

METHODOLOGY

Open Access



A benchmark for computational analysis of animal behavior, using animal-borne tags

Benjamin Hoffman^{1*}, Maddie Cusimano^{1†}, Vittorio Baglione², Daniela Canestrari², Damien Chevallier³, Dominic L. DeSantis⁴, Lorène Jeantet⁵, Monique A. Ladds⁶, Takuya Maekawa⁷, Vicente Mata-Silva⁸, Víctor Moreno-González², Anthony M. Pagano⁹, Eva Trapote², Outi Vainio¹⁰, Antti Vehkaoja¹¹, Ken Yoda¹², Katherine Zacarian¹ and Ari Friedlaender¹³

Abstract

Background Animal-borne sensors (“bio-loggers”) can record a suite of kinematic and environmental data, which are used to elucidate animal ecophysiology and improve conservation efforts. Machine learning techniques are used for interpreting the large amounts of data recorded by bio-loggers, but there exists no common framework for comparing the different machine learning techniques in this domain. This makes it difficult to, for example, identify patterns in what works well for machine learning-based analysis of bio-logger data. It also makes it difficult to evaluate the effectiveness of novel methods developed by the machine learning community.

Methods To address this, we present the Bio-logger Ethogram Benchmark (BEBE), a collection of datasets with behavioral annotations, as well as a modeling task and evaluation metrics. BEBE is to date the largest, most taxonomically diverse, publicly available benchmark of this type, and includes 1654 h of data collected from 149 individuals across nine taxa. Using BEBE, we compare the performance of deep and classical machine learning methods for identifying animal behaviors based on bio-logger data. As an example usage of BEBE, we test an approach based on self-supervised learning. To apply this approach to animal behavior classification, we adapt a deep neural network pre-trained with 700,000 h of data collected from human wrist-worn accelerometers.

Results We find that deep neural networks out-perform the classical machine learning methods we tested across all nine datasets in BEBE. We additionally find that the approach based on self-supervised learning out-performs the alternatives we tested, especially in settings when there is a low amount of training data available.

Conclusions In light of these results, we are able to make concrete suggestions for designing studies that rely on machine learning to infer behavior from bio-logger data. Therefore, we expect that BEBE will be useful for making similar suggestions in the future, as additional hypotheses about machine learning techniques are tested. Datasets, models, and evaluation code are made publicly available at <https://github.com/earthspecies/BEBE>, to enable community use of BEBE.

Keywords Machine learning, Bio-loggers, Animal behavior, Accelerometers, Time series, Self-supervised Learning

[†]B. Hoffman and M. Cusimano have contributed equally to this work.

*Correspondence:

Benjamin Hoffman

benjamin@earthspecies.org

Full list of author information is available at the end of the article



© The Author(s) 2024. **Open Access** This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

Background

Animal behavior is of central interest in ecology and evolution because an individual's behavior affects its reproductive opportunities and probability of survival [1]. Additionally, understanding animal behavior can be key to identifying conservation problems and planning successful management interventions [2], for example in rearing captive animals prior to reintroduction [3], designing protected areas [4], and reducing dispersal of introduced species [5].

One increasingly utilized approach for monitoring animal behavior is remote recording by animal-borne tags, or *bio-loggers* [6–8]. These tags can be composed of multiple sensors such as an accelerometer, gyroscope, altimeter, pressure, GPS receiver, microphone, and/or camera, which record time-series data on an individual's behavior and their *in situ* environment. Additionally, bio-logger datasets can include data from many-hour tag deployments on multiple individuals.

To give a behavioral interpretation to recorded bio-logger data, it is useful to construct an inventory of what types of actions an individual may perform [9]. This inventory, or *ethogram*, is then used to classify observed actions (Fig. 1A). Using an ethogram, one can quantify, for example, the proportion of time an animal spends in different behavioral states, and how these differ between groups (e.g., sex, age, populations), or change over time (e.g., seasonally), with physiological condition (e.g., healthy vs. sick) or across different environmental contexts (e.g. [10]).

For classifying the behaviors underlying bio-logger data, researchers are increasingly using supervised machine learning (ML) techniques [11]. In a typical workflow, a human annotates some of the recorded bio-logger data with the tagged individuals' behavioral states using a pre-determined ethogram, based on observations made simultaneously with data recording. These annotated data are used to train a ML model, which is then used to predict behavioral labels for the remaining un-annotated portion of the dataset. A test dataset, which is held out from the training stage, can be used to evaluate how well the trained model is able

to perform this behavior classification task. Using the predicted behavioral labels allows large datasets to be leveraged to address scientific questions, for example, through estimating activity budgets that vary by time or individual [12] or by environmental conditions [13]. In this manner, ML can help scientists to minimize manual effort required to ascribe behavioral labels to bio-logger data, or extend behavioral labels to data where manual ground-truthing is not possible.

Much research has applied ML to bio-logger data to establish its use with particular species, as well as to investigate the impact of different decisions made when using ML models for this purpose. Research questions include characterizing which methods are the most accurate, precise, sensitive, interpretable, or rapid (e.g., [14–21]), how to reduce the extent of ground-truthing necessary in species that are difficult to observe (e.g., [22–28]), and what kinds of behaviors are detectable with particular sensors and models (e.g., [15, 29–34]). However, the majority of studies focus on data from a single or a few closely-related species, making it difficult to identify patterns in how behavior classification methods are applied across multiple datasets.

A commonly used tool in ML for improving our understanding of analysis techniques is the *benchmark* (e.g. [42]). A benchmark consists of a publicly available dataset, a problem statement specifying a model's inputs and the desired outputs (a *task*), and a procedure for quantitatively evaluating a model's success on the task (using one or several *evaluation metrics*). In a common use-case for a benchmark, researchers report the performance of a proposed technique on the benchmark, helping the field to draw comparisons between different techniques and consolidate knowledge about promising directions. Developing benchmarks has been identified as an area of focus for ML applications in wildlife conservation [43] and animal behavior [44, 45].

For behavior classification from bio-loggers, a benchmark could assess model performance across a breadth of study systems in order to identify relevant patterns, such as how modeling decisions can influence classification performance. Indeed, researchers have analyzed multi-species datasets to identify best practices for other key

(See figure on next page.)

Fig. 1 **A** Examples of ethograms in BEBE. Left: gull ethogram with three behaviors. Right: a subset of the dog ethogram, with four behaviors. **B** BEBE consists of a supervised behavior classification task on nine annotated datasets, along with a set of metrics that compare model predictions with the annotations. Datasets and code are publicly available at <https://github.com/earthspecies/BEBE>. **C** Datasets in BEBE, with a photo of a representative individual and a 5-minute clip of annotated tri-axial accelerometer (TIA) data for each. Each accelerometer channel is min-max scaled for visualization. Top row: black-tailed gull (*Larus crassirostris*) [35], domestic dog (*Canis familiaris*) [29, 36], carrion crow (*Corvus corone*) [37] (see Methods). Middle row: western diamondback rattlesnake (*Crotalus atrox*) [17], humpback whale (*Megaptera novaeangliae*) [38], New Zealand fur seal (*Arctocephalus forsteri*) [39]. Bottom row: polar bear (*Ursus maritimus*) [22, 40], sea turtle (*Chelonia mydas*) [18], human (*Homo sapiens*) [41]. Gaps indicate that the behavior annotation is *Unknown*. For image attributions, see acknowledgments

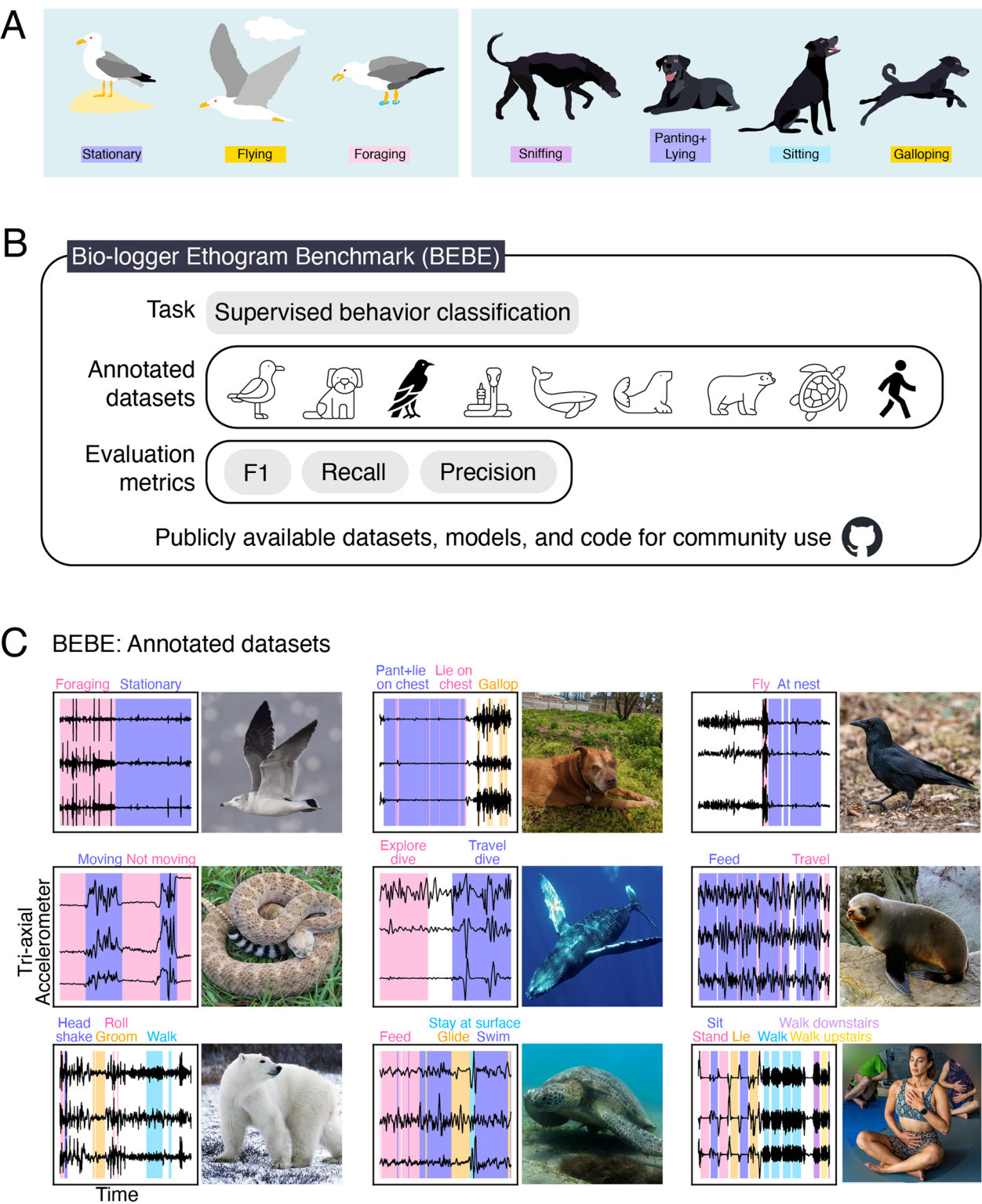


Fig. 1 (See legend on previous page.)

challenges in bio-logging, such as sensor calibration [46] and signal processing [47, 48]. Previous studies have applied one or more behavior classification techniques on multiple bio-logger datasets, with varying degrees of variability in the species and individuals included [20, 23, 49–55], but none have attempted to compile a managed, publicly available and diverse database that others could compare against as a benchmark.

In order to fill this gap, we present the Bio-logger Ethogram Benchmark (BEBE), designed to capture challenges in behavior classification from diverse bio-logger datasets. BEBE combines nine datasets collected by various research groups, each with behavioral annotations, as well as a supervised behavior classification task with corresponding evaluation metrics (Fig. 1B). These datasets are diverse, spanning multiple species, individuals, behavioral states, sampling rates, and sensor types (Fig. 1C), as well as large in size, ranging from six to over a thousand hours in duration. We focus on data collected from tri-axial accelerometers (TIA), in addition to gyroscopes, and environmental sensors. TIA are widely incorporated into bio-loggers because they are inexpensive and lightweight [6]. Additionally, the data they collect has been used to infer behavioral states on the order of seconds, in a wide variety of species [7].

As a first application of BEBE, we make and test several hypotheses about ML usage in bio-logger data (Table 1). We base these hypotheses on recent trends in ML, as well as based on their potential to influence the workflow of researchers using ML with bio-logger data. First, in many applications of ML, deep neural networks that make predictions based on raw data, out-perform classical ML methods such as random forests, which make predictions based on hand-crafted summary statistics or *features* [56]. Deep neural networks that operate on raw data, such as convolutional and recurrent neural networks, have previously been applied to behavior classification in wild non-human animals [49, 57–59], captive non-human animals [60] and in human activity

recognition [61–67]. We distinguish these from the multilayer perceptron, a type of neural network that utilizes hand-crafted features and has been used in several bio-logging studies (e.g., [14, 52, 68]). Studies differ on whether deep neural networks show performance benefits compared to classical methods (e.g., [49, 57]), and random forests remain the most commonly used ML methods for bio-logger data [69, Table 2], for which feature engineering is a key challenging step. In line with ML trends and results from [49], we predict that deep neural networks will outperform techniques using hand-chosen features (H1).

Second, we examine how data recorded for one species can be used to inform behavior predictions for a different species. Some previous works (e.g. [22, 23, 49]) have adopted a cross-species transfer learning strategy, by training a supervised model on one species and then applying this trained model to another related species. Our approach is different in that we use self-supervised learning. In self-supervised learning, a ML model (typically a deep neural network) is *pre-trained* to perform an auxiliary task on an unlabeled dataset. Importantly, training the model to perform this auxiliary task does not require any human-generated annotations of the data. In this way, self-supervised pre-training can make use of a large amount of un-annotated data that is easy to obtain (e.g., a species where a large dataset exists). Later, the pre-trained model can be trained (or *fine-tuned*) to perform the task of interest (such as behavior classification in a different species), using a small amount of annotated data. By learning to perform the auxiliary task the model learns a set of features, which often provide a good set of initial model parameters for performing the task introduced in the fine-tuning step. Inspired by the recent success of self-supervised learning in other domains of ML (e.g. language [70] and computer vision [71]), we predict that a deep neural network pre-trained on a large amount of human accelerometer data [61] will outperform alternative methods, after fine-tuning (H2). Moreover, the

Table 1 Hypotheses tested in this work

Hypothesis	Hypothesis confirmed?
(H1) Deep neural network-based approaches will outperform classical approaches based on hand-chosen summary statistics.	Yes, for the approaches we tested
(H2) Self-supervised pre-training using human accelerometer data will improve classification performance.	Partly
(H3) Self-supervised pre-training using human accelerometer data will improve classification performance when the amount of training data is reduced by a factor of four by removing individuals.	Yes
(H4) In terms of a single model's predictive performance, there is minimal improvement in some behavior classes when increasing the amount of training data by four times by adding individuals.	Yes

BEBE provides a means to identify patterns in behavior classification methods applied across multiple species and sensor types

Table 2 Summary of datasets in BEBE

	Human activity recognition (HAR) [41, 74]	Rattlesnake* [17]	Polar bear [22, 40]	Dog [29, 36]	Whale* [38]	Sea turtle [18]	Seals [21, 39]	Gull* [35]	Crow* (see methods)
Species	Human	Western diamond-back rattlesnake	Polar bear	Domestic dog	Humpback whale	Green turtle	Otarid spp.	Black-tailed gull	Carion crow
Typical body mass (kg)	60–70	1–3	150–300 (F), 300–800 (M)	5–70	40000	68–190	Fur seals: 40 (F), 140 (M) Sea lion: 100–300	0.4–0.6	0.4–0.6
Location	Genoa, Italy	Texas, USA	Arctic Ocean	Helsinki, Finland	Wilhelmina Bay, Western Antarctica Peninsula	Grande Anse d’Arjet, Martinique, France	Marine facilities on west coast of Australia	Kabushima Island, Japan	León, Spain
Tag position	Waist	Body	Neck	Back and neck	Dorsal surface or flank	Carapace	Back	Back or abdomen	Tail
# indiv.	30	13	5	45	8	14	12	11	11
Behavior classes	Lie down Stand Sit Walk downstairs walk upstairs walk	Moving not moving	Walk swim run roll rest pounce head shake groom eat dig	Walk trot stand sniff sit shake pant+stand dive	Feet dive Exploratory dive Travel Rest dive	Swim stay at surface scratch rest glide feed breathe	Feed groom rest travel	Fly forage stationary	In nest fly
Sample rate (Hz)	50	1	16	100	5	20	25	25	50
Data channels	TIA Gyroscope	TIA	TIA conductivity	2x TIA 2x gyroscope	TIA depth speed	TIA gyroscope depth	TIA depth	TIA	TIA
Hardware	Galaxy S II Samsung	AXY-3, AXY-4 (Technosmart)	Video collar: Exeye-Logger: TDR10-X-340D (Wildlife computers)	ActiGraph GT9X Link (ActiGraph LLC)	DTAG [75]	CATS (Customized animal tracking solutions)	CEFAS G6a+(CEFAS Technology Ltd)	TDK MPU-9250 (InvenSense)	miniDTAG [37], TIA: KX022-1020 (Kionix)
Attach. Method	Belt	Implant	Collar	Back: harness neck: tape	Suction	Suction	Fur seal: tape Sea lion: harness	Tape, harness	Elastic
Duration (hrs)	6.2	30.9	1108.4	29.5	184.6	77.1	14.0	85.7	114.6
Duration annotated (hrs)	4.2	30.9	196.1	16.9	114.1	67.8	11.6	85.0	3.4
Unknown behaviors (%)	33.4	0.0	82.3	42.7	38.2	12.1	17.0	0.8	95.7
Mean annot. dur. (sec)	17.5	721.8	127.2	15.5	119.8	47.2	24.8	2823.7	14.1
Annot. method	Direct obs., off tag video	Direct obs., off tag video	On tag video	Direct obs., Off tag video	Motion, on tag audio	On tag video	Direct obs., Off tag video	On tag video	On tag audio
License	Custom	Creative commons	Public domain	Creative commons	Creative commons	Public domain	Creative commons	Creative commons	Creative commons

Out of nine datasets, one comes from humans, three come from other terrestrial species, three come from aquatic species, and two come from flying species. We provide a summary of hardware, data collection, ethogram definition, and ground-truthing in Section 2.1.2; see the original papers for more details, including details on synchronization, calibration, and annotation validation. Datasets marked with an asterisk are publicly available for the first time in BEBE

Duration, the total duration of the dataset. Duration annotated: the duration of annotated data; Mean annot. dur., the average duration an individual spends in a behavioral state (described in Supplemental Information); Annot. method, the type of data that was used to make the behavioral annotations; Attach. method, manner of attaching the bio-logger to the organism

self-supervised pre-training step can reduce the amount of annotated data required to meet a given level of performance [72]. Therefore, we predict that this trend will hold when we reduce the amount of training data by a factor of four (reduced data setting; H3).

Finally, we investigate how the performance of our best models varies by behavior class and how this per-behavior performance scales with the amount of training data. In recent years, performance on some human behavior classification benchmarks has shown little improvement in spite of methodological advancements [73]. This suggests that the sensor data may not contain sufficient information to discriminate activities of interest. As found for human activity recognition and behavior classification in other animals [15, 29, 31–34], we expect that there will be a large degree of variation in per-behavior performance. Here, we test one possibility for improving classification performance: increasing the amount of training data. If some behaviours are not well discriminated by sensor data, we predict that increasing the amount of training data will show only minimal improvement (H4).

While we focus on a specific classification task in this study, all datasets, models, and evaluation code presented in BEBE are available at <https://github.com/earthspecies/BEBE> for general community use. Researchers may use the standardized task to test classification methods, or adapt BEBE datasets for their own research questions (see Discussion for examples). Given that one aim of BEBE is to improve our understanding of classification methods in bio-logger data, we are also seeking contributions to create an expanded benchmark with improved taxonomic coverage, a broader range of sensor types, additional standardization, and a wider variety of modeling tasks. Details about how to contribute in this way can also be found at our GitHub.

In summary, the main contributions of this study include:

1. Publicly available multi-species bio-logger benchmark dataset, centered on tri-axial accelerometers
2. Standardized evaluation framework for supervised behavior classification, with accompanying code and model examples
3. Demonstration of benchmark usage to investigate patterns in ML behavior classification performance (Table 1), including:
 - A comparison of deep learning and classical techniques on non-human bio-logger data
 - Successful cross-species application of a self-supervised neural network trained on human bio-logger data (based on [61])

Methods

Benchmark Datasets

We brought together nine animal motion datasets into a benchmark collection called the Bio-logger Ethogram Benchmark (BEBE) (Table 2). BEBE introduces a previously unpublished dataset (Crow); otherwise, these data were all collected in previous studies. Of the datasets included in BEBE, four are publicly available for the first time (Whale, Crow, Rattlesnake, Gull) and five were already publicly available (HAR, Polar bear, Sea turtle, Seals, Dog). We summarize datasets' hardware, data collection, ethogram definition, and ground-truthing in Sect. 2.1.2. For full details, including details on synchronization, calibration, and annotation validation, we refer readers to the original papers.

In each dataset, data were recorded by bio-loggers attached to several different individuals of the given species. Each dataset contains one species, except for the Seals dataset which contains four *Otariid* species. These bio-loggers collected kinematic and environmental time series data, such as acceleration, angular velocity, pressure, and conductivity (Fig. 2). While each dataset in BEBE includes acceleration data, different hardware configurations were used across studies. As a result, each dataset comes with its own particular set of data channels, and with its own sampling rate. We used calibrated data as provided by the original dataset authors.

In addition to the time series bio-logger data, each dataset in BEBE comes with human-generated behavioral annotations (Fig. 2, colored bars). Seven datasets were ground-truthed using either on or off tag video; the other two were ground-truthed using audio. In each dataset, each sampled time step is annotated with the current behavioral state of the tagged individual, which can be one of several discrete behavioral classes. At some time steps, it was not possible to observe the individual, or it was not possible to classify the individual's behavior using the predefined behavioral classes. In these cases, this time step is annotated as *Unknown*. These *Unknown* behavioral annotations are disregarded during model training and evaluation.

There are multiple time scales of behavior represented across the nine ethograms in BEBE, with some datasets including brief activities (e.g. shaking), and some including longer duration activities (e.g. foraging). In Table 2 we report the mean duration (in seconds) of an annotation in each dataset, as a rough estimate of the mean duration an individual spends in a given behavioral state.

For previously published datasets, the intentions were to validate an ethogram for use in free-ranging individuals (Polar Bear, Seals), use the ethogram to understand activity patterns (Whale, Sea Turtle, Rattlesnake), develop on-device algorithms to detect a specific rare

behavior of interest (Gull), or provide a publicly available dataset (HAR, Dog). The data used for annotation also varied from on-sensor (Polar Bear, Whale, Sea Turtle, Gull, Crow) to off-sensor (HAR, Seals, Rattlesnake, Dog). Given the range of purposes and data collection methods, the ethograms vary in how much of the animal's time is accounted for and how fine-grained the behavior categories are.

In addition to Fig. 2, we provide additional data visualizations in the Supplemental Information. Examples of each behavior class with the full set of channels reveal varying degrees of stereotypy in the behavioral classes (Supplemental Figs. S1–S10). For example, in the Sea Turtle dataset, *Stay at surface* appears more stereotyped than *Feed*. Summary statistics across different individuals for each behavior class suggest the presence of discriminative features for some datasets, as well as differences between individuals (Supplemental Figs. S11–S19). For example, in the Rattlesnake dataset, *Move* shows higher variance in raw accelerometer values compared to *Not Moving*, although to different degrees in different individuals.

Dataset collection

Datasets had to meet the following criteria to be included in BEBE:

1. Include animal motion data recorded by tri-axial accelerometer at ≥ 1 Hz;
2. Include annotations of animal behavioral states;
3. Comprise data recorded from tags attached to at least five individuals in order to reflect variation in sensor placement and individual motion patterns;
4. Contain over 100,000 sampled time steps with behavioral annotations;
5. Contribute to a diversity of taxa, as well as a balance among the categories of terrestrial, aquatic, and aerial species;
6. Have previously appeared in a peer-reviewed publication (with the exception of the Crow dataset, which is previously unpublished and described in more detail below);
7. Be licensed for modification and redistribution; or come with permission from dataset authors for modification and public distribution.

Four datasets were not previously publicly available and were collected by coauthors (Whale: A. Friedlaender; Crow: D. Canestrari, V. Baglione, V. Moreno-González, E. Trapote; Gull: T. Maekawa, K. Yoda; Rattlesnake: D. DeSantis, V. Mata-Silva). For these datasets, coauthors provided permission to publicly distribute the data. Through an informal literature search, we found five

publicly available datasets (HAR, Polar Bear, Dog, Sea Turtle, Seals). Of these, four were collected by coauthors (Polar Bear: A. Pagano; Dog: O. Vainio, A. Vehkaoja; Sea Turtle: L. Jeantet, D. Chevallier; Seals: M. Ladds). Finally, we assessed datasets from papers covered by a recent systematic literature review of automatic behavioral classification from bio-loggers [69, Page 12]. The supplemental material of [69] provides a table with the results of their systematic review, containing metadata on whether a paper used supervised learning, species, number of individuals, and number of timepoints. We looked exclusively at the supervised learning papers because these would require annotated datasets (criterion 2). Assessing criteria 1, 3, and 4 above resulted in twelve potential datasets out of 214. Of the twelve, two were already included in BEBE (Rattlesnake, Sea Turtle), nine studied terrestrial animals, a category which was already well-represented in BEBE, and one did not provide annotations. Therefore, no new datasets were added based on the results of the systematic literature review by [69].

Dataset summaries

In the following, we summarize the study design and data collection protocols for each of the datasets in BEBE. See also Table 2.

Human Activity Recognition (HAR) The study [41] was designed to provide a publicly available dataset of human (*Homo sapiens*) activities recorded by smartphone tri-axial accelerometers and gyroscopes. Thirty human subjects were instructed to perform a sequence of activities (Walking, Standing, Sitting, Lying Down, Walking Upstairs, and Walking Downstairs) while wearing a waist-mounted Samsung Galaxy S II smartphone. Behaviors were annotated based on video footage, and no information is provided about synchronization between video and motion sensor data. The ethogram used covers all behaviors performed by individuals, except transition periods between activities (e.g. moving from sitting to standing) which are treated as Unknown.

Rattlesnake The study [17] sought to quantify and evaluate variation in long-term activity patterns in free-ranging western diamondback rattlesnakes (*Crotalus atrox*) in the Indio Mountains Research Station, located in Texas. Individuals were implanted with Technosmart AXY-3 or Technosmart AXY-4 tri-axial accelerometers. Their behavior was directly observed, and recorded with a hand-held video camera. All recorded time steps were assigned to one of two behavior categories, Moving and Not Moving, in order to accommodate low-frequency recording and maximizing recording duration. Annotations were made based on field notes and recorded video.

Polar Bear The study [22, 40] was designed to validate the usage of tri-axial accelerometers and conductivity

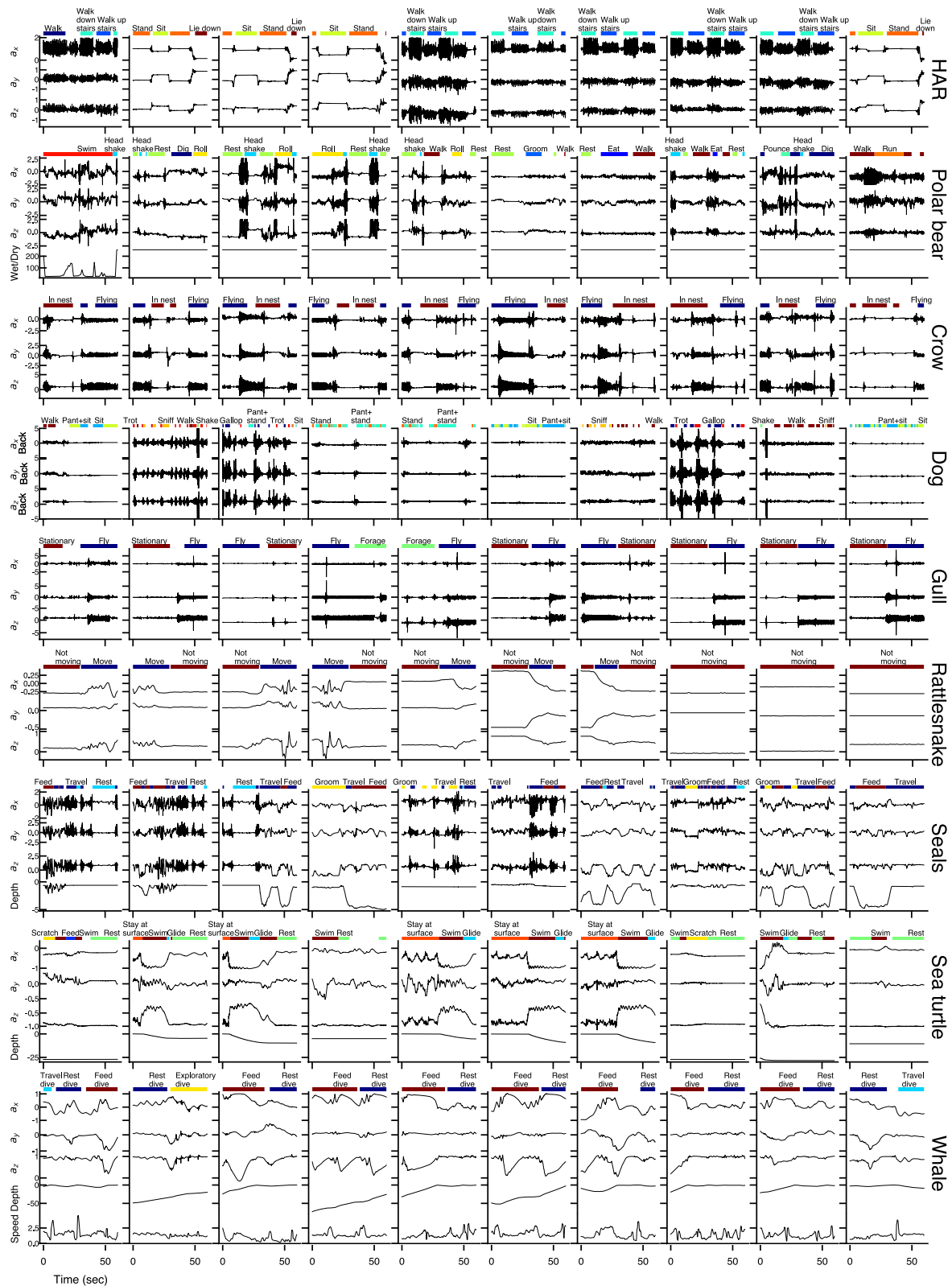


Fig. 2 Example data from BEBE. Each row displays ten 1-minute clips from one dataset, showing behavior labels, three tri-axial accelerometer channels (g), as well as speed (m/s), saltwater conductivity (wet/dry), and/or depth (m) if available. Examples were chosen to focus on transitions between behaviors. Acceleration traces for behavior classes range from highly stereotyped (e.g., *Sit* in HAR) to highly variable (e.g., *Feed* in Seals). For examples of each behavior in each dataset, with the full set of dataset channels, see Supplemental Figs. S1–S10

sensors, for the purpose of constructing daily activity budgets of polar bears (*Ursus maritimus*) on sea ice in the Beaufort Sea (Arctic Ocean). Bears were captured and equipped with Exeye video collars that also contained Wildlife Computers TDR10-X-340D motion loggers. Collars were retrieved after they fell off, or after the bear was re-captured. Motion data was annotated based on synchronized video footage. The ethogram was designed to cover all common behaviors observed in this footage, but excluded behaviors that were rare (such as fighting, breeding, drinking), extremely brief, or nondescript. Excluded behaviors were marked as Unknown. Additionally, the video camera was set to a 90-second duty cycle. Behaviors were marked as Unknown during time periods where the video camera was off.

Dog The study [29, 36] was intended to provide a dataset for developing methods that could be used to classify domestic dog (*Canis familiaris*) behaviors. Dogs were equipped with two ActiGraph GT9X Link loggers. One was placed on the neck using a collar, and the other was placed on the back using a harness. In an indoor arena, dogs were guided by their owners through a series of activities: sitting, standing, lying down, trotting, walking, playing, and treat-searching. Behaviors were annotated based on synchronized video footage. The ethogram was designed to reflect all the behaviors commonly performed during these activities. Ambiguous behaviors were recorded as Unknown.

Whale The study [38] characterized daily activity budgets of humpback whales (*Megaptera novaeangliae*) in Wilhelmina Bay, Antarctica, late in the feeding season. DTAG devices [75] were attached via suction cups to whales' dorsal surface or flank. They were programmed to release suction after 24 h, and were retrieved after release. The ethogram was designed to include common diving behavior, as well as resting. To identify different behaviors, whales' feeding lunges were first detected using an algorithm, based on recorded acoustic flow noise, as well as changes in the accelerometer signal. Then, dives were classified based on maximum dive depth, duration, and the presence and number of feeding events. These annotations were reviewed by two of the original study authors.

Sea Turtle The study [18] aimed to develop a machine learning method to compute activity budgets for green turtles from accelerometer and gyroscope data. CATS devices (Customized Animal Tracking Solutions, Germany) were attached using suction cups to the carapaces of free-ranging immature green turtles (*Chelonia mydas*) in Martinique, France. Behaviors were annotated based on synchronized on-device video footage. To design the ethogram, forty-six behaviors were initially identified in video footage. From these, seven frequent behavior

categories were identified, and the remaining behaviors (such as regurgitation, pursuit of other turtle) grouped into an 'other' category. The 'other' category was considered Unknown in the present study.

Seals The studies [10, 21] aimed to validate machine learning methods for behavior classification in otariids on captive seals, intended for eventual application in wild seals. Tags including CEFAS G6a+ accelerometers were attached to the backs of captive fur seals (*Arctocephalus forsteri* and *Arctocephalus tropicalis*) and sea lions (*Neophoca cinerea*). Behaviors were filmed in a swimming pool, by two or three underwater cameras (GoPro Hero 3 - Black edition) and one handheld camera above water (Sony HDRSR11E). In observation sessions, individuals either received a food item or were requested to perform behaviors learned through operant conditioning. The requested behaviors were chosen to reflect behaviors performed by wild seals. Twenty-six behaviors were identified in the video footage. Based on prior knowledge of wild seals, these initial behaviors were grouped into four categories or 'other' (such as direct feeding by trainer, seal out of sight of camera). The 'other' category was considered Unknown in the present study.

Gull The study [35] aimed to develop on-device machine learning algorithms to detect rare, ecologically important behaviors in accelerometer data (e.g., foraging) in order to control resource-intensive sensors like video cameras. The bio-loggers included a tri-axial accelerometer (TDK MPU-9250; InvenSense), integrated video camera, as well as other low-cost sensors which were not used in this study. Using waterproof tape and teflon harnesses, these were attached to either the back or abdomen of free-ranging black-tailed gulls (*Larus crassirostris*) in a colony on Kabushima Island near Hachinohe City, Japan. The main behavior of interest in this study was foraging, which included a variety of behaviors such as surface-dipping and plunging. Two other common non-foraging behaviors (flying and stationary) were also labeled. Behaviors were labeled based on video footage collected by the on-device video camera.

Crow The Crow dataset is presented for the first time here, and we provide complete details in the following section.

Crow dataset details

The data logger, called miniDTAG, was adapted from a 2.6 g bat tag integrating microphone, tri-axial accelerometer and tri-axial magnetometer [37] with changes that enable long duration recordings on medium-sized birds. The triaxial accelerometer (Kionix KX022-1020 configured for ± 8 g full scale, 16-bit resolution) was sampled at 1000 Hz and decimated to a sampling rate of 200 Hz before saving to a 32 GB flash memory. The 1.2 Ah

lithium primary battery (Saft LS14250) allowed continuous recording for about 6 days both in lab and field settings. Each miniDTAG was packaged with a micro radio transmitter (Biotrack Picopip Ag376) and attached to the two central tail feathers of carrion crows (*Corvus corone*) with a piece of the stem of a colored balloon following the procedure described in [76] (axes: x (backward-forward), y (lateral), z (down-up)). The thin rubber balloon material progressively deteriorated and finally broke, letting the miniDTAG fall to the ground, where it was radio-tracked using a Sika Biotrack receiver.

Accelerometer data were calibrated using Matlab tools from www.animaltags.org following standard procedures [75, 77]. The sensor channel was decimated by a factor of 4 before calibration, resulting in a sampling rate of 50 Hz. We normalized the tri-axial acceleration channels so the average magnitude of the acceleration vector was equal to 1.

For the present study, we tagged 11 individuals from 7 different territories near León, Spain, in a population that breeds cooperatively. Here crows live in stable kin group, in which a dominant breeding pair is assisted by subordinate helpers in raising the young [78]. Of these individuals, three were breeding males, two were helper males, five were breeding females, and one was a helper female, and all were attending an active nest. Data were collected in spring 2019, when all the birds were raising their nestlings. The miniDTAG plus battery (12.5 g) accounted on average ($\pm$ SE) for the $2.66 \pm 0.09\%$ of the crow body mass (range 2.29–3.15%). None of the crows abandoned the territory or deserted the nest after being tagged. From the recordings of these individuals, we selected 20 contiguous segments for annotation (average segment duration: 5.73 h), favoring segments where begging vocalizations and wing beats could be identified at multiple times during the recording (see Annotations below).

We divided the recorded data into five-second long non-overlapping segments (*clips*). Each accelerometer clip came with synchronized audio, which we used to assign behavioral annotations. If there were sounds of wingbeats for the entire duration of a clip, we annotated all sampled time steps in that clip as *Flying*. Additionally, if there were wingbeats followed by wind noise (interpreted as soaring), we also annotated all sampled time steps in that clip as *Flying*. Similarly, if a clip included sounds of chick begging calls, and no sounds of wingbeats, we annotated all sampled time steps in that clip as *In Nest*. Therefore, the label *In Nest* likely encompasses several behavioral states, such as resting, brooding, incubating, feeding chicks, and preening, which may occur at or near the nest. Crucially, these states do not include flying, and so there is no ambiguity between the two

labels. Clips that did not fit either of the criteria for being labeled *Flying* or *In Nest* were labeled as *Unknown*.

The two behaviours we chose for our ethogram (*Flying* and *In Nest*) are highly relevant for ethology research: Individual chick provisioning effort, measured as frequency of nest visits, is one of the key variables in the study of parental care behaviour in this species. It may be possible to infer nest visits based on alternating periods of *Flying* and *In Nest*. Counting nest visits typically requires either many hours of direct observations, which is time-consuming and difficult to carry out without interfering with the animals, or using video cameras at the nest, which requires costly equipment, high effort to install, and daily visits to change the battery.

Data pre-processing

For full implementation details, we refer the reader to the dataset preprocessing source code.¹ For all datasets, we used calibrated data and annotations provided by the original dataset authors; with the exception of the Crow dataset (described above and below), we refer the reader to the original publications for details. For two datasets (Sea Turtle, Gull), the average magnitude of the acceleration vector varied by more than 10% between tag deployments. To control for these differences, we normalized the tri-axial acceleration channels so that the average magnitude of the acceleration vector was equal to 1. For the Gull dataset, there were two possible tag placement positions (back or abdomen). To reduce data heterogeneity due to differences in tag placement, we rotated the calibrated data from some deployments by 180 degrees, around the axis parallel to the tagged individual's body. After performing this step, all deployments had, on average, positive acceleration in the vertical axis. While this step reduced heterogeneity between deployments, is unlikely to remove all differences between the data recorded by these different tag placements. We do not perform any additional special pre-processing steps on the datasets in BEBE, and we left each dataset in its original measurement units.

Annotations

In all datasets in BEBE, annotations indicated time intervals when the behavior class occurred (even when this time interval is only a few seconds long), rather than the occurrence of discrete behavioral events. For example, multiple discrete feeding events could occur within a time interval labeled as 'foraging'. The modeling task (per-time step classification) and evaluation procedure (per-time step classification precision and recall) are designed with this in mind. These annotations would be

¹ <https://github.com/earthspecies/BEBE-datasets/>.

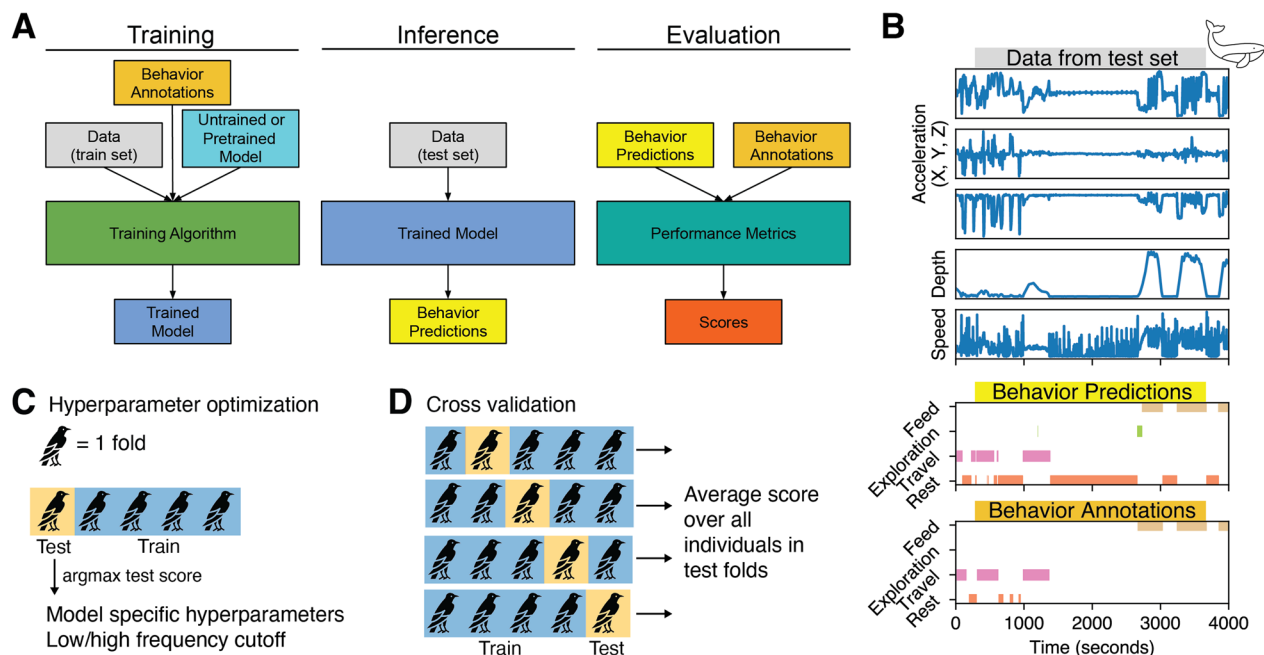


Fig. 3 **A** Summary of training and evaluation. Our process of data analysis follows the standard three steps of creating and evaluating machine learning models. In the first step (Training), the model learns from the train set of one dataset, including behavioral annotations. In the second step (Inference), the model makes predictions about the behavioral annotation for the test set data, which comprises data from a set of individuals distinct to those in the train set. In the third step (Evaluation), the model's predictions are evaluated based on their agreement with known behavioral annotations. **B** Example data from the Whale dataset [38], and predictions made by a CRNN model. The trained model is fed raw time series data, which it uses to make behavior predictions. These predictions are compared with annotations to arrive at performance scores. In this case, the model predicts the annotations well. Gaps in the behavior annotations indicate the behavior is *Unknown* at those samples; those samples are ignored in the evaluation metrics. **C** During hyperparameter optimization, we train a set of models with various hyperparameters and low/high frequency cutoffs. We obtain the model hyperparameters and low/high frequency cutoff from the model that maximizes the F1 score on the first test fold. **D** During cross-validation, we compute the test scores for the other four folds. The final score is averaged across all individuals in the test folds. The first test fold, used for hyperparameter optimization, is not used for testing

those required to describe the amount of time an animal spends performing different activities during a day.

With the exception of one dataset (Crow), the annotations in BEBE are derived from annotations made in the original studies. As a result, datasets in BEBE are annotated in a variety of ways (Table 2, row *Annotation method*). Datasets also vary in the specificity of their behavior classes: a behavior class may include several related behaviors and datasets vary in how much behaviors are grouped or split. For example, in the Dog dataset, there are fine distinctions between different behavioral classes (e.g. *Sit* vs. *Pant+sit*), relative to the Rattlesnake dataset, in which a behavior is only summarized as *Moving* or *Not Moving*.

For the remaining eight datasets, we used annotations as provided by the original dataset authors. For behaviors with few annotated timepoints in the original dataset, we treat these behaviors as *Unknown* (see dataset pre-processing source code for details). For all datasets, we used the time alignment between annotations and tag data that were produced by the original dataset authors.

Task and model evaluation

We provide a standard method for measuring different models' ability to classify behavior from bio-logger data (Fig. 3) that consists of a formal task, as well as a set of evaluation metrics. It reflects the following workflow. First, the researcher has defined the ethogram categories of interest and annotated the dataset. Then, the annotated dataset is split between the train set, used to train ML model, and the test dataset, which is used to evaluate the performance of the trained model.

Trained models are evaluated on their ability to predict behavior annotations. For each individual, we measure classification precision, recall, and F1 scores averaged across all sampled time steps from that individual and averaged across all behavioral classes. We disregard the time steps for which the annotation is *Unknown*. The entire pipeline, including training, inference, and evaluation, is repeated for each dataset in BEBE.

We split each dataset into five groups, or *folds*, so that no individual appears in more than one fold [28]. For evaluation, we use a cross validation procedure. During

cross validation, we train a model on the individuals from four folds, and test it on the individuals from the remaining fold. For all datasets, Figure S20 shows the proportion of behavior classes for each fold. This data partition reflects a common use case where researchers train a model from one set of individuals and then apply it to a separate unseen set of individuals (e.g., when behavior labels cannot be manually assigned for the latter).

Behavior classification task

Each dataset consists of a collection of multivariate discrete time series, where each time series $\{\mathbf{x}_t\}_{t \in \{1, 2, \dots, T\}}$ consists of samples $\mathbf{x}_t \in \mathbb{R}^D$. Here D is the number of data channels and T is the number of sampled time steps. Note that the number T may vary between different time series contained in a single dataset. Each time series is sampled from one bio-logger deployment attached to one individual and is sampled continuously at a fixed dataset-specific sampling rate (Table 2).

Each time series in a dataset also comes with a sequence of annotations $\{l_t\}_{t \in \{1, 2, \dots, T\}}$, where each $l_t \in \{Unknown, c_1, c_2, \dots, c_C\}$ encodes either the behavioral class c_j of the animal at time t , or the fact that the behavioral class is *Unknown*. Here C denotes the number of known behavioral classes in the dataset. The behavioral classes c_j vary between datasets in BEBE, and could be e.g. $c_j = \text{Foraging}$, $c_j = \text{Sniffing}$, or $c_j = \text{Flying}$.

The behavior classification task is to predict the behavioral annotation l_t of each sampled time step $\mathbf{x}_t$ (Fig. 3B). During training, models are given access to the behavioral annotations in the train set. We refer the reader to [79] for a review of studies with a similar task description. While behavior classification can also be formulated as a continuous time problem [80], we focus on a discrete time problem formulation in order to match the majority of prior studies.

Dataset Splits

A key part of a benchmark dataset is how it partitions the data used for model training (the *train set*) from the data used for model evaluation (the *test set*). This evaluation provides an estimate of how well a model performance generalizes outside of its train set. Therefore, the specific partition chosen determines what domains the ML model should generalize over.

In BEBE, we split each dataset into five groups (*folds*), which are used in a cross validation procedure. During cross validation, each time the model is trained, the train set consists of the data from four of these five folds, and the test set consists of the data from the remaining fold. For each dataset in BEBE, we divided the data so that no individual appears in more than one fold, and so that each fold has the same number of individuals represented

(± 1 individual). Therefore, during testing, a model's performance reflects its ability to generalize to new individuals, where effects such as tag placement [46] may influence model predictions.

Figure S20 displays the distribution of annotations across folds for all datasets in BEBE. Most datasets in BEBE have some behaviors with high representation (up to 92.4 percent of known behaviors, Rattlesnakes *Not Moving*), and some behaviors with very low representation (as little as 0.1 percent of known behaviors, Polar Bears *Pounce*).

Evaluation Metrics

Trained models are evaluated on their ability to predict the behavioral annotations of the test set. For each individual in the test set, we measure macro-averaged precision, recall and F1 scores of model predictions. By macro-averaging, performance on each behavioral class is weighted equally in the final metrics, regardless of their relative proportions in the test set. Finally, we average these scores across all individuals in the test set. In measuring these scores, we disregard the model's predictions for those time steps $\mathbf{x}_t$ for which $l_t = \text{Unknown}$. More precisely, for each individual in the test set we measure:

$$\text{Prec} = \frac{1}{C} \sum_{j=1}^C \text{Prec}_j, \quad \text{Rec} = \frac{1}{C} \sum_{j=1}^C \text{Rec}_j, \quad \text{F1} = \frac{1}{C} \sum_{j=1}^C \text{F1}_j, \quad (1)$$

where for each behavioral class index $j \in \{1, \dots, C\}$,

$$\text{Prec}_j = \frac{\text{TP}_j}{\text{TP}_j + \text{FP}_j}, \quad \text{Rec}_j = \frac{\text{TP}_j}{\text{TP}_j + \text{FN}_j}, \quad \text{F1}_j = 2 \cdot \frac{\text{Prec}_j \cdot \text{Rec}_j}{\text{Prec}_j + \text{Rec}_j}.$$

Here, TP_j , FP_j , and FN_j denote, respectively, the number of sampled time steps correctly predicted to be of class c_j (true positives), the number incorrectly predicted to be of class c_j (false positives), and the number incorrectly predicted to be not of class c_j (false negatives). Precision, recall, and F1 range between 0 and 1, with 1 reflecting optimal performance. After computing these scores for each individual, we calculate the average taken across all individuals in the test set. In addition to precision, recall, and F1 score, we compute confusion matrices for model predictions (see examples in Figures S26 and S27, with full set available at <https://zenodo.org/records/7947104>).

Hyperparameter Tuning and Cross Validation

All models we tested require the user to choose some parameters (known as *hyperparameters*) before training. To select hyperparameters for a given type of ML model and dataset, we performed an initial grid search across a range of possible values, using the first fold of the dataset as the test set and the remaining four folds of the dataset as the train set. We saved the hyperparameters which led

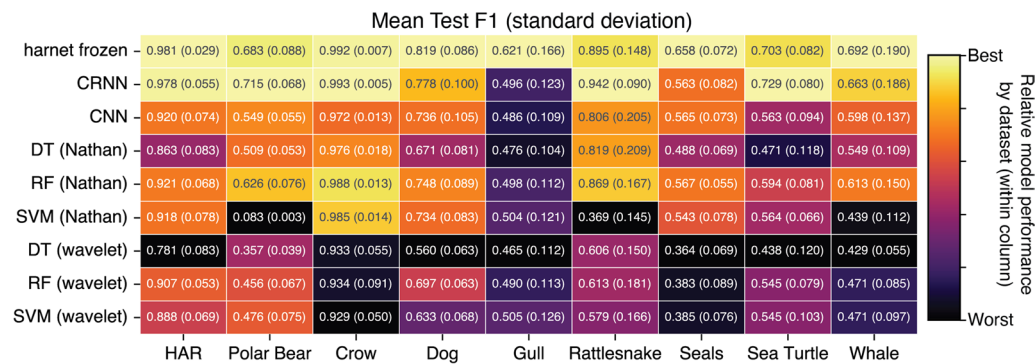


Fig. 4 F1 scores on the test set for supervised task. Here and elsewhere, the table is color-coded such that within a dataset (column), the brightest color indicates the best performing model for that metric, and the darkest color indicates the worst performing model. Numbers indicate the average score across individuals in the test folds, with the standard deviation in parentheses. The F1 score is macro-averaged across classes. Out of nine datasets, harnet does best on five datasets for F1, as indicated by the bright yellow entries in its row. CRNN does best on the other four datasets. For precision and recall results, see Figure S21

to the highest F1 score, averaged across individuals in the test set. The hyperparameter values included in the grid search are specified below, and the hyperparameter values that were saved for subsequent analyses are available at <https://github.com/earthspecies/BEBE>.

While it is common in the field of ML to use a single fixed train/test split of a dataset, we chose to use cross validation in order to capture the variation in motion and behavior between as many individuals as possible. After the initial hyperparameter grid search, we used the saved hyperparameters to train and test a model using each of the remaining four train/test splits of the dataset (which were not used for hyperparameter tuning). The final scores (precision, recall, and F1) we report are averaged across individuals taken from these four train/test splits.

All of the models we trained involve some randomness in the training process, which can introduce variance into model performance [81]. In addition, model performance varies between different individuals. Understanding the magnitude of this variation may be important when applying these techniques in new contexts.

To quantify variation in model test performance, for each model type we compute the standard deviation of each performance metric, taken across all individuals represented in the four test folds of the dataset that were not used for hyperparameter tuning. These values, given in parentheses in Figs. 4 and 5, therefore reflect variation in these scores due to differences in individual motion, as well as due to sources of variation in model training.

We do not perform significance tests using the variance in performance metrics computed through cross validation. In cross validation, data are reused in different train sets. The resulting metrics therefore violate the independence assumptions of many statistical tests, leading to underestimates in the likelihood of type I error [82].

Bootstrapping can produce better estimates of variance in model performance, but this involves high computational investment which may discourage future community use of a benchmark [81]. Therefore, as is typical for ML, we report variance in model performance in order to give a sense for its magnitude.

Low and high frequency components of acceleration

A common technique in analysis of acceleration data is to isolate acceleration due to gravity using a high- or low-pass filter [53], resulting in separate static and dynamic acceleration channels. It has been shown that the choice of cutoff frequency can have a strong effect on subsequent analyses [47]. Often, this frequency is chosen based on expert knowledge of an individual's physiology and typical movement patterns. As an alternative data-driven approach, we treated the cutoff frequency as a hyperparameter to be selected during model training. We use the terms “high frequency component” and “low frequency component” instead of “dynamic component” and “static component”, to reflect that in this data-driven approach, the cutoff frequency that leads to the best classification performance might not match the cutoff frequency that would isolate the acceleration due to gravity.

In more detail, for each raw acceleration channel, we apply a high-pass delay-free filter (using a linear-phase (symmetric) FIR filter with a Hamming window, followed by group delay correction) to obtain the high frequency component of the acceleration vector [83]. The high frequency component is then subtracted from the raw acceleration to obtain the low frequency component of the acceleration vector. The separated channels are then passed on as input for the rest of the model. For each dataset, the specific cutoff frequencies we selected from were 0 Hz (no filtering), 0.1 Hz, 0.4 Hz, 1.6 Hz, and

6.4 Hz. We omitted this step in the Rattlesnake dataset, where the high frequency components of the data had already been isolated. We also omitted this step for *harnet* models, since they were pre-trained using raw acceleration data, and for models using wavelet features, since the wavelet transform already decomposes a signal into different frequency components. For the experiments in Sect. 3.2 (Fig. 5), we used the cutoff frequency selected during the full data experiments in Sect. 3.1 (Fig. 4).

Model Implementation and Training Details

The methods we compared included the classical ML models Random Forests (RF), Decision Tree (DT), and Support Vector Machine (SVM), which are widely used to classify behavior recorded by bio-loggers (reviewed in [69, 79, Table 2]). These methods make predictions based on a set of pre-computed summary statistics, also known as *features*. For each classical method, we compared two different feature sets. The first feature set (denoted *Nathan*) was introduced in [14]. The second feature set (denoted *Wavelet*) consisted of spectral features, computed using a wavelet transform, inspired by [25] (see Methods).

In addition to classical methods, we compared methods based on deep neural networks that make predictions based on raw bio-logger data. We compared two types of models that are commonly used in other ML applications: a one-dimensional convolutional neural network (CNN), and a convolutional-recurrent neural network (CRNN). These types of methods have been employed by some recent studies focused on classifying animal and human behavior [60, 63, 84].

Additionally, we compared a convolutional-recurrent neural network which had been pre-trained with self-supervision, using human wrist-worn accelerometer data (*harnet* [61]). This network was pre-trained using over 700,000 days of un-annotated human wrist-worn accelerometer data recorded at 30 Hz [61]. The model was trained to predict whether these data had been modified, for instance by changing the direction of time, or by permuting the channels (Fig. 5A). We adapted this pre-trained model to our behavior classification task, which required small modifications of the network architecture (Fig. 5B; see details below). We refer to this modified model as *harnet frozen* (or *harnet* for short).

Models were implemented in Python 3.8, using PyTorch 1.12 [85] and scikit-learn 1.1.1 [86]. We used a variety of computing hardware depending on their availability through our computing platform (Google Cloud Platform). Deep neural networks (CNN, CRNN, *harnet*) used GPUs, and the rest of the models used CPUs. Our pool of GPUs included NVIDIA A100 and NVIDIA V100 GPUs. A single GPU was used to train each model. Our

pool of CPUs included machines with 16, 32, 64, 112 and 176 virtual CPUs.

For full implementation details, we refer readers to the source code,² which also contains the specific configurations that were evaluated during hyperparameter optimization. For all models we trained, we weighted the loss associated with each behavior class in inverse proportion to the frequency that that behavior occurred in the data; in initial experiments we found that this method for accounting for differences in behavior representation improved classification performance. For all the experiments presented, we trained over 2500 models for the purpose of hyperparameter tuning. For the hyperparameters that were then selected and used to obtain the reported results, we refer readers to our dataset repository.³

Supervised Neural Networks

All neural networks were implemented in PyTorch, with model-specific details given below. For each dataset in BEBE, we trained each type of model for 100 epochs using the Adam optimizer [87]. In each epoch, we randomly chose a subset of contiguous segments (*clips*) to use for training. The number of clips chosen per epoch was equal to twice the number of sampled time steps, divided by the clip length in samples. The clip length varied between datasets, and is specified below.

We used categorical cross-entropy loss, weighted in proportion to annotation imbalance. We applied cosine learning rate decay [88], and a batch size of 32. We masked all loss coming from sampled time steps annotated as *Unknown*.

CNN and CRNN models CNN consists of two dilated convolutional layers, a linear (i.e. width-1 convolution) prediction head, and a softmax layer. CRNN consists of two dilated convolutional layers, a bidirectional gated recurrent unit (GRU), a linear prediction head, and a softmax layer. In both CNN and CRNN, all convolutional layers are followed by ReLU activations and batch normalization. Each convolutional layer has 64 filters of size 7, and the GRU layer has 64 hidden dimensions. The outputs of these models are interpreted as class probabilities. CRNN models for behavior classification have previously been used by [49, 57].

We used a default clip length of 2048 samples (the time this represents will vary with the sampling rate of the dataset). However, two datasets include some deployments with fewer than 2048 recorded samples. For these, we used a shorter clip length (Rattlesnake, 64 samples; Seals, 128 samples).

² <https://github.com/earthspecies/BEBE/>.

³ <https://zenodo.org/record/7947104>.

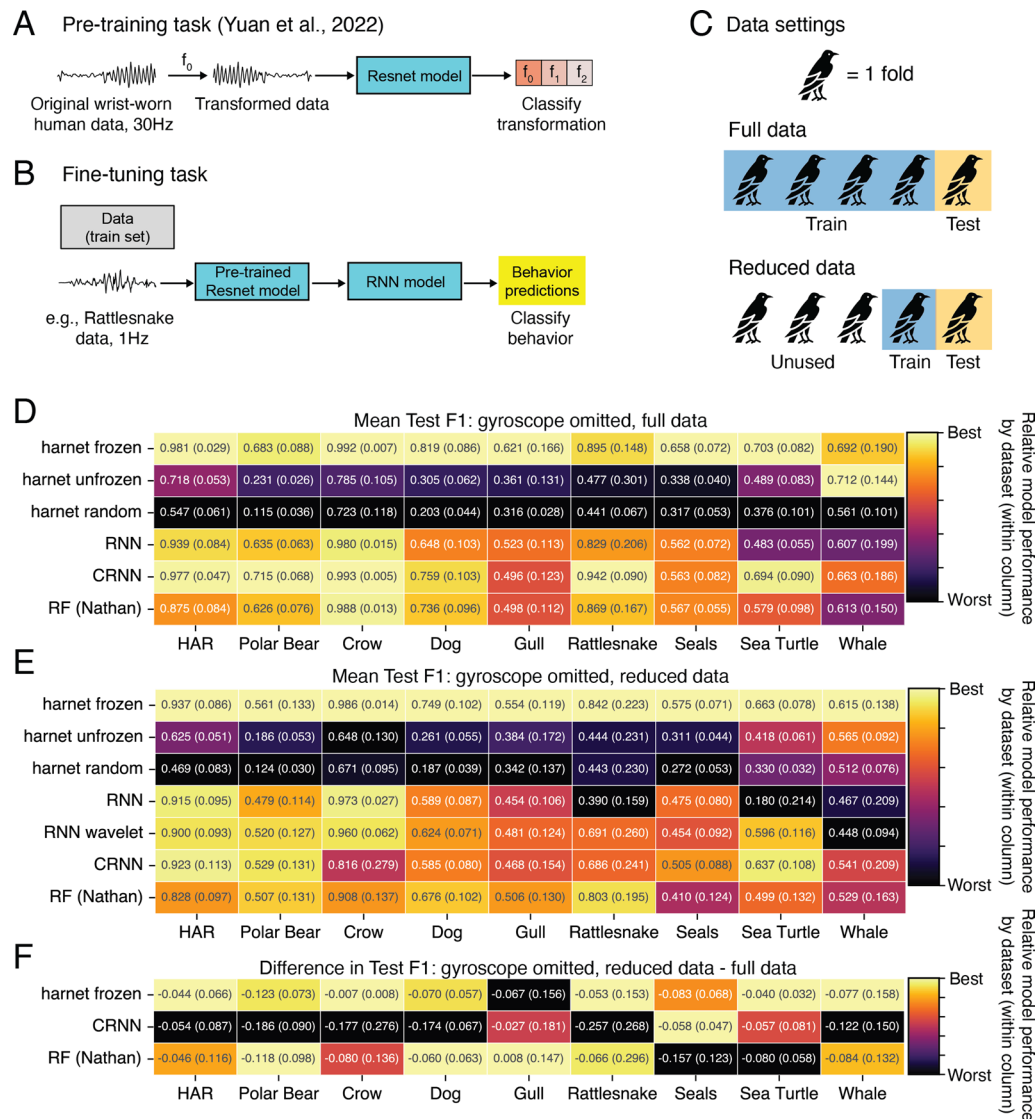


Fig. 5 Self-supervised pre-training and reduced data setting. **A** Pre-training task (performed in [61]): The main component of our *harnet* model has a Resnet architecture [89]. The Resnet was pre-trained with un-annotated human wrist-worn accelerometer data, which was modified with one of a set of signal transformations (e.g. f_0 = reversal in time). The network was trained to classify which transformation was applied to the original data. **B** In our *harnet* model, the input to the pre-trained Resnet was animal bio-logger data, without any modification to sampling rate. The outputs of the Resnet were passed to a recurrent neural network (RNN), which produced the behavior predictions. This full *harnet* model was then trained as shown in Fig. 3. **C** In the full data setting, four out of five folds are used to train the model in one-instance of cross validation. In the reduced data setting, only one fold is used for training while the test set is the same. In other words, approximately four times more individuals are included in the train set in the full data setting, than in the reduced data setting. **D** F1 scores for full data task. *harnet frozen* does best on five datasets and CRNN does best on three datasets. We omitted the RNN wavelet model from the full data experiments, due to high computational resources required for training, and its poor performance in the reduced data setting. **E** F1 scores for the reduced data task. *harnet frozen* does the best on all nine datasets. **F** Difference in F1 between reduced and full data tasks. For five datasets, *harnet frozen* shows the smallest decrease in F1 when using reduced data. For precision and recall results, see Figure S22

For our initial hyperparameter grid search, learning rate was selected from $\{1 \times 10^{-2}, 3 \times 10^{-3}, 1 \times 10^{-3}\}$, and convolutional filter dilation was selected from $\{1, 3, 5\}$.

Harnet and its variations For *harnet*, we use the pre-trained model described in [61], and obtained model weights from <https://github.com/OxWearables/ssl-wearables>. For a description of the pre-training setup using 700,000 h of un-annotated data from human wrist-worn

accelerometers, see Fig. 5 and [61]. As a default, we used the weights from the pre-trained model `harnet30` available at this repository. For the Rattlesnake and Seals datasets, we used the pre-trained model `harnet5` instead, since it was pre-trained on shorter duration clips. Both `harnet30` and `harnet5` consist of a sequence of convolution blocks, following [89]. Each pair of convolution blocks is separated by a pooling operation, which includes downsampling in the time domain. For our `harnet` model, we used the outputs of the second convolution block (of either `harnet30` or `harnet5`), which had been downsampled by a factor of 4. These outputs were then upsampled to their original temporal duration, before being passed through a bidirectional GRU with 64 hidden dimensions, and finally a linear (i.e. width-1 convolution) prediction head. The weights of the GRU and linear layers were randomly initialized before training.

Because the model trained in [61] operated on tri-axial accelerometer data, we only passed tri-axial accelerometer data through the convolutional blocks. To match the settings of the original pre-training setup, we normalized the acceleration channels so that the average magnitude of the acceleration vector was equal to 1. Additional data channels (depth, conductivity, and speed) were appended before being passed into the GRU. We omitted gyroscope channels; the data in these channels were relatively complex and outside the scope of the work of [61]. For our `harnet` model, we froze all weights during training except those in the GRU and prediction head. Our model `harnet_unfrozen` was identical to `harnet`, except we did not freeze any weights during training.

We evaluated several ablations of `harnet` (for results, see Fig. 5). The first, `harnet_random`, has the same architecture as `harnet` and `harnet_unfrozen`, except we used randomly initialized weights instead of the weights obtained by the pre-training procedure of [61]. This ablation was intended to disentangle the effect of pre-training from the effect of using this particular model architecture. The weights of `harnet_random` were all unfrozen during training.

The second ablation, `RNN`, omits the convolutional layers and passes the raw data directly into a GRU and linear prediction head. This ablation was intended to confirm that the improved performance of `harnet` was not due to the RNN architecture we used for the non-frozen part of the `harnet` model.

The third ablation, `RNN_wavelet`, replaces the convolutional layers of `harnet` with a wavelet transform. Each data channel is transformed using a Morlet wavelet transform, with 15 wavelets, using the `scipy.signal.cwt` module [86] (see details below). These wavelet transformed features are then passed into GRU

and linear layers, as in our `harnet` model. This ablation was intended to confirm that the improved performance of `harnet` could not be matched by computing spectral features, which would not require an intensive pre-training step. We tuned two hyperparameters for the wavelet transform, ω and C_{max} which are described in the following section. We selected ω and C_{max} from the same values as with the classical models.

For all the models described above, for our initial hyperparameter grid search, learning rate was selected from $\{1 \times 10^{-2}, 3 \times 10^{-3}, 1 \times 10^{-3}\}$. To match the pre-training setup of [61], we used a default clip length of 900 samples, and a clip length of 150 samples for Rattlesnake and Seals datasets.

Features for classical models

We tested two sets of features for the classical models (RF, DT, and SVM). The first set (`Nathan`) consists of summary statistics derived from [14], which have been used or adapted in a variety of behavior classification problems (e.g. [39, 68]). The second set are wavelet features (`Wavelet`) [25], which are commonly used to identify periodic motions like steps or tail beats (e.g. [16, 33]).

For the `Nathan` feature set, we first defined the feature set for tri-axial accelerometer channels in the same way as [14]. For each time step t , we first computed the root-mean-square amplitude q from the x , y , z -axes, as well as the high frequency component acceleration (as described above). Then, for the x , y , z -axes and q , we computed a *basic* set of features over a contiguous segment of data (*clip*) centered at t : mean, standard deviation, skew, kurtosis, maximum, minimum, 1-sample autocorrelation, and best fit slope. We also computed three pairwise correlations between the x , y , z -axes, the circular variance of inclination for q -axis, the circular variance of azimuth for q -axis, and a quantity analogous to the mean overall dynamic body acceleration (ODBA) using the high frequency component acceleration. If there were any tri-axial gyroscope channels, we computed the *basic* set of features as well as the pairwise correlations on the x , y , z -channels. For other channels (conductivity, depth and speed), we only computed the *basic* features. Because the datasets included different channels, they had a different number of input features. For our initial grid search, the duration (in seconds) of the clip used for feature computation was selected from $\{0.5, 1, 2, 4, 8, 16\}$ seconds. For the Dog dataset, this duration was selected from $\{0.5, 1, 2, 4, 8\}$ seconds due to memory limitations.

For the `Wavelet` feature set, we normalized each channel independently (z -score), and then for each channel we followed the procedure in [25]. We computed the continuous wavelet transform using `scipy.signal.cwt` with the complex Morlet wavelet function (`scipy.`

signal.morlet2), using 15 wavelets per data channel. We tuned two hyperparameters for the wavelet transform. The first hyperparameter, the dimensionless ω , controls the tradeoff between resolution in the time and frequency domains. The second, C_{max} controls the largest wavelength (in seconds) out of the 15 wavelets used. The parameter ω was selected from {5, 10, 20}, and C_{max} was selected from {1, 10, 100, 1000}. Once C_{max} was fixed, the wavelength of the k^{th} wavelet ($k = 0, 1, \dots, 14$) was equal to $C_{min} * \left(\frac{C_{max}}{C_{min}}\right)^{k/14}$. The minimum wavelength, C_{min} , was set to $2/sr$, where sr is the dataset-specific sampling rate.

Classical models

All classical models were implemented with Scikit-learn version 1.1.1, with model-specific details given below. We used loss functions weighted to account for annotation imbalance (corresponding to `class_weight = balanced`). During training, we did not include any sampled time steps which were annotated as *Unknown*.

Random Forest RF was implemented using `RandomForestClassifier` from the `sklearn.ensemble` package. For each tree we used 1/10 of the available training data (`max_samples = 0.1`). Other than `max_samples` and `class_weight`, we used the default settings. The model consists of 100 decision trees.

Decision Tree DT was implemented using `DecisionTreeClassifier` from the `sklearn.tree` package, using default settings except for `class_weight`.

Support Vector Machine SVM was implemented using `LinearSVC` from the `sklearn.svm` package, chosen for its scaling properties to large numbers of samples. We selected the algorithm to solve the primal optimization problem (i.e., `dual = False`) because `n_samples > n_features`. Other than `dual` and `class_weight`, we used the default settings.

Results

Deep neural networks improve classification performance

We used the datasets and model evaluation framework in BEBE to compare different methods for predicting behavior from bio-logger data. For each dataset, we compared the performance of three classical ML models with three deep neural networks. We predicted that neural network approaches would outperform the classical approaches we tested (Table 1, H1). The F1 scores for these models are given in Fig. 4, with precision and recall are presented in Supplemental Fig. S21.

In terms of classification F1 score, the methods we tested that were based on deep neural networks performed the best on all nine datasets in BEBE, confirming hypothesis (H1). The top performing model was always

either CRNN or `harnet`. The top performing deep neural net on a dataset achieved a F1 score that was 0.072 greater, on average, than the F1 score of the top performing classical model on that dataset. Deep neural networks achieved the best recall on all datasets and the best precision in seven out of nine datasets.

Self-supervised pre-training enables low-data applications

To evaluate the effectiveness of self-supervised pre-training for analysis of bio-logger data, we focus on the `harnet` neural network. This network was pre-trained using over 700,000 days of un-annotated human wrist-worn accelerometer data recorded at 30 Hz [61].

We predicted that `harnet` would out-perform the alternative methods we tested (Table 1, H2). We measured the performance of `harnet` on each dataset in BEBE, after fine-tuning (Fig. 5B). To better understand the contribution of the pre-training step on performance, we compared these results with various ablations of `harnet`: RNN, RNN wavelet, and `harnet random` (for justification, see Methods). Additionally, we compared `harnet` with CRNN and RF (Nathan), which were the best alternatives to `harnet` in Sect. 3.1. Finally, we also compared `harnet` with an alternative setup, `harnet unfrozen`, which had more tunable parameters (see Methods). The results of these comparisons are in Fig. 5D, with precision and recall scores in Supplemental Fig. S22. Gyroscope data were not included in the pre-training procedure of [61], and so we omitted these channels when obtaining these results in order to concentrate on the effect of the pre-training methodology. Other channels (e.g. depth) were still included, following the procedure detailed in Methods.

In terms of F1 score, `harnet` achieved the top score on five of the nine datasets, partly confirming hypothesis (H2). In three of the remaining four cases, CRNN achieved the top score, and in the remaining case, the `harnet unfrozen` variant achieved the top score. None of the ablations we tested approached the performance of `harnet`, indicating that the pre-training step, and not another design choice, was responsible for its high performance. The alternative setup for the pre-trained model, `harnet unfrozen`, achieved lower scores than `harnet` on eight of nine datasets (average F1 drop: .263).

To investigate the potential of self-supervised learning in low-data applications, we performed an additional set of computational experiments. For these, during cross validation, we trained each model using only one of five folds (rather than four of five) of the dataset (Fig. 5C). The folds used for testing remained the same. Because the folds partition the tagged individuals, this setting reflects a situation where the researcher can only annotate training data from a quarter of the tagged individuals. We

predicted that in this setting, *harnet* would outperform the alternative models we tested (Table 1, H3). The F1 scores of models in the reduced data setting are in Fig. 5E.

After reducing the amount of training data, the pre-trained *harnet* model dropped in F1 performance by 0.056 on average, as compared with a drop of 0.112 by CRNN and 0.069 by RF (Nathan) (Fig. 5F). Additionally, *harnet* achieved the top F1 score on all nine datasets in the reduced data setting (Fig. 5E) and across most behavior classes (Supplemental Fig. S23), and *harnet* achieved the best recall across all datasets. The ablations of *harnet* consistently performed poorly, relative to other models, in this reduced data setting. Taken together, these results confirm hypothesis (H3).

For some behavior classes, model performance improves minimally with increased training data

Using BEBE, we investigated the variation in F1 score across different behavioral classes within a single dataset. Using *harnet* in the full data setting, the inter-class range in F1 score ranged from small (Crow dataset, range: [0.990,0.995], max-min difference: 0.0055) to large (Gull dataset, range: [0.0768, 0.955], max-min difference: 0.878) (Fig. 6A, larger dots). This variation also existed in the reduced data setting, and for the CRNN and RF (Nathan) models (Fig. 6A, smaller dots; Supplemental Figs. S24–S25). We found that classes with relatively few training examples can perform as well as classes with many training examples (e.g., Sea turtle: *Stay at surface* vs. *Swim*).

Next, we examined the difference in performance between the reduced and full data settings. We predicted that model performance would improve minimally for some behavioral classes, when trained with four times as much training data (Table 1, H4; Fig. 6C). For *harnet*, the degree of improvement in per-class scores ranged from null (e.g., *Rest dive* in Whale dataset) to moderate (*Swim* in Polar bear dataset: 26 improvement in F1 score), with the median at 0.053 (Fig. 6B). Several classes with F1 scores lower than 0.9 in the reduced data setting showed small improvements (e.g., *Rest* in Seals dataset). This is consistent with what we would expect if predictive performance for these behaviors had reached an invisible ceiling, but does not conclusively demonstrate it: increasing the data even further or switching models may improve performance on these classes. Nevertheless, we expect that such classes are unlikely to become highly recognizable even going beyond a fourfold increase in training data.

Finally, we quantified the extent to which one could use model performance in the reduced data setting to predict performance in the full data setting. For the three models

we tested, there was a high degree of correlation between the per-class scores at these two data scales (*harnet*: $r(47) = 0.96, p < 0.001$, CRNN: $r(47) = 0.89, p < 0.001$, RF: $r(47) = 0.94, p < 0.001$; Fig. 6C). Last, we found that the rank ordering of classes within a dataset was conserved when the amount of training data was reduced. Within a dataset, the rank correlation between per-class scores in the reduced and full data settings was typically high (see Supplemental Table S1), indicating that similar classes performed better in the reduced and full data settings.

Conclusions

To support the development and application of methods for behavior classification in bio-logger data, we designed the Bio-logger Ethogram Benchmark (BEBE), a collection of nine annotated bio-logger datasets. BEBE is the largest, most diverse, publicly available bio-logger benchmark to date. As an example of how BEBE can be used by the community, we tested several hypotheses about ML methods applied to bio-logger data. Based on our results, we are able to make concrete suggestions for those designing studies that rely on ML to infer behavior from bio-logger data.

First, we found that methods based on deep neural networks outperformed the classical ML methods we tested (Table 1, H1; Fig. 4). While a similar trend has been observed in other applications of ML, the majority of studies involving bio-logger data still rely on classical methods such as random forests [69, Table 2]. Contrary to this trend, we suggest that researchers use methods based on deep neural networks in studies where the main objective for ML is to maximize the accuracy of the predicted behavior labels, and with datasets of comparable size to those in BEBE. In particular, for the behavior classification task we describe in this work, we suggest the use of a convolutional-recurrent architecture (which was used by both CRNN and *harnet*). In contrast to random forests, these deep neural networks learn their features directly from data and therefore do not require an intensive feature engineering step. Additionally, the dominance of CRNN over CNN across datasets demonstrates the importance of incorporating time scale as a learnable parameter (in contrast to RF and CNN where it is fixed).

While we found evidence that convolutional-recurrent networks tend to outperform feature-based methods like RF, they also tend to be more labor-intensive to implement and train. Thus, studies employing ML-based behavior classification may want to weigh the benefits of adopting these methods against their costs. Classification errors can propagate into downstream analyses, increasing the need to correct results for systematic bias [90].

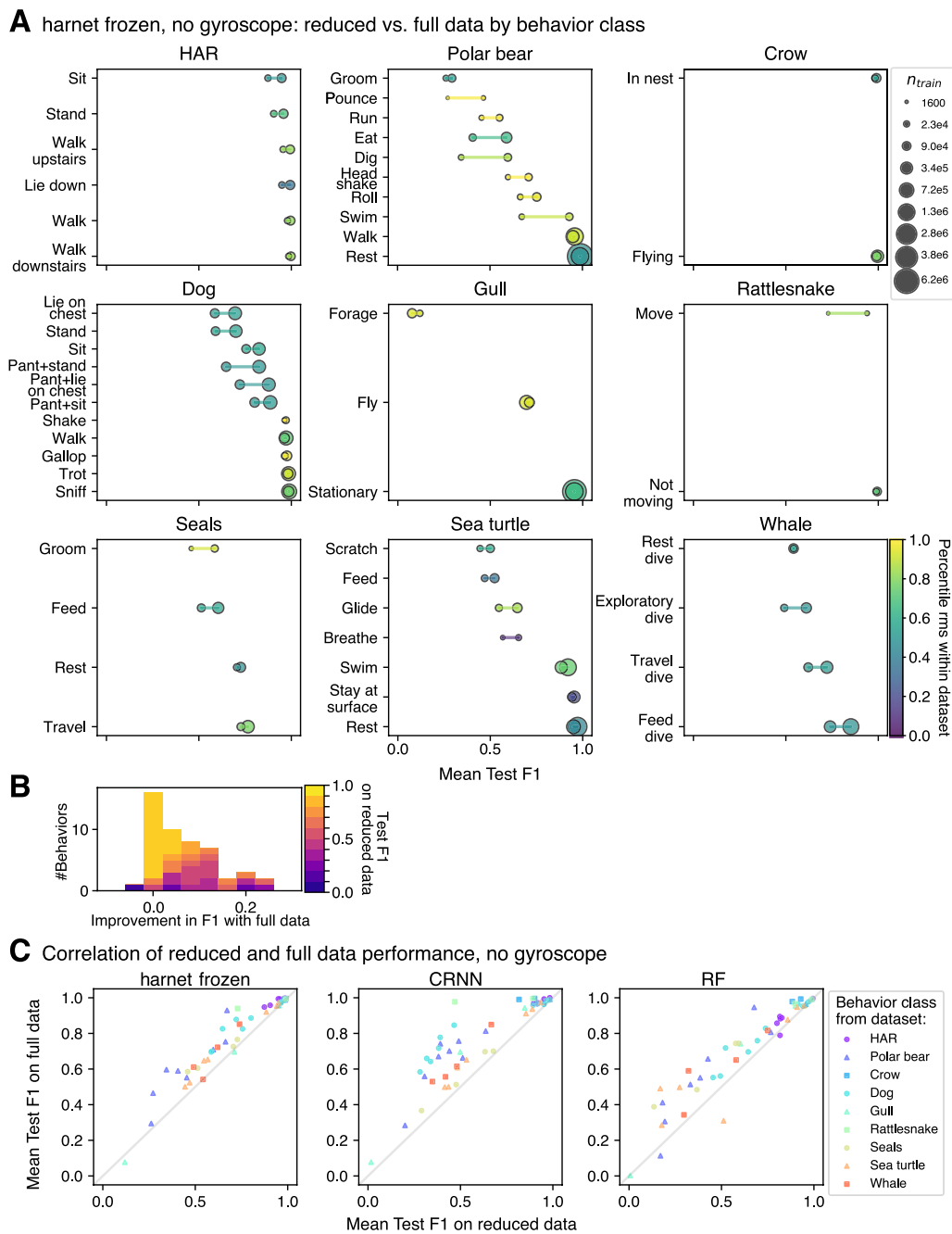


Fig. 6 Variation in model performance is conserved across data scales. **A** Performance of harnet in reduced and full data settings, by behavior class. Size of the marker indicates training dataset size (mean across folds). To help visualize the type of behavior, color indicates the percentile of the average root-mean-square (rms) amplitude of datapoints in that class, as compared to the root-mean-square amplitude of all labeled datapoints. Performance typically improves from the reduced to full data setting, but the rank order of behavior classes remains similar. **B** Histogram of per-behavior improvement in harnet's F1 score across all datasets, moving from reduced data to full data. Colors indicate F1 score in the reduced data setting. **C** Performance in reduced data setting versus performance in full data setting, for three models. Each point represents one behavioral class. If performance is the same in the reduced and full data settings, the point lies on the line. For all three models, there is a high correlation between the reduced and full data settings. CRNN shows the greatest improvements in F1 score with increased data (although CRNN and harnet have comparable performance in the full data setting; Fig. 5D)

Additionally, classification errors can increase uncertainty in downstream analyses. The uncertainty contributed by model classification errors can be quantified, for example, using bootstrap sampling from a confusion matrix [91]. Researchers may consider beginning with a feature-based model, such as RF, and adopting higher-performance methods if the feature-based method does not predict behaviors accurately enough to answer their ultimate scientific question.

We also found that a neural network pre-trained with self-supervision using a large amount of human wrist-worn accelerometer data achieved the best performance on just over half of the datasets (Table 1, H2; Fig. 4). This pattern became more pronounced when the amount of training data was reduced (Table 1, H3; Fig. 5F). Deep neural networks without self-supervision were not appreciably better than random forests in this reduced data setting (CRNN and `harnet` random vs. RF; Fig. 5E). Therefore, we suggest that studies examine adapting pre-trained neural networks (such as [61]), rather than training a neural network from scratch. This approach is especially promising for improving behavior classification in situations where a relatively small amount of annotated data is available, for example, due to the difficulty of obtaining ground truth behavior. Self-supervision could complement the use of surrogate species, which has had variable levels of success [22, 23, 92]. We believe this performance is particularly notable since the pre-training data came from only one species (humans, represented in only one dataset in BEBE), and came from only one attachment position (wrist, not represented in BEBE). Additionally, the pre-training data were recorded at 30 Hz (different from all datasets in BEBE), and we did not adjust for differences in sampling rate during fine-tuning. Last, one limitation of this approach was that the pre-training involved only data from accelerometers. A potential future direction would be to perform pre-training using data from more diverse taxa, using a wider variety of tag placements, sampling rates, and sensor types.

Finally, we found that model performance improves minimally for some behaviors, when increasing training data by adding individuals (Table 1, H4; Fig. 6B). Annotating a large train set may not provide sufficient benefit if some of the behaviors of interest are inherently difficult to identify from the available sensors. We also found that per-class model performance is correlated across data scales (Fig. 6C). This suggests that, as a part of designing a data analysis procedure, it may be worthwhile to attempt the analysis after annotating a small portion of the available data. This could provide a sense of the expected relative per-behavior model performance, and allow for adjustments to the ethogram if the performance is unlikely to reach the range necessary to address the study

question; for example, one might decide to reduce the number of behaviors in the ethogram in order to improve classification performance [15, 21]. With a diversity of bio-logging studies represented, BEBE may provide a shared resource for practitioners to identify which behaviors are easily discriminated from sensor data. While we cannot conclusively determine whether we reached ceiling performance for behaviors in BEBE with the experiments presented, we found that some behaviors showed minimal improvement with more training data (Fig. 6). Therefore, when using BEBE as a benchmark, we suggest that researchers track per-behavior performance improvements.

An important aspect of BEBE is that the data and evaluation code is openly available. This allows others to test hypotheses beyond those in Table 1, and to test future, yet to be developed ML methods. For example, future work could employ BEBE to systematically test different data augmentation techniques such as those suggested by [49], intended to improve performance with a small amount of annotated data. Additionally, it is possible that one could improve upon the conclusions in this article. For example, when testing (H1), for the classical models we tested, we used our own implementation of one of two types of generic hand-engineered features as model inputs [14, 25]. It is possible that other feature sets or model types would improve on the performance of the models we evaluated. For example, Wilson et al. [20], suggest using features tailored to the bio-mechanics of specific behaviors of interest. (For instance, to detect feeding in sea turtles, Jeantet et al. [18] pre-segmented data based on the variance in the angular speed in the sagittal plane.) Using BEBE as a common framework for comparison, others can interrogate the results presented here and improve suggestions for the bio-logging community.

A number of prior works have tested multiple ML algorithms for behavior classification, typically on a single dataset (e.g., [14, 16, 24]). We extend this trend by testing multiple ML algorithms on multiple datasets. Bio-logger datasets tend to be very heterogeneous, and can differ in study system, sensor type, sampling rate, ethogram definition and train-test split design. When comparing the results of disparate studies, it can be difficult to disentangle the effect of model design from the effects of the dataset. While we observed that there is a large amount of variability in model performance between datasets, we found that certain techniques performed consistently well (relative to alternatives). Therefore, by systematically testing techniques in a variety of settings, we are able to observe patterns in how ML methods are applied to bio-logger data, generally.

While this study focuses on supervised behavior classification, unsupervised behavior classification (which does

not rely on labeled data) is of interest where there is little knowledge of relevant behaviors or where behavior is difficult to observe [25, 50]. Based on an analysis of two seabird datasets, Sur et al. [24] argue that these methods perform worse than supervised methods in recovering pre-defined behaviors. They are also difficult to systematically evaluate. It is typically unclear what aspect of behavior they target and their outputs therefore require interpretation. To address this particular challenge, ethograms that are defined hierarchically, i.e., with behaviors composed of multiple more specific behaviors (e.g., [15, 18, 39, 53]), may provide a promising basis for evaluation, by providing insight into which aspects of behavior are modeled by a certain method.

We note limitations and potential improvements to our approach. First, the datasets in BEBE are primarily based on tri-axial accelerometers, which may not be able to represent motion accurately enough to distinguish all behaviors of interest [73]. Bio-loggers incorporating other sensor types, such as gyroscopes, audio, and video, will likely give researchers more complete characterizations of individuals' behaviors. Other benchmarks exist for animal behavior detection entirely from video [44, 45] but these do not focus on bio-logger data, which may present additional challenges such as intermittent video logging or a limited field of view [22, 35]. A future benchmark could include data types not examined in BEBE.

Second, bio-loggers can shed new light on conservation problems and interventions, as well as on patterns of animal behavior and energy expenditure [7, 11, 43, 93, 94]. In this study, we provide a standardized task for one extremely common analysis, behavior classification. Depending on the intended application, other analyses may be useful. These could include detecting unusual patterns in data [95] that may indicate changes in behavior or environmental conditions [96], or counting the rate at which a specific type of behavioral event occurs [9]. Future studies could use BEBE datasets, but formalize new tasks and evaluation metrics for use-cases that arise in these settings. In addition, studies based on BEBE could explore evaluation metrics to promote advances in on-device ML [35], such as device energy consumption metrics to assess on-device feasibility. This could additionally give insight into environmental impacts due to model usage [43, 97, 98]. Overall, we expect BEBE can support researchers by facilitating access to a breadth of study systems, which may involve using our standard task or creating different uses of BEBE datasets.

Third, there is variation in how the datasets were collected and annotated, and BEBE has limited taxonomic spread and no taxonomic replication. While some variation is desirable in order to promote generalizable methods development, it also complicates between-dataset

or between-behavior comparisons. These types of comparisons could illuminate how a model's predictive ability is related to biological factors, such as phylogeny or body size, and to non-biological factors, such as the choice of ethogram or data modalities included. It may be possible to quantify the effects of these factors using a benchmark with more datasets available and better data standardization.

Fourth, in our standardized evaluation framework, we exclude Unknown behaviors in the training objective and evaluation metric. The presented models assume that the provided known behavior labels are the only possible categories, and will apply one of them to all datapoints. This would be disadvantageous in applying supervised learning to bio-logger data where we usually know some behaviors have not been labeled. Approaches to accounting for unknown behaviors include using an "other" category [18, 39] and thresholding classification probabilities to make a prediction [99]. One potential future usage of BEBE would be to test these different methodologies for accounting for unknown behaviors, to elucidate their impact on recovering the behaviors of interest.

Finally, human-generated behavior annotations are susceptible to error, due to e.g. difficulty in observing a behavior, mistakes during data entry, or differences in human judgment. To mitigate this, behaviors can be annotated by multiple raters, and then checked against one another. However, this can be extremely time consuming, and in BEBE only the authors of the Rattlesnake, Whale, and Seal datasets performed this step. In spite of this potential annotation noise, we were able to observe consistent patterns in model performance across multiple datasets in BEBE. In the future, if more studies report between-rater agreement in annotations, it may become possible to quantify the magnitude of these errors.

Call for Collaboration The code repository includes instructions on how datasets outside of BEBE may be formatted for use with the methods in BEBE. Interested researchers may make their formatted datasets discoverable from the BEBE repository. These datasets would become easily available for others, but would not become part of the task in this paper, which must remain standardized.

However, it is typical for benchmarks to be updated when key challenges are sufficiently met [100]. In light of the preceding discussion, we seek community contributions that could lead to a more comprehensive benchmark, with three main objectives:

1. To provide researchers with means to understand how modeling decisions influence model performance,

2. To enable analyses which compare recorded behavior across taxa, and
3. To formalize tasks which reflect a variety of real-world applications, including conservation applications.

We expect these objectives will be best served by a benchmark with more diversity in its representation of taxa, data types, tag placement positions, sensor configurations, ethograms, and modeling tasks. Possible contributions include (1) annotated datasets to be made openly available to the research community (whether already available or not), (2) design of data and annotation standardization, and (3) design of benchmark tasks that reflect applications of ML and bio-logger technology. For any ensuing publications, contributors would have the option to co-author the manuscript. Interested researchers can follow the instructions at <https://github.com/earthspecies/BEBE>.

We have proposed that benchmarks can encourage the development and rigorous evaluation of ML methods for behavioral ecology. We envision many possible future outcomes for this line of research: for example, best practices for bio-logger data analysis, an ML-based toolkit that can be adapted to different study systems, or powerful species-agnostic tools that can be applied across taxa and sensor types. This could, in turn, inform more effective conservation interventions, as well as guide the development and testing of hypotheses about animal behavior.

Supplementary Information

The online version contains supplementary material available at <https://doi.org/10.1186/s40462-024-00511-8>.

Supplementary file 1.

Acknowledgements

This project was supported (in part) by a grant from the National Geographic Society. Compute resources were provided in part by Google Cloud Platform. We thank Phoebe Koenig for input on benchmark design. We thank Christian Rutz and Mark Johnson for critical contributions to an earlier version of the manuscript. We thank Felix Effenberger, Masato Hagiwara, Sara Keen, Jen-Yu Liu, and Marius Miron for constructive discussions. Any use of trade, firm, or product names is for descriptive purposes only and does not imply endorsement by the United States Government. Whale behavior data collection was funded by National Science Foundation Office of Polar Programs grants awarded to A. Friedlaender, and were collected under National Marine Fisheries Service Permits 14809 and 23095, Antarctic Conservation Act permits and UCSC IACUC Permit Friaa2004. Crow behavior data collection was funded by the Ministerio de Economía y Competitividad - España (Grant CGL2016 - 77636-P to V. Baglione). Rattlesnake behavior data collection was funded by a National Science Foundation Graduate Research Fellowship (NSF-GRFP) awarded to D. L. DeSantis and grants from the UTEP Graduate School (Dodson Research Grant) awarded to D. L. DeSantis. J. D. Johnson, A. E. Wagler, J. D. Emerson, M. J. Gaupp, S. Ebert, H. Smith, R. Gamez, Z. Ramirez, and D. Sanchez contributed to dataset development, and C. Catoni and Technosmart Europe srl. developed the accelerometers. Dog behavior data collection was funded by Business Finland, a Finnish funding agency for innovation, grant numbers

1665/31/2016, 1894/31/2016, and 7244/31/2016 in the context of “Buddy and the Smiths 2.0” project. Gull behavior data collection was funded by JSPS KAKENHI Grant Number JP21H05299 and JP21H05294, Japan. Sea turtle behavior data was collected within the framework of the Plan National d’Action Tortues marines des Antilles et the Plan National d’Action Tortues marines de Guyane Française, with the support of the ANTIDOT project (Pépinière Interdisciplinaire Guyane, Mission pour l’Interdisciplinarité, CNRS) and BEPHYTES project (FEDER Martinique) led by D. Chevallier. DEAL Martinique and Guyane, the CNES, the ODE Martinique, POEMM and ACWAA associations, Plongée-Passion, Explorations de Monaco team, the OFB Martinique and the SMPE Martinique provided technical support and field assistance. Numerous volunteers and free divers participated in the field operations to collect data. EGI, France Grilles and the IPHC Computing team provided technical support, computing and storage facilities for the original development of the Sea Turtle dataset. Polar bear behavior data collection was funded by the U.S. Geological Survey Changing Arctic Ecosystems Initiative and the Species and Land Management programs of the U.S. Geological Survey Ecosystems Mission Area. Seals behavior data collection was assisted by the marine mammal staff at Dolphin Marine Magic, Sealife Mooloolaba and Taronga Zoo Sydney. Animal icons in Fig. 1B and repeated throughout the paper have a FlatIcon license (free for personal/commercial use with attribution), attributions are as follows: (1) Gull: user *Smashicons*, (2) Rattlesnake: user *FreePik*, (3) Polar bear: user *Chanut-is-Industries*, (4) Dog: user *FreePik*, (5) Whale: user *The Chohans Brand*, (6) Turtle: user *FreePik*, (7) Crow: user *iconixar*, (8) Seal: user *monkik*, (9) Human: user *Bharat Icons*. Image attributions for Fig. 1C are as follows: (1) Gull: scaled and cropped from user Wildreturn (Flickr; CC BY 2.0), (2) Rattlesnake: scaled and cropped from user snakecollector (Flickr; CC BY 2.0), (3) Polar bear: scaled and cropped from user usfwshq (Flickr; CC BY 2.0), (4) Dog: Andrea Austin, (5) Whale: scaled and cropped image from user onms (Flickr; CC BY 2.0), (6) Sea turtle: scaled and cropped from user dominic-scaglioni on (Flickr; CC BY 2.0), (7) Crow: scaled and cropped from user alexislours (Flickr; CC BY 2.0), (8) Seal: scaled and cropped from volvob12b (Bernard Spragg) (Flickr; Public domain), (9) Human: Katie Zacarian.

Author contributions

M.C., A.F., and B.H. conceived the ideas; M.C. and B.H. designed methodology; V.B., D.C., D.C., D. D., A. F., L. J., M.L., T. M., V. M., V. M., A. P., E. T., O. V., A. V., K. Y., contributed the data; M. C., B. H., and K. Z. coordinated data contributions; M. C. and B. H. analysed the data; M. C. and B. H. led the writing of the manuscript. All authors contributed critically to the drafts and gave final approval for publication.

Availability of data and materials

The datasets generated and/or analysed during the current study are available in the Zenodo repository, (doi: 10.5281/zenodo.7947104). All model results used to create the figures are also available at the same repository. Code used to format the datasets is available at <https://github.com/earthspecies/BEBE-datasets/>. Code used to implement, train, and evaluate models is available at <https://github.com/earthspecies/BEBE/>.

Declarations

Ethics approval and consent to participate

All animal behavior datasets except the Crow dataset were reported in previous publications. Crow behavior data were collected in accordance with ASAB/ABS guidelines and Spanish regulations for animal research, and were authorized by Junta de Castilla y León (licence: EP/LE/681-2019).

Competing interest

The authors declare no competing interests.

Author details

¹Earth Species Project, Berkeley, CA, USA. ²University de León, León, Spain. ³CNRS Borea, Les Anses d’Arlet, France. ⁴Georgia College & State University, Milledgeville, GA, USA. ⁵African Institute for Mathematical Sciences, University of Stellenbosch, Stellenbosch, South Africa. ⁶Department of Conservation, Wellington, New Zealand. ⁷Osaka University, Osaka, Japan. ⁸University Texas El Paso, El Paso, TX, USA. ⁹US Geological Survey, Anchorage, AK, USA. ¹⁰University of Helsinki, Helsinki, Finland. ¹¹Tampere University, Tampere, Finland. ¹²Nagoya

University, Nagoya, Japan. ¹³University of California Santa Cruz, Santa Cruz, CA, USA.

Received: 27 June 2024 Accepted: 3 October 2024

Published online: 18 December 2024

References

- Davies NB, Krebs JR, West SA. An introduction to behavioural ecology. NJ: John Wiley & Sons; 2012.
- Berger-Tal O, Polak T, Oron A, Lubin Y, Kotler BP, Saltz D. Integrating animal behavior and conservation biology: a conceptual framework. *Behav Ecol*. 2011;22(2):236–9.
- Walters JR, Derrickson SR, Fry DM, Haig SM, Marzluff JMW Jr. Status of the California Condor (*Gymnogyps californianus*) and efforts to achieve its recovery. *The Auk*. 2010;127(4):969–1001.
- Thaxter CB, Lascelles B, Sugar K, Cook ASCP, Roos S, Bolton M, et al. Seabird foraging ranges as a preliminary tool for identifying candidate marine protected areas. *Biol Conserv*. 2012;156:53–61.
- Tingley R, Phillips BL, Letnic M, Brown GP, Shine R, Baird SJE. Identifying optimal barriers to halt the invasion of cane toads *Rhinella marina* in arid Australia. *J Appl Ecol*. 2013;50(1):129–37.
- Rutz C, Hays GC. New frontiers in biologging science. *Biol Lett*. 2009;5(3):289–92.
- Wilson R, Shepard E, Liebsch N. Prying into the intimate details of animal lives: use of a daily diary on animals. *Endanger Species Res*. 2008;01(4):123–37.
- Yoda K, Naito Y, Sato K, Takahashi A, Nishikawa J, Ropert-Coudert Y, et al. A new technique for monitoring the behaviour of free-ranging Adelie penguins. *J Exp Biol*. 2001;204(4):685–90.
- Bateson M, Martin P. Measuring behaviour: an introductory guide. Cambridge: Cambridge University Press; 2021.
- Ladds M, Salton M, Hocking D, McIntosh R, Thompson A, Slip D, et al. Using accelerometers to develop time-energy budgets of wild fur seals from captive surrogates. *Peer J*. 2018;10(6):e5814.
- Valletta JJ, Torney C, Kings M, Thornton A, Madden J. Applications of machine learning in animal behaviour studies. *Animal Behav*. 2017;124:203–20.
- Minasandra P, Jensen FH, Gersick AS, Holekamp KE, Strauss ED, Strandburg-Peshkin A. Accelerometer-based predictions of behaviour elucidate factors affecting the daily activity patterns of spotted hyenas. *Royal Soc Open Sci*. 2023;10(11): 230750.
- Studd E, Peers M, Menzies A, Derbyshire R, Majchrzak Y, Seguin J, et al. Behavioural adjustments of predators and prey to wind speed in the boreal forest. *Oecologia*. 2022;200(3):349–58.
- Nathan R, Spiegel O, Fortmann-Roe S, Harel R, Wikelski M, Getz WM. Using tri-axial acceleration data to identify behavioral modes of free-ranging animals: general concepts and tools illustrated for Griffon vultures. *J Exp Biol*. 2012;215(6):986–96.
- Studd EK, Landry-Cuerrier M, Menzies AK, Boutin S, McAdam AG, Lane JE, et al. Behavioral classification of low-frequency acceleration and temperature data from a free-ranging small mammal. *Ecol Evol*. 2019;9(1):619–30.
- Brewster L, Dale J, Guttridge T, Gruber S, Hansell A, Elliott M, et al. Development and application of a machine learning algorithm for classification of Elasmobranch behaviour from accelerometry data. *Mar Biol*. 2018;165(4):62.
- DeSantis DL, Mata-Silva V, Johnson JD, Wagler AE. Integrative framework for long-term activity monitoring of small and secretive animals: validation with a cryptic Pitviper. *Front Ecol Evol*. 2020;8:169.
- Jeantet L, Planas-Bielsa V, Benhamou S, Geiger S, Martin J, Siegwalt F, et al. Behavioural inference from signal processing using animal-borne multi-sensor loggers: a novel solution to extend the knowledge of sea turtle ecology. *Royal Soc Open Sci*. 2020;7(5):200139.
- Patterson A, Gilchrist HG, Chivers L, Hatch S, Elliott K. A comparison of techniques for classifying behavior from accelerometers for two species of seabird. *Ecol Evol*. 2019;9(6):3030–45.
- Wilson R, Holton M, Virgilio AD, Williams J, Shepard E, Lambertucci S, et al. Give the machine a hand: a Boolean time-based decision-tree template for rapidly finding animal behaviours in multisensor data. *Methods Ecol Evol*. 2018;9(11):2206.
- Ladds M, Thompson A, Kadar J, Slip D, Hocking D, Harcourt R. Super machine learning: improving accuracy and reducing variance of behaviour classification from accelerometry. *Animal Biotelem*. 2017;5:1–9.
- Pagano A, Rode K, Cutting A, Jensen S, Ware J, Robbins C, et al. Using tri-axial accelerometers to identify wild polar bear behaviors. *Endanger Species Res*. 2017;01:32.
- Campbell HA, Gao L, Bidder OR, Hunter J, Franklin CE. Creating a behavioural classification module for acceleration data: using a captive surrogate for difficult to observe species. *J Exp Biol*. 2013;216(24):4501–6.
- Sur M, Hall JC, Brandt J, Astell M, Poessel SA, Katzner TE. Supervised versus unsupervised approaches to classification of accelerometry data. *Ecol Evol*. 2023;13(5):e10035.
- Sakamoto KQ, Sato K, Ishizuka M, Watanuki Y, Takahashi A, Daunt F, et al. Can ethograms be automatically generated using body acceleration data from free-ranging birds? *PLoS ONE*. 2009;4(4):e5379.
- Leos-Barajas V, Photopoulou T, Langrock R, Patterson TA, Watanabe YY, Murgatroyd M, et al. Analysis of animal accelerometer data using hidden Markov models. *Methods Ecol Evol*. 2017;8(2):161–73.
- Hanscom RJ, DeSantis DL, Hill JL, Marbach T, Sukumaran J, Tipton AF, et al. How to study a predator that only eats a few meals a year: high-frequency accelerometry to quantify feeding behaviours of rattlesnakes (*Crotalus* spp.). *Animal Biotelem*. 2023;11(1):20.
- Ferdinandy B, Gerencsér L, Corrieri L, Perez P, Újváry D, Cszimadia G, et al. Challenges of machine learning model validation using correlated behaviour data: evaluation of cross-validation strategies and accuracy measures. *PLoS One*. 2020;15(7):e0236092.
- Kumpulainen P, Cardó AV, Somppi S, Törnqvist H, Väättäjä H, Majaranta P, et al. Dog behaviour classification with movement sensors placed on the harness and the collar. *Appl Animal Behav Sci*. 2021;241:105393.
- Sur M, Suffredini T, Wessells SM, Bloom PH, Lanzone M, Blackshire S, et al. Improved supervised classification of accelerometer data to distinguish behaviors of soaring birds. *PLoS ONE*. 2017;12(4):e0174785.
- Fehlmann G, O'Riain JM, Hopkins P, O'Sullivan J, Holton M, Shepard E, et al. Identification of behaviours from accelerometer data in a wild social primate. *Animal Biotelem*. 2017;03(5):1–11.
- Shamoun-Baranes J, Bom R, van Loon EE, Ens BJ, Oosterbeek K, Bouten W. From sensor data to animal behaviour: an oystercatcher example. *PLoS ONE*. 2012;7(5):e37997.
- Clarke TM, Whitmarsh SK, Hounslow JL, Gleiss AC, Payne NL, Huveneers C. Using tri-axial accelerometer loggers to identify spawning behaviours of large pelagic fish. *Mov Ecol*. 2021;9(1):26.
- McClune DW, Marks NJ, Wilson RP, Houghton JD, Montgomery IW, McGowan NE, et al. Tri-axial accelerometers quantify behaviour in the Eurasian Badger (*Meles Meles*): towards an automated interpretation of field data. *Animal Biotelem*. 2014;2(1):5.
- Korpela J, Suzuki H, Matsumoto S, Mizutani Y, Samejima M, Maekawa T, et al. Machine learning enables improved runtime and precision for bio-loggers on seabirds. *Commun Biol*. 2020;3(1):633.
- Vehkaoja A, Somppi S, Törnqvist H, Valdeoriola Cardó A, Kumpulainen P, Väättäjä H, et al. Description of movement sensor dataset for dog behavior classification. *Data Brief*. 2022;40:107822.
- Stidsholt L, Johnson M, Beedholm K, Jakobsen L, Kugler K, Brinkløv S, et al. A 2.6-g sound and movement tag for studying the acoustic scene and kinematics of echolocating bats. *Methods Ecol Evol*. 2019;10(1):48–58.
- Friedlaender A, Tyson R, Stimpert A, Read A, Nowacek D. Extreme diel variation in the feeding behavior of humpback whales along the western Antarctic peninsula during autumn. *Mar Ecol Progr Ser*. 2013;494:281–9.
- Ladds M, Thompson A, Slip D, Hocking D, Harcourt R. Seeing it all: evaluating supervised machine learning methods for the classification of diverse otariid behaviours. *PLoS ONE*. 2017;01(11):e0166898.
- Pagano A. Metabolic rate, body composition, foraging success, behavior, and GPS locations of female polar bears (*Ursus maritimus*), Beaufort Sea, spring, 2014–2016 and resting energetics of an adult female polar bear. US Geological Survey data release. 2018. Available from: <https://doi.org/10.5066/F7XW4H0P>.

41. Anguita D, Ghio A, Oneto L, Parra X, Reyes-Ortiz JL. A public domain dataset for human activity recognition using smartphones. *Comput Intell*. 2013;3:3.
42. Russakovsky O, Deng J, Su H, Krause J, Satheesh S, Ma S, et al. ImageNet large scale visual recognition challenge. *Int J Comput Vis*. 2015;115(3):211–52.
43. Tuia D, Kellenberger B, Beery S, Costelloe BR, Zuffi S, Risse B, et al. Perspectives in machine learning for wildlife conservation. *Nat Commun*. 2022;13(1):792.
44. Ng XL, Ong KE, Zheng Q, Ni Y, Yeo SY, Liu J. Animal kingdom: a large and diverse dataset for animal behavior understanding. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*; 2022; p. 19023–34.
45. Chen J, Hu M, Coker DJ, Berumen ML, Costelloe B, Beery S, et al. MammalNet: A large-scale video benchmark for mammal recognition and behavior understanding. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. 2023; p. 13052–61.
46. Garde B, Wilson RP, Fell A, Cole N, Tatayah V, Holton MD, et al. Ecological inference using data from accelerometers needs careful protocols. *Methods Ecol Evol*. 2022;13(4):813–25.
47. López LMM, de Soto NA, Madsen PT, Johnson M. Overall dynamic body acceleration measures activity differently on large versus small aquatic animals. *Methods Ecol Evol*. 2020;13(447):458.
48. Qasem L, Cardew A, Wilson A, Griffiths I, Halsey LG, Shepard EL, et al. Tri-axial dynamic acceleration as a proxy for animal energy expenditure; should we be summing values or calculating the vector? *PLoS ONE*. 2012;7(2):e31187.
49. Otsuka R, Yoshimura N, Tanigaki K, Koyama S, Mizutani Y, Yoda K, et al. Exploring deep learning techniques for wild animal behaviour classification using animal-borne accelerometers. *Methods Ecol Evol*. 2024;15(4):716–31.
50. Chimienti M, Cornulier T, Owen E, Bolton M, Davies IM, Travis MJ, et al. The use of an unsupervised learning approach for characterizing latent behaviors in accelerometer data. *Ecol Evol*. 2016;6(3):727–41.
51. Bidder OR, Campbell HA, Gómez-Laich A, Urgé P, Walker J, Cai Y, et al. Love thy neighbour: automatic animal behavioural classification of acceleration data using the k-nearest neighbour algorithm. *PLoS ONE*. 2014;9(2):e88609.
52. Yu H, Deng J, Nathan R, Kröschel M, Pekarsky S, Li G, et al. An evaluation of machine learning classifiers for next-generation, continuous-ethogram smart trackers. *Mov Ecol*. 2021;9(1):15.
53. Shepard E, Wilson R, Quintana F, Gómez-Laich A, Liebsch N, Albareda D, et al. Identification of animal movement patterns using tri-axial accelerometry. *Endanger Species Res*. 2008;10:47–60.
54. Hammond TT, Springthorpe D, Walsh RE, Berg-Kirkpatrick T. Using accelerometers to remotely and automatically characterize behavior in small animals. *J Exp Biol*. 2016;219(11):1618–24.
55. Resheff YS, Bensch HM, Zöttl M, Harel R, Matsumoto-Oda A, Crofoot MC, et al. How to treat mixed behavior segments in supervised machine learning of behavioural modes from inertial measurement data. *Mov Ecol*. 2024;12(1):44.
56. LeCun Y, Bengio Y, Hinton G. Deep learning. *Nature*. 2015;521(7553):436–44.
57. Aulsebrook AE, Jacques-Hamilton R, Kempnaers B. Quantifying mating behaviour using accelerometry and machine learning: challenges and opportunities. *Animal Behav*. 2024;207:55–76.
58. Jeantet L, Vigon V, Geiger S, Chevallier D. Fully convolutional neural network: a solution to infer animal behaviours from multi-sensor data. *Ecol Model*. 2021;450:109555.
59. Schoombie S, Jeantet L, Chimienti M, Sutton GJ, Pistorius PA, Dufourq E, et al. Identifying prey capture events of a free-ranging marine predator using bio-logger data and deep learning. *Royal Soc Open Sci*. 2024;11(6):240271.
60. Eerdekens A, Deruyck M, Fontaine J, Martens L, Poorter ED, Joseph W. Automatic equine activity detection by convolutional neural networks using accelerometer data. *Comput Electron Agric*. 2020;168:105139.
61. Yuan H, Chan S, Creagh AP, Tong C, Acquah A, Clifton DA, et al. Self-supervised learning for human activity recognition using 700,000 person-days of wearable data. *NPJ Digit Med*. 2024;7(1):91.
62. Zhang Y, Zhang Y, Zhang Z, Bao J, Song Y. Human activity recognition based on time series analysis using U-Net. 2018; [arXiv:1809.08113](https://arxiv.org/abs/1809.08113).
63. Ordóñez F, Roggen D. Deep convolutional and LSTM recurrent neural networks for multimodal wearable activity recognition. *Sensors*. 2016;16(1):115.
64. Chen K, Zhang D, Yao L, Guo B, Yu Z, Liu Y. Deep learning for sensor-based human activity recognition: overview, challenges and opportunities. 2021; [arXiv:2001.07416](https://arxiv.org/abs/2001.07416) [cs]
65. Hammerla NY, Halloran S, Ploetz T. Deep, convolutional, and recurrent models for human activity recognition using wearables. 2016; [arXiv:1604.08880](https://arxiv.org/abs/1604.08880) [cs, stat].
66. Saeed A, Ozcebe T, Lukkien J. Multi-task self-supervised learning for human activity detection. *Proc ACM Interact Mobile Wear Ubiquit Technol*. 2019;3(2):1–30.
67. Yang JB, Nguyen MN, San PP, Li XL, Krishnaswamy S. Deep convolutional neural networks on multichannel time series for human activity recognition. In: *Proceedings of the Twenty-Fourth International Joint Conference on Artificial Intelligence (IJCAI)*. vol. 15. Buenos Aires, Argentina; 2015. p. 3995–4001.
68. Resheff YS, Rotics S, Harel R, Spiegel O, Nathan R. AccelRater: a web application for supervised learning of behavioral modes from acceleration measurements. *Mov Ecol*. 2014;2(1):27.
69. Thiebault A, Huetz C, Pistorius P, Aubin T, Charrier I. Animal-borne acoustic data alone can provide high accuracy classification of activity budgets. *Animal Biotelem*. 2021;9(1):28.
70. Brown T, Mann B, Ryder N, Subbiah M, Kaplan JD, Dhariwal P, et al. Language models are few-shot learners. *Adv Neural Info Process Syst*. 2020;33:1877–901.
71. Caron M, Touvron H, Misra I, Jegou H, Mairal J, Bojanowski P, et al. Emerging properties in self-supervised vision transformers. In: 2021 IEEE/CVF International conference on computer vision (ICCV). 2021; p. 9630–40.
72. Chen T, Kornblith S, Swersky K, Norouzi M, Hinton GE. Big self-supervised models are strong semi-supervised learners. *Adv Neural Info Process Syst*. 2020;33:22243–55.
73. Tong C, Tailor SA, Lane ND. Are accelerometers for activity recognition a dead-end? In: *Proceedings of the 21st international workshop on mobile computing systems and applications*. New York, NY, USA: Association for Computing Machinery. 2020; p. 39–44.
74. Reyes-Ortiz JL, Oneto L, Samà A, Parra X, Anguita D. Transition-aware human activity recognition using smartphones. *Neurocomputing*. 2016;171:754–67.
75. Johnson MP, Tyack PL. A digital acoustic recording tag for measuring the response of wild marine mammals to sound. *IEEE J Ocean Eng*. 2003;28(1):3–12.
76. Rutz C, Troschianko J. Programmable, miniature video-loggers for deployment on wild birds and other wildlife. *Methods Ecol Evol*. 2013;4(2):114–22.
77. Martín López LM, Aguilar de Soto N, Miller P, Johnson M. Tracking the kinematics of caudal-oscillatory swimming: a comparison of two on-animal sensing methods. *J Exp Biol*. 2016;219(14):2103–9.
78. Baglione V, Marcos JM, Canestrari D. Cooperatively breeding groups of carrion crow (*Corvus corone*) in northern Spain. *The Auk*. 2002;119(3):790–9.
79. Wang G. Machine learning for inferring animal behavior from location and movement data. *Ecol Info*. 2019;49:69–76.
80. McClintock B, Johnson D, Hooten M, Ver Hoef J, Morales J. When to be discrete: the importance of time formulation in understanding animal movement. *Mov Ecol*. 2014;2:1–14.
81. Bouthillier X, Delaunay P, Bronzi M, Trofimov A, Nichyporuk B, Szeto J, et al. Accounting for variance in machine learning benchmarks. *Proc Mach Learn Syst*. 2021;3:747–69.
82. Dietterich TG. Approximate statistical tests for comparing supervised classification learning algorithms. *Neural Comput*. 1998;10:1895–923.
83. De Ruiter S, Johnson M, Harris C, Marques T, Swift R, Oh YJ, et al. The animal tag tools project. 2020; <https://animaltags.org/>.
84. Bohoslav JP, Wimalasena NK, Clausen KJ, Dai YY, Yarmolinsky DA, Cruz T, et al. DeepEthogram, a machine learning pipeline for supervised behavior classification from raw pixels. *eLife*. 2021;10:e63377.
85. Paszke A, Gross S, Massa F, Lerer A, Bradbury J, Chanan G, et al. Pytorch: An imperative style, high-performance deep learning library. *Adv Neural Info Process Syst*. 2019;32.

86. Pedregosa F, Varoquaux G, Gramfort A, Michel V, Thirion B, Grisel O, et al. Scikit-learn: machine learning in Python. *J Mach Learn Res*. 2011;12:2825–30.
87. Kingma DP, Ba J. Adam: A method for stochastic optimization. In: 3rd International conference on learning representations (ICLR); 2015.
88. Loshchilov I, Hutter F. SGDR: Stochastic gradient descent with warm restarts. In: 5th International conference on learning representations (ICLR); 2017.
89. He K, Zhang X, Ren S, Sun J. Deep residual learning for image recognition. In: Proceedings of the IEEE conference on computer vision and pattern recognition. 2016; p. 770–8.
90. Resheff YS, Bensch HM, Zöttl M, Rotics S. Correcting a bias in the computation of behavioural time budgets that are based on supervised learning. *Methods Ecol Evol*. 2022;13(7):1488–96.
91. Weinstein BG, Marconi S, Graves SJ, Zare A, Singh A, Bohlman SA, et al. Capturing long-tailed individual tree diversity using an airborne imaging and a multi-temporal hierarchical model. *Remote Sens Ecol Conserv*. 2023;9(5):656–70.
92. Dickinson ER, Twining JP, Wilson R, Stephens PA, Westander J, Marks N, et al. Limitations of using surrogates for behaviour classification of accelerometer data: refining methods using random forest models in Caprids. *Mov Ecol*. 2021;9(1):28.
93. Hussey NE, Kessel ST, Aarestrup K, Cooke SJ, Cowley PD, Fisk AT, et al. Aquatic animal telemetry: a panoramic window into the underwater world. *Science*. 2015;348(6240):1255642.
94. Kays R, Crofoot MC, Jetz W, Wikelski M. Terrestrial animal tracking as an eye on life and planet. *Science*. 2015;348(6240):aaa2478.
95. Pang G, Shen C, Cao L, Hengel AVD. Deep learning for anomaly detection: a review. *ACM Comput Surv*. 2021;54(2):1–38.
96. Jetz W, Tertitski G, Kays RW, Mueller U, Wikelski M, Åkesson S, et al. Biological earth observation with animal sensors. *Trends Ecol Evol*. 2022;37(4):293–8.
97. Henderson P, Hu J, Romoff J, Brunskill E, Jurafsky D, Pineau J. Towards the systematic reporting of the energy and carbon footprints of machine learning. *J Mach Learn Res*. 2020;21(248):1–43.
98. Raji D, Denton E, Bender EM, Hanna A, Paullada A. AI and the everything in the whole wide world benchmark. In: Proceedings of the neural information processing systems: Track on datasets and benchmarks. 2021; vol. 1.
99. Glass TW, Breed GA, Robards MD, Williams CT, Kielland K. Accounting for unknown behaviors of free-living animals in accelerometer-based classification models: demonstration on a wide-ranging mesopredator. *Ecol Info*. 2020;60:101152.
100. Dehghani M, Tay Y, Gritsenko AA, Zhao Z, Houlsby N, Diaz F, et al. The benchmark lottery. 2021; [arXiv:2107.07002](https://arxiv.org/abs/2107.07002).

Publisher's Note

Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.