



AGENT-RELATIVE RESTRICTIONS AND AGENT-RELATIVE VALUE

BY STEPHEN F. EMET

JOURNAL OF ETHICS & SOCIAL PHILOSOPHY

VOL. 4, NO. 3 | SEPTEMBER 2010

URL: WWW.JESP.ORG

COPYRIGHT © STEPHEN F. EMET 2010

Agent-Relative Restrictions and Agent-Relative Value

Stephen F. Emet

IN THIS PAPER, I POSE A CHALLENGE for attempts to ground all reasons in considerations of value. Some believe that all reasons for action are grounded in considerations of value. Some also believe that there are agent-centered restrictions, which provide us with agent-relative reasons against bringing about the best state of affairs, on an impartial ranking of states of affairs. Some would like to hold *both* of these beliefs. That is, they would like to hold that such agent-centered restrictions are compatible with a view that grounds all reasons for action in considerations of value. This is what I will argue is problematic.

My argument challenges a particular project, of showing that all ethical theories are broadly consequentialist. Proponents of this project claim that all ethical theories can be captured by the claim that what one ought to do is to perform that act which would bring about the optimal outcome.¹ Theories would merely differ on how they go about determining the ranking of outcomes. The idea is to take whatever other factors the deontological theory claims are relevant, and to work those factors into the evaluative ranking of outcomes. In this way, one has *consequentialized* the theory. An agent-neutral theory would claim that the ranking of outcomes is the same for everyone. But a theory that accorded more closely with common sense would have *agent-relative rankings*; how outcomes are ranked would vary from agent to agent. The attraction of this strategy, according to proponents, is that it would preserve what is compelling about consequentialism – its teleology and maximizing² – while also preserving more of commonsense morality than consequentialism does.³

I will call theories that ground all reasons in considerations of the good “teleological theories.” My claim is that agent-centered restrictions will not fit into a teleological theory, understood as one which grounds all reasons in considerations of the good. If the correct moral theory is a teleological one, then there are no agent-relative restrictions. If there are agent-relative restrictions, then teleology is false.⁴

¹ See, for instance, A. Sen, “Rights and Agency,” *Philosophy & Public Affairs* 11.1 (1982): 3-39, and “Evaluator Relativity and Consequential Evaluation,” *Philosophy & Public Affairs* 12.2 (1983): 113-132; J. Broome, *Weighing Goods* (Cambridge, MA: Blackwell, 1991); J. Dreier, “The Structures of Normative Theories,” *The Monist* 76 (1993): 22-40.

² S. Scheffler, “Agent-Centered Restrictions, Rationality, and the Virtues,” *Mind* 94.375 (1985): 409-419.

³ However, such an agent-relative approach does sacrifice the impartiality that is also often thought to be an appealing feature of consequentialism.

⁴ My purpose in this paper is not to argue for or against restrictions, but just to show their incompatibility with a certain conception of reasons for action.

1. Agent-relative Restrictions and Consequentializing

In this section, I show why those who seek to develop a teleological theory must accept agent-relative value if they are to capture agent-relative restrictions. Agent-relative restrictions are restrictions on what one can do to bring about the (agent-neutrally) best state of affairs.⁵ Agent-relative restrictions are thought to be an important element of commonsense morality that agent-neutral theories fail to capture – the other being agent-relative options, permissions to act in ways that do not bring about the best state of affairs.⁶ For instance, commonsense morality recognizes a restriction against the killing of innocent persons. It is impermissible, commonsensically, to kill an innocent individual in order to harvest his organs, even though doing so would enable us to save five other lives. Of course, it is open for a theory with a pluralistic axiology and no agent-centered restrictions to claim that such killings are particularly bad, worse than mere deaths, so that killing the innocent individual would *not* in fact produce the most good.⁷ But such a theory, although it captures the previous case, does not capture the following case: Suppose that killing one innocent individual is the only way to stop the killings of two other innocent individuals. It seems that commonsense morality still recognizes a restriction here – and many ethical theorists, for instance, Kantians, believe there to be a restriction here – but any agent-neutral teleological theory will not recognize such a restriction.⁸

The restriction against killing an innocent person says that at least sometimes one may not kill an innocent person, even if this would prevent the killings of multiple other innocent people, and there are no other morally relevant circumstances.⁹ Thus, (1) it is impermissible for John to kill one person, even if this will prevent George from killing two other people. And (2) it is impermissible for George to kill one person even if this will prevent John from killing two persons. For a teleological theory to capture (1), it must claim that the state of affairs in which John kills one is worse than the state of affairs in which George kills two. But to capture (2), a theory must claim that the state of affairs in which George kills one is worse than the state of

⁵ Scheffler, *Rejection*, pp. 2-3, 80-81.

⁶ Scheffler, *Rejection*; S. Kagan, *The Limits of Morality* (Oxford: Oxford University Press, 1989).

⁷ M. Schroeder, "Teleology, Agent-Relative Value, and 'Good,'" *Ethics* 117 (2007): 266.

⁸ An agent-neutral theory might claim that such killings are of infinite disvalue. On such a theory, I would then have an option. If I were faced with the choice of intentionally killing one to prevent the intentional killing of two, I could permissibly do either, or perhaps I should flip a coin. Such a view is rather implausible, since it would have to claim that the outcome in which a million are killed is no worse than the outcome in which one is killed.

⁹ Scheffler, *Rejection*, 80. Restrictions may come in various thresholds. If the restriction were absolute, then John must not intentionally kill an innocent person no matter what – even if it was the only way of preventing the annihilation of all sentient life in the universe. Such an infinite threshold seems implausible. Accounts may differ on where the threshold is, and why the threshold is where it is.

affairs in which John kills two. But if the theory is agent-neutral, then it applies the same ranking to all agents. Thus, an agent-neutral teleological theory cannot capture both (1) and (2).¹⁰ Thus, to capture such a restriction, which affirms both (1) and (2), within a value-based account, one must employ an agent-relative account of value. Then one may claim that it is worse, *relative to John*, for John to kill one than for George to kill two, and it is worse, *relative to George*, for George to kill one than for John to kill two.

Theories that incorporate a restriction on what one can do to bring about the state of affairs at the top of an agent-neutral evaluative ranking claim that one sometimes has decisive reason not to perform (agent-neutrally) optimal actions. This is an agent-relative reason. John has an agent-relative reason that *he* not kill, which is stronger than any agent-neutral reason he has to prevent George's killings. Likewise, George has special reason not to kill, which is stronger than any agent-neutral reasons he has to prevent John's killings. If these reasons are to be grounded in considerations of value, the evaluation must be agent-relative.

2. The Assumption

I now state a particular assumption, which I believe plausible, that I think undermines attempts to capture agent-relative restrictions in terms of agent-relative value:

- (A) For all possible outcomes in which there exist agents capable of having pro- and con-attitudes: (i) worthiness of favor is a strictly increasing function of goodness of outcomes; (ii) worthiness of disfavor is a strictly increasing function of badness of outcomes.¹¹

Thus, if some outcome O_1 is better than outcome O_2 , then O_1 is worthy of greater favor than O_2 . I will use "favor" to signify any sort of pro-attitude, and "disfavor" any sort of con-attitude. (A) is supposed to capture the apparent connection between (dis)value, and worthiness of (dis)favor. This assumption relates to the fitting-attitude analysis of value, but does not claim that worthiness of (dis)favor provides an *analysis* of value. Further, since my aim is not to provide an analysis of value, it is weakened further by restricting the class of outcomes it ranges over in order to avoid various challenges proposed for fitting-attitudes analyses of value.¹² As such, it is an assumption that both proponents and opponents of fitting-attitude analyses can share.

¹⁰ For similar arguments, on which this one is based, see Scheffler, *Rejection*, 87-90; Schroeder, "Teleology," 266-267; Portmore, "Consequentializing," 330-1.

¹¹ Thank you to an anonymous reviewer for this way of putting conditions (i) and (ii). This assumption is meant to be similar to FA*, in M. Zimmerman, "Partiality and Intrinsic Value," (forthcoming).

¹² See K. Bykvist, "No Good Fit: Why the Fitting Attitude Analysis of Value Fails," *Mind* 118 (2009): 1-30

The connection between value and worthiness of pro-attitudes is what is attractive about fitting-attitudes analyses.¹³ Even if Moore was right that the good is unanalyzable, we might still think that (A) is true. Further, some have proposed fitting-attitude analyses of agent-relative value.¹⁴ Indeed, given the challenges that other understandings of agent-relative value face, a fitting-attitude analysis may seem like the best hope of making sense of agent-relative value.¹⁵ And without some account of how agent-relative value relates to value *simpliciter*, the consequentialization project is rather uninteresting. So it is bad news for consequentializers if (A) is false, since they may need something stronger than (A). But (A) is itself, I argue, sufficient to generate problems for consequentializing agent-relative restrictions.

Recall that to capture agent-relative restrictions, we need to employ the notion of agent-relative value so that we can claim that it is worse, *relative to me*, for me to murder John than for George to murder John. By (A), however, this amounts to claiming that my murdering John is worthy of my greater disfavor than George's murdering John. But this strikes me as rather odd. I do not see how the outcome in which I murder John is *in itself* worthy of more disfavor, by me or by anyone, than the outcome in which George murders John.¹⁶

There may indeed be *reasons* to be especially concerned about our own wrongdoing – principally the fact that we have in general greater control over our own wrongdoing than the wrongdoing of others. Further, with wrongdoing come negative emotions like regret and guilt, and perhaps I have self-interested reason to be especially concerned about my own wellbeing. But even if there are reasons to be especially concerned with one's own wrongdoing, this alone does not imply that one's wrongdoing is worthy of greater disfavor.¹⁷

Suppose you have intentionally injured a patient. If you do not aid the patient, the patient will die because of your act. However, if you aid your patient you will be unable to aid two other patients, who have been similarly injured and will die if not aided. Traditionally, it is held that you ought not to murder one to prevent the murder of two others. If this is to be couched in terms of agent-relative value, by (A) it implies that the outcome in which your patient dies as a result of your intentionally injuring him is worthy of greater disfavor than the outcome in which two other patients die as a result of someone else's intentionally injuring them. This appears self-centered to me. Perhaps it is permissible for you to save your patient, because you would

¹³ Zimmerman, "Partiality."

¹⁴ E.g., T. Hurka, "Moore in the Middle," *Ethics* 113 (2003).

¹⁵ For strong criticisms of various attempts, including fitting-attitudes analyses, to make sense of agent-relative value, see Schroeder, "Teleology."

¹⁶ Note that worthiness of favor is not an exclusively agent-neutral notion. It seems to me that that a parent's child's suffering is worthy of the parent's greater disfavor than a distant child's suffering (even if, agent-neutrally, the sufferings are equally as bad).

¹⁷ Cf. Zimmerman's discussion of the Wrong Kind of Reasons problem, in "Partiality."

feel more regret and guilt about your patient than about the other two, and, even if irrational, this must be taken into account. Perhaps you are required to save your patient, since she has special claims on you in virtue of being *your* patient. But it is strange to say that your patient's being saved is worthy of greater favor. One could always claim that whereas non-relative value satisfies (A), agent-relative value is different, and that the relation between agent-relative value and worthiness of favor does not hold. But this is not an attractive response, for then the connection between agent-relative value and value simpliciter is very hard to see.

The worry I have is that I do not think that agent-relative value will support our intuitions about agent-relative restrictions. And this is because of the connection between value and worthiness of (dis)favor, and the fact that I find odd the idea that my wrongdoing is *worthy* of my greater disfavor than the wrongdoing of someone else. To say so suggests an overestimation of one's importance. Note that in standard cases of restrictions, one ought not to X even in order to prevent *very* many other Xings. One ought not to murder, even in order to prevent five or 50 other murders. But one would then need to claim that one's own wrongdoings are worthy of *vastly* greater disfavor than the wrongdoings of others. Even if worthy of *slightly* more disfavor, could one's wrongdoings really be worthy of such greater disfavor, as to support intuitions about agent-relative restrictions?

Here is another bit of evidence for the claim that one's own wrongdoings are not worthy of greater disfavor. The defender of the claim that one's own wrongdoing is worthy of special disfavor must presumably defend an asymmetry between wrongdoing and right-doing. That is, presumably there is, in general, no special reason to favor one's own good deed to a similar good deed of someone else.¹⁸ Someone who thought his own good deeds more worthy of favor than the good deeds of others would seem inappropriately self-centered. Suppose you must choose between a situation where you aid John (a stranger to you) and one where Jane (a stranger to you) aids John (who is also a stranger to Jane), and that John will be aided equally in both cases. You do not seem to have particularly good reason for choosing yourself. Now, that is not to say you would need particularly good reason here, since John will be aided equally in both cases. The point is simply that both outcomes seem worthy of similar favor. Suppose instead that, if Jane aids John, John will be made slightly better off. Here, it would seem inappropriate for you nevertheless to pursue your aiding John – you would seem like a glory-monger. There is nothing to be said for your aiding John that would warrant favoring it over Jane's aiding John to a slightly higher degree. But if your own right-doing was worthy of special favor, this would not be the case.

¹⁸ Special cases are perhaps parental obligations and other contractual obligations. Though, in the first case at least, the idea here is that others *cannot* do the good deed, that your child will be better off if *you* care for them, than if someone else cares for them, as this will further the parental relationship (which is perhaps good in and of itself) and will promote the child's own psychological wellbeing.

So your own right-doing is not worthy of special favor. So those who maintain that your own wrongdoing is worthy of special favor must provide some explanation as to why this is so when it comes only to wrongdoing.

Agent-relative restrictions already embody a similar sort of asymmetry. Such restrictions claim that we have special reason to avoid our own wrongdoing, but that we have no special reason to pursue our own right-doing. One ought not to harm others, even if doing so would prevent greater harms done by others. But no one claims that one ought to benefit others when this would prevent greater benefits done by others, except perhaps in cases involving parental and contractual obligations. And it might be hoped that, however one is to defend this asymmetry, such a defense will also account for the previous asymmetry.¹⁹

I am, however, not optimistic about this. Agent-relative restrictions are sometimes defended by appealing to the Kantian idea of inviolability.²⁰ As a person, one possesses a special kind of dignity, which gives one a degree of inviolability, and this inviolability makes it the case that one ought not to be used in various ways (for instance, as a mere means). Suppose then you are contemplating violating someone's rights, to prevent a greater number of rights violations. If this is permissible, then the very right (not to be used as a mere means, say) whose importance you are seeking to uphold by minimizing violations is negated. It would be permissible to use someone as a mere means of reducing the number of usings as mere means, which is just to say that one has no right against being used as a mere means. Regardless of whether or not this argument works, the point is that the agent-relative reason to be especially concerned about your own wrongdoing is only plausibly defended by referencing not some fact about *you*, the agent, but about the patient. It is in virtue of something about the patient, her dignity, her inviolability or the respect you owe her, or about the relationship between you and the patient, that the patient is making a particular claim on you, that generates the agent-relative reasons against wrongdoing.²¹

But such arguments need not, and should not, if they are to be plausible, reference anything about worthiness of disfavor. This merely invites the charge that agent-relative restrictions are objectionably self-centered and represent a misguided concern with one's own moral cleanliness. When couched in terms of agent-relative value, agent-relative restrictions seem to represent a kind of moral fastidiousness. But the defense of restrictions just considered does not appeal to any special importance you have as the agent, or any special value found in clean moral souls. Rather, it is supposed to reflect something we all have as persons. On such a defense, the "special concern" with

¹⁹ Thank you to an anonymous reviewer for suggesting this possible line of response.

²⁰ F.M. Kamm, "Non-Consequentialism, the Person as an End-in-Itself, and the Significance of Status," *Philosophy & Public Affairs* 21.4 (1992): 354-389.

²¹ Similarly, in cases of parental and contractual obligation, it is facts about the child, or the contractual party, and our relationship with that person, that generates the relevant agent-relative reasons.

wrongdoing, as opposed to right-doing, does not reflect some brute asymmetry between harm and benefit. Rather, the asymmetry between harm and benefit falls out of a prior notion of persons that the account seeks to reflect. Of course, violating someone's rights *is* worthy of disfavor. But this gets us no closer to seeing why the outcome in which you violate someone's rights is worthy of greater disfavor than the outcome with greater rights violations by others.

I am not claiming that the above defense of agent-relative restrictions succeeds – only that it will not account for the asymmetry between favoring and disfavoring. There are other conceptions of persons which do not support restrictions on harming. For instance, one might claim that, as persons, we are not inviolable, but *unignorable*.²² This is a less intuitive conception of persons; nevertheless, it is a competing conception, and we could debate which conception is better or more nearly captures whatever it is we are trying to capture.

Perhaps there is some other means of defending the asymmetry that one's wrongdoing is worthy of special disfavor while one's right-doing is worthy of no special favor. I do not see the prospects for such a defense. Further, even if such a defense could be provided, I question whether it would really amount to a *defense* of agent-centered restrictions. Rather, it seems to me that it would merely saddle such restrictions with a distasteful self-centeredness and would lead ultimately to the rejection of such restrictions. If one seeks to defend restrictions, rather than to merely make formal space for them, an account like the Kantian one mentioned here would be preferable.

The argument of this section, if correct, does two things. It shows that attempts to account for agent-relative restrictions in terms of agent-relative value cannot appeal to fitting-attitude analyses of value, since something weaker, namely (A), will not make sense of agent-relative restrictions. It also shows that such restrictions violate a plausible assumption about value, suggesting that the considerations which would generate restrictions are very different sorts of considerations, and that a theory which recognizes such restrictions must recognize reasons not grounded in considerations of value. A theory which does not admit those elements that would generate restrictions can approximate restrictions only with great artifice and awkwardness and can provide no defense of agent-centered restrictions, for their plausibility is lost when morphed into a claim about what is better or worse.

²² S. Kagan, "Replies to My Critics," *Philosophy & Phenomenological Research* 51.4 (1991): 920; K. Lippert-Rasmussen, "Moral Status and the Impermissibility of Minimizing Violations," *Philosophy & Public Affairs* 25.4 (1996): 340.

3. Objections

It might seem as if (A) leaves no room for anything but agent-neutral value. But, leaving aside agent-relative value for the moment, surely we understand claims about what is good and bad *for* a particular individual, even when such claims conflict with what is good and bad overall. It is bad for Jane if she gets fired from her job, even if it results in a better overall state of affairs – perhaps because Jane gets replaced by someone more efficient, who enables the company to generate more income and avoid greater lay-offs. And this might be taken to suggest either that (A) is false, or that (A) is of restricted scope, and so does not clearly apply to agent-relative value. But if (A) is false or does not apply to agent-relative value, then it is of no import whether agent-relative value violates (A) or not.

However, even if (A) were not to apply to what is good (or bad) *for* a particular individual, that alone would not show that it did not apply to agent-relative value – after all, what is good relative to me and what is good for me are supposed to be quite different. Second, I believe that (A) can be naturally extended to claims about what is good and bad for a particular person. In talking of what is good or bad for Jane, apart from what is good or bad overall, we take a different perspective, abstracting from all other relevant moral considerations, and focus just on Jane. Forgetting everyone else except Jane, how would this make things go? In other words, what would be worthy of favor or disfavor if Jane were the only entity with non-instrumental value? Alternately, suppose that my sole, overriding responsibility is to bring about the best of state of affairs for Jane that I possibly can. What then would be worthy of my favor or disfavor? By taking these more limited perspectives, where we focus on a single agent, we can see how what is good or bad for a particular person, rather than good or bad simpliciter, is still linked to the worthiness of favor or disfavor. But this extension of (A) to good-for will clearly not help those who wish to appeal to agent-relative value in order to consequentialize restrictions. For here we focus exclusively upon the agent and what would be best for her, but what is best for her might violate all sorts of agent-relative restrictions – for example, if an agent could kill her rich uncle, whom she has no tender feelings for, without getting caught. Perhaps we might instead construe our task as that of maximizing an individual's moral cleanliness to produce the best accounting we can, for that agent, in some moral ledger. But, again, this agent-focused perspective invites the charges of fastidiousness and self-centeredness discussed in the previous section.

A different sort of objection focuses on my claim that viewing my murder of John as worthy of greater disfavor than George's murder of John would be objectionably self-centered and would represent an overestimation of my own importance. This might seem ironic, for such charges of self-centeredness – and an obsession with one's own moral cleanliness – are often leveled against defenders of deontological restrictions. Is it not overly

self-centered, so the objection goes, to claim that you may not murder one, to prevent someone else from murdering five? What is so special about your murdering the one that gives you such great reason not to do it that you may not do it even to prevent five other similar murders? Succinctly, the objection is the following: if it is overly self-centered to *favor* killings by others more than killings by me, why is it not also overly self-centered to *allow* others to kill rather than myself killing (and thereby preventing their killings)? A consequentialist may grant the antecedent, which I have argued for in the previous section, but claim that this simply undermines deontological restrictions entirely.

What my argument suggests is that there may very well be circumstances in which we ought to favor an outcome which we may not bring about. There would thus be an asymmetry between what we ought to favor and what we ought to do. What could account for this asymmetry? The answer, I think, lies in the fact that, when others make claims on us, these are claims that we act or not act in some way or another. Others may have legitimate claims against me that I treat them in various ways,²³ but they do not have claims against me that I favor one outcome over another, except perhaps insofar as such a favoring attitude is instrumental to my acting in the way that satisfies their claim. The claim-rights of others bear directly on what we ought to do, because it is primarily through our actions that we run up against the claim-rights of others. What someone favors or disfavors is generally no business of mine, so long as that person acts toward me in an appropriate manner. If it is possible for me to favor one thing, yet do another, and if the claims of others can give me reason to act, then someone's claim may affect what I ought to do without affecting (or without affecting to the same degree) what I ought to favor.²⁴

In the case where I can prevent the killings of two people only by killing one person myself, my would-be victim's claim against me is that I not kill her. Suppose persons have a (non-absolute) right not to be used merely as a means. Then I may not kill one to prevent two others from being killed. In virtue of possessing this right, my would-be victim has a special claim against me – that I not use her merely as a means. If the claim against me not to be used is a genuine one, then it must give me special reason not to so act. In virtue of my having special reason not to so act in this way, I may acquire instrumental reason to disfavor my killing, if that will lead me to act in the way that I ought to. But my victim's claim is directly on what I do to her, not on what I favor or disfavor. There is no self-centeredness, since the reason why I may not kill one to save two has its source in the victim, not in me.

²³ Here I mean legitimate *moral* claims, not legal claims.

²⁴ If we always do what we most favor, then perhaps it will turn out that what we ought to favor and what we ought to do always coincide. But if there were creatures capable of intending one thing while favoring another, then situations could arise in which what they ought to do and what they ought to favor come apart.

A more complete response would outline all considerations which may differentially affect what we ought to do and what we ought to favor. I do not have such a complete response. Nevertheless, I hope to have sketched one rationale for the potential asymmetry between what we ought to do and what we ought to favor. More work would need to be done, but we should not assume, at the start, that any consideration which bears on what we ought to do bears on what we ought to favor, and vice versa.

Lastly, one might object that I have misunderstood the consequentializing project. I have construed this project as the attempt to reduce all reasons to considerations of value. D. W. Portmore, however, defends a consequentialization strategy which attempts to reduce all reasons for action not to *value*, but to *reasons to prefer* various outcomes.²⁵ However, my arguments apply against this project as well. One may indeed have self-interested reason to prefer the wrongdoing of others to one's own wrongdoing – if one acts wrongly, she may be punished or experience guilt, etc. But this can obviously not account for deontological restrictions, since sometimes violating such restrictions will very much be in one's interests. What one lacks is a reason *of the right kind* to prefer the wrongdoing of others over one's own wrongdoing. To have such a reason would be for one's wrongdoing to be worthy of greater disfavor than the wrongdoing of others, and this is just what I have argued is not the case.

Consider the agent-relative reason a parent has to prefer the outcome in which his child flourishes more than the outcome in which a stranger child flourishes. Assume that the agent-neutral reasons to prefer either outcome are roughly equivalent. It makes sense to say that, even if the outcome in which one's child flourishes is not better, agent-neutrally, one nevertheless has agent-relative reason to prefer that one's child flourishes.²⁶ And this is because part of the content of one's parental obligation is to be especially concerned with the wellbeing of one's children. So we can perhaps explain parental obligations in terms of agent-relative reasons to prefer. But this will not extend in the same way to deontological restrictions. Does the special reason I have to avoid wrongdoing really have anything to do with the outcomes I have reason to prefer? It seems not. Deontological restrictions do not require me to be especially concerned with the wellbeing of my potential victim. They just say that I cannot do certain things to people. Suppose that in virtue of my greater foresight, I see that if I break Jane's leg, she will be unable to go on a trip, on which she would suffer an accident and break both of her legs. I might try to alert Jane of this danger in some way, but it seems impermissible for me to break Jane's leg, against her will, even though it will

²⁵ D. Portmore, "Consequentializing Moral Theories," *Pacific Philosophical Quarterly* 88 (2007): 39-73; "Consequentializing," *Philosophy Compass* 4.2 (2009): 329-347.

²⁶ And this is not merely the instrumental reason that one may be in an especially good position to ensure the welfare of one's children. Rather, it seems plausible to me that the parental relation gives a parent non-instrumental reason to favor her child's flourishing over the flourishing of a distant child.

improve her wellbeing in the long run. The reason I have not to break Jane's leg is thus not something about the importance of Jane's wellbeing, but about her status as a person or her claim against me, or something of that sort.

In recent work, Portmore has moved away from talk of value and the desirability of outcomes.²⁷ So his project might appear immune from my arguments above. Portmore might then be able to accommodate agent-relative restrictions, by recognizing agent-relative reasons to prefer (agent-neutrally) suboptimal outcomes, and claim that this has nothing to do with such outcomes being better or more worthy of favor. This is a strength of Portmore's project. However, I see two difficulties with this. First, what sort of agent-relative reason to prefer is involved in the case of deontological restrictions? Given what has been said so far, it cannot be an egoistic reason, or a reason of the sort that arises in parental obligations, or the reason we have to prefer an outcome that is better or more favorable. So what is it? Second, even if such a move can formally accommodate restrictions, it would not answer worries about restrictions in the particular context in which they are thought to be puzzling. To see this, let me recapitulate some dialectic. It is thought that what is compelling about utilitarianism is its welfarism combined with its maximizing form.²⁸ The welfare of individuals is something that seems to be clearly valuable and something with motivational force – we care about the wellbeing of others. So welfare is good. Maximization now looks enticing – would not it be silly to bring about an outcome that is less good than some other outcome?²⁹ But then agent-relative restrictions look silly, since this is just what they require. One might subsequently reject welfarism, but whatever one's theory of the good, restrictions retain this “paradoxical” character, since they seem to require you to bring about a worse outcome. Defenders of agent-centered restrictions (and of commonsense morality generally) may respond that this air of paradoxicality is false, for there is no conflict between agent-centered restrictions and maximizing the good. Rather, the good is merely to be understood agent-relatively.³⁰

For this response to be successful, we need an elucidation of agent-relative value that shows how agent-relative value is intimately related to

²⁷ Portmore, “Consequentializing,” 332; see also his “Consequentializing Moral Theories,” 50.

²⁸ T.M. Scanlon, “Contractualism and Utilitarianism,” in *Utilitarianism and Beyond*, eds. A. Sen & B. Williams (Cambridge: Cambridge University Press, 1982), 108-111; “Sen and Consequentialism,” *Economics and Philosophy* 17 (2001): 40. For more on the appeal of utilitarianism and consequentialism generally, see Scheffler, “Agent-Centered Restrictions,” 414, and Portmore, “Consequentializing,” 332. However, I am unsure about Portmore's claim that what is compelling about consequentialism is the idea that “the reasons there are for and against performing a given act are wholly determined by the reasons there are for and against preferring its outcome to those of its available alternatives.” Rather, I think what is actually compelling about consequentialism (and egoism), despite more sophisticated versions of consequentialism, is really something about wellbeing, either general or individual wellbeing.

²⁹ Cf. P. Foot's “simple thought.” “Utilitarianism and the Virtues,” *Mind* 94.374 (1985): 198.

³⁰ Dreier, “Structures,” 24-5.

value simpliciter.³¹ After all, it is in a particular context that restrictions look paradoxical. We think we have a clear notion of the good, that it possesses clear motivational force, and so we would like to ground all other deontic concepts in the good.³² If this is one's aim, then a satisfactory dissolution of the paradoxical air of restrictions must show that they are grounded in considerations of the good as it is commonsensically understood. It is *because* we think we have some clear, commonsensical understanding of goodness that we want to reduce other deontic concepts to the good. As Schroeder has convincingly argued, *good* and *good-for* are concepts we have an intuitive grasp on (and are things it seems to make sense to want to maximize) – *good-relative-to* is not.³³ So, for the consequentialization of restrictions to be helpful here, the notion of agent-relative value it employs must be shown to be closely related to value simpliciter.

This is why I think that the project of grounding all reasons to promote in reasons to prefer will not answer the worry about restrictions, in the particular context in which it arose. The hope or interest, if there is any, in consequentialization, is to ground all reasons for action in something it is thought we have a clearer grip on and that is of obvious moral significance – the good. That is, what would be worthwhile would be if one could show how agent-relative restrictions can be ultimately grounded on considerations of value. But Portmore has moved away from talk of value, focusing on reasons to prefer. But it is unclear whether such agent-relative reasons to prefer relate clearly to, or reduce to, considerations of value. If they do, then much of my argument above – that agent-relative restrictions cannot be captured in terms of agent-relative value – would suggest that restrictions cannot be captured in terms of agent-relative reasons to prefer either. But if they do not, then the project of reducing all reasons to reasons to prefer is not deeply illuminating, since the hope of the reductive project is to reduce all reasons to something we have a better grip on. My first worry would then reemerge. Why do we have special reason to prefer the wrongdoing of others more than our own wrongdoing? It seems far more natural to me to say that we have *no* such reason – what we have special reason for is to not *act* in certain ways.

4. Conclusion

I have been arguing that, given a plausible assumption about the relationship between value and worthiness of (dis)favor, agent-relative restrictions cannot be captured by the notion of agent-relative value. The consequentializing project thus cannot succeed at capturing agent-relative restrictions. To the extent that such restrictions are thought paradoxical, the consequentializing

³¹ Schroeder, "Teleology."

³² Scanlon, "Contractualism," 108-111; "Sen and Consequentialism," 39-40.

³³ Schroeder, "Teleology," 291.

project cannot remove that stigma. So much the worse for consequentialization, some might say.³⁴ But perhaps others might be inclined to say, so much the worse for agent-relative restrictions! If agent-relative restrictions are resistant to consequentialization, does this call the coherency of such restrictions into doubt?

If one assumes that all reasons are grounded in value, then the failure of consequentialization surely would be problematic for restrictions. But this is a substantial assumption, and one not likely to be shared among many supporters of agent-relative restrictions.³⁵ Indeed, the notion that such restrictions are paradoxical might be thought to rest on just this presupposition.

We have here two strong, competing judgments. On the one hand, is it not obvious that I have special reason not to wrong someone, even if that would prevent similar wrongdoings by others? This is not to say that my wrongdoing is worse relative to me. Rather, it is just the thought that others can have special claims on me – claims that do not rest on considerations of value. On the other hand, is it not obvious that any reason for action must correspond to some value – namely the outcome the action would bring about? After all, what reason could there be for doing something, other than that it brings about something of value? These are both considered judgments that are held with high confidence by various parties. But without begging important questions, one cannot claim that the failure to consequentialize restrictions counts against the plausibility of restrictions.

In conclusion, I have argued that teleology and agent-relative restrictions are incompatible with one another. A consistent moral theory cannot both possess a teleological framework and incorporate agent-relative restrictions – such a theory will have to reject one of these elements.³⁶

Stephen F. Emet
Brown University
Department of Philosophy
Stephen_Emet@brown.edu

³⁴ E.g., McNaughton & Rawling, who argue against seeing agent-relative reasons as grounded on agent-relative value. “Value and Agent-Relative Reasons,” *Utilitas* 7 (1995): 31-47; “On Defending Deontology,” *Ratio* 11(1998): 37-54.

³⁵ For instance, Kamm, who rejects the attempt to ground restrictions in agent-relative value: “Non-Consequentialism,” 382.

³⁶ I would like to thank Rachel Singpurwalla and Alec Walen for helpful comments on earlier versions of this paper. I would especially like to thank the anonymous reviewers of this journal for many excellent comments that made this paper far better than it otherwise would have been.