

Title	知識の安全保障化とオープンサイエンスの共存可能性：問題圏の解明に向けて
Author(s)	岸本, 眞一郎
Citation	年次学術大会講演要旨集, 40: 339-341
Issue Date	2025-11-08
Type	Conference Paper
Text version	publisher
URL	https://hdl.handle.net/10119/20253
Rights	本著作物は研究・イノベーション学会の許可のもとに掲載するものです。This material is posted here with permission of the Japan Society for Research Policy and Innovation Management.
Description	一般講演要旨

1 E O 2

知識の安全保障化とオープンサイエンスの共存可能性：問題圏の解明に向けて

○岸本眞一郎（JAIST 支援機構）
Kishimoto.slr.rs@jaistso.or.jp

1. はじめに

オープンサイエンスは今日まで、科学技術および学術ソサエティーを活性化する方法として普及してきた。オープンサイエンスの潮流は、民間のデファクトスタンダードの普及のみならず、公共セクターの科学技術イノベーション（STI）政策の後押しを受けて成長してきた。たとえば、2015 年の OECD 報告 “Making Open Science a Reality” は、オープンアクセスと研究データの公開・再利用を政策目的として定式化した。その後、2021 年には UNESCO 総会が各国に実施を勧告する国際規范文書を採択した。欧州では 2018 年の Plan S と Rights Retention Strategy、英国では UKRI の OA ポリシー、米国では 2022 年 OSTP メモ（通称 Nelson Memo）がオープンアクセスを原則化するなど、民間のデファクト標準（プレプリント、OS ソフトウェア、生産性基盤）と公共セクターの STI 政策が相互補完的に普及を押し進めてきた。加えて、2020 年代以降の LLM 研究では、Transformer (2017)、BERT (2018)、GPT-3 (2020)、LLaMA (2023) など主要論文の多くが arXiv に先行公開されている。LLM 研究では、他者に先んじてオープンにし、迅速に priority を獲得することで、その業界でのポジショニング戦略を優位に進めようとする方略が定着している。

他方、オープンサイエンスが科学技術および学術ソサエティーにとって、常に正しく有益な方向に向かっているか否かについては、さほど議論されていない。管見の限り、オープンサイエンスが正しく有益な潮流であることを前提とした論調が多くみられる。

本論では、オープンサイエンスの潮流が自らの目的を破壊する方向にはたらく可能性について検討する。そのために、本論では “知識の安全保障化” という分析概念を導入する。この導入により、素朴なオープン思考のみでは解決しがたい、よりディープな潜在的問題を指摘することで、オープンサイエンスに係る STI 政策研究に新たな問題領域の提案を試みる。

2. 先行研究レビュー

2-1. 安全保障化とはなにか？

安全保障化（securitization）は、安全保障学・国際政治学におけるコペンハーゲン学派（Buzan／Wæver／de Wilde）が提示した理論枠組みである。安全保障化とは、政治主体（安全保障化の主体）がなにかしらの対象を生存的脅威（existential threat）」として言説化し、それを聴衆が受容することで、当該対象について、日常政治の手続きとは異なる例外的措置が正当化されることをさす。基礎文献として Buzan et al. (1998) と、その後の Balzacq (編) (2010) や Stritzel (2007) による言説・文脈・実践・受容過程の精緻化が広く参照されている。E-IR の入門解説（Eroukhmanoff, 2017/2018）では、「誰が／何を／いかにして脅威と称し／誰が受け入れるか」という相互行為の力学が俯瞰的に整理されている。

3. 概念化および問題化

概念化 -- 知識の安全保障化について： 本稿で暫定的に導入する「知識の安全保障化（Knowledge Securitization）」とは、科学研究・技術開発・イノベーション（STI）の個別領域が、技術安全保障や経済安全保障の語彙・制度・実務の下で再編される過程をスコープに収めようとする分析概念である。これにより：

- (a) 言説（脅威化の言語行為／政策文書の語彙変化）
- (b) 制度（研究セキュリティ・開示・輸出管理等の要件化）
- (c) 実務（大学・企業のデューディリジェンス，共同研究の管理，データ・サイバー体制）

これらを一体で把握でき、オープンサイエンスとの相互最適化やトレードオフの設計課題に光を当てることができる。

実際、米国では NSPM-33 実装ガイダンス（2022）および OSTP の研究セキュリティ・プログラム指針（2024）において、サイバー、海外渡航、トレーニング、輸出管理の 4 要素の実装が要求されている。EU では、外国干渉対処の実務文書（2022）と理事会勧告（2024）が整備されており、英国では NPSA/NCSC の “Trusted Research” の改訂ガイダンスにより大学・資金配分機関の行動基準を具体化している。

これらの枠組みは、FAIR 原則（2016）や DORA（2012-）等が推進したオープン化・評価改革のデファクト化（標準・慣行の定着）と、研究セキュリティ制度のディジュール化（要件・認証の法規範化）を同時に追跡でき、STI 政策研究における問いの創造の豊饒さ（enrichness）を増す道具だてのひとつとして有効であろうと考えられる。

概念化 -- 構成主義の採用： 知識の安全保障化の概念は、コペンハーゲン学派の系譜にある構成主義——すなわち、言説と受容が安全保障領域を構成するという考え方を前提に置く。たとえば、Balzacq（2010）は安全保障化の成功条件を「受容」「文脈」「実践」に分解し、Stritzel（2007）は発話だけでなく社会的埋め込みの寄与を強調している。構成主義のパースペクティブは、STI 政策におけるオープンネス、安全性、競争力といった価値の相互構成を検証するための鍵となる。

問題化： 知識の安全保障化という切り口で分析した場合、オープンサイエンスはオープンにされた研究開発データや言説のデファクト化に寄与する一方、そのデファクト化を狙う言語行為そのものが言論操作（Manipulation）を目的としておこなわれる可能性を指摘できる。

加えて、LLM による大量のテキストデータや図表データのスクレイピングと意味内容読取り、自動翻訳や自動要約は、もとの学習データのパターンによって biased された大量のテキストデータを生成する。すなわち、たとえ意図せずとも、LLM による Quasi-manipulative な生成が起こり得る。それらの知識が普及してしまえば、認識主体は普及された知識にもとづきテキストデータの真実性や正当性を判断するため、認識主体たちは質的なバイアス操作に気づきづらい。LLM 研究ではすでに、合成データ反復学習が性能崩壊（“model collapse”）を招き得ること、ベンチマーク汚染は評価の信頼性を損なうことが報告されている。

安全保障的パースペクティブが説得力をもつ場では、オープンサイエンスはこのような認識的リスクを自己産出し得る。公開性はオープンサイエンスの本来的な強みであるが、質的に biased された知識や偏ったテキスト生成物の反復的参照が社会に受容されると、受容関数（に相当する因子）が通増し、容易に閾値を超え、対抗言説や修正に必要な時間的余裕を奪う可能性がある。

STI 政策におけるこれらの潜在的リスクの影響は、国家的なイノベーションのコントローラビリティのみならず、科学技術の再現可能性の品質にも及び得る。今後、安全保障の 이슈がわが国および国際社会において重要性を増していくなかで、STI 政策形成の実務および政策研究においては、そこで取り扱われる STI 知識の再現可能性にフォーカスする必要があると考えられる。

4. 今後の研究展望

本論にて展開した概念化および問題化を精緻に議論し、社会的なリスクマネジメントのひとつとして知識安全保障の志向に関する合意形成をおこなっていくためには、今後、知識およびその安全保障についての哲学的あるいは学問論的基礎が必要になる。たとえば、Williamson（2000）の知識第一主義の認識論にもとづけば、知識と証拠（エビデンス）は同じ概念であり、密接な関係にある。今後は、エビデンス概念の明晰化と具体的実務の構成を含め、本論で提案した潜在的問題についてのより深い探求をお

こなっていきたい。

参考文献

- [1] David, Paul A. 2007. “The Historical Origins of ‘Open Science.’” SIEPR Discussion Paper 06-038.
- [2] David, Paul A. 2003. “The Economic Logic of Open Science and the Balance between Private Property Rights and the Public Domain.” SIEPR Discussion Paper 02-30.
- [3] Merton, Robert K. 1973 [1942]. “The Normative Structure of Science.” In *The Sociology of Science: Theoretical and Empirical Investigations*, ed. Norman W. Storer, 267–278. Chicago: University of Chicago Press.
- [4] OECD. 2015. *Making Open Science a Reality*. OECD Science, Technology and Industry Policy Papers 25.
- [5] UNESCO. 2021. *Recommendation on Open Science*. Paris: UNESCO.
- [6] UK Research and Innovation. 2021. *UKRI Open Access Policy*.
- [7] Office of Science and Technology Policy. 2022. *Ensuring Free, Immediate, and Equitable Access to Federally Funded Research*. Washington, DC: Executive Office of the President.
- [8] Fecher, Benedikt, and Sascha Friesike. 2014. “Open Science: One Term, Five Schools of Thought.” In *Opening Science*, ed. Sönke Bartling and Sascha Friesike. Springer.
- [9] DORA (San Francisco Declaration on Research Assessment). 2012. *San Francisco Declaration on Research Assessment*.
- [10] Hicks, Diana, Paul Wouters, Ludo Waltman, Sarah de Rijcke, and Ismael Rafols. 2015. “The Leiden Manifesto for Research Metrics.” *Nature* 520: 429–431.
- [11] European Commission. 2022. *Tackling R&I Foreign Interference—Staff Working Document*.
- [12] National Protective Security Authority (NPSA) and National Cyber Security Centre (NCSC). 2024–2025. *Trusted Research: Guidance*.
- [13] Vaswani, Ashish, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. “Attention Is All You Need.” arXiv:1706.03762.
- [14] Devlin, Jacob, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. “BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding.” In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL-HLT 2019) (ACL-N19-1423)*. See also arXiv:1810.04805 (2018).
- [15] Brown, Tom B., Benjamin Mann, Nick Ryder, et al. 2020. “Language Models Are Few-Shot Learners.” In *Advances in Neural Information Processing Systems 33*. See also arXiv:2005.14165.
- [16] Touvron, Hugo, Thibaut Lavril, Gautier Izacard, et al. 2023. “LLaMA: Open and Efficient Foundation Language Models.” arXiv:2302.13971.
- [17] 『「安全保障化」とは何か——脅威をめぐる政治力学』. 2024. ミネルヴァ書房.
- [18] Balzacq, Thierry, ed. 2010. *Securitization Theory*. Routledge.
- [19] Balzacq, Thierry. “Towards a Theory of Securitization: Copenhagen and Beyond.”
- [20] Executive Office of the President. 2022. *NSPM-33 Implementation Guidance*. Washington, DC.
- [21] European Commission. 2022. *Tackling R&I Foreign Interference*.
- [22] National Protective Security Authority (NPSA). n.d. *Trusted Research*. London: NPSA.
- [23] Wilkinson, Mark D., et al. 2016. “The FAIR Guiding Principles for Scientific Data Management and Stewardship.” *Scientific Data* 3: 160018.
- [24] DORA (San Francisco Declaration on Research Assessment). 2012. *San Francisco Declaration on Research Assessment*.
- [25] Bender, Emily M., Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmitchell. 2021. “On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?” In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAcT ’21)*.
- [26] Shumailov, Ilia, et al. 2023. “The Curse of Recursion: Training on Generated Data Makes Models Forget.” arXiv:2305.17493.
- [27] 2024. “Benchmark Data Contamination of Large Language Models: A Survey.” arXiv:2406.04244.
- [28] Williamson, Timothy. 2000. *Knowledge and Its Limits*. Oxford: Oxford University Press.