

Tara Smith's *Viable Values: A Study of Life as the Root and Reward of Morality*: A Discussion

Irfan Khawaja

I. Introduction

Tara Smith's *Viable Values* is a valuable addition (so to speak) to a newly-emerging literature on the philosophical thought of Ayn Rand. Its contribution, as I see it, is two-fold. On the one hand, it lays out the essential features of Ayn Rand's Objectivist Ethics in a clear and persuasive way; on the other hand, it presents a series of bold challenges, both direct and indirect, to a good deal of mainstream ethical theorizing in the "analytic" tradition.¹ The book as a whole adds up nicely to the sum of these parts.

Though there is much to say about virtually every section of the book, I'll confine myself here to discussion of four issues: questions of method; the conditionality of life as the basis of moral value; the content of moral prescriptions as egoistic rather than impersonal; and the conditional (as opposed to categorical) force of moral obligation. These four issues are central both to Smith's project and to Ayn Rand's distinctive ethical vision; they also mark a faultline that distinguishes the Objectivist conception of ethics from that of the mainstream. Focusing on them, I think, should bring into sharp relief the issues that divide Objectivism from mainstream ethical thought.²

¹ I use the term "mainstream" throughout to denote "the currently dominant mode of philosophical research and inquiry in the United States according to its avowed practitioners." As John Searle puts it, "The dominant mode of philosophizing in the United States is called analytic philosophy. Without exception, the best philosophy departments in the United States are dominated by analytic philosophers, and among the leading philosophers in the United States, all but a tiny handful would be classified as analytic philosophers....Analytic philosophy has never been fixed or stable, because it is intrinsically self-critical and its practitioners are always challenging their presuppositions and conclusions." John Searle, "Contemporary Philosophy in the United States," in *The Blackwell Companion to Philosophy*, (Cambridge, MA: Blackwell Publishers, 1996), edited by Nicholas Bunnin and E.P. Tsui-James, pp. 1-2. See also Bernard Williams, "Contemporary Philosophy: A Second Look," in the same volume, p. 26.

² "Objectivism" is the name Ayn Rand gave to her philosophy as a whole. Given Smith's claim to offer an elaboration and defense of Rand's account (*IV*, p. 83), it follows trivially (I think) that her book is a defense of the Objectivist Ethics. In saying this, I don't mean to imply that Smith is an official spokesperson for Objectivism, or that her book (or this paper) represent "official doctrine."

Reason Papers 26 (Summer 2000): 63-88. Copyright © 2000

II. Questions of Method

Viable Values begins with the age-old question, “Why should I be moral?” and with a provocative discussion of the failure of mainstream approaches to the question (VV, ch. 2). (For ease of reference, let’s call this *the Question*.) One of the unique features of her account, Smith says, is success at answering the Question where other theories have failed (VV, pp. 83, 187). Given the centrality of the Question to the book, it would be useful to know exactly what the Question is asking, and why it assumes such significance in the first place.

As many philosophers have noted, the Question “Why should I be moral?” is highly ambiguous; in fact, philosophers have often dismissed the question for just that reason. Smith ably refutes the most important of these criticisms (VV, pp. 16-28), but doesn’t in my view dwell at sufficient length on Rand’s unique way of posing the Question. In the first few paragraphs of her essay “The Objectivist Ethics,” Rand phrases her opening question as follows: The first question that has to be answered, as a precondition of any attempt to define, to judge or to accept any specific system of ethics, is: *Why* does man need a code of values?

Let me stress this. The first question is not: What particular code of values should man accept? The first question is: Does man need values at all—and why?³

Depending on how one counts them, there are anywhere between three and six questions here, each significantly different from (and more precise than) those typically discussed in the mainstream literature. They are:

1. Do we need to value? If so, why?
2. Given answers to (1) and (2), why do we need a *code* of values (and what is a “code of values”)?
3. Given affirmative answers to (1) and (2), *which* code of values do we need?

I’ll return to the answers to these questions in the next section. For now, I want to pause on three features of Rand’s method that distinguish it from contemporary analytic approaches to ethics, and which inform Smith’s method as well (VV, pp. 13-19).

First, note Rand’s emphasis on the proper formulation of the question that motivates ethical theorizing: we can’t, according to her, embark on ethical inquiry until we’re clear about the question that states its goal. Nor will just

In the space at my disposal, I’ve had to give a very compressed account of Smith’s argument; I’ve also felt free to move back and forth between Smith’s book and Rand’s original claims, without (I hope) distorting the content of either.

³ Ayn Rand, “The Objectivist Ethics,” in *The Virtue of Selfishness: A New Concept of Egoism*, (New York: Signet Books, 1964), pp. 13-14, italics in original.

any idiosyncratic question do. Only the right question gets you the right answer, and only the right answer so conceived will produce moral knowledge.

Secondly, note the insistence that the questions be posed and answered in a specific sequence. We're to begin with the most fundamental justificatory questions about value as such, questions that make no specifically *moral* presuppositions (questions 1 above). We then answer each of these morally-neutral questions before moving to derivative questions about the content of morality, taking the answers to the fundamental questions to be the basis for inquiry into the derivative ones. The result, if successful, should be an ethical theory with a foundationalist structure in which a basic normative principle(s) serves to support a superstructure of derivative ethical prescriptions.

Finally—and as a consequence of the preceding two points—note that the method as a whole forbids circularity and question-begging in ethical argument. Beginning with a morally neutral starting point rules out the assumption that our “considered moral beliefs” embody moral knowledge, and that the task of ethical theorizing is simply to “account” for these beliefs by appealing to higher-order principles. As Smith aptly puts it, the latter procedure merely allows one to seek “a rationalization for one’s existing beliefs rather than a justification of true beliefs” (*IV*, p. 15).

All of this flies in the face of received wisdom about method in analytic philosophy. As Dale Jamieson puts it, These days some version of coherentism is the dominant view of what constitutes proper method for theory construction in ethics. Coherentism can be roughly characterized as the view that beliefs can be justified only by their relation to other beliefs.

The most influential form of coherentism is Rawls’s method of reflective equilibrium. According to Rawls, proper method involves beginning with a set of considered beliefs, formulating general principles to account for them, and then revising both principles and beliefs in the light of each other, until an equilibrium is reached.⁴

Obviously, on this view, none of Rand’s methodological precepts applies. For the reflective equilibrium theorist, no specific question provides the goal of ethical inquiry; any question will do, and the questions can be posed and answered in any order, or in no particular order at all. Consequently, reflective equilibrium encourages rather than proscribes circular reasoning; in fact, circularity may well be the only stable and determinate feature of the method.

⁴ Dale Jamieson, “Method and Moral Theory,” p. 482 in *A Companion to Ethics*, (Malden, MA: Blackwell, 1993), edited by Peter Singer.

I think we can see the merits of Rand's approach to ethics if we compare the two methods' prospects for producing moral knowledge. Rand defines knowledge as "a mental grasp of a fact(s) of reality, reached either by perceptual observation or by a process of reason based on perceptual observation."⁵ On the mainstream view, knowledge is justified true belief with a proviso for handling so-called Gettier cases; alternatively, knowledge can be characterized as true belief that non-accidentally links truth and belief (as in Nozick's "truth-tracking" conception). On any plausible account, "knowledge" is a success term denoting a specific cognitive relationship between minds and truth, with truth conceived as correspondence to a mind-independent reality. What goes for knowledge generally goes for moral knowledge as well: moral knowledge is moral belief related in the right way to moral reality, where "moral reality" is discovered rather than constituted by mental activity.

By this standard, Rand's method does quite well. If we assume that question (1) is a legitimate and important question—an assumption that Smith indirectly defends (*IV*, pp. 16-28)—and we offer a genuinely responsive answer to that question (as Smith does, *IV* ch. 4), then if the answers are true, the result is a true foundation for moral knowledge, capable of transmitting truth to subsidiary principles via the answers to the subsidiary questions. It's a separate issue whether the resulting structure will provide the *whole* truth about moral reality, but it seems inarguable that it will constitute some important part of it by the very nature of the questions involved.

By contrast, reflective equilibrium does poorly by this standard, for three related reasons. For one thing, insofar as reflective equilibrium is coherentist, and coherentism countenances (or requires) circular reasoning, any theory based on reflective equilibrium will be circular. But if circularity is a fallacy, then all such theories will be fallacious—and no fallacious theory can count as knowledge.

Second, even if we waive the first objection, the method turns out to be vacuous. Reflective equilibrium tells us to achieve a coherentist "equilibrium" between our theoretical and pre-theoretical beliefs. But the method gives us (a) no determinate account of the criteria for "equilibrium," and (b) no actual method for achieving equilibrium, however we understand it.

Consider (a). "Equilibrium" is presumably a species of coherence. Unfortunately, it's unclear what species of coherence it is, because it's unclear what the genus "coherence" denotes. On the weakest interpretation, a coherent belief-set is one that is free of contradiction; on a stronger interpretation, every belief in such a belief-set entails every other. The difference between the two interpretations is stark, but the method gives us no way of discriminating between them. Since it doesn't, it neither gives us a determinate goal at which

⁵ Ayn Rand, *Introduction to Objectivist Epistemology*, 2nd ed., (New York: Meridian Books, 1990), p. 35.

to aim nor determinate prescriptions by which to achieve its aims. *A fortiori*, it gives us no goal or methods conducive to moral knowledge.

Now consider (b). The method of reflective equilibrium tells us to fit our pre-theoretical to our theoretical beliefs. The distinction between the two sorts of belief is obscure enough, but since the method offers no weighting of the one sort of belief to the other, the prescription turns out to be meaningless in any case: two agents starting with identical belief-sets could well follow the method but arrive at diametrically opposed belief-sets by following radically different procedures. Imagine that Abdul and Bilal begin with the same beliefs, both adopting the method of reflective equilibrium. Abdul brings his beliefs into equilibrium by taking a highly anti-theoretical strategy, keeping only his pre-theoretical beliefs and driving out all of his theoretical ones; Bilal does the reverse. By the terms of the method, neither Abdul nor Bilal has done anything wrong—even though each has done the reverse of what the other has done, and each has come to the reverse conclusions as the other has reached. This result seems to suggest that reflective equilibrium is less a “method” of ethical theorizing than the absence of one. For while we could describe both Abdul and Bilal as “following the method of reflective equilibrium,” we could with equal accuracy and greater simplicity describe them both as merely following their whims. A method this contentless hardly deserves the name.

Finally, even if we take the most charitable attitude toward the method, and waive the first two problems entirely, we’re still left with the fact that the method aims at no more than coherence. But coherence is a relation between beliefs, while knowledge is a relation between beliefs and a world that exists independently of beliefs. If so, the professed goal of the theory is irrelevant at best (and orthogonal at worst) to the goal of cognition, viz., knowledge. After all, if we begin with nothing but considered beliefs that are false, and merely put them into reflective equilibrium, we still end up with false beliefs, however elegantly expressed. Indeed, the very elegance of their expression might divert us even farther from moral knowledge than we might otherwise have been, since false beliefs elegantly presented are bound to attract us more than false beliefs without such theoretical embellishments.

So much for “theoretical” objections to reflective equilibrium. In practice, I think, the difference between Rand’s method and reflective equilibrium comes down to a difference between epistemic radicalism and conservatism. As we’ll see, Rand’s (and Smith’s) ethical views are quite radical, calling into question any number of commonplace assumptions about the status and content of morality. As Smith puts it:

The fact that certain moral conventions have been long or widely accepted carries no guarantee that they can withstand rational scrutiny. In order to arrive at a sound defense of morality, all prior suppositions about morality’s authority must be on the table. Any reader unwilling to suspend such suppositions is warned: Continued reading will be a waste of time. (*IV*, p. 15).

Reflective equilibrium, by contrast, encourages the complacent assumption that “we” already “have” moral knowledge, and that the task is simply to put it all in order. If the Rand-Smith argument is right, however, this assumption is wildly false. So let’s proceed now to the argument.

III. The Foundation: The Conditionality of Life

Let’s revisit the questions posed at the beginning of section II, and consider the Objectivist answers to them in turn. Question (1) was: Do humans *need* to value? Is there anything that entities like us *have* to do? As a placeholder for the notion of “value,” we can think of a value, as Rand does, as “that which one acts to gain and/or keep.” A “value,” in other words, is the object or goal of an action.

At first glance, our question may seem to drive us to a paradox. Start with the assumption that we have free will in the full, libertarian sense of that term. Free will, in turn, implies a power for alternate possibilities. “A course of thought or action is ‘free’,” writes Leonard Peikoff, “if it is selected from two or more courses of action possible under the circumstances. In such a case, the difference is made by the individual’s decision, which did not have to be what it is, i.e., which could have been otherwise.”⁶ The apparent paradox is this: on the one hand, moral obligation entails that I “have” to do what is moral, but on the other hand, if we have free will, and free will entails a power to select between alternate possibilities, there seems no sense in which I *have* to do anything.

An answer to Question (1) must resolve this apparent paradox. It must, in other words, identify an objective constraint on action—i.e., a principle about what we *have* to do—that is compatible with free will on a libertarian (rather than soft-determinist) interpretation.⁷ To put the point another way, an answer to Question (1) must identify a source of *practical necessitation consistent with metaphysical freedom*. To be free is to have alternatives; to value is to seek an object; to be obliged is to be required to act in a specific way for specific objects. What we need is a principle that preserves our freedom, while showing how we can be required to rank some alternatives over others.

The principle in question is what we might call *the principle of conditional necessity*. The first part of the principle governs the *causation* of

⁶ Leonard Peikoff, *Objectivism: The Philosophy of Ayn Rand*, (New York: Dutton, 1991), p. 55.

⁷ By “soft determinist,” of course, I mean “compatibilist,” but introducing the term in this context would only have created confusion. “Compatibilists” in analytic jargon take freedom to be compatible with determinism; by contrast, the compatibility I have in mind in the text is one between moral obligation and a *non*-deterministic form of freedom.

action. To be free is to be able to choose between alternative courses of action, or alternative ends. If I will a given end, the nature of goal-directed action implies that only some actions within the range of possibilities will qualify as effective means for achieving the end. Since the end can't achieve itself, it follows that insofar as I regard the end as my goal, I must take the most effective of the available means to it. Contrapositively, if I refuse to will the means, I must give up the end.

The second part of the principle governs the *consequences* of actions. In opening up alternatives, freedom makes me an agent: I initiate actions by breaking the chain of prior determinants affecting me. My action, in turn, has two aspects—one volitional or internal, the other existential or external, including causal consequences in the physical world and moral consequences on my character. An action that I initiate is irreducibly *mine*; it is my unique causal contribution to the world, brought into the world by my agency. Whatever action I take, then, I must be able to endorse my actions *as mine*, acknowledging my having caused them. To will an act is to will its foreseeable consequences; if I'm unwilling to endorse a consequence, I should be unwilling to take the action that would lead to it.

The senses of "should" and "must" involved here are simply what follows by acknowledging the facts of reality—specifically, the facts of causality in the external world, and the facts of causality arising from my own agency. To violate the principle of conditional necessity, whether as regards causes or consequences, is to flout my own awareness of causal reality. To will the end without willing the means is either to deny that the end requires specific means, or to deny that I've actually willed the end, both of which are plainly false. To refuse to acknowledge the consequences of my actions is implicitly to deny the efficacy of my own agency.⁸

Three additional points are worth making about this principle. First, the principle gets its name by contrast with the idea of *categorical necessity*. A categorical moral necessity is one that supersedes causality, both causally and consequentially: a categorical imperative binds intrinsically, out of relation to the agent's goals (as in Kant's claim that a good will is "good in itself") and it binds deontically, out of relation to the consequences of action (as in Kant's injunction to practice justice "though the heavens may fall"). By contrast, conditional necessity is conditional because its binding force is conditioned precisely on the way causality applies to moral agents. I *must* take measures literally *because* (and *only because*) my goals require it. Had I no goals, I would be subject to no constraints.

⁸ See Smith, *IV*, pp. 38-53; Rand, "Causality Versus Duty," pp. 95-101 in *Philosophy: Who Needs It*, (Signet, 1982). Cf. Thomas Hill, "The Hypothetical Imperative," pp. 17-37 in *Dignity and Practical Reason in Kant's Moral Theory*, (Ithaca: Cornell University Press, 1992).

Secondly, though the principle is grounded in facts about causality, it is perfectly consistent with freedom. While the principle provides practical constraints, it leaves me metaphysically free to set goals, and free to select the means to those goals; it also leaves me free *not* to set goals, and free to flout the principle itself in any number of ways. It does *not* leave me free to repeal the facts that the principle identifies, however. If I have a goal, only certain measures will suffice to enact it: my will has no say about that fact. If I refuse to enact the relevant measures, or refuse to set the goal, or refuse to heed the relevant consequences, causality will still operate in specific and determinate ways with specific and determinate consequences for me. My will has no say about *that* fact, either. So the principle identifies the background context within which free will operates. Freedom opens up alternatives; reality constrains them without encroaching on freedom itself. The principle of conditional necessity codifies that fact, and partly answers question (1).

The principle of conditional necessity does seem, at least *prima facie*, to provide an objective constraint on action consistent with freedom. But the way in which it does so provokes as many new questions as it answers. For one thing, the principle guides action on the assumption that the agent *has* goals. But the principle neither says that the agent *must* have goals, nor (*a fortiori*) identifies which goals she ought to have. Further, the principle provides no account of its own binding force: if goals are what bind agents, then the principle of conditional necessity only binds an agent if the agent regards adherence to the principle *itself* as a goal. But the principle doesn't explain why this must be so. Granted, violation of the principle would implicate the violator in self-deception, and might perhaps bring her adverse consequences. But a given agent may be willing to live with such self-deception, and in any case, by itself, the principle gives us no standard by which to distinguish beneficial from adverse consequences. Nor in any case does it give us a reason for preferring beneficial to adverse consequences. Clearly, then, the principle can't be the whole of the normative story.

In form, that story will have to identify *the ultimate condition that grounds the principle of conditional necessity*. If the principle of conditional necessity tells me what to do if I have a goal, we need to identify a principle that tells me why I must have goals in the first place—and tells me which ones to have, providing a way of ranking goals by their value. If the principle of conditional necessity only binds insofar as it *is* a goal, we need to identify the ultimate goal in relation to which adherence to the principle is a necessary means. If the consequence of violating the principle is self-deception, we need to identify a principle that explains what's wrong with self-deception. If violation of the principle brings the agent adverse consequences, we need a standard for distinguishing beneficial from adverse consequences, and one that explains why the beneficial is to be preferred to the adverse as such.

Smith's answer, which comes from Rand, turns on the conditional character of *life* as the ultimate (unconditioned) conditional necessity.

Following the Aristotelian tradition, Rand characterizes life as a conditional process of self-generated and self-sustaining action.

The existence of inanimate matter is unconditional, the existence of life is not: it depends on a specific course of action. Matter is indestructible, it changes its forms, but it cannot cease to exist. It is only a living organism that faces a constant alternative: the issue of life or death. Life is a process of self-sustaining and self-generated action. If an organism fails in that action, it dies; its chemical elements remain, but its life goes out of existence.⁹

In short, what this means is that a living being's very existence as the entity it is depends on its engaging in a complex repertoire of continuous action across its lifespan, generated and directed by the organism itself.¹⁰ The nature or identity of an organism determines the kind of action appropriate to it. Precisely because each organism *has* a determinate nature, each has needs, and because each has needs, each has a determinate mode of life appropriate to its capacities and environment. Given the determinacy of modes of life so conceived, there is such a thing as acting *in accordance with*, and acting *against*, the requirements of an organism's life—the most obvious requirements being those of physical health and functionality, the less obvious including proper mental functioning and psychological health. An organism's default on its mode of life undermines or impairs its functionality by degrees; severe default undermines its existence and identity altogether, and leads ultimately to death. Let's call this complex fact *the conditionality of life*, the thesis expressing it *the conditionality thesis*, and the entities to which it applies, *conditional entities*.

Contrary to a common misinterpretation, the central idea behind the conditionality thesis is not that life is a necessary condition for valuation, given that one can only value while alive (or no longer can if dead). Nor is the real point that life is a sufficient condition for valuation, in the sense that anything that's alive is bound to be valuing. Both claims are true, of course, but both miss the distinctive point of the thesis, which is fundamental explanatory.¹¹ We need an explanation for the conditional structure of valuation—for the fact that a goal necessitates the selection of means on its behalf, and that the consequences of an action matter to the agent bringing

⁹ Rand, "Objectivist Ethics," p. 16. Cf. Aristotle, *De Anima* II.2-4, and *Generation of Animals* II.1.

¹⁰ Harry Binswanger, *The Biological Basis of Teleological Concepts*, (Los Angeles: ARI Press, 1990), ch. 4.

¹¹ Peikoff makes this point as well, *Objectivism*, pp. 212-213.

them about. The conditionality of life explains both facts because conditionality is the property underlying both facts.¹² It's only living things qua living that *have* to set goals. It's only living things qua living that are affected by the goals they set, and by the manner in which those goals are pursued. A conditional entity has to set goals to meet its needs, and has to meet its needs to remain in existence, satisfying higher-order needs once the initial needs have been met, and yet higher-order needs once the intermediate needs have been met. Failure to satisfy needs undermines the entity's existence in proportion to the urgency of the need. Such an entity, whether rational or non-rational, conscious or non-conscious, faces a constant alternative, and the constant necessity of appropriate action in the face of that alternative. "Appropriate action" is action that takes life as the ultimate conditional necessity, and selects ends and means accordingly.

The conditionality thesis thus shows us that the principle of conditional necessity is not a freestanding thesis, but one rooted in a deeper phenomenon. If an organism is to preserve its existence and identity, it must act appropriately to its nature and circumstances. Being organisms, the thesis applies to us as well in the form of the principle of conditional necessity. It also explains what the principle of conditionality necessity by itself leaves unexplained. The principle of conditional necessity guides action on the assumption that the agent has goals, without explaining why she must, or which goals she ought to have. The conditionality thesis explains both: the agent must have goals to preserve her existence, and she must select those goals which would preserve the existence of a being like her. The principle of conditionality, we saw, provides no account of its own binding force; it leaves unanswered why one ought to adopt it as a goal. The conditionality thesis shows us that if an organism has a specific nature, its preservation can only be achieved by specific means—and that is what the principle of conditional necessity identifies. The principle of conditional necessity gives us no standard of beneficial and adverse, and no reason for preferring one to the other. The conditionality thesis does: survival qua human is the standard of value; self-preservation gives us a reason for preferring the beneficial to the harmful.

¹² Robert Nozick misses this point entirely in his much-celebrated "On the Randian Argument," (first published in *The Personalist*, vol. 52 (Spring 1971), pp. 282-304, reprinted in *Socratic Puzzles*, [Cambridge, MA: Harvard University Press, 1997]): "I would most like to set out the argument as a deductive argument and then examine the premises. Unfortunately, it is not clear to me what the argument is." (p. 249). Nozick's attempt to set out the "argument as a deductive argument" misses the point: "The Objectivist Ethics" is primarily a teleological explanation for why we *must* value, not (as Nozick tortuously interprets it) a "transcendental argument" for the *possibility* of value. Consequently, little of Nozick's discussion has any genuine relevance to Rand's argument. The oddity of Nozick's interpretive approach is heightened by his own insistence on the difference between explanations and deductions in the Introduction of his *Philosophical Explanations*, (Cambridge: Harvard University Press, 1981).

Finally, we saw that the principle of conditional necessity presupposes some version of a reality-orientation or reality principle; it only applies to the extent that the agent chooses not to flout the deliverances of her mind. The conditionality thesis explains why: contact with reality is a requirement of life itself.

That, in short, is the answer to Question (1). Ultimately, I must engage in valuation because as a living being, I'm a conditional being, and as a conditional being, my life depends on self-generated and self-sustaining action. Nor does it ever stop depending on it. If life requires continuous action, then every need I satisfy will bring another one in its wake, at a higher level of complexity. Whatever the variety and complexity of that structure, each element in it gets its ultimate *raison d'être* from its contribution to my life. Default on any element is default on the specifically human mode of living. I can ignore or evade that fact about myself, but I can't escape it—at least so long as I aim to exist. If so, I have to identify what to pursue, and how to pursue it; I have to find what's valuable, and resolve to track it.

What, then, about Question (2)? So far, we've learned that the conditionality of life requires constant valuation of us, and grounds one principle: the principle of conditional necessity. But even if we grant that, how does that give rise to a *code* of values?

First, a terminological detour: what is a "code"? The *Oxford English Dictionary* defines a "code" as a systematic collection or listing of statutes, rules, regulations, or signals. A code of *values*, we might infer, is a set of interconnected principles by means of which values are sought. This is clear enough in particular cases. A judge by definition values justice; she needs a legal code to secure it, and the judicial temperament to follow the code. A fire inspector by definition values fire-safe buildings; she has to insist that every building be "up to code," and cultivate the habits appropriate to her task. And so on, through any number of practices. The idea of a "code" in this context captures two things: (a) the *systematicity* of a collection of complex principles geared to realizing specific purposes in the world, and (b) the *traits* required to follow those principles in the concrete circumstances of the world. By induction, a moral agent would be someone who by definition valued moral ends, and sought those ends by adherence to the principles and habits constitutive of moral action.

Now let's return to our question. Human beings must value because we have needs, and we have needs because we're living, conditional entities. Conditional entities must value to remain in existence and to maintain their identities. Is having a code *itself* a value for a being in such a predicament?

The short answer is "yes." We've already seen how the conditionality of life underwrites a principle of conditional necessitation. The short answer for why we need a code of values is that the application of that principle is complex. Without a code for applying it, we would have no idea *how* to apply it. The fact that we need codes to engage in specific human practices—legal, professional, recreational, etc.—lends credence to the idea that we need a

meta-code that integrates and if necessary overrides the more particular ones. That meta-code is a code of moral values.

The longer answer requires that we look more carefully at how the conditionality of life gives rise to specifically moral value. In answering question (1), I noted that though the conditionality thesis applies to all organisms, it applies differently to each. The differences in application can be designated as differences of form and of content. As a matter of form, the conditionality thesis applies differently to human beings because we are the only organisms capable of following it on principle. As a matter of content, we differ because our needs and environment are different: it requires different actions and conditions of us than it requires of them.

As Ayn Rand first suggested,¹³ conditionality is an attribute of all and only organisms—from the lowliest prokaryote all the way across the biological taxa. Since humans are organisms, human beings are therefore conditional beings like all the others. Given the distinctiveness of human nature, however, the conditionality thesis (and the principle of conditional necessitation) applies in a special way to human beings—a way of particular importance to ethics. Our distinctive nature, for Rand as for Smith, arises from the special features of human consciousness, which differentiate us from all of the other organisms.

First consider the differences in *form*. In the other organisms, valuation is a thoroughly deterministic phenomenon, genetically coded by natural selection, and triggered by features of the environment. Whatever their simplicity or complexity, the lives of non-human organisms proceed more or less automatically—from metabolism, homeostasis and growth, through locomotion, learning, and cooperation. Non-human organisms are in short, automatic, deterministic value-trackers with life as their ultimate value, and an automatic awareness of, and propensity to act on, their needs. The problem of the conditionality of life is for them solved by *genetic coding*. They follow the principle of conditional necessitation because they can do no other.

By contrast, human beings are volitional, rational agents born *tabula rasa*. The fact that we're *tabula rasae* entails that we have no automatic means of discovering what our life-requirements are, or of maintaining a consistent commitment to them. Lacking innate knowledge, we have to discover our means of survival for ourselves; lacking innate behaviors or traits of character, we have to cultivate a willed commitment to tracking the good. Our distinctive moral task, then, is to employ volition and reason in the service of both ends. The problem posed by the conditionality of life is for us a matter of explicit, *volitional and rational* codification. We have to devise and live by a life-code on our own because we can willfully surrender our grip on the

¹³ For a more sustained argument, see Binswanger, *Biological Basis*, especially chs. 8-9.

conditionality of life, violating the principle of conditional necessitation, or any of its subsidiary principles.

Now consider issues of *content*. The basic issue here is that our nature differs radically from that of the other organisms. Consider first the sheer temporal extension of a human life. A full human life consists of the total expected length of an entire lifespan—say, eighty or ninety years. So we need a code that will last the distance, and traits that will do so as well. The second is the unity of a human life. A human lifespan is not just a series of disaggregated and unrelated sequences, but an integrated whole, each of whose parts contributes to the sum. So we need a code that will serve to integrate the parts into a whole, and traits that will do the same.¹⁴ The third crucial difference is that human beings are reflective organisms with a conception of, and need for, self-worth. The life we lead is one we reflect on, and one that we evaluate; the worth we ascribe to our purposes affects the quality of our lives. So we need a code that allows us to live lives we can endorse as worthy. Finally, human beings are social beings, but our mode of sociality differs fundamentally from that of the other social animals, precisely because of the differences of form and content in the application to us of the conditionality thesis. So a genuinely human code of social interaction will (at a minimum) have to reflect the preceding facts.

A code of values, then, is a system of principles, adopted by choice and internalized in character, which fully satisfies the uniquely human version of the problem of discovering and tracking values: *By what code does a long-lived, fully integrated, self-respecting social animal preserve itself?* The question is analogous to the questions that a biologist might ask of sunflowers, amoeba, wasps, lizards, lions, or dolphins. We need to preserve ourselves across a lifespan, as they do, but our mode of life differs from theirs to the extent that our life-requirements differ from theirs. We need morality to do for us what genetics does for them. That answers Question (2).

I can only touch here on Question (3): what code of values do human beings need? A moral code of values is a set of principles that aims to maintain an individual's existence by maintaining that individual's adherence to a human mode of life—i.e., to the way that the principle of conditional necessitation applies in the human case. The principles take the form of what Smith calls “fundamental prescriptive generalizations” (VV, p. 165), identifying the values that preserve and unify a full human life and enjoining action on their behalf. The traits take the form of self-beneficial virtues, geared to achieving the values in question.

The most fundamental principle on which the entire code is based is its ultimate standard of value—survival *qua* man. In Ayn Rand's words,

¹⁴ Cf. Christine Korsgaard, *The Sources of Normativity*, pp. 229-233.

Man's survival qua man means the terms, methods, conditions, and goals required for the survival of a rational being through the whole of his lifespan—in all those aspects of existence which are open to his choice.¹⁵

This standard, on Rand's view, gives rise to the need for three cardinal values: reason, purpose, and self-esteem, along with their corresponding virtues, rationality, productiveness, and pride. The cardinal virtues and values as Rand puts it are jointly "the means to and realization of" the ultimate value, life.¹⁶ They are, we might say, the actions and traits that make sustained adherence to the principle of conditional necessitation possible.

The standard of survival qua human also gives rise to other virtues, whose function in each case is to give us a constant, automatized (though not automatic) awareness of the facts required to pursue our values and maintain a commitment to them. Though the requirements of each of the virtues is complex, each has what might be called a "teleo-epistemic" function, facilitating in different ways a commitment to what I earlier called the reality-orientation. That commitment is volitional; it requires effort to maintain. Obviously, traits that influence and facilitate such effort will be assets to an entity aiming to do so. Virtue on this conception is therefore fundamentally a matter of applied epistemology. Among the major virtues Rand mentions are independence, integrity, honesty, justice, and benevolence,¹⁷ each described as adjuncts of rationality, and each described in this "teleo-epistemic" way.

This short outline of the Smith-Rand argument should alert us to the sheer distance between the Objectivist view of what ethics is about and the view we find in mainstream philosophy. In fact, a mainstream ethicist might well wonder whether I've been talking about *ethics* at all in this section. The reason for the puzzlement, I suspect, derives from the distance between the claims of the conditionality thesis and the standard operating concept of mainstream ethics, "the moral point of view." I turn to that issue next.

IV. Ethical Egoism Versus "the Moral Point of View"

The standard approach to morality in the mainstream literature is not to begin with the Question "Why be moral?" or anything like it, but to define morality

¹⁵ Ayn Rand, "The Objectivist Ethics," p. 26.

¹⁶ Ayn Rand, "The Objectivist Ethics," p. 27.

¹⁷ Benevolence is not on Rand's list of virtues in "The Objectivist Ethics," but she clearly endorses some such virtue in "The Ethics of Emergencies," (*Virtue of Selfishness*, essay 3) and in her remarks on the "benevolent sense of life" in *The Romantic Manifesto*, rev. ed., (New York: Signet, 1975), and in her novels.

circularly in terms of what is usually called “the moral point of view.”¹⁸ The moral point of view is a distinctive outlook on the world, which literally transforms the way it looks to the agent who adopts it. In doing so, one brings out the world’s morally salient features at the expense of its non-moral ones, filters out the latter, and adopts principles to govern one’s actions—usually in a social context.

One of the differentiating features of specifically moral principles on this view is their *impersonality*. In taking the moral point of view, one explicitly excludes considerations of self-benefit in assessing the moral worth of an action; one sees oneself as merely one member of the moral community, each one of whom is equally and agent-neutrally significant in one’s deliberations. Impersonality so conceived is typically thought to constitute morality as such, and is therefore neutral between conceptions of morality—neo-Kantian, Utilitarian, Contractarian, etc. The differences between these mainstream theories are differences *within* the moral point of view, not alternatives to it.¹⁹ On the mainstream conception, then, though self-interest is at least occasionally a permissible motive of action, it remains conceptually at odds with morality. Though we may make concessions to self-interest, it’s not itself a moral motive; to the extent that we insist on acting from a *specifically* moral motive, we must put self-interest aside.

¹⁸ The *locus classicus* is Kurt Baier’s *The Moral Point of View*, (Ithaca, NY: Cornell University Press, 1958); see also Baier’s “Moral Reasons and Reasons to be Moral,” in *Values and Morals*, edited by Alvin I. Goldman and Jaegwon Kim, (Dordrecht, Holland: Reidel, 1978).

Other notable accounts include R.M. Hare, *Freedom and Reason*, (Oxford: Oxford University Press, 1963); Thomas Nagel, *The Possibility of Altruism*, (Princeton, NJ: Princeton University Press, 1970), and *The View from Nowhere*, (Oxford: Oxford University Press, 1986), chs. 8-10; John Rawls, *A Theory of Justice*, (Cambridge, MA: Harvard University Press, 1971); William Frankena, *Ethics*, (Englewood Cliffs, NJ: Prentice Hall Press, 1973); J.L. Mackie, *Ethics: Inventing Right and Wrong*, (Harmondsworth, UK: Penguin Books, 1977), chs. 4-5; Alan Gewirth, *Reason and Morality*, (Chicago: University of Chicago Press, 1978); Peter Singer, *Practical Ethics*, (Cambridge, UK: Cambridge University Press, 1979); David Gauthier, *Morals by Agreement*, (Oxford: Oxford University Press, 1986); Kai Nielsen, *Why be Moral?* (Buffalo, NY: Prometheus Books, 1989); Christine Korsgaard, *The Sources of Normativity*, (Cambridge, UK: Cambridge University Press, 1996).

For a useful overview of the debate, see *The Definition of Morality*, (London: Methuen & Co., Ltd, 1970), edited by G. Wallace and A.D.M. Walker.

¹⁹ I think this remains largely true of so-called Virtue Theory as well, e.g., as described by Greg Pence in his entry “Virtue Theory” in *A Companion to Ethics*, pp. 249-258; see also the references he cites therein. At best, Virtue Theorists strike me as ambivalent about the impersonality constraint, not opposed to it.

Philosophers in the past few decades have begun to take issue with the stringency of the impersonality constraint, and with some of its particular demands, but none goes nearly as far as Smith in rejecting the constraint itself.²⁰ Against the impersonality thesis, Smith argues that the conditionality of life entails a commitment to *ethical egoism*: every moral obligation serves our objective self-interest; none fails to do so. More precisely, ethical egoism may be defined as the view that a moral agent should be the ultimate intended beneficiary of every action he takes.²¹ As Smith interprets it, ethical egoism entails that morality just *is* self-interest, and vice versa: “principled egoism is the only way to live” (in her words) because it’s the only way to be *moral*. This claim is perhaps the most controversial one in *Viable Values*, and is also the substantive ethical thesis for which Ayn Rand is most notorious.

How do we get from the conditionality thesis to ethical egoism? Smith puts the point succinctly. If life is the standard of value, self-benefit is the only rational and permissible motive: “Human beings survive *by* acting for their own benefit.” (*IV*, p. 155). Ethical egoism, then, is the first-personal expression of the conditionality thesis. Suppose I accept the conditionality thesis. In doing so, I make survival qua human my end; given the nature of human survival, the end in turn requires adherence to the principle of conditional necessitation. But that principle counsels that I adopt motivations consonant with the conditionality thesis, and reject all motivations at odds with it. It’s obvious, I think, why egoism is consonant with conditionality, while acting from the motives of duty or altruism are at odds with it. If my own good is my ultimate reason for acting, I have no reason to constrain my actions by norms positively subversive of my good (e.g., deontic side-constraints), and I have no reason to sacrifice my good to others who make no causal contribution to it. Neither altruism nor side-constraints could arise as a result of conditional necessities derived from the requirements of survival qua human, because adherence to neither brings me any survival-benefit. Hence no such action could be sanctioned by the standard of survival qua man, and every such action is to be ruled out. Following Smith, let’s call a life lived by this standard, a life of *flourishing*.

Flourishing so conceived involves a plan for an entire life, which itself involves a commitment to principled action internalized by self-

²⁰ Classic expressions of this view include Bernard Williams, “Persons, Character, and Morality,” in *Moral Luck: Philosophical Papers 1973-1980*, (Cambridge, UK: Cambridge University Press, 1981), pp. 1-19; Susan Wolf, “Moral Saints,” *Journal of Philosophy*, vol. 79 (August 1982), pp. 419-439; John Cottingham, “The Ethics of Self-Concern,” *Ethics*, vol. 101 (1991), pp. 798-817; and Jean Hampton, “Selflessness and the Loss of Self,” in *Self-Interest*, edited by J. Paul, F. Miller, and E. F. Paul, (Cambridge, UK: Cambridge University Press, 1993), pp. 135-165.

²¹ I owe this definition to Allan Gotthelf.

beneficial virtues. I mentioned some of the specific virtues earlier, but a general point about the role of virtue is worth stressing. The virtues, recall, give us an automatized grasp of the requirements of survival *qua* man. In doing so, they call for a principled commitment to virtue across the whole of one's lifespan, a commitment that requires that the agent make principled action an ineliminable part of his moral identity. As Smith argues, this commitment to virtue is strong enough to rule out the (apparent) value of "ill-gotten gains" derived from violations of moral principle: "Something is in a person's interest only if it offers a net benefit to the person's life. Since a person's life is not reducible to any isolated element of his condition, we cannot fasten on such elements to draw conclusions" about his interests (*IV*, p. 168).²² Note, however, that virtue is not "an isolated element" of one's condition, but rather something one needs across the whole of one's lifespan. The requirements of virtue, then, are always part of the calculation of "net benefit," and they override (rather than merely outweigh) vice. So *this* momentary departure from moral principle—this isolated, secret act of vice—no matter how tempting or apparently beneficial, is never in my interest.

Egoism so conceived is a "positive-sum game" (Smith's term) that holds out the possibility of a harmony of interests in which the virtue of each agent becomes an asset for that of every other. The harmony is underwritten by the fact that each virtuous agent sees his benefit in the virtues rather than the vices (or involuntary weaknesses) of others. A virtuous egoist benefits neither by sacrificing his interests to others, nor by preying on their vices or vulnerabilities, but by interacting with others in order to gain from their virtues

²² Though I generally agree with it, Smith's statement of the ill-gotten gains thesis is somewhat ambiguous on two counts.

- (1) Smith says that ill-gotten goods lack value, but does not explicitly tell us *for whom* they lack it. She is clearly committed to the thesis that ill-gotten goods lack value for the practitioner of wrongdoing. But we sometimes benefit from ill-gotten goods that others have acquired illegitimately without ourselves doing so; it's unclear to me what Smith would say about such cases. Would Smith say that we do not benefit from goods manufactured through slave labor, or goods that owe their existence to compulsory redistribution? Do I derive no benefit from the shirt I'm wearing if it was made (unbeknownst to me) in a slave labor camp in some Third World country? The problem because sharper given Smith's rejection elsewhere of compulsory redistribution (cf. *Moral Rights and Political Freedom*, [Lanham, MD: Rowman and Littlefield, 1995]). Are redistributed goods *never* beneficial? They sometimes at least appear so: think of the roads we all drive on, the electricity we all use, or the compulsory public health inspections that save us from illness. If Smith wants to deny that such "goods" are genuinely beneficial, we need a more detailed account of the matter.
- (2) Smith does not explicitly integrate her account of ill-gotten gains with the Objectivist thesis that moral principles apply differently in emergencies than they do in the rest of life. If I steal a boat to save someone drowning in a lake, is that use an ill-gotten gain?

and strengths, and by implication, from their self-interest. The rational pursuit of self-interest, then, is never the *cause* of conflicts of interest, because virtue constitutes our self-interest, and virtue can never be the cause of conflict. Conflicts arise not through the pursuit of objective interests, but through violations of the egoistic principles that secure our interests (*IV*, pp. 174-187).

This picture of egoism is very distant from the one we find in the mainstream literature. Kurt Baier's argument for excluding egoistic considerations from the moral point of view is typical of the mainstream conception of egoism:

Morality is designed to apply in just such cases, namely, those where interests conflict. But if the point of view of morality were that of self-interest, then there could never be moral solutions of conflicts of interest. However, when there are conflicts of interest, we always look for a 'higher' point of view, one from which such conflicts can be settled. Consistent egoism makes everyone's private interest the 'highest court of appeal'. But by the 'moral point of view' we mean a point of view which is a court of appeal for conflicts of interest. Hence it cannot (logically) be identical with the point of view of self-interest.²³

In a similar vein, James Rachels's polemic against egoism, which we can take to be representative of a type, is notorious. An egoist, Rachels argues, would have no reason not to burn down a building on whim, just for the spectacle of the fire and thrill of burning everyone inside to death. After all, if he *found* the action self-beneficial, what principle of egoism could one invoke to condemn him?²⁴

These arguments are the standard ones in the literature, repeated in virtually identical form wherever they are found. Despite their ubiquity, however, they are, I think, remarkably weak arguments, proceeding as they do from a common premise that none demonstrates—viz., that consistent egoism *does* in fact lead to conflicts of interest. The assumption that underwrites this claim is the belief that egoism involves a form of unconstrained preference-satisfaction. Since it is, "egoism" becomes a kind of black box, in both metaphorical senses of "black": inscrutable and nefarious. On this view, the agent's "self-interest" consists in doing whatever he wants, however he wants.

The problem with such accounts of egoism is not merely that the claim about conflicts of interests goes undefended; it's that we've been given no way of giving any content *at all* to the concept of "self-interest." Nothing in

²³ Kurt Baier, *The Moral Point of View*, p. 190.

²⁴ James Rachels, "Egoism and Moral Skepticism," p. 406 in *Vice and Virtue in Everyday Life: Introductory Readings in Ethics*, (San Diego, CA: Harcourt Brace Jovanovich, 1989), edited by Christina Hoff Sommers and Fred Sommers.

these discussions tells us *what* egoism prescribes (or why), or why anyone would think it rational, much less why it must lead to conflicts of interest. The ultimate court of appeal about the content of egoism reduces instead to linguistic consensus: “we” all speak as though consistent egoism led to conflicts of interest; since “we” generally think that “our” speech mirrors the way the world is, “we” should assume that ordinary language proves that egoism leads to conflicts of interest.

In the space I have available, let me simply offer a laundry list of reasons why I think these arguments fail, taking Smith’s account and discussion of egoism as my point of departure.

For one thing, it’s unclear why the claims of ordinary English language such assume such extraordinary epistemic powers in defining the nature of human self-interest. After all, self-interest is not so conceived in other languages—e.g., Attic Greek, ancient Arabic, or modern Urdu, for instance. In any case, if the conditionality thesis is true, our lives are structured by a determinate set of interests, whose nature is as discoverable as the requirements of health. We would (or at least should) surely be alarmed to discover that medical schools were teaching their students that the principles of clinical diagnosis and practice were based on nothing but the “linguistic intuitions of the medical community in reflective equilibrium.” Why should it be any different in ethics? Reflective equilibrium is no better a method for discovering methods of medical practice than it is for discovering methods of ethical practice. Indeed, it’s not a method of discovery at all.

Second, and relatedly, virtually nothing in the literature on egoism focuses on what actual egoists—e.g., Aristotle, Spinoza, Ayn Rand—have actually said about what our interests *are*. The preference is to concoct notional confrontations with notional egoists, bearing no relationship to actual defenders of the view. When this literature occasionally does make reference to historical egoists, it typically downplays their egoism; when it makes reference to contemporary egoists, like Ayn Rand, it typically distorts their claims in fantastic ways, or ignores important aspects of what they say. Perhaps Smith’s book will help to remedy this.

Third, there is little appreciation in this literature of the need or nature of planning an egoistic *life-plan*, taking one’s life as a single unit of long-term deliberation, as opposed to thinking of one’s life as a series of disaggregated and discrete moments. Consequently, the focus in the literature is always on problems of choice in the context of range-of-the-moment situations, as in Rachels’s arson example. But if Smith’s argument is right, an egoist is someone who adopts value-tracking principles and virtues so as *not* to live by the range-of-the-moment. Confronted by a supposed egoist with a yen for arson, as in Rachels’s example, the obvious question is: Why would arson promote his life as described in Smith’s or Rand’s theory? What life-promoting principles does it express? Such questions are never posed in the mainstream literature, even by implication.

Fourth, there is little or no discussion of the place of *virtue* in such a life or life-plan; when the egoist's values are not straightforwardly depicted as psychopathological, discussions of egoism almost always concern the unjust seizure of external goods (or exploitation of persons for this purpose). Such claims miss the point: on an egoistic view, virtue is in our interest, and the external goods are only valuable when they're gotten by virtuous means (*IV*, pp. 164-187).²⁵ So such examples are irrelevant unless they discuss the egoistic conception of virtue on its own terms.

Finally, there is little appreciation for how large a generalization the conflicts of interest claim is. Egoism, Baier tells us, is "among the *main causes* of our social problems."²⁶ But where does one get evidence of such a thing? None is ever provided, apart from adducing discrete examples of conflict-situations. Isolated examples, however, don't prove that something is "among the main causes" of something else; at best, they serve to confirm the initial impression that it is. If virtue as previously described is in our self-interest, how could it be among *any* of the causes of our social problems? This, too, is left unaddressed.

V. Categoricity Versus the Choice to Live

As suggested in the previous section, mainstream ethics is based on "the moral point of view." One of its differentiating features, as remarked there, is a commitment to impersonality. Another is the *categoricity* of moral obligation.

On the standard conception of morality, whatever their superficial grammatical form, moral principles are categorically binding, as opposed to conditionally so: morality consists of an important class of duties that one owes without having been incurred by any voluntary act. One just *has* them, period. As Rawls puts it,

[I]t is characteristic of natural duties that they apply to us without regard to our voluntary acts... Thus if the basic structure of society is just... everyone has a natural duty to do his part in the existing scheme. Each is bound to these institutions independent of his voluntary acts, performative or otherwise.²⁷

Likewise Stephen Darwall:

While other elements of Kant's theory of morals... perhaps reflect rather poorly any very widely shared view about morality, his insistence that moral

²⁵ Except in emergency situations, as discussed in "The Ethics of Emergencies," in *The Virtue of Selfishness*, essay 3.

²⁶ Kurt Baier, "Egoism," in *A Companion to Ethics*, p. 197.

²⁷ John Rawls, *A Theory of Justice*, pp. 114, 115.

requirements are categorical imperatives expresses our common sense about an important part of ethics.²⁸

And James Rachels:

Commonsense morality would say, then, that you should give money for famine relief rather than spending it on the movies.

This way of thinking involves a general assumption about our moral duties: it is assumed that we have moral duties *to other people*—and not merely duties we create, such as by making a promise or incurring a debt. We have natural duties to others *simply because they are people who could be helped or harmed by our actions*.²⁹

Such claims are widespread in mainstream ethics. The language of “binding” suggests the metaphor of a sort of cord that “ties” the agent inescapably to various duties. But the very description of a natural duty suggests that there is no way to explain how the cord came to be placed around us.

On the Objectivist view, by contrast, moral obligation is conditional all the way down. In fact, its conditionality rests on the individual’s choice. As Rand puts it, “To live is [man’s] basic act of choice. If he chooses to live, a rational ethics will tell him what principles of action are required to implement his choice. If he does not choose to live, nature will take its course.”³⁰ Smith writes, “the choice to live is prerational. It is a presupposition of the standards of rationality” (*IV*, 107).

The categorical conception of obligation poses a challenge to the Smith-Rand view of moral obligation, which makes moral obligation ultimately conditional on a “choice to live.” If the Objectivist view is really “objective,” how can morality’s binding force rest on a choice? Doesn’t it then collapse into subjectivity? We can make this objection more specific by subdividing it into two sub-objections. The first objection concerns the *escapability* of conditional obligations; the second concerns their apparent *arbitrariness*.

The escapability objection goes as follows. Either moral obligation is categorical or it’s merely conditional on a choice. If it’s categorical, the agent has obligations independently of her choices. If it’s conditional, by contrast, her obligations depend on her choice. But if obligation depends on choice,

²⁸ Stephen Darwall, *Impartial Reason*, (Ithaca, NY: Cornell University Press, 1983), p. 173.

²⁹ James Rachels, *The Elements of Moral Philosophy*, (New York: McGraw Hill, 1993), p. 76.

³⁰ Ayn Rand, “Causality Versus Duty,” in *Philosophy: Who Needs It*, p. 99.

what if the agent doesn't make the relevant choice? Then she's not obliged. But surely morality, and especially justice, must be "inescapable." Hence moral obligation must be categorical, not conditional.

The arbitrariness objection goes as follows. Either moral obligation is conditional or categorical. If it's conditional, it depends on a choice. But suppose someone decides not to make that choice. Then she's beyond the reach of morality. Imagine now that we try to persuade her to be moral. Clearly, to persuade her, we must appeal to *reasons* to make the relevant choice. But in the nature of the case, there are no reasons to give; the choice is "prerational." Hence the choice must be arbitrary, and must make morality arbitrary. Since morality is not merely arbitrary, its binding force must not rest on a choice.

It's clear that for both Smith and Rand, though the conditionality of life explains the overridingness of morality, its binding force for any given individual is conditional, not categorical. So if the two preceding objections were sound, they would apply to, and undermine, the Objectivist view. Though Ayn Rand discusses the issue in her essay "Causality Versus Duty," (where she articulated the view), she never addressed the preceding objections as such. In *Viable Values*, Smith couches her own account of the choice to live as a series of discrete responses to objections (*VV*, pp. 103-117) without, in my view quite offering a positive interpretation of the idea. In what follows, I offer my own thoughts on the issue, which I believe are an extension of Rand's view, and are compatible with Smith's.³¹

The key to understanding the "choice to live," as I see it, is to think of the binding force of an ultimate value by analogy with the binding force of a logical axiom, on the Aristotelian conception of an axiom (cf. *VV*, p. 107).³² So conceived, an axiom can be thoroughly conditional in its binding force without being either escapable or arbitrary in the senses implied by the objections just described. Though ideally, it would have been best to compare Smith's account of the choice to live with Ayn Rand's conception of "axiomatic concepts," to avoid a lengthy exposition of that concept, I'll use the more familiar example of Aristotle's defense of the Principle of Non-Contradiction (PNC).

At *Metaphysics* IV.3, Aristotle enunciates the PNC: "A thing cannot be and not-be at the same, in the same respect." This is an undeniable and foundational truth; its truth is merely re-affirmed in the attempt to doubt or deny it. Note, however, that the Principle is not a categorical injunction to engage in thought. In fact, it says nothing at all about thought, nor is it a

³¹ I owe the inspiration for my view to Allan Gotthelf's 1990 Ayn Rand Society paper, "The Choice to Value," but he is not responsible for my way of formulating the issues.

³² On axioms, see Aristotle, *Posterior Analytics*, I.2, *Metaphysics* IV.3, *Nicomachean Ethics* VI.2.

prescription of any kind. It merely states a fact about the world—one that becomes a guide for thought when and only when one *chooses* to think. In choosing to engage in thought, one sees in one's own case that *if* one is to do so successfully (i.e., at all), one *must* obey the PNC without exception. An isolated attempt to evade the Principle would be self-subverting. A wholesale attempt to evade it would render one entirely unable to think. And any attempt intermediate between these would be *both* self-subverting *and* render one incapable of thought, in proportion to the extent of the violation. In choosing to think, then, one is therefore bound by the Principle without ever having been forced or enjoined to do so. Indeed, given the conceptual relation between the PNC and thought, one can choose to be bound by it without ever having explicitly formulated it as a principle.

Is the PNC “escapable”? In one sense, yes; in another sense, not really. One *can* escape the PNC—if one is willing to pay the price. The PNC binds all thought; one way to evade it, then, is simply to stop thinking. And refusing to think *is* a way of escaping the PNC. The PNC applies *if* (and only if) one aims at thought. It *doesn't* apply to a non-thinker. On the other hand, its non-application to the non-thinker is hardly a threat to its logical or epistemic authority. A non-thinker can't constitute a problem for someone espousing the PNC, because she can't even raise whatever objection she might have to it. In fact, she can't even have an objection she refuses to raise, since even having a voiceless objection requires putting the objection in words, which calls the PNC back into operation. So the “problem” of escapability is no problem at all: the price of escape is such that it cannot, in principle, pose a threat to the Principle.

Is the PNC “arbitrary”? Not at all. To be sure, there's no *argument* for the PNC that proceeds from principles that are epistemically prior to and independent of it. But that's because there *are* no such principles. The PNC is an epistemically basic principle; all principles presuppose its truth, and none could in principle support it. The *Oxford English Dictionary* provides six relevant entries for the term “arbitrary”: “to be decided by one's liking,” “dependent on one's will or pleasure,” “at the discretion or option of anyone,” “derived from mere opinion or preference,” “capricious, uncertain, varying,” “unrestrained in the exercise of will.” I think it's clear from these entries that the concept of “arbitrariness” only applies in contexts where we face two or more alternatives and choose one while ignoring *the established principle that adjudicates between them*. In the nature of the case, this stricture couldn't apply to the PNC: putting aside the Law of Identity (which is equally basic),³³ the choice between accepting or rejecting the PNC couldn't be governed by our choosing some *other* principle which told us to accept or reject it. There is no such principle, and couldn't be one. So the PNC isn't arbitrary, because it applies prior to the proper application of the very concept of “the arbitrary.”

³³ I mean A=A, not Leibniz's Law.

The concept “arbitrary” only has meaning subsequent to accepting the PNC’s axiomatic status.

The analogy to moral obligation applies as follows. As a matter of non-prescriptive fact, life can only be kept in existence by a constant process of self-sustaining action. Moreover, life is unique in this respect: it’s the ultimate generator of practical requirements that explains why there are practical requirements at all. Let’s assume, on the basis of the argument in the third section of this paper, that life is the ultimate value in this sense.

Life’s conditional character, however, is not by itself a prescription. It’s simply a fact about the world. The fact by itself generates no categorical duty to keep one’s own life (or anyone else’s life) in existence, or indeed, to value or do anything at all. Life’s conditional character has prescriptive force for an agent when and only when the agent *chooses* to value and live, i.e., chooses to engage in goal-directed action, and chooses to promote her life by goal-directed action.

In choosing to live, one is conditionally bound by the requirements of life. The momentary attempt to evade that fact would lead one to practical inconsistency: if I will life as an end, I must will the means to it; if I refuse to will the means, I must give up the end. So a momentary act of evasion would stand condemned as a single violation of the principle of conditional necessitation. The wholesale attempt to evade the conditionality of life, of course, leads the evader to non-existence. And an intermediate attempt to live by combining rationality with irrationality or vice with a commitment to egoistic virtues would be futile: each commitment would either undercut the other, or else fail to yield the benefits that genuine commitment would bring. In choosing to live, then, one is bound to realize one’s life without ever having been required to do so by anything besides the choice to live. Note that as with the PNC, a person could well make the choice without making it under the description of “the choice to live” as described by the Objectivist literature (just as one can make the choice to be bound by the PNC without doing so in Aristotle’s words).

Is the choice to live escapable? Again, yes and no. It’s escapable in the sense that one can, in principle, fully opt out of the task of aiming at one’s self-preservation. But such a person is no more a threat to the authority of moral norms than the non-thinker is to the PNC. The person who denies the PNC is rendered incapable of thought; the person who denies the conditionality of life is rendered incapable (in the same sense of “incapacity”) of goal-directed action. The PNC-denier is rendered incapable of thought in proportion to her genuine commitment to denying the PNC. The life-denier is rendered incapable of action in proportion to her commitment to that denial.

Is the choice “arbitrary”? No. “Arbitrariness” is a matter of flouting some established principle. But once we take a regress of justifications back to an ultimate principle, there are no prescriptive principles left in terms of which to adjudicate the choice to adopt the ultimate principle. (This doesn’t mean that there is nothing to say about why the ultimate principle is ultimate.) *Ex*

hypothesi, the ultimate principle is survival qua human, so the ultimate choice is the choice to live. The choice, however, is hardly arbitrary. If life really *is* ultimate, then only *it* justifies normative standards and rational action. Whether or not it is ultimate is a matter of argument, and the argument of section II either succeeds or fails in this regard. If it succeeds, as I believe it does, then acceptance of life's ultimacy is the only thing that makes rational action possible. Since acceptance of that fact is ultimate, there is nothing arbitrary about accepting it in the only way it can be accepted. By contrast, if one rejects life, one has no legitimate claim to valuation based on standards at all. I conclude, then, that the "choice to live" is not arbitrary in any legitimate sense of that term.³⁴

VI. Conclusion

One of the great frustrations of anyone who has made his way through the literature on the Question "Why be moral?" is the fact that there are so few genuine attempts to answer it. What one too often gets in the literature is the admonition that the Question ought not to be asked because it's immoral to ask it—or the claim that what the Question ultimately requires is "the secular equivalent of faith in God" or recourse to "the sanctions of law" to persuade the recalcitrant.³⁵

One of the merits of *Viable Values* is that, in offering a clear and accessible exposition of Rand's theory, it offers us neither faith nor force as an answer to the question it poses, but a satisfying and responsive one. Our own lives, as Smith puts it, are "the sum of morality's claims on us." "Life sets the standard of value, life is the goal of morality, life is the reward of morality. What stronger answer can one imagine to the question of why we should be

³⁴ I take the preceding argument to answer the objections raised in Eric Mack's "The Fundamental Moral Elements of Rand's Theory of Rights," pp. 133-35 in *The Philosophic Thought of Ayn Rand*, (Urbana and Chicago: University of Illinois Press, 1986), edited by Douglas J. Den Uyl and Douglas B. Rasmussen.

Readers familiar with the debate between Philippa Foot and John McDowell on "morality as a system of hypothetical imperatives" may see a similarity between Rand's position and Foot's. Though Rand's position is closer to Foot's than to McDowell's, I think Foot's position falls prey to the "arbitrariness" objection discussed in the text in a way that Rand's does not. See Philippa Foot, "Morality as a System of Hypothetical Imperatives," in *Virtues and Vices*, (Berkeley: University of California Press, 1984), pp. 157-73; John McDowell, "Are Moral Requirements Hypothetical Imperatives?" *Proceedings of the Aristotelian Society*, vol. 52 (Suppl. 1978), pp. 13-29.

³⁵ The first quotation comes from Annette Baier, "Secular Faith," pp. 204-5 in *Revisions: Changing Perspectives in Moral Philosophy*, (Notre Dame: University of Notre Dame Press, 1983), edited by Stanley Hauerwas and Alasdair MacIntyre. The second comes from Peter Singer, *Practical Ethics*, p. 220.

moral?” (*IV*, p. 187). As I hope to have made clear in this paper, I don’t think there is one.³⁶

³⁶ This paper originated as a contribution to a book symposium on *Viable Values* arranged by the Ayn Rand Society, at the APA Eastern Division meeting in New York, December 28, 2000. I thank the symposium participants, Julia Driver, David Schmidtz, and Tara Smith, as well as the audience, for stimulating discussion on that occasion. I’d also like to thank Robert Hartford and Gregory Salmieri for detailed written comments on the paper, and Carolyn Ray for the opportunity to present it at the Enlightenment Online Conference, March 2001. Finally, a special thanks to Allan Gotthelf, who put the ARS program together, and whose comments on earlier drafts improved the paper immeasurably.