

Data architecture and visualization for a large-scale neuroscience collaboration

The International Brain Laboratory, Niccolò Bonacchi¹, Gaelle Chapuis², Anne Churchland³, Kenneth D. Harris², Cyrille Rossant², Maho Sasaki⁴, Shan Shen⁴, Nicholas A. Steinmetz⁵, Edgar Y. Walker⁴, Olivier Winter¹, Miles Wells²

¹Champalimaud Center for the Unknown, Av. Brasília, 1400-038 Lisboa, Portugal

²Institute of Neurology, University College London, London WC1N 3BG, UK

³Cold Spring Harbor Laboratory, Cold Spring Harbor, NY 11724, USA

⁴DataJoint Neuro, Houston, Texas, 77021

⁵Department of Biological Structure, University of Washington, Seattle, WA 98195, USA

Correspondence: kenneth.harris@ucl.ac.uk

Abstract

The International Brain laboratory (IBL) is a collaboration aiming to understand the neural basis of decision-making. Ten experimental labs use multiple neural recording modalities in diverse brain structures of mice making perceptual decisions. A primary requirement of IBL is to establish a data architecture that integrates data from all labs and modalities together. We have developed a system that allows users across 5 countries to automatically contribute data and metadata, search for relevant data, and share code for exploratory analysis. To accurately record metadata about subjects and experiments in a searchable and accessible way, we have developed a user-friendly, web-based electronic lab notebook system for colony and experiment management (Alyx). Alyx is a small relational metadatabase that records relevant information about each subject (age, strain, procedures), alongside information about experiments and their resulting data files. Once an experiment is completed, data registered in Alyx is automatically uploaded to a central server via Globus transfer. Users can search and access the data using a lightweight interface called the Open Neurophysiology Environment (ONE), implemented in Python and MATLAB. ONE defines a list of DatasetTypes, as arrays of predetermined shapes and units. Users search the database by running a command that returns an ID identifying experiments matching their criteria. A command `one.load` returns the requested DatasetTypes as a numerical array, downloading it from the central server and caching on the user's machine to avoid repeated downloads. Additional commands provide further ways to list experiments on the server, and the dataset types they each contain. This system allows users to search, load and process data that was collected in laboratories spanning multiple geographical locations. To visualize these data, we have developed a pipeline for automated analysis, based on the DataJoint framework.

Introduction

Three changes in the field of neuroscience have given rise to new challenges in data management. First, datasets are becoming vastly larger. For instance, advances in neural measurement technology have made it possible to measure neural activity from hundreds of neurons in daily experiments^{1,2}. These neural datasets are often accompanied by increasingly large behavioral datasets, such as video data tracking of the animal³⁻⁶. Second, many neuroscientists are now assembling into teams that include multiple researchers and universities and even spanning several countries⁷⁻⁹. This new team-approach to neuroscience necessitates that data are made accessible to team members at multiple sites, thus demanding a much more orderly data management system than is needed in an individual lab. Finally, the field has come to appreciate the need for public data sharing. Data sharing allows a single dataset to generate many insights into brain function, often far beyond those that were the focus of the original research. In order for a dataset to be successfully shared, it must be formatted so that it is readable by new users, must include relevant metadata, and must be stored reliably to ensure access.

Despite the widespread need for improved data management, many neuroscience labs lack the infrastructure needed to improve their current systems. Methods of data storage have not evolved to support the increasingly large datasets being generated. Newly established collaborations struggle to find affordable and efficient ways to pool and visualize data across a large team. Finally, standardized data formats^{10,11} have not been fully adopted, especially compared to data formats in other fields such as genomics and astronomy.

The International Brain Laboratory (IBL) is a collaboration of 21 neuroscience laboratories studying the computations supporting decision-making⁷. The collaboration is generating unprecedented amounts of behavioral and neural data which must be available immediately to all members of the collaboration, and later to the general public. IBL team members in experimental laboratories measure behavior and neural activity from multiple brain areas, and incoming data are available to collaboration members and to the public via a common data platform. Data management challenges within IBL are simpler than the general case, as the initial focus is only on a single behavioral task implemented with standardized instrumentation. Nevertheless, the data management system for the IBL is general enough to support a broad range of approaches. Here, we describe our approach to storing and organizing data for the IBL collaboration, and document the protocols, built on existing efforts of neurodata standardization, that allow users to search and access it. The standards are designed to be extensible, so they can be used for any behavior or recording. Common data types (such as those arising from electrophysiology recordings) are standardized, but users can add arbitrary data types (for example describing new behaviors) freely using their own custom “namespace”. The IBL approach can thus serve as a template that individual labs or collaborations can use to better

manage, access and learn from their growing datasets: standardizing what can be standardized, and allowing flexibility for what cannot.

Storing and organizing the data

The IBL data are generated via neural and behavioral measurements during a dynamic decision-making task for mice⁷. Data are collected daily at 7 experimental sites during both initial training and neurophysiological recordings in expert subjects. Data collected during decision-making experiments includes the output of sensors that subjects use to report choices, video recordings of the face and body from 3 high-speed cameras, sounds, and measurements from additional sensors that track ambient environmental parameters (e.g. temperature). Metadata on all subjects (e.g. genotype, age, surgical and training history, see Table 1) is also kept.

An initial goal of the IBL project is to collect a brain-wide activity map: sampling 270 recording sites twice each using Neuropixels probes, along with a single site that is sampled more heavily to assess repeatability across laboratories. This will generate recordings from ~200,000 neurons. In addition to this brain-wide map, individual IBL scientists will perform diverse experiments applying other recording modalities to the same task (e.g. two photon calcium imaging, fiber photometry of neuromodulators), or perform additional manipulations such as optogenetic stimulation or inactivation of specific cell populations. All these data must be put into a standard form, and collected into a single central location.

Public release of these data will require careful quality control and manual curation, and will only occur after sufficient time has allowed for this (planned for September 2020). In the meantime however internal IBL scientists require access to these data from multiple international locations. Different scientists require different data items, to perform tasks such as developing quality control metrics and analyzing behavioral learning curves. Data sharing solutions that require researchers to download the entire dataset to access specialized data items would therefore be impractical, adding unnecessary obstacles to the already difficult task of analyzing and exploring complex datasets. Furthermore, scientists need to be able to search for experiments and data items matching specified criteria. We have now established our core data pipeline and sharing architecture, which solves these problems, and will also form the basis of our external sharing system when the data are publicly released. As of November 2019, this system stores information on 12,250 behavioral sessions (currently 285,000 files).

Our data architecture is based on two components. A central relational database (hosted by Amazon Web Services) stores extensive metadata on all mice (e.g. genotype, lineage and history of surgery, weight, and water administration) experiments (e.g. experimenter, time of day), and data files. A bulk data server (250TB, currently hosted at the Flatiron Institute in New York but accessible worldwide) stores binary data files recorded by the experimental apparatus (Figure 1). The bulk data are generally stored as flat binary (.npy) files, however the backend format is irrelevant to users as access is provided via an API, which automatically downloads and delivers binary arrays of requested data directly to analysis software (Python or MATLAB), caching on the user's local computer to avoid repeated downloads. The files follow a standard

naming convention (described in Appendix 1) that allows relationships between data items to be easily encoded and understood by users.

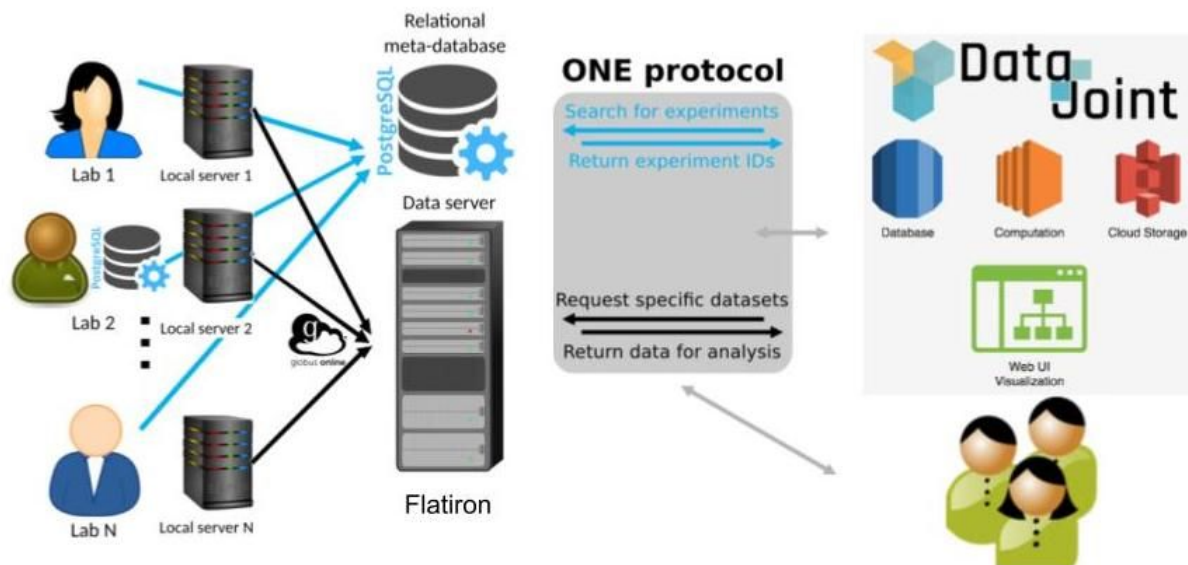


Figure 1. IBL data architecture. Data and metadata from each lab are integrated into central servers hosted at the Flatiron Institute via read-only connections. These data are then hosted to end users either directly through the Open Neurophysiology Environment protocol, or via additional protocols such as DataJoint or NeuroData without Borders.

All IBL labs run behavioral experiments using a common system of hardware and software. When an experiment starts, the experiment control computer contacts the central database to register the experiment together with metadata such as the subject ID. All experimental data files are registered with the central database, which automatically initiates a transfer of these files to the central bulk data server at 2am local time via Globus, a protocol for data transfer and sharing. Preprocessing steps such as video analysis and compression, spike sorting, and lossless compression of raw electrophysiology data are performed on the user's local machine prior to upload. Video is compressed using ffmpeg h264 codec with compression level crf 29 prior to upload, and tracking of body parts is performed with DeepLabCut⁴; raw electrophysiology is compressed by a factor of ~3 using a custom [lossless compression algorithm](#) we have developed specifically for this purpose (Appendix 3).

A key challenge is to ensure that metadata are comprehensive and accurate. While metadata about the experiments themselves can be collected by the recording hardware, metadata about the experimental subjects must be entered by lab members. Ensuring this happens reliably is primarily a problem of “social engineering” rather than software, and we have solved it by creating a user-friendly web-based client that connects from all IBL labs to our central database. This system functions as an electronic colony management / electronic lab notebook system storing details on all labs' mice, such as age and surgical history. These are critical because mouse behaviour is affected by a large number of factors¹². The system allows metadata to be entered at the time of collection, for example, recording each subject's weight before every experiment, which ensures data are entered more reliably than would occur by transcribing from paper lab notebooks. The system also performs other functions, such as generating email

notifications telling the experimenter how much supplementary water is needed after training, and recording delivery of these supplements. Other types of metadata (e.g. genotype results), are entered as soon as they are generated, which ensures their accuracy, in contrast to later entry of paper-based notes. The system, known as Alyx, is fully operational and is available as open source at <https://github.com/cortex-lab/alyx>.

Table 1: Terminology

	What is it?	What is an example?	What is it for?
Bulk data	Large-scale raw recordings or derived preprocessed results.	Raw electrophysiology signal, spike sorting results, behavioral outcomes	Bulk data are processed, analyzed, visualized, and is therefore critical for drawing scientific conclusions.
Metadata	Small data items that provide information on bulk data.	Mouse strain/sex, electrode location, experimenter ID, recording hardware configuration and lab	Metadata allows users to search for bulk datasets they want to analyze, and is essential to interpret results.
Relational Database	A system that stores complex information in easily-searchable and interdependent tables.	PostgreSQL, MySQL	Users search the database to find data they need. Searches are flexible, and include queries not originally conceived of by the designers, eg, "Find all data run on male mice on a Tuesday".
Data Standard	A set of rules specifying how data will be represented on a computer system.	Open Neurophysiology Environment, Neurodata Without Borders	A data standard allows users to read data from multiple providers without having to learn a new convention each time.
Analysis Language	A high-level programming language allowing data processing.	Python, MATLAB	Used by scientists to process and analyze bulk data.
Data Pipeline	A software tool that allows complex multi-step computations to be run automatically.	DataJoint , joblib , Dask	Allows a standard analysis to be run repeatedly on multiple standardized datasets, saving intermediate results and continuing after interruptions.

Accessing the data

To allow for individual scientists to work in different ways, IBL provides users with three protocols for searching, downloading, and sharing data. These are Open Neurophysiology Environment, DataJoint, and Neurodata Without Borders.

The [Open Neurophysiology Environment](#) (ONE) is a simple protocol developed by IBL, which employs a workflow that closely matches the way many neuroscientists currently operate.

Nearly all our scientists use one of two languages, MATLAB or Python, which they run on desktop workstations and use to access local data files. ONE is a lightweight interface that provides users with four functions allowing them to search for experiments of interest, and load required data from these experiments directly into MATLAB or Python (Figure 1, gray square). The user need not worry about underlying file formats or network connections, and data are cached on their local machine to avoid repeated downloads. The ONE system is well established and is used daily to access our data. Furthermore, it has been designed with a view to standardization, allowing a simple way for users to analyze data from multiple projects with the same codebase. All the information needed to understand and implement ONE is publicly available ([open source implementation](#), [demo](#)). ONE has been designed as an extensible open standard that can be adopted by anyone in the community. The ONE protocol can run with multiple backends, transparently interoperable to the user. For the main IBL data, we use a backend system that uses our central database. However, we have also provided [ONE light](#) (see Appendix 2), a backend that allows scientists to share data using ONE by simply uploading files to a web server, or to [figshare](#). The ONE system is extensible, allowing users to add non-standard data sets using their own namespace, and, for example, has been used to share large-scale electrophysiology recordings collected in a different task via figshare file upload¹³. Full details of the ONE protocol are given in Appendix 1.

The second access protocol used by the IBL is DataJoint. [DataJoint](#) is a workflow management system that integrates relational database (see Table 1) with automatic processing of analyses, that is seeing increasing use in neurophysiology (Figure 1, right). It provides a Python or MATLAB interface that directly operates on the database and allows intuitive and flexible queries. In IBL, DataJoint stores both experimental data and the results of analyses in a single relational database. A key capability of DataJoint is that it allows standard analyses to be run automatically on new data as the data comes in. For example, DataJoint creates daily summaries of animal behavior for all users in the collaboration to access via a browser (Figure 2). In collaboration with the team of DataJoint Neuro (the developers of this system), we have established a DataJoint system hosting IBL data on an Amazon cloud server, and automatic analysis pipelines process our behavioral data immediately after collection, with the results viewable on a web site. This web system is used by team members who are training mice to monitor progress, or to compare the performance of many animals, and also allows newcomers to the group to quickly get a sense of the behavioral data without needing to write their own analysis routines. A subset of these behavioral data are available on a [public website](#).

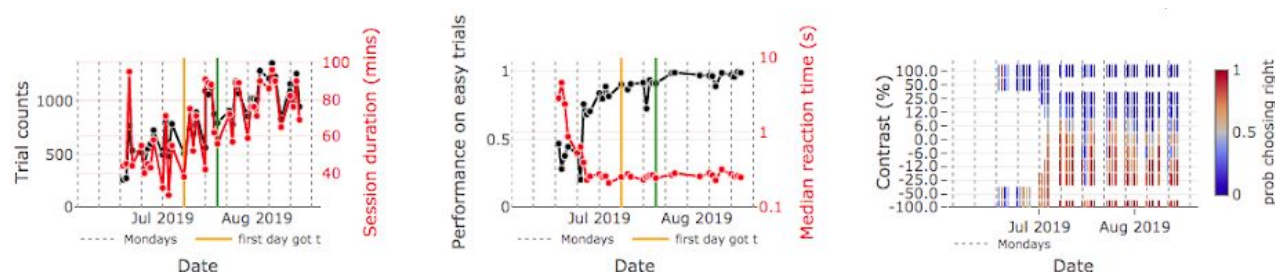


Figure 2. An analysis pipeline (created through DataJoint) automatically generates plots of behavioral data each time a subject is run. Plots generated by this pipeline are instantly available to all team members via a web browser. Data from a single mouse are shown. Left: daily trial counts and session

duration; middle: Proportion correct and reaction time on high contrast (easy) trials; Right: Accuracy for all contrast levels over a training period. Note that early training days (left) include only high coherence trials while later training days (right) include a mix of contrasts.

The final access protocol is the [NeuroData Without Borders](#) (NWB) protocol^{10,11}. NWB was created by neurophysiologists and software developers and aims to be a unified data standard (see Table 1) that is suitable for diverse neurophysiological and behavioral data. NWB is beginning to be widely adopted, and is the standard for the Allen Institute for Brain Sciences, for example. Version 2.0 of the NWB standard is based around a suite of data access functions that aims to allow users seamless access to neurophysiology data from multiple providers. The DataJoint Neuro team is working with the NWB team to provide NWB-compatible access to DataJoint databases, which will allow our data to be accessed via the NWB protocol. Importantly, using NWB for data sharing outside the collaboration does not preclude using a customized standard for data handling within the collaboration. Indeed, this is typical: for instance, the Allen institute employs the Laboratory Internal Management Systems (LIMS) and export data in NWB format when it is time for sharing, an approach individual laboratories within the collaboration have successfully piloted¹⁴.

Outlook

Although the IBL is still at an early stage, our experience thus far has shown it is possible to develop a reliable, easy-to-use data pipeline, that generates and organizes data suitable for sharing within a large collaboration, and beyond.

Data management, standardization and sharing has proven challenging within neuroscience. Datasets have grown in size, complexity and diversity in recent years, but few neuroscientists have prior experience or formal training in data management. The IBL, as a large collaboration, has been able to hire outstanding technical staff who work full time on data management, but most individual labs do not have the resources to do this. As a result, data management problems slow down projects, data sharing has not been universally adopted, and many shared datasets are lacking metadata. In this regard, neuroscience is substantially behind other scientific domains, such as genomics, astronomy, or particle physics.

We end with some thoughts on how the lessons we have learned in IBL could inform data management practices in individual labs. One possibility is for individual labs to nominate individuals to focus on data management. However, the time commitment required should not be underestimated: organizing, documenting, sharing, and supporting the data collected in even a medium-sized lab that uses modern neurophysiological methods is essentially a full-time job. An alternative possibility is for labs to rely on third-party companies that offer data management solutions (such as DataJoint Neuro, who have partnered with IBL as well as several individual labs outside the collaboration). In addition to the cost, however, this solution might not appeal to users hesitant to fully rely on a third party, fearing that this will limit flexibility and the ability to nimbly update workflows.

We hope that the open-source tools developed by IBL will help contribute to a solution allowing individual labs to curate and share data, without requiring dedicated full-time staff or external contractors. The ONE Light protocol allows data producers to “upload and forget” datasets at

the time of publication, while allowing users a standardized and searchable protocol to download data from multiple producers. Our colony management system requires minimal expertise to implement, and provides a solution to the problem of metadata collection and storage. We are keen to support the use of these tools by non-IBL labs, so we can continue to make them more useful and easy to apply by the general community. Our hope is that the IBL's data management plan can help pave the way forward to a new era of large scale data in neuroscience in which data from all labs is shared on a routine basis.

References

1. Jun, J. J. *et al.* Fully integrated silicon probes for high-density recording of neural activity. *Nature* **551**, 232–236 (2017).
2. Peron, S. P., Freeman, J., Iyer, V., Guo, C. & Svoboda, K. A Cellular Resolution Map of Barrel Cortex Activity during Tactile Behavior. *Neuron* **86**, 783–799 (2015).
3. Wiltischko, A. B. *et al.* Mapping Sub-Second Structure in Mouse Behavior. *Neuron* **88**, 1121–1135 (2015).
4. Mathis, A. *et al.* DeepLabCut: markerless pose estimation of user-defined body parts with deep learning. *Nat. Neurosci.* **21**, 1281–1289 (2018).
5. Stringer, C. *et al.* Inhibitory control of correlated intrinsic variability in cortical networks. *eLife* **5**, e19695 (2016).
6. Musall, S., Kaufman, M. T., Juavinett, A. L., Gluf, S. & Churchland, A. K. Single-trial neural dynamics are dominated by richly varied movements. *Nat. Neurosci.* **22**, 1677–1686 (2019).
7. Abbott, L. F. *et al.* An International Laboratory for Systems and Computational Neuroscience. *Neuron* **96**, 1213–1218 (2017).
8. Koch, C. & Jones, A. Big Science, Team Science, and Open Science for Neuroscience. *Neuron* **92**, 612–616 (2016).
9. Martin, C. L. & Chun, M. The BRAIN Initiative: Building, Strengthening, and Sustaining. *Neuron* **92**, 570–573 (2016).
10. Teeters, J. L. *et al.* Neurodata Without Borders: Creating a Common Data Format for Neurophysiology. *Neuron* **88**, 629–634 (2015).
11. Ruebel, O. *et al.* *NWB:N 2.0: An Accessible Data Standard for Neurophysiology*. <http://biorxiv.org/lookup/doi/10.1101/523035> (2019) doi:10.1101/523035.
12. Sorge, R. E. *et al.* Olfactory exposure to males, including men, causes stress and related analgesia in rodents. *Nat. Methods* **11**, 629–632 (2014).
13. Steinmetz, N. A., Zatka-Haas, P., Carandini, M. & Harris, K. D. Distributed coding of choice, action, and engagement across the mouse brain. *Nature* (in press).
14. Najafi, F. *et al.* *Dataset from "Farzaneh Najafi, Gamaleldin F Elsayed, Robin Cao, Eftychios Pnevmatikakis, Peter E. Latham, John P Cunningham, Anne K Churchland (bioRxiv, 2018); Excitatory and inhibitory subnetworks are equally selective during decision-making and emerge simultaneously during learning."*. <http://repository.cshl.edu/id/eprint/37693> (2019) doi:10.14224/1.37693.
15. Allen Institute for Brain Science. Allen Mouse Brain Atlas (2015) with region annotations (2017); download.alleninstitute.org/informatics-archive/current-release/mouse_ccf/annotatin

Appendix 1: Open Neurophysiology Environment details

The Open Neurophysiology Environment (ONE) user interface consists of four simple API functions that allow data consumers to search for experiments of interest and load data from them. This interface allows multiple backend instantiations, so data consumers can run the same exact code to process data from multiple data producers, even if these different producers store data in different locations. The main IBL data is provided via an ONE implementation that requires a backend SQL database, but we have also provided an “ONE light” implementation that allows data producers to share data simply by uploading files to a web server, or to figshare (a site offering free hosting for scientific data). These files can be uploaded in a variety of standard formats (numpy, csv, json, hdf5, mpeg, etc.), organized with one directory per experiment containing appropriately-named data files.

ONE clients have been implemented in Python and MATLAB, and are documented at <https://ibllib.readthedocs.io/en/develop/index.html>. Below, we describe how to use ONE in Python.

Importing ONE

Users can access data from different projects by loading and configuring appropriate implementations of the ONE library. For example, to access IBL’s main data in Python they would type

```
from oneibl.one import ONE
one = ONE()
```

whereas to access data from Steinmetz et al. via the ONE light figshare interface, they would type

```
from oneibl import onelight as one
one.set_figshare_url("https://figshare.com/articles/steinmetz/9974357")
```

In both cases, the user is returned an object named `one`, that allows the same commands to be run whichever ONE implementation is running behind the scenes.

Dataset types

The key to ONE’s standardization is the “dataset type”. When a user loads data from one of these types, they are guaranteed to be returned predictable information, organized in a predictable way, and measured in a predictable unit. For example, if a user requests the dataset `spikes.times` for a particular experiment, they are guaranteed to be returned the times of all extracellularly recorded spikes, measured in seconds relative to experiment start, and returned as a 1-dimensional column vector. The current list of ONE-standard dataset types was derived from the NWB data schema, and is shown in Appendix 2.

Datasets need not be two-dimensional numeric arrays – they can be arrays of any dimensionality, lists of strings, movies, or arrays of structures. In fact, a dataset can be anything

with the concept of a “row” – i.e. any data structure that can be addressed by an integer subscript, for example a leading dimension of a structure array, or frame number in a movie.

Dataset types are named in a specific way, which formalizes relations between data. Each dataset type has a two-part name, with the part before the period called the “object” and the part after called the “attribute”. When there are multiple dataset types with the same object (e.g. `spikes.times` and `spikes.clusters`), it is guaranteed that they will have the same number of rows, describing multiple attributes of the same object: in this case, the times and cluster assignments of each spike. If the attribute of one dataset matches the object of another, this represents a cross-reference. For example, `spikes.clusters` contains an integer cluster assignment for each spike, while `clusters.brainLocation` contains a structure array containing information on the physical location of each of these cells. The values in `spikes.clusters` can be used as indices to the rows of `clusters.brainLocation`.

Dataset types have specified measurement units. For example, locations in the brain are given in Allen CCF coordinates¹⁵, measured in mm. The units of measurement for all dataset types are specified in the documentation.

Some attributes have predefined meanings. For example it is guaranteed that any dataset type whose attribute is of the form `times` or `*_times` will represent the times of events, measured in seconds relative to experiment start. Datasets whose attribute is `intervals` or `*_intervals` are guaranteed to be two-column arrays giving start and end times of particular events in seconds relative to experiment start. Datasets whose attribute is `timestamps` contain timestamps for each sample of not-necessarily evenly sampled timeseries data, whose unit is again seconds relative to experiment start.

Not all data can be standardized. Data that are common across projects will be encoded in standard dataset types, and providers of data containing these types should return arrays of the specified dimension and measurement units. However, providers can add their own project-specific dataset types, by defining their own “namespace”. Names beginning with an underscore are guaranteed never to be standard. For example the dataset type `_ibl_trials.stimulusContrastLeft` contains information specific to the IBL project, and no other projects would be expected to use it. A list of standard dataset types will be maintained centrally, and will start small but increase over time as the community converges on good ways to standardize more information. The current list of standard and IBL-specific dataset types is in appendix 2.

Searching for data

Each experiment is characterized by a unique string known as the experiment ID (eID). The format of this string is determined by the implementation: ONE light uses directory pathnames, while the main IBL implementation uses hexadecimal UUIDs. Within an implementation, a specific eID is always guaranteed to refer to the same experiment, allowing data consumers to hard-code eID strings into their analysis software.

To search for experiments they want to analyze, a user runs the function `one.search` that returns a list of eIDs for all experiments matching the requested criteria. ONE currently allows searching for experiments by user (i.e. experimenter); by experimental subject; by date; and by

the dataset types associated with the experiment. For example, to search for all experiments on subject 'IBL_123' performed in the month of August 2019, for which spike-sorted extracellular electrophysiology and video eye tracking are available, one would type

```
eID_list = one.search(subject='IBL_123', date_range=['2019-08-02',
'2019-08-31'], dataset_types=['spikes.times', 'spikes.clusters',eye.xyPos'])
```

This command returns a list of eIDs for all experiments matching the specified criteria. Optionally, the search command also returns full metadata on each matching experiment in a structure list.

Loading data

Once a user knows the eID of an experiment they want to analyze, they can load data from this experiment in one of two ways. The first way is to load an individual dataset. For example, to load spike times, the user would type:

```
st = one.load_dataset(eID, 'spikes.times')
```

The second way is to load all datasets belonging to an object. For example the command

```
spikes = one.load_object(eID, 'spikes')
```

will return a dictionary with one entry for each dataset associated with that object (`spikes.times`, `spikes.clusters`, `spikes.depths`, `spikes.amps`, etc.).

Importantly, the user need not be concerned with the physical location or format of the raw data files - they just type these commands to load them into their analysis software. In practice, the data are cached on the user's local machine, so it need only be downloaded once for many uses; however details of the caching, as well as the underlying file formats are hidden from the user.

Listing available data

The fourth and final ONE function `one.contents(eID)` simply lists the dataset types available for a given experiment.

Appendix 2: ONE dataset types

The table below lists the dataset types currently used in the IBL's implementation of the ONE standard. The standard namespace contains data we believe can be standardized with other projects; these are largely adopted from the NWB data model. The `_ibl_` namespace contains data likely to be specific to our task or recording hardware. A live list of dataset types is maintained at <https://www.internationalbrainlab.com/resources>. The dataset types as of the time of writing are listed below.

We have provided a “ONE Light” interface, that allows any scientists to share data via ONE protocols, without needing to run a backend database. To do so, they create a data directory for each experiment, containing one data file for each of the dataset types they wish to provide. The files can be in multiple formats, but the standard is .npy files: flat binary files encoding numerical arrays, with a small header indicating the size and datatype. These files can be loaded and saved natively from numpy, and from MATLAB via [this toolbox](#). For example, to store spike times, the experiment directory would contain a file “spikes.times.npy”. Once a data sharer has created directories are creating containing these data files on their local machine, they run a command using [this library](#) to automatically upload them to figshare or to a web server. Data users can then search and download data using the standard ONE interface, though currently with simpler search capabilities than with the backend database.

ONE Dataset Type	Dimension	Description
spikes.times	[nspi]	Times of spikes (seconds, relative to experiment onset). Note this includes spikes from all probes, merged together.
spikes.clusters	[nspi]	Cluster assignments for each spike (integers counting from 0). Cluster assignment reflects the result of manual curation.
spikes.depths	[nspi]	Depth along probe of each spike (μm ; computed from waveform center of mass). 0 means deepest site, positive means above this.
spikes.amps	[nspi]	Peak amplitude of each spike (μV).
spikes.templates	[nspi]	Template ID for each spike as originally assigned by spike sorting software, prior to manual curation.
spikes.samples	[nspi]	Time of spikes, measured in units of samples in their own electrophysiology binary file.
templates.waveforms	[ntemp, nsw, nchSub]	Waveform of each template spike (stored as a sparse array, only for a subset of channels with large waveforms).
templates.waveformsChannels	[ntemp, nchSub]	Channels of the raw recording on which the template waveforms are defined.
clusters.metrics	[nc, nmetrics]	Quality control metrics for each cluster.

clusters.mlapdv	[nc, 3]	Estimated coordinates of the cell relative to bregma (mm) in Allen Common Coordinate Framework (CCF) ¹⁵ . See Appendix 4.
clusters.brainAcronyms	nc	Estimated cluster brain location acronym using Allen CCF notation. See Appendix 4.
clusters.waveforms	[nc, nsw, nchSub]	Mean unfiltered waveform of spikes in this cluster (NB: Neuropixels 1.0 data will have been hardware filtered).
clusters.waveformsChannels	[nc, nchSub]	Identities of the channels that are represented in clusters.meanWaveforms for each cluster sorted by amplitude.
clusters.depths	[nc]	Depth of mean cluster waveform on probe (μ m). 0 means deepest site, positive means above this.
clusters._phy_annotation	[nc]	Manual curation of spike cluster quality. 0 = noise, 1 = MUA, 2 = Good, 3 = Unsorted, 4-100 = manual quality score.
clusters._phy_ids	[nc]	Original cluster index assigned in Phy (counting from 0).
clusters.peakToThrough	[nc]	Trough to peak time of mean cluster waveform (ms).
clusters.amps	[nc]	Mean amplitude of each cluster (μ V).
clusters.channels	[nc, nch]	Channel which has the largest amplitude for this cluster. NB this counts all channels from all probes; to find this cluster's brain location index like so: channels.brainLocation[clusters.peakChannel[i],:]
clusters.probes	[nc, np]	Which probe this cluster came from (counting from zero).
probes.trajectory	[np, 7]	Trajectory coordinates of probe. See Appendix 4.
probes.description	[np]	Text description of probe: label (folder name), Model (3A, 3B1, 3B2), Serial Number, Original file name.
channels.probes	[nch]	Probe assignments for each channel (integers counting from 0). Can be used as direct indexing for the probes.* attributes.
channels.rawInd	[nch]	Array of indices in the raw recording file (of its home probe) that each channel corresponds to (counting from zero).
channels._phy_ids	[nch]	Array of indices in Phy that each channel corresponds to (counting from zero).
channels.brainAcronyms	[nch, 4]	Channel brain location acronym using Allen CCF notation.
channels.mlapdv	[nch, 3]	Channel location relative to bregma (mm) in Allen CCF. See Appendix 4.
channels.localCoordinates	[nch, 2]	Location of each channel relative to probe coordinate system (μ m): x (first) dimension is on the width of the shank; (y) is the depth where 0 is the deepest site, and positive above this.
channels.shank	[nch]	Array giving shank number (or other channel group) of

		each channel, used for spike sorting.
eye.timestamps	[nEyeSamples, 2]	Timestamps for pupil tracking timeseries: 2 column array giving sample number and time in seconds.
eye.raw	[nEyeSamples, nX, nY]	Raw movie data for pupil tracking.
eye.area	[nEyeSamples]	Area of pupil (pixels ²).
eye.xyPos	[nEyeSamples, 2]	Matrix with 2 columns giving x and y position of pupil (in pixels).
eye.blink	[nEyeSamples]	Boolean array saying whether eye was blinking in each frame.
licks.times	[nLicks]	Times of licks in seconds.
spontaneous.intervals	[nSpontInt, 2]	Times when no other protocol was going on for at least 30 seconds.
_ibl_wheel.position	[nWheelSamples]	Absolute linear displacement of wheel (cm).
_ibl_wheel.timestamps	[nWheelSamples, 2]	Times of position in seconds relative to session start, continuous (evenly spaced).
_ibl_wheel.velocity	[nWheelSamples]	Tangential velocity of the wheel (cm/s) where positive = CW.
_ibl_wheelMoves.intervals	[nWheelMoves, 2]	2 column array of onset and offset times of detected wheel movements in seconds relative to session start.
_ibl_wheelMoves.type	[nWheelMoves]	String array containing classified type of movement ('CW', 'CCW', 'Flinch', 'Other').
_ibl_trials.intervals	[nTrials, 2]	Start (i.e. beginning of quiescent period) and end (i.e. end of iti) times of each trial in seconds relative to session start.
_ibl_trials.included	[nTrials]	Boolean array of which trials to include in analysis, chosen at experimenter discretion, e.g. by excluding the block of incorrect trials at the end of the session when the mouse has stopped.
_ibl_trials.repNum	[nTrials]	The trial repetition number, i.e. how many trials have been repeated on this side (counting from 1).
_ibl_trials.goCue_times	[nTrials]	Time of go cue tone in seconds relative to session start.
_ibl_trials.goCueTrigger_times	[nTrials]	Time of go cue trigger command was sent in seconds relative to session start.
_ibl_trials.response_times	[nTrials]	Time in seconds relative to session start when a response was recorded (end of the closed loop state in bpod).
_ibl_trials.choice	[nTrials]	The response ID: -1 (turn CCW), +1 (turn CW), or 0 (nogo)
_ibl_trials.stimOn_times	[nTrials]	Times of stimulus onset in seconds relative to session start.

_ibl_trials.stimOnTrigger_times	[nTrials]	Times of stimulus onset trigger command in seconds relative to session start.
_ibl_trials.contrastLeft	[nTrials]	Contrast of left-side stimulus (0-1, nan if stimulus is on the other side).
_ibl_trials.contrastRight	[nTrials]	Contrast of right-side stimulus (0-1, nan if stimulus is on the other side).
_ibl_trials.feedback_times	[nTrials]	Time of feedback delivery (reward or noise) in seconds relative to session start.
_ibl_trials.feedbackType	[nTrials]	Whether feedback is positive or negative (-1 for negative, 1 for positive, 0 for no feedback).
_ibl_trials.rewardVolume	[nTrials]	Volume of reward given each trial (μ l).
_ibl_trials.itiDuration	[nTrials]	Inter-trial interval duration for each trial (seconds).
_ibl_trials.deadTime	[nTrials]	Time between state machine trial end and restart for every trial.
_ibl_trials.probabilityLeft	[nTrials]	Probability (0-1) that the stimulus will be on the left-hand side.
_ibl_passiveTrials.included	[nPassiveTrials]	Boolean suggesting which passive trials to include in analysis, chosen at experimenter discretion.
_ibl_passiveTrials.stimOn_times	[nPassiveTrials]	Times of stimulus onset in seconds relative to session start.
_ibl_passiveTrials.contrastLeft	[nPassiveTrials]	Contrast of left-side stimulus (0-1, nan if stimulus is on the other side).
_ibl_passiveTrials.contrastRight	[nPassiveTrials]	Contrast of right-side stimulus (0-1, nan if stimulus is on the other side).
_ibl_passiveValveClicks.times	[nPassiveTrials]	Times of valve opening during passive trial presentation in seconds relative to session start.
_ibl_passiveBeeps.times	[nPassiveTrials]	Times of the beep, equivalent to the go cue during task, in seconds relative to session start.
_ibl_passiveWhiteNoise.times	[nPassiveTrials]	Times of white noise bursts, equivalent to the negative feedback sound during the task, in seconds relative to session start.
_ibl_passiveNoise.intervals	[nPassiveTrials,2]	Passive noise trial start and end (i.e. end of iti) times in seconds relative to session start.
_ibl_sparseNoise.xy	[nSparseNoise,2]	2 column array giving x and y coordinates on screen of sparse noise stimulus squares.
_ibl_sparseNoise.times	[nSparseNoise]	Times of sparse noise stimulus onset in seconds relative to session start.
_ibl_extraRewards.times	[nExtraRewards]	Times of extra rewards in seconds relative to session start.
camera.dlc	[nframes, npoints x 3]	Coordinates of DeepLabCut (DLC) points (x position, y position, likelihood). Total points = 19 (fingers-8, nose-2, spout-2, tongue-2, eye-4).

camera.times	nframes	Time of each frame acquisition (leftCamera only for behavior, and right and bodyCameras for ephys rigs).
--------------	---------	--

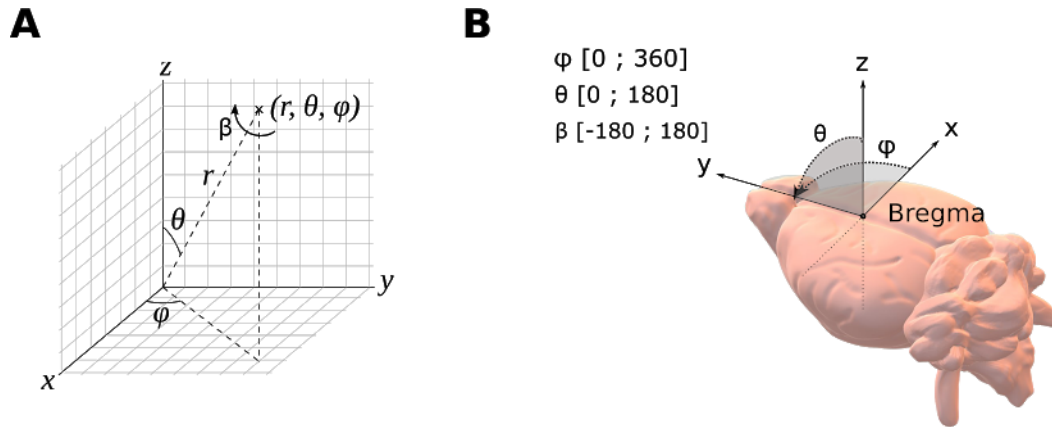
Appendix 3: Lossless compression algorithm

Our aim in developing the compression algorithm was to reduce data size, while maintaining not only the full signal available in the original data, but also maintaining the ease of random-access available with flat binary files. To do this, we took advantage of the temporal correlations in electrophysiological recordings, which show an approximate $1/f$ power spectrum.

The input to the algorithm is represented as a flat binary multiplexed file of 2-byte integers. Data are compressed independently in consecutive chunks of one second, which allows random access to any part of the recording without decompressing the whole signal. To compress a chunk, we first compute discrete time differences independently for each channel, which approximately whitens the signal. We then compress the result using the zlib lossless compression algorithm. The initial values for each chunk and compressed difference signals are then appended to a compressed binary file on a chunk-by-chunk basis, and a companion JSON file is saved storing the byte offset of every chunk. A decompression algorithm reads the JSON and binary files allowing random “slices” of the data to be retrieved on the fly without decompressing the whole file. The compression code is unit-tested with 100% coverage. In our benchmarking we could achieve a $\sim 3x$ compression ratio of our data. For our ~ 400 channel recordings, compression is $\sim 4x$ faster and decompression is $\sim 3x$ faster than real time on a Intel i9 10-core computer. To the best of our knowledge, this simple compression algorithm has not been previously described, and could be used in other applications that rely on multichannel time series of approximate $1/f$ spectrum, within and beyond neuroscience.

Appendix 4: Coordinate system within Allen CCF

Bregma is defined as Voxel ML-566, AP-540, DV-33 within the 10 μ m volume of the Allen CCF mouse Atlas¹⁵.



Coordinate System: **A)** Standard polar angles are used, alongside the standard $[x,y,z]$ base. In total, an electrode trajectory registers 7 entries (3 angles, 3 coordinates, depth) to define a probe insertion. **B)** The coordinate system is aligned to the ML (x , Right: positive), AP (y , A: positive), DV (z , D: positive) stereotaxic axis and zeroed at Bregma. The bounds for each angle are indicated.