

Moral Decision-Making via AI – deep ethics?

About shifting or losing responsibility

Tanja Henking

According to the latest statement of the German Ethics Council »Man and Machine – Challenges posed by Artificial Intelligence«, Vice-Chairman Julian Nida-Rümelin stated in a press release that »AI applications cannot replace human intelligence, responsibility and evaluation« (German Ethics Council 2023, press release). In addition to human intelligence, which may still need to be more clearly differentiated from Artificial Intelligence, this refers to two terms in particular, which will be examined closely in the following: evaluation/judgement and responsibility. In essence, it is about evaluation from an ethical perspective, which ultimately also has to stand up to legal scrutiny. This entails responsibility for making the evaluation, the evaluation process and, finally, the outcome of the evaluation/assessment process.

I

The terms »morality« and »ethics« are used in this article, but there is not enough space here to clarify them in greater detail. Put simply, morality is understood as the rules and values generally accepted in a society and ethics as the science of morality. It, thus, also addresses the question of right and wrong. Similarly, the formulation of ethical reflected decisions or actions is used. The question of whether a decision is right or wrong should be separated from the question of what rules and values we give ourselves as a society and which support our coexistence. In this context, the term »evaluation/judgement« mentioned at the beginning presupposes a moral capacity for judgement, which we assume in humans, but which can at best be trained into an Artificial Intelligence (AI). Whether an AI can also learn this is addressed by the question of a moral AI or the concept often already formulated of moral machines (or

only morally acting machines). The question that follows is whether an AI can ever acquire this ability or if it can only ever represent majority ideas. The hint that majority decisions are not always correct, and minorities should not be neglected does not really need any further mention. However, the question of a moral AI ignores a question that needs to be clarified beforehand: the understanding of ethics.

Formulations such as moral or ethical AI fall short. They pretend that the AI »experiences« an ethical conflict and has to solve it. This often involves different questions of morality, ethics and norms, and this on different levels. However, distinguishing between the different levels and tasks is essential if we are to approach the question of a moral AI at all. This begins with the issue of who determines the morals, ethics and norms that are fed into an algorithm and ends with the question of who evaluates the result at the end. Finally, it is also important to know whether humans are able to comprehend the decision-making process at all. This addresses a central (and largely unresolved) question, namely, the explainability of AI (Amann et al. 2020). Does the human consider the result found to be correct or does the AI evaluate it as correct – and what is the benchmark for evaluation? Should we not solve these questions first before we give in to the »temptation« of handing over these difficult questions to AI? And furthermore: Why should we actually want this? What benefits are promised, what risks are feared?

Do humans provoke ethical conflicts by using AI or do they solve them by employing AI? If they want to solve ethical conflicts by means of AI, this is based on an understanding of ethics that should be contradicted at this point. The ethical conflict relates to people and is bound to people. There can be no objective right and wrong, but a person-specific solution is required. The individual involved in an ethical conflict cannot allow himself to be relieved of responsibility by an AI (Sparrow 2021). This will be shown in the following.

This is because for an AI to »experience« a conflict, it must not only be able to decide between A and B, but also take into account the values that shape the situation and influence a decision, it must see what is in favour of A and in favour of B. In the case of a real ethical conflict or even dilemma, we see humans struggling with it, wrestling to make the right decision, and perhaps even suffering from the fact that there will always be vulnerabilities left in the decision-making process under any consideration. When humans have gone through this conflict, wrestled with a decision, we will often be able to accept

the decision made. An AI will not do this; it makes the decision according to a predetermined principle¹ (Iphofen and Kritikos 2019).

II

This raises the question of whether AI has morals or can be trained to have them. What moral concepts underlie the AI? What ethical values does the AI pursue? This question is anything but trivial. If an AI is to be used for questions that are primarily normative in nature, the search for an answer to this question is almost mandatory, because only then it is possible to make transparent according to which ideas the AI is operating, and which value system has been »written« into its training data set. If one takes a norm from a law, then the AI would at least have to be able to consider the legal methods of interpretation, and feed in case law and the literature of legal science. It would need to be able to classify possible »minority opinions« according to their significance for the debate. So, is the legal machine possible? Why law also needs common sense cannot be part of this contribution, which is to be limited to morality, ethics and responsibility. Nevertheless, some of the following considerations will be transferable.

However, first of all, let us return to moral values. This raises questions that have preoccupied philosophers and ethicists for centuries. According to which principles do we judge our actions? The deontological approach, which is predominant in the German tradition, and a utilitarian approach will be chosen for the following considerations in order to shed light on the problem of the criteria according to which a decision should be made. The ethics of principles, which is widespread in medical ethics, will be used again later to question whether the shifting of decisions to an AI-based system is convincing and which considerations should override the question of mere feasibility and, better still, be placed in front of it.

According to deontology, an action in itself can be good or bad, wrong or right. The action is not evaluated according to its consequences. Following its probably most prominent representative, Immanuel Kant, the action is judged according to whether it is in compliance with an obligatory rule and if the ac-

1 The discussion about the possible legal personality of the AI in the future is not intended to take place here.

tion is committed based on this obligation. Deontology, thus, focuses on the preconditions of action (Stanford Encyclopedia of Philosophy 2020).

By contrast, utilitarianism, which goes back to Jeremy Bentham (1780) and, thus, has its origins in the 18th century, pursues a different idea: action is evaluated according to its utility. Therefore, from a utilitarian perspective, the greatest possible benefit for as many people as possible is the standard for correct human decision-making behaviour (Stanford Encyclopedia of Philosophy 2017). *Prima facie*, it is easier to work with this approach in the context of AI (similar Iphofen & Kritikos 2019). It is important to keep this in mind because we might be abandoning our tradition.

III

Medical decision-making situations, as situations with normative impact and which could affect every human being, will be the focus of the following section. AI-based clinical decision support systems (CDSS) are increasingly finding their way into medicine. The Central Ethics Committee of the German Medical Association (ZEKO) dedicated a statement of its own regarding this issue in 2022, emphasising the role of the human being in addition to the possible benefit, whereby the human being should not be forced out of the decision (ZEKO 2022).

The concern about automated decisions is expressed in Art. 22 of the General Data Protection Regulation, where it states that »the data subject shall have the right not to be subject to a decision which produces legal effects concerning him or her or similarly significantly affects him or her and which is based solely on automated processing [...]«.

The CDSS can now help to find the right diagnosis, issue warnings, for example, if there is a risk of impending sepsis, or make therapy suggestions for a previously defined diagnosis.

IV

As part of the Check.App project funded by the Federal Ministry of Education and research (BMBF), the ethical, legal and social implications of so-called symptom checker apps that allow users to classify their symptoms by means of an app-controlled medical history are being investigated (Wetzel et al. 2022;

Müller et al. 2022). The apps work predominantly with the probabilities of a diagnosis. This is followed by the recommendation to consult a doctor or even go to an emergency room. These apps are currently being discussed primarily from the point of view of accuracy (Gilbert et al. 2022). However, it is interesting to see how the use of these apps could affect the doctor-patient encounter in the long term. As long as one sees this app operating in a social insurance system with free health care, a noticeable impact on the health-care system is probably not expected. The fact that the resources of even a fundamentally well-equipped health-care system are finite is shown by the current debate about the congestion of emergency rooms, with the result that work is being done on the use of algorithms for triage, the use of which, however, is not uncontroversial due to fears about patient safety (Elsner et al. 2018; Marburger Bund/DGINA/DIVI, Gemeinsame Pressemitteilung 2021).

The keyword »triage« addresses an area that exemplifies a real ethical dilemma. First of all, triage means sifting, sorting. The term originally comes from disaster medicine. The objective of triage is to save as many people as possible with limited resources. Whereas before the COVID-19 pandemic the term was probably only known to experts, the pandemic has turned it into a description of a scenario in the general population which people faced with worry and fear. While German hospitals – to the best of our knowledge – do not have to triage in the same way as hospitals in France, New York or Bergamo/Italy, debates have, nevertheless, been triggered that have, at least, called into question certain basic ethical assumptions (German Ethics Council 2020). The discussion was initially reduced to the question of who should receive a ventilator when only one is available, but two patients need it. The question is broadened to include the situation of the patient who arrives later and for whom no ventilator is available, but who may have a better survival situation than the patient who has already been ventilated. If we take the situation that is initially easy (or easier) to assess from a legal point of view: more patients arrive at the same time and there is not enough capacity available. Not all patients can be treated and saved (or given a chance of being saved). If the doctor saves patient A but not B, the doctor cannot be blamed from a legal point of view. This is because two obligations meet here, both of which the doctor cannot fulfil. Which patient he or she chooses, A or B, is ultimately irrelevant. In order to relieve physicians of the burden of this difficult and stressful decision-making situation, the medical society of intensive care physicians has provided them with criteria to help them make decisions. These criteria focus primarily on the prospects of a successful treatment (DIVI

2020). In other words, those who have a better chance of surviving the disease through treatment should receive it. Those who are less likely to survive, such as those who are older, overweight or have certain pre-existing conditions, will not receive the treatment. The idea of transferring this decision-making algorithm to an AI is not far-fetched, and it has sometimes been reported that such an algorithm has been used. If one considers the idea of triage according to the criteria of success, it quickly becomes apparent that the distribution of limited resources is based on the greatest possible benefit. Triage is utilitarian by its very nature. The medical likelihood of success may be a selection criterion, but it is not mandatory. Therefore, other criteria have been put up for discussion in the context of the debate on triage in the pandemic. These range from age and, thus, life expectancy (Hoven 2020), the juxtaposition of the young mother versus the older man with a previous illness, to the decision by lot (Walter 2022). If one were to (be able to) agree on one of these criteria, an AI would also have to take this into account if it is to decide morally in the sense of a normative consensus. However, this type of evaluative criteria, which can also be classified as utilitarian, has so far been consciously avoided in the German legal system. The decision of the Federal Constitutional Court on the Aviation Security Act (Luftsicherheitsgesetz), which was passed with the collapsing Twin Towers in New York in mind, is groundbreaking. Can shooting down a plane carrying 300 people be justified if it could save 3000 people? The Federal Constitutional Court has rejected the quantification of human lives and declared the authorisation to shoot down an aircraft unconstitutional and not justified (BVerfGE 115, 118 – 166). From a legal point of view, one thing is important to distinguish: the Aviation Security Act would have provided an authorisation basis for the selection decision to the detriment of the aircraft occupants. If the pilot themselves were to make a decision in the acute, specific situation that – regardless of how they decide – leads to the death of human lives, then they would not be blamed for making a particular decision. Similar to the case of medical triage, if all the people cannot be saved, the doctor cannot meet all the obligations, so cannot be blamed for their selection decision (Gerson 2020). That their decision remains a personally burdensome one, because their rescue decision is, at the same time, a non-salvation decision, is the essence of the real ethical dilemma. The person remains (morally) responsible for their decision.

Similar scenarios have recently been discussed for autonomous driving. Which group of people should the car steer into if a collision is unavoidable: the group of children or the group of pensioners? Here too, similar to the real triage

described above, a genuine ethical dilemma is constructed, which is, however, provoked by the use of autonomous vehicles. Therefore, it is demanded in some cases that the human being must remain (in the loop). Human beings can act neither rightly nor wrongly, because no matter how they decide, their decision will not be entirely without harm to one person or one group of people and, simultaneously, of benefit to the other person or group of people. Would chance, therefore, be the only ethically correct decision or should another ethically correct decision be taken into account within the framework of a programmed decision and would this decision then be dependent on the respective value system of a society. So should the car have been driven into group A in one country and into group B in another? This would be a normative question that would have to be decided in advance of an anticipated collision. However, this shows one thing quite clearly here: whereas the situation in the triage described above falls on the doctor and they have to make a decision, this triage is one that is calculated in advance. The algorithm will experience a moral stress and no decision burden. It is not a decision that affects or concerns the algorithm personally. It suffers from neither the decision to be made nor the consequences. For the algorithm, the decision remains impersonal. This makes one thing clear: ethical conflict is personal (Sparrow 2021). It is, therefore, not a question of an ethical AI or ethical decision-making in these situations, but rather of an anticipated, ethically problematic decision-making situation for which a decision has already been made in advance of the actual situation.

However, this raises the question of whether the use of algorithms/AI would reduce or even increase the decision-making burden and the moral distress associated with it. Can the person hand over the decision and, thus, relieve themselves, or would they experience it as a burden and ultimately feel powerless? Another difficulty in this context, however, should be highlighted. The starting point was the moral or ethical question underlying a decision. Although the law represents a basic moral consensus of a society, it can also raise ethical questions. When the legislator passes a norm, however, it is based on fundamental decisions of coexistence, on the one hand, and requires interpretation and is subject to further development, on the other hand, as advanced by the literature in the legal sciences and case law. Concepts in need of interpretation such as »well-being« come to mind. Well-being can be understood subjectively or objectively (Braun et al. 2022). While the »best interests« principle applies in other legal systems, German guardianship law, for example, does not recognise this, and it is not uncommon for translational inaccuracies and difficulties of interpretation to arise here in legal comparisons (relevant,

inter alia, in decisions concerning the end of life). Moreover, even the question of what is in a person's best interests is subject to interpretation and will be influenced by changing social understanding and Zeitgeist. Incidentally, the legislator decided to delete the term »well-being« in the law on guardianship in order to counter difficulties of interpretation and misunderstandings observed in practice (Schnellenbach et al. 2020; Henking, 2022). The legislator has now also enacted a law in the area of triage within the framework of the Protection against Infection Act. To the best of our knowledge, this regulation on triage is unique throughout the world. The legislature was called upon to act after the Federal Constitutional Court (BVerfG, decision of 16.12.2021, ref.: 1 BvR 1541/20) shared the concerns expressed in a constitutional complaint of discrimination against persons with disabilities in the case of triage. The German Protection against Infection Act was expanded by Sec. 5c of the act to include a prohibition of discrimination, according to which no one would be disadvantaged in the event of the insufficient availability of intensive care treatment capacities essential for survival because of, *inter alia*, their disability, age or degree of frailty. The allocation decision should only be made based on the current and short-term survival probability of the patient concerned. This already raises the question of how short »short-term« should be defined. Comorbidities may be taken into account, but only in the assessment of the current and short-term probability of survival if they significantly reduce the short-term probability of survival related to the current disease due to their severity or combination. If one were to make this decision (which, according to the law, has to be made by two doctors) with an AI-based decision support system, then all these normative preconditions, which are controversial in detail, would have to be taken into account. It would be feasible to limit the AI to a statement on the probability of survival, which only hints at another difficulty: the uncertainties of any prognosis – especially in the case of people, all of whom have a poor prognosis to begin with. Just consider that less than half of the patients suffering from COVID-19 who received invasive ventilation have survived (Karagiannidis et al. 2020). One can imagine AI-based decision support systems in various clinical settings where they are already partially in use (see above). But if one imagines this system in end-of-life decisions, then the question arises here, too, whether ethical, normative questions are excluded or frozen to a value concept at time X. And once again, the question arises as to which moral and value concepts influence the decision. This is because, according to our system of values, we do not follow a »best interest« principle (the content of which would also have to be clarified) which would

allow for an objective consideration, but what is required is a determination of the (presumed) will of the individual oriented towards their subjective interests. In other words, the value system of the person, of the individual, has to be used for the decision of further treatment or the limitation of therapy. This, in turn, makes the ethical decision a personal decision, both for the individual concerned and the person who may have to make it on their behalf. Let us assume that this decision has to be made by a relative (health-care proxy), then this decision remains a personal decision (Sparrow 2021). The relative has to determine and consider the values of the person concerned in their own individual case. They cannot remove themselves from this responsibility; they cannot transfer it to an AI system (Sparrow 2021). The AI will not be able to tell them what their relative would have wanted in concrete terms or whether they themselves are struggling with the interpretation and implementation of this will, because their decision also has consequences for their own life. Other considerations could possibly run the risk of giving (too much) weight to general values. The ethical complexity here arises from the challenge of determining the will of the person in the concrete situation and, thus, making the »right« decision for them.

The temptation to use AI for such decision-making situations may seem attractive. Meier et al. (2022) have presented a proof of concept for ethical decision-making in the clinic. This is intended to help with ethical decisions. Their main aim is to show whether a decision support for ethically difficult questions can be developed by means of an algorithm. The authors obtained their training data set from clinical cases of clinical ethics committees. The cases are categorised in the areas of abortion, decisions regarding minors, refusal of treatment and end-of-life decisions, among others. They used the principles of ethics according to Beauchamp and Childress (2019), the approach most frequently used in the clinical field to discuss ethical questions. They have tried to operationalise these principles of ethics and develop a corresponding algorithm. Their approach showed an astonishingly frequent agreement between the algorithm-based recommendations and those of experts or textbook cases. But is this now also proof that ethical questions can be solved in an AI-based way? And why would one want to do this? The authors speak, among other things, of limited time and personnel resources, which, in the worst-case scenario, could again mean triage, even if these case constellations did not play a role in the data set of the research design. However, the question remains: what are we actually hoping for (see critical statement by Gundersen and Baroe 2022)? What is the understanding of ethics behind these considerations (Spar-

row 2021)? Does the use of AI lead to a reduction in the workload of the staff responsible for the decision-making and care of the patient? This can only be convincing if the relinquishment of decision-making authority reduces stress. Can the doctor be relieved of responsibility? Can it shift the decision? However, the person remains the one who implements the decision in the end. They remain involved and, thus, it is still a personal ethical decision. The idea of ethical case consultation also aims to show the different values, the various interpretations and their points of contact that exist in a group and to work out these different viewpoints in a moderated process and develop a willingness to take on the perspective. This deliberation cannot be done by an AI. It remains a human task. The decisions that have to be made in the context of an ethical case conference represent ethical problems. Rarely do we encounter real ethical dilemmas as in triage. The authors are rightly confronted with the criticism that they have not sufficiently explained why the use of AI is necessary (Sauerbrei et al. 2022). Both the WHO warning not to use AI as a first resource and the EU Ethical Guideline for Trustworthy AI (HLEG 2019) have to be recalled in this context.

If we once again combine the examples of autonomous driving and decision-making in medical contexts, the question remains open as to whether and when we are prepared to leave or hand over the decision to an AI. At what level of automation does the person withdraw from the responsibility for the decision or is there a pressure to justify wanting to decide differently from the AI (the comparison of autonomous driving can be found at Eric Topol 2019; Henking, 2023; Duttge 2019; Katzenmeier 2019)?

V

While AI can be expected to bring new accuracy, an increase in knowledge and objectivity to the assessment of medical prospects of a successful outcome, in addition to the finding that the human body reacts individually, it is first necessary to address where and when implicit values play a role. Decisions at the end of life are a good example of this, especially because the increasing need as a person to be able to make decisions as independently as possible at the end of life has played an increasingly important role in recent decades. This also reflects the values of a society whose morals and ethics are changing. Whether this individuality and development can be depicted sufficiently quickly by an AI must, at least, be doubted. This is because it always requires sufficient reflec-

tion and a discourse on ethical ideas. Society cannot pass on the responsibility for this discourse to an AI, but must take it on itself as an ongoing task. This addresses part of the question of what constitutes ethics. Ethical conflicts can arise from the use of an algorithm, as in the inevitable collision in autonomous driving, because a situation is anticipated and the outcome is already determined. These considerations can be transferred to CDSS depending on how much the system displaces humans from the decision-making process. This suppression must find its limits when individual value concepts, and thus ethical questions, affect the person in terms of their own value system and ethical ideas. The ethical conflict is a personal one that can neither be left to an AI nor for which a proxy can absolve itself of responsibility (Henking, 2023; Sparrow 2019). He or she may even experience a hypothetical AI decision as powerlessness vis-à-vis the system.

VI

Before speaking of moral, ethical algorithms (in clinical use), we should reflect on what kind of ethics we are talking about. Is it the use of algorithms/AI where ethical problems can be anticipated? Should algorithms/AI be used to solve ethical problems? Should the algorithm itself act morally? Especially when answering the last question, it is necessary to think more deeply about what moral decision-making means and what competences, including the struggle for right and wrong, this requires. The evaluation of right and wrong cannot be left to an AI – it is the responsibility of humans and should remain so in the future. Those who call an AI ethical should always declare their understanding of ethics and morality.

References

- Amann, Julia, Alessandro Blasimme, Effy Vayena, Dietmar Frey, and Vince I. Madai. 2020. »Explainability for artificial intelligence in healthcare: a multidisciplinary perspective« *BMC Medical Informatics and Decision Making* (20): 310, DOI: 10.1186/s12911-020-01332.
- Alexander, Larry, and Michael Moore. 2021. »Deontological Ethics« *The Stanford Encyclopedia of Philosophy* (Winter 2021 Edition), edited by Edward N. Zalta.

- Accessed April 22, 2023. <https://plato.stanford.edu/archives/win2021/entries/ethics-deontological/>.
- Bentham, Jeremy. 1870. *An introduction to the principles of morals and legislation*. (1907 reprint of 1823 Edn.), Oxford: Clarendon Press.
- Beauchamp, Tom L., and James F. Childress. 2019. *Principles of biomedical ethics*. (8th ed.), New York: Oxford University Press.
- Braun, Esther, Jakov Gather, Tanja Henking, Jochen Vollmann, and Matthé Scholten. 2022. »Das Verständnis von Wohl im Betreuungsrecht – eine Analyse anlässlich der Streichung des Wohlbegriffs aus dem reformierten Gesetz« *Ethik in der Medizin* 2022 (34): 515–528.
- Deutscher Ethikrat. 2023. *Mensch und Maschine – Herausforderungen durch Künstliche Intelligenz*. Stellungnahme, Berlin. Accessed April 22, 2023. <https://www.ethikrat.org/fileadmin/Publikationen/Stellungnahmen/deutsch/stellungnahme-mensch-und-maschine.pdf>.
- Deutscher Ethikrat. 2020. *Solidarität und Verantwortung in der Corona-Krise*. Ad-hoc Stellungnahme, Berlin. Accessed April 22, 2023. <https://www.ethikrat.org/fileadmin/Publikationen/Ad-hoc-Empfehlungen/deutsch/ad-hoc-empfehlung-corona-krise.pdf>.
- DIVI. 2021. *Entscheidungen über die Zuteilung intensivmedizinischer Ressourcen im Kontext der COVID-19-Pandemie*. Version 3. Accessed April 22, 2023. <https://www.divi.de/joomlatools-files/docman-files/publikationen/covid-19-dokumente/211125-divi-covid-19-ethik-empfehlung-version-3-vorabfassung.pdf>.
- Driver, Julia. 2021. »The History of Utilitarianism«, *The Stanford Encyclopedia of Philosophy* (Winter 2022 Edition), edited by Edward N. Zalta and Uri Nodelman. Accessed April 22, 2023. <https://plato.stanford.edu/archives/win2022/entries/utilitarianism-history/>.
- Duttge, Gunnar. 2019. »Decision-Support-System« für Therapieentscheidungen am Lebensende?« *Medizinrecht* 37: 771–776.
- Elsner, Christian, Martin Blaschka, and Martin Kleehaus. 2018. »App-basierte Systeme im Bereich der medizinischen Notfallversorgung.« *Notfallmedizin up2date* 2018 13(03): 251–266.
- EU HIGH-LEVEL EXPERT GROUP ON ARTIFICIAL INTELLIGENCE. 2019. *Ethics Guideline for Trustworthy*. Accessed April 22, 2023. <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>.
- Gerson, Oliver Harry. 2020. »§ 3 Pflichtenkollision«. In *Pandemiestrafrecht*, edited by Esser, Robert and Michael Tsambikakis, C.H.Beck, München.

- Gilbert, Stephen, Alicia Mehl, Adel Baluch et al. 2020. »How accurate are digital symptom assessment apps for suggesting conditions and urgency advice? A clinical vignettes comparison to GPs.« *BMJ open* 2020 10 (12): e040269.
- Gundersen, Torbjørn, and Kristine Bærøe. 2022. »Ethical Algorithmic Advice: Some Reasons to Pause and Think Twice« *The American Journal of Bioethics* 22 (7): 26–28.
- Henking, Tanja. 2023. »Theorie- und evidenzbasierte Gesundheitspolitik in Zeiten der Digitalisierung und Künstlicher Intelligenz. Ethische und rechtliche Überlegungen.« In: *Smart Regulation: Theorie- und evidenzbasierte Politik*, edited by Wendland, Matthias, Iris Eisenberger, and Rainer Niemann, 1–13.
- Henking, Tanja. 2022. »Die Reform des Betreuungsrecht« In: *Der Nervenarzt* 2022 93: 1125–1133, DOI: 10.1007/s00115-022-01355-6.
- Hoven, Elisabeth. 2020. »Die »Triage«-Situation als Herausforderung für die Strafrechtswissenschaft«. In: *Juristenzeitung* 2020: 449.
- Iphofen, Ron and Mihalis Kritikos. 2019. »Regulating artificial intelligence and robotics: ethics by design in a digital society« *Contemporary Social Science* 16 (2): 170–184 DOI: doi.org/10.1080/21582041.2018.1563803.
- Karagiannidis, Christian, Carina Mostert, Corinna Hentschker et al. 2020. »Case characteristics, resource use, and outcomes of 10 021 patients with COVID-19 admitted to 920 German hospitals: an observational study« *Lancet Respir Med* September 2020 8(9): 853–862, DOI: 10.1016/S2213-2600(20)30316-7.
- Katzenmeier, Christian. 2019. »Big Data, E-Health, M-Health, KI und Robotik in der Medizin. Digitalisierung des Gesundheitswesens – Herausforderung des Rechts« *Medizinrecht* 37: 259–271.
- Koch, Roland, Anna-Jasmin Wetzels, Christine Preiser et al. 2022. »Ethical, legal, and social implications of symptom checker apps in primary health care (Check.App): Protocol for an interdisciplinary mixed methods study« *JMIR Research Protocols* 11 (5): e34026.
- Marburger Bund, DGINA, DIVI. 2021. *Gemeinsame Pressemitteilung, Keine Experimente mit der Patientensicherheit!* Accessed April 22, 2023. <https://www.marburger-bund.de/sites/default/files/files/2021-02/Gemeinsame%20Pressemitteilung%20MB%2C%20DGINA%2C%20DIVI%20-%20Keine%20Experimente%20mit%20der%20Patientensicherheit.pdf>
- Meyer, Lukas J., Alice Hein, Klaus Diepold, and Alena Buyx. 2022. »Algorithms for Ethical Decision-Making in the clinic: A proof of concept« *The American Journal of Bioethics* 22 (7): 4–20.

- Müller, Regina, Malte Klemmt, Hans-Jörg Ehni et al. 2022. »Ethical, legal, and social aspects of symptom checker applications: a scoping review« *Medicine, Health Care and Philosophy* 25 (4): 737–755.
- Sauerbrei, Aurelia, Nina Hallowell, and Angeliki Kerasidou. 2022. »Algorithmic Ethics: A Technically Sweet Solution to a Non-Problem« *The American Journal of Bioethics* 22 (7): 28–30.
- Sparrow, Robert. 2021. »Why machines cannot be moral« *AI & Society* 36: 685–693.
- Topol, Eric. 2019. *Deep Medicine. How Artificial Intelligence Can Make Healthcare Human Again*. New York: Basic Books.
- Walter, Tonio. 2022. »Keine Verpflichtung für ein Triage-Gesetz – und kaum Vorgaben dafür«. In: *Neue juristische Wochenschrift* 2022 (6): 363–366.
- Zentrale Ethikkommission der Bundesärztekammer (ZEKO). 2021. *Entscheidungsunterstützung ärztlicher Tätigkeit durch Künstliche Intelligenz. Stellungnahme*. Accessed April 22, 2023. https://www.zentrale-ethikkommission.de/fileadmin/user_upload/_old-files/downloads/pdf-Ordner/Zeko/ZEKO_SN_CDSS_Online_final.pdf.