

The Reversal Test: Eliminating Status Quo Bias in Applied Ethics*

Nick Bostrom and Toby Ord

I. INTRODUCTION

Suppose that we develop a medically safe and affordable means of enhancing human intelligence. For concreteness, we shall assume that the technology is genetic engineering (either somatic or germ line), although the argument we will present does not depend on the technological implementation. For simplicity, we shall speak of enhancing “intelligence” or “cognitive capacity,” but we do not presuppose that intelligence is best conceived of as a unitary attribute. Our considerations could be applied to specific cognitive abilities such as verbal fluency, memory, abstract reasoning, social intelligence, spatial cognition, numerical ability, or musical talent. It will emerge that the form of argument that we use can be applied much more generally to help assess other kinds of enhancement technologies as well as other kinds of reform. However, to give a detailed illustration of how the argument form works, we will focus on the prospect of cognitive enhancement.

Many ethical questions could be asked with regard to this prospect, but we shall address only one: do we have reason to believe that the long-term consequences of human cognitive enhancement would be, on balance, good? This may not be the only morally relevant question—

* For comments, we are grateful to Daniel Brock, David Calverley, Arthur Caplan, Jonathan Glover, Robin Hanson, Michael Sandel, Julian Savulescu, Peter Singer, Mark Walker, and to the participants of the “Methods in Applied Ethics” seminar at Oxford, the “How Can Human Nature Be Ethically Improved” conference in New York, the “Sport Medicine Ethics” conference in Stockholm, and the Oxford-Scandinavia Ethics Summit, where earlier versions of this article were presented. We are also grateful for the helpful comments from two anonymous referees and six anonymous members of the editorial board.

we leave open the possibility of deontological constraints—but it is certainly of great importance to any ethical decision making.¹

It is impossible to know what the long-term consequences of such an intervention would be. For simplicity, we may assume that the immediate biological effects are relatively well understood, so that the intervention can be regarded as medically safe. There would remain great uncertainty about the long-term direct and indirect consequences, including social, cultural, and political ramifications. Furthermore, even if (*per impossibile*) we knew what all the consequences would be, it might still be difficult to know whether they are on balance good. When assessing the consequences of cognitive enhancement, we thus face a double epistemic predicament: radical uncertainty about both prediction and evaluation.

This double predicament is not unique to cases involving cognitive enhancement or even human modification. It is part and parcel of the human condition. It arises in practically every important deliberation, in individual decision making as well as social policy. When we decide to marry or to back some major social reform, we are not—or at least we shouldn't be—under any illusion that there exists some scientifically rigorous method of determining the odds that the long-term consequences of our decision will be a net good. Human lives and social systems are simply too unpredictable for this to be possible. Nevertheless, some personal decisions and some social policies are wiser and better motivated than others. The simple point here is that our judgments about such matters are not based exclusively on hard evidence or rigorous statistical inference but rely also—crucially and unavoidably—on subjective, intuitive judgment.

The quality of such intuitive judgments depends partly on how well informed they are about the relevant facts. Yet other factors can also have a major influence. In particular, judgments can be impaired by various kinds of biases. Recognizing and removing a powerful bias will sometimes do more to improve our judgments than accumulating or analyzing a large body of particular facts. In this way, applied ethics could benefit from incorporating more empirical information from psychology and the social sciences about common human biases.

In this article we argue that one prevalent cognitive bias, status quo bias, may be responsible for much of the opposition to human en-

1. In parallel to affirming deontological side constraints, one might also hold that the value of a state of affairs depends on how that state was brought about. For instance, one might hold that the value of a state of affairs is reduced if it resulted from a decision that violated a deontological side constraint. When we discuss the consequentialist dimension of ethical or prudential decision making in this article, we mainly set aside this possibility.

hancement in general and to genetic cognitive enhancement in particular. Our strategy is as follows: first, we briefly review some of the psychological evidence for the pervasiveness of status quo bias in human decision making. This evidence provides some reason for suspecting that this bias may also be present in analyses of human enhancement ethics. We then propose two versions of a heuristic for reducing status quo bias. Applying this heuristic to consequentialist objections to genetic cognitive enhancements, we show that these objections are affected by status quo bias. When the bias is removed, the objections are revealed as extremely implausible. We conclude that the case for developing and using genetic cognitive enhancements is much stronger than commonly realized.

II. PSYCHOLOGICAL EVIDENCE OF STATUS QUO BIAS

That human thinking is susceptible to the influence of various biases has been known to reflective persons throughout the ages, but the scientific study of cognitive biases has made especially great strides in the past few decades.² We will focus on the family of phenomena referred to as status quo bias, which we define as an inappropriate (irrational) preference for an option because it preserves the status quo.

While we must refer the reader to the scientific literature for a comprehensive review of the evidence for the pervasiveness of status quo bias, a few examples will serve to illustrate the sorts of studies that have been taken to reveal this bias.³ These examples will also help delimit the particular kind of status quo bias that we are concerned with here.

The Mug Experiment.—Two groups of students were asked to fill out a short questionnaire. Immediately after completing the task, the students in one group were given decorated mugs as compensation, and the students in the other group were given large Swiss chocolate bars. All participants were then offered the choice to exchange the gift they had received for the other, by raising a card with the word “Trade” written on it. Approximately 90 percent of the participants retained the original reward.⁴

Since the two kinds of reward were assigned randomly, one would have expected that half the students would have got a different reward from the one they would have preferred *ex ante*. The fact that 90 percent of the participants preferred to retain the award they had been given

2. See, e.g., Thomas Gilovich, Dale W. Griffin, and Daniel Kahneman, *Heuristics and Biases: The Psychology of Intuitive Judgment* (Cambridge: Cambridge University Press, 2002).

3. For a good introduction to the literature on status quo bias and related phenomena, see Daniel Kahneman and Amos Tversky, *Choices, Values, and Frames* (Cambridge: Cambridge University Press, 2000).

4. Gilovich et al., *Heuristics and Biases*.

illustrates the “endowment effect,” which causes an item to be viewed as more desirable immediately upon its becoming part of one’s endowment.

The endowment effect may suggest a status quo bias. However, we have defined status quo bias as an inappropriate favoring of the status quo. One may speculate that the favoring of the status quo in the Mug Experiment results from the subjects forming an emotional attachment to their mug (or chocolate bar). An endowment effect of this kind may be a brute fact about human emotions and as such may be neither inappropriate nor in any sense irrational. The subjects may have responded rationally to an a-rational fact about their likings. There is thus an alternative explanation of the Mug Experiment which does not involve status quo bias.

In this article, we want to focus on genuine status quo bias that can be characterized as a cognitive error, where one option is incorrectly judged to be better than another because it represents the status quo. Moreover, since our concern is with ethics rather than prudence, our focus is on (consequentialist) ethical judgments. In this context, instances of status quo bias cannot be dismissed as merely apparent on grounds that the evaluator is psychologically predisposed to like the status quo, for the task of the evaluator is to make a sound ethical judgment, not simply to register his or her subjective likings. Of course, people’s emotional reactions to a choice may form part of the consequences of the choice and have to be taken into account in the ethical evaluation. Yet status quo bias remains a real threat. It is perfectly possible for a decision maker to be biased in judging the strength of people’s emotional reactions to a change in the status quo.⁵ Explanations in terms of emotional bonding seem less likely to account for the findings in the following two studies.

Hypothetical Choice Tasks.—Some subjects were given a hypothetical choice task in the following “neutral” version, in which no status quo was defined: “You are a serious reader of the financial pages but until recently you have had few funds to invest. That is when you inherited a large sum of money from your great-uncle. You are considering different portfolios. Your choices are to invest in: a moderate-risk company, a high-risk company, treasury bills, municipal bonds.” Other subjects were presented with the same problem but with one of the options designated as the status quo. In this case, the opening passage continued: “A significant portion of this portfolio is invested in a moderate risk company . . . (The tax and

5. Independent of the issue of status quo bias, there is evidence of a durability bias in affective forecasting, which leads people to systematically overestimate the duration of emotional reactions to future events; see, e.g., Gilovich et al., *Heuristics and Biases*, 292ff.

broker commission consequences of any changes are insignificant.)” The result was that an alternative became much more popular when it was designated as the status quo.⁶

Electric Power Consumers.—California electric power consumers were asked about their preferences regarding trade-offs between service reliability and rates. The respondents fell into two groups, one with much more reliable service than the other. Each group was asked to state a preference among six combinations of reliability and rates, with one of the combinations designated as the status quo. A strong bias to the status quo was observed. Of those in the high-reliability group, 60.2 percent chose the status quo, whereas a mere 5.7 percent chose the low-reliability option that the other group had been experiencing, despite its lower rates. Similarly, of those in the low-reliability group, 58.3 chose their low-reliability status quo, and only 5.8 chose the high-reliability option.⁷

It is hard to prove irrationality or bias, but taken as a whole, the evidence that has accumulated in many careful studies over the past several decades is certainly suggestive of widespread status quo bias. In considering the examples given here, it is important to bear in mind that they are extracted from a much larger body of evidence. It is easy to think of alternative explanations for the findings of these particular studies, but many of the potential confounding factors (such as transaction costs, thinking costs, and strategic behavior) have been ruled out by further experiments. Status quo bias plays a central role in prospect theory, an important recent development in descriptive economics (which earned one of its originators, Daniel Kahneman, a Nobel Prize in 2002).⁸ Psychologists and experimental economists have found extensive evidence for the prevalence of status quo bias in human decision making.⁹

6. William Samuelson and Richard Zeckhauser, “Status Quo Bias in Decision Making,” *Journal of Risk and Uncertainty* 1 (1988): 7–59.

7. Raymond S. Hartman, Michael J. Doane, and Chi-Keung Woo, “Consumer Rationality and the Status Quo,” *Quarterly Journal of Economics* 106 (1991): 141–62.

8. The work of Kahneman and Amos Tversky and their collaborators has convinced many economists that the standard economic paradigm, which postulates rational expected-utility maximizing agents, is, despite its simplicity and convenient formal features, not descriptively adequate for many situations of human decision making.

9. The exact nature and the psychological factors contributing to status quo bias are not yet fully understood. Loss aversion—the tendency to place a greater weight on aspects of outcomes when they are represented as “losses” (rather than, e.g., forfeited gains)—seems to be a significant part of the picture (James N. Druckman, “Evaluating Framing Effects,” *Journal of Economic Psychology* 22 [2001]: 91–101). It has also been suggested that omission bias may account for some of the findings previously ascribed to status quo bias. Omission bias is diagnosed when a decision maker prefers a harmful outcome that results from an omission to a less harmful outcome that results from an action (even in cases

Let us consider one more illustration of the empirical literature on status quo bias. One source of status quo bias is loss aversion, which can seduce people into judging the same set of alternatives differently depending on whether they are phrased in terms of potential losses or potential gains.

The Asian Disease Problem.—The same cover story was presented to all the subjects: “Imagine that the U.S. is preparing for the outbreak of an unusual Asian disease, which is expected to kill 600 people. Two alternative programs to combat the disease have been proposed. Assume that the exact scientific estimates of the consequences of the programs are as follows.” One group of subjects was presented with the following pair alternatives (the percentage of respondents choosing a given program is given in parentheses):

If Program A is adopted, 200 people will be saved (72 percent).
If Program B is adopted, there is a one-third probability that 600 people will be saved and a two-thirds probability that no people will be saved (28 percent).

Another group of subjects were instead offered the following alternatives:

If Program C is adopted, 400 people will die (22 percent).
If Program D is adopted, there is a one-third probability that nobody will die and a two-thirds probability that 600 people will die (78 percent).¹⁰

It is easy to verify that the options A and B are indistinguishable in real terms from options C and D, respectively. The difference is solely one of framing. In the first formulation, the outcomes are represented as gains (people are saved), while in the second formulation, outcomes are represented as losses (people die). The second formulation, however, assumes a reference state where nobody dies of the disease, and Program D is the only way to possibly avoid a loss. In the first formulation, by contrast, the assumed reference state is that nobody lives, and ordinary risk aversion explains why people prefer Program A (the safe bet).

The bias to avoid outcomes that are framed as “losses” is both pervasive and robust.¹¹ This has long been recognized by marketing

where presumably no moral deontological constraints are involved; Ilana Ritov and Jonathan Baron, “Status-Quo and Omission Biases,” *Journal of Risk and Uncertainty* 5 [1992]: 49–61).

10. Amos Tversky and Daniel Kahneman, “The Framing of Decisions and the Psychology of Choice,” *Science* 211, no. 4481 (1981): 453–58.

11. For a review of more recent confirmations of this framing effect, see, e.g., Druckman, “Evaluating Framing Effects.”

professionals. Credit card companies, for instance, lobbied vigorously to have the difference between a product's cash price and credit card price labeled "cash discount" (implying that the credit price is the reference point) rather than "credit card surcharge," presumably because consumers would be less willing to accept the "loss" of paying a surcharge than to forgo the "gain" of a discount.¹² The bias has been demonstrated among sophisticated respondents as well as among naive ones. For example, one study found that preferences of physicians and patients for surgery or radiation therapy for lung cancer varied markedly when their probable outcomes were described in terms of mortality or survival.¹³

Changes from the status quo will typically involve both gains and losses, with the change having good overall consequences if the gains outweigh these losses. A tendency to overemphasize the avoidance of losses will thus favor retaining the status quo, resulting in a status quo bias. Even though choosing the status quo may entail forfeiting certain positive consequences, when these are represented as forfeited "gains" they are psychologically given less weight than the "losses" that would be incurred if the status quo were changed.

Having noted that a body of data from psychology and experimental economics provides at least *prima facie* grounds for suspecting that a status quo bias may be endemic in human cognition, let us now turn to the case of human cognitive enhancement. Does status quo bias affect our judgments about such enhancements? If so, how can the bias be diagnosed and removed?

III. A HEURISTIC FOR REDUCING STATUS QUO BIAS

Many people judge that the consequences of increasing intelligence would be bad, even assuming that the method used would be medically safe. While enhancing intelligence would clearly have many potential benefits, both for individuals and for society, some feel that the outcome would be worse on balance than the status quo because increased intelligence might lead people to become bored more quickly, to become more competitive, or to be better at inventing destructive weapons; because social inequality would be aggravated if only some people had access to the enhancements; because parents might become less accepting of their children; because we might come to lose our "openness to the unbidden"; because the enhanced might oppress the rest; or

12. R. Thaler, "Toward a Positive Theory of Consumer Choice," *Journal of Human Behavior and Organization* 1 (1980): 39–60.

13. B. J. McNeil, S. G. Pauker, H. G. Sox, and A. Tversky, "On the Elicitation of Preferences for Alternative Therapies," *New England Journal of Medicine* 306 (1982): 1259–62.

because we might come to suffer from “existential dread.”¹⁴ These worries are often combined with skepticism about the potential upside of enhancement of cognitive and other human capacities:

Whether a general ‘improvement’ in height, strength, or intelligence would be a benefit at all is even more questionable. To the individual such improvements will benefit his or her social status, but only as long as the same improvements are not so widespread in society that most people share them, thereby again levelling the playing field. . . . What would be the status of Eton, Oxford and Cambridge if all could go there? . . . In general there seems to be no connection between intelligence and happiness, or intelligence and preference satisfaction. . . . Greater intelligence could, of course, also be a benefit if it led to a better world through more prudent decisions and useful inventions. For this suggestion there is little empirical evidence.¹⁵

In a recent article, another author opines: “Crucially, though, despite the fact that parents may want their children to be ‘intelligent’, where all parents want this any beneficial effect is nullified. On the one hand, intelligence could be raised to the same amount for all or, alternatively, intelligence could be raised by the same amount for all. In either case no one actually benefits over anyone else. . . . [The] aggregate effect, if all parents acted the same, would be that all their children would effectively be the same, in terms of outcome, as without selection.”¹⁶

Others have argued that the benefits of cognitive enhancement (for rationality, invention, or quality of life) could be very large and

14. See, e.g., Søren Holm, “Genetic Engineering and the North-South Divide,” in *Ethics and Biotechnology*, ed. A. Dyson and J. Harris (New York: Routledge, 1994), 47–63; Gregory S. Kavka, “Upside Risks: Social Consequences of Beneficial Biotechnology,” in *Are Genes Us? The Social Consequences of the New Genetics*, ed. C. Cranor (New Brunswick, NJ: Rutgers University Press, 1994), 155–79; George J. Annas, Lori B. Andrews, and Rosario M. Isasi, “Protecting the Endangered Human: Toward an International Treaty Prohibiting Cloning and Inheritable Alterations,” *American Journal of Law and Medicine* 28 (2002): 151–78; Francis Fukuyama, *Our Posthuman Future: Consequences of the Biotechnology Revolution* (New York: Farrar, Straus & Giroux, 2002); Leon Kass, *Life, Liberty, and the Defense of Dignity: The Challenge for Bioethics* (San Francisco: Encounter Books, 2002); Michael Sandel, “The Case against Perfection,” *Atlantic Monthly* 293 (2004): 51–62. For some of these, such as Leon Kass, it is sometimes difficult to discern to what extent the objection refers to the (narrow) effects of the intervention or to the mere fact that intervention and control is exercised. We partially address objections regarding the degree of control in Sec. V.

15. Holm, “Genetic Engineering and the North-South Divide,” 60 and n. 9.

16. Kean Birch, “Beneficence, Determinism and Justice: An Engagement with the Argument for the Genetic Selection of Intelligence,” *Bioethics* 16 (2005): 12–28, 24.

that many of the risks have been overstated.¹⁷ To proponents of this view, opinions like those expressed in the above quotation tend to seem puzzling: why think that greater mental faculties would be of no value if everybody shared in the improvement? Why be so suspicious of the consequences of the biological enhancement of intelligence when more familiar efforts to improve thinking ability (such as education) are met with near-universal approbation? To proponents, the idea that these negative judgments might derive partially from a bias against the new might seem plausible even without further argument. Opponents, of course, could return fire by charging proponents with a contrary bias in favor of the new. We need some way of adjudicating between the differing intuitions.

How can we determine whether the judgments opposing cognitive enhancement result from a status quo bias? One way to proceed is by reversing our perspective and asking a somewhat counterintuitive question: "Would using some method of safely lowering intelligence have net good consequences?"

The great majority of those who judge increases to intelligence to be worse than the status quo would likely also judge decreases to be worse than the status quo. But this puts them in the rather odd position of maintaining that the net value for society provided by our current level of intelligence is at a local optimum, with small changes in either direction producing something worse. We can then ask for an explanation of why this should be thought to be so. If no sufficient reason is provided, our suspicion that the original judgment was influenced by status quo bias is corroborated.

In its general form, the heuristic looks like this:

Reversal Test: When a proposal to change a certain parameter is thought to have bad overall consequences, consider a change to the same parameter in the opposite direction. If this is also thought to have bad overall consequences, then the onus is on those who reach these conclusions to explain why our position cannot be improved through changes to this parameter. If they are unable to

17. See, e.g., Ainsley Newson and Robert Williamson, "Should We Undertake Genetic Research on Intelligence?" *Bioethics* 13 (1999): 327–42; James Hudson, "What Kinds of People Should We Create?" *Journal of Applied Philosophy* 17 (2000): 131–43; Mark Walker, "Prolegomena to Any Future Philosophy," *Journal of Evolution and Technology* 10 (2002), <http://jetpress.org/contents.htm>; Nick Bostrom, "Human Genetic Enhancements: A Transhumanist Perspective," *Journal of Value Inquiry* 37 (2003): 493–506; see also Jonathan Glover, *What Sort of People Should There Be?* (Harmondsworth: Pelican, 1984); Anjan Chatterjee, "Cosmetic Neurology: The Controversy over Enhancing Movement, Mentation, and Mood," *Neurology* 63 (2004): 968–74; M. J. Farah, J. Illes, R. Cook-Deegan, H. Gardner, E. Kandel, P. King, E. Parens, B. Sahakian, and P. R. Wolpe, "Neurocognitive Enhancement: What Can We Do and What Should We Do?" *Nature Reviews Neuroscience* 5 (2004): 421.

do so, then we have reason to suspect that they suffer from status quo bias.

The rationale of the Reversal Test is simple: if a continuous parameter admits of a wide range of possible values, only a tiny subset of which can be local optima, then it is *prima facie* implausible that the actual value of that parameter should just happen to be at one of these rare local optima (fig. 1). This is why we claim that the burden of proof shifts to those who maintain that some actual parameter is at such a local optimum: they need to provide some good reason for supposing that it is so.

Obviously, the Reversal Test does not show that preferring the status quo is always unjustified. In many cases, it is possible to meet the challenge posed by the Reversal Test and thus to defeat the suspicion of status quo bias. Let us examine some of the possible ways in which one could try to do this in the case of medically safe, financially affordable, cognitive enhancement.

The Argument from Evolutionary Adaptation

For some biological parameters, one may argue on evolutionary grounds that it is likely that the current value is a local optimum. The idea is that we have adapted to live in a certain kind of environment, and that if a larger or a smaller value of the parameter had been a better adaptation, then evolution would have ensured that the parameter would have had this optimal value. For example, one could argue that the average ratio between heart size and body size is at a local optimum, because a suboptimal ratio would have been selected against. This argument would shift the burden of proof back on somebody who maintains that a particular person's heart—or the average human heart-to-body-size ratio—is too large or too small.

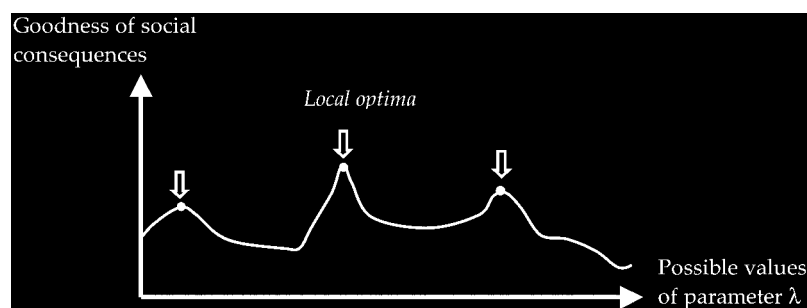


FIG. 1.—Only the points indicated with arrows are local optima. Typically, only a few points will be local optima, and most points will be such that a small shift in λ in the appropriate direction will increase the goodness of the social consequences.

The applicability of this evolutionary argument, however, is limited for several reasons. First, our current environment is in many respects very different from that of our evolutionary ancestors. A sweet tooth might have been adaptive in the Pleistocene, where high-calorie foods were scarce and the risk of starvation outweighed the health risks of a sugary diet. In wealthy modern societies, where a Mars bar is never far away, the risks of obesity and diabetes outweigh the risk of undernutrition, and a sweet tooth is now maladaptive. Our modern environment also places very different demands on cognitive functioning than did an illiterate life on the savanna: numeracy, literacy, logical reasoning, and the ability to concentrate on abstract material for prolonged periods of time have become important skills that facilitate successful participation in contemporary society.

Second, even if, say, a greater capacity for abstract reasoning had in itself been evolutionarily adaptive in the period of human evolutionary adaptation, there may have been trade-offs that made an increase in this parameter on balance maladaptive. For example, a larger brain might be correlated with greater cognitive capacity, yet a larger brain incurs substantial metabolic costs.¹⁸ These metabolic costs are no longer significant, thanks to the easy availability of food, suggesting that we may not be optimally adapted to the current environment. Similarly, the size of the birth canal used to place severe limitations on the head size of newborns, but this constraint is ameliorated by modern obstetrics and the possibility of cesarean section. An extended period of maturation was also vastly riskier ten thousand years ago than it is today.

Third, even if some trait would have been adaptive for our Pleistocene predecessors, there is no guarantee that evolutionary trial and error would have discovered it. This is especially likely for polygenic traits that are only adaptive once fully developed but that incur a fitness penalty in their intermediary stages of evolution. In some cases, the evolution of such traits may require an improbable coincidence of several simultaneous mutations that may simply not have occurred among our finite number of ancestors. An advanced genetic engineer, by contrast, may be able to solve some of the problems that proved intractable to blind evolution. She can think backward, starting with a goal in mind and working out what genetic modifications are necessary to attain it.

Fourth, there is no general reason for thinking that what evolution selects for—inclusive fitness—coincides with what makes our lives go well individually, much less collectively. The traits that would maximize our individual or collective well-being are not always the ones that maximize our tendency to propagate our genetic material. Evolution doesn't

18. Richard J. Haier, Rex E. Jung, Ronald A. Yeo, Kevin Head, and Michael T. Alkire, "Structural Brain Variation and General Intelligence," *Neuroimage* 23 (2004): 425.

care about human happiness. A capacity for rape, plunder, cheating, and cruelty might well have been evolutionarily adaptive, yet they have disastrous consequences for human welfare. Regarding intellectual faculties, we place a value on understanding, knowledge, and cognitive sensitivity that goes beyond the contribution these traits may make to our ability to survive and reproduce.

If we have reason for thinking that, for some human parameter, its role in contemporaneous society is identical to its role in Pleistocene human society, and that the trade-offs that this parameter involves have not changed, and that evolution would have had enough time to chance upon the optimum value, and that this parameter bears the same relation to human well-being as it did to reproductive success in the Pleistocene, then the argument from evolutionary adaptation would successfully meet the challenge posed by the Reversal Test. While these conditions may well hold for the heart-to-body-size ratio, they do not hold for all human parameters. In particular, they do not hold for human cognitive ability.

The Argument from Transition Costs

Consider the reluctance of the United States to move to the metric system of measurement units. While few would doubt the superiority of the metric system, it is nevertheless unclear whether the United States should adopt it. In cases like this, the transition costs are potentially so high as to overwhelm the benefits to be gained from the new situation. Those who oppose both increasing and decreasing some parameter can potentially appeal to such a rationale to explain why we should retain the status quo without having to insist that the status quo is (locally) optimal.

In the case of cognitive enhancements, one can anticipate many transition costs. Maybe school curricula would have to be redesigned to match the improved learning capacity of enhanced children. Tax codes and other regulations are often designed to strike a trade-off between how well they serve their intended function and how complex they are. (Complex regulations are harder to learn and enforce.) If people could learn complex rules more quickly, it may be appropriate to reevaluate these trade-offs and perhaps to adopt a more nuanced and complex set of social norms and regulations. Some games and recreational activities may likewise have to be modified to provide interesting levels of challenge to smarter participants. In the case of germline interventions, cognitively enhanced children might be raised by parents of normal cognitive ability, which could conceivably create some friction in such families and necessitate more preschool educational opportunities.

It is easy to overstate such transitional burdens. The cost would be

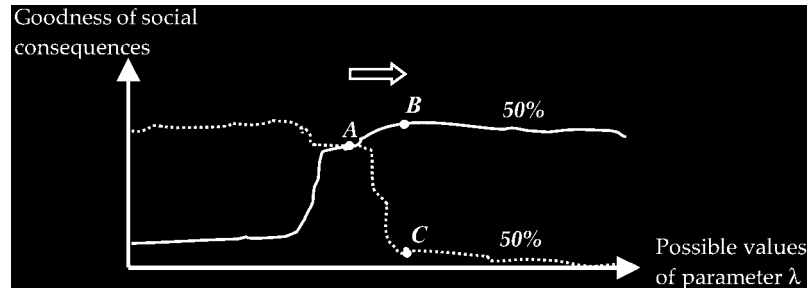


FIG. 2.—There is uncertainty whether the goodness of social consequences of a given value of a parameter λ is represented by the solid or the dotted line. A society that is currently in state A is not at a local optimum but may nevertheless resist a small shift in the parameter λ because of the risk that it would bring about state C rather than state B.

one-off while the benefits of enhancement would be permanent. School curricula are frequently rewritten for all sorts of trivial reasons. Modifying tax codes and regulations to fit a population with increased average intelligence would not be strictly necessary; it would simply be an opportunity to reap additional benefits of enhancement. Games and recreational activities are easy to invent, and we already have many games and cultural treasures that would presumably remain rewarding to people with substantially enhanced cognitive capacities. Even today, smart children are often raised by less smart parents, and while this might create problems in a few cases it certainly does not justify the conclusion that it would have been better, all things considered, if these children had been less talented than they are.

It would, however, be very difficult to exhaustively evaluate each possible transition cost against the permanent gains. Judgments about the balance between transition costs and long-term benefits will inevitably involve appeals to subjective intuitions. Such intuitions can easily be influenced by status quo bias. In Section IV, we will therefore present an extended version of our heuristic that takes account of transition costs.

The Argument from Risk

Even if it is agreed that we are probably not at a local optimum with respect to some parameter under consideration, one could still mount an argument from the risk against varying the parameter. If it is suspected that the potential gains from varying the parameter are quite low and the potential losses very high, it may be prudent to leave things as they are (fig. 2).

Uncertainty about the goodness of the consequences also means that results may be much better than anybody expected. It is not clear

that such uncertainty by itself provides any consequentialist ground whatsoever for resisting a proposed intervention. Only if the expectation value of the hypothetical negative results is larger than the expectation value of the hypothetical positive results does the uncertainty favor the preservation of the status quo.

The potential for unexpected gains should not be dismissed as a far-fetched theoretical possibility. In the case of cognitive enhancement, unanticipated consequences of enormous positive value seem not at all implausible. The fact that it may be easier to vividly imagine the possible downsides than the possible upsides of cognitive enhancement might psychologically join with other sources of human loss aversion to form a particularly strong status quo bias.

Imagine a tribe of *Australopithecus* debating whether they should enhance their intelligence to the level of modern humans. Is there any reason to suppose that they would have been able to foresee all the wonderful benefits we are enjoying thanks to our improved intellect? Only in retrospect did the myriad technological and social gains become apparent. And it would have been even less feasible for an *Australopithecus* to foresee the qualitative changes in our ways of experiencing, thinking, doing, and relating that our greater cognitive capacity have enabled, including literature, art, music, humor, poetry, and the rest of Mill's "higher pleasures." All these would have been impossible without our enhanced mental capacities; who knows what other wonderful things we are currently missing out on? It is an essential aspect of greater cognitive faculties that they facilitate new insights, inventions, and creative endeavors, as well as enabling new ways of thinking and experiencing. The uncertainty of the ultimate consequences of cognitive enhancement, far from being a sufficient ground for opposing them, is actually a strong consideration in their support.¹⁹

While some of the potential benefits might be hard to imagine, other benefits of greater cognitive faculties are quite plain. Diseases need cures, scientific questions need answers, poverty needs alleviation, and environmental problems need solutions. While a widespread increase in intelligence may not be sufficient to achieve all these goals, it could clearly help. Even the foreseeable benefits are very great.

One might object to this balancing of potential losses with potential gains by claiming that when it comes to the moral assessment of consequences, there is some normatively appropriate level of risk aversion which we must take into account. However, even if we accept such an account, and even if we completely disregard the possible gains mentioned above (both the unpredictable, and the more predictable ones),

19. Compare Nick Bostrom, "Transhumanist Values," in *Ethical Issues for the 21st Century*, ed. F. Adams (Charlottesville, VA: Philosophical Documentation Center, 2004).

it would still be difficult to make a case against cognitive enhancement. This is because while cognitive enhancement may create certain novel risks, it may also help to reduce many serious threats to humanity. In evaluating the riskiness of cognitive enhancement we must take into account both its risk-increasing and its risk-reducing effects. Mitigation of risk could result from a greater ability to protect ourselves from a wide range of natural hazards such as viral pandemics. There may also be threats to human civilization that we have not yet understood, but which greater intelligence would enable us to anticipate and counteract. The goal of reducing overall risk might turn out to be a strong reason for trying to develop ways to enhance our intelligence as soon as possible.²⁰

The Argument from Person-Affecting Ethics

Suppose that the cognitively enhanced would lead better lives. Does that give us a moral reason to enhance ourselves? Or to create cognitively enhanced people? It is possible to hold a person-affecting form of consequentialism according to which what we ought to do is to maximize the benefits we provide to people who either already exist or will come to exist independently of our decisions. On such views, there is no general moral reason to bring into existence people whose lives will be very good. By extension, there may be no moral reason to change ourselves into radically different sorts of people whose lives would be better than the ones we currently lead.²¹

Even if one accepts such a person-affecting ethics, one may still recognize moral reasons for supporting cognitive enhancement. In the case of bringing new people into existence, it would be difficult to deny that it would be a bad idea to deliberately select embryos with genetic disorders that cause severe retardation.²² This might indicate that one recognizes other types of moral considerations in addition to person-affecting ones or that one believes that selecting for mental retardation would adversely affect the existing population. Either way, the Reversal Test can be applied to put some pressure on those who hold these views to explain why the same kinds of reasons that make it a bad idea to

20. Compare Nick Bostrom, "Existential Risks: Analyzing Human Extinction Scenarios and Related Hazards," *Journal of Evolution and Technology* 9 (2002), <http://jetpress.org/contents.htm>.

21. See, e.g., Melinda A. Roberts, "A New Way of Doing the Best That We Can: Person-Based Consequentialism and the Equality Problem," *Ethics* 112 (2002): 315–50. Note that the idea of person-affecting ethics is not simply that what is good for one person may not be good for another. That the good for a person may partially depend on her preferences and other personal factors can of course be admitted by consequentialists who reject the person-affecting view.

22. Glover, *What Sort of People Should There Be?*

select for lower intelligence would not also make it a good idea to select for increased intelligence. For example, person-affecting reasons for bringing smarter children into the world could derive from many considerations, including the idea that present people might prefer to have such children or might benefit from being cared for in their old age by a more capable younger generation that could generate more economic resources for the elderly, invent more cures for diseases, and so on.

In the case of cognitively enhancing existing individuals, the consequentialist person-affecting reasons seem even stronger, at least for small or moderate enhancements. Practically everyone would agree that it would be the height of foolishness to set out to lower one's own intelligence, for instance, by deliberately ingesting lead paint. But if we think that becoming a little less intelligent would be bad for us, then we should either accept that becoming a little smarter would be good for us or else take on the burden of justifying the belief that we currently happen to have an optimal level of intelligence.

Very large cognitive enhancement for existing people is more problematic on a person-affecting view. A sufficiently radical enhancement might conceivably change an individual to such an extent that she would become a different person, an event that might be bad for the person that existed before. However, it is perhaps illuminating to make a comparison with children, whose cognitive capacities grow dramatically as they mature. Even though this eventually results in profound psychological changes, we don't think that it is bad for children to grow up. Similarly, it might be good for adults to continue to grow intellectually even if they eventually develop into rather different kinds of persons.²³

To summarize this section, we have proposed a heuristic for eliminating status quo bias, which transfers the onus of justification to those who reject both increasing and decreasing some human parameter. We illustrated this heuristic on the case of proposed cognitive enhancement. We considered four broad arguments by which one might attempt to carry the burden of justification, and we tried to show that, in regard to intelligence enhancement, these arguments do not succeed.

Our argument that status quo bias is widespread in bioethics thus proceeds in two steps. First, we note that the empirical literature shows that status quo bias affects many domains of human cognition, creating a *prima facie* reason for suspecting that it might affect some bioethical judgments in particular. Second, we apply the Reversal Test. Since the

23. However, in general, if the proposed change in a parameter is very large, the Reversal Test will tend to give a less definite verdict. This is because there is less *prima facie* implausibility in supposing that a larger interval of parameter values contains a local optimum than that a smaller interval does.

function of the Reversal Test is to remove whatever status quo bias is present, we infer that if our considered judgments change after the test has been applied, then our judgments prior to its implication were in fact affected by status quo bias. The four arguments we considered above are our best attempts, on behalf of the opponents of cognitive enhancement, to try to meet the burden of justification that the Reversal Test generates. We are not aware of any other arguments that have been advanced in the literature that could do this job. The test's challenge, of course, does not depend on its targets actually having offered these hypothetical arguments. If they have not and are not able to pass the test in some other way, then the indictment of status quo bias stands.

In order to further strengthen these conclusions, we will now present an extended version of the heuristic, which we call the Double Reversal Test. This version is especially useful in addressing the argument from transition costs and the argument from person-affecting ethics.

IV. THE DOUBLE REVERSAL TEST

Disaster! A hazardous chemical has entered our water supply. Try as we might, there is no way to get the poison out of the system, and there is no alternative water source. The poison will cause mild brain damage and thus reduced cognitive functioning in the current population. Fortunately, however, scientists have just developed a safe and affordable form of somatic gene therapy which, if used, will permanently increase our intellectual powers just enough to offset the toxicity-induced brain damage. Surely we should take the enhancement to prevent a decrease in our cognitive functioning.

Many years later it is found that the chemical is about to vanish from the water, allowing us to recover gradually from the brain damage. If we do nothing, we will become more intelligent, since our permanent cognitive enhancement will no longer be offset by continued poisoning. Ought we try to find some means of reducing our cognitive capacity to offset this change? Should we, for instance, deliberately pour poison into our water supply to preserve the brain damage or perhaps even undergo simple neurosurgery to keep our intelligence at the level of the status quo? Surely, it would be absurd to do so. Yet if we don't poison our water supply, the consequences will be equivalent to the consequences that would have resulted from performing cognitive enhancement in the case where the water supply hadn't been contaminated in the first place. Since it is good if no poison is added to the water supply in the present scenario, it is also good, in the scenario where the water was never poisoned, to replace that status quo with a state in which we are cognitively enhanced.

The argument contained in this thought experiment can be generalized into the following heuristic:

Double Reversal Test: Suppose it is thought that increasing a certain parameter and decreasing it would both have bad overall consequences. Consider a scenario in which a natural factor threatens to move the parameter in one direction and ask whether it would be good to counterbalance this change by an intervention to preserve the status quo. If so, consider a later time when the naturally occurring factor is about to vanish and ask whether it would be a good idea to intervene to reverse the first intervention. If not, then there is a strong *prima facie* case for thinking that it would be good to make the first intervention even in the absence of the natural countervailing factor.

The Double Reversal Test works by combining two possible perceptions of the status quo. On the one hand, the status quo can be thought of as defined by the current (average) value of the parameter in question. To preserve this status quo, we intervene to offset the decrease in cognitive ability that would result from exposure to the hazardous chemical. On the other hand, the status quo can also be thought of as the default state of affairs that results if we do not intervene. To preserve this status quo, we abstain from reversing the original cognitive enhancement when the damaging effects of the poisoning are about to wear off. By contrasting these two perceptions of the status quo, we can pin down the influence that status quo bias exerts on our intuitions about the expected benefit of modifying the parameter in our actual situation.

When this extended heuristic for assessing status quo bias can be applied, it accommodates a wider range of considerations than the simple Reversal Test. While the challenge posed by the Reversal Test can potentially be met in any of the several ways discussed above, the challenge posed by the Double Reversal Test already incorporates the possible arguments from evolutionary adaptation, transition costs, risk, and person-affecting morality into the overall assessment it makes. If we judge that, all things considered, it would be bad to reverse the original intervention when the natural factor disappears, this judgment already incorporates all these arguments.

The Double Reversal Test yields a particularly strong consequentialist reason for cognitive enhancement. While there may be a relevant difference between the two scenarios in terms of nonconsequentialist considerations (such as the distinction between acting and allowing), it is very difficult to find a difference in the expected consequences that could plausibly be thought of as decisive. Perhaps one could speculate that in the poisoning scenario, people would already have got used to

the idea of using a cognitive enhancement therapy, even though its effects were initially concealed by the presence of the natural factor. When the natural factor disappears, there might then be less psychological discomfort from allowing the enhancement to continue to operate. However, while such an effect is possible in principle, it seems unlikely that this speculative effect would be significant in realistic cases and utterly implausible to suppose that it could form a sufficient ground for opposing cognitive enhancement.²⁴

V. APPLYING THE REVERSAL TESTS TO OTHER CASES

We have illustrated the reversal heuristics on the hypothetical case of a medically safe and generally affordable enhancement of a population's cognitive capacity. The Reversal Tests, however, can be applied much more generally.

Consider a case where inequality and the distributional effects of an enhancement are concerns. Suppose, for example, that a cognitive enhancement could not be applied universally but only to some subset of the population. This might be because only the wealthy can afford to pay for it, or perhaps because certain groups decide not to avail themselves of the enhancement opportunity (perhaps on religious grounds). The development of such an enhancement would then potentially have negative consequences for social equality, and we may ask whether the benefits it would provide would be large enough to outweigh these potential inequality-increasing effects.

One way to approach this question would be to try to estimate the effects on social inequality that the development would have, come to some evaluative assessment of the severity of these effects, compare this assessment with an evaluation of the expected beneficial consequences that the enhancement technology would have, and then form a judgment of the overall expected goodness of the consequences based on

24. Alternatively, one could speculate that the direct enhancement of cognitive ability would set a different kind of precedent than either the "therapeutic enhancement" to compensate for a natural brain-damaging factor or the subsequent increase in cognitive ability that results when the natural factor disappears. But this speculation would have to be justified. If the idea is that direct cognitive enhancement would lead to further cognitive enhancement, it would have to be shown that (1) this is significantly more likely to result from direct cognitive enhancement than from therapeutic enhancement followed by a natural increase and (2) that further cognitive enhancement would be bad. But consider an iterated application of the Double Reversal Test: a series of disasters occur in which neurotoxins are released, each followed by a therapeutic enhancement to preserve the status quo and a subsequent elevation of cognitive ability when the neurotoxin disappears. At the end of the series, average cognitive ability is at a much higher level than it is today. Is there any point in this series where the brain damage ought not be compensated for by a therapeutic enhancement, or where it would be better to prevent the ensuing rise by preserving the brain-damaging factor?

this comparison. To consider the consequentialist grounds for enhancement means, of course, that one way or another we must make such a comparison. But realistically, there is no possibility of making this comparison in a completely scientifically rigorous way. Subjective intuitive judgment will inevitably enter into the assessment—both of what the likely consequences would be and of the goodness or badness of these consequences. We must therefore confront the possibility that these intuitions, which we perforce rely on, are biased in some way, and in particular the possibility that they are affected by status quo bias. This is where the Reversal Tests come in. Potential consequences that involve distributive concerns can be handled by the tests in the same way as other consequences.

In the case of cognitive enhancement, we can apply the simple Reversal Test and ask whether it would be a good thing if the treatment group (those who would be given the cognitive enhancement) instead had their cognitive capacity reduced. Are we prepared to claim that the status quo would be improved if the wealthy, say, suffered slight brain damage? If we are not prepared to make that claim, then the onus shifts to those who judge that the nonuniversal use of the cognitive enhancer would have on balance bad consequences: they need to explain why we should believe that the current cognitive ability of the potential enhancement users is at a local optimum such that both an increase and a decrease should be expected to make things worse on the whole.²⁵

We can also apply the Double Reversal Test. If the release of a hazardous chemical threatened to reduce cognitive ability among the potential enhancement users, would it be a good thing if they could use the permanent enhancement to stave off the impending decline? And if so, would it also be a good thing if, when the effects of the poison eventually started to wear off, the enhancement users refrained from taking steps to maintain their intellectual status quo (e.g., by injecting themselves with a neurotoxin)? If the answer to both these questions is yes, then there is a strong *prima facie* case for thinking that it would

25. Those (if any) who hold the opposite view should also address, e.g., whether the world would be better if nobody had access to expensive AIDS treatments, given that such treatments are not currently available to everybody. Or, to take a case more closely related to the one at hand, whether it would have been better in the past if nobody had been taught to read given that only elites had access to education. And considering that literacy is still far from universal, especially in the poorest countries, would it be better if nobody in those countries (or in developed countries?) were given this kind of cognitive enhancement unless and until everybody gets it? In such cases surely, *le mieux est l'ennemi du bien*.

be good overall—despite the assumed negative effect on equality—if the enhancement option is developed.²⁶

In a real-world situation, we are often interested in evaluating more than two alternatives. For example, we might conclude that even though it is better that the new enhancement option be allowed to reach the market than that it be banned, it is better still if its introduction be accompanied with inequality-reducing measures, for example, making the enhancement available via public health care at an affordable price. The Reversal Test could in principle be applied to evaluate each pair of policy options in turn. For instance, we could ask whether, given that the enhancement will be allowed on the market, it would be better if an inequality-*increasing* measure were implemented, and if the answer is no, we could place the onus on those who would maintain that neither inequality-increasing nor inequality-decreasing policies should be put in place to explain why we should think that the default degree of redistribution is optimal.²⁷ We can also apply the Reversal Test to cases of individual prudential decision making.

We have used the example of enhancement of cognitive ability, but the same considerations and the same heuristics can be applied to many other forms of human modification, such as proposed interventions to enhance the ability to concentrate, improve emotional well-being, reduce the need for sleep, or increase physical or sensory capacities. Those who deny that it would be a good thing for the healthy human life span to be extended may want to ask themselves whether they believe that it would be a good thing if health span were shortened, and if not, what reason there is for thinking that the current health span is optimal. They should also apply the Double Reversal Test and consider whether, for example, slowing the endogenous aging rate would be a good thing if it would serve to counterbalance some impending environmental factor that would otherwise shorten health spans; if this would be desirable,

26. There is a specific limitation when it comes to using the Reversal Test to address the issue of inequality. For the potential users who are already privileged in the status quo, the options of increasing the parameter and of decreasing it are opposites in terms of both equality and cognitive benefits. This allows our argument above to go through. However, for those who are at the average level of welfare in the society in the status quo this is not the case. While the options of increasing or decreasing cognitive ability will have opposite effects in the cognitive realm, both of these options will decrease social equality when applied to such a person. This is because in this situation any change to their well-being will create inequality: the equality of society is at a local optimum with respect to their welfare.

27. The Reversal Tests may sometimes appear to have less power to change opinion on matters of economic policy than on matters of human modification. If this is so, it might indicate that status quo bias is less pervasive in intuitions about economic policy, perhaps because we are more experienced in thinking about changes in economic policy than about changes in human nature.

then they should ask whether, were the environmental factor to be about to disappear, it would be desirable to take steps to preserve this damaging factor or else to adopt some alternative countermeasure (such as heavy smoking or an unhealthy diet) to retain the health-span status quo.

Beyond Therapy, the widely cited recent report produced by the President's Council on Bioethics, at one point comes tantalizingly close to considering a Reversal Test. After expressing many qualms and reservations about the consequences of using medical technology to extend the healthy human life span, the report reflects: "Yet if there is merit in the suggestion that too long a life, with its end out of sight and mind, might diminish its worth, one might wonder whether we have already gone too far in increasing longevity. If so, one might further suggest that we should, if we could, roll back at least some of the increases made in the average human lifespan over the past century."²⁸

In the next paragraph, the council makes clear that it does not favor such a rollback: "[Nothing] in our inquiry ought to suggest that the present average lifespan is itself ideal. We do not take the present (or any specific time past) to be 'the best of all possible worlds,' and we would not favor rolling back the average lifespan even if it were doable. Although we suggest some possible problems with substantially longer lifespans, we have not expressed, and would not express, a wish for shorter lifespans than are now the norm."²⁹

Having brought up the challenge, the council then unfortunately drops the subject after noting that while life expectancy (in the United States) has increased by thirty years in the last century, maximum life span has not changed much. But the reversal heuristic can be applied to hypothetical changes in either average or maximum life span. If the council believes that both shorter and longer maximum life spans would be worse than the present maximum life span, it owes us a convincing argument for why we should think this is so. It would have been interesting to know what conclusions the council members would have drawn if they had considered the reversal question more seriously. Would they have concluded that a shorter (maximum) life span would, after all, be preferable? Or that the current life span is exactly right? Or would they have changed their view and come out unambiguously in favor of a longer life span? Either way, the result would have been noteworthy and would have made it easier to assess the plausibility of the council's position.

The reversal heuristics do not indiscriminately favor all human en-

28. President's Council on Bioethics (U.S.), *Beyond Therapy: Biotechnology and the Pursuit of Happiness*, foreword by Leon Kass (Washington, DC: President's Council on Bioethics, 2004), 224.

29. Ibid.

hancements. For example, if we contemplate some intervention that would make everyone four inches taller, we may well come to the verdict that this would yield no net benefit. Clothes, buildings, vehicles, and so on, are designed for the current distribution of heights, so that changing average height would incur some costs. Since there are no obvious counterbalancing benefits from everybody being taller, these transition costs could justify the judgment that it is better to stick with the status quo. If we could safely and easily intervene to prevent an impending decrease in average height, say by administering growth hormone, we may have reason to do so; if whatever factor would otherwise have led to reduced height were to disappear, we might have reason to stop taking the growth hormone or to make some other intervention to prevent average height from increasing.³⁰

The Reversal Tests can be applied not only to choices affecting currently existing people but also to choices that affect what new types of people are brought into existence. Such choices, we may note, arise not only in the context of preimplantation genetic diagnosis, embryo screening, and possible future cases involving germ-line genetic modification but also in the contexts of maternal nutrition (e.g., whether to take a folic acid supplement to reduce the risk of neural tube defects), lifestyle (e.g., whether to abstain from heavy drinking during pregnancy), and timing (e.g., whether to conceive while suffering from rubella, or whether to postpone childbearing into one's forties). We can ask whether, from a consequentialist stance, it is better that a greater proportion of newborns are healthy or intelligent. Some critics of germ-line genetic enhancement have expressed doubt that it would be better if newborns had greater mental capacities. Applying the Reversal Test to this issue, we would ask whether it would be better if newborns had less intellectual capacity. If the answer is no, then we must ask for a strong justification for thinking that the current distribution of intellectual capacity in newborns is optimal.

Drawing moral conclusions about practices that influence what new types of people there should be may also require taking into account various deontological side constraints in addition to consequentialist considerations. Julian Savulescu, for example, has argued that parents have an obligation to select for the best children even if no net social

30. This example is not meant to be realistic. In the real world we have reason to celebrate the trend of increasing average height, as it is associated with beneficial developments resulting from improved nutrition. It is extremely implausible that the inconveniences of an increasing population height could ever be significant enough to outweigh the inevitable medical risks and costs of intervening to halt this trend (even setting aside important side constraints such as respect for individual autonomy).

benefit results.³¹ Others have opposed germ-line interventions on grounds that they involve an unjustifiable form of “tyranny” of the living over the unborn.³² While the heuristics offered in this article cannot fully address such deontological considerations, they may nevertheless be applied to check our intuitions for status quo bias insofar as consequentialist aspects feature in these deontological arguments.

For example, if the degree of control that present generations exercise over future ones is something that we should allegedly not increase by using germ-line therapy, we can apply the Reversal Test and ask whether we should instead reduce our control. Parents currently exert considerable influence over what new kinds of people there will be, through assortative mating, decisions whether to postpone pregnancies, the usage of preimplantation and embryonic screening, maternal nutrition, and child-rearing practices. If it is thought that it would be bad if parents had more influence over the traits of people-to-be, we should ask if it would be good if they had *less* influence. If this is not the case, we have reason for suspecting status quo bias. Michael Sandel, who argues against genetic enhancements on grounds that it is good for us to be “open to the unbidden,” seems to hold that it would be better if parents exerted less influence over their offspring than they currently do.³³ His view, therefore, may pass the Reversal Test.

The reversal heuristic is in principle applicable to any situation where we want to evaluate the consequences of some proposed change of a continuous parameter. However, its usefulness will vary from case to case. In many instances, it is possible to meet the challenge of the Reversal Tests: the method will certainly not always favor change over the status quo. The power of the heuristic lies in its ability to diagnose cases where status quo bias must be suspected and to challenge defenders of the status quo in these cases to provide further justification for their views. To what extent the example of cognitive enhancement generalizes to other issues remains to be seen, but the illustrations considered above suggest that the phenomenon is widespread. One might speculate that the popular intuition about the preferability of “the natural” might in part derive from a status quo bias. If so, then the manifestations of this bias may be endemic in human enhancement ethics and possibly in other parts of ethics as well.

A tool now exists for diagnosing status quo bias. While some reliance on intuitive judgment is unavoidable, there is no excuse for failing to test our intuitions with the most sophisticated methods available.

31. Julian Savulescu, “Procreative Beneficence: Why We Should Select the Best Children,” *Bioethics* 15 (2001): 413–26.

32. See, e.g., Jürgen Habermas, *The Future of Human Nature* (London: Blackwell, 2003).

33. Sandel, “The Case against Perfection.”