

THE MORAL IMPORTANCE OF LOW-WELFARE SPECIES

JAKOB LOHMAR
University of Oxford

Many species seem to have much smaller welfare ranges than we do. Which importance should we assign to the welfare of these *low-welfare species* when we have to decide whether to benefit members of these species or other humans? In particular, should we ever prioritize low-welfare species when some of the most significant human goods are at stake? It could be argued that benefits for low-welfare species are *irrelevant* to significant human benefits because they are much smaller in comparison. I begin by developing this argument but subsequently reject it based on the consideration of a hypothetical *high-welfare species* whose welfare range vastly exceeds the human welfare range. Arguments from empathy and fairness are given to support the opposite view on which some benefits for low-welfare species can outweigh even significant human benefits. Finally, I propose a principle that reconciles this view with the intuition that human headaches cannot outweigh human lives.

Keywords: low-welfare species; welfare ranges; relevance; aggregation; prioritization; animal ethics

Introduction

Many species seem to have welfare ranges that are much smaller than the welfare range of our species: the difference between the maximum and minimum level of lifetime well-being that members of these species can achieve is much smaller than in our case. Given that they have any welfare at all, this seems particularly plausible for invertebrates, but also many (small) mammals, fish, reptiles and birds plausibly have much smaller welfare ranges

Contact: Jakob Lohmar <jakob.lohmar@st-annes.ox.ac.uk>
 <https://orcid.org/0000-0003-1182-5780>

than we do.¹ There are various reasons this might be the case: members of these species have much shorter lives, they have less developed brains that might only allow for less intense pleasures and pains, and they might not be able to attain certain sophisticated goods (such as scientific understanding) that require a high level of intelligence or rationality. Generally speaking, the physical and psychological constitution of many species seems to limit their capacity for well-being in a way that does not allow them to be nearly as well off or badly off as we can be, at least under any realistic circumstances.² I will refer to these species whose welfare ranges are much smaller than (or ‘tiny’ compared to) the human welfare range as *low-welfare species*.³

In this essay, I will discuss the moral importance of low-welfare species. More specifically, I will ask whether we should always prioritize some significant human benefits over benefits for members of low-welfare species or whether some benefits for members of low-welfare species can, if numerous enough, outweigh even the most significant human benefits. For concreteness, I will often use the saving of a human life as an example for one of the most significant human benefits, where the saved life is assumed to be worth living and of substantial length. More specifically, we can assume that the life would be about as good and long as a human life can be and that the person living it has only had a short life so far. My question can then be stated as follows: Is there any number N and benefit B such that one ought to confer benefit B to N members of a low-welfare species rather than save one human life?

1. See Rethink Priorities’ pioneering research on the welfare ranges of other species, in particular Fischer (2023, 2024). Note that on their definition, welfare ranges measure individuals’ *momentary* welfare rather than their lifetime welfare.

2. The addition of ‘under any realistic circumstances’ is supposed to allow for the possibility that members of low-welfare species could be much better off in some circumstances that are fundamentally different from ours. Perhaps, for example, mice could live for many hundreds of years in extreme pleasure if there was a life-prolonging experience machine that is optimized for mice (cf. Nozick 1974: 42–45). While this might be a metaphysical (and perhaps nomological) possibility, it is not a realistic one. Hence, this possibility does not affect the welfare range of mice. That is not to say that it is always easy to specify whether a possible circumstance should be considered ‘realistic’ or ‘fundamentally different’. For example, should the (mere) possibility of radical enhancement of one’s welfare-relevant capacities count as a realistic circumstance? My tentative answer to this question is that such circumstances should be considered ‘fundamentally different’ given that they are not actually an available technology. See Vallentyne (2005) for an insightful discussion.

3. To be a bit more precise, I am thinking of welfare ranges that are smaller than the human welfare range *by at least several orders of magnitude*. A species with a welfare range that is about half as large as the human welfare range would not qualify as a low-welfare species; a species with a welfare range that amounts to less than 0.1 percent of the human welfare range would. I am not assuming that such comparisons could be made precisely, even in principle. But I do assume that welfare ranges are at least imprecisely comparable, so that we can say, for some N , that one welfare range is *at least* N times as large as another welfare range even if we cannot say, for *any* N , that it is *exactly* N times as large.

While it seems obvious that we should save a human life rather than benefit a single member of a low-welfare species (whatever the latter benefit may be), it is unobvious whether the same is true for *any* number of any benefits for members of low-welfare species. At the same time, this question is of high practical importance. As noted above, many actual species seem to be low-welfare species and, moreover, many of these species—in particular invertebrate species—contain incredibly many individuals. Of course, we cannot choose to benefit an arbitrarily large number of these “small animals” that are plausible candidates for members of low welfare species. But it seems that we can benefit *far* more of them than we can benefit humans. That goes in particular for the enormous numbers of small animals in factory farms (such as fish, shrimp, and farmed insects) whose living conditions we could greatly improve. But also the suffering of an even much larger number of small animals in the wild provides an opportunity for us to potentially do a lot of good. Even if the tractability of benefitting small wild animals is quite low, their astronomical numbers could render interventions that are aimed at their well-being the most effective (in terms of the increase in total well-being per resource).⁴ If benefits for low-welfare species could in principle outweigh all human benefits, it may well be, therefore, that they sometimes actually do outweigh some of the most significant human benefits.

This conclusion would be hard to accept. It would mean that in decision situations in which we can save either humans or members of low-welfare species but (due to scarce resources) not both of them, we sometimes ought to save the members of low-welfare species while allowing the humans to be severely harmed. Most of us are likely to face decisions of this kind. For example, when deciding where to donate some of our money, we could use that money to protect people in extreme poverty from deadly tropical diseases, or we could use the money to fight the inhumane treatment of factory-farmed salmon. Of course, we could split our donations between the two causes or look for a third option that allows us to benefit both humans and small animals. But usually, for every dollar we donate, there will be some trade-off between human and small animal welfare; there is no perfect correlation between what is best for humans and what is best for small animals. What goes for money also goes for other resources. For example, we could use some of our spare time to protest against the discrimination of certain human minorities, or we could use that time to advocate for improving the living conditions of bees. If benefits for low-welfare species can in principle outweigh all benefits for humans, it might well be (although this is far

4. See Horta and Teran (2023) for different promising ways of benefitting wild animals effectively. See Sebo (2023) for the argument that the sheer number of small animals makes it plausible that we sometimes ought to prioritize them on a utilitarian calculus.

from certain!) that we should advocate for the bees rather than for the discriminated people and that we should donate to help salmon rather than people in extreme poverty.⁵

As I said, these conclusions would be hard to accept; they may also seem unintuitive. But we should be wary of rejecting them without substantive argument. Not only could our intuitive rejection of these practical conclusions merely be the result of motivated reasoning that is based on our self-interest or speciesism; their unintuitiveness could also be explained by the underlying unintuitive empirical premises about the effectiveness of benefitting small animals, such as the (factually correct) premise that there are so incredibly many small animals. Even if we accept these empirical premises, they are hard to grasp intuitively and may therefore not be properly reflected in our judgment. We should therefore inquire whether there are good arguments for the view that some significant benefits for humans always outweigh benefits for low-welfare species.

I will present such an argument in the following section. Based on the consideration of species with welfare ranges that are much *larger* than our welfare range, I will, however, reject this argument (section 3). Subsequently, I will give two arguments for the opposite view on which some benefits for low-welfare species can outweigh even the most significant human benefits. The first of these arguments is based on the value of fairness (section 4), and the second is an argument from empathy (section 5). Finally, I discuss principles that could explain why benefits for low-welfare species, but not small benefits for humans, can outweigh human lives (section 6). I conclude in section 7.

The Argument from Irrelevance of Tiny Benefits

Here is a straightforward argument for the view that some benefits for humans cannot be outweighed by any benefits for low-welfare species: The welfare range of an individual's species restricts how much individuals of that species can be benefited—increasing their well-being from the minimum to the maximum well-being level of that species is the most that could be done for them. The welfare range of a low-welfare species is, by definition, tiny compared to the human welfare range. Hence, the largest benefits to individuals of a low-welfare

5. One important complication (besides many others) for reaching these conclusions is the nonidentity problem (Parfit 1984: ch. 17): When trying to benefit small animals on a large scale by improving their living conditions, we will probably also change their identities. We would therefore replace less well-off animals with better-off animals rather than make particular animals better off. This could be held to make a moral difference. But that moral difference should not be overstated: We also change the identities of humans by many interventions such as those that mitigate climate change.

species are tiny compared to the largest human benefits. Benefits to one individual cannot be outweighed by any number of comparatively tiny benefits to other individuals. Thus, no number of even the largest benefits to members of low-welfare species can outweigh any of the largest human benefits, such as the saving of a human life.

The crucial premise in this argument is the claim that benefits, no matter their number, cannot outweigh any benefit for another individual compared to which they are tiny. This claim, which I will refer to as *Irrelevance of Tiny Benefits*, gains plausibility from cases like *life for headaches* (see Norcross 1997): In a choice between saving a human life (which, as before, is assumed to be well worth living and of substantial length) and alleviating a huge number of mild headaches, it seems that we ought to save the human life, *no matter* the number of headaches.⁶ One way of expressing this intuition is to say that the alleviation of mild headaches is *irrelevant* when human lives are at stake (Scanlon 1998: 239). Irrelevance of Tiny Benefits provides a natural explanation for this intuition: The benefit of an alleviated headache is irrelevant to the benefit of a saved human life because the former benefit is tiny compared to the latter benefit. More generally, the comparative magnitude of benefits seems to explain our intuitions about whether one benefit is relevant to another benefit in many common cases. Intuitively, for example, cases of constant paralysis *are* relevant to the saving of a human life: for some number N , we should prevent N cases of constant paralysis rather than save a single human life (Scanlon 1998: 239). The reason for this seems to be that the benefit of preventing constant paralysis, unlike the benefit of preventing a headache, is sufficiently large compared to the benefit of saving a life. Irrelevance of Tiny Benefits can be considered an instance of the more general view that the comparative magnitude of benefits determines their relevance.⁷

A straightforward way of specifying this more general view is *Compare Absolute Magnitudes*: For all benefits B_1 and B_2 , B_1 is relevant to B_2 if and only if the magnitude of B_1 is sufficiently large compared to the magnitude of B_2 . Given that comparatively *tiny* benefits are not ‘sufficiently large’, Compare Absolute

6. Norcross himself, who introduced this case, ends up rejecting the intuition that we should prioritize the life over any number of headaches. Many others (e.g., Dorsey 2009) are convinced that our intuition about this case is correct.

7. Since this view implies that the numbers of benefits sometimes matter and sometimes don’t matter, it is *partially aggregative*. Alternatively, Irrelevance of Tiny Benefits could be motivated by the *nonaggregative* view that the numbers of benefits *generally* don’t matter (e.g., Taurek 1977). Given this view, it would only need to be argued that no *single* instance of a benefit can outweigh any benefit compared to which it is tiny; it would follow that the same holds for *any* number of comparatively tiny benefits. The nonaggregative view, however, has highly implausible implications (see Parfit 1978). For instance, it would follow that if we should save a life rather than prevent a case of constant paralysis, we should also save a life rather than prevent one *billion* cases of constant paralysis. But intuitively, the numbers do make a difference here.

Magnitudes implies Irrelevance of Tiny Benefits. The same is true for many variations of Compare Absolute Magnitudes that also reflect other factors besides comparative benefit magnitudes. For example, Compare Absolute Magnitudes could be supplemented by a further necessary condition for a benefit's relevance according to which the benefit may not be based on immoral preferences (such as in the case of sadistic pleasures). Since comparatively tiny benefits do not meet the first necessary condition of having a sufficiently large comparative magnitude, they would still count as irrelevant on this variation of Compare Absolute Magnitudes. Another popular variation gives extra weight to benefits for individuals who are less well off than the other potential beneficiaries (e.g., Voorhoeve 2014). Benefits that are too small to count as relevant on the simple Compare Absolute Magnitudes can be relevant on this variation of the view if the beneficiaries have comparatively low welfare. However, that does not have to hold for *tiny* benefits. Arguably, tiny benefits should be considered irrelevant even if they are conferred to individuals with much lower welfare. Otherwise, mild headaches could outweigh human lives after all if the headache patients are much less well off than the person whose life is at stake.

Overall, then, Irrelevance of Tiny Benefits is supported by case intuitions and is implied by several general views about the relevance of benefits that seem plausible. Nevertheless, I believe that Irrelevance of Tiny Benefits should be rejected. One reason to doubt this premise is, of course, the possibility that there are no irrelevant benefits at all: perhaps *all* benefits can, if numerous enough, outweigh any other benefit. This is a possibility that should be taken seriously.⁸ My rejection of Irrelevance of Tiny Benefits, however, does not rely on this contentious view. In what follows, I will argue that Irrelevance of Tiny Benefits should be rejected even on the assumption that there are irrelevant benefits.

The Problem of High-Welfare Species

Consider a hypothetical "high-welfare species" with a welfare range that is so large that the human welfare range is tiny in comparison. Perhaps individuals of that species live much longer or are able to experience much more intensely than humans, for example. Imagine that we could save the life of one individual of that species that is about as good as lives of that species can be. If we save its life, the individual's lifetime well-being will be close to the maximum well-being level of that species while its lifetime well-being would be close to neutral if we don't save it. This seems to be a morally weighty consideration, but

8. See, for example, Norcross (1997) and Horton (2018) for arguments for this *fully aggregative* view.

it also seems that the greatest goods for humans—such as our health, freedom, and flourishing—are still important moral factors that can be morally decisive. According to Irrelevance of Tiny Benefits, however, there is no number of any benefits for humans such that one ought to prioritize them since all human benefits are tiny compared to the benefit of saving the individual of the high-welfare species. This seems to be the wrong result. Surely, for example, there is *some* number of human lives such that we should save them rather than the one individual of the high-welfare species.

To see how plausible and weak this claim is, compare it to the intuition that we ought not to give all of our resources to a ‘utility monster’ that gains much more well-being from each resource than we do in total (Nozick 1974: 41). This widespread intuition that is frequently employed against utilitarianism is sufficient to argue against Irrelevance of Tiny Benefits: If we should sometimes benefit humans even though we could alternatively benefit the utility monster by a much larger magnitude, much smaller (or ‘comparatively tiny’) benefits can sometimes outweigh much larger benefits. While I take this to provide some reason to reject Irrelevance of Tiny Benefits, I don’t take it to be a decisive reason since the intuition that we sometimes ought to prioritize humans over the utility monster can be doubted. After all, the benefits for the utility monster are not only much larger than each individual human benefit but also larger than the *sum* of all human benefits. (Otherwise, utilitarianism would agree that we should not give all of our resources to the utility monster.) It could be held, therefore, that while human benefits could *in principle* outweigh the benefit to the utility monster, they do not actually outweigh it since they don’t add up to a sufficiently large total benefit.

This answer to the utility monster case is compatible with my claim that benefits for members of high-welfare species can be outweighed by some human benefits that are comparatively tiny. For my claim, it is sufficient that some human benefits could outweigh the benefit to the utility monster *in principle*. That is, I only need to claim that we should prioritize the human beneficiaries over the utility monster if the human beneficiaries were much more numerous (while the benefit to the utility monster remained unchanged). Given that each human beneficiary is benefited in a significant way, *this* seems close to indubitable. In particular, it seems clear that we should prioritize the human beneficiaries if they were so many that the sum of their benefits is *larger* than the benefit to the utility monster. Assume, for example, that the benefit to the utility monster is about a *million* times larger than the largest human benefit and that, with a given amount of resources, we could benefit either the utility monster or one *billion* humans in these ways. Benefitting humans would then be about one thousand times *more* effective (in terms of the total benefit per resource) than benefitting the utility monster. Of course, the ‘utility monster’ would then cease

to be a real utility monster since it would not derive more well-being from each resource than we do in total. But it would still be an example of an individual of a high-welfare species to which Irrelevance of Tiny Benefits accords priority: Due to its much larger welfare range, we can benefit it much more than we could possibly benefit an individual human being. This verdict of Irrelevance of Tiny Benefits seems clearly wrong. Even if we should give all of our resources to Nozick's utility monster, we should not spend all of our resources on such high-welfare species that we can benefit much less effectively than humans.

As I said, this claim is very weak. It is compatible with the view that we should prioritize one individual of a high-welfare species over a single human or even one million humans. It is also compatible with the view that, unit for unit, the well-being of high-welfare species matters more (e.g., due to their more developed psychological capacities) than our well-being (see Kagan 2019: 149). I only claim that we should not give *absolute* priority to high-welfare species whenever significant goods of theirs are at stake. It also seems that we don't need to be very suspicious about this claim. Not only does it not imply priority for humans but merely lack of absolute priority for other beings; we also have barely any self-interest in that claim being true since, as far as we know, there are no species with welfare ranges that are much larger than our welfare range.⁹ This is in contrast to the question at issue whether we sometimes have absolute priority over low-welfare species: those probably exist, and it is in our self-interest to have absolute priority over them.

I conclude that Irrelevance of Tiny Benefits should be rejected and that, accordingly, the above argument for the view that no benefits for low-welfare species can outweigh a human life is unsound. Moreover, the consideration of high-welfare species poses a general challenge for this view: to motivate it, we would need a principle according to which no benefits for low-welfare species can outweigh human lives while some benefits for humans *can* outweigh lives of high-welfare species. As argued above, Irrelevance of Tiny Benefits cannot meet this challenge. More generally, it is hard to see how the challenge could be met by *any* well-motivated (and similarly parsimonious) principle. An obvious response is that benefits for low-welfare species are much smaller than benefits for humans. While benefits for low-welfare species are small compared to human benefits and human benefits are small compared to benefits for high-welfare species, there is still the fact that some human benefits are large while all benefits for low-welfare species are small. In other words, some human benefits,

9. Interestingly, there could be high-welfare species in the future due to the possibility of sentient artificial intelligences (Shulman and Bostrom 2021). But this is a very unusual view that most people don't believe in. So this possibility plausibly does not affect our intuitions.

unlike any benefits for low-welfare species, are above a certain threshold above which all benefits are (intuitively) large and therefore relevant.

While this response gives us the desired results, it is also suspiciously convenient and anthropocentric. Our judgment that all benefits for low-welfare species are small seems to merely reflect that they are small compared to the range of human benefits. We may ask ourselves where members of high-welfare species would set the threshold above which benefits count as large enough to be relevant. Presumably, they would set the threshold at the level of benefits that are relatively large for *them* and hence higher than the largest human benefits, which *they* would consider small. Even if there is a threshold above which all benefits are relevant, we have no good reason to believe that our placing of the threshold is the right one. We should therefore not rely on this proposal to explain why some human benefits are relevant to much larger benefits for high-welfare species.¹⁰

In what follows, I will argue that, even when setting this challenge aside, there are good reasons to accept the opposite view on which some benefits for low-welfare species *can* outweigh human lives. I will present two arguments that attempt to make this view directly plausible. These arguments will also illuminate why it seems wrong to prioritize large benefits to members of high-welfare species over any number of any benefits for humans.

The Argument from Fairness

The first argument is based on the value of *fairness*. As far as this value is concerned, it seems plausible that we should at least sometimes benefit the members of a given species if we can benefit them effectively enough. More precisely, for some (arbitrarily large) sum of benefits *S*, if we never benefit the members of a species even though we can benefit so many of them in individually significant ways that the sum of their benefits exceeds *S*, we are treating them unfairly. Intuitively, low-welfare species are not an exception to this rule. If anything, fairness requires that we give their welfare *special* consideration to compensate for their lower natural capacity for welfare, so that it is *especially* unfair not to

10. Another possible answer to the challenge is that some human benefits are more than merely finitely many times larger than any benefit to (common) low-welfare species. This would break the symmetry given my assumption that the benefits for the high-welfare species are only finitely many times larger than human benefits. This answer is natural for those who believe that, even among human welfare goods, some 'lexically superior' goods contribute more to one's welfare than any amount of other goods (e.g., Dorsey 2009). However, these views about human well-being have highly implausible implications (see, e.g., Schönherr 2018). We may also wonder if there might be goods for high-welfare species that are lexically prior to any of our goods and whether we would be willing to assign absolute priority to them under this assumption.

benefit them if we can do so effectively.¹¹ Either way, fairness seems to require that, if we can benefit them effectively enough, we benefit low-welfare species at least sometimes.

This fairness requirement is in tension with the view that no number of even the most significant benefits for low-welfare species can outweigh one human life. This view implies that *no matter* how effectively we can benefit low-welfare species, we should not benefit them when we could save a human life alternatively. In a world like ours in which we could virtually always use our resources to save human lives with statistical certainty (e.g., by donating to effective charities), this would mean that we should *never* benefit low-welfare species even if we could benefit them much more effectively than we actually can. The only exception might be decisions about how to use small amounts of resources that do not suffice to help any human significantly, such as when we decide whether to use a minute of our time to help a bee that is trapped in a room. However, many small amounts of resources add up to a big amount with which a human *could* be benefited significantly. As long as we foresee that there will be sufficiently many situations in which we could decide to spend small amounts of resources on low-welfare species, it seems, therefore, that we should decide to always save the resources—at some point, these savings will add up to an amount that is sufficient to significantly benefit a human. Arguably, we should not even waste our time, then, to help bees escape our room.¹²

Even though we are used to prioritizing humans over (what we take to be) low-welfare species, this seems to be an extreme form of discrimination. Compare this to the fairness concern that a broadly utilitarian approach might require us to systematically prioritize the lives of healthy people over the lives of disabled people when the latter tend to contain less welfare (e.g., John et al. 2017; Kamm 2015). Whether or not it would be unfair to prioritize healthy lives over equally many disabled lives with lower welfare, which is a valid concern in itself, it *obviously* would be unfair to prioritize one healthy life over *any* number of disabled lives. The fairness claim that I am making with regard to low-welfare species is analogous to this extremely weak claim: It would be unfair to prioritize one human life over *any* number of lives of low-welfare species.

11. This fairness requirement could be spelled out in egalitarian or prioritarian terms. Alternatively, it could be argued that *compassion* requires that we give low-welfare species special consideration. Crisp (2003: 761) entertains the view that an impartial spectator might ‘feel compassion exactly in proportion to levels of overall welfare’.

12. Admittedly, it is often not possible to save nonmonetary resources in the straightforward way in which we can save money. However, it also seems plausible for nonmonetary resources that, *in expectation*, we will have more of these resources available for some valuable uses if we don’t occasionally spend small amounts of them on other things. In particular, it seems that, in expectation, we will have more time available to benefit humans if we decide to never spend any time on helping insects (except when this benefits humans).

While the discrimination against low-welfare species might seem less morally grave, it is important to see that there still is a fairness concern that we don't have when it comes to prioritizing human lives over human *headaches*: Intuitively, we do not discriminate against the headache patients by prioritizing someone's life over the alleviation of their headaches. On the contrary, it seems that it would be unfair against the person whose life is at stake if we prioritized the headache patients. These intuitions suggest that, in decisions in which human lives are at stake, we should not treat significant benefits for low-welfare species in the way that we treat the alleviation of headaches; even if the latter is irrelevant in such decisions, the former seem to remain a relevant moral factor.

The Argument from Empathy

The second argument can be considered an argument from *empathy*. Unlike the argument from fairness, this argument is not based on an interspecies welfare comparison but focuses on the perspective of members of the species with the smaller welfare range. The point is that from *their* perspective the largest goods attainable to them are significant even if they are much smaller than the goods that are attainable to us. Of course, this may not be their *actual* perspective. After all, low-welfare species might lack the cognitive capacities to contemplate the significance of their goods. But we can still ask which perspective would be *appropriate* to have for an individual with the welfare structure of a low-welfare species. From this (perhaps only hypothetical) perspective, it seems that some benefits for members of low-welfare species are significant. Perhaps, for example, the joy that they can experience is much less intense or much shorter than the joy that we are able to experience. Still, if that is the most joy that they can experience, experiencing it would constitute a highlight in their lives that is worth aspiring to. Or consider their lives: From our perspective they may seem insignificant because they have, for example, only a tiny fraction of the length of our lives. But for them, their lives are all that they have, and preventing their lives from being cut short is a worthy goal.

None of that is to say that their capacity for welfare is not in fact much lower than our capacity for welfare or that this difference is morally irrelevant. However, this change in perspective is supposed to show that the goods of low-welfare species can have a significance that makes it inappropriate to consider them irrelevant whenever important goods for us are at stake. It would be *arrogant* and perhaps *disrespectful* to consider their largest goods irrelevant whenever our largest goods are at stake.

Comparing Relative Magnitudes

I have argued that we should reject principles that ask us to prioritize a single human life over any number of any benefits for members of low-welfare species. How should we decide, then, in which cases some number of benefits of one kind can outweigh benefits of another kind? The arguments from fairness and empathy suggest that, rather than only considering the *absolute* magnitudes of benefits, we need to consider how large benefits are *relative* to the beneficiaries' welfare ranges or some related welfare quantity that differs between the beneficiaries. That is, we need to put each benefit in relation to this welfare quantity of the beneficiary, such as the beneficiary's welfare range, and then compare the resulting *relative benefit magnitudes*. If a benefit has a sufficiently large relative magnitude, it is relevant even if its absolute magnitude is small. This would allow us to maintain that each individual has some welfare goods that are never irrelevant, and that we should prioritize the members of each species in at least some hypothetical circumstances in which we could benefit them effectively enough.

In relation to which welfare quantity should we consider benefits? Since my discussion has been in terms of species with different welfare ranges, it is natural to assume that we should take the welfare range of the beneficiary's species. Although this would give us the right results in most cases, I think it is clear, however, that species' welfare ranges only matter derivatively. At least, they matter only insofar as they correlate with the *individual's* welfare range. This correlation is usually strong, but an individual's welfare range can also significantly come apart from the welfare range of its species, as in the case of some irremediable chronic diseases. Benefits to such individuals should not be considered irrelevant merely because they are small relative to the species' welfare range. As long as a benefit is large compared to the individual's restricted welfare range, it should be considered relevant.

It could be suggested, furthermore, that individuals' welfare ranges only matter insofar as they correlate with their *actual* welfare. This would mean that we need to put benefits in relation to the beneficiary's actual welfare, which may be much lower than the maximum level of the beneficiary's welfare range, and then compare *these* relative magnitudes. This is a natural suggestion for those who take the actual welfare of the beneficiaries to affect the relevance of benefits for reasons that are independent of the consideration of low-welfare species (e.g., Voorhoeve 2014).¹³ As noted before, however, if this view is to allow some

13. More precisely, Voorhoeve (2014) suggests giving more weight to benefits for beneficiaries who *would* have a lower level of well-being if they are *not* benefited. Whether we take the beneficiaries' well-being conditional on being benefited or conditional on not being benefited does not make a difference to my subsequent argument.

tiny benefits for individuals with low welfare to be relevant, such as in the case of benefits for low-welfare species, it also needs to accept that tiny benefits for humans with equally low welfare are relevant. When, additionally, their welfare *range* is very small, this is not implausible. For example, the alleviation of headaches should be considered relevant to the saving of a human life in the very unusual case where the headache patients have an extremely limited capacity for welfare (e.g., due to an extremely severe disability), so that the benefit of alleviating their headaches is significant compared to the welfare that they could ever realistically achieve. This case is analogous to the case of conferring significant benefits to members of low-welfare species, which, as I have already argued, should be considered relevant to the saving of human lives. However, should we also consider headaches to be relevant to human lives if the headache patients have a normal human capacity for welfare and only coincidentally have a very low welfare? The answer to this question decides whether we should choose the welfare range or the actual welfare of the beneficiary as the relational factor.¹⁴

Whether we take the welfare range or the actual welfare of beneficiaries, I think that a principle on which we can determine the relevance of benefits by comparing their relative magnitudes is promising. Such a principle would give us the same plausible results as Compare Absolute Magnitudes for common intraspecies comparisons in which the beneficiaries are assumed to have similar welfare ranges and to be (if benefited) similarly well off. In that case, the relational factor is roughly the same for all beneficiaries and therefore does not make a difference. We could thereby maintain, for example, the intuitive verdict that headaches cannot outweigh human lives when ‘all else is equal’. At the same time, comparing relative magnitudes avoids the problems discussed above for Compare Absolute Magnitudes that arise in interspecies comparisons. By comparing relative magnitudes of benefits, we would not discriminate against species with smaller welfare ranges and would do justice to their distinct perspectives. Most importantly, we could explain why the largest attainable benefits for a species, whether for humans or for low-welfare species, are always relevant: These benefits have the largest possible relative magnitude.

It is tempting to conclude that we *only* need to compare relative magnitudes to determine whether a given benefit is relevant: For all benefits B_1 and B_2 , B_1

14. A potential (third) solution would be to tie the relational factor to the notion of *fortune*, which is generally taken to be a relation between the individual’s well-being and some appropriate standard (e.g., McMahan 2002). The difficult question is what this standard is, where both the species’ and the individual’s potential or capacity for well-being are debated candidates (McMahan 2002; Vallentyne 2005). Given these parallels, it could be suggested that this standard, whatever it is, is the same as the factor relative to which benefits need to be considered when testing for their relevance. A principle on which we should compare relative benefit magnitudes could then be interpreted as a principle on which we should compare the increases in fortune that the benefits at issue would result in.

is relevant to B_2 if and only if the relative magnitude of B_1 is sufficiently large compared to the relative magnitude of B_2 . This simple principle, however, which we can refer to as *Compare Relative Magnitudes*, has a serious problem as well: Minor benefits for humans would be irrelevant compared to significant benefits for low-welfare species. Assume, for example, that alleviating a human headache has the same benefit magnitude as saving the life of a bee while the welfare range (as well as the actual welfare) of the human patients is much larger than the welfare range (and the actual welfare) of the bee. The *relative* benefit magnitude of alleviating the headaches is then much *smaller* than the relative benefit magnitude of saving the bee's life. Compare Relative Magnitudes would therefore consider the human headaches to be irrelevant, so that we should save the bee's life rather than alleviate *any* number of headaches. But surely we should spare the whole human population from suffering a headache rather than save the life of a single bee. Compare *Absolute* Magnitudes gives us the right result here: Per assumption, the absolute benefit magnitudes of alleviating a human headache and of saving a bee's life are roughly the same. Consequently, the human headaches count as relevant and can thus outweigh the bee's life if they are numerous enough.

My tentative proposal is therefore to combine Compare Absolute Magnitudes and Compare Relative Magnitudes as follows: Each principle provides a *sufficient* condition for a benefit's relevance, but it is only necessary that *one* of these two sufficient conditions is fulfilled. That is, a given benefit is relevant if and only if its absolute magnitude *or* its relative magnitude is sufficiently large. More precisely, for all benefits B_1 and B_2 , B_1 is relevant to B_2 if and only if the absolute magnitude of B_1 is sufficiently large compared to the absolute magnitude of B_2 *or* the relative magnitude of B_1 is sufficiently large compared to the relative magnitude of B_2 .¹⁵ This principle—*Compare Absolute and Relative Magnitudes*—gives us the intuitively right results in all considered cases: headaches cannot outweigh human lives (since both the absolute and the relative magnitude of headaches is too small); headaches can outweigh bee lives (since the absolute magnitude of headaches is large enough); human lives can outweigh lives of high-welfare species (since the relative magnitude of human lives is large enough); and lives of low-welfare species can outweigh human lives (since the relative magnitude of lives of low-welfare species is large enough).

15. As with Compare Absolute Magnitudes, we could alter this principle to give extra weight to benefits for beneficiaries who are less well off (if we choose welfare ranges as the relational factor), and we could add further necessary conditions to, for example, exclude benefits that are based on immoral preferences.

Conclusion

There is of course more to be said about Compare Absolute and Relative Magnitudes. For one thing, it is less parsimonious than the simple Compare Absolute Magnitudes and Compare Relative Magnitudes. However, I think that my discussion shows that there is at least a promising (type of) principle that can explain why some benefits for members of low-welfare species are relevant even when human lives are at stake. Since this verdict gains plausibility from considerations of fairness and empathy, while principles that deny it face the problem of high-welfare species, I conclude that the overall most plausible view is that some benefits for low-welfare species can in principle outweigh even the most significant benefits for humans. This conclusion is very preliminary, though. The arguments from empathy and from fairness are by no means knock-down arguments, and perhaps there is a principle that can respond to the challenge of high-welfare species in a satisfactory way. Further investigation of these issues is needed to reach a conclusion in which we can be more confident.

Acknowledgments

For helpful comments and discussion, I am very grateful to Ben Davies, Leonard Dung, Tomi Francis, Andreas Mogensen, Jeff Sebo, Teruji Thomas, and an audience at the SSEA colloquium series. While working on this paper, I have been generously supported with a grant from EA Funds and a scholarship from the Heinrich Böll Foundation.

References

- Crisp, Roger. 2003. 'Equality, Priority, and Compassion', *Ethics*, 113.4, pp. 745–763
- Dorsey, Dale. 2009. 'Headaches, Lives and Value', *Utilitas*, 21.1, pp. 36–58
- Fischer, Bob. 2023. 'Rethink Priorities' Welfare Range Estimates', *Rethink Priorities*, 23 January <https://rethinkpriorities.org/publications/welfare-range-estimates> [accessed July 29, 2025]
- Fischer, Bob. 2024. *Weighing Animal Welfare: Comparing Well-Being Across Species* (Oxford University Press)
- Horta, Oscar, and Dayron Teran. 2023. 'Reducing Wild Animal Suffering Effectively', *Ethics, Policy and Environment*, 26.2, pp. 217–230
- Horton, Joe. 2018. 'Always Aggregate', *Philosophy & Public Affairs*, 46.2, pp. 160–174
- John, Tyler M., Joseph Millum, and David Wasserman. 2017. 'How to Allocate Scarce Health Resources Without Discriminating Against People with Disabilities', *Economics and Philosophy*, 33.2, pp. 161–186, doi:10.1017/S0266267116000237
- Kagan, Shelly. 2019. *How to Count Animals, More or Less* (Oxford University Press)

- Kamm, Francis. 2015. 'Cost Effectiveness Analysis and Fairness', *Journal of Practical Ethics*, 3.1
- McMahan, Jeff. 2002. *The Ethics of Killing* (Oxford University Press)
- Norcross, Alastair. 1997. 'Comparing Harms: Headaches and Human Lives', *Philosophy & Public Affairs*, 26.2, 135–167
- Nozick, Robert. 1974. *Anarchy, State, and Utopia* (Basic Books)
- Parfit, Derek. 1978. 'Innumerate Ethics', *Philosophy & Public Affairs*, 7.4, pp. 285–301
- Parfit, Derek. 1984. *Reasons and Persons* (Oxford University Press)
- Scanlon, T. M. 1998. *What We Owe to Each Other* (Belknap Press)
- Schönherr, Julius. 2018. 'Still Lives for Headaches: A Reply to Dorsey and Voorhoeve', *Utilitas*, 30.2, pp. 209–218, doi:10.1017/S0953820817000152
- Sebo, Jeff. 2023. 'The Rebugnant Conclusion: Utilitarianism, Insects, Microbes, and AI Systems', *Ethics, Policy & Environment*, 26.2, pp. 249–264, doi:10.1080/21550085.2023.2200724
- Shulman, Carl, and Nick Bostrom. 2021. 'Sharing the World with Digital Minds', in *Rethinking Moral Status*, ed. by Steve Clarke, Hazem Zohny, and Julian Savulesc (Oxford University Press) doi:10.1093/oso/9780192894076.003.0018
- Taurek, John M. 1977. 'Should the Numbers Count?', *Philosophy & Public Affairs*, 6, pp. 293–316
- Vallentyne, Peter. 2005. "Of Mice And Men: Equality and Animals", *The Journal of Ethics*, 9, pp. 403–433, doi:10.1007/s10892-005-3509-x
- Voorhoeve, Alex. 2014. 'How Should We Aggregate Competing Claims?' *Ethics*, 125.1, pp. 64–87