

**STUDI KRITIK ETIKA SHANNON VALLOR
TERHADAP *ARTIFICIAL MORAL AGENT* (AMA)**

Skripsi

Diajukan untuk Memenuhi Sebagian

Syarat Memperoleh Gelar Sarjana Agama (S.Ag) dalam Program Studi Aqidah
dan Filsafat Islam



Oleh:

MUHAMMAD ABDUL ROSMI ALWI

NIM: E91219086

**PROGRAM STUDI AQIDAH DAN FILSAFAT ISLAM
FAKULTAS USHULUDDIN DAN FILSAFAT
UNIVERSITAS ISLAM NEGERI SUNAN AMPEL SURABAYA**

2026

PERNYATAAN KEASLIAN

Yang bertanda tangan di bawah ini, saya:

Nama : Muhammad Abdul Rosmi Alwi

NIM : E91219086

Program Studi : Aqidah dan Filsafat Islam

Dengan ini menyatakan bahwa skripsi ini secara keseluruhan adalah hasil penelitian atau karya saya sendiri. Kecuali pada bagian yang dirujuk sumbernya.

Trenggalek, 27 Juni 2026



Muhammad Abdul Rosmi Alwi

NIM: E91219086

PERSETUJUAN PEMBIMBING

Skripsi oleh *Muhammad Abdul Rosmi Alwi* ini telah disetujui untuk diujikan

Surabaya, 8 Juni 2026

Pembimbing,

A stylized, handwritten signature in black ink, likely belonging to Dr. Suhermanto, M.Hum.

Dr. Suhermanto, M.Hum

NIP: 196708201995031001

PENGESAHAN SKRIPSI

Skrispi berjudul “Studi Kritik Etika Shannon Vallor terhadap *Artificial Moral Agent* (AMA)” yang ditulis oleh Muhammad Abdul Rosmi Alwi ini telah diuji di depan Tim Penguji pada tanggal 8 Juni 2026:

1. Dr. Suhermanto, M. Hum (Penguji 1)
NIP: 196708201995031001



2. Dr. Isa Anshori, M. Ag (Penguji 2)
NIP: 197306042005011007



3. Dr. Fikri Mahzumi, M.Fil.I (Penguji 3)
NIP: 198204152015031001



4. Wildah Nurul Islami, M.Th.I (Penguji 4)
NIP: 198509232020122008



Surabaya, 8 Juni 2026

Dekan Fakultas Ushuluddin dan Filsafat



Prof. Abdul Kadir Rivadi, Ph.D.

NIP: 197008132005011003



**KEMENTERIAN AGAMA
UNIVERSITAS ISLAM NEGERI SUNAN AMPEL SURABAYA
PERPUSTAKAAN**

Jl. Jend. A. Yani 117 Surabaya 60237 Telp. 031-8431972 Fax.031-8413300
E-Mail: perpus@uinsby.ac.id

**LEMBAR PERNYATAAN PERSETUJUAN PUBLIKASI
KARYA ILMIAH UNTUK KEPENTINGAN AKADEMIS**

Sebagai sivitas akademika UIN Sunan Ampel Surabaya, yang bertanda tangan di bawah ini, saya:

Nama : Muhammad Abdul Rosmi Alwi
NIM : E91219086
Fakultas/Jurusan : Ushuluddin dan Filsafat/Aqidah dan Filsafat Islam
E-mail address : muhammadalwi323456@gmail.com

Demi pengembangan ilmu pengetahuan, menyetujui untuk memberikan kepada Perpustakaan UIN Sunan Ampel Surabaya, Hak Bebas Royalti Non-Eksklusif atas karya ilmiah :

☒ Skripsi ☐ Tesis ☐ Desertasi ☐ Lain-lain (.....)
yang berjudul :

STUDI KRITIK ETIKA SHANNON VALLOR

TERHADAP *ARTIFICIAL MORAL AGENT* (AMA)

berserta perangkat yang diperlukan (bila ada). Dengan Hak Bebas Royalti Non-Eksklusif ini Perpustakaan UIN Sunan Ampel Surabaya berhak menyimpan, mengalih-media/format-kan, mengelolanya dalam bentuk pangkalan data (database), mendistribusikannya, dan menampilkan/mempublikasikannya di Internet atau media lain secara **fulltext** untuk kepentingan akademis tanpa perlu meminta ijin dari saya selama tetap mencantumkan nama saya sebagai penulis/pencipta dan atau penerbit yang bersangkutan.

Saya bersedia untuk menanggung secara pribadi, tanpa melibatkan pihak Perpustakaan UIN Sunan Ampel Surabaya, segala bentuk tuntutan hukum yang timbul atas pelanggaran Hak Cipta dalam karya ilmiah saya ini.

Demikian pernyataan ini yang saya buat dengan sebenarnya.

Trenggalek, 27 Juni 2026

Muhammad Abdul Rosmi Alwi

ABSTRAK

Nama : Muhammad Abdul Rosmi Alwi
 Judul Skripsi : Studi Kritik Etika Shannon Vallor terhadap *Artificial Moral Agent* (AMA)
 NIM : E91219086

Latar belakang kajian berangkat dari problematika atribusi agensi moral pada mesin akibat reduksionisme komputasional, suatu kondisi yang berisiko mengaburkan kompleksitas ontologis serta kapasitas kognitif dan moral manusia. Evaluasi kritis terhadap postulat Shannon Vallor mengenai *Artificial Moral Agent* (AMA) dilakukan melalui pendekatan filosofis kualitatif berbasis kepustakaan, dengan memadukan analisis konseptual dan kritik normatif. Hasil analisis mengindikasikan adanya keterbatasan signifikan pada mesin, terutama terkait ketiadaan *phronesis* dan keterlibatan eksistensial autentik sebagai prasyarat kesadaran etis. Absennya dimensi ontologis tersebut mengindikasikan keterbatasan paradigma *value alignment* konvensional dalam menghadapi opasitas teknososial (*technosocial opacity*) sistem otonom modern. Kajian ini menegaskan tantangan etis utama bukanlah anomali komputasional, melainkan risiko *moral deskilling*, yakni penurunan kualitas penilaian moral akibat pendelegasian deliberasi etis secara tidak kritis kepada arsitektur non-sentien. Secara teoretis, kajian tersebut menegaskan urgensi pergeseran menuju etika kebajikan teknomoral (*technomoral virtue ethics*) demi mempertahankan kapasitas reflektif manusia dalam mengambil keputusan etis. Merespons hal tersebut, diperlukan integrasi gesekan deliberatif (*deliberative friction*) ke dalam arsitektur kecerdasan buatan. Penerapan mekanisme tersebut bertujuan mencegah pasivitas epistemik sekaligus mendorong sikap evaluatif terhadap kalkulasi probabilistik sistem, serta mempertegas batas antara simulasi sintaksis mesin dan pemahaman semantik manusia. Secara keseluruhan, dapat disimpulkan bahwa pelestarian ketajaman moral manusia merupakan prasyarat fundamental yang tidak dapat direduksi sekadar menjadi persoalan teknis dalam pengembangan kecerdasan buatan.

Kata kunci: *Artificial Moral Agent*, Gesekan Deliberatif, *Moral Deskilling*, *Phronesis*, Reduksionisme Komputasional

DAFTAR ISI

STUDI KRITIK ETIKA SHANNON VALLOR TERHADAP <i>ARTIFICIAL MORAL AGENT</i> (AMA)	
PERNYATAAN KEASLIAN	i
PERSETUJUAN PEMBIMBING	ii
PENGESAHAN SKRIPSI	iii
LEMBAR PERNYATAAN PERSETUJUAN PUBLIKASI	iv
ABSTRAK	v
DAFTAR ISI	vi
BAB I PENDAHULUAN	1
1.1 Latar Belakang	1
1.2 Rumusan Masalah	8
1.3 Tujuan Penelitian	8
1.4 Penelitian Terdahulu	9
1.5 Metode Penelitian	15
1.5.1 Jenis Penelitian	16
1.5.2 Sumber Data	17
1.5.3 Teknik Pengumpulan Data	20
1.5.4 Teknis Analisis Data	20
1.6 Kerangka Teoretis	21
1.7 Sistematika Pembahasan	23
1.7.1 Bab I: Pendahuluan	23
1.7.2 Bab II: Kajian Teori	24
1.7.3 Bab III: Metodologi Penelitian	24
1.7.4 Bab IV: Hasil dan Pembahasan	24
1.7.5 Bab V: Kesimpulan dan Saran	25

BAB II KAJIAN TEORI	26
2.1 Akar Genealogis: Etika Kebajikan (<i>Virtue Ethics</i>) Klasik hingga Kontemporer	26
2.1.1 Arsitektur Etika Aretaik Aristoteles	26
2.1.2 Reaktualisasi Neo-Aristotelian Modern.....	30
2.2 Magnis Opus Shannon Vallor: Kerangka <i>Technomoral Virtue Ethics</i>	34
2.2.1 Kondisi <i>Acute Technosocial Opacity</i>	34
2.2.2 Taksonomi 12 Kebajikan Teknomoral	38
2.2.3 Fenomena <i>Moral Deskillling</i>	42
2.3 Teori <i>AI Mirror</i> dan Dekonstruksi Komputasi.....	46
2.3.1 Konsep Cermin Sosio-Teknis (<i>The Speculum</i>)	46
2.3.2 Batasan Sintaksis Tanpa Semantik dalam Pemrosesan Bahasa	49
2.4 Konseptualisasi <i>Artificial Moral Agent</i> (AMA) dalam <i>Machine Ethics</i>	52
2.4.1 Pendekatan Fungsionalis Agensi Artifisial.....	52
2.4.2 Paradigma Rekayasa Etika Mesin (<i>Top-Down & Bottom-Up</i>).....	55
BAB III METODOLOGI PENELITIAN	60
3.1 Genealogi dan Transisi Arsitektural Kecerdasan Artifisial	60
3.2 Mekanika Sistem Probabilistik dan Opasitas Algoritmik	63
3.3 Arsitektur Large Language Models (LLM) dan Sintaksis Tanpa Semantik.....	65
3.4 Konseptualisasi Artificial Moral Agent dalam Diskursus Kontemporer	69
3.5 Pemaparan Objek Formal: Etika Kebajikan Teknomoral	72
3.6 Paradigma AI Mirror sebagai Reflektor Sosio-Teknis	74
BAB IV PEMBAHASAN.....	78
4.1 Konstruksi Kritik Etika Shannon Vallor terhadap <i>Artificial Moral Agent</i> (AMA)	78
4.1.1 Penolakan terhadap Reduksionisme Komputasional	78
4.1.2 Paradigma AI sebagai Spekulum Sosio-Teknis.....	81

4.2 Implikasi Teoretis: Dekonstruksi Paradigma dalam <i>Machine Ethics</i>	85
4.2.1 Kegagalan Etika Prediktif di Tengah Opasitas Teknososial.....	85
4.2.2 Demistifikasi Agensi Mesin	88
4.3 Implikasi Praktis: Kultivasi Kebajikan Teknomoral dan Tata Kelola AI.....	92
4.3.1 Ancaman <i>Moral Deskillling</i> dan Ilusi Kehidupan Tanpa Gesekan	92
4.3.2 Reaktualisasi Kebajikan Teknomoral sebagai Solusi Eksistensial.....	97
BAB V PENUTUP	102
5.1 Kesimpulan	102
5.1.1 Kritik terhadap Ilusi Moralitas Mesin	102
5.1.2 Kegagalan Etika Tradisional dan Jalan Buntu <i>Value Alignment</i>	102
5.1.3 Ancaman <i>Moral Deskillling</i> dan Pentingnya "Gesekan Deliberatif"	103
5.2 Saran.....	104
5.2.1 Saran Praktis (Untuk Pengembang dan Pembuat Kebijakan).....	104
5.2.2 Saran Akademis (Untuk Peneliti dan Filsuf Masa Depan).....	104

UIN SUNAN AMPEL
S U R A B A Y A

DAFTAR PUSTAKA

- Abakare, Chris. "The Origin Of Virtue Ethics: Aristotle's Views." *GNOSI: An Interdisciplinary Journal of Human Theory and Praxis*, Vol. 3, No. 1 (April 2020), 98–112.
- About Ali, Mohamad, Fadi Dornaika, and Jinan Charafeddine. "Agentic AI: A Comprehensive Survey of Architectures, Applications, and Future Directions." *Artificial Intelligence Review*, Vol. 59, No. 1 (November 2025). <https://doi.org/10.1007/s10462-025-11422-4>.
- Ahmad, Amar, Yvonne Vallès, and Youssef Idaghdour. "Bias in AI Systems: Integrating Formal and Socio-Technical Approaches." *Frontiers in Big Data*, Vol. 8(2026). <https://www.frontiersin.org/journals/big-data/articles/10.3389/fdata.2025.1686452>.
- Albous, Mohammad Rashed. "The Category Error of Human-like AI: Why We Hate the Mirror." *AI & SOCIETY*, Vol. 41, No. 4 (April 2026), 4173–4174. <https://doi.org/10.1007/s00146-025-02847-0>.
- Atakishiyev, Shahin, Housam K. B. Babiker, Jiayi Dai, Nawshad Farruque, Teruaki Hayashi, Nafisa Sadaf Hriti, Md Abed Rahman, et al. "Explainability of Large Language Models: Opportunities and Challenges toward Generating Trustworthy Explanations." In *arXiv [Cs.CL]* (20 Oktober 2025). <https://doi.org/10.48550/arXiv.2510.17256>.
- Bertolini, Andrea, and Francesca Episcopo. "Robots and AI as Legal Subjects? Disentangling the Ontological and Functional Perspective." *Frontiers in Robotics and AI*, Vol. 9 (2022). <https://www.frontiersin.org/journals/robotics-and-ai/articles/10.3389/frobt.2022.842213>.
- Blackman, Justin, and Richard Veerapen. "On the Practical, Ethical, and Legal Necessity of Clinical Artificial Intelligence Explainability: An Examination of Key Arguments." *BMC Medical Informatics and Decision Making*, Vol. 25, No. 1 (March 2025). <https://doi.org/10.1186/s12911-025-02891-2>.
- Brożek, Bartosz, and Bartosz Janik. "Can Artificial Intelligences Be Moral Agents?" *New Ideas in Psychology*, Vol. 54 (August 2019). <https://doi.org/10.1016/j.newideapsych.2018.12.002>.
- Brunozzi, Philippe, and Waldemar Brys. "Target-Centred Virtue Ethics: Aristotelian or Confucian?" *Inquiry*, Vol. 68, No. 10 (November 2025). <https://doi.org/10.1080/0020174X.2023.2220125>.
- Budnikov, Mikhail, Anna Bykova, and Ivan P. Yamshchikov. "Generalization Potential of Large Language Models." *Neural Computing and Applications*, Vol. 37, No. 4 (February 2025), 1973–1997. <https://doi.org/10.1007/s00521-024-10827-6>.
- Cardwell, Laurence. "Artificial Companionship: Moral Deskilling in the Era of Social AI." *Cambridge Journal of Artificial Intelligence*, Vol. 1, No. 2 (2024).

- Carel, Havi. "II—Virtue Without Excellence, Excellence Without Health." *Aristotelian Society Supplementary*, Vol. 90, No. 1 (June 2016). <https://doi.org/10.1093/arisup/akw006>.
- Cervantes, José-Antonio, Sonia López, Luis-Felipe Rodríguez, Salvador Cervantes, Francisco Cervantes, and Félix Ramos. "Artificial Moral Agents: A Survey of the Current Status." *Science and Engineering Ethics*, Vol. 26, No. 2 (April 2020), 501–532. <https://doi.org/10.1007/s11948-019-00151-x>.
- Chang, Yupeng, Xu Wang, Jindong Wang, Yuan Wu, Linyi Yang, Kaijie Zhu, Hao Chen, et al. "A Survey on Evaluation of Large Language Models." *ACM Transactions on Intelligent Systems and Technology*, Vol. 15, No. 3 (June 2024), 1–45. <https://doi.org/10.1145/3641289>.
- Chaudhari, Shreyas, Pranjal Aggarwal, Vishvak Murahari, Tanmay Rajpurohit, Ashwin Kalyan, Karthik Narasimhan, Ameet Deshpande, and Bruno Castro da Silva. "RLHF Deciphered: A Critical Analysis of Reinforcement Learning from Human Feedback for LLMs." *ACM Comput. Surv.*, Vol. 58, No. 2 (January 2026), 1–37. <https://doi.org/10.1145/3743127>.
- Chella, Antonio. "Artificial Consciousness: The Missing Ingredient for Ethical AI?" *Frontiers in Robotics and AI*, Vol. 10 (2023). <https://www.frontiersin.org/journals/robotics-and-ai/articles/10.3389/frobt.2023.1270460>.
- Chiarella, Salvatore G., Matteo Laurenzi, Marianna D'Onofrio, Shaun Gallagher, and Antonino Raffone. "Attribution of Consciousness to Non-Human Animals: Insights from AI and Multidimensional Frameworks." *Front. Psychol.*, Vol. 17, No. 1716363 (May 2026). <https://doi.org/10.3389/fpsyg.2026.1716363>.
- Coeckelbergh, Mark. "Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability." *Science and Engineering Ethics*, Vol. 26, No. 4 (August 2020), 2051–2068. <https://doi.org/10.1007/s11948-019-00146-8>.
- Colelough, Brandon C., and William Regli. "Neuro-Symbolic AI in 2024: A Systematic Review." In *arXiv [Cs.AI]* (9 January 2025). <https://doi.org/10.48550/arXiv.2501.05435>.
- Constantinescu, Mihaela, and Roger Crisp. "Can Robotic AI Systems Be Virtuous and Why Does This Matter?" *International Journal of Social Robotics*, Vol. 14, No. 6 (August 2022), 1547–1557. <https://doi.org/10.1007/s12369-022-00887-w>.
- Crespo, Ricardo F. "Aristotle on Agency, Habits and Institutions." *Journal of Institutional Economics*, Vol. 12, No. 4 (December 2016), 867–884. <https://doi.org/10.1017/s1744137416000059>.
- Cruz, Joe, and Patrick Lee Plaisance. "Virtue Ethics and a Technomoral Framework for Online Activism." *International Journal of Communication*, Vol. 15 (2021), 1330–1348.
- Darnell, Catherine, Liz Gulliford, Kristján Kristjánsson, and Panos Paris. "Phronesis and the Knowledge-Action Gap in Moral Psychology and Moral Education: A New Synthesis?" *Hum. Dev.*, Vol. 62, No. 3 (2019), 101–129. <https://doi.org/10.1159/000496136>.

- Donia, Joseph, and James A. Shaw. "Co-Design and Ethical Artificial Intelligence for Health: An Agenda for Critical Research and Practice." *Big Data & Society*, Vol. 8, No. 2 (July 2021). <https://doi.org/10.1177/205395172111065248>.
- Durán, Juan Manuel, and Karin Rolanda Jongsma. "Who Is Afraid of Black Box Algorithms? On the Epistemological and Ethical Basis of Trust in Medical AI." *Journal of Medical Ethics*, Vol. 47, No. 5 (May 2021). <https://doi.org/10.1136/medethics-2020-106820>.
- Dzwonkowska, Dominika. "Is Environmental Virtue Ethics Anthropocentric?" *Journal of Agricultural and Environmental Ethics*, Vol. 31, No. 6 (December 2018), 723–738. <https://doi.org/10.1007/s10806-018-9751-6>.
- Elden, Åke. "Epistemic Automation and the Deformation of the Human: Artificial Intelligence and the Reconfiguration of Theological Anthropology." *Religions*, Vol. 17, No. 5 (2026). <https://doi.org/10.3390/rel17050515>.
- Ferdman, Avigail. "AI, Deskillling, and the Prospects for Public Reason." *Minds and Machines*, Vol. 35, No. 3 (August 2025). <https://doi.org/10.1007/s11023-025-09737-w>.
- Floridi, LuciaNo. *The Ethics of Information*. London, England: Oxford University Press, (2013). <https://doi.org/10.1093/acprof:oso/9780199641321.001.0001>.
- Formosa, Paul, and Malcolm Ryan. "Making Moral Machines: Why We Need Artificial Moral Agents." *AI & Society*, Vol. 36 (2021).
- Frank, Lily Eva. "What Do We Have to Lose? Offloading Through Moral Technologies: Moral Struggle and Progress." *Science and Engineering Ethics*, Vol. 26, No. 1 (February 2020), 369–385. <https://doi.org/10.1007/s11948-019-00099-y>.
- Friedman, Cindy. "Artefacts of Change: The Disruptive Nature of Humanoid Robots Beyond Classificatory Concerns." *Science and Engineering Ethics*, Vol. 31, No. 2 (March 2025). <https://doi.org/10.1007/s11948-025-00532-5>.
- Fritz, Alexis, Wiebke Brandt, Henner Gimpel, and Sarah Bayer. "Moral Agency without Responsibility? Analysis of Three Ethical Models of Human-Computer Interaction in Times of Artificial Intelligence (AI)." *De Ethica*, Vol. 6, No. 1 (2020).
- Gabriel, Iason. "Artificial Intelligence, Values, and Alignment." *Minds and Machines*, Vol. 30, No. 3 (September 2020), 411–437. <https://doi.org/10.1007/s11023-020-09539-2>.
- García-Valdecasas, Miguel. "Are Large Language Models Intentional? The Limits of Referential Grounding." *Philosophy & Technology*, Vol. 39, No. 2 (March 2026). <https://doi.org/10.1007/s13347-026-01079-4>.
- Hacker-Wright, John. "Passions, Virtue, and Rational Life." *Philosophy & Social Criticism*, Vol. 46, No. 2 (February 2020), 131–140. <https://doi.org/10.1177/0191453718810928>.
- Hagendorff, Thilo. "The Ethics of AI Ethics: An Evaluation of Guidelines." *Minds and Machines*, Vol. 30, No. 1 (March 2020), 99–120. <https://doi.org/10.1007/s11023-020-09517-8>.

- Halbig, Christoph. "Virtue vs. Virtue Ethics." *Zeitschrift Für Ethik Und Moralphilosophie*, Vol. 3, No. 2 (October 2020), 301–313. <https://doi.org/10.1007/s42048-020-00078-0>.
- Hamilton, Kyle, Aparna Nayak, Bojan Božić, and Luca Longo. "Is Neuro-Symbolic AI Meeting Its Promises in Natural Language Processing? A Structured Review." *Semantic Web*, Vol. 15, No. 4 (October 2024). <https://doi.org/10.3233/SW-223228>.
- Hasan, Hanif. "Pengantar Metodologi Penelitian Kualitatif." In *Metode Penelitian Kualitatif*, edited by Rudy. Agam: Yayasan Tri Edukasi Ilmiah, 2025.
- Hayes, Paul, and Noel Fitzpatrick. "Narrativity and Responsible and Transparent Ai Practices." *AI & SOCIETY*, Vol. 40, No. 2 (February 2025), 605–625. <https://doi.org/10.1007/s00146-024-01881-8>.
- He, Ran, Jie Cao, and Tieniu Tan. "Generative Artificial Intelligence: A Historical Perspective." *Natl. Sci. Rev*, Vol. 12, No. 5 (May 2025). <https://doi.org/10.1093/nsr/nwaf050>.
- Héder, Mihály. "The Epistemic Opacity of Autonomous Systems and the Ethical Consequences." *AI & SOCIETY*, Vol. 38, No. 5 (October 2023), 1819–1827. <https://doi.org/10.1007/s00146-020-01024-9>.
- Hedfeld, Patrick. "AI as a Socio-Technical Actor: Rethinking Definitions for Ethics and Governance." *AI Ethics*, Vol. 6, No. 2 (April 2026). <https://doi.org/10.1007/s43681-026-01123-1>.
- Hennekes, Bram. "From the Philosopher's Stone to AI: Epistemologies of the Renaissance and the Digital Age." *Philosophies*, Vol. 10, No. 4 (2025). <https://doi.org/10.3390/philosophies10040079>.
- Hermann, Erik. "Leveraging Artificial Intelligence in Marketing for Social Good—An Ethical Perspective." *Journal of Business Ethics*, Vol. 179, No. 1 (August 2022), 43–61. <https://doi.org/10.1007/s10551-021-04843-y>.
- Hoipkemier, Mark. "Critical Realism and Common Goods." *Journal of Critical Realism*, Vol. 15, No. 1 (January 2016), 53–71. <https://doi.org/10.1080/14767430.2016.1124312>.
- Hossain, Delower, and Jake Y. Chen. "A Study on Neuro-Symbolic Artificial Intelligence: Healthcare Perspectives." In *arXiv [Cs.AI]*, (23 March 2025). <https://doi.org/10.48550/arXiv.2503.18213>.
- Hovd, Sigurd N. "Tools of War and Virtue—Institutional Structures as a Source of Ethical Deskillling." *Frontiers in Big Data*, Vol. 5-2022 (2023). <https://www.frontiersin.org/journals/big-data/articles/10.3389/fdata.2022.1019293>.
- Iizuka, Rie. "Situationism, Virtue Epistemology, and Self-Determination Theory." *Synthese*, Vol. 197, No. 6 (June 2020). <https://doi.org/10.1007/s11229-018-1750-7>.
- Jang, Eunhae, and Patrick Lee Plaisance. "Textual and Comparative Analyses on AI Policies: How Big Tech and Their Global Watchdogs Frame Virtues and Responsibility*." *Journal of Media Ethics*, Vol. 40, No. 4 (October 2025), 187–204. <https://doi.org/10.1080/23736992.2025.2580662>.
- Jarvis, Walter P., Danielle M. Logue, and Neilson Journals Publishing. "Cultivating Moral-Relational Judgement in Business Education." *Journal of Business*

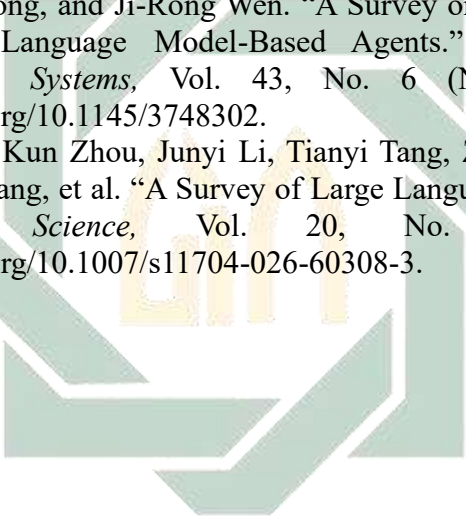
- Ethics Education*, Vol. 13 (2016), 349–372.
<https://doi.org/10.5840/jbee20161316>.
- Jeste, Dilip V., Sarah A. Graham, Tanya T. Nguyen, Colin A. Depp, Ellen E. Lee, and Ho-Cheol Kim. “Beyond Artificial Intelligence: Exploring Artificial Wisdom.” *Int. Psychogeriatr*, Vol. 32, No. 8 (August 2020), 993–1001.
<https://doi.org/10.1017/S1041610220000927>.
- Jeyaraman, Madhan, Sangeetha Balaji, Naveen Jeyaraman, and Sankalp Yadav. “Unraveling the Ethical Enigma: Artificial Intelligence in Healthcare.” *Cureus*, Vol. 15, No. 8 (August 2023). <https://doi.org/10.7759/cureus.43262>.
- Kirk, Hannah Rose, Bertie Vidgen, Paul Röttger, and Scott A. Hale. “Personalisation within Bounds: A Risk Taxonomy and Policy Framework for the Alignment of Large Language Models with Personalised Feedback.” In *arXiv [Cs.CL]*, (9 March 2023).
<https://doi.org/10.48550/arXiv.2303.05453>.
- Klimova, Blanka, Marcel Pikhart, and Jaroslav Kacetl. “Ethical Issues of the Use of AI-Driven Mobile Apps for Education.” *Front. Public Health*, Vol. 10 (2022). <https://doi.org/10.3389/fpubh.2022.1118116>.
- Knoll, Manuel. “TELEOLOGIA NA FILOSOFIA PRÁTICA DE ARISTÓTELES.” *Journal of Teleological Science*, Vol. 1, No. 1 (April 2022), 52–79.
<https://doi.org/10.59079/jts.v1i1.185>.
- Koch, Steffen. “Babbling Stochastic Parrots? A Kripkean Argument for Reference in Large Language Models.” *Journal Philosophy of AI*, Vol. 1 (2025).
<https://doi.org/10.18716/OJS/PHAI/2025.2325>.
- Ladak, Ali, Steve Loughnan, and Matti Wilks. “The Moral Psychology of Artificial Intelligence.” *Current Directions in Psychological Science*, Vol. 33, No. 1 (February 2024), 27–34. <https://doi.org/10.1177/09637214231205866>.
- Lawani-Luwaji Ebidor and Ilegbedion Ikhide. “Literature Review in Scientific Research: An Overview.” Articles. *East African Journal of Education Studies*, Vol. 7, No. 2 (2024). <https://doi.org/10.37284/eajes.7.2.1909>.
- Li, Shurui. “Perfect AI Mimicry and the Epistemology of Consciousness: A Solipsistic Dilemma.” In *arXiv [Cs.AI]*, (6 October 2025).
<https://doi.org/10.48550/arXiv.2510.04588>.
- Lim, Weng Marc. “What Is Qualitative Research? An Overview and Guidelines.” *Australasian Marketing Journal*, Vol. 33, No. 2 (2025).
<https://doi.org/10.1177/14413582241264619>.
- Lin, Ting-an, and Linus Ta-Lun Huang. “AI, Normality, and Oppressive Things.” *Minds and Machines*, Vol. 36, No. 2 (April 2026).
<https://doi.org/10.1007/s11023-026-09778-9>.
- List, Christian. “Probabilistically Coherent Credences despite Opacity.” *Economics and Philosophy*, Vol. 40, No. 2 (Cambridge Core: 2024), 497–506.
<https://doi.org/10.1017/S0266267124000038>.
- Lopez, Paola. “Bias Does Not Equal Bias: A Socio-Technical Typology of Bias in Data-Based Algorithmic Systems.” *Internet Pol. Rev*, Vol. 10, No. 4 (December 2021). <https://doi.org/10.14763/2021.4.1598>.
- Lumbreras, Sara. “The Limits of Machine Ethics.” *Religions*, Vol. 8, No. 5 (2017).
<https://doi.org/10.3390/rel8050100>.

- M. A. K. Raiaan, M. S. H. Mukta, K. Fatema, N. M. Fahad, S. Sakib, M. M. J. Mim, J. Ahmad, M. E. Ali, and S. Azam. "A Review on Large Language Models: Architectures, Applications, Taxonomies, Open Issues and Challenges." *IEEE Access*, Vol. 12 (2024). <https://doi.org/10.1109/ACCESS.2024.3365742>.
- McGregor, Tate. "The Philosophy of Artificial Wisdom." *Philosophy & Technology*, Vol. 38, No. 4 (September 2025). <https://doi.org/10.1007/s13347-025-00964-8>.
- McLoughlin, Shane, Verónica Fernández-Espinosa, and Marta Velázquez-Gil. "Practical Wisdom, Emerging Technologies, and AI." *Journal of Moral Education*, (17 March 2026), 1–31. <https://doi.org/10.1080/03057240.2026.2624863>.
- Mladin, Narcisa C., Dana Rad, Dumitru Ș. Coman, Miron G. Popescu, Maria I. Felea, Radiana Marcu, and Gavril Rad. "Algorithmic Habituation: A Neurocognitive and Systems-Based Framework for Human–AI Co-Adaptation." *Brain Sciences*, Vol. 16, No. 5 (2026). <https://doi.org/10.3390/brainsci16050473>.
- Mouta, Ana, Ana María Pinto-Llorente, and Eva María Torrecilla-Sánchez. "Uncovering Blind Spots in Education Ethics: Insights from a Systematic Literature Review on Artificial Intelligence in Education." *International Journal of Artificial Intelligence in Education*, Vol. 34, No. 3 (September 2024). <https://doi.org/10.1007/s40593-023-00384-9>.
- Naveed, Humza, Asad Ullah Khan, Shi Qiu, Muhammad Saqib, Saeed Anwar, Muhammad Usman, Naveed Akhtar, Nick Barnes, and Ajmal Mian. "A Comprehensive Overview of Large Language Models." *ACM Transactions on Intelligent Systems and Technology*, Vol. 16, No. 5 (October 2025), 1–72. <https://doi.org/10.1145/3744746>.
- Nawaz, Uzma, Mufti Anees-ur-Rahaman, and Zubair Saeed. "A Review of Neuro-Symbolic AI Integrating Reasoning and Learning for Advanced Cognitive Systems." *Intelligent Systems with Applications*, Vol. 26, No. 200541 (June 2025). <https://doi.org/10.1016/j.iswa.2025.200541>.
- Nyholm, Sven. "Responsibility Gaps, Value Alignment, and Meaningful Human Control over Artificial Intelligence." In *Risk and Responsibility in Context*, 1st ed., edited by Adriana Placani and Stearns Broadhead (New York: Routledge, 2023), 15–38.
- Omeljančiuk, Aleksej, Eimantas Peičius, Aušra Urbonienė, and Gvidas Urbonas. "Beyond Algorithmic Oversight: Internal Morality of Medicine and Meaningful Human Control in AI-Assisted Care." *Healthcare*, Vol. 14, No. 12 (2026). <https://doi.org/10.3390/healthcare14121638>.
- Palese, Rosario. "Artificial Truth: Algorithmic Power, Epistemic Authority, and the Crisis of Democratic Knowledge." *Societies*, Vol. 16, No. 3 (2026). <https://doi.org/10.3390/soc16030102>.
- Plaisance, Patrick Lee, and Martina Piantoni. "Humility: A Foundational Virtue for Digital Life." *Journal of Media Ethics*, Vol. 41, No. 1 (January 2026), 19–33. <https://doi.org/10.1080/23736992.2025.2536839>.

- Poznic, Michael, and Erik Fisher. "The Integrative Expert: Moral, Epistemic, and Poietic Virtues in Transformation Research." *Sustainability*, Vol. 13, No. 18 (2021). <https://doi.org/10.3390/su131810416>.
- Rahardjo, Mudjia. *Triangulasi Dalam Penelitian Kualitatif*. <https://uin-malang.ac.id/r/101001/triangulasi-dalam-penelitian-kualitatif.html>.
- Rani, Manecha, Bhupesh Kumar Mishra, and Dhavalkumar Thakker. "Advancing Symbolic Integration in Large Language Models: Beyond Conventional Neurosymbolic AI." In *arXiv [Cs.AI]*, (24 October 2025). <https://doi.org/10.48550/arXiv.2510.21425>.
- Reijers, Wessel. "Technology and Civic Virtue." *Philosophy & Technology*, Vol. 36, No. 4 (October 2023), 71. <https://doi.org/10.1007/s13347-023-00669-w>.
- S. Abdullahi, K. Usman Danyaro, A. Zakari, I. Abdul Aziz, N. Amila Wan Abdullah Zawawi, and S. Adamu. "Time-Series Large Language Models: A Systematic Review of State-of-the-Art." *IEEE Access*, Vol. 13 (2025). <https://doi.org/10.1109/ACCESS.2025.3535782>.
- Sabu, Fransiskus Xaverius, and Taufik Asmiyanto. "Rethinking Moral Agency: Luciano Floridi and the Ethics of AI in the Age of ChatGPT." *Khazanah Al-Hikmah*, Vol. 14, No. 1 (March 2026), 71–84. <https://doi.org/10.24252/kah.v14i1a5>.
- Sado, Fatai, Chu Kiong Loo, Wei Shiung Liew, Matthias Kerzel, and Stefan Wermter. "Explainable Goal-Driven Agents and Robots - A Comprehensive Review." *ACM Computing Surveys*, Vol. 55, No. 10 (October 2023), 1–41. <https://doi.org/10.1145/3564240>.
- Sambrotta, Mirco. "LLMs and the Logical Space of Reasons." *Minds and Machines*, Vol. 35, No. 4 (November 2025), 46. <https://doi.org/10.1007/s11023-025-09751-y>.
- Sandis, Constantine. "Virtue Ethics and Particularism." *Aristotelian Society Supplementary*, Vol. 95, No. 1 (July 2021), 205–232. <https://doi.org/10.1093/arisup/akab013>.
- Schuster, Nick. "The Skill Model: A Dilemma for Virtue Ethics." *Ethical Theory and Moral Practice*, Vol. 26, No. 3 (July 2023), 447–461. <https://doi.org/10.1007/s10677-023-10380-6>.
- Shool, Sina, Sara Adimi, Reza Saboori Amleshi, Ehsan Bitaraf, Reza Golpira, and Mahmood Tara. "A Systematic Review of Large Language Model (LLM) Evaluations in Clinical Medicine." *BMC Medical Informatics and Decision Making*, Vol. 25, No. 1 (March 2025). <https://doi.org/10.1186/s12911-025-02954-4>.
- Siala, Haytham, and Yichuan Wang. "SHIFTing Artificial Intelligence to Be Responsible in Healthcare: A Systematic Review." *Social Science & Medicine*, Vol. 296 (Maret 2022). <https://doi.org/10.1016/j.socscimed.2022.114782>.
- Steen, Marc, Martin Sand, Ibo Van de Poel, and Philosophy Documentation Center. "Virtue Ethics for Responsible Innovation." *Business and Professional Ethics Journal*, Vol. 40, No. 2 (2021), 243–68. <https://doi.org/10.5840/bpej2021319108>.

- Stenseke, Jakob. "Artificial Virtuous Agents in a Multi-Agent Tragedy of the Commons." *AI & Society*, Vol. 39, No. 3 (2024), 855–872. <https://doi.org/10.1007/s00146-022-01569-x>.
- Sulung, Undari, and Mohamad Muspawi. "Memahami Sumber Data Penelitian: Primer, Sekunder, Dan Tersier." *Edu Research*, Vol. 5, No. 3 (2024).
- Tollon, Fabio. "Technology and the Situationist Challenge to Virtue Ethics." *Science and Engineering Ethics*, Vol. 30, No. 10 (March 2024). <https://doi.org/10.1007/s11948-024-00474-4>.
- Tolmeijer, Suzanne, Markus Kneer, Cristina Sarasua, Markus Christen, and Abraham Bernstein. "Implementations in Machine Ethics." *ACM Computing Surveys*, Vol. 53, No. 6 (November 2021), 1–38. <https://doi.org/10.1145/3419633>.
- Tsai, Cheng-hung, and Hsiu-lin Ku. "Why AI May Undermine Phronesis and What to Do about It." *AI and Ethics*, Vol. 5, No. 3 (June 2025), 3079–3086. <https://doi.org/10.1007/s43681-024-00617-0>.
- Vallor, Shannon. "Moral Deskillling and Upskilling in a New Machine Age: Reflections on the Ambiguous Future of Character." *Philosophy & Technology*, Vol. 28, No. 1 (2015) 107–124. <https://doi.org/10.1007/s13347-014-0156-9>.
- . *Technology and the Virtues: A Philosophical Guide to a Future Worth Wanting*. New York: Oxford University Press, 2016.
- . *The AI Mirror: How to Reclaim Our Humanity in an Age of Machine Thinking*. New York: Oxford University Press, 2024.
- Vallor, Shannon, and Tillmann Vierkant. "Find the Gap: AI, Responsible Agency and Vulnerability." *Minds and Machines*, Vol. 34, No. 3 (June 2024). <https://doi.org/10.1007/s11023-024-09674-0>.
- Wallach, Wendell, and Colin Allen. *Moral Machines: Teaching Robots Right from Wrong*. Oxford University Press, 2009. <https://doi.org/10.1093/acprof:oso/9780195374049.001.0001>.
- Wan, Zishen, Che-Kai Liu, Hanchen Yang, Chaojian Li, Haoran You, Yonggan Fu, Cheng Wan, Tushar Krishna, Yingyan Lin, and Arijit Raychowdhury. "Towards Cognitive AI Systems: A Survey and Prospective on Neuro-Symbolic AI." In *arXiv [Cs.AI]* (2 January 2024). <https://doi.org/10.48550/arXiv.2401.01040>.
- Wang, Lei, Chen Ma, Xueyang Feng, Zeyu Zhang, Hao Yang, Jingsen Zhang, Zhiyuan Chen, et al. "A Survey on Large Language Model Based Autonomous Agents." *Frontiers of Computer Science*, Vol. 18, No. 6 (2024). <https://doi.org/10.1007/s11704-024-40231-1>.
- Wang, Shuhe, Shengyu Zhang, Jie Zhang, Runyi Hu, Xiaoya Li, Tianwei Zhang, Jiwei Li, Fei Wu, Guoyin Wang, and Eduard Hovy. "Reinforcement Learning Enhanced LLMs: A Survey." In *arXiv [Cs.CL]* (5 December 2024). <https://doi.org/10.48550/arXiv.2412.10400>.
- Weiner, Ellison B., Irene Dankwa-Mullan, William A. Nelson, and Saeed Hassanpour. "Ethical Challenges and Evolving Strategies in the Integration of Artificial Intelligence into Clinical Practice." *PLOS Digit. Health*, Vol. 4, No. 4 (April 2025). <https://doi.org/10.1371/journal.pdig.0000810>.

- Widłak, Tomasz. “Techno-Legal Virtues in Hybrid Adjudication: Judicial Excellence in AI-Augmented Courts.” *Minds and Machines*, Vol. 36, No. 1 (November 2025). <https://doi.org/10.1007/s11023-025-09756-7>.
- Wong, Pak-Hang. “Rituals and Machines: A Confucian Response to Technology-Driven Moral Deskillling.” *Philosophies*, Vol. 4, No. 4 (2019). <https://doi.org/10.3390/philosophies4040059>.
- Woodruff, Earl, and Jim Hewitt. “Epistemic Agency in the Age of Large Language Models: Design Principles for Knowledge-Building AI.” *AI*, Vol. 7, No. 3 (2026). <https://doi.org/10.3390/ai7030099>.
- Zafar, Mandy. “Normativity and AI Moral Agency.” *AI and Ethics*, Vol. 5, No. 3 (June 2025), 2605–2622. <https://doi.org/10.1007/s43681-024-00566-8>.
- Zhang, Zeyu, Quanyu Dai, Xiaohe Bo, Chen Ma, Rui Li, Xu Chen, Jieming Zhu, Zhenhua Dong, and Ji-Rong Wen. “A Survey on the Memory Mechanism of Large Language Model-Based Agents.” *ACM Transactions on Information Systems*, Vol. 43, No. 6 (November 2025), 1–47. <https://doi.org/10.1145/3748302>.
- Zhao, Wayne Xin, Kun Zhou, Junyi Li, Tianyi Tang, Zican Dong, Yupeng Hou, Beichen Zhang, et al. “A Survey of Large Language Models.” *Frontiers of Computer Science*, Vol. 20, No. 12 (May 2026). <https://doi.org/10.1007/s11704-026-60308-3>.



UIN SUNAN AMPEL
SURABAYA