

A Perspective on Neuroscience Data Standardization with Neurodata Without Borders

 **Andrea Pierré**,^{1,2*} **Tuan Pham**,^{1,2*}  **Jonah Pearl**,³  **Sandeep Robert Datta**,³  **Jason T. Ritt**,^{2*} and **Alexander Fleischmann**^{1,2}

¹Department of Neuroscience, Division of Biology and Medicine, Brown University, Providence, Rhode Island 02912, ²The Robert J. and Nancy D. Carney Institute for Brain Science, Brown University, Providence, Rhode Island 02912, and ³Department of Neurobiology, Harvard Medical School, Boston, Massachusetts 02115

Neuroscience research has evolved to generate increasingly large and complex experimental data sets, and advanced data science tools are taking on central roles in neuroscience research. Neurodata Without Borders (NWB), a standard language for neurophysiology data, has recently emerged as a powerful solution for data management, analysis, and sharing. We here discuss our labs' efforts to implement NWB data science pipelines. We describe general principles and specific use cases that illustrate successes, challenges, and non-trivial decisions in software engineering. We hope that our experience can provide guidance for the neuroscience community and help bridge the gap between experimental neuroscience and data science. Key takeaways from this article are that (1) standardization with NWB requires non-trivial design choices; (2) the general practice of standardization in the lab promotes data awareness and literacy, and improves transparency, rigor, and reproducibility in our science; (3) we offer several feature suggestions to ease the extensibility, publishing/sharing, and usability for NWB standard and users of NWB data.

Key words: big data neuroscience; collaborative science; data standardization with neurodata without borders; data science pipelines

Significance Statement

Neuroscience research generates increasingly large and complex data sets, requiring advanced data science tools. Neurodata Without Borders (NWB), an National Institutes of Health (NIH)-supported standard, is crucial for data management, analysis, and sharing. Despite its importance, adoption of NWB remains challenging for many labs. This article addresses practical, technical, social, and institutional aspects of data handling, and discusses common data science challenges in neuroscience. By sharing real-world experiences, we aim to stimulate discussion among key stakeholders such as tool developers, labs considering NWB adoption, and funding agencies promoting open science.

Received Feb. 27, 2024; revised July 24, 2024; accepted July 30, 2024.

Author contributions: A.P., T.P., J.P., S.R.D., J.T.R., and A.F. designed research; A.P., T.P., J.P., and J.T.R. performed research; A.P., T.P., J.P., and J.T.R. analyzed data; A.P., T.P., J.P., S.R.D., J.T.R., and A.F. wrote the paper.

We thank Simon Daste and Max Seppo for their input and provision of experimental data. We thank the Osmonauts (U19NS112953), Cindy Poo, Chris Rodgers, Rebecca Tripp, and Emily Ventriglia for helpful comments on earlier drafts. We thank the NWB, DANDI, and CatalystNeuro teams, including Cody Baker, Ben Dichter, Garrett Flynn, Satrajit Ghosh, Yaroslav Halchenko, and Oliver Rübél, for extensive discussions and helpful comments following posting of a preprint draft on arXiv. Work in the Fleischmann and Datta labs was supported by NIH Award U19NS112953. Work in the A.F. and J.T.R. labs was also supported by NIH award R01DC017437, and the Robert J. and Nancy D. Carney Institute for Brain Science. Carney Institute computational resources used in this work were supported by the NIH Office of the Director Award S100D025181. Parts of this article were processed using ChatGPT (<https://chat.openai.com/>) and Claude (<https://claude.ai/>) to improve language.

*A.P., T.P., and J.T.R. contributed equally to this work.

The authors declare no competing financial interests.

Correspondence should be addressed to Jason T. Ritt at jason_ritt@brown.edu or Alexander Fleischmann at alexander_fleischmann@brown.edu.

<https://doi.org/10.1523/JNEUROSCI.0381-24.2024>

Copyright © 2024 the authors

Introduction

Increasing complexity of neuroscience data

Over the past 20 years, neuroscience research has been radically changed by two major trends in data production and analysis. First, neuroscience research now routinely generates large data-sets of high complexity. Examples include recordings of activity across large populations of neurons, often with high-resolution behavioral tracking (Mathis et al., 2018; Steinmetz et al., 2019; Stringer et al., 2019; Siegle et al., 2021; Koch et al., 2022), analyses of neural connectivity at high spatial resolution and across large brain areas (Scheffer et al., 2020; Loomba et al., 2022), and detailed molecular profiling of neural cells (Callaway et al., 2021; Braun et al., 2022; Langlieb et al., 2023; Yao et al., 2023). Such large, multi-modal data sets are essential for solving major questions about brain function (Jorgenson et al., 2015; Brose, 2016; Koch and Jones, 2016).

Second, the collection and analysis of such datasets requires interdisciplinary teams, incorporating expertise in systems neuroscience, engineering, molecular biology, data science, and theory. These two trends are reflected in the increasing numbers of authors on scientific publications (Plume and Weijen, 2014) (<https://the-publicationplan.com/2016/12/13/new-investigation-reveals-the-number-of-authors-named-on-research-papers-is-increasing>) and the creation of mechanisms to support team science by the National Institutes of Health (NIH) and similar research funding bodies (Cooke and Hilton, 2015; Brose, 2016).

There is also an increasing scope of research questions that can be addressed by aggregating “open data” from multiple studies across independent labs. Funding agencies and publishers have begun to aggressively promote data sharing and open data, with the goals of improving reproducibility and increasing data reuse (Dallmeier-Tiessen et al., 2014; Tenopir et al., 2015; Pasquetto et al., 2017). However, open data may be unusable if scattered in a wide variety of naming conventions and file formats lacking machine-readable metadata.

Big data and team science necessitate new strategies for how to best organize data, with a key technical challenge being the development of standardized file formats for storing, sharing, and querying datasets. Prominent examples include the Brain Imaging Data Structure (BIDS; RRID:SCR_016124) for neuroimaging and Neurodata Without Borders (NWB; RRID:SCR_015242) for neurophysiology data (Teeters et al., 2015; Gorgolewski et al., 2016; Holdgraf et al., 2019; Rübél et al., 2022). The Open Neurophysiology Environment (ONE), best known from adoption by The International Brain Laboratory (IBL; [The International Brain Laboratory](https://www.international-brain-laboratory.org/), 2020, 2023), has a similar application domain to NWB, but a highly different technical design. There have been many standards over the years, including NIX (RRID:SCR_016196)—neuroscience information exchange format (Martone et al., 2020), neuroscience information framework (NIF; RRID:SCR_002894) (Gardner et al., 2008), minimum information about a neuroscience

investigation (MINI) for electrophysiology (Gibson et al., 2009), BRAINformat (Rübél et al., 2015), odML (RRID:SCR_001376; Grewe et al., 2011), Scalable Open Network Architecture Template for neural models (Dai et al., 2020), and the recent NeuroBlueprint <https://neuroblueprint.neuroinformatics.dev/>. Not all standards remain actively used, but NWB is one of them. These initiatives provide technical tools for storing and accessing data in known formats, but more importantly provide conceptual frameworks with which to standardize data organization and description in an (ideally) universal, interoperable, and machine-readable way.

Our labs' history in implementing NWB-based standardization

In 2019, the Fleischmann and Ritt labs initiated a collaboration to enhance the Fleischmann lab's data science and computational tooling and workflows. We expanded our team by hiring two research software engineers (RSEs), and by extending collaborations with data scientists and computational biologists. Similar efforts were underway in the Datta lab. An early common goal was the standardization of neurophysiology and behavioral data using a framework such as NWB. In this article, we provide our perspective on opportunities and challenges when adopting NWB data standardization.

Our labs investigate the functions of neural circuits for sensory processing and behavior in mice. Typical experiments include calcium imaging of neuronal activity in awake, head-fixed mice during odor presentation, with a number of behavioral readouts including sniffing, running, and facial movements (Fig. 1). In other experiments, mice are freely moving, with implanted GRIN lenses for miniscope imaging, odor and reward delivery in nose ports, and behavioral readouts including videographic tracking. Our experimental designs, data generation, and analyses are similar to many other labs investigating neural circuit mechanisms for sensory-motor transformations, learning, and memory (Box 1), though each lab has its own idiosyncrasies impinging on data management.

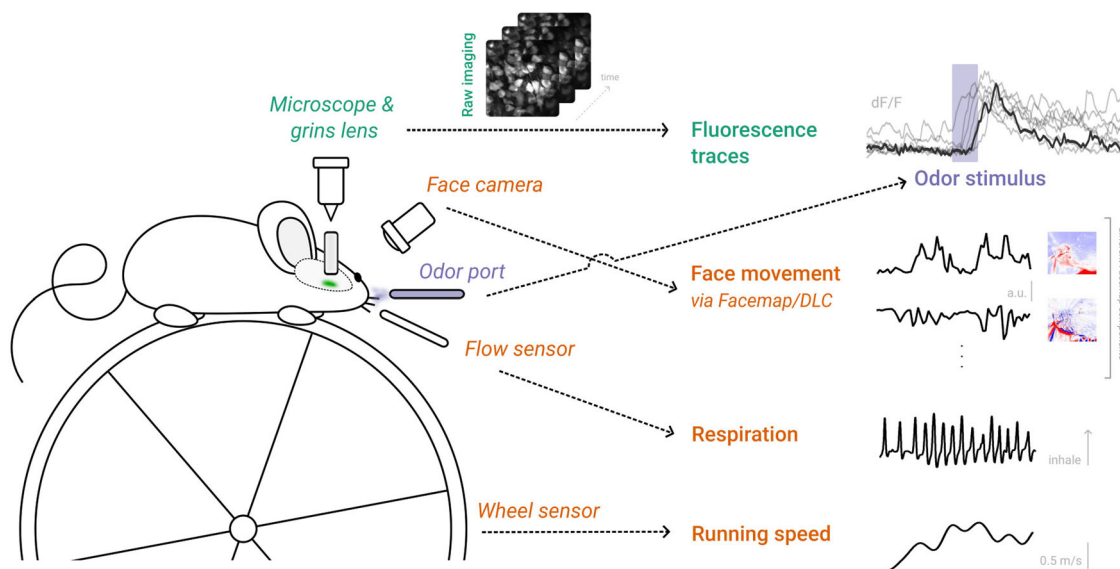


Figure 1. Setup of a typical Fleischmann lab experiment and resulting data streams. The left schematic illustrates in vivo head-fixed two-photon calcium imaging of a deep brain area (e.g., piriform cortex) through a GRIN lens. Throughout the paper, we use the following color scheme: green for neural activity, orange for animal behaviors, and purple for external variables (e.g., stimuli). Raw images from the microscope (top) are preprocessed to obtain fluorescence time series for each segmented neuron (top row, right). The animal receives odor stimuli through an odor port during a time window in each trial, marked by a light purple bar in the fluorescence time series plot. Several behaviors are tracked. A high-resolution camera captures facial movement, typically reduced using the application Facemap into principal components of image motion (middle), or through DeepLabCut into pose estimation or keypoints. Peri-nasal flow and wheel sensors, connected through a microcontroller, provide respiration and running speed estimates, respectively.

Choosing which standard to use could be the topic of its own paper. Here we chose NWB mainly because of our use case of neurophysiology experiments and popularity among our communities of researchers. Depending on the domain and community, researchers can review the guides and portfolio of endorsed standards maintained by organizations such as the International Neuroinformatics Coordinating Facility: <https://www.incf.org/resources/sbps>.

In this article, we first discuss our motivation and general considerations for implementing data standardization. We then describe the implementation of NWB data conversion pipelines, including domain-specific use cases and solutions for data sharing. We conclude by identifying opportunities for improving future user experience. We hope that by describing our experience, other labs planning to adopt NWB will benefit from comparisons with their own needs and capabilities. We also hope to provide a case study that may be informative for developers of NWB and similar data science toolboxes. Furthermore, we acknowledge that our experience with NWB was on the early side of the adoption curve, and that many issues we encountered have been resolved throughout the process. Box 2 summarizes the key insights from this article.

Box 1. Fleischmann Lab workflow

Data acquisition—experiments and systems We perform in vivo calcium imaging experiments in head-fixed (2-photon imaging) and freely moving (miniscope) mice. Experiments include multi-plane, multicolor, and/or multi-day recordings.

Data acquisition—tasks and stimuli In some experiments, animals receive pre-programmed odor stimuli independent of their behavior; in other experiments, sensory stimuli or an animal's behavior can trigger a reward. Behavior recording includes micro-controller-acquired time series (e.g., wheel speed, sniff rate, licks, and rewards) and video recordings of the animal's face or body motion.

Preprocessing Pipelines include conventional calcium imaging steps (e.g., motion correction, segmentation, deconvolution, multicolor or multi-day registration) using existing tools such as Suite2p (RRID:SCR_016434; Pachitariu et al., 2016) and Inscopix (<https://www.inscopix.com/>). Experiments with behavioral videos may also be preprocessed with toolboxes such as DeepLabCut (DLC, RRID:SCR_021391; Mathis et al., 2018) for pose estimation and Facemap (RRID:SCR_021513; Syeda et al., 2022) for facial motion extraction.

Conversion to standard format Raw and preprocessed data streams are integrated and stored in NWB files, using a custom tool, *calimag* (Pierré and Pham, 2023), developed in the Fleischmann lab.

Analyses Questions include stimulus or behavior tuning of single neuron or population activity, as well as how learning and experience shape neural activity.

Box 2. Key insights summary

Standardization with NWB requires non-trivial design choices (see section “How to organize data into a standard format”) **that have to balance between usability, feasibility, sustainability, and varying constraints of time, resources, and experience. This includes considerations at various stages of scientific work** (see sections “When to create and

use the standardized format” and “Pain points in the conversion workflow”). **More specific discussions touch on experiment setup** (e.g., section “Our experience with metadata capture”), **preprocessing** (e.g., section “How should different data types be stored?”), **conversion** (e.g., sections “How to organize data into a standard format” and “Creating NWB extensions allows fitting domain-specific use cases”), **analysis** (e.g., sections “Data access pain points,” “How should different data types be stored?,” “Should one standardize data from intermediate analysis stages?,” “Lab-specific metadata,” and “Odor stimulus metadata”), **and publishing** (sections “Considerations for sharing on DANDI,” “Public data sharing,” and “Where should raw data and supplemental information be stored?”). **This makes data sharing a distinct technique to be mastered by researchers, similar to scientific writing and scientific communication.**

In addition to technical considerations, there is a significant social component of adopting the NWB standard in neurophysiology research (sections “Within a lab,” “Our experience with metadata capture,” and “Social challenges in extension development”) **across a lab's different stakeholders** (section “Key stakeholders in adoption of a new lab standard”) **and research collaborators** (section “Collaboration”).

We observe a general practice of standardization that promotes more data awareness and literacy, and improves communication in our research activities (sections “Within a lab,” “Collaboration,” “Working with acquisition devices and software,” “An indirect benefit of using NWB is improved data awareness,” “Lab-specific metadata,” “Odor stimulus metadata,” and “Social challenges in extension development”), **and increases transparency, rigor, and reproducibility** (sections “Collaboration,” “Lab-specific metadata,” “Odor stimulus metadata,” “Social challenges in extension development,” and “Public data sharing”).

Compared to many neurophysiology standards, NWB is well-supported, well-maintained, and its developer team is very responsive to issues and requests (sections “NWB community” and “Documentation for extension development”). **The standard has a substantial and growing ecosystem of tools and resources** (sections “Public data sharing,” “NWB community,” “Neuroscience community,” “Off the shelf NWB conversion,” “Existing Neurodata Extensions,” and “Data exploration tool guidance”), **which is especially useful for new adopters.**

We encountered several unexpected challenges when working with and developing these open science and open source tools (sections “Our experience with metadata capture,” “Editing and merging of NWB files,” “Pain points in the conversion workflow,” “Timeliness of code contribution acceptance,” “Wishlist for NWB extensions,” “Potential surprises with data validation,” and “Modification of file organization”).

To increase usability and accessibility, we propose more centralized documentation and cataloged discussions on NWB standard development and the ecosystem around it (sections “NWB community” and “Documentation for extension development”). **Additionally, we offer several feature suggestions to ease NWB extensibility** (sections

“Framework extensions” and “Wishlist for NWB extensions”), **publishing** (sections “Considerations for sharing on DANDI” and “Alternatives to DANDI and general strategy with data repositories”), and **usability** (section “Data access pain points”).

Key Stakeholders in Adoption of a New Lab Standard

We first define, in high-level terms, three distinct personnel roles in a typical research lab, each of whom has their own needs and incentives surrounding data standardization. First, **Principal Investigators (PI)** and senior researchers manage research teams, labs, and projects. Second, **researchers** include research trainees (e.g., undergraduate and graduate students, postdoctoral associates), lab technicians, and data scientists, and more generally individuals collecting and/or analyzing data. Lastly, **Research Software Engineers (RSEs)** support researchers by developing and maintaining software, packages, and pipelines for data management, processing, and analysis.

PIs

Key desired outcomes for the adoption of lab-wide standardized data formats include improved efficiency, rigor, reproducibility, and ease of collaboration. Efficiency could follow from using common tools for saving, retrieving, analyzing, and sharing data; technical improvements by one member can have knock-on value for others. Rigor and reproducibility similarly benefit from increased access and scrutiny brought by all lab members being able to see each other's work, instead of working in isolation; data already in standard formats could ease communication and usage. An additional value for PIs is meeting the norms of their field for data management and sharing, including mandates from funding agencies such as the NIH, without requiring extensive ad hoc effort at the time of grant submissions or publication. Key objectives: efficiency, rigor, reproducibility, and collaboration.

However, there are several concerns when introducing standardized formats. PIs generally want to avoid major disruptions to scientific productivity in the lab. There is rarely a good time to slow or halt data collection and analysis in order to fully convert to new pipelines and workflows. On the other hand, a gradual transition can paradoxically lead to greater friction due to the simultaneous use of multiple incompatible systems. Adoption of a data standard can be much more than a point-and-click operation, requiring many decisions about the structure and use of the data not just as it is now, but also what the PI expects it to be in the future. One of the first decisions is the standard itself: it can be difficult to pick a “winner,” as standards may quickly become incompatible with the lab's evolving methods.

It is also uncommon to have institutional support, in the form of grant funding or university staffing allocated to the “low-level” task of revising data formats, or incentives such as promotion criteria that reward best practices in data management. While RSEs are increasingly recognized as valuable contributors to the research enterprise (Carver et al., 2022), most labs still do not have access to an RSE. This places the burden on students and postdocs, who are often enthusiastic to adopt new practices but are constrained by a need to make continual progress in their own careers. Moreover, lab members, including PIs, generally lack advanced training to know how to build automated systems that integrate multiple data streams into a single format with appropriate metadata, provide that data for analysis, and share data following community norms such as the findability, accessibility,

interoperability, and reusability (FAIR) guidelines (Wilkinson et al., 2016). Without support, adopting a standard is often a shared aspiration with little personal buy-in to do the needed work.

Researchers

The main motivation for researchers to adopt standardized data formats is to improve data analysis and shareability. Standardized data formats may support efficient and reproducible data processing and flexible, comprehensive data exploration and analysis. Efficient data analysis can, in turn, provide critical information for optimizing experimental design. Furthermore, standardized formats facilitate data sharing, which can yield new perspectives on datasets and increase their impact.

A main concern is that data standardization requires a significant increase in workload, whether researchers tackle it on their own or in collaboration with an RSE. The increased workload can happen at the experiment and data conversion stages, if data management standardization comes at the expense of experimental flexibility. At the stage of analysis, researchers may need to spend time to learn and adapt to the new standard in order to use the data. Researchers' diverse backgrounds, the availability/support of tools for standardized data, and the maturity of their projects further contribute to tradeoffs between making consistent experimental progress and standardizing experimental outputs. In particular, there are limited training opportunities in scientific computing as a topic in its own right, leaving most researchers without conceptual frameworks and technical knowledge to properly guide these choices. Additionally, researchers who decide to embrace standardization, open data, and reproducible workflows often lack recognition for the added work.

RSEs

RSEs directly support researchers in data management, analysis, sharing, and publication. Adopting standardized formats establishes predictability in the data that the researchers produce. This facilitates communication and makes it easier for RSEs to efficiently provide support in finding, using, and building appropriate systems to interact with the data. RSEs can also take advantage of such predictability to provide sufficient documentation and usable examples of the data for analysis, sharing, and reuse.

A core challenge is developing stable software implementations and workflows that are robust to small variations in experimental data, while still allowing flexibility to be useful to researchers engaged in rapid evolution of diverse experimental designs. Furthermore, choosing a new technology carries an elevated risk of bugs and missing features. Open source tools can be particularly unpredictable, and extensive in-house workarounds may be unsustainable and defeat the original purpose of standardization.

In addition, researchers and RSEs often come from different backgrounds. RSEs may not be familiar with scientific priorities and experimental constraints, and the expectations and timeline of research projects. Thus, diverging expectations and miscommunication between researchers and RSEs can lead to friction and delay in adopting the standards.

Social Scales of Working with the NWB Standard Within a lab

It is often desirable for members of a lab to share and use common technology, including analysis code, data conversion pipelines, and/or acquisition systems. This commonality allows members to jointly address technical problems, and build on

top of known solutions with some degree of prior validation, creating consistency across “generations” of graduate students and postdocs. For example, in our lab, researchers performing head-fixed two-photon calcium imaging share the same acquisition systems and data conversion pipeline, which allows them to get advice from their peers and to contribute their own solutions to common pain points.

A potential pitfall of sharing a common set of technologies may arise when the technology is not well maintained or kept up-to-date, forcing new projects to build on shaky ground. Another pitfall may come from the complexity of supporting a diverse enough set of use cases, and trying to make them all fit into the same technology.

On-boarding is key to encourage this economy of scale and self-regeneration of benefits, especially if a standard is not yet established. For example, rather than introduce NWB to researchers in new analysis notebooks, we tried to work backwards from the analysis pipelines they already used. That is, we refactored researchers’ existing code by replacing only file load operations, and converting from NWB structure to whatever variable names and data types the researchers already used (that often adopted suboptimal data conventions from the original raw file formats). Further experience with NWB might motivate changes to those conventions, but in this approach, initial learning is focused on practical steps whose value is innately recognized by the researcher, rather than on the generic NWB software interface. Naturally, it could be simpler for new lab members (or new projects) to start from a standardized “clean slate,” though our experience is that in practice there is usually still substantial inheritance of older code and procedures, at least in an established lab.

The Fleischmann lab uses lab-wide Git hosting (on GitLab), facilitating internal sharing and collaborative development of code. Combined with regular lab meeting discussion of data management and analysis topics, this culture of open communication and sharing helps disseminate technical progress across all lab members.

Collaboration

Our experience using NWB to send data to collaborators in other labs has been more mixed than for internal adoption. While standardization aims to establish a universal language for data, there can still be friction for recipients who have not already installed and used the necessary software, especially in the absence of good documentation and relevant working examples. We describe two cases with two different labs performing additional analyses on data we collected.

In the first case, we provided our collaborators with raw microscope images as TIFF (tag image file format) stacks and preprocessed calcium activity time series in NWB format. In contrast to our naive initial expectations, it was challenging for our collaborators to learn how to work with the NWB files. With hindsight, we should have included working example code that loaded and displayed data, which they could use as a starting template for their own work. However, there would still have been some friction, as their lab works primarily in Matlab, while we work almost entirely in Python. NWB provides Application Programming Interfaces (API) for both environments, but we would have needed to generate example code from scratch, and the two labs would have maintained two separate code bases. In the end, our collaborators used only the TIFF stacks, though partly in order to also work on novel preprocessing algorithms.

In the second case, our collaborators had previous experience with NWB. However, we were still refining our NWB conversion of that data, and were regularly making code breaking changes. Hence, we chose to create and send python “pickle” files that contained only a subset of the data, organized to simplify usage on their end and make it easier for us to create example code and documentation. As we continued to develop our internal pipelines, this approach hampered code interoperability between our labs. However, it was the more expedient choice to get the collaborators up and running. Since then, we have been working to improve the long-term stability of our NWB conversion pipeline, in order to converge on a collaboration strategy built entirely on NWB standardization. For this project, we and our collaborators have been working off the same NWB dataset, which is currently embargoed on Distributed Archives for Neurophysiology Data Integration (DANDI) Archive (<https://dandiarchive.org/dandiset/000785>). The ease of access to the data repository, accompanied by the availability of a minimal code repository for accessing and preprocessing (<https://gitlab.com/fleischmann-lab/calcium-imaging/projection-difference>), allows us to perform analysis of the same dataset independently and to more quickly share our findings with each other. This increases the reproducibility and interoperability of our and our collaborators’ analyses.

Our research has also become more rigorous and transparent as a result of the standard and the sharing of repositories. And the same reasons have allowed us to more easily discuss data internally and with collaborators while the data are actually accessible.

Public data sharing

Researchers are increasingly asked to publish their data on public archives. Apart from publication and funding requirements and opportunities for collaboration, these public data repositories increase chances of data reuse, e.g., for education, benchmarking new tools, computational modeling, or meta-analysis. Popular repositories include Figshare (RRID:SCR_004328; <https://figshare.com/>), Zenodo (RRID:SCR_004129; <https://zenodo.org/>), open science framework (OSF; RRID:SCR_003238; Foster and Dearnorff, 2017, and GIN G-Node (RRID:SCR_015864; <https://gin.g-node.org/>). These are more general repositories, with limited restrictions on data formats, though there sometimes could be other logistical/funding complications.

The DANDI (RRID:SCR_017571) is the recommended choice for public data sharing of NWB datasets (Halchenko et al., 2022) and is supported by both the BRAIN Initiative (Kaiser, 2022) and the Amazon Web Services (AWS) Public dataset programs. While it is more restrictive compared to other repositories (e.g., DANDI allows only standardized formats (<https://www.dandiarchive.org/handbook/about/policies/>), while Zenodo allows all formats (<https://about.zenodo.org/policies/>), the resulting rigor and consistency from DANDI may better facilitate reproducibility, modeling, meta-analysis, and tool development (Magland et al., 2024; https://github.com/catalystneuro/dandi_llms). We discuss our experience contributing a demonstrative calcium imaging dataset (Daste, 2022) on DANDI in “Considerations for sharing on DANDI.”

Apart from file format restrictions, researchers may need to take into account file size limits. DANDI has fairly generous limits, with 5 TB per file and no limit on dataset size, while some repositories have limits of less than 100 GB per file or dataset (some offer higher limits for a fee or other arrangement).

We discuss in more detail our experience contributing a demonstrative calcium imaging dataset (Daste, 2022) on DANDI in “Considerations for sharing on DANDI.” Generally, NWB

and DANDI have allowed our research to be more transparent. This dataset is also used in a recent publication (Srinivasan et al., 2023) and is available on DANDI and can be shared easily.

NWB community

During the process of developing our NWB data conversion pipeline, we had several opportunities to interact with the NWB development team. Some of these ways were the NWB/DANDI Slack for quick questions, GitHub issues for a technical question or bug, GitHub discussions for entry level questions, remote meetings with the NWB team for more in-depth substantial guidance, and organized events (hackathons, user days, and data re-hack) to meet others from the community and learn about the progress of the ecosystem. In general, our interactions with the NWB community were friendly, helpful, and responsive. For example, our questions on Slack usually received responses within the day. From our observation, this was also true for questions posed by other users.

As described in “Creating NWB extensions allows fitting domain-specific use cases,” we decided to design our own NWB extensions, which was technically challenging. Communication and assistance from the NWB team was very valuable in our design and implementation. Occasionally there were also helpful examples in GitHub issues or discussions on GitHub and Slack.

That said, many of these resources and communication channels are more familiar to computational scientists and software developers. The official documentation sometimes could be overwhelming to navigate (see, e.g., <https://github.com/NeurodataWithoutBorders/pynwb/issues/1482>), increasing a typical user’s need to find and access these discussions scattered around many channels. It could have been helpful to have a centralized, searchable resource that aggregated and archived these different issues and discussions across different forums, as a complement to the official documentation.

Neuroscience community

The advent of the open science movement, in parallel with standards development, has increased access to software tools and data that until recently was generally limited to high resource institutions. For example, the Allen Institute for Brain Science released a software development kit (SDK, RRID:SCR_018183) that simplifies retrieval of and interaction with extensive collections of NWB standardized data recorded with cutting edge electrophysiology and imaging tools. Such initiatives greatly expand opportunities to reuse data in education (Van Viegen et al., 2021; <https://training.incf.org/collection/neurodata-without-borders-nwb>), basic research (Deitch et al., 2021), and benchmarking of new computational models (Schneider et al., 2023).

However, given differences in cultures, priorities, resources, and incentives across different labs and institutions, adoption of NWB, and of open science practices more generally, remains challenging. Institutional policies like the recently updated NIH Data Management Policy (Martone and Nakamura, 2022; Contaxis et al., 2024; <https://oir.nih.gov/sourcebook/intramural-program-oversight/intramural-data-sharing/2023-nih-data-management-sharing-policy>) add new expectations for researchers, but without creating meaningful recognition and training to support and encourage changes in their practice. Individual institutions also have historically provided minimal support for adoption of data management best practices. We advocate for better funding for standardization as an essential practice in science in general, and particularly for NWB adoption. Some of this support could include partnerships with public resources such as `nwb4edu` (<https://nwb4edu.github.io>).

Building Our NWB-Based Data Conversion Pipeline: Experiences, Challenges, and Lessons Learned

How to organize data into a standard format

There have been many efforts at standardization of neuroscience data. NWB started as a pilot project to standardize neurophysiology data (Teeters et al., 2015), which then matured into NWB:N version 2.0 (NWB:N 2.0; Rübel et al., 2019).

However, NWB is not really a file format. The substantive outcome of the NWB development effort was an “ontology” that encapsulates the logical structure of neuroscience data at a high level, and schemas to translate these conceptual objects into precise computational objects. Unlike saving an image in joint photographic experts group (JPEG) format or a document in portable document format (PDF), to use NWB researchers must make a number of choices specific to their data, with both technical and conceptual implications. This makes data sharing a technique to be mastered by researchers, similar to scientific writing and scientific communication. These are an investment in time, require practice and engagement, but will pay off in the end for the scientific endeavor.

Figure 2 illustrates questions faced by researchers who may record multi-modal data scattered across different files and formats. The resulting data need to be organized, unified, and aligned in order to support analysis and collaboration. There can be different strategies to standardize these data, for example from a data lineage standpoint (the choice of the NWB team; Fig. 2, middle) or from a categorical standpoint (Fig. 2, right).

Our files mostly follow the default NWB internal structure for optical physiology, though we made our own extension to handle odor data (see “Odor stimulus metadata”) and argue researchers could benefit from alternative structures, perhaps using aliases or tags, that allow them to interact with their data files following categorical or other organization (see “Suggestion for better data access: tags and aliases”).

When to create and use the standardized format

Few experimental acquisition systems produce NWB files natively, so use of the standard requires researchers to choose a process and time to convert to NWB from some mixture of other data files. One strategy is to convert at the end of a project, perhaps to upload to a repository for sharing. This choice minimizes disruption to existing research workflows and preserves flexibility for intermediate analyses. However, this strategy may reduce reproducibility, as analysis is done on different files than are eventually shared. Also, shared code needs to be refactored at time of publication to account for these file differences.

Alternatively, conversion could occur prior to internal use. In the pipeline illustrated in Figure 3, conversion happens between preprocessing (using Suite2p and DeepLabCut) and analysis. Regardless of standardization, researchers typically reformat data before analysis, for example to compile information from multiple raw files into a convenient single data array or table. The key cost of standardization is to place *restrictions* on allowable output formats, in order to reap the benefit of harmonizing a particular dataset with common practice in the field. If data are converted early, then archival repositories can be used also as backups, possibly including data version control. Moreover, shared code does not need substantial rewriting at time of publication. However, if there is not already a robust conversion pipeline in place, this strategy introduces additional effort prior to progress of the scientific aims.

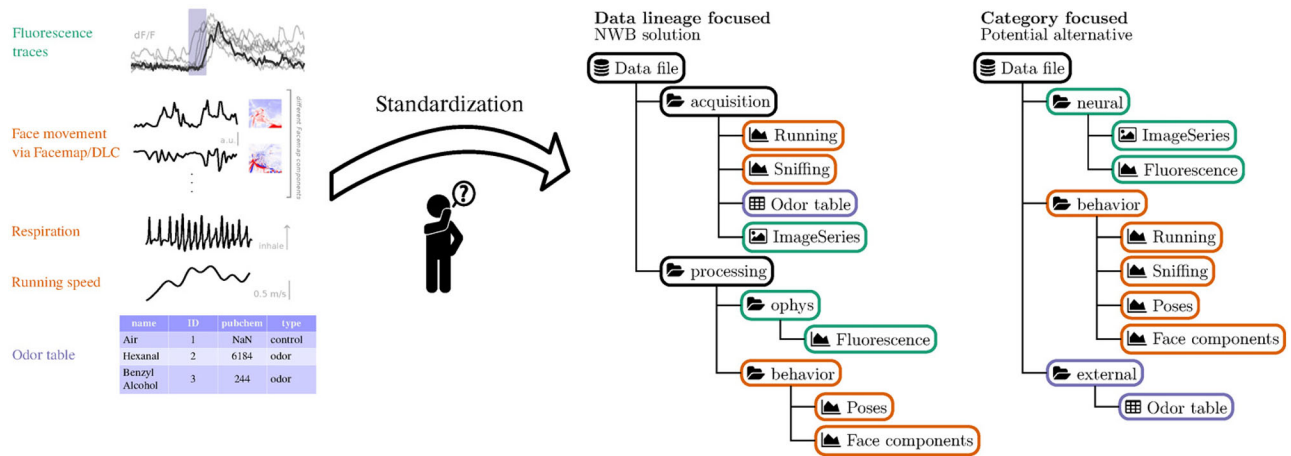


Figure 2. The issue of data standardization. Systems neuroscience data tend to be multi-modal, e.g., time series recorded from standalone sensors and extracted from neural imaging and behavioral videos, plus tables of stimulus or other events (left column). These data are usually scattered across different files in various formats. Researchers wanting a unified standard for ease of analysis and data sharing must choose between at least two possible organizational strategies: prioritizing the data lineage (chosen by NWB format; middle column) or prioritizing conceptual categories of data sources (right column). Color scheme: green for neural activity, orange for animal behavior, and purple for external variables (e.g., stimulus).

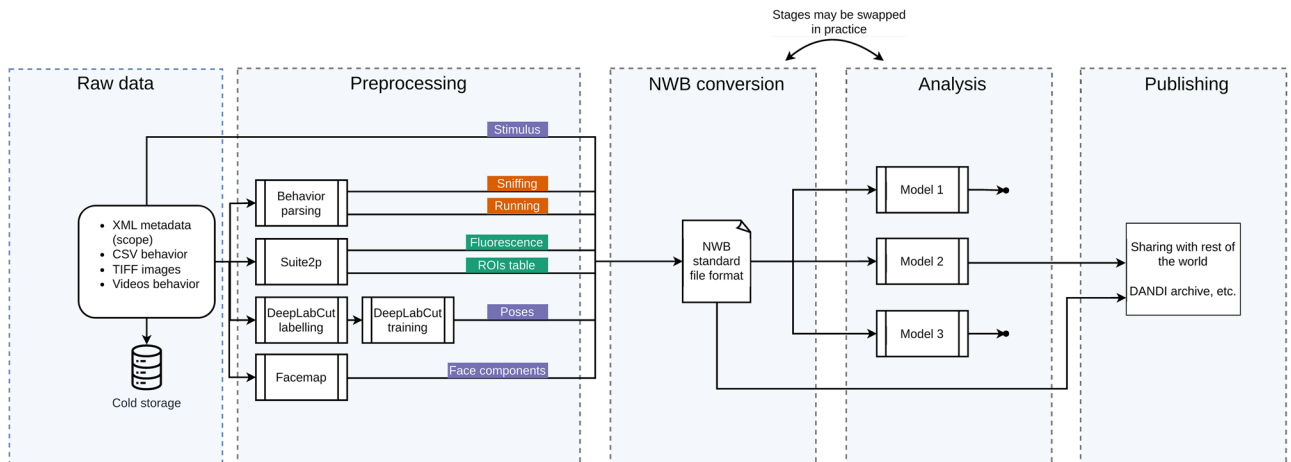


Figure 3. Our data pipeline. There are five primary stages in our data pipeline. *Raw data* acquired during experiments are archived in cold storage, and also fed to a *preprocessing* stage to be transformed into more directly usable information (e.g., fluorescence time series after cell segmentation). This stage uses a range of processing packages that produce multiple files, which are then combined during *NWB conversion* into a standardized format. Scientific *analysis* ideally is performed on the standardized data, but in practice may instead use individual files produced during preprocessing, in which case conversion and analysis stages are swapped. Standardized data are *published*, e.g., by uploading to a publicly accessible archive, in parallel with traditional journal publication.

Overall, our feeling is that the stages where NWB is most useful are integrating relatively stable preprocessed data, and archiving finalized data and analysis for publication.

Our experience with metadata capture

Metadata can be defined as “data about data”, for example, information about animal subjects (e.g., weight, sex, genetic line, age, and whether naive or trained), recording sessions (e.g., date, task type, experimenter name, manufacturer, and model of hardware), stimuli (e.g., chemical names, concentrations, and frequency of audio tones), supplemental text descriptions, and/or parameters used in data processing. Generally, metadata can aid in quality control, communicate contextual information to future users, and support cross-analyses of multiple data sets. Its use can extend beyond the lifetime of a project, including archiving, sharing, and reuse.

Quality of metadata capture

A benefit of moving data to NWB is that it encourages systematic handling of metadata. To convert into NWB format, some types

of metadata are required by the standard, while some are encouraged. Before moving to NWB, our metadata were scattered in several places. Now, all the relevant metadata are included in the NWB file, allowing consistent and easy access. This may help answer questions such as *What was the sex of animal X?*, *What imaging frame rate was used in experiment Y?*, or, when using our neurodata extension described in “Odor stimulus metadata,” *Which odor stimulus was used in trial Z?*, without having to go back to the raw data or to the experiment notebook.

Challenges to metadata capture

An obvious challenge to incorporating correct metadata in standardized files is that experimentalists do not always record metadata effectively. They may rapidly iterate an experimental design while piloting and record only “core” data for preliminary analyses, with a fuzzy boundary between these initial pilots and subsequent “real” data collection. Moreover, metadata often take unusual effort to document. Acquisition software may not support metadata capture at all. For example, mouse dates of birth or ages are

often not included in data files produced during an experiment, yet at least one of these values is needed to create NWB files that meet minimal upload requirements on DANDI (see “Considerations for sharing on DANDI”). Sometimes tools set incorrect metadata as a default; for example, we found the NWB conversion function within Suite2p defaulted to setting area of recording to be “V1” (<https://github.com/MouseLand/suite2p/blob/118901ac15c6881502c65e011a46fbc16e7a52d/suite2p/io/nwb.py#L346C27-L346C27>). Also, there is not always a clear purpose to recording metadata that go beyond the key variables in the original study design. Under the time pressure of the experiment, researchers may be induced either to use non-informative defaults or to enter random metadata to get underway.

This issue is exacerbated by a lack of accepted community standards of how to document for some types of metadata. For instance, in olfaction research, there is not yet consensus on how to document odor stimuli (though see [Castro et al., 2022](#) and “Odor stimulus metadata”).

More generally, metadata capture is needed not only during acquisition but also during preprocessing, analysis, and file conversion stages. Here again a lack of community consensus both motivates the need for detailed metadata capture and illustrates challenges in its implementation. For example, fluorescence is typically normalized, but there is wide variation in how that normalization is performed. Methods used to obtain so-called dF/F_0 can differ in parameter choices or the algorithm itself (e.g., global z-scoring, quantile normalization, or running normalization with additional filtering). Some methods may attempt to compute dF/noise instead (e.g., Inscopix CNMFE; [Boivin et al., 2021](#)). Often these choices are not apparent in publications and require careful inspection of code, if provided. Such nuances may affect how the data are used, the assumptions of tools that analyze such data, and efforts to replicate analyses.

Working with acquisition devices and software

In our labs, RSEs assist data conversion in part by working with researchers, equipment vendors, and others to determine what metadata are needed and how best to capture it.

Some commercial vendors put metadata in dedicated files, e.g., Bruker Microscope extensible markup language (XML) or ENV files (RRID:SCR_023608) while others integrate metadata into the same files as core data, e.g., Inscopix Miniscope (RRID:SCR_017407). However, some proprietary vendor files are poorly documented (and questions stayed unresolved after contacting support), such that we had to reverse-engineer files and make educated guesses as to the information in them. For example, some things we had to independently infer from Bruker XML files were where frame rates are recorded, what physical units different fields have, and what the reference frame coordinates are. Our inferences relied on field names, and were incomplete and possibly in error. More importantly, certain metadata can change the algorithm used to parse a file; for example, a flag indicating whether an experiment has multi-plane imaging affects the correct way to extract timestamps from the XML file. NeuroConv ([Baker et al., 2023](#)), the conversion tool from NWB developers (see “Off the shelf NWB conversion”), initially did not integrate Bruker metadata (see <https://github.com/catalystneuro/neuroconv/pull/390>), but we note support has been added during revision of this article.

Open source tools typically fill a space between commercial vendors and in-lab custom development. Some of these tools lack an ability to input metadata. An example is ArControl (RRID:SCR_021605; [Chen and Li, 2017](#)), which is an experiment control platform used with general purpose microcontrollers to present

stimuli and record behaviors. There is a project to convert its output into NWB format (<https://github.com/chenxinfeng4/ArControl-convert2-nwb>), but (as of this writing) still requiring post hoc metadata injection (see <https://github.com/chenxinfeng4/ArControl/issues/2#issuecomment-1416018374>).

We also develop custom scripts ourselves that generate comma-separated values (CSV)-like files on microcontrollers. This approach would ideally include informative headers, for example to give each data column an informative name, a plain text description, physical units, a data type, and possibly other metadata. We find that this step introduces friction and an increased chance of errors, especially as experimental designs change and researchers or software engineers need to keep code updated and documented. For now, metadata are often documented after acquisition. In an alternative approach, we implemented custom widgets in Jupyter notebooks used for data acquisition that allow experimenters to write in odor names. The notebook then saves the names in a YAML file along with separate core data files, and all files are integrated into an NWB file in a later conversion process. The widget was tedious to develop, but substantially improved the quality of metadata capture for odors at the time of the experiment.

Where should raw data and supplemental information be stored?

Researchers may want to store raw data in their NWB dataset. In our case, the raw data may contain calcium imaging TIFF stacks or behavior video recordings, both of which tend to be large. For example, a typical calcium imaging session in our lab generates a video of size around 40 GB, with associated behavioral videos around 3 GB. There has long been a question of what to do with videos (<https://github.com/NeurodataWithoutBorders/pynwb/issues/1647>, <https://github.com/NeurodataWithoutBorders/nwb-overview/issues/78>), contrasted with the much smaller data derived from them in preprocessing. Should raw videos be included in NWB files? If yes, how? If not, how should videos be handled when publishing to a repository?

The NWB team discourages writing videos in lossy compressed formats within NWB files (see the question *Why is it discouraged to write videos from lossy formats (mpg, mp4) to internal NWB datasets?* in the Frequently Asked Questions: <https://nwb-overview.readthedocs.io/en/latest/faq.html>). The main reason is an inability to decode the video without first copying the data to a standard file type (e.g., MP4) on the user’s computer; moreover, if the appropriate codec is not available, even a copied video would be unreadable. The preferred solution is to include videos in NWB files as an `ImageSeries` that has an external file reference (a relative path to, say, an MP4 file), [Rodgers \(2022\)](#) as an example. This solution also allows adding videos in published datasets on DANDI (see <https://www.dandiarchive.org/2022/03/03/external-links-organize.html>).

Often researchers may want to share explanatory content such as videos of experimental setups or down-sampled videos of calcium imaging registration aligned to behavior recording. Only a subset of recording sessions may have such associated content. A solution could be similar to storing raw data as external file references as described above, clearly labeled for demonstrative purposes to avoid confusion.

How should different data types be stored?

In NWB, neurodata types refer to different modalities of data and metadata, for example `DfOverF`, `PupilTracking`, or `SpikeEventSeries`. Each type has specific rules to fit

different use cases. If data belong to a standard neurodata type, there are usually clear examples and guidelines about where and how to store it in an NWB file. When it does not, non-trivial choices may be required, and variation across labs, each implementing their own conventions, may impact general reusability.

For each data source to be integrated into an NWB file, users must answer a number of questions about the data representation. Can the data be fit in a standard neurodata type? What metadata should be associated with it? Would an extension (see “Creating NWB extensions allows fitting domain-specific use cases”) add a more appropriate datatype? Does such extension exist? If not, is it worth the effort to develop one?

Additional questions concern where to place the data in the NWB hierarchy. The organization of the NWB standard is structured with data workflow stages at the top of the hierarchy: `acquisition` (usually raw), `processing`, and `analysis` (Fig. 2). While in theory preserving some element of data lineage, the semantics in practice are not always clear or observed, and can cause confusion when creating and using NWB files.

For example, should raw behavior time series acquired from microcontrollers be in `acquisition`, a module called `behavior in acquisition`, or in the same `behavior` module in `processing` that is often used to store post-experiment processing such as DeepLabCut pose estimation? From a data lineage point of view, it should be stored in `acquisition`. But from an analysis point of view, doing so spreads multiple fragments of behavior-related data across multiple hierarchical levels and modules.

As a detailed example of how small experimental variations can lead to non-trivial design choices in NWB files, we describe an experiment in the Fleischmann lab involving two color imaging of red (tdTomato) labeled cells in parallel with green (GCaMP) functional imaging. After using Suite2p for cell segmentation, the researcher classified each cell as expressing or not expressing the red fluorophore, producing a table of regions of interest (ROI) (cell) indices, Boolean values for whether a cell is red, and auxiliary data about the classification (average pixel intensity and a quality metric).

There are three levels of detail one might choose to keep in an NWB file (in addition to the functional imaging contained in a standard datatype): as the full table, as only the Boolean array, or as an array of indices of red cells. The last choice is the most compact, but does not preserve the auxiliary information that might be useful for quality control and reproducibility. Similarly, parameters of the classifier itself (e.g., intensity thresholds) should likely be saved as well. The choice of what information to retain both suggests and is constrained by what datatypes are available, or whether we would need to develop an extension (see “Creating NWB extensions allows fitting domain-specific use cases”). And a further decision is where to save the data in the file hierarchy (Fig. 2): as preprocessed data or an analysis result?

There is obvious value to saving the classification in the same place that stored the segmentation table from Suite2p output, essentially by adding more columns to that table. However, since the classification is not available at the time of Suite2p segmentation, and updating existing objects in the Suite2p NWB file was problematic (see “Editing and merging of NWB files”), we resorted to placing the classification table in another module called `cell_tag`. Given that the table came from Suite2p, whose outputs are in `processing`, we were unsure whether `cell_tag` should be considered `processing` or `analysis` in terms of lineage. However, in terms of usage, the tagging is not a useful result by itself, but is combined with

the calcium dependent activity. Hence, we decided to consider the table as processed data needed for analysis, and save it in `processing`.

As a second example of non-trivial design choices, the Datta lab records breathing signals with a temperature sensor implanted in the nose. An Arduino captures the signal, which is written into a CSV file in real-time. We developed a processing pipeline to clean and parse the breathing signal into individual breaths, and store the resulting data in an NWB file. There were a number of challenges along the way that highlight some limitations of the current NWB implementation.

Scipy’s `signal.find_peaks` function was the core of the breath processing pipeline; good results relied on choosing correct parameters to find true breaths while ignoring noise in the data. Sometimes we would update the defaults of those parameters based on new analyses, and it would have been helpful to traverse old files programmatically and update them. As it was, many key parameters ended up stored in the `description` property of the relevant `TimeSeries`, which may not be an obvious location to those looking at the data for the first time.

Also, there were a number of options for how to store information about each breath, which were difficult to differentiate ahead of time. It would have been ideal to choose based on, e.g., efficiency of storage or common practice, but in the end our decision was purely pragmatic. We first considered a tabular format like the `TimeIntervals` table, but adding data to the `TimeIntervals` table proved to be cumbersome (see <https://github.com/NeurodataWithoutBorders/helpdesk/discussions/30>). Then we considered an `IntervalSeries`, which would allow labeling onsets and offsets of inhales and exhales and convey the “interval” aspect of the data, but this did not lend itself to storing scalar descriptors for each breath, since the datatype stores only timestamps and not values. Finally, we settled on a simple solution: a `BehavioralTimeSeries`, containing many `TimeSeries` of length `number_of_breaths`. For example, inhale onset times, amplitudes, and peak flow rates each got their own `TimeSeries`. Inhales and exhales were paired in the preprocessing stage, and the `TimeSeries` that describe the inhales and exhales have the same length, thus implicitly pairing each inhale/exhale pair. We chose to save the `BehavioralTimeSeries` interface, called “breaths”, in the `processing` section of the NWB file.

Should one standardize data from intermediate analysis stages?

Research analysis pipelines typically have multiple stages, such as preprocessing, statistical modeling, simulation, or any computation whose inputs are the outputs of a previous stage. Those stages may also branch out to test a family of models, or vary analysis parameters. The NWB standard is limited in its handling of analysis parameters, for example as tables of metadata. Should intermediate results be appended to a single NWB file containing the entire history of analysis, each as their own “data source”? Should each analysis be stored in its own NWB file? Should all but the final published analysis be discarded?

Iterative analyses quickly become unwieldy without automated tracking of workflows (e.g., Renku; <https://renku.readthedocs.io/>) and/or data versions (e.g., DataLad; RRID:SCR_003931; Halchenko et al., 2021). NWB was not designed to compactly represent collections of results such as arise from parameter sweeping in an analysis. Similarly, NWB does not natively support tracking the partitioning of data, such as into “training” and “testing” subsets for cross validation (though there are possible *ad hoc* solutions under

the current standard, and new packages in late development—personal communication, NWB Developer Team—to support such functionality).

Editing and merging of NWB files

Early in our transition to NWB adoption, we needed to combine an NWB file exported from Suite2p with another NWB file produced by our own data pipeline. This turned out to be surprisingly difficult. Indeed, according to the PyNWB (RRID: SCR_017452) documentation, adding to files is supported, but removal and modifying of existing data is not allowed. We therefore tried two approaches to do this. In the first, we read the existing NWB file produced by Suite2p, added the missing data, and exported to a new NWB file. In the second, we looped over containers, i.e., HDF5 groups, in the existing NWB file, and copied each of them into a new NWB file, together with the new data.

The first approach produced an NWB file that, due to a bug in the underlying packages (which has since been fixed), caused crashes while reading with PyNWB (see <https://github.com/NeurodataWithoutBorders/pynwb/issues/1301>). Because of a different bug, the second approach failed to create a new NWB file with the new containers (see <https://github.com/NeurodataWithoutBorders/pynwb/issues/1297>). These unexpected errors in what seemed like intuitive workflows were frustrating both for the delay in switching over to NWB, and the additional effort needed to diagnose the bugs and find workarounds.

There are still limitations in copying containers from one NWB file to another. But compared to when we started working on this project, it is now more straightforward to copy datasets, i.e., a data array and its timestamps, from one file to another, and to read an existing NWB file, modify it, and export the modified file to a new file. It is also possible to append data to a file, in the sense of creating new datasets. However, to our knowledge, the only way to update metadata in an NWB file is to read the content of the existing file, use the NWB API to create an object with the correct metadata, and then export to a new file. In general, we have found that editing and merging NWB files can be a large source of confusion for users, and having a good tutorial or documentation

as proposed in <https://github.com/NeurodataWithoutBorders/pynwb/issues/1773> would be extremely useful.

We want to acknowledge that the bugs we encountered are not due to carelessness from the NWB developers, who have been very responsive to our feedback, but possibly underlie non-trivial technical issues that any HDF5-based open source software would have to address.

Pain points in the conversion workflow

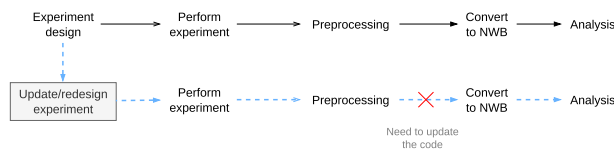
We encountered several pain points in our data conversion pipeline. One of the main pain points happens with branching experimental designs (Fig. 4a). Each time a design is updated, NWB conversion code may break and need to be updated. This is an issue especially early in project development, when many experimental details are undecided, but can continue far into a project's lifetime as researchers adjust their approach based on prior results.

Another pain point may arise when metadata are missing at conversion time (Fig. 4b). Researchers may be tempted to input nonsense values that need to be updated later, or the conversion may be blocked until the missing metadata are captured.

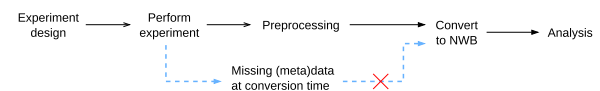
Sometimes, data in NWB files may need to be updated, e.g., to correct a previous entry, or to add data that becomes available later, such as histology (Fig. 4c). In this case, the pain point happens when the data conversion pipeline has to be run again on multiple already existing files. As discussed more in “Potential surprises with data validation,” a related issue can arise when sharing data in an archive such as DANDI (Halchenko et al., 2022). Validation to DANDI is stricter than requirements to build a file with the python API (PyNWB), requiring conversion code updates even after conversion was locally “successful” (Fig. 4d).

Timeliness of code contribution acceptance

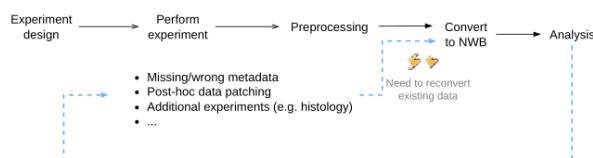
We discovered that Suite2p was dropping data from a second microscope channel in its NWB file output. The issue was that the NWB export function had been developed for only one microscope channel. Figure 5 shows the timeline of the issue until a fix was released. While fixing the issue internally took around



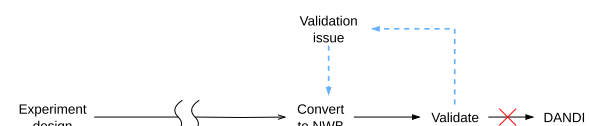
(a) Branching experiment



(b) Missing metadata at conversion time



(c) Data update/patching



(d) Validation issue before publishing

Figure 4. Pain point scenarios in the conversion workflow. This figure describes different scenarios adding burden to the research workflow. The red crosses represent a situation that breaks the existing workflow. The electric current symbol represents the location of a pain point. **a**, Branching from the main experiment, i.e., a redesign or update of the experiment, may break the current conversion code to NWB. **b**, If some metadata are missing at conversion time, it may force the researcher to come back to the experiment, to the original data, or to the conversion code. **c**, A scenario where existing NWB files need to be updated, e.g., when data from additional experiments like histology experiments become available, or when the NWB files have missing/wrong metadata, or if the NWB file has been found to have some data issues which need to be updated. **d**, A validation issue before publishing the data to DANDI which may force the researcher to update their conversion code to NWB and reprocess their NWB files.

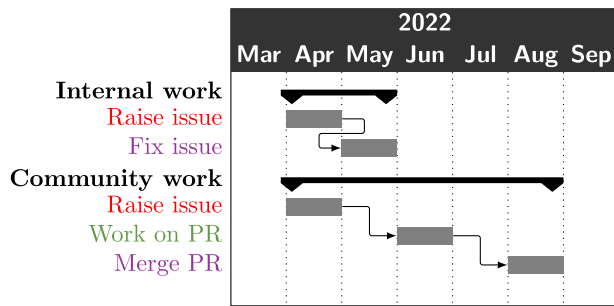


Figure 5. Example of a broader community issue resolution timeline. This figure illustrates the time taken to fix a Suite2p-related issue internally (i.e., two months), compared to the time it took to fix the issue for the broader community (i.e., five months).

two months, it took around five months (including time for us to complete a GitHub “pull request”) for the solution to be available to the Suite2p community. This is a long turnaround for what we considered to be a critical error, impacting all multicolor imaging analysis. We stress that we appreciate the Suite2p team’s review and acceptance of our code contribution. However, this experience illustrates a general problem for research software development in the open source community; researchers maintaining software may not have the bandwidth to address every issue or feature request in as timely a fashion as desired.

Off the shelf NWB conversion

Some friction during adoption of NWB can arise from the level of technical skill needed to be able to convert one’s data. When we started the process of adopting NWB, the options available were either to learn how to write our own data conversion pipeline, or hire a consultant to do the technical work. In the few years since, the NWB ecosystem has rapidly evolved. More recently introduced tools miss some areas of need (e.g., currently unsupported proprietary formats like Inscopix, or Suite2p output with multiple channels), but they solve many popular use cases.

NeuroConv (Baker et al., 2023) is a rapidly advancing Python package from core NWB developers to make it easier to convert from a variety of common neuroscience data formats. It is a flexible low-code solution for use in one-off conversion or as part of a lab pipeline. One benefit of NeuroConv is that it includes utilities to get metadata from proprietary formats with minimal effort. Additionally, it can combine files from multiple data sources with functionality to align timestamps, and contains utilities for file path inference to aid batch-conversion based on user-defined data organization. Coupled with the development of the NWB Graphical User Interface for Data Entry (NWB GUIDE; RRID:SCR_025467; <https://nwb-guide.readthedocs.io>; Flynn et al., 2024), which uses NeuroConv as a backend, NWB is considerably more accessible to newcomers than it was at the time we began our adoption.

These recent changes highlight a risk to early adopters of any standard that one may build features from scratch that quickly become obsolete after further developments from the community. If we started this project today, we would leverage these community projects, developing less custom code and using existing features from more widely tested projects used by the entire NWB community.

An indirect benefit of using NWB is improved data awareness

As a standard, NWB encourages good data practice. For example, each data array that is written in a file needs to have a timestamps vector attached to it. And ideally all the timestamps of the same

NWB file would be on a common axis, which can be quite challenging for experiments with multi-modal recording and has been discussed recently in Rodgers (2022). This includes the acquisition timezone, meaning an NWB file can easily be analyzed in different parts of the world without risking timestamps collision.

In our case, standardization encouraged better timestamping with custom instruments and sensors like Arduino and Teensy boards. For example, before we developed our own data pipeline, one lab researcher manually specified inter-trial intervals in their analysis code, as it was cumbersome to extract the (nearly constant) intervals from the recording system. Now they have access to the actual recorded timestamps for the inter-trial intervals and can catch and correct any system errors. Also, using NWB encouraged us to align timestamps across all data sources, simplifying downstream analysis work.

A general byproduct of moving our lab to NWB is increased awareness regarding data management itself. Lab members have become more familiar with general principles such as FAIR (Wilkinson et al., 2016) and emerging best practices. Although still harboring some skepticism of the direct usefulness to their research, lab members have become more welcoming to incorporating NWB into their workflows, and are supportive of the broader benefits, such as for data sharing.

Creating NWB Extensions Allows Fitting Domain-Specific Use Cases

An emerging standard with as broad a domain as NWB will naturally struggle to cover some applications, especially in less common experimental settings. Making the standard extensible creates a way for individual users or research groups to add functionality beyond what is created by the core developers. The NWB standard thus includes “neurodata extensions” to incorporate new data types. Extensions may be used individually, shared with the community, or, if the extension addresses a fundamental gap in NWB coverage, submitted for review to be added to the standard NWB data types. We have had some success using and creating NWB extensions to fit our specific research needs, though challenges and questions remain. motivation and definition of extensions

Existing neurodata extensions

Before deciding to create an extension, researchers should check the Neurodata Extensions Catalog (NDX Catalog; RRID:SCR_021340), a community led effort to create a central repository of contributions that, by design, arise from widely distributed effort. The NDX catalog includes extensions that support diverse types of data such as transistor–transistor logic pulses (<https://github.com/rly/ndx-events>), and popular acquisition systems such as miniscopes (<https://github.com/catalystneuro/ndx-miniscope>). However, not all Neurodata Extensions are listed on the NDX Catalog, since anyone can create and post an extension on lab websites, GitHub, or other sites.

Lab-specific metadata

One use case of NWB extensions is to record lab-specific metadata with greater flexibility than is supported in base NWB. We created `ndx-fleischmann-labmetadata` (Pham, 2023a) to store additional detail on recorded brain areas, and descriptions of the experiment and animals. Within our general type of experiment, we use many variations (Box 1), such as 1-photon or 2-photon calcium imaging, single or multicolor imaging, head-fixed or freely moving animals, and passively

presented or task-driven stimulation. NWB standard is missing fields to describe some of the complexity in these experiments; for example, we use multicolor imaging to retrograde label projections from the imaging site to distant brain regions, and there is no field to indicate this second (projection) area. Storing such additional experimental description as text in the top-level description field would be harder for quality control at time of entry, and less efficient to parse for queries at analysis time. With our extension, a subset of information ends up being repeated with standard locations in the NWB file; for example, imaging site is also stored under `ophys`, as suggested in the NWB documentation. However, we chose to centralize our metadata in one place to make querying, analysis, and aggregation of multiple data files easier.

Odor stimulus metadata

Another use case for extensions is to describe stimuli that do not fit within base NWB types. Our calcium imaging experiments use primarily odor stimuli, and some non-chemical stimuli such as sound. We are not aware of an extension to adequately describe these stimuli, and hence a year ago developed `ndx-odor-metadata` (Pham, 2023b). We characterize odor stimulus with standardized information automatically obtained from PubChem (Kim et al., 2023) using a PubChem Compound Identifier (chemical International Union of Pure and Applied Chemistry names, molecular formulas, and weights); dilution details such as concentration and solvent; metadata that are useful for analysis such as stimulus category (e.g., control or conditioned stimulus) and common chemical names; and identifiers to cross-reference with associated time series. The extension also allows non-odor stimuli to be described in plain text.

A major challenge with such extension development, although not an issue specific to NWB, is that there may not be community consensus or documentation to be used as starting points for extension design. For odor stimuli, it was not obvious what type and level of description would be necessary for both in-lab analysis and general reproducibility. Fleischmann lab RSEs used existing spreadsheets as starting examples, and learned only later that outside collaborators had independently created a package, `pyrfume` (Castro et al., 2022), for documentation of odorants. Future work could better harmonize these two efforts at stimulus metadata capture. More generally, the technical development of metadata capture can grow only in concert with the research community's understanding of what the standards for metadata ought to be.

Documentation for extension development

For most of the labs, we expect that extension development will be out of reach unless the lab has access to personnel with strong coding experience. A general challenge for us was that the available documentation could be confusing, and information was scattered across multiple sources, including documentation pages for PyNWB (<https://pynwb.readthedocs.io/>), hierarchical data modeling framework (HDMF; Tritt et al., 2019; RRID: SCR_021303; <https://hdmf.readthedocs.io/>), NWB Overview (<https://nwb-overview.readthedocs.io/>), and NWB Schema (<https://nwb-schema.readthedocs.io/>), and also in GitHub issues or examples on Slack. It would have been helpful in particular to have a larger set of use cases, examples, and/or tutorials. We stress that the NWB development team was highly responsive through GitHub, Slack, and emails, and their help was very valuable for our development work. In the future, we hope such support could be complemented by more comprehensive documentation.

Social challenges in extension development

One lesson learned from our experience is that creating the extension is only a technical part of a solution. Sustained engagement with researchers to choose, document, and record key information is the more fundamental requirement, especially if metadata standards motivating the extension are unsettled.

As a lab, we continue to refine what metadata we should track and how we should capture it. Some changes arise from variation in experiments conducted by different lab members. Some changes reflect interest in adding further types of information, such as water restriction details for experiments with behavioral training, as inspired by an International Brain Lab extension (<https://github.com/catalystneuro/ndx-ibl>). An extension may lower the technical barrier to metadata capture, but only if the extension is aligned with researchers' goals and practices, including changes over time.

A closely related challenge is that many metadata records must be captured post hoc instead of automatically during acquisition or preprocessing. Some acquisition systems lack features to enter metadata in machine-readable formats (necessary for software to correctly place that information in NWB files) during the experiments. Even where real-time capture is possible, the systems may be cumbersome to use, leading researchers to avoid comprehensive entry and checking of metadata. We usually need to work with researchers to collect metadata records in machine-readable formats after experiments and preprocessing are completed, leading to increased work and greater risk of errors and missing information.

We also have felt a tension between building minimal extensions that serve immediate needs versus investing in a longer development project that may have greater generalizability. For example, our odor stimulus extension provides for single odorant but not mixed odor stimuli. Though we generally do not use multi-component odors, they are used by some of our close collaborators (Wilson et al., 2017). We also designed our extension to build on PubChem standardization, which presents difficulties when studying custom-made or undocumented natural odors (Li et al., 2022). These limitations in our current implementation may become impediments as neuroscience tends toward more natural and ethologically relevant behaviors (Krakauer et al., 2017). However, surmounting these challenges will require substantial engagement from a broad section of the olfaction research community, before any technical contributions such as extensions can have a substantial impact.

Framework extensions

An extension is built on top of another NWB object. This object can be one of the four minimally structured objects (Groups, Attributes, Links, and Datasets) of the base NWB specification (<https://schema-language.readthedocs.io/>), but it is often better for an extension to build on a previously developed high-level data type that already captures much of the structure of the information being added. In addition to making it easier to develop the extension without starting from scratch, such inheritance can promote greater consistency by keeping almost all data organization the same as a "common" data type, except for the particular items added by the new extension. For example, a new fluorescence imaging data type might add beam path parameters to an existing fluorescence imaging type to provide for a scope that uses non-uniform laser scanning but otherwise collects standard data.

In cases where NWB is missing a more basic category of data, there is motivation to develop extensions intended to be used

specifically as building blocks for other extensions. We refer to these types of building blocks as “framework extensions.” In addition to facilitating development and serving as illustrative examples, framework extensions could add technical precision to discussions if a research community is working to converge to a consensus data standard.

For example, DeepLabCut and Facemap output time series of spatial locations of points on an animal’s body. While these outputs can be stored generically as simply behavior, they are both instances of a more specific concept of “pose,” and can be stored using the `ndx-pose` extension (<https://github.com/rly/ndx-pose>); the DLC developers offer the `DLC2NWB`—<https://github.com/DeepLabCut/DLC2NWB>—utility to ease conversion using this extension, but we are not aware of an analogous tool for Facemap).

An example framework extension that could have broad utility would store results from principal component analysis (PCA); one of the authors, T.P., participated in discussing this idea at a 2023 NWB Hackathon, but it is not yet implemented as far as we know). PCA is used widely as a simple data dimensionality reduction technique. There are several variants of PCA, such as joint PCA (jPCA) used to find low-dimensional structure in the activity of large neural ensembles (Churchland et al., 2012). Moreover, many analysis applications, including Facemap and MoSeq (Wiltschko et al., 2015, 2020; Sherry et al., 2023), use PCA as a preprocessing step. A general PCA extension could serve as a useful framework to incorporate these different uses within a consistent NWB format. The framework extension would define component eigenvalues, eigenvectors, and projections of the original time series.

As another example, BEADL (RRID:SCR_025464; <https://beadl.org/>) and ArControl (Chen and Li, 2017) model behaviors in a finite state machine framework. The extension

`ndx-structured-behavior` (<https://github.com/rly/ndx-structured-behavior>) is available for BEADL outputs, and it is possible to adapt the extension to handle ArControl output (see <https://github.com/chexinfeng4/ArControl-convert2-nwb>). However, as finite state machines are an important class of models for analysis, there could be value in establishing a more general framework extension, for example called `ndx-finite-state`, from which extensions for these specific analysis packages would inherit.

Wishlist for NWB extensions

Development, cataloging, and updating extensions could be more streamlined.

First, researchers may develop software using different repository hosting (e.g., GitLab instead of GitHub). It could be more inclusive for the `ndx-template` extension template (<https://github.com/nwb-extensions/ndx-template>) to not explicitly assume GitHub as the code repository. The template might also take into account both the Python Package Index (PyPI) and Anaconda as potential package repositories.

Second, currently, to be added to the NDX Catalog, new extensions are submitted via Pull Requests for review on GitHub. Some seem to be approved instantly while others are either stale (e.g., `ndx-pose`), or took around 2 months to be approved (Fig. 6). While the timeline for open source development is often highly variable, researchers and RSEs have to balance many priorities, and usually cannot dedicate much time to the approval process.

To simplify the review process, a bot could check critical requirements before asking for intervention from an NWB maintainer (taking some inspiration from the Conda-Forge community). For example, the bot could check if the package is already published on PyPI, if all the metadata fields in the

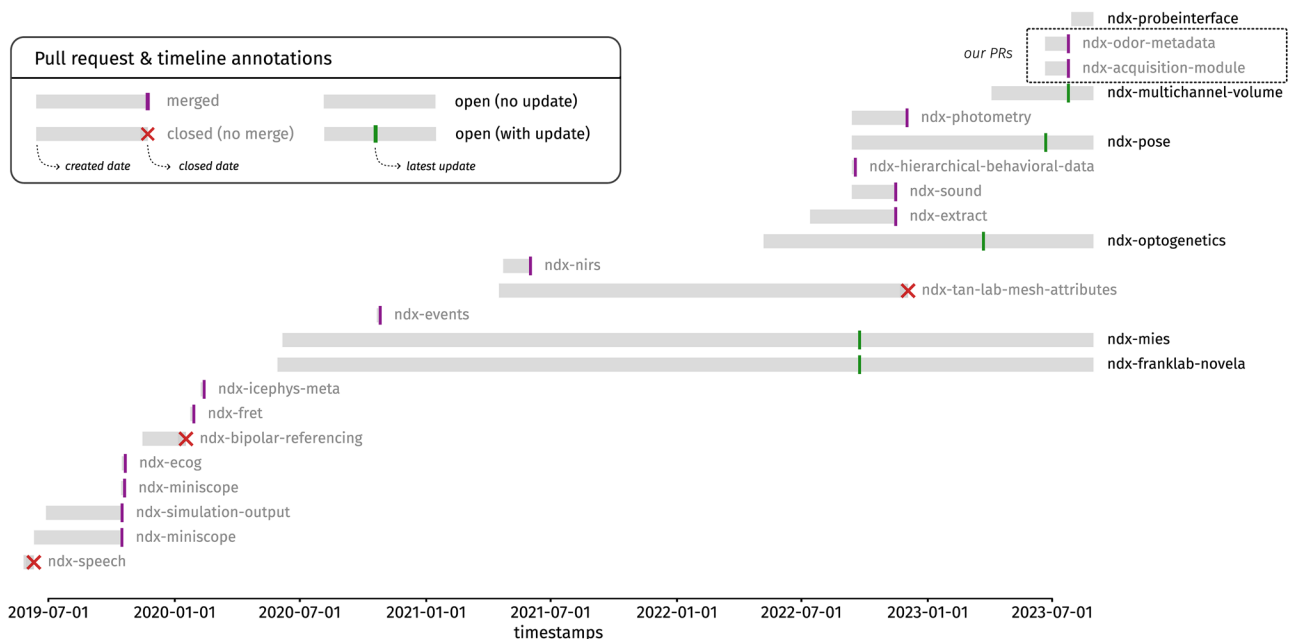


Figure 6. Pull requests (PRs) for publishing on the extension catalog may take a long time to be accepted. The data were obtained using GitHub API from [nwb-extensions/staged-extensions](https://github.com/nwb-extensions/staged-extensions) repository, on 2023-07-30. Out of 23 extension requests, about 61% (14/23) have been merged (bars ended with purple vertical sticks) and added to the catalog, while 13% (3/23) are closed without being added to the catalog (bars ended with red crosses). The review times for finished PRs vary, ranging between within a day to less than five months for most of them, with the exception being 1.6 years for the closed request for `ndx-tan-lab-mesh-attributes`. About 26% of the extension PRs (6/23) are still open, with 3 out of 6 being stale for more than a year. A notable one is `ndx-pose` for pose estimation extension (PR #31) which has been open for almost a year (Sept. 2022). Note: any closed/merged PR finished within less than 5 days is artificially extended to be 5 days for visibility.

`ndx-meta.yaml` file are filled in, and if all tests pass. Also, the bot could help for updating the extensions, say if the extension template or if some dependency has changed. Also, publishing to PyPI could be streamlined, for example by having a continuous integration (CI) job in the `ndx-template` extension template that supports automatic publishing to PyPI.

Additionally, we suggest adding some metadata to improve quality checks, centralization, and organization of extensions. To maintain quality control, the catalog could allow entries to be tagged to indicate whether an extension has been reviewed, similar to the distinction of preprint from peer-reviewed publications. To tackle fragmentation of extensions and tools, it might be helpful to also allow optional specification of the type and lineage of each entry, e.g., whether it is built upon another extension, and if it is a template extension for demonstration purposes. Additionally, we found it unclear whether the catalog submission policy welcomed lab-specific extensions (e.g., `ndx-ibl` for the IBL and ours `ndx-fleischmann-lab`), though in comments on an earlier draft, the NWB team clarified that they do encourage such submissions (personal communication). Although lab specific, these extensions could be useful examples or starting points for other labs to develop their own.

We hope to see depositing on the community catalog become more flexible and timely. A disadvantage is a potential reduction of quality control. However, more engagement, contribution, feedback, and discussion from the community are in general more likely to accelerate development of the standard. Extensions may serve as a starting point for such discussions, responding to community needs.

Considerations for Sharing on DANDI

In this section, we look at the last step of the data conversion workflow: data have already been converted to NWB and the researcher wants to share the data on a public repository, for

example to accompany a published paper. Here we look at DANDI (Halchenko et al., 2022), as the default solution recommended by the NWB team.

Potential surprises with data validation

One possible source of friction is validating the data before being able to push to DANDI. DANDI enforces a set of rules that NWB files have to meet before upload and publication as a “dandiset” is allowed, intended to promote adherence to consistent metadata standards and ensure the FAIRness (Wilkinson et al., 2016) of the archive. If files do not meet those requirements, researchers may need to (iteratively) redo their conversion with altered settings. This can be an unpleasant surprise, as one might have thought that having converted to NWB itself would be sufficient.

One solution could be to promote and describe the NWBInspector tool (RRID:SCR_025465; <https://nwbinspector.readthedocs.io>), used to validate NWB files, in the documentation and tutorials on how to create NWB files. It would also be helpful to be able to run NWBInspector from PyNWB to check files and get feedback at the time of initial conversion. This solution may soon be implemented when using no-code tools like NWB GUIDE (Flynn et al., 2024; <https://nwb-guide.readthedocs.io>; see also “Off the shelf NWB conversion”), though it did not exist when we started our projects.

Another point of friction can arise if a dandiset has already been published but needs to be updated later, for example (see <https://github.com/dandi/helpdesk/issues/98>). In our case, the validation rules changed after we first released the dandiset, and files that were already published became retroactively non-compliant. We had to go back to conversion from raw data. In general, if the cost to update a dandiset is too high, the risk is that researchers may decide not to correct stale or inaccurate information.

A potential solution would be to allow version-controlled inspection (Fig. 7). There could be at least two levels of

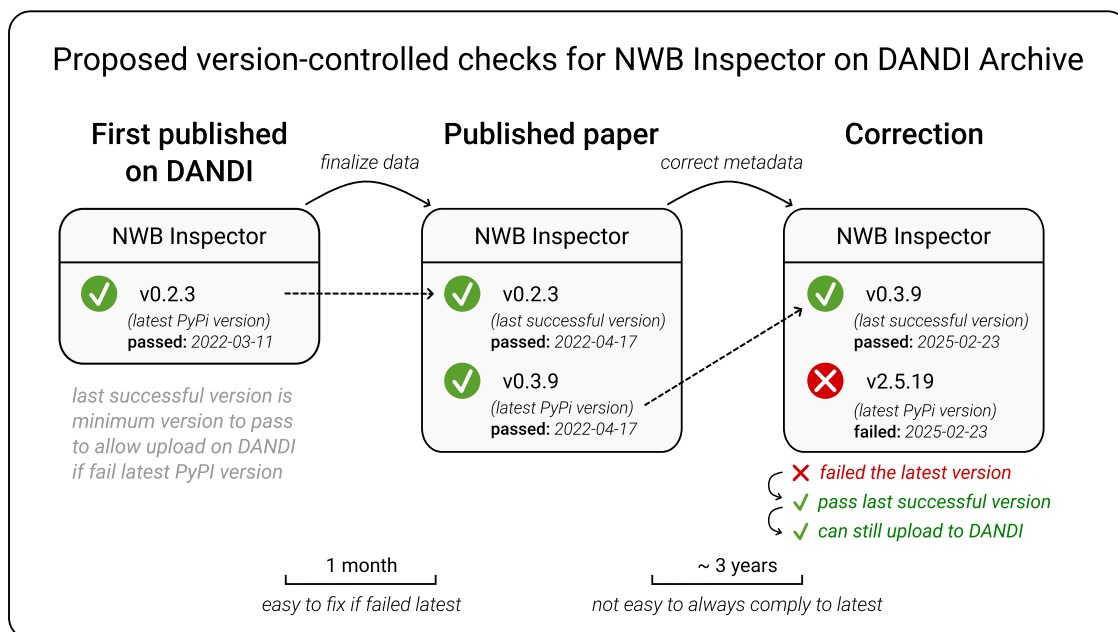


Figure 7. Proposed version-controlled checks for NWBInspector when uploading to DANDI Archive. To be published on DANDI Archive, datasets should always be checked and pass the latest version of NWB Inspector (first and second boxes) to maintain compliance with best practices. When existing datasets need to be updated, they may fail the latest version, for example 3 years after publication, to correct metadata (third box on left). The proposed solution is to allow for checking against the *last working version* for existing datasets, in cases of non-compliance with the latest version. This solution still allows researchers to disseminate updates and corrections, while maintaining transparency for the community in terms of non-compliance. This solution can be allowed a limited number of times, and failures can also be reported to DANDI Archive maintainers.

NWBInspector passing. Files that pass the most recent NWBInspector can always be uploaded. But if some files already on DANDI get updated and fail the most recent inspection, they could still be uploadable given they passed the previous working version of NWBInspector. Similar to CI systems, logs of fail/pass versions could be attached to the archive for developers and others to inspect. This approach would allow for researchers to flexibly upload corrections and updates, while still being transparent about compliance status. Failures could be reported to the DANDI team, allowing them to work with researchers to follow up-to-date best practices.

Modification of file organization

Another potential surprise is that the DANDI upload tool renames and reorganizes files into a “flatter” hierarchy. For example, one could have their NWB files organized by experiments with a nested directory structure organized by areas of recording, but DANDI refactors this structure to be organized only by subject directories, and moreover renames files by subject name and data type. DANDI also modifies external file links stored inside each NWB file to stay consistent with these file changes.

Changing the file structure may break existing analysis pipelines based on the original paths. Thus, it may be useful to think about data archiving from the start of a project. In that case, publishing the data to DANDI from the beginning of the project, with occasional updates, would make the researcher aware of this reorganization and account for it in their own code. In addition to saving effort at publication time, such a workflow would enhance analysis reproducibility. However, the cost is some increased overhead while data collection is still occurring.

Alternatives to DANDI and general strategy with data repositories

DANDI has strong restrictions on data file formats. While there is currently an exception on DANDI (Rodgers, 2022) that contains free-form source data (e.g., Python and NPY files), it is unclear whether this feature will officially be supported in the long run. Alternative repositories include Zenodo (<https://zenodo.org/>), Figshare (<https://figshare.com/>), GIN G-Node (<https://gin.g-node.org/>), OSF (Foster and Deardorff, 2017), or university data storage, potentially with Globus endpoints (RRID:SCR_011887; Foster and Kesselman, 1997, 1998; Foster, 2006). An alternative decentralized solution is Academic Torrents (Cohen and Lo, 2014; Lo and Cohen, 2016), which uses the BitTorrent protocol and leverages university bandwidth to avoid unsustainable data storage costs over the long term. These data archives can include NWB data and all related data such as raw data, pre-conversion data, analysis and summary data.

However, it may not always be feasible to centralize all data, and researchers might instead use a multi-site storage strategy. Large source data, including raw and pre-conversion data, could be deposited on university storage solutions, with Globus endpoints if possible, to take advantage of universities’ generally less restrictive quotas, assuming these data would rarely be accessed, updated, or used after conversion. Converted NWB files could then be deposited on DANDI, on which researchers can benefit from specialized software tools, as well as DANDI Hub, a Jupyter Hub with free computing resources on AWS. Lastly, along with code and documentation, researchers could continuously work on data with their analysis pipelines using solutions such as GIN G-Node, GitHub/GitLab with a DataLad (Halchenko et al., 2021), or data version control (DVC)

(Barrak et al., 2021; Kuprieiev et al., 2024) backend to manage aggregated and analyzed data and code. This helps with version-controlled code and data, without the restrictions from DANDI Archive.

We note that if researchers decide to follow a multi-site strategy, they would need to manually link these different archives together, preferably with DOI numbers and in machine-readable metadata on these different providers. The outlined example strategy separates the three archives (e.g., university storage, DANDI Archive, and GIN G-Node) by an assumed increasing update frequency, i.e., raw data files are less frequently updated compared to NWB files, and NWB files less than files with analysis or modeling results. With distributed storage, especially if these assumptions do not apply, researchers would need to manually keep track and link the updates regularly.

Suggestions to Streamline Data Reading and Writing

Data exploration tool guidance

The NWB ecosystem has many applications available for a researcher to quickly get a sense of what is inside an NWB file. As of writing, there are 4 general and 15 specialized data tools listed on the NWB Overview (<https://nwb-overview.readthedocs.io/>), and new tools continue to emerge. The number of active projects indicates a vibrant development community. However, new users may be overwhelmed by the choices, and not know how, except through brute force trials, to determine which tools are best for them. Moreover, consolidation around a few key applications could help channel valuable developer efforts into refining and improving existing tools, some of which still exhibit rough spots like freezing on large files or frequent crashes.

This situation is common in open source development ecosystems (for example, there are many partially redundant but not interchangeable python plotting packages). A difference here is that the NWB standard was created and continues to be maintained through a somewhat centralized development team, with an explicit agenda to be adopted as a ubiquitous standard for neurophysiology. There is thus a stronger case that innovation arising from widely dispersed development should be balanced by centralized advising over third party tools.

For example, primary NWB documentation could maintain a section with some (automatically scraped) metrics for each tool (e.g., number of GitHub stars and number of downloads on PyPI) next to accessible summaries of the features of each tool, and descriptions of who their target users are. At time of writing, several of these changes are in process or planned (NWB Team, personal communication; https://nwb-overview.readthedocs.io/en/latest/tools/analysis_tools_home.html).

A more assertive approach would select recommended tools, on the basis of features, robustness (e.g., resolution of bugs and handling of large file sizes), and probable longevity. For data exploration, some natural candidates could be NWBWidgets (RRID:SCR_021154; <https://nwb-widgets.readthedocs.io/>), which is also integrated with DANDI Hub, and relatively newer NeuroSift (RRID:SCR_025466; Magland et al., 2024), which is an interactive visualization tool that works directly in the user’s browser. In our experience, NeuroSift is highly accessible, without requiring installation, and offers strong visualization functionality out of the box. Both tools support streaming data from the DANDI Archive. Again, the goal would be to provide soft incentives that encourage contributors to focus primarily on existing tool refinement, while still leaving space for new specialized projects in early development.

Data access pain points

Figuring out where data are

We find that new NWB users often struggle to find and access information, with confusion arising from where the information is in the internal hierarchy, or because the datatype of a particular object does not intuitively describe what it is. Many scientists look first for modules based on source of data (e.g., fluorescence, behavior, and stimuli). But access under the NWB schema runs first through stage of processing (e.g., acquisition, preprocessing, and analysis) and then descends through multiple levels of hierarchy to data source. That is, researchers may employ a mental sequence of *where is my behavior* (say) then *what processing has been applied*, which is the opposite ordering from what NWB currently uses (Fig. 2).

An outlier is that `stimulus` is at the top of the hierarchy, with `acquisition` and `processing`. However, stimulus time series sometimes need additional processing, for example, to transform raw digital outputs recorded by behavior control devices into a semantically useful tabular format. Should such stimuli be saved within `stimulus` (with processing stage indicated in name or `description` attributes) or in a module inside `processing`? Additionally, tables cannot be saved inside `stimulus`, and only limited metadata can be associated. It is recommended to use dedicated modules or objects designed to save metadata, for example `devices` for recording or `lab_metadata` for lab-specific metadata. This again runs into the potential issue of categorically similar objects being widely separated.

Cumbersome syntax to extract data

A challenge for new users that is parallel to understanding object locations is confusion over the addressing syntax, i.e., when to use dot syntax, `object1.object2`, or Python dictionary syntax, `object1["object2"]`. The syntactic variation derives from the structure of the HDF5 file specification and the NWB schema, both of which are generally unknown and opaque to users.

Two obvious alternative possibilities for API syntax would simply make one or the other access method universal (e.g., through a Python `DataClass`). Either choice would obscure the real differences between types of objects in the NWB implementation (e.g., a fluorescence object including metadata attributes, vs a numpy array just of the dF/F_0 values), but we are not convinced that most users benefit from having these differences encoded in syntax.

Another possibility that is both general and convenient for programmatic access would support universal reference via “path strings”, such as `nwbfile[pathstr]` where `pathstr = 'object1/object2/object3'`.

Listing 1. Retrieving data using the PyNWB API

```
# 1D array of timestamps
t = nwb_file.processing["behavior"]["interpd_500"] [
    ↪ therm_highpassed"].timestamps[:]

# 1D array of data
therm = nwb_file.processing["behavior"] ["interpd_500"
    ↪ ] ["therm_highpassed"].data[:]
```

Lab-specific wrapper workaround

In its current state, long hierarchies in NWB files (e.g., processing → behavior → interpolated → position → data) are slow to type and hard to remember, and tend to clutter code. A common method to hide complexity in an individual user's analysis code is to first create “wrappers” (Fig. 8). For example, a wrapper may define simple “`get()`” methods that automatically skip parts of the object path, e.g., `data=nwb_wrapper.get("dFF0")`. Wrappers can also add convenience features, such as aggregating different time series into a single data frame, and wrappers can be stored in dictionaries for easy looping over multiple files.

On the other hand, wrappers may be complex to design and may introduce a maintenance burden if they aim to work across the usually wide range of experiments and data streams that arise even within a single lab. In practice, then, individual researchers often end up partially or completely rewriting similar helper code with each new project.

Suggestion for better data access: tags and aliases

A potential solution for better data access is a feature we call “fluid NWB” (Fig. 9), allowing for a list of tags for each object, including “flat” objects such as timeseries, tables, and modules. Users could add annotations and categories as they see fit, and specialized communities could evolve their own norms for “virtual” file organization, without confounding the underlying standard. Aliases, to our knowledge, are currently not possible, but the integration of such a feature may allow for users to have easier and quicker access, and could also aid documentation. For example, the AllenSDK has a dedicated dictionary for metadata field mapping to NWB/HDF5 locations; this shares some similarity with aliasing and illustrates a place for annotation usage (https://alleninstitute.github.io/AllenSDK/allensdk.core.brain_observatory_nwb_data_set.html#allensdk.core.brain_observatory_nwb_data_set.BrainObservatoryNwbDataSet.FILE_METADATA_MAPPING). Supporting custom tags for neurodata types is currently an open GitHub issue (see <https://github.com/NeurodataWithoutBorders/nwb-schema/issues/531>).

Tags and aliases would be a “decorative layer” on top of the NWB standard, allowing for more “fluid” data structures, which researchers and developers could exploit for usability and discoverability. However, in the absence of convergence on naming norms within a given research area, overlapping tags, complex tag formatting, and tag relations could proliferate to the point of no longer being useful. For example, should cardiac recordings (electrocardiogram), saccades, and arena locations all carry a common *behavior* tag? Should muscle recordings (EMG) be tagged both as *neural* and *behavior* in a brain-machine-interface study? The added flexibility of an alias or tag system would

Listing 2. Retrieving data through a custom wrapper

```
# Use wrapper to create an alias to the data
interpd = MyCustomWrapper(
    nwb_field="processing",
    path_to_interface=["behavior", "interpd_500"],
    nwb_file=nwb_file
)

# After one-time setup, simpler data retrieval
t, therm = interpd.get("therm_highpassed")
```

Figure 8. Code snippet comparison showing how to retrieve data from an NWB file using the “raw” PyNWB API (left) compared to using a custom wrapper (right). After a one-time setup, retrieving the data through a custom wrapper reduces the cognitive load for the user.

(a)

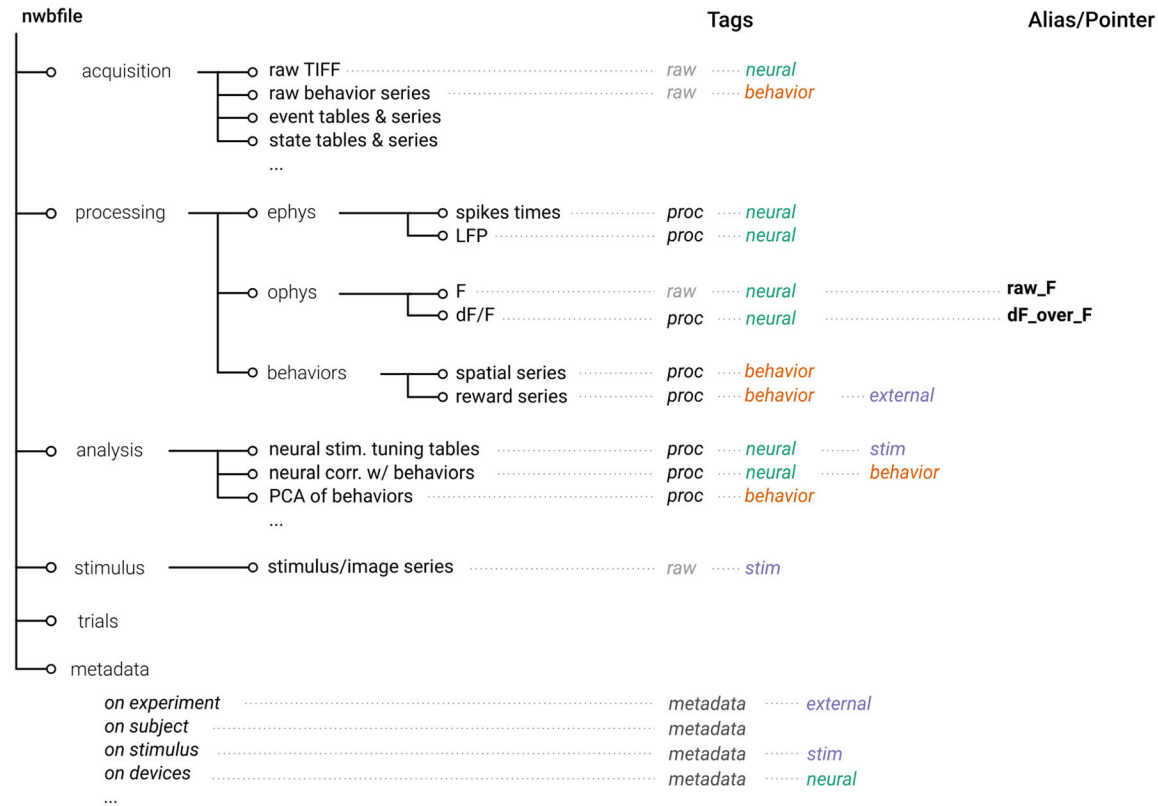
Current NWB

*hierarchical, processing-stage focused
standard, more structured, defined
uses hdf5 and/or pynwb*

(b)

Proposed "decorative layer" for *fluid* NWB

*tag & alias support to assist query, exploration & analysis
less structured, more user/lab/community-defined
but needs a special API*



(c)

User: unclear where all neural-related data are, need to look through everything

-> `fluid_nwb.get(nwbfile, tags="neural")`

Narrow down with

-> `fluid_nwb.get(nwbfile, tags=["neural", "proc"])`

(d)

`pynwb: nwbfile.processing["ophys"]["Fluorescence"]["RoiResponseSeries"].data[.]`

-> `fluid_nwb.get(nwbfile, alias="raw_F")`

Figure 9. A proposed design layer for the NWB standard to assist with data retrieval and organization. The nature of the current NWB structure is hierarchical and tends to be organized by processing stages; panel (a) shows an example of this structure. Accessing relevant data requires knowledge of where it is located, which may be multiple levels deep, see for example bottom box (d) to access raw fluorescence data with `PYNWB`. The proposed "decorative layer" allows for more "fluid" interaction with NWB via additional specifications in NWB objects to assist querying, exploration and analysis with more user/lab/community's control and customization, without breaking the existing hierarchical NWB structure. Panel (b) illustrates examples of adding tags and aliases. Tags can be more specific, multi-faceted, and customized to concepts of recording/analysis that users tend to look for (e.g., *neural*, *behavior*, *stim*, and *external*), as well as higher level details such as processing stages (e.g., *raw* and *proc*). Aliases and/or pointers allow users to add names for objects that are most frequently accessed, or expected to be so. Taking advantage of this "decorative layer," users and developers may design a `fluid_nwb` API to interact with NWB files in a more flexible and less verbose manner, for example with tags in box (c) and aliases in box (d).

produce the greatest benefit if complemented by a process to secure community consensus around tagging conventions.

Discussion

Standardization is an essential component of modern data management, analysis, and sharing, and NWB has introduced a comprehensive and versatile data science ecosystem for neuroscience

research. However, our experience suggests that implementation of NWB workflows at the level of individual labs or research collaborations still requires significant effort and commitment. Furthermore, given the rapid pace of technology development in neuroscience research, we expect that the development and implementation of adequate data science tools will continue to pose new challenges for some time. Solutions to these challenges will likely require a reorganization of neuroscience research to

facilitate interdisciplinary collaborations, including additional institutional support not just for the creation of new tools but also for their adoption by research labs at all levels of technical capability.

Our intention with this article was to describe our experience in implementing our own NWB pipeline. It is important to note that some of the rough edges that we experienced have already been improved, thanks to the NWB team being very responsive both during our pipeline development and since the preprint version of this article has been published online. We are convinced that researchers should use a standard to share their data, and we believe NWB is a serious contender for the job. We want to acknowledge that the issues we encountered are not specific to NWB per se and would have needed to be solved by any alternative standard. These problems and considerations are very likely to be universal for standardizing, managing, and sharing neurophysiology datasets.

For researchers starting their data standardization journey today, we would recommend using and building on top of projects such as NWB GUIDE (Flynn et al., 2024) and NeuroConv (Baker et al., 2023). Such tools were not readily available when we started our data standardization pipeline, but they are nowadays a good choice for newcomers to benefit from accessible tools, tested by the entire NWB community. The implementation and use of NWB standards in our labs has consistently increased over the past five years. We anticipate that the benefits of standardized data science pipelines will be a major asset for future generations of trainees in the labs.

References

- Baker C, Mayorquin H, Weigl AS, Tauffer L, Buccino AP, Sharda S, Dichter B (2023) NeuroConv. original-date: 2022-07-19T16:49:38Z.
- Barrak A, Eghan EE, Adams B (2021) On the co-evolution of ML pipelines and source code—empirical study of DVC projects. In: *2021 IEEE International Conference on Software Analysis, Evolution and Reengineering (SANER)*, pp 422–433, Honolulu, HI: IEEE.
- Boivin B, Adil N, Neufeld S (2021) Inscopix CNMF-E. original-date: 2021-02-09T01:17:56Z.
- Braun E, et al. (2022) Comprehensive cell atlas of the first-trimester developing human brain. Pages: 2022.10.24.513487. Section: new results.
- Brose K (2016) Global neuroscience. *Neuron* 92:557–558.
- Callaway EM, et al. (2021) A multimodal cell census and atlas of the mammalian primary motor cortex. *Nature* 598:86–102.
- Carver JC, Weber N, Ram K, Gesing S, Katz DS (2022) A survey of the state of the practice for research software in the United States. *PeerJ Comput Sci* 8:e963.
- Castro JB, Gould TJ, Pellegrino R, Liang Z, Coleman LA, Patel F, Wallace DS, Bhatnagar T, Mainland JD, Gerkin RC (2022) “Pyrume: A Window to the World’s Olfactory Data.” Preprint, Neuroscience.
- Chen X, Li H (2017) ArControl: an arduino-based comprehensive behavioral platform with real-time performance. *Front Behav Neurosci* 11:244.
- Churchland MM, Cunningham JP, Kaufman MT, Foster JD, Nuyujukian P, Ryu SI, Shenoy KV (2012) Neural population dynamics during reaching. *Nature* 487:51–56.
- Cohen JP, Lo HZ (2014) Academic torrents: a community-maintained distributed repository. In: *Proceedings of the 2014 Annual Conference on Extreme Science and Engineering Discovery Environment, XSEDE '14*, pp 1–2. New York: Association for Computing Machinery.
- Contaxis N, et al. (2024) Building on NIH’s data sharing policy. *Science* 384:747–748.
- Cooke NJ, Hilton ML, editors (2015) *Enhancing the effectiveness of team science*. Washington, DC: National Academies Press.
- Dai K, et al. (2020) The SONATA data format for efficient description of large-scale network models. *PLoS Comput Biol* 16:e1007696.
- Dallmeier-Tiessen S, Darby R, Gitmans K, Lambert S, Matthews B, Mele S, Suhonen J, Wilson M (2014) Enabling sharing and reuse of scientific data. *New Rev Inf Networking* 19:16–43.
- Daste S, Pierré A (2022) Two photon calcium imaging of mice piriform cortex under passive odor presentation (Version 0.220928.1306) [Data set]. DANDI archive.
- Deitch D, Rubin A, Ziv Y (2021) Representational drift in the mouse visual cortex. *Curr Biol* 31:4327–4339.e6.
- Flynn MAGM, et al. (2024) Neurodata Without Borders/nwb-guide: 0.0.15.
- Foster I (2006) Globus toolkit version 4: software for service-oriented systems. *J Comput Sci Technol* 21:513–520.
- Foster ED, Deardorff A (2017) Open science framework (OSF). *J Med Libr Assoc* 105:203–206.
- Foster I, Kesselman C (1997) Globus: a metacomputing infrastructure toolkit. *Int J Supercomput Appl High Perform Comput* 11:115–128.
- Foster I, Kesselman C (1998) The globus project: a status report. In: *Proceedings Seventh Heterogeneous Computing Workshop (HCW'98)*, Orlando, FL, 1998, pp 4–18.
- Gardner D, Goldberg DH, Grafstein B, Robert A, Gardner EP (2008) Terminology for neuroscience data discovery: multi-tree syntax and investigator-derived semantics. *Neuroinformatics* 6:161–174.
- Gibson F, et al. (2009) Minimum information about a neuroscience investigation (MINI): electrophysiology. *Nat Precedings* 1–7.
- Gorgolewski KJ, et al. (2016) The brain imaging data structure, a format for organizing and describing outputs of neuroimaging experiments. *Sci Data* 3:160044.
- Grewe J, Wachtler T, Benda J (2011) A bottom-up approach to data annotation in neurophysiology. *Front Neuroinform* 5:16.
- Halchenko Y, et al. (2021) DataLad: distributed system for joint management of code, data, and their relationship. *J Open Sour Software* 6:3262.
- Halchenko Y, et al. (2022) dandi/dandi-cli: 0.46.2.
- Holdgraf C, et al. (2019) iEEG-BIDS, extending the brain imaging data structure specification to human intracranial electrophysiology. *Sci Data* 6:102.
- Jorgenson LA, et al. (2015) The BRAIN initiative: developing technology to catalyze neuroscience discovery. *Philos Trans R Soc B: Biol Sci* 370:20140164.
- Kaiser J (2022) NIH’s BRAIN Initiative Puts \$500 Million into Creating Most Detailed Ever Human Brain Atlas. 22 Sept. 2022. Available at: <https://www.science.org/content/article/nihs-brain-initiative-puts-dollar500-million-creating-detailed-ever-human-brain-atlas>.
- Kim S, et al. (2023) PubChem 2023 update. *Nucleic Acids Res* D51:D1373–D1380.
- Koch C, et al. (2022) Next-generation brain observatories. *Neuron* 110:3661–3666.
- Koch C, Jones A (2016) Big science, team science, and open science for neuroscience. *Neuron* 92:612–616.
- Krakauer JW, Ghazanfar AA, Gomez-Marín A, MacIver MA, Poeppel D (2017) Neuroscience needs behavior: correcting a reductionist bias. *Neuron* 93:480–490.
- Kupriev R, et al. (2024) DVC: Data Version Control - Git for Data & Models. Zenodo.
- Langlieb J, et al. (2023) The cell type composition of the adult mouse brain revealed by single cell and spatial genomics. Pages: 2023.03.06.531307. Section: New results.
- Li B, Kamarck ML, Peng Q, Lim F-L, Keller A, Smeets MAM, Mainland JD, Wang S (2022) From musk to body odor: decoding olfaction through genetic variation. *PLoS Genet* 18:e1009564.
- Lo HZ, Cohen JP (2016) “Academic Torrents: Scalable Data Distribution.” arXiv:1603.04395 [cs].
- Loombs S, et al. (2022) Connectomic comparison of mouse and human cortex. *Science* 377:eabo0924.
- Magland J, Soules J, Baker C, Dichter B (2024) Neurosift: dANDI exploration and NWB visualization in the browser. *J Open Sour Software* 9:6590.
- Martone M, Gerkin R, Moucek R, Das S, Goscinski W, Hellgren-Kotaleski J, Kennedy D, Leergaard T, Boline J, Abrams M (2020) NIX—neuroscience information exchange format. *F1000Research* 9:358.
- Martone ME, Nakamura R (2022) Changing the culture on data management and sharing: getting ready for the new NIH data sharing policy. *Harvard Data Science Review* 4:3.
- Mathis A, Mamidanna P, Cury KM, Abe T, Murthy VN, Mathis MW, Bethge M (2018) DeepLabCut: markerless pose estimation of user-defined body parts with deep learning. *Nat Neurosci* 21:1281–1289.
- Pachitariu M, Stringer C, Dipoppa M, Schröder S, Rossi LF, Dalgleish H, Carandini M, Harris KD (2016) “Suite2p: Beyond 10,000 Neurons with Standard Two-Photon Microscopy.” preprint, Neuroscience.

- Pasquetto IV, Randles BM, Borgman CL (2017) On the reuse of scientific data. *CODATA Data Sci J* 16:8.
- Pham T (2023a) ndx-fleischmann-labmetadata.
- Pham T (2023b) ndx-odor-metadata.
- Pierré A, Pham T (2023) calimag. Language: eng.
- Plume Andrew, van Weijen Daphne (2014) Publish or perish? The rise of the fractional author ... *Research trends* 1:16–18.
- Rodgers CC (2022) A detailed behavioral, videographic, and neural dataset on object recognition in mice. *Sci Data* 9:620.
- Rübel O, et al. (2019) NWB:N 2.0: an accessible data standard for neurophysiology. Pages: 523035. Section: new results.
- Rübel O, et al. (2022) The neurodata without borders ecosystem for neurophysiological data science. *eLife* 11:1–48.
- Rübel O, Prabhat M, Denes P, Conant D, Chang E, Bouchard K (2015) BRAINformat: a data standardization framework for neuroscience data. *bioRxiv* 024521.
- Scheffer LK, et al. (2020) A connectome and analysis of the adult *Drosophila* central brain. *eLife* 9:e57443.
- Schneider A, Azabou M, McDougall-Vigier L, Parks DF, Ensley S, Bhaskaran-Nair K, Nowakowski T, Dyer EL, Hengen KB (2023) Transcriptomic cell type structures in vivo neuronal activity across multiple timescales. *Cell Rep* 42:112318.
- Sherry Z, Jaggi A, Brann D, Lovell J, Weinreb C (2023) dattalab/moseq2-app: v1.3.1.
- Siegle JH, et al. (2021) Survey of spiking in the mouse visual system reveals functional hierarchy. *Nature* 592:86–92.
- Srinivasan S, Daste S, Modi MN, Turner GC, Fleischmann A, Navlakha S (2023) Effects of stochastic coding on olfactory discrimination in flies and mice. *PLoS Biol* 21:e3002206.
- Steinmetz NA, Zatzka-Haas P, Carandini M, Harris KD (2019) Distributed coding of choice, action and engagement across the mouse brain. *Nature* 576:266–273.
- Stringer C, Pachitariu M, Steinmetz N, Reddy CB, Carandini M, Harris KD (2019) Spontaneous behaviors drive multidimensional, brainwide activity. *Science* 364:eaav7893.
- Syeda A, Zhong L, Tung R, Long W, Pachitariu M, Stringer C (2022) Facemap: a framework for modeling neural activity based on orofacial tracking. Pages: 2022.11.03.515121. Section: new results.
- Teeters JL, et al. (2015) Neurodata without borders: creating a common data format for neurophysiology. *Neuron* 88:629–634.
- Tenopir C, Dalton ED, Allard S, Frame M, Pjesivac I, Birch B, Pollock D, Dorsett K (2015) Changes in data sharing and data reuse practices and perceptions among scientists worldwide. *PLoS One* 10:e0134826.
- The International Brain Laboratory, et al. (2020) Data architecture for a large-scale neuroscience collaboration. Pages: 827873. Section: new results.
- The International Brain Laboratory, et al. (2023) A modular architecture for organizing, processing and sharing neurophysiology data. *Nat Methods* 20:403–407.
- Tritt Andrew J, Rbel Oliver, Dichter Benjamin, Ly Ryan, Kang Donghe, Chang Edward F, Frank Loren M, Bouchard Kristofer 2019 “HDMF: Hierarchical Data Modeling Framework for Modern Science Data Standards” 2019 IEEE International Conference on Big Data Big Data Los Angeles CA USA 2019 pp. 165–179.
- Van Viegen T, et al. (2021) Neuromatch academy: teaching computational neuroscience with global accessibility. *Trends Cognit Sci* 25:535–538.
- Wilkinson MD, et al. (2016) The FAIR guiding principles for scientific data management and stewardship. *Sci Data* 3:160018.
- Wilson CD, Serrano GO, Koulakov AA, Rinberg D (2017) A primacy code for odor identity. *Nat Commun* 8:1477.
- Wiltshko AB, Johnson MJ, Iurilli G, Peterson RE, Katon JM, Pashkovski SL, Abaira VE, Adams RP, Datta SR (2015) Mapping sub-second structure in mouse behavior. *Neuron* 88:1121–1135.
- Wiltshko AB, Tsukahara T, Zeine A, Anyoha R, Gillis WF, Markowitz JE, Peterson RE, Katon J, Johnson MJ, Datta SR (2020) Revealing the structure of pharmacobehavioral space through motion sequencing. *Nat Neurosci* 23:1433–1443.
- Yao Z, et al. (2023) A high-resolution transcriptomic and spatial atlas of cell types in the whole mouse brain. Pages: 2023.03.06.531121. Section: new results.