

A SIMPLE SOLUTION TO THE SCOPE PROBLEM

JACOB BARRETT
Vanderbilt University

According to the desire-satisfaction theory of welfare, something is good for me to the extent that I desire it. This theory faces the “scope problem”: many of the things I desire, intuitively, lie beyond the scope of my welfare. Here, I argue that a simple solution to this problem is available. First, I suggest that it is a general feature of desires that they can differ not only in their objects but also in their “targets,” or for the sake of whom one has the desire. For example, I can desire that my child win an award either for their sake or for my own sake. Second, I show that we can use this idea to solve the scope problem by holding that something is good for me to the extent that I desire it for my own sake. Despite first appearances, this solution is not ad hoc, incomplete, or circular.

Keywords: desire satisfaction; preference satisfaction; welfare; well-being; the scope problem

1. The Scope Problem

We desire all sorts of different things. We desire our own pleasure, but we also desire the pleasure of others. We desire various personal achievements and luxuries—a rewarding career, meaningful relationships, to win prizes, good food and drink, to travel, to frolic in the sun—and, sometimes, we desire that others achieve these things too. We desire to be respected, to live in a just world, that our favorite sports team wins. We desire to do the right thing, that others do the right thing, and—at least if we have a retributivist streak—that those who don’t do the right thing suffer the consequences. The list goes on.

According to the desire-satisfaction theory of welfare, something is good for someone to the extent that they desire it. The sheer range of desires we have, however, puts considerable pressure on this theory. It seems plausible enough that when I satisfy desires for my own experiences or achievements this con-

Contact: Jacob Barrett <jacob.a.barrett@vanderbilt.edu>

tributes to my welfare. But is it also good for me that others experience pleasure and obtain various achievements, simply because I desire this? For example, if I strike up a conversation with a stranger on a train, and develop a desire that their illness is cured, am I thereby made better off when they are cured, even if I never hear of this or meet the person again (Parfit 1984: 494)? Many find this difficult to believe. Similarly, is it always on balance good for me to do the right thing, or to promote the benefits of others, if I on balance desire to do this? If, for example, I on balance desire to perform actions that benefit others at what intuitively seems like a great personal sacrifice (for example, seriously injuring myself or donating a huge sum of money to charity), does doing so really on balance promote my own welfare—such that this is not really a sacrifice after all (Overvold 1980: 108)? Or, to take another sort of case, is it good for me that morally better states of affairs that intuitively have nothing to do with me obtain, if I desire this—for example, that the war in Ukraine ends (to update an example from Sumner 1996: 134)? Again, this seems implausible.

These and similar examples generate the “scope problem” for the desire-satisfaction theory of welfare (Darwall 2002: 27). The problem is that many of our desires seem to have nothing to do with our welfare; the desire-satisfaction theory includes far too much in its scope. Or, as James Griffin eloquently puts it: “The trouble is that one’s desires spread themselves so widely over the world that their objects extend far outside the bound of what, with any plausibility, one could take as touching one’s own well-being” (1987: 17).

There are, broadly speaking, three ways to respond to the scope problem. The first is to reject the desire-satisfaction theory of welfare. Since I wish to defend this theory here, I set this response aside. The second is to bite the bullet and maintain that it is in fact good for me to satisfy the various desires that seem to generate the scope problem, our initial intuitions notwithstanding. This is widely regarded as too large a bullet to bite (though see Lucas 2010), and I will assume as much here. The third is to narrow the scope of the desires that count toward my welfare by claiming that only a proper subset of my desires matter in this way. The problem with this response, however, is that no particular account of how to narrow this scope has met with much success.

My goal in this paper is to argue that this third response is easier to pull off than has been appreciated. Specifically, I will suggest that we can rescue the desire-satisfaction theory of welfare simply by claiming that something is good for me when I desire that thing *for my own sake*. I cannot claim much originality for this solution. Stephen Darwall (2002), for example, makes a similar move when defending his alternative rational care theory of welfare, on which a person’s welfare consists in what we should desire for that person for their sake—rather than, as on the view I am suggesting, on what a person in fact desires for their own sake. And I suspect that many who have encountered the scope problem

have come to a similar idea: after all, it seems like precisely what has gone wrong in the above examples generating the scope problem is that the things I desire are not things I desire for my own sake. But, for whatever reason, this solution hasn't caught on or even been widely discussed (though see Dorsey 2012: 421), most likely because it seems ad hoc, incomplete, or circular—like the first step toward a solution, rather than a solution itself. Don't we owe some further story about what it is for me to desire something for my own sake, and isn't this precisely what an adequate solution to the scope problem would need to provide?

When the above solution first occurred to me, I had the same response, but I have since come to believe that the solution is perfectly adequate. Here, I develop it in three steps. First, I defend an intuitive but overlooked claim about the structure of desires. Second, I show that accepting this claim allows us to easily solve the scope problem. Third, I respond to some objections. Throughout, I make no attempt to provide a positive argument for the desire-satisfaction theory or to demonstrate its preferability to its rivals. Nor, for that matter, do I argue that my solution to the scope problem is superior to others in the literature. Instead, I aim only to show that my solution works, such that there is at least one viable solution to the scope problem available—and a very simple one at that. The scope problem therefore provides no grounds on which to reject the desire-satisfaction theory of welfare.

Before going on, let me clarify the scope of the scope problem. There are many other problems for the desire-satisfaction theory of welfare, and I cannot address them all here. But one is worth singling out since it is easily conflated with the scope problem. This is the problem of what we might call “worthless” desires: certain desires have objects that do not seem worth pursuing at all, and whose achievement correspondingly does not seem to contribute to one's welfare. Standard examples here include someone with the sole desire to count blades of grass (Rawls 1971: 434), to knock down as many icicles as possible (Kraut 1994: 40), or to pursue another similarly pointless project.

The problem of worthless desires is related to the scope problem, since both involve desires whose satisfaction does not, intuitively, benefit one. But the two problems are distinct, since the scope problem arises even when it comes to worthwhile desires, of the sort discussed above, whose objects are very much worth pursuing but whose satisfaction just doesn't seem to make one better off. And likewise, the problem of worthless desires can arise even for desires that do not intuitively fall outside the scope of one's welfare: for example, the desire that I have exactly 92 friends seems worthless, not worth pursuing, but not because its object lies *beyond* my welfare. As a result, standard solutions to the problem of worthless desires—for example, the idea that we shouldn't trust intuitions about what is good for people with psychologically atypical desires, or that we should focus not on all actual desires but on idealized, informed, or rational desires (e.g., Brandt 1979), or on desires we can justify with reasons (Bruckner

2016)—do not solve the scope problem. After all, even psychologically typical, idealized, informed, and rational desires, which we can justify with reasons, can still intuitively fall outside the scope of our welfare. Consider again the desire that someone else's life go well.

In what follows, I therefore focus exclusively on worthwhile desires that generate the scope problem, leaving to the side the question of which solution to the problem of worthless desires is best. The scope problem, then, is the problem that certain desires, *which otherwise seem perfectly sensible and worthwhile*, seem to have objects falling outside the scope of one's welfare. Put this way, we can also see that the scope problem excludes various other problems or puzzles for the desire-satisfaction theory, concerning whether we should put further constraints not on which desires count as welfare-relevant, but on what it takes to *satisfy* a desire in a welfare-relevant way.¹ For example, must I be *aware* that a desire is satisfied, and must I have a desire *at the same time* that it is satisfied, in order for its satisfaction to benefit me (e.g., Forrester 2023)? Although some standard cases used to illustrate the scope problem run these issues together, we should not be distracted by these features. For example, it still doesn't seem like I am made better off simply in virtue of satisfying my desire that a stranger's illness is cured, even if we specify that I am made aware of their recovery when it happens and that I have maintained a desire for their recovery since I met them once, many years ago, on a train.

For my purposes here, then, I remain neutral about whether the desire-satisfaction theory should embrace further constraints either on when desires count as worthwhile or on when satisfying a desire contributes to one's welfare—though I will generally speak as if no such constraints exist to avoid adding needless complexity. Indeed, as we will see, it is a virtue of my solution to the scope problem that it is not only simple but modular, permitting a wide range of answers to other questions faced by the desire-satisfaction theory of welfare.

2. The Desire Relation

I now begin the first step of my argument, which involves defending an intuitive but overlooked claim about the structure of desires. It is commonly assumed that desires involve a three place relation: an agent *A* desires object *x* with strength *s*, or $D(A, x, s)$. That agents have desires is uncontroversial, but it is somewhat more controversial what to say about desires' objects or strengths. The standard view is that desires take propositions as their objects. For example, when we say that Angela desires orange juice, this is shorthand for the claim that she desires

1. Thanks to an anonymous referee for suggesting this clarification.

that some relevant proposition obtain—for example, Angela desires *that she drink orange juice*. A competitor view holds that agents desire concrete objects directly, such that when we say that Angela desires orange juice, this isn't shorthand for anything: Angela's desire takes *orange juice* as its object (Thagard 2006). This debate is irrelevant to my argument, so I will simply assume that desires have propositional objects, since this is the more common and, I believe, the more plausible view. The desire relation $D(A, x, s)$ should therefore be interpreted as claiming that agent A desires that proposition x obtain with strength s .

There is also disagreement about desire strength. On one view, which seems more common among economists, desires lack primitive strengths. What agents really have, at the most fundamental level, are ordinal (or “strengthless”) preferences, of the form A prefers x to y , or xP_Ay —as preferences are conventionally represented. If an agent has preferences not only over propositions, but also over probability distributions (or “lotteries”) of propositions, and if these preferences meet certain consistency conditions, the agent can then be described as having desires with numerically representable strengths (von Neumann & Morgenstern 1944: ch. 3). So, on this view, when we say that the difference between the strength of Angela's desires for orange juice and apple juice is twice the difference between the strength of her desires for apple juice and water, this is shorthand for making some claim about Angela's ordinal preferences over probability distributions of orange juice, apple juice, and water. On another view, which seems more common among philosophers, desires simply come with psychologically basic strengths, such that when we make the above claim about Angela's desire strength, this isn't shorthand for anything. Since we make claims about desire strength all the time, we can think of this as a “primitivist” solution to the problem of desire strength: rather than explaining facts about desire strengths in terms of something more fundamental or basic, we simply hold that desire strength is psychologically primitive or basic itself.

It does not matter for our purposes here whether we think that desire strengths are primitive, but I will assume that they are, since this is the simpler and, as I have discussed elsewhere, the more plausible view (Barrett 2019; 2022). But the distinction between “shorthand” and “primitivist” explanations of intuitive features of desires will be useful. Specifically, in addition to making claims about which agents desire things, what they desire, and how strongly they desire them, we often also make claims about what I will call desires' *targets*, or “for the sake of” whom one desires certain things (Darwall 2002: 47). For example, parents often desire not merely that certain propositions obtain that relate to their children; they desire that these things obtain *for the sake of* their children. In fact, a parent might desire the exact same object either for their own sake or for their child's sake: Imagine a woman up for a prestigious award accusing her father of only desiring that she wins the award for his sake, rather than that she wins

the award for her sake. Or, to take another example, suppose you are working on a group project, whose success will be credited equally to both you and your partner. Here, it seems perfectly intelligible to distinguish the case where you want the project to succeed for your own sake from the case where you want it to succeed for your partner's sake. Indeed, we can imagine a spectrum of possibilities, from one in which you exclusively desire that the project succeeds for your own sake, to one in which you strongly desire that it succeeds for you own sake but weakly desire that it succeeds for your partner's sake, all the way until the case where you don't at all desire that the project succeeds for your own sake but desire this only for your partner's sake.

At least on its face, then, desires appear to involve a four-place relation. An agent *A* desires that a proposition *x* obtain with strength *s* *for the sake of* target *t*, or $D(A, x, s, t)$. And my suggestion is that we take this appearance at face value: we should take it as primitive, or psychologically basic, that desires have this fourth argument place, corresponding to the target of a desire, or for the sake of whom one has the desire. The reason for this is straightforward: I know of no other plausible way to explain the intuitive difference between a case where a father desires that his daughter succeed for his own sake and the case where he desires that she succeed for her sake, or to explain the above range of desires one might have about the success of a group project. Take the former pair of cases. We cannot explain the difference in the father's desires by reference to which agent *A* has the desire (in each case it is the father doing the desiring), what object *x* the desire takes (in each case, it is that the daughter win the award, so both are equally "about" the daughter), or the strength *s* of the desire (which, we may stipulate, is equally strong in each case). But there really is an intuitive difference between these cases, so, to explain this, we must posit that desires involve four-place relations with argument places not only for agents, objects, and strengths, but also for targets. Of course, some desires might be untargeted, in the sense that we don't desire them for anyone in particular—consider a desire that some abstract value like beauty obtains in the world. But that is consistent with my claim here. My proposal is that some desires have targets, not necessarily that all of them do. Or, if one prefers, one can think of all desires as having targets, and intuitively "untargeted" desires as those that take some perfectly generic object as their target—for example, they may be desires one has for the sake of "the universe" or "the world."

One objection to this proposal draws on the above analogy with desire strength. Even though, intuitively, it might seem that we can't explain desire strength except by claiming that it is primitive, it turns out that we can indeed provide an alternative and non-obvious account of desire strength in terms of preferences over probability distributions or "lotteries." So, perhaps it is similarly possible to provide some alternative and non-obvious account of desires'

targets, or for the sake of whom one desires things, in terms of something else: perhaps we should adopt a “shorthand” rather than a “primitivist” account of desires’ targets. A deflationary possibility is that, in the cases in question, what is really going on is that the agent has different intrinsic desires in the background, whereas the differentially targeted desires in question are only instrumental. The difference between the father desiring that the daughter wins the award for his sake and the father desiring that the daughter wins for her sake, on this account, is that in the former the father intrinsically desires, say, the warm glow of vicarious parental success, whereas in the latter he intrinsically desires that the child feels the warm glow of personal success. The desire that the child win the award is only instrumental to, or a means to, the satisfaction of one of these intrinsic desires, and it is this difference in intrinsic desire that explains the intuitive difference between what I have been calling the “target” of the father’s two desires. Or perhaps in the former case, the father intrinsically desires to feel a warm glow and only instrumentally desires that his daughter wins the award, whereas in the latter he intrinsically desires that his daughter wins.

This account may seem plausible in the case at hand, but it fails to generalize, since it cannot explain the difference between desires’ targets in cases where such desires are both intrinsic. Suppose a friend’s child is in danger of dying or being seriously harmed, and you intrinsically desire that this doesn’t occur. It seems perfectly intelligible that this desire might have different targets: you might desire that the child is okay for their sake, or you might desire that the child is okay for their parent’s sake, or perhaps you might desire that the child is okay in an untargeted or impersonal way—for no one’s sake in particular (or for the sake of the world) (Darwall 2002: 64). Similarly, if it is your own child, you might intrinsically desire that they are okay for your own sake, or intrinsically desire this for their sake. Unless there is something incoherent in claiming that such desires can indeed be intrinsic, which I see no reason to think, this is enough to defeat the account in question.

To be clear, we do sometimes use the “for the sake of” locution to refer to instrumental relations, but that is not what is going on here. In the instrumental sense, we desire one proposition “for the sake of” some other proposition—for example, we desire that someone wins an award for the sake of *or as a means to* the realization of the proposition that they feel a warm glow. Similarly, we often refer to intrinsic desires as desires that we have for objects *for their own sake*, by which we mean that we desire certain objects non-instrumentally or finally: we desire that they obtain *full stop*, not as a means to anything else. But in the sense that is relevant when talking about desires’ targets, we desire a proposition not “for the sake of” some further proposition, but rather for the sake of some concrete object, paradigmatically, some person. So while it might sound odd to say, about some targeted intrinsic desire, that I desire an object for its own sake for

the sake of some person, this is only because the “for the sake of” locution is being used in two different senses. Desires have propositional objects, but they also have non-propositional targets (Darwall 2002: 69). And both, I am suggesting, are equally primitive.

It would therefore be a mistake to analyze desires’ targets in terms of instrumental relations: intrinsic desires, too, can differ in their targets. But perhaps, one might still object, some other non-primitivist or shorthand account of (intrinsic) desires’ targets is available. Indeed, one might think that various proposed solutions to the scope problem might similarly provide accounts of desires’ targets. For example, one proposal is that satisfying a desire only contributes to my welfare if the desire is for something that, as a matter of logical necessity, can only occur if I exist at the time it is satisfied (Overvold 1980: 118). Analogously, we might say that my desire takes some person as a target when I desire something that, as a matter of logical necessity, can only occur if they exist at the time it is satisfied. This account fails to solve the scope problem since it is extensionally inadequate: For example, it labels the satisfaction of one’s on balance desire to do one’s duty as contributing to one’s welfare, even in cases where this comes at great personal sacrifice (Velleman 2002: 30). And the analogous account similarly fails to explain how to make sense of a desire’s target: It necessarily labels a desire that someone else do their duty as a desire for their sake, when this needn’t always be the case. Nevertheless, if some better version of such an account were available (perhaps involving some careful tweaks to the above proposal), one can see how it might explain how desires can differ in their targets without positing targets as primitive.

My own view is that no such account is available, and that we should therefore hold that it is indeed a psychologically basic, primitive fact about desires that they can differ in their targets. But note that we can solve the scope problem either way. For if I am wrong about this, and there is some non-primitive account of desires targets available, well, then we can use this account to solve the scope problem, in the way I am about to explain. And if I am right about this, and it is a primitive psychological fact that desires can differ in their targets, then, as I will now argue, we can use this account to solve the scope problem too. Going forward, I therefore proceed largely on the assumption that desires’ targets are indeed primitive, but it is important to remember that my argument doesn’t strictly speaking rely on this. Rather, the solution I will develop in the next section only requires that (intrinsic) desires can differ in their targets, not necessarily that such targets are primitive—at least as long as one’s non-primitive analysis of desires’ targets does not generate a certain sort of circularity or explanatory deficiency that I will address in due course.

I recognize that I am proposing an unusual account of desire, but I believe this is simply because this feature of desires has been overlooked. Once we think

about it, it is perfectly intuitive that desires have different targets—that I can desire something for my own sake, or for someone else’s sake, or perhaps, say, for my country’s sake—and it doesn’t seem at all unintuitive to think that this is a primitive psychological fact about desires. So I submit that desires involve a four-place relation, between an agent *A*, an object *x*, a strength *s*, and a target *t*. Crucially, while *A* and *t* often refer to the same agent—we often desire things for our own sake—this isn’t always the case.

3. Solving the Scope Problem

With this account of desire in the background, I now offer a simple solution to the scope problem. According to the standard desire-satisfaction theory of welfare, the realization of some proposition *x* is good for an agent *A* to the extent that *A* (intrinsically) desires that *x*. According to my proposed modification, the realization of some proposition *x* is good for an agent *A* to the extent that *A* (intrinsically) desires that *x* for *A*’s own sake. This modification, it seems to me, can handle all the standard cases where the scope problem arises, since all such cases involve an agent desiring an object that, intuitively, they don’t desire for their own sake. For example, if I develop a desire that another agent’s life goes well (say, a stranger on a train), then I presumably desire that their life goes well for their sake, and this explains why it does not make me better off if their life goes well. At the same time, however, it does seem intuitive that parents or romantic partners, say, can indeed be made somewhat better or worse off by things going well for their children or partners. And my proposed account can explain this too: one remarkable feature of being in a close relationship is that one comes to desire that the other’s life go well not only for their sake, but also for one’s own sake.

This modification can handle all other cases I am aware of as well, for example, cases of self-sacrifice: even if one most strongly desires to do the right thing, one typically desires this for others’ sake, or in an untargeted or impersonal way. So one can do what one on balance most desires and yet set back one’s own welfare, if this sets back the desires one has for one’s own sake. And, again, if I desire some morally good event occurs that intuitively doesn’t affect me, such as the war in Ukraine ending, this doesn’t benefit me since I don’t desire this for my own sake. Now, to be clear, the account in question will label the satisfaction of some very similar desires as contributing to my welfare. For example, if I desire that I be “morally upright” for my own sake, then becoming morally upright will come out as benefiting me (Dorsey 2012: 421). But this seems perfectly intuitive. Insofar as I desire that I be morally upright (or, say, that I do my duty) for my own sake, satisfying this desire plausibly contributes to my welfare, much

like achieving any other personal project does. But insofar as I desire that I be morally upright (or that I do my duty) for others' sake, or for no one's sake in particular, satisfying this desire doesn't plausibly contribute to my welfare. And that is exactly what the account I am proposing predicts.

My solution to the scope problem therefore appears to handle all the standard problem cases for the desire-satisfaction theory of welfare. But can we modify these cases to generate new problems? Suppose that, as before, you meet a stranger on the train—only this time, you form a desire that their illness is cured *for your sake*. My simple solution suggests that if the stranger recovers, you are thereby made better off. Isn't this as counterintuitive as the original case?²

My response to this concern is twofold. First, I do *not* think this is similarly counterintuitive, once we fill in the details in any way that makes the desire psychologically realistic. For, as I have noted, in the context of a close relationship, we do think that you can be made better or worse off by what happens to another, and it is an advantage of my account that it neatly explains this in terms of the idea that you can desire that things go well for those you stand in close relations to *for your own sake*. Just consider a parent saying to a child, "I know you don't care about your safety, but I'm not being paternalistic here. I want you to be okay *selfishly*; please wear a helmet *for me*." Of course, it does seem odd to form desires like this about a stranger one has just met, but to make this plausible we can suppose, for example, that the stranger had a significant impact on you, that you felt a real connection, that you felt a kindred spirit in them, and so you formed desires about them, for their sake, of the sort you would normally only form with a close relation. In this case it no longer seems unintuitive to think that when their illness is cured, you are made better off, just as you might be if your child or spouse's illness is cured. The oddity of the case is just that we do not normally form desires in this way—as if by a sort of (platonic) love at first sight.

Of course, one might insist that psychological realism is not what matters: we can conceive of cases where you desire that a stranger's illness is cured for your own sake, despite forming no connection with them, and without any of the other details making the case psychologically plausible. Since satisfying such a desire doesn't intuitively make you better off, isn't this enough to defeat my account? Now, I must admit that my own intuitions about this case are rather indecisive. But if one does have the strong intuition that satisfying this desire fails to make you better off, this brings me to the second part of my reply: if we push the issue in this direction, then we find ourselves back at the problem of worthless desires from earlier. Absent some further story making it psychologically plausible, desiring a stranger do well *for your sake* seems pointless and

2. Thanks to an anonymous referee for raising this objection.

bizarre, much like desiring that you spend all day counting blades of grass or that you have exactly 92 friends. And in this case, all the standard replies to the problem of worthless desires become available. Perhaps we lack reliable intuitions about the welfare of people with atypical desires, or perhaps we should restrict our focus to rational, informed, or idealized desires, or to desires one can justify with reasons. As I have said, it is not my goal to defend a particular solution to the problem of worthless desires here. Rather, my point is that *if* we tell a story making your desire that the stranger's illness is cured for your sake seem worthwhile, *then* we lose the intuition that satisfying it fails to contribute to your welfare. The problem my solution to the scope problem leaves unresolved, to the extent that it arises, reduces to the problem of worthless desires.

A final complication involves posthumous desires.³ Suppose I desire that, after I die, my friends visit my grave. Am I made better off if they do so? Certain alternative solutions to the scope problem rule this out—for example, the above (failed) solution, on which satisfying a desire only contributes to my welfare when that desire is for something that can only occur if I exist at the time it is satisfied. But while my account does not similarly rule out the possibility of the satisfaction of posthumous desires benefiting one, it also does not rule this in—which is all the better, as intuitions vary widely about such cases. After all, the core issue with posthumous desires is that such desires can only be satisfied at a time when agents no longer have them, since such agents no longer exist. This, as we have seen, is a separate issue from the scope problem, as it concerns what further conditions must be in place for the *satisfaction* of a desire to qualify as benefiting one, and not which desires qualify as *in scope*. Some think that satisfying a desire only benefits one if one has the desire at the same time it is satisfied; others think one can be retroactively made better off if a desire one has at one time is satisfied at a later time, even if one no longer has that desire then. My own inclination is to go for the first option and to conclude that satisfying posthumous desires does not benefit one. But I am glad that my solution to the scope problem is compatible with either approach, as this makes my solution not only simple but also modular, and so more widely acceptable.

4. Is the Solution Too Easy?

So far, my solution to the scope problem seems to check all the boxes: it handles all the standard problem cases generating the scope problem, handles modifications to these cases insofar as they do not reduce to the problem of worthless desires, and is compatible with different approaches to other puzzles, for

3. Thanks to an anonymous referee for raising this issue.

example, about the timing of desires. Still, I recognize that one might have some residual feeling that there is something “too easy” about it. How might we make sense of this feeling?

One possible worry is that my solution is ad hoc, since the idea of desire’s targets or a “for the sake of” relation might itself seem rather ad hoc. But note that the argument I gave for the existence of this relation, for desires involving a four-place rather than a three-place relation, made no reference to the scope problem. We had to appeal to this relation to draw the distinctions we intuitively want to draw between different cases of desires, for example, the different desires a parent might have in relation to their child. So there is nothing ad hoc about the solution after all.

Another worry we might have is that the solution is incomplete. It is all very well to say that something is good for me if I desire that thing for my own sake, but don’t we need a further explanation of what it is to desire something for my own sake? I have suggested, however, that no such explanation is available or needed, since it is plausibly a primitive fact about desires that we can not only desire different propositional objects with different strengths, but that we can also desire such objects for the sake of different non-propositional targets. If that’s right, then there is nothing incomplete about this solution after all. And if it’s not right, and one insists on some alternative non-primitive or “shorthand” account of desire’s targets, then we can use that account to complete our solution, and the charge of incompleteness will dissipate once more.

Finally, and perhaps most significantly, one might think that there is something circular about my proposal. We are after an account of when something is good for someone, and we have claimed that something is good for some agent *A* when *A* desires that thing for *A*’s own sake. But I have admitted that I cannot rule out the possibility that desiring something for someone’s sake does not involve a primitive relation, and this opens the window to a new problem: maybe to desire something for someone’s sake is to desire it *insofar as it is good for them*, rendering the account problematically circular. Connie Rosati raises a similar objection against Darwall’s (2002) rational care theory of welfare, and perhaps the objection works against Darwall given his further commitments (Rosati 2006: 620–626; see also Arneson 1999: 124). But in the present context, this objection fails too—though it takes a bit more work to see why, as the objection comes in two main varieties.

On the first version of the objection, the worry is that we must analyze the distinction between desires that I do and do not have for my own sake in terms of the idea that only satisfying the former sort of desire is good for me. This would render my account obviously circular, as I would now be trying to explain the distinction between desires whose satisfaction is and is not good for me by appeal to the very same distinction: I would be explaining what it is for

something to be good for me in terms of me desiring it for my own sake, while also explaining what it is to desire something for someone's sake in terms of it being good for them. However, my account just as obviously does not face this objection, as it claims that whether we desire something for someone's sake is a psychological fact about our desires, not a normative or evaluative fact about whether what we desire is in fact good for someone. Indeed, my account implies the possibility of desiring things for others' sakes that are *not* good for them—for example, I might desire something for your sake, even though you do not desire it for your own sake, such that, on my account, it is not good for you. So desiring something for someone's sake can come apart from it being good for them, we cannot analyze the (psychological) “for the sake of” relation in terms of the (normative or evaluative) “good for” relation, and this first version of the objection is easily dismissed.

The second version of the objection psychologizes it. Here, the idea is that we must analyze me desiring something for your sake in terms of me desiring it while thinking it is good for you, or perhaps more carefully, in terms of me desiring it *under the description* of it being good for you. Now, an initial problem for this version of the objection is that it does not obviously generate any circularity. For suppose we grant the analysis in question, so that my account says: something is good for you if you desire it for your own sake, you desire something for your own sake if you desire it under the description of it being good for you, so something is good for you if you desire it under the description of it being good for you. Although this account of what is good for someone makes reference to the concept “good for,” there need not be anything circular about this since the reference is embedded in people's attitudes. That is, we are providing an account of the normative or evaluative *property* of welfare, of what *is* good for people, not in terms of the same property, but rather in terms of psychological facts about how the *concept* of welfare or “good for” arises in people's desires. By analogy, consider a simple version of moral relativism on which something is wrong if most people believe it is wrong. For all its flaws, this account does not seem guilty of a problematic form of circularity either. It is not problematically circular to explain a *property* like wrongness or welfare in terms of people's cognitive or conative attitudes involving the *concept* “wrongness” or “welfare.” So there seems to be no genuine charge of circularity on this second interpretation of the objection.

However, I recognize that there might nevertheless seem something odd, something explanatorily deficient—even if not strictly speaking circular—about combining the account of desires' targets in question with the solution to the scope problem I have proposed. So let me raise a further, and more decisive reason to dismiss this objection. This is simply that it is not plausible to analyze desires' targets in terms of the idea that when you desire something for someone's sake, you desire it under the description of it being good for them. This is not a plau-

sible alternative, that is, to my own primitivist account. Now, I am happy to grant that—at least setting aside cases of worthless desires—when we desire something for someone’s sake, we do generally think that it is good for them in the very broad sense that it would provide them with something of *positive valence*. However, it is important not to conflate this broader sense of “good for” with the narrower sense of “good for” at issue, where something is good for someone insofar as it promotes their *welfare*. And it is not plausible that desiring something for someone’s sake necessarily involves desiring it while thinking that it is good for them in this narrower sense, as the objection in question presupposes.

To see why, return to our earlier example of desiring for someone’s sake that they become morally upright, or perhaps that they become a good or virtuous person. It seems clearly possible to desire this for someone, for their sake, without thinking that this would make them better off in the sense that it promotes their welfare. For instance, a parent can sincerely and intelligibly say to their adult child, “Look, I don’t deny that you are very well off with a very high welfare level. But at the same time, you are living a life of debauchery and villainy, and I don’t want that life for you. I want you to be a good person, to be morally upright, to be virtuous—I want that *for your sake*, even if it makes you worse off.” Or suppose, to take a less usual but perhaps even clearer example, that the parent is a devout adherent of the view that even satisfying worthless desires makes someone better off, but their child solely desires to count blades of grass and spends all day doing so. Here, the parent might desire, for their child’s sake, that the child develops other desires and lives differently—not because they think this will make the child better off, but because they think it will make the child’s life more *meaningful*, and they want a more meaningful life for their child, for their child’s sake. Or, as a final example, I can think of lots of cases where I desire, for a friend’s or loved one’s sake, that they succeed in achieving something that they desire, even though they do not desire that thing for their own sake. For example, I might desire, for my friend’s sake, that they succeed in their campaign to pass some local ordinance that they care a lot about passing for the sake of other less advantaged members of their community, even if I personally think the ordinance is pointless and won’t help anyone. Since I do in fact accept the view I am defending in this paper, this involves a clear case where I desire something for someone’s sake despite not thinking that it contributes to their welfare, because I know they don’t desire it for their own sake.

These examples demonstrate that desiring something for someone’s sake need not involve desiring it under the description of it being good for them, in the sense that achieving it would contribute to their welfare. So, much like the first version of the incompleteness objection, we can dismiss the second. Not only does the account of desires’ targets it presupposes fail to generate any strict form of circularity, but the account is not plausible on its face.

I conclude that my solution to the scope problem, though simple, is not “too easy”: it solves the problem without being problematically ad hoc, incomplete, or circular (or otherwise explanatorily deficient). It does, of course, require us to accept a revisionary account of desire, but this account is independently plausible. Indeed, I think we should accept that desires involve a four-place rather than a three-place relation, with an extra place for their targets, regardless of what theory of welfare we favor. It is a nice bonus that accepting as much allows us to solve the scope problem for the desire-satisfaction theory of welfare.

Acknowledgements

For helpful comments and discussion, I’d like to thank Scott Aikin, Luca Hemmerich, Christopher Howard, Sarah Raskoff, and two anonymous referees.

References

- Arneson, Richard J. (1999). Human Flourishing Versus Desire Satisfaction. *Social Philosophy & Policy*, 16(1), 113–142.
- Barrett, Jacob (2019). Interpersonal Comparisons with Preferences and Desires. *Politics, Philosophy & Economics*, 18(3), 219–241.
- Barrett, Jacob (2022). Subjectivism and Degrees of Well-Being. *Utilitas*, 34(1), 97–104.
- Bruckner, Donald (2016). Quirky Desires and Well-Being. *Journal of Ethics and Social Philosophy*, 10(2), 1–34.
- Darwall, Stephen (2002). *Welfare and Rational Care*. Princeton University Press.
- Dorsey, Dale (2012). Subjectivism without Desire. *The Philosophical Review*, 121(3), 417–442.
- Forrester, Paul (2023). Concurrent Awareness Desire Satisfaction. *Utilitas*, 35(3), 198–217.
- Griffin, James (1987). *Well-Being: Its Meaning, Measurement and Moral Importance*. Clarendon Press.
- Kraut, Richard (1994). Desire and the Human Good. *Proceedings and Addresses of the American Philosophical Association*, 68(2), 39–54.
- Lukas, Mark (2010). Desire Satisfactionism and the Problem of Irrelevant Desires. *Journal of Ethics and Social Philosophy*, 4(2), 1–24.
- Overvold, Mark C. (1980). Self-Interest and the Concept of Self-Sacrifice. *Canadian Journal of Philosophy*, 10(1), 105–118.
- Parfit, Derek (1984). *Reasons and Persons*. Oxford University Press.
- Rawls, John (1971). *A Theory of Justice*. Belknap Press.
- Rosati, Connie (2006). Darwall on Welfare and Rational Care. *Philosophical Studies*, 130(3), 619–635.
- Sumner, L. W. (1996). *Welfare, Happiness, and Ethics*. Oxford University Press.
- Thagard, Paul (2006). Desires Are Not Propositional Attitudes. *Dialogue*, 45(1), 151–156.
- von Neumann, John and Oskar Morgenstern (1944). *Theory of Games and Economic Behavior*. Princeton University Press.