

Political Ideals and the Feasibility Frontier

David Wiens

Abstract. Recent methodological debates regarding the place of feasibility considerations in normative political theory are hindered for want of a rigorous model of the feasibility frontier. To address this shortfall, I present an analysis of feasibility that generalizes the economic concept of a production possibility frontier and then develop a rigorous model of the feasibility frontier using the familiar possible worlds technology. I then show that this model has significant methodological implications for political philosophy. On the Target View, a political ideal presents a long-term goal for morally progressive reform efforts and, thus, serves as an important reference point for our specification of normative political principles. I use the model to show that we cannot reasonably expect that adopting political ideals as long-term reform objectives will guide us toward the realization of morally optimal feasible states of affairs. I conclude by proposing that political philosophers turn their attention to the analysis of actual social failures rather than political ideals.

Political philosophers have increasingly turned their attention to questions regarding the proper place of feasibility considerations in normative (i.e., action-guiding) political theory. Much of this debate takes place under the rubric of “ideal vs. nonideal theory” (see Valentini 2012*a*). Can moral and political ideals be action-guiding despite their apparent infeasibility? To what extent should normative political principles be sensitive to feasibility concerns? These are important methodological issues, on which much remains to be said. But we can’t answer these questions until we have a firm grip on the concept of feasibility and the logic of feasibility assessments. On these issues, philosophers have had much less to say.

Author’s note. For helpful discussion, thanks to audiences at the Australian National University, the University of Michigan, and the University of Sydney. Thanks in particular to: Elizabeth Anderson, Christian Barry, David Braddon-Mitchell, Geoff Brennan, Rachael Briggs, Sarah Buss, Mark Colyvan, Billy Dunaway, Jason Konek, Holly Lawford-Smith, Seth Lazar, Leon Leontyev, David Lusby, John Maier, Sarah Moss, David Plunkett, Peter Railton, Paolo Santorio, Wolfgang Schwarz, Alex Silk, Dan Singer, Nic Southwood, Leif Wenar, and several anonymous referees. Thanks also to the Editor, Richard Bradley, for helpful suggestions. Support from the Australian Research Council (Discovery Grant DP120101507) is gratefully acknowledged.

My first aim is to develop a *restricted possibility account* of feasibility to address this shortfall. I start with a brief discussion of the rival *conditional probability account* (Brennan and Southwood 2007; Lawford-Smith 2013), showing that this view faces enough difficulties to warrant developing an alternative. I start my positive proposal in section 2 with the familiar but overly narrow economic concept of a production possibility frontier. To render this notion applicable to the concerns of political philosophers, I extend it to a more general concept of an all-purpose production possibility frontier. This yields an analysis of feasibility in terms of all-purpose resource availability. Roughly, realizing a state of affairs is *feasible* only if it can be realized given facts about the (financial, technological, motivational, institutional, etc.) resources at our disposal.¹ In section 3, I refine this analysis by developing a model of the feasibility frontier using a possible worlds framework. The main point of this model is to make precise the core components of the analysis and demonstrate their interplay. This gives us a rigorous base from which to investigate the broader implications of the analysis.

My second aim is to show that the restricted possibility account has significant implications for the aforementioned methodological disputes. On the conventional *Target View*, a political ideal presents a more or less well-defined target for real world political reform and, thus, serves as an important reference point for our specification of normative political principles (among others, see Buchanan 2004; Gilabert 2012*a*; Robeyns 2008; Simmons 2010; Valentini 2012*b*). In section 4, I first use core features of the model developed in section 3 to differentiate between two distinct specifications of the Target View (each interprets the nature of the target differently). This imposes some structure on the debate. I then show that the Target View (in both guises) is mistaken — we cannot reasonably expect that adopting political ideals as long-term reform objectives will guide us toward the realization of morally optimal feasible states of affairs. This conclusion might provoke pessimism about the prospects for action-guiding political philosophy. To combat such pessimism, I suggest that political philosophers can provide action-guiding normative principles with critical purchase on the status quo by turning their attention to the analysis of actual social failures rather than political ideals (cf. Wiens 2012).

1 I should note at the outset that I only present a necessary condition for feasibility here; hence, my analysis is not fully reductive. I originally thought that this necessary condition was sufficient too. For various reasons, I am inclined to resist this now. Alas, further discussion of the issue is beyond the scope of this paper; in any case, my thoughts regarding additional necessary conditions are unsettled. Yet my analysis remains significant for at least three reasons: it identifies and makes precise an analytic condition that resonates with ordinary usage of the concept; it fixes a somewhat restrictive outer bound on the feasible set (certainly more restrictive than Lawford-Smith's rival account); and it offers a baseline model for understanding the relationship between the relevant facts and the feasible set. Put differently, if correct, my analysis presents a framework within which any further analysis of feasibility must operate.

1. MOTIVATING A POSSIBILITY-BASED ANALYSIS

Philosophers have offered little besides preliminary suggestions for an analysis of the concept of feasibility (e.g., Brennan and Pettit 2005; Brennan 2013; Cohen 2009; Cowen 2007; Räikkä 1998).² The only candidate analysis worked out in any detail is the *conditional probability account*, first suggested by Brennan and Southwood (2007) and developed more fully by Lawford-Smith (2013). Since Brennan and Southwood do not discuss their proposal in much detail, I focus here on Lawford-Smith's analysis (cf. Gilabert and Lawford-Smith 2012; Gilabert 2012*b*).

Lawford-Smith starts with three intuitive parameters for her analysis. The first is that the concept of feasibility should be fit to rule out as invalid any normative theory that makes infeasible demands, much as the familiar “ought implies can” proviso is meant to do (2013: 245). The second parameter is that feasibility is more restrictive than broad notions of possibility (such as logical or metaphysical possibility) (2013: 256). Our notion of feasibility must be sensitive to (among other things) economic, political, technological, and cultural constraints. But, third, our notion of feasibility should not be unduly sensitive to these latter constraints, otherwise it will “rule out oughts that shouldn't be ruled out” (2013: 254).

To accommodate these intuitions, Lawford-Smith offers a bifurcated analysis. First, she defines a binary notion of feasibility that is insensitive to economic, political, cultural (and so on) constraints to play the role of ruling theories out: realizing a state of affairs is feasible if and only if there exists an agent (individual or collective) with a nonzero probability of realizing the state given that it attempts an action that is not “logically, conceptually, metaphysically and nomologically impossible” (2013: 252–253).³ Put simply, realizing a state of affairs is feasible if and only if it is possible (broadly understood) for an agent to bring it about. However, this weak notion of feasibility is “unlikely to do much work” in normative political analysis because it fails to accommodate the second intuition (2013: 252). On Lawford-Smith's view, “soft constraints” (such as: the limits of human ability, resource availability, technological limitations, institutional constraints, and cultural factors) are ill-suited to rule out normative theories because they are “malleable” or

² One common (though unscrutinized) suggestion is that stability is a necessary condition for feasibility: realizing a state of affairs is feasible only if it is *stably* realized (e.g., Cohen 2009: 56f; Gilabert 2012*b*: sec. 4.3). This suggestion is mistaken. To wit, it is feasible to balance a spinning top on its point despite the manifest instability of that position (a nudge will send the top to its side). Stability is best conceived as a potentially desirable property of a target state of affairs (which is typically left implicit).

³ Lawford-Smith departs from Brennan and Southwood here; for the latter, realizing a state of affairs is feasible if and only if it is *sufficiently likely* given that the relevant agents try to realize it (Brennan and Southwood 2007: 9f).

“dynamic” limitations. Instead, soft constraints “make an outcome *less likely to obtain*” (2013: 254, original emphasis). She defines a continuous (or “scalar”) notion of feasibility accordingly: in view of the relevant “soft constraints”, realizing one state of affairs is more feasible than realizing another just in case the former is more likely than the latter to be realized by the relevant agent’s actions conditional upon the agent undertaking the relevant actions (2013: 255). This second notion of feasibility is meant to rank normative theories rather than exclude them from consideration.

Lawford-Smith’s analysis admirably enables us to negotiate the narrow road between “cynical realism” and “impotent idealism” (Gilabert and Lawford-Smith 2012: 815). But I think we can do better. For one thing, the conditional probability view treats the complete set of feasibility constraints in a disjointed way—logical, conceptual, metaphysical, and nomological constraints play an excluding role in the analysis, which is wholly distinct from the ranking role played by economic, political, technological, or institutional constraints. I don’t deny that “soft” constraints differ from “hard” constraints in their dynamic malleability, so there is a case to be made for treating them differently in some respects. Yet, it seems to me that there are hard limits to the economic, institutional, and cultural resources we can obtain (although these limits are difficult if not impossible to discern), thereby placing hard limits on what we can achieve with such resources. This is a feature such resources share with so-called “hard” constraints, one that our analysis of feasibility should account for.

Ordinary usage further suggests that the binary notion of feasibility is sensitive to (among others) economic, political, and technological constraints. Consider:

- Nigerian President Goodluck Jonathan’s statement that “2017 is not feasible for the take-off of the Continental Free Trade regime in Africa as the environment in the continent would not be conducive for its smooth operation” (Usigbe 2012);
- Former Indian civil servant Prodipto Ghosh’s claim that “it is not feasible [for India] to do anything [about climate change] at this stage [beyond voluntary reductions of greenhouse gases]. With the present state of technology development, we are likely to encounter severe constraints to our growth” (Ghosh 2012).

These are paradigmatic uses of the binary notion, which one would be hard pressed to deny are sensitive to economic, institutional, or technological limits. Certainly, they are not saying that an African Free Trade agreement or greater Indian efforts on climate change are “logically, conceptually, metaphysically and nomologically impossible”.

Finally, the conditional probability analysis appears to be less sensitive to motivational constraints than ordinary usage seems to be. While the motivational limitations of other

people are accounted for as “part of the context in which an agent acts”, the conditional probability analysis leaves aside the motivational limitations of the agent in question by assuming that the agent tries to perform an action that potentially conduces to the realization of the target state of affairs (“tries to realize” for brevity) (Lawford-Smith 2013: 256; cf. Estlund 2011). Lawford-Smith’s reason for assuming trying is plainly moral theoretic: “we don’t want to let agents off the moral hook. The fact that a person won’t do something isn’t enough for us to retract an imperative that she ought to” (2013: 254f). That is, our analysis of feasibility shouldn’t deliver a concept that rules out certain actions as obligatory simply because agents refuse to perform them, especially when they refuse for morally dubious reasons. But this reasoning presupposes that an adequate analysis of feasibility should render the concept useful for delimiting our moral obligations. I am sympathetic to this thought. But it is an open question whether feasibility is fit to play such a delimiting role (cf. Brennan and Southwood 2007; Brennan 2013). Accordingly, our conceptual analysis shouldn’t presuppose that it can.

Some ordinary feasibility claims seem duly attentive to at least *some* motivational constraints that prevent the agent in question from trying to realize the target state. The Ghosh quote above is a case in point, taking the Indian government’s reluctance to decrease economic growth for the environment’s sake as a feasibility constraint. Importantly, these motivational limitations need not be morally dubious. To wit, virtuous agents can be prevented from trying to realize desirable states of affairs due to *moral problems of assurance*—when the stakes are high and assurance of others’ cooperation is absent, morally motivated agents with a healthy dose of fear and risk aversion can lack the motivation required to realize a cooperative outcome (James 2012: ch. 4). The breakdown of international economic cooperation between World Wars I and II seems to be such a case. My point is that, by assuming trying, the conditional probability analysis excludes morally unimpeachable motivational limitations from among the full range of feasibility constraints, which seems in tension with ordinary use of the concept.

Nothing I have said here gives us conclusive reason to reject the conditional probability account. Instead, I take the preceding difficulties to form a case for developing a possibility-based alternative to Lawford-Smith’s probability-based account. My aim in the next two sections is to do just that: to rigorously develop a *restricted possibility account* of (binary) feasibility that is: (1) more restrictive than logical, conceptual, metaphysical, and nomological possibility; (2) treats all relevant types of constraints (logical and metaphysical, economic and technological) in a unified way; (3) is a (defeasible) candidate for a feasibility requirement that delimits our moral responsibilities. Once we have presented the restricted possibility alternative in some detail, we can then contrast our

options, seeing more clearly their comparative advantages and disadvantages. I leave the task of making a detailed comparison between the conditional probability and restricted possibility accounts for another time.⁴

2. PRODUCTION POSSIBILITY AND THE FEASIBILITY FRONTIER

We can derive an intuitive notion of feasibility from the economist's concept of a *production possibility frontier*. As the name suggests, a production possibility frontier delimits the set of commodity bundles (including goods and services) we can produce given our production functions (one for each commodity) and a schedule of constraints on the production inputs (labor, capital, raw material, and so forth). Every commodity bundle within the (closed) set defined by the possibility frontier is within the *feasible production set*. The basic idea to be carried forward is this: the feasible set is circumscribed by the set of bundles we⁵ can realize given the limitations on production processes and the available resources.

Although intuitive, the economist's notion of feasibility is too narrow for the purposes of political philosophy. Political philosophers are interested in social and political possibilities more generally, not simply economic production possibilities. Further, the set of relevant inputs extends beyond economic production inputs to include things like institutional capacity and political will. Fortunately, we can render this narrow concept applicable to these more general concerns by extending it to a more general notion of an all-purpose production possibility frontier.

Before I do this, let me fix ideas by briefly extending the discussion from section 1 about the general types of constraints any sensible feasibility assessment must consider. In general, feasibility assessments must attend to any fact that constrains our capacity to alter the status quo. We can make our assessment more tractable by grouping the

4 To be clear, Lawford-Smith (2013) takes her primary contribution to be showing that scalar feasibility is "more interesting" than binary feasibility (cf. Cowen 2007). Let me stress that my focus here is on the conceptual challenges faced by her analysis of *binary* feasibility. I aim to present an alternative analysis of binary feasibility that (unlike Lawford-Smith's view, by her own admission) is rich enough to have important implications for normative political theory. I conjecture that, once the details are in view, my analysis diminishes at least one motivation for analyzing scalar feasibility as Lawford-Smith does — namely, to enrich the theoretical interest of the conditional probability analysis by accounting for economic, technological, and institutional constraints (and so on). I set aside further discussion of a scalar notion of feasibility, as nothing in the remainder of the paper turns on it.

5 In the context of discussing the realization of states of affairs, "we" (and other collective pronouns) should be read as referring to the relevant set of agents, whoever they are. "We" is not meant to suggest a collective agent that may not, in fact, exist. Thanks to Holly Lawford-Smith for pointing out the potential confusion.

relevant facts into analytically useful general categories. Some straightforwardly rigid constraints are *logical consistency*, the *laws of nature*, and *human biology*. But we should also attend to less rigid, more malleable constraints. I only present a few salient examples here, leaving development of a full list as a practical exercise: *ability constraints*, which comprise facts about human abilities; *cognitive constraints*, which comprise facts about our cognitive capacity, including cognitive biases and computational limitations; *economic constraints*, which — if taken broadly — comprise facts about possible allocations of money, labor power, and time; *institutional constraints*, which include facts about institutional structure and capacity (for example, the number and distribution of veto points in a collective decision procedure and the ways in which political officials are selected); and, *technological constraints*, which include facts about the tools, techniques, and organizational schemes available for bringing about new states of affairs. I depart from the conditional probability account in also including *motivational constraints*, which identify the limits of what people can be motivated to do given intrinsic features of human agents that affect motivation (including affective biases, prejudices, and fears), as well as the extrinsic features of an agent's environment that interface with her intrinsic motivational capacities (including social norms and incentives).⁶ These categories are meant to group together broadly related facts about the means at our disposal for altering the status quo. In other words, they identify limits on the production inputs required to realize possibilities.

We can use this list of general types of constraints to illuminate the notion of an all-purpose production possibility frontier. Generally, the limits of our present technological, motivational, institutional, financial (and so on) resources define a multidimensional constraint set, which determines the total means available for realizing states of affairs. For example, motivational constraints are defined by (the limits of) our motivational resources, namely, human agents' capacity to be motivated to behave in certain ways given their affective biases, prejudices, fears, and the like, as well as the social norms and incentives to which they are subject. Our *current total stock of all-purpose resources* is composed of the technological, institutional, motivational (and so on) means we have on hand. Additionally, we acknowledge that these different types of resources can be more or less fungible; for example, we can convert some portion of our existing technological and institutional resources into a greater stock of motivational resources

⁶ Note that these motivational constraints are not taken to be so tight as to exclude all possible options but those actually chosen. Rather, they are meant to exclude options that people find themselves motivationally incapable of choosing, leaving open a range of possibilities besides those actually chosen. Thanks to Geoff Brennan for drawing my attention to this point.

by using technological and institutional means to cultivate desirable dispositions. But we should acknowledge that using these technological and institutional resources in any particular way has opportunity costs—they cannot be expended elsewhere to realize competing ends. Given the limited fungibility of these different types of resources, we can change the composition of our total resource stock through a (limited) series of conversions. We can even grow our current stock of resources by “investing” some portion of it. For example, we might presently invest some of our current motivational and institutional “capital” to reform our current institutions in ways that improve their capacity to help us overcome future collective action problems. So we need not treat the size of our total stock of resources as fixed over time. Let’s say that an *attainable total stock of all-purpose resources* is a configuration of (institutional, cultural, financial) resources that could emerge from a transformation of our current resource stock by a series of conversions or investments that is available to us (more on this below).

Note that our means for altering the status quo are finite—which is to say that we do not have unlimited capacity for realizing possible states of affairs. Moreover, each type of resource has a limited number of uses. Money or technology alone is not enough to realize many target states of affairs; we must also be motivated to use the available money or technology effectively. (Nor, it should be said, is motivation enough; we also require certain financial or institutional means to get the job done.) There are also limits on the ways in which we can convert one type of resource into another. For example, we can’t exchange all of our technological assets for improved institutional capacity, but we can use some combination of technology, political will, and cultural resources to reform our institutions. Finally, the opportunity costs associated with using our resources in any particular way sets limits on the allocations that are possible. All these limitations on our allocation of the resources at our disposal restricts the states of affairs we can realize.

The foregoing illuminates two senses in which the feasible set is circumscribed by constraints on all-purpose resource availability. The world sets the resource requirements for realizing our objectives (e.g., how much steel and concrete we need to build a stable bridge under various conditions). This restricts the possible resource allocations that can effectively realize our objectives. Further, the world—in particular, the composition of our current resource stock and the permitted conversions—also restricts which possible resource allocations can be attained (e.g., how much steel and concrete we can acquire). In both these ways, the world places limits on our capacity to alter the status quo.

This notion of an all-purpose production possibility frontier highlights the fact that our analysis of the feasible set must attend to the relevant constraints *together*. We are, after all, limited to realizing those states of affairs whose input requirements can be *jointly* satisfied. Suppose realizing a target state of affairs requires that we have 12

units of institutional capacity, 20 units of motivational capacity, 5 units of technological assets, and \$15 million.⁷ Suppose further that we can fulfill each of these requirements separately given our current stock of resources. That is, we can transform our current stock into one that includes the required institutional capacity, or into one that includes the required motivational capacity, and so on. None of this implies that we can *jointly* satisfy these demands. The limits on the ways in which we can convert one type of resource into another and the opportunity costs associated with any particular allocation might prevent us from doing so. The significance of this point is heightened once we acknowledge that our attempts to realize distinct moral values (e.g., freedom and equality) can place competing demands on our resources. Accordingly, the feasible set includes only those states of affairs whose production input requirements can be *jointly* satisfied given our current stock of resources.

We are now in a position to analyze this general notion of feasibility more precisely. Start with this: realizing a state of affairs is *feasible* only if we possess the resources required to realize it — that is, realizing the target state has production input requirements that can be met by our current stock of all-purpose resources. This is not quite right, since our analysis should make explicit the fact that the size and composition of our total resource stock can change and grow. Recalling our notion of an attainable resource stock, we should say that realizing a target state of affairs is *feasible* only if there is an attainable resource stock that enables us to realize it. (I deal with the remaining modals — *attainable* and *enable* — in section 3.) Colloquially, feasible states of affairs have production input demands that we can satisfy (now or in the future) given the all-purpose resources available to us; infeasible states are “too expensive” given the set of all-purpose resource stocks that are attainable.

It is worth pausing to note that this analysis of feasibility is void of any moral content. Accordingly, feasibility assessments do not incorporate our judgments about which states of affairs are *worth* realizing from a moral standpoint. This might seem counterintuitive (cf. Buchanan 2004: 61; Räikkä 1998). But if we conceive of normative (i.e., action-guiding) political analysis as “a confrontation of the feasible with the [morally] desirable” (Brennan 2007: 119), then this is the right way to go. One might be skeptical here that we can separate our feasibility assessments from our judgments about which possibilities are morally desirable. The same set of facts can be relevant for both feasibility and desirability judgments after all. This is most obvious when we consider causal history. If greater socioeconomic equality is feasible but only by authoritarian means, it thereby becomes

⁷ I use the idealization of precisely quantifiable institutional, motivational, etc., capacity solely for ease of exposition.

less desirable.⁸ Nothing I have said here is meant to deny this dual relevance of facts. I only mean to assert that a single set of facts can enter our feasibility and desirability analyses to different effect. If the relevant facts are such that greater equality arises only by authoritarian means, then greater equality may be feasible but not desirable. A non-moralized analysis of feasibility gives us a clear view of the role played by particular facts in our feasibility assessments and permits us to see clearly at which points in our normative analysis our feasibility judgments, as opposed to our desirability judgments, are doing the work. A moralized notion of feasibility precludes this analytic separation.

3. MODELING THE FEASIBILITY FRONTIER

The analysis presented in the last section, although intuitive, is still too imprecise to give us a firm grip on the proposed account's wider implications. In this section, I refine the analysis by developing a fairly rigorous model of the all-purpose production possibility frontier ("feasibility frontier" hereafter) using a possible worlds framework. This modeling exercise is important for several reasons. First, it provides an analysis of the key modal notions (*attainable*, *enable*) left unanalyzed in the last section. Second, a rigorous model clarifies the units of analysis and the variables to which our feasibility assessments are sensitive. Relatedly, the model specifies the ways in which resources, states of affairs, agents, and times enter the analysis. Finally, the model facilitates clear reasoning about the implications of the analysis for normative political theory, as I show in section 4.

On the standard approach to modeling modal relations, possibility expresses existential quantification over a contextually restricted set of possible worlds (Kratzer 1991; cf. Lewis 1973; Stalnaker 1968). I propose that we model the feasibility frontier similarly (whence the moniker "restricted possibility account"). The task in this section is to develop this proposal.⁹

We start with the set of all possible worlds, each of which represents a complete way the world might be. Following the analysis in the last section, we are here interested in how worlds differ with respect to four properties: (1) the total resource stock; (2) the processes for altering the composition of our resource stock, including the conversion rates (simply, "conversion processes"); (3) the causal processes by which we use our resource stock to realize social and political states of affairs (simply, "causal processes"); and (4) the social and political states of affairs of analytical interest. These are the properties of worlds to

⁸ Thanks to Geoff Brennan for raising this point.

⁹ I introduce the model informally in the main text. A more precise formulation is given in a formal appendix.

which our feasibility analysis is particularly sensitive. Accordingly, we represent each possible world as a summary description of these properties.¹⁰ The social and political states of affairs that obtain at possible worlds play a particularly important role: we are interested in assessing the feasibility of realizing these properties at the actual world. To avoid confusion, states of affairs are less comprehensive than possible worlds. Here, a *state of affairs* is represented by a description of a social or political feature of a world that can serve as a sensible object of political reform. Examples include: the constitutive features of a country's legislature; the level and distribution of public goods provision; or the level of ethnic violence in a region. Recall that realizing a state of affairs is feasible if and only if there exists an attainable resource stock that enables us to realize that state. Put in terms of the model developed here, we say (roughly) that realizing that state is feasible only if it obtains at a world that is consistent with the total resource stock, the conversion processes, the causal processes, and the state(s) of affairs that obtain at the actual world at the time of evaluation. Equivalently, a state is feasible only if it obtains at a world that represents a possible future trajectory from the state of the actual world at the time of evaluation.¹¹

Put this way, we can readily see that the concept of feasibility as analyzed in the last section can be modeled as a binary accessibility relation on possible worlds. On the standard view, a world w' is accessible from w (if and) only if certain specified facts that obtain at w do not rule out w' (Kratzer 1991). My analysis of feasibility implies that a world is a member of the feasible set only if it is compatible with the facts pertaining to the composition of the total resource stock, the conversion processes, the causal processes, and the state(s) of affairs that obtain at the actual world at the time of evaluation. In the

10 It is worth putting this precisely here to avoid confusion hereafter. We treat worlds at a time as a quadruple: $w^t = \langle r, v, c, s \rangle^{w,t}$, where r denotes the particular resource stock that obtains at w , v represents the conversion processes at w , c represents the particular causal processes at w , and s represents the states of affairs of interest at w . A complete world-history is thus represented as a sequence of quadruples, one quadruple for each time t .

11 Note that states of affairs are not the objects of direct choice even though they are the objects of our feasibility analyses. As a matter of choice, we focus on actions or policies, which yield probability distributions over states of affairs. (Thanks to Geoff Brennan for raising this point.) This is a place where probability and feasibility come apart on my account: realizing a state of affairs is feasible just in case it is a possible outcome of (i.e., has positive probability given) an action or policy that is consistent with the total resource stock, the conversion processes, the causal processes, and the state(s) of affairs that obtain at the actual world at the time of evaluation. (In contrast, Lawford-Smith says that realizing a state of affairs is feasible just in case it has positive probability given an action or policy that is logically, conceptually, metaphysically, and nomologically possible.) Given this analysis, probability distributions are relevant for choosing some action or policy (because we care about *expected* outcomes) but are irrelevant for assessing the feasibility of realizing a state of affairs. Hence, the complications introduced by probability distributions can be safely left aside here.

standard terminology, the feasible set is thus delimited by a *circumstantial modal base*: it includes only those possible worlds that are consistent with the salient circumstances at the world of evaluation. My point here is to note that we can draw on a familiar apparatus to model the feasibility frontier — the outer bound of the feasible set — and refine the analysis from section 2.¹² Using this apparatus, we can say that a possible world is a member of the feasible set only if that world is circumstantially accessible from the actual world, in view of the facts that constitute the feasibility modal base. It follows that realizing a state is feasible only if there is at least one world at which the state is realized that is circumstantially accessible from the actual world; realizing the state is otherwise infeasible.

We should unpack this circumstantial accessibility relation in greater detail. Departing from convention, we evaluate accessibility in two steps here. This is to highlight the fact that all-purpose resources are (potentially) spent at two points leading up the realization of the target state. Resources are most clearly spent to realize target states of affairs: we must spend technological, institutional, and motivational capital to realize a new form of government or implement a scheme that relieves global poverty. Perhaps less obviously, resources are spent when we undertake the conversions and investments required to make any necessary changes to the composition of our current resource stock. Put simply, we must use current resources to “buy” a new resource stock if our current resource stock does not meet the input requirements of realizing our target state of affairs. Hence, we can think of our accessibility relation on worlds as embedding two accessibility relations.¹³

The first pertains to attainable resource stocks: a resource stock is *attainable* if and only if it is a property of a world that is consistent with the resource stock and conversion processes that obtain at the actual world at the time of evaluation. This first step yields the set of worlds with attainable resource stocks. The second embedded accessibility relation pertains to those states of affairs that can be realized by attainable resource stocks. However, notice that the states of affairs that can be realized with a resource stock by one causal process might differ from those that can be realized by an alternative process. Thus, a particular resource stock might enable agents to realize different states of affairs at two different worlds, since the sets of causal processes that obtain at these worlds might differ. At this second stage, then, we must narrow the set of worlds with attainable resource stocks to the set of worlds with causal processes that are consistent with the actual world

¹² “Outer bound” because, as I noted at the outset, I only present a necessary condition for feasibility here.

¹³ Although what follows is precise enough for now, this is a part of the discussion where particularly important complicating factors arising from temporal issues had to be left aside in the interest of tractability. Please see the appendix for discussion of these complications.

at the time of evaluation. This restriction yields what we might call the *constraint set*, the set of worlds such that, for each member world, the resource stock, conversion processes, and causal processes that obtain at the world at the time of evaluation are consistent with the resource stock, conversions processes, causal processes, and states of affairs that obtain at the actual world at the time of evaluation. This constraint set fixes an outer bound for the set of feasible worlds. But we are ultimately interested in the feasibility of realizing states of affairs, which are properties of worlds. We now say that realizing a state of affairs is feasible only if it obtains at some world in the constraint set. This entails the analysis offered earlier: realizing a state of affairs is feasible only if it is a property of a world that is consistent with, or accessible given, the salient properties of the actual world at the time of evaluation.

Our ordinary notion of feasibility accommodates both agent- and time-relative feasibility judgments — realizing a state of affairs is said to be feasible *for* a particular agent *at* or *before* some particular time (Gilabert and Lawford-Smith 2012; Lawford-Smith 2013). It is thus important to indicate how the model I’ve developed can index feasibility judgments to agents and times. Let’s start with times. Recall that the basic constituents of our model are possible worlds. We represent worlds as a description of their salient properties: total resource stock; conversion processes; causal processes; and states of affairs. What we left implicit earlier for the sake of simplicity (but see footnote 10) is that we represent worlds as sequences of these descriptions, one for each time. So the worlds in our model are constructed from *time-slice descriptions*, which come together to form a complete world-history. This allows us to index feasibility judgments to times quite easily: realizing a state of affairs is feasible at or before a particular time t only if it obtains at a world in the constraint set at or prior to t , which is just to say that it is a property of a time-slice $t' \leq t$ that is consistent with the salient properties of the actual world at the time of evaluation.¹⁴

Accommodating agent-relative claims is no more difficult. This becomes apparent once we recognize that the causal processes we are interested in — the processes by which our resource stock is used to realize social and political states of affairs — are

14 Complications arising from time indexing raise numerous interesting issues, which are left aside here for simplicity. But one issue in particular is worth noting. Let t denote the time of evaluation. My analysis implies that the only states of affairs that are feasible at t are those that are realized at the actual world at t . But this should not be taken to deny that other counterfactual states of affairs could have been realized at t (and, so, were feasible prior to t) had different actions been taken at some time t' earlier than t . This point is worth noting because such counterfactual judgments might be used to ground certain kinds of moral judgments; for instance, that some agent is currently responsible for perpetuating harm because she was involved in bringing about an institutional scheme that is less just than a counterfactually feasible scheme (cf. Pogge 2008). (Thanks to an anonymous reviewer for raising this point.)

social causal processes. They are causal processes that are constituted by *the actions of agents*. This isn't to deny the (perhaps critical) involvement of impersonal environmental forces (for example, when we harness hydro or wind power to provide energy for human enterprise). But the states of affairs in which we are interested here are not realized by wholly impersonal causal processes. The causal processes by which social and political states are realized may be simple (they involve a single agent acting in isolation) or complex (they aggregate the actions of several agents). The key point is that they centrally involve agents. Consequently, we can say that realizing a state of affairs is *feasible for* an agent, *a*, only if it obtains at a world in the constraint set as a result of a causal process that involves actions taken by *a*. Realizing a state of affairs is feasible only if there exists an agent for whom realizing that state is feasible.

I've aimed to refine the analysis presented in section 2 by modeling the feasibility frontier—the outer bound of the feasible set—with a high degree of rigor. But we should not expect the model to deliver the precise outputs obtained from economists' models of the production possibility frontier. In fact, given the model, I'm skeptical that we estimate the feasibility frontier with any confidence whatsoever, especially if we are looking beyond the medium-term (this point will become important in the next section). This estimation difficulty isn't due to any imprecision in the model, though. It is because our estimates of the required inputs—of the required resource conversions, of the states of affairs causally realized by distinct causal processes, of the relative costliness of particular stock transformations—are bound to be radically imprecise. There's no way around this. Imprecision is an inevitable part of counterfactual comparisons—especially those involving distant counterfactuals—which are integral to feasibility analysis. We'll take up the implications of these estimation difficulties below. The model presented here is nonetheless helpful despite the noted problem. By indicating the salient variables and clarifying the role that these variables play in our feasibility assessments, the model brings much needed discipline to philosophers' feasibility talk. This is more than enough to serve as a basis for reasoning about the relationship between feasibility and moral considerations.

4. IMPLICATIONS FOR NORMATIVE ANALYSIS

I now show that the restricted possibility account of feasibility has significant implications for normative political theorizing. On the conventional Target View of normative political philosophy, a political ideal presents a more or less well-defined target for real world political reform and, thus, serves as a reference point for our specification of normative political principles. In this section, I use the model developed above to show that

proponents of the Target View are unjustified in asserting that political ideals present appropriate long-term goals for morally progressive reform efforts. I conclude by suggesting that normative political theory can nonetheless provide normative guidance for morally progressive reform by analyzing actual social failures rather than political ideals.

One familiar aim of political philosophy is to enumerate evaluative principles; understood functionally, these are principles that rank possible worlds according to some standard of moral desirability.¹⁵ Evaluative principles give an interpretation of our basic evaluative criteria (e.g., well-being, freedom, equality, security, community) and establish interconnections between these criteria, specifying their relative importance and the respects in which one criterion's (e.g., welfare) contribution to overall moral desirability is conditioned by other criteria (e.g., freedom or equality) (cf. Hamlin and Stemplowska 2012). Yet political philosophers are generally not satisfied with providing a moral ranking of possible worlds; they often have the practical aim of saying something about what we should *do*. They aspire, in other words, to specify general moral directives for action in view of our evaluative judgments. These directive principles are meant to serve a deontic function — they demarcate the lines between obligatory, permissible, and impermissible actions or institutional arrangements (for instance, by enumerating a general schema of rights and obligations). So understood, directive principles advise (perhaps command¹⁶) us to realize particular (sets of) possible worlds — namely, those worlds that are consistent with the content of the principle.

To illustrate the distinction drawn here, take a principle of strict distributive equality. Understood evaluatively, this principle ranks worlds that are strictly equal as morally better than worlds that are not strictly equal; understood directly, this principle advises us to realize a world at which strict equality is satisfied. Of course, our political theories are generally more complex than this, so the exercise of ranking or advising will require us to balance a number of competing criteria. The key point here is to distinguish two ways in which we might understand a single principle.

Directive principles are typically thought to be subject to a feasibility requirement whereas evaluative principles are not. Assume this is correct for the sake of argument (but

15 I do not mean to commit here to any view about the relative priority of “the good” and “the right”, or of the evaluative and the deontic more generally. All I need is that there be a moral standard according to which we evaluate ways the world might be. If that standard is supplied by some account of basic value, then the ranking will be in terms of “overall goodness” or some such; if that standard is supplied by basic deontic facts, then the ranking will be in terms of “overall rightness” or “overall consistency with the deontic ideal” or some such. “Moral desirability” is meant to be neutral between these.

16 For lack of a better term, I stick with “advise” hereafter to remain neutral on the stringency of directive principles.

see Estlund 2014; Gheaus 2013). Then a normative theory is plausible only if its principles for action do not require agents to realize infeasible ends. Accordingly, philosophers widely concede that directive principles must be sensitive to empirical facts in some sense.¹⁷ But how sensitive, and in what way?

On the *Target View*, our normative theorizing starts by characterizing a political ideal. Here, a “political ideal” more or less defines the constitutive features of an ideally just world — the highest ranked world according to our evaluative criteria (cf. Gilibert 2012b: 123).¹⁸ Political ideals are often understood in directive terms: an ideal provides “a ‘compass’ for action... a regulative end at which our actions should aim” (Valentini 2012b: 38). In other words, an ideal is typically thought to identify a reference point for normative theorizing by characterizing the states of affairs we should aim to realize (perhaps incrementally) or by delivering a set of normative principles that we should aim to satisfy, even if only approximately. This familiar picture is sustained by the intuitive plausibility of a handful of simple but vague metaphors — the political ideal as a compass or a map. (Proponents of the Target View so understood are legion; see, e.g., Buchanan 2004; Gilibert 2012b; Robeyns 2008; Simmons 2010; Valentini 2012b.)

To facilitate critical scrutiny, we can differentiate two versions of the Target View, neither of which entails the other (although these are rarely distinguished by proponents of the view). First, we might say that a political ideal is supposed to provide a general description of the states of affairs we should aim to realize. On this specification, a political ideal is represented by the ideal world — the highest ranked world according to our evaluative standard — and we should aim to realize (by incremental steps) the states of affairs that obtain at the ideal world. Accordingly, directive principles are specified with this aim in view.¹⁹ Second, we might say that a political ideal is meant to enumerate the directive principles we should implement to regulate our political affairs. On this specification, a political ideal is represented by the set of directive principles that govern interpersonal transactions and institutional design at an ideal world and we should strive to (eventually) satisfy those principles within our own context. I discuss each of these

17 Cohen (2003) might be an exception here, although it seems clear to me (and others, e.g., Valentini 2012b: 30) that Cohen’s “fact-free principles” are not directive but evaluative principles, as I distinguish them. Thus, in my terms, Cohen’s view is that evaluative principles must be fact-insensitive and that directive principles must be grounded in evaluative principles.

18 Perhaps a theory of *justice* does not attend to all our evaluative criteria, but only to a proper subset of them (Rawls 1999: 8f; cf. Cohen 2003: 244f). Then an ideal of a *just* society is sensitive only to those evaluative criteria that are relevant for realizing justice.

19 Perhaps the ideal is represented by a set of possible worlds, all of which attain the highest ranking according to our moral standard. As nothing I say turns on this complication, I set it aside.

interpretations in turn.²⁰

According to the first specification, a political ideal is represented by the highest ranked world according to our standard of moral desirability.²¹ An ideal provides a target for normative theorizing in the sense that our reform efforts should aim to realize the states of affairs that obtain at the ideal world — the “ideal states”.²² Recalling the feasibility requirement asserted earlier, I assume (for the sake of argument) that we have a standing obligation to realize the morally best, or optimal, feasible world. Hence, a normative theory can plausibly require us to realize the ideal states only if the ideal states obtain at the optimal feasible world. By definition, the optimal feasible world is in the constraint set. Given our model of the feasibility frontier, we know that all states of affairs that obtain at the optimal feasible world arise from a sequence that starts with the initial conditions that obtain at the actual world. (“Initial conditions” refers to the resource stock, conversion processes, causal processes, and states of affairs that obtain at a world at the time of evaluation).

Suppose for a moment that the ideal world is not in the constraint set. It follows that the initial conditions at the ideal world do not match the initial conditions at the actual world; thus, they do not match the initial conditions at the optimal feasible world. Put differently, the sequence that leads to the realization of the ideal states has a different starting point than the sequence that leads to the realization of the optimal feasible states. Given this, we are justified in believing that the ideal states obtain at the optimal feasible world only if we are justified in believing that both sequences, with their disparate sets of initial conditions, are such that they eventually converge on the ideal states. Thus, if the ideal world is not in the constraint set, we ought to aim to realize the ideal states only if we are justified in believing that the sequence by which the ideal states are realized converges with the sequence by which the optimal feasible states are realized, despite their divergent starting points. This is a controversial assumption; it is certainly unwarranted without a

20 Instead of adopting either version of the Target View, someone might endorse a third view of the role of political ideals in normative theorizing: namely, that an ideal itself is not the ultimate target for political reform but is meant to provide an evaluative standard by which we rank feasible alternatives (see, e.g., Gilabert 2012a: 45f, Jubb 2012; Sangiovanni 2008; Stemplowska 2008). Call this the *Benchmark View*. On this view, a political ideal does not provide straightforward guidance for specifying directive principles; instead, it helps us figure out which among our feasible alternatives we should aim to realize. As this view raises issues that are quite distinct from those raised by the Target View, I deal with it in Wiens (N.d.).

21 A more precise formulation of the following line of reasoning is presented in section 6.2 of the appendix.

22 Presumably, on this view, the normative principles with which we ought to comply are those that would, given the expected level of compliance, realize the ideal states. Note that these principles might differ from the normative principles with which people generally comply at the ideal world, given the differences in external circumstances between the ideal world and the actual world. More on this in a moment.

fairly thorough examination of both the ideal world and the optimal feasible world.

Now suppose the ideal world is in the constraint set. Then it follows that the optimal feasible world is identical to the ideal world. It also follows that the sequence that culminates in the realization of the optimal feasible states leads to the realization of the ideal states. However, we are justified in believing that we should realize the ideal states only if we have shown, at least, that the ideal world is not likely to be excluded from the constraint set. In practice, we can't simply assume that the ideal world is the optimal feasible world; this is an unwarranted assumption without a fairly thorough feasibility assessment of the ideal world.

Thus, whether the ideal world is within the feasibility frontier or not, the first specification of the Target View is unjustified in asserting that we should aim to realize the political ideal without an extensive exploration of the limits of the feasibility frontier. Proponents of the Target View can justify their assertion only by conducting a thorough feasibility assessment of the ideal world prior to specifying the states of affairs we ought to realize. Keep this point in mind; we'll return to it in a moment.

According to the second specification, a political ideal is represented by a set of "ideal principles", the directive principles that regulate interpersonal transactions and institutional design at an ideal world.²³ An ideal provides a target for normative theorizing in the sense that we should strive to (eventually) implement ideal principles within our own context. Given the feasibility requirement, we have an obligation to realize the optimal feasible world. Hence, we should implement ideal principles only if general compliance with them eventually leads to the realization of the optimal feasible world.

It is worth highlighting at this point that implementing a directive principle can yield different outcomes depending on the circumstances in which that principle is implemented. As a simple illustration, suppose we regulate our affairs by a distributive principle requiring strict material equality. If people are generally motivated such that they do not require differential material rewards as an incentive to maximize their productive contribution, then general compliance with this principle should have little effect on total social production (and, thus, little effect on total welfare, among other things). If instead people are generally motivated to maximize their productive contribution only by a scheme of differential rewards, then general compliance with this principle will keep total social production (and, thus, total welfare) lower than it might otherwise be. So a single principle can yield differing levels of welfare or freedom depending on the external circumstances. Generally, this is because compliance with directive principles, in conjunction with other factors, shapes (at least) the ends people pursue and the means

23 A more precise formulation of the following line of reasoning is presented in section 6.2 of the appendix.

they deploy in their pursuits. In terms of the model, we treat directive principles as a variable that, assuming general compliance, influences the way people use their stock of all-purpose resources and shapes the causal processes that operate at a world. In this light, we should implement ideal principles only if general compliance with them yields causal processes that eventually realize the optimal feasible world.

Suppose for a moment that the ideal world is not in the constraint set. It follows that the initial conditions at the ideal world differ from the initial conditions at the optimal feasible world. (Recall that “initial conditions” refers to the resource stock, conversion processes, causal processes, and states of affairs that obtain at a world at the time of evaluation.) By assumption, general compliance with the ideal principles obtains at some time at the ideal world and, consequently, yields causal processes that realize the ideal states. Given this, we are justified in believing that general compliance with the ideal principles realizes the optimal feasible states only if we are justified in believing that general compliance with the ideal principles is (i) part of a sequence that yields the ideal states and (ii) part of a sequence that yields the optimal feasible states, despite the two sequences’ distinct initial conditions. Thus, if the ideal world is not in the constraint set, we ought to implement the ideal principles only if we assume (as above) that, despite distinct starting and end points, the two sequences — the one that yields the ideal states and the one that yields the optimal feasible states — converge somewhere along the way. This time, we must assume convergence on the same causal processes as shaped by general compliance with ideal principles. This, too, is a controversial assumption, one that is surely unwarranted without a fairly thorough examination of both the ideal world and the optimal feasible world.

Now suppose the ideal world is in the constraint set. Then it follows that the optimal feasible world is identical to the ideal world. It also follows that general compliance with the ideal principles culminates in the realization of the optimal feasible states. However, we are justified in believing that we should implement the ideal principles only if we have shown, at least, that the ideal world is not likely to be excluded from the constraint set. As above, assuming that the ideal world is the optimal feasible world is unwarranted without first conducting a fairly thorough feasibility assessment of the ideal world. Thus, whether the ideal world is feasible or not, the second specification of the Target View is also unjustified in asserting that we should aim to realize the political ideal without extensive exploration of the feasibility frontier. As previously, proponents of the Target View can justify their assertion only by conducting a thorough feasibility assessment of the ideal world prior to specifying the directive principles with which we ought to comply.

What have we learned from the preceding discussion? On both specifications, a

political ideal identifies a more or less well-defined target for political reform. However, on both views, our adoption of an ideal as our reform target is unjustified in the absence of a thorough assessment of the feasibility frontier. If normative theories are subject to a feasibility requirement,²⁴ then we are not justified in adopting the ideal as a target for real world political reform simply because it is shown to be morally best. We must also show that the ideal is unlikely to be infeasible.

This point might seem rather banal. Indeed, philosophers are increasingly attuned to feasibility issues in the course of arguing for their proposed ideals (e.g., Gilibert 2012*b*; James 2012; Ypi 2012). The first thing to note in reply is that philosophers have generally fallen far short in their attempts to analyze the feasibility of their proposed ideals. I'm not merely saying that philosophers haven't offered enough evidence to warrant belief that their ideals are feasible; the point is not just about the burden of proof philosophers bear. Rather, philosophers typically fail to offer the *right kind* of evidence to warrant such belief. An adequate feasibility analysis must at least investigate the causal mechanisms that generate the status quo and the ways in which the causal mechanisms required to realize the ideal are likely to interface with the status quo (Wiens 2013). Yet philosophers frequently satisfy themselves with optimistic speculations about the causal processes that *could* support their proposed ideal at some counterfactual world, typically assuming key changes to core features of the status quo at the actual world — assuming, for instance, that political leaders are more morally motivated than they seem to be. Few philosophers (at least in the global justice literature with which I am familiar) have credibly engaged with social scientific explanations of the status quo or shown how their proposed ideal could emerge by a process that is consistent with the salient features of the status quo.²⁵ So we are not *yet* justified in adopting many extant ideals as targets for reform.

This shortcoming seems readily fixed though — we simply do more thorough feasibility analyses before endorsing any particular ideal as a target for reform. Here's where the restricted possibility analysis bites. The model presented in section 3 highlights the extreme difficulties of estimating the feasibility frontier with any confidence, especially if we are looking beyond the medium-term. The problem has two sources: the sheer complexity of the calculations required in any feasibility assessment and the fact that political ideals typically constitute fundamental (perhaps revolutionary) departures from the status quo. (I suppose some might propose more limited political ideals, but the general tendency among philosophers is to propose deep and far-reaching changes as the ultimate aim.)

24 Which I have assumed for the moment for the sake of argument. I consider the possibility of dropping this assumption below.

25 Amartya Sen (e.g., 1981; 1999) is an obvious exception.

First, given the number of variables to which our feasibility assessments must be sensitive, the complexity of their interactions, and the potential for path-dependence,²⁶ determining whether any particular long-range objective is feasible is beyond human cognitive capacity. Determining whether limited reforms to the status quo are feasible in the short- to medium-term might be tractable, since the limited time horizon and the limited departure from the status quo reduce the number of variables in play and the range of possibilities to be considered. But — and this is the second point — political ideals are rarely limited in the relevant ways. Philosophers' ideals typically represent social and political arrangements that are fundamentally (perhaps radically) distinct from the status quo as a long-term aim.²⁷ To think that we can estimate with any confidence whether profound and far-reaching changes to our current social and political arrangements are feasible strains credulity (cf. James 2012: 116). The upshot of these epistemic limitations is that, given the restricted possibility account, we simply cannot determine with any confidence whether particular long-range objectives are feasible, let alone with sufficient confidence to justify adopting a political ideal as a reform target.

This uncertainty might not be worrisome if we can reasonably expect that our efforts to realize an infeasible ideal would get us close to the morally optimal feasible world. Indeed, the standard reply to the claim that an ideal might be infeasible is to assert that we should still aim to approximate it as far as possible (e.g., Gilabert 2012*b*: 243). This would be a compelling reply if we had any reason to expect that the optimal feasible world were a fair approximation of the ideal world. But the arguments above imply that we have no such reason, certainly not without identifying the optimal feasible world and thoroughly examining the states of affairs it eventually realizes.²⁸ Given the epistemic limitations noted above — and given that we are a long way from the optimal feasible world — identifying the optimal feasible world is no less difficult than determining whether any particular ideal is feasible. So it is not just an open question whether we are duty-bound to aim at the ideal, but an *unresolvable* question. Worse, Lipsey and Lancaster's (1956) "general theory of second best" gives us a strong reason to think that the optimal feasible world is unlikely to approximate an infeasible ideal. They demonstrate that, if one of our ideal directive principles can't be satisfied, then the optimal feasible world does not necessarily satisfy the remaining ideal principles. So trying to approximate an infeasible ideal carries a substantial risk of leading us away from the optimal feasible world and toward morally

²⁶ "Path-dependence" refers to the phenomenon of future option sets being strongly determined by present choices to the extent that some choices can irrevocably close off certain future possibilities.

²⁷ Cf. Gilabert's (2012*b*: 241) comment about a five hundred year or longer time horizon.

²⁸ I make the argument for this skeptical conclusion explicit in section 6.3 of the appendix.

inferior states of affairs.²⁹ Without any reasonable expectation that aiming at a political ideal will get us close to the optimal feasible states of affairs, we should abandon the Target View.

5. HOW TO AVOID PESSIMISM

We can respond to this pessimistic conclusion in one of three ways. Two responses reject key premises of the argument in the last section: first, that normative political theories face a feasibility requirement; second, the restricted possibility account of feasibility. I briefly consider these two responses before turning to a third, which accepts the conclusion and proposes a strategy for overcoming pessimism about the prospects for action-guiding political philosophy.

The fact that we cannot determine whether political ideals are feasible poses a problem only if normative theories are required to prescribe feasible reform targets. If this feasibility requirement is mistaken, we need not estimate whether the reform target represented by some political ideal is feasible. At this point, one might say “So much worse for the feasibility requirement!” (e.g., Estlund 2014; Gheaus 2013). I don’t wish to deny this response, since, in section 1, I left it an open question whether the concept of feasibility is suited to rule out normative theories along the lines of an “ought implies can” proviso. Instead, I simply note that rejecting the feasibility requirement will require a compelling explanation (vindicating or debunking) for the apparent validity of certain familiar contrapositive inferences, like the following: “Strict egalitarian socialism is infeasible; therefore, we are not duty-bound to realize a strict egalitarian socialist regime.” Inferences such as these are naturally explained by some sort of feasibility requirement. Given this requirement’s wide acceptance, I submit that the burden of proof rests with those who wish to deny all relevant versions of the feasibility requirement.³⁰

29 Also see Brennan’s (2013) instructive discussion about the need to balance the likelihood of morally disastrous outcomes arising from a particular series of reform policies that may be recommended in pursuit of (optimistically) anticipated desirable outcomes.

30 A related move one might make is to claim that, strictly speaking, we don’t have duties to realize infeasible reform targets, but that we have “transitional” or “dynamic” duties to make reform targets that are currently infeasible feasible (Gheaus 2013; Gilibert 2012*b*). But this move makes no sense given the restricted possibility account. The plausibility of a transitional duty rests on the claim that there are states of affairs that might not be accessible from our current circumstances yet might be accessible from different circumstances. (Our transitional duty, then, is to bring about the circumstances from which the currently inaccessible reform target can be accessed.) But this is just to say that that the proposed reform target is currently accessible, albeit by a complicated and protracted series of steps. On the restricted possibility account, a reform target that can be accessed in this way satisfies the proffered necessary condition for

A second response is to reject the restricted possibility account of feasibility. This account (or something like it) generates the epistemic problem because it makes the feasibility of realizing a state of affairs an intractable function of myriad facts about the status quo: facts about our current economic, institutional, technological, and motivational resources (among others); facts about how these resources interact to bring about changes to the status quo; facts about how these changes to the status quo affect our subsequent resource stock. In addition, if we are to estimate the limits of the feasibility frontier at any reasonable temporal distance from the status quo, we must account for the fact that the effects of certain changes compound over time and interact in complex ways with later changes. But we can avoid the need to make all these intractable calculations if we adopt a simpler analysis of feasibility. Lawford-Smith's account is an option here. On her view, to determine whether a reform target is feasible, we need only settle whether it is logically, conceptually, metaphysically, and nomologically possible. This is much simpler because political philosophers are unlikely to propose social and political arrangements that are logically inconsistent, conceptually incoherent, or that violate fundamental physical laws. Accordingly, we can be sufficiently confident that any intelligible political ideal is feasible.

I don't wish to deny this response either. This may well be a point in favor of Lawford-Smith's alternative. I make two simple points in reply. First, as noted in section 1, this alternative is not without its disadvantages. (It bears noting here that Lawford-Smith's account seems at odds with the sorts of apparently valid contrapositive inferences mentioned above. Strict egalitarian socialism is almost certainly logically, conceptually, metaphysically, and nomologically possible, and extreme improbability does not license the inference.) Which account we should adopt should turn on a full consideration of their relative merits given our theoretical purposes. Perhaps we should even concede that we should adopt different notions of feasibility for different theoretical purposes. I leave a full discussion of the relative advantages and disadvantages of the two accounts for another time. Second, it's not at all clear to me how much positive weight we should give to the fact that an account like Lawford-Smith's makes it relatively easy to estimate the feasibility frontier (or, conversely, how much negative weight we should assign to the restricted possibility account in virtue of the epistemic challenges it presents). Should we favor an analysis that makes feasibility analysis easier? If we want an analysis that fits ordinary usage, then I don't see why ease of calculation should be a desideratum. Ordinary (binary) feasibility judgments are vague, at least from an epistemic perspective. The restricted possibility account accommodates that vagueness, while Lawford-Smith's

feasibility. Notice that it remains an issue whether we can determine with any confidence whether the reform objective is accessible in this way.

account of binary feasibility does not (although it can capture the vagueness of scalar feasibility judgments).

I stress that I do not wish to deny these replies; I simply point out the (nontrivial) work that must be done to sustain them. I now want to outline a third response for those who are sympathetic to both the feasibility requirement on normative political theories and the restricted possibility account of feasibility. The basic thought is to drop the forward-looking notion of moral progress underlying the Target View and adopt a backward-looking notion of moral progress instead. Let me explain.

The Target View is summed up in the thought that a long-term objective is needed to guide progressive political reform efforts. The predominant rationale for this thought is twofold. First, we need a long-term objective in view to help us chart a progressive transitional path. If we are to chart a *transitional* path, we must be transitioning *toward* something; if that path is to be *progressive*, our transitional path should end at a state of affairs that is morally superior to the here and now. If we have our choice of ends, why not choose a politically ideal end? Simmons expresses the intuition well:

We can hardly claim to know whether we are on the path to the ideal of justice until we can specify in what that ideal consists. . . . The requirement that nonideal policies be “likely to be successful” requires that we know how to measure success; and that measure makes essential reference to the ultimate target, the ideal of perfect justice. (Simmons 2010: 34)

Second, and relatedly, we need a long-term aim in view to ensure that our reform efforts don’t land us at a local rather than global optimum, perhaps unnecessarily closing off the possibility of realizing states of affairs that are morally superior to the local optimum (Gilbert 2012a: 47; Simmons 2010: 24).

Underlying this twofold rationale is a notion of moral progress as *progress toward* a goal—a political ideal, in this case. The arguments in the last section raise difficulties for a normative methodology based on such a forward-looking notion of progress. The basic problem is that our efforts to characterize a distant objective that represents fundamental departures from the status quo are too fraught to justify asserting political obligations to work toward the realization of any particular long-range target. No doubt our evaluative judgments remain untouched by my argument; we can continue to assert that realizing some political ideal is morally desirable in some sense. But, as argued above, we cannot reasonably expect reform efforts aimed at a distant political ideal to realize the optimal feasible states of affairs.

We might concede that the Target View holds out an uncertain and risky goal as a

standard of moral progress yet continue to hold the Target View for lack of a better alternative. Unless we can formulate a method for specifying normative principles that holds greater promise for guiding us toward morally progressive reform, we might nonetheless persist in thinking that political ideals offer the best guidance for our efforts to realize greater justice. Fortunately, an alternative exists.

The key step is to reorient normative political theory around an alternative conception of moral progress—a backward-looking notion of *progress from* actual injustice rather than a forward-looking notion of *progress toward* ideal justice (cf. Schmidtz 2011). Thinking of progress in this way gives normative political theory a new impetus. Instead of trying (against all odds) to chart an uncertain transitional path toward a risky goal, we reorient normative theory to focus on concrete *social failures* rather than political ideals.³¹ Here, “social failure” refers to a state of affairs that is morally inferior to (known) feasible alternatives according to our evaluative criteria. Reoriented toward the analysis of failures, normative political philosophy undertakes the following clinical tasks: identifying states of affairs that are suboptimal with respect to some evaluative criteria³²; diagnosing the causes of the suboptimal states; evaluating potential remedies for circumventing or mitigating the identified failure, assessing the moral costs and benefits of alternatives; prescribing remedies that are (as far as we can tell) likely to leave open possibilities for future progress. Detailed discussion of the mechanics of this “failure analysis approach” is beyond the scope of this paper (see Wiens 2012; cf. Anderson 2010; Knight and Johnson 2011). My point here is that reorienting normative theory around an alternative notion of progress circumvents the challenges faced by the Target View. Unlike an ideal-oriented approach to normative theory, the failure analysis approach has no motivation to chart an uncertain path toward a risky goal in the face of severe epistemic and practical challenges. Yet the failure analysis approach can deliver truly progressive normative guidance, pressing us to redress existing injustices in ways that do not close off (as far as we can tell) possibilities for further improvement.

There remains a worry that, by adopting a largely backward-looking approach like failure analysis, we might land ourselves at a local optimum or perhaps a suboptimal

31 This isn't to say that there is no place for ideal analysis in political philosophy; indeed, ideal analysis can even serve the aims of a reoriented normative political philosophy. A well-developed ideal models a counterfactual world, which we can investigate to determine the necessary and sufficient conditions for realizing the modeled states of affairs. Analysis of this sort can be useful for exploring the feasibility frontier. All my argument implies is that we cannot justifiably derive moral directives for real world reform efforts from the analysis of ideals.

32 One might press here that we need political ideals to provide the evaluative standard by which we identify social failures. As I noted in footnote 20, the view that political ideals can serve as an evaluative standard (as defined here) raises issues that I cannot deal with here. See Wiens (N.d.) for a skeptical argument.

state of affairs with few avenues of escape. I don't deny this worry. But it does not arise because we seek guidance by looking to failures rather than ideals. The worry arises because of our limited foresight, which I've shown to be a problem for the Target View too. Quite simply, we can't see where particular paths of reform lead very far into the future. Even if we have a clear picture of a political ideal, we can't chart even a hazy path from the status quo toward that ideal because it typically lies beyond the limits of our foresight. So this worry about getting stuck at a local rather than global (feasible) optimum applies to the Target View as much as it applies to a failure analysis approach. Ironically, the Target View is less able to cope with this concern than the failure analysis approach. As I have argued, we have no reasonable confidence of making moral progress by aiming to realize an ideal. Further, by focusing primarily on ideals and only indirectly on failures, political philosophers leave themselves susceptible to undue optimism about the potential for their proposed ideals to realize moral improvements, often neglecting to consider the ways in which their ideals risk creating new failures (Wiens 2014; cf. Petroski 2006). Finally, by neglecting to analyze concrete failures rigorously, we leave ourselves ill-equipped to provide trenchant normative guidance for the fraught task of overcoming actual injustices — a task that, if done well, promises to bring real (if incremental) moral gains.

6. APPENDIX

This appendix has three aims. Section 6.1 formulates more precisely the core elements of the model developed informally in section 3. Section 6.2 formulates more precisely the reasoning underlying the arguments at the beginning of section 4. Section 6.3 makes explicit a key premise in the argument at the end of section 4.

6.1. The feasibility frontier. Let W denote the set of possible worlds. Let R be the set of possible resource stocks, V the set of possible conversion processes, C the set of possible causal processes, and S the set of possible states of affairs. Since R , V , C , and S are the only properties of interest here, we represent worlds at a time t as a quadruple: for every $w \in W$, $w^t = \langle r, v, c, s \rangle^{w,t}$. r is a vector indicating the amount on hand of each component of the stock at the indicated world and time; $r^{w,t} = (r_1, \dots, r_n)$. (We could also represent v , c , and s as sets if necessary; for instance, $v^{w,t} = \{v_1, \dots, v_n\} \subseteq V$.) We represent the actual world at t as $\alpha^t = \langle r, v, c, s \rangle^{\alpha,t}$. A world-history is represented as a sequence of quadruples, $\{w^{t_i}\}_{i=-\infty}^{\infty}$. t_0 denotes the time of evaluation throughout.

We treat conversion processes as functions, $v: R \rightarrow R$; $r' \in v(r)$ if and only if, for some $w \in W$ and $t_n \geq t_k$, there is a $\{w^{t_i}\}_{i=k}^{t_n}$ such that $w^{t_k} = \langle r, v, \cdot, \cdot \rangle$ and $w^{t_n} = \langle r', \cdot, \cdot, \cdot \rangle$. We

treat causal processes as functions too, $c : R \times S \rightarrow S$; $s' \in c(r, s)$ if and only if, for some $w \in W$ and $t_n \geq t_k$, there is a $\{w^{t_i}\}_i$ such that $w^{t_k} = \langle r, \cdot, c, s \rangle$ and $w^{t_n} = \langle \cdot, \cdot, \cdot, s' \rangle$. (I replace a variable with a center dot $(\cdot)$ whenever the particular value of the variable is irrelevant.)

The embedded accessibility relations in detail. Recall that realizing s at t is feasible only if s obtains at a world that is accessible from α^{t_0} . We evaluate accessibility in two steps (to make explicit that resources might be spent at two distinct points). The first step pertains to attainable resource stocks. Let v^{α, t_0} be the conversion process that obtains at α at t_0 ; let r^{α, t_0} be the resource stock that obtains at α at t_0 . A resource stock $r \in R$ is attainable if and only if $r \in v^{\alpha, t_0}(r^{\alpha, t_0})$. Let W^R denote the set of worlds with attainable resource stocks; $W^R = \{w \in W : w^{t_0} = \langle r^{\alpha, t_0}, v^{\alpha, t_0}, \cdot, \cdot \rangle\}$

The second embedded accessibility relation pertains to those states of affairs that can be realized by attainable resource stocks. We note that, given a particular resource stock r and initial state of affairs s , the states of affairs that can be realized by $c(r, s)$ might differ from those that can be realized by $c'(r, s)$. Let c^{α, t_0} be the causal process(es) that obtain at α at t_0 and s^{α, t_0} be the state(s) of affairs that obtain at α at t_0 . We say that realizing s is feasible only if $s \in c^{\alpha, t_0}(r^{\alpha, t_0}, s^{\alpha, t_0})$. This second step thus generates the constraint set, denoted $W^C = \{w \in W^R : w^{t_0} = \langle r^{\alpha, t_0}, v^{\alpha, t_0}, c^{\alpha, t_0}, s^{\alpha, t_0} \rangle = \alpha^{t_0}\} \subseteq W^R$.

Recall that W^C circumscribes the feasibility frontier, i.e., the outer bound of the feasible set. We are ultimately interested in the feasibility of realizing states of affairs, which are properties of worlds. We say that realizing s at $t^* \geq t_0$ is feasible only if there is a $w \in W^C$ and $t_n \geq t_0$ such that $w^{t_n} = \langle \cdot, \cdot, \cdot, s \rangle$ and $t_n = t^*$. This implies the analysis offered earlier: realizing s at t is feasible only if s is a property of a world that is accessible from α^{t_0} .

Agent relativity made explicit. Let A be the set of agents; let c_a be a causal process involving actions taken by $a \in A$. Let $c_a : R \times S \rightarrow S$ be a function from resources to states of affairs, such that $s' \in c_a(r, s)$ if and only if, for a w and $t_n \geq t_0$, there is a $\{w^{t_i}\}_i$ such that $w^{t_0} = \langle r, \cdot, c_a, s \rangle$ and $w^{t_n} = \langle \cdot, \cdot, \cdot, s' \rangle$. Then: realizing s' is feasible for some agent a only if $s' \in c_a(r^{\alpha, t_0}, s^{\alpha, t_0})$. Realizing s' is feasible *simpliciter* only if realizing s' is feasible for some $a \in A$.

This completes the formal presentation of the model.

6.2. Feasibility in normative theory. Let PI denote the ideal world, the highest ranked world according to our standard of moral desirability. Let OF denote the optimal feasible world, the highest ranked *feasible* world according to our evaluative standard. Assuming

(for the sake of argument) a feasibility requirement on normative theories, I assume throughout that we have a standing obligation to realize OF .

First specification of the Target View. PI constitutes the political ideal and our reform efforts should aim to realize s^{PI,t^*} for some $t^* \geq t_0$ (where $s^{w,t}$ denotes the states of affairs that obtain at w at t). Given our obligation to realize OF , we have an obligation to realize $s^{OF,t}$ for some $t \geq t_0$. Thus, we should aim to realize s^{PI,t^*} only if the sequence $\{OF^{t_i}\}_i$ is such that $OF^{t_0} = \langle r, v, c, s \rangle^{\alpha, t_0}$ and $OF^{t_n} = \langle \cdot, \cdot, \cdot, s^{PI,t^*} \rangle$ for some $t_n \geq t_0$. In words, we ought to realize s^{PI} only if the sequence that constitutes OF culminates in the realization of s^{PI} at some t .³³

Suppose $PI \notin W^C$. Then $PI^{t_0} \neq \langle r, v, c, s \rangle^{\alpha, t_0} = \langle r, v, c, s \rangle^{OF, t_0}$. Thus, $\{PI^{t_i}\}_i \neq \{OF^{t_i}\}_i$. Given this, we are justified in believing that $\{OF^{t_i}\}_i$ will culminate in the realization of s^{PI,t^*} only if we are justified in believing that $\langle r, v, c, s \rangle^{PI, t_0}$ and $\langle r, v, c, s \rangle^{OF, t_0}$ are precisely tuned such that both $\{OF^{t_i}\}_i$ and $\{PI^{t_i}\}_i$ eventually converge on the same states of affairs. Thus, if $PI \notin W^C$, we ought to realize s^{PI,t^*} only if we are justified in believing that both $\{OF^{t_i}\}_i$ and $\{PI^{t_i}\}_i$ converge to s^{PI,t^*} despite (perhaps radically) different initial conditions. Without fairly thorough examination of both $\{OF^{t_i}\}_i$ and $\{PI^{t_i}\}_i$, this is an unwarranted assumption.

Suppose $PI \in W^C$. Then $OF = PI$ and $\{OF^{t_i}\}_i$ is such that $OF^{t_n} = \langle \cdot, \cdot, \cdot, s^{PI,t^*} \rangle$ for some $t_n \geq t_0$. Thus, we should realize s^{PI,t^*} . However, notice that we are justified in believing that we should realize s^{PI,t^*} only if we have shown, at least, that PI is not likely to be excluded from W^C , the constraint set. Thus, our belief that we should realize s^{PI,t^*} is unwarranted until we have done a thorough feasibility assessment of PI .

33 This way of formulating the point seems to posit a static target for political reform — the states of affairs that obtain at a target world (PI or OF) at a particular time t^* . This is a somewhat simplistic, and thus unsatisfactory, picture. But the model presented here can easily accommodate dynamic targets instead. For example, the Target View might identify an ideal as the sequence of states that obtains at PI starting at some critical time t^* ; or we might say that we have a standing obligation to realize the sequence of states that obtains at OF starting at some critical time t^* (perhaps the earliest possible time at which we can “lock on” to the OF sequence from the actual world, or the time at which we can “lock on” to the OF sequence while incurring the least moral costs). Alternatively, instead of identifying a (timeless) optimal feasible world, we might identify the optimal feasible states for each time period, so that w_{OF}^i denotes the morally optimal world at t_i , $i = 1, \dots, N$. (Note that this permits the possibility that $w_{OF}^i \neq w_{OF}^j$ for any $i \neq j$.) The possibility of accommodating dynamic reform targets allows us to analyze issues related to time inconsistency or path dependence when setting reform targets at different time horizons. For example, on the latter proposal, there’s no reason to assume a priori that w_{OF}^0 will be accessible from w_{OF}^5 , even though both are (by assumption) accessible from α^0 . Such a case thus raises intricate questions about the tradeoffs we must make when adopting different time horizons. Alas, such issues must be left for another time. (Thanks to an anonymous reviewer for raising these important issues.)

In either case, we are not justified in believing that we should realize s^{PI, t^*} unless we have thoroughly explored the feasibility frontier.

Second specification of the Target View. The political ideal is constituted by the normative principles that, given general compliance, sustain s^{PI} and our reform efforts should aim to implement those principles, denoted p^* . Since different normative principles sustain different states of affairs (see discussion in main text), we treat normative principles as a variable that conditions the causal processes that operate at a world; c_p denotes the causal processes that operate given general compliance with a set of normative principles, p .

Given our obligation to realize s^{OF, t_n} for some $t_n \geq t_0$, we should (eventually) implement p^* only if $\{OF^{t_i}\}_i$ is such that $OF^{t_j} = \langle \cdot, \cdot, c_{p^*}, \cdot \rangle$ for at least one $0 \leq j \leq n$. In words, we should implement p^* only if general compliance with those principles (at some time) yields a causal process at OF that leads to the realization of s^{OF, t_n} .

Suppose $PI \notin W^C$. Then $PI^{t_0} \neq \langle r, v, c, s \rangle^{\alpha, t_0} = \langle r, v, c, s \rangle^{OF, t_0}$. Thus, $\{PI^{t_i}\}_i \neq \{OF^{t_i}\}_i$. By assumption, $\{PI^{t_i}\}_i$ is such that, for some $j \geq 0$, $PI^{t_k} = \langle \cdot, \cdot, c_{p^*}, \cdot \rangle$ for all $k \geq j$. Given this, we are justified in believing that $\{OF^{t_i}\}_i$ is such that $OF^{t_j} = \langle \cdot, \cdot, c_{p^*}, \cdot \rangle$ for at least one $0 \leq j \leq n$ only if we are justified in believing that, for some $j \leq n$, $s^{OF, t_n} \in c_{p^*}(r^{OF, t_j}, s^{OF, t_j})$ and $s^{PI, t_n} \in c_{p^*}(r^{PI, t_j}, s^{PI, t_j})$. Thus, we ought to implement p^* only if we are justified in believing that $\langle r, v, c, s \rangle^{PI, t_0}$ and $\langle r, v, c, s \rangle^{OF, t_0}$ are precisely tuned, such that both $\{OF^{t_i}\}_i$ and $\{PI^{t_i}\}_i$ converge on $\langle \cdot, \cdot, c_{p^*}, \cdot \rangle^{t_j}$ for at least one $0 \leq j \leq n$, despite distinct starting and end points. Without fairly thorough examination of both $\{OF^{t_i}\}_i$ and $\{PI^{t_i}\}_i$, this is an unwarranted assumption.

Suppose $PI \in W^C$. Then $OF = PI$ and $\{OF^{t_i}\}_i$ is such that $OF^{t_j} = \langle \cdot, \cdot, c_{p^*}, \cdot \rangle$ for at least one $0 \leq j \leq n$. Thus, we should implement p^* . However, notice that we are justified in believing that we should implement p^* only if we have shown, at least, that PI is not likely to be excluded from W^C , the constraint set. Thus, our belief that we should implement p^* is unwarranted until we have done a thorough feasibility assessment of PI .

In either case, we are not justified in believing that we should implement p^* unless we have thoroughly explored the boundaries of the feasible set.

6.3. Claim. *We have no reason to believe that the optimal feasible world approximates the political ideal.*

Proof. Suppose $PI \notin W^C$. Then $PI^{t_0} \neq \langle r, v, c, s \rangle^{\alpha, t_0} = \langle r, v, c, s \rangle^{OF, t_0}$. Thus, the political ideal is not identical to the optimal feasible world; $\{PI^{t_i}\}_i \neq \{OF^{t_i}\}_i$. Let S^{PI} be the set of states of affairs that are similar enough to s^{PI} (the ideal states realized at PI) to count as

approximations to s^{PI} . Then OF approximates PI only if $\{OF^{t_i}\}_i$ converges to $s \in S^{PI}$ at some time. We have no reason to believe this to be true without thoroughly examining both $\{PI^{t_i}\}_i$ and $\{OF^{t_i}\}_i$.

REFERENCES

- Anderson, Elizabeth. 2010. *The Imperative of Integration*. Princeton: Princeton University Press.
- Brennan, Geoffrey. 2007. Economics. In *A Companion to Contemporary Political Philosophy*, ed. Robert E. Goodin, Philip Pettit and Thomas Pogge. 2nd ed. Malden, MA: Blackwell.
- Brennan, Geoffrey. 2013. "Feasibility in Optimizing Ethics." *Social Philosophy and Policy* 20(1/2):314–329.
- Brennan, Geoffrey and Nicholas Southwood. 2007. Feasibility in Action and Attitude. In *Hommage à Wlodek. Philosophical Papers Dedicated to Wlodek Rabinowicz*, ed. T. Rønnow-Rasmussen, B. Petersson, J. Josefsson and D. Egonsson. Available online at www.fil.lu.se/hommageawlodek.
- Brennan, Geoffrey and Philip Pettit. 2005. The Feasibility Issue. In *Oxford Handbook of Contemporary Philosophy*, ed. Frank Jackson and Michael Smith. Oxford: Oxford University Press.
- Buchanan, Allen. 2004. *Justice, Legitimacy, and Self-Determination: Moral Foundations for International Law*. New York: Oxford University Press.
- Cohen, G. A. 2009. *Why Not Socialism?* Princeton: Princeton University Press.
- Cohen, G.A. 2003. "Facts and Principles." *Philosophy & Public Affairs* 31(3):211–245.
- Cowen, Tyler. 2007. "The Importance of Defining the Feasible Set." *Economics and Philosophy* 23(1):1–14.
- Estlund, David. 2011. "Human Nature and the Limits (If Any) of Political Philosophy." *Philosophy & Public Affairs* 39(3):207–237.
- Estlund, David. 2014. "Utopophobia." *Philosophy & Public Affairs* 42(2):113–134.
- Gheaus, Anca. 2013. "The Feasibility Constraint on the Concept of Justice." *The Philosophical Quarterly* 63(252):445–464.

- Ghosh, Prodipto. 2012. "Climate change: Those creating problems least willing to take action." Interview in *ZeeNews.com*. 2 February 2012. http://zeenews.india.com/news/exclusive/climate-change-those-creating-problems-least-willing-to-take-action_756232.html.
- Gilabert, Pablo. 2012a. "Comparative Assessments of Justice, Political Feasibility, and Ideal Theory." *Ethical Theory & Moral Practice* 15(1):39–56.
- Gilabert, Pablo. 2012b. *From Global Poverty to Global Equality*. New York: Oxford University Press.
- Gilabert, Pablo and Holly Lawford-Smith. 2012. "Political Feasibility: A Conceptual Exploration." *Political Studies* 60(4):809–825.
- Hamlin, Alan and Zofia Stemplowska. 2012. "Theory, Ideal Theory and the Theory of Ideals." *Political Studies Review* 10:48–62.
- James, Aaron. 2012. *Fairness in Practice: A Social Contract for a Global Economy*. New York: Oxford University Press.
- Jubb, Robert. 2012. "Tragedies of Non-Ideal Theory." *European Journal of Political Theory* 11(3):229–246.
- Knight, Jack and James Johnson. 2011. *The Priority of Democracy: Political Consequences of Pragmatism*. Princeton: Princeton University Press.
- Kratzer, Angelika. 1991. Modality. In *Semantics: An International Handbook of Contemporary Research*, ed. Arnim von Stechow and Dieter Wunderlich. Berlin: de Gruyter pp. 639–650.
- Lawford-Smith, Holly. 2013. "Understanding Political Feasibility." *The Journal of Political Philosophy* 21(3):243–259.
- Lewis, David. 1973. "Counterfactuals and Comparative Possibility." *Journal of Philosophical Logic* 2(4):418–446.
- Lipsey, R.G. and Kelvin Lancaster. 1956. "The General Theory of Second Best." *The Review of Economic Studies* 24(1):11–32.
- Petroski, Henry. 2006. *Success Through Failure: The Paradox of Design*. Princeton and Oxford: Princeton University Press.

- Pogge, Thomas W. 2008. *World Poverty and Human Rights*. 2nd ed. Malden, MA: Polity Press.
- Räikkä, Juha. 1998. "The Feasibility Condition in Political Theory." *Journal of Political Philosophy* 6(1):27–40.
- Rawls, John. 1999. *A Theory of Justice*. 2 ed. Cambridge, MA: Harvard University Press.
- Robeyns, Ingrid. 2008. "Ideal Theory in Theory and Practice." *Social Theory and Practice* 34(3):341–362.
- Sangiovanni, Andrea. 2008. Normative Political Theory: A Flight From Reality? In *Political Thought and International Relations: Variations on a Realist Theme*, ed. Duncan Bell. Oxford: Oxford University Press pp. 219–239.
- Schmidtz, David. 2011. "Nonideal Theory: What It Is and What It Needs to Be." *Ethics* 121(4):772–796.
- Sen, Amartya. 1981. *Poverty and Famines: An Essay on Entitlement and Deprivation*. New York: Oxford University Press.
- Sen, Amartya. 1999. *Development as Freedom*. New York: Anchor Books.
- Simmons, A. John. 2010. "Ideal and Nonideal Theory." *Philosophy & Public Affairs* 38(1):5–36.
- Stalnaker, Robert C. 1968. A Theory of Conditionals. In *Studies in Logical Theory*, ed. Nicholas Rescher. Oxford: Oxford University Press pp. 98–112.
- Stemplowska, Zofia. 2008. "What's Ideal About Ideal Theory?" *Social Theory and Practice* 34(3):319–340.
- Usigbe, Leon. 2012. "Continental free trade regime not feasible in 2017 - Jonathan." *Nigerian Tribune*. 30 January 2012. <http://tribune.com.ng/index.php/news/35122-continental-free-trade-regime-not-feasible-in-2017-jonathan>.
- Valentini, Laura. 2012a. "Ideal vs. Non-ideal Theory: A Conceptual Map." *Philosophy Compass* 7(9):654–664.
- Valentini, Laura. 2012b. *Justice in a Globalized World: A Normative Framework*. New York: Oxford University Press.
- Wiens, David. 2012. "Prescribing Institutions Without Ideal Theory." *The Journal of Political Philosophy* 20(1):45–70.

Political Ideals and the Feasibility Frontier

- Wiens, David. 2013. "Demands of Justice, Feasible Alternatives, and the Need for Causal Analysis." *Ethical Theory & Moral Practice* 16(2):325–338.
- Wiens, David. 2014. Achieving Global Justice: Why Failures Matter More Than Ideals. In *Making Global Institutions Work: Power, Accountability and Change*, ed. Kate Brennan. New York: Routledge.
- Wiens, David. N.d. "Against Ideal Guidance." Manuscript, University of California, San Diego, <www.dwiens.com>.
- Ypi, Lea. 2012. *Global Justice & Avant-Garde Political Agency*. New York: Oxford University Press.