

**PARTICULARISM AND GENERALISM REVISITED: TOWARDS A PRINCIPLED
PARTICULARISM OF CONTINGENCY**

GEORGES THÉRIAULT

Thesis submitted to the University of Ottawa
in partial Fulfillment of the requirements for the M.A. degree in Philosophy

Department of Philosophy
Faculty of Arts
University of Ottawa
Unceded Algonquin Territory

© Georges Thériault, Ottawa, Canada, 2020

Abstract:

This thesis will revisit the debate between moral particularists and moral generalists in the field of meta-ethics. The general aim of this project will be to come to a better understanding of the status and role of moral principles in a reasons-holistic moral landscape. The specific aim will be to develop a viable position within the particularism-generalism debate that will combine elements from both theories. My central argument will be threefold: (a) I will argue that, in a reasons-holistic moral landscape, exceptionless moral principles are not sufficient to ground the possibility of moral thought and judgement; (b) that the possibility of moral thought and judgement depends in part on a determinate set of exceptionless moral principles and an indeterminate set of defeasible moral principles; and (c) that moral principles are insufficient to codify all or most moral truths in finite and manageable terms. My position therefore is a modified version of Principled Particularism. In this thesis, despite defending a version of particularism, I will not shy away from employing generalist terminology or from accepting certain generalist assumptions. Furthermore, unlike some particularists, I will stress the necessity and utility of moral principles. This thesis will also incorporate research about moral thought and judgement from the fields of moral psychology, cognitive science, and neuroscience.

Keywords:

Particularism, generalism, holism about reasons, moral thought, moral judgement, moral principles, moral psychology, moral learning, contingency, defeasibility, codification.

Acknowledgements

This thesis would not have been possible without the help and support of many people. First of all, I am thankful to my supervisor Andrew Sneddon for reading my drafts multiple times and providing very helpful feedback on my arguments, my successful *and* unsuccessful thought experiments, and on the overall structure of the thesis. It definitely takes another mind to help spot both the subtle and obvious mistakes in one's argument! I am also thankful for Andrew's overall support of this project as a worthwhile and engaging endeavour.

Many thanks also to David Hyder and Nigel DeSouza for being part of the evaluation committee. Both of them provided healthy and constructive criticisms during my Thesis Project Seminar held in January 2020. With their help, I hope to have managed to refine my explanations of reasons-holism, to clarify the necessity of moral principles despite their theoretical and practical limitations, and to provide some workable answer to the problem of moral relativism.

Thank you to Paul Rusnock for being the President at the Project Seminar. Paul's course on Russell and Carnap was doubtless invaluable for my philosophical development. Thank you as well to Vincent Bergeron for being the President at the Thesis Defence.

Thank you to Caitlin for always supporting me as a classmate, a fellow co-editor, and for being present at the Project Seminar.

This thesis certainly would not have been possible without the excellent team of professors at Laurentian University and UofS. So thank you again to Alain Beaulieu, Michael Yeo, Gillian Crozier, Paolo Biondi, Rachel Haliburton, Lucien Pelletier, Normand Raymond, and Réal Fillion. Many thanks to Réal for being present at the Thesis Defence and generally for always being supportive!

Huge thanks go out also to my family. My parents, my sister, and my grandparents on both sides of the family have supported me throughout my entire education. In their own ways, they have all provided me with incredible encouragement throughout my university journey.

Final thanks go to my partner Victoria. Without you, none of this would be possible. You have always been there for me in both the good times and the hard times, and you have positively shaped me into the person I am today. Thank you for your boundless love and support.

This thesis is dedicated to all of you.

Table of Contents

<u>Title Page</u>	
<u>Abstract</u>	ii
<u>Acknowledgements</u>	iii
<u>Table of Contents</u>	iv
<u>Introduction</u>	1
<u>Chapter 1: Setting the Stage</u>	
<u>1.1. Particularism vs. Generalism: The Spectrum of the Debate</u>	4
<u>1.2. Holism and Atomism About Reasons</u>	14
<u>1.3. Different Types of Moral Principles</u>	19
<u>Chapter 2: Are exceptionless moral principles sufficient to ground the possibility of moral thought and judgement?</u>	
<u>2.1. Introduction</u>	22
<u>2.2. The Categorical Imperative</u>	24
<u>2.3. The Principle of Utility</u>	37
<u>2.4. Causing Gratuitous Pain on the Unwilling</u>	47
<u>2.5. The Principle of Special Reason to do Harm</u>	49
<u>2.6. Richard Holton's <i>That's it</i> Principle</u>	51
<u>2.7. Conclusion</u>	56
<u>Chapter 3: Does the possibility of moral thought and judgement depend on the provision of a suitable supply of moral principles?</u>	
<u>3.1. Introduction</u>	59
<u>3.2. Particularism and Defeasible Principles</u>	60
<u>3.3. Determinate and Indeterminate Sets of Principles</u>	61
<u>3.4. Moral Psychology and Moral Learning</u>	64
<u>3.5. Relativism</u>	84
<u>3.6. Conclusion</u>	90
<u>Chapter 4: Can moral principles codify all or most of morality in finite and manageable terms?</u>	
<u>4.1. Introduction</u>	92
<u>4.2. Holton's Principled Particularism</u>	92
<u>4.3. Rejection of Transcendental Generalism</u>	94
<u>4.4. Rejection of Contextual Generalism</u>	103
<u>4.5. Rejection of Deep Moral Contextualism</u>	107
<u>4.6. Rejection of Anti-Transcendental Particularism</u>	113
<u>4.7. Rejection of Principle Eliminativism</u>	118
<u>4.8. Conclusion: Principled Particularism of Contingency</u>	121
<u>Conclusion</u>	124
<u>Bibliography</u>	127

Introduction

This thesis will revisit the debate between moral particularists¹ and moral generalists² in the field of meta-ethics. The general aim of this project will be to come to a better understanding of the status and role of moral principles in a reasons-holistic moral landscape. The specific aim will be to develop a viable position within the particularism-generalism debate that will combine elements from both theories. My central argument will be threefold: (a) I will argue that, in a reasons-holistic moral landscape, exceptionless moral principles are not sufficient to ground the possibility of moral thought and judgement; (b) that the possibility of moral thought and judgement depends in part on a determinate set of exceptionless moral principles and an indeterminate set of defeasible moral principles; and (c) that moral principles are insufficient to codify all or most moral truths in finite and manageable terms. My position therefore is a modified version of Principled Particularism.³ In this thesis, despite defending a version of particularism, I will not shy away from employing generalist terminology or from accepting certain generalist assumptions. Furthermore, unlike some particularists,⁴ I will stress the necessity and utility of moral principles.⁵

¹ The strongest case for moral particularism can be found in Jonathan Dancy (2004). Other notable particularists: Walter Sinnott-Armstrong (1999), Roger Crisp (2000), Margaret Little (2000), David McNaughton (2000), Garrett Cullity (2002), Richard Holton (2002), Andrew Gleeson (2007), Alan Thomas (2011), Amelia Hicks (2013), Daniel Sharp (2016).

² The strongest case for moral generalism can be found in Sean McKeever and Michael Ridge (2006). Other notable generalists: Brad Hooker (2000, 2008), Luke Robinson (2006), Manuel Hernández-Iglesias (2006) Pekka Väyrynen (2006, 2008, 2009), Robert Audi (2008), Christine Korsgaard (2008, 2009), Maike Albertzart (2013), Jennifer Zamzow (2015).

³ A view developed in Holton (2002).

⁴ E.g. principle eliminativists such as Garrett Cullity (2002), principle-abstinence particularists such as David McNaughton (2000), or anti-transcendental particularists such as Jonathan Dancy (2004). These different forms of particularism will be defined and explained in chapter 1.

⁵ Particularists who are more open to moral principles include: Little (2000) and Richard Holton (2002). It is also worth looking into what Lance and Little have produced. In both their (2004) and their (2008), they attempt to provide a middle-ground position between particularism and generalism. In the latter, they call it “deep moral contextualism”.

Chapter 1 will serve the purpose of setting the stage. First, I will discuss the various positions in the debate, both particularist and generalist, as well as my own. Second, I will define and explain holism and atomism about reasons. *Contra* McKeever and Ridge (2006) and Albertzart (2013), I will argue that the holism-atomism debate is not orthogonal to the particularism-generalism debate.⁶ Third, with the help of McKeever and Ridge, I will briefly discuss different types of moral principles such as hedged and unhedged principles, standards, guides, and action-guiding standards. Fourth, I will provide an important distinction between exceptionless and defeasible principles that will come back again and again in this thesis.

Chapter 2 will propose an answer to the following question: are exceptionless moral principles sufficient to ground the possibility of moral thought and judgement? I will argue that they are not. I will argue that many of the strongest candidates for invariant reasons are not truly invariant because they can be shown to vary in certain cases. It will be argued that, since the majority of reasons (including the strongest candidates for invariant reasons) can be shown to vary, the variability of reasons is the norm, not the exception. Consequently, exceptionless principles are rare, so rare that they are insufficient to ground the possibility of moral thought and judgement. This means we need to turn to defeasible principles to see whether they can complete the picture.

Chapter 3 will propose an answer to the following question: does the possibility of moral thought and judgement depend on the provision of a suitable supply of moral principles? In this thesis, this question will be called the “dependence question”. I will argue that the possibility of moral thought and judgement *does* depend *in part* on the provision of a suitable supply of moral principles. More specifically, I will argue that it depends in part on (1) a determinate set of exceptionless principles and (2) an indeterminate set of defeasible ones. In regards to the first

⁶ See McKeever and Ridge (2005: 93-103; 2006: 27-45) and Albertzart (2013: 342, footnote 6).

claim, the idea is that there is a specific set of universal and necessarily true principles that must figure in any account that seeks to explain the possibility of moral thought and judgement. In respect to the second claim, the idea is that there are principles that are context-dependent and contingent that also play a role in making moral thought and judgment possible. The catch, with this second claim, however, is that any set of them can contribute to making moral thought and judgement possible. In other words, there is no need for any specific set of them to be present in a given socio-cultural environment to make moral thought and judgement possible for all individuals. In this vein, my guiding assumption is that different sets of defeasible principles can contribute to the emergence and development of moral thought and judgement across different cultures and individuals. In this chapter we will also turn to the moral psychology, cognitive science, and neuroscience literatures to support our claims. A core assumption throughout this thesis will be that there are multiple factors that make moral thought and judgement possible (e.g. moral teachers and exemplars, evolutionarily entrenched moral intuitions, instincts, and domain-general learning capacities), not just moral principles.

Chapter 4 will seek to answer the following question: can moral principles codify all or most of morality in finite and manageable terms? This question will be called “the codification question”. I will argue that moral principles are insufficient to codify all or most of morality in finite and manageable terms. Still, I will argue that *some* codification is possible. First, I will argue that there are at least two exceptionless moral principles that meet all of the criteria for codification. Second, there are defeasible principles that capture default reasons that appear to be universal qua *ceteris paribus*. This chapter will also seek to reject other views such as Transcendental Generalism, Anti-Transcendental Particularism, and Deep Moral Contextualism, among others.

Chapter 1

1.1. Particularism vs. Generalism: The Spectrum of the Debate

Particularism is often described as the view that agents do not need moral principles to think and judge in moral terms. Generalism is often understood as the opposite claim. But these general markers for defining particularism and generalism fall short of the mark, and they leave many other forms unmentioned. For instance, Principled Particularism accepts the claim that agents need principles to think and judge in moral terms, and to distinguish right from wrong. There are also different kinds of moral principles (e.g. default, defeasible, exceptionless) that may be accepted or rejected by both particularists and generalists. So how do we appropriately delineate between particularism and generalism? On the basis of an overview of the scholarly literature, there seem to be two central questions that help to distinguish particularists from generalists: (1) does the possibility of moral thought and judgement depend on the provision of a suitable supply of moral principles?; and (2) can moral principles codify the moral landscape in finite and manageable terms?. If one answers yes to both questions, one is some kind of generalist; whereas if one answers no to both questions, one is some kind of particularist. The tricky part in this debate is when one answers yes to the first but no to the second. Here, one can be either a particularist or a generalist. This will of course depend on specific details such as, for instance, the types of principles one wants to defend and what kind of codification one has in mind. With these

considerations in the background, I will now delineate the various forms of particularism and generalism.⁷

McKeever and Ridge (2006) have provided a very useful taxonomy for the different forms of particularism. First, the most extreme variety of particularism is Principle Eliminativism (PE). PE claims “there are no true moral principles”.⁸ PE holds that moral thought and judgement have nothing to do with the application of moral principles to moral situations, and that the moral exemplar should not be understood as the person of principle.⁹ Second, Principle Skepticism Particularism (PSP) claims “there is no good reason to think there are any true moral principles”.¹⁰ The idea here is that, given the complexity of the moral landscape, the constant discovery of new exceptions, and perhaps the generalized nature of moral principles, moral principles are just not cut-out to maintain their truth-status. Third, Principled Particularism (PP) claims that “any finite set of moral principles will be insufficient to capture all moral truths”.¹¹ This is a claim that is defended by Holton (2002). In this thesis, I will also be defending this claim. The strength of this view is that it takes into account the complexity of the moral landscape and the limitation of moral principles, while also recognizing the necessity and utility of moral principles insofar as they do capture *some* moral truths. Fourth, we have Principle Abstinence Particularism (PAP) which claims that “we ought not rely upon moral principles”.¹² PAP evokes the idea that principles may distort our perception of moral cases, and for this reason, should not be relied on. Proponents of

⁷ If one answers no to the first and yes to the second, one is defending a rather unlikely view. That is, if moral thought and judgement do not depend on moral principles, then it seems unnecessary to codify moral truths (with admittedly unnecessary principles) that can be discovered in other, more natural or efficient ways (e.g. by means of moral instincts, intuitions, and/or a special kind of discernment). Thus, answering no to the first question takes away a central reason for answering yes to the second.

⁸ McKeever and Ridge, *Principled Ethics*, 15.

⁹ “Moral Particularism.” In *The Stanford Encyclopedia of Philosophy*, Winter Edition 2017.

¹⁰ McKeever and Ridge, *Principled Ethics*, 15.

¹¹ *Ibid.*, 16.

¹² *Ibid.*, 17.

PAP may be defenders of PE and PSP, but not necessarily. A naïve criticism of PAP is that it contradicts itself. A critic might argue that no matter how much we would want to rely on the claim “we ought not rely upon moral principles”, we would inevitably be relying on at least one moral principle. But this assumes that the claim is itself a moral directive when really it is an epistemic and practical one. And finally, fifth, Anti-Transcendental Particularism (ATP) claims “the possibility of moral thought and judgement does not depend on the provision of a suitable supply of moral principles”.¹³ This is Jonathan Dancy’s central thesis in his *Ethics Without Principles*.

As it stands in Dancy (2004), this claim has, in my view, three weaknesses, all of which consist in failures to make key distinctions. First, the claim does not adequately distinguish between exceptionless and defeasible moral principles. Second, it does not distinguish between determinate and indeterminate sets of principles. Here, the question at issue is: does the possibility of moral thought and judgement depend on a specific/determinate set of moral principles, or can a variety of different sets make moral thought and judgement possible for different individuals in different moral cultures during the complicated and culturally-dependent process of moral learning? This question, of course, treads on moral psychology and cultural relativism, both of which will be discussed in chapter 3. Third, Dancy’s claim does not say what the possibility of moral thought and judgement *does* depend on. The questions at issue here are: (1) what does the possibility of moral thought and judgement depend on?; and (2) doesn’t the possibility of moral thought and judgement depend on more than just moral principles? These three lacunae in Dancy’s claim will be explained in more detail in chapter 3.

McKeever and Ridge’s taxonomy is also useful for situating other particularists in the debate. On the extreme end, McKeever and Ridge seem to think that John McDowell (1998)

¹³ Ibid., 19.

endorses a form of Principle Eliminativism, for he suggests that moral cases resist capture in any general set of moral principles.¹⁴ Margaret Olivia Little is also an important scholar in this debate for many reasons. First, some of what she has claimed, according to McKeever and Ridge, could be interpreted as an endorsement of Principle Eliminativism. For example, she claims “we have reason to doubt that there are any moral principles, even very complicated ones, capable of codifying the moral landscape”.¹⁵ Second, the fact that she doubts that there are any moral principles and that the moral landscape can be codified by means of such principles could make her a Principle Skeptic. However, at the end of her (2000) she makes a very interesting remark, one that I think captures her central point: “While particularism dethrones moral generalizations, though, it hardly exiles them. We will think it does only when the lesson of holism is misidentified; the enemy its objections target is not generalization, but codification”¹⁶. The suggestion here is that holism does not rule out the possibility of moral principles (moral generalizations). Thus, particularism, which is supported by holism, cannot eliminate moral principles, but can only dethrone them. *Contra* McKeever and Ridge, Little is clearly not a Principle Eliminativist, for she admits the existence of functional moral generalizations. What she and other particularists are critiquing is how generalists use those moral generalizations, i.e. to codify the moral landscape. Little’s point is that moral principles do exist, they just don’t occupy the central throne of morality.

Lance and Little (2004) have also made important contributions by developing what they call “defeasible generalizations”, in their attempt to find a middle ground between particularism and generalism. Defeasible generalizations, they propose, can be revised and reformulated to

¹⁴ John McDowell. *Mind, Value and Reality*. Cambridge, MA: Harvard University Press, 1998, 149, quoted in McKeever and Ridge, *Principled Ethics*, 15.

¹⁵ Margaret Olivia Little. “Moral Generalities Revisited.” In *Moral Particularism*, edited by Brad Hooker and Margaret Little, 276–304. New York: Oxford University Press, 2000, 277 quoted in *Ibid.*, *Principled Ethics*, 15.

¹⁶ Little, “Moral Generalities Revisited”, 304.

include exceptional cases in their scope. Both Dancy (2004) and McKeever and Ridge (2006) have responded to Lance and Little (2004). Dancy argues that their defeasible generalizations fail to determine the nature of exceptional cases, to tell us why they are the way they are, and to act as reliable guides to action when confronted with them.¹⁷ McKeever and Ridge, for their part, argue that Lance and Little's defeasible generalizations only avoid "flattening the normative landscape" insofar as they endorse generalist terminology.¹⁸ Lance and Little (2008) is another attempt at finding a middle ground between particularism and generalism. They call their view "Deep Moral Contextualism". Here, they maintain a view of defeasible generalizations, and reject exceptionless generalizations altogether. Lance and Little reject exceptionless moral principles and instead endorse defeasible ones. I will discuss their views in more detail in chapter 3. David McNaughton (2000), for his part, is a Principle Abstinence Particularist, for he is renowned for having claimed that "moral principles are at best useless, and at worst a hindrance, in trying to find out which is the right action"¹⁹. The view being implied here is that the focus on rational principles to interpret moral cases may actually distort our perception of those cases. Richard Holton (2002), for his part, is a Principled Particularist, for he argues that no finite set of moral principles will be able to codify the moral landscape.²⁰ He is not a Principle Eliminativist because he believes there are some true moral principles. He argues for what he calls "*That's it*" principles. In essence, these are principles that specify cases where there is a reason to do something, and where there are no other morally relevant features to disable that reason. In other words, there is a reason to *x* and since *That's it* (since there are no other morally relevant features to disable it), it is a reason in favour of *x*.²¹ In

¹⁷ Dancy, *Ethics Without Principles*, 117.

¹⁸ McKeever and Ridge, *Principled Ethics*, 64.

¹⁹ McNaughton, *Moral Vision*, 191.

²⁰ Richard Holton, "Particularism and Moral Theory: Principles and Particularisms." *Aristotelian Society: Supplementary Volume* Supp. no. 76 (2002): 191–209.

²¹ Holton, "Principles and Particularisms.", 199.

these situations, Holton believes these principles are what we might call exceptionless generalizations. However, he reminds us that there are simply not enough of these to codify the moral landscape.

On the generalist side, there are, I would say, three primary positions: Transcendental Generalism (TG), Contextual Generalism (CG), and Principle Absolutism (PA). TG includes McKeever and Ridge's brand of generalism which they call "generalism as a regulative ideal". It is a direct response to Dancy's particularism in that it argues that the possibility of moral thought and judgement *does* depend on the provision of a suitable supply of moral principles. More specifically, they argue that ordinary moral thought, insofar as it aspires to count as knowledge, presupposes the availability of a suitable supply of moral principles—namely a suitable supply of default moral principles.²² They go further by arguing that "morality also presupposes the availability of unhedged (exceptionless) moral principles, principles whose defeating and enabling conditions have been fully and finitely articulated in descriptive terms".²³ With these unhedged or exceptionless moral principles, they believe the entire moral landscape can be codified in finite and manageable terms. CG, as I conceive it, accepts the rarity of exceptionless principles, but grounds itself primarily on defeasible ones. It would be close to Lance and Little's Deep Moral Contextualism, except it would allow for some exceptionless generalizations. In addition, it would be able to perform context-dependent codification of the moral landscape. PA, for its part, is the extremist position on the generalist end of the spectrum. The idea here is that there is a specific set of exceptionless moral principles—or perhaps just one ultimate principle—that make moral thought and judgement possible and that codifies the entire moral landscape. In PA, it is *only* one specific set of principles that makes moral thought and judgement possible and that codifies the

²² McKeever and Ridge, *Principled Ethics*, 113.

²³ *Ibid.*, 177, the parenthesis is my own addition.

moral landscape, and nothing else. Constitutive Generalism, for instance, as it is described and rejected in McKeever and Ridge,²⁴ is a form of PA.

Besides McKeever and Ridge, it is worth mentioning two more scholars. First, we have Luke Robinson (2006) who argues that holism is itself a form of moral generalism.²⁵ This is an interesting but also radical claim, for it seems to conflate principles about how moral reasons behave with principles that tell us about right and wrong. Robinson also argues for what he calls “moral dispositionalism”. He argues that we have moral dispositions based on general moral facts such as “homicide is wrong”, and that these facts can explain why *particular* moral facts such as “*that* homicide is wrong” are true. He argues this is so because “for homicide to be wrong is for it to be dispositionally wrong and for a particular homicide to be occurrently wrong is for it to manifest a disposition to be wrong”²⁶. Thus, the moral disposition expressed by the general fact “homicide is wrong” has the same explanatory power as the general moral principle “homicide is wrong”. Ultimately, I will be arguing against Robinson’s conception of moral facts. I will argue that it is the particular moral facts that explain whether a moral generalization is true, rather than vice versa. In other words, our a priori knowledge that an act is dispositionally wrong is dependent on our having accumulated a sufficient range of a posteriori knowledge of particular wrong-making features that make that act generally wrong.²⁷

²⁴ See McKeever and Ridge (2006): “According to one dimension of the generalist tradition, moral principles are built into the very meaning of moral predicates. They are analytic truths, and thus anyone who is in fact competent with a given moral concept is (perhaps implicitly) committed to the associated principle that spells out the object to which the concept applies. On this view, certain moral principles are constitutive of moral thought and judgment; this view is called ‘constitutive generalism’.” (Abstract of chapter 5, Oxford Scholarship Online). McKeever and Ridge go on to say that a more direct strategy for arguing for Constitutive Generalism would be to argue for a “specific set of principles as constitutive of moral judgment” (97) wherein “moral predicates like ‘morally required’ can be analyzed into purely descriptive language” (97).

²⁵ Luke Robinson, “Moral Holism, Moral Generalism, and Moral Dispositionalism.” *Mind: A Quarterly Review of Philosophy* 115, no. 458 (2006): 353.

²⁶ Robinson, “Moral Holism, Moral Generalism, and Moral Dispositionalism.”, 352.

²⁷ See section 4.5, pages 111-113 for our arguments in favour of a view of moral reasons entitled “Core Holism”.

There is also Maike Albertzart (2013) who provides an original defense of Transcendental Generalism. He holds that the person of moral judgement is best understood as a person of principle.²⁸ To defend this claim, he develops a conception of moral principles called “internalised long-term commitments”. He critiques the polarisation of moral judgement and moral principles assumed by particularists, and argues that moral judgement actually depends on moral principles understood as internalised long-term commitments, for they are what shape our moral character and sensitivities.

In sum, the debate is wide-ranging. I am indebted to all the scholars mentioned above, including many more. Each of them will serve a purpose in this thesis. Our primary players, however, are Dancy (2004, 2008, 2017), Little (2000), and Holton (2002) on the particularist side; Lance and Little (2004, 2008) in the middle; and McKeever and Ridge (2006, 2008, 2016), Robinson (2006), and Albertzart (2013) on the generalist side. Below I have included a spectrum for this debate.

Principle Absolutism
Transcendental Generalism
Contextual Generalism
Deep Moral Contextualism
Principled Particularism
Anti-Transcendental Particularism
Principle Eliminativism

As we go up the spectrum, principles have a more prominent role in morality; as we go down, they progressively become less so.²⁹ Deep Moral Contextualism is the centrist position. This is the position where principles are officially dethroned from morality, and context becomes

²⁸ Maike Albertzart. “Principle-Based Moral Judgement.” *Ethical Theory and Moral Practice: An International Forum* 16, no. 2 (2013), 340.

²⁹ I have intentionally left out Principle Skepticism Particularism and Principle Abstinence Particularism for obvious reasons. The former is more of an attitude than a position, and the latter is an epistemic and practical directive, rather than a complete and workable theory. What is more, both of these are compatible with Principle Eliminativism and Anti-Transcendental Particularism.

the focal point. Below I have included a table of the different theories and how they answer key questions:

<u>Meta-ethical Theory</u>	<u>What view of reasons does the theory assume or claim to be compatible with?</u>	<u>Do moral thought and judgement depend on principles?</u>	<u>What types of principles do moral thought and judgement depend on?</u>	<u>Do moral thought and judgement depend on a determinate or indeterminate set of principles?</u>	<u>Do moral thought and judgement depend on factors other than principles?</u>	<u>What type of codification is advocated by the theory?</u>	<u>How much codification does the theory claim to achieve?</u>
Principle Absolutism	Atomism	Yes	Exceptionless	Determinate	No	Transcendental	Complete
Transcendental Generalism	Atomism or Hybridism	Yes	Exceptionless and Default	Determinate	Yes	Transcendental	Complete
Contextual Generalism	Hybridism or Holism	Yes	Exceptionless and Default	Determinate	Yes	Context-Dependent	General
Deep Moral Contextualism	Holism	Yes	Defeasible	?	Yes	Context-Dependent or N/A?	Particular, or codification is not possible?
Principled Particularism of Contingency	Holism	Yes	Exceptionless and Defeasible	Both	Yes	Context-Dependent	Particular
Holton's Principled Particularism	Holism	Yes	Exceptionless	?	Yes	?	Particular
Anti-Transcendental Particularism	Holism	No	N/A	N/A	Yes	N/A or Context-Dependent	Codification is not possible, or very limited
Principle Eliminativism	Holism	No	N/A	N/A	Yes	N/A	Codification is not possible

The terms and questions in this table will become clearer as we go along. For the moment, notice the lacunae in Deep Moral Contextualism (DMC) and Holton's Principled Particularism (PP). It is unclear whether Lance and Little (2008) argue that moral thought and judgement depend on a determinate or indeterminate set of defeasible principles. It is also unclear whether they believe codification is even possible, or if they support a view according to which particular codification is possible. As for Holton, we know he advocates for particular codification of the moral landscape when he claims that "different moral verdicts will be entailed by different sets of principles; but there is no one set that will entail them all".³⁰ However, it is unclear whether he

³⁰ Holton, "Principles and Particularisms", 195.

advocates for *transcendental* codification or *context-dependent* codification given that he supports holism (which, at face-value, lends itself better with context-dependent codification) and exceptionless “*That’s it*” principles (which lend themselves better to a transcendental theory).

My own position is a version of PP in which I combine elements from Holton’s version and Lance and Little’s DMC. By doing so, I hope to explain away the aforementioned lacunae. In my version of PP, I have bifurcated two considerations. First, that which codifies the moral landscape is not necessarily that which makes moral thought and judgement possible. Indeed, if one is a generalist, one either believes that moral principles codify a general part or the whole moral landscape, but this doesn’t necessarily imply that they make moral thought and judgement possible in the first place. My assumption throughout this thesis will be that there are multiple factors that make moral thought and judgement possible (e.g. moral teachers and exemplars, evolutionarily entrenched moral intuitions, instincts, an agent’s always-developing ability of discernment, and domain-general learning capacities, to name a few), not just moral principles. Second, that upon which moral thought and judgement depend do not automatically codify all moral truths. As I mentioned earlier, the dependence question and the codification question should be treated individually, for an affirmative answer to the first does not necessarily entail an affirmative answer to the second, and vice versa; and same thing with negative answers. As such, my brand of PP will make two claims: (1) that moral principles are necessary but insufficient for making moral thought and judgement possible; and (2) that moral principles are capable of codifying particular parts (but not the entirety) of the moral landscape in finite and manageable terms.

With these ideas in the background, I believe the central criterion for what makes one a particularist or a generalist is one’s stance on the breadth of moral principles’ codification. If one

believes that complete or general codification of the moral landscape is possible by means of principles, one is *necessarily* a generalist; whereas if one believes that only particular parts (or no parts) of the moral landscape are codifiable, then one is *necessarily* a particularist. Indeed, if one believes that moral thought and judgement do not depend on principles, then one is very likely a particularist. But this is not a guarantee. On the other hand, if one believes that moral thought and judgement *do* depend on principles, it is not at all clear whether one is a generalist or a particularist. For example, Holton claims that moral thought and judgement depend on *That's it* principles, but he is a particularist. Hence, one's answer to the dependence question is insufficient for determining with certainty whether one is a generalist or a particularist. Therefore, *contra* Dancy, I do not believe the answer to the dependence question is what necessarily divides generalists and particularists. As I see it, what divides generalists and particularists is what they believe moral principles can do, how important or unimportant they are for moral theory, deliberation, and action-guidance, and how much of the moral throne they occupy. In short, what seems to divide generalists and particularists is how much of the moral landscape they believe moral principles can codify.

1.2. Holism and Atomism About Reasons

Much of the particularism-generalism debate revolves around the holism-atomism debate (although some scholars do not fully agree with this, as we will see). Whether holism leads to particularism or generalism, or whether atomism leads to generalism, or whether the holism-atomism debate is irrelevant to the particularism-generalism debate, is much debated in the scholarship. Despite these disagreements, both holism and atomism about reasons are nevertheless

omnipresent in the particularism-generalism literature. Therefore, for our purposes, it is worthwhile to define and explain them.

According to Dancy, reasons-holism is the claim that “a feature that is a reason in one case may be no reason at all, or an opposite reason, in another”.³¹ Dancy distinguishes reasons-holism from reasons-atomism according to which “a feature that is a reason in one case must remain a reason, and retain the same polarity, in any other”.³² To better highlight the difference between holism and atomism about reasons, consider the following illustration.

Suppose that a consideration (C) may be a reason (R) that favours an action (A). Reasons-holism holds that whether C is an R to A depends on more than C and A; it depends on other features of the situation in which C and A figure. Reasons-atomism, by contrast, claims that whether C is an R to A depends only on the nature of C and A.³³

For example, suppose that saving a life may be a reason that favours lying. Reasons-holism holds that whether saving a life is a reason to lie depends on more than the fact that lying will save a life and the act of lying; it depends on other features of the situation in which saving a life and lying figure. Reasons-atomism, by contrast, claims that whether saving a life is a reason to lie depends only on the nature of saving a life and the nature of lying. For an atomist, life-saving is always right-making and lying is always wrong-making. The question is: which consideration is more important? Suppose we are saving the life of Adolph Hitler. Indeed, it might very well be better to tell the truth and let Hitler die. In such a case, lying would arguably be wrong-making because it would save the life of a man who intends on murdering innocent people. The holist would argue that saving his life would also be wrong-making for that same reason. The atomist, however, would argue that saving his life is still right-making (because life-saving is *invariantly*

³¹ Dancy, *Ethics Without Principles*, 7.

³² Ibid., 7.

³³ Thank you to my thesis supervisor Andrew Sneddon for this more comprehensive explanation of reasons-holism.

right-making) but that this consideration is trumped or outweighed by the wrong-makingness of lying and the right-makingness of telling the truth. But now suppose the man whose life might be saved is Jewish and he is hiding in our house, and the person we are about to lie or tell the truth to is a Nazi officer. The atomist holds that lying is always a reason against, and that life-saving is always a reason in favour. Therefore, the atomist would argue that while lying is always wrong-making, saving the Jewish man's life against a tyrannical and genocidal government is more important. Thus, the latter consideration outweighs the former, and lying becomes in this case morally permissible even though it is always morally wrong. The holist, however, would argue that lying is right-making in this case. First, the holist might argue that lying to a Nazi officer takes considerable courage since we would be acting against the law and putting ourselves at risk of imprisonment or even death. Second, lying in this case is strictly intended to produce a positive outcome, namely to save the life of a human being against a genocidal government. These two reasons *make* the act of lying right, not wrong or neutral. So for the holist, while lying might be wrong-making *by default*, it is not always or *invariantly* wrong-making, for there may always be other features in the situation that can cause its moral status to change. The same thing can be said about all other reasons, including life-saving. The atomist, by contrast, argues that moral reasons do not change valences; rather, they always remain the reasons that they are (either for or against). For the atomist, since all reasons have the same moral contribution, we have to determine which reason outweighs the others in a given case. For the holist, however, although seeing which reason outweighs the others in a given case is always important, a deeper look at context is also important

to see in which ways these reasons are contributing to the case at hand; for they may not always contribute in the same way.³⁴

Holism and atomism are defined along similar lines by other scholars as well. According to Katie McShane, reasons-holism is a claim that explanation goes from whole to parts, and reasons-atomism a claim that explanation goes from parts to whole. McShane juxtaposes reasons-holism with meaning-holism in the philosophy of language. Just as the meaning-holist holds that the meaning of words or sentences is determined by their linguistic context, so the reasons-holist claims that the normative status of features in a given case is determined by the context in which they appear. These are ontological claims because they assert that there is no such thing as meaning or normative force independently of context. They are also exclusive claims because they tell us to look to sentences-in-languages and features-in-situations, not merely to sentences or features in isolation. Finally, they are also assertions about whole-to-part explanations because the implication is that the linguistic/situational context *explains why* a particular sentence/feature has the meaning/normative force that it does.³⁵ According to McKeever and Ridge, reasons-holism holds that “whether some consideration is a reason may depend upon the presence (or absence) of various background conditions”;³⁶ while reasons-atomism, by contrast, claims that “a consideration that is a reason in one setting will similarly be a reason in any other setting”.³⁷

³⁴ There are certain cases where a given consideration can provide reasons holistically while also always being invariant in its contribution to a given case. This is often the case with thick moral concepts such as “justice”. Justice can generate both reasons for and against, but it will always generate a reason for even if that reason may be outweighed by a reason against in a given case. This possibility is consistent with reasons-holism. See Andrew Sneddon (2010) for a detailed discussion on the matter. For a discussion on what are called “contributory reasons”, see Dancy (2017).

³⁵ Katie McShane. “Holism.” In *International Encyclopedia of Ethics*. Hoboken, New Jersey: John Wiley & Sons, Ltd., 2013, 3.

³⁶ Sean McKeever and Michael Ridge. “Reasons: Holism and Atomism.” In *International Encyclopedia of Ethics*. Hoboken, New Jersey: John Wiley & Sons, Ltd., 2013, 1.

³⁷ McKeever and Ridge. “Reasons: Holism and Atomism.”, 1.

Simply put, under the holistic view, reasons are context-sensitive, meaning that other features in the situation may change their moral valence; whereas under the atomistic view, reasons behave according to their nature, such that they will have the same moral valence in all contexts in which they appear. In the holistic world, the vast majority of reasons are variant, though there may be some invariant ones,³⁸ whereas in the atomistic world, all reasons are *necessarily* invariant. If holism is true, then the landscape of moral reasons is not very hospitable for exceptionless moral principles to the extent that such principles attempt to specify invariant reasons. In chapter 2, I will defend reasons-holism by examining the strongest candidates for invariant reasons and demonstrating why most of them are not truly invariant.

Most scholars in the particularism-generalism debate either support holism³⁹ or they believe that the holism-atomism debate is orthogonal to the particularism-generalism debate.⁴⁰ The most prominent generalists have rejected the inference from reasons-holism to particularism on the grounds that the former leaves open the possibility that the behaviour of reasons, enablers, disablers, intensifiers, and attenuators is codifiable.⁴¹ In short, generalists claim that holism leaves open the possibility for a context-dependent codification of the moral landscape by means of moral principles. For this reason, they argue that holism cannot support particularism which claims that the complete codification of the moral landscape by means of principles is impossible. In my view, this argument has an important flaw. It assumes that particularists are against all forms of codification, and that codification can only be achieved under a generalist theory. Both of these assumptions are untrue. Holton's Principled Particularism, for example, allows for a partial

³⁸ According to Dancy, although the vast majority of reasons in the holistic world are variant, it is not inconsistent with holism to allow for the possibility of some invariant reasons. Some reasons are necessarily invariant because of the particular reasons that they are (i.e. their specific content). See Dancy (2004: 77-78).

³⁹ Dancy (2004, 2017), Little (2000), Lance and Little (2004, 2008) to name a few.

⁴⁰ See McKeever and Ridge (2005: 93-103; 2006: 27-45) and Albrecht (2013: 342, footnote 6).

⁴¹ Väyrynen (2004), McKeever and Ridge (2005, 2006, 2016), and Albrecht (2013).

codification of the moral landscape. PP simply admits that principles are insufficient for codifying the *entire* moral landscape.⁴² So while holism rejects the *complete* codification of the moral landscape to the extent that such an endeavour would require exceptionless moral principles—principles that seldom exist in a reasons-holistic world—it is not against moral generalizations themselves, nor is it against partial codification.⁴³ In this thesis, I will argue that holism may not lead directly to particularism, but it does seem to reject Principle Absolutism and Transcendental Generalism (McKeever and Ridge’s brand of generalism) insofar as these latter are founded on exceptionless moral principles that can supposedly codify all moral truths. In chapter 2, we will demonstrate how many of the strongest supposedly invariant reasons (specified by exceptionless principles) are not truly invariant insofar as they can be shown to vary in certain contexts. The only version of generalism that might be compatible with holism seems to be Contextual Generalism, which of course calls for a context-dependent codification of the moral landscape. However, in chapter 4, we will provide arguments against CG. On the other side of the spectrum, holism is compatible with all forms of particularism. On the whole, then, holism is certainly not entirely orthogonal to the particularism-generalism debate. First, it rejects two forms of generalism; and second, given that it rejects two forms of generalism and is compatible with all forms of particularism, it seems to favour particularism over generalism.

1.3. Different Types of Moral Principles

McKeever and Ridge distinguish between what they call “hedged” and “unhedged” principles. Hedged principles include certain qualifications in their formulations, such as *ceteris*

⁴² See Holton (2002: 193).

⁴³ See Little (2000: 304).

paribus, *prima facie*, and *pro tanto* clauses.⁴⁴ For example, “all things being equal, lying is wrong” is a hedged principle.⁴⁵ Hedged principles are incapable of codifying morality because they leave out exceptional cases from their scope. Consequently, they also cannot guide the moral agent in novel circumstances. McKeever and Ridge defend a specific type of hedged principle they call a “default principle”. Default principles claim “that a given action has a given moral property only if no further feature of the situation explains why it does not”.⁴⁶ Unhedged principles, for their part, do not have such qualifications. They are exceptionless generalizations expressed by necessarily true propositions. They are intended to include all possible exceptions to a given norm into their scope. As such, they are, in theory, capable of codifying all of morality.⁴⁷ Hedged or default principles, however, do not specify features that might serve as potential disabling conditions or countervailing reasons. According to McKeever and Ridge, we move from hedged to unhedged principles when we begin adding potential disabling conditions and countervailing reasons in the antecedent of the principle.⁴⁸

McKeever and Ridge also distinguish between principles qua standards, guides, action-guiding standards, algorithmic decision-procedures, and moral truth-makers.⁴⁹ Let’s briefly examine the first two of these. On the one hand, principles qua standards are defined as “entirely exceptionless generalizations that provide the truth-conditions for the application of a moral concept”.⁵⁰ Principles qua standards claim to provide explanations of *why* certain actions are right

⁴⁴ McKeever and Ridge, *Principled Ethics*, 21.

⁴⁵ Sean McKeever and Michael Ridge. “Moral Particularism and Moral Generalism.” In *The Stanford Encyclopedia of Philosophy*. Winter 2016 Edition, see section 2 ““Particularism” and “Generalism” are said in many ways”.

⁴⁶ McKeever and Ridge, *Principle Ethics*, “Abstract and Key Words”, 113. For the abstract, see the version of the book on Oxford Scholarship Online.

⁴⁷ *Ibid.*, 21.

⁴⁸ McKeever and Ridge, *Principled Ethics*, 118-119.

⁴⁹ *Ibid.*, 5-14.

⁵⁰ *Ibid.*, 7.

or wrong, why certain moral considerations are reasons with a specific valence and weight, and why a given character trait is a virtue or a vice.⁵¹ According to McKeever and Ridge, these principles should be understood as necessary truths because “any contingent principle would leave open the possibility of a deeper standard that explains why the contingent principle holds true in our world but not in others”.⁵² On the other hand, principles qua guides “purport to be well suited to guiding action”.⁵³ But according to McKeever and Ridge, they “need not provide entirely accurate application conditions for moral concepts”.⁵⁴ This is because they might not cover all cases and they might even lead the moral agent to error. The only criterion for a principle qua guide is that it provides useful direction to a “conscientious moral agent”.⁵⁵

If we juxtapose these two types of principles, the primary difference is that standards are necessary truths, whereas guides are only contingently true because the valence and weight of the moral reasons they specify depend on contextual features. This difference is so because the former principles attempt to specify invariant reasons, whereas the latter attempt to specify variant ones. The reason why principles qua guides might have false implications and thus may lead the moral agent to error is because they are generalizations attempting to specify reasons that may vary from case to case.⁵⁶ This brings me to an important distinction that will come up often in this thesis: that between exceptionless and defeasible moral principles:

An exceptionless moral principle is a moral generalization expressed by a necessarily true proposition that specifies an invariant reason or set of reasons.

⁵¹ McKeever and Ridge, “Moral Particularism and Moral Generalism”, section 2 ““Particularism” and “Generalism” are said in many ways”.

⁵² McKeever and Ridge, *Principled Ethics*, 7.

⁵³ McKeever and Ridge, “Moral Particularism and Moral Generalism”, section 2 ““Particularism” and “Generalism” are said in many ways”.

⁵⁴ McKeever and Ridge, *Principled Ethics*, 8.

⁵⁵ *Ibid.*, 8.

⁵⁶ Note that if a principle qua guide specifies an invariant reason it is because it is also a standard; hence, a principle qua action-guiding standard.

A defeasible moral principle is a moral generalization expressed by a contingently true proposition that specifies a variant reason or set of reasons.

It is important to make this distinction when trying to determine whether moral principles make moral thought and judgement possible. We ought to ask: what types of principles make moral thought and judgement possible? Do we need exceptionless principles? Or are defeasible ones sufficient for the task? This distinction is also important when considering whether reasons-holism or reasons-atomism is true. If reasons-holism is true, then most moral reasons are variant. If reasons-atomism is true, then all reasons are necessarily invariant. In our case, we will be defending reasons-holism, therefore exceptionless principles will be hard to come by. Indeed, a reasons-holistic world in which the variability of reasons is the norm is largely inhospitable to exceptionless principles that seek to specify invariant reasons. Nonetheless, as it turns out, in chapter 2 we will discover at least two exceptionless principles. Here, we will argue that they are necessary but insufficient for making moral thought and judgement possible. This means that we will need to turn to defeasible principles to see whether they can complete the picture of morality. Defeasible principles, as we shall see, will constitute the bulk of morality. In chapter 3, I will argue that defeasible principles are also necessary but insufficient for making moral thought and judgement possible. Such things as moral exemplars, moral teachers, and evolutionarily entrenched moral intuitions, instincts, and domain-general learning capacities are also necessary.

Chapter 2

2.1. Introduction

Are exceptionless moral principles sufficient to ground the possibility of moral thought and judgement? I will argue that they are not because of reasons-holism. Since this question depends

on the nature of moral reasons, we will have to explain why reasons-holism is true and reasons-atomism is false. In order to do so, we will need to determine whether reasons in the moral landscape are mostly variant or invariant. If variance is the norm and invariance the exception, it is very likely that exceptionless principles (which purport to specify invariant reasons) will not be able to make moral thought and judgement possible on their own.

Reasons-atomism, which claims that all reasons are necessarily invariant, already seems untrue at first glance. For instance, that lying is always wrong-making seems clearly untrue when we consider for example the game of Diplomacy in which lying is the point of the game; or again when we are protecting Jews and are suddenly visited by the Nazi Gestapo to whom the truth is not owed.⁵⁷ That stealing is always wrong-making also seems untrue in cases where our military steals weapons of mass destruction from an enemy who intends to use them to commit the mass murder of innocent lives. That killing is always wrong also seems untrue in cases of self-defence. Finally, that telling the truth is always right also seems untrue when we take into account those cases where telling the truth causes more harm than good. At first glance, most reasons do seem to be variant rather than invariant. But in what follows, I will take on the strongest candidates for invariant reasons and see whether or not they are truly invariant. The ones just listed are strong, but they are clearly not the strongest. In this chapter, I will argue that most supposedly invariant reasons are not truly invariant by pointing to exceptional cases where they can be shown to vary. As it turns out, we will land upon two invariant reasons, but argue that they are nevertheless insufficient for making moral thought and judgement possible.

In what follows, I will critically examine five supposedly universal principles and demonstrate why the reasons they specify are not truly invariant. They are: (1) Kant's Categorical

⁵⁷ Examples taken from Lance and Little (2008: 53).

Imperative; (2) Bentham and Mill's Principle of Utility; (3) Causing Gratuitous Pain on the Unwilling; (4) the belief that one must have a good reason to do actions that cause harm (specified by a principle we will call The Principle of Special Reason to do Harm); and (5) Richard Holton's *That's it!* Principle. The supposedly invariant reasons specified in the first three of these principles will be shown to vary. As for the last two, I will concede that they are invariant reasons, but argue that their specific contents make them incapable of grounding the possibility of moral thought and judgement on their own. Therefore, by the end of this chapter, I will argue that exceptionless moral principles are not sufficient to make moral thought and judgement possible.

2.2. The Categorical Imperative

As is well-known, the Categorical Imperative (CI) has multiple formulations. For our purposes, we will only concern ourselves with the formula of universality and the formula of humanity. They read as follows:

Act only in accordance with that maxim through which you can at the same time will that it become a universal law.⁵⁸

Act that you use humanity, whether in your own person or in the person of any other, always at the same time as an end, never merely as a means.⁵⁹

One way of interpreting the first formulation (and this will be the interpretation that we will adhere to throughout)⁶⁰ is that it specifies the following supposedly invariant reason: *an act is right if it can become a universal moral law*. The second formulation is more straightforward. It

⁵⁸ Immanuel Kant. *Groundwork of the Metaphysics of Morals: Translated and Edited by Mary Gregor*. United Kingdom: Cambridge University Press, 1997, 31 (4:421).

⁵⁹ Kant. *Groundwork of the Metaphysics of Morals*, 38 (4:429).

⁶⁰ We won't explore other valid interpretations since it is not necessary for the purposes of this thesis.

contains two supposedly invariant reasons, both of which imply each other: *treating people as ends in themselves is right*; and *treating people merely as means is wrong*. In what follows, I will demonstrate how none of these three supposedly invariant reasons are invariant.

The idea behind the first is that whenever you do something, you should make sure that it would at least be *possible* for everybody else to act in the same way. Consider the act of breaking a promise along with the maxim “if you want joy, break a promise”. Kant’s thought is that if we universalize this maxim, it will lead to contradiction. The contradiction here is a conceptual one. If one’s false promise maxim were universalized, then simultaneously the grounds of possibility of thinking about promises and making or breaking them would disappear. Thus, the conceptual implications yield a contradiction. In the world of actions, if everyone starts breaking their promises, the very idea of promising will lose its credibility, and no one will trust anyone. In the final analysis, the institution of promising would collapse, and it would be impossible for anyone to make promises. Hence, the original goal of wanting joy by breaking a promise would not be possible because the institution of promising would not exist. The same logic applies to other acts such as cheating and stealing. The point is that the very possibility of a person’s doing these acts requires that most people, most of the time, do not act in these ways; this is what makes them wrong acts.

But we do not here have an invariant reason. First, the fact that the institution of promising depends on most people keeping and not breaking their promises most of the time is not that which makes keeping a promise right or breaking a promise wrong. What makes keeping a promise right are the right-making features that count in its favour, such as building trust with someone, making them happy, and so on. On the other hand, what makes breaking a promise wrong is the potential and/or actual harm being done both to the deceiver and to the people being deceived. Therefore,

that the institution of promising depends on most people keeping and not breaking their promises most of the time does not yield a universal moral law, but only a general truth.

Second, there are certain cases where breaking a promise is the right thing to do—cases whose consequences will not lead to the collapsing of the institution of promising if the acts that produced them were universalized. This is because we would be universalizing the exceptions to the rule without contradicting its general truth. For example, suppose I promised to attend an important meeting with my boss, but on my way there I get a call that my mother was rushed to the hospital. Many such examples can be related. So long as there are good reasons to break a promise the institution of promising will not collapse.

Third, there seems to be a fallacious underlying assumption at work here, namely that “if it is wrong for someone to *always* break their promises, then it is right for someone to *always* keep their promises”. Indeed, it is true that always breaking one’s promises is wrong. But this is unsurprising. If you always break your promises—even if in some cases you have good reasons to do so—arguably you will be wrong in most cases. Suppose that in your lifetime you broke all one million of your promises, and that you were only right to break ten thousand of them. Here, approaching the world with the imperative “always break your promises” is certainly not a good general rule-of-thumb. Therefore, it would be wrong for one to adhere to it. But this does not mean that the opposite imperative (viz. always keep your promises) is universally right. Indeed, you would be wrong to do so ten thousand times. Undoubtedly, our principles must be informed by the context in which we seek to apply them. And surely, following a principle we consider to be universally true *merely* because of its specific content (and without taking context into account) constitutes its own form of immorality.

Finally, this supposedly invariant reason for action—namely that *an act is right if it can become a universal moral law*—is by its own admission insufficient to maintain its status as an invariant reason without there already existing a set of invariant reasons for action in the moral landscape. This is because the truth of this principle depends on there being other invariant reasons for action in the field of moral possibilities to confirm its validity. In other words, in order to confirm that it is in fact invariantly true, we must study individual actions to see if there are invariant reasons that support them. Thus, this principle's supposed invariance is contingent upon proving the existence of other invariant reasons. But, insofar as it does not itself prove this, to claim that it is itself an invariant reason would be to beg the question.

Moving on now to the next supposedly invariant reasons, namely: *treating people as ends in themselves is right* and *treating people merely as means is wrong*. These supposedly invariant reasons are also not invariant. There are cases where treating people merely as means is right-making, and where treating people as ends in themselves is wrong-making.

Let's discuss Kant's arguments in regards to the classic "murderer at the door" scenario.⁶¹ Should we lie or tell the truth about the whereabouts of our friend to someone who intends on murdering them? For Kant, given that we must always treat others as ends in themselves and never merely as means, we are morally required to always tell the truth. Lying, under any circumstance, is morally wrong because it treats others merely as means. Although lying in this case would stem from altruistic motives, it may have unexpectedly harmful consequences that we would be responsible for, and liable to, by law. For example, if we lie by saying our friend is not at home,

⁶¹ Immanuel Kant. "On a Supposed Right to Lie from Altruistic Motives", in *Kant's Critique of Practical Reason and Other Writings in Moral Philosophy*, trans. and ed. Lewis White Beck, 346-350. Chicago: University of Chicago Press, 1949.

but it happens that our friend slipped out without our knowing, the murderer might very well find our friend later on and murder them anyway. In such a case, Kant argues, we might be rightly accused of the wrongful death of our friend; whereas if we had told the truth, we could have then called for help and the murderer might have been apprehended by the neighbours.⁶² According to Kant, if we tell the truth and the murderer manages to murder our friend, we are entirely blameless because we are following our universally binding duty. Thus, for Kant, we should only act based on what we know is right, and always avoid what we know is wrong; we should never act based on a known wrong for the sake of a supposedly better outcome. For Kant, even if we could be certain that lying would produce an extremely positive outcome, this fact would be irrelevant to whether or not the act in question is right or wrong.⁶³ What makes lying wrong is the nature of the act of lying itself: it involves treating people merely as means and not as ends in themselves. If we start making exceptions to our universal principles based merely on inductive and contingent possibilities rather than deductive and necessary certainties, our principles will become “uncertain and useless”,⁶⁴ for it is written into the definition of a principle that it ought to be universal.⁶⁵ For these reasons, according to Kant, we should always tell the truth and treat others as ends in themselves and never merely as means.

⁶² Kant, “On a Supposed Right to Lie”, 347.

⁶³ To be clear, Kant’s moral argument (i.e. that we have a duty to always tell the truth) does not depend on and is not reinforced by the unpredictability of the outcomes of lying or telling the truth. That lying is wrong does not depend on the negative consequences it might or will produce. By the same token, that telling the truth is right does not depend on the positive consequences it might or will produce. Rather, for Kant, lying is always wrong in itself regardless of its consequences in the world. Kant only refers to the possible outcomes that lying or telling the truth might produce in order show that they are often unpredictable and, therefore, should never be relied on as justification to act against our duties. This is why, according to Kant, we should stick with what we know to be necessarily, universally, and deductively true (i.e. our universal duty to tell the truth and to treat everyone as ends in themselves), and simply avoid dabbling in the analysis of consequences and outcomes altogether.

⁶⁴ Ibid., 347.

⁶⁵ Ibid., 350. We will return to this claim shortly.

For Kant, not only is treating people as ends in themselves always right-making, it is always morally right *overall* for all people in all cases and, therefore, is always morally required. By the same token, treating people merely as means is not only always wrong-making, it is wrong universally and we are absolutely forbidden from doing it. This makes Kant a strong atomist about reasons. The weak atomist would claim that while treating people as ends in themselves is always right-making and treating people merely as means is always wrong-making, they are not always right and wrong overall; for there may be some cases where there is an overriding moral consideration at play. Thus, for the weak atomist, features that are reasons for belief and action always *contribute* in the same way to a given case, but they may be defeated by other competing features. For example, the weak atomist might argue that we should indeed lie in order to save a life. By contrast, for the strong atomist, such as Kant, all reasons have a universal moral valence and therefore can never be overridden. As I explained in section 1.2, I believe both of these views are false.

In response to atomism, all we have to show is that there are cases in which other factors override our duty to tell the truth, and that lying can itself be right-making. The “murderer at the door” example is such a case. There are at least two reasons to support the act of lying in this case. First, assuming we have good epistemic reasons to believe the individual at the door is in fact a murderer⁶⁶ who intends on murdering our friend,⁶⁷ these facts are themselves reasons for us to believe that the murderer will act as he intends to. Of course, we cannot be *completely certain* that he intends to murder our friend, but we do have *strong reasons to believe* that he does. This inductive knowledge arguably generates a moral obligation to lie in order to prevent the murderer

⁶⁶ In the scenario Kant is discussing, the person at the door is in fact a murderer. It is considered to be a known fact, rather than a probable conjecture. See p. 346.

⁶⁷ It is stated in the scenario that the murderer is “pursuing” our friend. See p. 346.

from achieving his intended act and outcome—both of which are of negative moral valence.⁶⁸ Second, the fact that we are lying in order to save a life—an act of positive moral valence⁶⁹—is also a reason in favour of lying. Even though lying is, other things being equal, inherently wrong, that it might very well save a life arguably trumps the alternative of telling a truth that will result in death. The assumption that truth-telling is always morally superior to life-saving (i.e. insofar as the former is a perfect duty and the latter is an imperfect duty) is unwarranted. In what follows, I will provide a case in support of this claim.

The problem with Kant's view is that even in cases where it is overwhelmingly reasonable to believe that the individual in question will do what he intends to do (e.g. murder the intended victim), we still can never be certain. Understood as such, Kant's view is that we should only act based on what we know is (supposedly) necessarily true (e.g. telling the truth is universally right). So, for Kant, any inductive moral argument in favour of some imperfect duty—no matter how strong and compelling—can never override any of the perfect duties because these latter are based on necessary, deductive truths. Consequently, Kant's position, insofar as it considers certain moral duties as universally and necessarily true and binding, leaves us with intuitively implausible answers to scenarios such as the “murderer at the door”. Kant's view requires us to follow our perfect duty in all circumstances—even in cases in which we have overwhelming evidence to predict a specific negative outcome we can successfully avoid—and thus argues that breaking this duty is absolutely prohibited in all circumstances. Not only does Kant's view produce intuitively implausible solutions to these types of scenarios, it is also neglecting an important part of rationality, namely its connection to the empirical. In my view, both our abstract principles and

⁶⁸ By “negative moral valence” I mean acts that are by default, other things being equal, inherently wrong.

⁶⁹ By the same token, by “positive moral valence” I mean acts that are by default, other things being equal, inherently right.

the empirical context in which they are applied should be considered simultaneously and viewed as equally important. If rationalism and empiricism should be synthesized in the theory of knowledge, a similar synthesis should occur in the moral realm as well.

In my view, all relevant facts, both moral and non-moral, should be seriously considered in all cases. When we have, for example, overwhelming reasons to believe that a particular negative outcome will occur, actions to avoid such an outcome—especially ones that will bring about a positive outcome—should be seriously considered as viable contenders. I think a holistic sensitivity to intentions, the intrinsic nature of our actions and duties, as well as to external factors and probable outcomes, is important for any ethical theory. Theories that neglect one or more of these facets, such as deontology in regards to outcomes, tend to give implausible and often counterintuitive answers to certain moral dilemmas.⁷⁰ In my view, lying to the murderer at the door is right-making for three reasons. First, because we have overwhelming evidence to believe the murderer at the door is in fact a murderer and will likely do what he intends to do (i.e. murder our friend). Second, because of the fact that lying will be done with good intentions (i.e. for the specific purpose of saving a life). Third, because of the fact that lying to the murderer will likely produce the positive outcome of saving a life.

Kant's claim that if we admit of any exceptions to our principle it will become "uncertain and useless" is also unwarranted. It appears to be predicated on the logical fallacy of the slippery slope according to which other similar exceptions will continually be made until we reach the point where our principles become self-contradictory:

If one is asked whether he intends to speak truthfully in a statement that he is about to make and does not receive the question with indignation at the suspicion it expressed that he might be a liar, but rather asks permission to consider possible exceptions, that person is already potentially a liar. That is because he shows that he does not acknowledge

⁷⁰ We will discuss our holistic view in the coming paragraphs.

truthfulness as an intrinsic duty but makes reservations with respect to a rule which does not permit any exception, inasmuch as any exception would directly contradict itself. All practical principles of right must contain rigorous truth, and the so-called "mediating principles" can contain only the more accurate definition of their application to actual cases (according to rules of policy), but they can never contain exceptions from the former. Such exceptions would nullify their universality, and that is precisely the reason that they are called principles.⁷¹

In my view, if we adopt reasons-holism and the notion of a default principle, we can avoid this problem. It seems to me entirely possible to have exception-laden principles while simultaneously holding on to our intrinsic duties. Reasons-holism, at least the version that I adopt, claims that the normative status of features that are reasons for belief or action depend on *both* the internal status of those features and the other external features of the context. Before delving into this view, there are a few other views to mention, some of which we have already discussed.

First, the claim that the normative status of features depends *only* on their internal status is the atomistic view. The strong version, as we saw, is that the features in question are always right or wrong *overall* in all cases; whereas the weak version is that they are always right-or-wrong-*making*, but not necessarily right or wrong overall. Second, the claim that the normative status of features depends *only* on other external features of the context is the strong holistic view, which I reject. The strong holistic view posits a world without principles at all, whether exceptionless or default, and thereby lends itself to Principle Eliminativism.⁷² Indeed, Kant is right that if we want to put forward a principled morality, we cannot reject intrinsic duties altogether. However, he misses out on an important possibility. The weak version of holism, which is the view that I adopt, holds that the normative status of reasons depends *both* on their internal default status and on other external features of the context. For example, lying is by nature wrong, other things being equal,

⁷¹ Kant, "On a Supposed Right to Lie", 350.

⁷² See section 4.7 for my discussion of Principle Eliminativism. Also, cf. strong holism with absolute holism in section 4.5.

and we have an abstract duty not to lie; but when considered in context, its default wrongness might vary due to other features of the case, such as the likely possibility that it will save a life.

I believe this weaker version of holism is the superior view. It synthesizes the information from our rationalistic and abstract moral principles with the empirical context in which they are applied, and views our intrinsic moral duties as equally important as the consequences of our actions. It accounts for the fact that some features have an inherent normative valence, while simultaneously recognizing the influence of context. This view of holism helps us provide more reasonable and compelling answers to complicated dilemmas such as “the murderer at the door”. Under this view of holism, lying is still considered both wrong and wrong-making by default, but in the “murderer at the door” case, it is *made* right due to more important moral considerations such as saving a life. Thus, in the abstract, the default principle “lying is wrong” is still affirmed as a true principle qua *ceteris paribus*; however, the moral reason it specifies, when considered in context, is nonetheless a variant one given the fact that its status may change due to other features present in the situation.

Let’s take another example to demonstrate how the reasons specified in the formula of humanity are not truly invariant. During WWII, allied soldiers had a duty to go to war against German soldiers. Here we have a case where allied soldiers had a duty specifically *not* to respect German soldiers as ends in themselves. The Nazi Regime, to whom German soldiers owed their allegiance, did not hesitate to implement systems that treated specific peoples merely as means (e.g. concentration camps). The WWII scenario is an example where the world had to use self-defence against tyrannical, imperialistic, and genocidal governments. Arguably, not all German soldiers were “Nazis”, nor were all of them directly implicated in the genocide of the Jews; but insofar as they went to war in the name of their country which was run by the Nazi Regime, they

were nevertheless implicated in the Nazi War Machine as a whole. It would have been arguably impossible and counterproductive for the Allies to treat them as ends in themselves while *only* treating the *real* Nazis merely as means. The state of war is complicated and in order to achieve victory the enemy must be clearly identified. The Allies had a choice: either they respect the German soldiers as ends in themselves, or their innocent victims; but they could not do both.

We have here the possibility of a generalizable argument against invariant reasons: for any supposedly invariant reason, imagine a dilemmatic situation in which we are torn between incompatible acts that involve fulfilling and breaking the same moral requirement (e.g. treating people as ends in themselves in a state of war).⁷³ When it is justifiable to do so, as arguably it often will be (e.g. WWII, murderer at the door, etc.), then the moral reasons in question are proven variant rather than invariant. This clash between incompatible acts does not imply that the moral principle is rejected altogether, but it does reject the supposed invariance of the moral reason it specifies. In the example above, in order to treat the innocent as ends in themselves, the Allies had to treat the perpetrators merely as means. Treating the innocent as ends in themselves and the perpetrators merely as means was right-making; whereas treating the perpetrators as ends in themselves would have been arguably wrong-making, as it would imply absolute pacifism as a necessary response to a mass murdering enemy—a response that would produce a significantly worse outcome. Consequently, when considered in context, treating others as ends in themselves is not always right-making, and treating others merely as means is not always wrong-making. Treating German soldiers and Nazis as ends in themselves would imply not going to war against them, in which case they would have a free pass to continue their invasions and mass

⁷³ It is worth looking into some of the arguments made by Samuel Scheffler in his book *The Rejection of Consequentialism*. Scheffler uses similar dilemma structures to argue against absolute prohibitions. We will discuss Scheffler in the coming paragraph.

exterminations. Indeed, going to war against them involved treating them merely as means, but it was the only thing the Allies could do in the circumstances to protect the greater number of innocent people. Reasons-holism, insofar as it forces us to study the context as opposed to relying solely on abstract principles, enables us to see the importance of both deontic and consequentialist reasons for action, rather than seeing one view as completely right and the other as completely wrong. Moreover, reasons-holism is compatible with default principles. These are principles that are true in the abstract qua *ceteris paribus* clauses, but that, inasmuch as they only specify the internal status of reasons, their truth, as applied to cases, is contingent upon external features of the contexts in which they are applied.

Samuel Scheffler appeals to a dilemma similar to the one we have laid out above in his rejection of the standard brand of non-consequentialism. According to the standard non-consequentialist view, “if you can either prevent the murder of one innocent person or the murder of five innocent people, it is rational to prevent the murder of the five, since you owe only a duty of positive aid to all six potential victims, and if you can either provide more or less positive aid, it is rational to provide more. But if *you* would have to murder one innocent person in order to prevent the murder of five by someone else, then you must not do so”.⁷⁴ Non-consequentialist views such as Kant’s deontological theory, hold certain duties to maximum esteem, thereby prescribing absolute prohibitions to act against them. But there are certain situations, such as the one Scheffler describes, that involve *personally* doing harm (rather than merely preventing harm or letting harm occur) in order to save more innocent lives. In the first case, one is preventing harm to five people, and letting one person be harmed; whereas in the second case, one is *personally* harming one person to save five. In the latter case, the Kantian must argue that one must not harm

⁷⁴ Samuel Scheffler. *The Rejection of Consequentialism: A Philosophical Investigation of the Considerations Underlying Rival Moral Conceptions*. New York: Oxford University Press, 1994, 24-25.

anyone under any circumstances. Thus, the Kantian must do no harm and let the five people be murdered.

According to Scheffler, non-consequentialists have yet to provide a satisfactory, non-question-begging answer to this dilemma. The answer that is typically given is that not inflicting harm is a stronger duty than preventing harm. But why is the former stronger than the latter? Appealing to either duty's internal content is hardly sufficient for concluding which of them trumps the other in a given case, let alone which of them lends appropriate support to agent-centered restrictions. Indeed, we seem forced to investigate the particular case at hand rather than focus merely on the abstract individual strengths of these respective duties. Given that there are cases where preventing the mass murder of many people would depend on our murdering a significantly lower number of people (especially those cases in which the latter are the would-be murderers of the former, such as in the WWII scenario), a more plausible answer to such a dilemma should depend on a contextual analysis that considers both our duties *and* their potential outcomes. Hardly anyone would argue that the Allies should not have gone to war against the Axis. Arguably, not going to war would have gone against our moral intuitions. Now, of course, this is not a sufficient reason as to why the Allies *should* have gone to war. Furthermore, to be clear, I am not saying that everything the Allies did was morally praiseworthy or permissible. However, the alternative view here, i.e. the Kantian or non-consequentialist view of absolute prohibition, seems significantly inferior to the weak holistic view I have described and defended in this section. The Kantian view not only goes against our intuitions and reasons in these aforementioned exceptional cases, but it also neglects the moral importance of context while at the same time blinding itself to empirical evidence that supports the likelihood of specific, morally relevant outcomes. The weak holistic view, by contrast, retains the intrinsic moral valence of our duties while simultaneously

considering the outcomes of those duties in the contexts in which they are enacted. The solution, then, is not to throw Kantianism out the window, but rather to reject the reasons-atomism upon which it is based.

With these considerations in mind, I think the Categorical Imperative, as applied to cases, should be viewed as a default rather than an exceptionless principle. Qua *ceteris paribus*, it still holds true, but when considered in context, the reason it specifies may vary. I believe that adopting reasons-holism and the notion of a default principle allows us to avoid acting in ways that are contrary to our intuitions, such as telling the truth to someone we have good reasons to believe will murder our friend, or remaining at home while imperialistic and genocidal regimes wreak havoc on other continents and plot the take-over of the world. The weak version of holism allows us to maintain the thesis that treating people as ends in themselves is intrinsically right and that treating people merely as means is intrinsically wrong (other things being equal), while simultaneously accepting the possibility that external features present in the context may change the otherwise default normative status of these abstract moral considerations. For these reasons, I believe the adoption of default principles in accordance with the aforementioned weak version of holism is the more compelling view when it comes to the behaviour of moral reasons.

2.3. The Principle of Utility

The Principle of Utility (or the Greatest Happiness Principle) is also said to specify an invariant reason, or a set of invariant reasons. Although the Principle of Utility (PU) is very useful in a wide array of situations, I will argue that it does not specify a set of truly invariant reasons. Below I have included its more popular formulations:

1. Actions are right in proportion as they tend to promote happiness, wrong as they tend to produce the reverse of happiness.⁷⁵

2. An act is right only if it produces the greatest amount of happiness for the greatest number of people.⁷⁶

To begin, the first formulation does not take into account cases where acts that tend to produce happiness are wrong, and cases where acts that tend produce the reverse of happiness are right. Let's begin with the former case. Suppose we live in a world of blue people and green people. The blue people's population is 5 billion, and the green people's population is 100 million. The green people, despite being a minority, have been in power for nearly a century. Suppose now that the blue people revolt and take over the government. Their leader is a tyrant who convinces them that they are racially superior, and that the green people have been the cause of all their social and economic problems. The blue tyrant calls for the extermination of the green people and receives complete support. In this example, the act of genocide will make more people happy. Thus, statistically speaking, this act tends to promote more happiness than the reverse of happiness. But of course, the act of genocide is clearly wrong. Hence, the act that tends to promote happiness is wrong-making in this case. Therefore, this supposedly invariant reason is not truly invariant.

On the other hand, sometimes acts that tend to produce the reverse of happiness are right. Suppose a man has murdered a mother's son. In court, the murderer is deemed insane and is given a lighter jail sentence along with a rehabilitation plan. The mother, doomed to live an unhappy life as a result of her son's violent death, is very angry at the verdict. Only the death penalty or life in imprisonment would have reduced some of her suffering. The current verdict, unfortunately, only

⁷⁵ John Stuart Mill. *Utilitarianism*. Mineola, New York: Dover Publications, Inc., 2007, 6.

⁷⁶ James H. Burns and Herbert L. A. Hart. *The Collected Works of Jeremy Bentham: A Comment on the Commentaries and A Fragment on Government*. Oxford University Press, 1977, 393: "...it is the greatest happiness of the greatest number that is the measure of right and wrong."

serves the function of worsening her suffering. But the court process was fair. The man was a paranoid schizophrenic, had no criminal record, and was honest and very remorseful during the proceedings. The man had no family or friends, but his co-workers said he was always respectful and kind to others. The man explained that he was struggling financially and could no longer afford his medications. As a consequence, his schizophrenia worsened and he became depressed. His episodes became very intense, and they eventually led him to murder someone he thought was trying to kill him. The murder was not intended or premeditated, and he was only defending himself in the moment to a threat that he thought was real. In the end, the judge made what most of us would probably call a fair judgement. Let's suppose he received 5 years in prison with a lifetime supply of his medications paid for by the state (to insure such a scenario doesn't happen again), followed by 3 years in a rehabilitation facility and 5 years of community service.

In this situation, perhaps handing out the sentence the mother was hoping for would have reduced *her* suffering / increased *her* happiness. But the court's primary duty is not to make people happier but to hand out sentences the perpetrators of a crime deserve. In this case, the ruling would have a negative impact on the mother's happiness, yet it would be a morally right ruling. Let's also not forget that, since the man had no family and friends, he had no one in his life that would feel great pain if he would have been sentenced to death or had received life in prison. So the only people whose happiness is relevant is the mother (and her friends and family), the man, and the general community. We can assume that the mother and her circle would be mostly unhappy with this ruling. We can probably also assume that the general community would also be mostly unhappy. This is because those who would be opposed to the ruling would most likely feel the reverse of happiness (anger, fear, sadness, etc.) more intensely than those in favour of the ruling who, for their part, would arguably feel a comparatively milder feeling of happiness. So, overall

this ruling would arguably create more of the reverse of happiness than happiness itself; yet, it would be entirely reasonable and fair. Thus, there are some cases where we have good reason to think that the act that tends to produce the reverse of happiness is right. Hence, the reason identified in this version of the PU can, in some cases, be no reason at all or even an opposite reason. Therefore, this supposedly invariant reason is not truly invariant.⁷⁷

Moving on now to the second formulation of the PU. To this supposedly invariant reason, I will provide three counterexamples: (1) the miserable coma patient; (2) the unwanted Aunt Suzie; and (3) the suicidal child molester under the knife.

Let's begin with our first counterexample: the miserable coma patient. Suppose an elderly woman of 78 years has been in a coma for 2 months. Given the type of coma she's in, her alcoholism, her age, and other factors, the doctors believe she may only wake up after several months or years. The odds are not good, but they are also not hopeless. As it happens, there is no written or recorded testament from the elderly woman stating that she would want to be kept alive if ever she was to fall into a coma. She has five children, and they are asked if their mother wanted extreme measures taken to preserve her life. Although there is no official proof, the children know

⁷⁷ A utilitarian might object here by claiming that having fair trials maximizes happiness overall, despite causing unhappiness at times (e.g. in the case of the mother whose son was murdered). That is, increasing the mother's happiness here would not maximize happiness (or tend to produce happiness) overall when we take a longer and broader view of the relevant consequences of a fair trial. Relatedly, this type of utilitarian could say that increasing the mother's happiness would be right, but that since fair trials tend to cause more happiness, the fair trial is more right. Here, though, there is a crucial difference between maximizing overall happiness and the tendency to produce happiness. This objection would be an excellent response to someone arguing against the invariance of formulation 2. For formulation 1, however, it does not seem to hold. All things considered, it might be true that fair trials tend to produce more overall happiness, despite causing unhappiness at times. But if we isolate the trials having to do with mothers whose sons were murdered and where the majority of the people affected are unhappy about the verdict, we can say that *those* trials tend to produce more unhappiness. So in the end, it's all in how far we want to go in terms of related consequences (i.e. short, medium, and long-term) and whose happiness we want to concern ourselves with. Also, it is not evident that fair trials, all things considered, tend to produce more happiness. They do seem to produce more justice, but we shouldn't just assume that the production of justice tends to produce more happiness. As a final point, formulation 1 would be significantly stronger if it were "Actions are right in proportion as they tend to promote *more overall* happiness, wrong as they tend to produce *more unhappiness*".

without a doubt that their mother would have wanted the extreme measures taken. In fact, this was something she always talked about after she woke up from a previous coma many years ago after overdosing on alcohol. But the children despise their mother for all the torment she has caused in their lives. She cheated on her husband, forcing a divorce; she was and is still an alcoholic; and she is generally a miserable person who makes everyone around her miserable. So the children lie to the doctors by telling them that their mother would have wanted them to pull the plug after 3 months of inactivity, which also happens to be the law. The children figure that with their mother out of the picture they would all be significantly happier, their mother would be put out of her misery, and the world would be a better place.

In this scenario, pulling the plug seems to be the action that would cause the greatest amount of happiness for the greatest number of people. The mother's perpetual reduction of everyone's happiness would be stopped, and the children would be happier. However, this reason for action is arguably overridden by more important considerations such as the respect for life and the respect for a person's autonomy. That is, the formula of humanity version of the CI seems to trump the PU in this case. Pulling the plug against their mother's wishes in order to make everyone happier would be to use her merely as a means to attain that end. Arguably, with their mother in a coma the children should be able to live their lives more freely and happily now. Moreover, they would barely need to visit. For these reasons, then, the act of pulling the plug in this case, no matter how much happiness it produces for the children, is wrong. Not only would it be done by means of a lie, it is also defeated by more important moral considerations such as the respect for the mother's life and the respect for her autonomy.

Moving on to our second counterexample: the unwanted Aunt Suzie. Suppose that everyone in the family hates Aunt Suzie, and that Aunt Suzie hates everyone in the family. Suppose further that, despite this reciprocal hatred, the family still invites Aunt Suzie for Thanksgiving dinner every year out of mere politeness, and that Aunt Suzie shows up every year also out of mere politeness. If Aunt Suzie were to decide one day that she would not be showing up at family gatherings anymore, everyone in the family would be happier, including herself. Her decision, it seems, would promote the most overall happiness. Suppose now that the family and Aunt Suzie agree that Aunt Suzie will no longer be showing up to family gatherings precisely for this reason.

In my view, the voluntary and collective decision of having Aunt Suzie removed from the family—which in this case promotes the most overall happiness—is arguably morally wrong because it is a show of cowardice and a lack of ambition by both parties (the family and Aunt Suzie).⁷⁸ Families should work together to maintain and build relationships; they shouldn't just give up by taking the easy way out. First of all, getting along is not impossible; and secondly, the adults should act as mature, responsible, and virtuous role models for their children. Yes, working with someone who has a polar-opposite personality is hard, but who said doing the right thing was going to be easy? Arguably, this voluntary and collective decision that promotes the most overall happiness is not only defeated by more important considerations, it is also itself wrong-making.

But a utilitarian might argue that a society where family members try to get along, despite all the unhappiness that may cause, would, in the long-term, promote a better society than one where people take the easy way out. More unhappiness might indeed result, but the fact that family members are actively trying to get along could lead to significant improvement in family relations

⁷⁸ Cowardice and lack of ambition are moral vices. See Aristotle's *Nicomachean Ethics* Books II-V for a detailed study of virtue and vice. For courage and cowardice see Book III, 6-9. For ambition and lack of ambition see Book IV, 4.

and, consequently, to more overall happiness in the long-term. Furthermore, if every family would do the same, then we probably *would* live in a better society in which the greatest amount of happiness for the greatest number of people is achieved. This type of utilitarian would value the striving towards the greatest possible society even if some unhappiness may result along the way.

The problem with this objection is that this “better society” will not necessarily be a happier one. It might be a more virtuous one, and/or a more meaningful one, but it is not evident that it would be a happier one. I do agree that we should aim for the greatest possible society, and I do agree that families should strive to get along despite all the unhappiness this may inevitably cause, but this aim will not necessarily lead to a happier society. Isn’t it true that an important part of being virtuous sometimes involves doing the right thing even if it makes most people unhappy? In the situation above, maintaining and building family relationships and showing the example to the kids does seem to be the right thing to do. To be clear, if Aunt Suzie were an alcoholic brute who frequently swore in front of the kids and who insulted everyone every chance she got, then of course the situation would be different. But here, she is capable of acting politely during gatherings which demonstrates that she is at least capable of some decency and self-improvement.

Another utilitarian could argue that it is not *necessary* to build a relationship with Aunt Suzie. The kids have plenty of other adult role models, and the absence of Aunt Suzie will probably barely be felt at all. Although I agree that the kids have a *sufficient* number of role models and that Aunt Suzie’s absence will likely not negatively affect the trajectory of the kids’ moral development, I am still unconvinced that her removal—which would promote the most overall happiness—is the right thing to do. What is not necessary is not necessarily morally irrelevant, and what is perhaps morally permissible is not necessarily the best course of action available. Although the removal of Aunt Suzie would cause the most overall happiness, it is a sub-optimal attempt at

a better family because it involves the exclusion of a member. A true attempt at a better family is generally one that *includes* everyone in the family—not one that excludes some members for the betterment of the majority simply because it is easier than trying to get along.⁷⁹ So while excluding Aunt Suzie may have its benefits, it is arguably not the best moral option available. For these reasons, then, the reason specified in the second formulation of the PU does seem to be defeated by more important considerations such as the virtue of proper ambition, the virtue of sacrifice, and the value of the family.

Moving on to our third counterexample: the suicidal child molester under the knife. Suppose an unknown man is undergoing brain surgery. He attempted suicide by a self-inflicted gunshot to the head. The suicide note said that he had no family or friends and that he had made terrible mistakes in his life. He also mentioned that he has no desire to contribute to society or to continue living. At first, the prognosis is not very good. But early into the operation, the neurosurgeon becomes quite confident that he can remove the bullet and that the man will make a slow but full recovery. However, right before he's about to inform his team that the man's chances at recovery are quite good, one of the scrub nurses walks in with a report on the man's background. She says the man under the knife is a serial child molester who is under house arrest. The neurosurgeon begins to question himself: should I really save this man?

The man has done terrible things in his life, and the fact that he is under house arrest means that he can barely contribute to society. In fact, the suicide note specifically mentions that he has no desire to contribute in any way to the betterment of society. What is more, his mere presence in his house likely reduces the happiness of the people in his neighbourhood who know who he is,

⁷⁹ I say “generally” because, of course, there are circumstances where a family member is abusive and violent and should indeed be excluded from the family.

especially families who have children. Finally, his attempt at suicide clearly indicates that he no longer wants to live. Given these facts, if the neurosurgeon saves him the man will likely reduce the overall happiness of the people around him for the rest of his life. On the other hand, the neurosurgeon could easily choose not to save him. He knows he could easily fabricate a convincing lie in the report by saying that the bullet was lodged at a slightly different location in the brain making a safe removal impossible. He could also purposely remove the bullet less carefully than he otherwise would specifically in order to cause an internal bleed that would kill the man.

In this case, killing this man would arguably cause more overall happiness. Furthermore, cases such as these are so rare that if this particular act of killing for the greater happiness of all were universalized (i.e. if all neurosurgeons in this situation would kill the man for the greatest happiness of all) the institutions of happiness and justice would not be at any risk of collapsing; in fact, more overall happiness would be produced. But of course the act of killing for the promotion of the greatest overall happiness is clearly wrong in this case. The respect for the man's life, the duty of the surgeon to save lives, and his duty of non-maleficence as a health care professional clearly defeat the PU in this case.

A utilitarian might object and argue that a society in which surgeons practice impartiality, treat all patients equally, and fulfil their duties is a better society than one where exceptions such as these are made. While I agree that such a society would be a better one, it would not necessarily be a happier one. Again, cases such as these are quite rare; and universalizing the act of killing in these particular circumstances will not cause the institutions of happiness and justice to collapse. In fact, doing so will reduce both actual and potential harm to others and consequently increase overall happiness. For these reasons, it appears the utilitarian who wants to defend the universal and invariant truth of the PU is forced to defend the act of killing in this case.

To escape this objection, the utilitarian would have to admit that his principle is subject to exceptions (in which case it does not specify an invariant reason and our argument is correct). To maintain invariance, the utilitarian would arguably have to modify the formulation of his principle to “an act is right only if it promotes the greatest amount of *good* for the greatest number of people”. The term “good” here would include all beneficial things (e.g. happiness, meaningfulness, pleasure, self-overcoming, empowerment of individuals and communities, security, etc.).⁸⁰ The issue with this version of the PU, however, is that while its scope becomes larger and its reason perhaps invariant, its resolution is significantly reduced. For one, we would have to find a way to rank-order the various ends/goods (happiness, meaningfulness, empowerment, etc.). That is, we would have to determine which is superior overall and which is superior in contexts x, y, and z. Furthermore, we would have to account for contexts where a combination of these ends are pursued. In such cases, we would have to determine which combinations are superior. These are no easy problems to solve. So although the reason specified in this version of the PU may be invariant in its *theoretical* formulation, its resolution is so unclear and unspecified that its *practical* application will almost always be variant. Simply put, it is so vague that it will be too difficult for the moral agent to apply it in the same way (the objectively best way) across different situations. Thus, it will likely lead the moral agent to apply it in a variety of different, sub-optimal ways.

In conclusion, The Greatest Overall Happiness Formulation of the PU does not specify an invariant reason because it can be shown to vary in certain situations. Therefore, while it is a very strong *defeasible* principle, it is not an exceptionless one. On the other hand, The Greatest Overall

⁸⁰ This is the route taken by George Edward Moore with his idea of “Ideal Utilitarianism” in the last chapter of *Principia Ethica* entitled “The Ideal”. See George E. Moore. *Principia Ethica*. (Revised Edition with “Preface to the second edition” and other papers, T. Baldwin (ed.), Cambridge: Cambridge University Press, 1993). Cambridge: Cambridge University Press, 1903.

Good Formulation, although possibly invariant and exceptionless, is very difficult to apply invariantly due to issues of vagueness. Indeed, it is arguably a principle that is implicitly assumed by all moral agents striving to do what is right; but while it might be necessary for moral thought and judgement, it is undoubtedly insufficient.

2.4. Causing Gratuitous Pain on the Unwilling

The causing of gratuitous pain on unwilling victims is often considered to be invariantly wrong. But suppose the world is on the brink of destruction because a tech-savvy psychopath has gained control over several nuclear weapons. Suppose further that you have been kidnapped by the psychopath, and he tells you that if you want to save the lives of millions of people, you must pull a switch that will cause pain to five unwilling victims tied to electric chairs. In this case, you have strong reasons to believe that the psychopath does indeed have control over weapons of mass destruction. For example, suppose he pressed 3 buttons and you subsequently saw videos of 3 explosions in three different countries on the news along with very graphic and realistic images of the wreckage and of the dead bodies; suppose further that a military or secret intelligence official has contacted you to confirm that the psychopath does in fact have these capabilities; and suppose again that the psychopath persuasively explained to you step by step how he managed to orchestrate this whole affair. Insofar as you have strong reasons to believe that he does have the ability to eradicate millions of people, it seems you are morally required to pull the switch, despite how terrible the act is. In this situation, not pulling the switch is far more terrible an act than pulling it, for the psychopath can press more buttons and continue to eradicate more people. If the situation were such that you did not have strong reasons to believe the psychopath's capabilities (or if you weren't very certain) then it would be permissible not to pull the switch, perhaps even preferable.

But this is not the situation here. In this situation, you have very strong evidence that the psychopath does have such capabilities. Therefore, it seems here, that we have a strong counterexample against this supposedly invariant reason.

Dancy also argues that there are such cases out there. Dancy writes: "...can the particularist admit the existence of some invariant reasons? The obvious examples are things such as the causing of gratuitous pain on unwilling victims. Surely, it is commonly urged, this is always for the worse, even if overall we might in some case be morally forced to do it".⁸¹ But suppose you, your family, and a fat man are in a cave quickly filling up with water. You and your family are stuck in between the fat man, who is blocking the only exit. But there is a way out: you just so happen to have some dynamite with you. Do you cause gratuitous pain on the blameless and unwilling fat man by blowing him up and escaping to safety? Or is causing gratuitous pain on this unwilling victim still a reason against doing it in this case? It seems we are morally required to do it. If it were merely preferred over the alternative it would simply diminish its power as a reason. But since it is morally required, it defeats it. In either case, the supposedly invariant reason has varied. Of course, many of these hypothetical cases are silly and unlikely to occur, but if we are going to test the limits of moral principles that purport to specify reasons that supposedly *never* vary, this is to be expected. Thus, although causing gratuitous pain on unwilling victims is generally wrong by default, it is still not *invariantly* wrong.

⁸¹ Dancy, *Ethics Without Principles*, 77.

2.5. The Principle of Special Reason to do Harm

The idea behind this principle is that one *must* have a good reason to do actions that cause harm, or to believe that it is right to do them.⁸² Is this an invariant moral reason? The Principle of Special Reason to do Harm states that one needs a good reason either to do an act that is default immoral such as killing, stealing, or lying, or to believe that it is right to do them. This principle came to mind after reading a passage from the eighth chapter of *Ethics and the Limits of Philosophy* by Bernard Williams. Williams writes: “No doubt there are some ethical beliefs, universally held and usually vague (“one has to have a special reason to kill someone”), that we can be sure will survive at the reflective level. But they fall far short of any adequate, still less systematic, body of ethical knowledge”.⁸³ Taken in the abstract, killing someone is default wrong, but if there is a feature in the situation that would make it right, then that feature would serve as the special reason for why it would be right to kill someone. This principle seems to apply itself invariantly in all situations where there is an exception to the rule—that is, situations where there is a defeater or a disabling condition that would change the default status of the belief or act in question. Indeed, this principle specifies an invariant reason, but this should not come as a surprise, for it is one of the fundamental tenets of morality itself. That is, by definition, we cannot do anything otherwise wrong without a special reason to do so. The reason why we can’t come up with a counterexample is because it is an integral part of the specific content of morality. In other words, this principle must be true, for if it wasn’t, nihilism would ensue and everything would be permitted.

However, although this principle makes moral justification necessary and specifies an invariant reason, it does not by itself make moral thought and judgement possible. It is a necessary

⁸² Here, we mean harm of any kind whether it be social, physical, mental, or emotional.

⁸³ Bernard Williams. *Ethics and the Limits of Philosophy*. Abingdon, United Kingdom: Routledge, 2006, 148.

condition for the possibility of moral thought and judgement, but it is not a sufficient one. For it is a principle that recognizes (but does not specify) a set of reasons that permit us to do acts such as killing, stealing, and lying in certain exceptional situations. It does not mention any specific reasons for doing those acts—it simply admits that there are some. Simply put, it is a true but vague ethical belief. So even if it is an ethical belief that is both necessarily true and a necessary condition for the possibility of moral thought and judgement, it remains insufficient for the task because it is an abstract consideration without any contextual content. In order to be both necessary and sufficient to make moral thought and judgement possible, it would also need to explain: (a) why specific acts have a certain status; and (b) why the status of those acts may change. If it cannot meet these conditions, it would have to be considered in combination with other invariant reasons that *can* provide these explanations.

Another thing to note is that this principle is open to the possibility that the status of certain reasons can vary across contexts. That is, the principle assumes that there are special features in certain situations that are reasons to believe or do the opposite of the default moral belief or action. Thus, it presupposes the inherent variability of default reasons, and therefore implicitly accepts reasons-holism. So, although this principle specifies an invariant reason, its specific content is such that it actually endorses the variability of reasons. We have here an exception that proves the rule.

There might be more of these invariantly true abstract considerations, but insofar as they do not explain why specific acts have a certain status, or why the status of those acts may change in specific patterns, they are insufficient for making moral thought and judgement possible, let alone codifying all moral truths in finite and manageable terms. Still, so far it seems that a proper view of reasons-holism admits of at least one invariant reason. One always needs a special reason to do something that would cause harm, otherwise everything would be permitted.

2.6. Richard Holton's *That's it* Principle

Richard Holton's *That's it* principle is also a strong candidate that may provide an invariant reason. *That's it* is a conjunct clause that can be added to the antecedent of the central premise of a moral argument. The *That's it* clause stands for "there are no further relevant moral principles and non-moral facts; i.e. there is no true moral principle and set of true non-moral sentences which supersede those which appear in this argument".⁸⁴ In simpler terms, a *That's it* clause is a way of saying: "so far, from what we know, there are no other relevant features in this case that could change our moral verdict". Consider the following moral argument provided by Holton himself:

P1 This is a killing
P2 $\forall x ((x \text{ is a killing} \ \& \ \textit{That's it}) \rightarrow \text{you shouldn't do } x)$
P3 *That's it*
C You shouldn't do this.⁸⁵

In this example, we have an action that is default wrong (i.e. killing), and there is nothing else to suggest that it might have been right in this context. Therefore, it remains wrong. According to Holton, *That's it* clauses are implicitly present in all moral arguments and situations. No matter how many conjuncts there are in the antecedent, the *That's it* clause is always implicitly assumed.⁸⁶ This is because when we revisit a case and see that there were, in fact, other relevant features we had originally missed, it is *That's it* that is denied. Suppose we revisit the case of the killing and find that it was done in self-defence. The new argument would be formulated as follows:

P1 This is a killing
P2 This is done in self defence
P3 $\forall x ((x \text{ is a killing} \ \& \ x \text{ is done in self defence} \ \& \ \textit{That's it}) \rightarrow \text{you may do } x)$
P4 *That's it*

⁸⁴ Holton, "Principles and Particularisms", 199.

⁸⁵ Ibid., 199.

⁸⁶ Ibid., 208-209.

C You may do this.⁸⁷

The *That's it* clause in the original argument has been denied because we have found a relevant feature (i.e. the fact that the killing was done in self-defence) that we had previously overlooked. This new fact is thus added to the argument, and a new *That's it* clause is affirmed. So far, from what we know, there was a killing and it was done in self-defence, and that's it. But this knowledge is not fool-proof, for new information can always appear and lead us to change our verdict. For example, suppose that, although the killing was done in self-defence, we find evidence that the killer unnecessarily went above and beyond and even enjoyed it.

The point is that the *That's it* clause seems to be implicitly assumed in all moral arguments and situations. This seems to indicate that it provides us with an invariant moral reason for belief. The invariant reason it seems to be putting forward might be rendered as follows:

For all moral situations in which there is a moral/immoral belief or action, and no other relevant features to disable or defeat that belief or action, then it is right/wrong.

For example, suppose there is a killing (an act that is default wrong) and we can confirm that there are no other relevant features that make it right (i.e. it was not done out of self-defence) because we have clear video footage of the murder as well as other convincing evidence (e.g. unanimous testimonies, clear intent and premeditation in the killer's personal diary, etc.). In this case, the act of killing is wrong. The fact that there are no other relevant features (alongside all the other evidence) is a reason to believe/confirm that the killing is wrong. Of course, for this to be a reason we need to know beyond a reasonable doubt that there are, in fact, no other relevant features.

⁸⁷ Ibid., 199.

We can imagine many other cases where *That's it* cannot be denied. Consider the following example. Suppose my mother asks me if I want to go the movies on Saturday. She feels we haven't spent quality time together in a while and really wants to spend time with her son. I promise to go. But then on Friday my buddy Dave asks me if I want to go fishing all-day with him on Saturday. I really like Dave and spending time fishing with him is always a blast. So I decide to lie to my mother by telling her that I can't go to the movies. I tell her I need to finish an online assignment for school that I had forgotten was due Saturday night by 11:59pm. I tell her I will need to go to the library and access some books put on reserve by my professor—but the truth is that I'm going fishing with Dave. This is an example where I have no good reasons to break my promise to my mom, or to lie to her. The moral argument for why my action is wrong would be as follows:

P1 This is breaking a promise by lying.

P2 $\forall x ((x \text{ is a breaking a promise by lying} \ \& \ \textit{That's it}) \rightarrow \text{you shouldn't do } x)$

P3 *That's it*

C You shouldn't do this.

In this example, I am clearly in the wrong, and there does not seem to be any way that I can justify breaking my promise with my mom. The idea behind *That's it* is that, for all situations in which we have compelling evidence to conclude beyond a reasonable doubt that the *That's it* clause is not denied, we have an invariant reason to believe that the conclusion of our moral argument is true. Indeed, many moral situations are too complex to conclude beyond a reasonable doubt that *That's it* is not denied; but there do appear to be some situations where we can effectively make this conclusion, such as the example about breaking a promise with my mom. The amount of joy I would have on my fishing trip with Dave simply does not make it right for me to break my promise to my mom by lying to her. Certainly, the disfavourers or wrong-making features outweigh any potential favourers in this situation. A fishing trip can be planned another

time, and it is of course wrong to deceive others unnecessarily (especially one's own mother) and potentially cause them harm in the process. So given that there are no other relevant features that could disable or outweigh the wrong-making features in this case, and given that we can be reasonably sure of this, the *That's it* clause stands firm and our verdict is not overturned. The same line of argument seems to hold in the case involving the act of killing. The killing that was not done out of self-defence, and which was video-taped in such a way as to confirm the description of what happened, and which is further supported by other compelling evidence, is clearly wrong. The wrong-making features of the act, including those in the situation surrounding it, clearly outweigh any existing right-making features, if there are any. The only possible exception where the right-making features might outweigh the wrong-making features would be if the murderer was coerced into murdering in order to save the lives of multiple people. But such a scenario would be quite a complicated affair and it would be difficult to determine exactly what would be the best course of action. Consequently, we certainly would not be able to conclude *beyond a reasonable doubt* which course of action would be best. Thus, such a scenario simply falls outside the scope of the invariant reason we are here considering and does not serve as a working counterexample.

So far, this supposedly invariant reason seems quite strong. But is it really an invariant reason specified by an exceptionless principle? As I understand it, this invariant reason does seem to be necessarily true, but only in specific contexts whose conditions are already outlined in the formulation of the principle itself. In other words, the structure of the principle is already laying out the perfect set of conditions for invariance. If we're in a situation in which all enablers and disablers have been accounted for, and wherein we know beyond a reasonable doubt all the favourers and disfavourers relevant in the case, then of course our moral conclusion will be certain: the right/wrong belief or action will depend on whether the favourers outweigh the disfavourers

and vice versa. Therefore, it is unsurprising that the reason specified in *That's it* is necessarily invariant—its very structure is built for invariance. Still, there is a catch: although this reason is necessarily invariant, its truth depends on its being informed by contextual experience. In other words, we need to know a posteriori that there are no other relevant features that could supersede those which appear in our moral argument.

The great thing about this principle is that it is not infinitely or overly long. Rather, it is finite, relatively short in formulation, theoretically graspable, and practically applicable. As such, it certainly meets the criteria of an action-guiding standard. But despite these promising characteristics, it is still by itself insufficient for making moral thought and judgement possible, for it only captures a small portion of the moral landscape. That is, it only covers the contexts in which we effectively *can* conclude, beyond a reasonable doubt, that there are no further relevant moral principles and non-moral facts which supersede those which appear in our moral arguments. Thus, although it allows for a partial context-dependent codification, it does not codify all moral truths a priori. There is still a huge part of the moral landscape that it does not cover. First, it does not apply in cases in which the conclusion—that there are no further relevant features that supersede those which appear in our moral argument—is too difficult to determine due to the complexity of said cases. It also does not cover those cases where the *That's it* clause is denied and reaffirmed indefinitely as a result of the acquisition of new information.

Interestingly, this whole process of denying and re-affirming *That's it* actually proves the variability of reasons. Just like with the Principle of Special Reason to do Harm, we have yet another exception (an invariant reason) that proves the rule (the variability of reasons). Furthermore, insofar as we are constantly revising and reformulating our principles after each denial of *That's it*, we are admitting the *defeasibility* of our principles, not their supposed

exceptionlessness. And, as we saw in chapter 1, defeasible principles are generalizations expressed by contingent propositions that specify *variant* reasons, not invariant ones.

But although this moral reason is invariant, its scope is far too limited to make moral thought and judgement possible, let alone codify all moral truths. It is therefore with right reason that Holton argues that no finite set of finite principles can axiomatize morality.⁸⁸ This is in fact the central claim of Principled Particularism—a view we not only agree with in this thesis, but one upon which we would like to build. We will discuss this important claim in the next two chapters.

2.7. Conclusion

So far, I have argued that there are good reasons to doubt the supposed invariance of the vast majority of moral reasons, given that the majority of the strongest candidates for invariant reasons have been shown to vary in certain cases. In the end, we only found that two of them were truly invariant: (1) the true ethical belief that one must have a good reason to do actions that cause harm (specified by the Principle of Special Reason to do Harm); and (2) Richard Holton's *That's it* principle. In the first case, we noted that the invariance should not come as a surprise because if such a principle/belief was not invariant, then morality would not be possible in the first place and nihilism would ensue. One of the reasons why we have ethical debates and adhere to certain ethical standards in our day-to-day lives is at least partly because we assume that this principle/belief is a given. In essence, it prevents the justification of the immoral. In the second case, we observed that Richard Holton's *That's it* principle does indeed specify an invariant reason, but only in specific contexts whose conditions are already outlined in the formulation of the principle itself. Here, we noted that this is because the principle is already laying out the perfect set of conditions for

⁸⁸ Ibid., 193.

invariance. Finally, in both cases, we showed that both principles implicitly assume the variability of reasons.

So to return to this chapter's research question: are exceptionless moral principles sufficient to ground the possibility of moral thought and judgement? I do not believe they are. Exceptionless moral principles are supposed to specify invariant moral reasons, but given that invariant reasons are few and far between, there are simply not enough of them to ground the possibility of moral thought and judgement. Indeed, one of them is necessary to avoid nihilism, and the second to confirm that a moral verdict is not overridden by other considerations. However, they are both nevertheless insufficient to make moral thought and judgement possible. If reasons-atomism was true and all reasons were necessarily invariant, then exceptionless principles would be able to make moral thought and judgement possible a priori. Here, the contingencies of context such as an individual's socio-cultural environment, time, place, people present in x situation, moral teachers, moral exemplars, and so on, would only serve the function of activating a moral knowledge that was already there. But since reasons-atomism is false given that most reasons are in fact variant, this simply cannot be the case. Rather, we should say: since reasons-holism is true, an individual *must* have at least some sense of the context-sensitive nature and behaviour of variant moral reasons in order to think and judge in moral terms. Therefore, the contingencies of context *must* play an important role in making moral thought and judgement possible, since they are the things that can cause a reason to vary from context to context (i.e. they can serve as enabling or disabling conditions that can affect the status of a reason, or as intensifiers or attenuators that can affect the strength or weight of a reason). Thus, the contingencies of context are necessary not only to help activate moral knowledge, but first and foremost to understand the nature and behaviour of moral reasons. Hence, this means that, in a reasons-holistic moral landscape, principles are not the only

things that play a role in making moral thought and judgement possible. Consequently, Principle Absolutism is false because it claims that a specific set of exceptionless moral principles is sufficient to make moral thought and judgement possible.⁸⁹ Such things as evolutionarily entrenched moral intuitions, instincts, and domain-general learning capacities should also be examined in relation to moral thought and judgement, for they too may prove to be necessary for their possibility. We should study these because it is not at all evident that our moral principles perfectly represent the nature of, say, our moral instincts, nor is it obvious that our moral instincts can be at all expressed in propositional form. For these reasons, it seems to me that it is necessary to discuss the moral psychology literature. These considerations will comprise the bulk of our next chapter. However, before moving on to chapter 3, let's summarize our argument for chapter 2:

1. The majority of the strongest candidates for invariant moral reasons can vary in certain contexts (proof: our counterexamples).
2. Therefore, the majority of the strongest candidates for invariant reasons are not invariant.
3. Since the majority of the strongest candidates for invariant reasons are not actually invariant, then it is highly likely that the vast majority of reasons are variant.
4. Therefore, reasons-holism is true.
5. Since exceptionless principles purport to specify invariant reasons, and invariant reasons are few and far between, exceptionless principles are also few and far between.
6. Since variant reasons constitute the majority of the moral landscape, exceptionless principles are incapable of codifying the entirety of the moral landscape (because they will only be able to codify the invariant reasons, which again are very few in number).⁹⁰
7. Since reasons-holism is true and the vast majority of reasons are variant, an individual must have at least some sense of the context-sensitive nature and behaviour of variant reasons in order to think and judge in moral terms.
8. Therefore, the contingencies of context (e.g. socio-cultural environment, time, place, people present, moral teachers, moral exemplars, etc.) *must* play an important role in making moral thought and judgement possible; this is because they are the things that can cause a reason to vary from context to context (i.e. they can serve as enabling or disabling conditions that can affect the status of a reason; or as intensifiers or attenuators that can affect the strength or weight of a reason).

⁸⁹ This may put a dent in Transcendental Generalism insofar as it vouches for exceptionless moral principles alongside default ones. For confirmation that Transcendental Generalism vouches for exceptionless moral principles, see McKeever and Ridge (2006: 7, 64-65) and (2016: "Epistemological Arguments").

⁹⁰ This certainly puts a dent in Transcendental Generalism, because it claims that exceptionless principles, at least in theory, can codify the entire moral landscape. We will discuss this in the next two chapters.

8. Therefore, in a reasons-holistic moral landscape, principles are not the only things that play a role in making moral thought and judgement possible.
9. Therefore, exceptionless moral principles by themselves are insufficient to make moral thought and judgement possible.
10. Consequently, Principle Absolutism—which claims that a specific set of exceptionless moral principles is sufficient to make moral thought and judgement possible—is false.
11. But even though exceptionless principles are not sufficient to make moral thought and judgement possible, some of them appear to be necessary such as, for example, the Principle of Special Reason to do Harm and *That's it!*.
12. Special Reason to do Harm is necessary to avoid the justification of the immoral and explains why default reasons can be defeated or disabled; and *That's it!* is implicitly assumed in all moral arguments because its clause is always either affirmed (i.e. when we know that there are no further relevant features that could supersede the ones which appear in our argument) or denied (i.e. in cases where we discover features we had overlooked that *do* supersede the ones which appear in our argument).
13. Therefore, moral thought and judgement necessarily depend on a *determinate* set of exceptionless principles, even though this set is insufficient for making them possible.
14. Therefore, Anti-Transcendental Particularism and Principle Eliminativism—which both claim that moral thought and judgement do not depend on moral principles of any kind—are false.⁹¹

Chapter 3

3.1. Introduction

Does the possibility of moral thought and judgement depend on the provision of a suitable supply of moral principles? Yes. I will argue that moral thought and judgement depend on a *determinate* set of exceptionless principles and an *indeterminate* set of defeasible ones, but that moral principles as a whole are *insufficient* for making them possible. By examining research in the fields of moral psychology, cognitive science, and neuroscience, I will argue that there are other factors that figure in the emergence and development of moral thought and judgement.

⁹¹ There will be more arguments provided against ATP and PE in chapter 4 in the sections entitled “Rejection of Anti-Transcendental Particularism” and “Rejection of Principle Eliminativism”. Still, I believe our arguments in chapter 2 are sufficient for their rejection.

3.2. Particularism and Defeasible Principles

In chapter 2, we already argued that moral thought and judgement depend on a determinate set of at least two exceptionless principles: the Principle of Special Reason to do Harm and Holton's *That's it* principle. But that is only a part of the answer; we also have to consider *defeasible* moral principles. We also have to clarify that an affirmative answer to the dependence question does not automatically entail that we are advocating for generalism. As we argued in chapter 1, one's answer to the dependence question is not the determining criterion for one's stance in the particularism-generalism debate. It is possible to imagine a form of particularism that claims that moral thought and judgement depend on the provision of a suitable supply of moral principles, but that also claims that such a provision will nevertheless be insufficient for codifying the moral landscape.⁹²

Particularism, as Margaret Little puts it, is “not against moral generalizations, but against a certain picture of what those generalizations have to be like in order to do their work”⁹³. Even though particularists are generally skeptical of exceptionless generalizations expressed by necessarily true propositions, they are not necessarily against them. What is more, they are not necessarily against defeasible principles expressed by contingent propositions. Thus, particularists are not against the idea or use of certain types of moral generalizations. Rather, they are opposed to the way generalists apply them. Particularists are unanimously against the *complete codification* of the moral landscape by means of a *specific set* of moral principles.⁹⁴ In this vein, Little writes: “While particularism dethrones moral generalizations, though, it hardly exiles them. We will think

⁹² In fact, this is what we will argue in chapter 4.

⁹³ Brad Hooker and Margaret Olivia Little. “Introduction.” In *Moral Particularism*, edited by Brad Hooker and Margaret Little, New York: Oxford University Press 2000, xi.

⁹⁴ However, we will leave open the possibility of there being different ways of codifying the moral landscape which might be compatible with certain forms of particularism.

it does only when the lesson of holism is misidentified; the enemy its objections target is not generalization, but codification”⁹⁵. Holism is an attack on the assumption that things *must* be a certain way, it is not an attack on moral generalizations per se. Particularism is not incompatible with the possibility of moral principles derived from inductive generalizations based on contextual experience. Defeasible principles, for example, are generalizations expressed by contingent propositions that function perfectly well in a reasons-holistic world, inasmuch as they seek to specify variant reasons. Hence, in this chapter, I will argue that moral thought and judgement also depend on an *indeterminate* set of defeasible principles. In my view, an individual needs a suitable amount of defeasible principles in her socio-cultural environment to work with before she can learn to think and judge in moral terms. The fact that some of us *can* judge in moral terms implies that there is a suitable supply of principles available for us in our socio-cultural environment that made our moral judgement possible in the first place. Allow me to defend these claims.

3.3. Determinate and Indeterminate Sets of Principles

Basic moral facts may be contingent a priori,⁹⁶ but in order for moral thought and judgement to function, an agent must have a posteriori knowledge of a sufficient amount of moral principles in order to think and judge based on those facts. Thus, the a priori knowledge of those basic moral facts depends upon there being a suitable supply of moral principles that one can learn a posteriori. That is, in order for one’s a priori knowledge of contingent moral facts to be activated and developed over time, there needs to be principles out there in one’s socio-cultural environment that one can observe, come to know, and use. Most of these principles will be defeasible such as

⁹⁵ Little. “Moral Generalities Revisited.”, 304.

⁹⁶ See Dancy (2008) “Are Basic Moral Facts Both Contingent and a Priori?” for a detailed discussion on the matter.

“tell the truth”, “don’t kill”, “everything in moderation”, “equality of opportunity for all”, and so on. How many of these defeasible principles are necessary to activate moral thought and judgement is hard to tell. Arguably, to narrow them down to a specific number, or to some specific set, is far-fetched. A more reasonable assumption is that the activation of moral thought and judgement is different for every individual. Presumably it would depend on the individual’s natural ability to learn, the efficacy and competency of her moral teachers and exemplars, and the moral principles available in her socio-cultural environment. For each individual learns at a different pace, each teacher teaches differently and showcases different levels of moral competence, and each socio-cultural environment will have different sets of defeasible principles available for one to learn.⁹⁷ Of course, different environments will have many of the same moral principles, especially if there are biologically innate ones,⁹⁸ but significant differences will presumably be present as well.

So moral thought and judgement do depend (at least in part⁹⁹) on the provision of a suitable supply of defeasible moral principles. But since defeasible principles are contingent, there is no need for them to be a part of determinate set. The claim that there *must* be a *determinate* set of defeasible moral principles in order for moral thought and judgement to be possible is problematic for at least two reasons: it assumes (a) that the same set of defeasible principles is present in all socio-cultural environments; and (b) that these principles cannot be affected/modified by the environments in which they are articulated and applied. As I see it, this is highly implausible given the wide diversity of moral cultures across the world. There are principles that simply do not exist

⁹⁷ Following Peter Railton (2017: 187): “As in the case of implicit learning of a language or of causal relations, the implicit moral understanding at which one arrives will be hostage to the quality of one’s learning environment, and all particular learning environments have limitations and biases—just as no method of learning is free of liabilities”. One’s culture, teachers, and exemplars all play a role in the quality of one’s learning environment, and each environment will have different principles, norms, prejudices, and biases.

⁹⁸ We will return to the possibility of biologically innate principles shortly.

⁹⁹ It may, for instance, also depend on the provision of a suitable supply of moral virtues, or on certain evolutionarily entrenched emotions, instincts, intuitions, and domain-general learning capacities.

in some socio-cultural environments yet exist in others such as the principles of freedom of speech and gender equality. Moreover, different cultures often have different ways of justifying certain exceptions to certain rules. Therefore, to say that there *must* be a *specific* set of defeasible moral principles to make moral thought and judgement possible is not only implausible given the wide diversity of moral cultures in the world, but also because it assumes there is an *innate* and *universal* moral structure that is unaffected by the various features and activities of any particular socio-cultural environment—an assumption that contradicts the definition of a defeasible principle. A defeasible principle is expressed by a contingent proposition, not a necessarily true one. If there is an innate and universal moral structure composed of a specific set of principles that make moral thought and judgement possible, these principles would be exceptionless ones, not defeasible ones.

Dancy, however, argues against defeasible generalizations.¹⁰⁰ According to him, a principled ethic must meet four necessary conditions: coverage, reasons, epistemology, and applicability. First, “the moral status of every action must be determined by the principles, in one way or another”.¹⁰¹ For if principles cannot fulfil this function they will fail to cover the moral landscape. Second, for each action that has a moral status, the principle must explain why it has that status. Third, principles must be learnable, either from experience or from each other by testimony. Fourth, principles must be capable of working as guides to action in new situations. For having learned them, one must then be able to apply them correctly.¹⁰²

Dancy then argues that Lance and Little’s (2004) defeasible generalizations fail to meet three of these necessary conditions, and therefore that the morality they have built around them is not really a principled one.¹⁰³ In my view, Dancy’s criteria for a principled ethic work for

¹⁰⁰ Dancy, *Ethics Without Principles*, 111-117.

¹⁰¹ Ibid., 116.

¹⁰² Ibid., 116-117.

¹⁰³ Ibid.

exceptionless principles, but are too strict for defeasible ones. For we can imagine certain types of defeasible principles that do not need to meet all of these criteria in order to contribute to the emergence and development of moral thought and judgement. There is no need for principles to explain the nature of the exceptional cases nor guide the moral agent through them in order to contribute to the emergence and development of moral thought and judgement. This is why McKeever and Ridge's taxonomy of principle-types is so useful, for there are principles that are not action-guiding standards but which still play an important role in moral learning. Dancy's criteria target only exceptionless principles, not defeasible ones. This is because he is assuming that a principled ethic must be composed of a *specific* set of moral principles that not only make moral thought and judgement possible, but that codify all moral truths. But as we saw earlier, (a) a principled ethic does not need to codify all moral truths to be considered a principled ethic;¹⁰⁴ and (b) defeasible moral principles cannot come in a specific set that is unaffected by context, for this would contradict the very definition of defeasible moral principles. Thus, by assuming that a principled ethic must be composed of a specific set of moral principles, Dancy prevents himself from considering the credibility of defeasible moral principles.

3.4. Moral Psychology and Moral Learning

So far, our argument is that moral thought and judgement depend on a determinate set of exceptionless principles, and an indeterminate set of defeasible ones. In this section, we want to provide evidence from the moral psychology literature which lends support to our argument. The first thing we need to show is that moral principles are *necessary* for making moral thought and

¹⁰⁴ See chapter 1.1. We saw that Holton's Principled Particularism is a principled ethic that claims to only codify *some* moral truths.

judgement possible. We will provide two reasons to support this claim: a philosophical one and a biological/psychological one. The evidence for the second is not set in stone, but it is worthwhile to discuss it. In short, we will point to findings that suggest that some innate moral principles exist.¹⁰⁵ As we will see, some of these innate principles might be exceptionless and universally true (e.g. the Principle of Special Reason to do Harm). In other cases, we will see that the default clauses of some defeasible principles (e.g. the principle of non-maleficence: “other things being equal, inflicting harm is wrong”) might be innate. The second thing we need to show is that the diversity of moral cultures demonstrates the *insufficiency* of a determinate set of universal and exceptionless moral principles to make moral thought and judgement possible. This will involve a demonstration of how individuals from different moral cultures reason, cognitively speaking. The third thing we need to show is that moral principles as a whole are *insufficient* for making moral thought and judgement possible. This will involve a demonstration that there are other factors that contribute to the emergence and development of moral thought and judgement such as evolutionarily entrenched moral instincts, intuitions, and domain-general learning capacities.

a) The Necessity of Moral Principles

Beginning with the first point, there are at least two reasons that show that moral principles are necessary for moral thought, judgement, and learning. The first is philosophical and the second is biological/psychological. First, following Lance and Little (2008), explanation of why a consideration is a reason to believe or do something must be somehow linked to generality.¹⁰⁶ That is, when we adduce reasons to explain why a belief or an action is right or wrong, we are

¹⁰⁵ To be clear, this is not the same as claiming that there is an innate “moral grammar” or module (see Mikhail 2011). We will argue against this theory later in this section in our discussion of moral learning.

¹⁰⁶ Lance and Little (2008: 60-61).

necessarily committed to a generalization, just as we are in other forms of reasoning (e.g. logical, epistemic, scientific etc.). It would be quite a stretch to say that moral reasoning is the only form of reasoning that doesn't depend on generalizations. Hence, just as the assumption that moral reasons behave differently than other types of reasons is ill-founded (as Dancy himself argues),¹⁰⁷ so too the assumption that moral reasoning is the only type of reasoning in which explanatory generalizations have no part is misplaced. Moral generalizations in a reasons-holistic landscape may not meet all of Dancy's criteria for a principled ethic, but this does not entail that moral agents are not committed to them or that moral thought and judgement do not depend on them. Just as we are committed to explanatory generalizations in other domains such as logic, science, and epistemology, so we are committed to them in ethics. This is because explanatory generalizations are necessary for *reasoning in general*. Although the content of what we're reasoning about will differ in these aforementioned domains, all types of reasoning involve the grouping of particular facts under certain kinds of generalizations. For this reason, moral principles understood as explanatory moral generalizations are necessary for moral thought and judgement.

The second reason for believing that moral principles are necessary for moral thought and judgement is a biological/psychological one. On the one hand, if there are innate moral principles, then moral principles are clearly necessary for moral thought and judgement. However, the question of innate principles has yielded a diversity of views on the matter. John Mikhail (2011) has argued for the existence of an innate and universal moral faculty analogous to Noam

¹⁰⁷ See Dancy (2017: "What the Particularist Believes"): "The particularist suggests that there is no reason to suppose that moral reasons function in a radically different way from other reasons. Indeed, there is a sort of presumption that they don't. That presumption is partly grounded on the fact that nobody is able to say with any confidence just which reasons are moral ones and which are not. This means that providing a radical difference between the way in which reasons of the two sorts function should seem rather peculiar. But the presumption is also partly grounded in the fact that the difference suggested by the generalist is very radical, since it affects what one might call the very logic of moral thought. To suppose that moral thought has a different logic from other thought is to adopt a bifurcated conception of rationality".

Chomsky's "universal grammar".¹⁰⁸ Jesse Prinz (2007), however, argues that morality, like all human capacities, "depends on having particular biological predispositions, but none of these deserves to be called a 'moral faculty'". He argues that "morality is a by-product—accidental or invented—of faculties that evolved for other purposes".¹⁰⁹ Finally, scholars such as Paul Bloom (2013) find themselves in the middle position. They are skeptical about the idea of an innate and universal "moral grammar" but seem convinced that there are strong candidates for innate moral principles.¹¹⁰

Beginning with the latter, Paul Bloom in his book *Just Babies: The Origins of Good and Evil*, argues that infants possess an innate—albeit rudimentary and underdeveloped—knowledge of good and evil. He argues that humans have at least four "natural endowments":

1. a moral sense—some capacity to distinguish between kind and cruel actions
2. empathy and compassion—suffering at the pain of those around us and the wish to make this pain go away
3. a rudimentary sense of fairness—a tendency to favor equal divisions of resources
4. a rudimentary sense of justice—a desire to see good actions rewarded and bad actions punished.¹¹¹

These evolutionarily entrenched endowments can be said to presuppose certain moral principles such as the principles of non-maleficence, justice, and equality. Indeed, these endowments involve emotional capacities, instincts, and intuitions. But Bloom's idea is that they might operate according to certain moral principles such as the ones just mentioned. For example, when an infant feels empathy or compassion at the suffering of others, or when it manages to

¹⁰⁸ John Mikhail. *Elements of moral cognition: Rawls's linguistic analogy and the cognitive science of moral and legal judgment*, Cambridge: Cambridge University Press, 2011.

¹⁰⁹ Jesse Prinz. "Is Morality Innate?" In *Moral Psychology: The Evolution of Morality: Adaptations and Innateness*, 367-406, Cambridge, MA: MIT Press, 2007, 368.

¹¹⁰ Paul Bloom. *Just Babies: The Origins of Good and Evil*. Crown Publishing Group: New York, 2013, cf. 18, 120. Another scholar that might fall under this view is Owen Flanagan (2017) in his book *The Geography of Morals: Varieties of Moral Possibility*.

¹¹¹ Bloom, *Just Babies*, 5.

distinguish kind actions from cruel ones, Bloom's suggestion is that there is more at play than just a feeling or an ability. According to Bloom, infants have a rudimentary knowledge of these things. Infants *know*, at a low-resolution understanding, that kind actions are good, and cruel ones are evil. They *know* that good actions should be rewarded and bad actions punished, and so on.

Greene and others (2001) have demonstrated that human beings have, evolutionarily speaking, developed a revulsion to physically and personally harming an innocent person.¹¹² Mikhail (2011) has built on this research and has argued further that acts such as intentionally hitting someone without their permission (i.e. harming an innocent and unwilling person) has an immediate negative moral valence that all humans respond to. This and other acts can, he argues, reveal an innate and universal moral faculty in human beings.¹¹³ Bloom is skeptical about the idea of a universal moral grammar, but he considers such acts as being strong candidates for reflecting moral rules that “transcend space and time”. Bloom writes: “if you punch someone in the face, you’d better have a damn good reason for it”.¹¹⁴ Other scholars, however, do not think there are innate moral principles. Prinz (2007), for example, argues that morality depends on specific biological predispositions, but that none of these presuppose a universal moral grammar. For Prinz, morality is a by-product of biological faculties that evolved to serve other functions. Under this view, Bloom's four natural endowments of our moral sense would not constitute predispositions that *presuppose* innate moral principles. Rather, for Prinz, our moral principles are accidental or invented to articulate in more precise terms the moral features these endowments are concerned with. Bloom might very well agree with this view, for he has not argued in favour of a complete,

¹¹² Joshua R. Greene, Sommerville, Leigh Nystrom, John Darley, and Jonathan Cohen. “An fMRI Investigation of Emotional Engagement in Moral Judgment.” *Science (Washington)* 293, no. 5537 (2001): 2105–2108. <https://doi.org/10.1126/science.1062872>.

¹¹³ Mikhail. *Elements of moral cognition*, 2011.

¹¹⁴ Bloom (2013: 18).

innate, and universal grammar.¹¹⁵ He simply believes that there may be strong candidates for *individual* innate moral principles.¹¹⁶ Arguably, the fact that our moral sense (which is comprised of emotional capacities, instincts, and intuitions) is biologically innate is entirely compatible with the claim that our moral principles are learned socially over time to improve said moral sense and to better articulate the behaviour of features in the moral landscape.

For our part, we find Mikhail's theory to be quite a stretch. There are many reasons to believe that the linguistic analogy is flawed or at least limited.¹¹⁷ The debate as to whether there are innate moral principles at all is a complicated one that goes beyond the scope of this thesis. For our purposes, it does not seem necessary to find the answer to this question. If it is true that only our emotions, instincts, and intuitions are innate and that moral principles are socially learned over the course of our moral development as individuals, it is still the case that they are necessary for moral thought, judgement, and learning. Admittedly, under Prinz's view, moral principles might not provide the grounds for the *initial emergence* of moral thought and judgement, but they are arguably still necessary for the *full development* of these latter. Without moral principles we would only have a rudimentary understanding of right and wrong. Indeed, following Bloom, it does seem as though we would be *just* babies. Moreover, Jonathan Baron (1998) has shown that when people act *only* based on their moral intuitions, they often produce sub-optimal, counterproductive, or even morally catastrophic consequences in matters of public health, public policy, and the tort system.¹¹⁸ Therefore, even if moral principles are accidental or invented by-

¹¹⁵ Ibid., 120.

¹¹⁶ Ibid., 18.

¹¹⁷ Later in this section, we will provide evidence from the moral psychology, cognitive science, and neuroscience literatures that disfavours the idea of an innate moral faculty or "module".

¹¹⁸ Jonathan Baron. *Judgment misguided: Intuition and error in public decision making*. New York: Oxford University Press, 1998.

products, they are still necessary for the full and proper development of moral thought and judgement.

For these reasons, then, moral principles are necessary for moral thought and judgement. First, explanations about why a consideration is a reason to believe or do something must somehow be linked to generalizations. This is so in all domains of reasoning, whether it be in epistemology, logic, science, or ethics. The assumption that morality has a unique form of reasoning that does not utilize generalizations is unfounded. Thus, since moral principles are generalizations that tell us right from wrong, they are necessary for moral thought and judgement. Second, if there are innate moral principles, then they are certainly necessary for moral thought and judgement. If there aren't, moral principles *in general* are still necessary for their full and proper development.

b) The Insufficiency of a Determinate Set of Universal and Exceptionless Moral Principles

Let us move on to the second point outlined in the introductory paragraph for this section. I will argue that the diversity of moral cultures demonstrates the insufficiency of a determinate set of universal and exceptionless moral principles to make moral thought and judgement possible. In addition to a determinate set of exceptionless moral principles, I will argue that moral thought and judgement also depend on an *indeterminate* set of defeasible ones. There are at least three reasons to support this claim. First, individuals are not exposed to the same set of moral principles during their moral learning (even though there will be some overlap between different individuals and cultures). Second, there are multiple different sets of principles that are capable of contributing to the emergence and development of moral thought and judgement. Third, while most, if not all, cultures will have principles that specify default clauses such as “other things being equal, lying is wrong”, they differ in the ways in which they justify certain exceptions to these defaults.

The first two reasons are self-evident. Regarding the first one: someone born in Toronto with European roots will not be exposed to the same set of moral principles as a Melanesian born in New Guinea. Indeed, there will be some commonalities (e.g. the principle of Special Reason to do Harm, and several default principles), but there will be significant differences in principles as well due to differences in culture and history. These commonalities will be greater if we compare, say, someone born in Toronto with European roots with someone born in the New York also with European roots. With respect to the second point: there are multiple different cultures in the world, each of which has their own set of moral principles, but yet the vast majority of individuals from all cultures are perfectly capable of thinking and judging in moral terms. This can only mean that *various* sets of defeasible principles can contribute to the emergence and development of moral thought and judgement in human beings, not just one determinate set of them.

Let us move on now to the third reason. Indeed, there are some default principles that are found in most, if not all, cultures such as “telling the truth is right”, “harming others is wrong”, “stealing is wrong” and a host of others. However, these principles only constitute commonalities in their simple formulations, not necessarily in their complex formulations.¹¹⁹ When we get down to details and specific cases, larger differences begin to emerge across cultures. For instance, in one culture killing might be wrong except in some cases of self-defence, but in another culture it might be wrong except in cases of self-defence *and* in cases of the death penalty. Furthermore, the fact that stealing is generally wrong is also not necessarily made right in the same exceptional

¹¹⁹ By the “simple formulation” of a principle I mean principles whose propositions only include default clauses without further conditions, qualifications, or exceptions. For example, “other things being equal, lying is wrong”, “killing is wrong”, “avoid excess and deficiency” and so on. By the “complex formulation” of a principle I mean just the opposite: principles whose propositions *do* include one or more conditions, qualifications and/or exceptions in their antecedents as conjuncts to the default clause. For example, “lying is wrong except in cases where lying could save a life”, “killing is wrong except in cases of self-defence”, and so on. Complex principles can, in theory, be quite long. We will discuss the final logical consequences of complex principles in chapter 4.

cases for different cultures. In this regard, we only have to refer to the diversity of copyright laws across different countries (compare the USA to China).

When we get down to details, specific cases, and exceptions, the above principles will often differ across cultures. Indeed, the default clauses of some of those principles might be commonalities, but this alone is not enough to conclude that such principles (i.e. in their simple, default formulations) act as a determinate set that make moral thought and judgement possible across cultures. First, given the wide diversity of moral cultures, it seems more likely that moral thought and judgement will be made possible in different ways for different individuals in different cultures. Second, the complex formulations of principles are just as important as the simple ones in the process of moral learning. We need to know that there are exceptions to defaults in order to fully understand why a particular consideration has the default valence that it has. That is, the concept of a default implies that its valence may change, and the precondition for knowing this is knowing about the possibility of exceptional features. For these reasons, the possibility of moral thought and judgement for a given individual in a given culture will depend on the simple and complex formulations of his culture's defeasible moral principles. More specifically, it will depend on the specific contextual justifications that are used in his culture and by his moral teachers to support certain beliefs and actions, and how well those justifications are taught and learned.¹²⁰

These three reasons can be supported by the moral psychology, cognitive science, and neuroscience literatures. Henrich et al. (2010) have compiled a large body of research which

¹²⁰ It is worth noting that the simple formulations of our principles are the ones we are most familiar with in our day-to-day lives, and they are the ones we seem to use more often. In my view, the simpler the principles are, the better they are at guiding us across multiple situations. It is only when we arrive at specific situations with more unusual features that we often feel compelled to justify why the simple formulation of our principle is true or applicable, or why it needs to be modified by adding certain conditions, or why it must simply be rejected altogether. In my view, the possibility of moral thought and judgement depends both on the complex formulations of our defeasible moral principles, and on their simpler formulations. For the former are needed to guide us in specific and difficult situations, whereas the latter serve as general action-guiding standards across multiple situations.

suggests that, cross-culturally, there is tremendous diversity among humans in domains such as visual perception, fairness, cooperation, spatial reasoning, categorization and inferential induction, moral reasoning, reasoning styles, self-concepts and related motivations, and the heritability of IQ.¹²¹ Their thesis attacks the assumption that we can generalize across human populations based on a sample of “standard subjects”. Indeed, there may be many commonalities between diverse populations, but the breadth of the differences casts great doubt on the possibility of standardization, especially if the standard subject being put forward or assumed is WEIRD (Western, Educated, Industrialized, Rich, and Democratic).¹²² These differences in the aforementioned domains should already have us consider the possibility that there will be some degree of difference in moral principles, justifications, and reasoning across different cultures. If this is indeed the case, then moral thought and judgement will not only depend on a determinate set of exceptionless moral principles, but also on an indeterminate set of defeasible ones. For our purposes, we will discuss the findings compiled by Henrich et al. that pertain to moral reasoning. Their findings, we argue, support our argument that various sets of principles can contribute to the emergence and development of moral thought and judgement cross-culturally.

John Snarey (1985) carried out a meta-analysis with data from 27 countries to see whether there was any cross-cultural evidence for Kohlberg’s three levels in the development of moral reasoning.¹²³ While all three levels were found in WEIRD populations, very little evidence was

¹²¹ Joseph Henrich, Steven J. Heine, and Ara Norenzayan. “The Weirdest People in the World?” *Behavioral and Brain Sciences* 33 (2010): 61–135.

¹²² Henrich et al., *The Weirdest People in the World?*, 1.

¹²³ Kohlberg proposed that humans, in the evolutionary development of their moral reasoning, progressed through the same three levels. First, he argued that children begin at the “pre-conventional level”, where they interpret right and wrong as grounded in internal principles about the physical effects or outcomes of actions. Second, they progress to the “conventional level”, where morality is founded on external norms that are interpreted as upholding the social order of their group, community or society. Third, some advance even further to the “post-conventional level”, where they no longer depend on external norms for judging right from wrong, but instead begin to utilize abstract moral principles about justice and individual rights. See Kohlberg (1976, 1981).

found for the post-conventional level in non-WEIRD populations from small-scale societies.¹²⁴ Also, formal education was not found to be a prerequisite for attaining Kohlberg's post-conventional level. There are highly educated non-Western populations that do not demonstrate something akin to post-conventional reasoning. According to Al-Shehab (2002), faculty members at Kuwait University scored significantly lower on Kohlberg's schemes than the typical standards for Western educated adults.¹²⁵ Thus, moralities based on a set of abstract moral principles about justice and individual rights are far from being universal across different populations.

To be clear, this doesn't mean that there aren't any principles underlying moralities that do not use post-conventional reasoning. Just because a particular culture's morality is primarily communal, emotive, or virtue-based, doesn't mean individuals in that culture aren't implicitly using what we would call moral principles. Moral principles, as we conceive of them, are simply generalizations expressed by propositions that specify certain moral features. These generalizations need not be explicitly expressed by a culture to prove that they are necessary for moral thought and judgement. Arguably, they are at work implicitly during an individual's virtue-exhibition, and assumed or presupposed by a culture's traditions and practices.

Still, it is clear that there are great differences in moral reasoning across cultures. But these findings do not prove that there aren't any universal, exceptionless, and innate moral principles. However, they do demonstrate that different cultures have different ways of reasoning about what is right and wrong. This fact implies that, cross-culturally, moral thought and judgement depend at least *in part* on an indeterminate set of principles. These principles can be explicitly expressed (as in the case of cultures that utilize post-conventional reasoning) and/or implicitly assumed or

¹²⁴ John R. Snarey. "The cross-cultural universality of social-moral development: A critical review of Kohlbergian research." *Psychological Bulletin* 97, no. 2 (1985): 202-232. Cited in Henrich et al. (2010).

¹²⁵ A. J. Al-Shehab. "A cross-sectional examination of levels of moral reasoning in a sample of Kuwait University faculty members." *Social Behavior and Personality* 30 (2002): 813-820. Cited in Henrich et al. (2010).

presupposed (as in the case of cultures whose reasoning is based on an ethic of community, an ethic of divinity, a set of virtues, or any other ethic where abstract principles might not be verbally expressed). Simply put, our argument is that *various* sets of principles can contribute to the emergence and development of moral thought and judgement in human beings, not just one specific set that it is found universally across all moral cultures.¹²⁶

c) The Insufficiency of Moral Principles

There is plenty of evidence to show that moral principles *as a whole* are insufficient for making moral thought and judgement possible. Research suggests that other factors are needed. In addition to moral thought and judgement depending on a determinate set of exceptionless principles and an indeterminate set of defeasible ones, we will demonstrate that they also depend on evolutionarily entrenched moral instincts, intuitions, and domain-general learning capacities.

To begin, there has been a long-standing rationalist tradition in moral psychology that has pedestaled the role of reason in the workings of moral thought and judgement. Only in recent decades has the role of emotions been taken more seriously. Arguably, both reason and emotions play important roles in the emergence and development of moral thought and judgement. Our own claim is that moral thought and judgement depend not only on generalizations obtained by rationalist means, but also on other things such as instincts, emotions, and intuitions.

Consider the infamous trolley problems. The most common versions are (a) the Bystander at the Switch and (b) the Footbridge. These two versions of the trolley problem imply the same lives lost / lives saved outcomes, but they differ in a very important way.

¹²⁶ Again, we do believe that there are some universal principles (e.g. Principle of Special Reason to do Harm).

1. Bystander at the Switch

You are at the wheel of a runaway trolley quickly approaching a fork in the tracks. On the tracks extending to the left is a group of five railway workmen. On the tracks extending to the right is a single railway workman. If you do nothing the trolley will proceed to the left, causing the deaths of the five workmen. The only way to avoid the deaths of these workmen is to hit a switch on your dashboard that will cause the trolley to proceed to the right, causing the death of the single workman. Is it appropriate for you to hit the switch in order to avoid the deaths of the five workmen?

2. Footbridge

A runaway trolley is heading down the tracks toward five workmen who will be killed if the trolley proceeds on its present course. You are on a footbridge over the tracks, in between the approaching trolley and the five workmen. Next to you on this footbridge is a stranger who happens to be very large. The only way to save the lives of the five workmen is to push this stranger off the bridge and onto the tracks below where his large body will stop the trolley. The stranger will die if you do this, but the five workmen will be saved. Is it appropriate for you to push the stranger on to the tracks in order to save the five workmen?¹²⁷

In studies replicated multiple times, most people say yes to the first but no to the second. But both problems involve saving five lives and losing one. So why then do people answer differently? Joshua Greene and others (2001) suggest that human beings have, evolutionarily speaking, developed a revulsion to the idea and act of personally and physically harming an innocent person. According to them, this instinct will typically overpower actions based on a rational calculation that would favour saving more lives.¹²⁸ Greene and his team used functional M.R.I to look for signs of a conflict between those areas of the brain linked with instinct and emotions, and those devoted to calculation and rational analysis. When participants thought about dilemmas that required killing someone by physically putting their hands on them (such as in the Footbridge version of the trolley problem), several brain networks lit up. The first included the

¹²⁷ The formulations of both dilemmas are taken directly from Greene et al. (2001: "Figures & Data"), however they originate from Judith Jarvis Thompson. "The Trolley Problem." *The Yale Law Journal*, 94 (1985): 1395-1415 and Philippa Foot. "The Problem of Abortion and the Doctrine of the Double Effect." *Oxford Review*, no. 5 (1967): 1-5.

¹²⁸ Joshua Greene, Sommerville, Leigh Nystrom, John Darley, and Jonathan Cohen. "An fMRI Investigation of Emotional Engagement in Moral Judgment." *Science (Washington)* 293, no. 5537 (2001): 2105–2108. <https://doi.org/10.1126/science.1062872>

medial parts of the frontal lobes, which are associated with emotions about other people such as empathy and compassion. The second was the dorsolateral surface of the frontal lobes, which is implicated in continuous mental calculations which include, for example, non-moral reasoning such as deciding whether to get somewhere by plane or by train. The third area that lit up was the anterior cingulate cortex, which records a conflict between an impulse coming from one part of the brain and an advisory coming from another.¹²⁹

However, when participants were thinking about more impersonal and hands-off dilemmas such as Bystander at the Switch, their brains responded differently. Only the region with rational analysis lit up. This difference in brain responses suggests that our moral thought and judgement can be influenced by both rational and emotional faculties depending on the situation we find ourselves in. According to Steven Pinker (2008), “other studies have shown that neurological patients who have blunted emotions because of damage to the frontal lobes become utilitarians: they think it makes perfect sense to throw the fat man off the bridge”.¹³⁰ Together, these findings support Greene’s theory that our non-utilitarian instincts derive from an emotional impulse that overpowers the rational calculation of outcomes.¹³¹

Our mission here is not to solve the trolley problems, but simply to show that moral thought and judgement depend on more than just rationalistic moral principles. That we evolved to feel a sense of revulsion at the idea of personally and physically harming an innocent person has arguably done us much good in our evolution. Arguably, on the whole, it has probably saved many lives throughout human history despite the fact that in particular cases wherein harming the innocent might be required, it urges us to abstain. It is in these cases, arguably, that rationalistic moral

¹²⁹ Steven Pinker. “The Moral Instinct.” *The New York Times*, sec. Magazine, 2008. Accessed May 15, 2020. <https://www.nytimes.com/2008/01/13/magazine/13Psychology-t.html>.

¹³⁰ Pinker, “The Moral Instinct”, 2008.

¹³¹ Ibid.

principles should intervene. Maybe the right thing *is* to push the fat man off the footbridge; it's hard to tell. Regardless, the instinct in question undoubtedly has been and continues to be beneficial on the whole, and it clearly influences our moral thought and judgement in important ways. Interestingly, it might be the case that the principle of non-maleficence emerged from this very instinct or from a family of similar ones. It might also be the case that our own principle of Special Reason to do Harm entered our moral consciousness precisely as a result of a conflict between this instinct and situations that demand that we act against it. Although we will not delve into the connections between moral instincts and moral principles, given the aforementioned evidence, our hypothesis is that there definitely appears to be a vital connection, and that perhaps the latter evolved from the former.

With this compelling evidence in the background, it seems clear to me that moral thought and judgement necessarily depend on evolutionary entrenched moral instincts and emotions. What is more, these findings not only demonstrate that moral thought and judgement cannot merely depend on moral principles, but also that they do not merely depend on a *determinate* set of moral principles. Returning to the trolley problems, Greene and others admit as well that a solution is not obvious. The explanation they provide is very much in line with what we are claiming here:

If a solution to this problem exists, it is not obvious. That is, there is no set of consistent, readily accessible moral principles that captures people's intuitions concerning what behavior is or is not appropriate in these and similar cases. This leaves psychologists with a puzzle of their own: How is it that nearly everyone manages to conclude that it is acceptable to sacrifice one life for five in the trolley dilemma but not in the footbridge dilemma, in spite of the fact that a satisfying justification for distinguishing between these two cases is remarkably difficult to find?¹³²

¹³² Green et al. "An fMR Investigation of Emotional Engagement in Moral Judgment." 2106.

Solutions to moral cases are not always easy to come by. Furthermore, given the complexities of the moral landscape and the holism of reasons, it is seldom the case that we can access relevant and exceptionless moral principles to solve such cases. The multifarious ways in which our reason, instincts, emotions, and intuitions compete and cooperate are difficult to capture in a set of consistent and fully articulated principles. These considerations, therefore, make it highly unlikely that moral principles by themselves are capable of making moral thought and judgement possible, let alone provide us with the complete code to morality. So far, then, we have convincing evidence that moral principles are insufficient for making moral thought and judgement possible. However, let us provide more evidence.

Some scholars have gone further arguing that intuitions comprise the bulk of our moral judgements. Jonathan Haidt is renowned in the moral psychology literature for advancing what he calls the “Social Intuitionist Approach to Moral Judgement”.¹³³ His argument, in summary, is twofold: he argues (1) that moral judgment is mostly the result of quick and automatic evaluations that involve the triggering of intuitions by certain stimuli in the environment; and (2) that moral reasoning is often a post hoc construction after a judgement has already been reached.¹³⁴ Consider the following scenario devised by Haidt:

Julie is traveling in France on summer vacation from college with her brother Mark. One night they decide that it would be interesting and fun if they tried making love. Julie was already taking birth-control pills, but Mark uses a condom, too, just to be safe. They both enjoy the sex but decide not to do it again. They keep the night as a special secret, which makes them feel closer to each other. What do you think about that — was it O.K. for them to make love?¹³⁵

¹³³ Jonathan Haidt. “The Emotional Dog and Its Rational Tail: A Social Intuitionist Approach to Moral Judgement.” *Psychological Review* 108, no. 4 (2001): 814–834.

¹³⁴ Haidt. “The Emotional Dog and Its Rational Tail”, 814.

¹³⁵ *Ibid.*, 814.

According to Haidt's findings, most participants in the study immediately state that the act is wrong and are quick to provide explanations. Many were quick to mention the possibility of birth defects if a child were conceived. However, they were reminded that Julie was already using a contraceptive. Participants were also quick to mention that the siblings will be hurt emotionally by this act. However, they were reminded that the story makes it clear not only that they weren't hurt emotionally, but in fact that it made them feel closer to each other. Many also suggested that the act would offend their parents and the community. However, they were reminded that it was kept a secret. Eventually, most participants admitted, "I don't know, I can't explain it, I just know it's wrong."¹³⁶ Haidt's argument is that most people typically respond intuitively to a moral/immoral act, and only justify their arguments when they are asked questions. According to his findings, people start with a moral conclusion (which they've arrived at by some emotional response), and then they work backwards and search for reasons to support their conclusion. This is clearly not an adequate form of moral reasoning, but a judgement resulting from an intuitive or emotional response followed by a post hoc construction. With these well-established findings, we can safely say that moral thought and judgement depend in a very important way on moral intuitions, not solely on moral principles.

To be clear, Haidt mentions that these findings do not imply that moral reasoning is never a direct cause of moral judgement: "It must be stressed at the outset that the social intuitionist model is an antirationalist model only in one limited sense: It says that moral reasoning is rarely the direct cause of moral judgment. That is a descriptive claim, about how moral judgments are actually made. It is not a normative or prescriptive claim, about how moral judgments ought to be

¹³⁶ Ibid., 814.

made”.¹³⁷ The claim being made is not that moral reasoning never causes moral judgement, but simply that it is rare. Furthermore, even though most people most of the time might judge beliefs and acts based primarily on emotions and intuitions, this does not mean that moral reasoning is not what they *should* be doing. Arguably, we should aim to improve our moral judgement by learning not only to control our emotions and intuitions, but also to develop rational arguments to support our judgements. Finally, moral reasoning is also important in children’s moral learning. For example, when our child hits another child on the playground, our reaction as parents might very well be intuitive and emotional (e.g. yelling at our child), and the judgement we make as to which punishment to enforce might also be emotionally driven. However, the rational explanations we provide to our child as to why their act of hitting was wrong (if we do a good job of it, and if our child is receptive enough) will teach our child the *reasons why* it is wrong to hit other people. Moreover, the next time around our child will hopefully be able to control their impulses precisely because they have received these important explanations. For these reasons, then, moral thought and judgement—although they certainly depend on moral intuitions—also depend on moral reasoning which, for its part, involves the rational articulation of moral principles.

Let us now discuss the topic of moral learning. Recent evidence suggests that moral learning and the development of moral understanding in infants does not depend on an innate moral module.¹³⁸ Peter Railton (2017) has done a good job of gathering much of the supporting evidence for this thesis. He writes:

Research on infants provides evidence that they form non-perspectival expected-value representations of agents and actions... which help them to navigate the human

¹³⁷ Ibid., 815.

¹³⁸ See Moll (2001, 2002), Greene and Haidt (2002), Nichols (2005), Prinz (2007), Greene (2014), and Railton (2017) to name a few.

environment. Such representations can be formed by highly-general mental processes such as causal and empathic simulation, and thus afford a foundation for spontaneous moral learning and action that requires no innate moral faculty and can exhibit substantial autonomy with respect to community norms.¹³⁹

Railton argues that if domain-general learning processes lie at the core of the development of moral understanding, then it should not depend upon the “triggering of an innate moral module or external socialization into a set of norms”.¹⁴⁰ Instead, the infant’s capacities for implicit discrimination and choice along morally-relevant dimensions appear to develop integrally with causal, conceptual, intentional, and evaluative *learning in general*, not because of a unique and self-contained moral domain consisting of a specific set of principles. Greene and Haidt (2002) concur on this point in their own way. They argue that multiple regions of the brain play a role in giving rise to moral judgement. For instance, if we were to factor out emotional processing, mental imagery, and abstract reasoning in an attempt to uncover a distinctive part of the brain dedicated specifically to moral judgement (as was attempted by Moll *et al.* (2001; 2002)),¹⁴¹ there would likely be no phenomena left to study on the imaging scans. This is because there is a variety of brain processes involved in giving rise to moral judgement, both cognitive and affective ones.¹⁴²

The key insight here is not that moral principles are rendered unnecessary for moral learning, but that these various capacities needed for learning of *all* types (whether it be social, moral, or linguistic learning) *precede* moral principles and norms. In other words, the acquisition of moral principles by the infant presupposes the infant’s ability to form non-perspectival expected-value representations of agents and actions that help it navigate the human environment

¹³⁹ Peter Railton. “Moral Learning: Conceptual Foundations and Normative Relevance.” *Cognition* 167 (2017): 172.

¹⁴⁰ Railton, “Moral Learning”, 176.

¹⁴¹ Jorge Moll *et al.* “Frontopolar and anterior temporal cortex activation in a moral judgment task: preliminary functional MRI results in normal subjects.” *Arq. Neuropsiquiatr.* 59 (2001): 657–664; “Functional networks in emotional moral and nonmoral social judgments.” *Neuroimage* 16 (2002): 696–703.

¹⁴² Joshua Green and Jonathan Haidt. “How (and where) does moral judgement work?” *Trends in Cognitive Sciences* 6, no. 12 (2002): 522-523.

as a whole. That is, the infant must already have begun to develop discriminative abilities to factor experience and behavior into such categories as harm, benefit, risk, cooperation, and intention vs. accident, *before* being able to grasp and understand moral principles and norms. Railton insists that this must be so otherwise transgressions of moral norms would not be properly detected by the infant, and principles that imply helping others or punishing violators would not be properly understood or applied.¹⁴³

Our own account fits very well with this research in moral learning for three reasons. First, our account does not claim that there is an innate “moral module” consisting of a specific or determinate set of principles *sufficient* for moral learning. This is because we accept that moral thought and judgement can be made possible not only by different sets of principles depending on one’s socio-cultural environment, but also by other factors such as instincts, intuitions, and domain-general learning capacities. In regards to the claim that moral thought and judgement depend on an indeterminate set of defeasible principles, the evidence for this is quite simply the cultural and moral diversity in the world. As we saw earlier, different moral cultures have different moral principles (with some commonalities of course), yet the vast majority of people are perfectly capable of thinking and judging in moral terms. This can only imply that multiple sets of principles can contribute to making moral thought and judgement possible.

Second, although we have argued that there are some exceptionless moral principles necessary for moral thought and judgement, these need not be part of some innate, universal, necessary, and sufficient moral module known a priori. As we have already argued, exceptionless principles are not sufficient for making moral thought and judgement possible; we also need a sufficient range of defeasible principles, an innate set of domain-general learning capacities, moral

¹⁴³ Railton (2017: 176).

instincts, and moral intuitions. The exceptionless principles we have defended in chapter 2 such as the Principle of Special Reason to do Harm and Holton's *That's it* principle, may very well be innate moral principles. However, because of their insufficiency for making moral thought and judgement possible on their own, they cannot, by definition, be part of an innate moral *module* understood as a self-contained set of principles that is independent from a more general learning module. In fact, in chapter 4, we will argue that the knowledge of these exceptionless principles is dependent upon the knowledge of defeasible ones.¹⁴⁴

Third, our account assumes the need for an a posteriori dimension for the acquisition of moral principles and for the possibility of moral thought and judgement (for instance, our claim that defeasible moral principles are necessary for moral thought and judgement) while at the same time not rejecting the credibility of innate learning capacities.

3.5. Relativism

How, then, does our account avoid the problem of relativism? One way in which we avoid it is to the extent that we accept that there are at least a few exceptionless and universally held moral principles. To be clear, we are not claiming that moral thought and judgement depend *solely* on an indeterminate set of defeasible principles. We are not saying that each culture has different principles, and therefore that there are no universal principles. Indeed, that would be to commit the error of relativism. Our argument rather is that moral thought and judgement depend on both a

¹⁴⁴ For example, the Principle of Special Reason to do Harm—which is actually grounded in a deeper principle that we might call “special reason to do something that is default immoral”—is dependent on our knowledge of some sufficient range of wrong-making features about certain actions. For instance, lying, killing, and stealing are default wrong because we have collected large amounts of evidence over thousands of years that suggests that they cause significantly more harm than good. As a result, we have developed defeasible principles that specify these default valences. Now, the exceptionless principle that states that one needs a special reason to perform actions that oppose these defaults, albeit necessarily true and exceptionless, is clearly dependent on the a posteriori knowledge we have gathered about these defaults.

determinate set of exceptionless principles and an indeterminate set of defeasible ones. In other words, we are arguing that, across different cultures, there is some degree of overlap in principles. Some of that overlap is necessarily in exceptionless principles, but some of it is also in default principles. For instance, the default principle “other things being equal, lying is wrong”, even though it is only the simple formulation of a defeasible principle, is clearly present in most if not all cultures. Of course, there will be differences in what one culture considers a valid exception to the rule as compared to another, but the default clause does seem to be present in most (probably all) cultures. So we might say then that there are also default clauses that are universal, and therefore that we avoid relativism in this way as well.

How, then, does our account adjudicate moral disagreements between different cultures? The first thing to stress is that failure to answer this question completely is not to fall into relativism. That we apparently *can* adjudicate *all* moral disagreements between different cultures assumes that there is a complete moral theory that can fully articulate and codify all possible moral truths. In chapter 4, we will argue against the validity of this assumption. To give a brief answer: our claim is that, in a reasons-holistic moral landscape, it is not possible to capture all possible moral truths and fully articulate the ways in which reasons may vary across contexts in a way that is theoretically and practically manageable. Consequently, it is impossible to adjudicate all possible moral disagreements. There will always be some situations where we will not be able to present the right principle for the right disagreement that fully articulates why one culture is right and the other is wrong. However, to be clear, we are not arguing that *no* moral disagreements can be adjudicated. That indeed *would be* to commit the error of relativism. Our view rather is that it is possible to adjudicate *some* disagreements *some* of the time.

First, there are obvious cases such as when one culture views a particular group of people as meriting death due to their skin colour. Clearly, this is wrong, and cultures who believe in the opposite view according to which no one merits death based on race are in the right. Here, the disagreement can be perfectly adjudicated. But of course that is an easy case. In regards to more complex cases, let's take the example of retirement homes for seniors.

In Canada, retirement homes are very common and we typically don't seem ethically bothered by them. But in China, retirement homes seldom exist largely for traditional and ethical reasons. The duty of "filial piety",¹⁴⁵ for instance, is deeply entrenched in Chinese history. The idea of sending one's parents away to be taken care of by strangers when said parents could easily be taken care of at home by their children certainly does not honour the duty of filial piety. Indeed, there are certainly cases where we cannot take care of our parents due to severe illness, and as result we must send them to *nursing homes*. But nursing homes are quite different from retirement homes. In China, while the former are necessary and morally acceptable, the latter are largely an infringement upon the duty of filial piety. Despite this moral infringement, in the last decade China has had to develop various elderly care policies to keep up with its aging population. In fact, the proportion of its population over the age of 65 is expected to double from 2010 to 2030.¹⁴⁶ It is only recently that elderly care centers other than nursing homes have been a subject of consideration.

The Chinese Government elderly care policy is governed by a 90/7/3 formula, meaning it aims for 90 per cent of seniors to remain at home, 7 per cent to stay at intermediate facilities and 3 per cent at nursing homes. Bearing in mind the Chinese tradition for people to stay

¹⁴⁵ For a philosophical explication of filial piety, see Confucius' *Analects* in Robert Eno. *The Analects of Confucius : An Online Teaching Translation*, 2015. [http://www.indiana.edu/~p374/Analects_of_Confucius_\(Eno-2015\).pdf](http://www.indiana.edu/~p374/Analects_of_Confucius_(Eno-2015).pdf). For a more historical and encyclopedic explication, see Aris Teon. "Filial Piety (孝) in Chinese Culture." *The Greater China Journal*, 2016. <https://china-journal.org/2016/03/14/filial-piety-in-chinese-culture/>.

¹⁴⁶ Elena Glinskay and Feng, Zhanlian, *Options for Aged Care in China*, World Bank, 2018, accessed 12th September 2019, <<http://documents.worldbank.org/curated/en/171061542660777579/Options-for-Aged-Care-in-China-Building-an-Efficient-and-Sustainable-Aged-Care-System>>

at home as long as possible, there is substantial potential for providing home care services in various forms. Those services could be on a one-on-one basis, but also for instance through neighbourhood day-care centres, where the elderly can stay during the daytime and return to their home or that of their children at night. For children, this might be a way of fulfilling their ‘filial piety’ by allowing them to bring their parent(s) to and from the centres.¹⁴⁷

Let’s apply the Chinese concept of filial piety to retirement homes in Canada. The Chinese moral argument seems to be that, on account of the fact that people in retirement homes can be taken care of by their children, the children should fulfil their responsibility to their parents by putting in the necessary work rather than taking the easy way out in an attempt to avoid being inconvenienced. So the question is: are retirement homes an abdication of our moral responsibility to our parents?

As Canadians we might disagree and argue that retirement homes are generally better at taking care of the elderly because of their resources and trained staff. We might also argue that the elderly who are in retirement homes have usually voluntarily agreed to be there. We might also add that many of our elderly are happy because (a) they have some degree of independence; (b) the daily activities they have a harder time doing can be done by trained staff; (c) they are in a home with people of their own generation and perhaps even many of their friends; and (d) they can take part in events organized by the home they otherwise wouldn’t have been able to take part in. On the face of it, these reasons seem rather convincing.

However, a Chinese person might argue that retirement homes are not only an abdication of our moral responsibility to our parents, but also that all of the benefits outlined in a, b, c, and d can be achieved at home with the children or elsewhere in the community. Independence can easily

¹⁴⁷ EURObiz. “China and Its Elderly Care System,” October 17, 2019. <https://www.eurobiz.com.cn/china-and-its-elderly-care-system/>.

be fostered at home; the activities that are harder to manage can be done by the children; our elderly parents can meet with their friends and other people their age at other locations in town such as the theater, the church, the community center, the shopping center, the arts centre, the fitness centre, and so on; and finally, residency at a retirement home is not a prerequisite for participation in local and community events.

So who's right? It's hard to tell. Both arguments make strong cases. Can our account of Principled Particularism adjudicate this moral disagreement? And if it cannot, can we still avoid invoking what we might call "bad relativism"? Sometimes the answer to a clear-cut question is not so clear-cut; and sometimes a clear-cut answer simply cannot be provided. I would argue that both arguments can be morally acceptable in their own respective countries/cultures. Both countries have political processes that involve both explicit and implicit contracts with their citizens—processes by which their citizens can meaningfully participate in the creation and operation of retirement and nursing homes. There are also important possibilities to consider. In some cases, there may be Canadians who are taking the easy way out by sending their parents to retirement homes; in others, the parents have not only agreed but are excited to go, even when their children have vowed to take care of them. On the other side, in some cases, there may be Chinese people who are taking care of their parents and who are doing a poor job of it; while in other cases, they're doing a better job than most retirement homes ever could. Hence, when it comes to sending or not sending our elderly parents to a retirement home, there seem to be multiple right ways of acting depending of course on the particular situation we find ourselves in. Undoubtedly, free voluntary agreement is a crucial moral consideration, but so is the duty to one's parents and the variability of cultural values. The assumption that there *must* be or *is* only One Way to adjudicate a moral disagreement should be treated with great skepticism, even though there may be some cases in

which this assumption is true. Therefore, we should look at specific cases rather than the moral validity of retirement homes as a whole.

Suppose we have a case where the parents want to go to a retirement home. In China, arguably, the duty of filial piety would require that the children honour their parents' decision. Indeed, the children can, and probably should, try to convince their parents to stay at home with them so that they may fulfil their duty of filial piety. But if their parents insist that they absolutely want to live in a retirement home, the duty of filial piety changes sides and it becomes quite clear that the children should respect their parents' wishes. In Canada, the children should also respect their parents' wishes, even though they should certainly try to convince their parents to stay at home with them if they are capable of taking care of them. On the other hand, in a case where the elderly parents want to stay home with their children, and their children are perfectly capable of taking care of their parents, arguably the children should respect their parents' wishes and take on the responsibility rather than taking the easy way out. In this case, compassion, honouring one's parents, and the virtue of sacrifice seem to trump the children's preference and convenience. Of course, if such a situation would render the children extremely unhappy, then it may be morally acceptable for them to respectfully tell their parents to find somewhere else to live. But arguably, in most cases, the children would do and be just fine, even if some of their freedom would be comprised and some of their responsibilities increased.

Our account, then, adjudicates moral disagreements by examining the particularities of contexts and seeing which moral considerations come out on top in those contexts. Sometimes culture A has it right, other times culture B does. Sometimes culture A has one part right, and culture B has the other part right, and all that is needed is a kind of synthesis. Sometimes both cultures are acting rightly in their own ways, but perhaps not in the best possible way. Sometimes

the best possible way is discoverable, other times it is unclear whether it can be discoverable at all. But just because we cannot discover the best possible way, does not necessarily imply that we are acting wrongly or even sub-optimally (for perhaps, in a given case, we are acting in the best possible way without knowing it). Moreover, just because we may not be able to adjudicate a particular moral disagreement between two cultures once-and-for-all by means of an objective point of view that exists beyond said cultures, does not mean we cannot make improvements through intercultural dialogue. Indeed, sometimes the objectively best way is simply not known to us. Sometimes all we have left is cultural and intercultural dialogue, and the strategy of weighing the conditions, pros, and cons of our moral beliefs and actions in particular situations.

3.6. Conclusion

We have argued that moral thought and judgement necessarily depend on a determinate set of exceptionless moral principles and an indeterminate set of defeasible ones. We saw that the Principle of Special Reason to do Harm is necessary to avoid the justification of the immoral, and therefore must be exceptionlessly true and universally held cross-culturally. We also saw that defeasible principles are perfectly compatible with reasons-holism to the extent that they specify variant reasons whose valences are contingent on contextual features.

We then provided two reasons to believe that moral principles are necessary for the possibility of moral thought and judgement. First, we argued that explanations about why a consideration is a reason to believe or do something must be somehow linked to generalizations. Thus, since moral principles are generalizations that tell us right from wrong, they are necessary for moral thought and judgement. Second, we argued that moral principles *in general* are necessary

for the *full and proper* development of moral thought and judgement, even if there aren't innate principles capable of igniting the *initial emergence* of moral thought and judgement.

We then provided three reasons to support our argument that the diversity of moral cultures demonstrates the insufficiency of a determinate set of universal and exceptionless moral principles to make moral thought and judgement possible. In this regard, we argued that moral thought and judgement also depend on an *indeterminate* set of defeasible moral principles. First, we clarified that individuals are not exposed to the same set of moral principles during their moral learning (even though there will be some overlap across different individuals and cultures). Second, we noted that there are multiple different sets of principles capable of contributing to the emergence and development of moral thought and judgement. Third, we indicated that while most, if not all, cultures have principles that specify default clauses such as "other things being equal, lying is wrong", they differ in the ways in which they justify certain exceptions to these defaults.

We then provided substantial evidence to show that moral principles *as a whole* are insufficient for making moral thought and judgement possible. In this regard, we cited research suggesting that other factors are needed. In addition to moral thought and judgement depending on a determinate set of exceptionless principles and an indeterminate set of defeasible ones, we demonstrated that they also depend on evolutionarily entrenched moral instincts, intuitions, and domain-general learning capacities.

Finally, we discussed the issue of moral relativism and explained why our account does not fall prey to it. We argued that insofar as our account admits of the existence of exceptionless and universally held moral principles, it consequently does not fall prey to what we might call "bad relativism" according to which objective truth is impossible. We also noted that while our account does not, in theory, adjudicate *all* possible moral disagreements, *it is* capable of adjudicating *some*

moral disagreements *some* of the time. Despite this inherent incompleteness, we noted that this too is not equivalent to falling prey to so-called “bad relativism”.

Chapter 4

4.1. Introduction

Can moral principles codify all or most of morality in finite and manageable terms? In this chapter, I will argue that they cannot. However, I will seek to demonstrate that some *particular* codification is possible. The first section will begin by describing Richard Holton’s Principled Particularism in order to give the reader a general idea of what PP is all about (i.e. its central claims and assumptions). The last section will set out to describe our own modified version of PP. The sections in between will comprise our arguments against the other positions in the debate. Throughout these middle sections we will gradually build an answer to the codification question.

4.2. Holton’s Principled Particularism

According to Richard Holton, Principled Particularism is the thesis that: “There is no finite set of finite principles that serves to axiomatize ethical evaluation: that is, no finite set of finite principles, such that, given any action fully described in non-moral terms, the principles and the description entail a given moral verdict if and only if it is true”.¹⁴⁸ This is quite a difficult thesis to wrap our heads around so we will need to unpack it in more detail.

¹⁴⁸ Holton, “Principles and Particularisms”, 193.

A simpler characterization of PP would be: *the thesis that there is no finite set of finite principles that can fully codify morality*. That is, no matter what set of principles we come up with, there will always be moral truths that will escape their grasp. To help us understand this in more precise terms, Holton provides us with two clarifications. First, the Principled Particularist denies that there is any set of principles from which we can *deduce* which action is right given its description in non-moral terms. Second, the Principled Particularist also denies the possibility of deducing whether an action is right given its characterization in terms of thick ethical concepts.¹⁴⁹

But now, what does Holton mean by “there is no *finite* set of *finite* principles”? According to Holton, since particularists accept the supervenience of the moral on the descriptive (i.e. “the thesis that any two situations that are identical in their descriptive properties will be identical in their moral properties”),¹⁵⁰ it would be possible, in theory, to give a complete and consistent axiomatization of morality by listing every possible action and every true moral verdict. But since there is an infinite number of possible actions, this list would also have to be infinite.¹⁵¹ But now, since a finite list of moral truths could, in theory, be captured in a single principle (so long as that principle is infinitely long), Holton’s position is that “particularists will have to insist that their claim is that there is no finite set of *finite* principles”¹⁵² that can axiomatize morality. In other words, only an infinite principle would be able to capture all moral truths; but of course, such a principle could never be grasped by finite minds, thereby making it impractical for the moral agent. Thus, according to Holton, “particularists are committed to thinking that there is no one set of true principles that entails, and hence justifies, each true moral verdict”.¹⁵³ But to be clear, this does

¹⁴⁹ Ibid., 193, footnote 2.

¹⁵⁰ Ibid., 192.

¹⁵¹ Ibid., 193.

¹⁵² Ibid., 193.

¹⁵³ Ibid., 194.

not mean that we are thereby committed to the claim that *no* true moral verdicts are entailed by *any* moral principles.¹⁵⁴ That is, the mere fact that a finite set of finite principles does not capture all moral truths does not mean that some moral truths cannot be captured by *any* set of moral principles. Rather, for principled particularists, the idea is that “different moral verdicts will be entailed by different sets of principles; but there is no one set that will entail them all”.¹⁵⁵ So PP argues that we can only achieve a particular and limited codification of the moral landscape.

Other scholars disagree and argue that most or all of morality can be codified. But the question is: what would it take to achieve general or complete codification without invoking a set of unmanageable principles? The generalist needs to develop a finite set of finite principles that can codify most or all moral truths in manageable terms. In what follows, we argue that generalists have failed to develop such a view.

4.3. Rejection of Transcendental Generalism

According to Dancy (2017), there are two strategies the generalist can try to use in order to ensure that his principles remain both exceptionless and capable of codifying all moral truths.¹⁵⁶ First, the generalist can complicate his principles by including exceptional features into the accounts of the invariant reasons they seek to specify. McKeever and Ridge have themselves alluded to this strategy: “Even if moral reasons are context-dependent, it might well be possible for us to track those reasons only if we had access to a moral principle which articulated exactly how those reasons vary from one context to the next”.¹⁵⁷ McKeever and Ridge go on to argue in

¹⁵⁴ Ibid., 194-195.

¹⁵⁵ Ibid., 195.

¹⁵⁶ Jonathan Dancy. “Moral Particularism.” In *The Stanford Encyclopedia of Philosophy*, Winter Edition 2017, see “The Generalists’ Reply”.

¹⁵⁷ McKeever and Ridge, *Principled Ethics*, 37-38.

chapter 7 of *Principled Ethics* that we can and *need* to develop exceptionless principles “whose defeating and enabling conditions have been *fully* and finitely articulated in descriptive terms”.¹⁵⁸ If successful, this first strategy of principle-complication would allow the generalist to develop exceptionless, necessarily true generalizations capable of codifying all of morality. However, if this first strategy turns out to be unsuccessful, Dancy notes that the generalist can try to use another strategy. Second, the generalist can develop a core-periphery view according to which a core of invariant reasons can account for the existence and status of a periphery of variant ones.¹⁵⁹ Dancy argues that both of these strategies are incapable of saving the generalist, and we agree.

First, the strategy of complicating one’s principles in order to guarantee that they remain necessarily true will fail to codify morality in finite and manageable terms. This is because, in order to effectively do so, one would have to render principles so complex that moral agents would not be able to understand them as theoretical standards, let alone use them as guides to action in normal and novel cases. To illustrate this point, suppose one promised to do something. If the fact that one promised is sometimes not a reason in favour of doing it, we can expect that there will be an explanation for this. Suppose the explanation is that what one promised to do was immoral. Now the reason in normal cases is that one promised to do it and it is not immoral. But of course, one can object here by claiming that this reason is not truly invariant, for there is always the possibility of being coerced into promising. So one would have to complicate the principle further in order to ensure it remains exceptionlessly and necessarily true. Thus, in order to account for this possibility, the reason in normal cases now reads as follows: one promised to do it, what one promised is not itself immoral, and one was not coerced into making it. As we can see, the strategy here is to include all possible exceptional features into one’s account of the reason in question. But

¹⁵⁸ Ibid., 177. Emphases in italics are my own.

¹⁵⁹ Dancy (2017), “The Generalists’ Reply”.

this seems like a hopeless task, as there is no obvious end to this process of complication.¹⁶⁰ Arguably, this process seems to reveal the variability of reasons, not their invariability. Despite this, the generalist might still argue that, eventually, one will in fact reach this end point in finite and manageable terms and that it will be impossible to provide a counter-example to the fully articulated and exceptionless principle. Our position is that even if we reach this end point in finite terms, manageability is a whole other thing. The complication of principles can get so long and complex that moral agents will not only be unable to apply them usefully as guides to action, but will also fail to grasp them appropriately as standards.

For example, if we want to guarantee the exceptionless status of the simple principle “lying is wrong”, one option would be to include all the possible exceptions into the formulation of the principle. Consider the following principle:

Lying is wrong except in cases where telling the truth would cause more harm, or where the only way to tell the truth would be to do something more immoral, or where it would save a life, or save someone from physical injury or psychological pain, or extend someone’s life, or prevent unnecessary conflict between people, or if it was done under coercion.

In essence, this process of complication may go on indefinitely. What is more, for each exception to the default, there would need to be certain specifications to ensure that the exception in question is in fact relevant in this case. So for example, suppose that in a given case the default clause “lying is wrong” is disabled by the exceptional feature “the only way to tell the truth *here* is to do something more immoral than lying”. First, we would have to specify how that disabler was enabled (this would include explaining why lying is, in this case, the lesser of all evils; or, in other words, why the other possibilities are more immoral than the act of lying), and why that

¹⁶⁰ Ibid.

enabler is not itself disabled by some other feature. Indeed, we would have to do this for each possible exception outlined in the above principle. Also, the total number of possible exceptions is arguably much higher than what we have here.

Despite this complicated process, McKeever and Ridge have recommended a similar strategy as a way to track the variance of moral reasons from context to context.¹⁶¹ They argue that unhedged (i.e. exceptionless) principles are capable of codifying all of morality and that we have good reasons to articulate them.¹⁶² First, let us recall that their default, “hedged” principles claim “that a given action has a given moral property only if no further feature of the situation explains why it does not”.¹⁶³ Conceived as such, the principle “all other things being equal, lying is wrong”, for instance, is an example of a hedged principle. Hedged principles are not capable of codifying all of morality because they are not exceptionless generalizations. McKeever and Ridge admit themselves that “default principles are insufficient to guide action without an independent ability to judge whether relevant defeaters are present”.¹⁶⁴ McKeever and Ridge go on to argue that we can and *need* to move from hedged to unhedged moral principles “whose defeating and enabling conditions have been *fully* and finitely articulated in descriptive terms”.¹⁶⁵ They argue further that “the person of practical wisdom is a person in command of a set of finite and unhedged principles which codifies all of morality (or at least, the entire portion of morality of which knowledge is possible) in descriptive terms. So if practical wisdom is possible then morality can be captured in a finite set of unhedged principles with purely descriptive antecedents”.¹⁶⁶

¹⁶¹ McKeever and Ridge, *Principled Ethics*, 37-38.

¹⁶² *Ibid.*, 192-196.

¹⁶³ McKeever and Ridge, *Principle Ethics*, “Abstract and Key Words”, 113. For the abstract, see the version of the book on Oxford Scholarship Online.

¹⁶⁴ McKeever and Ridge, *Principled Ethics*, 138.

¹⁶⁵ *Ibid.*, 177. Emphases in italics are my own.

¹⁶⁶ *Ibid.*, 140.

But as we have seen above, this strategy of *fully* articulating the exceptional features that may defeat or disable the default clause of the principle is a fool's errand. The principle becomes so complex that it becomes unmanageable from the point of the view of the moral agent. Furthermore, this process of complication or of "full articulation" to ensure the necessary truth of the principle actually seems to suggest that holism is the correct view of moral reasons rather than atomism or some hybrid view of reasons. As we have argued in chapter 2, given that the majority of the strongest candidates for invariant reasons can be shown to vary, why should we expect all other reasons to be invariant? Given the fact that we can almost always come up with counterexamples to supposedly invariant reasons—even their strongest candidates—it seems highly unlikely not only that atomism is true, but also that holistic reasons can be fully captured by principles. Therefore, this process of complication in order to guarantee the exceptionless status of principles doesn't seem very promising. The principle as an action-guiding standard will become too complicated to guide the moral agent, and eventually too complicated to even grasp as a standard. Maike Albertzart (2014), a moral generalist himself, agrees with us:

McKeever and Ridge's account of unhedged moral principles is too demanding. So far no satisfactory set of unhedged principles has been presented and it is telling that McKeever and Ridge themselves do not offer even an example of an unhedged moral principle. (...) even if there were a set of unhedged principles which codifies all of morality, yet to be discovered, it would be highly complex. In order to cover all different circumstances such principles would have to include various exception clauses that would render their antecedents unmanageably long.¹⁶⁷

We agree that reformulating our principles in order to accommodate the information gathered from new contextual experiences is worthwhile, but only to the extent that we admit the defeasible nature of our principles and keep them relatively simple and easy to grasp. I think the

¹⁶⁷ Maike Albertzart. *Moral Principles*. London, UK and New York, USA: Bloomsbury, 2014, 74.

lesson here is that we should not attempt to *guarantee* the necessary truth-status of principles. This is why I have suggested that *defeasible* principles expressed by contingent propositions (which can include hedged or default principles) should be preferred to exceptionless ones expressed by necessarily true propositions. Although defeasible moral principles will not be able to codify the entire moral landscape, this does not mean that we cannot make moral progress by means of them.

The second strategy the generalist can employ is to adopt a core-periphery view of moral reasons. That is, he can save his generalism by restricting himself to a specific group of invariant reasons (a moral core) that somehow explains the existence, status, and behaviour of the variant reasons on the periphery. More specifically, because (a) atomism is likely to be false inasmuch as the majority of the strongest candidates for invariant reasons can be shown to vary by counterexample; and (b) the strategy of complicating one's principles to guarantee their exceptionless status is a fool's errand; therefore (c) the transcendental generalist, if he wants his generalism to achieve complete codification of the moral landscape, seems forced to abandon both atomism and holism and to adopt a hybrid view of moral reasons.¹⁶⁸ Unfortunately, as we will see, this hybrid view is problematic. But first, let us provide an example of such a view.

William David Ross distinguishes between derived and underived *prima facie* duties. The underived duties are fidelity, reparation, gratitude, the promotion of the general good, and non-maleficence.¹⁶⁹ Other duties are derived from these core, underived ones. For instance, the duty of non-maleficence is *prima facie* more fundamental and binding than the duty of beneficence,

¹⁶⁸ If this hybrid view fails to achieve complete codification, the generalist will have no option left to achieve complete codification. As we will soon see, the hybrid view will fail to do so, and the generalist will be forced to pursue a *general* codification of the moral landscape instead.

¹⁶⁹ William David Ross. *The Right and the Good*. Edited by Philip Stratton-Lake. New York: Clarendon Press Oxford, 2002, 21.

because otherwise acts that are generally wrong but made out of beneficence such as killing someone to save another, or stealing from one to give to another, would be *prima facie* right; and, of course, we should not see such acts as generally right, even though there may be some exceptional circumstances in which they are. Therefore, to avoid these immoral justifications, one should first have a duty not to harm others, and then a duty to be benevolent. For these reasons, beneficence is derived from non-maleficence.¹⁷⁰ A similar logic applies to the other duties. For Ross, the point is that there is a core of underived and invariant *prima facie* duties from which all other variant duties are derived. Dancy provides his own example in regards to promising:

I might have a duty to go up to London today to see my son Hugh. But this duty is derived from a general duty to do what I have promised to do. As we might put it, that Hugh is expecting to see me today sometimes gives me a reason to go up to London and sometimes does not; it is a derived, and therefore variable, reason. If it does give me a reason, it will be because it is keyed in some way into an unvariable, underived reason. So derived reasons are variable, and underived ones invariant.¹⁷¹

Clearly then, counterexamples will only be effective at refuting generalism if they target the supposed *underived* invariant reasons upon which it rests, not the variant ones that are derived from them.¹⁷² The problem with this hybrid view of reasons, however, is that while there may exist some invariant reasons, they are so few and far between that they will not be able to fulfil their ‘core’ role (no pun intended). The core of underived and invariant reasons needs to be able to explain the existence, status, and behaviour of the variant reasons on the periphery. As we discussed in chapter 2, many of the strongest candidates for invariant reasons were shown to vary in certain contexts. So it is unlikely that these rare and scattered invariant reasons will be able to do the work required of them. There are two relevant points to make here.

¹⁷⁰ Ross, *The Right and the Good*, 21-22.

¹⁷¹ Dancy (2017), “The Generalists’ Reply”.

¹⁷² Ibid. Cf. McNaughton and Rawling (2000: 266-272) who concur on this point as well.

First, although some of these candidates for invariant reasons may be underived, that they are does not necessarily entail that they are invariant. In other words, the fact that a reason is underived is insufficient for demonstrating that it is invariant. The moral reasons specified in Ross' underived prima facie duties, or in the categorical imperative, or in the principle of utility, can all be shown to vary in certain contexts. Second, just because there are a few invariant reasons in the moral landscape does not necessarily entail that *all* variant reasons can be derived from them, or that all moral truths can be codified transcendentally or context-dependently by means of them.

Take for example Holton's *That's it* principle. Recall that we agreed that it specifies an invariant reason for belief. *That's it* holds that, for all situations in which we have compelling evidence to conclude beyond a reasonable doubt that there are no further relevant features that supersede those which appear in our moral argument, we have an invariant reason to believe that the conclusion of our moral argument is true. However, we noted that it is usually quite rare that we know this fact beyond a reasonable doubt; and since *That's it* only covers the cases in which we do, its coverage of the moral landscape is quite limited. We also noted that the continual denial of the *That's it* clause actually proves the variability of reasons described in holism. Here, a generalist who accepts *That's it* might be tempted to say that it therefore codifies this variant periphery and thereby fulfils its 'core' role under the hybrid account of reasons. Unfortunately, this leap in logic is unwarranted. Although the continual denial of *That's it* proves the *existence* of variant reasons, it does not explain their status and behaviour across contexts. That is, it does not in any way 'track' or fully articulate the behaviour of variant reasons on the periphery. Thus, although the reason for belief specified in *That's it* is underived and invariant, it does not by itself explain the status and behaviour of variant reasons, let alone codify most or all of morality.

The same is true of the Principle of Special Reason to do Harm. Even though it specifies an underived and invariant reason, it is by itself insufficient to explain the behaviour of the variant reasons on the periphery. In chapter 2, we noted that this principle is open to the possibility that the status of certain reasons can vary according to context. In other words, the principle assumes that there are special features in certain situations that are reasons to believe or do the opposite of the default. We then argued that this principle presupposes the variability of default reasons, and therefore implicitly accepts holism. In this regard, we had noted that it is an invariant exception that proves the rule of variability. But an exception that proves the rule is far from being a core that proves the periphery. Already, this notion sounds unlikely and counterintuitive. The principle that one needs a special reason to do harm may help explain the existence and status of default reasons related to harm-causing, but it will not explain the various ways in which they may vary across contexts, nor will it be able to codify other default reasons unrelated to harm-causing. So even though the Principle of Special Reason to do Harm may provide more coverage than *That's it*, its coverage is still insufficient for general or complete codification of the moral landscape.

So it seems the transcendental generalist has nowhere to go. First of all, it is highly ill-advised for him to support his generalism with atomism, given that the majority of the strongest candidates for invariant reasons can be shown to vary. Second, if holism is true, it is impossible for him to codify all or most of morality in finite and *manageable* terms by means of exceptionless principles. And third, it is also ill-advised for him to seek support from a hybrid view of reasons given that there are not enough invariant reasons that can work together as a determinate core to explain the existence, status, *and* behaviour of the variant reasons on the periphery. It seems, then, that we have strong reasons to believe that Transcendental Generalism is false.

4.4. Rejection of Contextual Generalism

Contextual Generalism (CG) is the position McKeever and Ridge would have to defend if holism turns out to be true. Since we have provided strong reasons to believe that holism is true and that Transcendental Generalism is false, McKeever and Ridge's "generalism as a regulative ideal" does seem forced to adopt a contextual view of codification. As we saw earlier, McKeever and Ridge do believe that the codification of context-dependent or holistic reasons is possible.¹⁷³ CG must therefore prove that it is compatible with holism while simultaneously providing a determinate set of principles that can codify most of morality. Insofar as it endeavours to build upon holism, it can, in theory, admit of some invariant reasons and therefore of some exceptionless principles. CG needs to show that some combination of moral principles can account for the status and behaviour of variant reasons across contexts *without* assuming atomism or hybridism.

In order for the contextual generalist to meet these demands he must formulate a dual account of principles: (1) an account of defeasible principles that can codify all default reasons; and (2) an account of exceptionless principles that can articulate, in finite and manageable terms, *common patterns* according to which those defaults may vary from context to context. A general starting question the contextual generalist might ask is: how are default reasons disabled or defeated most of the time? A more specific question would be: how is the default wrong-making status of lying disabled or defeated most of the time? Other questions the contextual generalist might ask are: (a) which disabling conditions and defeaters are the most common across defaults?; (b) which features most commonly enable the disabling conditions and defeaters that typically

¹⁷³ McKeever and Ridge (2006: 37-38): "Even if moral reasons are context-dependent, it might well be possible for us to track those reasons only if we had access to a moral principle which articulated exactly how those reasons vary from one context to the next".

change the status of default reasons?; and (c) are there enablers, disablers, and/or defeaters that typically appear together as a group across multiple contexts?

CG is a step down from TG which theoretically attempts to achieve *complete* codification. This form of CG only seeks to achieve a general codification of morality in order to guide the moral agent. Other moral tools such as moral instincts, intuitions, discernment, and an account of the virtues could help trim the hedges and fill in the missing gaps. The main problem with CG, however, is that common patterns are not that easy to come by, let alone codify. Let's take up again our example of the complication of the principle "lying is wrong".

Lying is wrong except in cases where telling the truth would cause more harm, or where the only way to tell the truth would be to do something more immoral, or where it would save a life, or save someone from physical injury or psychological pain, or extend someone's life, or prevent harmful conflict between people, or if it was done under coercion.

As we noted earlier, there may be many more exceptions. Recall that for the transcendental generalist, incorporating all possible exceptions into the formulation of the principle is not enough to achieve complete codification. TG needs to fully articulate the various ways in which the default reason in question may vary from context to context. That is, for each exception to the default, TG must provide the necessary specifications to ensure that the exception in question is in fact relevant to the case at hand. First, TG must specify how the disabler in question was enabled; and second, why the enabler is not itself disabled by some other feature. TG would have to do this for each possible exception outlined in the above principle. The problem here is that the principle would become unmanageably long. CG, on the other hand, might be able to avoid this.

CG would ask the following question about the lying principle outlined above: can we discern commonalities in the exceptional features such that we may render the principle more

manageable? The commonalities seem to be the following: “doing so in this case is less immoral” and “refraining from telling or telling the truth in this case is more immoral”. As such, the overly complicated principle above could be reformulated into two more manageable principles: (1) “lying is wrong except in cases where doing so is less immoral than the other options”; and (2) “lying is wrong except in cases where refraining from telling or telling the truth is more immoral”. In terms of practical action-guidance, this is definitely an improvement. However, these two principles are far from exceptionless; they are clearly defeasible. Consequently, although these principles might capture the common ways in which the default wrong-making status of lying is disabled or defeated, CG still needs to develop a second set of higher-order exceptionless principles that can explain in finite and manageable terms the common ways in which these disablers and defeaters are themselves enabled. These principles would be formulated along the following lines: “A typically disables B when C is present” and “G typically defeats B when H is present”. Suppose there are many more cases where B can be disabled or defeated by another feature. Suppose further that there is a finite and manageable group of features that can do this—let us suppose a hundred more aside from A and G. If such a group is in fact finite and manageable, then when C or H is present in a case where B is morally relevant, B’s default status will *generally* be disabled or defeated. The presence of C or H in cases where B is morally relevant is therefore a good general indicator that the status of B will change. This is because they typically enable other features to do so—features we are able to group together in finite and manageable terms to the extent that they share something in common: namely that they are enabled by C or H.

The first problem with CG is that it is not evidently compatible with holism—a compatibility CG sets out to achieve in the first place. That is, the whole point of developing CG was for it to be compatible with the context-sensitivity of reasons described in reasons-holism.

Unfortunately, in order to guarantee that common patterns in the behaviour of reasons are captured, CG had to develop a second set of higher-order exceptionless principles. In other words, CG had to develop a core of invariant considerations to explain and track the variance of the main preponderance of reasons. Arguably, this is a hybrid view of reasons, not a holistic one.

The second problem with CG is the assumption that features such as disablers, defeaters, and enablers behave in a finite number of ways that is *manageable*. I highly doubt this. For the record, there are three possibilities: either features behave in a finite number of ways that is manageable, or in a finite number of ways that is unmanageable, or in an infinite number of ways that is unmanageable. CG of course is only true if the first of these is true. In order to refute CG, all we need to show is that these features behave in ways that are too complex to yield manageable exceptionless principles. Whether these features are finite or infinite in number is an interesting question, but one that is unnecessary to answer in order to refute CG.

Our above example is just a taste of what these higher-order exceptionless principles might look like. We have not considered all possible exceptional features that may disable or defeat the wrong-making default status of lying, nor have we considered other types of features such as attenuators, intensifiers, and alternators. We also have not discussed how sets or groups of features might disable or defeat each other. Still, without delving into these complexities, in our example above, we can already see that these principles are very far-removed from ordinary moral practice and deliberation. Shouldn't CG be more concrete and practical? We remain quite skeptical as to whether these principles are very useful for ordinary moral agents. Arguably, if we really tried to develop these higher-order exceptionless principles that supposedly fully articulate the common patterns of the behaviours of moral and non-moral features, they would be unmanageable from the point of the view of the moral agent. Consequently, it is highly doubtful that such principles

contribute to the emergence and development of moral thought and judgement given that they are practically never used during moral learning. If one wants to develop a theory that captures common patterns in the behaviours of moral and non-moral features while simultaneously avoiding the problem of unmanageability, then perhaps one should leave exceptionless principles behind them and focus on defeasible ones. The downside of a position built solely on defeasible principles is that it may not be a properly generalist one if it cannot achieve complete or general codification. But this is only a downside if one absolutely wants to maintain their generalism. Arguably, such a position is what Lance and Little (2008) put forward and they call it “Deep Moral Contextualism”.

4.5. Rejection of Deep Moral Contextualism

Deep Moral Contextualism (DMC) is an attempt at a middle-ground position between generalism and particularism. It is made up entirely of defeasible generalizations as it is firmly against exceptionless ones due to their unmanageability.¹⁷⁴ For the sake of clarity, let us revisit the table from chapter 1.

<u>Meta-ethical Theory</u>	<u>What view of reasons does the theory assume or claim to be compatible with?</u>	<u>Do moral thought and judgement depend on principles?</u>	<u>What types of principles do moral thought and judgement depend on?</u>	<u>Do moral thought and judgement depend on a determinate or indeterminate set of principles?</u>	<u>Do moral thought and judgement depend on factors other than principles?</u>	<u>What type of codification is advocated by the theory?</u>	<u>How much codification does the theory claim to achieve?</u>
Deep Moral Contextualism	Holism	Yes	Defeasible	?	Yes	Context-Dependent or N/A?	Particular, or codification is not possible?

¹⁷⁴ See Lance and Little (2008: 73).

For DMC, do moral thought and judgement depend on a determinate or indeterminate set of defeasible principles? As we argued in chapter 3, cross-culturally, defeasible principles contribute to the emergence and development of moral thought and judgement in different sets of principles. We argued that the claim that moral thought and judgement *must* depend on a *specific* set of defeasible principles is implausible for two reasons: (1) because of the wide diversity of moral cultures in the world; and (2) because it assumes there is an *innate* and *universal* moral structure that remains intact no matter the various features and activities of any particular socio-cultural environment. To support these claims we cited research which suggests that human beings do not all reason (cognitively speaking) about right and wrong in the same way across cultures;¹⁷⁵ and that the existence of an innate and universal moral module in the brain is highly unlikely.¹⁷⁶ Although Lance and Little do not specify whether they support the notion of a determinate set of principles, I think we can safely assume that they do not given their contextualism and their strong opposition to exceptionless, necessary, universal principles.

Another question left unanswered by Lance and Little: does DMC vouch for particular or no codification? That is, can defeasible principles codify a part of the moral landscape in finite and manageable terms? DMC is clearly against transcendental codification of any kind (whether complete, general, or particular) given that it rejects explanatory, exceptionless principles. DMC is also clearly against complete context-dependent codification given that it rejects exceptionless principles that would seek to track the variance of individual reasons. DMC should also be against *general* context-dependent codification because it argues that acknowledgement of supervenience functions expressed by exceptionless generalizations that explain the moral as supervening on the natural is not equivalent to the acknowledgement of explanatory and action-guiding moral

¹⁷⁵ Henrich *et al.* (2010).

¹⁷⁶ Moll (2001, 2002), Greene and Haidt (2002), Railton (2017).

principles. Following John McDowell, Lance and Little argue that such exceptionless generalizations have simply been relocated to the supervenient level.¹⁷⁷ Contextual Generalism, in its attempt at *general* context-dependent codification, does just this. It develops a second set of higher-order principles that attempt to capture the ways in which the moral supervenes on the natural. The problem that emerges from it is just the one alluded to here: the exceptionless generalizations in question do not yield explanatory and action-guiding moral principles that are *manageable*. But the problem with CG goes deeper. Following McDowell, Lance and Little write: “supervenience can be admitted so readily because doing so admits to so little: it doesn’t mean that there are any useful patterns to the way in which the dependencies line up”.¹⁷⁸ Grouping together the situations in which, say, an action is cruel, may “yield groupings that would simply look gerrymandered to anyone who does not have independent competency with the moral concepts”.¹⁷⁹ In other words, the grouping of patterns in which an act has a given moral valence, even if this can be done at the higher-order level of supervenience, will not be usefully action-guiding unless the moral agent has the appropriate competency with the relevant moral concepts at play. For these reasons, we can safely conclude that Lance and Little reject general context-dependent codification.

The question therefore is: do they support *particular* context-dependent codification, or no codification at all? It is unclear which of these Lance and Little support. Luckily, the answer to this question might not matter all that much for our purposes. As we argued in chapter 1, the criterion of whether one is a generalist or a particularist is one’s view on codification. Even Little

¹⁷⁷ Lance and Little, “From Particularism to Defeasibility in Ethics”, 56.

¹⁷⁸ John McDowell, “Virtue and Reason.” *Monist: An International Quarterly Journal of General Philosophical Inquiry* 62 (1979): 331–50 quoted in Lance and Little (2008: 56).

¹⁷⁹ Lance and Little (2008: 56).

herself has admitted this.¹⁸⁰ Since Lance and Little support either particular or no codification, DMC is necessarily a form of particularism. When they say they are not “anti-generalist”, this is because they recognize that moral thought and judgement depend on principles. But as we argued in chapter 1, answering the dependence question affirmatively is not the deciding factor of whether one is a generalist or a particularist, even though it does place one farther away from the stronger forms of particularism. If recognizing a dependence on principles were the deciding factor, the debate would be over-simplified when it is actually much more complex. If this were so, one could simply say: “generalists support principles, particularists are against them”. But here one would be leaving out the valid position of Principled Particularism which claims that moral thought and judgement do depend on principles, but that only a particular part of the moral landscape can be codified by means of them. Hence, Lance and Little are falsely assuming that both the dependence question and the codification question are of equal import. If this were true, then there would certainly be grounds for a valid middle-position such as contextualism. But insofar as it is not, it cannot be used as a way to justify the validity of this middle-position.

Contextualism, therefore, might not be a valid and distinct position in the generalism-particularism debate. Rather, it seems to be a defining characteristic of particularism insofar as it differs very little from reasons-holism. In a certain sense, holism *is* contextualism. It is in this respect that we begin to question whether DMC really is a distinct view at all. For instance, all forms of particularism presuppose holism, therefore they adopt most, if not all, of the core contextualist tenets. What is more, the defeasible principles put forward in DMC can fit perfectly well under our version of PP. Our view, however, differs from DMC insofar as it admits of some

¹⁸⁰ Again, we refer to Little (2000: 304): “While particularism dethrones moral generalizations, though, it hardly exiles them. We will think it does only when the lesson of holism is misidentified; the enemy its objections target is not generalization, but codification”.

exceptionless principles such as the Principle of Special Reason to do Harm and *That's it*. Another way in which our version of PP differs from DMC is our view of reasons. Below I have outlined a few possible views one might have about reasons, some of which we have already discussed:

Atomism: All reasons are necessarily invariant.

Hybridism: A core of invariant reasons explains the variant reasons on the periphery.

Core Holism: A core of variant reasons explains invariant reasons on the periphery.

Disjointed Holism: Invariant reasons are possible but not explained by a core of variant reasons.

Absolute Holism: All reasons are necessarily variant.¹⁸¹

Principle absolutists are atomists; principle eliminativists are absolute holists; transcendental generalists can be atomists or hybridists, but not holists (as we argued in our rejection of TG in this chapter); contextual generalists are arguably hybridists although they claim to be holists; and anti-transcendental particularists are disjointed holists. Disjointed Holism is opposed to core-periphery views of any kind. That is, according to Disjointed Holism, the invariance that may exist does not explain the variance that exists, and the variance that exists does not explain the invariance that may exist. For the disjointed holist, most reasons are of course variant, but there may be some invariant ones that are “cosmic accidents”. Core Holism, on the other hand, is the view that there is a core of variant reasons that explains the existence, status, and behaviour of invariant reasons on the periphery.

DMC arguably supports Disjointed Holism because it rejects Core Holism. We can infer this for two reasons. First, Lance and Little’s view is that we do not need exceptionless generalizations to do the explanatory work we need them to do.¹⁸² But, as we have argued already,

¹⁸¹ For the sake of slightly better symmetry, we could also call Atomism “Absolute Atomism” and Hybridism “Core Atomism”. However, we will stick to the terminology above given that we have already discussed atomism and hybridism and referred to them by those names.

¹⁸² Lance and Little (2008: 54).

we need principles like *That's it* and Special Reason to do Harm (alongside an indeterminate number of defeasible principles) to make moral thought and judgement possible. Second, although Lance and Little develop a semantics and epistemology of defeasible generalizations “in an attempt to show that it is possible for them to play explanatory roles without being reducible to, replaceable by or ultimately beholden to exceptionless generalizations”,¹⁸³ they have not considered the possibility that exceptionless generalizations can be dependent on or beholden to defeasible ones. In other words, they have not considered the possibility of Core Holism. Recall that our view is that there are some invariant reasons and therefore some exceptionless principles; and that moral thought and judgement depend on at least two of them. In what follows, we will try to explain why the knowledge of some exceptionless principles is dependent upon the knowledge of defeasible ones. In other words, we will seek to show that there is a core of default reasons specified by defeasible principles that can explain the existence, status, and behaviour of *some* invariant reasons specified by exceptionless principles. This will be in an effort to show that Lance and Little's picture is incomplete, and that Core Holism is both true and compatible with Disjointed Holism.

Our view of PP presupposes both Core Holism and Disjointed Holism. We believe there are default variant reasons that explain the existence, status, and behaviour of *some* invariant reasons, but not all. For there may be some invariant reasons that are cosmic accidents and/or that are explained in other ways. First, recall the Principle of Special Reason to do Harm. This principle is arguably one of a few other principles subsumed under a more fundamental principle: *the Principle of Special Reason to do something that is default immoral*. This principle tells us that we need a special reason to lie, to kill, to cheat, to steal, to break a promise, or to do anything that is default wrong. So, on the one hand, we have default or defeasible principles such as “all other

¹⁸³ Ibid.

things equal, lying is wrong” and “all other things equal, killing is wrong” that specify variant reasons; and on the other hand, we have the exceptionless principle “one needs a special reason to do something that is default immoral” that specifies an invariant reason. Core Holism argues that the core of default variant reasons explains the existence, status, and behaviour of invariant ones. This appears to be the case here. The fact that we know that we need a special reason to do something that is default, defeasibly, or generally wrong, is because we have already internalized a sufficient range of wrong-making features about lying, killing, stealing, and all those other acts that are deemed wrong by default. Thus, our knowledge of the exceptionless principle that tells us that we need a special reason to do something that is default immoral is dependent on our knowledge of some sufficient range or core of default wrong-making features. Therefore, Core Holism is true to some degree. However, second, this is only the case for *some* invariant reasons. Since holism is by definition a possibility thesis, it would be inconsistent to say that *all* possible invariant reasons are *necessarily* explained by a core of variant ones. Our view, then, is that while *some* invariant reasons can be explained by a core of variant ones, there may very well be other invariant reasons that are cosmic accidents and/or that are explained in other ways. Therefore, our Core Holism is entirely compatible with Disjointed Holism. For these reasons, DMC—which assumes Disjointed Holism but rejects Core Holism—is in the wrong.

4.6. Rejection of Anti-Transcendental Particularism

Anti-Transcendental Particularism (ATP) makes two claims: (1) that moral thought and judgement do not depend on the provision of a suitable supply of moral principles;¹⁸⁴ and (2) that morality cannot be codified completely or generally in finite and manageable terms by means of

¹⁸⁴ Dancy (2004: 7).

such principles. In chapter 6 of *Ethics Without Principles*, Dancy gives us a set of criteria for a principled ethic: coverage, reasons, epistemology, and applicability.¹⁸⁵ Principles must be able to determine the moral status of every action, explain why each action has the status it has, be learnable, and be capable of guiding the moral agent to action in ordinary, exceptional, and novel cases.¹⁸⁶ What Dancy has in mind then when he thinks of a principled ethic is one that is composed of a *determinate* set of exceptionless principles. As we saw in chapter 3, according to Dancy, defeasible principles will fail on three of the aforementioned criteria, and therefore will not be able to come together to form a principled morality.¹⁸⁷

But as we saw in chapter 2, moral thought and judgement do seem to depend on at least two exceptionless principles. The Principle of Special Reason to do Harm is necessary to avoid the justification of the immoral, and *That's it!* is implicitly assumed in all moral arguments because its clause is always either affirmed (i.e. when we know that there are no further relevant features that could supersede the ones which appear in our argument) or denied (i.e. in cases where we discover features we had overlooked that do supersede the ones which appear in our argument). This alone is enough to reject ATP, but we will provide additional reasons.

Something Dancy has not considered is the possibility that moral thought and judgement depend on an *indeterminate* set of principles. Even if his set of criteria is correct, and even if an indeterminate set of principles does not amount to a true and complete morality, this does nothing to show that moral thought and judgement do not depend on an indeterminate set of principles. In chapter 3, we argued that, given the wide diversity of moral cultures, teachers, and exemplars, individuals will likely be exposed to different sets of moral principles during the course of their

¹⁸⁵ Ibid., 116-117.

¹⁸⁶ Ibid.

¹⁸⁷ See pages 111-117 of Dancy (2004).

moral learning. We then argued that, given the fact that the vast majority of people can think and judge in moral terms, it is very likely that different sets of moral principles are capable of contributing to the emergence and development of moral thought and judgement cross-culturally. We also provided evidence from the moral psychology, cognitive science, and neuroscience literatures to support these claims. For these reasons, moral thought and judgement do seem to depend, at least in part, on moral principles. Therefore, the first claim of ATP is false.

Another problem with Dancy's view is that his set of criteria for a principled ethic is too narrow. What he has in mind is complete codification by means of a specific set of exceptionless principles; but as we have seen throughout this thesis, there are many other ways in which a principled ethic may be constituted. Complete or even general codification can hardly be said to be prerequisites for a principled ethic since Principled Particularism—which vouches for a particularly limited codification—is widely recognized as a principled ethic in the scholarship. Instead, the deciding factor in determining whether a theory should count as a principled ethic should be how it answers to the dependence question, not to the codification question. That is, if one holds that moral thought and judgment depend on moral principles, one should be viewed as advancing a principled ethic. Whereas if one holds that they do not depend on moral principles of any kind, one should be viewed as negating the possibility of a principled ethic and therefore as advancing an unprincipled ethic. So while one's answer to the codification question determines whether one is a particularist or a generalist, one's answer to the dependence question determines whether one is advancing a principled or an unprincipled ethic. That principles are "dethroned" in one's theory does not necessarily entail that one is advancing an unprincipled ethic, but perhaps simply that principles are not as central and powerful as they were traditionally thought to be. Dancy's set of criteria is too narrow because it only accepts principled ethics that place principles

on the throne without considering other possibilities in which, for example, principles are necessary for moral thought and judgement despite being dethroned and insufficient for codification.¹⁸⁸

Now, because Dancy's set of criteria for a principled ethic is too narrow, he prevents himself from properly considering the credibility of defeasible moral principles. Following Lance and Little (2008), explanation of why a consideration is a reason to believe or do something must be somehow linked to generality.¹⁸⁹ That is, when we adduce reasons to explain why a belief or an action is right or wrong, we are necessarily committed to generalization. Our generalizations may not be exceptionless the vast majority of the time, and a suitable provision of them may never be capable of codifying all or most morality, but we are still committed to them in moral reasoning just as we are in other forms of reasoning (e.g. logical, epistemic, scientific etc.). Hence, as we argued in chapter 3, just as the assumption that moral reasons behave differently from other kinds of reasons is ill-founded (as Dancy himself argues),¹⁹⁰ so too the assumption that moral reasoning is the only kind of reasoning in which explanatory generalizations have no part is misplaced. The correct lesson to draw from reasons-holism is twofold: (1) the articulation of exceptionless generalizations is perhaps mostly ill-advised albeit sometimes valid; and (2) defeasible generalizations are perfectly capable of providing functional explanations to guide the moral agent

¹⁸⁸ Dancy's criteria are, however, useful for determining whether or not a principled ethic achieves complete codification. This we do not deny.

¹⁸⁹ Lance and Little (2008: 60-61).

¹⁹⁰ See Dancy (2017: "What the Particularist Believes"): "The particularist suggests that there is no reason to suppose that moral reasons function in a radically different way from other reasons. Indeed, there is a sort of presumption that they don't. That presumption is partly grounded on the fact that nobody is able to say with any confidence just which reasons are moral ones and which are not. This means that providing a radical difference between the way in which reasons of the two sorts function should seem rather peculiar. But the presumption is also partly grounded in the fact that the difference suggested by the generalist is very radical, since it affects what one might call the very logic of moral thought. To suppose that moral thought has a different logic from other thought is to adopt a bifurcated conception of rationality".

through the landscape of variant reasons. The answer then is not to eschew generalizations, but to “change our picture of what they must look like to do the work they need to do”.¹⁹¹

Finally, just because invariant reasons and exceptionless principles are few and far between does not mean that we cannot codify some particular parts of the moral landscape in finite and manageable terms. For example, the Principle of Special Reason to do Harm—which presupposes an even deeper principle that we need a special reason to do anything that is default immoral—appears to help us codify cases in which doing harm is permissible or required. Now, this does not entail that we can *fully* articulate *all* cases in which this is so, but it does demonstrate that there are some such cases out there that can be sufficiently articulated, explained, and justified to conscientious moral agents. This principle also works in unison with the sum total of defeasible principles that specify default reasons, insofar as it covers the cases in which the statuses of default reasons are altered. Although we cannot fully articulate the ways in which variant reasons may vary, or explain the status of every possible action, or guide the moral agent through all possible exceptional and novel cases, we can still confirm/know that there are some features that are right-making or wrong-making by default. This is because (1) we *know* we need special reasons to believe or act in the opposite way of what the defaults prescribe; and (2) we *know* that the defaults maintain their status across cases *unless* they are altered by some other features. So here too, it seems that principles are capable of codifying at least some particular parts of the moral landscape. For these reasons, the second claim of ATP is false.

¹⁹¹ Lance and Little (2008: 61).

4.7. Rejection of Principle Eliminativism

There are at least three main branches of Principle Eliminativism. First, there are Principle Eliminativists who accept the existence of some true principles, but who seek to eliminate them from their theory to the extent that they distort the moral agent's perception of most moral truths.¹⁹² Second, there are Principle Eliminativists about exceptionless principles, but not about other types of principles.¹⁹³ Third, there are Principle Eliminativists who believe that there are no true principles of any kind.¹⁹⁴

Proponents of the first version are also defenders of Principle Abstinence Particularism (PAP). The problem with this view is that it relies on a fact that does not seem to be generally true, namely that principles always or mostly distort moral judgement. While we are ready to accept that some or even all principles *some* of the time might distort our moral judgement, we are skeptical that distortion is *inevitable* all or most of the time. We also remain unconvinced that distortion is necessarily always counterproductive to moral thought, judgement, and learning. The distortion of judgement by one principle in a given situation might positively lead the moral agent to perceive this distortion and to think of another non-distorting principle to guide him appropriately. Thus, the experience of distortion could enhance moral learning to the extent that it might force the moral agent to re-think his application of certain principles to certain cases. Finally, the distortion that principles may cause does not necessarily seem to depend ultimately on their specific formulation, but rather more so on their practical application by the moral agent. In this regard, although we certainly agree that an exceptionless principle with an unmanageable amount of clauses in its antecedent could easily distort the moral agent's view on what is morally salient

¹⁹² See McNaughton (2000).

¹⁹³ See Little (2000) and Lance and Little (2004, 2006a, 2006b, 2008).

¹⁹⁴ See Cullity (2002) and Hicks (2013).

in a given situation,¹⁹⁵ defeasible principles with manageable clauses are entirely different. A defeasible principle such as “All other things equal, lying is wrong” can hardly be said to distort unless the moral agent misuses it. Thus, the distortion that would occur would be a result of the moral agent’s failure to apply the principle correctly, not a result of the principle’s faulty or unmanageable formulation. For these reasons, this first version of PE is unlikely to be true.

Moving on to the other two versions of PE. Principle Eliminativists about exceptionless or all principles are in the wrong if they support their view with Absolute Holism—a view of reasons that is built on the logical fallacy of proving non-existence. Although one can be skeptical of most supposedly invariant reasons given holism, one cannot logically rule them out. We cannot prove the non-existence of invariant reasons by logic alone just as we cannot prove the non-existence of God or of Russell’s flying teapot, even though we can provide strong reasons to believe in the unlikelihood of there being very many of them out there. The burden of proof is on those who claim that invariant reasons exist, not on their opponents to prove that they do not exist. Yet Principle Eliminativists who support their view with Absolute Holism argue that we can prove the non-existence of invariantly true reasons. This is not only to commit the logical fallacy of proving non-existence, but it is also inconsistent with a proper view of holism. We have to remember that holism is a possibility thesis; that is, it is not positing what *must* happen, but only what *may* happen. Thus, although it posits a world where variability is the norm, it does not in any way rule out the possibility of some invariability. Simply put, holism properly understood is not claiming that there *must* be no invariance.

¹⁹⁵ This is essentially what we argued in this chapter in our discussion of the two strategies for complete codification. Exceptionless principles with unmanageable clauses cannot serve as graspable theoretical standards nor as practical guides to action.

This critique may be enough to refute some Principle Eliminativists, but not all. For there may be Principle Eliminativists who do not support their view with Absolute Holism. Other Principle Eliminativists might argue that the very nature of a moral generalization is not built to capture moral truths. The same proponents might also argue that moral truths are just too complex to capture with generalizations. These defenders of PE would have to support some version of Contractualism (C) or Virtue Ethics (VE) combined with some concept of moral discernment or perception. Without objecting to the credibility of either of these views (C and VE) or to the validity of moral discernment as a practical and theoretical skill, we can reject this version of PE by simply pointing to functional moral generalizations that can do the work required of them to contribute to the possibility of moral thought, judgement, and learning.

First, the Principle of Special Reason to do Harm and Holton's *That's it* principle, as we have argued, are necessary for moral thought and judgement. As we saw in chapter 2, the former is needed to prevent the justification of the immoral, while the latter cannot *not* be assumed in moral deliberation for it is constantly being either affirmed (i.e. when we know that there are no other relevant features that could supersede the ones which appear in our moral argument) or denied (i.e. when we discover features we had overlooked that do supersede the ones which appear in our argument). Second, following Lance and Little (2008), any explanation in moral reasoning of why a consideration is a reason to believe a certain thing or act in a certain way must somehow be linked to generality. When we adduce moral reasons for belief or action—that is, when we think and judge in moral terms—we cannot help but employ generalizations. This claim is predicated on the assumption that moral reasoning is no different in its appeal to generalizations as other types of reasoning. Therefore, since moral principles are expressed by generalizations, moral principles are necessary for moral thought and judgement, *contra* PE.

Another reason to reject this version of PE is because it is founded on a very narrow view of truth. It seems to only accept necessary, invariant, absolute, and universal truths. It then argues that since principles cannot capture such truths, that there are no true principles. This is a narrow view of truth because it leaves out important *contingent* truths that are characteristic of the holistic landscape. Hence, this version of PE is founded on a view of truth that is incompatible with holism. But since holism is true (as we argued in chapter 2), this form of PE must be false.

4.8. Conclusion: Principled Particularism of Contingency

We are therefore left with the view that claims (a) that moral thought and judgement depend on both a determinate set of exceptionless moral principles (such as the Principle of Special Reason to do harm and Holton's *That's it* principle) *and* on an indeterminate set of defeasible moral principles; and (b) that the sum total of moral principles can only codify *particular parts* of morality in finite and manageable terms, not most or all parts. Our view gains particular attraction insofar as it does not claim that moral principles are the be-all and end-all of morality, while simultaneously recognizing their necessity and utility for moral thought, judgement, action-guidance, and learning. Just because moral principles do not codify all or most of morality does not mean they do not play an essential role in making moral thought and judgement possible. Nor does it show that they are incapable of enhancing moral learning, or of guiding individuals towards performing right actions, exhibiting socially desirable virtues, or pursuing righteous ends beneficial to themselves and to the greater good of others.

Our view is particularist and not generalist because it does not advance the claim that morality can be completely or generally codified. Our view is that only some specific parts of morality can be codified by principles, and that the rest of the picture must be obtained by other

means such as an account of the virtues, and/or some account of moral perception or discernment. But in these latter accounts (which have not been provided here¹⁹⁶), codification will likely mean something entirely different; and it might not even make good sense to view such accounts as providing anything akin to codification at all. In addition, although our view is particularist, it is not a form of Anti-Transcendental Particularism because we recognize that moral thought and judgement depend at least in part on moral principles. For these reasons, it makes more sense to call our view a version of Principled Particularism. Finally, our view is not a form of contextualism because we recognize the possibility of exceptionless principles that can do important explanatory work, codify certain parts of morality, and be procedurally useful for the moral agent. Having said this, these exceptionless principles, such as the Principle of Special Reason to do Harm and Holton's *That's it* principle, although necessary for the possibility of moral thought and judgement, are by no means sufficient. Indeed, as we argued earlier, the possibility of moral thought and judgement also depends on an indeterminate set of defeasible principles, innate moral instincts, intuitions, and domain-general learning capacities.

Our central claim in chapter 3 was that moral thought and judgement depend on a determinate set of exceptionless principles and an indeterminate set of defeasible principles. That is, for each individual, moral thought and judgement will be made possible by a different set of principles. However, we specified that certain exceptionless principles such as the Principle of Special Reason to do Harm and *That's it* must be present in all possible sets, otherwise moral thought and judgement would not be possible at all. In each different set, there will be

¹⁹⁶ Rebecca Stangl (2010) and Christine Swanton (2015) have both offered an account of the virtues compatible with moral particularism. The former calls her view "Asymmetrical Virtue Particularism" and the latter "Particularist Virtue Ethics". These accounts are worth looking into as they provide convincing reasons to believe that particularism can and should marry itself with some account of the virtues. Swanton's theory is especially interesting insofar as her "target-centered" view emphasizes the necessity for a proper account of moral perception and discernment.

commonalities (the exceptionless principles and perhaps some default clauses), but also differences due to the diversity of moral cultures and socio-cultural environments. The only necessity regarding defeasible principles is that individuals be conditioned by a sufficient range of them. Which defeasible principles will be present will depend on the individual's socio-cultural environment which includes her moral teachers, exemplars, and cultural values among other things. Our view in chapter 3 is particularly attractive to the extent that it recognizes that there are some exceptionless principles that are necessarily true cross-culturally, while also recognizing that contingently true defeasible principles relative to one's own culture also play an essentially role in the emergence and development of moral thought and judgement. In other words, our view adequately recognizes the plain and obvious truth of the relativity of cultural values, while simultaneously avoiding the problem of "bad relativism" according to which objective truth and fundamental morality are impossible.

Our central claim throughout this fourth chapter has been that moral principles can only codify *particular parts* of morality in finite and manageable terms, not most or all parts. On the one hand, there are at least four reasons to believe that complete and general codification are both impossible. First, in a reasons-holistic landscape, there just aren't enough invariant reasons and exceptionless principles to codify most or all of morality. Second, moral and non-moral features (whether they are finite or infinite in number) cannot all be captured by a manageable set of principles. Third, there may always be an indefinite number of novel and exceptional situations that will escape complete capture by means of principles. Fourth, there is no way to capture the common patterns in which moral and non-moral features may vary from case to case in a manageable way and without assuming a hybrid view of reasons that would contradict our holism.

On the other hand, there are at least two reasons to believe that some particular codification is possible. First, there are at least two true exceptionless principles, one of which prevents the justification of the immoral and explains why default reasons can be defeated or disabled, while the other covers all cases in which there are no further relevant moral and non-moral features which supersede the ones which appear in our moral arguments. Second, there are true defeasible principles that cover all cases in which the default reason is *not* defeated or disabled. These defeasible principles are capable of (1) determining the moral status of all actions that are not defeated or disabled by other considerations, (2) explaining why each of these actions have the status they have, (3) being learned, and (4) guiding the moral agent to action in ordinary, exceptional, and novel cases. Now, of course, these principles can, at times, lead the moral agent to error. But this should not come as a surprise since defeasible principles can always be challenged, modified, or annulled to the extent that they specify reasons that may vary in different contexts. If these principles could successfully guide the agent through all possible situations without failure, they would be exceptionless principles and they would codify all of morality—which is what we have been arguing against throughout this chapter. In both chapters 3 and 4, we noted that the experiences of error and distortion are key components to the development of moral thought, judgement, and learning.

Conclusion

The scholarly goal of this thesis has been to make a modest contribution to the particularism-generalism debate. I hope to have built upon Richard Holton's Principled Particularism by adding elements from Lance and Little's Deep Moral Contextualism, Dancy's

Anti-Transcendental Particularism, and McKeever and Ridge's Generalism as a regulative ideal. I also hope to have clarified the spectrum of the debate by asking more precise questions and drawing key distinctions. For example, we saw that Dancy's central claim—that the possibility of moral thought and judgement does not depend on the provision of a suitable supply of moral principles—does not adequately distinguish between exceptionless and defeasible principles. Accordingly, one goal of this thesis has been to show that one can be a particularist while also believing that defeasible principles play an important role in moral thought, judgement, and learning. We also bifurcated two important considerations into two distinct questions: the dependence question and the codification question. The former asks whether moral thought and judgement depend on moral principles, whereas the latter asks whether morality (all moral truths) can be codified in finite and manageable terms. In this regard, we argued that that upon which moral thought and judgement depend is not necessarily capable of codifying all moral truths. In chapter 1, we argued that one's answer to the codification question determines whether one is a particularist or a generalist; and later, in chapter 4, we argued that one's answer to the dependence question determines whether one is advocating for a principled or an unprincipled ethic. This bifurcation, as we saw, allowed for the possibility of more middle positions between Dancy's brand of particularism and McKeever and Ridge's brand of generalism. Our position, for instance, is a principled ethic that does not codify all moral truths. These distinctions are not found in the scholarship, and so we hope to have provided a useful contribution in this regard.

The general aim of this project has been to come to a better understanding of the status and role of moral principles in a reasons-holistic moral landscape. The specific aim has been to develop a viable position within the particularism-generalism debate that combined elements from both theories. By asking and answering three specific research questions, I have developed an account

of Principled Particularism. First, I argued that, in a reasons-holistic moral landscape, exceptionless moral principles are necessary but insufficient to ground the possibility of moral thought and judgement. Second, I argued that the possibility of moral thought and judgement depends on the provision of a suitable supply of moral principles. More specifically, I argued that they depend on a determinate set of exceptionless moral principles and an indeterminate set of defeasible ones. Third, I argued that moral principles are insufficient to codify the moral landscape in finite and manageable terms. In this vein, I argued that moral principles can only codify particular parts of morality, not most or all parts.

In summary, moral principles are only a piece of the pie. In terms of the emergence and development of moral thought and judgement, they have been shown to be necessary but nevertheless insufficient. After doing a survey of the moral psychology, cognitive science, and neuroscience literatures in chapter 3, we saw that moral thought and judgement also depend on innate moral instincts, intuitions, and domain-general learning capacities. In terms of the codification of morality, we argued in chapter 4 that a finite set of principles can only provide us with a partial and blurry code despite the fact that principles are necessary for moral thought and judgement. Indeed, the reasons-holistic landscape is complex, ever-changing, messy, and blurry. Arguably, theories that claim to fully articulate all moral truths—i.e. that attempt to render the picture of this complicated and blurry moral world sharp—should be met with great skepticism.¹⁹⁷

¹⁹⁷ “... imagine having to draw a sharp picture ‘corresponding’ to a blurred one. In the latter there is a blurred red rectangle; you replace it with a sharp one. Of course – several such sharply delineated rectangles could be drawn to correspond to the blurred one. – But if the colours in the original shade into one another without a hint of any boundary, won’t it become a hopeless task to draw a sharp picture corresponding to the blurred one? Won’t you then have to say: “Here I might just as well draw a circle as a rectangle or a heart, for all the colours merge. Anything – and nothing – is right.” – And this is the position in which, for example, someone finds himself in ethics or aesthetics when he looks for definitions that correspond to our concepts”. Ludwig Wittgenstein. *Philosophical Investigations: The German Text with an English Translation* by G. E. M. Anscombe, P. M. S. Hacker and Joachim Schulte. 4th edition. Wiley-Blackwell, 2009: page 40, section 77.

Bibliography

- Albertzart, Maïke. *Moral Principles*. London, UK and New York, USA: Bloomsbury, 2014.
- . “Principle-Based Moral Judgement.” *Ethical Theory and Moral Practice: An International Forum* 16, no. 2 (2013): 339–54.
- Al-Shehab, A. J. “A cross-sectional examination of levels of moral reasoning in a sample of Kuwait University faculty members.” *Social Behavior and Personality* 30 (2002): 813–820.
- Aristotle. “Nicomachean Ethics: Translated by W.D. Ross.” In *The Basic Works of Aristotle: Edited by Richard McKeon*. New York: Modern Library, 2001.
- Audi, Robert. “Ethical Generality and Moral Judgment.” In *Challenging Moral Particularism*, edited by Mark Norris Lance, Matjaž Potrč, and Vojko Strahovnik, 31–52. New York: Routledge, 2008.
- Baron, Jonathan. *Judgment misguided: Intuition and error in public decision making*. New York: Oxford University Press, 1998.
- Bloom, Paul. *Just Babies: The Origins of Good and Evil*. Crown Publishing Group: New York, 2013.
- Burns, James H., and Herbert L. A. Hart. *The Collected Works of Jeremy Bentham: A Comment on the Commentaries and A Fragment on Government*. Oxford University Press, 1977.
<https://www.oxfordscholarlyeditions.com/view/10.1093/actrade/9780199553471.book.1/actrade-9780199553471-book-1>.
- Crisp, Roger, 2000, “Particularizing Particularism”, In *Moral Particularism*, edited by Brad Hooker and Margaret Little, 23–47. New York: Oxford University Press, 2000.
- Cullity, Garrett. “Particularism and Moral Theory: Particularism and Presumptive Reasons: Garrett Cullity.” *Aristotelian Society Supplementary Volume* 76, no. 1 (2002): 169–90.
<https://doi.org/10.1111/1467-8349.00094>.
- Dancy, Jonathan. “Are Basic Moral Facts Both Contingent and a Priori?” In *Challenging Moral Particularism*, edited by Mark Norris Lance, Matjaž Potrč, and Vojko Strahovnik, 116–22. New York: Routledge, 2008.
- . “Defending Particularism.” *Metaphilosophy* 30, no. 1–2 (1999): 25–32.
- . *Ethics Without Principles*. New York: Oxford University Press, 2004.
- . “Moral Particularism.” In *The Stanford Encyclopedia of Philosophy*, Winter Edition 2017. <https://plato.stanford.edu/archives/win2017/entries/moral-particularism>.

- . *Necessity, Universality, and the A Priori in Ethics*. Oxford University Press, 2007.
- . “Review of: Principled Ethics: Generalism as a Regulative Ideal by Sean McKeever and Michael Ridge.” *Mind* 116, no. 462 (2007): 462–67.
- . “The Particularist’s Progress.” In *Moral Particularism*, edited by Brad Hooker and Margaret Little, 130–56. New York: Oxford University Press, 2000.
- . *What Do Reasons Do?* Clarendon Press, 2006.
- Eno, Robert. *The Analects of Confucius : An Online Teaching Translation*, 2015.
[http://www.indiana.edu/~p374/Analects_of_Confucius_\(Eno-2015\).pdf](http://www.indiana.edu/~p374/Analects_of_Confucius_(Eno-2015).pdf).
- EURObiz. “China and Its Elderly Care System,” October 17, 2019.
<https://www.eurobiz.com.cn/china-and-its-elderly-care-system/>.
- Fiala, Andrew. “Pacifism.” In *The Stanford Encyclopedia of Philosophy*, edited by Edward N. Zalta. Metaphysics Research Lab, Stanford University, Fall Edition 2018.
<https://plato.stanford.edu/archives/fall2018/entries/pacifism/>.
- Flanagan, Owen. *The Geography of Morals: Varieties of Moral Possibility*. New York: Oxford University Press, 2017.
- Flynn, Jennifer. “Moral Particularism.” *Analysis* 70, no. 1 (2010): 140–48.
- Foot, Philippa. “The Problem of Abortion and the Doctrine of the Double Effect Philippa Foot.” *Oxford Review*, no. 5 (1967): 1-5.
- Garfield, Jay. “Particularity and Principle: The Structure of Moral Knowledge.” In *Moral Particularism*, 178–204. New York: Oxford University Press, 2000.
- Gleeson, Andrew. “Moral Particularism Reconfigured.” *Philosophical Investigations* 30, no. 4 (2007): 363–80.
- Glinskaya, Elena, and Feng, Zhanlian, *Options for Aged Care in China*, World Bank, 2018, viewed 12th September 2019
<http://documents.worldbank.org/curated/en/171061542660777579/Options-for-Aged-Care-in-China-Building-an-Efficient-and-Sustainable-Aged-Care-System>
- Greene, Joshua. “Beyond point-and-shoot morality: Why cognitive (neuro)science matters for ethics.” *Ethics*, 124 (2014): 695–726.
- Greene, Joshua, and Jonathan Haidt. “How (and where) does moral judgement work?” *Trends in Cognitive Sciences* 6, no. 12 (2002): 517-523.
- Greene, Joshua, R. Sommerville, Leigh Nystrom, John Darley, and Jonathan Cohen. “An fMRI

- Investigation of Emotional Engagement in Moral Judgment.” *Science (Washington)* 293, no. 5537 (2001): 2105–2108. <https://doi.org/10.1126/science.1062872>.
- Haidt, Jonathan. “The Emotional Dog and Its Rational Tail.” *Psychological Review* 108, no. 4 (2001): 814–834.
- Henrich, Joseph, Steven J. Heine, and Ara Norenzayan. “The Weirdest People in the World?” *Behavioral and Brain Sciences* 33 (2010): 61–135.
- Hernández-Iglesias, Manuel. “Generalism without Foundations.” *Acta Analytica: International Periodical for Philosophy in the Analytical Tradition* 21:2, no. 39 (2006): 71–86.
- Hicks, Amelia. “Moral Particularism: The Nonexistence of Moral Principles and the Nature of Proper Moral Deliberation.” Ph.D. Thesis, University of Notre Dame, 2013.
- . “Particularism Doesn’t Flatten.” *Journal of Moral Philosophy: An International Journal of Moral, Political and Legal Philosophy* 13, no. 3 (2016): 339–62.
- Holton, Richard. “Particularism and Moral Theory: Principles and Particularisms.” *Aristotelian Society: Supplementary Volume* Supp, no. 76 (2002): 191–209.
- Hooker, Brad. “Moral Particularism and the Real World.” In *Challenging Moral Particularism*, edited by Mark Norris Lance, Matjaž Potrč, and Vojko Strahovnik, 12–30. New York: Routledge, 2008.
- . “Moral Particularism: Wrong and Bad.” In *Moral Particularism*, edited by Brad Hooker and Margaret Little, 1–22. New York: Oxford University Press, 2000.
- Kant, Immanuel. *Groundwork of the Metaphysics of Morals: Translated and Edited by Mary Gregor*. United Kingdom: Cambridge University Press, 1997.
- . “On a Supposed Right to Lie from Altruistic Motives”, in *Kant’s Critique of Practical Reason and Other Writings in Moral Philosophy*, trans. and ed. Lewis White Beck, 346–350. Chicago: University of Chicago Press, 1949.
- Kohlberg, Lawrence. “Moral stages and moralization.” In *Moral development and behavior: Theory, research and social issues*, 31–53. Rinehart and Winston, 1976.
- . *The philosophy of moral development*. Harper & Row, 1981.
- Korsgaard, Christine. *The Constitution of Agency: Essays on Practical Reason and Moral Psychology*, Oxford: Oxford University Press, 2008
- . *Self-Constitution: Agency, Identity and Integrity*, Oxford: Oxford University Press, 2009
- Lance, Mark Norris, and Margaret Olivia Little. “Defeasibility and the Normative Grasp of

- Context” *Contextualisms in Epistemology* 61, no. 2/3 (2004): 535–55.
- . “From Particularism to Defeasibility in Ethics.” In *Challenging Moral Particularism*, edited by Mark Norris Lance, Matjaž Potrč, and Vojko Strahovnik, 53–74. New York: Routledge, 2008.
- . *Particularism and Antitheory*. New York: Oxford University Press, 2006.
- . *Defending Moral Particularism*. Blackwell Publishing, 2006.
- Little, Margaret Olivia. “Moral Generalities Revisited.” In *Moral Particularism*, edited by Brad Hooker and Margaret Little, 276–304. New York: Oxford University Press, 2000.
- . “On Knowing the ‘Why’: Particularism and Moral Theory.” *Hastings Center Report* 31, no. 4 (2001): 32–40.
- McDowell, John. *Mind, Value and Reality*. Cambridge, MA: Harvard University Press, 1998.
- . “Virtue and Reason.” *Monist: An International Quarterly Journal of General Philosophical Inquiry* 62 (1979): 331–50.
- . “Wittgenstein on Following a Rule.” *Synthese: An International Journal for Epistemology, Methodology and Philosophy of Science* 58 (1984): 325–64.
- McKeever, Sean, and Michael Ridge. “Moral Particularism and Moral Generalism.” In *The Stanford Encyclopedia of Philosophy*. Winter 2016 Edition, <https://plato.stanford.edu/archives/win2016/entries/moral-particularism-generalism>.
- . “Particularism and the Contingent a Priori.” In *Challenging Moral Particularism*, edited by Mark Norris Lance, Matjaž Potrč, and Vojko Strahovnik, 107–15. New York: Routledge, 2008.
- . *Principled Ethics: Generalism as a Regulative Ideal*. New York: Oxford University Press, 2006.
- . “Reasons: Holism and Atomism.” In *International Encyclopedia of Ethics*. Hoboken, New Jersey: John Wiley & Sons, Ltd., 2013. <https://doi.org/10.1002/9781444367072.wbiee507>.
- . “Review of: Ethics Without Principles by Jonathan Dancy.” *The Philosophical Review* 116, no. 1 (2007): 124–28.
- . “The Many Moral Particularisms.” *Canadian Journal of Philosophy* 35, no. 1 (2005): 83–106.
- . “What Does Holism Have to Do with Moral Particularism?” *Ratio: An International*

- Journal of Analytic Philosophy* 18, no. 1 (2005): 93–103.
- McNaughton, David. *Moral Vision*. Oxford: Blackwell, 1988.
- McNaughton, David, and Piers Rawling. “Holism about Value.” In *Challenging Moral Particularism*, edited by Mark Norris Lance, Matjaž Potrč, and Vojko Strahovnik, 166–84. New York: Routledge, 2008.
- . “Unprincipled Ethics.” In *Moral Particularism*, edited by Brad Hooker and Margaret Little, 256–75. New York: Oxford University Press, 2000.
- McShane, Katie. “Holism.” In *International Encyclopedia of Ethics*. Hoboken, New Jersey: John Wiley & Sons, Ltd., 2013.
- Mikhail, John. *Elements of moral cognition: Rawls’s linguistic analogy and the cognitive science of moral and legal judgment*, Cambridge: Cambridge University Press, 2011.
- Mill, John Stuart. *Utilitarianism*. Mineola, New York: Dover Publications, Inc., 2007.
- Moll, Jorge, *et al.* “Frontopolar and anterior temporal cortex activation in a moral judgment task: preliminary functional MRI results in normal subjects.” *Arq. Neuropsiquiatr.* 59 (2001): 657–664.
- . “Functional networks in emotional moral and nonmoral social judgments.” *Neuroimage* 16 (2002): 696–703.
- Moore, George E. *Principia Ethica*. (Revised Edition with “Preface to the second edition” and other papers, T. Baldwin (ed.), Cambridge: Cambridge University Press, 1993). Cambridge: Cambridge University Press, 1903.
- Nichols, Shaun. “Innateness and moral psychology.” In *The innate mind: Structure and contents*, 353–370, Oxford: Oxford University Press, 2005.
- Potrč, Matjaz. “Particularism and Resultance.” *Acta Analytica: International Periodical for Philosophy in the Analytical Tradition* 19, no. 33 (2004): 163–87.
- Prinz, Jesse J. “Is Morality Innate?” In *Moral Psychology: The Evolution of Morality: Adaptations and Innateness*, 367–406, Cambridge, MA: MIT Press, 2007.
- Railton, Peter. “Moral Learning: Conceptual Foundations and Normative Relevance.” *Cognition* 167 (2017): 172–90.
- Robinson, Luke. “Moral Holism, Moral Generalism, and Moral Dispositionalism.” *Mind: A Quarterly Review of Philosophy* 115, no. 458 (2006): 331–60.
- Ross, William David. *The Right and the Good*. Edited by Philip Stratton-Lake. New York: Clarendon Press Oxford, 2002.

- Sharp, Daniel. "The Rule-Following Considerations and the Case for Moral Particularism: A Critical Examination." *Ethical Perspectives: Journal of the European Ethics Network* 23, no. 3 (2016): 411–43.
- Samuel Scheffler. *The Rejection of Consequentialism: A Philosophical Investigation of the Considerations Underlying Rival Moral Conceptions*. New York: Oxford University Press, 1994.
- Sinnott-Armstrong, Walter. "Some Varieties of Particularism." *Metaphilosophy* 30, no. 1–2 (1999): 1–12.
- Snarey, John R. "The cross-cultural universality of social-moral development: A critical review of Kohlbergian research." *Psychological Bulletin* 97, no. 2 (1985): 202–232.
- Sneddon, Andrew. "Thick Concepts and Holism about Reasons." *Journal of Value Inquiry* 44, no. 4 (2010): 461–68.
- Stangl, Rebecca. "Asymmetrical Virtue Particularism." *Ethics: An International Journal of Social, Political, and Legal Philosophy* 121, no. 1 (2010): 37–57.
- Strahovnik, Vojko. "Introduction: Challenging Moral Particularism." In *Challenging Moral Particularism*, edited by Mark Norris Lance, Matjaž Potrč, and Vojko Strahovnik, 1–11. New York: Routledge, 2008.
- Swanton, Christine. "A Particularist but Codifiable Virtue Ethics." In *Oxford Studies in Normative Ethics*, 5:38–63. New York: Oxford University Press, 2015.
- Teon, Aris. "Filial Piety (孝) in Chinese Culture." *The Greater China Journal*, 2016. <https://china-journal.org/2016/03/14/filial-piety-in-chinese-culture/>.
- Thomas, Alan. "Another Particularism: Reasons, Status and Defaults." *Ethical Theory and Moral Practice: An International Forum* 14, no. 2 (2011): 151–67.
- Thompson, Judith Jarvis. "The Trolley Problem." *The Yale Law Journal*, 94 (1985): 1395–1415.
- Väyrynen, Pekka. "A Theory of Hedged Moral Principles", in Russ Shafer-Landau (ed.) *Oxford Studies in Metaethics*, 4 (2009): 91–132.
- . "Moral Generalism: Enjoy in Moderation." *Ethics* 116, no. 4 (2006): 707–41.
- . "Particularism and Default Reasons." *Ethical Theory and Moral Practice: An International Forum* 7, no. 1 (2004): 53–79.
- . "Usable Moral Principles." In *Challenging Moral Particularism*, edited by Mark Norris Lance, Matjaž Potrč, and Vojko Strahovnik, 75–106. New York:

Routledge, 2008.

Williams, Bernard. *Ethics and the Limits of Philosophy*. Abingdon, United Kingdom: Routledge, 2006.

Wittgenstein, Ludwig. *Philosophical Investigations: The German Text with an English Translation by G. E. M. Anscombe, P. M. S. Hacker and Joachim Schulte*. 4th edition. Blackwell Publishing, 2009.

Zamzow, Jennifer L. "Rules and Principles in Moral Decision Making: An Empirical Objection to Moral Particularism." *Ethical Theory and Moral Practice: An International Forum* 18, no. 1 (2015): 123–34.