because isomorphic structures can and do validate different formulations in the language of set theory, and as such, structural identity is different from actual identity. But these examples carry over almost verbatim to the ETCS. Does this mean that category theory is not a structuralist foundation for mathematics? Not in the sense used by Tsementzis, but it should be noted that there are valid reasons to think that his sense is a bit too strict. First, the ETCS is at least structuralist, in the sense that the objects of the theory - the objects that everything else is defined in terms of - are indeed objects as they are specified only up to isomorphism. Arguably, being structuralist in this sense is sufficient for a philosophical case for structuralism.

Secondly; it is not entirely obvious to me that (Tsementzis 2017) is right to claim that a structuralist system must *always* consider all isomorphic objects as elementarily equivalent. Tsementzis is of course right to point out that isomorphism sometimes results in a loss of information in ETCS, and that this does seem unexpected in a structuralist system, but I believe that it is worth pointing out that these cases are restricted to situations where a *reduct* of the language has taken place. What I mean by this is that we know from model theory that if two $\mathcal{L}$ -structures $\mathcal{M}, \mathcal{N}$ are isomorphic, then they are elementarily equivalent, i.e. for any $\mathcal{L}$ -sentence φ we have $\mathcal{M} \models \varphi$ if and only if $\mathcal{N} \models \varphi$. In other words, when the two structures in the same language are isomorphic, they preserve information as we would expect them to, with the caveat that it is only the information *expressible* in $\mathcal{L}$. In other words, the SFOM is recoverable for ETCS if we add this single stipulation. I believe this

is a rather small concession. In summary; we have two promising candidates for a structuralist foundation for mathematics. One constructivist, the other of the same proof theoretic strength as ZFC. One structuralist in a quite strong sense, where isomorphic objects are literally indistinguishable, the other in a slightly more fine grained sense that allows for some language specific discrepancies, but still specifies its primitive objects only up to isomorphism. Both of which would lend themselves to quite a strong case for mathematical structuralism.

BIBLIOGRAPHY

- Awodey, Steve. 2004. "An Answer to Hellman's Question: "Does Category Theory Provide a Framework for Mathematical Structuralism?"" *Philosophia Mathematica* 12 (1): 54–64. <https://doi.org/10.1093/philmat/12.1.54>.
- Benacerraf, Paul. 1965. "What Numbers Could Not Be." *The Philosophical Review* 74 (1): 47–73. <http://www.jstor.org/stable/2183530>.
- . 1973. "Mathematical Truth." *Journal of Philosophy* 70 (19): 661–79. <https://doi.org/10.2307/2025075>.
- Colyvan, Mark. 2012. *An Introduction to Philosophy of Mathematics*. Cambridge Introductions to Philosophy. Cambridge university Press.
- Hellman, G. 1989. *Mathematics Without Numbers*. Oxford University Press.
- Hellman, Geoffrey. 2003. "Does Category Theory Provide a Framework for Mathematical Structuralism?" *Philosophia Mathematica* 11 (2): 129–57. <https://doi.org/10.1093/philmat/11.2.129>.
- Lawvere, F. William. 1964. "An Elementary Theory of the Category of Sets." *Proceedings of the National Academy of Sciences of the United States of America* 52 (6): 1506–11. <http://www.jstor.org/stable/72513>.
- Leinster, Tom. 2012. "Rethinking Set Theory." at <https://arxiv.org/abs/1612.09375>. <https://arxiv.org/pdf/1212.6543>.
- Linnebo, Øystein, and Richard Pettigrew. 2011. "Category Theory as an Autonomous Foundation." *Philosophia Mathematica* 19 (3): 227–54. <https://doi.org/10.1093/philmat/nkr024>.
- Reck, Erich, and Georg Schiemer. 2025. "Structuralism in the Philosophy of Mathematics." In *The Stanford Encyclopedia of Philosophy*, edited by Edward N. Zalta and Uri Nodelman, Fall 2025. <https://plato.stanford.edu/archives/fall2025/entries/structuralism-mathematics/>; Metaphysics Research Lab, Stanford University.
- Shapiro, S. 1997. *Philosophy of Mathematics: Structure and Ontology*. Oxford University Press.
- Tsementzis, Dimitris. 2017. "Univalent Foundations as Structuralist Foundations." *Synthese* 194 (9): 3583–3617. <https://doi.org/10.1007/s11229-016-1109-x>.

REAL NEWS: ABSTRACTOPHOBIA

HUGO RENNINGS | MATHEMATICS

We've all been there: it's grandma's birthday so you're at her apartment in Roosendaal with your whole family. Your uncle is talking about how he's using *the A.I.* at his work and says that all the kids these days are using *chapgpt* for their assignments. He looks at you with a smirk, 'Right?' You say 'Hmmm, well, I don't know', and mumble something about neuro-symbolic bla bla explainable bla bla until he loses interest. A few minutes later: Bingo. Your grandpa asks you (like every time) 'Wat is dat ook alweer, dat logica?'¹. Hahaha, how to start? Well maybe you say something about first order logic (maybe model theory? No, actually, that's way too much), or perhaps you try to explain Frege's puzzle. Yeah, that's a good idea! But before you've even referenced the meaning of modality he's nodding like that was the end of your story and offering the next person some tea. Well, whatever.

Scientists have recently coined the term abstractophobia for the psychological condition characterised by 'A repulsion to (discussion of) abstraction or abstract concepts, especially without the explicit mention of a practical, real world, hands-on, impactful, results-driven application and/or example'. Early research suggests that it affects nearly everyone you talk to about your mathematics degree. Symptoms depend on the situation and interlocutor, but mathematicians are often greeted with an 'I actually *hated* maths in high school' or a 'So you're like *super* smart'

1. What is it again, that logic?

followed quickly by a 'What do you want to do with that?' Chronic abstractophobia leads to hyperfixation on the tangible (as perceived by the patient). This causes the irresistible urge to ask 'Is that useful?' or, 'Does that have any applications?'

'Abstractophobes are not just uninterested in set theory or algebra', explains Dr Zorn Slemma, 'They have a deep-seated disdain for the concept of abstraction'. One indication of this is the lack of consistency in the discourse around abstraction. A recent study found that patients associated the word 'abstract' with both 'That art that is just splatters of paint and actually my one-month-old infant could also make that' as well as with 'That stuff that is way too confusing for my one-month-old infant'. Additionally, researchers found that they could reduce objections of abstractophobes to theorems of real analysis by stating that they could be 'Used for googoo gaga' and 'Help optimise leeleee-lala'. In fact, the words 'use', 'apply', 'optimise' and 'improve' were all found to significantly reduce discomfort in patients, even when used in a completely incomprehensible way. These results motivated scientists to designate abstractophobia as a legitimate condition.

Scientists are still working to determine exactly the causes of abstractophobia, but here is what they know: Individuals are usually infected with abstractophobia by others, especially when someone they trust starts to say that 'It can't hurt

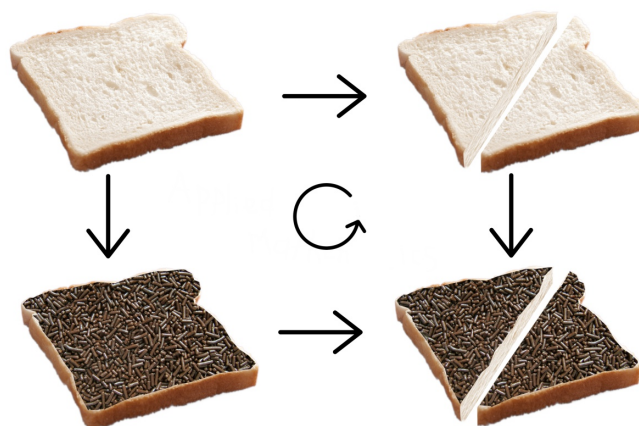
to earn a bit of money' and 'I don't think that boolean algebra homomorphisms are the next big thing'. Children of parents who suffer from the condition are also significantly more prone to it themselves, and are less likely to recover. Epidemiologists are thus concerned at the increase in abstractophobia among toddlers, with one study reporting that nearly half of three-year-olds showed a disinterest in placing wooden shapes in the correctly shaped holes of a box. One participant is quoted 'When will I ever need to use this in real life?' On the other hand, experts found that the popular internet character 'Tung Tung Tung Sahur', a leading figure in online abstract and surreal culture, is particularly popular under the younger generation, so there may be some hope yet.

Now, our readers might think that they are not in any danger themselves, since they view the world as a commutative diagram and haven't thought about an actual natural number in months, but be warned! The government has, for a while now, been concerning itself with an evil plan known as 'E.C. Onomy' and to succeed they need to convince you to stop thinking about math and start thinking about cash. That is why universities try at every

opportunity to turn the pure mathematics student into an engineer. Career day? More like 'Ca-n't-you-study-something-to-make-us-rich-er day'. Applied maths? More like 'Appl-ease-help-us-ensure-more-profit-is-realised maths'. In fact, even those 'students' who ask the lecturer to 'Make this more concrete' and to 'Give an example' are actually agents of the state! And this strategy works. One day your buddies are with you, laughing and having fun while learning about algebraic logic, next day they're sending Phillip-Morris their CV saying that they have 'Ultra knowledge of proper filters'. Watch out. Nobody is safe.

By describing abstractophobia, researchers hope to create awareness for the condition, which might eventually lead to a cure. In fact, a plan to develop a vaccine has already been drawn up. The idea is to produce an injection $\iota : \text{People} \hookrightarrow \text{People Who Like Maths}$, which is sufficient to eradicate abstractophobia. We hope to bring you good news on this front in the near future.

Next week in *Real News*: $\aleph_1$ reasons your morphisms are not epic.



Applied Mathematics

MATTEO CELLI, JOSJE VAN DER LAAN, GIANNIS RACHMANIS

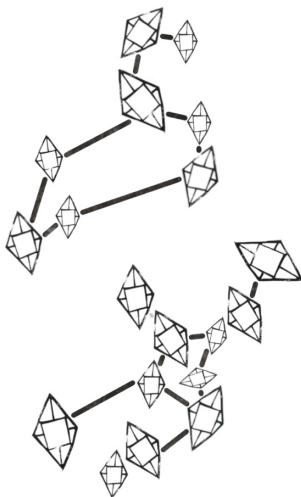
Aries (Mar 21 – Apr 19)

You have attitude, propositional and beyond. The problem is, it makes your life hard in embedded contexts. But do not let this prevent you from expressing yourself: Frege and Geach do not understand who you really are and never will. Don't scope over them much, you can substitute as you freely will. If you have to choose between the law of the excluded middle and double negation elimination, look towards the stars: both and neither are written in there, but at least you will see the light.

With the sun moving to Pisces, the other water zodiacs will not be unaffected. It is written in the stars: a trip to the blue-book-blue sea in Zandvoort is the right step . The first moves towards a new journey are always at the start and usually uneven but quite liberating. You need to make a colist since your comrades and colleagues cook for you, and this time, it is not only consistency that matters. The recipe is to start from the final case which is evident by definition and proceed coinductively to the previous step until the beginning.

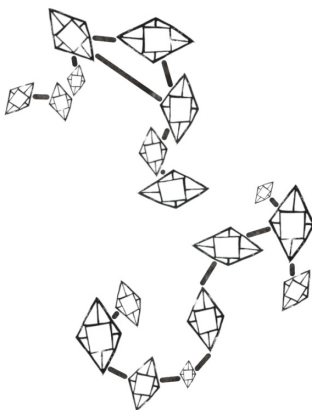
Leo (July 23 – August 22)

You try to find satisfaction in life by trying out all the possible options, and you hope that the stars can give you polynomial bits of advice. But beware: by the Karp-Lipton theorem this would make the polynomial hierarchy collapse to the second level. The darkness is only a reminder of the light you know how to experience, so do not commit to model theory just yet. Accept that the best things in life happen when you least expect it, but with nonzero probability of being incorrect.



Libra (September 23 – October 22)

New year, new you? Who knows? The stars have a finite algorithm to compute. Regardless of the answer, this is the time in which you force yourself to actually sit down and learn what forcing is. You will finally learn to experience yourself that the independence of the Continuum Hypothesis from ZFC is not because of Mercury retrograde, but because of the unexplanatory nature of axioms. Every time you finish a bottle of wine, a new one appears: whose fault that is, is an answer that the first type theorist that you meet on the way home can question.



Virgo (August 23 – September 22)

Free spirit, freely-generated subalgebras, free borders, free jazz, free choice, but you keep working in the MoL room: the stars are telling you to listen, and enjoy some freedom as well. By construction, you need to cycle around Oosterpark, but this path might never end (unless the sun hints you into the correct direction). You may go to the park or to the beach, but then does it mean that you may go to both? Ok, maybe one more hour in the MoL room to wrap your head around this.

Scorpio (October 23 – Nov. 21)

The end of the academic year has been a bit of a cocomplete mess, but don't be too harsh on yourself! In some toposes, even the world's simplest Axiom of Choice fails. And do not overcomplicate matters: a coconut is after all just a nut. Remember to sprinkle constructivist propaganda on the minds of your fellows to save yourself from yet another mistake. What the future brings, is a question for tomorrow: now, you can only overthink about the past.

Sagittarius (Nov 22 – Dec 21)

You spent a year trying to justify a new constructive, non-wellfounded, paraconsistent, predicative, vegan and gluten-free set theory, but now self-membered sets are and are not subsets of themselves, and you are running out of arguments to uphold that this is actually “natural”. But don’t worry: the Moon and Uranus in the house of Leo will bring you a fresh perspective on the matter.

Aquarius (Jan 20 – Feb 18)

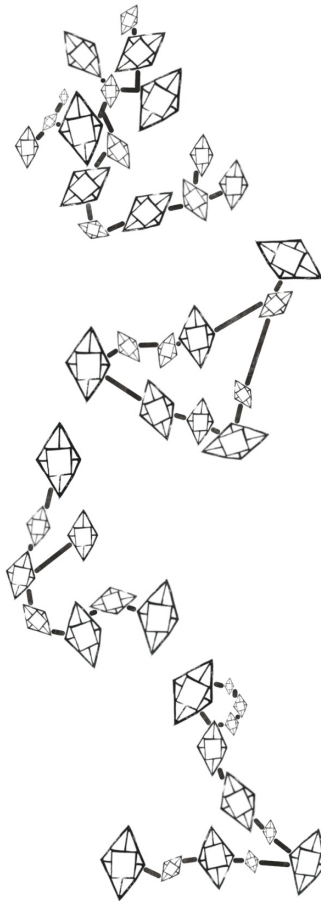
You are busy trying to understand what “sun moving to sagittarius” means, to procrastinate on the dynamics of modal logic. But fear not: if you carefully look at the stars, you will notice that there is a linear order with PDL-interpolation. Fate is on your side, trust the universe and keep amalgamating - believing in the stars is not much different from believing in the axiom of choice anyway.

Capricorn (Dec 22 – Jan 19)

The sun rises and the moon takes a break from the heat in the house of Aquarius. This is a good sign, as it means that all your ideas will come to life. Unless you are a generality absolutist, and you truly mean all ideas. In that case, your head will explode and all ideas will let loose, inspiring more logicians to have an absolute generality of ideas and in the end you will have a proof of infinity.

Pisces (Feb 19 – Mar 20)

They recommend that you cut on spending as mental Mercury clashes with deceptive Neptune in your finance zone. But you believe in eliminating cuts, and opt for a reckless life. But what they do not get about you is that cut-free also means consistent: go girl! Remember you are your own creating subject: you can prove whatever you want to show your selfworth, even trivial statements, if that is what makes you complete.



ILLUSTRATIONS: HUGO RENNINGS



♥ Diversity ♥ Committee

WHO WE ARE

Our mission is to provide a welcoming and supportive environment to all members of its community regardless of background or identity. Our tasks:

1. To advise, support and act as contact for the ILLC community on matters of diversity, inclusivity and appropriate behaviour
2. To promote diversity and inclusivity in our community via information, organizing events and trainings

We would like to become more active again in the next year. For this we need to know about your needs and would like to hear your suggestions. Because of this we ask you to take part in this short, *anonymous* questionnaire:



QUESTIONNAIRE



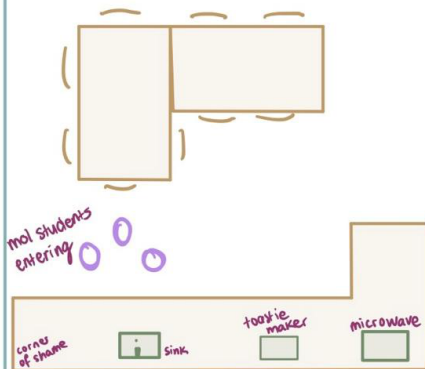
WEBSITE

Andy | Alex | Anna | Balder | Filip | Jelke | Karolina | Ronald

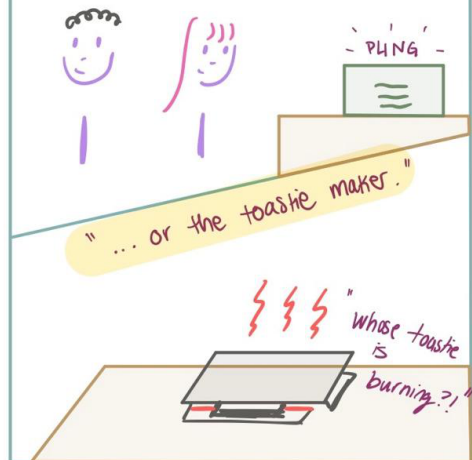
WWW.TINYURL.COM/DIVERSITY-AT-ILLC

1 LUNCHTIME?

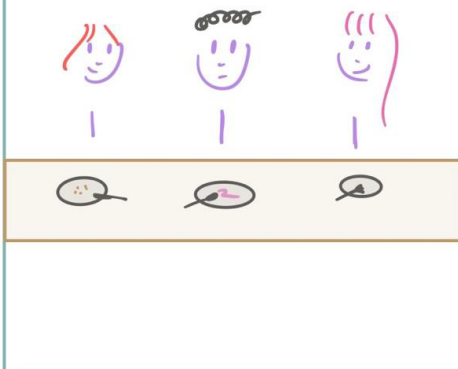
(David Attenborough voice-over)
"The common room at 13:00 on a week day, also known as lunchtime. We'll witness an important ritual here!"



"The mol students wait in line for the microwave..."



"And then, when lunch is ingested, it is time for the most important part of the ritual."



"coffee break?"

