

**Algorithmic care: Peter Singer's ethics and the challenge of AI
in the end-of-life medicine****Louise Batôt¹ & Alessio Belli²****Abstract**

This article applies Peter Singer's ethical framework to evaluate the impact of predictive AI systems in end-of-life care. Singer's work is particularly fitting for this analysis because it focuses on the normative variables that AI technologies influence: the moral significance of suffering, the formation of reflective preferences, and the assignment of responsibility for foreseeable outcomes. His perspective on personhood, emphasis on autonomy as rational self-determination, and principle of equal consideration of interests provide a solid foundation for assessing how algorithmic models affect clinical timing, deliberation, and risk allocation. By examining two predictive tools used in end-of-life care, the article demonstrates how Singer's framework clarifies the ways in which AI can enhance procedural rationality in end-of-life decisions. Additionally, it identifies the risks that AI poses to autonomy, responsibility diffusion, and the reproduction of structural inequalities. Ultimately, Singer's framework offers a conceptual means to differentiate algorithmic interventions that support ethical decision-making from those that undermine the non-negotiable ethical conditions he advocates.

Keywords: Peter Singer, bioethics, palliative care, artificial intelligence, AI, suffering

Introduction

In the ethically sensitive field of end-of-life medicine, AI systems are increasingly utilized for diagnosis, information management, and care allocation (Davenport & Kalakota, 2019; Stafie et al., 2023). The literature highlights several advantages: AI can help anticipate patient suffering, facilitate earlier and more structured discussions between clinicians and patients, and, in some cases, contribute to the rehumanization of medical practice. However, current research emphasizes the need for a theory-based ethics of AI that can effectively address the challenges posed by these tools. AI systems may significantly threaten patients' decision-making autonomy and complicate the distribution of responsibility for clinical decisions—issues that are particularly critical in end-of-life care.

To address these emerging concerns, we propose applying Peter Singer's ethical framework. As one of the leading philosophers in contemporary practical ethics, Singer approaches sensitive issues such as abortion and euthanasia with a perspective free from religious or outdated beliefs, grounded instead in a rational concern for individual agency. His work provides a solid foundation for developing a theory-based ethics of AI that can assess the technology's impact in medical contexts. In this article, we argue that Singer's framework clarifies how algorithmic models in end-of-life care can enhance both the procedural and consequentialist aspects of clinical practice, while also emphasizing that such systems cannot exercise conscious or accountable moral agency.

While algorithmic tools can optimize outcomes, anticipate suffering, and support rational deliberation, they cannot fulfill the essential requirement of a Singerian ethical framework: the exercise of deliberative and morally responsible agency. This is a key contribution of Singer's ethics to the current transdisciplinary discussions regarding normative structures for AI in end-of-life care. We emphasize this aspect of Singer's theory because it provides a clear analytical tool for assessing how AI impacts moral agency and deliberative responsibility, while avoiding absolutist positions by recognizing both the potential benefits and the challenges that these technologies present. The paper is structured as follows: it first reconstructs the Singerian ethical framework, then examines the use of algorithmic applications in end-of-life care, and

¹ University of Luxembourg (Luxembourg); email: batotlouise@outlook.fr; ORCID: 0009-0008-5013-4230

² University of Luxembourg (Luxembourg); email: alessiobelli705@gmail.com; ORCID: 0009-0007-7908-9480

finally evaluates their normative adequacy according to Singer's criteria.

What it means to be a person

To outline the ethical framework guiding our analysis, we begin by reconstructing the key philosophical shifts introduced by Singer in contemporary bioethics. Throughout his work, Singer aims to reformulate the foundations of moral judgment by challenging what he views as the limitations of traditional Christian ethics—particularly its opposition to practices such as abortion, assisted suicide, euthanasia, and prostitution (Androne, 2017, p. 37). He argues that technological and medical advances, especially in end-of-life care, reveal the inadequacy of inherited moral doctrines to address the complexities of contemporary cases.

Singer contends that ethical reflection should abandon metaphysical assumptions rooted in absolute truths. Instead, he advocates for understanding ethics as an extension of our natural moral capacities—an evolved faculty that guides our intuitive judgments of right and wrong (Singer, 2016, p. 17). These intuitions, however, are not definitive; they should be subjected to rational scrutiny. According to Singer, the aim of moral philosophy is to examine and refine these judgments while reassessing how we assign value to life and death, bringing ethical deliberation back to a secular and empirically informed domain (Singer, 2016, p. 18). To live ethically, he argues, means adopting standards that we can rationally justify and defend.

By grounding his work in a critical analysis of moral thought, Singer has sought since *Practical Ethics* to dismantle one of the central tenets of Christian moral doctrine: the principle of the sanctity of life. This principle, foundational in many modern legal codes and implicit in the right to life, asserts that human life holds a special and superior value compared to all other forms of life. According to this view, killing a human being is always wrong, and no meaningful moral comparison can be made between ending a human life and ending any other kind of life.

To address this, Singer distinguishes between being human—biologically belonging to the species *Homo sapiens*—and being a person, defined as an individual possessing characteristics such as self-awareness, the capacity to feel pain and pleasure, the ability to form interests and desires, and an awareness of temporal continuity across past, present, and future. This distinction introduces a qualitative criterion whereby personhood is not restricted to adult, rational human beings but may also extend to nonhuman animals that exhibit these morally relevant capacities.

The value of a life—and the moral implications of ending it—hinges on the quality of the being's conscious experience and its capacity for pleasure, pain, interests, and preferences. According to Singer, the significance of violating the right to life does not depend on species membership, which he views as morally irrelevant and a secular remnant of religious doctrines. Instead, it should be assessed based on characteristics such as rationality, self-consciousness, awareness, autonomy, and the ability to experience pleasure and pain (Singer, 2011, p. 135). Consequently, the right to life is not absolute; it is a characteristic of persons, and its applicability depends on whether the being in question possesses the qualities that contribute to a life's value.

Singerian utilitarianism and the Principle of equal consideration

These ethical considerations exist not within a deontological framework but within a consequentialist one. This distinction is important for our understanding of how AI influences end-of-life decision-making; our goal is not solely to evaluate its morality based on strict adherence to established rules. Singer, drawing from the British utilitarian tradition, has thoroughly explored its variants (including classical and hedonistic utilitarianism) and has developed a distinctive position (Singer, 2011, p. 135; Singer, 2003). As previously mentioned, for Singer, “to live ethically is to live according to standards one can defend” (Buckle, 2005, p. 177). However, defending an ethical position cannot rest on partial or group-based interests.

The Kantian demand for universalization, within Singer's framework, transforms into the requirement that "I must formulate my standards as universal judgments that apply to me because they apply equally to all" (Buckle, 2005, p. 177).

To act ethically means giving equal consideration to the interests of everyone affected and choosing "the course of action most likely to maximize the interests of those affected ... that has the best consequences, on balance, for all" (Singer, 2011, p. 13). This principle of equal consideration asserts that ethical reflection must attribute equal value to the interests of all parties involved, serving as the central tenet of Singer's utilitarianism. Within this framework, a preference is defined as what an individual selects after careful, informed, and thoughtful reflection aimed at promoting their own interests (Pauer-Studer, 1993). Consequently, a preference cannot be replaced or adequately represented by statistical measures, automated processes, or algorithmic proxies without losing the subjective and deliberative aspects that provide it with moral authority.

These insights, pivotal to Singer's philosophy, inform our analysis as we assess the impact of AI on end-of-life decision-making. From this perspective, taking a person's life is wrong because it undermines that individual's established preference to continue living. However, when a competent person determines that their future will involve more suffering than joy and autonomously concludes that dying would be preferable, the argument against killing transforms into a reason to honor that person's wish (Singer, 2003).

This view assumes the individual possesses the qualities that define a 'person' in Singer's interpretation. Indeed, "the principle of equal consideration of interests does not protect (in the same way) the lives of people whose rationality or self-awareness has been destroyed by illness or damage to the body" (Synowiec, 2019, p. 54). Moreover, it does not fully apply to potential capacities that have not yet developed, such as those of fetuses or newborns. According to Singer, such beings lack the rational, self-aware interests that support a person's claim to continued life and thus cannot be regarded as persons in the morally relevant sense.

This does not imply that it is inherently justified to kill a being that has not yet developed the characteristics of a person or that has lost them fully or partially. To evaluate the moral status of killing a being incapable of forming rational preferences, Singer introduces an additional argument based on classical utilitarianism. Utilitarianism aims to maximize the overall balance of pleasure over pain, suggesting that it is wrong to kill a being capable of experiencing pleasure if its expected future pleasures outweigh its expected pains. In elaborating this framework, Singer also considers other factors, including wishes that establish a right to life, respect for autonomy, and the indirect consequences that killing may have on others (Pauer-Studer, 1993). In general, he identifies two interconnected criteria for assessing the morality of killing: the preference to continue living and the capacity to experience pleasure.

When applied to end-of-life situations, this leads to the following interpretation: when a being is self-conscious and can articulate long-term preferences, the frustration of those preferences weighs heavily on the moral implications of death. The hedonic aspect—the expected balance of pleasure and pain—plays a role in those preferences. However, when reflective agency is absent or diminished, the ethical assessment shifts to the hedonic dimension alone as a proxy for well-being. This shift maintains continuity between rational and merely sentient forms of moral consideration. At the same time, certain cases may present hybrid situations where drawing a clear line between preference-based and hedonistic judgments becomes challenging. Therefore, these two levels of assessment will guide our analysis of how AI systems in end-of-life care engage with, approximate, or at times conflate Singer's criteria.

Conceptual framework of Singer's position on end-of-life care

With Singer's theoretical framework in mind, we can now examine his views on end-of-life care. As Singer states, "euthanasia originally meant a 'gentle and easy death', but it is now used

to describe the killing of individuals who are incurably ill and suffering, in order to spare them further pain or distress” (Singer, 2011, p. 157). He identifies three forms of euthanasia: 1) *Voluntary euthanasia* occurs when the decision is made at the competent and informed request of the person who is to die; 2) *Involuntary euthanasia* takes place when the person is capable of consenting but is not asked or is asked and chooses to continue living; 3) *Non-voluntary euthanasia* refers to cases involving individuals who cannot comprehend the choice between life and death and who have not previously expressed a preference relevant to their current condition (Singer, 2011, p. 158). Regarding voluntary euthanasia, Singer argues that it is morally permissible under specific conditions. These conditions include: the individual making a free and rational decision to request death; confirmation by two physicians that the person is suffering from an illness that causes, or is certain to cause, severe and unrelievable suffering; and the submission of a written request at least thirty days prior to the intended act, witnessed by two individuals (Pauer-Studer, 1993).

Singer’s perspective requires us to move beyond the strict belief that it is always wrong to take the life of an innocent human being. For those who possess full cognitive abilities, the primary focus should be on upholding their autonomy and responsibility regarding their own lives, “allowing them to decide whether their lives are worth living” (Singer, 2003, p. 529). If they determine that their lives are not worth living, they should have the option to end them. From this viewpoint, the role of AI should be to support and enhance an individual’s agency, rather than to replace it or serve as a mere substitute.

Singer does not consider involuntary euthanasia to be morally permissible, but he recognizes non-voluntary euthanasia as a complex issue. To clarify this complexity, he distinguishes between passive euthanasia—allowing someone to die—and active euthanasia—directly intervening to bring about death. He argues that “there is no intrinsic moral difference between killing and allowing to die” (Singer, 2011, p. 183). However, this assertion does not imply simple moral equivalence; Singer acknowledges “extrinsic differences – especially differences in the time it takes for death to occur” (Singer, 2011, p. 185)—which carry moral significance and must be considered. Therefore, if passive euthanasia can be morally justified, then “active euthanasia should also be accepted as humane and proper in certain circumstances” (Singer, 2011, p. 183). This distinction is particularly relevant in cases of non-voluntary euthanasia, which may involve individuals in irreversible comas or vegetative states, or those with severe, untreatable impairments who are being kept alive by artificial means. In such instances, Singer deems active euthanasia permissible, as it not only alleviates irremediable suffering but also ends a life that, according to his criteria, is no longer worth living.

Access to palliative care is essential in this context. However, for Singer, the central point is that seriously ill individuals must have the freedom to make rational choices about how they wish to approach the final stage of life, and physicians should support them in these choices. Autonomy, as a governing ethical principle, necessitates that clinicians give equal consideration to all available options and their consequences, thereby protecting and strengthening the patient’s agency. Palliative care should therefore be presented alongside options like medically assisted suicide and voluntary euthanasia, rather than serving as a basis for excluding them or stigmatizing the choice to die. Failure to do so would undermine autonomy and risk causing—or prolonging—suffering that the individual perceives as unbearable.

In the final chapter of *Rethinking Life and Death*, Singer distills his effort to go beyond the limitations of traditional ethics into five core propositions. He argues that an ethical theory that is adequate to contemporary challenges must:

- 1) Recognize that the worth of human life varies according to morally relevant characteristics such as self-awareness, the capacity for physical, emotional, and social interaction, and the ability to form conscious preferences (Singer, 1994, p. 219).

- 2) Take responsibility for the consequences of one’s decisions by acknowledging that

practitioners are equally accountable for actions and omissions when outcomes are foreseeable (Singer, 1994, p. 195). This responsibility should not be diluted through standardized protocols; clinicians must continually question whether a decision they foresee will end a patient's life is, all things considered, the right one (Singer, 1994, p. 196).

3) Respect a person's desire to live or die, placing autonomy and rational preference at the center of ethical deliberation—particularly when distinguishing persons from humans more generally (Singer, 1994, pp. 197–200).

4) Bring children into the world only if they are wanted, challenging the assumption that fetuses or newborns possess an intrinsic right to life independent of personhood (Singer, 1994, p. 200).

5) Reject discrimination based on species and affirm that moral worth derives from capacities rather than biological classification (Singer, 1994, p. 202).

These principles are historical and evolving, and their application cannot be automatic; it must be responsive to the specifics of each case and the balances those particulars demand. Not all five principles are equally relevant in the empirical cases discussed below. Principles 4 and 5 play only a peripheral role in end-of-life medicine as mediated by AI, while the first three principles suffice to construct the normative framework guiding our assessment of AI-assisted decision-making in this context. This distinction is particularly important because AI is not just a technical tool; its moral significance is influenced by how it is used. AI acts as an intermediary within a decision-making process that involves both the patient and the clinician, and its impact on personal preferences and interests requires careful examination.

In the following sections, we will explore how certain algorithmic tools may seem to operationalize Singer's core ethical commitments, while also evaluating the tensions these tools may introduce. Throughout our discussion, we emphasize that, according to Singer, moral weight is given to individuals' capacities and interests, not their social value or group membership. Our argument focuses on the ethics of decision-making concerning costly interventions, rather than making judgments about which lives hold value in a broader sense.

Promises of AI in end-of-life care

Singer's consequentialist framework expands moral responsibility to include any process that can produce foreseeable outcomes, such as procedural or algorithmic decision-making. This means that systems used for triage or terminal sedation protocols are also subject to moral evaluation. Ethically, what matters is not whether a system acts or refrains from acting—since, as Singer points out, “there is no difference which depends solely on the distinction between an act and an omission” (Singer, 2011, p. 183)—but whether its operation predictably alleviates or prolongs suffering and respects or frustrates the patient's considered preferences. This perspective guides our analysis of how algorithmic tools in end-of-life care may implement or sometimes distort Singer's principles.

The rise of predictive algorithms in end-of-life care shows a significant alignment with several key aspects of Singer's ethics, including the focus on a patient's capacity to experience suffering and the gradual development of rational preferences through clinical dialogue. Predictive systems have demonstrated the ability to “outperform traditional, clinically used predictive models”³ (Rajkomar et al., 2018, p. 4). Tools that engage with complex ethical considerations often align with Singer's two macro-criteria: respecting the rational preferences of individuals capable of self-reflection and minimizing suffering for those whose ability to

³ This observation is based on results from a peer-reviewed study that developed and validated deep-learning models using data from over 216,000 hospitalizations across two U.S. academic centers (UCSF and Chicago Medicine). These models achieved predictive accuracy of up to 95% for in-hospital mortality, highlighting a significant improvement in distinguishing high-risk patients from low-risk patients compared to traditional clinical scoring methods, which typically achieve around 91% accuracy.

form such preferences is impaired. They also address the intermediate processes through which preferences are shaped and revised during medical encounters. Indeed, “AI significantly transforms the traditional health care paradigm toward an evidence-based and patient-centered model” (Abejas et al., 2025, p. 2), which presents a promising development from a Singerian perspective.

For example, the short-term mortality prediction model developed by Jvion CORE analyzes a combination of clinical, behavioral, and socio-demographic variables to identify oncology patients at a high risk of mortality within thirty days. At Northwest Medical Specialties, the algorithm produces recommendations, such as redirection to palliative care or hospice, which are sent to a care coordinator. This coordinator can then schedule or initiate discussions about end-of-life care with the patient (Gajra, 2021).⁴ Each alert is accompanied by essential risk factors and explanatory indicators, enabling physicians to understand the reasoning behind the prediction and to prioritize the timing of conversations and care planning. This statistical anticipation intervenes early in the clinical workflow, even before patients may have expressed explicit preferences.

By identifying patients at significant risk of pain and deterioration before they can articulate their preferences, the algorithm temporarily substitutes predictive reasoning for preference satisfaction when rational agency is diminished (Van den Beuken-van Everdingen, 2007). In Singer’s framework, the ability to suffer is the minimal condition for moral consideration: unarticulated suffering still indicates the fundamental goal of minimizing pain and promoting well-being. Preferences further refine this orientation through reflective and temporally extended judgment. From this perspective, the algorithmic model serves as a hybrid space between predicting suffering and tracing preferences, pinpointing cases where the anticipated trajectory of suffering is significant enough to justify earlier clinical engagement. This anticipatory function can help establish the temporal and informational conditions necessary for forming well-considered preferences later on.

At this stage, the first Singerian principle—that the moral weight of a life varies based on characteristics relevant to suffering, self-awareness, and agency—seems to be partially reflected in the algorithm’s focus on predicted suffering. This alignment supports Singer’s emphasis on qualitative considerations over strictly quantitative ones. The statistical predictions can also serve as a preliminary indicator of the urgency of preference satisfaction, signaling where morally significant interests (such as avoiding severe pain) are most at risk. This supports clinical prioritization in a manner consistent with consequentialist reasoning. However, the alignment remains procedural rather than substantive, as the model highlights conditions under which preferences might matter without accessing or replacing the preference itself. The individualized construction of patient trajectories further reinforces this procedural alignment.

The incorporation of behavioral and socio-demographic indicators enables a strong contextual assessment, resonating with Singer’s view that ethical judgment must consider the specific characteristics of the sentient subject rather than rely solely on general rules. Even the 30-day predictive window can be viewed as ethically significant, delineating a period in which agency and reflective preference formation may still be exercised. Finally, the third principle is indirectly honored: by providing doctors with explanatory factors behind each alert, the system can facilitate a deliberative process in which information is shared, options are clarified, and preferences can begin to take shape. In this sense, its contribution lies less in determining

⁴ For each patient identified as high- or medium-risk, the system highlights the key clinical and socioeconomic factors influencing the risk estimate, along with five personalized intervention suggestions (e.g., palliative care, pain management, symptom control, or social support). The primary variables that shape the risk assessment include functional decline, comorbidity burden, cancer type, pain intensity, and indicators of socioeconomic vulnerability such as income, education, and living conditions.

outcomes and more in structuring the conditions for informed and autonomous moral deliberation—an essential component of Singer’s ethical framework.

Another machine-learning algorithm, utilized at the University of Pennsylvania Health System (UPHS), was prospectively deployed across 18 outpatient oncology departments (Manz et al., 2020).⁵ The model generates real-time short-term mortality risk estimates, allowing clinicians to identify high-risk patients earlier and initiate timely discussions about end-of-life goals and preferences. Compared to conventional prognostic scores such as ECOG and Elixhauser, the algorithm significantly improved the predictive accuracy of high-risk alerts (Manz et al., 2020, pp. 1727–1728).⁶ In practice, this meant fewer unnecessary alerts and a more efficient allocation of clinical attention to patients nearing the end of life. This approach maximizes the value of end-of-life conversations while minimizing avoidable anxiety and interventions.

This situation closely aligns with Singer’s principle of respect for considered preferences, which asserts that moral action requires recognizing and acting upon an individual’s informed and rational choices. Within this framework, the UPHS model acts as a tool that promotes considered preferences: decisions about living or dying remain in the hands of fully informed and autonomous agents. By improving prognostic accuracy, the system allows both patients and clinicians to engage in more meaningful discussions about continuing or interrupting treatment, as well as when to initiate palliative care or (where legally permitted) voluntary euthanasia in line with the patient’s reflective preferences. Thus, the algorithm supports moral autonomy by reducing prognostic uncertainty and enhances the clinician’s ability to facilitate self-determined choices—another contribution consistent with Singer’s emphasis on promoting and respecting considered preferences.

Overall, these cases highlight why some experts claim that AI in medicine can ‘rehumanize’ clinical practice. Several commentators contend that “the greatest opportunity offered by AI is not reducing errors or workloads, or even curing cancer, but the opportunity to restore the precious and time-honored connection and trust—the human touch—between patients and doctors” (Topol, 2019, p. 18). This restoration enables clinicians to focus on “the tasks that are uniquely human: building relationships, exercising empathy, and using judgment to guide and advise” (Fogel & Kvedar, 2018, p. 3) and allows them “spend more time with patients, hence re-humanizing clinical practice” (Jotterand & Bosco, 2020, p. 2457). In fields such as oncology—where patients “frequently receive care near the end of life that is not concordant with their wishes and worsens their quality of life” (Manz et al., 2020, p. 1727)—this shift could represent a significant improvement.

From this perspective, AI can enhance moral reasoning by improving the timing and conditions under which patient preferences are identified and acted upon. Together, these two models illustrate the range of Singer’s criteria for moral assessment. Jvion CORE functions as a hedonic-preference hybrid that is suitable for cases where agency is weak or latent, while UPHS offers a more straightforward preference-based application grounded in autonomy and deliberation. Both models demonstrate how Singer’s dual criteria can inform the ethical interpretation of algorithmic tools in end-of-life care. However, this convergence has its limits. The procedural efficiencies introduced by algorithmic systems cannot address the deeper moral asymmetry between algorithmic optimization and the reflective, accountable moral reasoning that Singer’s framework requires.

⁵ This peer-reviewed study prospectively validated a machine learning model that predicts 180-day mortality among 26,525 oncology outpatients across 19 clinics in the University of Pennsylvania Health System.

⁶ The baseline performance was established using ECOG and Elixhauser scores, with the machine learning model outperforming both: it showed an improvement of 0.17 over ECOG and 0.20 over Elixhauser. This improvement corresponds to roughly one-third fewer false alerts and a positive predictive value of 45%.

From promises to tensions

Despite some promising alignments, the use of AI in end-of-life care does not fully meet the essential requirements of a Singerian ethical framework. While certain AI models demonstrate a form of consequentialist reasoning, the literature reveals significant limitations in truly recognizing patient preferences, supporting the assumption of full moral responsibility, and treating each instance of suffering with equal moral weight—elements that are crucial in Singer’s ethics.

A closer examination raises doubts about the alignment with Singer’s third principle. When an algorithm predicts a patient’s suffering, it may influence care decisions without the patient having clearly expressed or confirmed their preferences. As several scholars point out, “without taking into account value plurality, there is a real risk of the AI’s decisions undermining the patient’s autonomy” (Sauerbrei et al., 2023, p. 8). In clinical environments, algorithmic recommendations often lack transparent justifications, making it difficult to explain how a particular conclusion was reached (Nikoloudi & Mystakidou, 2025). This lack of transparency can weaken the clinician’s ability to act responsibly and diminish the patient’s capacity for autonomous, informed decision-making (Abejas et al., 2025).

Singer emphasizes the importance of conscious and rational preferences—not mere approximations. Ethical action must be grounded in what individuals know, want, and choose as situated subjects. In medical decision-making, preferences develop and evolve through dialogue, information exchange, and reflection. In end-of-life care, this process typically occurs gradually, from early prognostic discussions to revisable goals-of-care and advance planning. This gradual unfolding allows patients to test, adjust, and stabilize their choices as their condition and understanding change.

This dynamic process is what gives preferences their ethical significance in Singer’s view: they stem from self-awareness, deliberation, and informed consent. Predictive systems can potentially enhance this process by improving the timing of conversations and clarifying uncertain prognostic horizons. However, by suggesting options, framing outcomes, or establishing defaults, these systems can also subtly shape the decision-making landscape and influence how preferences are formed or revised. The ethical question that arises is not only whether algorithms make accurate predictions but also whether they maintain the reflective space in which authentic preferences can develop. When an algorithm infers a preference from statistical associations—such as treatment history, socioeconomic profile, or other proxies—it effectively creates a substitute for valuation that bypasses the subjective process of judgment. In doing so, it risks stripping the preference of its moral weight and treating the patient not as a person but as a collection of signals.

In the end, the system operates based on a preference that was never truly articulated. A Singerian preference is a deliberate, context-specific expression of interest by a self-aware individual, which cannot be captured by an algorithm that only predicts behaviors. While the algorithm’s output may resemble a preference statistically, it lacks the intentional and moral depth that confers ethical significance. This reduction risks undermining the very concept of moral agency that Singer aims to uphold—especially when clinical teams treat algorithmic suggestions as morally authoritative without critical examination.

Such situations can arise because algorithms assume a three-dimensional position of authority: epistemic, due to the credibility of their scores (Mittelstadt et al., 2017); institutional, through their integration into clinical protocols (Floridi & Cowls, 2019); and structural, because they create a monoculture of decision-making that standardizes clinical judgment around algorithmic outputs, thereby diminishing epistemic and moral diversity (O’Neil, 2016). In these cases, the algorithm no longer aids ethical deliberation; instead, it begins to replace it—an outcome that Singer’s framework is designed to prevent. From a Singerian perspective, end-of-life AI systems must maintain deliberative autonomy: their recommendations should be

interpretable to clinicians, serving as prompts for dialogue rather than substitutes for informed choice. Only under these conditions can patients articulate—and, when necessary, revise—their preferences through informed discussion.

The second principle states that ethical responsibility applies equally to both action and inaction when outcomes are foreseeable. This also highlights structural flaws in how algorithms operate. AI systems often create a grey area of responsibility, where errors in prediction (like false negatives) or failures to follow up (such as inaction following an alert) result in moral harm without a clearly identifiable agent (Ferlito et al., 2024). The study of the UPHS algorithm illustrates this issue: even with a 40% risk threshold, the algorithm's sensitivity was only 27% (Manz et al., 2020, p. 1729). This means that fewer than one-third of patients who died shortly afterward were accurately identified. Many of these patients, often those with atypical cancer trajectories, did not receive timely referrals to palliative care. From a Singerian perspective, such omissions constitute moral failures, as they hinder the prompt recognition and expression of preferences that could have been formed—and respected—had the risks been anticipated.

Algorithmic opacity exacerbates the challenges within healthcare systems. It limits clinicians' ability to justify their decisions in clear and shareable ways. When clinicians cannot explain how a risk score was generated or defend the reasoning behind an algorithmic recommendation, their decisions lose the moral quality associated with acting on sound reasoning, as suggested by Singer. Furthermore, this lack of clarity disperses accountability across various entities—datasets, designers, and clinical protocols—resulting in no single agent taking full ownership of the consequences. As a result, responsibility becomes diluted within the system's framework.

This creates a moral vacuum: a situation where actions occur without clear justification and where it is difficult to identify who is accountable. In Singer's view, both the erosion of rational agency and the diffusion of institutional responsibility indicate that the expectation to take responsibility for consequences is not being met. To restore moral responsibility in algorithmic medicine, transparency and governance safeguards are essential. Clinicians need to be able to comprehend and justify the recommendations they use, as rational accountability is necessary for ethical action in a Singerian sense. Institutions, on their part, should establish traceable chains of responsibility that connect data design, model deployment, and clinical application, ensuring that every predictable outcome has a clear moral owner. Implementing these measures would translate Singer's second principle into specific requirements for the design and oversight of AI systems.

Additionally, the first principle is compromised when algorithmic models fail to consider structural disparities. Research indicates that AI tools trained on biased datasets often underrepresent or misclassify patients from racial minorities, lower socioeconomic backgrounds, or rural areas (Davis et al., 2017; Brajer et al., 2020).⁷ The well-known study by Obermeyer et al. provides a clear example: using prior healthcare spending as a proxy for clinical need, the algorithm consistently underestimated the health risks of Black patients—not because their needs were lower, but because unequal access to care had reduced prior utilization (Obermeyer et al., 2019). This demonstrates how data infrastructures can skew moral equality. When institutional or economic factors are mistaken for true indicators of need, bias can reinforce existing inequalities rather than address them. Additionally, bias may occur at the measurement stage, as symptoms are recorded with less accuracy, and these disparities may be

⁷ For instance, studies by Davis et al. revealed that predictive models exhibit “calibration drift” over time: while discrimination levels might remain stable, models tend to overestimate risk as patient populations and care practices change. Similarly, a study by Bajer et al. validated an in-hospital mortality prediction model across three independent hospitals. Although the model's performance remained high, it highlighted that models trained in one location may lose some accuracy when applied elsewhere, emphasizing the difficulty of maintaining external validity and equitable performance across different institutions.

further exacerbated by AI technologies (Norori et al., 2021). Bias might also arise during the deployment of models, where those trained in one context operate under mismatched thresholds or resource constraints, creating new imbalances.

From a Singerian perspective, rectifying these imbalances is a moral obligation. Equal consideration of suffering demands systematic bias audits, diversification of datasets to reflect the full range of patient experiences, and adjustments to predictive thresholds across different populations. These actions would help ensure that algorithmic systems do not reinforce institutional inequities but instead uphold Singer's principle that every instance of suffering deserves equal moral consideration. Incorporating these requirements into model evaluation would transform Singer's ethic of impartial concern into concrete, measurable design standards for AI in end-of-life care.

These safeguards demonstrate how an ethic informed by Singer can guide the responsible design and governance of AI in end-of-life medicine. By translating Singer's three commitments—reflective preference, accountability for consequences, and equal consideration of suffering—into operational criteria like transparency, deliberative space, and bias correction, this framework bridges Singer's ethical principles with the practical needs of clinical decision-making.

Conclusion

This article illustrates that while artificial intelligence (AI) can enhance procedural and consequentialist aspects of Singerian ethics—such as anticipating clinical outcomes, improving deliberation conditions, and supporting rational decision-making—it cannot act as a moral agent. Predictive models may serve as useful tools that contribute to Singer's criteria, but they ultimately remain technical aids that shape moral judgment rather than exercise it. Our analysis uncovers some ambivalences regarding their role in assisting human deliberation. On one hand, the Jvion and UPHS models demonstrate how AI can initiate clinical conversations earlier, improve the recognition of impending suffering, and help articulate patient preferences. On the other hand, their reliance on statistical functioning can limit deliberative autonomy, diffuse responsibility, and perpetuate morally unacceptable asymmetries from a Singerian perspective. These limitations risk creating an ethics devoid of subjects, rationality that lacks deliberation, and humanism that fails to recognize moral agency.

Based on these observations, we propose a set of normative safeguards to guide the governance of AI in end-of-life care. These include interpretable and transparent recommendations, traceable chains of responsibility, and systematic audits for bias paired with efforts to diversify datasets. These measures embody the essence of Singer's ethical legacy: a preference is never just a statistical datum, an interest cannot be merely a correlation, and a life should never be reduced to a variable.

Acknowledgment

The Sekyra Foundation funds the open access publication of the article in the journal.

References

- ABEJAS, A. G., SANTOS, D. G., LEITE COSTA, F., CORDERO BOTEJARA, A., MOTA-FILIPPE, H. & SALVADOR VERGÉS, À. (2025): Ethical Challenges and Opportunities of AI in End-of-Life Palliative Care: Integrative Review. In: *Interactive Journal of Medical Research*. [online] [Retrieved May 14, 2025]. Available at: <https://pubmed.ncbi.nlm.nih.gov/40302210/>
- ANDRONE, M. (2017): Some Considerations on Peter Singer's Practical Ethics. In: *LUMEN Proceedings*, 1, pp. 34–43.
- BRAJER, N., COZZI, B., GAO, M., NICHOLS, M., REVOIR, M., BALU, S., FUTOMA, J., BAE, J., SETJI, N., HERNANDEZ, A. & SENDAK, M. (2020): Prospective and External

Evaluation of a Machine Learning Model to Predict In-Hospital Mortality of Adults at Time of Admission. In: *JAMA Network Open*, 3(2). [online] [Retrieved February 5, 2020]. Available at: <https://pubmed.ncbi.nlm.nih.gov/32031645/>

BUCKLE, S. (2005): Peter Singer's Argument for Utilitarianism. In: *Theoretical Medicine and Bioethics*, 26(3), pp. 175–194.

DAVENPORT, T. & KALAKOTA, R. (2019): The potential for artificial intelligence in healthcare. In: *Future Healthcare Journal*, 6(2), pp. 94–98.

DAVIS, S. E., LASKO, T. A., CHEN, G., SIEW, E. D. & MATHENY, M. E. (2017): Calibration drift in regression and machine learning models for acute kidney injury. In: *J. Am. Med. Inform. Assoc.*, 24(6), pp.1052–1061.

FERLITO, B., SEGERS, S., DE PROOST, M. & MERTES, H. (2024): Responsibility Gap(s) Due to the Introduction of AI in Healthcare: An Ubuntu-Inspired Approach. In: *Science and Engineering Ethics*, 30(34). [online] [Retrieved August 1, 2024]. Available at: <https://doi.org/10.1007/s11948-024-00501-4>

FLORIDI, L. & COWLS, J. (2019): A Unified Framework of Five Principles for AI in Society. In: *Harvard Data Science Review*, 1(1). [online] [Retrieved July 1, 2019]. Available at: <https://hdr.mitpress.mit.edu/pub/10jsh9d1/release/8>

FOGEL, A. L. & KVEDAR, J. C. (2018): Artificial intelligence powers digital medicine. In: *NPJ Digit. Med.*, 1(5). [online] [Retrieved March 14, 2018]. Available at: <https://www.nature.com/articles/s41746-017-0012-2>

GAJRA, A., ZETTLER, M. E., MILLER, K. A., FROWNFELTER, J. G., SHOWALTER, J., VALLEY, A. W., SHARMA, S., SRIDHARAN, S., KISH, J. K. & BLAU, S. (2021): Impact of Augmented Intelligence on Utilization of Palliative Care Services in a Real-World Oncology Setting. In: *JCO Oncol. Pract.*, 19(1), pp. 80–88.

JOTTERAND, F. & BOSCO, C. (2020): Keeping the “Human in the Loop” in the Age of Artificial Intelligence. In: *Science and Engineering Ethics*, 26(5), pp. 2455–2460.

MANZ, C. R., CHEN, J., LIU, M., CHIVERS, C., REGLI, S. H., BRAUN, J., DRAUGELIS, M., HANSON, C. W., SHULMAN, L. N., SCHUCHTER, L. M., O'CONNOR, N., BEKELMAN, J. E., PATEL, M. S. & PARIKH, R. B. (2020): Validation of a Machine Learning Algorithm to Predict 180-Day Mortality for Outpatients with Cancer. In: *JAMA Oncology*, 6(11), pp. 1723–1730.

MITTELSTADT, B., ALLO, P., TADDEO, M., WATCHER, S. & FLORIDI, L. (2017): The Ethics of Algorithms: Mapping the Debate. In: *Big Data & Society*, 3(2). [online] [Retrieved November 1, 2016]. Available at: <https://ssrn.com/abstract=2909885>

NIKOLOUDI, M. & MYSTAKIDOU, K. (2025): Artificial Intelligence in Palliative Care: A Scoping Review of Current Applications, Challenges and Future Directions. In: *American Journal of Hospice and Palliative Medicine* (Online ahead of print), [online] [Retrieved July 1, 2025]. Available at: <https://pubmed.ncbi.nlm.nih.gov/40598879/>

NORORI, N., HU, Q., AELLEN, F. M., FARACI, F. D. & TZOVARA, A. (2021): Addressing bias in big data and AI for health care: A call for open science. In: *Patterns*, 2(10). [online] [Retrieved October 8, 2021]. Available at: <https://pubmed.ncbi.nlm.nih.gov/34693373/>

OBERMEYER, Z., POWERS, B., VOGELI, C. & MULLAINATHAN, S. (2019): Dissecting racial bias in an algorithm used to manage the health of populations. In: *Science*, 336(6464), pp. 447–453.

O'NEIL, C. (2016): *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. New York: Crown Publishing.

PAUER-STUDER, H. (1993): Peter Singer on Euthanasia. In: *The Monist*, 76(2), pp. 135–157.

RAJKOMAR, A., OREN, E., CHEN, K., DAI, A. M., HAJAJ, N., HARDT, M., LIU, P. J., LIU, X., MARCUS, J., SUN, M., SUNDBERG, P., YEE, H., ZHANG, K., ZHANG, Y., FLORES, G., DUGGAN, G. E., IRVINE, J., LE, Q., LITSCH, K., MOSSIN, A., TANSUWAN, J., WANG,

- D., WEXLER, J., WILSON, J., LUDWIG, D., VOLCHENBOUM, S. L., CHOU, K., PEARSON, M., MADABUSHI, S., SHAH, N. H., BUTTE, A. J., HOWELL, M. D., CUI, C., CORRADO, G. S. & DEAN, J. (2018): Scalable and accurate deep learning with electronic health records. In: *NPJ Digit. Med.*, 1(18) [online] [Retrieved May 8, 2018]. Available at: <https://www.nature.com/articles/s41746-018-0029-1>
- RAJKOMAR, A., DEAN, J. & KOHANE, I. (2019): Machine Learning in Medicine. In: *New England Journal of Medicine*, 380(14), pp. 1347–1358.
- SAUERBREI, A., KERASIDOU, A., LUCIVERO, F. & HALLOWELL, N. (2023): The impact of artificial intelligence on the person-centred, doctor-patient relationship: some problems and solutions. In: *BMC medical informatics and decision making*, 23(1). [online] [Retrieved April 20, 2023]. Available at: <https://pmc.ncbi.nlm.nih.gov/articles/PMC10116477/>
- SINGER, P. (1994): *Rethinking Life and Death. The Collapse of Our Traditional Ethics* (First U.S. Edition). New York: St. Martin's Press.
- SINGER, P. (2003): Voluntary Euthanasia: A Utilitarian Perspective. In: *Bioethics*, 17(5–6), pp. 526–541.
- SINGER, P. (2011): *Practical Ethics* (Third Edition). New York: Cambridge University Press.
- SINGER, P. (2016): *Ethics in the Real World*. Princeton: Princeton University Press.
- STAFIE, C. S., SUFARU, I.-G., GHICIUC, C. M., STAFIE, I.-I., SUFARU, E.-C., SOLOMON, S. M. & HANCIANU, M. (2023): Exploring the Intersection of Artificial Intelligence and Clinical Healthcare: A Multidisciplinary Review. In: *Diagnostics*, 13(12), pp. 1–37.
- SYNOWIEC, J. (2019): Metaethical grounds for changing the understanding of values in Peter Singerian ethics (in the context of his reflection on the value of human life). In: *Kultura i Wartości*, (28), pp. 47–64.
- TOPOL, E. (2019): *Deep Medicine: How Artificial Intelligence Can Make Healthcare Human Again*. New York: Basic Books.
- VAN DEN BEUKEN-VAN EVERDINGEN, M. H. J., DE RIJKE, J. M., KESSELS, A. G., SCHOUTEN, H. C., VAN KLEEF, M. & PATIJN, J. (2007): Prevalence of pain in patients with cancer: a systematic review of the past 40 years. In: *Ann. Oncol.*, 18(9), pp. 1437–1449.