



Delft University of Technology

A unificationist approach to wrongful pure risking

Maheshwari, Kritika

DOI

[10.1080/0020174X.2024.2339512](https://doi.org/10.1080/0020174X.2024.2339512)

Publication date

2024

Document Version

Final published version

Published in

Inquiry (United Kingdom)

Citation (APA)

Maheshwari, K. (2024). A unificationist approach to wrongful pure risking. *Inquiry (United Kingdom)*.
<https://doi.org/10.1080/0020174X.2024.2339512>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

A unificationist approach to wrongful pure risking

Kritika Maheshwari

Department of Values, Technology, and Innovation, Delft University of Technology, Delft, Netherlands

ABSTRACT

What makes cases of pure risking sometimes wrong? There is a strong intuition that the wrongness of pure risking stands in an explanatory relationship with the wrongness of the non-risky act. Yet, we cannot simply take this for granted insofar as in cases of wrongful pure risking, the risked outcome fails to materialize. To this end, I motivate and develop an under explored approach in the literature that I call Unificationism. According to the Unificationist account that I defend, the fact that pure risking φ is *pro tanto* wrong is grounded by a general moral fact that φ -ing is *pro tanto* or all-things-considered wrong, other things being equal. This relationship holds even if and when an agent's risky conduct fails to transpire or culminate into φ -ing *ex post*. I argue that this Unificationist account captures our explanatory intuition, avoids problems of extensional and explanatory inadequacy that existing alternative faces, and most importantly, renders Unificationism as a plausible view within the ethics of pure risking.

ARTICLE HISTORY Received 24 May 2023; Accepted 21 March 2024

KEYWORDS Pure risking; wrongness; Unificationism

1. Introduction

Recent theorizing in the ethics of risk has paid great attention to cases of *wrongful pure risking* where the risk fails to materialize.¹ Here's one example.²

Russian Roulette: Bill finds Joe asleep and decides to put a loaded gun against Joe's head, just for fun. There is one in six chance that if he pulls the trigger, the

CONTACT Kritika Maheshwari  kritika136@gmail.com  Delft University of Technology, Faculty of Technology, Policy and Management, Jaffalaan 5, 2628 BX Delft

¹We can contrast such cases from cases of impure risking where the risk materializes (Thomson 1986, p.173). Note that for the purpose of this paper, I will understand risk as the probability of some bad outcome (e.g. harm).

²Variations of this example are discussed in, amongst others, Thomson 1986, Cripps 2013, James 2016, Oberdiek 2012, 2017, Quong 2020, Frowe 2021, Rowe 2022, and Bowen 2022.

© 2024 The Author(s). Published by Informa UK Limited, trading as Taylor & Francis Group

This is an Open Access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited. The terms on which this article has been published allow the posting of the Accepted Manuscript in a repository by the author(s) or with their consent.

gun will fire and kill Joe. No bullet fires when Bill pulls the trigger, and Joe is left unharmed.

Most would agree that Bill acts wrongly in risking killing Joe. This judgment of Bill's conduct does not seem to depend on, nor is it determined by whether the risk materializes for Joe. Tom Parr & Adam Slavny (2019) capture this thought nicely when they say, '[r]iskers commit wrongs when they cast the dice, not just when they turn up snake eyes' (p.76).³ However, things are not so clear when considering *why* rather than whether (or when) Bill's pure risking is *pro tanto* morally wrong.⁴

There is a strong intuition that if Bill acts wrongly, then the wrongness of his conduct must be explained in relation to the fact that killing Joe is wrong, other things being equal. Yet, we cannot simply take this for granted. Since Bill (luckily) fails to actually kill Joe, it can be hard to articulate precisely how and why the wrongness of killing Joe stands in an explanatory relationship with why pure risking killing him is also wrong. This paper explores how we can best address and capture this explanatory intuition.

To this end, I motivate and argue for what I call the Unificationist approach. In a nutshell, Unificationism aims to unify and explain the morality of act type pure risking φ in relation to that of act type φ -ing. According to the Unificationist account that I defend in this paper, there exists an explanatory relationship between the fact that pure risking φ is wrong and a *general* moral fact that φ -ing is *pro tanto* wrong. This explanatory relationship holds even if and when an agent's risky conduct fails to transpire or culminate into φ -ing *ex post*. Call this the Simple Account. My defense of the Simple Account builds on a recent Unificationist account as defended by Parr & Slavny (2019).⁵

According to their Buck Passing Account, an explanatory relationship exists between the wrongness of pure risking φ and *particular* wrong-making factors of φ -ing. While our accounts differ subtly with respect to the *explanans* of the wrongness of pure risking (general moral fact versus particular wrong-making factors), this difference proves substantive. In particular, it has significant implications for capturing our initial

³In discussing cases of intentionally imposed risks, Adriana Placani (2016) voices a similar thought when she writes that '[j]udgments of wrongfulness for incomplete actions are not dependent on whether or not the actions succeeded with respect to their aims. Both wrongfulness and culpability can be established prior to act-completion and in the absence of outcome materialization' (p.85).

⁴Henceforth, I will drop the term *pro tanto* for simplicity.

⁵Note that in their paper, Parr & Slavny don't defend Unificationism, nor do they explicitly qualify their account as Unificationist. However they admitted to it in personal communication.

explanatory intuition and defending a general Unificationist account that is both extensionally and explanatorily adequate, or so I'll argue.

My discussion is organized into five sections. In §2, I introduce and motivate Unificationism. I also discuss how it both contrasts with and complements dominant approach in recent literature which I call Isolationism.⁶ Unlike Unificationists, Isolationists prefer to locate and explain the wrongness of act type pure risking φ in isolation from the morality of act type φ -ing. In §3, I argue why Unificationists should explicate the explanatory relationship between them in terms of grounding. In §4, I present Parr & Slavny's Buck Passing Account and discuss why it fails as a general Unificationist account despite its intuitive appeal. In §5, I propose and defend my Simple Account as a plausible alternative and conclude in §6.

2. Two approaches: isolationism and unificationism

The existing literature is ripe with accounts that explain the relational or non-relational wrong of pure risking in distinct ways. For instance, those who are attracted to a harm-based account think that Bill acts wrongly in *Russian Roulette* when and because pure risk is harmful in and of itself. When Bill imposes pure risk on Joe, he harms and thereby wrongs Joe by frustrating his preference not to be exposed to risk (Finkelstein 2003), or because he disadvantages Joe by affecting his future functionalities (Wolff & de-Shalit 2018), or because he sets back Joe's dignitary interest (Placani 2016) or his interest in controlling his exposure to risk (James 2016).

Those who are attracted to an option- or opportunity-based account argue that Bill acts wrongly when and because he interferes with Joe's autonomy by reducing his safe and choice-worthy options (Oberdiek 2017 & 2012), or because he diminishes Joe's overall negative freedom by reducing his compossible set of options for actions (Ferretti 2016). Amongst those who lean towards rights-based thinking, some think that Bill acts wrongly by virtue of infringing Joe's right not to be exposed to risk grounded in a distinctive interest that risk affects (Oberdiek 2009, Perry 2003 & 2014).⁷ Some others have appealed to the idea

⁶I thank John Oberdiek for suggesting this name to me.

⁷On Oberdiek's view, it's an agent's interest in autonomy, whereas on Perry's view, it's an agent's second-order interest in not being subject to risk. Note that others who have also defended a right against risk deviate from Oberdiek and Perry's view and ground this right in an agent's interest in not being harmed (See McCarthy 1997, Bowen 2022). Thanks to an anonymous reviewer for bringing this to my attention.

that in risking killing Joe, Bill acts wrongly because he frustrates Joe's legitimate expectations of how he should have conducted himself (Kumar 2003) or because he exercises a kind of dominating power over him (Maheshwari & Nyholm 2022).

Evidently, these accounts are quite diverse in their explanatory scope and focus. Yet, there is (at least) one common thread connecting them. These accounts explain the wrongness of Bill's conduct in terms of facts that are over and beyond any facts about the wrongness of what Bill was risking, namely, actually killing Joe.⁸ In other words, facts about risk itself being a kind of harm or a form of interference with Joe's freedom, rights, autonomy, welfare, and so on are contingent or constitutive wrong-making facts that directly relate to or are effects of Bill's act of pure risking.

If and when they these substantive wrong-making facts obtain, their obtainment constitutes or gives rise to a *pro tanto* moral reason against Bill's conduct, other things being equal. Call this general approach of explaining the wrongness of act type pure risking φ by way of isolating and referring to substantive wrong-making facts (or features) that are distinct from facts about the wrongness of the risked act type φ -ing, *Isolationism*. By contrast, call the alternative approach of referring to facts (or features) about the wrongness of the risked act type φ -ing in explaining the wrongness of act type pure risking φ , *Unificationism*.

Although Isolationists don't explicitly deny or refute any potential connection between the morality of pure risking φ and that of φ -ing, their search for some distinct substantive wrong-making feature of Bill's risky conduct is driven by the predicament that the risked harm failed to materialize. Given this, their own accounts 'attempt to prescind as much as possible from the material harm that risks can ripen into' (Oberdiek 2012, p.343).

For Unificationists, by contrast, the fact that the risk failed to materialize does not hinder appealing to the morality of killing Joe for capturing why Bill acts wrongly. The theoretical impetus for their approach stems from the strong explanatory intuition stated earlier, namely, that if and when there is a *pro tanto* moral reason against Bill imposing a pure risk of killing Joe, this reason must have its explanatory roots or relations

⁸Maria-Paola Ferretti (2016), for instance, clarifies this move in her development of a freedom-based account for cases of pure risking harm by noting that, '[T]he question I address is not about the 'outcome harm' that can follow from risk imposition (that is, the harm inflicted when the more or less probable negative outcomes are actually realized), but what, if anything, is morally problematic with risk imposition as such' (p.262).

with why killing him is wrong, other things being equal. A handful of theorists who have signaled their attraction towards Unificationism appeal to this intuition, albeit in different ways.

For instance, in addressing the debate on the rights-based approach to cases of risking harm, David McCarthy (1997) notes that ‘[I]t would be very surprising if facts about the morality of imposing risks of harms did not connect importantly with the morality of harming’ (p.205). In arguing against Isolationist accounts, Parr & Slavny (2019) similarly hold that ‘[A]ny account of the wrongness of pure risking should be able to explain this general relationship between wrongness in non-risking cases’ (p.81). More recently, in her critique of non-consequentialist approaches to risk, Barbara Fried (2020) also observes that ‘[It] is hard to see how deontologists can avoid the conclusion that what is wrong with risky conduct, finally, is the risk of *ex post* harm it poses.’ (p.247).

However, even theorists like McCarthy, Parr & Slavny, and Fried do not explicate why the morality of risking harm connects with that of harm in cases where the harm fails to materialize. Nor do they clarify or discuss in depth what kind of connection or relationship exists, if any. Where Isolationism has received a lot of attention in the literature, as indicated by the sheer variety of accounts on offer, Unificationism has remained largely under-explored and under-motivated. What, then, are some initial reasons for thinking that an explanatory relation exists between the wrongness of pure risking and that in non-risky cases? In other words, why take Unificationism seriously?

There are at least five independent reasons why. First, there seems to be a presumption in favor of something like Unificationism in our everyday deliberation and reasoning about cases like *Russian Roulette*. From an *ex post* perspective, when asked why it was wrong for Bill to play Russian Roulette with Joe, we might respond something like, ‘Well, because he could have actually killed Joe’. Alternatively, from an *ex ante* perspective, when asked why it would be wrong for Bill to play Russian Roulette, we might respond something like, ‘Well, because he *could* actually kill Joe’. One distinctive feature of these everyday explanations is that they explicitly and directly appeal to the potential for harm and/or the reasons against killing Joe even if and when Bill never does such a thing *ex post*.

Second, part of our existing philosophical discussion around risk seems pre-committed to something like Unificationism without always making it explicit as such. Consider, for instance, Shelly Kagan’s (1991) famous problem of paralysis for the deontological constraint against harming. The problem of paralysis arises considering the idea that ‘it cannot only

be *actually* causing harm to another which is forbidden [by constraint against harming]; some forms of merely *risking* harm to another must be prohibited as well' (p.87, my addition in square brackets).⁹ Accordingly, if it's impermissible to harm others, then it is also impermissible to risk harm, other things being equal.

Kagan's discussion suggests that potentially all our everyday, mundane actions are rendered impermissible by virtue of carrying even a small risk of harm. Keeping aside the significance and implications of this problem for our everyday conduct, it seems that the problem of paralysis is itself rooted in and suggestive of a prior connection between the morality of harming someone and that of risking harm, just as Unificationism contends. If the wrongness of (pure or impure) risking harming were normatively irrelevant to, or independent of the wrongness of harming, then the constraint against the latter would not apply to cases of risk and thereby generate the problem of paralysis in the first place.

Third, consider the following conditional. If and when it is *pro tanto* wrong for someone to φ , it is *pro tanto* wrong for them to (purely) risk φ -ing, other things being equal. For instance, if it is *pro tanto* wrong for me to injure someone, then it is *pro tanto* wrong for me to risk injuring them. Conversely, if and when it is not *pro tanto* wrong for someone to φ , then it is not *pro tanto* wrong for them to (purely) risk φ -ing, other things being equal.¹⁰ For instance, if it is not *pro tanto* wrong for me to injure someone to save their life, then it is not *pro tanto* wrong for me to risk injuring them to save their life. Of course, these different relations don't hold uniformly in all cases.

For instance, we routinely confront cases where φ -ing is *pro tanto* or all-things-considered wrong, but pure risking φ -ing is still considered all-things-considered permissible. Consider driving to work. The individual and the social benefits involved in driving seem to offset or outweigh the risks involved and, subsequently, outweigh the reasons against driving. Notwithstanding, we can grant that the obtainment of the conditional(s) stated above in some or enough number of cases motivates the Unificationists' thesis. The motivation comes from the idea that on some, if not every occasion that the conditional holds, the reason why

⁹Note that Kagan's problem of paralysis is distinct from Hayenhjelm & Wolff's (2011) discussion of the problem of paralysis. The latter concerns the issue of assigning people individual rights not to be exposed to risks insofar as virtually everything we do imposes some risk on others.

¹⁰Alternatively, we might say that if and when it is permissible for someone to φ , it is permissible for them to risk φ -ing, other things being equal (See Lackey 1986, p.647-8; Thomson 1986, p.177). See Peterson & Seidel (2021) for a discussion of this conditional (and its many other variations), which they call the deontic transfer principle.

purely risking φ is *pro tanto* wrong refers to the reason(s) why φ -ing is *pro tanto* or all-things-considered wrong, other things being equal. This is not unfamiliar territory once we consider the same conditional as it applies to other distinct act types.

For instance, consider the conditional if and when lying to your parents (act_1) is wrong, then lying to your mother (act_2) is also wrong. Or consider the conditional if and when killing (act_1) is wrong, then attempting (or trying) to kill (act_2) is also wrong. It seems intuitively plausible that these conditionals obtain in some cases, and at least on some occasions when they obtain, the reasons against performing (act_2) refer to (or simply are) the reasons against performing (act_1), other things being equal. *Mutatis mutandis* for act type φ -ing and act type pure risking φ -ing.¹¹

Fourth and relatedly, Unificationism is also suggested by a relation of co-variation between the wrongness of risking φ -ing and the wrongness of φ -ing even in cases where φ -ing fails to materialize. This can be illustrated by considering variations in the *degree* of wrongness of φ -ing and that of pure risking φ -ing. Other things being equal, if there is an increase in the former, then the latter seems to vary accordingly, where this variation goes in one direction (from reasons against φ -ing to reasons against pure risking φ , not *vice versa*). If the two were not related in this way, we would not see some monotonously increasing variation or difference between the two.¹² It is hard to think of cases where there is variation in the degree to which φ -ing is wrong, but the degree of wrongness of pure risking φ remains constant, other things being equal.

Lastly, pure risking is a four-place relational property in the cases that are of interest. It holds between an agent who subjects risk onto someone or something, the probability of some unwanted act (or outcome) obtaining, and the risked act (or outcome) itself. Because of this relationality, it seems natural to think that if pure risking is morally significant, then its significance is partly captured by or explained by facts about what it stands in relation to, such as facts about the wrongness of the risked act or outcome that failed to realize. This, too, suggests the Unificationist thesis.

These reasons make a good case for the *prima facie* plausibility of Unificationism, independently of whether (and to what extent) they also support Isolationism. For instance, consider Isolationist accounts that

¹¹ I thank an anonymous reviewer for pressing me on this point.

¹² Thanks to Andreas Schmidt for suggesting this point.

claim that pure risking is wrong because risking is harmful (Finkelstein 2003, Oberdiek 2012, Placani 2016). If this is correct, then the constraint against harming applies to cases of pure risking in virtue of pure risking itself being harmful. Moreover, we observe co-variation between the wrongness of harming and the wrongness of pure risking in virtue of pure risking itself being harmful. The fact that these considerations seem to favor isolationism equally may lead one to doubt the attraction and motivation of Unificationism as a stand-alone alternative to Isolationism.¹³

Notwithstanding, there are compelling reasons why anyone, especially Isolationists, should still take Unificationism seriously. First, the harm cited in both our everyday intuitions and language practices of explaining why purely risking harm is sometimes wrong concerns the risked ‘outcome-harm’ (alternatively, the potential or expected harm) and not the ‘risk-harm’ (the harm of risking itself) that the aforementioned Isolationist account identify.¹⁴ This suggests that if Isolationists want to make sense of our ordinary folk explanation of why purely risking harm is wrong, then they would still need to appeal to the Unificationist approach.

Second, as some have recently argued, belief- or evidence-relative risks are not harmful because they cannot interfere with one’s interests (Rowe 2022) or have any tangible effect on the risk victims. After all, such risks ‘exist only in the mind of the risk-taker’ (Parry & Slavny 2019, p.83). Yet, most would agree that it is sometimes wrong to impose them (even if and when the objective risk is zero). However, if wrongful imposition of belief- or evidence-relative risking cannot be harmful, then without retorting to Unificationism, Isolationism alone proves insufficient for explicating why the constraint against harming applies to these cases or why reasons or features that make harming wrong can sometimes ground a right against being subject to these risks (Bowen 2022).¹⁵

Third and relatedly, if we accept the intuitive claim that the more serious the harm risked, the more wrongful it is to impose belief- or evidence-relative risk of that harm (co-variation), then we should also accept that the obtaining of this relation is neither motivated by nor motivates the Isolationist’s position on the basis that wrongful imposition of belief- or evidence-relative risks are harmful. It does, however, naturally

¹³Thanks to an anonymous reviewer for pressing me on this point.

¹⁴The terminological distinction between ‘outcome-harm’ and ‘risk-harm’ comes from Finkelstein (2003).

¹⁵This limitation proves especially problematic once we acknowledge that many instances of wrongful pure risking concern belief- or evidence-relative risks in light of an agent’s epistemic inability to access or acquire knowledge about objective risk (Rowe 2022) or due to practical limits agents face in quantifying objective risks (Burri 2022).

motivate Unificationism insofar as the wrongfulness of imposing belief- or evidence-relative risk of some harm varies with how wrongful it is to impose that harm.

Note that the above discussion should not be taken to imply that Unificationism is the *only* correct approach to developing an explanatory account of wrongful risking.¹⁶ Nor should it be taken to imply that Isolationism is misguided. On the picture I'm advancing, we can treat Isolationism and Unificationism as distinctively plausible approaches that don't compete with or contradict one another. We can attest this by considering an Isolationist account of our choice, say, one that explains the wrongness of risking in terms of its impact on the risk-bearer's overall freedom (Ferretti 2016).

Suppose that at the moment when Joe faces the risk of being shot by Bill, he has fewer safe and choice-worthy options at his disposal than he did before. If diminishing Joe's overall freedom is a substantive wrong-making feature of Bill's conduct, then we can say that Bill acts wrongly when he risks killing Joe because in doing so, he diminishes his overall freedom. This, to me, strikes as one (but not the only) plausible explanation for why Bill's risk is wrong that is distinct from, but also compatible with a Unificationist explanation that cites or appeals to the morality of actually killing Joe.

However, if we were committed to rejecting Isolationism altogether, then we would have to give up the aforementioned Isolationist explanation of why Bill acts wrongly.¹⁷ This, however, seems suspect without further argumentation. Accordingly, if both Unificationism and Isolationism are complementary approaches, then we might say that a complete theory of wrongful pure risking specifying a set of *totality* of wrong-making facts would often include wrong-making facts identified by different Isolationist accounts (if and when these facts are obtained) and ones identified by the Unificationist account. This naturally raises the following question. What precisely is the Unificationist story about which wrong-making facts explain why Bill acts wrongly in *Russian Roulette*?

3. Unificationism and grounding

The Unificationist idea presented so far, namely that there is an explanatory relationship between the wrongness of pure risking φ and the

¹⁶Contra Parr & Slavny (2019 p.83).

¹⁷*Mutatis Mutandis* for other Isolationist explanations of wrongful pure risking that I don't discuss here due to limitations of space.

wrongness of φ -ing is under-specified in at least two respects. Firstly, we haven't yet consolidated our Unificationist talk of the explanatory relationship. That is, it is unclear what exactly the nature of this relationship is. Secondly, we haven't yet clarified the relata of this relationship. That is, what is the *explanans* of the wrongness of pure risking φ according to Unificationism? Before tackling this latter question, I will first argue in favor of positing grounding as the operative explanatory relation.

Getting clear on this relation is important for ultimately motivating my Simple Account (§5) but also for explaining why the Buck Passing Account falls short as a successful Unificationist account (§5). Moreover, as we'll see later, appealing to certain formal properties of grounding will also help explicate why it seems difficult to connect the morality of pure risking φ -ing and that of φ -ing in cases like *Russian Roulette*. Those those sympathetic to something akin to Unificationism have suggested that there is a 'parasitic' (Fried 2012) or a 'close' (Parr & Slavny 2019) relationship between the two. Yet, they don't clarify what this close or parasitic relationship amounts to.¹⁸

Some have attempted to make this relationship precise by hinting at reduction, for instance, by noting that the wrongness of pure risking φ is 'derived recursively' from that of φ -ing (Perry 2007). Others refer to 'transference' (Oberdiek 2012, p.342-3) and 'inheritance' (Peterson & Seidel 2021, p.1186) to drive home the same point. The question of how the wrongness of pure risking φ is derived recursively, transferred, or inherited from the morality of φ -ing is, however, left unaddressed by these theorists. Moreover, it remains unclear whether these theorists take these relations to be explanatory in nature (or kind), or whether they treat them as mere placeholders for some other explanatory relation.

Notwithstanding this, the relation under consideration can be made more concrete by acknowledging that it cannot be causal, conceptual, or a hybrid of the two. It is not a causal relationship because given their causal inefficacy, moral facts about the wrongness of φ -ing do not *cause* the wrongness of pure risking φ .¹⁹ The relationship is also not a conceptual one because it is simply not a matter of conceptual truth that the wrongness of φ -ing explains the wrongness of pure risking φ . This leaves us in the domain of non-causal relations, amongst which there are several options that we can immediately set aside. These include, but are not

¹⁸Although Parr & Slavny appeal to grounding in their formulation of the Buck Passing Account, they don't motivate or clarify this relation in any detail.

¹⁹See Cornell realism for an exception to the view that moral facts are morally inefficacious (e.g. Sturgeon 1984, Railton 1986).

limited to, composition, identity, logical entailment, and supervenience. I'll briefly consider these relations in turn.

Consider composition. Some classic examples of composition include cases like 'a chair is composed of atoms', or 'a battle is a part of a war'. The explanation for why Bill acts wrongly in pure risking, however, is not composed of moral facts about what Bill is risking, or vice versa. It cannot be an identity relation because moral facts about what Bill is risking are not semantically or conceptually identical to the wrongness of pure risking killing Joe. Moreover, we are also not dealing with a relation of logical entailment. If moral facts about φ -ing pertain to actions that are an instantiation of some act type φ -ing, then it is not a matter of logical entailment that these same facts explain, capture, or apply to actions that are an instantiation of a distinct act type pure risking φ .

Next on our list is supervenience. On a general understanding, the relation of supervenience would hold that facts about the wrongness of pure risking φ supervene on facts about the morality of φ -ing.²⁰ However, supervenience, too, seems like the wrong candidate. Take risking harm as an example. A supervenience relation between the wrongness of risking harm and that of harming would mean that if it would be permissible for me to risk harm on someone, then it would also be permissible for me to harm them, other things being equal.

However, this isn't right insofar as it is sometimes permissible to risk harm to someone without it being permissible to harm them, other things being equal. For instance, while it is permissible for me to offer a starving individual some food despite the trivial but non-negligible risk that they might get food poisoning, it is not permissible for me to do so when it is certain or highly likely that consuming my food will give them food poisoning. Supervenience, however, seems to ignore that facts about the wrong of risking harm depend on facts at the level of risk too (e.g. the magnitude, degree, or nature of risk), amongst other things.²¹ Where does this leave us?

I think the most promising and appropriate way for the Unificationist to interpret the explanatory relation is in terms of grounding.²² This is for several reasons. First, grounding is commonly thought to bear an intimate relation to explanation. Some theorists consider it central to the very

²⁰As I understand it here, supervenience is simply the view that *A-properties* supervene on *B-properties* if and only if a difference in *A-properties* *requires* a difference in *B-properties* (McLaughlin & Bennett 2021).

²¹Thanks to Joseph Conrad for this suggestion.

²²Some philosophers distinguish metaphysical grounding from normative grounding (See Fine 2012). For the purposes of this paper, I leave it open whether these are two distinct kinds of grounding and simply adhere to the language of grounding.

notion of grounding that it picks out the 'in virtue of,' 'because of,' or 'make it the case' explanatory relation that is neither causal, strictly logical, nor conceptual.²³ On this way of understanding grounding claims, to say that A grounds B is to say A stands in a non-causal dependence relation with B to the effect of A providing an explanation of B.

Second, grounding is thought to express and fix the direction of priority between facts.²⁴ For instance, a fact is grounded when there are more fundamental facts that it non-causally derives from, such as when facts about wrongness are said to be grounded in facts about goodness. In the context of Unificationism, the parallel thought is that when we cite the grounds for the wrongness of pure risking φ , we cite facts that are strictly prior to it, namely, facts that pertain to the morality of φ . Grounding thus gives traction to the intuitive idea that facts about the morality of φ -ing are in some sense fundamental to facts about the morality of pure risking φ . It is the obtaining of the former facts that play a role in making it the case that facts about the latter obtain, if and when they do.

Third, within moral philosophy, discussions surrounding the moral explanation of something being right in virtue of something else are quite naturally and intuitively interpreted in terms of grounding. For instance, when consequentialists insist that the relevant facts about rightness obtain in virtue of the corresponding facts about goodness, they are not making a claim about type identity, set membership, or any kind of determinate–determinable relation. Rather, they insist that facts about rightness obtain in virtue of or are *grounded* by certain facts about goodness (Berker 2017). By parity of reasoning, we can understand Unificationism too as an approach concerning the grounding of moral facts (namely, the wrongness of pure risking φ) by other facts (namely, ones pertaining to the morality of φ -ing).²⁵

Finally, my contention that the explanatory relation is one of grounding is further supported by various hallmark features of grounding, such as asymmetry. Grounding relations are supposedly asymmetric (owing to its irreflexivity and transitivity). That is, if 'A grounds B', then this is made true by a single grounding relation running from A to B. This dependency relation is unidirectional. There is no inverse grounding relation that runs from B to A, which makes 'B is grounded in A' true. The same insight is at work once we acknowledge that it is facts about the wrongness of φ -ing

²³See Schaffer 2009, Audi 2012, Raven 2012.

²⁴See Wilson 2016.

²⁵Note that so far, I've purposely refrained from specifying whether the relevant facts that pertain to the morality of φ -ing are normative, natural, or both. I will return to this in my discussion in §3 and §4.

that ground facts about the wrongness of risking φ (rather than the other way around).

In light of these considerations, then, I suggest we reformulate the general Unificationist thesis as follows: when pure risking φ -ing is wrong, the fact that it is wrong is *grounded* in certain facts about the wrongness of φ -ing. Our task now is to answer the second question posed earlier, namely, identifying the *explanans* of what makes pure risking φ wrong according to the Unificationist approach.

4. An existing unificationist proposal: the buck passing account

In recent work, Parr & Slavny (2019) have offered us a Unificationist account for capturing what they consider to be a ‘close relationship’ between the *factors* that make an act wrong and the factors that make risking that act wrong. According to what they call the Buck Passing Account (henceforth, BPA),

‘When it is wrong for an agent to risk [φ -ing] on someone, the fact that it is wrong for him to risk [φ -ing] is grounded directly in the fact that he increases the probability of a set of facts {f} obtaining that would make it wrong for him to [φ]’ (Parr & Slavny 2019, p.83, my modification in square brackets).²⁶

Call the members in the set of facts {f} that would make φ -ing (the non-risked act such as Bill killing Joe) wrong, if they obtained, the *explanans* or grounding facts of the wrongness of φ -ing. These are simply the morally salient wrong-makers of φ -ing. Just why the facts that ground or explain the wrongness of pure risking φ are probabilistic facts regarding wrong-makers of φ -ing if they were to obtain is easily clarified by looking into how facts about the wrongness of some act are, in general, grounded.

According to a widely accepted way of understanding the grounding of moral (or normative or evaluative) properties of an act, no action is considered *brutely* wrong (Väyrynen 2013). Selim Berker (2019) helpfully illustrates this thought by noting that ‘[w]hen I do something wrong, there are certain properties of that action (and, perhaps, of other things in the world as well) that make what I do wrong’ (p.904). For instance, the fact that an act is a telling of a lie (a natural or a non-normative fact) grounds the moral fact that the act is morally wrong. Similarly, the fact that an act is a violation of someone’s right may ground the fact that the act is morally wrong.

²⁶I’ve replaced the author’s use of ‘v-ing’ to ‘ φ -ing’ to keep its usage consistent with my own discussion.

Importantly, the wrongness of a particular token-act in question is grounded by grounding facts instantiated at time t in world w . Accordingly, if the wrongness of Bill's killing Joe [p] is grounded by the obtaining of the salient wrong-making features in $\{f\}$ at time t in world w , then the wrongness of Bill's *risking* killing Joe is grounded, not by the actual obtainment of wrong-making features in $\{f\}$ at t in w , but rather, by the *likely* or probabilistic obtainment of wrong-making features in $\{f\}$ when Bill risks killing Joe at t in w .

In line with the Unificationist thesis, then, BPA unifies the grounds of the wrongness of pure risking φ with that of particular wrong-makers of φ -ing by simply passing the 'explanatory buck' to the latter.²⁷ At first glance, BPA seems promising and intuitively plausible as an account of what makes pure risking wrong. However, it is open-ended whether it is also theoretically appealing as a *general* Unificationist account. To test this, I propose the following desiderata.

First, the Extension Desideratum. It inquires whether the Unificationist account under consideration is extensionally adequate, that is, whether it has the right scope. As Parr & Slavny (2019) note, '[T]he correct view of wrongful risking must incorporate *exactly* the factors that can make risked acts wrong' (p.83). Accordingly, to be extensionally adequate in this way, BPA must capture the grounds of the wrongness of pure risking by successfully incorporating the full range of factors, whatever they are, that make φ -ing wrong.

Second, the Explanation Desideratum. It inquires whether the Unificationist account under consideration is explanatorily adequate. As the authors note, '[A]ny account of the wrongness of pure risking should be able to explain this general relationship between wrongness in non-risking cases and the wrong of risking' (p.81). Accordingly, to be explanatorily adequate in this way, BPA should capture and make sense of *why* pure risking φ is wrong in relation to increasing the likelihood of the particular wrong-making features of φ -ing obtaining.

4.1. Extensional adequacy

Consider, first, whether BPA fulfills the Extension Desideratum. According to its proponents, the intuitive plausibility of BPA 'does not depend on

²⁷Elsewhere, the authors note that the BPA passes the explanatory buck to the wrongness of the risked act (p.82-83) as opposed to passing the buck to the 'the wrong-makers' of the risked act. As the remaining discussion will make clear, these two are distinct *explanans* of what makes pure risking φ -ing wrong.

any particular account of the facts that ground moral wrongness. It holds regardless of whether the wrongness of the risked act is based on harm, autonomy, impersonal value, or any other wrong making property' (p.79). In other words, if the wrongness of Bill pure risking killing Joe is grounded in some set of facts $\{f\}$, then every given fact in $\{f\}$, regardless of what this fact is, plays some role in making it the case that Bill acts wrongly in *Russian Roulette*.²⁸

Presumably, there is a plurality of wrong-making facts in $\{f\}$ that ground why a given act is wrong. As the quote above already indicates, such facts may include violation of one's rights, diminishment of one's autonomy, reduction of one's overall freedom, the frustration of one's preferences, bad motives, intentions, negligence, recklessness, and so on. Let's suppose that violating Joe's right not to be killed is one member of the set of facts $\{f\}$ that grounds the wrongness of Bill actually killing Joe. In playing Russian Roulette, if Bill wrongly ends up shooting and actually killing Joe, then the instantiation of the said wrong-making fact grounds why Bill acts wrongly, other things being equal.

Accordingly, BPA holds that when Bill acts wrongly in pure *risking* shooting and killing Joe, the fact that he acts wrongly is directly grounded in the fact that he increases the probability of violating Joe's right not to be killed, other things being equal. To take another example, suppose that shooting and killing Joe would be wrong because it frustrates Joe's preference not to die. According to BPA, Bill's risking that act is wrong because doing so increases the likelihood of frustrating Joe's preference not to die. Likewise, if shooting and killing Joe would be wrong because it diminishes his autonomy, then according to BPA, risking that act is wrong because doing so increases the likelihood of diminishing Joe's autonomy, other things being equal. All this, then, just strictly follows from the logic of the account as formulated by its proponents.

For the aforementioned examples, it seems that BPA gets Unificationism right. But if BPA is the correct *general* Unificationist account, then for the sake of extensional adequacy, we should expect it to generate structurally similar explanations for why Bill's risking killing Joe is wrong for *any* and *every* other wrong-maker in $\{f\}$, regardless of what it is. This means that given any wrongful act φ -ing, if we picked out at random some wrong-making fact in $\{f\}$, we should expect BPA to explain why purely risking φ -ing is wrong in virtue of increasing the probability of *that*

²⁸This seems to be generally true of grounding. As Rosen (2010) notes, if $[x]$ is grounded in some set of facts S , then every fact in S 'plays some role in making it the case that $[x]$ ' (p.116).

wrong-making factor (or feature) obtaining. However, BPA falls short in this regard for at least some members in $\{f\}$. Let me illustrate this with two concrete examples.

Suppose that, for the sake of the argument, intentions can sometimes explain why some non-risky act is wrong. We needn't assume a particular story of what intentions are or what imparts them the status of a wrong-making fact or moral permissibility-determining feature of an act, assuming they enjoy this status.²⁹ For now, we can work with the following simple picture. An agent's intention, understood here as the agent's aim in performing a particular act and how he sees that act as promoting his objective, can sometimes make that act wrong or impermissible (Liao 2012).

Now, as Parr & Slavny (2019) themselves note, '[I]n cases where an act is wrong because of an agent's intentions, we also explain the wrongness of risking these acts with reference to *these* intentions' (p.84, my emphasis in italics). This seems correct insofar as there is no principled objection against appealing to Bill's intentions in explaining the wrongness of risking killing Joe *if* we also appeal to his intentions in explaining the wrongness of actually killing Joe. However, if and when intentions are the salient wrong-makers of a token act of killing someone, BPA fails to generate the correct Unificationist explanation of what makes Bill's pure risking killing Joe wrong.

Recall that the BPA connects the explanatory grounds of the wrongness of pure *risking* φ with probabilistic facts regarding particular wrong-makers of φ -ing if they were to obtain. Following this, BPA tells us that if intentions are (one of) the salient wrong-maker, then Bill's purely risking killing Joe is wrong insofar as in doing so, he raises the probability of that wrong-maker obtain, namely, his intention to kill him. Now, it might be that *sometimes*, in risking killing an individual, an agent acts wrongly by virtue of increasing the probability of forming an intention to kill that individual. Or it might even be that sometimes, in risking killing an individual, an agent acts wrongly in virtue of increasing the probability of forming *and* also acting upon their intention to kill that individual.

To use one of Parr & Slavny's own (slightly modified) example, it might be that Bill decides to flip a coin such that if it comes up heads, he will *then* form and act on the intention to kill Joe and inflict harm on him;

²⁹Parr & Slavny (2019) acknowledge that it is debatable whether intentions enjoy the status of wrong-makers or whether they affect the permissibility of actions. For my purposes, I will assume, like the authors, that intentions sometimes enjoy this status.

whereas if it comes up tails, he will leave Joe alone.³⁰ In such a case, Bill's flipping the coin is itself a risky act. BPA explains the wrongness of performing this act by appealing to a probabilistic fact about the obtainment of Bill's intention to kill Joe (that is, the wrong-maker of his non-risky act were it to actually obtain).

While it appears that BPA captures cases like this and ones alike correctly, our original case of *Russian Roulette* is different from them. By stipulation, Bill plays Russian Roulette with Joe for the sheer fun of risking. In doing so, he neither intends to kill Joe nor does he increase the likelihood of forming an intention to kill him when he holds the gun against Joe's head for fun (at least this seems true within the confines of the original example). Accordingly, if BPA takes intentions to be one of the salient wrong-makers of Bill's non-risky act, then it is hard to see how it captures and explains cases like *Russian Roulette*.

More importantly, even if we stipulated that in *Russian Roulette*, Bill acts riskily with the intention to kill Joe, then the fact that he acts wrongly in risking Joe's death is *not* grounded or explained by the obtainment of the *probabilistic* fact that he increases the likelihood of intending to kill him, as BPA would say. Rather, it is directly and straightforwardly grounded or explained by the obtainment of the *non-probabilistic* fact that he intends to kill him. BPA, however, does not vindicate this intuitively plausible Unificationist explanation.³¹

Consider another example. According to Tim Scanlon's contractualism, an act is wrong when and *in virtue of* being prohibited by a moral principle that no one can reasonably reject.³² Suppose, then, that violation of a non-rejectable moral principle (say, 'treat others with respect') is a salient wrong-making feature in {f}. If we accept it as a salient wrong-making feature of a non-risky act, then BPA holds that in purely risking killing Joe, Bill acts wrongly because in doing so, he increases the likelihood of the instantiation of an act being prohibited by a non-rejectable moral principle ('treat others with respect'). Yet, as most contractualist discussions on the topic illustrate, the wrongness of pure risking is directly and straightforwardly grounded or explained by the obtainment of the

³⁰Here, the chances of forming the intention to kill are 50–50 by stipulation.

³¹Some think that besides intentions, an act can sometimes be wrong by virtue of the motivation with which it was performed (Tadros 2011). Suppose an agent's motivation (say, to have fun) is relevant to explaining why killing someone is wrong. In that case, BPA holds that risking killing Joe is wrong because in doing so, Bill raises the probability of his motivation (to have fun) Joe obtaining. This, too, illustrates that BPA gets the Unificationist explanation wrong if and when an agent's motivations are the salient wrong-making feature of an act.

³²For a discussion of whether Scanlon's contractualism describes one way or the only way in which acts are wrong, see Frick 2015.

non-probabilistic grounding fact that risking is prohibited by a non-rejectable moral principle (based on either *ex ante* or *ex post* complaints).³³

As a general result, then, various wrong-makers of a non-risky act do not always relate to the wrongness of pure risking as a matter of being grounded by a probabilistic fact, *contra* BPA. When Bill acts wrongly in risking killing Joe, he does not increase the probability of instantiation or obtainment of wrong-makers of killing Joe such as intentions, motivational state, violation of some non-rejectable moral principle, legitimate expectation or social norm, and so on. Instead, in explaining why Bill's risking killing Joe is wrong, these aforementioned wrong-makers feature directly in our *explanans* as a matter of a non-probabilistic grounding fact (except perhaps in some exceptional circumstances).

Thus, investigating further into *which* particular facts ground the moral wrongness of φ -ing shows that there is no *one* general relationship between these facts and the grounds of the wrongness of risking φ . Instead, there are (at least) two distinct relationships. Some wrong-makers are explanatorily relevant in the way BPA takes them to be, namely, as a matter of probabilistic grounding fact, whereas others appear to be explanatorily relevant as a matter of non-probabilistic grounding fact. The upshot, then, is that we end up with a watered-down Unificationist account of why pure risking φ is wrong, insofar as we can only appeal to a subset of wrong-makers of φ -ing in $\{f\}$ in explicating the wrongness of pure risking φ . Thus, BPA proves to be extensionally limited.

4.2. Explanatory adequacy

Consider, next, whether BPA fulfills the Explanatory Desideratum. Stated as such, it is not immediately clear or obvious *why* pure risking φ -ing is wrong in virtue of an agent increasing the likelihood of some wrong-making features of φ -ing obtain. In this regard, we might wonder whether BPA is explanatorily adequate by way of *explaining* why the relevant explanatory relation holds. But as the authors themselves admit, their account 'relies on a basic normative fact and therefore it is difficult to defend through independent argument'(p.84). Notwithstanding this, it is still doubtful whether BPA is explanatorily adequate by way of *capturing* the fact that Bill 'increases the likelihood of wrong-

³³This is, of course, a very simplistic way of presenting a complex debate on the contractualist morality of wrongful pure risking. See Kumar (2003& 2015), James (2012), Frick (2015) Suikkanen (2019).

making features of φ -ing obtaining' as the correct *explanans* of why his conduct is wrongful.

If increasing the likelihood of the particular wrong-maker killing Joe is simply what constitutes risking killing Joe, or in other words, it is simply a description of what it is for Bill to *risk* killing Joe, then on one reading, BPA merely restates or reaffirms what risking *involves* as an explanation for why it is wrong when it is. This is problematic insofar as just saying that 'Bill's risk is wrong because it involves risking' is to beg the question.

On a different reading, perhaps the BPA explains why Bill playing Russian Roulette with Joe is wrong by identifying risking itself as a salient wrong-making feature of his conduct. Understood this way, the worry about question-begging does not arise insofar as our initial explanandum (why risking is wrong) now becomes the explanans (that some action constitutes pure risking and hence wrong). This, too, however, does not get us any closer to answering *why* pure risking killing is wrong.

While its plausible that an action is wrong because it constitutes risking φ , this is not a brute fact insusceptible of any further explanation. It is instead a fact that obtains in virtue of further facts, and on the Unificationist approach, facts about the wrongness of φ -ing. It is thus not unintelligible to ask (and this is indeed the Unificationist's way of thinking) why the fact that performing some action that would involve risking count against it in the first place. That is precisely the question the Unificationist wants to address. However, it appears that BPA falls short in answering this if, as a Unificationist account, it merely captures or identifies risking as a wrong-making feature of some action. Besides, it's hard to reconcile this with Parr & Slavny's Unificationist idea that BPA 'passes the explanatory buck by referring to the wrongness of the risked act' (p.83), insofar as the buck is not passed if the fact of risking φ is what makes Bill's conduct wrong. Thus, the BPA does not satisfy the Explanatory Desideratum.³⁴

5. A new unificationist proposal: the simple account

Taking stock, the discussion of BPA is insightful in at least two ways. First, as I've argued, it captures the Unificationist's explanatory relationship with some wrong-makers of φ -ing correctly, but arbitrarily leaves out some others. It is an unattractive feature of any Unificationist account

³⁴This is especially problematic insofar as Parr & Slavny propose BPA as an alternative to the Isolationist approach and reject the latter on the grounds of being explanatory inadequate.

to exclude certain wrong-makers of φ -ing in this way. This means that a satisfactorily broad Unificationist account of the wrongness of risking φ cannot fully or only concern probabilistic grounding facts. Second, the discussion also explicates that for our Unificationist account to be explanatory in nature, we should go beyond simply redescribing pure risking φ or identifying risking φ as a salient wrong-making feature of one's conduct under consideration.

Moving forward, I suggest that we appropriately modify and reformulate BPA in a way that accommodates or subsumes the distinct explanatory relationships that hold between different grounds of the wrongness of φ -ing and that of pure risking φ , whilst also adequately capturing the Unificationists' intuition that the former stands in an explanatory relation with the latter. In line with this suggestion, I now propose the following as a general Unificationist account. According to what I call the Simple Account,

When it is wrong for an agent to purely risk φ -ing, the fact that it is wrong for the agent to purely risk φ -ing is grounded directly in a general moral fact that φ -ing is *pro tanto* (or all things considered) wrong *simpliciter*.³⁵

To understand exactly how the Simple Account captures cases like *Russian Roulette*, think of Bill as someone who acts similarly to an agent who tries to φ but fails.³⁶ When Bill points his loaded gun at Joe, he is engaged in doing something that *could* have, in more propitious circumstances, become, or culminated into φ -ing if and when certain sufficient and/or necessary conditions for the fruition or completion of the act of φ -ing had obtained. Accordingly, Bill's act of risking φ , similar to some agent's act of trying to φ , may be thought of as an *incipient* act of φ -ing.

As Matthew Hanser (2014) argues, the wrongness of incipient acts, such as trying to do something, can be evaluated with reference to that which they *could* become or in which they could culminate, even if and when trying to φ ultimately proves unsuccessful. As he notes, '[W]e can judge that an agent *acted* wrongly insofar as he φ -ed, or that he

³⁵As the Simple Account appeals to the general moral fact that it is *pro tanto* wrong to φ *simpliciter*, it can be interpreted suitably to extend to cases where the agent risks something that is *pro tanto* wrong to bring about (like an event, or outcome, or state of affair) as well as cases where the agent risks doing something (action) that is *pro tanto* wrong to do.

³⁶Note that I don't mean to say that risking and trying are the same type of actions or that risking is an instance of trying. I only mean that in some respect, risking is similar to trying. I acknowledge that both are distinct types of action in at least one following way: trying necessarily involve an agent's being motivated via an intention towards bringing about what he is trying to bring about; whilst risking does not (necessarily) involve an agent's being motivated via the intention to bring about what he is risking bringing about.

was acting wrongly insofar as he was (trying) φ -ing. Both judgments are supported by the same underlying fact: that it was, in the circumstances, wrong to φ ' (2013, p.149).

Something similar, I contend, can be said more generally of cases of pure risking φ , even if and when pure risking φ is not strictly always an act of trying, nor an incipient act of φ -ing. In particular, we could say of the agent that he has done something wrong in *pure* risking φ -ing in virtue of the *general* moral fact that φ -ing is wrong because he was, at a time, risking that which *could* have culminated into or become something that was, in the circumstances, *pro tanto* wrong. As such, if it is *pro tanto* wrong to φ , then this general moral fact licenses both judgments of the form that 'the agent acted wrongly insofar he was φ -ing' and 'the agent was acting wrongly insofar as he was risking φ '.

Many theorists seem to accept the idea that general moral facts, also referred to as moral principles, like 'killing is wrong' and 'keep your promises,' play an important fundamental role in morality. At least some who accept this hold that one theoretical function of moral principles is that they are explanatory in nature (Fogal & Risberg 2020; Väyrynen 2013 & 2018).³⁷ That is, general moral facts or principles play a special role in the *explanation* of moral phenomena and accordingly, they are considered by some theorists as explanation-serving and not just explanation-involving.³⁸

Yet, despite its initial plausibility, it can be puzzling just how the Simple Account gets around the predicament noted earlier that has motivated Isolationists. If it is *pro tanto* wrong for me to kill someone, for instance, then this is presumably because my successful act of killing someone possesses some morally salient wrong-making feature that gives me a reason against performing such an act. This seems to be motivated by the idea that the moral fact that it is wrong for me to kill is explained by the particular non-moral (or natural) fact about my act, namely, that it is an instance of killing, together with a general explanatory moral principle, namely, that it is wrong to kill (Fogal & Risberg 2020).

But if it is a characteristic feature of instances of *pure* risking that they are unsuccessful or unrealized token acts of φ -ing, then how can the wrongness of an agent's conduct be grounded by the general fact that

³⁷In making this objection, I am siding with those theorists who affirm that general moral principles play some fundamental role in morality. By contrast, moral particularists reject this position. See Väyrynen 2018 for the debate between moral generalism and moral particularism.

³⁸See Berker 2019 for an opposing view.

φ -ing is wrong given that the act type φ -ing is never instantiated?³⁹ The force of this objection comes to shine in Judith Thomson's (1986) discussion of *Russian Roulette*. As she writes, '[I]t is preferable to say (as I suggested) that if I play Russian Roulette on you, and there is no bullet under the firing pin when I fire, then I do not harm you. So, *we do not have available the fact that my doing so* is my harming you as a ground for thinking that I infringe a right of yours when I do so – for there is no such fact' (1986, p.165, my emphasis in italics).

There is some kernel of truth to Thomson's point, at least if we understand her as saying that the *particular* moral fact of *my* harming you cannot ground the wrongness of my risking harm upon you in cases like *Russian Roulette*. We cannot appeal to the former to explain why the latter holds insofar as the wrongness of *an agent* harming (or φ -ing) is instantiated when the token action of an agent harming someone obtains. How, then, do we reconcile this difficulty and the Unificationist intuition captured by the Simple Account? The two appear to be in tension with each other.

Note that Thomson considers the non-obtainment of a *particular* moral fact, namely, the wrongness of 'my killing someone' (or my φ -ing), as a hindrance to grounding the wrongness of 'my risking killing someone' (or my risking φ) when I (luckily) fail to kill someone. The mistake, however, is to overlook the aforementioned distinction between *particular* and *general* moral facts.⁴⁰ A particular moral fact is a moral fact about a particular (dated, non-repeatable) thing, such as a particular action, person, or state of affairs. It grounds *only* particular concrete token instances of actions of that type when it obtains (e.g. *Bill's* killing Joe is a concrete token action that obtains when Bill actually kills Joe).

By contrast, general, non-particular facts, including facts such as pain is bad or that killing is pro tanto wrong, don't ground in a time-bound (nor in a world-bound) way. These facts are fundamental (and thus ungrounded) and are widely assumed to obtain as a matter of 'metaphysical necessity.' Given this, we can explain why Bill acted wrongly in pure risking φ by directly appealing to the *general* moral fact that it is wrong to φ , insofar as obtainment of this fact is not dependent on or requires that Bill φ s.

³⁹As noted earlier, this apparent difficulty seems to have been the driving force behind some Isolationist accounts.

⁴⁰In discussing the case of Bill recklessly drunk driving and risking harming an agent, Parr & Slavny (2019) also seem to make incorrectly note that we explain Bill's wrong by direct appeal to a *particular* moral fact, namely, wrongness of harming that agent (p. 83).

This holds even if pure risking, by definition, doesn't result in the instantiation of the particular moral fact of Bill acting wrongly in φ -ing *ex post* either by luck (perhaps because the gun falls out of his hand or it simply fails to fire a shot) or due to some other reason (maybe because Bill abandons his plan to play Russian Roulette mid-way). It suffices that Bill merely risks φ to say that he acts wrongly in doing so because φ -ing is wrong, where the latter is a general moral fact in the sense explained here.

Thus, if a general moral fact like killing is wrong *itself* plays a role in explaining particular moral facts involving instances of Bill's killing Joe, then *that* general fact is also part of the set $\{f\}$ that encompasses all wrong-makers of killing that are explanatory relevant for capturing why Bill's pure risking killing Joe is wrong.⁴¹ In this way, the Simple Account attends to the 'close relationship between *factors* that make an act wrong and factors that make risking that act wrong' that Parr & Slavny (2019) set out to capture by proposing the BPA. Yet, the Simple Account differs from the BPA in a very important way.

The BPA takes the explanans for why risking φ is wrong to include the particular wrong-making features that make φ -ing wrong as a matter of a *probabilistic* grounding fact. Whereas, the Simple Account takes the explanans for why risking φ is wrong to include the general fact that what one is risking is *pro tanto* wrong to do as a matter of a *non-probabilistic* grounding fact. This is because such principles do not stand in *any* causal or non-causal relation with how one acts or behaves, let alone standing in a probability-raising relationship. They are often regarded as obtaining as a matter of necessity rather than probability (Väyrynen 2018).

However, this does not mean that the Simple Account treats particular wrong-makers of φ -ing identified by the BPA as irrelevant. Rather, the account allows us to derive a number of distinct *reductive* explanations of the wrongness of risking φ at the level of particular wrong-makers of φ -ing. This is because the fact that an act is *pro tanto* wrong is itself reducible to particular facts or features in virtue of which that act is *pro tanto* wrong. One of these reductive explanations will replicate the BPA and accordingly, the Simple Account subsumes the BPA and retains Parr & Slavny's Unificationist idea that some wrong-makers of φ -ing ground the wrongness of risking φ -ing as a matter of a probabilistic grounding fact.

⁴¹Those who accept this view endorse a tripartite structure of moral explanation of what makes an act wrong.

Other reductive explanations will take the following form: if φ -ing is wrong, then that which grounds the wrongness of φ -ing also grounds the wrongness of pure risking φ . Accordingly, the Simple Account retains my Unificationist idea that fundamental, general moral facts (as well as agent-relative wrong-making features of φ -ing) have an explanatory role to play, as a matter of non-probabilistic fact, in grounding the wrongness of pure risking φ , given some wrongful act type φ . I consider it a theoretical virtue of the Simple Account that it allows us to derive these different Unificationist explanations without pre-committing us to one or the other, unlike the BPA.

As an upshot, the Simple Account avoids concerns around extensional adequacy insofar as the account does not arbitrarily exclude any particular wrong-makers of φ -ing in capturing the wrong of pure risking. The Simple Account is pluralistic in this regard. Moreover, in its general formulation, the Simple Account does not commit itself to a particular kind of relationship (whether it is probabilistic or non-probabilistic) that holds between the morality of φ -ing and that of pure risking φ , unlike the BPA.

Besides, the account also bypasses problems of extensional inadequacy in cases where the particular grounds of the wrongness of φ -ing are unknown or where there is disagreement over the particular grounds of some non-risky act.⁴² Additionally, it also captures cases that some Isolationist accounts fail to accommodate. Consider the following version of *Russian Roulette*:

Russian Roulette (Fake): Everything as *Russian Roulette*, except this time, unbeknownst to Bill, the gun is fake. There is a one in six belief-relative risk that the gun will shoot a bullet and injure him. No bullet fires when Bill pulls the trigger and Joe is left unharmed.

Some argue that Isolationist accounts that locate the wrongness of pure risking φ in contingent effects cannot capture cases like *Russian Roulette (Fake)*, where the pure risk of killing Joe is purely epistemic and there is no objective or fact-relative risk of dying that befalls Joe.⁴³ As a result, these accounts are explanatorily limited; they only explain cases of pure risking where Bill acts wrongly in a fact-relative sense but not in a belief- or evidence-relative sense.⁴⁴ On the Simple Account, the fact that the risk imposed is epistemic (and thus an instance of belief- or evidence-relative

⁴²For instance, there could be cases where we don't know the wrong-making features of action but know that it is wrong to act that way.

⁴³See Parr & Slavny (2019, p.77-78 & 81-82) for discussion of these cases.

⁴⁴For an argument along these lines, see Parr & Slavny's (2019) Determination Objection and Rowe's (2022) Argument from Interference.

wronging) is not a deterrent to generating a Unificationist explanation of why Bill acts wrongly.

The account holds that the wrongness of evidence-relative pure risking in cases like *Russian Roulette (Fake)*, similar to wrongful cases of fact-relative pure risking, is explained by the general moral fact that killing Joe is wrong. As an upshot, the Simple Account unifies distinct cases, irrespective of whether it involves wronging others in an evidence- or belief- or fact-relative sense. This way, the Simple Account not only retains but also extends the scope of the Buck-Passing Account and thus fulfills the Extensional Desideratum. Finally, insofar as the Simple Account passes the explanatory buck to the wrongness of φ -ing without identifying risking itself as a wrong-maker of some acts and without appealing to that which it is trying to explain as part of the explanation itself, it also fulfills the Explanatory Desideratum.

6. Conclusion

Cases of wrongful pure risking raise an interesting puzzle for the ethics of risk: While it seems that the morality of pure risking stands in a close explanatory relationship with that of non-risky acts, it remains unclear and non-obvious how to make sense of this in cases where the risk fails to materialize. By motivating Unificationism and explicating this explanatory relationship in terms of grounding, I've argued that a plausible distinction between general and particular moral facts allows us to capture the explanatory intuition that risking harm (or some other unwanted outcome) can be sometimes wrong in virtue of the wrongness of the harm itself, even if and when the latter fails to realize. In doing so, I've defended Unificationism (and, in particular, the Simple Account) as a plausible approach within the ethics of pure risk, one that deserves its own place in existing discussions that are currently dominated by Isolationist accounts.

Acknowledgments

I would like to thank audiences at the 1st Munich Graduate Conference in Ethics (2020), work in progress PhD meeting at Faculty of Philosophy Groningen, and the OZSW Annual Philosophy Conference (2020) for their helpful feedback. A special thanks to Adam Slavny, Alec Walen, Andreas Schmidt, Frank Hindriks, John Oberdiek, Sven Nyholm, and two anonymous reviewers for their valuable written comments that helped improve the paper. I would also like to extend my gratitude to Conrad Bakka, Daan Evers, Lukas Wolf, and Michael Gregory for multiple stimulating discussions on the topic.

Disclosure statement

No potential conflict of interest was reported by the author(s).

Funding

This work was supported by Rijksuniversiteit Groningen: [Grant Number PhD Scholarship Bursary].

Bibliography

- Audi, P. 2012. "Grounding: Toward a Theory of the *In-Virtue-of* Relation." *Journal of Philosophy* 109 (12): 685–711. <https://doi.org/10.5840/jphil20121091232>.
- Berker, S. 2017. "The Unity of Grounding." *Mind; A Quarterly Review of Psychology and Philosophy* 127 (507): 729–777. <https://doi.org/10.1093/mind/fzw069>.
- Berker, S. 2019. "The Explanatory Ambitions of Moral Principles." *Noûs*, <https://doi.org/10.1111/nous.12246>.
- Bowen, J. 2022. "Robust Rights and Harmless Wronging." In *Oxford Studies in Normative Ethics Volume 12*, edited by Mark Timmons, 31–53. 1st ed. Oxford: Oxford University Press.
- Burri, S. 2022. "Conceptualising Morally Permissible Risk Imposition without Quantified Individual Risks." *Synthese* 200 (415). <https://doi.org/10.1007/s11229-022-03888-4>.
- Cripps, E. 2013. *Climate Change And The Moral Agent*, 27–57. 1st ed. Oxford: Oxford University Press.
- Ferretti, M. 2016. "Risk Imposition and Freedom." *Politics, Philosophy & Economics* 15 (3): 261–279. <https://doi.org/10.1177/1470594X15605437>
- Fine, K. 2012. "Guide to Ground." In *Metaphysical Grounding: Understanding the Structure of Reality*, edited by Fabrice Correia, and Benjamin Schnieder, 37–80. 1st ed. Cambridge: Cambridge University Press.
- Finkelstein, C. 2003. "Is Risk a Harm?" *University of Pennsylvania Law Review* 151 (3): 963–1001. <https://doi.org/10.2307/3312883>
- Fogal, D., and O. Risberg. 2020. "The Metaphysics of Moral Explanations." In *Oxford Studies in Metaethics Volume 15*, edited by R. Shafer-Landau, 170–194. 1st ed. Oxford: Oxford University Press.
- Frick, J. 2015. "Contractualism and Social Risk." *Philosophy and Public Affairs* 43 (3): 175–223. <https://doi.org/10.1111/papa.12058>
- Fried, B. 2012. "The Limits of Non-Consequentialist Approach to Tort." *Legal Theory* 18 (3): 231–262. <https://doi.org/10.1017/S1352325212000183>.
- Fried, B. 2020. *Facing Up to Scarcity: The Logic and Limits of Nonconsequentialist Thought*. 1st ed. Oxford: Oxford University Press.
- Frowe, H. 2021. "Risk Imposition and Liability to Defensive Harm." *Criminal Law & Philosophy*, <https://doi.org/10.1007/s11572-021-09588-3>.
- Hanser, M. 2014. "Acting Wrongly by Trying." In *Oxford Studies in Normative Ethics Volume 4*, edited by M. Timmons, 138–158. 1st ed. Oxford: Oxford University Press.

- Hayenhjelm, M., and J. Wolff. 2011. "The Moral Problem of Risk Impositions: A Survey of the Literature." *European Journal of Philosophy* 20: E26–E51.
- James, A. 2012. "Contratualism's (Not So) Slippery Slope." *Legal Theory* 18 (3): 263–292. <https://doi.org/10.1017/S135232521200002X>.
- James, A. 2016. "The Distinctive Significance of Systemic Risk." *Ratio Juris* 30 (3): 239–258. <https://doi.org/10.1111/raju.12150>.
- Kagan, S. 1991. *The Limits of Morality*. 1st ed. Oxford: Oxford University Press.
- Kumar, R. 2003. "Who Can Be Wronged?" *Philosophy & Public Affairs* 31 (2): 99–118. <https://doi.org/10.1111/j.1088-4963.2003.00099.x>.
- Kumar, R. 2015. "Risking and Wrong." *Philosophy and Public Affairs* 43 (1): 27–51. <https://doi.org/10.1111/papa.12042>
- Lackey, D. P. 1986. "Taking Risk Seriously." *Journal of Philosophy* 83 (11): 633–640.
- Liao, S. M. 2012. "Intentions and Moral Permissibility: The Case of Acting Permissibly with Bad Intentions." *Law and Philosophy* 31 (6): 703–724. <https://doi.org/10.1007/s10982-012-9134-5>.
- Maheshwari, K., and S. Nyholm. 2022. "Dominating Risk Impositions." *Journal of Ethics* 26 (4): 613–637. <https://doi.org/10.1007/s10892-022-09407-4>
- McCarthy, D. 1997. "Rights, Explanation, and Risks." *Ethics* 107 (2): 205–225. <https://doi.org/10.1086/233718>
- McLaughlin, B., and K. Bennett. 2021. "Supervenience." In *The Stanford Encyclopedia of Philosophy*, edited by Edward N. Zalta. <https://plato.stanford.edu/archives/fall2021/entries/supervenience/>.
- Oberdiek, J. 2009. "Towards a Right Against Risking." *Law and Philosophy* 28 (4): 367–392. <https://doi.org/10.1007/s10982-008-9039-5>
- Oberdiek, J. 2012. "The Moral Significance of Risking." *Legal Theory* 18 (3): 339–356. <https://doi.org/10.1017/S1352325212000018>
- Oberdiek, J. 2017. "Imposing Risk: A Normative Framework." In *Oxford Legal Philosophy*. 1st ed. Oxford: Oxford University Press.
- Parr, T., and A. Slavny. 2019. "What's Wrong with Risk?" *Thought: A Journal of Philosophy* 8 (2): 76–85. <https://doi.org/10.1002/tht3.407>
- Perry, S. 2003. "Harm, History and Counterfactuals." *San Diego Law Review* 40 (4): 1283–1313.
- Perry, S. 2007. "Risk, Harms, Interests, and Rights." In *Risk: Philosophical Perspectives*, edited by T. Lewens, 190–201. 1st ed. New York: Routledge.
- Perry, S. 2014. "Torts, rights, and risk." In *Philosophical Foundations of Tort Law*, edited by J. Oberdiek, 38–64. Oxford: Oxford University Press.
- Peterson, M., and Seidel S. 2021. "The Deontic Transfer Principle." *Erkenntnis* 86 (5): 1185–1195. <https://doi.org/10.1007/s10670-019-00149-8>
- Placani, A. 2016. "When the Risk of Harm Harms." *Law and Philosophy* 36 (1): 77–100. <https://doi.org/10.1007/s10982-016-9277-x>
- Quong, J. 2020. *The Morality of Defensive Force*. New York: Oxford University Press.
- Railton, P. 1986. "Moral realism." *The Philosophical Review* 95 (2): 163–207. <https://doi.org/10.2307/2185589>
- Raven, M. 2012. "In Defence of Ground." *Australasian Journal of Philosophy* 90 (4): 687–701. <https://doi.org/10.1080/00048402.2011.616900>.

- Rosen, G. 2010. "Metaphysical Dependence: Grounding and Reduction." In *Modality: Metaphysics, Logic, and Epistemology*, edited by R. Hale, and A. Hoffman, 109–136. Oxford: Oxford University Press.
- Rowe, T. 2022. "Can a Risk of Harm Itself be a Harm?" *Analysis* 81 (4): 694–701. <https://doi.org/10.1093/analys/anab033>
- Schaffer, J. 2009. *On What Grounds What*. In *Chalmers & Wasserman Metametaphysics: New Essays on the Foundations of Ontology*, 347–383. 1st ed. Oxford: Oxford University Press.
- Sturgeon, N. 1984. "Moral Explanations." In *Morality, Reason, and Truth*, edited by David Copp, and David Zimmerman, 49–78. Totowa, NJ: Rowman and Allanheld.
- Suikkanen, J. 2019. "Ex Ante and Ex Post Contractualism: A Synthesis." *The Journal of Ethics*, <https://doi.org/10.1007/s10892-019-09282-6>.
- Tadros, Victor. 2011. *The Ends of Harm: The Moral Foundations of Criminal Law*. Oxford: Oxford University Press.
- Thomson, J. 1986. *Rights, Restitution, and Risk: Essays in Moral Theory*. 1st ed. Cambridge, MA: Harvard University Press.
- Väyrynen, P. 2013. "I – Grounding and Normative Explanation." *Aristotelian Society Supplementary Volume* 87 (1): 155–178. <https://doi.org/10.1111/j.1467-8349.2013.00224.x>.
- Väyrynen, P. 2018. "Reasons and Moral Principles." In *The Oxford Handbook of Reasons and Normativity*, edited by D. Star, 839–862. 1st ed. Oxford: Oxford University Press.
- Wolff, J., and A. De-Shalit. 2007. *Disadvantage*. Oxford: Oxford University Press.