

BETWEEN THE SPECIES

Towards a Coherent Theory of Animal Equality

ABSTRACT

In this article I want to construct in a simple and systematic way an ethical theory of animal equality. The goal is a consistent theory containing a set of clear and coherent universalized ethical principles that best fits our strongest moral intuitions in all possible morally relevant situations that we can think of, without too many arbitrary elements. I demonstrate that impartiality with a level of risk aversion and empathy with a need for efficiency are two different approaches that both result in the same consequentialist principle of prioritarianism. Next, I discuss that this principle can be trumped by an ethic of care principle of tolerated partiality, and a deontological principle of basic right. These three principles represent different kinds of equality that can be applied to animal ethics. Finally, the predation problem leads to the introduction of a triple-N-principle that is related to the value of biodiversity.

STIJN BRUERS
Ghent University
stijn.bruers@ugent.de

Volume 17, Issue 1

Jun 2014

© *Between the Species*, 2014
<http://digitalcommons.calpoly.edu/bts/>

Introduction

In this article my ambitious goal is to construct a coherent ethical system that is capable of dealing with all relevant issues in principle-based animal ethics. The basic line of reasoning of this construction goes as follows: I start with a factual property of the world, which ignites a moral intuition or emotion, i.e. a direct moral response or judgment that has no further rational justification. Then, in a process of reflection, this intuition is translated into a universalized ethical principle, where “universalized” means: “relevant to all morally similar situations.” Sometimes different moral intuitions will mutually support each other, resulting in a set of coherent universalized ethical principles. But sometimes we encounter a new fact or situation that again ignites another moral intuition or emotion, which might be in contradiction with our constructed set of universalized ethical principles. To solve this conflict or moral dilemma, we can either change the ethical principles, or introduce a new ethical principle that trumps the previous ethical principles in that particular situation. This new ethical principle needs to be universalized as well to all relevantly similar situations.

This process continues: we again test the constructed coherent set of universalized principles in new situations, and if we encounter a moral dilemma, we look for further refinements. Eventually, all situations and all facts that ignite moral intuitions should be covered, and we move to a consistent ethical system of hierarchical universalized principles, where some principles trump others. In other words, we reach a theory in “reflective equilibrium” (Rawls, 1971), which means that our strongest moral intuitions and ethical principles are coherent (mutually supporting each other).

This approach can be compared with solving a crossword puzzle. The descriptions of the words are the analogues of objective facts in the world. The white boxes refer to the possible situations, the individual letters represent the intuitive moral judgments in particular situations. The words correspond with the universalized ethical principles (applied to all similar situations), and these words mutually support each other and form a coherent solution to the puzzle. So let's derive a coherent ethic of animal equality, starting from the most basic, indisputable objective facts and moral intuitions.

The construction of a coherent system

Fact 1: All sentient beings have a well-being and they value their own well-being (and everything that contributes to well-being). Sentient beings are beings that have and can subjectively feel interests. Things subjectively matter to them, meaning first of all that the individual has a mechanism (i.e. a complex functioning nervous system) that enables the individual to have representations of their bodies and environments. These representations can have intentionality, resulting in qualitative experiences (phenomenological sensations or qualia). For example: through my fingers I can feel this book. I know the difference between this feeling and an absence of feeling, for example when my fingers are anaesthetized. However, just before I paid attention to this feeling of touch, I was not aware of it. There was an unconscious neural activity (but no anesthesia). Only after I focused on my fingertips, it became a conscious experience or "quale" of touch. Now, qualia are often neutral. I don't feel an urge to avoid touching books. But other qualia are affective in nature; they are evaluated as being positive or negative. A needle in my finger generates a quale that I wish to avoid. This quale is called pain and it generates an urge in me. Once a quale becomes an affective mental state (i.e. a positive

or negative feeling or emotion such as pain, distress, joy,...), well-being comes into play. These feelings are related to interests or needs: they are nothing but subjective experiences of (un)satisfied interests. Fear, pain, and frustration indicate that the needs for respectively safety, bodily integrity, and freedom are not satisfied.

Moral intuition 1: Impartiality is morally important. We can consider a two-step process to increase impartiality, from rational egoism to extended contractualism. A rational egoist would strive for a contractarian ethic (cf. Thomas Hobbes), where all rational beings (i.e. beings with whom one can negotiate) of equal power will become part of the moral community, because those rational egoists gain mutual advantages through the social contract. However, in a first step to extend impartiality, Rawls (1971) used the method of the veil of ignorance to delete the second condition of equality of power. He arrives at a contractualist ethic that also includes rational people in dependent or weaker positions (minorities, future generations,...). The veil of ignorance is a thought experiment, whereby you imagine that you will be born as a rational agent, but you don't know who you will be. You can determine the moral and political laws, based on your knowledge of the natural laws. I would suggest a second step to extend impartiality, whereby we delete the condition of rationality. Imagine that you might be any object or entity in the world, but you don't know who or what you might be. For complete impartiality, you have to imagine you could be a planet, an electron, a pig in the year 3000 or anything you can think of. How would you like that entity to be treated? If you were non-sentient, this question would not matter to you, because nothing done to you will influence your well-being (you would not have a well-being). The kind of treatment becomes important only for those beings whose

well-being can be influenced by moral agents. Non-sentient entities should not be taken into account in this moral evaluation. So the least arbitrary and most impartial thing to do is to delete both conditions (of rationality and equality of power), which is what Rowlands (1998) argued, from which it follows that well-being still remains important.

Universal ethical principle 1: All moral agents should strive towards impartiality in all situations, and should take everyone's well-being into consideration in an impartial way. Moral agents are people who are able to understand the notion of impartiality.

Fact 2: Empathy is meaningful for all and only for sentient beings (feeling empathy for non-sentient beings such as teddy bears would be a kind of projection of emotions). Empathy is the capacity to experience or sample the emotions of others. This emotional response occurs when the perspective (frame of reference) of the other is taken.

Moral intuition 2: Compassion (empathy plus the desire to alleviate the suffering of the other) is a virtue.

Universal ethical principle 2: All moral agents should develop compassion in all situations (hence also towards all sentient beings). Moral agents are people who are able to develop compassion, are able to understand the virtue of compassion, and are able to help others. Those moral agents should try to improve the well-being of others.

The above two universal ethical principles are coherent with each other, and give a rational and emotional basis of the moral importance of sentience. They are based on contractualism, consequentialism, and virtue ethics. The coherence gets even

stronger when we consider the following moral intuitions. A) Mental capacities (self-consciousness, rationality,...) are morally important. They are very special, complex, and vulnerable, hence worth protecting. B) Babies and mentally disabled humans have rights because they have something morally important. They have a higher moral status than human egg cells, skin cells, dead human bodies, plants, or stones. Together with the fact that sentience is the only mental capacity that mentally disabled persons have in common with other humans, it follows from A and B that sentience is important. The link between rights and sentience is also not farfetched: rights protect interests; feelings detect interests.

This gives us a strong coherent case for the moral relevance of sentience. It is a scientific question (i.e. a matter of fact) what entity has a well-being and how its well-being can be influenced. We can briefly compare this moral relevance of sentience with the moral irrelevance of a criterion such as the species *Homo sapiens*. First, the species is one of the many biological classifications, thus it is arbitrary to pick a specific species and not a specific population, genus, family, order, class,... Second, the definition of a species is very complicated. One of the definitions refers to a set of individuals who could get fertile offspring. But reference to fertility and offspring is very artificial and farfetched when it comes to determining who has rights. Third, science will never be able to determine whether a human-chimpanzee hybrid, a human-animal chimera, an ancestor (*Australopithecus*, *Homo habilis*,...) or a genetically modified humanoid should still be called *Homo sapiens*. The boundaries are fuzzy. Fourth, all species are temporally related to all other species in a similar way, as populations can be spatially related in a ring species (a ring species consists of a spatial spreading of populations, where A can get fertile off-

spring with B, B with C, but C not with A). Fifth, if the moral status of a species is determined by genes or bodily appearance, then it is also very arbitrary to pick out those genes or bodily characteristics and not others (such as skin color,...). We are not responsible for our genes, so it would be a violation of the desert principle if we based moral status on genes. In summary, the species boundary is too arbitrary, artificial, and abstract to be morally relevant.

So far, our ethic is not yet unambiguous and clear. We observe that there are different sentient beings, and multiple ways to influence their well-being (for example: increasing everyone's well-being a little bit versus increasing the well-being of one individual a lot). So what is a just distribution of well-being? First of all, we value parsimony and simplicity. One simple solution would be to add the levels of well-being of all sentient beings for a specific time interval, and then take the sum over all times. Then we could try to maximize this sum. This is sum-utilitarianism. But there are also other simple options, such as trying to maximize the well-being of the worst-off sentient being (the one with the lowest level of well-being). This is maximin-utilitarianism. However, according to many people, both sum-utilitarianism and maximin-utilitarianism have some counterintuitive implications. With sum-utilitarianism, it is morally good to sacrifice one individual in order to increase the well-beings of others, or to kill one individual and replace him with another sentient being, or to keep on breeding sentient beings in order to increase the sum of well-being. The latter is known as the "repugnant conclusion" (Parfit, 1984): an overpopulated world with a trillion individuals with a well-being slightly above zero, might be better than a world with only a thousand individuals who have a satisfying high level of well-being. Our moral intuitions go against these conclusions.

These conclusions can be avoided by introducing a level of risk aversion.

Fact 3: There are many sentient being, and some beings can be worse-off than others. This fact implies that from behind the impartial veil of ignorance, how to maximize your well-being becomes a game of chance. Mathematically, sum-utilitarianism implies that the expectation value of your well-being will be maximized. But you have to be aware that there is a risk that you might be born as one of the worst-off individuals. For example: two individuals might have well-beings equal to 10 and 100, so the expectation value will be equal to 55 (the average). In sum-utilitarianism, this situation would be equal to the situation where those two beings both have a well-being of 55. The problem is that in the first situation, you might end up as the person with level 10. When much is at stake, most moral agents have a risk aversion (need for safety—to play it safe), and in this game of chance, this means that they would not opt for sum-utilitarianism, but to some kind of prioritarianism: giving priority to increases of well-being of the worst-off positions. Therefore they prefer the second situation. If you have maximum risk aversion (a maximum need for safety), you would take the maximin-utilitarian strategy (maximizing the minimum/lowest well-being), giving all priority to the worst-off position, because you are so worried at becoming this worst-off individual. If you have zero risk aversion, you are a sum-utilitarian. A high but not maximum level of risk aversion would result in a prioritarianism that is in between maximin-utilitarianism and sum-utilitarianism. We could call this ‘quasi-maximin prioritarianism’.

Moral intuition 3: A (high) level of risk aversion is good (especially when much of your well-being is at stake; then most people are risk averse).

Universal ethical principle 3: Quasi-maximin prioritarianism should be applied in all situations. Mathematically, this principle can be expressed as the maximization of the expectation value of a weighted average of qualities of life of all sentient beings. Let's clarify this expression. The maximization runs over all available choices from behind a veil of ignorance. Each choice gives a different world-history. In each choice, we only consider the sentient beings that exist in that world-history (in the present or the future), and only those sentient-beings whose well-being can be influenced by our choice. The expectation value is in case the outcomes of qualities of life are uncertain (then it becomes a double game of chance: first, from behind the veil of ignorance you don't know which one of the possible beings you will be, and second, you don't know exactly what qualities of life those possible beings will get). The weighted average is over all sentient beings that exist in the present or will exist in the future, whereby the lowest qualities of life get the highest weight factors, resulting in a higher priority to maximize those lowest positions. So it is a priority averaged well-being that matters. (When the lowest position gets the weight 1 and the others 0, then we obtain maximin-utilitarianism. When all qualities of life get an equal weight factor, we arrive at sum-utilitarianism.) The quality of life refers to the total preferred well-being of an individual over his/her complete lifespan. This preferred well-being is the value that one would ascribe to living the complete life of that individual, when looking from the most impartial point of view, e.g. from behind a veil of ignorance. The quality of life contains everything that

would matter to you, all the preferences that you would have, if you would live the life of that sentient being.

Quasi-maximin prioritarianism has some elegant features. It avoids the abovementioned objections against sum-utilitarianism, and also a lot of objections against animal ethics. First, consider the idea of painlessly killing someone (for example in his sleep). From behind the veil of ignorance, you cannot prefer such killing, even if you are not aware that you will be killed. This means that a sentient being should now be defined as a being that has already developed the capacity to feel and has not yet permanently lost this capacity. Indeed, quality of life starts from the first feeling and ends at the last feeling.

Next, take the problem of replaceability. Is it allowed to kill a sentient being (painlessly), and then let another sentient being be born? This happens when we breed and slaughter cows. If we kill a sentient being, his quality of life will be e.g. 5, whereas it would have been 10 otherwise (when he lives a full life). So in a first option, one individual will have a life with total well-being equal to 5, and a second one will have at most level 10. In a second option, we will have only one being, with level 10. From behind the veil of ignorance, in the first option you have probability $\frac{1}{2}$ to have a low quality of life equal to 5. In the second option, you are sure you will have level 10. A sum-utilitarian would say that the first option is better, because the total quality of life equals 15, which is higher than 10 in the second situation. But being risk averse, I would prefer the second situation, and that's also what our prioritarian theory says. Therefore, sentient beings are not replaceable. Also the repugnant conclusion (the idea to keep on breeding sentient beings until their qualities of life are about to drop below zero), can be avoided, by simply noting that behind the veil of ignorance

you would not prefer an overpopulated world where everyone has a very low quality of life. So quasi-maximin prioritarianism avoids the often heard argument that breeding livestock animals is good, because they owe their lives to the breeders, and it is better to live a life on a farm than not to be born at all. According to our prioritarianism, the choice is not between an existing life on a farm versus a non-existing life, because as said above: in each choice, we only consider the sentient beings that exist in that world-history.

Another famous problem in animal ethics is the lifeboat dilemma (e.g. Regan, 1983). Suppose there are different sentient beings in a lifeboat, but we cannot save everyone. Those beings can have different expected life expectancies, but they can also differ in complexity (richness) of emotions, the amounts of needs, the levels of satisfaction when needs are satisfied,... This means that the potential qualities of life can differ amongst the different sentient beings in the lifeboat. The potential qualities of life between a (mentally normal) human, a dog, or a frog can differ. This influences our choices whom to rescue. As Regan argued, it might be required to sacrifice the dogs first, because they experience a less rich life than the humans. However, Regan said that the life of one human would trump the lives of a million and more dogs. According to our prioritarianism (the veil of ignorance with a high but not maximum level of risk aversion), there would be a number of dogs, above which the loss of that amount of dogs would be worse than the loss of one human life.

The quasi-maximin principle is coherent with a lot of our moral intuitions. And there is a second way to arrive at this principle.

Fact 4: There might be situations where we can decrease someone's well-being with a huge amount (drive him into extreme poverty) in order to increase the worst-off position with a negligible small amount.

Moral intuition 4: Efficiency is important to some degree. Empathy might have a tendency to give absolute priority to improving the worst-off individual, which results in a maximin strategy. But if we value efficiency, we would not sacrifice too much well-being.

Universal ethical principle 4: This equals quasi-maximin prioritarianism (principle 3). We should maximize the quality of life of all sentient beings, giving a strong priority to increase the lowest values of well-being. In other words: we should maximize the quality of life of the worst off individuals, unless this is at the expense of much more well-being of others.

In summary: a rational approach of impartiality (the veil of ignorance) with a high but not maximum risk aversion (need for safety) coheres with an emotional approach of empathy with a low but non-zero need for efficiency. These are two approaches resulting in the same quasi-maximin prioritarian principle.

This principle has two disadvantages. As a first problem, the qualities of life are very difficult to measure and compare. All we have is our empathy, our scientific knowledge, and our imagination. We have to try placing ourselves in the position of others, by using empathy, or by imagining that we could be the other individual, with all his or her needs and feelings. So the "emotional" method of empathy and the "rational" method of the veil of ignorance are actually two rules of thumb to make educated guesses about the order of the qualities of life of dif-

ferent individuals. Empathy and imagination are virtues to be developed and already allow us to move quite far.

A second disadvantage is that the level of priority given to the worst-off (in other words: the level of risk aversion or the need for efficiency), is in some sense arbitrary. The level is somewhere between 0 (sum-utilitarianism with zero risk aversion) and 1 (maximin-utilitarianism with maximum risk aversion). However, I believe our coherent picture is strong enough to withstand this objection. The arbitrariness is less worse than overriding a coherent set of strong moral intuitions. The good thing is that no-one has a strong preference to a sharp level of priority. No-one says the value should be 0,748. It's more like a fuzzy range that we prefer. So we can and should be a bit tolerant to the levels of priority that other moral agents would prefer, and this means we can be flexible and could come to a democratic or mutual consensus between all moral agents. But once we have set a level of priority, we should apply it consistently in all relevantly similar cases.

The quasi-maximin prioritarianism is the basic framework of a coherent ethical system of animal equality. All sentient beings are in some sense equal from an impartial perspective such as behind a veil of ignorance. It is a consequentialist ethic, because it only looks at outcomes of qualities of life. Giving a level of priority for the worst-off positions, some people (true consequentialists) might prefer to stop the construction of a coherent ethical system here. However, there are some more intuitions that do not nicely fit in the prioritarian ethic. We first discuss an intuition related to an ethic of care, and next an intuition related to an ethic of rights.

Fact 5: There is a possible situation where I have to choose between a sentient being I hold dear and one or more other unknown sentient beings, e.g., in a burning house dilemma, where I have to choose between saving my child or other individuals from the flames

Moral intuition 5: I am allowed to help the person I hold dear.

Universal ethical principle 5: It is allowed to be partial in all situations where someone is involved whom you hold dear (with whom you have a personal relationship or strong feelings of empathy), as long as we tolerate similar levels of partiality of everyone else. This principle of tolerated partiality trumps the above prioritarian principle to some degree, but not too much.

Burning house dilemmas such as “Your child or the dog?” (Francione, 2000) are often used to counter animal equality. But here we introduced a new principle of tolerated partiality, which hides a new kind of equality. In the burning house, I would save my child instead of someone else. But all individuals in the house are equal if I would tolerate your choice to save someone else. A white racist would say that it is immoral to save black children from the house instead of white children. A speciesist would say that it is immoral to save someone belonging to another species. But if someone has an emotional connection with a dog, we should tolerate his choice to save the dog. Saving a dog instead of a human, saving a mentally disabled orphan instead of a mentally normal child, or saving your lover instead of two unknown persons, might be violations of the quasi-maximin principle. But I think we are allowed to violate this principle to some degree. Also here we could try to reach a democratic or mutual consensus between all moral

agents, about the degree of violation that is allowed. We should apply this degree of partiality consistently in all situations.

Fact 6: The organ transplantation problem. There is a possible situation, where five patients in a hospital would die unless we sacrifice an innocent person against his will and use five of his organs for transplantations. This would be allowed according to prioritarianism.

Moral intuition 6: I (and most people) feel emotional distress and restraint to sacrifice this one person against his will. We should not sacrifice someone, even if prioritarianism is violated and even if someone I hold dear is one of the patients in the hospital. So this intuition trumps both prioritarianism and tolerated partiality.

There are a lot of other moral dilemmas where we can use someone without his/her consent as merely means to save others. Torturing someone in order to gain information about a bomb, throwing someone (a sentient being such as a mentally disabled human) in front of a trolley in order to block the trolley that is about to kill other people, using someone as a shield against bullets, using someone as a slave, using someone in medical experiments, terror bombing civilians in order to demoralize the enemy, raping someone, killing and eating someone (cannibalism), trafficking,... All these situations generate moral intuitions that are very coherent if we translate them into the following deontological principle (an extension of Kantian ethics).

Universal ethical principle 6: All sentient beings have a basic right not to be used as merely means to someone else's ends. A victim is used as merely means, when two conditions are met. 1) A moral agent causes the victim a "disrespectful

harm” against its will. A disrespectful harm means a treatment as property or commodity (see Francione, 2000) or a violation of bodily integrity. 2) The presence of the victim is required in order to reach the ends. The latter is an important criterion because there are moral dilemmas whereby you are allowed to cause harm to someone in order to save others. In those dilemmas, the presence of the victim was not required in order to save the others.

This principle is coherent with the notion of respect, which is next to empathy an important moral virtue, and it is coherent with the notion of intrinsic value (the opposite of instrumental value) as well.

The ethical principle of the basic right trumps both the principle of priority and the principle of tolerated partiality. But the basic right is not absolute: if the principle of priority is strongly violated (if thousands of sentient beings will die), then a basic right might be violated (this corresponds with a need for efficiency). As with the above principles, this level of violation can be determined on the basis of a democratic or mutual consensus among moral agents. And here we have flexibility as well: there are different levels of harm, there is a morally relevant gradation in someone’s ends (from the vital needs of many sentient beings to the luxury ends of one individual), and there is a gradation in the level of sentience and mental capacities. These gradations could be coupled. For example: a being with higher levels of morally relevant mental capacities has a stronger claim to this basic right.

Let’s briefly apply this principle to the “least harm” objection against veganism (Davis, 2003). Suppose that a meat eater can kill and eat one cow, whereas a vegan needs a crop field to

get the same amount of nutrients. Suppose using that crop field accidentally kills five mice. The meat eater causes least harm, but he violates the basic right of the cow, which is worse. The mice are not used as merely means, so therefore veganism remains the morally better choice. (For further criticism on the least harm argument of Davis, see Matheny, 2003, and Lamey, 2007).

We now arrive at an ethical system with three principles of equality. The first is based on impartiality (interchangeability of sentient beings) and results in a form of prioritarianism. The second is a tolerated partiality, whereby we tolerate the choices of others to save those they prefer. From this tolerated partiality, the individuals in a burning house inherit a “tolerated choice equality.” This principle weakly trumps the first principle. The third principle is a basic right equality, and this trumps the two former principles to a strong but not absolute degree. The three principles are related to respectively a consequentialist ethic of well-being and justice, a feminist ethic of care, and a deontological ethic of rights.

These three principles imply veganism. Consider a dairy cow in the livestock industry and a human who likes to eat cheese. Start with the veil of ignorance. In one situation, dairy cows are not bred, so we can only be a human being, who has a quality of life equal to 10. In the second situation, this human enjoys the cheese (his quality of life increases to 11), but the cow has a miserable life (suffering in the livestock industry, early death,...). So her quality of life equals 3. According to quasi-maximin prioritarianism, the first situation is preferred. If you’d choose the second situation, from behind the veil of ignorance, you have probability $\frac{1}{2}$ to end up in the worst-off position. (According to sum-utilitarianism, the second situa-

tion is better.) Tolerated partiality is also violated: if we prefer the enjoyment of cheese above the use of the cow, we should also tolerate the other option: breeding women and using their breast milk to make cheese for cows (suppose the cow likes human cheese). This we would not tolerate. The third principle is also violated, because the cow in the livestock industry is used as merely means (her bodily integrity is violated and she is treated as property).

With these three principles, we arrive at a coherent system that best fits our strongest moral intuitions. Some intuitions based on speciesist judgments are not compatible with this system of animal equality. These intuitions are too weak and cannot be incorporated without introducing highly arbitrary and artificial constructions, so we have to dismiss these speciesist intuitions as being moral illusions. Although our theory implies veganism, it still allows for some partiality (the tolerated partiality meets our intuitive preference for some individuals). However, there is one serious problem remaining.

Fact 7: Predators need meat in order to survive. If predators cannot use other sentient beings as merely means, they will all become extinct. If principles 4, 5, and 6 were universalized to predator animals, this would imply that they have to become extinct.

Moral intuition 7: Predators are allowed to hunt and hence violate the basic rights and well-being of prey. It would be a tragedy if they became extinct. It is not easy to formulate a clear principle that is coherent with this intuition as well as with the intuitions that we encountered before. If we suppose that biodiversity has a moral value, then we have the following option.

Universal ethical principle 7: If a sufficiently large group of sentient beings became by an evolutionary process dependent on the use of other sentient beings for their survival, they are allowed to use other sentient beings for that purpose (until feasible alternatives, that don't violate basic rights, are found).

If we suppose that biodiversity has moral (intrinsic) value, and if we define biodiversity as the diversity of everything that is the direct product of evolutionary processes, then this seventh principle becomes coherent with the value of biodiversity. So predators and their behavior contribute to biodiversity and we should not destroy that biodiversity.

This principle is also coherent with a “triple-N-principle,” which refers to the three values “natural, normal and necessary” of a carnist ideology (Joy, 2009). This connection works if we define natural as: behavior that is a direct consequence of a process of evolution (genetic mutation and natural selection). So it refers to an “evolutionary process.” Normal means that the behavior happens a lot, so it refers to a “sufficiently large group.” And necessary means that those beings would die if they no longer exhibit that behavior. This refers to “dependency for survival.”

Putting the three criteria together, natural+normal+necessary, means that a lot of biodiversity would be lost when the behavior stopped. And a lot of biodiversity has a lot of moral value, sufficiently enough to trump the basic right. Predation is normal, natural, and necessary, so it is allowed (as long as there are no feasible alternatives), even if it violates the basic right. For humans, eating animal products is not necessary (according to the Academy of Nutrition and Dietetics), so we are not allowed to violate the basic rights of animals. Organ transplantation (by sacrificing a sentient being against his will) is not allowed ei-

ther, because it is a violation of the basic right and it is not normal and natural (it is necessary though).

Note that this value-of-biodiversity principle is completely unrelated to the value-of-sentience principles discussed before, although we could compare biodiversity as an intrinsically valuable property of ecosystems, with sentience as an intrinsically valuable property of sentient beings. In itself, the biodiversity principle seems arbitrary, but it is coherent with lots of our intuitions. For example: moving around and killing insects (by accident) is considered allowed, even if scientists are able to demonstrate that insects are sentient. But the 3-N-principle says that moving around is natural, normal, and necessary behavior of animals.

Finally, we also have situations where predators attack us or beings that we hold dear. Our intuition say we are allowed to defend ourselves and others, and we have a stronger duty to protect some individuals with whom we have special relationships. All sentient beings have the right to defend themselves or others, they have the right to be partial in such decisions, as long as they respect similar kinds of partiality of others (see principle 5) and as long as biodiversity is not threatened. If we wish, we could also add that we have a duty to protect all beings that have (or will develop) moral agency or rationality. Those rational beings not only feel their interests, but they also know and understand their interests. This rationality applies to most human beings, except, for example, seriously mentally disabled human orphans. This satisfies people's intuitions that we have a duty to protect humans from predators. (But if we say that we have a duty to protect mentally disabled humans whereas we do not have a duty to protect non-human animals, because all humans have a higher moral status than non-humans, then

we become too partial. This kind of speciesism, like racism or sexism, is a kind of partiality and arbitrariness that we cannot tolerate.)

This completes the process. We now have a theory of animal equality, with clear and coherent universalized ethical principles that best fit our strongest moral intuitions, and without too many arbitrary elements.

References

- Academy of Nutrition and Dietetics - American Dietetic Association. 2003. "Position of the American Dietetic Association and Dietitians of Canada: Vegetarian Diets." *Journal of the American Dietetic Association* 103 (6).
- Davis, S. 2003. "Least Harm." *Journal of Agricultural and Environmental Ethics* 16 (4).
- Francione G. 2000. *Introduction to Animal Rights: Your Child or the Dog?* Philadelphia: Temple University Press.
- Joy, M. 2009. *Why We Love Dogs, Eat Pigs and Wear Cows: An Introduction to Carnism*. Conari Press.
- Lamey, A. 2007. "Food Fight! Davis Versus Regan on the Ethics of Eating Beef." *Journal of Social Philosophy* 38 (2).
- Matheny, G. 2003. "Least Harm: A Defense of Vegetarianism from Steven Davis's Omnivorous Proposal." *Journal of Agricultural and Environmental Ethics* 16: 505–511.
- Parfit, D. 1984. *Reasons and Persons*. Oxford: Clarendon Press.

Rawls, J. 1971. *A Theory of Justice*. Cambridge: Harvard University Press.

Regan T. 1983. *The Case for Animal Rights*. Berkeley: University of California Press.

Rowlands, M. 1998. *Animal Rights: A Philosophical Defence*. Macmillan/St Martin's Press.