

The Modal Vocabulary: What Possible, Plausible, Probable, and Desirable Should Mean?

Veli Virmajoki

Introduction

Assume, for the sake of argument, that a scenario team gathers around a whiteboard. Someone has sketched four futures for the energy system in 2050, and the facilitator asks the group to sort them: which of these are *possible*? Which *plausible*? Which *probable*? And which *desirable*?¹ Imagine the room's opinion are split. One participant says that she thinks that total grid decarbonisation by 2050 is plausible, since no known physical law rules it out. Another objects, correctly it seems, that plausibility asks for more than the absence of impossibility – it requires we can identify a causal pathway, mechanism, that could actually get us there. A third participant cynically adds that the scenario is not even *possible* in any sense that matters, given current political realities; a fourth answers that political realities are precisely the kind of thing scenarios are met to challenge, so referring to political reality is not an argument against the scenario as such – desirability affects what can be made possible. Twenty minutes on, the team has learned a great deal about one another's assumptions. However, they have learned nothing at all about what the words they use – *plausible*, *possible*, *desirable* – mean.

This is not a failure of facilitation, people learn. It still is a failure of vocabulary, and that does matter. Let us explain in this text.

The modal notions central to futures studies practices and communications – *possible*, *plausible*, *probable*, *preferable* – are at the same time (i) the field's most indispensable concepts and (ii) its most dangerously imprecise. These notions are expressed in nearly every canonical framework. For example, Voros's (2003) *futures cone* (which already exist at least in Hancock and Bezold's 1994 work), sorts futures into nested zones – and these are zones of possibility, plausibility, probability, and preferability – maybe we can add desirability, and layer of preposterous. Amara's (1981) three commitments – that the future is not predetermined, the future is not predictable, and future outcomes

¹We use the term *desirable* whenever possible, i.e., when we are not citing author whose use of *preferable* is not at the very core of their writing.

can be influenced by present choices – presuppose that one can draw meaningful distinctions between what is predetermined, what is predictable, and what is open to influence. Finally, Dator’s (2009) four archetype futures – *Continued Growth, Collapse, Discipline, Transformation* – sorts, at least implicitly, futures in terms of their relation to present trajectories and ask us to think type of *alternatives* – modally loaded term.

And in none of these frameworks are the terms given definitions sharp enough to settle the disagreement that opens this text. The words are living the life of their own in actual work in futures studies done every day but rarely anyone has said what the words mean with philosophical clarity.

Behind these terms lie metaphysical questions about the reality of the future, the openness of what is to come, and the extent of the space of possibility. These are questions about the furniture of reality – the metaphysical architecture that makes talk of alternative futures (a modal notion) coherent at all – and they are bracketed in what follows. This text asks a downstream question: given some space of possibilities, how should the vocabulary for moving in that space work? What does it mean, exactly, to call a future *possible* rather than *plausible*, or *plausible* rather than *probable*? What kind of reasoning is associated with the modal notions in the field?

Due to the various and even contradictory use of the modal notions in the field, we do not claim to track their usage to some ultimate point of departure or their “real meaning”. The tool we use in this text is what in philosophy is called *conceptual engineering* – the deliberate “sharpening” of concepts so that they are fit better to certain practices, such as thinking about futures. The goal is not legislation of meanings from outside. Rather, by connecting the use of a concept to its meaning through “engineering”, and asking what each term must mean if it is to do what futures studies does. We proceed term by term: discuss *possible*, then *plausible*, then *probable*, and touch on *desirable*. Along the way we run into a form of reasoning that underwrites much of the practice (or so we argue): counterfactual reasoning, the disciplined “what if things had been different?” thinking, which is can be associated with the engine of scenario construction.

Conceptual Engineering and the Futures Vocabulary

Before sharpening any single term, we must steer clear about the kind of work is to follow. Conceptual engineering, as it is practised in recent analytic philosophy, begins the observation concepts we “inherit” are not always fit for purposes we use them. They are *vague* where we need precision, *ambiguous* where we need clarity, or *loaded with connotations* that quietly affects research and its

use (see e.g., Cappelen, 2018; Haslanger, 2012). The term *conceptual engineering* itself was, as is fitting to its own meaning, coined more than once – independently within “Carnap scholarship” (Carus, 2007; Creath, 1991) and within “metaphilosophy more broadly” (Blackburn, 1999; Floridi, 2011). The two traditions have since converged, at least somewhat, on a shared programme of kind of concept-improvement (Brun, 2016; Isaac, Koch, & Nefdt, 2022). Conceptual engineering does not to describe how a concept happens to be used – that is conceptual analysis, the older and more “descriptive” tradition – but to ask how it *should* be used, given what we need it for (Nado, 2021; Chalmers, 2025). Philosophy here becomes a “translational”, problem-solving thinking aimed at workable solutions to conceptual problems (Nado, 2021; Blackburn, 1999). The question shifts: not “what do people mean by X?” but “what should X mean, if it is to do perform certain functions?” – a move that treats concepts not as fixed objects of contemplation but as *devices* whose functional efficacy can, and should, be improved (Nado, 2021; Isaac et al., 2022; Woodward 2003).

The process has a roughly the following nature. At a minimum it runs through four phases: *describing* the target concept, *assessing* how it currently performs its function, *improving* where it fails or causes confusion, and *implementing* the revised version in the community (like futures studies scholars and foresight practitioners) that uses it (Isaac et al., 2022; Cappelen & Plunkett, 2020). Drawing on Carnap’s (1950) “method of explication”, we may treat these phases not as forming a tidy sequence but, rather, we move back and forth between the phases until we achieve what we want (e.g., Brun, 2016; Chalmers, 2025; Isaac, Koch, and Nefdt, 2022; Thomasson, 2025).

One lesson for our purposes can be pinpointed for our purposes: the goals of futures studies as field and practice are plural. In conceptual engineering, one may wish to change how people classify things, or what inferences they draw, or the truth-conditions of key sentences, or the norms of usage – and each of these aims, it seems, requires a different method (Isaac et al., 2022; Koch, 2021; Nado, 2023). The framework of conceptual engineering is, however, “meta-normatively neutral”: it works with whatever substantive purpose one may have, but it insists the purpose has to be made explicit during the engineering process so it works (see e.g., Dutilh Novaes, 2020; Haslanger, 2005).

Futures studies is an exceptionally good candidate for the work of conceptual engineering. The field draws its modal vocabulary from several traditions, sources, and intuitions. Often uses of “possible” and “probable”, for example, differ from people to people, from project to project, from one futures studies tradition to another. Consider this: when a *logician* says something is “possible”, she means, more or less, that it is not self-contradictory. When a *strategist* in a company says a scenario is

“possible”, she means something far richer – something along the lines that the scenario consistent with what is known about the world and is reachable by at least a rough causal (that can be at least sketched) pathway from here. When an *ordinary person* says “it is possible that tomorrow rains”, she may be expressing uncertainty (given weather forecasts), not making a claim about logical consistency or causal routes. Three concepts are thus travelling under one word – and we are in the danger of moving between them, more or less, without warning. This is the kind of situation conceptual engineering was tailored to diagnose and repair (Cappelen, 2018; Isaac et al., 2022).

The stakes are not merely academic, that we do consistent research reporting. When a foresight deliverable calls a future *plausible* to an audience – a board of directors, a government ministry, a community planning body – the stakeholders deserve to know what kind of meaning the word has in that case and what is the evidence and reasoning that are taken to warrant *plausibility*-claim. For example, if *plausible* would only mean “not obviously self-contradictory”, the warrant is thin. If it means “an identifiable chain of causes could lead from here to there”, the warrant is considerably heavier (assuming the chain is well-reasoned). If it means “some expert assigns it non-trivial probability”, it is heavier again but different in kind (no causal reasoning behind the claim but trust to experts). The current vocabulary is such that all of these uses are allowed. And here is where conceptual engineering is needed: such ambiguities are not innocent (Cappelen, 2018; Queloz & Bieber, 2022). When one word is chewed with three different epistemic standards in three mouths, the risk is not merely that people talk past one another. It is that decisions based on “plausibility” get murky and contradictory even.

There is an interesting discussion about the history of the most loaded word in the lot. Ramírez and Selin (2014), tracing the etymology of “plausibility” and “probability” in English, show that the two ran together until the seventeenth century and pulled apart “only” over the following three hundred years. “Probable”, once tied increasingly to mathematical and statistical reasoning and representations, gained a kind of scientific nobility. “Plausible” kept, from the start, an association with mere appearance. This can be understood once we notice that it is a cousin to *applause*, act which indicates the winning of public approval. In its earliest English uses it carried the hint of a “false appearance of veracity”, and – no matter how ironical it may sound in our current world – this a lineage that ends, more or less, in the Cold-War coinage *plausible deniability*: a polite name for institutional lying.

This means that a term the field of futures studies has at its core, *plausibility*, and associates with certain presentations of future, is loaded with certain associations – looking good rather than being true, winning people over, sounding convincing while being empty. The word, one might say, has a odd past when mirrored to its current use in futures studies. And that is precisely the type of a mismatch between inherited meaning and required function that conceptual engineering was built to repair.

In what follows, we ask: (i) what work does this or that concept do, or what function it has, in futures studies, and (ii) how could we provide a meaning to the concept that allows to do that work or satisfy the function? How we could and should use a concept will not likely correspond its current meaning-bundle, that is the point, but rather we want to improve the rigour, the clarity, and the decision-relevance of the concept. We begin discussion concerning the concepts by very rough mapping of how the term is used across the field's representatives (description of usage comes first, as Isaac et al. 2022 argue) before saying how it should be perhaps be refined. And we keep one ethical eye open throughout: conceptual engineering is tied into conceptual ethics (Burgess & Plunkett, 2020). Conceptual engineering, as M. Shields (2021) and Cappelen (2018) warn, can become akin to *conceptual domination* if its ends are hidden and its proposals are shielded from criticism from the very people they are meant to serve. This means that the proposals we provide are offered as an ongoing, participatory conversation, not handed down as orders.

“Possible”: More Than Merely Conceivable

Of all the modal terms, *possible* may appear as the least controversial. It is not. The metaphysics of possibility includes Lewis's (1986) concrete modal realism, Plantinga's (1974) abstractionism, and Armstrong's (1997) combinatorialism. Lewis treats possible worlds as concrete realities, Plantinga as maximal possible states of affairs, and Armstrong as recombinations of actual objects and properties. Each provides a different account of what possible worlds are and how vast is the space of possibility. These accounts are ontological: about what exists. The question now is conceptual: about what practitioners mean, and should (could) mean, when they stamp a future “possible”.

In philosophical usage, several well-separated varieties of possibility have been distinguished, and the differences matter enormously for futures studies – this is perhaps the most crucial place where philosophy should be studied by those who work on futures. Modal epistemology fixes mostly on the so-called *alethic* or “objective” modalities – possibility is grounded on of *how the world is*, not of

how it is believed to be. Among the principal three alethic possibilities are *logical*, *physical*, and *metaphysical* possibility. (Roughly: logical = not self-contradictory conceptually; metaphysical = how reality could actually have been deep down; physical = allowed by the laws of nature).

The “textbook picture” is to nest these: physical possibility inside metaphysical, metaphysical inside logical (Kripke, 1980; Hale, 2013). The boundaries, it must be noted, are not settled. According to Chalmers (2002) metaphysical and logical (conceptual) modality coincide; according to Shoemaker (1998) metaphysical modality deflates so that it is co-extensive with physical modality; and Priest (2021) questions whether there is a notion of metaphysical necessity that is distinct from conceptual necessity or physical necessity. While the picture can thus be contested, we can use it anyway.

For futures studies the practical question is a blunt one: which variety of possibility draws the outer boundary of the space of possibilities we must deal with?

If the answer is logical possibility, the space is vast, too vast, – anything that does not contradict itself is on the table. Escaping to metaphysical possibilities does not help much. Even those futures that break every known law of physics are included. A world in which people can fly are metaphysically possible. Obviously, these are mostly too wide to be much use. However, issue is not that simple. These logical possibilities that are far-fetched may earn their place in the “most speculative” exercises, for example, at the mythic floor of Inayatullah’s (1998) Causal Layered Analysis; where myth and metaphor probe the very edge of what we can think at all. To call these “speculative” is not criticism, rather, we wish to point out that sometimes even logical possibilities form the vast space we must take into account. It is not easy, but nothing valuable is easy, some would say.

Then again, if the answer is physical possibility, the space is narrower but still enormously vast: the laws of physics license a staggering range of social, political, technological, and ecological arrangements (for example, people have no names but are all identified by number, to give a toy-example). The answer must, then, be some sort of a “practical possibility” – consistent not only with natural law but with the world as it actually now stands and work. The space narrows again, to the futures one could actually reach from where we stand. It is this last boundary that everyday futures studies more or less lives inside of. This is, then, what futures studies scholars and people thinking about mostly mean when they talk about *possible*.

Practical possibility is a peculiar thing, and a few philosophers approach it almost head-on. The argument is that we grasp what is possible not through elaborate reasoning but more or less directly.

One strand of this “head-on” approach is *perceptual*. Strohminger (2015) argues that perception itself can provide knowledge of possibilities; seeing a low branch we in some sense *see* that it could be climbed. A second strand is *agentive* rather than perceptual. Vetter (2023) argues that our sense of what is possible is rooted in the experience of our own action: we know what we can do, and the wider idea of what is possible grows out from this experience.

Of course, philosophers try to make sense of possibilities, like practical possibilities going to their roots – nobody would say (credibly) that they can see the future or feel how they change it; that would be plain silly. However, the arguments may be more relevant to foresight than they sound. When one tracks weak-signal or does horizon-scanning, one is trying to peak or have some feel what might be going on now and, thereby, in the future. When Ansoff (1975) brought weak signals into strategic management, he was pointing at something we might call “a perception of practical possibility”: an early indication that a development is possible, before anyone can say whether it is plausible or probable – these are separate issues. Hiltunen (2008), referring to Peirce, can be seen as sharpening the point: a *future sign* (distinguished from *weak signal* due to the inflation of meanings of the latter) consist of three parts – a signal, an issue, an interpretation. This would mean that the “perception” of practical possibility is interpretation-laden from the start. This, in itself is not a problem: even scientific observations are theory-laden. So there are ideas where practical possibilities are “seen” and form a seed for understanding futures.

A more generally used route to knowledge about possibility runs through the epistemology of conceivability. In brief, Yablo (1993) took what we can coherently imagine to count in favour of – though never to settle – what is possible: imaginability is evidence, but evidence that can be overturned. This is, near enough, what a brainstorming session does; it uses imagination as an access way to understand what is possible. But the route is not safe, to be sure, and we know that. The move from “I can imagine this” to “this is possible” is problematic, at least when we talk about practical possibilities. But the distinction has even philosophical problems. Vaidya and Wallner (2021) discuss “modal epistemic friction”. There must be constraints that keep the leap from “imagine” to possible” away from being arbitrary. Something can seem conceivable and turn out, on a closer look, to be impossible. One can imagine me having different parents but I could not – different genes, different person. And something can fail to seem conceivable to a particular group and be perfectly possible all the same; this would be a failure of imagination, not a fact about the world – we cannot conceive a science that is not there – describing it would be, paradoxically, to have it – but it is possible that

science changes. Van Inwagen (1998), in short, compares modal judgment to seeing: at the far edge of the visual field we withhold verdicts, and the same withholding is due once modal claims stretch past what our minds handle with any reliability. In terms of future: The more exotic the future, the less our being able to picture it tells us. We can picture, imagine, freely. The world is under no obligation to comply or follow.

What this means for practice is clear enough. Calling a future “possible” is not just something one can say when nothing else can be argued for. It is an *epistemic claim* – that, at minimum, so far as we know, nothing rules this future out – and its epistemic weight (how seriously we take it in decision-making) depends entirely on the variety of possibility in play – logical, metaphysical, physical, “practical”. A future that is “logically possible” carries less decision-relevant weight than one a future that is “physically possible”, which carries less than a future that is “practically achievable from where we stand” (how we ever can know this; in this text we talk about vocabulary). One must be case-by-case engineer here: when one speaks about “possible”, one has to tell what one *means* by this and be consistent. If one is shown one’s vision for “practically possible” future is not really practically possible, one cannot back down and say that the future, at least, non-contradictory, one commits intellectual crime. The future may be lifted to a table in another context where brainstorming is more welcomed, but that does not mean mistake is erased. *Possible* is tricky, but it can somewhat be tamed through consistency.

“Plausible”: The Case for Causal Pathways

If *possible* is, when taken in the widest sense, the outer boundary of what cannot be ruled out, *plausible* is meant to form a more interesting zone of futures: the futures that earn serious analytical attention. In Voros’s (2003) *futures cone* the plausible sits inside the possible and outside the probable. In practice, it is perhaps the most used but least consistently defined word in the field; in a sense, futures studies is the study of plausible (not just trend extrapolation, no mere speculation, something in-between). What the discipline actually demands when it talks about plausibilities is rarely obvious.

Consider two quite different things a practitioner might mean by *plausible*. According to the first reading, a scenario is plausible if nobody in the room can name a reason it could not happen. Call this *negative plausibility*: plausibility as the absence of identified constraints. According to the second reading, a scenario is plausible only if someone can lay out a causal pathway – a sequence of events,

mechanisms, and conditions – running from the present to that “plausible” future. Call this *positive plausibility*: plausibility as the presence of an identifiable causal route.

Needless to say, the gap between the two is enormous in practice. Negative plausibility is cheap: It depends on ignorance too much, as it refers to us not knowing something – and we know very little, really. The bar is low. The result, if we chose to use this definition of plausibility is that a great many futures count as *plausible* –all not-yet-refuted. For example, a future in which cold fusion becomes commercially viable by 2040 is negatively plausible in just this sense: nobody has proven it impossible, but nobody can say how it would come about – research is ongoing. Positive plausibility costs more, because it asks the scenario builder to construct, or at least in sketch, a causal path from the present to the described future. Fewer futures can be called plausible: no sketch of the path, no plausibility. However, the ones that are plausible under the definition of positive plausibility may well be worth more to a decision-making for the very reason that the causal pathway gives her something to build, monitor, test, and respond to.

The philosophical resources that are to be used here are, we suggest the theory of counterfactuals and the interventionist account of causation. Woodward’s (2003) interventionist account gives a natural way to think about a causal pathway. On Woodward’s account, to identify a cause is to identify a relationship that stays stable under at least some range of interventions. Changing taxation changes consumption, for example. The so-called *structural-equation tradition* made this concrete (and inspired the interventionist way of thinking about causality). In the causal models of Spirtes, Glymour, and Scheines (1993) and, above all for our purposes, Judea Pearl (2000, 2009), a system is a set of variables tied together by functional relationships, and to evaluate a counterfactual (i.e., *contrary-to-fact* claim where we say something about something that *did not happen*) – “what would happen if X took value x?” – one *intervenes*: Instead of waiting to see what value X happens to take, you reach in, set it to x yourself, disconnect it from whatever normally controls it, and follow the effects down the chain.

Pearl calls this surgery, and a simple example shows why it matters. Suppose I want to know what happens when I turn on a garden sprinkler. I do not worry about why the sprinkler is usually on (maybe a timer, maybe the gardener, whatever). I just switch it on, cut it off from those usual causes (for example, keep the gardener from touching it, no matter how much the person hates to see grass getting too wet), and watch what follows: the grass gets wet, (maybe too wet). This is the difference between *doing something* and merely *seeing it*. A wet roof and rain tend to come together, giving us

a hunch of causal relationship between the two, but turning on the sprinkler myself tells me what changed as a consequence.

To rely on these very specific (for historical reasons *causality* has been under extreme stress in philosophy) theories of causality is *not to say* that every plausible scenario come with a fully specified causal model. That would be unrealistic in futures studies. However, relying on these specific theories tell us that plausibility can and even *must* carry a *causal expectation*: that, in principle, a chain of causes joining present to future could be sketched. Otherwise, plausibility is cheap. When a scenario cannot even satisfy this demand for causal sketch that – when nobody can what is being called “plausible” is really “not yet shown to be impossible”. That is a much weaker claim as we saw, too weak and cheap. Honesty in foresight practice means that we strip the name “plausibility” off these claims. One can then go deeper and demand more: that the causal description is not mere sketch, that it is something more – and what this “more” means can be read off from the theories described above about causality. Details, a lot of them are to be found there; but that is not our main topic now.

Adopting *positive plausibility* as the standard has a cost, then, but it is something we may need to pay. Luckily, there is futures studies literature where work that getting *plausibility* right, can be seen in the form described above. *Cross-impact balance analysis* (see Weimer-Jehle, 2006) can be interpreted as operationalising something close to *positive plausibility*, when it looks for configurations of variables are internally consistent given a specified network of promoting and inhibiting influences, that is, interdependencies that can reasonably, although not uncontroversially, be read as causal ones.

However, positive plausibility is not just a nice term that happens to match some method in futures studies. While *general morphological analysis* (see Ritchey, 2018) does something similar as cross-impact balance analysis. However, its so-called cross-consistency assessment eliminates combinations that violate known logical, empirical, or normative constraints. In this way, it screens for what can coexist rather than focuses on causal structure as such. Ritchey (2018, 85) is, in fact, explicit that the assessment does not have to refer to causality (even if causal arguments happen to inform a particular consistency judgment). There is no causality here, at least automatically. Thus, there is no plausibility here, at least not automatically. That is: if we follow the definition given for plausibility in this text.

The two methods thus sit, in a sense, on different sides; but they are close at least when it comes to one crucial issue: What they seem to share is a refusal to let mere conceivability settle the question

of plausibility or even relevant futures. A concept of plausibility that requires causal pathways would therefore seem not to import a foreign standard but to follow the field's own reasoning. While plausibility is one thing – it is difficult to achieve – there are still methods that are at least close to finding plausible futures (even when they bracket causality) and these methods are far from mere speculation, so to say. There is a long continuum between the search for (positive) plausibility and just guessing, that is for sure.

Finally, we wish to pre-empt one worry and possible misreading: that requiring causal pathways biases foresight toward incremental, predictable futures and against the genuinely transformative scenarios we do not know yet how to reach. The worry is important but, in the end, misplaced. First, a causal pathway need not be conventional. It can run through mechanisms that are novel and even unprecedented. What matters is not that the mechanism be familiar but that it be *articulable*: that the builder of a scenario can say, even roughly, “here is how this could happen.” Second, we are now discussing *plausibility*. We bite the bullet in the sense that we admit that not all futures are plausible, and only incremental, predictable futures may belong to this group (which is not necessarily the case, see the first reply). To say something is *plausible* is not to say it is only thing we must take into account when we think about the future. We only wish to give *plausibility* some posture here – if everything is plausible, then nothing is (and what would then be transformative or unexpected, if plausibility does not serve as contrast?)

Finally, we may notice that Ramírez and Selin (2014) propose an alternative to *plausibility* itself as a measure of quality of a scenario: productive discomfort and the surfacing of ignorance. They argue that the most useful scenarios sit at an intermediate distance from the familiar – neither so obvious as to be dismissed nor so strange as to be rejected outright – and it is, the argument goes on, there that learning is greatest. This is a high-quality example of a genuine challenge to any framework, ours included, that *appears* to treat plausibility as a straightforwardly virtuous, at least implicitly. *However, we need to distinguish between plausible and positive.* One can accept the view about productive discomfort and the surfacing of ignorance as real criteria, but the view does not exclude (positive) plausibility. Plausibility is a requirement on a scenario's intelligibility – that one can say how it might be causally reached – and says relatively little about the strangeness of the path itself (see also above – paths can be novel). Moreover, a scenario can be causally describable and deeply discomfiting at the same time. One could even argue that disciplining oneself to construct a causal pathway to an unwelcome future is a “good” way to make discomfort productive rather than

paralysing. In other words: a scenario being plausible, positively plausible, does not mean the scenario has to be nice or make one satisfied. *Plausibility*, as a concept, guarantees nothing about positivity. In fact, most plausible futures, like climate changing faster, are least positive.

“Probable”: Interpretation and lessons from the Philosophy of Climate Science

If *plausible* has been obscure term in futures studies, *probable* causes even trouble, given the nature of what we have learned about futures thinking. While many like to talk about *probable futures*, the term has no natural home in this or that tradition. It is used by almost all, even if the notion of *probability* is far from clear in the phrase *probable future*.

When we look more closely, we see differences in explicit orientations towards *probability*. The term *probable* implies quantification – probabilities, likelihoods, numbers – and futures studies has always had an uneasy relationship with providing exact numbers, given the uncertainties and the basic fact that often the features of the future must be imagined before they can be count. However, some branches of futures studies are more intimately connected to numbers: the Delphi method (see origins in Dalkey & Helmer, 1963; Linstone & Turoff, 1975) often asks panellists at least for probability estimates; cross-impact analysis (Gordon & Hayward, 1968; Helmer, 1981) reasons over conditional probabilities, modeling how the occurrence of one event (re)shapes the likelihood of other events. Other traditions – such as *la prospective*, CLA, critical futures work in general – are not in the business of providing probability assignments; probability, in futures studies, implies spurious precision, deep uncertainty dressed up as rigour. Both sides have a point about *probability* as notion in futures studies.

The philosophy of probability can be used to explain both the use and mistrust. The first and most important lesson is that *probability* is not one concept. An anecdote tells how famous philosopher Russell called it the most important concept in modern science – and then added that nobody has the slightest idea what it means. He was not, if the anecdote is to be trusted, exaggerating. It is a *family of concepts*. In this family, each concept of *probability* is grounded in a different conceptual interpretation. When we go further, we notice that each interpretation is appropriate to a different context of application.

The problem is that futures studies rarely pauses to ask which one of the concepts or interpretations of probability is in play in this or that talk about *probable futures*.

Let's go to the interpretations, then. The *classical* interpretation, on the table since Laplace (1814/1999), defines a probability as favourable outcomes divided by all equally possible ones. This works for a dice and shuffled cards, but only because their physical *symmetry* grounds the fact that the outcomes really are equally likely. Without that type of symmetry, there is no reason to call outcomes equal. Consider the problem through a famous example (see van Fraassen, 1989; after Bertrand, 1889): A cube factory makes cubes up to 2 cm a side, and you ask whether a random dice has a side under 1 cm. Measured by length, the answer looks like half – 50% probability; however measured in terms of volume, it is one part in eight (12.5% probability). Same cubes, same question – but you may receive two answers, shaped only by which yardstick you picked. Futures studies obviously cannot satisfy the symmetry requirements built in this interpretation of *probability*. A question like “a major disruption in industry and workforce in region X by 2040” has no natural list of outcomes to count, so the favourable-over-total ratio never gets off the ground. And even if you force a list, nothing really makes its entries equally likely – you have simply chosen the items on the list, forced them arbitrarily. This arbitrariness that Bertrand and van Fraassen exposed in their “puzzle”.

The *frequentist* interpretation treats probability as frequency: how often something does happen as a how often something happens, divided by how many times it could have happened. The finite version counts *actual* cases: how many members of a reference class have the attribute (Venn, 1876). 7 heads in 10 tosses means a probability of 0.7. The hypothetical version imagines an endless run of trials and takes the frequency it would settle towards (Reichenbach, 1949; von Mises, 1957). The finite version runs straight into what Hájek (1996) calls *the problem of the single case*, which torments futures studies almost everywhere. Toss a coin a single time and heads comes up either always or never, whatever the coin's real bias, and therefore the intuitively “real probability” is. events futurists care about (the fall of a particular regime, the arrival of a particular technology, a pandemic from a new pathogen) happen once, if at all – and that is exactly what makes them interesting. This means the frequentist interpretation hardly is meant when futures studies mention *probable futures*. Futures are tossed at most once: there is only one history. Both versions (the finite and the hypothetical one) also share the reference-class problem discussed above: any event belongs to many classes at once, each with its own frequency. For example, a coming energy transition, the probability of which futures studies may be interested in, might be grouped together with past energy transitions, with infrastructure overhauls in general, or with policy-driven technology shifts (or what have you – the

problem becomes greater, the more items you list) and each grouping gives a different number for probability, if taken as frequentist concept.

The *subjective* or *Bayesian* interpretation treats probability as a rational agent's degree of confidence. The tradition was shaped decisively by Ramsey (1926/1990) and de Finetti (1937/1980), who "measured belief" through kind of betting: your degree of belief in an event is the price at which you would buy or sell a bet on it. However, the so-called *Dutch Book argument* shows the danger in confidence levels that are not coherent. Suppose, in a two-horse race, you call each horse 60% likely to win: the figures each feel reasonable, yet they sum to 120% for an event that must come to exactly 100%. Incoherence of this kind (and it takes subtler forms too, of course, above we have a toy-example) lets a bookie hand you a set of bets you would each accept as fair but that together guarantee a loss (Skyrms, 1984).

The issue is that such incoherencies can be already quietly at work every time a Delphi panellist clicks a percentage button – take, for example 10 claims; who really counts their total percentage? Even if the events (that the Delphi claims present) are then argued to be not mutually exclusive, this is *ad hoc* and does not really remove the conceptual core of the problem. Moreover, *Orthodox Bayesianism*, in de Finetti's style, asks only two things of rational confidence: obey the probability calculus, and update by conditionalisation. This is rather permissive. Two equally rational panellists, given the same evidence, can each assign very different probabilities for the same event. If *probable* in futures studies means "believed likely by some rational agent," the notion of *probable* carries less weight than one might hope. Many rational agents, inflation of probable futures – we guess no one would happy about such inflation.²

The *propensity* interpretation treats probability as an objective physical *disposition* or *tendency* (Popper, 1957, 1959; Peirce, 1957), and it might seem more promising for futures studies: it grounds probability in the items of the world. A society characterized by inequality might be said to have a propensity toward political instability, in the same way a loaded die has a propensity to land on certain faces. But propensity interpretations can be slippery. For example, Hitchcock (2002) notes that calling something a "propensity" does little to specify what the *property actually is*. Moreover, Humphreys

²Bayesians have, of course, added further constraints such as matching credences (confidence) to observed frequencies (van Fraassen, 1983), or the Principal Principle, which says that when you know the real, objective odds of something, your personal confidence should simply match them (Lewis, 1980). However, these constraints tame subjective probability by tying it to something outside the individual mind – to frequencies, or to chances handed down by our best models – and share similar issues as the interpretations based directly to them.

(1985) shows that propensities do not fit Bayes' theorem, the rule for revising a probability in the light of new evidence (for example, you have a hunch, something happens, and you revise the belief behind the hunch in light of what happened – Bayes' theorem is simply the exact arithmetic for that revision). The theorem works in both directions, linking the chance of A given B to the chance of B given A, but a propensity only points one way. A viral mutation can have a tendency to cause a pandemic, but it makes no sense to say a pandemic has a tendency to have been caused by that particular mutation – causes push forward to their effects, not backward. The forward probabilities are there; the backward ones the theorem needs are not. This may seem something we can simply ignore as a formal game, but if you buy an interpretation of probability, you must take all of it. And the problem above does matter to futures studies: wherever one reasons backward from a future outcome to its likely cause (how probably a coming pandemic would trace to a lab accident, say), and this is exactly the type of inference propensities cannot make sense of.

What, then, should futures studies do with *probability*? Maybe the best way to approach the issue is by analogy. The philosophy of climate science offers a good model for handling probability under conditions of deep uncertainty – and it is worth focusing on, because climate scientists face problems also discussed in futures studies. They must assign probabilities to outcomes – global temperature rise, sea-level change, forest fires, heat-waves – that are unique, that arise from complex systems characterized by deep uncertainty, and that are extremely urgent for policies. Climate scientists cannot run the climate system repeatedly to search frequencies. They cannot lean on naive subjective estimates, because the stakes demand intersubjective accountability. And they cannot avoid probability altogether, because decision-makers demand estimates of probability (to count the cost of something against the risks, for example).

In this context a more disciplined form of expert judgment has emerged, one that distinguishes among the sources of uncertainty rather than collapsing them into either false precision or unhelpful vagueness. Its most visible public expression is the IPCC's calibrated language (see Mastrandrea et al. 2011), which pairs graded probability terms with a separate scale for the strength of evidence and the degree of expert agreement. Rather than asking "how likely is this outcome?" in the abstract, researchers ground their assessments in causal models. In philosophy of climate science Hannart et al. (2016) have built a causal counterfactual framework for the attribution of weather and climate-related events, drawing explicitly on Pearl's (2000; 2009) interventionist semantics (see above). Stott et al. (2016) argue that pairing physical models with observational constraints allows the factual

climate to be compared against a counterfactual world without human “intervention”, which provides estimates of the human contribution to the risk of an event. In this way, the question shifts from “what is the *probability* of this outcome?” to “how much more, or less, *likely* does this outcome become *given a specific causal intervention*?” The move connects probability to the interventionist causal reasoning that, we argued above, should underwrite plausibility. Probability and plausibility are not the same notion, but we do not need to be exact how their difference is formulated. Both can be translated to the language of counterfactual causality. This not only makes the notions easier to grasp but supports our argument that counterfactual reasoning is the “hidden engine” of futures thinking, at least in plausibility and probability.

The lesson for futures studies is not that the field should import climate science’s specific methods as such – we build on an analogy. The lesson is that futures studies should take seriously the distinction between different *grounds* for a probability assignment. A probability derived from a well-specified causal model, even a rough one, carries more epistemic weight than a probability derived from an someone’s “gut feeling” or “intuition” – even when both are expressed with the same numerical value.

The scenario-planning literature is, in a sense, parallel to this at places at least. Ramírez and Selin (2014) argue that, whatever role *probability* may play within scenario work, it cannot legitimately serve as something that crowns one scenario the “most likely” to the extent that the scenario is to become a type base case (a practice Ramirez and Selin find in some prominent consulting guidance). This is because the very situations that call for scenarios are situations of *Knightian uncertainty* (i.e., risk without measurable or quantifiable probabilities; Knight, 1921), where the requisite data and stable distributions simply do not exist. To stamp *probability* to something under those conditions is not rigour. This is the same caution our distinction between different epistemic grounds (see above) is meant to argue for: a number that no model or adequate real-world data can back up is just a personal estimate dressed up to look like a measurement – and being honest about which it really is makes all the difference. Personal estimates can be extremely useful and “correct” but the difference, on a conceptual level, must be made.

The upshot for the aims of conceptual engineering is the following: When futures studies project or task uses the term *probable*, it should be clear about which kind of *probability* is supposed to be at play. At a minimum, practitioners should distinguish between *model-based probabilities* (derived from specified causal or structural models, bracketing for now their different qualities), *expert-*

judgment probabilities (collected through structured processes like Delphi, with the reasoning documented, not just “clicked”), and *subjective credences* (individual degrees of confidence, acknowledged as such). These are not equally authoritative, epistemically speaking, and pretending otherwise is where the trouble starts. A futures studies project or task that reports a future as *probable* while specifying which kind of probability is in question within the claim already represents a real advance in rigour. And when no *probability* can responsibly be defined – the uncertainty too deep, the models too inadequate, the situation too novel – the honest response is to say so, rather than to disguise ignorance in terminology of *probability* or *likelihood*.

“Desirable” and the Hidden Modalities of Preference

The fourth term that deserves attention here and is included in the standard vocabulary is *desirable*, (often *preferable*)³ – is usually treated differently than the notions discussed in this text thus far. Possibility, plausibility, and probability look like ontic or epistemic concepts: they describe the world or our knowledge of it and confidence. Desirability is an axiological (value-related) concept: it describes what we find good or just. Bell (1997) dedicated the entire second volume of his *Foundations of Futures Studies* to normative questions- Bell argued that value judgments are an inescapable part of futures studies. This is true to a great extent. We often ask about desirability. Moreover, to mention some examples, backcasting (Robinson, 1982) takes a desirable future as its starting point and works backward to find the steps to reach it; and there is the “desirable” zone in the *futures cone* (Voros, 2003) – an area of futures that can be mapped.

However, *desirable* does not separate from the other modal concepts in its core. To judge a future desirable is, under the surface, to judge it, first of all, possible (assuming one does not rationally desire what one believes to be truly impossible). Second, there are several desirable futures (hopefully) and this means not all of them can actualize, thereby, some are merely possible. Moreover, it seems that to select a desirable future as the target of a backcasting process is to judge it, at minimum, plausible (given our definition of plausible: something that can be reached). In backcasting, the desirable future must be reachable via some pathway from the present. Desirability, in other words, carries hidden modal commitments. Every desirable future implicitly refers possible, sometimes at least *plausible*, and sometimes *achievable, given the right actions*.

³As noted, we use the term desirable whenever possible, i.e., when we are not citing author whose use of preferable is not at the very core of their writing.

This means that the quality of a normative futures studies research more than likely depends on the quality of the modal reasoning behind it, at least to the extent the modal reasoning constitutes the researcher at hand. A desirable future that turns out to be physically impossible is not really desirable in any *action-guiding sense*. Rather, it is a fantasy. A desirable future that lacks any identifiable causal pathway from the present is an ideal not tied to real strategy. When conceptually engineering *desirable futures* we should therefore drag these hidden commitments into the light: when future studies discusses desirable futures, it should be able to say not only *why* that future is valued but in what way it is possible, plausible, probable, and so on. The take away lesson is that one cannot simply talk about desirability without committing some other modal concepts. Desirability cannot fly free from any other considerations of what we consider as modal in any sense of that word. There is no one-to-one connection between *desirability* and some another modal category; sometimes desirable futures may be limited to plausible ones, sometimes even “merely” possible count, and so on. But one has to explicate what connection one means by her use of the notion desirability.

We may finally notice that desirability is connected to agency. In Belnap et al.’s (2001) STIT framework, roughly, each choice available to an agent at a moment narrows the branching futures down to a subset, and an agent “sees to it that” something happens when the outcome is guaranteed by that choice rather than coming about anyway, independently of it. Notice, again, that there is dependency-relations in play here. A *desirable future*, then, in STIT terms, is one whose realisation depends on some agent’s choices that lead to a better option that would be reached without that choice. If no identifiable agent has the capacity to bring about certain “desirable future”, calling it “desirable” is misleading, wishful thinking. At least in cannot be setting a target. So discussion of ontology of possibilities and modal vocabulary are, not surprisingly, intimately connected. However, we cannot map be the connections here to their full extent.

Counterfactual Reasoning: The Hidden Engine of Scenario Construction

In this writing, we have argued that *plausibility* should require an identifiable causal pathway and that *probability* should, where it can, be grounded in causal models – or at least then the probability judgements become most robust. Both requirements share the same type of reasoning that runs through futures studies: *counterfactual reasoning*. Counterfactual reasoning is disciplined asking of “what-if-things-had-been-different” questions (in its simplest form: what would happen if this variable changed, for example if this or that event occurred, if this policy were adopted, and so on). In what follow, we make more explicit why there are good reasons to consider it as the hidden

cognitive engine that drives scenario construction, modal assessment (that is, assessments of possible, plausible, probable, and so on), and most of the analytical work that separates competent foresight from speculation.

The idea is that knowledge of what is possible grows out of our capacity for counterfactual reasoning, i.e., reasoning of the form “if X were the case, then Y would be the case.” We can now put this to work on what scenario builders are actually doing when they construct and evaluate scenarios. A useful guide here is Virtajoki’s (2023) account of causal explanation in history, which discusses, in its final chapter explicitly the similarities between counterfactual based causal explanations in history and futures studies.

The starting point is a simple conception of *cause*: a cause of an outcome is something *that makes a difference* to an outcome. A long tradition running back to at least Max Weber (1949), sees causal analysis as the work of separating the factors that made a difference from those that did not. This is also what futures studies does when it the driving forces and critical uncertainties that affect the future – what would change, if this or that was to change? If nothing, there is no causal relevance. If something changes, we have identified something valuable. A good scenario, like a good historical explanation, isolates the variables that are relevant to historical trajectories (including trajectories towards the future; they will become history at some point).

A counterfactual account of (explanation-seeking) reasoning can be based on Woodward’s *interventionist* idea that an explanation answers *what-if-things-had-been-different* questions (2003). Explanations are *contrastive*: they answer “why X *rather* than Y?” by identifying factors such that, had things been different in some specific way, Y *rather* than X would have been the case. Scenarios share this superstructure. We may ask “what if oil prices will double rather than remains the same?” – a what-if question – and answer by providing a scenario (or set of scenarios) where we present the consequences in the form that makes clear how the world could have been different from how (the “rather” part) it may be when the prices do not double. In general, a scenario-builder considers a change and traces what would follow.

The interventionist account also gives a clean solution to the problem that is present whenever we discuss the future: which features of the present we *hold fixed* and what do we let change? If “everything can change” is on the table (meaning that we are allowed to consider that *everything* changes at once), then chaos follows (of course, we do not argue against the possibility that everything changes, but, we wish to point out that assuming that everything changes suddenly leaves little room

to stay on track what the future would be like). In the interventionist (counterfactual) framework, we fix certain factors and ask what follows, if they are thus fixed. To simplify a bit, we stipulate that something is in way or another (or will be) and then follow through. This idea (that certain factors are fixed) is standard in causal explanation literature, and should generate no further trouble for futures studies either. We must specify what changes we allow in any given exercise, and may bracket others for later use.

Consider now the fact, discussed through this text, that how we understand *plausible*, *probable*, and so on, differ greatly in their consequences of how seriously we should take the associated judgements concerning *plausible* or *probable* (or what have you) *futures*. This is where a link from history to the future, already present in the earlier work on history-future links (see Staley, 2002; Bradfield et al., 2016 – we do not claim to have discovered the link), becomes most directly relevant to understanding what makes some futures studies better than others. The factors that give the *deepest* counterfactual explanations – the factors that would have produced similar outcomes across many alternative pasts – are exactly the *factors robust enough* to matter across many different cases, including possible futures. For example, if alliance systems and geopolitical tensions would have led to large-scale war down numerous alternative historical paths in early 20th century, we have reason to expect that alliance systems and geopolitical tensions to shape the difference between war and peace also in the future. Past *contingencies* do not repeat in detail, but *certain structural causal features* seem to repeat themselves in certain forms. It also follows that the sheer number of things that *could* change a future outcome should not make us unable to say anything about the future: as with the past, not all of those “could”-changes are equally relevant (what if asteroid hit the earth is different than from asking what if alliances were different in early 20th century, to give a toy example), and the work is to separate *plausible* or *probable* departures of things from business-as-usual from far-fetched ones. What counts as *plausible* or *probable*, of course, depends on how we define these – and the relevance of *probable* and *plausible* departures may vary accordingly. Still some departures, some what “what if”-s are more relevant than other; some are plausible, some are merely noncontradictory.

Finally, we may notice how counterfactual reasoning as the engine of scenario construction has several positive consequences for understanding futures studies. First, it explains what makes some scenarios better than others: better scenarios are based on better what-if reasoning. A sharp and robust identification of difference-makers is shared by good counterfactual thinking and futures thinking. Good what-if thinking, in turn, is based on the fact that sometimes we, or some group, have better

knowledge of certain domain that the what-if reasoning is about. Experienced practitioners produce better scenarios not because they are more imaginative but because they reason about dependencies with more precision. Or so it seems natural to argue.

Second, it clarifies the relationship between evidence and scenarios. As Virtajoki (2023) notes, historical explanation and evidence are developed hand-in-hand. Explanations cannot be simply extracted from pre-given facts, since deciding what would count as relevant evidence is itself shaped by the explanation under construction. We can know what evidence to look for only once we attempt to sketch an explanation. Scenarios work the same way: The futures we sketch as possible shape which evidence about the present (and past) we treat as evidence for the futures we sketch – and from where we look for the evidence. In general, we should be aware of the relationship between (re)presentations of the future and the evidence they rely on, and how the two are developed and found hand-in-hand.

Third, and most importantly, the idea of counterfactual reasoning as the engine of futures thinking connect the field's modal vocabulary (*possible*, *probable*, *plausible*, *desirable*, and so on) in ordinary cognition. We can judge a future *possible* because we can coherently develop the idea of a condition that *would* lead to it; we can judge future *plausible* when we can specify a causal pathway (series of counterfactual dependencies) connecting it to the present; and more or less *probable* by asking how robust the chain is to our knowledge is – what we consider could change it. A future is *desirable* if it *would* be good or just and at least possible. The more disciplined the reasoning in any given case, the better the futures studies. Asking what-if-things-had-been-different questions is at the heart of futures thinking, our argument goes. The issue runs even deeper: To add an icing on the cake, according to Williamson (2007), our knowledge of not just ordinary possibilities but also about metaphysical possibility and necessity is not the product of a special faculty but is simply a by-product of our ordinary capacity to evaluate counterfactual conditionals. Counterfactual reasoning pierces through our minds.

From Vocabulary to Practice

The arguments of this text have put on the table and suggested certain proposals for how the field's core modal terms should be understood. These are not through-and-through one-size fits for all, imposed, *definitions* but ideas (building on the practice of conceptual engineering) about refinements and added rigour to what is needed from modal vocabulary of the field. We argued as follows:

Possible should be defined whenever used: at minimum, futures studies should distinguish *logical possibility* (no contradiction), *physical possibility* (a future situation that is consistent with laws of nature), and *practical possibility* (a future situation consistent with laws of nature and other facts we know about how world works); things can be “possible” in many ways, and this matter above all when futures studies is connected decision-making. In the case of *plausible*, we have taken a bit more explicit stance on how it should be understood: we should require *positive plausibility*, which means the identification, or at least of a sketch, of a causal pathway from the present to the described future. *Negative plausibility*, which means the mere absence of identified constraints, is too weak to carry the weight we wish to give to *plausible futures* in futures studies. *Probable* can be understood in many ways, and it is difficult to pinpoint the “best” interpretation of the concept. We argued that transparency about the grounds of the assignment: model-based, expert-judgment, or subjective credence (“confidence”). Finally, one must understand that *desirable* implicit modal commitments that need to be made explicit – one has to tell if a desirable future at hand at any moment is associated with possible, plausible, or probable reading of the modality of that *desirable* future.

How do these proposals connect to existing futures studies methods? This question is what ties, in conceptual engineering, concepts to practice. There are clear connections between the definitions above and the methods, already discussed (but, of course, we cannot explicate all of them and gaps are possible, to be honest). Horizon scanning and weak-signal (or future sign) detection (Ansoff, 1975; Hiltunen, 2008) can be associated with forms of *possibility*: they look for and catch signals that indicate that something previously unthought may be possible. The positive plausibility standard does not constrain these methods, since horizon scanning is precisely an attempt to widen the boundary of what we recognise as possible, not to find paths to the possibilities. It is after the signal that the standard bites, if it is to bite: moving from “this might be possible” to “this is plausible” is the moment a causal pathway is sketched.

Scenario planning (see e.g., Wack, 1985a, 1985b; Schwartz, 1991; Schoemaker, 1995; Bradfield et al., 2005; Amer et al., 2013) is perhaps best understood at in terms of *plausibility* and, in some versions, *probability*. In our analysis. futures studies thinking is embedded in counterfactual reasoning: the scenario builder considers changes to key variables and traces the consequences. The *positive plausibility* standard requires that each scenario is delivered with at least an outline of the causal pathway from the present to the future the scenario describes. Skilled practice already tends to do this, but making the requirement concerning causal pathways explicit gives a quality criterion to

scenario planning that can be taught, assessed, and improved. Again, we are not saying that something has to be done; we only attempt to set the ground for what could be done.

The Delphi method (long tradition in futures studies, starting from e.g., Dalkey & Helmer, 1963; Linstone & Turoff, 1975;) is perhaps best understood at in terms *probability* and should be explicit about which interpretation of *probability* it uses. Notice a trap here: probable is not always the “this is the most automatic to happen; more than plausible (no matter what items like the *futures cone* presents). The weight of probability assignment depends on how probability is understood and operationalized – that is exactly our argument and where philosophy becomes useful; probability is one topic which is difficult and confusing, and *probable* cannot be used without care. Be that as it may, the Delphi method is usually built to extract expert-judgment about probabilities. However, an expert who merely clicks “percentage buttons” with no documented reasoning delivers a subjective credence at best (and this subjective credence can often be incoherently structured, strictly speaking; see above). Recording the reasoning, not just the number, lets a Delphi exercise separate well-grounded judgments from guesses and weight them accordingly. Interestingly, Landeta’s (2006) review makes a similar point: Delphi can be taken as valid when applied with methodological rigour, but we must remain cautious if consensus is mistaken for truth and its measurement done without sufficient care. We argue that flagging the grounds of each assessment is one cheap way of keeping the two apart.

We provide, as an example, the observation that cross-impact balance analysis (Weimer-Jehle, 2006) aligns somewhat naturally with positive plausibility, since it looks for configurations of variables that are internally *consistent*, *given* certain network of influences; interdependencies that can reasonably be read as causal (causality is tricky even if we accept it as something living inside a method, see e.g., Woodward 2003. In this context, *plausibility* means the scenario corresponds to a self-consistent configuration. Morphological analysis (Ritchey, 2018) does something similar, but what it calls *cross-consistency assessment* diagnoses what can *coexist* rather than causal structure as such. As we saw, Ritchey (2018, 85) is explicit in that the method does need not refer to causality. This means that surviving cross-consistency analysis is not, on its own, *positive plausibility* in the sense we defended in this text. The conceptual study of this text makes visible what these methods do while steering clear about where consistency stops and causality begins; where *plausibility* is bracketed, where it is not.

Finally, we may notice how Causal Layered Analysis (CLA) (Inayatullah, 1998) operates at several levels of depth at once, and modalities are intertwined with several issues and work in several ways. No single modal category applies across the levels of CLA uniformly; the demand each level places on a modal claim shifts as one moves through the levels. At the litany level of trends and headlines, the relevant modality comes to close to *practical possibility* and, where numbers attach, *probability*: claims here are continuous with business-as-usual and tied to what people think world turns out to be, given what appears obvious. At the systemic level of structures and causes, we move closer to *positive plausibility*: the analyst identifies causal relations that would have to be made intervention to in order to world to change. As we move onto the level of discourse and worldview, what shifts are the framings and claims of legitimacy that determine which futures possible to think, given one has certain worldview. For example, futures of science are seen differently given different philosophies of science, possibilities open and close as one accepts this or that philosophical worldview on science (see Virmajoki 2022). At the deepest level of myth and metaphor, the modalities reach furthest, probing the *logical* and *metaphysical* possibilities that look toward the edge of what we can think, and the standards of modality rightly relax (which does mean this level is easy to map, far from it). Relaxation is not abolition: we can still be precise what type of possibilities we talk about – even plausibilities, given we find mechanism where myth produces visible outcomes or commitment (see discussion in Virmajoki 2022). CLA thus does not escape the modal vocabulary. Rather, it *stratifies* it, and tracking down the relationship between the levels of CLA and modalities in futures studies would be a worthwhile exercise for that very reason. Unfortunately, that must wait treatment outside the scope of this text.

Conclusion

This writing has argued that the modal vocabulary at the heart of futures studies – the terms practitioners use to classify, weigh, and communicate futures – is in need of *engineering*, not just *interpretation*, as the usage is, to put it bluntly, wild. *Possible*, *plausible*, *probable*, and *desirable* are not mere labels to be applied however one wishes, at least if we wish to maintain some posture in the unity of futures studies and how it communicates to society out there. These modal notions are associated with epistemic claims, and they carry both practical and philosophical weight. That weight can only be assessed, if modal vocabulary is communicated clearly. *Possibility* should specify what kind of possibility. *Plausibility* should say whether it is associated with causal pathways – if not, we argue, it is term used too loosely. *Probability* should perhaps not deeply disclose its exact

interpretation but still explain what grounds when one is saying something is *probable*. *Desirability* claims should make explicit their hidden commitments to other forms of modality discussed in this text.

We have also argued that counterfactual reasoning is the hidden cognitive engine of scenario construction and, in general, of the futures studies when it moves in the space of possibilities – *what could have happened, had this or that changed* is structurally similar to *what may happen, given this or that change*. The claim about the hidden engine ties well together, when looked closely, with the work in futures studies – asking “what if?”, tracing consequences, evaluating pathways – which is why we make the argument in the first place. Understanding conceptual core is necessary, and our claim about counterfactual reasoning is a claim about this conceptual core, when it comes to modality. The link is not decorative, then. Counterfactual reasoning is to foresight what calculation is to accounting; it is not the whole of the practice, but it is operation without which the practice could not get off the ground. Counterfactual armchair pondering is no substitute for real methods.

But sharpening the vocabulary raises a further question we have raised but not really answered. If plausibility requires causal pathways and probability requires transparency about grounds, what counts, in detail, as relevant evidence for such claims? How do we know that a causal pathway is genuine and not just a plausible-sounding story – causality is tricky, as noted (and often loosely used as a concept in futures studies, unfortunately – where are, say, interventionist causal models or equations beyond forecasting, where talk about *causality* is still there)? How do we assess whether probability assignments are well calibrated and coherent? What are the limits of our ability to know what lies ahead at all – what really determines what we can think about possibility even in its widest sense? These are *epistemological* questions – questions about the nature and the limits of the knowledge that futures studies produces and can produce. This raises another challenge: what, exactly, can be known about the future, and how? No matter how much effort we spend to conceptually engineering modal vocabulary, the work is empty without associated treatment of the knowledge that is transmitted with and through modal vocabulary. Again, this is out of the scope of this writing.

Bibliography

Amara, R. (1981). The futures field: Searching for definitions and boundaries. *The Futurist*, 15(1), 25–29.

- Amer, M., Daim, T. U., & Jetter, A. (2013). A review of scenario planning. *Futures*, 46, 23–40.
- Ansoff, H. I. (1975). Managing strategic surprise by response to weak signals. *California Management Review*, 18(2), 21–33.
- Armstrong, D. M. (1997). *A world of states of affairs*. Cambridge University Press.
- Bell, W. (1997). *Foundations of futures studies: Vol. 2. Values, objectivity, and the good society*. Transaction Publishers.
- Belnap, N., Perloff, M., & Xu, M. (2001). *Facing the future: Agents and choices in our indeterminist world*. Oxford University Press.
- Bertrand, J. (1889). *Calcul des probabilités*. Gauthier-Villars.
- Blackburn, S. (1999). *Think*. Oxford University Press.
- Bradfield, R., Wright, G., Burt, G., Cairns, G., & van der Heijden, K. (2005). The origins and evolution of scenario techniques in long range planning. *Futures*, 37(8), 795–812.
- Bradfield, R., Derbyshire, J., & Wright, G. (2016). The critical role of history in scenario thinking: Augmenting causal analysis within the intuitive logics scenario development methodology. *Futures*, 77, 56–66.
- Brun, G. (2016). Explication as a method of conceptual re-engineering. *Erkenntnis*, 81(6), 1211–1241.
- Burgess, A., & Plunkett, D. (2020). On the relation between conceptual engineering and conceptual ethics. *Ratio*, 33(4), 281–294.
- Cappelen, H. (2018). *Fixing language*. Oxford University Press.
- Cappelen, H., & Plunkett, D. (2020). Introduction: A guided tour of conceptual engineering and conceptual ethics. In A. Burgess, H. Cappelen, & D. Plunkett (Eds.), *Conceptual engineering and conceptual ethics* (pp. 1–26). Oxford University Press.
- Carnap, R. (1950). *Logical foundations of probability*. University of Chicago Press.
- Carus, A. W. (2007). *Carnap and twentieth-century thought*. Cambridge University Press.

Chalmers, D. J. (2002). Does conceivability entail possibility? In T. S. Gendler & J. Hawthorne (Eds.), *Conceivability and possibility* (pp. 145–200). Oxford University Press.

Chalmers, D. J. (2025). What is conceptual engineering and what should it be? *Inquiry*, 68(9), 2902–2919.

Creath, R. (1991). Every dogma has its day. *Erkenntnis*, 35, 347–389.

Dalkey, N., & Helmer, O. (1963). An experimental application of the Delphi method to the use of experts. *Management Science*, 9(3), 458–467.

Dator, J. (2009). Alternative futures at the Manoa School. *Journal of Futures Studies*, 14(2), 1–18.

de Finetti, B. (1980). Foresight: Its logical laws, its subjective sources (H. E. Kyburg Jr., Trans.). In H. E. Kyburg Jr. & H. E. Smokler (Eds.), *Studies in subjective probability* (pp. 55–118). Robert E. Krieger. (Original work published 1937)

Dutilh Novaes, C. (2020). Carnapian explication and ameliorative analysis. *Synthese*, 197(3), 1011–1034.

Floridi, L. (2011). *The philosophy of information*. Oxford University Press.

Gordon, T. J., & Hayward, H. (1968). Initial experiments with the cross-impact matrix method of forecasting. *Futures*, 1(2), 100–116.

Hájek, A. (1996). “Mises redux”—Redux: Fifteen arguments against finite frequentism. *Erkenntnis*, 45, 209–227.

Hale, B. (2013). *Necessary beings*. Oxford University Press.

Hancock, T., & Bezold, C. (1994). Possible futures, preferable futures. *Healthcare Forum Journal*, 37(2), 23–29.

Hannart, A., Pearl, J., Otto, F. E. L., Naveau, P., & Ghil, M. (2016). Causal counterfactual theory for the attribution of weather and climate-related events. *Bulletin of the American Meteorological Society*, 97(1), 99–110.

Haslanger, S. (2005). What are we talking about? The semantics and politics of social kinds. *Hypatia*, 20(4), 10–26.

- Haslanger, S. (2012). *Resisting reality: Social construction and social critique*. Oxford University Press.
- Helmer, O. (1981). Reassessment of cross-impact analysis. *Futures*, 13(5), 389–400.
- Hiltunen, E. (2008). The future sign and its three dimensions. *Futures*, 40(3), 247–260.
- Hitchcock, C. (2002). Probability and chance: Philosophical aspects. In *International encyclopedia of the social & behavioral sciences* (Vol. 18, pp. 12089–12095). Elsevier.
- Humphreys, P. (1985). Why propensities cannot be probabilities. *The Philosophical Review*, 94(4), 557–570.
- Inayatullah, S. (1998). Causal layered analysis: Poststructuralism as method. *Futures*, 30(8), 815–829.
- Isaac, M. G., Koch, S., & Nefdt, R. M. (2022). Conceptual engineering: A road map to practice. *Philosophy Compass*, 17(10), Article e12879.
- Knight, F. H. (1921). *Risk, uncertainty, and profit*. Houghton Mifflin.
- Koch, S. (2021). There is no dilemma for conceptual engineering: Reply to Max Deutsch. *Philosophical Studies*, 178(7), 2279–2291.
- Kripke, S. A. (1980). *Naming and necessity*. Harvard University Press.
- Landeta, J. (2006). Current validity of the Delphi method in social sciences. *Technological Forecasting and Social Change*, 73(5), 467–482.
- Laplace, P. S. (1999). *A philosophical essay on probabilities* (A. I. Dale, Trans.). Springer. (Original work published 1814)
- Lewis, D. (1980). A subjectivist's guide to objective chance. In R. C. Jeffrey (Ed.), *Studies in inductive logic and probability* (Vol. 2, pp. 263–293). University of California Press.
- Lewis, D. (1986). *On the plurality of worlds*. Blackwell.
- Linstone, H. A., & Turoff, M. (Eds.). (1975). *The Delphi method: Techniques and applications*. Addison-Wesley.

- Mastrandrea, M. D., Mach, K. J., Plattner, G.-K., Edenhofer, O., Stocker, T. F., Field, C. B., Ebi, K. L., & Matschoss, P. R. (2011). The IPCC AR5 guidance note on consistent treatment of uncertainties: A common approach across the working groups. *Climatic Change*, 108(4), 675–691.
- Nado, J. (2021). Conceptual engineering, truth, and efficacy. *Synthese*, 198(Suppl. 7), 1507–1527.
- Nado, J. (2023). Classification procedures as the targets of conceptual engineering. *Philosophy and Phenomenological Research*, 106(1), 136–156.
- Pearl, J. (2000). *Causality: Models, reasoning, and inference*. Cambridge University Press.
- Pearl, J. (2009). *Causality: Models, reasoning, and inference* (2nd ed.). Cambridge University Press.
- Peirce, C. S. (1957). Notes on the doctrine of chances. In V. Tomas (Ed.), *Essays in the philosophy of science* (pp. 74–84). Bobbs-Merrill.
- Plantinga, A. (1974). *The nature of necessity*. Oxford University Press.
- Popper, K. R. (1957). The propensity interpretation of the calculus of probability and the quantum theory. In S. Körner (Ed.), *Observation and interpretation* (Colston Papers, Vol. 9, pp. 65–70). Butterworths.
- Popper, K. R. (1959). The propensity interpretation of probability. *The British Journal for the Philosophy of Science*, 10(37), 25–42.
- Priest, G. (2021). Metaphysical necessity: A skeptical perspective. *Synthese*, 198(Suppl. 8), 1873–1885.
- Queloz, M., & Bieber, F. (2022). Conceptual engineering and the politics of implementation. *Pacific Philosophical Quarterly*, 103(3), 670–691.
- Ramírez, R., & Selin, C. (2014). Plausibility and probability in scenario planning. *Foresight*, 16(1), 54–74.
- Ramsey, F. P. (1990). Truth and probability. In D. H. Mellor (Ed.), *Philosophical papers* (pp. 52–94). Cambridge University Press. (Original work published 1926)
- Reichenbach, H. (1949). *The theory of probability*. University of California Press.

Ritchey, T. (2018). General morphological analysis as a basic scientific modelling method. *Technological Forecasting and Social Change*, 126, 81–91.

Robinson, J. B. (1982). Energy backcasting: A proposed method of policy analysis. *Energy Policy*, 10(4), 337–344.

Schoemaker, P. J. H. (1995). Scenario planning: A tool for strategic thinking. *Sloan Management Review*, 36(2), 25–40.

Schwartz, P. (1991). *The art of the long view: Planning for the future in an uncertain world*. Doubleday.

Shields, M. (2021). Conceptual domination. *Synthese*, 199(5–6), 15043–15067.

Shoemaker, S. (1998). Causal and metaphysical necessity. *Pacific Philosophical Quarterly*, 79(1), 59–77.

Skyrms, B. (1984). *Pragmatics and empiricism*. Yale University Press.

Spirtes, P., Glymour, C., & Scheines, R. (1993). *Causation, prediction, and search*. Springer-Verlag.

Staley, D. J. (2002). A history of the future. *History and Theory*, 41(4), 72–89.

Stott, P. A., Christidis, N., Otto, F. E. L., Sun, Y., Vanderlinden, J.-P., van Oldenborgh, G. J., Vautard, R., von Storch, H., Walton, P., Yiou, P., & Zwiers, F. W. (2016). Attribution of extreme weather and climate-related events. *WIREs Climate Change*, 7(1), 23–41.

Strohming, M. (2015). Perceptual knowledge of nonactual possibilities. *Philosophical Perspectives*, 29(1), 363–375.

Thomasson, A. L. (2025). Conceptual engineering: When do we need it? How can we do it? *Inquiry*, 68(9), 2993–3018.

Vaidya, A. J., & Wallner, M. (2021). The epistemology of modality and the problem of modal epistemic friction. *Synthese*, 198(Suppl. 8), 1909–1935.

van Fraassen, B. C. (1983). Calibration: A frequency justification for personal probability. In R. S. Cohen & L. Laudan (Eds.), *Physics, philosophy and psychoanalysis: Essays in honor of Adolf Grünbaum* (pp. 295–319). D. Reidel.

van Fraassen, B. C. (1989). *Laws and symmetry*. Clarendon Press.

van Inwagen, P. (1998). Modal epistemology. *Philosophical Studies*, 92(1–2), 67–84.

Venn, J. (1876). *The logic of chance* (2nd ed.). Macmillan.

Vetter, B. (2023). An agency-based epistemology of modality. In D. Prelević & A. Vaidya (Eds.), *The epistemology of modality and philosophical methodology* (pp. 44–69). Routledge.

Virmajoki, V. (2022). Understanding futures of science: Connecting causal layered analysis and philosophy of science. *Journal of Futures Studies*, 27(2).

Virmajoki, V. (2023). *Causal explanation in historiography*. Palgrave Macmillan.

von Mises, R. (1957). *Probability, statistics and truth* (rev. English ed.). Macmillan.

Voros, J. (2003). A generic foresight process framework. *Foresight*, 5(3), 10–21.

Wack, P. (1985a). Scenarios: Uncharted waters ahead. *Harvard Business Review*, 63(5), 72–89.

Wack, P. (1985b). Scenarios: Shooting the rapids. *Harvard Business Review*, 63(6), 139–150.

Weber, M. (1949). *The methodology of the social sciences* (E. A. Shils & H. A. Finch, Eds. & Trans.). Free Press.

Weimer-Jehle, W. (2006). Cross-impact balances: A system-theoretical approach to cross-impact analysis. *Technological Forecasting and Social Change*, 73(4), 334–361.

Williamson, T. (2007). *The philosophy of philosophy*. Blackwell.

Woodward, J. (2003). *Making things happen: A theory of causal explanation*. Oxford University Press.

Yablo, S. (1993). Is conceivability a guide to possibility? *Philosophy and Phenomenological Research*, 53(1), 1–42.