

Rethinking Rationality

To make its research readily available to a broad audience, the Institute for Philosophy and Public Policy publishes a quarterly newsletter: *QQ—Report from the Institute for Philosophy and Public Policy*. Named after the abbreviation for “questions,” *QQ* summarizes and supplements Institute books and working papers and features other selected work on public policy questions. Articles in *QQ* are intended to advance philosophically informed debate on current policy choices; the views presented are not necessarily those of the Institute or its sponsors.

* * *

In this issue:

Recent research on the scope and limits of rationality triggers worries about the wisdom of our public policies governing risk p. 1

Does the scientific basis for human genetic engineering justify faith that it will not create more problems than it can solve? p. 6

Does nuclear deterrence work? For an answer we need to take a closer look at the nature of deterrence itself p. 9

How good a person do I have to be? (How good a person do I *want* to be?) p. 12

The first volume in a new book series is announced p. 15

Just when you were getting used to the idea that alfalfa sprouts cause cancer and exercise causes infertility comes word about the possible threat posed by radon in your home. Every night of late local news stations in Washington, D.C., advertise their evening report by a teaser promising more information on the threat posed by this invisible form of indoor air pollution. The estimated chance of death by radon, even on the most pessimistic accounts, is far smaller than the chance of dying in a car accident; yet few news broadcasts try to attract viewers by headlines promising defensive driving tips. Everyone is afraid of getting AIDS; far fewer dread diabetes, a much more prolific killer. One survey shows that over 90 percent of us think we are better than average drivers.

It is well known that people's worries about various risks correlate poorly with the actual dangers they pose. Our judgments are seriously flawed, and it seems that, at the very least, our attitudes about risk are often inconsistent, if not perverse.

Of course, those charged with inconsistency in such matters are free to respond, “So I’m inconsistent. Big deal!” Many would join with Emerson in holding that “consistency is the hobgoblin of little minds,” or with William Allen White, that “consistency is a paste jewel that only cheap men cherish.” We may have our own reasons to pick and choose our own fears, whether or not our choices line up with somebody else’s dispassionate assessment of what is truly fearsome.

Yet consistency in some form defines what it is to be rational, and we are less complacent about the charge

of irrationality. Certain principles of consistency in our choices—so obvious, on first reflection, that they hardly need stating—have been taken to be constitutive of what rationality is. If risk is what you care about, and automobiles are riskier than radon in the home, isn't it irrational not to adjust your concerns about the two accordingly?

People talk about risk, and experts talk about risk, but the two groups may be talking past each other, for "risk" often seems just a convenient label to slap on a broad family of other concerns.

If it turns out that people are indeed irrational in their attitudes toward risk, this has troubling implications for many of our public policies governing risk. In a democracy public attitudes are the cornerstone for policies; they are embodied in the laws we pass and reflected in regulations that give those laws substance. If public opinion on regulatory issues is shaped by frivolous or confused considerations, this bodes ill for our prospects of establishing rational public policies.

Perhaps, however, apparent inconsistencies in our attitudes about risks can be given some other explanation. Perhaps people focus concern on certain risks because they are responding to factors that researchers are simply failing to measure. Or perhaps the standard model of rationality cannot accommodate the complexities of human reasoning. How rational or irrational are we? And how much does this matter?

Who Cares About Risk Anyway?

A first, sympathetic explanation of why people worry disproportionately about certain risks—that is to say, out of proportion to the riskiness of the risk—is that the category of risk hardly exhausts the full range of what most of us care about in our daily lives. People talk about risk, and experts talk about risk, but the two groups may be talking past each other, for "risk" often seems just a convenient label to slap on a broad family of other concerns.

Risk analysts compute the chance of dying from a given activity, but most people care not only about their chance of dying but about what life is like while they are living it. Trade-offs between quality and quantity of life are made all the time in personal decisions about health and safety: Mark Twain, told that he could add five years to his life by giving up smoking and drinking, reportedly quipped that five years without smoking and drinking weren't worth living. Similar trade-offs are relevant in the policy arena as well.

Thus people may prefer one technology to another for reasons other than the actual risks to life and health associated with it. They tend to care about the form

of social organization it encourages (solar power lends itself to decentralization, while nuclear power is by its nature highly centralized); about the control they feel they have over its risks (one, perhaps specious, reason for worrying less about driving than flying); about whether the risks are assumed voluntarily or imposed by others, whether these occur now or later, affect many or few, and so forth.

These factors help to explain some of the results the experts find so puzzling. If risk analysts are preoccupied exclusively with risk, while our concerns are more diverse and wide ranging, it is not surprising that their research should deliver a verdict of irrationality. But here the fault lies not in how human beings think about risk, but in an overly narrow and restrictive focus of measurement.

Other findings are less easily explained, however, by pointing to a richness in our values overlooked by risk analysts. They suggest that people are often driven by indefensible cognitive processes; at least some of the time we simply process information in crazy ways.

Failing Grades in Probability

Measurements of risk have two components: an assessment of the probability of some outcome and an assessment of whether that outcome would be good or bad, and how good or bad it would be. Thus opportunities arise for people to make two different kinds of mistakes.

First, we can make erroneous judgments about probabilities. Many of the laws of probability have a counter-intuitive flavor, and temptations to fallacy are common. Most of us have to struggle to resist the so-called gambler's fallacy, the wistful belief that after enough successive losses the odds have to favor a win. Only reluctantly do we abandon our gut feeling that after five coins in a row have come up heads, surely the next coin will come up tails, despite the mathematician's insistence that the odds on any fair coin toss remain fifty/fifty.

Most of us, who can hardly balance our checkbooks without a pocket calculator, wisely eschew complicated calculations of probabilities in favor of simple rules of thumb, or heuristics. . . .

The gambler's fallacy is an obvious and familiar example of a cognitive failure. More interesting is research on how poorly people do at probabilistic reasoning even when they appear to be doing it quite well.

Most of us, who can hardly balance our checkbooks without a pocket calculator, wisely eschew complicated calculations of probabilities in favor of simple rules of

thumb, or heuristics, that help us assess probabilities in a rough and ready way. The pioneering psychologists Daniel Kahneman and Amos Tversky have studied many of these heuristics. One is "salience," or the ease with which we can call examples of a certain kind of occurrence to mind. Salience will be a generally reliable guide to likely patterns in the world around us: we can think of fewer redheads than brunettes in our acquaintance because brunettes indeed outnumber redheads in the population. But such heuristics can lead us to make mistakes, as when the salience of some risk is reinforced by undue media attention (another possible reason why airplane crashes are more feared than automobile accidents). On balance, however, the heuristics Kahneman and Tversky identify, unlike the gambler's fallacy, are promising strategies for coping efficiently with the uncertainties we face. It is not surprising that even the best heuristics will not always give the same results one would get if one took the trouble to figure in the specifics of the case, but it is surely reasonable to sacrifice some accuracy for convenience.

What is surprising, however, is the extent to which these heuristics dominate even obviously relevant countervailing information. Tversky and Kahneman have shown that when heuristics are triggered, people let themselves ignore completely information that may be far more important but less salient. In one study, Tversky and Kahneman asked people to judge whether an individual selected from some population was more likely to be an engineer or a lawyer, where the population consisted of 70 percent engineers and 30 percent lawyers. Such "base rate data" about the background population is crucial to assessing probabilities correctly. They found that in the absence of any description of the individual's characteristics and traits, people's judgments were based on what they knew about the percentage of lawyers and engineers in the population, but "prior probabilities were effectively ignored when a description was introduced, even when this description was totally uninformative." When asked to decide

if some undescribed Tom was more likely to be a lawyer or an engineer, subjects guessed he was 70 percent likely to be an engineer, but when they were asked about Dick, who was described simply as a "30 year old man, married with no children," these subjects opted for a fifty/fifty probability.

The conclusion to draw from such research seems fairly straightforward. We not only rely on heuristics but are dominated by them, even when we shouldn't be. When it comes to judging probabilities people are, with some frequency, simply wrong.

Knowing What We Value

A more controversial and troubling realm of error lies in the way that people value outcomes. We may expect people to be poor at assessing probabilities, but we want to give them credit for at least knowing their own minds when it comes to assigning values to the outcomes of their choices. They can confidently judge which of two alternatives is the more attractive.

Or can they? Another striking finding of Kahneman and Tversky is that people's value judgments are notoriously influenced by the way in which various choices are framed. Every retailer knows that customers consider \$99 a far more attractive price than one a mere dollar higher, while the difference between, say, \$97 and \$98 makes no difference at all. Similarly, customers tend to be delighted by a discount for cash payment but bristle at a surcharge for using credit cards—although identical policies can usually be described either way.

What Tversky and Kahneman have done is to describe the systematic nature of how preferences are affected by the way a problem is framed. They have scientifically grounded the suspicion that people prefer the "half full" cup over the "half empty" cup, showing the extent to which framing does indeed affect judgment, even on some important policy issues.

Kahneman and Tversky ask us to suppose that the United States is preparing for the outbreak of an

A cab was involved in a hit-and-run accident. Two cab companies serve the city: the Green, which operates 85 percent of the cabs, and the Blue, which operates the remaining 15 percent. A witness identifies the hit-and-run cab as Blue. Tests show that under circumstances similar to those on the night of the accident, witnesses correctly distinguish Blue cabs from Green cabs 80 percent of the time and misidentify them the other 20 percent. What's the probability that the cab involved in the accident was Blue, as the witness stated?

Most subjects conclude that it is 80 percent likely that the cab will be Blue. If there were 85 Green cabs and 15 Blue ones in the city, however, a witness with an 80 percent accuracy rate would incorrectly identify 17 Green cabs as Blue (20 percent of 85 = 17), and he would correctly identify 12 of the Blue cabs (80 percent of 15 = 12). He would thus identify 29 cabs as Blue and be correct in only 12 cases, an error rate of almost 60 percent. The base rate—the preponderance of Green—makes the odds 60 to 40 that he has misidentified a Green cab rather than correctly identified a Blue one.

(Example from Tversky and Kahneman)

Linda is 31, outspoken, and very bright. She majored in philosophy in college. As a student, she was deeply concerned with discrimination and other social issues and participated in anti-nuclear demonstrations. Which statement is more likely:

- a. Linda is a bank teller.**
- b. Linda is a bank teller and active in the feminist movement.**

87 percent of respondents ranked (b) as more likely than (a), but it is a law of probability that a conjunction cannot be more likely than one of its constituents.

(Example from Tversky and Kahneman)

unusual Asian disease, which is expected to kill 600 people, unless action is taken. Two alternative programs to combat the disease have been proposed. If program A is adopted, 200 people will be saved. If program B is adopted, there is a $\frac{1}{3}$ probability that 600 people will be saved and a $\frac{2}{3}$ probability that no people will be saved. When the alternatives were posed in these terms in a test survey, 72 percent of respondents opted for Program A, only 28 percent for program B. A second group was given the same options, but described in this way: If program A is adopted, 400 people will die; if program B is adopted, there is a $\frac{1}{3}$ probability that nobody will die, and a $\frac{2}{3}$ probability that 600 people will die. This time only 22 percent opted for the first program, while 78 percent opted for the second—a clear preference reversal. The framing of the question proved decisive in eliciting a response. When program A was seen as involving a gain of 200 lives it was rated far more favorably than when it was seen as involving a loss of 400.

This research, moreover, reveals not only the extent to which framing affects what is perceived as a loss or a gain, but a deep and abiding asymmetry in how losses and gains themselves are regarded. In study after study, people show themselves more concerned to avoid a loss than to receive an equivalent gain. In one experiment, for example, researchers Jack Knetsch and J.A. Sinden gave half their subjects tickets to a lottery and the other half \$3.00. When the first group was given an opportunity to sell their tickets for \$3.00, 82 percent kept them, but when the second group was allowed to buy lottery tickets for their \$3.00, only 38 percent wanted the tickets. Human beings seem strongly disposed, for whatever reason, to defend their own personal status quo, to hang on to what they've got. Tversky and Kahneman call this tendency "loss aversion."

Framing effects and loss aversion together explain an important class of preference reversals. People value gains and losses differently, and the framing of a problem determines a "reference point" from which outcomes are viewed as gains or losses. In the Asian

disease example, the description of the problem determines whether people see the alternatives as lives saved or lives lost, which in turn determines preferences among the alternatives.

Is loss aversion irrational? On one popular model of economic rationality it is. According to this conception, what matters is the bottom line, where you end up, whether you got there by avoiding a loss or forgoing a gain. But for others, the choice process itself may legitimately matter to us, as well as its ultimate outcome. Douglas MacLean, director of the Institute for Philosophy and Public Policy's project on risk and rationality, argues that "we have a common set of reactive attitudes, like regret and reproach, which are often provoked by the decisions or choices we make and not just by the choiceless outcomes that result from our decisions. The very same outcome might provoke delight in one context—some small gain you were lucky to realize—but regret in another, for example if you could have done far better by choosing differently." It is not always reasonable to be affected by such attitudes, of course, but given their persistence and pervasiveness, we may do well to factor them into our decisions. For example, MacLean suggests, "if your acquaintances, your boss, or your constituents are more likely to hold you responsible for the losses that result from your risky decisions than for the lost opportunities that result from choosing to avoid risks, then choosing to avoid a loss rather than maximize the expected outcome would seem to be an eminently prudent strategy." An adequate conception of rational choice may have to assess the value of decision processes themselves and our reasonable reactive attitudes and aversions. This would involve a more complex understanding of alternatives than one that exclusively focuses on outcomes in the measurement of risk.

We might accept loss aversion, but what about our

Imagine you have operable lung cancer and must choose between two treatments: surgery and radiation therapy. Of 100 people having surgery, 10 die during the operation, 32 are dead after one year, and 66 after five years. Of 100 people having radiation therapy, none die during treatment, 23 are dead after one year, and 78 after five years. Which treatment would you prefer?

When this question was posed to a group of physicians, framed in terms of mortality rates, 50 percent of respondents favored radiation therapy. But when the same information was presented in terms of survival rates, radiation was preferred by only 16 percent. Surgery is somewhat riskier at the outset, but has better odds of long-term survival.

(Example from Kahneman and Tversky)

susceptibility to framing effects? Surely preference reversals in cases like the Asian disease example are irrational. After all, the two programs of disease control remain the same, however described. Nevertheless, the issue of framing effects might be more complicated. MacLean acknowledges that "if two different descriptions describe the same choices, rationality requires that we make the same decision in each case. Stated in this way, the principle seems hardly exceptionable." But "the problem is that there may be no clear way to determine what counts as a redescription of the same prospect, rather than descriptions of different prospects." Is the situation of someone who gained some money and lost it gambling the same or different as that of someone who never had any money to begin with? While framing can determine the reference point from which different possible outcomes are viewed as gains or losses, there may be no general way to determine what the correct reference point should be. One's current position is not always an appropriate benchmark. A raise in salary, for example, might trigger delight by comparison to one's actual salary, but disappointment compared to one's legitimate expectations.

Much of the moral fabric of our lives depends on finding the appropriate descriptions of objects and events that, from some detached perspective, could equally well be described in some other way. A person's arm can be seen as a human limb or as meat and bones. Assessments of guilt and responsibility, in law and in morality, often turn on determining what description of an event is true or most appropriate. And our culture is currently divided most deeply on whether a human fetus is a person or merely a human organism.

Finally, some recent research in this area suggests further complexities, that people's value judgments apparently depend not only on how the outcomes of a choice are framed, but also on the framing of the choice itself, the process by which the judgments of gain or loss themselves are elicited. Paul Slovic, of Decision Research, Inc., has shown that we will get a different picture of the value people place on, say, tea or coffee, according to whether we ask them whether they *prefer* tea to coffee, whether they *judge* tea to be *better* than coffee, or whether they would be *willing to pay* more for tea than for coffee. This research also reveals some remarkable preference reversals. Some of Slovic's examples focus on a choice between two bets with roughly equal expected payoffs: one has a higher probability of winning, while the other involves a smaller chance of a larger gain. People prefer to take the higher probability bet, but they are willing to pay more for the chance to take the other. Again, according to one popular conception of rationality, values, preferences, and choices should reflect the same consistent ordering, so these reversals would be irrational. But some of Slovic's results might be taken instead to show how values can be expressed in importantly different ways. We might value some things deeply, for example, but think it inappropriate or wrong to pay for them at all.

What Should We Conclude?

The conclusions from this research are mixed. We find some striking examples of human fallibility and some clear limits to human rationality. But other cases show plain folks exhibiting reasoning that may not square with expert assessments of risk but nonetheless does not seem obviously faulty. On the one hand, people do poorly at judging probabilities and are persistently led astray by framing effects. On the other hand, their reasoning often shows both good sense and a sensitive appreciation of the complex nature of difficult choices.

In any case, it is doubtful that people will ever change to become more fully rational. Tversky and Kahneman's research findings, confirmed in countless studies over many years, show that people's patterns of choice are amazingly tenacious and persist at all levels of education and technical sophistication. In a country where most adults would be hard pressed to find a least common denominator between two fractions, education in the laws of probability and in the subtleties of risk analysis is going to produce limited results. We are going to have to deal with people as they are.

On the one hand, people do poorly at judging probabilities and are persistently led astray by framing effects. On the other hand, their reasoning often shows both good sense and a sensitive appreciation of the complex nature of difficult choices.

Nor does increased reliance on expert risk analysis seem to provide a better foundation for public policy. Our irrationality shows the need for expert guidance in risk management, but the complexities in our values also suggest that the expert's analytic techniques for revealing public values may be flawed. Nobody ever said that democracy, with its reliance on the popular will, would produce the best results all of the time, but only results that were preferable on balance to those generated by alternative arrangements. We seem to be more rational than animals, less rational than angels, or computers. In other words, human.

The sources quoted in this article are: Amos Tversky and Daniel Kahneman, "The Framing of Decisions and the Psychology of Choice," *Science*, vol. 211 (1981); Tversky and Kahneman, "Rational Choice and the Framing of Decisions," *Journal of Business*, vol. 59 (1986); Jack L. Knetsch and J.A. Sinden, "Willingness to Pay and Compensation Demanded: Experimental Evidence of an Unexpected Disparity in Measures of Value," *The Quarterly Journal of Economics* (August 1984); Douglas MacLean, "Risk, Regret, and Rationality," Institute for Philosophy and Public Policy Working Paper RR-5 (forthcoming), and unpublished manuscripts; Paul Slovic, "Preference Reversals," Institute for Philosophy and Public Policy Working Paper RR-2 (March 1987).

Does Human Genetic Engineering Have a Scientific Basis?

New technologies have traditionally started out on a small scale; if they are based on sound scientific principles they can take hold and flourish, and eventually provide the basis for large-scale enterprises. The Anasazi Indians of the American Southwest would probably not have ventured to build and occupy hundred-room dwellings on the faces of mile-high cliffs had they not first tried out their construction techniques on small buildings closer to firm ground.

In certain cases, however, particularly during this century, technologies have been instituted on a large scale almost from their inception. Nuclear power plants by their nature service huge consumer grids; the proposed Strategic Defense Initiative must shield the entire country if it is to work at all. In a rational world, the more people put at risk by an enterprise, the firmer should be its scientific foundation. If the theories by which one predicts the consequences of a large-scale technology are perceived to be flawed or seriously incomplete, the desirability of implementing the technology altogether may be questioned.

The use of molecular biology to modify human characteristics, so-called human genetic engineering, also represents a technological enterprise of potentially enormous scale. Genetic manipulations confined to the non-reproductive cells of individual persons ("somatic" gene modification), if found to be effective for rare life-threatening disorders, would almost certainly come to be advocated as a therapeutic procedure for more common health-threatening conditions, such as obesity, diabetes, hypertension, and depression. And genetic modification of the reproductive cells, or early embryo, for the purpose of prospectively correcting inherited defects or predispositions to disease ("germline" genetic engineering), while not on the short-term medical agenda, is considered by many to be a reasonable prospect, particularly because of recent dramatic results along these lines in animal experiments. The widely touted idea that organisms are "programmed by their genes," and that scientists are now learning how to rewrite the programs, fosters confidence in the inevitable progress of this approach.

Research and development, particularly when accompanied by promises of important health benefits, can create large scientific, medical, and commercial infrastructures, and corresponding ambitions to apply the fruits of this labor as soon as possible. Therefore, it is essential, even at this early stage, to examine in

greater detail the idea that organisms are programmed by their genes and can therefore be reprogrammed by genetic manipulation. Does the purported scientific basis for this technology justify faith that human genetic engineering will not create more problems than it can solve?

Are There Genetic Programs?

A typical expression of the theoretical perspective behind the genetic engineering enterprise can be found in a 1978 article by the molecular biologists Alexander Rich and S.H. Kim: "It is now widely known that the instructions for the assembly and organization of a living system are embodied in the DNA molecules contained within the living cell." Taken together with our remarkable technical capacity to cut, splice, rearrange, synthesize, and decipher the sequence of DNA, this would appear to provide the required prescription for the makeover of all sorts of organisms, including our descendants.

Most biologists and knowledgeable lay people would agree with the quoted statement; equivalent declarations can be found in virtually every textbook, magazine article, and television program on modern biology to appear in the past twenty-five years. Nonetheless, it is reasonable to ask if this viewpoint really has any content at all. For scientists it may simply represent a shorthand form for a patchwork of results in cellular and molecular biology, none of which really add up to any formal set of "instructions for the assembly and organization of a living system." However, the view in question may actually be taken to mean that DNA "programs" an organism's activities much the way a floppy disk programs the tasks of a microcomputer. As incisive an analyst as the physicist Freeman Dyson has somehow concluded from his survey of molecular biology that "hardware processes information; software embodies information. These two components have their exact analogues in the living cell; protein is hardware and nucleic acid is software." Thus, there seems to be a major and potentially dangerous gap between the state of molecular biology as a science and the perception of it by analysts, policymakers, and the public being asked to bear the costs and risks of its application.

Considering the wealth of data available on the cellular role of DNA, it is difficult to understand why the facile metaphor that likens this molecule to a com-

puter program has persisted for so long. In effect, DNA constitutes a list of ingredients, not a recipe for their interactions. Whereas all somatic cells in a given individual contain the same genes, these are not equally expressed in each differentiated cell type. Different cell types in the pancreas, for instance, produce insulin and digestive enzymes. But the genes which specify these proteins are common to all pancreatic cells, as well as to the numerous other cell types of the body, which make no pancreas-specific products. In a formal sense, the "program" for expressing genes resides among the various components of the cell and its environment. Biological information is therefore not uniquely located in any one structure or molecule.

Numerous studies have been conducted in which genes are mutated, moved, or increased in number, and the consequences for whole organisms or their cells are recorded. These studies show that small genetic changes can have effects on a cell or organ that may be extensive or negligible, depending on the particular protein primarily altered, and on the interactions that protein has with other cellular components. But simply because a small physical or chemical change in a complex system results in changes throughout that system does not mean it does so by issuing instructions. Such a line of thought, followed consistently, is equivalent to maintaining that a defective gasket "programmed" the space shuttle Challenger to explode.

If a cell's DNA indeed embodied a computer-type program, it should be feasible to decipher the language in which the program is written. I refer here not to the genetic code: the correspondence between nucleotide triplets and amino acids is uncontested; it is the grammar of "the instructions for the assembly and organization of a living organ" that we seek. If a cellular programming language existed it would be expected to have signifiers whose meanings transcend the immediate context. The word "house," or the nucleotide triplet GAT, have meanings that are constant in the appropriate systems, provided that they are used according to certain grammatical rules. However, no such rules exist in the cell that permit an assignment of the "programmatic" meaning of a change from one nucleotide to another, for example, or the insertion of a perfectly normal gene into a new place in the chromosome. The visible outcome of such changes will be different in different situations.

It is true that we can presume a unique correspondence between the genetic makeup, or genotype, and the physical properties, or phenotype, of an organism under a specific environment. But the "environment" of the genome includes not only externally controllable factors like temperature and nutrition, but also numerous maternally provided proteins present in the egg cell at the time of fertilization. These proteins influence gene activity, and variations in their amounts and spatial distribution in the egg can cause embryos even of genetically identical twins to develop in uniquely different ways. Thus, the

same genotype, under different conditions, can develop into a host of different phenotypes. Our genetic endowment need not be our destiny.

Is Genetic Engineering Possible?

Even if the assumption of unique correspondence between genotype and phenotype for a given environment were granted, this is not the sort of programming that can be put to effective use by genetic engineers. Unless a normal gene can be inserted in an exactly normal location in the cells targeted for gene therapy (a procedure that would necessitate precisely removing the abnormal genes without otherwise damaging the cells), an abnormal genotype will be constructed, whose corresponding phenotype will be unpredictable.

The chromosomes of each cell type are generated in sequential steps that occur during embryonic development, including chemical modification of the DNA itself. The result of these steps is that some genes wind up in chromosomal regions that allow their information to be used, while other genes are packed away into non-functional chromosomal regions. Different genes are allocated to "active" and "inactive" chromosomal domains in different cell types. It is consequently a far from straightforward matter to insert foreign DNA into precise locations within these structures. Furthermore, the specific removal of defective genes is impossible with current or foreseeable technologies.

In the case of somatic gene modification these difficulties would be reflected in failure of the procedure if the normal protein was either not produced or produced in amounts too small to correct the disorder. If the protein was overproduced it could be physiologi-



"We hope to make antibiotics, interferon and diagnostic products, but just to be on the safe side, we're starting out with a line of shampoo."

cally disruptive. In the worst case the genetically modified cells could become cancerous and eventually kill the patient. This is not to say that expression of a desired gene cannot occasionally be achieved in target cells, and that implantation of such cells into a patient's body could not ameliorate the symptoms of a gene-related disability. Such therapy may offer hope in some rare, desperate cases. But scientific principles that

Gene therapy is the prestige field of current biomedical research, but the vague promise that it may someday be the treatment of choice for many health disorders is a bad guide for public health policy.

would allow one to predict the long-term behavior of bodily tissues that express foreign genes simply do not exist.

The introduction of foreign DNA into the egg or sperm prior to fertilization, or into early embryos, presents an additional set of problems. These procedures are collectively referred to as "germ-line genetic engineering" because they have the probable result that the altered genes will become incorporated into the embryo's own reproductive cell precursors, and thence conveyed to subsequent generations. The success that investigators have achieved in introducing functional genes into the eggs and embryos of experimental animals might appear to be a result of a deep and thorough understanding of the process of gene expression during early development. Quite the opposite is true.

Early embryos have long been known to have the capability of enduring major traumas and insults and still developing into normal-looking organisms. Embryos of sea urchins, frogs, and even mammals can be experimentally dissociated into their constituent cells; if this is done at an early enough stage, each cell gives rise to a fully formed individual. If two unrelated early mouse embryos are jumbled together, the constituent cells will readjust their fates to yield a single individual with four parents. Phenomena of this sort are more a tribute to biological prodigiousness than to human ingenuity.

Therefore it was interesting, but not unprecedented, to learn that foreign genes injected into fertilized mouse eggs could become expressed in the resulting embryo, which can then develop into a recognizable animal. Because of the extraordinary homeostatic capacities of the embryo, physically normal, or even "improved," development can occur in embryos that have been rendered genetically abnormal by these procedures. In such cases, however, the cryptic genetic defect could show up in subsequent generations. A recent study,

for example, found that normal-looking offspring of one genetically modified mouse developed cancer by the middle of their lives at more than 40 times the rate of the unmodified strain.

The Human Species as an Experimental System

This scientific critique of genetic engineering is not in any way intended to shift the focus of the debate to mere technical issues. Even were the technology entirely capable of delivering on its promises to prospectively correct genetic disorders with no additional risk to health, there would still be numerous extra-scientific reasons to be wary of attempts to modify humans and even nonhuman species to reflect current tastes and needs. And even were the program of retrospectively curing sick persons using somatic methods accepted as being a natural extension of current drug therapies, debate would still occur on whether such treatments, which cannot be withdrawn in the face of adverse reaction, actually represent an ethical medical option in situations that are not life-threatening.

On the other hand, if the popular conception of the state of a science like molecular biology is based on false analogies and expectations, an informed discussion of the ethics and desirabilities of new technologies based on that science cannot occur. Gene therapy is the prestige field of current biomedical research, but the vague promise that it may someday be the treatment of choice for many health disorders is a bad guide for public health policy. Just because brain surgery and the transmutation of the elements eventually found a scientific basis does not mean that Aztec shamans and medieval alchemists were on the right track.

Success in germ-line modification in humans will likely be judged by a favorable outcome in the original groups of embryos treated. However, a scientifically

Just because brain surgery and the transmutation of the elements eventually found a scientific basis does not mean that Aztec shamans and medieval alchemists were on the right track.

rigorous assessment of the impact of genetic manipulation would actually require a program of controlled breeding over several generations, an option that would be unacceptable with regard to humans. Failing this, errors introduced into the human gene pool could affect countless persons in the future, and their spread throughout the population curtailed only at great social cost.

In the United States, statutory prohibition on human embryo experimentation extends only to the provision

of financial support by the Department of Health and Human Services for in vitro fertilization studies. The Recombinant DNA Advisory Committee of the National Institutes of Health unanimously voted down a proposal to ban human germ-line modification. If the public in this country continues to accept a laissez-faire attitude toward human germ-line gene modification, in the space of a few years we may be confronted with human experimental products disavowed by their parents, partial successes whose characteristics are only approximately human, and unfortunate lineages susceptible to early cancers and exotic new genetic disorders. The timetable for embarking on this dubious program will depend, then, not on the state of scien-

tific understanding, but rather on how widely certain false notions of genetic perfectability come to be accepted, and how rapidly narrow commercial and professional interests can succeed in exploiting the resulting demands.

—Stuart A. Newman

Stuart Newman is a professor at New York Medical College and a member of the executive council of the Committee for Responsible Genetics, a Boston-based public interest group. This article is substantially condensed and adapted from his paper "Human Genetic Engineering: Who Needs It, and Does It Have a Scientific Basis?", prepared for the Institute for Philosophy and Public Policy's Working Group on Moral and Social Issues in Biotechnology.

Does Nuclear Deterrence Work?

Does nuclear deterrence work? This is the central question in any examination of the fundamentals of nuclear weapons policy. If nuclear deterrence works well, there is strong prudential reason, in terms of national self-interest, to retain it, and if it does not, there is strong prudential reason to abandon it. But this question is central for a moral appraisal of nuclear deterrence as well, on any moral view that takes the consequences of our actions and policies seriously.

The question of whether nuclear deterrence works well can be taken in at least two different ways. The broader question is about the absolute deterrent value of nuclear threats, that is, their effectiveness in comparison with a policy of no military threats at all; the narrower question is about the marginal deterrent value of nuclear threats, that is, their effectiveness in comparison with nonnuclear military threats. The latter question is the one most of us care about in asking whether nuclear deterrence works.

How should this question be answered? Most commonly, an empirical answer is attempted, as in the argument that nuclear deterrence must work well because there has not been a war between the United States and the Soviet Union since 1945. But such an argument exhibits the error made by the person who is sure that the amulet she wears keeps elephants away because she has not seen any elephants since beginning to wear it. No effort is made to show any connection beyond mere concurrence between nuclear deterrence policy and the absence of war.

A comparison between nuclear deterrence and conventional military deterrence might better be made not empirically, but by a closer look at the nature of deterrence itself. After all, deterrence is a pervasive relation among persons and institutions at all levels of

social groupings, from the family to the nation to the world order. If we try to analyze what conditions are necessary for deterrence threats in general to be effective, and then examine how these are satisfied in a paradigmatic case of effective deterrence, such as legal deterrence, this can help us understand how the comparison between general military deterrence and nuclear deterrence should be made.

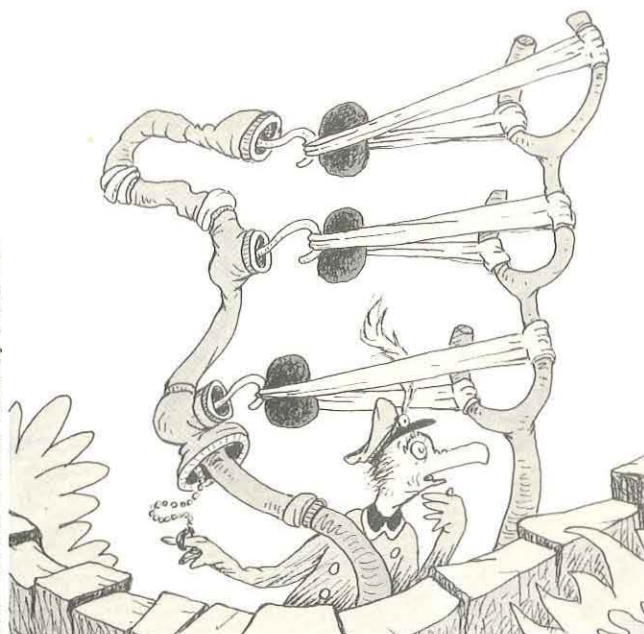
Legal Deterrence

Effective deterrent threats satisfy two conditions. First, they create in the minds of the parties threatened the belief that should they fail to conform their behavior to the required standards it is likely that the threatened harm will be imposed upon them; that is to say, that the threatening party has the capability and the willingness to carry out the threats. Second, these beliefs result in the threatened parties conforming their behavior to the required standards.

How well do legal deterrent threats satisfy these conditions? For legal deterrent threats, the general ability of the state to carry out its threats is not in doubt, given the state's effective monopoly on the use of force. But there is some room for doubt about the state's ability to impose the threatened harm for particular cases of nonconforming behavior, since it is often difficult to detect who is responsible for such behavior. Thus, the grounds for the belief that the state is able to carry out its legal deterrent threats are not completely firm; particular potential violators may have some reason to doubt the state's ability to punish them.

What about the basis for the belief that the state is willing to carry out its legal threats? Mere declarations of an intention to carry out a threat are normally not a sufficient basis for a belief that the threat will be car-

from *The Butter Battle Book*, by Dr. Seuss



© 1984 by Dr. Seuss and A. S. Geisel

ried out. Thus the basis for the belief that the state is willing to carry out legal threats is to be found in the state's past behavior of legal threat executions. To create the belief that it is willing to carry out legal threats, the state must have a history of having done so. To put the point paradoxically, the general success of legal deterrence is dependent on its occasional failure. Legal deterrence does not merely tolerate failures, maintaining its overall success despite them, but it actually makes use of them, and even requires them, for its overall success.

Now, if the belief that the threatener is able and willing to carry out its threat is to lead the threatened parties to conform their behavior to the required standard, they must be assured that if they do so conform, the harm threatened for nonconformity will not be inflicted upon them; otherwise they would have no incentive to conform in order to avoid this harm. Moreover, the harm threatened must be sufficiently severe in relation to the expected gain from nonconforming behavior: nonconforming behavior must not pay. Our system of legal deterrence generally satisfies these conditions.

Finally, for a threat to work, the party threatened must know what behavior will avoid the infliction of the threatened harm, which means that the required standards of behavior, in addition to being public, must be sufficiently clear and precise. Legal systems make use of devices, in addition to careful legislative drafting, to minimize vagueness in the law, such as relying on precedent in judicial interpretation of laws. But vagueness in the standards, coupled with the fear of the threatened harm, may actually lead the threatened parties to restrict their behavior more than they would if the standards were clear and precise, to "err on the side of safety." Vagueness in the required standards of

behavior may then either undercut or enhance deterrence effectiveness.

Military Deterrence

General military deterrence differs from legal deterrence in several important respects bearing on its effectiveness.

For military deterrence, the belief that the threatener has the general ability to carry out its threat comes less easily. Military deterrence, in most cases, is mutual: each party is both the threatener and the threatened. As opposing states approach parity in military force levels, the ability of each to carry out its military threats against the other becomes more and more doubtful.

Military deterrence has important implications as well for the belief that the threatener is willing to carry out its threats. The history of particular nations executing their military threats is usually short and sometimes nonexistent. Engaging in military action, whether by challenging a threat or executing a threat, is usually very costly. In addition, whatever history of threat executions there is is likely to be ambiguous. The earlier threat executions may have been by a different regime, against a different state with different military capabilities and a different relation to the threatener, and in response to a different sort of challenge.

Since history is largely inadequate to provide a demonstration of a state's willingness to carry out its military threats, the alternative is to find a measure for attributing a presumption of willingness. In the absence of evidence to the contrary, it may be presumed that a state is willing to execute its threats if and only if it would be in its perceived self-interest to do so. The main reason for the execution of a military threat, and hence the main factor in assessing its rationality, is the military prospect for denying the opponent his objective in aggression, considered in conjunction with an appreciation of the importance of the interest compromised by the aggression.

The mutuality of military deterrence also undermines the assurance that the threatened party can avoid harm through conformity to the threatener's standards, since mutuality creates motivations for, and consequent fears of, preemption. The possibility of a preemptive attack means that neither side can be assured that conforming behavior on its part will make it immune from attack. But military deterrence is made more robust by the fact that if two states are at least close to parity in military capability, the harm that is threatened is certainly severe enough to outweigh whatever the aggressor might hope to gain. (States are notorious, however, for their lack of foresight about the cost of war.)

The threats and standards in military deterrence are, finally, more vague than in legal deterrence. The threatened party will be uncertain about what range of aggressive behavior on its part would lead to threat execution because this depends, in the context of military deterrence, on a highly speculative assessment of what responses the threatener would perceive to be

rational. While such uncertainty might lead the threatened party to be especially cautious, it may also weaken deterrence by fostering risk-taking behavior.

The purpose in comparing legal and military deterrence is not to show that military deterrence is not effective, but to show what factors are relevant to the assessment of its effectiveness. Military and legal deterrence are not alternatives to each other, since they operate in different realms, one at the domestic level and the other at the international level. To show that legal deterrence is more effective than military deterrence does not show that there is anything more effective than military deterrence at the international level (short of the domestic law of a world government). Our purpose is simply to set the stage for drawing some conclusions about a specific form of military deterrence, namely, nuclear deterrence.

Nuclear Deterrence

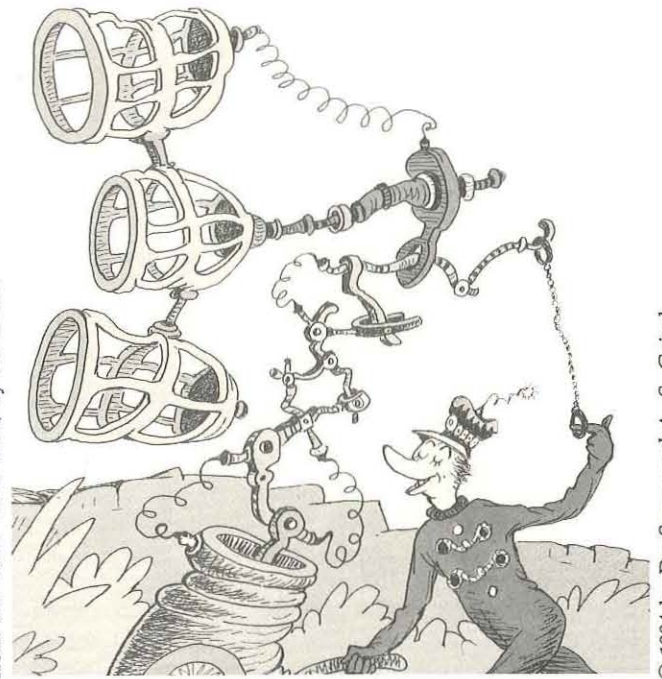
How does nuclear deterrence fare in comparison with general military deterrence in terms of the factors just discussed?

Nuclear deterrence, like general military deterrence, is mutual, but, as the label "mutual assured destruction" suggests, it is mutual in a stronger sense. The nuclear situation is one in which each side has the military ability to destroy the other side no matter who strikes first. Thus, it is a form of mutuality which should raise no doubts about the ability of the threatener to execute its threat, unlike the mutuality of other forms of military deterrence. The mutuality of nuclear deterrence thus should promote rather than undermine the effectiveness of deterrence. As a result, nuclear deterrence has the additional advantage of increasing assurances that the party threatened can avoid the threatened harm by conformity to the standards: motivations for preemption should be eliminated under nuclear deterrence, because with each side able to destroy the other no matter what, there would be nothing to gain and much to lose by striking first. (Doctrines of nuclear war fighting are, however, increasingly bringing this into question.)

Even more clearly for nuclear deterrence than for conventional military deterrence, the willingness of the threatener to execute its threats must be attributed presumptively; the history of nuclear deterrence provides no instances of threat executions. But such a presumption fails completely for nuclear deterrence. There are no reasons sufficient to make rational the execution of a nuclear threat between the superpowers in the context of the nuclear situation. There can be no interest of sufficient importance to outweigh the potential losses from military conflict when these losses carry a serious risk of amounting to the destruction of the society.

The severity of the threatened harm is a distinctive feature of nuclear threats that would seem to have a strong impact on their effectiveness. Combined with the certain ability each side has to carry out its threat, the severity of the threat creates what has been called

from *The Butter Battle Book*, by Dr. Seuss



© 1984 by Dr. Seuss and A. S. Geisel

"the crystal ball effect." Any leader of a superpower contemplating aggression against the other can foresee clearly, as if in a crystal ball, the likely outcome of total ruin. This may, as well, have an effect on whether the vagueness inherent in nuclear threats (as in other military threats) increases or decreases their deterrent effect. The crystal ball effect may lead to a decided tendency toward greater caution.

Our comparison shows that nuclear deterrence has both striking advantages and striking disadvantages when compared with general military deterrence. On the one hand, nuclear threats seem superior in deterrent value due to the distinctive character of the mutuality of nuclear threats and to the special restraining potential of the crystal ball effect. On the other hand, the deterrent value of nuclear threats seems to suffer considerably due to the failure of the presumption of willingness to execute nuclear threats. Focusing on the former set of factors would lead one to be very optimistic about the effectiveness of nuclear deterrence, while focusing on the latter set of factors would lead one to extreme pessimism. Attempting an overall comparative evaluation, taking into account both sets of factors, given their significance and the sharpness of the opposition between them, may well leave one in a great state of uncertainty.

There remains an argument, however, which may show that nuclear threats do *not* have marginal deterrent value. Threats cannot have marginal deterrent value, however much they may appear to do so, unless they have some absolute deterrent value. If there is some special feature of nuclear threats, in distinction from other military threats, that raises the possibility that they have *no* absolute deterrent value, then whatever marginal deterrent value they may seem to have would be illusory.

There is such a special feature of nuclear threats. The point is often made that nuclear deterrence can tolerate no failures, that is, no instances of nonconforming behavior (assuming that such an instance would be or would lead to nuclear war). If there was an instance of nonconforming behavior, the likely result of destruction of society would mean that the system of deterrence as a whole had failed. The system would have proved ineffective in an absolute sense, because, as one might put it, it would not have brought about fewer instances of nonconforming behavior than what is necessary to allow the social order to continue. (Likewise, a system of legal deterrence would have no absolute deterrent value if it did not succeed in avoiding complete social chaos.) For nuclear deterrence to have absolute deterrent value, we may say, the probability of its failing (say, per year) must be so low that it is very unlikely that a failure would occur over decades or even centuries. If nuclear deterrence cannot guarantee that it is very unlikely to fail over an extended number of years, it must be regarded as ineffective in an absolute sense. Unless nuclear deterrence can do a substantially better job at deterring aggression than history has

shown general military deterrence has been able to do, then nuclear deterrence is absolutely ineffective, because general military deterrence can tolerate a much higher rate of failure without social breakdown (or destruction) than nuclear deterrence can.

If nuclear deterrence is substantially more effective over the long haul than general military deterrence, it must have a substantial, not merely minimal, marginal deterrent value in comparison. But our argument does not support this: the advantages in effectiveness of nuclear deterrence over conventional military deterrence are matched by its disadvantages. As a result, there is reason to believe that nuclear deterrence has no absolute deterrent value and so no marginal deterrent value over conventional military deterrence. Having nuclear deterrence would then be worse than having no system of military deterrent threats at all.

—Steven Lee

Steven Lee is professor of philosophy at Hobart and William Smith Colleges. During 1986–87 he was a Rockefeller Resident Fellow at the Institute for Philosophy and Public Policy. This article is drawn from "The Logics of Deterrence," a chapter of a work in progress, *Morality, Prudence, and Nuclear Weapons*.

How Good a Person Do I Have to Be?

How good a person am I? This is a question most of us care about being able to answer, and most of us, I suspect, know what we would like the answer to be. I, for one, would like to think that I'm a pretty good person, not saintly, but as good as most and better than some of my colleagues and friends. I'd like to be able to pat myself on the back but also to point a finger at others; I want the requirements of morality to be lenient enough so I'll score high, but tough enough so that my high score still means something. It would be nice if it turned out that however good I am is just about exactly how good a person ought to be. Contemporary moral philosophers have given a good deal of attention to these questions. What kind of a moral report card would the rest of us get from them?

Being a Good Person

If my household's gross income last year was \$40,000 and I gave away \$2000 to charity, am I (a) a splendid person; (b) a good person; (c) an okay person; (d) a bad person? A 1984 Rockefeller Brothers Fund survey reports that the average family with an income in that range donated \$1,060, but while that lets me know that my level of giving is somewhat higher than average,

it doesn't tell me whether it is enough to satisfy the demands of morality. When I think of friends who probably gave less, I feel like leaving my income tax return casually lying around so they can see what a good person I am. But when I think about the amount of human suffering to be alleviated in the world, my contribution suddenly seems much less generous.

Most of us tend to decide what we owe others in part by seeing what's left over after we've secured a moderately comfortable—but not extravagant—standard of living for ourselves. We feel we shouldn't be faulted for aspiring to a middle class lifestyle; the wealthy, on the other hand, have a good deal to apologize for. Not surprisingly, surveys show that 95 percent of Americans consider themselves middle or working class. The rich—those who ought to be "soaked" to provide benefits for the rest of us—are invariably those who have \$10,000 a year more than we have.

I confess that I recently bought a \$1500 stereo system. That sum of money, as the UNICEF literature reminds me, can purchase a lot of vials of penicillin at 25 cents each. But, I hasten to explain, this is the first nice stereo I've ever bought in my whole entire life. And I know

for a fact that many of my friends have stereos worth five times as much. *That's* selfish, in a world where millions of people go to bed every night hungry. Of course, someone who sold her stereo to raise money for the homeless will view me much as I view my stereophile colleagues. The difference between their lifestyle and mine is far less than the difference between my lifestyle and that of starving children sleeping on the streets in Calcutta. True enough. But a St. Francis who gives all that he has to the poor is just that: a saint. I'm not expected to be morally magnificent, but only morally decent, a middling sort of person, who does my share and then goes home to listen to *Don Giovanni* on my compact disc player. To do more is to qualify for moral extra credit.

Some philosophers define the boundaries of what we have a right to keep for ourselves more narrowly,

however. Utilitarians are notorious for denying the whole category of moral extra credit. On their view, not to do a good thing is just the same morally as doing a bad thing, so every time we pass up an opportunity for a good deed, we not only forgo canonization but invite moral criticism. Peter Singer, for one, argues that "if it is in our power to prevent something very bad happening, without thereby sacrificing anything of comparable moral significance, we ought to do it." The trouble is that given how many very bad things are happening in the world, and how very bad they are, little else is of comparable moral significance, which means we may be called upon morally to give up a *lot*. Singer gives a partial list: "color television, stylish clothes, expensive dinners, a sophisticated stereo system, overseas holidays, a (second?) car, a large house, private schools for our children. . . ."



Drawing by Handelman; © 1986
The New Yorker Magazine, Inc.

"John Stuart Mill taught that the happiness of the individual is paramount. He didn't name names, but I suspect that you and I are the sort of individual he had in mind."

Nor are utilitarians the only ones who saddle the rest of us with burdensome obligations. Rights theorists are likely to think that if somebody has a right, somebody else has a duty, and the somebody else, as often as not, turns out to be us. Henry Shue, writing in *Basic Rights*, maintains that "One is required to sacrifice, as necessary, anything but one's basic rights in order to honor the basic rights of others." It sounds suspiciously like that involves sacrificing our stereotypes.

...the project of going around all day making other people happy can only get off the ground if there are at least some "first-order projects," if there is at least somebody being happy.

These arguments suggest, moreover, that we not only ought to *have* less but we ought to *do* more. It's true that a poet earning \$10,000 a year can give only so much before cutting into her own basic needs. But maybe she should consider a new career on Wall Street. Maybe the moral high road doesn't involve earning as little as you can, but earning as much as you can, so as to have more to give away. Philosophers who would make the rest of us feel guilty for not giving more may be at fault themselves for remaining in a non-lucrative profession when they could be more gainfully if less happily employed elsewhere. Whoever said you had a right to do just what you please?

Well, a number of philosophers have said, not that we have a right to do whatever we please, but that it's important to be able to make certain key choices that give meaning and coherence to our lives. In a famous argument against utilitarianism, Bernard Williams rejects the view that one is required to abandon one's own deepest commitments whenever they fail to advance the greatest good of the greatest number. To require someone to surrender his chosen life's work or his passionate creative pursuits would be "to alienate him in a real sense from his actions and the source of his actions in his own convictions. . . . It is thus, in the most literal sense, an attack on his integrity." Moreover, the project of going around all day making other people happy can only get off the ground if there are at least some "first-order projects," if there is at least somebody *being* happy. After the Revolution there has to be somebody left to do the *living*.

One problem here is that while we want some selfishness to turn out to be justified, we don't want *all* selfishness to turn out to be justified. The challenge, according to Shelly Kagan, philosopher at the University of Illinois at Chicago Circle, is a complex one: we want to explain why "it is sometimes permissible to refuse to perform an optimal act. But all those unwilling to embrace egoism must at the same time avoid arguments that rule out the possibility of there being

any moral requirements at all. Thus the explanation must also account for the fact that sometimes a given optimal act *is* required by morality." If we try to establish some protected zone of self-interest, we will have trouble showing that this zone can ever be legitimately encroached upon. And if we argue that reasons of self-interest sometimes outweigh moral reasons, then we will have to say that moral sacrifice is sometimes actually *unjustified*. But certainly we want it to be permissible, even if not required, for the moral saint to go the extra mile. Kagan concludes that it is difficult to set principled limits to what morality may demand of us.

In any case, some cold comfort may be derived from the fact that almost without exception philosophers who call for moral sacrifices fail to practice what they preach. They themselves are not rushing off to sign up with Mother Teresa. Some of them drive very nice cars. And insofar as they propose any specific guidelines for moral behavior, these tend to be calculated to reassure. Singer advocates (at minimum) giving "a round percentage of one's income like, say, 10%—more than a token donation, yet not so high as to be beyond all but saints." *That* isn't too bad, we might think. Shue points out that worldwide poverty could be reduced significantly with relatively modest sacrifices by affluent nations: "The affluent are expected not to enjoy less, but only to acquire more at a somewhat slower rate than they would if they maximized their own interests, narrowly construed." One feels uneasy, however, at Kagan's challenge: why this much and no more?

Being a Sainly Person

Despite our fond wishes otherwise, it may turn out that morality is uncomfortably demanding. Its bare minimum may look a lot like our maximum: maybe Mother Teresa herself should be doing more than she is. But those of us who don't want, frankly, to be a whole lot more moral than we are may want to reply instead: All right, morality demands a lot. But do we have to do *everything* it demands?

Despite our fond wishes otherwise, it may turn out that morality is uncomfortably demanding. Its bare minimum may look a lot like our maximum: maybe Mother Teresa herself should be doing more than she is.

Susan Wolf, a philosopher at The Johns Hopkins University, suggests that being as morally good as we can be isn't in fact an admirable goal. Glad that she herself and her loved ones are not "moral saints," she argues that "moral perfection . . . does not constitute a model of personal well-being toward which it would be particularly rational or good or desirable for a

human being to strive." In a moral saint, she argues, the moral virtues (all present and all to an extreme degree) "are apt to crowd out the nonmoral virtues, as well as many of the interests and personal characteristics that we generally think contribute to a healthy, well-rounded, richly developed character." Someone who devotes all his time to raising money for Oxfam "necessarily is not reading Victorian novels, playing the oboe, or improving his backhand." Thus his is "a life strangely barren." Nor can the moral saint, it would seem, encourage in himself otherwise delightful characteristics that go "against the moral grain," such as a cynical or sarcastic wit. A moral saint, Wolf observes, "will have to be very very nice." Nice, and dreary company.

Wolf cautions, however, that "the fact that models of moral saints are unattractive does not necessarily mean that they are unsuitable ideals. Perhaps they are unattractive because they make us feel uncomfortable—they highlight our own weaknesses, vices, and flaws. If so, the fault lies not in the characters of the saints, but in those of our own unsaintly selves." But she notes that some of the qualities the moral saint necessarily lacks are *good* qualities, qualities we *ought* to admire, "virtues, albeit nonmoral virtues, in the unsaintly characters who have them."

Thus, Wolf concludes that "moral ideals do not, and need not, make the best personal ideals. . . . we have sound, compelling, and not particularly selfish reasons to choose not to devote ourselves univocally to realizing [our] unlimited potential to be morally good." On Wolf's view, we needn't be defensive about the fact that our lives are not as morally good as they might be, because "a person may be *perfectly wonderful* without being *perfectly moral*." It is not always better to be morally better.

Conclusion

If Wolf is right, we can concede that morality is demanding, but we can devote our lives at least in part to other pursuits than making ourselves maximally moral. Her view is not a rationalization of selfishness, however, not the view that if God hadn't meant for us to grab as much as we could he wouldn't have given us two hands to grab it with. Instead, it's a call for a broader and more diverse ideal of human excellence, for an opportunity to cultivate in ourselves a rich array of both moral and nonmoral excellences.

There is little immediate danger, of course, that most of us will knock ourselves out being too good. Whatever the optimal balance between moral and nonmoral excellences (and Wolf leaves it open how this balance should be struck), most of us err on the side of selfishness pure and simple. Nothing Wolf says gives any reason not to adopt, say, a policy of tithing.

But perhaps we have reason not to become overly obsessed with moral report cards. What's tiresome is not so much being good, but harping on goodness from morning to night. We could all probably stand to be a lot better morally than we are, and a lot better nonmorally as well, but maybe a first start toward progress would be simply to do more and to talk less.

—Claudia Mills

Sources quoted in the article are: Peter Singer, *Practical Ethics* (Cambridge: Cambridge University Press, 1979), Chapter 8 ("Rich and Poor"); Henry Shue, *Basic Rights* (Princeton, N.J.: Princeton University Press, 1980); Bernard Williams, "A Critique of Utilitarianism," in J.J.C. Smart and Bernard Williams, *Utilitarianism: For and Against* (Cambridge: Cambridge University Press, 1973); Shelly Kagan, "Does Consequentialism Demand Too Much?" *Philosophy & Public Affairs*, vol. 13, no. 3 (Summer 1984); and Susan Wolf, "Moral Saints," *Journal of Philosophy*, vol. 79, no. 8 (August 1982).

Announcing the first volume in the
Cambridge Studies in Philosophy and Public Policy

Mark Sagoff

The Economy of the Earth:
Philosophy, Law, and the Environment

\$29.95

SPECIAL 20% DISCOUNT PRICE FOR QQ READERS: \$23.96

Prof. Sagoff debunks an influential approach to environmental policy that reduces social values to benefits and costs incurred by individuals. He presents an analysis based on ethical, cultural, and political concerns that reveals how environmental legislation can be sensitive to political reality and responsive to economic costs.

"It is timely in terms of the issues and problems discussed. . . . The examples, which are typically rendered in personal, anecdotal forms, are often both humorous and telling, resonating with and illuminating features of everyday experience and action."—Richard Wasserstrom

To order, send checks or money orders to Cambridge University Press, 32 East 57th Street, New York, NY 10022. (All orders must be prepaid; residents of New York and California include sales tax.) To order by VISA or Mastercard call 800-872-7423. ISBN 0-521-34113-2. To obtain the special discount, use Order #796.

The Institute for Philosophy and Public Policy was founded in 1976 to conduct research into the conceptual and normative questions underlying public policy formulation. This research is conducted cooperatively by philosophers, policymakers and analysts, and other experts from within and without the government.

All material copyright ©1988 by the Institute for Philosophy and Public Policy, unless otherwise acknowledged. For permission to reprint articles appearing in *QQ*, please contact the editor.

Editor: Claudia Mills

STAFF:

Douglas MacLean, Director
Robert K. Fullinwider, Research Scholar
Nancy Jecker, Rockefeller Resident Fellow
Judith Lichtenberg, Research Scholar
David Luban, Research Scholar
Mark Sagoff, Research Scholar
Jerome Segal, Research Scholar
Robert Wachbroit, Research Scholar
Claudia Mills, Editor
Kathleen Wiersema, Assistant to the Director

ADVISORY BOARD

Brian Barry / European University Institute, Florence, Italy
Hugo Bedau / Professor of Philosophy, Tufts University
Richard Bolling / Retired Member, U.S. House of Representatives
Peter G. Brown / School of Public Affairs, University of Maryland
Daniel Callahan / Director, The Hastings Center
David Cohen / Roosevelt Center
Joel Fleishman / Vice Chancellor, Duke University
Virginia Held / Professor of Philosophy, City University of New York
Charles McC. Mathias, Jr. / Retired Member, U.S. Senate
Michael Nacht (ex officio) / Dean, School of Public Affairs, University of Maryland

Institute for Philosophy and Public Policy
University of Maryland
College Park, Maryland 20742

Address correction requested.

THIRD CLASS BULK Non-Profit Organization U.S. Postage PAID Permit No. 10 College Park, Md.
--