

ARTICLE

Welfarism and continuity in ethical theory: a formal comparison of prospect utilitarianism vs. sufficientarianism

Susumu Cato¹  and Hun Chung² 

¹Institute of Social Science, University of Tokyo, 7-3-1 Hongo, Bunkyo-ku, Tokyo, Japan 113-0033 and

²Department of Data and Decision Sciences, Emory University, 36 Eagle Row, Atlanta 30322, USA

Corresponding author: Hun Chung; Emails: hun.chung@emory.edu, hunchung1980@gmail.com

(Received 02 August 2023; revised 27 June 2025; accepted 08 July 2025)

Abstract

This paper offers a formal analysis of continuity, welfarism, value satiability, lifeboat cases, along with their interconnectedness with sufficientarianism, with particular attention to the recent defences of sufficientarianism by Ben Davies and Lasse Nielsen in response to Hun Chung's Prospect Utilitarianism (PU). It demonstrates how precise formal definitions help resolve conceptual ambiguities and sharpen philosophical argumentation in distributive ethics. Without such precision, one risks misidentifying or mischaracterizing important normative concepts and theories, leading to confusion or strawman critiques. By highlighting these risks, the paper underscores the methodological importance of precise definitions and formal analysis in ensuring clarity, consistency, and rigor in ethical theorizing.

Keywords: Prospect utilitarianism; sufficientarianism; welfarism; continuity; value satiability; lifeboat cases

1. Introduction: Prospect Utilitarianism and its Recent Critiques

In his 2017 paper, “Prospect Utilitarianism: A Better Alternative to Sufficientarianism”. Hun Chung proposed a theory of distributive justice called ‘Prospect Utilitarianism (PU)’.¹ PU combines the utilitarian social welfare function²

¹More recently, Chung (2023a) has shown how PU can be defended from the original position. See also Chung 2020, 2021, 2022, 2023b) for further discussion of the original position as it pertains to utilitarianism and Rawls.

²Let $N = \{1, \dots, n\}$ be the set of individuals. Let X be the set of social alternatives. For each $i \in N$, let $u_i: X \rightarrow \mathbb{R}$ be individual i 's utility function. Then, for any social alternative $x \in X$, the utilitarian social welfare function $U: X \rightarrow \mathbb{R}$ is defined as: $U(x) = u_1(x) + \dots + u_n(x) = \sum_{i=1}^n u_i(x)$. Hence, according to utilitarian social welfare function, for any social alternatives $x, y \in X$, x is socially preferred to y iff $U(x) > U(y)$ iff $\sum_{i=1}^n u_i(x) > \sum_{i=1}^n u_i(y)$. That is, social alternative x is socially preferred to social alternative y iff the total sum of individual utility generated by x is greater than that generated by y .

© The Author(s), 2025. Published by Cambridge University Press. This is an Open Access article, distributed under the terms of the Creative Commons Attribution-NonCommercial licence (<https://creativecommons.org/licenses/by-nc/4.0/>), which permits non-commercial re-use, distribution, and reproduction in any medium, provided the original article is properly cited. The written permission of Cambridge University Press must be obtained prior to any commercial use.

with individual utility functions characterized in accordance with Kahneman and Tversky's (1979) prospect theory.³ The main selling point of PU is that "it has elements that can appeal to utilitarians, sufficientarians, prioritararians, and egalitarians at the same time"; it is, according to Chung, a "total package" (Chung 2017: 1932). To summarize:

- (1) PU "recommends a distribution that always maximizes the instances of sufficiency" (Chung 2017: 1932), and, thereby, accommodates sufficientarianism's *positive thesis*.⁴
- (2) As a result, PU gives right answers to 'lifeboat' scenarios⁵ (Chung 2017: 1923–25).
- (3) PU "explains why we are not usually morally disturbed by the relative inequality between the rich and super-rich" (Chung 2017: 1923–25) and thereby (partially) accommodates sufficientarianism's *negative thesis*.⁶
- (4) PU always generates *continuous* ethical judgements⁷ (Chung 2017: 1925–27).
- (5) Despite being a version of utilitarianism, PU is "immune to a standard objection to utilitarianism – namely, that it may justify vastly unequal distributions" (Chung 2017: 1932).
- (6) This is so because "[a]fter maximizing sufficiency, prospect utilitarianism prioritizes the individual who is worse-off in terms of his/her welfare level" (Chung 2017: 1927, Proposition 4). In other words, although PU does not 'directly' aim to equalize welfare, because of the specific way that individual utility functions are characterized in PU, PU has a built-in bias towards equality and is, like prioritarianism, "derivatively egalitarian" (Benbaji 2005: 312; Chung 2017: 1929).

³Kahneman and Tversky propose that each individual's utility is: "(i) defined on deviations from the reference point; (ii) generally concave for gains and commonly convex for losses; (iii) steeper for losses than for gains" (Kahneman and Tversky 1979: 279). From this, Chung defines a reference utility function that meets these three characteristics and defines each individual's utility function as a horizontal translation of the reference utility function so that it generates a utility of 0 at their respective critical sufficiency threshold of resource levels (Chung 2017: 1922–23). The reference utility function is designed to measure the extent to which providing a specific benefit to an individual situated at a certain distance from their reference point contributes to the overall value of the entire distribution.

⁴Sufficientarianism's 'positive thesis' claims that it is morally important for people to have enough material resources.

⁵A 'lifeboat situation' is characterized by the fact that some people will necessarily fall below a critical sufficiency threshold no matter how we distribute the remaining resources" (Chung 2017: 1913).

⁶Sufficientarianism's 'negative thesis' claims that once everybody has enough material resources, whether somebody has more or less material resources than others has no moral significance. According to Chung, PU implies a slightly weaker version of the negative thesis called "the weak negative thesis" according to which "As two people have more and more material resources, the ethical significance of their relative inequality matters less and less" (Chung 2017: 1930).

⁷Intuitively, ethical judgements are continuous if "there are no 'jumps' in the ethical preference order. . . . it says that two social states that are almost the same, in terms of the welfare levels of society's members must be viewed as almost ethically indifferent" (Roemer 2004: 272). We may think of continuity as formally expressing Aristotle's principle to treat 'like cases alike and different cases differently'. We will discuss this in more detail in the next section.

As a version of utilitarianism, PU appeals to utilitarians by default. In addition, (1) and (3) make PU appealing to sufficientarians; (2) and (4) are independent moral desiderata that any principle of distributive ethics should ideally satisfy; and (5) and (6) make PU appealing to both egalitarians and prioritarists.

According to Chung, the two central problems of sufficientarianism (including many of its recent non-headcount variants⁸) that make PU a better alternative is that it fails to meet (2) and (4): that is, “[sufficientarianism] cannot provide right answers to lifeboat situations ... [failing (2)] and ... it fails to provide continuous ethical evaluations ... [failing (4)]” (Chung 2017: 1916). In his 2017 paper, Chung uses these two failures of sufficientarianism as a foil to motivate PU as a distributive ethical principle that retains all major attractions of sufficientarianism while avoiding its drawbacks.

In a recent paper titled, “The Prospects for ‘Prospect Utilitarianism’”, Ben Davies (2022) attempts to defend sufficientarianism from these two charges. With respect to desideratum (2), Davies (2022: 336–38) argues that recent non-headcount versions of sufficientarianism can indeed give right answers to lifeboat cases by considering “benefit size”. With respect to desideratum (4), Davies (2022: 339–41) argues that sufficientarianism can satisfy a more general concept of continuity even if it fails to satisfy the narrower (and, hence, more questionable) concept of continuity (which Davies calls “welfarist continuity”) on which Chung’s critique of sufficientarianism purportedly relies. In a related vein, Lasse Nielsen has recently argued that critiquing sufficientarianism on grounds of discontinuity is based on what he calls “the numbers fallacy”. According to Nielsen, this fallacy arises from the utilization of numerical examples featuring “empty numbers” that fail to accurately capture the underlying value framework that may possess satiability and/or range properties (Nielsen 2019: 802–9; 2023: 8).

This paper provides a formal analysis of continuity, welfarism, value satiability, as well as their interconnectedness with sufficientarianism, with particular attention to the recent defences of sufficientarianism by Davies and Nielsen. In the end, this paper underscores the importance of providing precise formal definitions and rigorous treatments of normative concepts. This methodological commitment is especially critical given that certain mathematical concepts – such as continuity – are often invoked in philosophical literature without a precise mathematical definition. Such imprecise usage permits multiple interpretations of continuity to arise, depending on the context or ethical framework, leading to discrepancies in both argument and conclusion. As this paper will demonstrate, without clearly and precisely defining key concepts, one can easily be misled into thinking they are proposing or criticizing a particular normative concept or theory, when in fact they are engaging with a very different one altogether. This likely gives rise to unnecessary conceptual confusion and the risk of strawman arguments. A precise formal approach is therefore essential to ensuring clarity, consistency, and robustness in theoretical debates within distributive ethics.⁹

⁸See, for example, Crisp (2003), Brown (2005), Casal (2007), Huseby (2010), Hirose (2016), Bossert *et al.* (2022).

⁹For a related discussion of how formal models can contribute to normative inquiry in political theory, see Chung and Kogelmann (2024).

2. What is Continuity? General Continuity vs. Welfarist Continuity

Let us first probe the issue of continuity. Intuitively, we say that ethical judgements are continuous if “there are no ‘jumps’ in the ethical preference order. . . . two social states that are almost the same . . . must be viewed as almost ethically indifferent” (Roemer 2004: 272). We may conceptualize continuity as a formalization of Aristotle’s principle that requires us to “treat like cases alike and different cases differently”. Here, what counts as *like cases* or *different cases* depends on the specific context and this is not an issue that continuity addresses. Continuity is not a theory of what is alike or dissimilar; rather, it asserts that once we have established what counts as similar and dissimilar (determined by the relevant topology), the same moral evaluation should apply to any two states that are almost the same or sufficiently similar.

Why is continuity important for a distributive ethical theory? Some scholars regard continuity as having little theoretical or practical significance. We disagree. One key reason continuity is essential for a distributive ethical theory is that we currently lack a precise scale for measuring and comparing individuals’ well-being with pinpoint accuracy. Continuity serves as a robustness requirement in the face of inherently rough and imprecise welfare measurements. Without continuity, ethical judgements could fluctuate unpredictably due to simple measurement error or minor perturbations in the ethical data. If we wish to avoid making vastly inconsistent ethical judgements over nearly identical situations merely due to measurement error or imprecision, then we should endorse continuity and recognize its violation as both a theoretical and practical drawback for any distributive ethical theory.¹⁰

Simply put, continuity requires us to apply the same moral evaluation to any two states that are almost the same or sufficiently similar according to some well-defined notion of similarity or proximity. Here, we must ask: two social states that are almost the same *with respect to what*? The definition of continuity presumes a space (a domain) over which moral evaluations of ethical preferability are made. If the domain under consideration is a space of resources, then continuity claims that there is a continuous relationship between the distribution of resources and ethical preferability. If the domain under consideration is a space of welfare, then continuity claims that there is a continuous relationship between distribution of welfare and ethical preferability.

Many theories of distributive ethics that work under a welfarist framework are either explicitly or implicitly committed to the view that there is a continuous relationship between the distribution of welfare and ethical preferability. (See Kaplow and Shavell 2001; Roemer 2004; Fleurbaey 2015: 207; Chung 2017;

¹⁰An anonymous reviewer raised the important question of whether continuity should be treated as a normative requirement if the concept of well-being is itself essentially ambiguous (cf. Sen 1992: 49). We agree that this kind of conceptual ambiguity gives reason to doubt the *completeness* of our comparative ethical judgements. However, it is less clear that such ambiguity undermines *continuity*. That said, we acknowledge that once completeness is abandoned, the significance of continuity becomes more technically complex. In particular, standard topological definitions of continuity – viz., the open-set and closed-set definitions – are no longer equivalent. These are subtle issues we prefer to bracket and leave open at this stage.

Adler and Holtug 2019: 108; etc.¹¹) As Fleurbaey (2015) notes, “[t]he only sensible ranking that this mild condition [viz., continuity] rules out is the leximin, but [even] the leximin can be treated as the limit of a family of continuous social rankings” (Fleurbaey 2015: 207, footnote 4).

However, endorsing continuity does *not* require one to be a welfarist (which we will soon define and discuss more formally). One might think that factors beyond welfare – such as primary goods, autonomy, etc. – are also ethically relevant and, at the same time, think that our judgements of ethical preferability should vary continuously with the distribution of both welfare and non-welfare considerations. This seems to be the basic intuition that led Davies (2022) to distinguish between “welfarist continuity” from “continuity *per se*” (Davies 2022: 339).

Continuity *per se* is an attractive feature of an ethical view if understood as saying that, where is only a slight difference *across all ethically relevant features* between two outcomes, this should not make a big difference to our ethical preferences. But we can clearly embrace a principle of continuity understood in this way, while rejecting *welfarist* continuity, since two outcomes may differ only slightly in welfare, but differ more significantly on some other ethically relevant factor. (Davies 2022: 340)

... I conclude that welfarist continuity is not axiomatic. (Davies 2022: 341)

Here, Davies is asserting two claims:

Claim 1: Although continuity in general (i.e. continuity *per se*) is an “attractive” feature of an ethical view; *welfarist continuity* is not (i.e. welfarist continuity is not axiomatic).

Claim 2: An ethical view can endorse this more *general notion of continuity* (i.e. continuity *per se*) defined over *all ethically relevant features*, while rejecting the narrower notion of *welfarist* continuity (Davies 2022: 340).

To better understand Davies’s argument, let us try to understand these two claims more precisely and formally. To do so, we would need to distinguish between the distribution of *welfare* and the distribution of *ethically relevant non-welfare features*. Let $N = \{1, \dots, n\}$ be a set of individuals ($n \geq 2$). Let $W = \mathbb{R}^n$ denote a space of welfare distributions and let $X = X_1 \times \dots \times X_m = \mathbb{R}^{nm}$ denote a space of

¹¹Among ethical theorists and axiologists, continuity has usually been discussed in the context of determining the value ordering of risky lotteries in accordance with ‘expected utility theory’. This is called the lottery framework. Intuitively, the continuity axiom under this framework states that for any triple of outcomes x , y , and z , such that x is strictly preferred to y , which, in turn, is strictly preferred to z , there exists some probability $p \in (0, 1)$ that makes the decision-maker indifferent between getting y for sure on the one hand, and playing a gamble (or lottery) that gives x with probability p and z with probability $1 - p$ on the other hand. Temkin (2001) has argued that such a continuity axiom is implausible when x is very good while z is extremely bad [‘continuity of extreme cases’], but this is implied by the more plausible continuity axiom that is restricted to adjacent cases, where the differences in the desirability of three outcomes x , y , and z are small, together with ‘substitutivity of equivalents’ and ‘transitivity’. Temkin concludes that, therefore, there must be something wrong with ‘transitivity’ (see also Temkin 1996). Binmore and Voorhoeve (2003), Arrhenius and Rabinowicz (2005), and Stefansson (2022) point out that Temkin’s (1996, 2001) arguments are technically flawed.

distributions of *all ethically relevant non-welfare features*.¹² For instance, X_1 might be the space of resource distributions, X_2 might be the set of distributions of autonomy, X_3 might be the set of distributions of individual rights, and so on. The specific assignments or interpretations of each X_i as well as the number m of ethically relevant non-welfare features will depend on one's particular ethical theory or theory of the good. Together, $W \times X$ will denote the space of all possible distributions of *all ethically relevant, both welfare and non-welfare, factors*. So, any distribution $(u, x) \in W \times X$ will consist of two components: (i) a distribution of welfare $u \in W$, and (ii) a distribution of ethically relevant non-welfare features $x \in X$ among the n individuals.

Let our ethical preference relation $\succeq$ be defined over all possible distributions of all ethically relevant factors in $W \times X$. For any two distributions of all ethically relevant factors $(u, x), (v, y) \in W \times X$, $(u, x) \succeq (v, y)$ means that distribution (u, x) is at least as morally good as distribution (v, y) . Let $\succ$ and $\sim$ be the asymmetric and symmetric parts of $\succeq$. So, $(u, x) \succ (v, y)$ means that distribution (u, x) is strictly ethically preferred to distribution (v, y) , while $(u, x) \sim (v, y)$ means that the two distributions are equally morally good or ethically indifferent.

Note that this framework is general enough to accommodate both welfarist and non-welfarist ethical theories. For instance, a welfarist theory would disregard the distribution of ethically relevant non-welfare features and base its judgements of ethical preferability solely on the distribution of welfare, whereas a non-welfarist theory would not. Here is the formal definition of welfarism.

Welfarism: *There exists a relation $\succeq^*$ on W such that for any $(u, x), (v, y) \in W \times X$, $u \succeq^* v \Leftrightarrow (u, x) \succeq (v, y)$.*¹³

In words, we say that an ethical view is *welfarist* if there exists an ethical preference relation defined on the space of welfare distributions that completely coincides with a corresponding ethical preference relation defined on the space of all ethically relevant factors such that one welfare distribution u is ethically preferred to another welfare distribution v if and only if the distribution (u, x) of all ethically relevant factors containing the welfare distribution u is ethically preferred to the distribution (v, y) regardless of what the non-welfare features x and y happen to be. Welfarism essentially says that *all non-welfare information is irrelevant*, and the only ethically relevant consideration is welfare. That is, the distribution of welfare completely and solely determines our ethical preferences. A non-welfarist ethical theory will simply deny that this kind of ethical preference relation defined on the space of welfare distributions exists.

By contrast, continuity is a different concept. Let us first define general continuity (or what Davies calls “continuity per se”) over all ethically relevant features.

General Continuity (over all ethically relevant features): *For any $(u, x), (v, y) \in W \times X = \mathbb{R}^n \times \mathbb{R}^{nm}$ and any sequence $\{(v^k, y^k)\}$ such that*

¹² X can be more general. For example, we can get the same results for most cases when X is a partially ordered set.

¹³Our definition is compatible with the definition of welfarism, which has been used by many welfare economists since the work of Sen (1979). A concise argument on this topic can be found in Blackorby *et al.* (2005). Our definition is particularly similar to that of Kaplow and Shavell (2001); see also Fleurbaey *et al.* (2003).

$(v^k, y^k) \rightarrow (v, y)$, if $(u, x) \succ (v, y)$ (respectively, $(v, y) \succ (u, x)$), then there exists a $k^* \in \mathbb{N}$ such that $(u, x) \succ (v^k, y^k)$ (respectively, $(v^k, y^k) \succ (u, x)$) for all $k > k^*$.¹⁴

In words, general continuity requires that if two distributions are almost the same in all ethically relevant features, then the two distributions should be almost ethically indifferent in the sense that if one distribution is ethically preferred to another distribution with respect to all ethically relevant features, then the former distribution should remain ethically preferred to any distribution that is sufficiently close to the latter distribution in all of its ethically relevant features. General continuity formally captures Davies's intuitive idea that "[if there] is only a slight difference across all ethically relevant features between two outcomes, this should not make a big difference to our ethical preferences" (Davies 2022: 340).

General continuity is not a theory about which features of a distribution are ethically relevant. Rather, it asserts that once we have identified the ethically relevant features of a distribution by our prior ethical theory, there should be a continuous relationship between these features and our judgements of ethical preferability. Proposition 1 formally demonstrates that general continuity and welfarism are indeed logically independent, meaning that neither implies the other.

Proposition 1. *General Continuity and Welfarism are independent.*

Proofs of Proposition 1 and all subsequent propositions are relegated to the Appendix. Proposition 1 shows that endorsing general continuity does not require one to be a welfarist, nor does being a welfarist require one to endorse general continuity. The two concepts are independent.

Then, how does this notion of general continuity differ from what Davies calls "welfarist continuity"? Since Davies does not provide a precise formal definition of welfarist continuity, we must infer its formal definition from his informal explanations. Consider the following:

we can clearly embrace a principle of [general] continuity understood in this way, while rejecting *welfarist* continuity, since two outcomes may differ only slightly in welfare, but differ more significantly on some other ethically relevant factor. (Davies 2022: 340)

¹⁴Alternately, General Continuity can be equivalently stated in the following way:

- **Open Set Definition of General Continuity (over all ethically relevant features):**
For any $(v, y) \in W \times X = \mathbb{R}^n \times \mathbb{R}^m$ the sets $\{(u, x) \in W \times X | (u, x) \succ (v, y)\}$ and $\{(u, x) \in W \times X | (v, y) \succ (u, x)\}$ are open.

There is an alternative definition that uses closedness:

- **Closed Set Definition of General Continuity (over all ethically relevant features):**
For any $(v, y) \in W \times X = \mathbb{R}^n \times \mathbb{R}^m$ the sets $\{(u, x) \in W \times X | (u, x) \succeq (v, y)\}$ and $\{(u, x) \in W \times X | (v, y) \succeq (u, x)\}$ are closed.

These definitions are logically equivalent as long as the weak ethical preference relation $\succeq$ is complete. This applies to other definitions in the rest of this paper.

... sufficientarians can plausibly reject welfarist continuity while embracing the less specific principle of continuity I have suggested. For instance, Shields (2018: 44–81) suggests that one candidate for a sufficientarian threshold is autonomy, and that we have especially weighty reasons to secure for people a level of autonomy sufficient to develop and pursue their own view of the good. ... On this view, it is possible to have two identical distributions of welfare, one paternalistically imposed and the other autonomously chosen, and prefer the latter for non-welfarist reasons. (Davies 2022: 340–41)

... If welfarist continuity is an axiom, this seems to rule out by fiat any view that is interested in autonomy except insofar as it impacts welfare. (Davies 2022: 340)

Here, we can see that Davies is asserting an additional claim:

Claim 3: Sufficientarianism is compatible with general continuity, even though it rejects welfarist continuity.

For the record, we state that Claim 3 is false. This is an important issue that we will return in the next section. In the meantime, let us try to understand the notion of welfarist continuity that Davies has in mind in the passage above. The underlying idea seems to be that anyone committed to welfarist continuity must treat any two distributions with identical welfare distributions as ethically indifferent, *regardless* of how much the two distributions differ in their *non-welfare* features. Call this Welfarist Continuity, which may be formally defined as follows:

Welfarist Continuity: For any $(u, x), (v, y) \in W \times X = \mathbb{R}^n \times \mathbb{R}^{nm}$ and any sequence $\{(v^k, w^k)\}_{k \in \mathbb{N}}$ such that $v^k \rightarrow v$, if $(u, x) \succ (v, y)$ (respectively, $(v, y) \succ (u, x)$), then there exists a $k^* \in \mathbb{N}$ such that $(u, x) \succ (v^k, w^k)$ (respectively, $(v^k, w^k) \succ (u, x)$) for all $k > k^*$.¹⁵

In words, Welfarist Continuity requires that if two distributions of all ethically relevant features have nearly identical *welfare distributions*, then the two distributions should be almost ethically indifferent in the sense that if one distribution is ethically preferred to another distribution, then the former distribution should remain ethically preferred to any distribution whose *welfare distribution* is sufficiently close to that of the latter *regardless* of how much the two distributions differ in their distributions of *non-welfare features*.

As already noted, continuity, by itself, is not a theory of what is similar or dissimilar, nor does it specify which features of a given distribution should be regarded as ethically relevant. Rather, it asserts that once we have determined how to measure the distance or proximity between any two distributions and have identified, through our prior ethical theory, which features are ethically relevant, there should be a continuous relationship between those ethically relevant features and our judgements of ethical preferability.

¹⁵We are grateful to an anonymous reviewer for suggesting this formulation of Davies's welfarist continuity.

In this sense, Welfarist Continuity is not a pure conception of continuity under the current framework that also includes ethically relevant non-welfare features. Welfarist Continuity augments the idea of continuity with a particular theory of the good, specifying which features of a distribution are ethically relevant. Then, what particular theory of the good does Welfarist Continuity combine to the notion of continuity? It turns out that Welfarist Continuity *presupposes* Welfarism; more precisely, Welfarist Continuity is a conjunction of General Continuity and Welfarism, which are independent of each other (Proposition 1).

Proposition 2. *An ethical preference relation $\succsim$ satisfies Welfarist Continuity if and only if it satisfies General Continuity and Welfarism.*

Since Welfarist Continuity is the conjunction of General Continuity and Welfarism, rejecting it must stem from rejecting either General Continuity or Welfarism (or both). Davies claims that General Continuity (i.e. continuity extended to all ethically relevant factors) is an “attractive” feature of an ethical theory. If so, then we can say his rejection of Welfarist Continuity is really a rejection of *welfarism*, rather than a rejection of *welfarist continuity*.

When Roemer (2004) described continuity as an “ethically attractive” *axiom* for ethical preferences (Roemer 2004: 272), it is important to understand that he was operating *within* a welfarist framework *only* for the sake of argument.

Welfarism is the view, first, that everything of value about a person’s life can be summed up in a number that measures his or her welfare and that, second, a distributional ethic need only rank possible distributions of welfare, in a population, to be complete. As I here adopt a welfarist framework, I am not concerned with the ‘equality-of-what’ debate, which focuses upon what features of the human condition should be the objects of distributional concern. (Roemer 2004: 267–8)

Roemer makes it clear that he is *not* a welfarist, yet he explains why he, nonetheless, adopts a welfarist framework in the article.

Some may wonder why, in this article, I take a welfarist approach, whereas in other work I have been quite critical of welfarism. The answer is that it is unwise to fight all one’s battles at the same time. While I remain critical of welfarism as a political philosophy, that is not the focus of my concern here. Rather, I focus upon a kind of dogmatism; that it is here illustrated with respect to the welfarist tradition is not particularly important. A similar dogmatism can be found in non-welfarist theory, and I would make a similar critique in that case. As nonwelfarist ethics are generally more complex than welfarist ones, making the critique there would be somewhat more complicated, and pedagogically less transparent, than making it in the welfarist context. (Roemer 2004: 268)

In other words, Roemer’s argumentative strategy was to simply bypass the ‘equality-of-what’ debate – i.e. the question of which features of a distribution are ethically relevant – in order to more transparently illustrate that there may not be a

one-size-fits-all ethical theory that can universally apply to all contexts and situations. Instead, he argues that an “attractive ethical view is contextual, eclectic, and pluralistic” (Roemer 2004: 281). His adoption of a welfarist framework was not an endorsement of welfarism but rather a methodological choice, made purely for the sake of argument and pedagogic simplicity.

However, once we extend our ethical framework to include all ethically relevant non-welfare features – i.e. when we no longer bypass the ‘equality-of-what’ debate and adopt the current framework, which incorporates all ethically relevant, both welfare and non-welfare, considerations (a framework Roemer acknowledges as “more complex”) – Roemer would likely not find Welfarist Continuity ethically attractive. This is because Welfarist Continuity effectively implies that all non-welfare features of a distribution are ethically irrelevant – an implication Roemer, as a non-welfarist, would clearly reject.

As already noted, continuity itself is neutral regarding which features of a given distribution should be considered ethically relevant. Since Welfarist Continuity presupposes the truth of welfarism, to the extent that one has reservations about welfarism, Welfarist Continuity is an inappropriate way to define continuity within our current framework, which includes all ethically relevant, both welfare and non-welfare, considerations.

Suppose one rejects welfarism and accepts that both welfare and non-welfare factors are ethically relevant in evaluating the moral goodness of distributions. Suppose further that one holds that there should be a continuous relationship between welfare distributions and our judgements of ethical preferability. What, then, would be the appropriate conception of continuity for such a non-welfarist? We propose the following:

Restricted Welfare Continuity: *For any $(u, x), (v, y) \in W \times X$ and any sequence $\{(v^k, y)\}$ such that $(v^k, y) \rightarrow (v, y)$, if $(u, x) \succ (v, y)$ (respectively, $(v, y) \succ (u, x)$), then there exists a $k^* \in \mathbb{N}$ such that $(u, x) \succ (v^k, y)$ (respectively, $(v^k, y) \succ (u, x)$) for all $k > k^*$.*

Restricted Welfare Continuity states that, given that welfare is just one among many ethically relevant considerations, if one distribution of all ethically relevant features is ethically preferred to another, then the former should remain ethically preferred to any other distribution that closely approximates the latter in terms of welfare while maintaining *identical* non-welfare features. This conception of continuity does not presume welfarism; rather, it asserts that, once non-welfare features are fixed, our ethical judgements should vary continuously with changes in welfare distributions without encountering abrupt ‘jumps’. The next result, Proposition 3, confirms that Restricted Welfare Continuity and Welfarism are indeed logically independent of each other.

Proposition 3. *Restricted Welfare Continuity and Welfarism are logically independent.*

We believe that Restricted Welfare Continuity provides a plausible definition of continuity in welfare space, particularly once we acknowledge that ethically relevant

non-welfare factors can meaningfully influence our ethical judgements. However, defined in this way, it turns out that general continuity (defined over all ethically relevant features) necessarily implies Restricted Welfare Continuity.

Proposition 4. *General Continuity implies Restricted Welfare Continuity.*

In other words, it is *logically impossible* to endorse General Continuity without also endorsing Restricted Welfare Continuity; or equivalently, it is logically impossible to reject Restricted Welfare Continuity without rejecting General Continuity as well.

Recall that Davies's argument was that it is possible to have two identical distributions of welfare, but ethically prefer one distribution over the other on grounds that one distribution has been autonomously chosen while the other has been paternalistically imposed. This is a valid point. But it is a mistake to think that this violates continuity (in the sense of Restricted Welfare Continuity). Davies claims: "If welfarist continuity is an axiom, this seems to rule out by fiat any view that is interested in autonomy except insofar as it impacts welfare" (Davies 2022: 341). This is incorrect. The claim that autonomy matters only insofar as it affects welfare is a logical implication of *welfarism*, which holds that welfare is the sole ethically relevant consideration. However, it is *not* a logical implication of continuity in welfare space defined by Restricted Welfare Continuity, which merely requires that ethical judgements be continuous over the space of welfare. Recognizing ethically relevant non-welfare factors does not preclude ethical judgements from exhibiting continuity within welfare space.

Going back to Davies's example of autonomy, consider two identical distributions of welfare, where one is generated autonomously, while the other is imposed paternalistically. Since the two distributions are identical in terms of welfare, welfarism implies that the two distributions should be ethically indifferent. By contrast, Restricted Welfare Continuity says that if the autonomously generated distribution is ethically preferred to the paternalistically imposed distribution for non-welfare reasons (despite the fact that the two distributions contain identical distributions of welfare), then the autonomously generated distribution *should remain ethically preferred* to the paternalistically imposed distribution even if the autonomously generated distribution contained *slightly less welfare* or the paternalistically imposed distribution contained *slightly more welfare*. The two concepts are clearly different, but it is not easy to recognize the theoretical difference between the two concepts if one conflates two different questions that should be kept separate. The two questions are:

Question 1: Should welfare be *the only* relevant consideration in our ethical evaluations?

Question 2: Should our ethical evaluations be *continuous*?

Welfarism says 'yes' to the first question. Utilitarians say 'yes' to both questions. However, the answers to the two questions, as we have seen, are independent (Proposition 1); so, a 'yes' or 'no' answer to the first question does not imply either a

‘yes’ or ‘no’ answer to the second question, and vice versa. Davies wants to say ‘yes’ to the second question and ‘no’ to the first question (i.e. there are morally relevant considerations other than welfare). If so, then Davies’s criticism is a criticism against welfarism, *not* continuity. As a matter of fact, we can now see (after defining welfarism and continuity in welfare space precisely) that Davies’s purported counterexamples and criticisms are all targeting welfarism, not continuity.

This is seemingly the root of Davies’s error. What truly troubles Davies is welfarism, a commitment shared by all versions of utilitarianism. At the same time, Davies aims to defend sufficientarianism. However, Chung (2017) did not criticize sufficientarianism on the grounds of welfarism; rather, he challenged it for generating discontinuity. Thus, to defend sufficientarianism against Chung’s critique, Davies needed to show that sufficientarianism does not violate continuity – a task that, as we will soon see in the next section, is essentially impossible. To achieve this, he artificially split the concept of continuity, arguing that he rejects not continuity *per se*, but rather a much narrower version, which he labels “welfarist continuity”, allegedly the foundation of Chung’s criticism. (To be clear, Chung has never distinguished between general continuity and welfarist continuity, and his criticisms against sufficientarianism were based on the fact that it violates continuity *simpliciter*, not a specific form of continuity.) Davies then claims (without formal argument) that although sufficientarianism fails to satisfy this narrow version of (welfarist) continuity (which is unproblematic since welfarist continuity is implausible and non-axiomatic), it can still satisfy a more general principle of continuity defined over all ethically relevant features. (This is Claim 3.) In the next section, we will explain why such a strategy cannot succeed.

Here is a summary of our discussion in section 2. According to Davies, while continuity is generally an attractive axiom for an ethical theory, welfarist continuity is not (Claim 1). Moreover, he argues that an ethical theory can endorse General Continuity – continuity defined over all ethically relevant features – while rejecting the narrower notion of Welfarist Continuity (Claim 2). However, it turns out that Welfarist Continuity is simply the conjunction of General Continuity and Welfarism (Proposition 2), and that these two components are logically independent of one another (Proposition 1). Hence, if anyone – such as Davies – wishes to endorse General Continuity while rejecting Welfarist Continuity, this implies that their objection is, in effect, directed at Welfarism rather than at continuity itself. Indeed, as explained, all of Davies’s purported counterexamples to Welfarist Continuity turn out to be critiques of Welfarism – not of continuity. Moreover, once we move beyond a strictly welfarist framework and allow for the possibility that non-welfare features may also be ethically relevant, Welfarist Continuity (as defined above) becomes an *inappropriate* conception of continuity, as it presupposes the truth of Welfarism – something the expanded ethical framework explicitly denies. In this expanded framework, the appropriate way to understand continuity in welfare space is what we call Restricted Welfare Continuity. This principle holds that, once non-welfare features are fixed, ethical preferability judgements should vary continuously with changes in the welfare distribution. Restricted Welfare Continuity neither presupposes Welfarism nor excludes the relevance of other ethical factors. Once we adopt Restricted Welfare Continuity as our working conception of continuity within this expanded framework, it not only emerges as a plausible axiom (thereby contradicting Claim 1), but is also

logically implied by General Continuity – making it impossible to endorse General Continuity while rejecting Restricted Welfare Continuity (thereby contradicting Claim 2).

3. Sufficientarianism and Continuity

Now, let us move on to assess Claim 3. Recall:

Claim 3: Sufficientarianism is compatible with general continuity, even though it violates welfarist continuity.

As we have seen, continuity restricted to welfare space, when precisely defined over all ethically relevant features, is different from welfarism, and is, therefore, not implausible for reasons that somebody might find welfarism implausible. But even if it were implausible for other valid reasons, Davies’s argumentative strategy would still not work because as long as sufficientarianism presumes the existence of *some* critical sufficiency threshold of *some* ethically relevant feature (be it autonomy, welfare, reasonable contentment, or whatever) over which the ethical value of the distribution suddenly ‘jumps’, it will necessarily be discontinuous in whatever domain that our working notion of continuity (be it (restricted) welfare continuity, general continuity, etc.) is defined over.

To see this, assume that there are two individuals (i.e. $N = \{1, 2\}$) and that welfare and autonomy are our two main ethically relevant considerations. Let $W = \mathbb{R}^2$ be a space of welfare distributions and let $A = \mathbb{R}^2$ be a space of distributions of autonomy. Suppose $\underline{a} > 0$ is the level of autonomy “sufficient to develop and pursue [one’s] own view of the good” (Davies 2022: 340) and suppose our sufficientarian theory claims that we have “especially weighty reasons to secure” each individual an autonomy level of $\underline{a}$. Suppose our ethical preference relation ranks each distribution of autonomy and welfare *lexicographically*: specifically, it first looks at the levels of autonomy the two individuals enjoy and ranks a distribution with more individuals above the autonomy threshold $\underline{a}$ over any distribution with fewer individuals above the autonomy threshold $\underline{a}$; next, when the number of individuals above the critical autonomy threshold $\underline{a}$ are the same, then it applies utilitarianism and prefers the distribution that yields a larger total sum of individual well-being levels in their welfare distributions.

Suppose $u > v$. Then, we have $((u, \underline{a}), (u, \underline{a})) \succ ((v, \underline{a}), (v, \underline{a}))$. Now, consider a sequence $\{((u^k, \underline{a}^k), (u^k, \underline{a}^k))\}$ where $u^k \rightarrow u$ and $\underline{a}^k = \underline{a} - \frac{1}{k}$. Then, we have $\{((u^k, \underline{a}^k), (u^k, \underline{a}^k))\} \rightarrow ((u, \underline{a}), (u, \underline{a}))$. However, we have $((v, \underline{a}), (v, \underline{a})) \succ ((u^k, \underline{a}^k), (u^k, \underline{a}^k))$ for all $k \in \mathbb{N}$. Therefore, our ethical preference relation that incorporates sufficientarian considerations about individual autonomy is *discontinuous* from the perspective of *general* continuity defined over *all ethically relevant considerations*, which, in our working example, are welfare and autonomy. This argument can be easily generalized to any number of ethically relevant considerations (besides welfare and autonomy).

For this purpose, let us consider a generalized version of headcount sufficientarianism. Consider any continuous and monotonic function $h: \mathbb{R}^{1+m} \rightarrow \mathbb{R}$ and some critical level $\theta^* \in \mathbb{R}$ such that for each $i \in N$, $h(u_i^*, x_i^*) = \theta^*$ for some $(u_i^*, x_i^*) \in \mathbb{R}^{1+m}$. Then, according to generalized headcount sufficientarianism, for any $(u, x), (v, y) \in \mathbb{R}^{(1+m)n}$,

$$(u, x) \succeq (v, y) \Leftrightarrow \#\{i \in N | h(u_i, x_i) \geq \theta^*\} \geq \#\{i \in N | h(v_i, y_i) \geq \theta^*\},$$

where (u_i, x_i) is a vector that consists of individual i 's welfare level and the distribution of all ethically relevant non-welfare features. By monotonicity, we mean that $h(u_i, x_i) > h(v_i, y_i)$ if every coordinate in (u_i, x_i) is larger than that in (v_i, y_i) – i.e. if $(u_i, x_i) \gg (v_i, y_i)$ holds. Here, $h(u_i, x_i)$ works as a general aggregation measure for sufficiency. That is, if its value for individual i is not lower than θ^* , then this individual is considered to have met their critical sufficiency threshold and has enough. Our notion of generalized headcount sufficientarianism generalizes sufficientarianism by allowing us to regard sufficientarianism based on autonomy as a special case. From this, we are able to prove the following general result that essentially shows that expanding the number of ethically relevant considerations does not save sufficientarianism from the problem of discontinuity.

Proposition 5. *Generalized headcount sufficientarianism violates General Continuity.*¹⁶

In contrast, utilitarianism satisfies this stronger notion of general continuity.

Proposition 6. *Utilitarianism satisfies General Continuity.*

As Chung's PU is utilitarian, it follows from Proposition 6 that it satisfies general continuity defined over all ethically relevant considerations.

Corollary of Proposition 6. *Prospect Utilitarianism satisfies General Continuity.*

Proposition 5 shows that Davies's continuity claim on sufficientarianism was simply wrong. Continuity is not a concept that can be arbitrarily stretched to be made consistent with sufficientarianism. If one is a sufficientarian, then one cannot avoid discontinuous ethical judgements defined over whatever domain of ethically relevant considerations; discontinuity is simply a cost that sufficientarians must stomach.

4. Value Satiability and Continuity

To this, some sufficientarians may acknowledge the discontinuity inherent in sufficientarianism but dismiss its severity. Recently, Lasse Nielsen (2023) has

¹⁶Some sufficientarian theories satisfy general continuity. For example, Benbaji (2005, 2006) proposes an ordering that does not exhibit 'absolute priority' for people below the threshold. However, this is not the type of sufficientarian theory that Chung and Davies are targeting.

contended that the three standard objections raised against sufficientarianism, namely the problem of indifference,¹⁷ the problem of outweighing priority,¹⁸ and the problem of discontinuity,¹⁹ can all be viewed as specific manifestations of a more broader and general criticism termed “the threshold abruptness objection” (Nielsen 2023: 3–6). According to Nielsen, all instances of the threshold abruptness objection (including the objection of discontinuity) are based on the construction of numerical counterexamples, whose intuitive force draws from “empty numbers” that lack any substantive moral content, and, therefore, falls prey to what he calls “the numbers fallacy” (Nielsen 2019: 812; 2023: 12).

The reason why the numbers used in these numerical examples are, according to Nielsen, “empty” is that they “falsely . . . assume the Archimedean properties of real numbers in their underlying value-theory” (Nielsen 2019: 812). This assumption suggests that the underlying value being represented is *insatiable*, implying that regardless of the existing level of value, it is in principle always possible to introduce an additional unit of the underlying value’s currency to further augment its magnitude. However, Nielsen argues that various fundamental values such as autonomy (Raz 1986), capability (Nussbaum 2000), and reasonable contentment (Huseby 2010), among others, exhibit the capacity for satiation or may exhibit range properties. In the latter case, these values are characterized by a binary nature wherein individuals either possess them or do not, and once obtained, further accumulation is impossible (Nielsen 2023: 8, 15). Consequently, if sufficientarians were to establish their framework of sufficientarianism by incorporating the relevant thresholds based on such satiable foundational values, the apparent implausibility of the threshold abruptness objection could be readily explained away. Nielsen refers to this resultant form of sufficientarianism as “value-satiability sufficientarianism”.

Value-satiability sufficientarianism assumes that the relevant value is (or relevant values are) satiable and that distributive justice is fulfilled if and only if the distribution is such that the [sic] everyone is sated in regard to the relevant value(s). In other words, satiable-value sufficientarianism identifies the threshold as the point above which any person will not become better-off in terms of the relevant value by having more of whatever can be allocated to her. Here, the upper limit is founded on that it is impossible for any human person to be relevantly better-off than this in terms of justice-relevant values. (Nielsen 2019: 806–7)

The thought seems to be that if we accept that many important foundational normative values (such as autonomy, capability, freedom, and so on) can be sated, then the fact that a certain principle of distributive ethics [such as sufficientarianism, whose critical sufficiency thresholds are defined on the basis of such satiable foundational values (Nielsen 2019, 2023)], displays discontinuity may not be that problematic. In other words, value satiability may *explain* why our ethical judgements can plausibly be discontinuous.

¹⁷This is the objection that claims that sufficientarianism “implies indifference about inequality between the superrich and people who are barely above the threshold” (Nielsen 2023: 3).

¹⁸This is the objection that claims that sufficientarianism “implies that we should allow marginal benefits given to people below the threshold to outweigh very large benefits given to people above the threshold” (Nielsen 2023: 4).

¹⁹This is the objection that claims that sufficientarianism “fails to take into consideration the moral significance of ‘being almost there’” such that “any state below the cut-off point will be significantly deficient, and thus there cannot be any (even almost) ethical indifference” (Nielsen 2023: 5).

Davies essentially makes the same point when he distinguishes between ‘ethical attractiveness’ and the ‘demands of justice’, and argues, following Nielsen, that “the values that are relevant to justice are *satiabile*; once the relevant value has been sated, people can become better off, but not in a way that is relevant to justice” (Davies 2022: 341). From this, Davies argues that

What justice demands – and sufficientarianism is a theory of *justice* – may be different from the ‘moral value’ of a total distribution. Thus, even if small differences in welfare necessarily implied only small differences in the ethical preferability of an outcome ... such a relationship may not hold between welfare and the demands of justice. (Davies 2022: 341)

What Davies is saying is that if demands of justice are *satiabile* (because they are based on satiable foundational values), then this implies that demands of justice display *discontinuity*. This is incorrect. Continuity says that small differences in our ethical data should not generate big differences in our ethical evaluations; however, continuity does *not* say that big differences in our ethical data should always generate big differences in our ethical evaluations (which is the inverse statement of continuity, and inverse statements are generally not logically equivalent to their original statements).

For the sake of argument, let us accept Nielsen’s point that many foundational values pertinent to justice, or relevant to defining sufficientarian thresholds, exhibit the property of satiation. Would this make our evaluations of justice *discontinuous*? Not necessarily. Rather, value-foundational satiability merely implies that our evaluations of justice may not be *strictly monotonic* across the whole domain.

Let us illustrate this with a simple example. Suppose that a person’s autonomy is satiable: that is, autonomy increases with a person’s resource levels up to a point after which the person’s autonomy is completely fulfilled and sated. Suppose that individual i ’s autonomy function α_i , which measures the degree of autonomy individual i enjoys, is defined as follows:

$$\alpha_i(x) = \begin{cases} x & \text{for } x < \underline{x} \\ \underline{x} & \text{for } x \geq \underline{x} \end{cases}$$

where $x \in \mathbb{R}$ is the amount of resources i has. So, individual i ’s autonomy increases linearly in the amount of resources up to $\underline{x}$ after which i ’s autonomy is fully sated and remains constant. Nevertheless, the autonomy function α_i that measures individual i ’s degree of autonomy is *continuous* in resources.²⁰ Contrast this with another autonomy function β_i :

$$\beta_i(x) = \begin{cases} x & \text{for } x < \underline{x} \\ \underline{x} + 100 & \text{for } x \geq \underline{x} \end{cases}$$

²⁰**Proof.** Since both a linear function and a constant function are continuous, all we need to show is that $\alpha_i(x)$ is continuous at $\underline{x}$. Pick any $\varepsilon > 0$, and let $\delta = \varepsilon$. Then, for all $x \in \mathbb{R}$ such that $|x - \underline{x}| < \delta$, we have $|\alpha_i(x) - \alpha_i(\underline{x})| < \varepsilon$. ■

Here, individual i 's autonomy increases linearly by the amount of resources that i receives up to *right before* $\underline{x}$; but *at* $\underline{x}$, individual i 's autonomy *abruptly jumps* to $\underline{x} + 100$ and stays constant thereafter. Clearly, the autonomy function β_i is *discontinuous* at $\underline{x}$.

Note that both autonomy functions α_i and β_i are consistent with Nielsen's value-foundational satiability. Specifically, both autonomy functions "identif[y] the threshold [namely, $\underline{x}$] as the point above which any person will not become better-off in terms of the relevant value [in this case, autonomy] by having more of whatever can be allocated to her [in this case, resources]" (Nielsen 2019: 806–7). But, unlike β_i , which is *discontinuous* at the critical sufficiency threshold $\underline{x}$, α_i is *continuous* at the critical sufficiency threshold $\underline{x}$.

What these examples show is that endorsing value-foundational satiability does not necessarily commit us to endorse discontinuity. We used autonomy as an example, but the same thing can be said to any other satiable foundation value such as freedom, capability, reasonable contentment, and so on. Hence, even if one successfully builds sufficientarianism or any theory of justice on the basis of value-foundational satiability as Nielsen proposes, such fact does not explain away why the demands of justice may be discontinuous. As a consequence, it is a mistake to dismiss the problem of discontinuity raised against sufficientarianism merely as an instance of *the numbers fallacy*.

Then, the question remains. If discontinuity is indeed a *cost*, and if PU can avoid the cost while retaining all the main attractions of sufficientarianism, why not accept PU instead of sufficientarianism?

5. Sufficientarianism and Lifeboat Cases

Let us now move on to lifeboat cases, which were originally introduced by Frankfurt (1987: 30).²¹ In a typical lifeboat case, we are faced with a choice of two options:

- **Some Survive:** Save some people by allocating resources unequally, or
- **All Die:** Allocate resources equally and let everyone die.

(Telic) Egalitarianism claims that (undeserved) inequality is bad in itself. Prioritarianism claims that benefiting people matters more the worse off these people are (Parfit 1997: 213). Both egalitarianism and prioritarianism imply an equal distribution of resources. In so far as Crisp's (2003) and Huseby's (2010) recent non-headcount versions of sufficientarianism endorse prioritarianism below the critical sufficiency threshold,²² they also imply an equal distribution of resources

²¹See Chung (2016) for a critical discussion of Frankfurt's seminal work on sufficientarianism.

²²According to Crisp, "absolute priority is to be given to benefits to those below the threshold at which compassion enters. Below the threshold, benefitting people matters more the worse off those people are, the more of those people there are, and the greater the size of the benefit in question. Above the threshold, or in cases concerning only trivial benefits below the threshold, no priority is to be given" (Crisp 2003: 758). Similarly, Huseby claims that "First, individuals below the maximal sufficiency threshold should have absolute priority over individuals above this threshold. . . . Between the minimal and maximal sufficiency thresholds, I propose that we should apply a constrained and inverse form of prioritarianism. . . . Second, strong priority should be given to those below the minimal sufficiency threshold. By strong priority, I intend

below the critical sufficiency threshold.²³ Since everybody in a lifeboat scenario is clearly below the critical sufficiency threshold (however this may be defined), egalitarianism, prioritarianism, and Crisp's and Huseby's non-headcount sufficientarianisms all imply that we should choose *All Die* instead of *Some Survive*. Insofar as we agree that choosing *Some Survive* is the "right answer",²⁴ lifeboat cases pose a problem for all these views. This is the gist of Chung's (2017) original argument.

If that was too brief, here is a more formal argument that demonstrates why telic egalitarianism, prioritarianism, and recent non-headcount versions of sufficientarianism endorse *All Die* in lifeboat situations.

Firstly, telic egalitarianism claims that (undeserved) inequality is bad in itself, and, hence, reducing inequality will always be in at least one respect good. This commits all telic egalitarians to say that *All Die* is, at least in one respect, morally better than *Some Survive*. A *strong* telic egalitarian would further argue that everybody dying equally (i.e. *All Die*) would be strictly morally preferable to some people (who are no more deserving than others) surviving (i.e. *Some Survive*).

Secondly, prioritarianism can be formally characterized as a distributive ethical view that maximizes $\sum_{i \in N} f(u_i)$: that is, for any two welfare distributions $u, v \in W$, $u \succ v$ according to prioritarianism if and only if $\sum_{i \in N} f(u_i) > \sum_{i \in N} f(v_i)$, where f is a strictly increasing, strictly concave function. Note that a function f is strictly concave (i.e. has a decreasing slope) if and only if for any $x, y \in \mathbb{R}$ and any $\tau \in (0, 1)$,

$$f(\tau x + (1 - \tau)y) > \tau f(x) + (1 - \tau)f(y).$$

Because of the strict concavity of f , prioritarianism is committed to the following principle:

- **The Pigou–Dalton Equalization Principle:** For any $u, v \in W$, if $u_i > v_i = v_j > u_j$, $u_i + u_j = v_i + v_j$, and $u_k = v_k$ for all $k \neq i, j$, then $v \succ u$.

To see this, pick any $u, v \in W$, such that $u_i > v_i = v_j > u_j$, $u_i + u_j = v_i + v_j$, and $u_k = v_k$ for all $k \neq i, j$. Then, strict concavity of f implies that for any $u, v \in W$ and any $\tau \in (0, 1)$,

$$f(\tau u_i + (1 - \tau)u_j) > \tau f(u_i) + (1 - \tau)f(u_j).$$

In particular, the inequality holds when $\tau = \frac{1}{2}$. Substituting $\tau = \frac{1}{2}$ into the inequality, we obtain:

something less than absolute priority, but something more than straightforward weighted aggregation" (Huseby 2010: 184–185).

²³See also Brown (2005), Casal (2007), Hirose (2016) and Bossert *et al.* (2022, 2023) for related accounts of sufficientarianism.

²⁴We understand that some sufficientarians would not even consider that *Some Survive* is the right answer to lifeboat situations (indeed, one of the authors of this paper, Susumu Cato, takes this position). For example, if there are two individuals and 10 units of resources, these sufficientarians may think that (5, 5) is ethically acceptable or even ethically preferable even when the sufficiency threshold is 10. However, this is not Davies's stance: unlike these sufficientarians, Davies agrees that *Some Survive* is the right answer to lifeboat situations, and, instead, tries to argue that sufficientarianism can prescribe *Some Survive* in lifeboat situations by considering benefit size.

$$f(v_i) = f(v_j) = f\left(\frac{u_i + u_j}{2}\right) > \frac{f(u_i)}{2} + \frac{f(u_j)}{2}.$$

This implies that $f(v_i) + f(v_j) > f(u_i) + f(u_j)$ and $f(v_k) = f(u_k)$ for all $k \neq i, j$. As a result, we have $\sum_{i \in N} f(v_i) > \sum_{i \in N} f(u_i)$, and, hence, $v \succ u$ according to prioritarianism, which is precisely what the Pigou–Dalton Equalization Principle requires. (Note that this holds independently of the position of the critical sufficiency threshold θ .)

The Pigou–Dalton Equalization Principle essentially says that it is always morally preferable to distribute the same total sum of resources/welfare equally among the individuals. Hence, if the total amount K of distributable resources/welfare is less than the total amount of resources needed to meet everybody's critical sufficiency threshold θ (i.e. $K < n\theta$), prioritarianism requires the resources to be allocated equally, which results in each individual receiving $\frac{K}{n}$. Note that $\frac{K}{n} < \theta$, i.e. the amount allocated to each individual is below their critical sufficiency threshold θ . This shows that prioritarianism will prescribe *All Die* instead of *Some Survive* in the lifeboat scenario.

Thirdly, in so far as Crisp's (2003) and Huseby's (2010) recent non-headcount versions of sufficientarianism endorse prioritarianism *below* the critical sufficiency threshold, they also imply an equal distribution of resources *below* the critical sufficiency threshold θ by the same argument. So, the crucial point is whether these recent non-headcount versions of sufficientarianism are committed to the following:²⁵

- **The Pigou–Dalton Equalization Principle (Across the Threshold θ):** For any $u, v \in W$, if $u_i \geq \theta > v_i = v_j > u_j$, $u_i + u_j = v_i + v_j$, and $u_k = v_k$ for all $k \neq i, j$, then $v \succ u$.

The Pigou–Dalton Equalization Principle (Across the Threshold θ) states that it is morally better to distribute resources/welfare equally even when this results in pulling someone below their critical sufficiency threshold θ . As long as one believes that benefitting those below the critical sufficiency threshold must be *absolutely prioritized* (as suggested by Crisp), recent non-headcount versions of sufficientarianism have no theoretical resources to reject the Pigou–Dalton Equalization Principle (Across the Threshold θ).

We will later consider how a modified version of non-headcount sufficientarianism (which we call *Sufficientarianism S*) can address this problem. For now, we conclude that telic egalitarianism, prioritarianism, and (Crisp's and Huseby's) recent non-headcount sufficientarianism all imply that we should choose *All Die* instead of *Some Survive* in the above lifeboat scenario.

Note that the welfare levels of those who die in *All Die* are higher than those who die in *Some Survive*. Therefore, we can say that facing death in *All Die* is better than facing death in *Some Survive*. However, this does not change the fact that everybody eventually dies in *All Die* and that this is what all three views (i.e. telic egalitarianism, prioritarianism, non-headcount sufficientarianism) prescribe. So, again, insofar as we agree that choosing *Some Survive* is the “right answer”, lifeboat cases pose a problem for all these views.

²⁵This principle is formally proposed and examined by Bossert *et al.* (2022, 2023). In particular, Bossert *et al.* (2002) show how this principle is satisfied by what they call critical-level sufficientarian orderings.

Of course, Chung makes it clear that

This is not to deny that there are ways for both egalitarians and prioritarians to [choose *Some Survive*] in such scenario. For instance, as long as the egalitarian and the prioritarian do not give *absolute* priority to the worst off, it could be possible for them to recommend [*Some Survive*] by allowing inequalitarian or anti-prioritarian distributions in exceptional cases in which sufficiently large gains for a sufficient number of people who aren't worst-off will outweigh a gain to the worst-off. However, the point is that both egalitarianism and prioritarianism will be able to give right answers to lifeboat situations only allowing exceptions, not as a matter of principle. Note that the same criticism would apply to both Crisp's and Huseby's versions of sufficientarianism, as both versions endorse prioritarianism below the critical sufficiency threshold. (Chung 2017: 1914)

Davies argues that, unlike Chung's accusations, the two recent *non-headcount* versions of sufficientarianism by Crisp (2003) and Huseby (2010) are able to provide right answers to lifeboat cases by considering *benefit size* and that "the idea of allowing benefit size to outweigh the importance of benefits going to the worst off is not an *ad hoc* exception, but rather a principle that applies in all cases" (Davies 2022: 337).

We would first like to note that there are several places where both Crisp and Huseby appear to acknowledge that their respective non-headcount versions of sufficientarianism do not select *Some Survive* in lifeboat cases – even when the size of the benefit is taken into account. For instance, Crisp explains that "[o]ne possible problem with [his] view is . . . that the view will prefer the smallest nontrivial benefit to any number of individuals below the threshold to any benefit, no matter how large, to any number of individuals above the threshold" (Crisp 2003: 758). Similarly, Huseby explains that a potential problem with his view is that it may allow "giving a small benefit to a person who is below the sufficiency threshold" at the expense of giving "a much larger (or extremely much larger) benefit to many (or extremely many) people above the sufficiency threshold" (Huseby 2010: 186). Huseby calls this broader issue – of which lifeboat failures are a special case – "the problem of waste" (Huseby 2010: 186). (This is essentially what Nielsen called the problem of outweighing priority.)

Unlike what Davies thinks, neither Crisp nor Huseby thought that considerations of benefit size will help their non-headcount versions of sufficientarianism choose *Some Survive* in lifeboat cases. Instead, Crisp simply bites the bullet and argues that the problem of waste, "may not be as implausible as it seems once we give proper recognition to the fact that the threshold is the point at which compassion no longer applies" (Crisp 2003: 758). Similarly, instead of saying that there is a way to overcome the problem of waste by considering benefit size, Huseby tries to dilute the severity of the objection by pointing out that other well-known distributive principles all face the same problem: "Egalitarianism, absolute prioritarianism, and the difference principle all demand waste also in situations *where everyone is insufficiently well off*" (Huseby 2010: 187 emphasis his). His response is not that

there is a way to solve the problem by considering benefit size; rather, he finds such a problem a *cost* of sufficientarianism that he is willing to bear:

In my view, sufficiency is, all things considered, a more defensible principle than its main alternatives. The principle is grounded in the concern for the badly off. This concern has as its cost the problem of waste. I find this cost acceptable. (Huseby 2010: 187)

In sum, both Crisp and Huseby acknowledged that their non-headcount versions of sufficientarianism may fail to provide right answers to lifeboat cases and they also did not think that considerations of benefit size can be used to overcome such shortcomings.

But maybe both Crisp and Huseby were not perfectly aware of the true potential of their theoretical framework in coping with lifeboat cases. Hence, let us be maximally charitable and consider what formal conditions are required to construct a version of (non-headcount) sufficientarianism that delivers the correct verdicts in lifeboat cases – specifically in the way Davies envisions, by taking *benefit size* into account.

The key is to define sufficientarianism's *social welfare function* in a way that accounts for the (huge) benefit that accrues to somebody meeting their sufficiency threshold. We may do this by assuming two critical values of transferable resources (or utility): (a) a critical sufficiency threshold $\theta \geq 0$, and (b) a critical level $\alpha \geq \theta$. Now, take any continuous, concave, and strictly increasing function $g: \mathbb{R} \rightarrow \mathbb{R}$ that is not bounded from above and define the sufficientarian social welfare function (SWF) as follows:

$$S(x) = \sum_{i \in N: x_i < \theta} [g(x_i) - g(\alpha)],$$

where $x = (x_1, \dots, x_n)$ denotes any distribution and x_i denotes the amount of resources that individual i receives in distribution x .²⁶ Then, for any two distributions $x = (x_1, \dots, x_n)$ and $y = (y_1, \dots, y_n)$, our sufficientarian theory will strictly ethically prefer x to y if and only if:

$$S(x) > S(y) \Leftrightarrow \sum_{i \in N: x_i < \theta} [g(x_i) - g(\alpha)] > \sum_{i \in N: y_i < \theta} [g(y_i) - g(\alpha)].$$

Since g is concave, our sufficientarian SWF has prioritarian tendencies when everybody is below the sufficiency threshold and is unable to reach it. Let us now decompose our sufficientarian SWF into two components:

²⁶This is a simplified version of the sufficientarian SWF introduced by Bossert *et al.* (2023) under the name of generalized critical-level sufficientarianism. In a variable-population setting, these authors consider a “Paretian” sufficientarianism, where a critical level is different from the threshold. By contrast, the sufficientarian SWF proposed in this paper does *not* satisfy the Pareto principle because all changes in utilities/well-beings above the threshold do not increase the sufficientarian SWF and are not taken into account. This implies that our sufficientarian SWF $S(x)$ is congruent with Nielsen's value satiability assumption and satisfies sufficientarianism's negative thesis.

$$S(x) = \underbrace{\sum_{i \in N: x_i < \theta} [g(x_i) - g(\theta)]}_{\text{First Component}} + \underbrace{\sum_{i \in N: x_i < \theta} [g(\theta) - g(\alpha)]}_{\text{Second Component}}$$

If we use $m < |N| = n$ to denote the number of those below the critical sufficiency threshold θ [i.e. $m = \#\{i \in N | x_i < \theta\}$], this can be written as follows:

$$S(x) = \underbrace{\sum_{i \in N: x_i < \theta} [g(x_i) - g(\theta)]}_{\text{Shortfall Component}} + \underbrace{m[g(\theta) - g(\alpha)]}_{\text{Headcount Component}}.$$

Then, how should we interpret these two terms? Note that the first term (i.e. the Shortfall Component) is negative unless the set of those below the critical sufficiency threshold θ is empty, while the second term (i.e. the Headcount Component) is negative unless the critical level α coincides with the critical sufficiency threshold θ or m is zero. Together, the value of our sufficientarian SWF $S(x)$ cannot be positive.

Note that for each individual $i \in N$, $[g(x_i) - g(\theta)]$ represents the disvalue of i falling short of meeting their sufficiency threshold θ . The disvalue of this gap $[g(x_i) - g(\theta)]$ decreases as x_i approaches the critical sufficiency threshold θ and becomes zero when x_i reaches or exceeds θ . The first term of our sufficientarian SWF sums the disvalues of such a gap across all individuals who are below the sufficiency threshold θ : i.e. $\sum_{i \in N: x_i < \theta} [g(x_i) - g(\theta)]$. Therefore, we can say that the first term (i.e. the Shortfall Component) of our sufficientarian SWF represents the overall negative impact on society generated by the cumulative shortfall of individuals who are below the critical sufficiency threshold θ in a given distribution. As noted above, we can see that our sufficientarian SWF has a prioritarian tendency when everybody is below their sufficiency threshold. This is because, unless the second term changes, the overall disvalue of the sufficientarian SWF can most effectively be reduced by distributing any available resources or utility to the worse off.

Then, what about the second term, $m[g(\theta) - g(\alpha)]$, the Headcount Component? When somebody who was previously below the sufficiency threshold θ successfully reaches or surpasses it, the total number of individuals who are below the sufficiency threshold θ decreases from m to $m - 1$. As a result, the value of the sufficientarian SWF increases by an increment of $[g(\alpha) - g(\theta)]$ – the difference between the value assigned to the critical level α and that assigned to the critical sufficiency threshold θ – for each additional person who reaches their critical sufficiency threshold θ . We might think of $[g(\alpha) - g(\theta)]$ as representing the *size* of the *moral benefit* of letting somebody reach their critical sufficiency threshold. Or, to put differently, we might think of $[g(\theta) - g(\alpha)]$ as representing the negative moral value of somebody being below the critical sufficiency threshold θ , and $m[g(\theta) - g(\alpha)]$ as representing the cumulative moral disvalue of having m individuals below the critical sufficiency threshold θ .

For instance, suppose $x_i < \theta$, that is, individual i is below the sufficiency threshold, and suppose we move individual i from x_i to x'_i (where $x_i < x'_i$). If $x_i < x'_i < \theta$, then the increased moral value of the distribution becomes: $g(x'_i) - g(x_i)$. However, if $x_i < \theta < x'_i$, the increased moral value of the distribution becomes: $[g(x'_i) - g(x_i)] + [g(\alpha) - g(\theta)]$.

That is, there is an additional moral value that gets added to the distribution by allowing somebody to not merely reduce their gap while still falling short of the critical sufficiency threshold θ , but allowing that person to successfully reach or surpass the critical sufficiency threshold θ .

We can see that our sufficientarian SWF $S(x)$ is designed to find a distribution that minimizes the negative moral disvalue of the distribution. This is in accordance with Nielsen's (2023) recent proposal:

This invites the idea that sufficientarianism should be interpreted negatively in the sense of focusing on elimination of deficiency rather than on securing enough of some given currency. Thus, we arrive at the following generic principle:

(S) A distribution is just if, and only if, no one suffers deficiencies from a justice-relevant threshold (Nielsen 2023: 17).

Our sufficientarian SWF $S(x)$ is also in line with value-satiability sufficientarianism as it assumes that the moral value that accrues to the distribution by giving more to a given individual is sated once that individual reaches their sufficiency threshold θ . Following Nielsen, let us call our formalization of sufficientarianism, *sufficientarianism S*: *Sufficientarianism S* generates a moral order over the set of distributions in accordance with the sufficientarian SWF $S(x)$. Would *sufficientarianism S* be able to now choose *Some Survive* to lifeboat situations?

Proposition 7. *In any lifeboat situation, Sufficientarianism S always maximizes the number of instances meeting their critical sufficiency threshold θ as long as the critical level α is sufficiently large.*

Proposition 7 implies that *sufficientarianism S* can provide right answers to lifeboat situations by recommending *Some Survive* instead of *All Die* as long as it assumes that the 'size' of the moral benefit (i.e. $[g(\alpha) - g(\theta)]$) that accrues to the distribution *in addition to* the gain in individual welfare of making somebody reach their critical sufficiency threshold θ is sufficiently large.²⁷

Defined in this way, we can see that there is a sense in which *sufficientarianism S* begs the very question it tries to answer: specifically, it *presupposes* that the ethical significance of ensuring individuals reach their critical sufficiency threshold θ is substantial enough to demonstrate that it will endorse a distribution that maximizes the number of individuals attaining the said threshold. In this context, the determination of what qualifies as assigning an appropriately high moral significance to meeting the critical sufficiency threshold θ is defined

²⁷Interestingly, the pure headcount approach can be seen as a limit of this sufficientarian theory. More precisely, sufficientarianism *S* approaches to headcount sufficientarianism as $[g(\alpha) - g(\theta)]$ approaches to infinity. It is noteworthy that our sufficientarianism can be regarded as a hybrid theory because the first term corresponds to prioritarianism below the threshold, and the second term essentially corresponds to the headcount approach to sufficientarianism. In this sense, $[g(\alpha) - g(\theta)]$ is a "weight" on headcount sufficientarianism. This offers an intuitive reason why headcount sufficientarianism can be obtained as a limit.

endogenously by its capacity to yield morally correct resolutions in lifeboat scenarios. That is, *sufficientarianism S* is being adjusted for the sole purpose of generating a predetermined correct answer – *Some Survive* – in lifeboat scenarios.

Note that in our *sufficientarianism S*, the size of the moral benefit associated with reaching the critical sufficiency threshold θ is given by the difference $g(\alpha) - g(\theta)$, where $g(\alpha)$ is the value of the critical level α and $g(\theta)$ is the value of the critical sufficiency threshold θ . To give right answers to lifeboat scenarios, the critical level α must be sufficiently greater than the critical sufficiency threshold θ . However, unless the critical level α is identical to the critical sufficiency threshold θ , the moral ordering generated by *sufficientarianism S* will necessarily display *discontinuity* at the critical sufficiency threshold θ .

Proposition 8. *The moral ordering of Sufficientarianism S is discontinuous at the critical sufficiency threshold θ whenever the critical level α exceeds the critical sufficiency threshold θ , i.e. whenever $\alpha > \theta$.*

Hence, even if we are able to formulate sufficientarianism (negatively) to provide right answers to lifeboat situations by considering “benefit size” as Davies proposes, we can only do so by completely giving up continuity.²⁸

It is worth noting that there is an important distinction between the *negative* formulation of sufficientarianism – as requiring the *elimination of deficiency* – and the *positive* formulation – as requiring the *maximization of sufficiency*. These two formulations tend to align with different versions of sufficientarianism: the negative formulation aligns more naturally with non-headcount sufficientarianism, while the positive formulation aligns with headcount sufficientarianism. This distinction yields divergent distributive implications. For example, if the critical sufficiency threshold is set at 10, a positive formulation of sufficientarianism that emphasizes the maximization of sufficiency (i.e. headcount sufficientarianism) will favour the distribution (10, 2, 1) over the distribution (9, 8, 8), since more people meet the threshold. In contrast, the negative formulation of sufficientarianism that emphasizes the elimination of deficiency (i.e. non-headcount sufficientarianism) may regard (9, 8, 8) as preferable, because it yields fewer severe shortfalls, even

²⁸One potential way for sufficientarians to ensure that their theory is mathematically consistent with continuity is to measure all sufficientarian values (e.g. autonomy, resources, etc.) on a finite discrete scale. The core idea is that each ethically relevant feature would be assessed using a finite set of values – such as $\{0, 1\}$, where 0 represents “insufficient” and 1 represents “sufficient” for that particular feature. The resulting sufficientarian social welfare function would also take on a finite set of values – for instance, $\{0, 1, \dots, n\}$, where each numeric value indicates the number of individuals who meet the sufficiency threshold. Since it can be mathematically demonstrated that any function from a finite set to another finite set is continuous, this approach would render sufficientarianism continuous, at least in a formal mathematical sense. However, adopting this strategy would prevent sufficientarians from making meaningful distinctions between individuals who are closer to or further from the sufficiency threshold. In other words, it would eliminate the ability to account for degrees of insufficiency. As a result, this approach is only viable for *pure headcount sufficientarians*, who are concerned solely with the number of individuals who meet the threshold. For any *non-headcount version of sufficientarianism* – where the ethical assessment depends not just on whether people meet the threshold but also on how far below it they fall – discontinuity remains unavoidable. Therefore, it remains true that, contrary to Davies’s claim, any non-headcount version of sufficientarianism remain discontinuous.

though no one reaches the threshold.²⁹ Such divergent distributive implications arise even in the class of variable-population orderings. For instance, consider a comparison between the populations (10, 10, 5, 5) and (10, 5), assuming again a sufficiency threshold of 10. The *positive* formulation of sufficientarianism (i.e. headcount sufficientarianism) would judge the former to be better, as it includes more individuals who meet the sufficiency threshold. In contrast, the *negative* formulation of sufficientarianism (i.e. non-headcount sufficientarianism) would prefer the latter, as it involves a smaller total shortfall from sufficiency – that is, less cumulative deficiency across the population.

Such a distinction is formally reflected in our sufficientarian social welfare function $S(x)$, which incorporates two sources of moral disvalues: (a) the moral disvalue of failing to meet the critical sufficiency threshold, and (b) the moral disvalue of being farther away from the critical sufficiency threshold when below it. Accordingly, an implication of our Proposition 7 is that whenever the critical level α is *insufficiently* high, sufficientarianism S , represented by the sufficientarian SWF $S(x)$, may prefer (9, 8, 8) to (10, 2, 1) – despite the fact that no one meets the sufficiency threshold in (9, 8, 8). This indeed is the main reason why non-headcount sufficientarianism – especially in forms that incorporate prioritarian weighting below the threshold – generally fail to generate correct answers to lifeboat situations. Propositions 7 and 8 show that one can reformulate non-headcount sufficientarianism in a clever way to avoid this problem – but only at the cost of giving up continuity.

Thus, the question remains: if offering appropriate solutions to lifeboat situations is indeed a desired outcome and if generating discontinuous ethical judgements is indeed a cost, and if PU can circumvent such costs while preserving the key merits of sufficientarianism, why not embrace PU instead of sufficientarianism?

6. Concluding Remarks

In this paper, we have argued that the respective defences of sufficientarianism by both Davies and Nielsen are not entirely successful. When defined precisely, continuity emerges as a much more plausible normative requirement than is often assumed. While a discontinuous ethical theory might still be defensible all things considered, this does not eliminate the fact that discontinuity constitutes a theoretical cost. Crucially, continuity and welfarism are distinct concepts and should not be conflated. Confusing the two risks misattributing reasons for rejecting one to the other.

Moreover, any plausible principle of distributive ethics should ideally be capable of delivering the correct verdicts in lifeboat scenarios. We have shown that sufficientarianism can be formally adapted to yield the correct verdicts in such scenarios by introducing two critical values: (a) a critical sufficiency threshold θ , and (b) a critical level α that lies sufficiently above this threshold. However, this adaptation necessarily renders sufficientarianism discontinuous at the critical sufficiency threshold. None of these theoretical insights would have been possible without a careful and precise formal analysis.

²⁹We thank an anonymous reviewer for making this point.

So, we will ask one last time: if the violation of continuity is indeed a *cost*, and if PU can avoid the cost while preserving all the principal attractions of sufficientarianism – including providing right answers to lifeboat situations – why not accept PU instead of sufficientarianism?

Acknowledgements. We are grateful to Roger Crisp and Ben Davies for their helpful comments and constructive criticisms on an earlier version of this paper.

Competing Interests Declaration. The authors have no competing interests to declare.

References

- Adler M.D. 2018. Prioritarianism: room for desert? *Utilitas* **30**, 172–197.
- Adler, M.D. 2019. *Measuring Social Welfare: An Introduction*. Oxford: Oxford University Press.
- Adler M.D. and N. Holtug 2019. Prioritarianism: a response to critics. *Politics, Philosophy & Economics* **18**, 101–144.
- Arrhenius G. and W. Rabinowicz 2005. Millian superiorities. *Utilitas* **17**, 127–146.
- Benbaji Y. 2005. The doctrine of sufficiency: a defence. *Utilitas* **17**, 310–332.
- Benbaji Y. 2006. Sufficiency or priority? *European Journal of Philosophy* **14**, 327–348.
- Binnmore K. and A. Voorhoeve. 2003. Defending transitivity against Zeno's paradox. *Philosophy & Public Affairs* **3**, 272–279.
- Blackorby C., W. Bossert, and D. Donaldson 2005. *Population Issues in Social Choice Theory, Welfare Economics, and Ethics*. Cambridge: Cambridge University Press.
- Bossert W., S. Cato and K. Kamaga 2022. Critical-level sufficientarianism. *Journal of Political Philosophy* **30**, 434–461.
- Bossert W., S. Cato and K. Kamaga 2023. Thresholds, critical levels, and generalized sufficientarian principles. *Economic Theory* **75**, 1099–1139.
- Brown C. 2005. Priority or sufficiency... or both? *Economics & Philosophy* **21**, 199–220.
- Casal P. 2007. Why sufficiency is not enough. *Ethics* **117**, 296–326.
- Chung H. 2016. Is Harry Frankfurt's 'Doctrine of Sufficiency' sufficient? *Organon F* **23**, 50–71.
- Chung H. 2017. Prospect utilitarianism: a better alternative to sufficientarianism. *Philosophical Studies* **174**, 1911–1933.
- Chung H. 2020. Rawls's self-defeat, a formal analysis. *Erkenntnis* **85**, 1169–1197.
- Chung H. 2021. On choosing the difference principle behind the veil of ignorance: a reply to Gustafsson. *Journal of Philosophy* **118**, 450–462.
- Chung H. 2022. Chain-connection, close-knitness, and the difference principle. *Journal of Politics* **84**, 2305–2311.
- Chung H. 2023a. Prospect utilitarianism and the original position. *Journal of the American Philosophical Association* **9**, 670–704.
- Chung H. 2023b. When utilitarianism dominates justice as fairness: an economic defense of utilitarianism from the original position. *Economics and Philosophy* **39**, 308–333.
- Chung H. and B. Kogelmann 2024. Formal Models in Normative Political Theory. *Journal of Theoretical Politics* **36**, 256–274.
- Davies B. 2022. The prospects for 'prospect utilitarianism'. *Utilitas* **34**, 33–43.
- Fleurbaey M. 2015. Equality versus priority: how relevant is the distinction? *Economics and Philosophy* **31**, 203–217.
- Fleurbaey M., B. Tungodden and H.F. Chang 2003. Any non-welfarist method of policy assessment violates the Pareto principle: a comment. *Journal of Political Economy* **111**, 1382–1385.
- Hirose I. 2016. Axiological sufficientarianism. In *What is Enough? Sufficiency, Justice, and Health*, eds C. Fourie and A. Rid, 51–68. Oxford: Oxford University Press.
- Huseby R. 2010. Sufficiency: restated and defended. *Journal of Political Philosophy* **18**, 178–197.
- Kahneman D. and A. Tversky 1979. Prospect theory: an analysis of decision under risk. *Econometrica* **47**, 263–292.
- Kaplow L. and S. Shavell 2001. Any non-welfarist method of policy assessment violates the Pareto principle. *Journal of Political Economy* **109**, 281–286.
- Nielsen L. 2019. Sufficiency and satiable values. *Journal of Applied Philosophy* **36**, 800–816.

- Nielsen L. 2023. The numbers fallacy: rescuing sufficientarianism from arithmeticism. *Inquiry* 1–22. <https://doi.org/10.1080/0020174X.2023.2186950>.
- Nussbaum M. 2000. *Women and Human Development*. Cambridge: Cambridge University Press.
- Parfit D. 1997. Equality and priority. *Ratio* 10, 202–221.
- Raz J. 1986. *The Morality of Freedom*. Oxford: Clarendon Press.
- Roemer J. 2004. Eclectic distributional ethics. *Politics, Philosophy and Economics* 3, 267–281.
- Sen A. 1979. Personal utilities and public judgements: or what's wrong with welfare economics? *Economic Journal* 89, 537–558.
- Stefánsson H.O. 2022. Continuity and catastrophic risk. *Economics and Philosophy* 38, 266–274.
- Suzumura K. 1983. *Rational Choice, Collective Decisions, and Social Welfare*. Cambridge: Cambridge University Press.
- Temkin L.S. 1996. A continuum argument for intransitivity. *Philosophy & Public Affairs* 25, 175–210.
- Temkin L.S. 2001. Inequality: a complex, individualistic, and comparative notion. *Philosophical Issues* 11, 327–353.

APPENDIX Proofs of Propositions

Proposition 1. *General Continuity and Welfarism are independent.*

Proof. First, we show that General Continuity does not imply Welfarism. Take any continuous, strictly monotone, numerical function $g: X \rightarrow \mathbb{R}$. Define $\succsim_g^*$ by letting: $(u, x) \succsim_g^*(v, y)$ iff $g(x) \geq g(y)$. We claim that $\succsim_g^*$ satisfies General Continuity. To see this, pick any $(u, x), (v, y) \in W \times X$ and suppose $(u, x) \succ_g^*(v, y)$. This implies that $g(x) > g(y)$. Take any sequence $\{(v^k, y^k)\}$ such that $(v^k, y^k) \rightarrow (v, y)$. Since g is continuous, there exists $k^* \geq 1$ such that $g(x) > g(y^k)$ for all $k > k^*$. By the definition of $\succsim_g^*$, we obtain $(u, x) \succ_g^*(v^k, y^k)$ for all $k > k^*$. Thus, General Continuity is satisfied. Take any $u \in W$ and $x, y \in X$ such that $x \gg y$. By the definition of $\succsim_g^*$, $(u, x) \succ_g^*(u, y)$, which violates Welfarism because Welfarism requires these to be indifferent. Second, we show that Welfarism does not imply General Continuity. Pick any critical sufficiency threshold $\theta \geq 0$ and consider the welfarist headcount sufficientarian ordering, defined as follows: $(u, x) \succ_{WH}^*(v, y)$ iff $\#\{i | u_i \geq \theta\} \geq \#\{i | v_i \geq \theta\}$. This satisfies Welfarism, but violates General Continuity; this follows from Proposition 8, which shows the violation of General Continuity of a general version of the headcount ordering. ■³⁰

Proposition 2. *An ethical preference relation $\succsim$ satisfies Welfarist Continuity if and only if it satisfies General Continuity and Welfarism.*

Proof. ‘Only If’ direction: Assume that $\succsim$ satisfies Welfarist Continuity. First, we show that it satisfies General Continuity. Let $(u, x), (v, y) \in W \times X$ such that $(u, x) \succ (v, y)$. Take a sequence $\{(v^k, y^k)\} \rightarrow (v, y)$. Since $v^k \rightarrow v$, Welfarist Continuity implies that there exists $k^* \in \mathbb{N}$ such that $(u, x) \succ (v^k, y^k)$ for all $k > k^*$. The other part in the parentheses can be similarly proved. Hence, General Continuity implies Welfarist Continuity.

Second, we show that $\succsim$ satisfies Welfarism. As an auxiliary step, we show that $u = v \Rightarrow (u, x) \sim (v, y)$. To the contrary, assume that $u = v$ but $(u, x) \succ (v, y)$. Take a sequence $\{(v^k, z^k)\}$ such that (i) $v^k = v$ for all $k \geq 1$ and (ii) $z^1 = y$ and $z^k = x$ for all $k \geq 2$. Then, Welfarist Continuity implies that there exists a $k^* \in \mathbb{N}$ such that $(u, x) \succ (v^k, z^k)$ for all $k > k^*$. This implies that $(u, x) \succ (v, x) = (u, x)$. This is a contradiction. Thus, we obtain $u = v \Rightarrow (u, x) \sim (v, y)$. Now, define $\succsim^*$ by letting $u \succsim^* v$ iff there exist $x^*, y^* \in X$ such that $(u, x^*) \succ (v, y^*)$. We show that, for any $(u, x), (v, y) \in W \times X$, $u \succsim^* v \Leftrightarrow (u, x) \succsim (v, y)$. To show $u \succsim^* v \Rightarrow (u, x) \succsim (v, y)$, take any $(u, x), (v, y) \in W \times X$ and assume that $u \succsim^* v$. Since $u \succsim^* v$, there exist $x^*, y^* \in X$ such that

³⁰Consider the welfarist leximin $\succsim_{WL}^*$, which is a lexicographic ordering on W ; see Suzumura (1983) and Adler (2019) for its precise definition. Now, define $\succsim_{WL}$ by letting $u \succsim_{WL} v \Leftrightarrow (u, x) \succ_{WL} (v, y)$. This is another example that satisfies Welfarism, but violates General Continuity.

$(u, x^*) \succ (v, y^*)$. To the contrary, assume that $(v, y) \succ (u, x)$. By the claim established in the auxiliary step above, we have $(u, x) \sim (u, x^*)$. Transitivity implies that $(v, y) \succ (v, y^*)$, contradicting the claim established in the auxiliary step. Hence, $(u, x) \succeq (v, y)$. To show that the converse direction is true, suppose that there exist $(u, x), (v, y) \in W \times X$ such that $(u, x) \succeq (v, y)$, but $v \succ^* u$. Then, by the definition of $\succeq^*$, there exist $x^*, y^* \in X$ such that $(v, y^*) \succ (u, x^*)$. By the claim established in the auxiliary step above, we have $(u, x^*) \sim (u, x)$ and $(v, y) \sim (v, y^*)$. By transitivity, this leads to $(v, y^*) \succ (v, y)$ and $(u, x) \succ (u, x^*)$, another contradiction. Thus, Welfarism holds.

If direction: Assume that $\succeq$ satisfies General Continuity and Welfarism. To prove that Welfarist Continuity holds, let $(u, x), (v, y) \in W \times X$ such that $(u, x) \succ (v, y)$. Take a sequence $\{(v^k, y^k)\}$ such that $v^k \rightarrow v$. By Welfarism, there exists an ordering $\succeq^*$ on W such that, for any $(u, x^*), (v, y^*) \in W \times X$, $u \succeq^* v$ iff $(u, x^*) \succeq (v, y^*)$. Now, we consider another sequence $\{(v^k, \hat{y}^k)\} \rightarrow (v, y)$. That is, $\hat{y}^k \rightarrow y$ holds, as well as $v^k \rightarrow v$. By General Continuity, there exists $k^* \in \mathbb{N}$ such that $(u, x) \succ (v^k, \hat{y}^k)$ for all $k > k^*$. Since Welfarism holds, $u \succ^* v^k$, which implies $(u, x) \succ (v^k, y^k)$ for all $k > k^*$. The other part in the parentheses can be similarly proved. Thus, Welfarist Continuity holds. ■

Proposition 3. *Restricted Welfare Continuity and Welfarism are logically independent.*

Proof. First, we show that Welfarism does not imply Restricted Welfare Continuity. Let $\underline{u} > 0$ and for any $(u, x) \in W \times X$, let $f(u, x) = \#\{i \in N | u_i \geq \underline{u}\}$. Define an ethical preference relation $\succeq$ on $W \times X$ as follows: for all $(u, x), (v, y) \in W \times X$, $(u, x) \succeq (v, y)$ if and only if $f(u, x) \geq f(v, y)$. We can confirm that $\succeq$ satisfies Welfarism. To see this, we can define a relation $\succeq^*$ on W by letting: $u \succeq^* v$ if and only if $\#\{i \in N | u_i \geq \underline{u}\} \geq \#\{i \in N | v_i \geq \underline{u}\}$. Note that for all $(u, x), (v, y) \in W \times X$, $u \succeq^* v$ if and only if $(u, x) \succeq (v, y)$. Thus, Welfarism holds. Now, let $u = (\underline{u}, 0, \dots, 0)$ and $v = (\underline{u}, \dots, \underline{u})$. Then, we have for all $x, y \in X$, $f(u, x) = 1$ and $f(v, y) = n$. Hence, we have $(v, y) \succ (u, x)$. Consider the sequence $\{(v^k, y)\}$, where $v^k = (\underline{u} - \frac{1}{k}, \dots, \underline{u} - \frac{1}{k})$ for all $k \in \mathbb{N}$. Then, we have $(v^k, y) \rightarrow (v, y)$. However, we have $(u, x) \succ (v^k, y)$ for all $k \in \mathbb{N}$. Therefore, the ethical preference relation $\succeq$ violates Restricted Welfare Continuity. Second, we show that Restricted Welfare Continuity does not imply Welfarism. Now, define another ethical preference relation $\succeq'$ such that for all $(u, x), (v, y) \in W \times X$, $(u, x) \succeq' (v, y)$ if and only if $\sum_{i \in N} \sum_{\ell=1}^m x_{\ell i} \geq \sum_{i \in N} \sum_{\ell=1}^m y_{\ell i}$. Suppose $(u, x) \succ' (v, y)$. Pick any sequence $\{(v^k, y)\}$ such that $(v^k, y) \rightarrow (v, y)$. Then, we have $(u, x) \succ' (v^k, y)$ for all $k \in \mathbb{N}$. Hence, the ethical preference relation $\succeq'$ satisfies Restricted Welfare Continuity. However, $\succeq'$ violates Welfarism because we have $(w, x) \succ' (w, y)$ for all $w \in W$. ■

Proposition 4. *General Continuity implies Restricted Welfare Continuity.*

Proof. Suppose that an ethical preference relation $\succeq$ on $W \times X$ satisfies General Continuity. Pick any $(u, x), (v, y) \in W \times X$, any sequence $\{(v^k, y)\}$ such that $(v^k, y) \rightarrow (v, y)$, and suppose $(u, x) \succ (v, y)$. Define a new sequence $\{(v^k, \hat{y}^k)\}$ such that $v^k = v^k$ and $\hat{y}^k = y$ for all $k \in \mathbb{N}$. Then, we have $(v^k, \hat{y}^k) \rightarrow (v, y)$. By General Continuity, there exists a $k^* \in \mathbb{N}$ such that $(u, x) \succ (v^k, \hat{y}^k) = (v^k, y)$ for all $k > k^*$. (The proof for the case where $(v, y) \succ (u, x)$ is analogous.) Hence, our ethical preference relation $\succeq$ on $W \times X$ satisfies Restricted Welfare Continuity. ■

Proposition 5. *Generalized headcount sufficientarianism violates General Continuity.*³¹

Proof. For each $i \in N$, let (u_i^*, x_i^*) be such that $h(u_i^*, x_i^*) = \theta^*$. Let $\mathbf{1}_s$ be the s -dimensional vector composed of s ones. Let (u^*, x^*) be a profile such that each individual $i \in N$ obtains (u_i^*, x_i^*) . Also, let (v, y) denote a profile such that the first individual obtains (u_1^*, x_1^*) and all of other individuals $i \in N \setminus \{1\}$ obtains $(u_i^*, x_i^*) - \mathbf{1}_{(1+m)}$. Note that $n = \#\{i \in N | h(u_i^*, x_i^*) \geq \theta^*\} > \#\{i \in N | h(v_i, y_i) \geq \theta^*\} = 1$, and, hence, $(u^*, x^*) \succ (v, y)$. Now, take the following sequence $\{(u, x)^k\}$ of profiles: $(u, x)^k = (u^*, x^*) - \frac{1}{k} \mathbf{1}_{n(1+m)}$.

³¹Some sufficientarian theories satisfy general continuity. For example, Benbaji (2005, 2006) proposes an ordering that does not exhibit 'absolute priority' for people below the threshold. However, this is not the type of sufficientarian theory that Chung and Davies are targeting.

This sequence converges to (u^*, x^*) . However, $(v, y) \succ (u, x)^k$ holds for all $k > 0$. Therefore, this violates General Continuity. ■

Proposition 6. *Utilitarianism satisfies General Continuity.*

Proof. Let the ethical preference relation $\succsim$ defined on $W \times X$ be utilitarian: that is, for any $(u, x), (v, y) \in W \times X$, $(u, x) \succsim (v, y)$ if and only if $\sum_{i \in N} u_i \geq \sum_{i \in N} v_i$. Pick any $(u, x), (v, y) \in W \times X$ and suppose $(u, x) \succ (v, y)$. Let $\delta = (\sum_{i \in N} u_i - \sum_{i \in N} v_i)$ and pick any sequence $\{(v^k, y^k)\}$ such that $(v^k, y^k) \rightarrow (v, y)$. Since $(v^k, y^k) \rightarrow (v, y)$, there exists a $k^* \in \mathbb{N}$ such that for all $i \in N$, $|v_i^k - v_i| < \frac{\delta}{n}$ for all $k > k^*$. Then, by the triangle inequality, we have for all $k > k^*$:

$$\begin{aligned} \left| \sum_{i \in N} v_i^k - \sum_{i \in N} v_i \right| &= |(v_1^k + \dots + v_n^k) - (v_1 + \dots + v_n)| \\ &= |(v_1^k - v_1) + \dots + (v_n^k - v_n)| \\ &\leq |v_1^k - v_1| + \dots + |v_n^k - v_n| < \frac{\delta}{n} + \dots + \frac{\delta}{n} = \delta. \end{aligned}$$

This implies that for all $k > k^*$, $\sum_{i \in N} u_i > \sum_{i \in N} v_i^k$, and, hence, $(u, x) \succ (v^k, y^k)$, as desired. ■

Proposition 7. *In any lifeboat situation, Sufficientarianism S always maximizes the number of instances meeting their critical sufficiency threshold θ as long as the critical level α is sufficiently large.*

Proof. Consider any lifeboat situation where the available resources are $k\theta$, where $k < |N| = n$. So, we have just enough resources to allow $k < n$ people meet their critical sufficiency threshold θ . Let $x^k = (x_1^k, \dots, x_n^k)$ be any distribution of the transferable goods that makes k (the maximum number of) individuals meet their critical sufficiency thresholds θ , and, without loss of generality, rearrange the individuals so that $x_i^k = \theta$ for $i = 1, \dots, k$ and $x_j^k = 0$ for $j = k + 1, \dots, n$, i.e.,

$$x^k = \left(\underbrace{\theta, \theta, \theta, \dots, \theta}_{(k \text{ individuals})}, \underbrace{0, 0, 0, \dots, 0}_{(n-k \text{ individuals})} \right).$$

We claim that distribution x^k maximizes our sufficientarian SWF. To show this, we will show that there is no way to increase the value of our sufficientarian SWF by moving to an alternate distribution. Obviously, moving to any other distribution that makes a different set of k individuals meet their critical sufficiency thresholds θ will generate the same value for our sufficientarian SWF, and, hence, will not increase its value. Now, consider moving to another distribution $x^a = (x_1^a, \dots, x_n^a)$ under which $a < k$ individuals meet their critical sufficiency thresholds θ . Without loss of generality, let the individuals from 1 to a be the individuals who meet their sufficiency threshold θ in distribution x^a . Then, for individuals $i \in \{a + 1, a + 2, \dots, k\}$, we have $x_i^k = \theta > x_i^a$. Since g is increasing and concave, the best distribution that would generate the highest value for the sufficientarian SWF, while having $a < k$ individuals meet their critical sufficiency thresholds θ , would give $x_i^a = \theta$ for $i = 1, \dots, a$ and give $x_j^a = \frac{k-a}{n-a}\theta < \theta$ for $j = a + 1, \dots, n$, i.e.

$$x^a = \left(\underbrace{\theta, \dots, \theta}_{(a \text{ individuals})}, \underbrace{\frac{k-a}{n-a}\theta, \dots, \frac{k-a}{n-a}\theta}_{(n-a \text{ individuals})} \right).$$

According to our sufficientarian SWF, for individuals $i \in \{a + 1, a + 2, \dots, k\}$, moving from $x_i^k = \theta$ to $x_i^a = \frac{k-a}{n-a}\theta$ generates the following total social loss:

$$- \underbrace{(k-a)}_{\text{(a total of } k-a \text{ individuals)}} \left\{ \underbrace{\left[g(x_i^k = \theta) - g\left(x_i^a = \frac{k-a}{n-a}\theta\right) \right]}_{i\text{'s welfare loss}} + \underbrace{[g(\alpha) - g(\theta)]}_{\text{moral disvalue of failing to meet threshold } \theta} \right\} \quad (1)$$

For individuals $j = k+1, k+2, \dots, n$, moving from $x_j^k = 0$ to $x_j^a = \frac{k-a}{n-a}\theta$ generates the following total social gain:

$$(n-k) \underbrace{\left[g\left(x_j^a = \frac{k-a}{n-a}\theta\right) - g(x_j^k = 0) \right]}_{j\text{'s welfare gain } (j=k+1, \dots, n)} \quad (2)$$

A move from distribution x^k to distribution x^α would *not* be morally preferable, if and only if,

$$\begin{aligned} (1) + (2) &\leq 0 \\ \Leftrightarrow (n-k) \left[g\left(\frac{k-a}{n-a}\theta\right) - g(0) \right] + (k-a) g\left(\frac{k-a}{n-a}\theta\right) &\leq (k-a) g(\alpha) \\ \Leftrightarrow \frac{n-a}{k-a} g\left(\frac{k-a}{n-a}\theta\right) - \frac{n-k}{k-a} g(0) &\leq g(\alpha) \end{aligned}$$

which is true whenever the critical level α is sufficiently high. Hence, in any lifeboat situation, moving from a distribution that makes the maximum number of individuals meet their critical sufficiency thresholds θ to another distribution that makes a lesser number individuals meet their critical sufficiency thresholds θ will never increase the value of the sufficientarian SWF as long as the critical level α is sufficiently high. ■

Proposition 8. *The moral ordering of Sufficientarianism S is discontinuous at the critical sufficiency threshold θ whenever the critical level α exceeds the critical sufficiency threshold θ , i.e. whenever $\alpha > \theta$.*

Proof. According to our sufficientarian SWF, a distribution x is strictly morally preferred to another distribution y if and only if

$$S(x) = \sum_{i \in N: x_i < \theta} [g(x_i) - g(\alpha)] > \sum_{i \in N: y_i < \theta} [g(y_i) - g(\alpha)] = S(y).$$

If $\alpha = \theta$ (i.e. if the critical level α is equal to the critical sufficiency threshold θ), then, for any distribution $x = (x_1, \dots, x_n)$, our sufficientarian SWF can be re-written as:

$$S(x) = \sum_{i \in N} \min\{g(x_i) - g(\theta), 0\}.$$

We note that both $g(x_i) - g(\theta)$ and $h(x) = 0$ (i.e. a constant function) is continuous. Therefore, $\min\{g(x_i) - g(\theta), 0\}$ is also continuous. Moreover, the sum of continuous functions yields a continuous function. We thus conclude that the moral ordering induced by our sufficientarian SWF, S , is continuous when $\alpha = \theta$.

Now, suppose $\alpha > \theta$. Consider a distribution $x = (\theta, \dots, \theta)$ where everybody meets their critical sufficiency threshold θ . Note that $S(x) = 0$. Now, consider a sequence of distributions $\{x^k\} = (\theta - \frac{1}{k}, \dots, \theta - \frac{1}{k})$ for each $k \in \mathbb{N}$. Note that $x^k \rightarrow x$ (i.e. x^k converges to x). Then, we have:

$$\begin{aligned} \lim_{k \rightarrow \infty} S(x^k) &= \lim_{k \rightarrow \infty} \left\{ n \left[g\left(\theta - \frac{1}{k}\right) - g(\theta) \right] + n [g(\theta) - g(\alpha)] \right\} \\ &= [g(\theta) - g(\theta)] + n [g(\theta) - g(\alpha)] \\ &= n [g(\theta) - g(\alpha)] \neq 0 = S(x). \end{aligned}$$

Hence, our sufficientarian SWF is discontinuous at $x = (\theta, \dots, \theta)$. As an illustration, consider a two-person case and consider two distributions: $u = (u_1, u_2) = (\theta, \theta)$ and $v = (v_1, v_2) = (\theta, 0)$ and suppose $g(\alpha) > 2g(\theta) - g(0)$. Then, since $S(u) = 0 > g(0) - g(\alpha) = S(v)$, we have $u \succ v$. Now, consider a sequence $\{u^k\} = (u_1^k, u_2^k) = (\theta - \frac{1}{k}, \theta - \frac{1}{k})$. Note that $u^k \rightarrow u$.

Then, from $g(\alpha) > 2g(\theta) - g(0)$, we have, for all $k \in \mathbb{N}$:

$$\begin{aligned} g(\alpha) &> 2g(\theta) - g(0) \\ \Rightarrow g(0) - g(\alpha) &> 2[g(\theta) - g(\alpha)] > 2[g(\theta - \tfrac{1}{k}) - g(\alpha)] \\ \Rightarrow g(0) - g(\alpha) &> 2[g(\theta - \tfrac{1}{k}) - g(\alpha)] \\ \Rightarrow S(v) &> S(u^k) \end{aligned}$$

and, hence, $v \succ u^k$. Hence, there exists no $k^* \in \mathbb{N}$ such that $u^k \succ v$ for all $k > k^*$. So, the moral ordering of sufficientarian S is discontinuous at the critical sufficiency threshold θ . ■

Susumu Cato is Professor of Economics at the Institute of Social Science, University of Tokyo. His research interests include decision theory, social choice theory, axiology, political philosophy and organizational economics. Email: cato@iss.u-tokyo.ac.jp, susumu.cato@gmail.com. For more information, see: <https://sites.google.com/site/susumucato/>

Hun Chung is Associate Professor in the Department of Data and Decision Sciences at Emory University. His research spans formal ethics/political philosophy, game theory/social choice theory, and the interdisciplinary field of PPE (Philosophy, Politics, and Economics). Prior to joining Emory, he was Professor in the Faculty of Political Science and Economics at Waseda University Email: hun.chung@emory.edu, hunchung1980@gmail.com. For more information, see: <https://sites.google.com/site/hunchung1980>.