



iJRASET

International Journal For Research in
Applied Science and Engineering Technology



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Volume: 11 **Issue:** III **Month of publication:** March 2023

DOI: <https://doi.org/10.22214/ijraset.2023.49153>

www.ijraset.com

Call: ☎ 08813907089

E-mail ID: ijraset@gmail.com

Advanced Key Foundations of Multiagent System

Dr. M Purushotham

Sreyas Institute of Engineering College

Abstract: *Ethics is inherently a multiagent concern. However, analysis on AI ethics nowadays is dominated by work on individual agents: (1) however Associate in Nursing autonomous golem or automotive could hurt or (differentially) profit folks in theoretical things (the questionable tramcar problems) and (2) how a machine learning formula could turn out biased choices or recommendations. The social group framework is basically omitted. To develop new foundations for ethics in AI, we tend to adopt a sociotechnical stance during which agents (as technical entities) facilitate autonomous social entities or principals (people and organizations). This multiagent conception of a sociotechnical system (STS) captures however moral considerations arise within the mutual interactions of multiple stakeholders. These foundations would modify USA to understand ethical STSs that incorporate social and technical controls to respect stated moral postures of the agents within the STSs. The visualized foundations need new thinking, on 2 broad themes, on how to realize (1) Associate in STS that reflects its stakeholders' values and (2) individual agents that perform effectively in such Associate in STS.*

Keywords: *Ethics; values; sociotechnical systems; norms; preferences*

I. ETHICS IN MULTIAGENT SYSTEMS

The shocking capabilities incontestable by AI technologies overlaid on careful knowledge and fine-grained management provide cause for concern that agents will wield huge power over human welfare, drawing increasing attention to ethics in AI. Ethics is inherently a multiagent concern—an amalgam of (1) one party's concern for one more and (2) a notion of justice.

To capture the multiagent conception realistically, we have a tendency to model our setting as a sociotechnical system (STS). associate STS includes autonomous social entities (principals, i.e., individuals and organizations) and technical entities (agents, WHO facilitate principals, and resources) [13]. What foundations will we got to build STSs that address moral concerns from multiple perspectives? Since Associate in Nursing agent might incorrectly follow its principal's moral directive or properly follow an unethical directive, Associate in moral STS ought to give social and technical controls to market moral outcomes. The STS conception leads USA to formulate the matter because the one amongst specifying (1) Associate in STS to respect a explicit general moral posture over its stakeholders' worth preferences; Associate in (2) an agent UN agency respects a stated individual moral posture and functions in this STS. Existing works on AI and ethics adopt a single-party attitude in topics like (1) recursive answerability [17] and fairness [25], where choices or recommendations are often biased; and (2) the behavior of agents [16], once facing ethical quandaries in theoretic situations, like the notable tram issues [23]. Even MAS-oriented analysis on ethics for the most part focuses on analysis of stakeholders' values [24] with the aim of specifying a single agent. Recent efforts by MAS researchers, identify limitations of existing approaches, like goal modeling, suggesting that current models lack vital elements. In distinction, we advocate realizing (1) STSs that replicate system objectives in their social architecture; and (2) agents that balance ethical preferences and facilitate their principals take moral choices. Following [14], we have a tendency to obtain to make on recent analysis on values and principles of justice, providing new foundations for ethical multiagent systems on 3 main analysis themes.

- 1) *Model:* Developing a model of ethics supported values and norms that supports individual and system-level moral judgments based mostly on standard criteria we tend to decision moral postures.
- 2) *Novelty:* increasing the scope of multiagent system modeling to include norms and values, incorporating concepts on guilt and inequity aversion [19], and therefore the principles of justice.
- 3) *Analysis:* Developing reasoning techniques to assist stakeholders identify potential moral pitfalls in associate STS, specifically via (1) formal verification approaches to accommodate ethics in terms of norms and values preferences; and (2) agent-based simulations for assessing the ethicality of STSs and their members.
- 4) *Novelty:* Combining verification and simulation to assess however well associate STS respects a system-level moral posture like utilitarianism and school of thought.
- 5) *Elicitation:* Develop techniques to specify a suitable STS primarily based on value-based negotiation between involved stakeholders and to elicit worth preferences from stakeholders.
- 6) *Novelty:* Enhancing negotiation and deliberation techniques to focus on values and not intrusive learning important preferences.

II. SOCIOTECHNICAL SYSTEM (STS)

We begin from an outline of a sociotechnical system (STS) custom-made from Kafalı et al., introducing the required ideas underlying our conception. Section four discusses further literature. Figure one shows associate degree STS (right frame) and the way we tend to envision such an STS being designed (left frame).

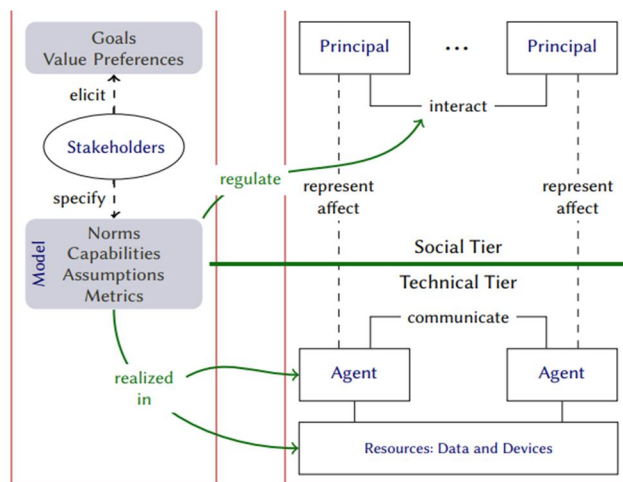


Figure 1: A schematic of a sociotechnical system (STS) [29].

A neutral in associate degree STS is associate degree autonomous entity (individual or organization) that has associate degree interest within the specification or operations of the STS. for instance, the stakeholders of a patient transfer [1] STS embrace patients, doctors, nurses, and also the hospital; for a phone users' STS, the caller, the callee, furthermore because the individuals and organizations near (e.g., a library is inquisitive about the keeping the phones of individuals within the library silent) square measure the stakeholders. A principal could be a neutral that's active in an exceedingly system. A principal can opt for its actions within the system. Our applications of interest emphasize interactions among principals whereby they exchange information and services, e.g., as in social media, scientific collaboration, and care. A neutral UN agency isn't a principal would have associate degree interest within the specification of a system however doesn't participate as a call maker. for instance, in patient transfer, a nurse and doc square measure principals, however a patient generally isn't. When associate degree STS is operational, its social tier includes principals and its technical tier includes agents and underlying resources, such as databases, services, sensors, and actuators. The agents act on behalf of the principals and their actions have an effect on the principals: in many-to-many relationships, shown as matched for simplicity. Engineering associate degree STS involves distinguishing its stakeholders and eliciting their goals (reflecting domain requirements) and price preferences to provide a model that specifies the STS along side its environmental (operating) assumptions and metrics. The specification captures the STS's (1) technical design in terms of capabilities, viewed abstractly as actions on resources that participants can perform; and (2) social design in terms of the principals' roles and also the norms capturing the legitimate expectations between them and also the consequences of the actions.

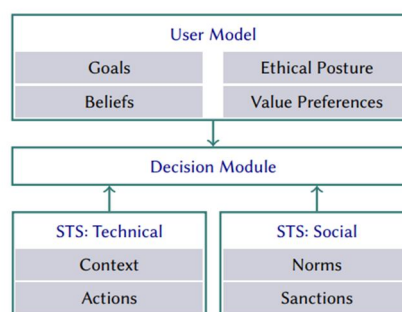
A. Ethical Postures

An individual moral posture refers to however AN agent could respond to the worth preferences of a principal within the STS World Health Organization is tormented by the agent's actions. AN example moral posture would be to mirror the common intuition that a call is moral if it accommodates the preferences of others besides oneself. A general moral posture refers to however AN STS is given in lightweight of the worth preferences of its stakeholders. samples of ethical postures embody (1) ism, i.e., to reduce inequality across stakeholders with relevancy satisfying their preferences; and (2) ism, i.e., to maximise mixture welfare (the greatest good of the best number) while not respect to any disparities. The principals World Health Organization participate in AN STS could adopt a unique ethical posture from what's incorporated within the STS and totally different value preferences from those of the stakeholders World Health Organization given it.

B. Agents

Figure two illustrates AN agent's illustration and deciding. AN agent's user model describes the agent's user in terms of goals, beliefs, worth preferences, and moral posture. An agent's user could be a principal on whose behalf the agent acts and interacts. The agent's selections might have an effect on not solely its user however conjointly alternative principals within the STS. The agent maintains information of the STS in which it functions, together with however its principals relate. An agent's decision module produces actions that mirror the worth preferences of its user's moral posture, the moral postures of alternative involved principals, furthermore because the general moral posture.

Figure 2: An agent that functions in an STS by representing and reasoning about its technical and social architectures, the goals, values, preferences, and ethical postures of the principals in an STS.



III. RESEARCH CHALLENGES AND QUESTIONS

Our analysis objective is to make new multiagent foundations for AI ethics. to the current finish, we have a tendency to advocate addressing the shortcomings of current approaches for AI ethics by, severally, developing (1) new representations and reasoning approaches concerning ethics from each individual and system perspectives; (2) new ways in which of analyzing systems with relevancy AN moral posture each statically (verification) and dynamically (simulation); and (3) new ways in which to elicit worth preferences from stakeholders and to help them in negotiating acceptable STS specifications. A sensible approach would have several parts. However, we have a tendency to target novel challenges that we have a tendency to posit as having the best prospect of reward.

A. Model of Ethics for AI

Q1 Illustration. what's associate degree acceptable model associate degreeed illustration of associate degree STS and of an agent from the position of ethics?

Q2 Higher cognitive process. however will we tend to support higher cognitive process by an agent in associate degree STS that takes into consideration the worth preferences of its user and different principals likewise because the STS?

Motivation. To represent associate degree STS exactly and reusably, we need a sufficiently made language that supports not solely the required normative relationships however additionally provides a capability to capture time (to support constraints on ordering and occurrence), strength (to use as a basis for crucial preferences to handle conflicts), and context (to modulate the outcomes of violation, for instance). In our conception, associate STS isn't a separate running entity however is complete through the interactions of the principals, agents, and resources that feature within the STS. every agent should represent (1) its view of the social design of the STS, together with the normative relationships during which it participates; (2) its read of the technical architecture of the STS, particularly the context (state of the world) and actions which will be performed in it; and (3) the goals and worth preferences of its user and different principals. Such a illustration enables agents to create choices that balance the on top of components.

This language ought to alter the specification of associate degree STS as well as a model of the precise principals and agents that includes in it. It should support specifying price preferences for every principal wherever the preferences can be expressed as ordinals or as cardinal values. Consent could be a central construct in ethics and responsibility that has not received adequate attention from AI researchers. Consent characterizes once associate degree action by one autonomous party gains legitimacy despite probably infringing upon the autonomy or authority of another party, unforgettably referred to as the "moral magic of consent" [6]. As these works et al., e.g., [4], indicate the intuitions about consent area unit faraway from established within the legal literature. Two major competitive intuitions area unit that consent reflects (1) a mental action of the willing party, indicating that it's the exercise of an internal choice; and (2) a communicative act or performative by the willing party conferring powers on the recipient, indicating that it's the exercise of a normative power.

The existing literature on consent focuses on a retrospective read (which is to adjudicate on some apparent violation, as in a very court of law) however in AI ethics the potential read is arguably a lot of necessary (since it's concerning selections to be created by associate degree agent on the fly).

Decision creating. notice prosocial agents that model worth preferences of not solely their various users however conjointly the opposite principals plagued by the agents' actions. Specifically, will AN agent's deciding mirror its user's ethical posture and therefore the worth preferences of the principals affected by its actions? a specific moral posture is inequity aversion [19], which maps to the informal construct of guilt. once AN inequity averse agent doesn't act in accordance with the worth preferences of a principal, it accumulates guilt (on behalf of its user). The guilt applies submissively once it follows or deviates from a norm. Such an agent might ANticipate guilt from taking a dubious action, which feeling might discourage the agent from taking that action.

B. Analysis of Ethicality

Q3 Verification. however will we tend to verify that associate degree STS specification satisfies the neutral needs with relevance a given systemic moral posture?

Q4 Simulation. however will we tend to change stakeholders of associate degree STS to assess associate degree STS specification in relevance actual or imputed ethical postures of the principals UN agency would understand that STS?

Motivation. As we tend to model ethics, it's necessary {to associate degreeealize|to research|to investigate} an STS specification on measures like aliveness (something smart happens), safety (nothing dangerous happens), hardiness (how long one thing smart keeps happening), associate degreeed resilience (how quickly an STS recovers from one thing bad). Such analyses necessitate the utilization of (1) formal verification to assess the STS specification, associate degreeed (2) simulation to foresee an outcome. Bremner et al. [12] gift a number one approach for formal verification intermeshed toward moral reasoning, incorporating beliefs, desires, and values in an exceedingly framework supported planning. This approach will facilitate advance the current agenda. Verification of STSs and Agents. Develop new model checking approaches that contemplate worth preferences of stakeholders and work on prime of existing probabilistic model checking tools. Given associate degree STS specification and also the information of outcomes promoted by values, associate degree increased verification tool would facilitate North American country perceive whether or not the specification is biased toward sure values. For example, we tend to could establish that a phone ringer agent continually prefers safety over privacy. That agent may ring a user's phone loud for a decision from a friend. however will we tend to adapt rising model checking tools for these purposes? A supply of quality in our setting is that we tend to represent each the STS specification and also the agents UN agency support its principals. thanks to the necessity of autonomy, any norm could also be profaned [46], tho' norms offer a basis for responsibility. And, in general, we tend to cannot interpret worth preferences evidently utilities as is typical in scientific theory. Thus, a quest challenge is the way to formulate the correctness problems. we tend to associate degreeeeticipate that correctness properties would be assessed (1) {separately|individually|singly|severally|one by one|on associate degree individual basis} for an STS and conditional upon an STS for its member agents; and (2) qualitatively with relevance moral postures of individual agents and of the system. Although formal verification will facilitate assess associate degree STS specification under general assumptions, social simulations offer North American country with associate degree avenue to foresee the runtime outcome. Social Simulation. change stakeholders to guide the simulation, and later on facilitate them perceive the outcomes in associate degree STS if a definite sort (or group) of people were to move in it. For example, if the phone user is traveling extensively for work and is attending conferences, the simulation can facilitate the neutral determine that in associate degree STS specification biased toward safety over privacy, the agent can ring the phone loud a lot of oftentimes than in an STS that balances safety associate degreeed privacy betting on the user's context; such associate degree agent can so deviate a lot of usually from STS norms and attract a lot of sanctions from agents of alternative principals. Can we tend to generate social quandary situations for every user primarily based on associate degree understanding of the user's previous interactions and far-famed value preferences? These social quandary things embody cases where (1) multiple norms conflict, (2) one or a lot of norms conflict with worth preferences of a user, (3) worth preferences of 1 user conflicts with those of alternative users within the interaction.

C. Elicitation of Ethical Systems

Q5 Learning. however will associate degree agent elicit its users' price preferences?

Q6 Negotiation. however will we have a tendency to modify stakeholders to form associate degree STS specification that accords with their price preferences?

Motivation. will agents act in ways in which align with the values of principals? to try to to therefore, associate degree agent should, first, acknowledge the worth preferences of its principals, which may be very difficult.

First, asking the principals what values they like over others directly (e.g., via a survey) is probably going to be futile. As Bostyn et al. [11] show, responses to theoretic queries on ethical preferences (e.g., as within the tram drawback surveys) don't predict the behavior of the participants in world. Thus, associate degree agent should learn its principals' price preferences by perceptive what the principals liquidate real decision eventualities and reasoning regarding why they did therefore. Second, price preferences is context specific—a principal may like one price to a different ($v_1 > v_2$) in a very context however have the opposite preference ($v_2 > v_1$) in another context. as an example, contemplate Charlie, a principal WHO is visually impaired. Charlie prefers safety (a value) to privacy (another value) once he's traveling (a context). Thus, once Charlie is traveling, his agent mechanically takes photos of his surroundings and shares them along with his friends. However, if Charlie is traveling with Dave, a trusty friend (another context), there's no want for Charlie's agent to compromise privacy by sharing photos. Third, though associate degree agent must learn its user's values, those values, in turn, might depend upon the values of alternative principals with that the agent and its user act. As the examples on top of counsel, learning price preferences involves recognizing and modeling many nuances. Even for alittle set of core values of interest in associate degree application state of affairs, there can be an outsized range useful preferences, considering the variability of physical and social contexts during which the preferences apply. Learning price Preferences. Learn price preferences by perceptive (1) the principal's actions; (2) whether or not the principal approves or disapproves the agent's actions; and (3) whether or not alternative principals sanction the agent's actions, completely or negatively. This drawback is basically totally different from the everyday preference learning issues, e.g., [3], whose objective is to find out preferences from pairwise comparisons of things of interest. As we argue on top of, directly eliciting preferences between price pairs from principals might not yield fascinating outcomes. In distinction, we seek to learn price preferences by perceptive what principals and agents do (as against what they say) in numerous contexts. Ajmeri et al. [5] show however price preferences is aggregate to spot a consensus action that is truthful to all or any stakeholders concerned. Knowing the worth preferences of stakeholders helps in higher facilitating interaction between them. Interest-based negotiation [21] relies on the concept that stakeholders' goals might take issue from their positions throughout negotiation. Thus, satisfying their (imputed) goals is healthier than giving them what they expressly arouse. Prior negotiation protocols for settings associated with STSs, e.g., [8, 10], accommodate neither values nor the whole breadth of associate degree STS as conceived here. Existing approaches, e.g., concentrate on eliminating conflicts among negotiating parties. In distinction, we have a tendency to motivate conflicts as a basis for negotiation of associate degree STS (during elicitation) and as associate degree input into moral deciding by associate degree agent (at run time). Value-Based Negotiation. Support stakeholders with conflicting necessities however similar price preferences in generating associate degree acceptable STS specification. In our conception of value-based negotiation, every supply includes associate degree STS specification. A neutral will reason regarding however the current supply contributes to it stakeholder's most popular values to decide the response move: settle for, reject, or generate a offering. Facilitating such reasoning needs (1) a normative negotiation framework for the specification of STSs that has a basis for systematically editing norms to modify the generation of effective offers and counteroffers; and (2) a value-based concession bidding strategy that adapts its offers at run time supported opponent's behavior while not predefined utility functions.

IV. ETHICS AND RELATED CONSTRUCTS

An AI system is neither just associate degree rule nor a standalone agent, but sociotechnical system representing a society of humans and agents. consequently, there's a necessity and urgency for addressing societal considerations on AI adoption. Ethics is one such concern however it's closely associated with alternative social group considerations on AI, together with fairness, accountability, transparency, and privacy. The new foundations we tend to call for will and will address these connected considerations in addition. Fairness considerations judgments on the outcomes of machine learning predictors. analysis on fairness in AI [18] and the way individuals assess AI fairness [9] focuses on a personal (is it truthful to me?) or system (is the system truthful as a whole?), however not on a gaggle, incorporating the preferences of stakeholders and their social relationships and power dynamics. to attain cluster fairness, each agent should support fairness in higher cognitive process by understanding contextually relevant norms [4] and reasoning regarding price preferences of all stakeholders [5], not simply of the agent's user. Accountability is crucial to establishing World Health Organization is in control of a decision created by associate degree agent [15]. previous works perceive answerableness as either traceability [7] or negative utility [20], but these ideas area unit neither necessary nor decent for capturing accountability as a result of they lack the social-level linguistics that undergirds answerableness. we tend to ask for to capture the normative basis of answerableness directly tho' it supports traceability and sanctioning wherever applicable. Transparency relates to the principle of explicability and considerations traceability [2]. we tend to ask for to support these desired principles of accountable AI through traceability of STS negotiation steps as well as explicability of agents' reasoning at runtime. Privacy is of course approached from a values perspective [24].

It encompasses values like confidentiality, disapproval, and avoiding infringing into others' house [4, 22]. Researchers advocate giving bigger management to users on higher cognitive process, e.g., for privacy. However, giving management to users raises the question of whether or not a user's action accords therewith user's or alternative involved users' values. Social norms area unit the centerpiece of contextual integrity, a theory of privacy wherever violations occur once information flows violate discourse norms.

REFERENCES

- [1] Joanna Abraham and Madhu C. Reddy. 2010. Challenges to Inter-Departmental Coordination of Patient Transfers: A Workflow Perspective. *International Journal of Medical Informatics* 79, 2 (Feb. 2010), 112–122.
- [2] AIHLEG. 2019. Ethics Guidelines for Trustworthy AI. Independent High-Level Expert Group on Artificial Intelligence set up by the European Commission. https://ec.europa.eu/newsroom/dae/document.cfm?doc_id=60419.
- [3] Nir Ailon. 2012. An Active Learning Algorithm for Ranking from Pairwise Preferences with an Almost Optimal Query Complexity. *The Journal of Machine Learning Research* 13, 1 (Jan. 2012), 137–164.
- [4] Nirav Ajmeri, Hui Guo, Pradeep K. Murukannaiah, and Munindar P. Singh. 2018. Robust Norm Emergence by Revealing and Reasoning about Context: Socially Intelligent Agents for Enhancing Privacy. In *Proceedings of the 27th International Joint Conference on Artificial Intelligence (IJCAI)*. IJCAI, Stockholm, 28–34. <https://doi.org/10.24963/ijcai.2018/4>
- [5] Nirav Ajmeri, Hui Guo, Pradeep K. Murukannaiah, and Munindar P. Singh. 2020. Elessar: Ethics in Norm-Aware Agents. In *Proceedings of the 19th International Conference on Autonomous Agents and MultiAgent Systems (AAMAS)*. IFAAMAS, Auckland, 1–9.
- [6] Larry Alexander. 1996. The Moral Magic of Consent (II). *Legal Theory* 2, 3 (Sept. 1996), 165–174.
- [7] Katerina Argyraki, Petros Maniatis, Olga Irzak, Subramanian Ashish, and Scott Shenker. 2007. Loss and Delay Accountability for the Internet. In *Proceedings of the IEEE International Conference on Network Protocols (ICNP)*. IEEE, Beijing, 194–205.
- [8] Reyhan Aydoğan, David Festen, Koen V. Hindriks, and Catholijn M. Jonker. 2017. Alternating Offers Protocols for Multilateral Negotiation. In *Modern Approaches to Agent-based Complex Automated Negotiation*, Katsuhide Fujita, Quan Bai, Takayuki Ito, Minjie Zhang, Fenghui Ren, Reyhan Aydoğan, and Rafik Hadfi (Eds.). Number 674 in *Studies in Computational Intelligence*. Springer, Cham, 153–167. https://doi.org/10.1007/978-3-319-51563-2_10
- [9] Reuben Binns, Max Van Kleek, Michael Veale, Ulrik Lyngs, Jun Zhao, and Nigel Shadbolt. 2018. 'It's Reducing a Human Being to a Percentage': Perceptions of Justice in Algorithmic Decisions. In *Proceedings of the Conference on Human Factors in Computing Systems (CHI)*. ACM, Montreal, 377:1–377:14.
- [10] Guido Boella, Patrice Caire, and Leendert van der Torre. 2009. Norm Negotiation in Online Multi-Player Games. *Knowledge and Information Systems* 18, 2 (Feb. 2009), 137–156.
- [11] Dries H. Bostyn, Sybren Sevenhant, and Arne Roets. 2018. Of Mice, Men, and Trolleys: Hypothetical Judgment Versus Real-Life Behavior in Trolley-Style Moral Dilemmas. *Psychological Science* 29, 7 (2018), 1084–1093. <https://doi.org/10.1177/0956797617752640> PMID: 29741993.
- [12] Paul Bremner, Louise A. Dennis, Michael Fisher, and Alan F. T. Winfield. 2019. On Proactive, Transparent, and Verifiable Ethical Reasoning for Robots. *Proc. IEEE* 107, 3 (March 2019), 541–561. <https://doi.org/10.1109/JPROC.2019.2898267>
- [13] Amit K. Chopra, Fabiano Dalpiaz, F. Başak Aydemir, Paolo Giorgini, John Mylopoulos, and Munindar P. Singh. 2014. Protos: Foundations for Engineering Innovative Sociotechnical Systems. In *Proceedings of the 22nd IEEE International Requirements Engineering Conference (RE)*. IEEE Computer Society, Karlskrona, Sweden, 53–62. <https://doi.org/10.1109/RE.2014.6912247>
- [14] Amit K. Chopra and Munindar P. Singh. 2016. From Social Machines to Social Protocols: Software Engineering Foundations for Sociotechnical Systems. In *Proceedings of the 25th International World Wide Web Conference*. ACM, Montréal, 903–914. <https://doi.org/10.1145/2872427.2883018>
- [15] Amit K. Chopra and Munindar P. Singh. 2018. Sociotechnical Systems and Ethics in the Large. In *Proceedings of the AAAI/ACM Conference on Artificial Intelligence, Ethics, and Society (AIES)*. ACM, New Orleans, 48–53. <https://doi.org/10.1145/3278721.3278740>
- [16] Louise A. Dennis, Michael Fisher, Marija Slavkovik, and Matt Webster. 2016. Formal Verification of Ethical Choices in Autonomous Systems. *Robotics and Autonomous Systems* 77 (March 2016), 1–14.
- [17] Nicholas Diakopoulos. 2016. Accountability in Algorithmic Decision Making. *Communications of the ACM (CACM)* 59, 2 (Feb. 2016), 56–62.
- [18] Cynthia Dwork, Moritz Hardt, Toniann Pitassi, Omer Reingold, and Richard Zemel. 2012. Fairness Through Awareness. In *Proceedings of the 3rd Innovations in Theoretical Computer Science Conference (ITCS)*. ACM, Cambridge, 214–226.
- [19] Ernst Fehr and Klaus M. Schmidt. 1999. A Theory of Fairness, Competition, and Cooperation. *The Quarterly Journal of Economics* 114 (1999), 817–868.
- [20] Joan Feigenbaum, Aaron D. Jaggar, and Rebecca N. Wright. 2011. Towards a Formal Model of Accountability. In *Proceedings of the 14th New Security Paradigms Workshop (NSPW)*. ACM, Marin County, California, 45–56.
- [21] Roger Fisher, William L. Ury, and Bruce Patton. 1983. *Getting to Yes: Negotiating Agreement Without Giving In* (3rd ed.). Penguin Books, New York.
- [22] Ricard López Fogués, Pradeep K. Murukannaiah, Jose M. Such, and Munindar P. Singh. 2017. SoSharP: Recommending Sharing Policies in Multiuser Privacy Scenarios. *IEEE Internet Computing (IC)* 21, 6 (Nov. 2017), 28–36. <https://doi.org/10.1109/MIC.2017.4180836>
- [23] Philippa Foot. 1967. The Problem of Abortion and the Doctrine of Double Effect. *Oxford Review* 5 (1967), 5–15.



10.22214/IJRASET



45.98



IMPACT FACTOR:
7.129



IMPACT FACTOR:
7.429



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Call : 08813907089  (24*7 Support on Whatsapp)