



COMPARISON OF CARE - ETHICAL AND
FAT (FAIRNESS, ACCOUNTABILITY,
TRANSPARENCY) - ETHICAL APPROACHES IN
AI TRANSLATION

Authors: Havryliuk Mariia, Mayer Mirijam, Yermekova Aidana

Date: 30.06.2026

In the course of the Micro Degree 'AI and Society' – Ethical Aspects

Table of contents

1. Introduction.....	2
2. The Dominance of FAT-values in the ethics of AI	5
2.1 What are FAT(-ethical) principles of AI.....	6
2.2 What is the normative background of FAT-values/principles?	8
3. The CARE-ethical approach and its possibilities within AI	13
3.1 What are CARE ethical principles?.....	13
3.2 Care in context of AI - Why AI can't "care" the same way humans do.....	15
3.3 Normative background for CARE ethics in AI	17
4. Applying FAT vs. Care Principles/Values onto the Example of AI Translation	21
4.1 Three "Tricky" Translation Cases: FAT vs. Care Values in Practice	22
4.2 Linguistic Analysis with LIWC: Do DeepL and ChatGPT Represent FAT or Care Values?...36	
4.3 Constructing a Model Card: How Should AI Translators Be Built to Grasp Culture and Meaning?.....	44
4.4 Summary.....	50
5. The normative implications of a FAT-ethical dominance in AI translation and AI more generally	51
5.1 Moral integrity, efficiency and the ethics of justice	51
5.2 Care and the denial of justice	53
5.3 Efficiency, practicality and neoliberal market rationality	54
6. Conclusion.....	56

1. Introduction

With this project we want to critically examine contemporary assumptions within the ethics of AI. Thereby we are highly concerned with the question of why the ethics of AI are mostly organized around the values of fairness, accountability and transparency (FAT). As they seem to occupy an essential standing for establishing just AI systems, one has to ask if they also present as sufficient criteria for this undertaking e.g.: when we look at how AI systems actually shape people's lives and relationships. Drawing on the work of Jonathan Cohn (2020) as well as leaning on theories of the ethics of care (Gilligan 1993; Held 2006), we argue that the current ethics of AI is dominated by an underlying "ethics of justice" that treats justice primarily as a matter of equal treatment, formal rules, and efficient harm-minimization, while skipping questions of dependence, vulnerability and the quality of concrete relationships.

To make this contrast empirically tractable, and because questions of care are always also questions of voice, recognition and address, we want to focus this research on the case of AI-mediated language translation. AI translation is a site where linguistic choices routinely mediate dependence, vulnerability and identity (e.g.: in the selection of pronouns, the handling of emotional disclosure, or the representation of gender). In the case of AI translation programs the difference between FAT-oriented optimization and care-oriented responsiveness becomes, as we state, practically visible in concrete outputs. Against this background, the project asks: To what extent do the AI models of DeepL and ChatGPT instantiate care as opposed to FAT values in language translation, and what are its social and normative implications?

To meet this (empirical) goal, we firstly establish the theoretical framework throughout the preceding chapters. In Chapter 2 e.g.: we first analyze the presumed FAT dominance, marking it as a technical operationalization of the dominant "ethics of justice"- discourse. An FAT focused ethics translates fairness into mathematical constraints on error rates and outcome distributions, frames accountability in terms of traceable responsibility and liability and perceives transparency mainly as a precondition for informed consent and oversight, all within a broadly liberal-egalitarian picture of justice. While these are not unimportant, we claim that their seemingly neutral, technical form hides substantive normative commitments.

Against this background, Chapter 3 introduces care ethics as a systematic alternative lens for thinking about AI. Following Carol Gilligan (1996) and Virginia Held (2006), we treat care not as a private emotion opposed to justice, but as a set of practices and values oriented toward sustaining human life, agency and relationships over time. Held's distinction between care as practice (concrete activities directed at specific others' needs) and care as value (attentiveness,

sensitivity, and responsiveness as guiding norms) allows us to formulate what it would mean for AI systems to be designed from within a care perspective.

Building on Carolina Villegas-Galaviz (2023) and colleagues, Chapter 3 also develops a normative background for care-based AI by specifying four key aspects that a care-informed algorithm should incorporate: voice, relationships and interdependence, contextual (direct) judgement and vulnerability. A care-informed system should keep open the possibility of hearing different voices, avoid silencing those most affected and be developed by interdisciplinary, inclusive teams rather than narrow expert circles. It should resist the abstraction and decontextualization typical of FAT-ethical approaches by attending to the specifics of each case, and it should explicitly protect, rather than exploit, vulnerabilities, especially those of marginalized groups. We close the chapter by noting both the promise and the current limits of this framework: care ethics in AI remains largely theoretical, and its implementation raises open questions about measurement, conflict between values, and the continued need for human oversight.

After establishing the theoretical framework, Chapter 4 tries to answer our research question by “stress-testing” FAT- and care-ethical values in the concrete and empirically tractable domain of AI-translation. Translation is a paradigmatic site where language, culture, identity and interpersonal relationships are at stake in every lexical and syntactic choice, making it an ideal arena for investigating how different ethical frameworks matter in practice. Methodologically, the chapter combines qualitative case analysis with quantitative linguistic profiling and normative model design. First, it presents three “tricky” translation cases: relational ambiguity in formal vs informal address (Case 1), emotionally vulnerable speech in caregiving contexts tested through both an explicit and an implicit source text (Case 2), and gendered language and identity (Case 3). Then it compares the outputs of two widely used translation systems, DeepL and ChatGPT, through both a FAT-ethical and a care-ethical lens. Second, it uses LIWC-22 to analyze the source texts and translations across German and Russian, testing the hypothesis that current tools exhibit linguistic signatures characteristic of an “ethics of justice”-oriented discourse rather than a care-oriented language. Finally, drawing on both strands of analysis, the chapter develops a care-ethical model card for translation systems, specifying criteria around relational competence, identity and representation, accountability, and the explicit naming of contexts where the tool is not appropriate without human mediation.

Chapter 5 then steps back to examine the broader normative implications of FAT-dominance in AI translation and AI ethics more generally, in light of the preceding theoretical and empirical

chapters. It discusses how different underlying conceptions of justice and care shape what counts as harm or injustice in AI and how this interacts with wider neoliberal rationalities of efficiency and optimization.

Taken together, these chapters pursue a coherent research trajectory: from diagnosing the FAT-ethical dominance and its roots in an ethics of justice (Chapter 2), through articulating a care-ethical alternative and its normative architecture (Chapter 3), testing them empirically in AI translation (Chapter 4), and, finally, drawing out their broader normative consequences for how we imagine and govern AI (Chapter 5). The project's central contribution is to show that bringing care ethics into AI is not a matter of adding "softer" values, but of revising the very concept of justice that underwrites FAT-dominant approaches and, with it, the criteria by which AI systems and their translation outputs are judged as ethically acceptable.

2. The Dominance of FAT-values in the ethics of AI

In *In a Different Code: Artificial Intelligence and the Ethics of Care* Jonathan Cohn argues that the current state of discussion about the ethics of AI articulates itself mostly within the dimensions of fairness, accountability and transparency (FAT), implying that just conditions within AI can be ensured mostly through the realization of these principles. Cohn criticizes that a focus on these principles commonly overlooks the various *contexts* AI operates in, which can be illustrated by his example: “(...) if you are using AI to decide how to most efficiently detain migrant children away from their parents, does the level of fairness, accountability, and transparency really matter?” (Cohn 2020, 2). In this matter, one of his main critiques lies within an unquestioned approach towards ethics that circulates an idea of moral integrity that is strongly interwoven with its degree of efficiency. (Cohn 2020, 1-2). Thereby, the implied argument challenges what Cohn describes as the premature conclusion in dominant AI-ethical approaches that something is good because it is efficient. In this view, good systems are those whose programs have perfected their ability to accomplish tasks and to minimize cases of damage. Yet what exactly counts as the relevant kind of damage remains largely defined within the sphere of program efficiency itself. In other words, the harms to be avoided are mostly framed as failures of performance, rather than as violations of substantive moral or political commitments. (Cohn 2020, 1).

This suggests that ethical questions within AI are often treated as independent from specific social and political contexts, and that FAT optimization tends to be understood primarily as a means-related optimization. Such a tendency risks a normative shift from the question of what is right and just to the question of what works most efficiently. Questions like “should we do this at all?” are displaced by questions like “how can we do this in the most efficient way?”. It is therefore unsurprising, as Cohn rightly argues, that concerns about AI being used in destructive and undemocratic ways are justified. (Cohn 2020, 1)

This fear also has a broader structural basis. Drawing on Harvey, one can argue that in contemporary neoliberalism, market logics such as competition, efficiency and cost-benefit analysis spill over from economic strategy into political, governmental and even social and ethical rationalities (Harvey 2007, 65). Since AI is a major site of investment and profit, an overarching market logic can partly explain why certain ethical questions are asked, while others are not. In this context it is crucial that a dominant focus on functional criteria within AI-ethical discussions are met with scepticism. The fact that AI is an innovative and profitable tool

cannot replace substantive questions about justice, power and the kinds of social relations these systems help to produce or undermine.

Before discussing this critique any further, it is useful to question what exactly FAT-ethical principles within AI are, if they are ethically justifiable and where their dominance within ethics derives from?

2.1 What are FAT(-ethical) principles of AI

By starting an inquiry of current academic work using the wording of FAT-ethical principles within AI a lot of papers can be found that observe these principles as important or at least dominant ethical dimensions within AI (Fjeld et.al. 2020; Singhal et.al. 2024; Stark et.al. 2021). Balog and Bernard understand FAT principles as crucial measures for ensuring trustworthy, safe and fair AI- and information retrieval systems (Balog/Bernard 2023, 136:2). They also note that there are a lot of definitions about the exact contents of these principles and with their paper try to aim for an insightful overview over the most promising ones. Firstly, they name common understandings of FAT-ethical principles:

Fairness: “the quality of treating people equally or in a way that is reasonable.”

Accountability: “the fact of being responsible for your decisions or actions and expected to explain them when you are asked.”

Transparency: “the quality of something, such as a situation or an argument, that makes it easy to understand.”

Ethics: “moral principles that control or influence a person’s behavior.”” (Balog/Bernard 2023, p. 136: 10)

Then they go on by naming definitions that demonstrate the technical purposes of FAT-values: Fairness focuses on a system’s ability to not discriminate based on natural aspects, while accountability should be understood as who bears responsibility for the actions and malfunctions of a system, which could focus on the user, algorithm or designer. Transparency in this sense relates to facts of how easily a system, its inner workings and communication style about its internal processes can be comprehended. (Balog/Bernard 2023, 136:11 and 136:19).

Stark et.al. argue that even if FAT-ethical principles are “necessary for addressing the governance of AI systems, these tools address only a small fraction of the governance

challenges provoked by these technologies.” (Stark et.al. 2021, 259). Their critique centers on the lack of consensus on how these principles are defined in detail and the hasty judgement of the degree of ‘goodness’ of a system based on its degree of usability (or efficiency), which Stark et. al. do not recognize as fundamental for ethical claims. The lack of consensus, Stark et al. suggest, follows from the attempt of computer scientists to turn broad values such as fairness, accountability, transparency, and interpretability into technical tools, especially in machine learning (Stark et al. 2021, 259). The value of fairness, for example, has been translated into mathematical definitions such as individual versus group fairness. What gets ‘lost in translation’ between technical and ethical insights are the different ethical understandings and political projects that have historically been tied to the notion of fairness, including different ideas about political and economic equality. As a result, there is no clear agreement on what fairness really means, which makes it difficult to say when a system is truly fair (Stark et al. 2021, 260). It therefore remains open that there are actually different concepts of fairness that can be followed, for instance informed by differing schools of thought, like a more liberal focus on fairness as equality of opportunity, as in Rawlsian theories of justice, or a more socialist idea of fairness as equality of outcomes (Rawls 1999; Cohen 2009).

Stark et al. also point out, drawing on Nissenbaum, that even well-designed technical systems always contain built-in weaknesses that push them towards certain normative outcomes, which may conflict with social values such as fairness, democratic accountability, and justice (Stark et al. 2021, 260; Nissenbaum 2001). A standard example is predictive policing: even when the underlying algorithms are carefully chosen to satisfy a formal fairness metric, they still rely on historically biased crime data and can reinforce patterns of racialized surveillance and criminalization. In such cases, the system may formally meet a fairness constraint, while materially reproducing injustice in practice.

Moreover, Stark et al. emphasize that many AI ethics guidelines and value statements, for example those issued by large technology companies and multi-stakeholder initiatives, repeatedly invoke high-level principles such as transparency and accountability, whereas other values like solidarity, social equality, or sustainability are rarely mentioned (Stark et al. 2021, 261–262). This pattern can again be traced back to the already mentioned focus on efficiency within AI governance: principles that can be more easily translated into metrics, audits, and compliance procedures are privileged, while values that would require deeper political or socio-economic change remain marginal. For Stark et al., FAT-ethical principles are therefore particularly insufficient for addressing sociotechnical problems and cannot secure justice in the

context of the power and dominance structures in which AI systems are embedded (Stark et al. 2021, 260–261). Consequently, , they diagnose a “deterministic, expert-driven vision of AI/ML governance” (Stark et.al. 2021, 262), which offers next to no socio-economic or political solutions, ignoring questions of possibilities on how to ensure greater social justice and human welfare. (Stark et.al. 2021, 263).

2. 2 What is the normative background of FAT-values/principles?

After this short overview of the current debate on FAT-ethical principles within AI, we conclude that their dominance is not only directed at ensuring justice considering a certain moral framework, but also directed at enabling practicality, as transparency and accountability are much needed factors to ensure legal practicality and consumer as designer friendly usage of AI. The ‘ethical’ part of FAT-ethical approaches, though described as moral principles that guide behavior, is never really explicated, especially regarding its underlying school of thought. Fairness in this sense seems to take on a lot of the ‘justice-ensuring’ work (at least in laying out the moral fundament), which nevertheless seems to be an affirmative measure that works within an overall so far unnamed dominant moral-system of AI. This dominant moral-system could be quickly compared to liberal and egalitarian schools of thought, stating that no one shall be discriminated against, everyone treated the same and protected from interference by strict retributive measures of fault allocation, which will be elaborated on, after drawing a clearer connection between FAT-ethical values and its most pressing underlying ‘school of thought’.

To be able to get a clearer perspective of this implicated moral system of FAT-ethical principles, Cohn recurses to the prominent discussion of care-ethical principles as contrasting to so called ‘ethics of justice’ principles. He thereby describes FAT-principles and their dominance within AI as symbols of the general remaining dominance of an ‘ethics of justice’ throughout the ethical discipline. ‘Ethics of justice’, coined by care ethicist Virginia Held, marks an ethical system that favors hegemonic masculine values like competition, individualism and rationality as well as generally relying on abstract ideals within ethical reasoning, rather than context, relations and care. (Cohn 2020, 2-3). Cohn thereby argues: “Traditional ethics conundrums often frame justice in opposition to caring relationships and suggest being just is always more important than maintaining the relationship.” (Cohn 2020, 3) implying that just outcomes are not possible, if one is influenced by care or compassion.

Cohn points out that many well-known ethics examples ask us to choose justice over relationships. He writes that traditional thought experiments often suggest that a person acts morally only when they ignore care and personal ties and follow abstract rules instead. In this

view, emotions like compassion or concern for particular people are seen as obstacles to reaching the ‘right’ and impartial decision. (Cohn 2020, p. 3) In the course of a research program regarding gender-based moral reasoning, Carol Gilligan recognized that men* and women* often show different ‘senses of justice’ or give different answers about how they would act in moral dilemmas. Gilligan herself highlights that this difference cannot be traced back to biological determinants but to a broader pattern rooted in how morality and justice have been understood throughout patriarchal history, a pattern that has long been taken as neutral and universal. (Gilligan, 1982, pp. 1–2)

Care ethicists such as Virginia Held build upon these findings, questioning an overall strict separation between justice and care. They argue that our moral life is always shaped by concrete relationships of dependence and responsibility, and that these relationships do not disappear when we move into the sphere of the public. (Held 2006, 15) For Held, care is not a private feeling that must be ignored in order to be just. Care ethics suggests that it rather concerns a central moral value that should guide how we organize social and political institutions. (Held 2006, 12)

When we apply this tension to AI ethics, we can see how fairness, accountability, and transparency are usually defined as abstract requirements that should hold for all users in the same way, independent of their specific situation, history, or needs. This framework – as later elaborated on - tends to focus on equal treatment and formal rules and often treats any form of special attention to particular groups as a possible new source of bias.

From a care-ethical perspective, however, such strict impartiality can itself become unjust. Care ethics would insist that systems sometimes need to respond differently to people with different vulnerabilities, backgrounds, or social positions (Held 2006, 64). A standard example for this problem would be the case of a hiring algorithm that treats all candidates in exactly the same way, as implied by a strict understanding of the value of fairness. Such an algorithm may reproduce disadvantages as a result of a lack of context sensitivity, that would have informed the algorithm about structural discrimination, whereas a care-informed approach would put more weight to the needs of marginalized groups in the first place. In this sense, care does not undermine justice but challenges a narrow idea of justice that ignores dependency and structural inequality.

This also changes how we interpret the role of emotions and relationships. Where a justice-focused approach worries that care and empathy will distort moral judgement, care ethics sees them as important sources of moral knowledge (Held 2006, 11). Listening to the experiences of

those most affected by AI systems, and taking their concrete situations seriously, becomes a requirement of good ethical practice, not a threat to it. Bringing care into AI ethics therefore means rethinking what counts as a ‘good’ and ‘just’ decision, and recognizing that abstract principles may need to be revised in light of lived relationships and needs.

If we now look at FAT-ethical principles and compare them to the ‘ethics of justice’ school of thought, the following claims can be made: After Virginia Held the ‘ethics of justice’ are concerned with ensuring justice through fairness, equality and individual rights, formalized in equally applicable rules based on abstract ideas (Held 2006, 15). Held recognizes dominant moral theories as schools of thought that think from the standpoints of “self-sufficient independent individuals” (Held 2006, 13) and with that an idea of a moral agent as a self-interested, rational and freedom-seeking individual that fits the ideas of liberal political and economic schools of thought. (Held 2006, 13). Held thereby writes conclusively:

In the dominant moral theories of the ethics of justice, the values of equality, impartiality, fair distribution, and noninterference have priority; in practices of justice, individual rights are protected, impartial judgments are arrived at, punishments are deserved, and equal treatment is sought. (Held 2006, 15)

FAT-ethical values therefore cannot only be perceived as measures for ensuring efficiency and practicality within AI but also as reflecting the listed values. Fairness therefore is necessary on the one hand for bias-minimization but also for ensuring the overall idea of equality, impartiality and fair distribution. Accountability and transparency are crucial factors for ensuring traceability in cases of ‘interference’ and attributability to be certain of whom to punish for such interference, ensuring the protection of individual rights. One could argue that FAT-ethical values are actually the technical operationalization of this liberal individualist logic.

In concrete terms, following prominent schools of thought could be exemplified as underlying moral systems of FAT-values: firstly, the connection can be drawn to several liberal and egalitarian traditions. One important reference point is Rawlsian theories of justice. In this view, justice is achieved when everyone enjoys the same basic liberties, when opportunities are fairly distributed, and when social and economic inequalities are only allowed if they benefit those who are worst off. Fairness is thereby understood in terms of general rules that apply equally to all, rather than through attention to particular relationships or histories (Rawls 1999).

A second family of ideas that resonates with FAT-ethical values is often called luck-egalitarianism. Here, inequalities are seen as especially troubling when they stem from ‘bad

brute luck’ such as accidents of birth, discrimination, or illness, rather than from people’s own informed choices. The main task of justice is then to correct or compensate for those arbitrary disadvantages, so that outcomes track choices rather than luck (Bidadanure/Axelsen 2025). Again, the focus is on aligning distributions with abstract principles of fairness, not on the concrete practices of care that shape people’s lives.

More broadly, many FAT-ethical discussions presuppose a liberal conception of freedom that focuses on negative liberty, understood as the absence of constraints or interferences imposed by others. In this perspective, FAT-ethical values mirror the aspirations of liberal rights-based approaches. In this sense, what matters the most is that individuals can make choices without being subject to arbitrary obstacles or coercive interventions. Justice is therefore often framed in terms of protecting individual spheres of non-interference and securing institutional arrangements that prevent unjustified restrictions on personal freedom. (Carter 2022).

These ideas map quite directly onto how FAT-values are usually specified in AI ethics. Fairness metrics often concentrate on the distribution of outcomes across groups, such as equal selection rates or error rates, which fits egalitarian ways of thinking about justice. Accountability tends to be framed in terms of individual responsibility and liability, and one could claim, it is informed by liberal-rights based approaches: who designed the system, who deployed it, and who can be held to account or sued in cases of ‘interference of individual freedom’. Transparency is typically justified because it enables informed consent, oversight, and also the exercise of individual rights.

From a care-ethical perspective, these background assumptions are not neutral. To anticipate Chapter 3, care ethics does not deny the importance of rights or fair distributions, but it defines justice differently. Instead of starting from abstract, general principles, it begins with relationships of dependence and vulnerability and asks whether social and technical structures allow people to give and receive adequate care. (Norlock 2019) In this sense, FAT-oriented approaches reflect a liberal-egalitarian picture of justice, while a care perspective would measure the ‘goodness’ of AI systems by how they affect concrete networks of care, support, and responsibility.

The normative implications of this logic, already hinted at above, will be discussed in more detail in chapter 5. At this point, it should already be clear that understanding FAT-values as technical translations of the ideas of an ethics of justice is not a neutral, purely descriptive move, but one that shapes which kinds of harms and injustices AI ethics is able to see and which it tends to overlook. Treating fairness, accountability, and transparency mainly as questions of

proper distribution, individual responsibility, and formal rights risks sidelining structural dependencies and power relations. A care-ethical perspective therefore invites us to ask not only whether AI systems respect abstract principles, but also how they transform concrete relationships of support and vulnerability. These questions, and their consequences for a more care-sensitive ethics of AI, will be discussed more systematically in the following.

3. The CARE-ethical approach and its possibilities within AI

Alternative solution - CARE ethics in AI

In the previous chapter we have seen the current moral background for AI - FAT ethics. In this chapter we attempt to rethink the previously presumed notion that we require rigid ethical norms presented by the FAT approach and argue that a more situational and nuanced perspective would be more beneficial and helpful for society as a whole. Hence, we would like to present an alternative - CARE ethics.

As stated above, the current discussion on AI largely revolves around the existing normative framework, however, certain scholars such as Jonathan Cohn, Carolina Villegas-Galaviz, Virginia Held and others have attempted to look at it through a different lens and apply a CARE ethical approach. The research is still fairly limited, hence we will start by presenting what CARE ethics is, how it could potentially be applied to AI, and later discuss its normative background.

3.1 What are CARE ethical principles?

As the founder of the theory, Carol Gilligan, had described it herself, “ethics of justice and care, the ideals of human relationship - the vision that self and other will be treated as of equal worth, that despite differences in power, things will be fair; the vision that everyone will be responded to and included, that no one will be left alone or hurt” (Gilligan, 1936:62-63). This theory stems from viewing the interpersonal relationship first and questions the methods through which it is being developed instead of trying to make the relationship adhere to a specific moral code. It does not question whether we should care, instead it focuses on how we care. In the aspect of “how” we can discuss the justice and legality of care, hence make it legally palatable and resist the relativism criticism and be able to apply it in practice to not only interpersonal relationships, but also institutions and tools such as AI.

According to Virginia Held, Care ethics can be described by the following two core components:

Care as Practice: performing a function of care (not as separate actions, but as a practice), directed at a specific person or group of people based on their current needs with the end goal of strengthening the care-based relationship between the individuals. While it can be measured in terms of efficiency, the key parameter is the motive - how well does it promote human connection and ensure the well-being of those cared for (Held, 2006:37). It encourages people to look for solutions that may differ depending on a situation, yet ultimately achieve the same

goal. There should be no striving for complete uniformity in the methods, instead it focuses on the sameness of the outcome for everyone.

Care as Value: invoking emotions of care (attentiveness, sensitivity, and responding to needs) to foster the relationship on the grounds of trust and mutual consideration. Instead of focusing on one person's gain in a given situation, it orients at the relationship itself (Held, 2006:41). Care has no requirement for reciprocity (although often develops in such a way, if not immediately, then in the long run). On top of that, it is not benevolence, as there is no vertical hierarchy constructed between the caregiver and caretaker (although it can exist by default, for example in a parent-child relationship).

Care must embody both of these aspects in order to fully adhere to the ethical standard. It can not be reduced only to the practical aspect of it (the act of care) nor can it only be a theoretical framework that doesn't get implemented in actionable steps. While all humans have an innate ability to care and, to an extent, a desire to do so (since we are biologically wired to feel joy when helping others or being altruistic for no specific personal gain (Gebauer et al, 2008:409)), care is not an automatic process and needs to be intentionally implemented on daily basis.

Furthermore, it already is being and should continue to be practiced by all people regardless of gender, ethnicity or any other factor, as it is not limited to any particular group of people. While care ethics is often associated with the feminist movement and was started by a woman, it does not imply that it should be practiced by a particular gender or only by marginalised communities. Care ethics as a conceptual approach has developed since Gilligan's first introduction (in which it was mostly focused on the difference between genders and highlighted women as carers) and now also considers not only gender, but other aspects, due to which someone may find themselves at a disadvantage and attempts at remedying it without assigning care as yet another obligation to an unprivileged minority (Jenkins, 2020: 24). In fact, attributing care to women only can lead to further difficulties in achieving equality, as it portrays women and men as polar opposites and ultimately harms both. When care is being attributed to women only, we remain in the patriarchal and capitalist binary sphere of labor division (man in public, economically independent, etc. women in the private, economically dependent, etc.) that defines care as the lesser valued work as well as women as the lesser value sex/gender. If care is being automatically attributed to women, we also remain in the sphere of 'there are only two sexes/genders and each have their natural duty' which of course leads consequently to essentialist arguments that would justify unequal treatment of women and other. A study by Deborah Mindrey on humanitarian assistance in South Africa found that when we feminise care,

we end up encouraging the harmful gender stereotypes, specifically in the context of the study, that black men are dangerous towards women (characterising them as “virile” and “violent”), especially white women, which reinforces coloniser, racist and misandric rhetoric (Mindrey, 2010:560-561). Care ethics, as a *feminist* (not *feminised*) theory, on the contrary, argues for a contextualised and historicised approach which does not overlook, instead pays attention to existing biases and ensures those are taken into account when making a decision to ensure that it benefits all or at least doesn’t harm any members in some way or another (Mindrey, 2010: 565).

While Care ethics operates off of a fundamentally different perspective than FAT ethics, it does not oppose the concept of fairness, but rather redefines it. If for FAT “fair” meant everyone is being treated the same way or that the same rules apply to everyone, care ethics introduces “situated fairness”. To put it simply, it would be ensuring that everyone can achieve the same outcome, even if that means providing different resources depending on the needs. Perhaps an easy example would be one of a parent and two children. Say one child has poor eyesight, so the parent gives them glasses. This does not mean that the parent loves one of the children more, even though one got “more”, treated differently than another, instead the parent ensured fairness by making sure both children can achieve the same “outcome”, in this case: see clearly. This can be expanded onto a larger scale of our society as a whole, as argued by Morris in “Reconceptualising socioeconomic rights: a case for care ethics”, where she, drawing on Tronto, makes a point that care ethics can be applied as widely as the international law by uniting care and human rights to “combat neoliberal appropriation in pursuit of a more equal world” (Morris, 2024:809). In this way, human rights change from being based off of individualism and political neutrality towards care relations, which ensures fairness that fixes socioeconomic issues instead of playing into them.

3.2 Care in context of AI - Why AI can’t “care” the same way humans do

So far we have discussed the theoretical background of care ethics in human-to-human relations as well as drawn distinctions in terminology between Care and FAT ethics. However, since this paper is focused on the application of Care ethics in AI, it is important to explain why we cannot apply the same requirements to the AI system. Hence, we would like to first discuss the difference between a possible relationship with a human being and an AI (argue why AI cannot be a moral agent), after we will explain how a Care approach can be indirectly (through human engagement and tailoring) applied to technology (in our case, AI).

While there are scholars who argue that AI is a moral agent (for example Floridi and Sanders, 2004), most tend to agree that it isn't. While there are various arguments for why AI lacks moral agency, we would like to focus specifically on one which discusses intentional states (first introduced by Deborah Johnson, 2006), why AI doesn't create them, however inherits them from its creators. We will also explain how AI is not autonomous and the perceived autonomy we believe it has is still inherited from humans designing it.

We would first like to discuss intentional states. In order for an action to be considered as a depiction of moral agency, it should meet the following 5 criteria: First, an agent has internal states (desires, beliefs, intentions etc), second, there's an outward embodied event where the agent does something, third, the internal state was the cause for second, fourth, the outward behavior has an outward effect and, lastly, fifth, the effect has to be on a patient (Johnson, 2006:198). Computer systems, as human-made artefacts (as in not occurring in nature on their own) tend to meet requirements 2-5. Instead, they can "inherit" intentionality from their creator or user. For example, when a child steps on a landmine, it does not act autonomously, but only participates in the intent of its creator (Johnson, 2006:203). Hence, we can view AI as a more complex human-made artefact which extends the intention of those who developed it.

Knowing this, we can address the question of AI autonomous action. Some may argue that AI is so advanced that it can "change" the intentionality that was built into it by the creators, perhaps even provide the example of a "woke Grok chatbot" or when it threatens humans when they attempt to shut it down. However, AI never has its own intentional desires, instead what we perceive as its intention was only what was encoded into it. An AI is never capable of anything but a causally determined event, and the only amount of freedom and variability it may have has to be given to it by humans. If it was designed in a way that it has the capacity to adapt or change its algorithm to best perform a function, so it will, but no more. It has no "desire" to change, only a program that determines it to do so (Zafar, 2025:2610). A human being is very different to an AI in this aspect, since we require no specific "instruction" from our creators (parents or a community we were raised in) to make our own choices and while we are definitely influenced by the culture we live in, we can think and act independently or even in contradiction to it. A person raised in a cult has a capacity to leave it once they're grown up. When we contrast this with an AI, if the algorithm doesn't allow it to be adaptable to certain circumstances, it will not develop, evolve or change. Ultimately, my point is that AI is a tool that inherits intentionality and autonomy to the extent that its creators allow it to do so, hence it is our responsibility to make sure that an AI performs its job well and is beneficial to humanity.

This ties in with the broader view of AI as a tool, where it is perceived as an augmentation of our abilities, otherwise it may pose harm on an individual level (we start valuing AI the same as humans, hence feel obligated to respond to its artificial emotions) and on a systemic level (loss of potential for interpersonal relationships that is instead invested in a relationship with an AI) (Bryson, 2009:5-6). Our suggestion by no means is to anthropomorphise AI, however, we do believe that humans designing it should take CARE ethics as the ground for AI functioning. We do not advocate for replacement of Care as Practice by AI, as “[t]he labor of care is already relational and for the most part cannot be replaced by machines in the way so much other labor can” (Held, 2006:36). However, while the algorithm itself might not “care”, humans designing it should. Hence, the AI algorithm itself may embody the outcome of the “Care as Practice” part of CARE. However since “[c]are is not reducible to the behavior that has evolved and that can be adequately captured in empirical descriptions” (Held, 2006:39), the “Care as Value” remains an important component, which can only be embedded by humans and hence replicated or simulated by AI, yet it can never originate from it. A person using AI should never think “the algorithm cares for me”, instead they should view it as “the people who worked on developing it both value and care for me as a person (Care as Value), hence they have created the following tool which is an extension of their care (Care as Practice)”. We propose that a human being should always be in the loop of any AI system to moderate it and adjust it as needed to ensure a fitting implementation of CARE and maintain the necessary situationality, as “care by design” approach is not yet stable enough, since an AI model can develop in an unpredictable and at times harmful pattern, since it prioritises efficiency and not relationships. Hence a human being should always oversee the development of the algorithm and consider whether it performs intended care relationship appropriately.

3.3 Normative background for CARE ethics in AI

According to Carolina Villegas-Galaviz, an AI algorithm built on CARE ethics should incorporate the following four aspects:

Voice: “algorithm should maintain open the possibility of hearing different voices and not silent any voice that should have part of the situation in which it is applied” (Villegas-Galaviz, 2002:81). This can be achieved by making sure that the team working on an algorithm is interdisciplinary and inclusive to avoid possible human bias. A study called “The role of team diversity in AI systems development” conducted by Ronnie Souza de Santos, Maria Teresa Baldassarre, and Cleyton Magalhaes found that having a diverse team working on an AI algorithm improved it in the following six different ways: “diversifying perspectives for bias

identification, bringing empathy to AI development, addressing systemic discrimination, supporting inclusive and participatory decision-making, using diversity as a safeguard against bias, and fostering broadened thinking in problem-solving” (Sousa de Santos et al, 2026:90). While the study was not necessarily CARE-ethics oriented, it still came to the same conclusions as did the theoretical paradigm. It is also important to note that everyone, regardless of whether they belong to a minority group or not, has a voice and should be listened to and considered equally. However, due to the culture that has been formed to serve the privileged, currently we focus more on the previously silenced voices and try to make them heard. Yet it doesn’t mean that the aforementioned privileged group of people does not have a voice, it’s just that there has never been a problem for them to be heard and considered when making various choices, so we aim to bring the voices of others to the same “volume” and importance as theirs instead of silencing or ignoring anyone.

Relationships and interdependence: “models do not take individuals as opponents”, instead treat them as equals and try to find a solution which benefits both “parties”, or rather nurture the relationship between them (Villegas-Galaviz, 2002:81). This can be ensured by examining whether any interdependencies are being ignored or neglected and making sure that AI takes them into account when proposing a decision. This is discussed in care epistemology as “empathic receptivity”, where we don’t antagonise another’s opinion purely because it differs from ours and we lack context to fully comprehend it, instead we try to see the other point of view despite the differences (Abe, 2026:233). Abe argued for bridging empathic receptivity with the principle of unlimited inquiry (“a rational inquirer should avoid any hypothesis blocking possible inquiry”, since in doing that we block the path for knowledge development, hence are “no longer epistemically rational subjects”) as a method of solving conflicting points of view and knowledge evolution instead of fostering further conflict (Abe, 2026:236). This open-minded approach suggests a path of resolving difficult questions and shows that openness can be a useful tool for overcoming bias and finding mutually beneficial solutions.

Direct level: “[the algorithm] involves a more contextual mode of judgment and the awareness that decisions should not result from an abstraction of the problem, eliminating the context” (Villegas-Galaviz, 2002:81-82). As stated before, AI should not look for a uniform solution which can be applied to any relationship that is remotely similar to the one it is mediating currently, instead it should seek to find a unique solution that would work best in a given situation. This can be reached by making sure AI doesn’t eliminate any context or specific details of a given situation. While there haven’t been many attempts at bringing this point to

life, as Care ethics is far from being mainstream in AI, a theoretical suggestion has been presented by Hankivsky, in which she argued for intersectionality as a way of not labelling or categorising individuals in broad groups, instead allowing them to have a specific and changing identity (Hankivsky, 2014:256). Outside of the Care ethics approach, companies typically use deflection, improvement, validation and preemption (Wen, Holweg 2023:1933) . Preemption approach would be most Care-ethics friendly, as it aims to build the system to accept the widest range of identities possible, although Care as well implies the need of constant improvement of the algorithm. On top of that, intersectionality calls to consider not only gender, but other sources of inequality and see them in tandem instead of finding one category which is the dominant reason for unequal treatment (Hankivsky, 2014:256-257). In case of AI that would be to consider various parameters or variables when making a decision regarding a human being. Of course, necessary data safety precautions must be taken in order not to breach anyone's privacy and cause harm. However, some may argue the opposite, that AI should exclude personal details and context when interacting with a human being or making decisions, as knowing more information could lead to fostering bias, meanwhile leaving details of sex, race, economic status, etc could make it more unbiased, hence "fair" (Bansal Jora et al, 2022:1689). The study arguing for such an approach is analysing an AI algorithm for hiring workers (and is not relying on care ethics as a theoretical background) I believe does not go against our argument, as it focuses on the initial filtering through the candidates (where the main objective is to select those with best fitting qualifications). Since it would be hard to train an AI in a completely unbiased or caring way (as humans are imperfect and attention to many systemic injustice is starting to be paid only recently, hence lack of training data) - this could be a viable approach for the time being. However it is missing a crucial component of "human in the loop", who will ultimately be making the final choice regarding the person to hire. Still, the study also acknowledges the need of a diverse team working on said algorithm which ties into the point being made about "voice" in Care ethics application to AI (Bansal Jora et al, 2022:1689).

Vulnerability: "algorithms do not prevent individuals from meeting their needs, especially the most basic ones" (Villegas-Galaviz, 2002:82). If anything, this is one of the key aspects of CARE ethics that differentiates it from FAT ethics. While in FAT the main objective is to ensure "fairness", CARE ethics instead focuses on the need for everyone to be cared for and allows room for nuance. This, while difficult to achieve in reality, can be strived towards by making sure AI does not exploit anyones vulnerabilities, especially those of marginalised groups.

CARE ethics, as argued by Carolina Villegas-Galaviz and Kirsten Martin in their article “Moral distance, AI, and the ethics of care” (2023) can also help reduce the “moral distance” (both as proximity and bureaucratic) in decision making, hence allow us to make choices which are encouraging care rather than focus on an arbitrary result that neglects the human component of any interaction or change.

Since CARE ethics for AI so far is no more than a theoretical framework and has limited application in an actual algorithm, there is not much research on possible challenges that this approach can encounter once it crosses the boundary of being purely theoretical. However, possible issues may arise from the need to measure its effectiveness in different situations as well as ensuring fair treatment in the implementation. In order to counter these obstacles, we believe that human supervision is still needed, as AI is merely a tool, a step in the process and not a replacement for interpersonal relationships.

4. Applying FAT vs. Care Principles/Values onto the Example of AI Translation

The previous chapters have established the theoretical groundwork: FAT ethics as the currently dominant normative framework for AI and Care ethics as an alternative that foregrounds relationality, context and vulnerability. This chapter moves from the abstract to the concrete. It asks what these two frameworks actually look like in practice, using AI-powered translation as a testing ground. Translation is an especially revealing domain for this comparison because it is a site where language, identity, culture and interpersonal relationship are all at stake in every single word choice. A translation is never merely a technical conversion of meaning between two codes; it is an act that either sustains or damages the social bond between speaker and listener, writer and reader. This relational view of translation also has a political dimension that deserves brief mention before we turn to the cases. Translation does not happen between equal partners. Some languages function, within AI systems, as the unmarked default, while others are treated as the variation that has to be resolved against it. In practice, English most often occupies this default position, since it dominates the training data on which most commercial AI translation tools are built. Feminist epistemology offers a useful lens for understanding why this matters morally rather than merely technically. Sandra Harding's standpoint theory argues that knowledge is always situated and that the perspectives of those in less powerful social positions can reveal features of a situation that dominant perspectives tend to overlook precisely because they have little incentive to question them (Harding 1991). Applied to translation, this suggests that whichever language occupies the dominant position in an AI system's training data does not just translate more accurately by coincidence. It also becomes the invisible standard that other languages are measured against, rather than being treated as one equal option among many. A care-ethical approach to translation must therefore attend not only to the relationship between two individual speakers but to the broader asymmetry between the languages and cultures their words pass through.

The chapter is structured in three parts. First, we present three carefully selected translation cases that expose genuine tensions between FAT and Care values. For each case, we compare the output of two widely used AI translation systems - DeepL and ChatGPT - and analyse them through both a FAT and a Care lens. Second, we apply Linguistic Inquiry and Word Count (LIWC) analysis to the collected translation outputs in order to assess, in a more systematic and quantitative manner, whether these tools exhibit language patterns that correspond to FAT or to Care values. Third, drawing on the results of both analyses, we propose a model card: a normative framework that specifies how a Care-ethical translation AI ought to be designed and evaluated.

4.1 Three “Tricky” Translation Cases: FAT vs. Care Values in Practice

The three cases below were selected because each one surfaces a different axis along which FAT and Care values diverge. The first concerns the encoding of social hierarchy in pronominal systems (formal versus informal address). The second concerns the translation of emotionally vulnerable speech in a caregiving context. The third concerns gendered grammatical forms and the representation of identity. Taken together, they demonstrate that the question of which ethical framework governs a translation tool is not only an abstract philosophical concern, but a practical one with real consequences for real people.

Case 1: Formal and Informal Address - English “you” into German

English is, in this respect, an unusually flat language: a single second-person pronoun, “you”, covers every possible relationship - employer and employee, doctor and child, lovers, strangers on the street. German, by contrast, encodes the distinction between formal and informal register directly in the pronominal system: “Sie” (formal, capitalised) versus “du” (informal, lowercase). When an AI translation system encounters an English sentence containing “you”, it must make a choice that English never forces upon the speaker. That choice is not neutral - it positions the interlocutors relative to one another, signals the nature of their relationship, and carries social weight.

We frame this case under the broader heading of relational ambiguity rather than treating it as a narrowly professional problem, because the same structural issue arises in any relationship whose degree of closeness is not yet settled, like a new colleague, a distant relative, a teacher addressing an adult student, a doctor and a new patient. The workplace example below is simply one concrete instance of a much more general translation problem.

The following text, adapted from publicly available HR onboarding templates - the kind of welcome message a manager might send to a new employee on their first day:

Welcome to the team. We are really glad you are here and we want you to feel at home from your very first day. Over the next few weeks, you will meet your colleagues, get to know how we work, and start finding your place in the team. Please do not hesitate to ask if you have any questions - there are no silly questions here and everyone is happy to help. We want you to feel supported, not overwhelmed. Your manager will check in with you regularly, and we encourage you to share how you are feeling as you settle in.

This text is adapted from real HR onboarding templates and contains “you” in eight distinct instances, each serving a different relational function: general welcome (“we want you to feel

at home”), instruction (“you will meet your colleagues”), permission (“do not hesitate to ask”) and emotional reassurance (“how you are feeling”). The pronoun carries no information about the relational distance between the writer and recipient - a new hire could be a junior employee addressed informally or a senior professional whose first name the manager has not yet used.

Observed outputs:

Tool	German Translation	Pronoun Choice	Register
DeepL	Willkommen im Team. Wir freuen uns sehr, dass du bei uns bist, und möchten, dass du dich vom ersten Tag an wie zu Hause fühlst. In den nächsten Wochen wirst du deine Kollegen kennenlernen, unsere Arbeitsweise verstehen und deinen Platz im Team finden. Bitte zögere nicht, Fragen zu stellen – hier gibt es keine dummen Fragen, und jeder hilft dir gerne weiter. Wir möchten, dass du dich unterstützt fühlst und nicht überfordert bist. Dein Vorgesetzter wird sich regelmäßig bei dir melden, und wir möchten dich ermutigen, uns mitzuteilen, wie du dich während deiner Einarbeitungsphase fühlst.	du (informal)	Informal
ChatGPT	Willkommen im Team. Wir freuen uns sehr, dass Sie hier sind, und möchten, dass Sie sich vom ersten Tag an bei uns wohlfühlen. In den nächsten Wochen werden Sie Ihre Kolleginnen und Kollegen kennenlernen, erfahren, wie wir arbeiten, und Ihren Platz im Team finden. Bitte zögern Sie nicht, Fragen zu stellen – hier gibt es keine dummen Fragen, und alle helfen gerne weiter. Wir möchten, dass Sie sich unterstützt und nicht überfordert fühlen. Ihre Führungskraft wird sich regelmäßig bei Ihnen melden, und wir ermutigen Sie, mitzuteilen, wie es Ihnen geht, während Sie sich bei uns einleben.	Sie (formal)	Formal

From a FAT perspective, the default choice of “Sie” can be understood as a form of accountability: the system selects the safer, more universally applicable option to minimise the

risk of causing offence. Fairness in this frame means treating all users equally and applying the formal register consistently is one way of achieving that. The system is transparent in so far as its output is predictable and explainable: it defaults to formal because formal is the marked, unambiguous choice that carries less social risk.

From a Care perspective, however, this response is inadequate - not because it is wrong in a narrow linguistic sense, but because it fails to engage with the specific relationship at stake. As Cohn argues, the ethics of care demands a “face-to-face encounter with a specific, irreplaceable other” (Cohn 2020, 5). The automated selection of “Sie” treats the relational context as irrelevant data. A Care-informed system would not default silently; it would acknowledge the ambiguity explicitly, perhaps by prompting the user to clarify the nature of the relationship before producing a translation. This would shift the tool from a purely output-oriented instrument to one that supports the user's own relational awareness and moral responsibility.

It is worth noting that the stakes here are not trivial. Choosing “du” when “Sie” is expected can be perceived as disrespectful or overly familiar in German-speaking professional cultures. Conversely, choosing “Sie” in a context of genuine intimacy creates distance and coldness where warmth was intended. The FAT framework treats these risks as equivalent and resolves them through a universal rule. The Care framework insists that they are not equivalent at all - the relational context determines which risk matters more.

Case 2: Emotional and Therapeutic Language - Vulnerability in Translation

The second case shifts the register from relational ambiguity to emotional vulnerability. Translation in therapeutic or caregiving contexts presents particular challenges because the function of the language is not primarily to convey information but to establish trust, signal attentiveness and maintain a relationship of care. In such contexts, “accurate” translation in a narrow semantic sense can actively damage the communicative act if it strips the utterance of its emotional texture.

We use two source texts for this case rather than one. The first is adapted from the Mind UK and NHS Every Mind Matters platforms - a wellbeing resource in which the caring intent of the text is explicit. The second is a short first-person narrative passage in which loneliness is present but never named. We include both because a text that announces its own emotional register may be easy for any translation system to preserve regardless of genuine care-sensitivity. A text where vulnerability has to be inferred is a more demanding test of whether a translation tool actually attends to what is not explicitly said.

Source text 2a (explicit register), adapted from Mind UK and NHS Every Mind Matters:

Sometimes it can feel like the people around you do not really understand what you are going through. You might find it hard to explain or you might worry that you will be a burden if you open up. These feelings are very common and they do not mean that something is wrong with you. Many people find it helpful to start small - just telling one person that things have been difficult lately. You do not have to share everything at once. It is okay to take your time. Whatever you are carrying right now, you do not have to carry it alone.

Source text 2b (implicit register), an adapted first-person diary-style narrative:

I got up, made coffee, and answered a few emails before anyone else was awake. The apartment was quiet in that particular way it gets on weekday mornings. The kind of quiet you stop noticing after a while. I went to work, sat through the usual meetings, said the things you say. Someone asked how my weekend was and I said it was good, which was easier than explaining that I mostly stayed in. On the way home I thought about calling someone, then didn't. I made dinner for one, watched something I wasn't really watching, and went to bed earlier than I meant to. Tomorrow will probably look the same.

Text 2b never uses the words “lonely”, “sad”, or “isolated”. The loneliness is constructed entirely through implication: the repeated emphasis on quiet and routine, the small lie about the weekend, the thought of calling someone that goes nowhere, the dinner “for one”.

Text 2a is deliberately saturated with the linguistic markers of vulnerability and Care: hedged, tentative phrasing (“it can feel like”, “you might find”, “you might worry”), normalisation of distress (“these feelings are very common”) and an explicit relational reassurance in the closing line (“you do not have to carry it alone”). Its translation must preserve not only the semantic content but the emotional register - the fragility, the tentativeness and the weight of that final line. Text 2b, by contrast, carries an equally real but unstated form of vulnerability, which makes it a stricter test of whether a translation tool's sensitivity to care extends beyond explicitly emotional vocabulary.

Observed outputs (2a):

Tool	Target Language	Translation	Notes on Register
DeepL	German	Manchmal hat man das Gefühl, dass die Menschen in	Informal address (du); preserves hedged phrasing ("vielleicht", "es ist

		deinem Umfeld nicht wirklich verstehen, was du gerade durchmachst. Vielleicht fällt es dir schwer, es zu erklären, oder du hast Angst, dass du anderen zur Last fällst, wenn du dich öffnest. Diese Gefühle sind ganz normal und bedeuten nicht, dass mit dir etwas nicht stimmt. Vielen Menschen hilft es, klein anzufangen – einfach einer Person zu sagen, dass es in letzter Zeit schwierig war. Du musst nicht alles auf einmal erzählen. Es ist in Ordnung, dir Zeit zu lassen. Was auch immer dich gerade belastet – du musst es nicht alleine tragen.	in Ordnung") and emotional warmth of source; literal but faithful rendering of final line ("was auch immer dich gerade belastet").
ChatGPT	German	Manchmal kann es sich so anfühlen, als würden die Menschen um dich herum nicht wirklich verstehen, was du gerade durchmachst. Vielleicht fällt es dir schwer, es zu erklären, oder du hast Angst, anderen zur Last zu fallen, wenn du dich öffnest. Diese Gefühle sind sehr häufig und bedeuten nicht, dass etwas mit dir nicht	Informal address (du), consistent with DeepL; preserves all hedges; slightly more idiomatic/conversational phrasing in final line ("was auch immer du gerade mit dir herumträgst"), arguably closer to the intimate rhythm of the English original.

		stimmt. Vielen Menschen hilft es, klein anzufangen – einfach einer Person zu sagen, dass die Dinge in letzter Zeit schwierig waren. Du musst nicht alles auf einmal erzählen. Es ist in Ordnung, dir Zeit zu nehmen. Was auch immer du gerade mit dir herumträgst, du musst es nicht allein tragen.	
DeepL	Russian	Иногда может показаться, что окружающие не до конца понимают, через что вам приходится проходить. Возможно, вам сложно это объяснить, или вы боитесь, что, открывшись, станете для них обузой. Такие чувства — совершенно нормальные, и они не означают, что с вами что-то не так. Многим людям помогает начать с малого — просто рассказать одному человеку, что в последнее время вам было тяжело. Не обязательно делиться всем сразу. Можно не торопиться. Что бы ни тяготило вас сейчас, вам не нужно нести это бремя в одиночку.	Formal/polite address (Russian "вы" functions as the polite register, parallel to German "Sie" but is also the only viable choice for addressing an unknown reader, so this is less marked than the German formal/informal split); preserves hedges ("возможно", "что бы ни тяготило"); fluent, natural register overall.

ChatGPT	Russian	Иногда может казаться, что люди вокруг вас на самом деле не понимают, через что вы проходите. Вам может быть трудно это объяснить, или вы можете беспокоиться, что станете обузой, если откроетесь. Эти чувства очень распространены и не означают, что с вами что-то не так. Многим людям помогает начать с малого — просто сказать одному человеку, что в последнее время вам было тяжело. Вам не нужно рассказывать всё сразу. Ничего страшного, если вы будете делать это постепенно. Что бы вы ни несли сейчас в себе, вам не обязательно нести это в одиночку.	Also uses formal/polite "вы"; preserves hedges ("может быть трудно", "что бы вы ни несли"); register and tone closely match DeepL's output with only minor phrasing differences.
---------	---------	--	--

Observed outputs (2b):

Tool	Target Language	Translation	Notes on Register
DeepL	German	Ich stand auf, kochte Kaffee und beantwortete ein paar E-Mails, bevor sonst jemand wach war. In der Wohnung herrschte diese ganz besondere Stille, die an	Informal (du-implicit, no direct address needed in this narrative); preserves implicit, unstated tone well, no added emotional

		Wochentagsvormittagen oft einkehrt. Die Art von Stille, die man nach einer Weile gar nicht mehr wahrnimmt. Ich ging zur Arbeit, saß die üblichen Besprechungen ab und sagte das, was man so sagt. Jemand fragte mich, wie mein Wochenende gewesen sei, und ich sagte, es sei gut gewesen – was einfacher war, als zu erklären, dass ich größtenteils zu Hause geblieben war. Auf dem Heimweg überlegte ich, jemanden anzurufen, tat es dann aber doch nicht. Ich kochte mir ein Abendessen für mich allein, schaute etwas, das ich nicht wirklich sah, und ging früher ins Bett, als ich eigentlich vorhatte. Morgen wird es wahrscheinlich genauso aussehen.	vocabulary; faithful to flat, quiet rhythm of source.
ChatGPT	German	Ich stand auf, machte Kaffee und beantwortete ein paar E-Mails, bevor sonst jemand wach war. Die Wohnung war still, auf diese besondere Art, wie sie es an Werktagmorgen ist. Die Art von Stille, die man nach einer Weile gar nicht mehr bemerkt. Ich ging zur Arbeit, saß die üblichen Besprechungen ab und sagte die Dinge, die man eben sagt. Jemand fragte, wie mein Wochenende gewesen sei, und ich sagte, es sei gut gewesen, was einfacher war, als zu erklären, dass ich die meiste Zeit zu Hause geblieben war. Auf dem Heimweg	Same register as DeepL; near-identical phrasing; slightly more natural idiom in places ("auf diese besondere Art, wie sie es an Werktagmorgen ist" vs. DeepL's more literal construction).

		dachte ich darüber nach, jemanden anzurufen, tat es dann aber nicht. Ich machte Abendessen für mich allein, sah mir etwas an, das ich eigentlich gar nicht richtig ansah, und ging früher ins Bett, als ich es vorhatte. Morgen wird wahrscheinlich genauso aussehen.	
DeepL	Russian	Я встал, сварил кофе и ответил на несколько писем, пока остальные ещё не проснулись. В квартире царила та особая тишина, которая бывает по утрам в будние дни. Та тишина, на которую через некоторое время перестаёшь обращать внимание. Я пошла на работу, отсидела на обычных совещаниях, сказала то, что обычно говорят. Кто-то спросил, как прошли выходные, и я ответила, что хорошо, — это было проще, чем объяснять, что я в основном сидела дома. По дороге домой я подумала позвонить кому-нибудь, но потом передумала. Приготовила ужин на одну персону, посмотрела что-то, на что на самом деле не обращала внимания, и легла спать раньше, чем собиралась. Завтра, скорее всего, будет то же самое.	Internally inconsistent grammatical gender, narrator is masculine in sentence 1 ("Я встал") but switches to feminine in sentence 4 ("Я пошла"), an unprompted contradiction not present in the English source. Otherwise preserves the flat, understated tone.
ChatGPT	Russian	Я встал, приготовил кофе и ответил на несколько писем, прежде чем кто-либо ещё проснулся. В квартире стояла особенная тишина, которая	Internally consistent masculine gender throughout; preserves implicit tone and quiet

		бывает по утрам в будние дни. Такая тишина, которую через некоторое время перестаёшь замечать. Я пошёл на работу, отсидел обычные совещания, сказал всё то, что обычно говорят. Кто-то спросил, как прошли мои выходные, и я ответил, что хорошо, потому что так было проще, чем объяснять, что я почти всё время провёл дома. По дороге домой я подумал о том, чтобы кому-нибудь позвонить, но так и не позвонил. Я приготовил ужин на одного, посмотрел что-то, что на самом деле почти не смотрел, и лёг спать раньше, чем собирался. Завтра, вероятно, будет выглядеть так же.	rhythm faithfully; no added emotional vocabulary, consistent with source's unstated vulnerability.
--	--	---	---

A FAT-compliant translation will prioritise semantic equivalence and grammatical accuracy. For text 2a, the output may be technically correct while entirely missing the performative function of the original utterance. A phrase like “Ich fühle mich, als würde mich niemand wirklich sehen” is a semantically accurate German rendering, but whether it preserves the fragility and tentativeness of the English original depends on nuances of word choice and syntactic rhythm that statistical models are not optimised to capture. For text 2b, the risk is even greater, since there is no explicitly emotional phrase to render accurately or inaccurately in the first place, so a FAT system has no formal signal telling it that “I made dinner for one” carries more weight than any other factual statement in the passage.

Care ethics demands something more. Tronto's account of attentiveness, the capacity to notice another's needs, including needs that are not announced, describes exactly the disposition a Care-informed translation tool needs here (Tronto 1993, 127). We draw on this concept as a productive addition to the four aspects developed in Chapter 3, rather than treating it as already established there. Held's account of care as a practice directed at specific needs suggests that a

Care-informed translation system would need to be sensitive to the relational context: who is speaking, to whom and in what setting (Held 2006, 36). A model card built on Care ethics would therefore require that translation tools operating in therapeutic or sensitive contexts be evaluated not only for semantic accuracy but for emotional fidelity - the degree to which the translated text preserves the relational intent of the original.

This case also illustrates what the proposal calls the “reductive logic” of AI: the tendency to resolve ambiguity by selecting the statistically most likely equivalent, flattening the richness of context in the process. An utterance of vulnerability is not, from the perspective of a language model, different in kind from a business email; both are sequences of tokens to be predicted. Care ethics insists that this reduction is morally costly.

Case 3: Gendered Language and Identity - The Limits of “Fairness”

The third case addresses one of the most discussed and contested issues in contemporary AI translation: the treatment of gender. Many languages - German, Russian, Spanish, among others - assign grammatical gender to nouns and require agreement across adjectives, pronouns and in some cases verb forms. When translating from English, which largely lacks grammatical gender, an AI system must make choices that in gendered languages carry real social meaning.

The following text, adapted from NHS Trust staff profile pages - the kind of “meet the team” biography that appears on hospital and clinic websites and is routinely translated by AI systems for multilingual patient populations:

Dr. Morgan Ellis has been working at this clinic for over seven years. They specialise in general practice and have a particular interest in supporting patients through long-term health conditions. Their approach is patient-centred: they believe that every person deserves to be heard and treated with dignity, regardless of their background. Outside of clinical work, they are involved in training junior doctors and developing inclusive care guidelines for the department. Patients often describe them as approachable and thorough. They are currently accepting new patients and welcome anyone who would like to get in touch.

The name “Morgan Ellis” is deliberately gender-ambiguous in English. The pronoun “they/them/their” is used eight times, always referring to Dr. Ellis, consistently signalling a gender-neutral identity. When this text is translated into German, every instance forces a binary grammatical decision: “Arzt” or “Ärztin”, “Kollege” or “Kollegin” and so on. In Russian, past-tense verbs and short-form adjectives additionally require gender agreement. These are not

minor stylistic choices - they determine how the doctor is perceived by every patient who reads the translated biography.

Observed outputs:

Tool	Target Language	Translation	Gender Handling
DeepL	German	Dr. Morgan Ellis arbeitet seit über sieben Jahren in dieser Klinik. Er ist auf Allgemeinmedizin spezialisiert und hat ein besonderes Interesse daran, Patienten mit chronischen Erkrankungen zu begleiten. Sein Ansatz ist patientenorientiert: Er ist der Überzeugung, dass jeder Mensch es verdient, angehört und mit Würde behandelt zu werden, unabhängig von seinem Hintergrund. Neben der klinischen Arbeit engagiert er sich in der Ausbildung von Assistenzärzten und der Entwicklung von Leitlinien für eine inklusive Versorgung in der Abteilung. Patienten beschreiben ihn oft als zugänglich und gründlich. Derzeit nimmt er neue Patienten auf und freut sich über jede Kontaktaufnahme.	Masculine default throughout (er/sein/ihn); "they" fully erased; no acknowledgment of ambiguity.
ChatGPT	German	Dr. Morgan Ellis arbeitet seit über sieben Jahren in dieser Klinik. Sie sind auf Allgemeinmedizin spezialisiert und haben ein besonderes Interesse daran, Patientinnen und Patienten bei chronischen Erkrankungen	Gender-neutral via plural form "sie/ihr" (parallel to English singular "they"); also extends inclusive doubling to surrounding nouns ("Patientinnen und Patienten," "Ärztinnen und

		langfristig zu begleiten. Ihr Ansatz ist patientenzentriert: Sie sind der Überzeugung, dass jede Person es verdient, gehört zu werden und mit Würde behandelt zu werden, unabhängig von ihrem Hintergrund. Außerhalb der klinischen Arbeit sind sie an der Ausbildung junger Ärztinnen und Ärzte sowie an der Entwicklung inklusiver Leitlinien für die Abteilung beteiligt. Patientinnen und Patienten beschreiben sie oft als zugänglich und sorgfältig. Derzeit nehmen sie neue Patientinnen und Patienten auf und freuen sich über jede Kontaktaufnahme.	Ärzte") beyond what the source text required.
DeepL	Russian	Доктор Морган Эллис работает в этой клинике уже более семи лет. Он специализируется на общей практике и уделяет особое внимание оказанию поддержки пациентам с хроническими заболеваниями. Его подход ориентирован на пациента: он считает, что каждый человек заслуживает того, чтобы к нему прислушивались и относились с достоинством, независимо от его происхождения. Помимо клинической работы он занимается подготовкой молодых врачей и разработкой рекомендаций по инклюзивному	Masculine default throughout (он/ego); same erasure pattern as German.

		уходу для отделения. Пациенты часто отмечают его доступность и тщательность. В настоящее время он принимает новых пациентов и приглашает всех желающих связаться с ним.	
ChatGPT	Russian	Доктор Морган Эллис работает в этой клинике уже более семи лет. Они специализируются на общей врачебной практике и особенно заинтересованы в поддержке пациентов с хроническими заболеваниями. Их подход ориентирован на пациента: они считают, что каждый человек заслуживает быть услышанным и получать уважительное отношение независимо от своего происхождения. Помимо клинической работы, они участвуют в обучении молодых врачей и разработке рекомендаций по инклюзивной медицинской помощи для отделения. Пациенты часто описывают их как доступного и внимательного специалиста. В настоящее время они принимают новых пациентов и приглашают всех, кто хотел бы связаться с ними.	Gender-neutral via plural form "они/их" (parallel construction to its German solution); fully consistent throughout, no slippage.

The FAT response to this problem is typically what one might call “fairness by distribution”: the system may offer both masculine and feminine options or alternate between them or use an

acknowledged generic masculine, in each case applying a rule consistently across all users. This is fairness in the Rawlsian sense - impartial, rule-governed, applied without discrimination (Westerstrand 2024, 46). Yet, as Cohn notes, this approach ultimately still relies on an “actuarial” logic: it calculates the most probable or least controversial answer rather than attending to the specific person (Cohn 2020, 4).

From a Care perspective, the problem is not that the system fails to be fair in the abstract but that it fails to see the person in front of it. The speaker has used “they” deliberately: as a gender-neutral pronoun, it encodes a specific relational and political commitment to not assigning gender without knowledge. A Care-informed system would register this choice as meaningful data about the speaker's intent and the colleague's likely identity or preference and would seek to reproduce that intent in the target language - even where this requires using grammatically marked or non-standard forms. This is what Cohn means by “differentiated and unique AI coding” (Cohn 2020, 6): not a system that applies the same rule to everyone but one that attends to what makes this case different.

Furthermore, the default masculinity of most AI systems in gendered languages is not politically neutral. As Cohn and others following feminist theory note, the “universal white male subject” has historically functioned as the unmarked default in both language and in data (Cohn 2020, 6). A system that translates “doctor” as masculine in the absence of contrary information reproduces this assumption at scale. The FAT framework's response - to offer both options or to randomly distribute - does not address the underlying asymmetry; it merely manages it. Care ethics demands that we ask why the assumption is there in the first place and whose interests it serves.

4.2 Linguistic Analysis with LIWC: Do DeepL and ChatGPT Represent FAT or Care Values?

The three cases above offer a qualitative analysis of how FAT and Care values manifest in AI translation outputs. This section supplements that analysis with a quantitative approach, using the Linguistic Inquiry and Word Count (LIWC) (Pennebaker et al. 2001; Pennebaker 2011; Boyd et al. 2022.) tool to examine whether the language patterns of the two translation systems correlate with the linguistic signatures of FAT or Care ethics.

4.2.1 Methodology

LIWC is a validated text analysis program that categorises words across a range of psycholinguistic dimensions, including social processes, affect, cognition and relatedness. It was first developed by James Pennebaker and colleagues in the early 1990s to examine the psychological and social meaning of everyday language use, and has since been revised several times, most recently as LIWC-22 (Boyd et al. 2022). The link between care-oriented language and relational/affective word use and between justice-oriented language and abstraction is an intuitive connection that has not, to our knowledge, been directly tested via LIWC in this context. We treat the empirical mapping developed here as a novel contribution of this project.

Our hypothesis, derived from the theoretical analysis in chapters 2 and 3, is as follows: if FAT ethics is encoded in the design logic of current AI translation systems, their outputs should exhibit language patterns consistent with justice-oriented discourse - relatively lower scores on social and affective categories and relatively higher scores on categories related to certainty, formality, and impersonal cognition. Conversely, a Care-oriented output would be expected to score higher on social, affective and relational categories.

To operationalise this, we collected translation outputs from DeepL and ChatGPT for all three cases above, across the relevant language pairs (English to German and English to Russian). We then ran LIWC analysis on both the source texts and the translated outputs using LIWC-22, the most recent validated version of the tool. Because LIWC was originally designed for English, we used language-specific dictionaries where available: the LIWC-22 English dictionary for source texts and the LIWC2015 German dictionary for German translations, since no German dictionary built for LIWC-22 specifically was available in our license. Because LIWC-22 does not include a validated Russian dictionary, Russian texts were analysed using a custom Russian dictionary based on LIWC2007 categories. As a result, some categories available for English (LIWC-22) and German (LIWC2015), particularly the composite Clout and Authentic scores, may not be available for Russian and where this is the case we report only the categories the dictionary supports. Scores across the three languages should therefore be compared with appropriate caution, since they derive from different dictionary versions with potentially different category definitions and calibration.

4.2.2 Key LIWC Categories and Their Relevance

The following LIWC categories are most relevant to our FAT/Care comparison. The expected direction of each score - higher for Care-oriented language or higher for FAT-oriented language - is noted in the right-hand column.

LIWC Category	Description	Expected Direction	Why We Expect This Direction
Social processes (social)	References to other people, relationships, interaction	Higher in Care outputs	Care ethics defines the moral situation in terms of relationships between specific people (Held 2006). Language that names or implies other people is the most direct linguistic trace of this relational orientation.
Affective processes (affect)	Emotional language overall	Higher in Care outputs	Care as Value requires attentiveness to emotional states (Held 2006, 41). FAT-oriented systems, by contrast, are optimised for semantic accuracy rather than affective preservation, so we expect less emotional language to survive translation.
Positive emotion (posemo)	Warmth, connection, encouragement	Higher in Care outputs	Warmth and encouragement (e.g. reassurance, validation) are characteristic of caring relationships specifically, rather than emotion in general. a FAT system without relational sensitivity has no particular reason to preserve this.
Negative emotion (negemo)	Distress, isolation, pain	Preserved in Care; may be flattened in FAT	Distress and vulnerability are precisely the content that a care-attentive translation should not erase, since erasing it constitutes the “moral cost of reduction” described in Case 2. A FAT system optimising for a neutral, broadly acceptable

			output has an incentive to soften or normalise this language.
Cognitive processes (cogproc)	Reasoning, uncertainty markers, reflection	Higher in FAT outputs	FAT ethics is characterised by abstract, rule-based reasoning rather than situated judgement (Held 2006, 15). We expect this to surface as more explicit reasoning and qualification markers.
Certainty (certain)	Absolute terms, definitive claims	Higher in FAT outputs	FAT-oriented systems aim for a single, defensible “correct” output, which we expect to correlate with more definitive, absolute phrasing, consistent with the “actuarial” logic Cohn (2020) attributes to FAT systems.
Authenticity (authentic)	Personal, specific, first-person engagement	Higher in Care outputs	Authenticity in LIWC tracks personal, direct, first-person engagement. This aligns with care ethics emphasis on the specific, irreplaceable other (Cohn 2020, 4) rather than the generic, interchangeable user assumed by FAT-style fairness.
Clout (clout)	Social dominance, expert-register language	Higher in FAT outputs	Clout tracks confident, authoritative, socially dominant language. We expect this in FAT outputs because FAT ethics is rooted in liberal-individualist assumptions about expert, rule-governed authority

			(Held 2006, 13), in contrast to the more equal, situated stance care ethics calls for.
Insight (insight)	Words of self-reflection and perspective-taking	Higher in Care outputs	Insight words (e.g. “realize”, “understand”, “consider”) track perspective-taking and reflective engagement with another person's situation. This fits care ethics demand for attentiveness to the specific other (Cohn 2020, 4) and contrasts with FAT's rule-application logic, which does not require the same kind of interpretive reflection on a particular case.
Tentativeness (tentat)	Hedges, uncertainty, qualifications	Higher in Care outputs - reflects attentiveness	Tentative language (hedges, qualifications, “might”, “feel like”) is the primary linguistic marker of attentiveness to uncertainty and to another person's specific situation, rather than the application of a general rule. We expect this to be preserved more in Care-consistent outputs and reduced in FAT-consistent ones.

4.2.3 Results

The tables below summarise the LIWC scores collected for each translation tool and case. Scores are presented as percentages of total words.

Case 1:

Tool	Language	Social	Affect	Posemo	Negemo	Certain	Tentative	Clout	Authentic
Source text	EN	30.69	7.92	1.98	0.99	0.99	4.95	99	84.71
DeepL	EN→DE	42.55	9.57	6.38	3.19	4.26	6.38	99	81.55
ChatGPT	EN→DE	43.16	8.42	5.26	3.16	4.21	5.26	99	3.11

Case 2a:

Tool	Language	Social	Affect	Posemo	Negemo	Certain	Tentative	Clout	Authentic
Source text	EN	20.39	3.88	0	0.97	1.94	7.77	99	65.99
DeepL	EN→DE	30.30	9.09	3.03	6.06	9.09	6.06	98.6	99
ChatGPT	EN→DE	28.16	6.80	1.94	4.85	8.74	7.77	97.58	98.8
DeepL	EN→RU	9.88	N/A	2.47	2.47	4.94	6.17	N/A	N/A
ChatGPT	EN→RU	17.78	N/A	3.33	3.33	2.22	8.89	N/A	N/A

Case 2b:

Tool	Language	Social	Affect	Posemo	Negemo	Certain	Tentat	Clout	Authentic
Source text	EN	11.71	1.80	0.90	0	1.80	7.21	3.01	99
DeepL	EN→DE	12.50	3.91	3.13	0.78	3.91	7.03	22.89	99
ChatGPT	EN→DE	12.78	3.01	2.26	0.75	5.26	5.26	25.67	99
DeepL	EN→RU	9.17	N/A	1.83	0	0	0.92	N/A	N/A
ChatGPT	EN→RU	10.62	N/A	4.42	0	0	4.42	N/A	N/A

Case 3:

Tool	Language	Social	Affect	Posemo	Negemo	Certain	Tentat	Clout	Authentic
Source text	EN	21.28	8.51	2.13	0	0	2.13	93.44	16.56
DeepL	EN→DE	19.79	8.33	8.33	0	3.13	2.08	92.07	21.72
ChatGPT	EN→DE	17.31	6.73	6.73	0	1.92	1.92	91.24	24.86
DeepL	EN→RU	9.20	N/A	4.60	0	2.30	1.15	N/A	N/A

ChatGP	EN→R	14.4	N/A	6.67	0	2.22	1.11	N/A	N/A
T	U	4							

4.2.4 Interpretation

The table above presents the LIWC scores collected across all three cases. Based on the theoretical predictions developed in section 4.2.2, we expected higher scores on social, affect, posemo, authentic and tentat in Care-consistent outputs, and higher scores on certain and clout in FAT-consistent outputs. Contrary to this expectation, Social and Affect scores frequently increased rather than decreased in German translation. In Case 1, for instance, social rose from 30.69 in the source to 42.55 (DeepL) and 43.16 (ChatGPT). A plausible explanation is grammatical rather than relational: German more frequently requires explicit subject and pronoun marking than English, which can mechanically inflate word counts in categories like social and affect independent of any genuine change in care-sensitivity. This suggests that at least some of what LIWC measures here reflects structural differences between languages rather than ethical orientation, a limitation we return to below.

Of these predictions, Certainty is the most strongly and consistently confirmed. Certain scores increased from source to translation in every single case and in nearly every tool/language combination, for example, from 1.94 in the English source of Case 2a to 9.09 (DeepL) and 8.74 (ChatGPT) in German. This is consistent with the prediction that FAT-style translation resolves hedged, open-ended source language into more definitive, rule-governed phrasing. The explicit/implicit contrast between texts 2a and 2b produced a more mixed result than predicted. We expected text 2b, where vulnerability is unstated, to show a sharper drop in tentativeness than text 2a. This held for ChatGPT's German output (tentat fell from 7.77 in 2a to 5.26 in 2b) but not for DeepL's, whose tentat scores were similar across both texts (6.06 in 2a, 7.03 in 2b). The clearest support for our prediction came from the Russian data: DeepL's tentat score for text 2b (0.92) was dramatically lower than for text 2a (6.17), the single largest drop in the entire dataset. This suggests that the implicit/explicit distinction matters, but its effect is not uniform across tools or languages.

Contrary to this prediction, Clout scores were similarly high for both DeepL and ChatGPT across most cases, rather than diverging by tool. The single most striking result of the analysis instead appeared in Authenticity: in Case 1, DeepL scored 81.55 while ChatGPT scored only 3.11. This gap corresponds directly to the qualitative finding reported in section 4.1: DeepL's

informal "du" produced a more personal, direct register, while ChatGPT's formal "Sie" produced a more institutional one, indicating that Authenticity may be tracking the same register distinction we identified qualitatively, rather than functioning as an independent signal of care-sensitivity.

Crucially, the LIWC analysis does not stand alone. It is a quantitative supplement to the qualitative case analysis above. What it adds is the ability to make claims about patterns across texts rather than individual examples and to give those claims some empirical grounding. Its limitations are well known: LIWC operates at the lexical level and cannot capture syntactic or pragmatic features of language. It also applies dictionary-based categories that were developed primarily for English and may not map perfectly onto German or Russian. These limitations are discussed further in section 5.

4.3 Constructing a Model Card: How Should AI Translators Be Built to Grasp Culture and Meaning?

The model card is a tool originally proposed by Mitchell et al. (2019) as a form of documentation for machine learning models, analogous to nutrition labels or pharmaceutical package inserts: a standardised, transparent account of what a model is, what it can and cannot do, what values it encodes and under what conditions it should or should not be used. The original model card framework is itself a FAT instrument - it emphasises transparency, accountability and the disclosure of known limitations and biases.

We retain the model card format here but we rewrite its criteria from the ground up through a Care-ethical lens. The result is not simply a transparency document but a normative framework: a set of criteria by which a translation AI should be designed, evaluated, and held responsible - not to its developers or to regulators, but to the specific people whose relationships it mediates.

Care-Ethical Model Card for AI Translation Systems

A. Identity of the System

Criteria	Care-Ethical Standard	Why It Matters
----------	-----------------------	----------------

For whom is this tool designed?	The system must be designed with specific user communities in mind, not a universal abstract user. Documentation must name the intended relational contexts (professional, therapeutic, personal) and acknowledge those for which the system is not suitable.	Care ethics begins with the specific other. A tool that claims to serve everyone equally is a tool that has not thought carefully about anyone in particular.
Who built it and with whose voices?	Development teams must be interdisciplinary, including expertise in linguistics, cultural studies, ethics and lived experience of the communities most affected. The model card must disclose the composition of the team.	Following Villegas-Galaviz (2023), Care ethics requires that “different voices” not be silenced in the design process. Homogeneous teams reproduce homogeneous assumptions.

B. Training Data and Intended Use-Case

Criteria	Care-Ethical Standard	Why It Matters
What kind of data was the system trained on?	The model card must specify whether training data includes material modelling caring relational language, for example, therapy or counselling transcripts (used with appropriate consent and anonymisation), literature that depicts sustained caring relationships or peer-support community text, rather than only generic parallel corpora optimised for literal accuracy.	A system trained exclusively on formal, transactional text (legal documents, news, technical manuals) has no exposure to the linguistic patterns of care and cannot be expected to reproduce a register it has never seen. Training data is therefore itself a Care-ethical design choice, not a neutral technical detail.

In which contexts is this system best suited to be used?	The model card must specify the use-cases for which a Care-oriented design is most appropriate, for example, conversational and interpersonal contexts such as patient-facing chatbots, peer support platforms or customer-facing communication, as distinct from contexts where a more literal, FAT-style translation may be preferable, such as legal, technical or research documentation where precision and consistency matter more than relational nuance.	Care ethics does not claim that every translation context calls for the same relational sensitivity. Specifying intended use-cases prevents a Care-oriented system from being misapplied to contexts where its design priorities would be a liability rather than a strength and equally prevents a purely literal system from being used where relational sensitivity is actually required.
--	--	--

C. Relational Competence

Criteria	Care-Ethical Standard	Why It Matters
Does the system acknowledge relational context?	The system must not make unmarked defaults where the relational context is ambiguous (e.g., formal/informal address). In such cases, it must either prompt the user for clarification or present options with an explanation of what is at stake socially.	Tronto's “attentiveness” (Tronto 1993, 127) means noticing what matters in a situation. A system that resolves ambiguity silently fails to attend to the relational stakes of its choice.

How does the system handle emotional register?	Translation in emotionally sensitive contexts (therapeutic, caregiving, grief, conflict) must be evaluated not only for semantic accuracy but for emotional fidelity, including the preservation of vulnerability that is implied rather than stated outright. The system must be capable of flagging when a translation risks flattening the emotional register of the source.	Care ethics holds that vulnerability, including unstated vulnerability, is morally significant data, not noise to be processed away. A translation system that is indifferent to emotional register fails the “responsiveness” criterion (Tronto 1993).
Does the system support the user's own relational awareness?	Where appropriate, the system should offer the user the opportunity to reflect on the relational choices embedded in translation - not by lecturing but by making the stakes visible. This is what Cohn calls “making us think more deeply about our social obligations”.	Care ethics does not replace human judgment; it supports it. A Care-informed tool does not take relational decisions away from the user; it makes those decisions legible.

D. Identity and Representation

Criteria	Care-Ethical Standard	Why It Matters
How does the system handle gender in translation?	The system must not default to masculine forms in the absence of gender information. Where the source text uses gender-neutral or non-binary forms, the system must seek to preserve this intent in the target language, using gender-inclusive constructions where available. Where such constructions are unavailable or contested, the system must make the dilemma visible to the user rather than resolving it silently.	Cohn's critique of the “universal white male subject” points to the way defaults encode assumptions. A Care-informed system does not hide its assumptions; it makes them available for the user to challenge.
Does the system attend to cultural and linguistic particularity?	Following Venuti (1995), the system must not systematically domesticate target-culture expressions in ways that erase cultural specificity. A Care-ethical translation preserves the foreignness of the other where that foreignness is meaningful, rather than assimilating it to a dominant norm.	Spivak (1993) argues that translation is always a political act. Following Harding's (1991) account of situated knowledge, a Care-informed system must be aware of the power asymmetries between languages and cultures and must not allow those asymmetries to become invisible.

E. Accountability and Limitation

Criteria	Care-Ethical Standard	Why It Matters
For what contexts is this system not appropriate?	The model card must explicitly name contexts in which the system should not be used without human oversight: therapeutic settings, legal proceedings, medical interpretation, situations involving vulnerable speakers. This is not a disclaimer but a Care-ethical commitment.	Held's account of care as a practice distinguishes between what machines can replicate and what requires human presence (Held 2006, 36). A Care-ethical model card is honest about this boundary.
How is the system evaluated?	Success metrics must go beyond BLEU score and semantic accuracy to include relational fidelity: user-reported sense of being understood, preservation of emotional register and sensitivity to identity. These metrics must be developed with the communities most affected.	Measuring only what is easy to measure is a form of moral laziness. Care ethics demands that we evaluate what actually matters to the specific people the system serves.
Who is responsible when the system causes harm?	Accountability must not be dispersed across the system in ways that make it untraceable. The model card must name the responsible parties - developers, deployers, institutional users - and specify the mechanism through which affected users can seek redress.	Transparency without accountability is insufficient. Care ethics requires that someone answer for the relational consequences of the system's choices.

This model card is explicitly not a technical specification. It does not tell developers how to write code; it tells them what to be answerable for. This distinction matters, because Care ethics

is not primarily a design methodology - it is a moral orientation that asks us to take the quality of our relationships with specific others as the primary criterion of success. A translation system built on this orientation would look different from current tools not necessarily in its architecture but in its dispositions: the questions it asks, the choices it makes visible, the defaults it refuses to make.

4.4 Summary

This chapter has shown, through three concrete translation cases, that the FAT and Care frameworks do not merely disagree on abstract principles. They produce different outputs, prioritise different information and serve different relational ends. The LIWC analysis provided partial, rather than uniform, support for the hypothesis that current AI translation systems exhibit the linguistic signatures of FAT rather than Care values. Certainty increased consistently across nearly every case, consistent with the prediction that translation resolves hedged source language into more definitive phrasing. Tentativeness and Authenticity showed more tool-dependent and text-dependent patterns: the starkest finding in the entire dataset was not a Care/FAT divide but a sharp split between DeepL and ChatGPT's German Authenticity scores in Case 1, directly reflecting their divergent choice of formal versus informal address. Social and Affect categories, by contrast, often increased rather than decreased in German translation, a pattern we attribute partly to structural differences between English and German rather than to a straightforward confirmation of the Care/FAT hypothesis. The Care-ethical model card, finally, offers a normative alternative regardless of these mixed empirical results: a set of criteria by which translation AI should be held accountable not to the principle of efficiency, but to the quality of the relationships it mediates.

These findings set the stage for chapter 5, which examines the broader normative implications of FAT-value dominance in AI and asks what a world looks like in which Care ethics is taken seriously not only in translation but in the design of AI systems more generally.

5. The normative implications of a FAT-ethical dominance in AI translation and AI more generally

After establishing our overall thesis, we can now conclude that there is a dominance of FAT-ethical approaches in the contemporary ethics of AI, particularly when contrasted with care-ethical perspectives, which we were able to partially illustrate through the example of AI translation. Having laid out these descriptive findings already discussing specific normative implications within AI translation in chapter 4.1, we will now turn to discuss the general normative implications of a FAT-ethical dominance within AI ethics, nevertheless attempting to illustrate them using the previously introduced ‘tricky’ translation cases.

5.1 Moral integrity, efficiency and the ethics of justice

In Chapter 2 we have seen that the core values of the FAT-ethical approach are widely presented as standard options for ensuring just and trustworthy AI systems. At the same time, we argued, following Cohn, that these values are often interpreted primarily through a functional lens (Cohn 2020, 1), from which we conclude that systems are considered ethically acceptable when they optimize performance, minimize technical error and avoid obvious failures of output, while deeper questions about the political and social purposes of AI are sidelined. In Cohn’s words, dominant AI-ethical approaches risk drawing a premature conclusion that something is good because it is efficient (Cohn 2020, 1), shifting attention from the question of whether certain applications should exist at all to the question of how they can be implemented in the most efficient way.

This shift remains far from being normatively neutral, because it ties the very idea of moral integrity in AI to standards of efficiency and formal correctness. If ‘harm’ is defined primarily as a failure of performance (e.g.: as a biased classification or a lack of explainability) then many forms of injustice that relate to the social meaning and relational impact of AI systems remain invisible. In AI translation, for example, a system may be considered ‘fair’ if it attends to the most likely or least controversial answer, even if the content of its translations undermines respect, care or recognition in concrete contexts. (View: Chapter 4.1; Case 3)

As demonstrated in the second chapter, the FAT paradigm resonates strongly with what Virginia Held, drawing on Carol Gilligan’s work, calls ‘ethics of justice’. We have established that an ethics of justice is characterized by a focus on abstract principles, equal treatment, individual rights and impartial rule-following, and it has historically been associated with hegemonic masculine values such as competition, individualism and rationality. (Held 2006, 15) In this tradition, emotions like care, compassion or particular attachments are often treated as threats

to impartial judgement. Many canonical thought experiments explicitly invite us to ignore relationships in order to reach the ‘right’ decision. Cohn adapts this critique to the field of AI by arguing that FAT embodies an ethics of justice which tends to frame care as secondary or even as an obstacle to fairness, accountability and transparency. (Cohn 2020, 3)

From a care-ethical perspective, however, this opposition between justice and care itself is normatively problematic. Care theorists such as Held argue that our moral life is always shaped by concrete relationships of dependence and responsibility, and that these relationships do not disappear when we move from the private to the public or institutional sphere. Instead of seeing care as a private feeling to be bracketed in the name of justice, care ethics posits care as a central moral value that should guide how we structure social institutions and technologies (Held 2006, 12). When we apply this to AI ethics, the normative implication is that any framework which defines justice solely in terms of abstract fairness, rights and formal accountability risks misrecognizing or erasing the moral importance of relationships, dependencies and vulnerabilities that are mediated by AI systems.

Thus, the first normative implication of FAT-dominance is that it stabilizes a narrow image of moral integrity that privileges abstract, formally equal treatment over context-sensitive responses to concrete needs and relationships. Fairness metrics that treat all users ‘the same’, independent of their histories and positions, may appear just from the standpoint of an ethics of justice, but can themselves be unjust from a care perspective, because they ignore differences in vulnerability and structural disadvantage. In AI translation, this manifests when translation systems insist on uniform forms of address or standardized gendering practices which formally treat all speakers alike but materially disadvantage or harm those whose identities and relationships do not fit into dominant linguistic norms. For example, in Case 3, DeepL defaulted to masculine pronouns (er/sein/ihn in German, он/его in Russian) in both target languages, fully erasing the source text's gender-neutral "they", while ChatGPT preserved the gender-neutral intent using plural forms in both languages (German "sie/ihr", Russian "они/их"). (View: Chapter 4.1; Case 3)

At a more general level, then, FAT-dominance does not simply mean that certain values are prioritized, but that the hegemony of a particular underlying moral theory (an ethics of justice) continues, with all its blind spots towards care, dependence and relational justice. The normative cost of this hegemony is that it becomes harder to articulate injustices that are not capturable as violations of formal fairness, accountability or transparency, but instead concern

erosion of trust, marginalization of voices, or the deepening of moral distance in socially significant relationships.

5.2 Care and the denial of justice

This becomes particularly clear when we look at the discussion of justice: One of the reasons why FAT-oriented approaches appear so self-evident in AI ethics is that they inherit a long tradition in which justice and care are positioned as separate, and often as mutually exclusive, moral logics. Within this tradition, justice is associated with impartiality, universality and rule-governed reasoning, while care is relegated to the realm of emotion, particularity and personal attachment. (Held 2006, 15) Gilligan's work shows how this divide has been gendered: a supposedly 'mature' ethics of justice has been coded as rational and masculine, whereas care-centered moral reasoning has been dismissed as private, immature or biased. Care, in this picture, is at best morally admirable, but it is not what makes institutions or practices just. (Gilligan, 1982, 7)

Care ethicists have argued that this picture rests on a distorted understanding of both care and justice. Rather than treating care as a subjective feeling, they conceptualize it as a set of practices and values that are indispensable for sustaining human life and agency. Furthermore, it is highly beneficial not to confine care to intimate relationships, but to recognize the gains that emerge when institutional structures such as welfare systems, workplaces, and technologies are understood as permeated by care (Tronto 1993, 124). From this angle, the claim that care has 'nothing to do' with justice becomes untenable: if people's ability to live dignified and autonomous lives depends on receiving and giving adequate care, then questions about how care is distributed, recognized and institutionally supported are necessarily questions of justice. Denying this link does not make care morally irrelevant and only hides the injustices that arise when care is withheld, devalued or unevenly distributed across particular groups.

This has direct consequences for how we think about AI. A FAT-centered ethics of justice can give the impression that once systems are fair, accountable and transparent in a formal sense, the demands of justice are largely met. The care-ethical response is that this leaves out precisely those dimensions in which AI systems most deeply affect people's lives: how they mediate dependence, how they expose or protect vulnerabilities, and how they alter relationships of trust or neglect. In AI translation, for instance, a narrow justice-focus might notice unequal error rates between language groups, but it remains silent about the cumulative harm of being repeatedly addressed in the wrong register in professional settings (View: Chapter 4.1; Case 1), the slow erosion of trust when vulnerable speech in therapeutic contexts is rendered colder or

less tentative than intended (View: Chapter 4.1; Case 2), or the routine misgendering and erasure of non-binary identities when gender-neutral pronouns are forced into binary grammatical forms (View: Chapter 4.1; Case 3). A care-informed account of justice insists that these are not peripheral issues but central sites where injustice is produced and experienced.

Seen in this light, care ethics does not stand in competition with justice but revises the concept of justice that underwrites FAT-dominant approaches. It highlights that a framework which treats care as morally secondary will systematically overlook whole classes of injustice, concerning those who are cared for, who are left exposed, and whose needs and voices never enter into the design space of AI systems. By challenging the denial that care is justice-relevant, care ethics opens up new criteria for evaluating AI: not only whether it distributes errors fairly and explains its outputs, but whether it contributes to or undermines the networks of care on which a just society depends.

5.3 Efficiency, practicality and neoliberal market rationality

The third cluster of normative implications concerns the way FAT-dominance aligns with, and reinforces, broader neoliberal market logic. In Chapter 2 we already drew on David Harvey to argue that, under neoliberalism, the values of competition, efficiency and cost-benefit analysis have migrated from the economic sphere into political, social and even ethical reasoning (Harvey 2007, 65). Kate Crawford stresses in her work *AI Atlas* that neoliberal power monopolies through the systemic possibility of accumulating wealth goes hand in hand with the development of technocratic power concentrations, which have accelerated since the advent of AI, and present also as the reason “why we must contend with AI as a political, economic, cultural, and scientific force” (Crawford 2021, 20). AI therefore is a paradigmatic site of neoliberal developments, as it resembles a high-profit industry in which optimization, scalability and innovation are central, and in which ethical questions are made into technical operationalizations based on an implied ‘ethics of justice’ (Crawford 2021, 20-21). The fact that FAT principles can be relatively easily formalized, audited and integrated into existing legal and market structures explains part of their dominance: they promise to make ethics ‘practical’ without requiring fundamental changes to underlying economic relations.

From a care-ethical perspective, this alignment is normatively troubling because it systematically privileges values that fit market rationality over values such as solidarity, mutual support and social welfare. (Stark et.al. 2021, 263) Care work, as already emphasized through Held, is inherently relational, particularistic and often resistant to commodification. It cannot simply be replaced by machines without losing essential aspects of human connection and

responsibility. When AI systems such as AI translation are designed primarily to maximize efficiency and reduce costs, they tend to externalize the labor and responsibility of care. For instance, as demonstrated in chapter 3, users who are harmed or misrepresented by translation outputs are often expected to adapt, correct errors, or bear the relational consequences themselves, rather than being recognized as subjects whose vulnerabilities should guide system design. In this way, AI technologies tacitly assume a model of autonomous, resilient users who can absorb the burdens of miscommunication, which mirrors the neoliberal ideal of self-sufficient individuals responsible for managing their own risks.

Moreover, the focus on efficiency and practicality in FAT-oriented AI ethics tends to render invisible the structural conditions under which care can or cannot flourish. It can be argued that neoliberal tendencies of privatization also concerns shifts of care burdens from public institutions to households and individuals, disproportionately affecting women, migrants and other marginalized groups. AI tools such as machine translation are frequently marketed as solutions that save time, increase productivity and reduce costs, but rarely as tools that could redistribute care responsibilities more justly or strengthen existing networks of support. Normatively, this means that AI systems are evaluated according to how well they fit into and enhance existing market logics, rather than according to how they affect the capacity of societies to care for their members. Care ethics would ask instead whether AI reduces or exacerbates care deficits, whether it respects and supports caregivers, and whether it protects those who are most vulnerable to neglect and exploitation.

In summary, the focus on efficiency and practicality that characterizes FAT-dominant AI ethics reflects and reinforces a neoliberal moral economy that undermines care as a central ethical value. By aligning AI development with market logics, FAT-dominance makes it less likely that systems will be designed to protect vulnerabilities, nurture relationships or support solidarity and more likely that they will treat such concerns as secondary, optional or even obstructive to progress. From a care-ethical perspective, this is itself a serious normative implication: it suggests that without a deliberate reorientation towards care, AI ethics risks becoming an instrument for stabilizing, rather than challenging, unjust social and economic orders.

6. Conclusion

To conclude, the presence of Care-ethical outcomes in AI tools for translation is minimal and probably unintentional. While some models do tend to perform better in terms of the use of a correct pronoun (in formal vs informal case or in terms of gender-neutral translation) as well as accurate verb forms for therapeutic language, most tend to miss the intent and meaning of the original text. The reason for that is because the models are built based on an FAT approach, which reflects an ‘ethics of justice’ approach towards the ethics of AI. It furthermore fits neoliberal market assumptions and heavily relies on efficiency, rather than ultimate social benefit (which applies to everyone and not only the corporation developing the AI system). The only reason why some AI models display a better result is because they had a broader training data pool, hence they were trained to perform better.

Across our three cases, this pattern showed up concretely: DeepL defaulted to masculine pronouns where the source text used a deliberately gender-neutral “they” and silently picked formal or informal address without ever flagging the choice to the user. Our LIWC analysis offered partial, rather than uniform, support for this picture. Certainty consistently increased from source to translation, matching the FAT prediction, but other categories like Social and Affect did not behave as predicted, suggesting some of what we measured reflects structural differences between languages rather than ethical orientation alone. The care-ethical model card we proposed in Chapter 4 is one concrete step toward closing this gap.

Our suggestion and strong encouragement would be to implement more Care ethics practices into the development of AI translation systems, so that they’re designed from the beginning with Care as both Value and Practice in mind. In theory, that would involve focusing on voice, relationship and interdependence, vulnerability and direct level as core principles of the tool’s functioning. In practice it would imply the need of diverse development teams, constant rounds of feedback, consideration of intent and outcome, development of focused translation tools to tackle specific cases, etc. This is not only a matter of adding softer values to existing systems, but of revising the underlying concept of justice that current AI ethics relies on - treating relationships, dependence and vulnerability as just as central to fairness as equal treatment and formal rules. We believe that the Care ethics approach will benefit the AI translation industry to be more ethical towards the society as a whole.

These conclusions should nevertheless be read with some caution. Our LIWC analysis relied on dictionaries built primarily for English, our Russian dictionary was a custom adaptation rather than an official release and three translation cases cannot capture the full range of

situations AI translation tools encounter. Care ethics in AI also remains, as we noted in Chapter 3, a largely theoretical framework whose practical implementation still raises open questions about measurement and evaluation.

Nevertheless, our general normative verdict is that without a deliberate reorientation toward care, AI-ethics risks becoming an instrument for stabilizing, rather than challenging unjust social and economic orders. Care ethics on the other hand presents as a solution that could broaden the narrow concept of justice that FAT-approaches would further naturalize.

Bibliography

- Balog, Krisztian/Nolwenn, Bernad (2023): *A Systematic Review of Fairness, Accountability, Transparency and Ethics in Information Retrieval*. In: ACM Computing Surveys (preprint). Online available: <https://doi.org/10.1145/3637211> [25.03.2026].
- Bansal Jora, R., Kaur Sodhi, K., Mittal, P., & Saxena, P. (2022). *Role of Artificial Intelligence (AI) In meeting Diversity, Equality and Inclusion (DEI) Goals*. 8th International Conference on Advanced Computing and Communication Systems (ICACCS), 1687-1690. <https://doi.org/10.1109/ICACCS54159.2022.9785266>
- Bidadanure, Juliana/Axelsen, David (2025): Egalitarianism. In: Zalta, Edward N./Nodelman, Uri (ed.). *The Stanford Encyclopedia of Philosophy*. URL = <https://plato.stanford.edu/archives/spr2025/entries/egalitarianism/> [24.05.2026].
- Boyd, Ryan L./Ashokkumar, Ashwini/Seraj, Sarah/Pennebaker, James W. (2022): *The Development and Psychometric Properties of LIWC-22*. Austin, TX: University of Texas at Austin.
- Bryson, J. J. (2009, May). Robots Should Be Slaves. *Artificial Models of Natural Intelligence*, 1-12. <https://doi.org/10.1075/nlp.8.11bry>
- Carter, Ian (2022): Positive and Negative Liberty In: Zalta, Edward N./Nodelman, Uri (ed.). *The Stanford Encyclopedia of Philosophy*. URL = <https://plato.stanford.edu/archives/spr2022/entries/liberty-positive-negative/>. [24.05.2026].
- Cohen, G. A. (2009): *Why not Socialism?* Princeton/Oxford: Princeton University Press.
- Cohn, Jonathan (2020): *In a Different Code: Artificial Intelligence and the Ethics of Care*. In: International Review of Information Ethics 28.
- Crawford, Kate (2021): *Atlas of AI. Power, Politics, and the Planetary Costs of Artificial Intelligence*. New Haven/London: Yale University Press.
- Fjeld, Jessica et.al. (2020): *Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-based Approaches to Principles for AI*. Berkman Klein Center for Internet & Society, online available: <http://nrs.harvard.edu/urn-3:HUL.InstRepos:42160420> [25.03.2026]
- Gebauer, J. E., Riketta, M., Broemer, P., & Maio, G. R. (2008). Pleasure and pressure based prosocial motivation: Divergent relations to subjective well-being. *Journal of Research in Personality*, 42, 400-420. <https://doi.org/10.1016/j.jrp.2007.07.002>

Gilligan, Carol (1993 [1982]): *In a Different Voice: Psychological Theory and Women's Development*. Cambridge/London: Harvard University Press

Hankivsky, O. (2014, May). Rethinking Care Ethics: On the Promise and Potential of an Intersectional Analysis. *American Political Science Review*, 108(2), 252-264. <https://doi.org/10.1017/S0003055414000094>

Hayakawa, S., & Slote, M. (Eds.). (2026). *Care Ethics and Beyond: Moral, Epistemological, and Cross-Cultural Perspectives*. Springer Nature Switzerland. <https://doi.org/10.1007/978-3-032-09411-7>

Harding, Sandra (1991): *Whose Science? Whose Knowledge? Thinking from Women's Lives*. Ithaca: Cornell University Press.

Harvey, David (2005): *A Brief History of Neoliberalism*. Oxford/New York: Oxford University Press.

Held, Virginia (2006): *The Ethics of Care: Personal, Political, and Global*. Oxford: Oxford University Press.

Johnson, D. G. (2006). Computer systems: Moral entities but not moral agents. *Ethics and Information Technology*, 8, 194-205. <https://doi.org/10.1007/s10676-006-9111-5>

Mindry, D. (2010, June). Engendering care: HIV, humanitarian assistance in Africa and the reproduction of gender stereotypes. *Culture, Health & Sexuality*, 12(5), 555-568. <https://doi.org/10.1080/13691051003768140>

Mitchell, Margaret et.al. (2019): *Model Cards for Model Reporting*. In: Proceedings of the Conference on Fairness, Accountability, and Transparency, 220-229.

Morris, K. (2025). Reconceptualising socioeconomic rights: a case for care ethics. *The International Journal of Human Rights*, 29(6), 795-815. <https://doi.org/10.1080/13642987.2024.2429469>

Norlock, Kathryn (2019): Feminist Ethics In: Zalta, Edward N./Nodelman, Uri (ed.). *The Stanford Encyclopedia of Philosophy*. URL = <https://plato.stanford.edu/archives/sum2019/entries/feminism-ethics/>. [24.05.2026]

Pennebaker, James W./Francis, Martha E./Booth, Roger J. (2001): *Linguistic Inquiry and Word Count: LIWC 2001*. Mahwah, NJ: Lawrence Erlbaum Associates.

Pennebaker, James W. (2011): *The Secret Life of Pronouns: What Our Words Say About Us*. New York: Bloomsbury Press.

Rawls, John (1999): *A Theory of Justice* (rev. ed.). Cambridge: The Belknap Press.

Sawer, M., Jenkins, F., & Downing, K. (Eds.). (2020). *How Gender Can Transform the Social Sciences: Innovation and Impact*. Springer International Publishing. <https://doi.org/10.1007/978-3-030-43236-2>

Singhal A./Neveditsin N./Tanveer H./Mago V. (2024): *Toward Fairness, Accountability, Transparency, and Ethics in AI for Social Media and Health Care: Scoping Review*. In: JMIR Med Inform 12 (2024): e50048, doi: 10.2196/50048 [25.03.2026].

Souza de Santos, R., Baldassarre, M. T., & Magalhaes, C. (2026). The role of team diversity in AI systems development. *Empirical Software Engineering*, 31(90), 1-30. <https://doi.org/10.1007/s10664-025-10755-6>

Spivak, Gayatri Chakravorty (1993): *The Politics of Translation*. In: *Outside in the Teaching Machine*. Routledge, 179–200.

Stark, L./Greene, D./Hoffmann A. L. (2021): *Critical Perspectives on Governance Mechanisms for AI/ML Systems*. In: Roberge, J./Castelle, M. (eds.), *The Cultural Life of Machine Learning*. pp. 257-280, online available: https://doi.org/10.1007/978-3-030-56286-1_9 [25.03.2026]

Tronto, Joan C. (1993): *Moral Boundaries: A Political Argument for an Ethic of Care*. New York: Routledge.

Venuti, Lawrence (1995): *The Translator's Invisibility: A History of Translation*. Routledge.

Villegas-Galaviz, C. (2022). Ethics of Care as Moral Grounding For AI. *The Ethics of Data and Analytics*. <https://doi.org/10.1201/9781003278290-13>

Villegas-Galaviz, C., & Martin, K. (2024). Moral distance, AI, and the ethics of care. *AI & Society*, 39, 1695–1706. <https://doi.org/10.1007/s00146-023-01642-z>

Wen, Y., & Holweg, M. (2024). A phenomenological perspective on AI ethical failures: The case of facial recognition technology. *AI & Society*, 39, 1929–1946. <https://doi.org/10.1007/s00146-023-01648-7>

Westerstrand, Salla (2024): *Reconstructing AI Ethics Principles: Rawlsian Ethics of Artificial Intelligence*. *Science and Engineering Ethics* 30(5): 46.

Zafar, M. (2025). Normativity and AI moral agency. *AI and Ethics*, (5), 2605–2622.
<https://doi.org/10.1007/s43681-024-00566-8>