



Responsibility assignment won't solve the moral issues of artificial intelligence

Jan-Hendrik Heinrichs^{1,2}

Received: 21 October 2021 / Accepted: 31 December 2021
© The Author(s) 2022

Abstract

Who is responsible for the events and consequences caused by using artificially intelligent tools, and is there a gap between what human agents can be responsible for and what is being done using artificial intelligence? Both questions presuppose that the term ‘responsibility’ is a good tool for analysing the moral issues surrounding artificial intelligence. This article will draw this presupposition into doubt and show how reference to responsibility obscures the complexity of moral situations and moral agency, which can be analysed with a more differentiated toolset of moral terminology. It suggests that the impression of responsibility gaps only occurs if we gloss over the complexity of the moral situation in which artificial intelligent tools are employed and if—counterfactually—we ascribe them some kind of pseudo-agential status.

Keywords Responsibility · Artificial intelligence · Ethical terminology

1 Introduction

Contemporary debates about the moral issues of generating and using artificially intelligent systems strongly focus on the question of who is responsible for these system's behaviour. A core suggestion is that there is a gap between what moral agents can be responsible for [24], or what they can suffer retribution for [9] and the behaviour of AI systems. On the other hand, several authors either claim that there is no gap between our practice of responsibility ascription and the behaviour of AI [20, 38] or make suggestions on how to modify our practice of responsibility ascription to bridge these gaps [27]. This article argues that we are looking in the wrong direction: responsibility is not a term suited to solve moral problems surrounding AI. Componential analysis of ‘responsibility’ as provided by many in the debate should be a reason to drop the unanalysed, general term and use the components instead, at least in ethical theorising.

Attempts to tackle questions regarding the use of artificially intelligent systems exist in a variety of contexts,

ranging from autonomous driving [28, 29] and the military use of robots [36] to epistemic responsibility for artificially intelligent diagnostic systems [14]. The results are equally diverse. They range from a complete rejection of the use of artificially intelligent systems due to the impossibility of reliably attributing responsibility, to a rejection or significant revision of the concept of responsibility because it is not suitable for clarifying questions of robot ethics [23].

The diagnosis that technological developments challenge our legal and moral practise, especially that of ascribing responsibility, has predecessors in several previous technologies, starting with the steam engine. More recent examples include the car, and the issue has become most prominent in nuclear energy and genetic technologies. Gifford for example discusses how early automation in automated looms and railroads prompted “the replacement of the preexisting strict liability tort standard with the negligence regime” ([11], p. 1) and how the widespread introduction of automobiles made it necessary to generate “»financial responsibility laws« that required automobile owners to either purchase insurance or to provide proof that they had sufficient financial resources to pay claims” ([11], p. 41 f.). It has been pointed out in these and other contexts that there is a discrepancy between our established practice of attributing responsibility and the moral requirements of current technical developments. This discrepancy or gap is caused by at least two factors: It is caused firstly by the increasingly

✉ Jan-Hendrik Heinrichs
j.heinrichs@fz-juelich.de

¹ Institut für Ethik in den Neurowissenschaften (INM-8),
Forschungszentrum Jülich, Jülich, Germany

² Faculty for Philosophy, RWTH Aachen, Aachen, Germany

complex social and institutional contexts in which actions and their consequences can hardly be attributed to a single human actor [18]. And it is caused secondly by the development of technical systems that increasingly determine the conditions, courses and consequences of actions ([37], p. 107 ff.). The extreme point of both factors can be found in the use of artificially intelligent systems. AI systems are obviously technical support systems for action that strongly influence said action's course. However, they are also—less obviously—integrated into complex social and institutional contexts, insofar as their production, training and even the use of AI systems involves multiple actors in different institutional dependencies.

Saying that AI is a technical support system for human action narrows down the term's meaning for the current purpose. It is intended to refer to learning systems of narrow or at most moderate generality. Such systems can and often do show superhuman performance in a limited number of particular tasks, but typically do not generalize across tasks [7]. Such systems heavily depend on human intervention during design, training, and validation. They can, however, vary wildly in the need for human intervention in the way they are employed. Some systems such as language transformers heavily depend on human input, others such as automated picture recognition systems can work with automated sensor data acquisition and without human intervention for quite some time. While both remain support systems for human action, one is a close support system, whereas the other one is a distant one. In contrast, systems with human-level ability to generalize and potential superintelligences will probably have to be treated as genuine authors of actions in the full philosophical meaning of the term and thus cease to be a mere support system to human action [2]. Thus, they require a different treatment than the one I propose in this paper.

2 Why responsibility for artificial intelligence is controversial

Controversies about the responsibility for actions with the help of AI systems currently abound, a prime example being the discussion about the use of military robots. These applications are the focus of political attention in the campaign to outlaw killer robots and the core topic of the by now canonical example [27, 38] of philosophical articles on the responsibility for the behaviour of AIs: Robert Sparrow's landmark article, 'Killer Robots'. In this text he asks "who we should hold responsible when an autonomous weapon system is involved in an atrocity of the sort that would normally be described as a war crime" ([36], p. 62). Sparrow's killer robot thus is a case of distant support systems, i.e., a system comprising robotic parts which has little to no human intervention during its employment. He tries to show that of

all possible candidates for responsibility, none are fit to take responsibility for actions carried out with the help of semi-autonomous military robots.

Sparrow seems to assume that responsibility always belongs to an individual, and that it depends on the ability to predict or control the event in question. Because of the lack of control over the actions of an autonomous weapon system, Sparrow claims that neither the manufacturer, nor the programmer, nor the commanding officers are in a position to take responsibility for the system's behaviour if it amounts to a war crime. Because responsibility can only be assumed by those who can be punished, the machine itself is not suitable for assuming responsibility. Because it neither suffers nor can its behaviour be corrected by punishment, let alone by praise or blame, it is simply not possible or appropriate to hold the machine responsible. There is what has later come to be called a gap between the events that require a moral evaluation and the actions that allow no ascription of responsibility, a responsibility gap.

The term 'responsibility gap' has probably been coined in 2004 by Andreas Matthias, who argued that AI systems would create responsibility gaps that could not be bridged by the established notion of responsibility: "If we want to avoid the injustice of holding men responsible for actions of machines over which they could not have sufficient control, we must find a way to address the responsibility gap in moral practice and legislation" ([24], p. 183).

Sparrow makes use of the concept—if not of the term—of a responsibility gap. Holding anyone responsible for the action of an autonomous weapon system would, as introduced by Matthias, be unjust because of their lack of control over systems behavior. Sparrow's overall conclusion is straightforward: because it is not possible to reasonably hold anyone responsible for war crimes committed by an autonomous weapon system, but the possibility of holding specific individuals responsible for war crimes is a necessary condition for a just war, the use of autonomous weapons systems in war cannot be justified ([36], p. 66). They should not be used and probably not be built in the first place.

A similar argumentative structure can be observed in a completely different type of document, namely in industry guidelines for AI engineers. The British Standards and the IEEE Guidelines both note that there is uncertainty in the prediction and control of some AI-based systems. Both argue that this uncertainty results in a gap between the responsibility of engineers (and users) and the real consequences of the machines' behaviour. In this, they concur with Sparrow. But unlike him—obviously—they do not call for abandoning the technology. Neither do they, however, simply accept the epistemic gap or call for abandoning our practice of ascribing responsibility whenever a person cannot predict and control a device. Rather, they call on engineers to design the systems in question in a way which allows for prediction

or at least retrodiction of the system's behaviour (British Standard for Robots and robotic devices, p.5; IEEE Global Initiative for Ethical Considerations in Artificial Intelligence and Autonomous Systems, p. 90 ff.)

The guidance of these institutions serves not only the purpose of creating the best products for the user but also to protect manufacturers and engineers from legal and moral blame and liability for undesired behaviour of their products. Philosophical authors as well as stakeholder organisations seem to doubt that our contemporary practice of attributing responsibility can solve the moral and legal problems that arise from the use of AI systems. Either unpredictable or uncontrollable AI systems must be dispensed with altogether, or they must be adapted to such an extent that they fit back into our practice of responsibility.

In the following, I will argue that these authors are completely right in claiming that our contemporary practice of responsibility ascriptions does not solve the moral issue of how to respond to morally problematic actions which have been committed with the support of or by AI systems. However, this is not because artificially intelligent systems change the structure of action or the requirements of moral evaluation in some fundamental way, but because the term 'responsibility' is not suited for detailed moral analysis of complex actions.

The claim that these moral problems of AI-assisted actions cannot be solved given our current practices of ascribing responsibility is due to the fact that 'responsibility' is a bundle term with too many internal tensions. The internal tensions of the term are adumbrated by the fact that it allows for a certain dismantling of distinctions, as can be seen in a prominent application to artificial intelligence: „There are many ways agents are held responsible for their actions within legal systems, corresponding to different available means of punishment. Human agents have historically been punished in a variety of ways: through the infliction of pain, social ostracism or banishment, fines or other confiscation of property, or deprivation of liberty or life itself. Debates over whether it makes sense to hold (ro) bots accountable for their actions often center on whether any of these traditionally applied punishments makes sense for artificial agents.” ([41], p. 208, my emphasis).

This short passage showcases how several distinct moral properties and relations—here liability, punishment, and accountability—are gathered under the term 'responsibility'. By trying to solve moral issues of artificial intelligence with this term, we try to solve several highly specialised tasks with one rather a crude tool.

3 'Responsibility'—etymological structure and functional role

The concept 'responsibility' has originally been modelled in analogy to the interaction between a judge and a defendant. Its first philosophical occurrences in Reid¹ and Kant's *Metaphysics of Morals* refer to the person giving an account of her actions. This situation at first appears to be morally pretty straightforward. One person has to answer for her previous actions to another. This prototypically clear distribution of obligations and entitlements is what guides the use of the term. As such, 'responsibility' is a term used for the coarse-grained moral task of identifying who has to give reasons for his or her actions to whom [4, 15, 33].

Contrary to appearance, even this seemingly straightforward moral situation is characterised not by one but by a variety of moral relations between the persons involved. It involves mutual recognition, mutual attribution of moral status and states, and several concrete moral demands and obligations such as duties that one person has towards the other claims that she can make on him, the justification that she owes him, the reparation that may be due, and so on. These moral relations can and need to be analysed with the complex conceptual toolbox of ethical theory, involving terms such as attributability, accountability, liability, culpability, individual and group agency, etc.

Using the term 'responsibility' is common in many contemporary moral discourses about AI as exemplified above in Sparrow's article, Wallach and Allen's book, the British standards, and the IEEE guidelines. The following analysis will try to show that this is using the term beyond its capability, i.e. beyond the coarse grained task of identifying a judge and a defendant.

4 Prototype and variations of attribution of responsibility

In the current literature, 'responsibility' is differentiated into different meanings of the term and into different relations it refers to. According to the first type of differentiation going back to Hart's seminal work, being responsible for something can for example mean that someone has a duty to do something, or that he or she is liable for something [12, 13]. The second differentiation distinguishes the relation of responsibility, i.e. who can be responsible for what to whom, according to which norm and for whose sake [22, 30]. These two differentiations are drawn together in

¹ Reid in his *d's Essays on the Active Powers of the Mind* uses the term 'accountability', but his argument is taken up by Mill who uses 'responsibility' as a synonym ([25], p. 7).

newer analysis ([39, 40]. The term 'responsibility' thus refers to a broad cluster of practices that answer the very general question "In what way (1) is who (2) responsible for what (3) to whom (4) according to what norm (5) and for whose sake (6)?" Each interrogative term in this question can have different answers depending on the context, but some answers form the core of our way of speaking, others are variations or even deviations.² In the following, I will provide some detailed support for the thesis that the prototypical answer to all six interrogative words is modelled on the simplified dialogical situation and refers to a specific component of the bundle term 'responsibility'. I'll show that, in contrast, analysing moral issues of artificial intelligence requires a different model than the dialogical situation and the employment of a multitude of moral relations, not just one. To show how this is the case I will focus on the answers to the first three interrogative terms.

There is another tradition of analysing the concept of responsibility, which dominates the writings of Köhler and colleagues [20] as well as Tigard [38]. Tigard in particular, following David Shoemaker, distinguishes the concept of responsibility into accountability, attributability, and answerability, rejecting a gap in the respective practices of responsibility: "Like our accountability practices toward fellow humans, we can hold AI to account by imposing sanctions, correcting undesirable behavioral patterns acquired, and generally seeing that the target of our responses works to improve for the future—a bottom-up process of reinforcement learning." ([38], p. 604).

Both are versions of pluralism concerning moral responsibility, and they overlap in some of the components of responsibility. However, there are mismatches in nomenclature—e. g. what Tigard calls 'answerability' is here subsumed under 'accountability'—and in depth of analysis. While Tigard uses a tripart distinction of the types of responsibility into attributability, accountability and answerability, the present differentiation goes beyond an analysis into types and identifies slightly more types, too.

4.1 The different answers to "Responsible in what way?"

"Responsible in what way?" can have several different descriptive and normative answers. In the following, I will only discuss the normative versions. It should, however, be mentioned that the prototypical descriptive answer refers to being a cause: rain is responsible for the road being wet. This

is a common way of speaking and is presupposed in many (not all) normative judgements of responsibility.

The normative dimensions of the term 'responsibility' are given by three retrospective dimensions accountability, praise- and blameworthiness, liability on which I'll focus in the following and two prospective ones, obligation and virtue [39, 40]. Most of these picks out a separate dimension of the dialogical model of 'responsibility'. If one person approaches another on grounds of her actions, she can ask for justification, praise, or blame the person or assign some kind of liability.³

The prototypical answer to 'responsible in what way' is, in my opinion, 'accountable'. Here is why: Being accountable is a necessary precondition of the further moral relations subsumed under 'responsibility'. To be responsible in the sense of being accountable means to be able and to be expected to give reasons for, i.e., to justify past actions. Someone who is accountable in this sense can be asked about their actions in direct interaction and answer, i.e. give an account why they said or did something (cf. [3]). Being asked for and giving reasons for one's behaviour is the entrance ticket to every form of rule-governed interaction, a sufficient (and maybe necessary) condition for taking part in the discourse. Being able to give justifications and getting challenged to do so is, moreover, the core characteristic of rational agents, i.e. of beings who can act for reasons [4, 33]. This is why responsibility in the sense of accountability and ability to justify oneself/one's actions is sometimes also equated with the ability to act, sometimes even with personhood. Moral discourse necessarily requires the ability for justification, for giving and taking reasons [15], the ability to receive or give praise and blame or to assign or accept liability fully depend thereon.

There is a danger of confusing accountability and explainability. Being accountable implies being the kind of person

² A pragmatic prototype conception of concepts will be used here, i.e. I will understand concepts as linguistic tools and their meaning as determined by a prototypical use and variations [21, 31, 43].

³ Another prominent approach to further analyse the 'in what way' bundle of 'responsibility', the one used by Tigard, is introduced in [42]. Watson distinguishes between accountability and attributability, with accountability being used more or less as above, attributability as the ability to be subjects of moral appraisal (p. 240). Shoemaker [34] in turn focusses on the distinction between attributability, answerability and accountability. According to him attributability is the ability to get attributed a predicate which expresses one's practical attitudes, answerability is the ability to give the reasons with which one justified one's conduct and accountability is the "capacity to recognize and appreciate the demands defining the various relationships as reason-giving" (p. 631). For the present purpose, I will leave attributability to one side. This is not to say it is not an important dimension of our moral practice and the reactive attitudes we have towards others. However, it is more of a precondition of holding someone responsible in any meaning of the term rather than one of the components of responsibility. Thus, it is not an answer to the question in what way someone can be held responsible, rather it is an answer to the question who can be held responsible and thus will be discussed below.

who can provide reasons for one's behaviour, not explaining its causal antecedents. When asked "Why did you do that?" the explanation "the combination of internal states $x_{1..n}$, input $y_{1..n}$ and the mechanisms $z_{1..n}$ operating on the former caused this behaviour as an output" is not a valid reply. It might at best be the start of one but lacks any justificatory relevance. It fails to place the action in question in a net of mutually accepted norms of conduct.

As mentioned, the moral relation of praise- or blameworthiness closely depends on that of accountability. It does so because praise and blame have—at least among other things—the function of correcting behaviour. If a being cannot respond to praise or blame with changes in behaviour, it is not a candidate for them (apart from providing an outlet for the blamer's anger). Modifying one's own behaviour on the basis of praise and blame, in turn, requires the ability to act for reasons. This is different from the mere ability to modify one's behaviour on the basis of punishment and reward. The former is possible for discursive beings only, the latter for discursive as well as non-discursive beings. Responsibility in the sense of praiseworthiness or blameworthiness makes a being an active member of the moral community in the sense that it can apply moral reasons to its actions.

Related is the understanding of 'responsibility' as liability, i.e., that individuals can be expected to compensate or apologise for the consequences of their behaviour [6]. The extent of an agent's liability typically depends on particular norms and their goal. It will for example make a difference whether norms are designed for the aim of retribution or for the aim of betterment. In the former case, the dialogical situation might be a suitable model, in the latter case a more complex social arrangement needs to stand in as a model. In any case, however, taking an agent to be liable presupposes a certain set of mental capacities, without which the application of the norm becomes pointless ([40], p. 25). This presupposition of certain mental capacities is what Sparrow referred to when he claimed that autonomous weapon systems are not suited to take responsibility because they do not suffer and thus cannot be punished. Even liability is closely connected to the dialogical model for the term 'responsibility'. We usually talk about agents being liable only if they can in some minimal way grasp the norms for the breach of which they are liable. That implies that these norms have at least to be possible reasons for action for this agent. They must be able to relate their own behaviour to these norms when asked to justify their conduct. It does not suffice if the agents are merely conditioned according to these norms.

In the debate about how someone—anyone—could be responsible for the behaviour of artificially intelligent agents, two answers dominate: first, a deviant form of accountability for the system itself. As already mentioned, industry standards recommend that artificially intelligent systems be

designed in such a way that their actions can be traced at any time. The same idea lies at the heart of the explainable AI movement. However, the system does not give reasons for actions, but is designed to make the causes for its behaviour transparent—or worse: allows for another—often equally opaque—system to make the causes for the original behaviour transparent (cf. [38], p. 602, [45]). As noted above, the explanation of causes is not a justification with reasons. At best, the people who created or used the system are enabled to consult its logs in their reasons for continuing to operate, modify or shut down the system [32]. Thus, this deviant case runs into a dilemma: Either the information about the causes of behaviour is not an answer to the normative "Why did you do this?" question. Then making the system explainable tracks causation, which is at best a nonnormative version of 'responsibility', not a case of accountability. Alternatively, the information is taken to be an answer to a normative question, but then it is a question aimed at the engineers, providers, and users of the system: "Why was the system trained on these data?", "Why were these outputs during training considered correct or acceptable?", "Why was the system employed in this environment and for this purpose?" etc. Then this is not accountability of the system but of the people creating, selling, or employing it.

Second, there is a regular discussion about who is liable for the events in which artificially intelligent systems are involved and their consequences. One proposal is a form of group liability. The idea is to generate a joint liability by producers, distributors, users, etc. of AI systems, by generating a money fund from which damages caused by the system—and possibly fines incurred—are to be compensated [1]. This is a borderline case of responsibility as liability for two reasons. First, it is an extremely narrow and purely legal scope of liability. It lacks any regard for genuinely moral reactions and for behavioural change. Second, and this point is closely related to the next interrogative term, it involves a deviant answer to the question "Who is liable for actions involving artificially intelligent systems?" It suggests a form of group responsibility, but even that would be a deviant case, simply because the group in question does not fit the standard conditions for any type of group agency and responsibility. There is barely a joint context of action, no self-identification of an acting body or anything similar [8]. The only connection between the members of this liable group is some involvement with the product.

From the description, it should have become clear that applying the term 'responsibility' in both cases obscures more of the moral relations than it reveals. This is not to say that the solutions to specific moral questions of dealing with AI-assisted actions are bad. They are not. But they are oversimplified or misdescribed by the term 'responsibility'. Neither is the pursuit of explainable AI simply a case of accountability or of any other normative version of

responsibility nor is the complex legal constellation of joint liabilities for consequences of AI employment adequately described by the undifferentiated ‘responsibility’. Both questions, that concerning accountability and that concerning liability, involves not just one way of being responsible, but several; and they both refer to a situation which simply cannot be modelled on a dialogical situation but always involves more parties. That leads us to the second interrogative term, the ‘who’ of responsibility.

4.2 The different answers to "Who is responsible?"

The subject of responsibility which the interrogative clause ‘who is responsible?’ inquires about is the responsible person from whom a reaction or response is required. The prototypical subject of responsibility is a typical interlocutor in rule-governed discourses. These are the individuals we usually turn to when we ask why someone did or said something, blame them or even demand compensation or punishment.

This characteristic of being an interlocutor is mirrored in what has come to be called attributability, i.e. the ability to be attributed moral properties independently of external ascription [38, 42]. Another human being can correctly be attributed practical reasons, reactive attitudes, or virtues and vices. The moral properties of algorithms on the other hand are not attributes of the models themselves, but of their being designed in a certain way and being placed in a specific use context for a particular purpose. It is to the combination of (human) design, employment, and purpose that moral properties can be attributed, not to the algorithm itself.

The prototype of typical interlocutors obviously does not exhaust the range of possible subjects of moral relations. The range of possible subjects of responsibility—in any meaning of the term—is typically described by a set of characteristics which need to be correctly attributable to the system in question, such as rationality, receptivity to reasons, receptivity to correction through praise and blame or punishment, etc. Since many of these properties come in degrees, it is plausible that being able to stand in the moral relation in question is either itself a gradual or a threshold phenomenon. Nor must all variations of responsibility be alike in this regard, for example, the legal practice seems to know cases of full responsibility (liability, ability to give informed consent) given a certain threshold of competency, while at least educational practice sees responsibility (accountability) as a gradual phenomenon that adolescents acquire successively. Depending on which characteristics at which threshold are constitutive for the ability to stand in a given moral relation, it is a matter of dispute whether, for example, some animals, human infants, people with severe mental disabilities or artificially intelligent systems can be part of the relation

in question. On the other hand, moral relations can tie together more than just two individuals. For example, one agent can be morally obligated towards a number of different individuals within the same relation, e.g., the obligation to respect one’s parents, both of them. There is also the special case of group responsibility. Admittedly it is an as yet undecided question in philosophy whether moral rights or obligations of groups, which are firmly established in the legal context, have a moral equivalent that cannot be traced back to the rights and obligations of individual group members (cf. [8]).

In the debate about attributing responsibility for the behaviour of artificially intelligent systems, two marginal answers to "Who is responsible?" dominate. The first of these has been discussed in science fiction scenarios since at least Samuel Butler’s novel *Erewhon* [5]: the artificially intelligent system itself. As seen above in Sparrow’s example, some philosophers mention this option, only to dismiss it immediately. This seems to be the reaction of most authors in this field: artificially intelligent systems are rarely seriously discussed as bearers of responsibility. Others take it slightly more serious and metaphorically talk of the duty of a programme to solve certain tasks, or to adhere to certain limits. This usage is even found in the prominent philosophical contribution by Wallach and Allen, who speak of designing systems in such a way that they do not transgress concrete norms of action. But in the end, they admit that artificially intelligent systems are at best deviant cases of responsibility [41]. Some very few pick a special aspect of the term ‘responsibility’, namely the descriptive use as in ‘causally responsible’ and on this basis ascribe responsibility to artificially intelligent systems [10].

The second major strand of attributing responsibility for AI systems constructs a form of group or systems responsibility for the consequences of actions carried out with the help of artificially intelligent systems. As mentioned above, Beck for example constructs a rough analogy to the concept of a legal person and proposes to generate a special legal status for artificially intelligent systems. It is intended to consist of the partial responsibilities of all parties involved in the creation and use of the artificially intelligent system: "This legal person for robots would only be the bundling of all the legal responsibilities of the different parties (users, sellers, producers, etc.). This bundling is actually the main reason why a new classification for these machines is necessary." ([1], p. 479). In philosophy, a similar suggestion has been made by Sven Nyholm, who argues that not much of a responsibility gap remains if we understand responsibility as a relation property of the network or system of agents and tools involved in a certain behaviour or event, what he calls a human–machine collaborations [27]. But as Nyholm himself admits, his suggestion involves a significant re-conceptualisation of ‘responsibility’ and ‘agency’. In the terms introduced here, it drives the use of the term further

apart from its etymological origin instead of trying to salvage the dialogical constellation.

Nyholm introduces several illuminating differentiations in the terms of agency and responsibility. In particular, he differentiates agency into individual versus collaborative, into domain specific versus domain independent as well as into basic versus principled and into supervised versus deferential and responsible ([27], p. 1207 f.). Contemporary and near future machine learning systems which Nyholm exemplifies with automated cars and military robots have “supervised and deferential agency [...] of a collaborative type” ([27], p. 1211). Given that the agency of artificially intelligent systems is collaborative, defers to and is supervised by a human in the collaboration, Nyholm can assign responsibility to the pair or group of collaborators but place the locus of responsibility with the human part. Admittedly, Nyholm can thereby show that there is not much of a responsibility gap, but he does so by showing that the humans designing, training, selling, using artificially intelligent systems are the responsible parties: “The most difficult questions here instead concern what humans are most responsible for any potential bad outcomes caused by their robot collaborators.” ([27], p. 1214).

What is more, Nyholm succeeds in bridging the gaps in moral relations because he further differentiates the concept of responsibility. The humans designing, training, selling, using artificially intelligent systems are not just plainly responsible, but rather we should try “to determine who is responsible for what aspects of the actions the car performs” ([27], p. 1214). Thus, Nyholm’s project to vindicate the use of the concept of ‘responsibility’ for artificially intelligent systems only succeeds by (a) differentiating types of responsibility and (b) assigning it to the humans involved with the systems. The underlying idea of Nyholm’s solution to dealing with moral issues raised by the use of AI-systems is quite convincing, namely, to consider them either collaborative tools or scaffolds of human agency [16, 17]. But it is, contrary to appearance, not a solution which relies on the bundle concept of ‘responsibility’.

In general, analyzing moral issues of artificial intelligence with the term ‘responsibility’ which is modelled on a dialogical situation tends to obstruct the view of the complex contexts in which, as Nyholm demonstrates, these issues occur. The moral issues of artificially intelligent systems typically involve more than just two persons, and most of the time there is no neat grouping of the affected and involved persons into collective agents, much less into two parties.

4.3 The different answers to “What is someone responsible for?”

The object of responsibility is the issue addressed in the demand for response, what the subject is responsible for. In

everyday language, very different things can take the place of the object in ‘responsible for ...’, for example, physical objects, actions, states of affairs. A person can be responsible for his car, for the care of his children or for the safety of a workplace. However, many authors think that these different kinds of things a person can be responsible for can be reduced to only one or two kinds, namely actions and possibly the success of actions. The undisputed prototype of the object of responsibility is a previous action.

The person responsible for his or her car is, according to reductive positions, actually responsible for certain actions concerning her car: for regular maintenance, for driving it in a certain, safe way, for parking it only in marked parking spaces and so on. It is an open question whether people can be responsible for the success of actions or only for the action itself: Does a person who is responsible for the safety of a workspace have to remove obstacles, provide guidance, etc., and has fulfilled his responsibility when he has performed all these tasks? Or is she responsible for ensuring that the workspace is safe beyond these actions and has neglected her responsibility if someone is harmed in this workplace, even though all specific measures have been carried out to make the workplace safe? The answer to this question depends on the exact meaning of ‘responsibility’ as introduced above. For example, a person may be legally (and even morally [6]) liable for events that have occurred, even though they have performed all the actions they were obligated to perform.

The discourse of responsibility for behaviours of artificially intelligent systems is fundamentally at odds with this conceptual prototype for two reasons. First, because it moves events into the focus, which definitely are not actions in any but a very loose sense, namely the data processing and output of computer systems. An action is behaviour carried out on the basis of practical reasons, which in turn are composed of some pro-attitude like a wish or desire and some instrumental belief. Even if it is convenient to describe the behaviour of AI-systems as actions, this is merely borrowing the psychological terminology. It might from an intentional stance make sense to ascribe beliefs to AI-systems. However, ascribing them pro-attitudes is at best a derivative use of the term, if not just a figure of speech.

This derivative use has become deeply entrenched in our description especially of reinforcement learning and tree search algorithms. We talk about them maximising a reward function or a preference function, respectively. The term has been imported from behavioural sciences, be it animal behaviour studies or behavioural economics where organisms’ behaviours are explained in the same terms, namely as reward and preference functions. We should, however, differentiate between the notion of maximising a reward or preference function and that of being susceptible to rewards [35] or having preferences,

though these are easily confused. While it might be possible to describe animal behaviour within a given domain as guided by the maximisation of reward functions for the purpose of scientific explanation, this abstraction leaves out quite a bit of relevant psychological information. The same is true for describing behaviour, especially human behaviour, as guided by a preference function. This might be a useful abstraction, but it captures only a minor part of what it means to have a preference. The mere fact that AI-systems can be designed making use of these explanatory models of psychological and behavioural science does not mean that these systems have the same type of state as the organisms so described. It is one of the scientifically exciting aspect of AI-models in science that they do not need to imitate the structure of the phenomenon which they explain [26]. This holds equally true for models of the human mind.

It is exactly for this reason that ascriptions of cognitive and conative states or of agency to AI-systems are carried out with great care to what exactly is being ascribed. Take the examples of Nyholm: “Domain-specific principled agency: pursuing goals on the basis of representations in a way that is regulated and constrained by certain rules or principles, within certain limited domains.” [27], p. 1208, my emphasis) or Floridi: “Agent =_{def.} a system, situated within and a part of an environment, which initiates a transformation, produces an effect, or exerts power on it over time.” ([10], p. 140, my emphasis). Only a few authors would insist that a human agent’s pro-attitudes and beliefs together with the resulting actions are of the same kind as an AI-system’s goals or objectives and representations together with the resulting behavior [44].

The second reason why the discourse of responsibility for behaviours of artificially intelligent systems is fundamentally at odds with the conceptual prototype is because it rarely refers to the real actions involved in the moral issues raised by artificially intelligent systems, namely the actions of individual human agents. This is probably the most disconcerting diagnosis of the current debate. The question whether a software engineer is accountable for the way he programmed and trained an AI system, a manager liable for ordering or selling it, or an officer blameworthy for the order to use an autonomous weapon system is rarely asked, although these are the basic moral phenomena we should inquire about. In contrast, the question people in fact ask is whether any of these individuals can be responsible tout court for an artificially intelligent system or the events brought about by such a system. Several authors wonder whether it is the software engineer who is responsible for an artificially intelligent system or whether the commanding officer is responsible for the behaviour of an autonomous weapon system or to what degree—as a contrast to in what way—each of these is responsible.

5 What does the concept of responsibility do for the moral analysis of the use of artificially intelligent systems?

Two things should have been particularly noticeable in the exposition of the concept of responsibility and its particular use in discourses on artificially intelligent systems: first, that the concept of responsibility contains an internal tension, between the simplified dialogical prototype and the very differentiated practice of attributing responsibility. Second, that the practice of using the term, i.e., of attributing responsibility, draws on the entire set of instruments for answering normative questions, at least in ethics and law. It tries to reconstruct the different moral relations of the dense network of obligations, duties, expectations, liabilities, etc. and model them on the dialogical situation between typical moral subjects. Identifying different components or versions of ‘responsibility’, such as accountability-responsibility or responsibility as duty, is first and foremost a task of isolating distinguishable moral relations between different constellations of agents. Analysing a complex moral situation such as that of actions involving AI-tools according to these distinguishable moral relations is the logical next step. Unsurprisingly, this works better for some novel moral relations and worse for others. However, these analyses are at best misnamed as investigations of responsibility for the actions of AI, when in fact they show that the undifferentiated term ‘responsibility’ itself does no analytical work in this regard.

Let me put it into a rough analogy: the original endeavour (moral analysis) is characterised by a complex set of sub-tasks (identifying different moral relations) for which we need different specialised (terminological) tools. These tools are amongst others: ‘right’, ‘duty’, ‘justification’, ‘liability’, ‘recompensation’, etc. It turned out that we can package a relevant number of these tools together in the multitool ‘responsibility’—much like a swiss army knife. And much as in real multitools, we do not assemble a full version of the specialised tool, but one that fits into the casing. Here, the casing is the dialogical construction of the term ‘responsibility’. Now we analyse that multitool and generate differentiations such as ‘liability responsibility’. But these primarily refer to the multitool with a certain sub-tool extended. Liability responsibility is nothing over and above liability and responsibility as the obligation is nothing but obligation. When our original endeavour takes a critical and partially unfamiliar turn (in the analysis of moral issues of artificial intelligence), we try to stick to our multitool instead of opening the full toolbox, although we encounter gaps in the multitool’s applicability: responsibility gaps.

At no point in the above analysis of the attribution of responsibility did it turn out that the term does more than

the differentiated set of instruments. The moral and legal questions regarding the production and use of artificially intelligent tools by persons in their institutional embeddings can be solved with the more differentiated conceptual tools of the relevant disciplines, it cannot with the multitool ‘responsibility’.

Deborah Johnson has made a similar observation, when she approached our practice of ascribing responsibility to engineers: “Although there is broad consensus that engineers have social responsibilities, what is owed in the name of social responsibility is not well understood.” ([19], p. 85). After analyzing the codes of several engineering societies which predominantly talk in terms of obligations and duties, she suggests settling for an analysis of engineers’ responsibility as accountability, which in turn consists of obligations according to shared norms. Instead of sticking to the blanket identification of the involved parties as ‘engineers’ and ‘the public’ she disentangles who exactly is involved in the social practice of accountability. If this is the path set out for our analysis of responsibility for artificially intelligent systems and their use, we would do better so skip the detour and use the more specialized terminology from the outset. To return to Sparrow’s example: Instead of talking about the responsibility for the events caused by an autonomous weapons system, we should ask for the diverse moral rights and obligations of everybody involved. Here are some: Does a commanding officer have the duty to prevent any action carried out by his troops and with his equipment, which would amount to a war crime? Is she blameworthy if her troops or their equipment are involved in a violation of the codes of just war? Should she feel regret or blame herself? Can she be excused, or her guilt be mitigated if she fulfilled all her obligations of oversight? Can victims or their relatives make claims against her in that case? Is a company liable for damages caused by their goods if there is no fault on the side of the user, in this case the military unit? Is the company obligated to inform their user about the possibility of unforeseen events and the risks thereof? Is a company which produces equipment that is likely to be involved in events that count as war crimes morally blameworthy? Is a military or political leader morally blameworthy for equipping a military with tools which can even under good conditions cause events which would have to be described as war crimes? I have little doubt that all these questions can be answered—without gaps—given some specifics of the cases under scrutiny.

Under which circumstances would the analysis with the full toolbox of ethical terminology result in gaps? I think this question is diagnostic for the debate about responsibility for artificially intelligent systems. The answer seems to be that gaps would occur, if there were a subject of some moral relation, which for some reason cannot stand in this relation, e.g., a subject of obligation which cannot have obligations. The thought behind this is, that artificial intelligence can

be understood as an agent, but it cannot be understood as a subject of moral relations that come with being an agent. I think the problem is with the former and not with the latter part of this thought. Only if we think of artificial intelligence systems as pseudo-agent, do responsibility gaps occur. The disconcerting fact, however, is that at the moment an artificially intelligent system becomes an agent, there are no responsibility gaps anymore. An agent is a person who can be involved in all the moral relations packaged under the term ‘responsibility’. Pseudo- and non-agents, however, simply do not stand in any type of moral relation, they can merely be the subject matter of moral relations between agents. The latter are typically not analyzed with the more general term ‘responsibility’, but with the more specialized components. There is no reason to do otherwise if artificially intelligent tools are being used by these agents.

Funding Open Access funding enabled and organized by Projekt DEAL and the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) – 491111487.

Declarations

Conflict of interest The author declares that there is no conflict of interest.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

1. Beck, S.: The problem of ascribing legal responsibility in the case of robotics. *AI Soc.* **31**(4), 473–481 (2016). <https://doi.org/10.1007/s00146-015-0624-5>
2. Bostrom, N.: *Superintelligence: paths, dangers, strategies*. Oxford University Press, Oxford (2014)
3. Bovens, M.: Analysing and assessing accountability: a conceptual framework1. *Eur. Law J.* **13**(4), 447–468 (2007). <https://doi.org/10.1111/j.1468-0386.2007.00378.x>
4. Brandom, R.: Action, norms, and practical reasoning. *Nous* **32**(Supplement 12), 127–139 (1998)
5. Butler, S.: Erehwon. Trubner & Co, London (1872)
6. Capes, J.A.: Strict moral liability. *Soc. Philos. Policy* **36**(1), 52–71 (2019). <https://doi.org/10.1017/S0265052519000220>
7. Chollet, F.: On the measure of intelligence. <https://arXiv.org/1911.01547>. (2019)

8. Crone, K.: Foundations of a we-perspective. *Synthese* (2020). <https://doi.org/10.1007/s11229-020-02834-6>
9. Danaher, J.: Robots, law and the retribution gap. *Ethics Inf. Technol.* **18**(4), 299–309 (2016). <https://doi.org/10.1007/s10676-016-9403-3>
10. Floridi, L.: The ethics of information, 1st edn. Oxford University Press, Oxford (2013)
11. Gifford, D.G.: Technological triggers to tort revolutions: steam locomotives, autonomous vehicles, and accident compensation. *J. Tort Law* **11**(1), 71–143 (2018). <https://doi.org/10.1515/jtl-2017-0029>
12. Hart, H.L.A.: Punishment and responsibility: essays in the philosophy of law. Oxford University Press, Oxford (1968)
13. Haydon, G.: On being responsible. *Philos. Quar.* **28**(110), 46–57 (1978)
14. Heinrichs, B., Eickhoff, S.B.: Your evidence? Machine learning algorithms for medical diagnosis and prediction. *Hum. Brain Mapp.* **41**(6), 1435–1444 (2020). <https://doi.org/10.1002/hbm.24886>
15. Heinrichs, B., Knell, S.: Aliens in the space of reasons? On the interaction between humans and artificial intelligent agents. *Philos. Technol.* (2021). <https://doi.org/10.1007/s13347-021-00475-2>
16. Heinrichs, J.-H.: Artificial intelligence in extended minds: intrapersonal diffusion of responsibility and legal multiple personality. In: Beck, B., Kühler, M. (eds.) *Technology, anthropology, and dimensions of responsibility*, pp. 159–176. J.B. Metzler, Stuttgart (2020)
17. Hernández-Orallo, J., Vold, K.: AI extenders: the ethical and societal implications of humans cognitively extended by AI. Paper presented at the Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society, Honolulu, HI, USA (2019) <https://doi.org/10.1145/3306618.3314238>
18. Horwitz, M.J.: The transformation of American Law, 1780–1860. Harvard University Press (1977)
19. Johnson, D.G.: Rethinking the social responsibilities of engineers as a form of accountability. In: Michelfelder, D.P., Newberry, B., Zhu, Q. (eds.) *Philosophy and engineering: exploring boundaries, expanding connections*, pp. 85–98. Springer International Publishing, Cham (2017)
20. Köhler, S., Roughley, N., Sauer, H.: Technologically blurred accountability? Technology, responsibility gaps and the robustness of our everyday conceptual scheme. In: Ulbert, C., Finkbusch, P., Sondermann, E., Debiel, T. (eds.) *Moral agency and the politics of responsibility*, pp. 51–67. Routledge, London (2017)
21. Lakoff, G.: Women, fire, and dangerous thing. What categories reveal about the mind. University of Chicago Press (1987)
22. Lenk, H., Maring, M.: Deskriptive und normative Zuschreibungen von Verantwortung. In: Lenk, H. (ed.) *Zwischen wissenschaft und ethik*. Suhrkamp, Frankfurt (1991)
23. Loh, J.: Responsibility and robot ethics: a critical overview. *Philosophies.* **4**(4), 58 (2019). Retrieved from <https://www.mdpi.com/2409-9287/4/4/58>
24. Matthias, A.: The responsibility gap: ascribing responsibility for the actions of learning automata. *Ethics Inf. Technol.* **6**(3), 175–183 (2004). <https://doi.org/10.1007/s10676-004-3422-1>
25. McKeon, R.: The development and the significance of the concept of responsibility. *Revue Internationale de Philosophie*, **11**(39 (1)), 3–32 (1957). Retrieved from <http://www.jstor.org/stable/23940271>
26. Napoletani, D., Panza, M., Struppa, D.C.: Is big data enough? A reflection on the changing role of mathematics in applications. In: Mircea, P. (ed.) *The best writing on mathematics 2015*, pp. 293–304. Princeton University Press (2016)
27. Nyholm, S.: Attributing agency to automated systems: reflections on human-robot collaborations and responsibility-loci. *Sci. Eng. Ethics* **24**(4), 1201–1219 (2018). <https://doi.org/10.1007/s11948-017-9943-x>
28. Nyholm, S.: The ethics of crashes with self-driving cars: A roadmap. I. *Philosophy Compass* **13**(7), e12507 (2018). <https://doi.org/10.1111/phc3.12507>
29. Nyholm, S.: The ethics of crashes with self-driving cars: a roadmap. II. *Philos. Compass* **13**(7), e12506 (2018). <https://doi.org/10.1111/phc3.12506>
30. Ropohl, G.: Das risiko im prinzip verantwortung. *Ethik Und Sozialwissenschaften* **5**(1), 109–120 (1994)
31. Rosch, E.: Cognitive reference points. *Cogn. Psychol.* **7**(4), 532–547 (1975). [https://doi.org/10.1016/0010-0285\(75\)90021-3](https://doi.org/10.1016/0010-0285(75)90021-3)
32. Rudin, C., Radin, J.: Why Are we using black box models in AI when we don't need to? A lesson from an explainable AI competition. *Harvard Data Science Review* (2019). <https://doi.org/10.1162/99608f92.5a8a3a3d>
33. Sellars, W.S.: Empiricism and the philosophy of mind. Harvard University Press, Cambridge and London (1997)
34. Shoemaker, D.: Attributability, answerability, and accountability: toward a wider theory of moral responsibility. *Ethics* **121**(3), 602–632 (2011). <https://doi.org/10.1086/659003>
35. Silver, D., Singh, S., Precup, D., Sutton, R.S.: Reward is enough. *Artif. Intell.* **299**, 103535 (2021). <https://doi.org/10.1016/j.artint.2021.103535>
36. Sparrow, R.: Killer robots. *J. Appl. Philos.* **24**(1), 62–77 (2007). <https://doi.org/10.1111/j.1468-5930.2007.00346.x>
37. Sussman, H.: Victorian technology: invention, innovation, and the rise of the machine. Praeger Publishers, Santa Barbara, Denver and Oxford (2009)
38. Tigard, D.W.: There is no techno-responsibility gap. *Philos. Technol.* **34**(3), 589–607 (2021). <https://doi.org/10.1007/s13347-020-00414-7>
39. van de Poel, I.: The relation between forward-looking and backward-looking responsibility. In: Vincent, N.A., van de Poel, I., van den Hoven, J. (eds.) *Moral responsibility: beyond free will and determinism*, pp. 37–52. Springer, Netherlands, Dordrecht (2011)
40. Vincent, N.A.: A structured taxonomy of responsibility concepts. In: Vincent, N.A., van de Poel, I., van den Hoven, J. (eds.) *Moral responsibility: Beyond Free Will and Determinism*, pp. 15–35. Springer, Netherlands, Dordrecht (2011)
41. Wallach, W., Allen, C.: Moral machines. Teaching robots right from wrong. Oxford University Press, Oxford and New York (2009)
42. Watson, G.: Two faces of responsibility. *Philos. Top.* **24**(2), 227–248 (1996). Retrieved from <http://www.jstor.org/stable/43154245>
43. Wittgenstein, L.: Philosophische Untersuchungen. In: *Schriften* (Vol. 1). Frankfurt am Main: Suhrkamp (1997)
44. Wooldridge, M.: Reasoning about rational agents. MIT Press (2003)
45. Zeiler, M. D., Krishnan, D., Taylor, G. W., Fergus, R.: Deconvolutional networks. Paper presented at the 2010 IEEE Computer Society Conference on Computer Vision and Pattern Recognition, 13–18 June 2010 (2010)

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.