

A Transformative Intuitionist Logic for Examining Negation in Identity-Thinking

Rebecca Kosten

Please cite definitive version once available, forthcoming in
the *Australasian Journal of Logic*.

Abstract

Negation often reinforces problematic habits of othering, but rethinking negation can make good on feminist hopes for logic as a transformative space for inclusion. As Plumwood argues in her 1993 paper, not all uses of negation in the context of social identity are inherently problematic, but the widespread implicit use of classical negation has limited our options with respect to representing difference, ultimately reinforcing dualisms that essentialize social differences in problematic ways. In response to these limitations, I take inspiration from Dembroff's recent work on the metaphysics of genderqueer identity to build models of social identity using the *Heyting-Brouwer* logic developed by Rauszer in her 1974 paper. Ultimately, I argue that these models demonstrate both how classical negation reinforces problematic habits of othering and how alternative forms of negation can transform our treatment of social identity altogether.

1 Introduction

The use of logic to represent, model, or discuss socially engaged phenomena is fraught with concerns. Feminists who are critical of logic have argued that logic is not suitable for engaging with social phenomena because it abstracts away from lived experience and reinforces existing hierarchical relationships between social categories.¹ In light of these critiques, Valerie Plumwood

¹See, for example, (Nye 1990). For an analysis of this and other critiques, see (Plumwood 1993).

argues that the problem lies not with logic as a whole, but with a certain set of assumptions guiding its function. In particular, she suggests that the culprit in many examples of problematic or ineffective logical representations of social phenomena is the widespread use of classical negation.

I agree with Plumwood that the features of classical negation, and many of the assumptions guiding classical logic, limit our options with respect to representing difference, ultimately reinforcing dualisms that essentialize social differences in problematic ways. However, since the tendency to produce and reinforce dualisms is based on the assumption of immutable, exhaustive, and exclusive categories, this tendency can be mitigated through using a logic that does not rely on these assumptions.² Based on a close examination of the problematic features of dualisms, Plumwood recommends criteria for a negation that can be used for feminist purposes. These include being able to recognize the contribution and significance of each individual category, reconceiving relations in more integrated ways, reclaiming denied areas of overlap, and affirming positive or independent sources of definition (Plumwood 2002).

For example, if we assume a gender binary with immutable, exhaustive, and exclusive categories, this situates the relevant social categories as a dualism. In their paper, Robin Dembroff identifies the binary axis as a framework which situates the categories of men and women in this way (Dembroff 2020). In this case, the only options for self-identification with regard to gender become *man* and *woman*. If these categories are exhaustive, then each agent must be categorized as either a man or a woman. Owing to the hierarchical relationship between these categories in the relevant social contexts, the category of man becomes more prominent and being a woman becomes simply a matter of failing to be a man. Furthermore, if these categories are exhaustive, then no agent can be both a man and a woman. As a result, when individuals regulate the boundaries between these social categories,

²Whether or not a given logic relies on these assumptions will, of course, depend on how these assumptions are articulated. While exhaustiveness and exclusivity can be directly captured by particular logical sentences, immutability is more difficult to capture. Depending upon how we decide to capture immutability, a classical approach could likely mitigate some of its harmful effects. In the models I offer here, immutability occurs when an agent does not have an available path for transitioning between social categories. Lacking such a path impacts not only the agent's future-related options, but also the way in which they are conceiving of their current identity. See section 4.2 for an example of how assuming immutability limits not only an agent's future related options, but also an agent's understanding of the social categories in question.

being a man will require rejecting anything associated with the category of woman. This rejection ultimately creates problematic habits of othering whereby agents in the dominant social category consider themselves to have nothing in common with the “other” less dominant social categories. And if these categories are immutable, then those in the dominant social category (in this case that of *man*) are secure in maintaining their social position and its associated privileges. In this way, the immutability of social categories results in a reinforcement of social hierarchy: no matter the social value and contribution of a non-dominant social category, agents who self-identify this way are trapped in a social position that is characterized by its failure to approximate the dominant and is systematically rejected through practices of othering.

Plumwood argues that not all uses of negation in the context of social identity are inherently problematic. I agree – my models in this paper provide a way of visualizing the problematic features of classical negation and imagining alternative ways to conceive of the relationships between social categories. While classical negation does often reinforce problematic habits of othering, rethinking negation can make good on feminist hopes for logic as a transformative space for inclusion.

However, while Plumwood recommends rethinking negation by utilizing *relevance* logic, I utilize the *Heyting-Brouwer* logic developed by Cecylia Rauszer instead, which is a type of *intuitionist* logic. *Heyting-Brouwer* logic is an excellent fit for modeling social identity because it allows us to isolate, represent, and ultimately abandon the assumptions of immutability, exhaustiveness, and exclusivity.

Nevertheless, this paper is intended as a friendly supplement to Plumwood’s groundbreaking work, rather than a criticism of it. By demonstrating that Plumwood’s concerns regarding negation can be accommodated in *Heyting-Brouwer* logic, this paper expands our available options for effectively modeling social identity. In doing so, this paper contributes to an emerging literature that revisits Plumwood’s concerns regarding negation and seeks to provide better models of social identity.³ While it may well be the case that some logical systems produce more effective models of social identity than others, it is worthwhile to start by examining our various options for modeling social identity using a variety of non-classical logics prior to any such determination. Hopefully, indulging in such examinations will

³See, for example, (Eckert Forthcoming).

ultimately allow for better critical engagement with a more diverse set of social phenomena.

Through combining Plumwood’s examination of the problematic features of dualisms with the flexible, alternative forms of negation in Rauszer’s *Heyting-Brouwer* logic, my analysis here allows for a radical transformation of our logical treatment of social identity. By working in a logic that allows me to selectively introduce and remove problematic assumptions about the way social identities are related, it is indeed possible to recognize the contribution and significance of individual categories, re-conceive relations in more integrated ways, reclaim denied areas of overlap, and affirm positive or independent sources of definition.

Contrary to the critique that logic is hopelessly mired in abstraction away from lived experience, my models provide a guide for individuals who are working to question and interrogate their own role in reinforcing problematic dualistic conceptions of social identity. With my models, I demonstrate that our self-identifications with individual social categories reflect far more than our assumptions about ourselves. In fact, our self-identifications tend to demonstrate how we consider ourselves to be in relation to other people. If we are not careful, the assumptions guiding these relations can reinforce existing social hierarchies and problematic habits of othering. However, as this paper shows, with a bit of caution we can reclaim logic as a tool for transforming our treatment of social identity altogether.

2 Setting Up Models

As we engage with the social world, we are constantly relying upon various assumptions about how we are situated in the given social context. These assumptions help us to create various pictures, or models, that help us to navigate the social world more efficiently through understanding ourselves in relation to the social categories which are available to us.

In this section, I introduce the features of the logical models I use to represent how an agent might situate themselves in a given social context. By rendering these pictures using the tools of formal logic, I aim to illuminate how some of the most common social contexts are riddled with assumptions that can be unduly harmful to individuals as they engage in the process of self-identification.

2.1 Like a Map

A map represents a physical place. It simplifies the place, making certain assumptions about what to include and what to leave out. The specific information included on each map will depend on what is most important or helpful for individuals navigating in that particular context. A map of a busy downtown area is more likely to highlight prominent streets, bus routes, and businesses, than it is to highlight topographical features such as elevation, vegetation, and water levels. Certainly a map of a busy downtown could include information about topographical features (many modern digital maps do). But since this information is less helpful for typical trips in this area, like taking the bus to a local coffee shop, it makes sense to leave out this information when creating the map. Conversely, a map of a hiking trail is much more likely to highlight topographical features. In this case, topographical features are more relevant for common trips in this area, like hiking to a natural landmark.

For this project, I am mapping a social phenomenon rather than a physical one, and I use logical tools rather than cartographic ones. But, the process of simplification is similar. Like maps of physical locations, models of social phenomena must simplify the situation somewhat.

These models focus on details that can help agents navigate contexts where the demonstrated understanding of the given social categories is prominent. Just as the choices about what to include on a map of the physical world are not random, the choices I make in the construction of these models are not random either. I've made certain choices because these choices are useful in helping me represent how an agent is understanding themselves in relation to social categories. Just as a map of the busy downtown area might differ from a map of a hiking trail, what I include in a given model is shaped by what is most useful for the present analysis.

In practice, these models are a collection of dots and arrows. Roughly speaking, each dot will represent a way that the agent might self-identify at a given time and each arrow represents a way that the agent might change their thinking about their relationship to the social categories which are available to them.

2.2 Building Blocks of Models

So far, I have argued that it is best to think of these models as maps, or visualizations, of the social world and have given a brief overview of the phenomenon I aim to represent. With this context in mind, I now introduce the four major building blocks of the models I will use.

To begin, the whole picture is the model. In my diagrams, each model is presented enclosed in a box and labeled $\mathcal{M}_n$ where n is a number indicating which model we’re talking about.

Any given model is a snapshot of how an agent is thinking of their self-identity with respect to the given social space. By situating the agent at a particular point in the model and understanding how all of the pieces of the model interact, we can isolate the options that the agent takes themselves to have with respect to their identification with the available social categories.

Each model will have the following pieces:

Σ : The Set of Specifications

In a model, each dot or “stopping point” is a specification. Each of these represents a way that the agent might self-identify at a given time.⁴

For reference, each specification will be labeled S_n where n is a number indicating which specification we’re talking about.⁵ Collectively, all of the specifications in a model are included in the set of specifications for that model. This set of specifications is called Σ .

When defining or introducing a model, I will list the members of Σ to indicate how many specifications are included in the given model. For example, if a model were to have just two specifications, labeled S_1 and S_2 respectively, the set of specifications would be described like this: $\Sigma = \{S_1, S_2\}$.

In order to understand how the specifications are positioned with respect to one another, we will need to consider how the specifications are related. This is done using the next piece of the model.

⁴Note that this implies that self-identifying is a bit more than just claiming a predicate. Rather it is claiming a predicate given the rules you take to be operative regarding that predicate at the particular point.

⁵There are infinitely many specifications available: $S_n, n \in \mathbb{N}$.

R : The Accessibility Relation

The pattern of arrows connecting the individual specifications in a model is given by the accessibility relation, R . This relation tells us how the individual specifications are positioned with respect to one another. Since each specification represents a way that the agent might self-identify at a given time, the relations between specifications indicate the available paths for the agent to change their thinking about their identification with available social categories. Essentially, we can trace the various ways that an agent might travel along arrows in a model from one particular self-identification, or specification, to another in order to understand how their thinking about their social identity might change over time.

In terms of notation, R yields a set of ordered pairs that tells us which dots are connected by arrows in which direction. For example, if there is an arrow from S_1 to S_2 , then the pair $\langle S_1, S_2 \rangle$ is in the set of pairs given by R . When talking about a particular arrow or path, the notation $\mathcal{R}(S_n, S_m)$ will be used to name the path between S_n and S_m for any two such specifications. The arrow from S_1 to S_2 would thus be described as $\mathcal{R}(S_1, S_2)$.

One interesting feature of the accessibility relation is that the properties of the pattern of arrows reflect the assumptions that are operative in the agent's thinking about their relationship to social categories.⁶ I'll talk more about the significance of this in Sections 3 and 4. For now, it is important to note that in all models I'll use in this paper, R is both **reflexive** and **transitive**.

An accessibility relation is **reflexive** if and only if for any specification S_n it is the case that $\mathcal{R}(S_n, S_n)$. Or, in other words, an accessibility relation is reflexive just in case there is an arrow from each specification to itself. Practically speaking, this means that in all the models I discuss, the agent always has the option to "stay put" or to maintain their current self-identification.

An accessibility relation is **transitive** if and only if for any specifications S_i , S_j , and S_k it is the case that if $\mathcal{R}(S_i, S_j)$ and $\mathcal{R}(S_j, S_k)$, then also $\mathcal{R}(S_i, S_k)$. Essentially, transitivity guarantees that every time there is a step-wise path following forwards arrows from S_i to S_k via stopping at S_j , there

⁶These assumptions can also be used to identify which logic is reflective of the agent's thinking. While *Heyting-Brouwer* logic is extremely flexible, adding certain assumptions in the construction of a model can approximate the function of other logical systems.

is also a direct path from S_i to S_k that skips S_j .⁷ This is important because, in practice, models typically include unique specifications to indicate each possible step-wise change. While agents do sometimes make changes to their self-identification gradually, transitivity means that long paths along arrows in the same direction can be made shorter. In this way, an agent can make several changes in their self-identification at once using a direct path or make the same changes gradually through following a step-wise path.

In practice, it will often be convenient for the accessibility relation to be **antisymmetric**. An accessibility relation is **antisymmetric** if and only if for any specifications S_l and S_m : if $\mathcal{R}(S_l, S_m)$ and $\mathcal{R}(S_m, S_l)$, then $l = m$. Essentially, antisymmetry states that arrows are uni-directional between distinct specifications. If there is an arrow from S_l to S_m and an arrow from S_m to S_l , then l and m are the same specification. This is a good idea in practice because it reduces redundancy.⁸ However, including two specifications that can see one another is sometimes useful in rendering complex models.

Δ : Set of Agents

The individuals in question are the agents. It turns out that modeling how one individual is thinking about their identity with respect to a given social space is already quite complex, so all of the models I’ll discuss here have only one agent, whom I’ll refer to using the lowercase letter a .⁹

However, there is no significant reason (technically-speaking) why we couldn’t add more agents. To allow for this, it is useful to generalize somewhat. Therefore, each model will have a set of agents, labeled Δ . Using

⁷Note that, as will be important later, transitivity is directional. It applies to patterns of arrows going in the same direction. This means that not all step-wise paths will collapse to a single path. For example, if the pattern of arrows was such that $\mathcal{R}(S_2, S_1)$ and $\mathcal{R}(S_2, S_3)$, then moving from S_1 to S_3 would require first moving from S_1 to S_2 , traveling in the reverse direction along the $\mathcal{R}(S_2, S_1)$, and then moving from S_2 to S_3 , traveling forwards along the $\mathcal{R}(S_2, S_3)$ arrow. Since this step-wise path involves both forward and reverse movement, there is no guarantee that it can be accomplished in a single step.

⁸This is often convenient because two specifications that can see one another will have the same truths, so it is not helpful to list them separately in most applications.

⁹I use the pronouns “they” and “them” to refer to the agent throughout the paper. This is done for uniformity and clarity in models where the agent is considering, committing to, and rejecting identification with multiple gendered social categories in rapid succession and is not intended to reflect real-world practices regarding pronoun use.

this terminology, the category of agents that only has agent a and no other agents would be described like this: $\Delta = \{a\}$. This will be the case for all models I discuss in this paper.

σ : Function to Handle Truth Values

The final piece of each model is a function that handles truth values by stipulating which basic claims about identity hold of each agent at each specification. This function is called σ .

At their most basic level, claims about identity can be expressed using the language of predicates. The formal language I use here includes predicate symbols,¹⁰ which represent the different social categories with which an individual might self-identify. Informally, for ease of reading, I'll use capital letters that relate to the social category in question to represent these predicates.

Essentially, what σ does is take each predicate-specification pair and provide a set of individuals of whom the given predicate is true at the specification in question. Typically, the predicates which are true of an agent at a given specification is a matter of what the particular model is intended to show. If, for example, we want to model an individual who self-identifies with category X , we might choose to situate the agent a at a given specification, say S_1 for instance, and stipulate that the predicate X is true of a and only a at that specification. If we did so, this would be written as follows: $\sigma(\langle X, S_1 \rangle) = \{a\}$.

In the *Heyting-Brouwer* logic I use here, the σ function is subject to the following constraint:

Monotonicity:

For any $S_1, S_2 \in \Sigma$, any predicate P , and any agent $\xi_n \in \Delta$,
 if ξ_n is in the set of agents assigned to P by σ at S_1 and $\mathcal{R}(S_1, S_2)$,
 then ξ_n is in the set of agents assigned to P by σ at S_2 .

In other words, monotonicity requires that once a predicate is true of an agent at a specification, it must remain true of them at all specifications that can be reached via following arrows forwards from that point. Typically, in

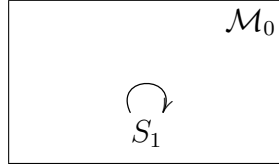
¹⁰There are infinitely many of these non-logical n -ary predicate symbols: $P_m^n(\xi_1, \xi_2, \xi_3, \dots \xi_n)$, $m \in \mathbb{N}$. When it is necessary to talk about these predicates collectively, they can be categorized together in Pred , the set of all such non-logical predicates.

models like this, the only type of change allowed is forwards movement.¹¹ This may initially seem problematic, since agents do sometimes retract their commitments to social categories.

To address this concern, I allow both forwards and reverse movement in my models. This change allows the models to more effectively capture dynamic engagement with social identity.¹² Viewed this way, monotonicity does not place a constraint directly upon an agent's identity-thinking, but rather upon how the specifications are positioned in the model: specifications that can be reached by via following forward paths are those where an agent retains their existing commitments, whereas specifications that can be reached by following reverse paths are those where an agent retracts at least one commitment.

Putting the Pieces Together for a Model

Thus, a model is specified as follows, using an ordered tuple to organize the four pieces given above: $\mathcal{M}_n = \langle \Sigma, R, \Delta, \sigma \rangle$. For example, the model would be described as follows, assuming that the predicate X is true of a at S_1 :^{13,14}



$$\mathcal{M}_0 = \langle \Sigma = \{S_1\}, \mathcal{R} = \{\langle S_1, S_1 \rangle\}, \Delta = \{a\}, \sigma(\langle X, S_1 \rangle) = \{a\} \rangle$$

¹¹For traditional presentations of models like this, see (Priest 2008).

¹²This choice is explained in more detail in section 3.2 where I demonstrate how both forwards and reverse movement function in a model.

¹³Note that the picture representing the model does not indicate directly which claims are true at S_1 , even though σ does indicate that X is true of a at S_1 . For this purpose, I'll make use of some additional notation introduced in the next section.

¹⁴I draw the reflexive arrow here and for all one-specification models because it helps to emphasize that despite there being only one specification, there is still a path for forwards and reverse movement. In more complicated models, I omit both reflexive and transitive arrows for ease of reading.

3 Modeling Basic Identity-Thinking

This initial model $\mathcal{M}_0$ provides an excellent starting point for modeling basic identity-thinking. In particular, if we take agent a to be situated at S_1 in the model $\mathcal{M}_0$, this provides a way to represent the situation described in the following general case:

Pressure to Commit

People constantly keep labeling me as X . As a result, I've internalized label X and now self-identify as X . Except, when I think about it, X doesn't feel right for me, even though X seems like it works for others. I'd rather not identify one way or the other regarding X .

In this case, the agent is pressured into identifying as X but wants to retract this commitment. This pressure is, in many cases, an unreasonable restriction on the agent's exploration of their own social identity.

In order to see what leads the agent to be "trapped into" identifying this way, it will be helpful to isolate various basic claims an agent might make with regard to a particular social category. This is the focus of the first half of this section. In the second half of this section, I turn to what changes in this identity-thinking might look like. After demonstrating what possible changes might look like, I use this perspective to clarify how the restrictions on this agent's exploration of their identity are functioning.

3.1 Claims about Identity-Thinking

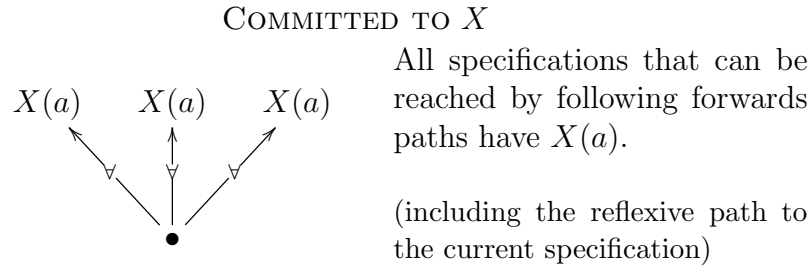
I've already mentioned that each specification represents some way that the agent self-identifies at a given time and that the most basic claims about identity-thinking can be expressed using predicates. These basic claims about social identity are the direct claims that an agent might make about their self-identification with a particular social category.

Recall, above, the initial model for an individual who self-identifies with category X . To capture this self-identification, we situated agent a at a specification and stipulated that the predicate X was true of a at that specification. Identifying as a member of a particular social category is the first basic claim about an agent's self-identity that can be expressed using these predicates. In this section, I will introduce our first negation operator to expand this list of basic claims about an agent's self-identity to include four

options: being committed to X , being uncommitted to X , rejecting X , and being purely uncommitted to X .

Each of these basic claims affects the shape of the overall model. For example, as we have already seen with the definition of monotonicity, if a claim is true at the current specification, it remains true at all specifications that can be reached via following a forwards path from the current specification. This, and the nature of the semantic clauses for operators in *Heyting-Brouwer* logic, means that each claim about an agent's self-identity reflects not only what is currently true of them but also the options that the agent takes themselves to have with respect to changing this identification.

In order to demonstrate how these basic claims influence the shape of the model as a whole, I use diagrams to represent what the claim in question requires or allows for in the structure of paths. In these diagrams, $\bullet$ represents the current specification. It is assumed that the claim in question is asserted at $\bullet$. When a claim requires that all specifications reached via following paths in a particular direction have a certain characteristic, I indicate this by including multiple arrows labeled with $\forall$, each of which terminates in a label that gives the characteristic in question. When a claim requires the existence of a path in a particular direction leading to a specification with a certain characteristic, I indicate this by including a single arrow labeled with $\exists$ which terminates in a label that gives the characteristic in question. When a claim allows for a path with a certain characteristic, but does not require it, I indicate this by including a single arrow labeled with $\diamond$ which terminates in a label that gives the characteristic in question. These diagrams are not proper models in the sense described in the previous section.¹⁵ They are simply intended to be illustrative of the definition of each of these basic claims about self-identity.



If an agent is committed to X , then X is true of them at the given

¹⁵As such, they do not show reflexive or transitive arrows.

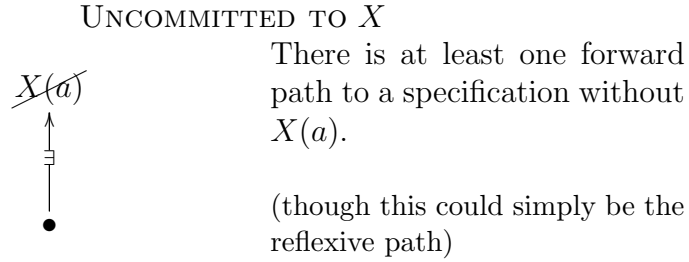
specification. Since these models are monotonic, once a claim is true at a specification, it remains true at all specifications accessible forwards from this point. This means that committing to X is not just the claim that one currently identifies with the social category. Rather, it is a commitment to continue identifying as X along all forward paths.

To describe this, it is useful to have an additional piece of notation.¹⁶ To state that any arbitrary claim Φ is true at a specification, it is said that: $\mathcal{M}_n, S_m \models \Phi$. Thus, the claim that X is true of a at the specification S_1 can be expressed as follows:

$$\mathcal{M}_0, S_1 \models X(a)$$

This can be interpreted as saying that a has committed to claiming, affirming, or identifying as X at S_1 in the model $\mathcal{M}_0$. This reflects the situation described in the initial model constructed above.

Of course, a predicate may also fail to be true of an agent at a given specification.¹⁷ If this is the case, then the agent is said to be uncommitted to X .



¹⁶We could instead, as we did above when describing the construction of the model, make use of σ , the function that keeps track of which claims are true at each specification, to indicate that X is true of a at S_1 :

$$\sigma(\langle X, S_1 \rangle) = \{a\}$$

However, σ can only indicate whether a given predicate is true of a given agent at a particular specification according to the construction of the model. It cannot directly capture, for example, a situation where a rejects X . Additionally, this notation presumes a is the only agent who identifies with X at S_1 . Since there is only one agent in all the models I discuss, σ will always return either the empty set or the singleton of a and this extra assumption doesn't do much work. However, in more general applications it will be helpful to claim that X is true of a without making any claims about other agents. Thus, in order to represent more claims about identity, a different notation is needed.

¹⁷Notice that I say “fail to be true” – in this logic a claim that fails to be true is not thereby false, where falsity is understood as the truth of the negation. More on this below.

Similar to above, in order to state that any arbitrary claim Φ is not true at a specification, it is said that: $\mathcal{M}_n, S_m \not\models \Phi$. Thus, the claim that X is not true of a at the specification S_m can be expressed as follows:

$$\mathcal{M}_n, S_m \not\models X(a)$$

This indicates that a has not committed to X at S_m , which means that not all forward paths have $X(a)$. Note that some forward paths could still have $X(a)$, it simply must be the case that not all of them do.

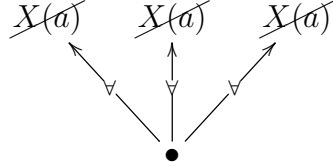
In one sense, being uncommitted to X is the straightforward opposite of being committed to X , since one cannot be both committed and uncommitted at the same time. At minimum, being uncommitted to X requires that X is not true of a at S_m , the current specification. If X is not true of a at S_m , then the reflexive path $\mathcal{R}(S_m, S_m)$ is a forward path that lacks $X(a)$. If X were true at S_m , then it would have to be true at all specifications forward from that point and there would be no forwards path that lacks X . Indeed, for any predicate P and any agent $\xi_n \in \Delta$, the agent ξ_n is either committed to P or uncommitted to P at each specification.¹⁸

It might be tempting to think of being uncommitted to X as a kind of indecision, but this would not be precisely accurate. While being uncommitted to X does in fact imply a lack of commitment to X , being uncommitted to X is compatible with other attitudes towards a given social category. In particular, rejecting X implies being uncommitted to X .¹⁹ To see how, let's turn to what rejecting X looks like.

¹⁸This is because it must be the case that either $\mathcal{M}_n, S_m \models P(\xi_n)$ or $\mathcal{M}_n, S_m \not\models P(\xi_n)$.

¹⁹The reverse implication does not hold. This is a helpful feature of the *Heyting-Brouwer* logic I work with here. When modeling an agent's thinking about their self-identification, it is often useful to distinguish between merely being uncommitted to a given social category at a particular time and rejecting identification with a given social category. If necessary, it is still possible to restrict this in specific models so that being uncommitted to X implies a rejection of X . Indeed, I give an example of this at the end of Section 3.2.

REJECTING X



All specifications along forward paths do not have $X(a)$ / No specification along a forwards path has $X(a)$.

(including the reflexive path to the current specification)

If an agent rejects identifying as X , this is a bit more than just being uncommitted to X . Essentially, rejecting X is a way of ruling out options that involve any forward movement towards X . Just as with making a commitment to X , rejecting X is persistent among all paths that can be reached forwards from this point. This denial or distancing from the social category is more than just a claim about what is true now. Rather, it is a refusal to identify as X along all forwards paths.²⁰

To describe this, I use the first of our two negation operators.²¹ Here is the semantic clause:

$$\mathcal{M}_n, S_m \models \neg\Phi \text{ iff for all } S_k \text{ accessible from } S_m, \text{ it is the case that } \mathcal{M}, S_k \not\models \Phi.$$

Notice that when we apply this semantic clause to $X(a)$, it captures the exact situation depicted above.²² All S_k accessible from S_m are the specifications that can be reached following a forwards path. If all such specifications are such that $\mathcal{M}_n, S_k \not\models X(a)$, then a has rejected X at S_m .

Thus, the claim that a has rejected X at the specification S_m can be expressed as follows:

$$\mathcal{M}_n, S_m \models \neg X(a)$$

²⁰To see why, recall that R is transitive. Any specification that can be reached via following a path of only forwards arrows is reachable via a single step.

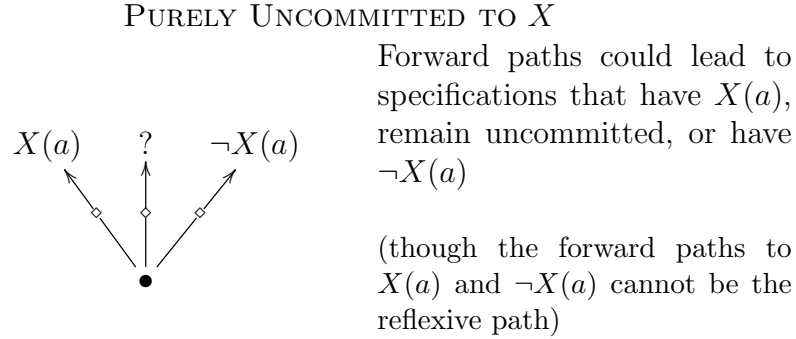
²¹Readers who are familiar with *intuitionistic* logic will notice that this is the intuitionist negation. Intuitionist negation requires a bit more than a claim simply failing to be true, hence why it is perfect for modeling rejection.

²²Here, I describe what it looks like to reject a predicate. It is possible to reject more complex claims about identity through applying the semantic clause to the claim in question.

Now we can see why rejecting X implies being uncommitted to X . Since the accessibility relation R is reflexive, S_m is one of the specifications accessible from itself. Thus, if $\mathcal{M}_n, S_m \models \neg X(a)$, then, according to the semantic clause for $\neg$ it will be the case that $\mathcal{M}_n, S_m \not\models X(a)$ which means that if a rejects X , then a is also trivially uncommitted to X .

Above, I noted that an agent is uncommitted to X whenever X fails to be true of them at the given specification. However, just because a predicate X fails to be true of an agent at a specification does not thereby mean that the agent rejects X . In order to reject X , the agent must also rule out any forward movement towards X . Thus, while rejecting X implies being uncommitted to X , being uncommitted to X does not imply rejecting X .

In light of this, it is useful to distinguish between being uncommitted to X and being purely uncommitted to X .



When a is uncommitted to X and also does not reject X , this reflects being purely uncommitted to X .²³ Such cases represent a situation where the agent has not made a commitment one way or the other about the social category at the specification. This can be expressed as follows:

$$\mathcal{M}_n, S_m \not\models X(a) \text{ and } \mathcal{M}_n, S_m \not\models \neg X(a)$$

This preserves a kind of “undecided” or “neutral” option, which is needed for modeling much of our thinking about social identity. Even if we ultimately decide that it is not possible to be purely uncommitted to a social category

²³This is the first example of how basic identity thinking can be combined. Below, I’ll define semantic clauses for the remaining operators in *Heyting-Brouwer* logic to expand this list of combinations. I mention this one here because it requires making use of the meta-theoretical notation.

in a given context, it is still useful to be able to model this way of thinking about social identity. It is, at the very least, a common attitude that people tend to have in their self-identification.

When having an “undecided” or “neutral” option is not desirable, it is possible to restrict the model so that being uncommitted to X requires rejecting X . This is one way of expressing the idea that it is not possible to be fully neutral with regard to one’s attitude towards a given social category. If any failure to claim membership in a social category is a rejection of it, then it is not possible to be purely uncommitted to any social category.

But this approach is not without its problems. Indeed, restricting a model so that being uncommitted to X requires rejecting X will re-introduce the harmful effects of dualisms that Plumwood is so concerned about. To see how, see the model I develop in the next section.

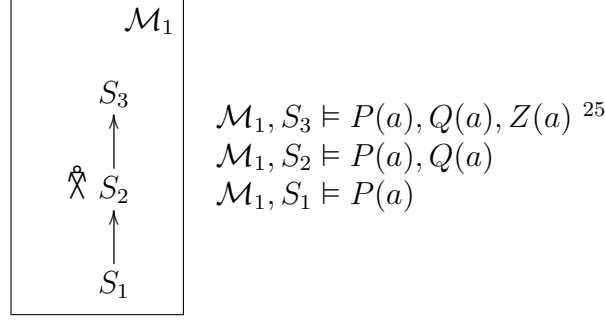
Thankfully, this approach is not the only way to demonstrate that a purportedly neutral position towards an identity can never be fully neutral. While I have presented these attitudes towards X as distinct, individual attitudes that one might have, the broader picture of an agent’s thinking about their relationship to available social categories may be significantly more complicated than this, as I demonstrate below.

Self-identifying as X is not as simple as just asserting “I am X .” Rather, it is to situate one’s understanding of oneself with regard to a complex set of commitments and rejections. Self-identifying as X is to assert a more complex claim of the form: “I am X , where I take X to operate in the following way...” As is evident from the definition of committing to X , the ways in which an agent self-identifies affect more than what is currently true of them. Through committing to X , an agent positions themselves a particular point in a broader context that reflects a specific understanding of available paths. When an agent is considering their relationship with more than one social category at a time, the interactions among these individual claims about self-identification make the picture significantly more complex.

3.2 Changes in Basic Identity-Thinking

I’ve already mentioned that each arrow represents a way that the agent might change their thinking about their relationship to available social categories. Essentially, each move along an arrow reflects one option the agent has for changing their self-identification. Let’s take a closer look at what these changes might look like.

In order to talk concretely about these changes, the following model will be useful.²⁴ Suppose that agent a is considering P , Q , and Z in the following way:



Let's start by situating a at S_2 in $\mathcal{M}_1$. From this point, a has two options: forwards movement to S_3 or reverse movement to S_1 . In this section, I argue that it makes sense to allow for both forwards and reverse movement and that we should think of moving forwards along arrows as stabilizing and moving in the reverse direction along arrows as destabilizing.²⁶ Roughly speaking, stabilizing allows for adding additional claims while keeping all existing claims, whereas destabilizing retracts acceptance of claims we already have.²⁷

²⁴As is standard, in models with more than one specification, I omit reflexive and additional transitive arrows for ease of reading throughout. The model pictured above can be defined as follows:

$$\mathcal{M}_1 = \left\langle \Sigma = \{S_1, S_2, S_3\}, \mathcal{R} = \left\{ \langle S_n, S_n \rangle \text{ for all } S_n \in \Sigma, \right. \right. \\ \left. \left. \langle S_1, S_2 \rangle, \langle S_2, S_3 \rangle, \langle S_1, S_3 \rangle \right\}, \Delta = \{a\}, \right. \\ \left. \sigma = \{a\} \text{ for the following pairs: } \langle P, S_3 \rangle, \langle Q, S_3 \rangle, \langle Z, S_3 \rangle, \right. \\ \left. \langle P, S_2 \rangle, \langle Q, S_2 \rangle, \langle P, S_1 \rangle \text{ and } \sigma = \{ \} \text{ otherwise} \right\rangle$$

²⁵Here I use commas to separate the individual claims that are true at each specification. In this case, $P(a)$, $Q(a)$, and $Z(a)$ are all true at S_3 in $\mathcal{M}_1$.

²⁶This language is inspired by Dembroff, who uses the language of destabilizing in relation to social identity in their definition of a critical gender kind (Dembroff 2020).

²⁷Claims here is meant in the most general sense. In this particular example, the claims in question are all commitments to the available social categories P , Q , and Z . This is done for the sake of simplicity. In practice, the agent could stabilize by adding other types of claims (such as rejecting a social category, for instance).

Forwards Movement: Stabilizing Change

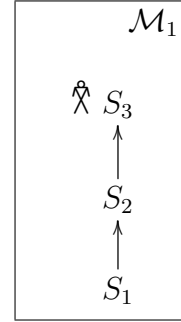
If agent a is situated at S_2 in $\mathcal{M}_1$, then they have currently committed to P and Q . One option for changing their self-identification would be to make an additional commitment, in this case by committing to Z .

In order to commit to Z , agent a moves from S_2 to S_3 by following the $\mathcal{R}(S_2, S_3)$ arrow in the forwards direction. When agent a does this, they retain their commitment to P and Q . Indeed, since S_3 is accessible from S_2 and the accessibility relation R is monotonic, all of the claims that were true at S_2 remain true at S_3 . Still, this action of moving from S_2 to S_3 does represent a significant change for agent a . At S_1 , agent a has made a commitment to Z that they did not have at S_2 . Now, at S_3 , agent a has made a commitment to P , Q , and Z .

Any action that travels forwards along an arrow is stabilizing. Just as in the example above, all moves in the forwards direction along an arrow will preserve an agent's existing commitments. This movement also allows for new claims to be made, so long as they are compatible with what was already true at the initial specification.²⁸

As I mentioned above, in models like this, forwards movement is typically the only type of movement possible. On this understanding of models, one typically begins by situating the agent at the bottom of a model, considering the various choices they might make to move up a branch. Hence, in this case, the model suggests that the agent might have begun at S_1 by identifying with only P , then made a commitment to Q by moving to S_2 , and finally made another commitment to Z by moving to S_3 . At a first glance, this is sensible: agent a initially knows less about their identity and then over time gradually comes to know more as they consider (and ultimately make) additional commitments.

However, given the constraint of monotonicity, only allowing forwards movement is incredibly limiting. A common change in an individual's identity-thinking is retracting a commitment to a social category so that one no longer self-identifies in the same way as before. While stabilizing change is an important type of change in self-identification, it is by no means the only type



²⁸Note that, in practice, specifications are only listed separately if there is at least one claim that distinguishes them. Technically, the move from S_2 to itself forwards along the reflexive arrow would count as a stabilizing action, but only trivially. The interesting cases of stabilization are ones where the agent moves to a new specification as a result.

of change in identity-thinking and logical models of identity should be able to reflect this.

If models only allow for forwards movement and we think of each arrow as one potential step that an agent might take in changing their identity-thinking, then once agent a has committed to P and Q , they can never retract this commitment. While a is free to make additional claims and move further up in the model, if they can only move forwards, they will never be able to reach a specification without P or without Q .

This would be an unreasonable limitation for all models of identity-thinking. Models of identity-thinking need to allow for flexibility in cases where agents do retract their commitments to a given social category. Happily, it is possible to allow for this flexibility. As I demonstrate next, the dual negation operator in *Heyting-Brouwer* logic provides an excellent mechanism for modeling reverse movement along arrows.

Reverse Movement: Destabilizing Action

Above, agent a made a stabilizing change by moving from S_2 where they had committed to P and Q to a new specification S_3 where they are committed to P , Q , and Z . Now suppose that, upon reflection, agent a decides that they no longer want to commit to Z .²⁹

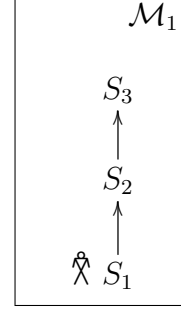
In order to retract this commitment, agent a needs to move back to S_2 . If we allow for movement in the reverse direction, then such a change is possible. Indeed, this is relatively easy: there is already an arrow $\mathcal{R}(S_2, S_3)$ which agent a can follow in the reverse direction back to S_2 .

Generally speaking, as an agent moves in the reverse direction along an arrow, they rethink or reassess claims to which they had already assented. This reflects a different type of change in identity-thinking. Rather than adding additional claims, in a destabilizing action the agent changes thinking

²⁹This case is an example of what Dembroff refers to as existential destabilizing because the destabilizing is based on the agent's felt or desired social role, rather than on the agent's social or political commitments regarding social norms. We might, alternatively, imagine an agent who wants to articulate or embrace an ideology which provides for an option where one does not identify with a particular social category. To do this, the agent might destabilize merely in order to create another available option, even if they still want to commit to the given social category. This would then be an example of principled destabilizing (Dembroff 2020).

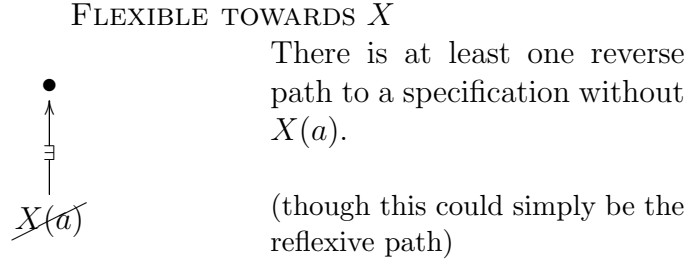
by retracting claims to which they had already assented.³⁰

In general, as long as there is a specification that can be reached via following arrows in the reverse direction where agent a has not committed to the predicate, agent a can retract their commitment to the predicate. In the example above, then, agent a can also retract their commitment to Q by moving in the reverse direction along the $\mathcal{R}(S_1, S_2)$ arrow from S_2 to S_1 .³¹



However, in this case, once agent a reaches S_1 , there are no more options for reverse movement in the model $\mathcal{M}_1$. Indeed, all specifications in this model are such that agent a commits to P . Based on agent a 's available options in this model, a must always commit to P .

This is where another attitude or basic claim in identity-thinking comes into play. The reason that agent a must always commit to P in the model above is that agent a is not flexible towards P at any specification in $\mathcal{M}_1$ in the following sense:



An agent is only flexible towards X when they have destabilized their commitment towards X .³² This means that there is at least one reverse path

³⁰Just as with stabilizing actions above, there is always a trivial destabilizing option available where the agent merely follows the reflexive arrow in the reverse direction back to the current specification. The interesting cases remain ones where the agent destabilizes to a new specification.

³¹Since the accessibility relation is transitive, agent a could also retract their commitment to both Z and Q in a single step by moving in the reverse direction along the $\mathcal{R}(S_1, S_3)$ arrow directly from S_3 to S_1 .

³²Note that we might also think of being purely uncommitted to a social category as a type of forward-looking flexibility. For my purposes here, I assign reverse-looking flexibility this label because this is the type of flexibility that allows an agent who already claims something to move to a position where they can be purely uncommitted to that claim.

available where they do not commit to X . Essentially, being flexible towards X guarantees that even if the agent does eventually commit to X , there will be a reverse path available where they do not commit to X .³³ Thus, as long as we allow reverse movement, the agent will always be able to reach a point without X .

This is expressed with the dual negation operator.³⁴ The semantic clause is as follows:

$$\mathcal{M}_n, S_m \models \neg_D \Phi \text{ iff } \text{there exists } S_k \in \Sigma \text{ such that } \mathcal{R}(S_k, S_m) \text{ and } \mathcal{M}_n, S_k \not\models \Phi$$

An agent has destabilized their claim that Φ , or is currently flexible towards Φ just in case there is a reverse path to a specification where they do not claim Φ .

Thus, in the example above, agent a is flexible towards Q and Z at all three specifications because from each of these specifications, it is possible to follow an arrow in the reverse direction to S_1 where agent a does not commit to Q or commit to Z . But, as we just saw, there is no specification in $\mathcal{M}_1$ where agent a has not committed to P .

Notice that claims require work to destabilize once they're asserted. Especially in the case where there is no existing specification where the agent a doesn't claim the predicate in question, this work can be substantial. In such cases, the agent currently does not have a flexible attitude towards the predicate at any point in the model. As a result, opening up a new destabilization option forces a change in the model that the agent is using to capture their identity-thinking. The definition of being flexible towards a predicate indicates exactly which changes we'd need to make in order to create some flexibility towards P : we'd need to open up a destabilization option where agent a does not claim P . In this case, we might modify the model by adding a new specification S_4 below S_1 where a has not yet committed to P .³⁵

³³I'll talk about how flexibility towards a predicate can be combined with other attitudes in the next section. These combinations represent more complex types of flexibility, like a flexible commitment or a flexible rejection.

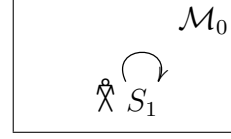
³⁴This "reverse-looking" negation has other names as well. I have chosen to refer to it as the dual negation throughout because this reflects its function and is the terminology used by Rauser in her initial formulation of *Heyting-Brouwer* logic.

³⁵To do this, the new model $\mathcal{M}_1^*$ will be identical to $\mathcal{M}_1$ with the following additions:

In this way, destabilization is a method for forcing flexibility towards a predicate. Once a claim has been destabilized and the reverse path created, this option remains available to the agent in a modified model. The agent need not avoid committing to the predicate indefinitely to preserve this flexibility. As long as they take the new specification seriously as an available reverse path, flexibility is maintained.

Now that we have seen the radical potential of destabilizing change, we have a nice way to capture the restrictive problem in the **Pressure to Commit** case. In this case, agent a has been pressured into committing to X , but no longer feels that this is right for them and would prefer not to identify one way or the other regarding X . Above, I mentioned that $\mathcal{M}_0$ provides an excellent model of the pressure the agent is under in this case.

Recall that in this model a has made a commitment to X at S_1 : $\mathcal{M}_0, S_1 \models X(a)$. Since S_1 is the only specification in $\mathcal{M}_0$, there is no destabilization option available for agent a regarding X and a cannot reach a specification where they do not claim X .



Thus, the pressure on agent a operates through a denial of flexibility. As the agent realizes that X doesn't feel right for them, they recognize that the lack of flexibility is problematic for their articulation of their social identity. In this case, the agent states that they would prefer "not to identify one way or the other regarding X ." To accomplish this, agent a would need to be purely uncommitted to X . But since they are currently at a specification where they claim X and there is no destabilization option available, they cannot become purely uncommitted to X unless they have some flexibility with regard to X . Through destabilizing their commitment to X , agent a could recover some flexibility towards X . In this way, destabilization allows the agent to open up additional paths that are not dependent upon the commitments that may have been given to them or forced upon them.

Add S_4 to Σ , add the pairs $\langle S_4, S_n \rangle$ for all $S_n \in \Sigma$ to R , and update σ such that it returns $\{ \}$ for the following pairs: $\langle P, S_4 \rangle, \langle Q, S_4 \rangle, \langle Z, S_4 \rangle$. If this is the case, then $\mathcal{M}_1^*, S_4 \not\models P(a)$ and S_4 serves as the destabilization option for all specifications, since there is a reverse path to S_4 from every specification in $\mathcal{M}_1^*$.

4 Modeling Complex Identity-Thinking

Above, I argued that denying flexibility or refusing to allow an agent to destabilize their commitment to a given predicate is often inappropriate and harmful. While this is certainly the case, many of the more common limitations on identity-thinking arise when we think about claims involving multiple predicates. Consider the following case:

Forced Association

I self-identify as W . But people say that if I'm W , then I must also be Y . Claiming W feels right for me and I still want to identify as W , but I don't want to end up being Y as a result. It seems to me that there should be a way to be W without also being Y .

In this case, the agent self-identifies as W and expresses a commitment to W . However, the claims about social categories in their context suggest that claiming W requires also claiming Y , which the agent does not want to do. In order to understand cases like this one, more resources are needed. Throughout this section, I discuss some of the options available for putting together thinking about self-identity. I then use these tools to analyze an example of **Forced Association** related to toxic masculinity.

4.1 Putting Together Identity-Thinking

As we move towards more complex models, it will be useful to talk about the connections between basic claims about identity. To do so, I introduce the two more operators in *Heyting-Brouwer* logic: $\wedge$ and $\vee$.

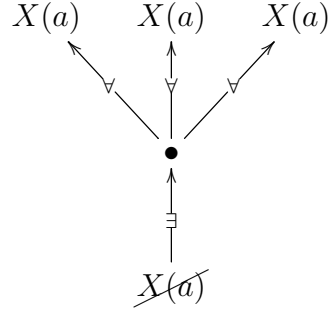
“And”

$$\mathcal{M}_n, S_m \models \Phi \wedge \Psi \text{ iff } \mathcal{M}_n, S_m \models \Phi \text{ and } \mathcal{M}_n, S_m \models \Psi$$

Straightforwardly, the claim $\Phi \wedge \Psi$ is true at S_m in $\mathcal{M}_n$ just in case agent a claims Φ at S_m in $\mathcal{M}_n$ and claims Ψ at S_m in $\mathcal{M}_n$. Taking monotonicity into account, asserting a conjunctive claim like this means that every specification that stabilizes the current specification S_m is such that both Φ and Ψ are true. In order to get to a non- Φ or non- Ψ option the agent would first need to destabilize.

Using this operator, it is possible to capture some interesting nuances of the identity-related attitudes discussed above:

FLEXIBLE COMMITMENT TO X

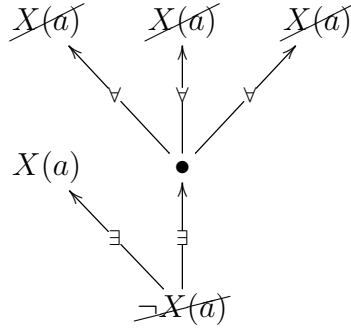


All specifications along forward paths have $X(a)$ (including the reflexive path to the current specification), but there is at least one reverse path to a specification without $X(a)$.

$$\mathcal{M}_n, S_m \models X(a) \wedge \neg_{\mathbf{D}} X(a)$$

If agent a makes a flexible commitment to X , then this means that they claim X , but recognize that they do not have to. Essentially, a flexible commitment is a less assertive commitment to X . While every stabilization option available is still such that agent a claims X , it is also the case that every stabilization option available still allows for a path back to a specification where agent a does not claim X .

FLEXIBLE REJECTION OF X



No forward paths lead to a specification with $X(a)$ (including the reflexive path to the current specification), but there is at least one reverse path to a specification without $\neg X(a)$.

$$\mathcal{M}_n, S_m \models \neg X(a) \wedge \neg_{\mathbf{D}} \neg X(a)$$

If agent a flexibly rejects X , then this means that they reject X , but recognize that they do not have to. Essentially, a flexible rejection is a less assertive rejection to X . While agent a has still ruled out forwards movement towards X , every stabilization option still allows for a path back to a specification where a does not reject X . If there is a reverse path to a

specification where $\neg X(a)$ fails to be true, then according to the semantic clause for $\neg$, there must also be a forwards path to a specification where a claims X . One way of understanding this is that a flexible rejection must be such that the agent takes themselves to be able to retract their rejection and move to a place where they still have the option of committing to X .³⁶

“Or”

$\mathcal{M}_n, S_m \models \Phi \vee \Psi \text{ iff } \mathcal{M}_n, S_m \models \Phi \text{ or } \mathcal{M}_n, S_m \models \Psi$
--

Straightforwardly, $\Phi \vee \Psi$ is true at S_m in $\mathcal{M}_n$ just in case agent a claims Φ at S_m in $\mathcal{M}_n$, claims Ψ at S_m in $\mathcal{M}_n$, or both. Taking monotonicity into account, asserting a disjunctive claim like this means that the agent a has narrowed their available options such that every specification that stabilizes S_m is such that either Φ or Ψ is true. Essentially, they’ve limited their forwards options to Φ options and Ψ options. In order to expand their options beyond these two cases, they would first need to destabilize.

Of course, traditional presentations of *Heyting-Brouwer* logic also include implication, dual implication, and quantifiers.³⁷ However, the operators I have explained thus far provide more than enough to work with for an initial foray into modeling the ways that individual agents self-identify.

4.2 Complex Identity-Thinking in Action

To begin modeling how individual agents often self-identify in complex ways, consider the following articulation of toxic masculinity in the context of a gender binary:

Anything But Feminine: Being an American boy is a setup. We train boys to believe that the way to become a man is to objectify

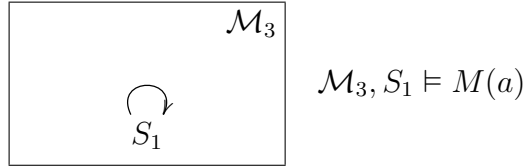
³⁶One might expect that flexible commitment and flexible rejection would be morphologically similar. However, as is evident above, the diagram for flexible rejection is significantly more complex. This is due to the way negation works in *Heyting-Brouwer* logic. As noted above in the initial explanations of commitment and rejection, more conditions must be met for a negated claim to fail to be true than for a basic predicate claim to fail to be true.

³⁷For more information on the dual operators, see (Ferguson 2014), (Priest 2009), and (Priest 2011).

and conquer women, value wealth and power above all, and suppress any emotions other than competitiveness and rage... We tell them, “Don’t be these things, because these are feminine things to be. Be anything but feminine!” (Doyle 2020)

To represent this case, I construct the following model. Model $\mathcal{M}_3$ is such that the only specification is S_1 , there is a path from S_1 to S_1 , a is the only agent, and M is true of a at S_1 .

$$\mathcal{M}_3 = \langle \Sigma = \{S_1\}, \mathcal{R} = \{\langle S_1, S_1 \rangle\}, \Delta = \{a\}, \sigma(\langle M, S_1 \rangle) = \{a\} \rangle$$

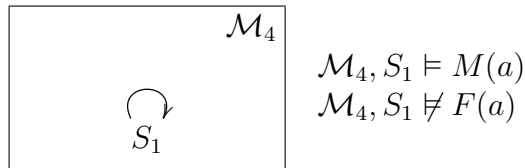


So far, all I’ve done is assign “being an American boy/becoming a man” to the predicate M and made it true of the agent in a specification. I’ve started with the smallest possible model, that of only one specification.³⁸ Since the accessibility relation must be reflexive, there is a path from this specification to itself. Then I assumed that M was true of a at this specification.

But this doesn’t yet capture the description of toxic masculinity. To capture this description, we’ll need to talk about the relationship between “being an American boy/becoming a man” and rejecting “femininity.” This means that we’ll need a predicate to represent this second social category of “being feminine.” Let’s use F for this purpose.

Now, modify the previous model so that F is true of no agents at S_1 . Note that this is the exact same model as above, except that it has an additional predicate, F , which does not apply to any agents.

$$\mathcal{M}_4 = \left\langle \Sigma = \{S_1\}, \mathcal{R} = \{\langle S_1, S_1 \rangle\}, \Delta = \{a\}, \begin{array}{l} \sigma(\langle M, S_1 \rangle) = \{a\} \\ \sigma(\langle F, S_1 \rangle) = \{ \} \end{array} \right\rangle$$



³⁸This is often a good way to start building a model, subsequently adding pieces as needed to represent the situation in question. As I demonstrate in this section, though, the one-specification model turns out to be an excellent model for the kind of restrictive assumptions present in the articulation of toxic masculinity given above.

At first glance, $\mathcal{M}_4$ might not seem very different from $\mathcal{M}_3$. However, the structural features of this model (namely, the fact that it has only one specification) produce some additional consequences regarding the interaction of the predicates M and F .

Since the only specification that can be reached by following a forwards path is S_1 itself and $\mathcal{M}_4, S_1 \models M(a)$, this means that agent a is committed to M at S_1 . And since S_1 is also such that $\mathcal{M}_4, S_1 \not\models F(a)$, this means that F fails to be true of agent a at all specifications that can be reached along a forwards path from S_1 . Thus, at S_1 , agent a rejects F . Hence, $\mathcal{M}_4, S_1 \models \neg F(a)$.

This representation of agent a 's options is pretty limiting. Indeed, since the only specification is S_1 , there is no destabilization option available for any of the claims asserted at S_1 . As a result, there is no flexibility for agent a with regard to any of their assertions.

The one-specification model $\mathcal{M}_4$ defined above effectively demonstrates how toxic masculinity affects agent a 's engagement with an immutable, exhaustive, and exclusive gender binary in a mutually reinforcing way. Observe that the following complex claims are true at S_1 in $\mathcal{M}_4$:

$$\begin{array}{ll} \text{(exhaustiveness)} & M(a) \vee F(a) \\ \text{(exclusivity)} & \neg(M(a) \wedge F(a)) \end{array}$$

As agent a internalizes the ideology of toxic masculinity, they are confronted with a choice between becoming a man by rejecting femininity altogether (thereby inflexibly committing to M and inflexibly rejecting F) or failing to become a man. If agent a fails to become a man, then they will be grouped into category F . And if they are grouped into category F , they cannot subsequently commit to M because becoming a man requires rejecting femininity.

This model satisfies all the principles of classical logic. For any predicate whatsoever, either it will be true of agent a at S_1 , meaning that agent a has claimed membership in the relevant social category, or it will fail to be true of agent a and they will consequently be said to reject membership in the relevant social category. Indeed, in all models with this structure, any failure to claim membership in a social category results in a rejection of it.

Thus, this model provides another way of seeing how classical logic reinforces problematic habits of othering. In this context, being a man requires an inflexible commitment to M combined with an inflexible rejection of all

“other” social categories. In order to be a man, agent a must be anything but feminine.

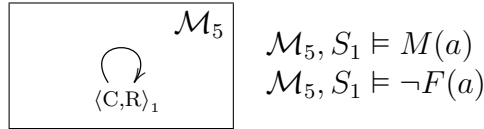
5 Modeling Systems of Identity-Thinking

Above, with the **Pressure to Commit** case, we saw that one common problematic feature of identity thinking is restriction of movement or denial of destabilization. Let’s explore how allowing for additional movement might help destabilize some of the attitudes of toxic masculinity in the **Anything But Feminine** case.³⁹

To do so, it will be helpful to add an additional convention to the diagrams the following models. Instead of using S_n to label each node in the diagram, I use a modified ordered pair notation with a subscript to indicate the specification number. The first item of the ordered pair will indicate the attitude that agent a has towards M at the specification and the second item of the ordered pair will indicate the attitude that agent a has towards F at the specification. The chart below indicates the possible options for indicating the agent’s attitude in the ordered pair.⁴⁰

C	a is COMMITTED TO the predicate
R	a REJECTS the predicate
P	a is PURELY UNCOMMITTED TO the predicate
C_f	a is FLEXIBLY COMMITTED TO the predicate
R_f	a FLEXIBLY REJECTS the predicate

Thus, for example, the attitudes described at S_1 in $\mathcal{M}_4$ above would yield the following modified model:



³⁹Note that while I explore options for allowing additional movement in terms of the **Anything But Feminine** case, the same strategy can be applied to the more generic **Forced Association** case from the beginning of Section 4.

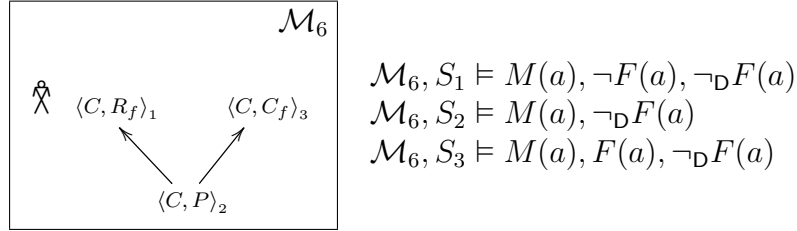
⁴⁰Notice that I do not list UNCOMMITTED TO the predicate in this list. This is to reduce redundancy: any time an agent either REJECTS or is FLEXIBLY COMMITTED TO the predicate, they will also be UNCOMMITTED TO it.

This model, $\mathcal{M}_5$, indicates a case where agent a is committed to M and rejects F at S_1 because C is in the first position of the ordered pair, which indicates the attitude that agent a has towards M at the specification, and R is in the second position of the ordered pair, which indicates the attitude that agent a has towards F at the specification. The subscript indicates that this pair denotes the specification S_1 .

5.1 Ways to be M : Destabilizing Toxic Masculinity

As I argued above, trapping agent a into an inflexible rejection of F merely because they claim M is unreasonably restrictive. In this section, I expand the previous model $\mathcal{M}_5$ to allow for other ways of being or becoming a man which do not require an immediate, inflexible rejection of F .

Using the new labeling conventions, let's now expand the model to include an additional destabilizing option, S_2 , and a new stabilizing option from that point, S_3 . Situate the agent at S_1 .⁴¹



With agent a situated at S_1 in $\mathcal{M}_6$, their situation is quite similar to what it was in at S_1 in $\mathcal{M}_5$ with one key difference: agent a has now destabilized their rejection of F . To see how, observe that the new model now makes the following true: $\mathcal{M}_6, S_1 \models \neg F(a) \wedge \neg_{\mathbf{D}} \neg F(a)$.

At S_1 , agent a still rejects F , since there is no forward path to a specification where $F(a)$, but they do so flexibly, since there is a destabilization option (S_2) from which a could stabilize such that they commit to F (at S_3).

⁴¹The model $\mathcal{M}_6$ can be described as follows:

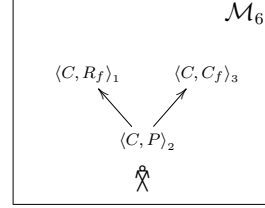
$$\mathcal{M}_6 = \left\langle \Sigma = \{S_1, S_2, S_3\}, \mathcal{R} = \left\{ \begin{array}{l} \langle S_n, S_n \rangle \text{ for all } S_n \in \Sigma, \\ \langle S_2, S_1 \rangle, \langle S_2, S_3 \rangle \end{array} \right\}, \Delta = \{a\}, \right.$$

$$\left. \sigma = \{a\} \text{ for the following pairs: } \langle M, S_1 \rangle, \langle M, S_2 \rangle, \langle M, S_3 \rangle, \langle F, S_3 \rangle \right\rangle$$

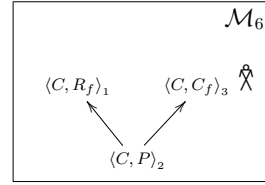
$$\text{and } \sigma = \{ \} \text{ otherwise}$$

Note that recognizing this destabilization option is enough to create flexibility regarding F . Thus, S_1 in $\mathcal{M}_6$ already represents a significantly different way of self-identifying as or becoming a man than S_1 in $\mathcal{M}_5$ represented.

However, we might also imagine agent a retracting their rejection of F by moving to S_2 , where they are still committed to M but are purely uncommitted to F . In order to have such an attitude towards F , agent a must take seriously the options to either reject F by moving forwards to S_1 or commit to F by moving forwards to S_3 .



We might further imagine that agent a takes this stabilization option by moving to S_3 . If they do so, then they would be flexibly committed to F . Notice that while it took a great deal of effort for agent a to destabilize their rejection of F and create a new model with a destabilization option that allows for a more flexible rejection of F at S_1 , this work is not needed to destabilize agent a 's commitment to F at S_3 because a suitable destabilization option already exists.



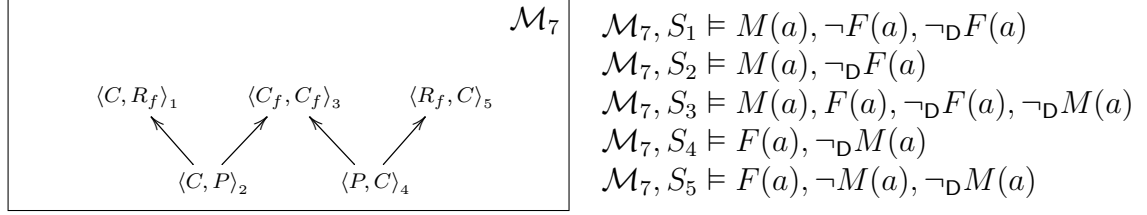
In this model, S_3 provides a counterexample to exclusivity because $\mathcal{M}_6, S_3 \not\models \neg(M(a) \wedge F(a))$. But the categories M and F are still exhaustive (the claim $M(a) \vee F(a)$ is true at all specifications). And since all specifications in $\mathcal{M}_6$ are such that agent a is committed to M , this model does not allow for any flexibility regarding M . This is what I go on to add in the following section.

5.2 Maximizing Options for Identity-Thinking

The previous model, $\mathcal{M}_6$, maintained agent a 's commitment to M at every specification. This provided a good way of seeing the various options that agent a has with respect to self-identifying with category M . In order to model more options for agent a 's identity-thinking, expand the model to include an additional destabilizing option, S_4 , and a new stabilizing option from that point, S_5 .⁴²

⁴²The model $\mathcal{M}_7$ can be described as follows:

$$\mathcal{M}_7 = \left\langle \Sigma = \left\{ \begin{matrix} S_1, S_2, S_3 \\ S_4, S_5 \end{matrix} \right\}, \mathcal{R} = \left\{ \begin{matrix} \langle S_n, S_n \rangle \text{ for all } S_n \in \Sigma, \\ \langle S_2, S_1 \rangle, \langle S_2, S_3 \rangle, \\ \langle S_4, S_3 \rangle, \langle S_4, S_5 \rangle \end{matrix} \right\}, \Delta = \{a\} \right\rangle$$



Imagine agent a at S_3 where they claim both M and F . From their vantage point at S_3 , agent a has two similar paths available: they could retract their recent commitment of F and return to S_2 or they could retract their commitment to M by moving to S_4 . Let's focus on the latter option for the moment.

If agent a destabilizes their commitment to M , this takes them to a specification where they no longer self-identify with social category M . At S_4 , agent a remains committed to F , but is purely uncommitted to M .

Notice that S_4 is the mirror image of S_2 . Just as at S_2 , a has now committed one way or another regarding only one of the two social categories they're considering. This time, instead of being committed to M and being purely uncommitted to F , a has committed to F and is purely uncommitted to M .

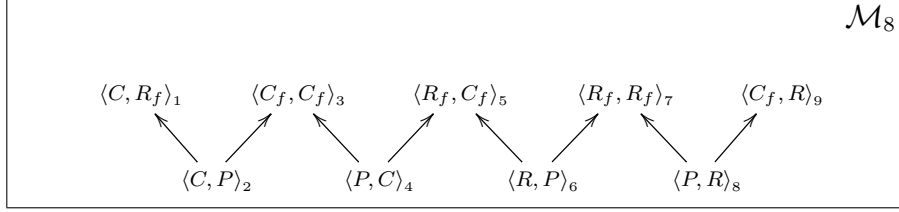
From S_4 , agent a has yet another option for stabilizing: reject identity M . This would take agent a to S_5 , which is the mirror image of S_1 . At S_5 , agent a flexibly rejects M and is committed to F .

While agent a is flexible towards both M and F at S_3 because they have the option to retract their commitment to either of them by moving to S_4 or S_2 , this flexibility is not present at both S_5 and S_1 . At these specifications, agent a commits to a social category inflexibly (M at S_1 and F at S_5). To allow for more flexibility, we can continue to build out the model.

For instance, we could expand the model to include two additional destabilizing and re-stabilizing steps, as we have been doing throughout this section. This adds an additional destabilizing option, S_6 , and a new stabilizing option from that point, S_7 . Additionally, this adds an additional destabilizing option from S_7 , denoted as S_8 , and a new stabilizing option from that

$$\sigma = \{a\} \text{ for the following pairs: } \langle M, S_1 \rangle, \langle M, S_2 \rangle, \langle M, S_3 \rangle, \langle F, S_3 \rangle, \langle F, S_4 \rangle, \langle F, S_5 \rangle \text{ and } \sigma = \{ \} \text{ otherwise}$$

point, S_9 .⁴³



In $\mathcal{M}_8$, the following claims hold at each specification:

$$\begin{array}{ll}
\mathcal{M}_8, S_1 \models M(a), \neg F(a), \neg_{\mathbf{D}} F(a) & \mathcal{M}_8, S_2 \models M(a), \neg_{\mathbf{D}} F(a) \\
\mathcal{M}_8, S_3 \models M(a), F(a), \neg_{\mathbf{D}} F(a), \neg_{\mathbf{D}} M(a) & \mathcal{M}_8, S_4 \models F(a), \neg_{\mathbf{D}} M(a) \\
\mathcal{M}_8, S_5 \models \neg M(a), F(a), \neg_{\mathbf{D}} F(a), \neg_{\mathbf{D}} M(a) & \mathcal{M}_8, S_6 \models \neg M(a), \neg_{\mathbf{D}} M(a), \neg_{\mathbf{D}} F(a) \\
\mathcal{M}_8, S_7 \models \neg M(a), \neg F(a), \neg_{\mathbf{D}} F(a), \neg_{\mathbf{D}} M(a) & \mathcal{M}_8, S_8 \models \neg F(a), \neg_{\mathbf{D}} M(a), \neg_{\mathbf{D}} F(a) \\
\mathcal{M}_8, S_9 \models M(a), \neg F(a), \neg_{\mathbf{D}} F(a), \neg_{\mathbf{D}} M(a) &
\end{array}$$

Using the list above, we can check the status of complex claims regarding M and F . For example, by the time that agent a reaches S_7 , they (flexibly) reject both M and F . This is a counterexample to exhaustiveness because $\mathcal{M}_8, S_7 \not\models M(a) \vee F(a)$.

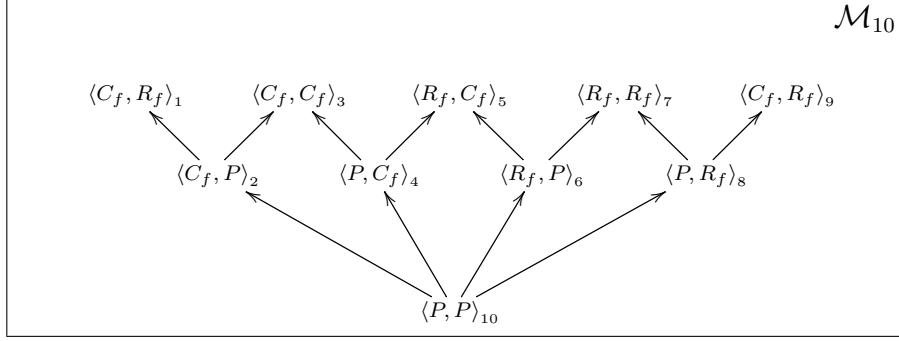
At this point, there still remain two specifications in the top row which have an inflexible attitude towards a predicate: S_1 , where agent a is inflexibly committed to M , and S_9 , where agent a inflexibly rejects F .

It is also the case that each commitment or rejection in the second row is inflexible. While these specifications served as destabilization options for specifications in the top row, these specifications themselves do not have access to a destabilization option.

⁴³The model $\mathcal{M}_8$ can be described as follows:

$$\mathcal{M}_8 = \left\langle \Sigma = \left\{ \begin{array}{l} S_1, S_2, S_3, S_4, S_5 \\ S_6, S_7, S_8, S_9 \end{array} \right\}, \mathcal{R} = \left\{ \begin{array}{l} \langle S_n, S_n \rangle \text{ for all } S_n \in \Sigma, \langle S_2, S_1 \rangle, \\ \langle S_2, S_3 \rangle, \langle S_4, S_3 \rangle, \langle S_4, S_5 \rangle, \\ \langle S_6, S_5 \rangle, \langle S_6, S_7 \rangle, \langle S_8, S_7 \rangle, \langle S_8, S_9 \rangle \end{array} \right\}, \Delta = \{a\}, \right. \\
\left. \sigma = \{a\} \text{ for the following pairs: } \langle M, S_1 \rangle, \langle M, S_2 \rangle, \langle M, S_3 \rangle, \right. \\
\left. \langle F, S_3 \rangle, \langle F, S_4 \rangle, \langle F, S_5 \rangle, \langle M, S_9 \rangle \text{ and } \sigma = \{ \} \text{ otherwise } \right\rangle$$

We can address this by expanding the model to include a radical destabilization option, S_{10} , where agent a is purely uncommitted towards both M and F .⁴⁴



The inclusion of S_{10} results in all attitudes towards social categories becoming flexible. Since there is a path $\mathcal{R}(S_{10}, S_n)$ for all $S_n \in \Sigma$ and at S_{10} agent a is purely uncommitted to all available predicates, this means that from any point it is possible to destabilize any commitment or rejection towards a predicate.

This model shows all of the paths that could be created by destabilizing, and then subsequently stabilizing differently, agent a 's attitudes with respect to the social categories M and F . At the bottom of the model with S_{10} , agent a makes no commitments with respect to M and F . In the second row of the model, with specifications S_2 , S_4 , S_6 , and S_8 , agent a has taken a stance (either a commitment or a rejection) with respect to exactly one of M and F . At each specification in the top row, agent a has taken a stance (either a commitment or a rejection) with respect to both of M and F .

Displayed this way, any stabilizing move brings us further up the diagram whereas any destabilizing move brings us further down the diagram. Importantly, agent a need not start at the bottom of the diagram where they are

⁴⁴The model $\mathcal{M}_{10}$ can be described as follows:

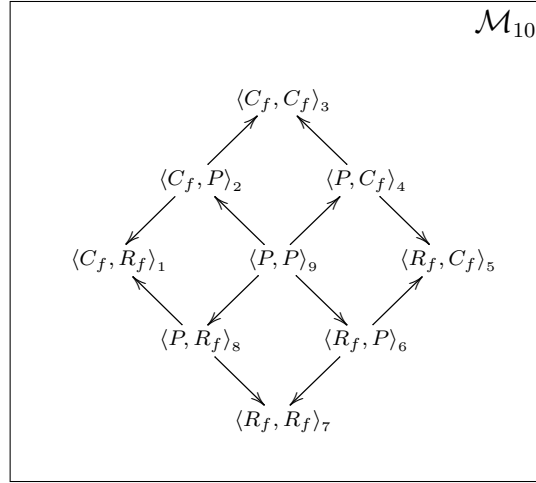
$$\mathcal{M}_{10} = \left\langle \Sigma = \left\{ S_1, S_2, S_3, S_4, S_5 \right\}, \mathcal{R} = \left\{ \begin{array}{l} \langle S_n, S_n \rangle \text{ for all } S_n \in \Sigma, \langle S_2, S_1 \rangle, \\ \langle S_2, S_3 \rangle, \langle S_4, S_3 \rangle, \langle S_4, S_5 \rangle, \\ \langle S_6, S_5 \rangle, \langle S_6, S_7 \rangle, \langle S_8, S_7 \rangle, \langle S_8, S_9 \rangle, \\ \langle S_{10}, S_n \rangle \text{ for all } S_n \in \Sigma \end{array} \right\}, \Delta = \{a\}, \right.$$

$$\left. \begin{array}{l} \sigma = \{a\} \text{ for the following pairs: } \langle M, S_1 \rangle, \langle M, S_2 \rangle, \langle M, S_3 \rangle, \\ \langle F, S_3 \rangle, \langle F, S_4 \rangle, \langle F, S_5 \rangle, \langle M, S_9 \rangle \text{ and } \sigma = \{ \} \text{ otherwise} \end{array} \right\rangle$$

purely uncommitted to both M and F . Few individuals start exploring their relationships to available social categories from a radically neutral position like S_{10} .

Furthermore, agent a is not “stuck” at any uppermost point in the diagram either. While each of these specifications represent claims with respect to both M and F , agent a always has the option to destabilize their commitment or rejection to any social category. This, too, makes sense: identities are not static. Individuals can and do alter their attitudes towards different social categories throughout their lives.

Another perspective on this model can be obtained by wrapping the specifications around, so that S_1 and S_9 overlap, as follows:⁴⁵



Viewed this way, stabilizing actions move outwards and destabilizing actions move inwards. Each edge of the diamond is a place where agent a commits to or rejects one of the available social categories. Hence, on the top left edge, agent a commits to M whereas on the top right edge agent a

⁴⁵The model $\mathcal{M}_{10}$ can be described as follows:

$$\mathcal{M}_{10} = \left\langle \Sigma = \left\{ \begin{matrix} S_1, S_2, S_3, S_4, S_5 \\ S_6, S_7, S_8, S_9 \end{matrix} \right\}, \mathcal{R} = \left\{ \begin{matrix} \langle S_n, S_n \rangle \text{ for all } S_n \in \Sigma, \langle S_2, S_1 \rangle, \\ \langle S_2, S_3 \rangle, \langle S_4, S_3 \rangle, \langle S_4, S_5 \rangle, \\ \langle S_6, S_5 \rangle, \langle S_6, S_7 \rangle, \langle S_8, S_7 \rangle, \langle S_8, S_1 \rangle, \\ \langle S_9, S_n \rangle \text{ for all } S_n \in \Sigma \end{matrix} \right\}, \Delta = \{a\}, \right.$$

$$\left. \sigma = \{a\} \text{ for the following pairs: } \langle M, S_1 \rangle, \langle M, S_2 \rangle, \langle M, S_3 \rangle, \right.$$

$$\left. \langle F, S_3 \rangle, \langle F, S_4 \rangle, \langle F, S_5 \rangle \text{ and } \sigma = \{ \} \text{ otherwise} \right\rangle$$

commits to F . Each bottom edge is a place where agent a rejects the social category that they claimed on the opposite side of the diamond.

As any agent moves through this space of specifications, committing to and rejecting different social identities, their movement traces through the various possibilities that are open to them. As shown above with agent a , an agent's destabilizing and stabilizing actions might serve to open up options that were not previously in the model, but ultimately those options will be constrained by any simplifying assumptions.

For example, if we assume exhaustiveness, specifications S_9 , S_8 , S_6 , and S_7 would drop out of the diamond model above because for each of them, $\mathcal{M}_{10}, S_n \not\models M(a) \vee F(a)$. Adding in an assumption of exclusivity would remove S_3 because $\mathcal{M}_{10}, S_3 \not\models \neg(M(a) \wedge F(a))$. Without the alternate stabilization path that S_3 provides for both S_2 and S_4 , it would not be possible for agent a to remain purely uncommitted towards either social category at these specifications. As a result, the constraints of exhaustiveness and exclusivity taken together yield a disconnected model with one specification (S_1) where agent a is inflexibly committed to M and inflexibly rejects F and another specification (S_5) where agent a is inflexibly committed to F and inflexibly rejects M .

Thus, this model $\mathcal{M}_{10}$ provides both a way to envision maximal options for exploring identification with two social categories in *Heyting-Brouwer* logic and a way to envision how the constraining assumptions can unduly restrict these options.

In this case, I chose to focus only on two identities and one individual, so we end up with a diamond pattern of specifications that trace out the total possibilities for articulating social identity in this space. The same pattern of maximal options would hold for any two social categories, though we may want to restrict this in specific applications. If we want to move away from a binary, it is also possible to focus on more than two social categories at a time, though this does increase the complexity of the model and the number of available options.⁴⁶

⁴⁶For example, the maximal options when focusing on three social categories creates a cube structure, where each face of the cube is a diamond that holds either commitment or rejection towards one of the available categories constant, in much the same way that each edge of the diamond holds commitment or rejection towards one of the two social categories constant.

6 Conclusion

As I have demonstrated in this paper, the use of logic to represent, discuss, or model socially engaged phenomena is indeed fraught, but it is not hopeless. Many of the models in this paper provide a way of visualizing the limitations that logic, and classical logic in particular, can face when attempting to represent an agent's interaction with social categories.

However, these limitations do not mean that we should abandon the task of modeling socially engaged phenomena. Rather, my analysis suggests that we should use caution when modeling socially engaged phenomena and actively seek clarity regarding the limitations of the models we build. Examining these limitations will often involve a pluralistic approach: through examining how another logic expands the space of available options, the limitations of our initial approach become more evident.

In just such a way, the models I have built using the *Heyting-Brouwer* logic developed by Rauszer help illuminate the limitations of classical logic when representing an agent's process of self-identification. Through destabilizing the assumptions fueling the binary axis, I argue, we can reclaim a greater degree of flexibility in our models of social categories.

My models are practical: they demonstrate the structure of various ways of thinking about social categories so that it is possible to identify and isolate the problematic assumptions that may be hampering our analysis. For those individuals who have accepted a limited model of social categories with very little flexibility, envisioning a more expansive model of social categories can open up options for self-identification that were previously inconceivable. And for those of us aiming to encourage others to tolerate, understand, and ultimately accept patterns of identification that lie outside of the available space granted by limited models, being able to demonstrate how a more expansive model unfolds will likely prove useful.

Through building models which allow for areas of overlap, situate individual categories equitably with respect to one another, meaningfully acknowledge the interrelations among social categories, and affirm positive or independent sources of definition, my analysis provides an initial step towards a radical transformation of our logical treatment of social identity. As Plumwood's criteria for a negation that can be used for feminist purposes indicate, this radical transformation is no easy task.⁴⁷ The following would

⁴⁷As shown here, Plumwood's criteria for a negation for feminist purposes provide a

be fruitful paths for further work on this topic:

- Test the effectiveness of this approach through applying it to more real-world cases.
- Utilize more combinations of the operators of *Heyting-Brouwer* logic to include more complex claims as part of the analysis.⁴⁸
- Build models with more than one agent or consider how the models built by different agents might interact.
- Analyze the impact of disconnected models, where a path from one self-identification to another is unavailable.
- Discuss how an agent’s position in the model reflects their positionality. This could be helpful for an evaluation of *Heyting-Brouwer* logic’s ability to represent intersectionality.

As more work emerges in this area, it is my hope that we will collectively improve our ability to see, acknowledge, and discuss the suitability of individual approaches to logically representing socially engaged phenomena. While some attempts will certainly have flaws – the limitations fueling the available options for self-identification in cases like **Pressure to Commit** and **Anything But Feminine** are clear evidence of this – even incorrect models can help us to see where additional flexibility or an alternative presentation might be needed.

Rather than abandoning logical tools in our efforts to avoid problematic habits of othering, we ought to embrace the unique ability of these tools to showcase how individual assumptions contribute to these problematic habits. Only then will we be able to radically transform our treatment of social identity to match the needs of the present moment.⁴⁹

helpful guide for projects seeking to address the limitations of classical logic (Plumwood 2002).

⁴⁸For example, see (Cook R., R. Kosten, R. A. Rakacolli, and I. Valasquez Manuscript) for an exploration of how strings of both negations would work in this system.

⁴⁹Acknowledgments

References

- Cook R., R. Kosten, R. A. Rakacolli, and I. Valasquez. (Manuscript), “Heyting-Brouwer Logic, Duality, and Counting No-Dalities”
- Dembroff, R. (2020), “Beyond Binary: Genderqueer as a Critical Gender Kind.” *Philosopher’s Imprint* 20 (9): 1 – 23.
- Doyle, G. (2020), *Untamed*. The Dial Press.
- Eckert, M. (Forthcoming), “Re-Centering and Genderqueering Val Plumwood’s Feminist Logic.” *Feminist Philosophy and Formal Logic*. Ed. A. Yap and R. Cook.
- Ferguson, T. (2014), “Extensions of Priest-da Costa Logic”, *Studia Logica* 102: 145 – 174.
- Plumwood, V. (1993), “The Politics of Reason: Toward a Feminist Logic”. *Australasian Journal of Philosophy* 71 (4): 436 – 462.
- Plumwood, V. (2002), “Feminism and the Logic of Alterity”. *Representing Reason: Feminist Theory and Formal Logic* Ed. R.J. Falmagne and M. Hass, Rowman and Littlefield: 45 – 70.
- Nye, A. (1990), *Words of Power: A Feminist Reading of the History of Logic*. Thinking Gender. New York: Routledge.
- Priest, G. (2008), *An Introduction to Non-Classical Logic, Second Edition: From If to Is*. Cambridge University Press.
- Priest, G. (2009), “Dualising Intuitionistic Negation”, *Principia* 13(2): 165 – 184.
- Priest, G. (2011), “First-order da Costa Logic”, *Studia Logica* 97(1): 183 – 198.
- Rauszer, C. (1974), “A Formalization of the Propositional Calculus of H-B Logic”, *Studia Logica* 33(1): 23 – 34
- Rauszer, C. (1974), “Semi-Boolean Algebras and their Application to Intuitionistic Logic with Dual Operators”, *Fundamentae Mathematicae* 83(1): 219 – 249.

- Rauszer, C. (1977), “Applications of Kripke Models to Heyting-Brouwer”,
Studia Logica 36(1/2): 61 – 71.
- Rauszer, C. (1977), “Model Theory for an Extension of Intuitionistic Logic”,
Studia Logica 36(1/2): 73 – 87.