

Machine learning for ecosystem services

Simon Willcock^{a,b,*}, Javier Martínez-López^c, Danny A.P. Hooftman^{d,e}, Kenneth J. Bagstad^f, Stefano Balbi^c, Alessia Marzo^g, Carlo Prato^g, Saverio Sciandrello^g, Giovanni Signorello^g, Brian Voigt^h, Ferdinando Villa^{c,i}, James M. Bullock^e, Ioannis N. Athanasiadis^j

^a School of Environment, Natural Resources and Geography, Bangor University, United Kingdom

^b Biological Sciences, University of Southampton, United Kingdom

^c Basque Centre of Climate Change, Spain

^d Lactuca: Environmental Data Analyses and Modelling, The Netherlands

^e NERC Centre for Ecology and Hydrology, Wallingford, United Kingdom

^f Geosciences & Environmental Change Science Center, U.S. Geological Survey, Denver, CO, USA

^g Centre for the Conservation and Management of Nature and Agroecosystems (CUTGANA), University of Catania, Catania, Italy

^h Gund Institute, Rubenstein School of Environment and Natural Resources, University of Vermont, Burlington, VT, USA

ⁱ IKERBASQUE, Basque Foundation for Science, Bilbao, Bizkaia, Spain

^j Information Technology Group, Wageningen University, The Netherlands

ARTICLE INFO

Article history:

Received 29 September 2017

Received in revised form 1 March 2018

Accepted 11 April 2018

Available online 5 May 2018

Keywords:

ARIES

Artificial intelligence

Big data

Data driven modelling

Data science

Machine learning

Mapping

Modelling

Uncertainty, Weka

ABSTRACT

Recent developments in machine learning have expanded data-driven modelling (DDM) capabilities, allowing artificial intelligence to infer the behaviour of a system by computing and exploiting correlations between observed variables within it. Machine learning algorithms may enable the use of increasingly available 'big data' and assist applying ecosystem service models across scales, analysing and predicting the flows of these services to disaggregated beneficiaries. We use the Weka and ARIES software to produce two examples of DDM: firewood use in South Africa and biodiversity value in Sicily, respectively. Our South African example demonstrates that DDM (64–91% accuracy) can identify the areas where firewood use is within the top quartile with comparable accuracy as conventional modelling techniques (54–77% accuracy). The Sicilian example highlights how DDM can be made more accessible to decision makers, who show both capacity and willingness to engage with uncertainty information. Uncertainty estimates, produced as part of the DDM process, allow decision makers to determine what level of uncertainty is acceptable to them and to use their own expertise for potentially contentious decisions. We conclude that DDM has a clear role to play when modelling ecosystem services, helping produce interdisciplinary models and holistic solutions to complex socio-ecological issues.

© 2018 The Authors. Published by Elsevier B.V. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

1. Introduction

Many scientific disciplines are taking an increasingly integrative approach to planetary problems such as global climate change, food security and human migration (Baziliana et al., 2011; Bullock et al., 2017). To address such challenges, methods and practices are becoming more reliant on large, interdisciplinary data

repositories often collected in cutting-edge ways, for example via citizen scientists or automated data collection (Isaac et al., 2014). Recent developments in information technology have expanded modelling capabilities, allowing researchers to maximise the utility of such 'big data' (Lokers et al., 2016). Here, we focus on one of these developments: data-driven modelling (DDM). DDM is a type of empirical modelling by which the data about a system are used to create models, which use observed systems' states as inputs for estimating some other system state(s), i.e., outputs (Jordan and Mitchell, 2015; Witten et al., 2016). Thus, DDM is the process of identifying useful patterns in data, a process sometimes previously referred to as knowledge discovery in databases (Fayyad et al., 1996). This process consists of five key steps: 1) understanding the research goal, 2) selecting appropriate data, 3) data cleaning, pre-processing and transformation, 4) data mining

* Corresponding author at: School of Environment, Natural Resources and Geography, Bangor University, United Kingdom.

E-mail addresses: s.willcock@bangor.ac.uk (S. Willcock), javier.martinez@bc3research.org (J. Martínez-López), danny.hooftman@lactuca.nl (D.A.P. Hooftman), kjbagstad@usgs.gov (K.J. Bagstad), stefano.balbi@bc3research.org (S. Balbi), alessia.marzo@unict.it (A. Marzo), s.sciandrello@unict.it (S. Sciandrello), g.signorello@unict.it (G. Signorello), Brian.Voigt@uvm.edu (B. Voigt), ferdinando.villa@bc3research.org (F. Villa), jmbul@ceh.ac.uk (J.M. Bullock), ioannis@athanasiadis.info (I.N. Athanasiadis).

(creating a data driven model), and 5) interpretation/evaluation (Fayyad et al., 1996) (Fig. 1). A variety of methods for data mining and analysis are available, some of which utilise machine learning algorithms (Witten et al., 2016; Wu et al., 2014) (Fig. 1). A machine learning algorithm is a process that is used to fit a model to a dataset, through training or learning. The learned model is subsequently used against an independent dataset, in order to determine how well the learned model can generalise against the unseen data, a process called testing (Ghahramani, 2015; Witten et al., 2016). This training–testing process is analogous to the calibration–validation process associated with many process-based models.

In general, machine learning algorithms can be divided into two main groups (supervised- and unsupervised-learning; Fig. 1), separated by the use of explicit feedback in the learning process (Blum and Langley, 1997; Russell and Norvig, 2003; Tarca et al., 2007). Supervised-learning algorithms use predefined input-output pairs and learn how to derive outputs from inputs. The user specifies which variables (i.e., outputs) are considered dependent on others (i.e., inputs), which sometimes indicates causality (Hastie et al., 2009). The machine learning toolbox includes several linear and non-linear supervised learners, predicting either numeric outputs (regressors) or nominal outputs (classifiers) (Table 1). An example of supervised machine learning that is familiar to many ecosystem service (ES) scientists is using a general linear model, whereby the user provides a selection of input variables hypothesised to predict values of an output variable and the general linear model learns to reproduce this relationship. The learning process needs to be fine-tuned through a process, as for example in the case of stepwise selection where an algorithm selects the most parsimonious best-fit model (Yamashita et al., 2007). However, note that

Table 1

A simplified summary of machine learning algorithms (categorised as supervised and unsupervised).

Category	Task	Common algorithms
Unsupervised learning (learning without feedback from a trainer)	Clustering Associations Dimensionality reduction	k-means Apriori PCA
Supervised learning (learning past actions/decisions with trainer)	Classification (learning a categorical variable) Regression (learning a continuous variable)	Bayesian Networks, Decision Trees, Neural Networks Linear Regression, Perceptron

stepwise functions may also be used in unsupervised learning processes when combined with other methods. Within unsupervised-learning processes, there is no specific feedback supplied for input data and the machine learning algorithm learns to detect patterns from the inputs. In this respect, there are no predefined outputs, only inputs for which the machine learning algorithm determines relationships between them (Mjolsness and DeCoste, 2001). An example unsupervised-learning algorithm, cluster analysis, groups variables based on their closeness to one another, defining the number and composition of groups within the dataset (Mouchet et al., 2014). Within the supervised- and unsupervised-learning categories, there are several different varieties of machine learning algorithms, including: neural networks, decision trees, decision rules and Bayesian networks. Others have described the varieties of machine learning algorithms (Blum and Langley, 1997;

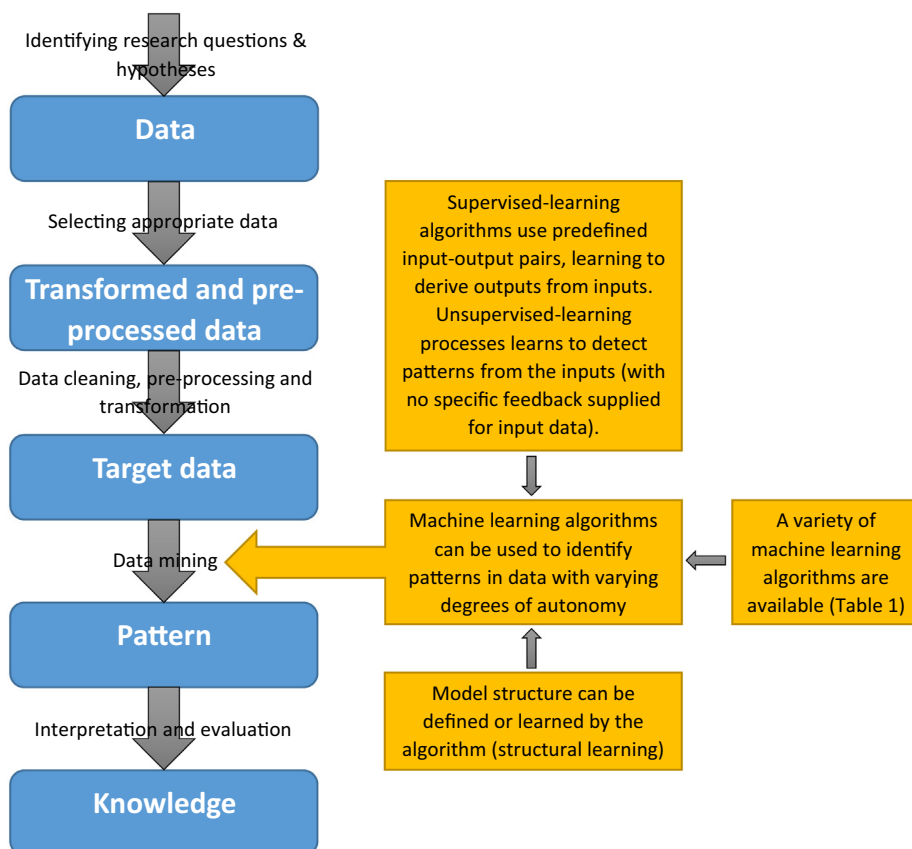


Fig. 1. A schematic outlining how machine learning algorithms (yellow) can contribute to the data-driven modelling process (blue) (Fayyad et al., 1996). (For interpretation of the references to colour in this figure legend, the reader is referred to the web version of this article.)

Mjolsness and DeCoste, 2001; Russell and Norvig, 2003; Tarca et al., 2007) and so we only provide a brief summary here, leaving out more advanced methods such as reinforcement learning, and deep learning (see Table 1).

DDM undoubtedly has a role to play when modelling socio-ecological systems and assessing ES. DDM can give useful predictive insight into areas where understanding of the underlying processes is limited. However, as with many statistical methods, DDM requires adequate data availability. The level of data required is determined on a case-by-case basis, depending of the research question being asked. For example, to use machine learning algorithms, data must be able to be divided into training and testing subsets (Smith and Frank, 2016). Machine learning algorithms assume considerable changes in the modelled system have not taken place during the time period covered by the model (Ghahramani, 2015; Jordan and Mitchell, 2015), though machine learning can also be used for identifying change, i.e., detecting concept drift (Gama et al., 2004). Model validation/testing, which has yet to become standard practice within the ES modelling community (Baveye, 2017; Hamel and Bryant, 2017), is an integral part of the machine learning process within DDM. This is vital as DDM can result in overfitting, which occurs when the model learns the training data well (i.e., a close fit to the training data), but performs poorly on independent test data (Clark, 2003).

To assess the quality of the learning process, machine learning algorithms use various methods (summarised in Witten et al. (2016)) to ensure that the results are generalizable and avoid overfitting. For example, *k*-fold cross validation allows for fine-tuning of model performance (Varma and Simon, 2006; Wiens et al., 2008). This approach maximises the data availability for model training by dividing the data into *k* subsets and using *k*-1 subsets to train the model whilst retaining a subset for independent validation. This process is repeated *k* times so that all available data have been used for validation exactly once. The results of the *k*-folds are then combined to produce metrics of quality for the machine learning process, often accompanied with an estimation of the model uncertainty (i.e., the cross-validation statistic). Whilst the goodness-of-fit parameter used varies within DDM (e.g., root mean square error is used extensively within regression models, but the standard error is more commonly used in Bayesian machine learning (Cheung and Rensvold, 2002; Uusitalo, 2007)), it provides the user with a transparent estimate of model uncertainty. Whilst estimates of uncertainty are useful, users of DDM should be aware that such models do not represent the underlying processes within socio-ecological systems, but instead capture relationships between variables (Ghahramani, 2015). However, for some datasets and model applications (see Section 4 for further details), DDM can produce more accurate models than process-based models, as the latter may suffer from an incomplete representation of the socio-ecological processes (Jordan and Mitchell, 2015; Tarca et al., 2007). Finally, as with any modelling, DDM depends on the quality of the training and testing datasets used; whilst some extreme cases or outliers might get ignored during DDM, the quality of the information supplied to the machine learning algorithms should be verified beforehand (Galelli et al., 2014).

The aim of this paper is to demonstrate the utility of DDM to the ES community. We present two examples of DDM using Bayesian networks (a supervised learning technique), as implemented in the Waikato Environment for Knowledge Analysis machine learning software (Weka; <http://www.cs.waikato.ac.nz/ml/weka/>; Frank et al. (2016); Hall et al. (2009)), used both standalone and as part of the Artificial Intelligence for Ecosystem Services (ARIES; <http://aries.integratedmodelling.org/>; Villa et al. (2014)) modelling platform. We chose Bayesian network methods as uncertainty metrics describing both the model fit and the grid-cell uncertainty can be calculated (Aguilera et al., 2011; Landuyt et al., 2013; Uusitalo,

2007). Our Weka example focusses on firewood use in South Africa, and is comparable to conventional ES models recently published by Willcock et al. (submitted for publication). Using ARIES, we model biodiversity value within Sicily, and demonstrate how DDM can make use of volunteered geographical information by incorporating data from Open Street Maps into the machine learning process. In both examples, we highlight how model structure and uncertainty computed in the machine learning process supplement and enhance the value of the results reported to the user.

2. Methods

For the first example, we used Weka, an open-source library of machine learning algorithms (Frank et al., 2016; Hall et al., 2009), to create a model capable of identifying the upper quartile of sites for firewood use in South Africa. We chose this example as: 1) firewood use is of high policy relevance in sub-Saharan Africa (Willcock et al., 2016); 2) robust spatial data on firewood use are available within South Africa and may, for some municipalities, provide a comparable context to other parts of sub-Saharan Africa, which are often more vulnerable but data deficient (Hamann et al., 2015); 3) models ranking the relative importance of different sites were rated as useful to support ES decision-making by nearly 90% of experts in sub-Saharan Africa (Willcock et al., 2016); and 4) multiple conventional models have recently been run for this ES covering this spatial extent (see Willcock et al. (submitted for publication) for full details).

The firewood use data are freely available (Hamann et al., 2015) and are based on the South African 2011 population census, which provides proportions of households per local municipality using a specific ES (similar data are available for a set of other ES; see www.statssa.gov.za for all 2011 census output). For this paper we used the proportion of households that use collected firewood as a resource for cooking (Hamann et al., 2015). To derive a measure of total resource use, we multiplied the proportion of use by the 2011 official census municipal population size (from www.statssa.gov.za) as: $[(\% \text{ households using a service}) \times (\text{municipal population size})]$. We then divided this value by the area of each local municipality to provide an estimate of firewood use density, ensuring that model inputs are independent of the land area of the local municipality.

To utilise Bayesian networks, the decision variable (firewood use density) had to be converted into a categorical (nominal) attribute; note, the categories created during this process are unordered. The goal of this task was to predict the areas in the upper quartile, reflecting demand from decision-makers for identification of the most important sites for ES production and, once identified, enabling these areas to be prioritised for sustainable management (Willcock et al., 2016). Thus, the firewood use density data were categorised within the highest 25% (Q4) and the lowest 75% (Q1–Q3) quartiles using Weka's *Discretize* filter to create ranges of equal frequencies (four in our case). Out of the generated quartiles, the three lower ones were merged with the *MergeTwoValues* filter. To ensure like-for-like comparisons between our DDM and conventional models, we provided the machine learning algorithms with the same user supplied input data used to model firewood within Willcock et al. (submitted for publication) (Table 2). Since most Bayesian network inference algorithms can use only categorical data as inputs, the input data were discretised by grouping their values in five bins of equal frequencies. Selecting the number of bins is a design choice and may impact model output (Friedman and Goldszmidt, 1996; Nojavan et al., 2017). As such, the sensitivity of the modelled output to variable bin numbers warrants future investigation, but is beyond the scope of this first-order introduction to machine learning for ES.

Table 2

The municipal-scale inputs into the Weka machine learning algorithms to estimate firewood use in South Africa. Overfitting is avoided by first training the algorithm on subset of these data and then testing against the remaining data.

Attribute	Description
LCAgriculture	The proportion of agricultural land area, derived from GeoTerralimage (2015)
LCForest	The proportion of forested land area, derived from GeoTerralimage (2015)
LCGrassland	The proportion of grassland land area, derived from GeoTerralimage (2015)
LCUrban	The proportion of urban areas, derived from GeoTerralimage (2015)
LCWater	The proportion of water bodies area, derived from GeoTerralimage (2015)
COFirewood	The proportion of area on which firewood can be produced (Forest, Woodland, Savanna), derived from GeoTerralimage (2015)
OProtected	The proportion of protected natural areas, derived from the World Database on Protected Areas (www.protectedplanet.net)
MOCarbon	Mean amount of carbon stored per hectare, as calculated in Willcock et al. (submitted for publication)
OGrowthDay	Average number of growing days in the area as driven by the relationship between rainfall and evapotranspiration, as calculated in Willcock et al. (submitted for publication)
ZScholesA	A metric of the nutrient-supplying capacity of the soil (Scholes, 1998)
ZScholesB	A metric of the nutrient-supplying capacity of the soil (Scholes, 1998)
ZScholesD	Scholes (1998) land use correction, as calculated in Willcock et al. (submitted for publication)
ZSlope	This is the mean slope in the area, based on the global 90-m digital elevation model downloaded from CGIAR-CSI (srtm.csi.cgiar.org/SELECTION/inputCoord.asp)
Population_density	The municipal population based on the South African 2011 census (www.statssa.gov.za)
Firewood_density	Observed firewood use for cooking from the South African 2011 census (Hamann et al., 2015)

We used the *BayesNet* implementation of Weka to train our DDM. The machine learning algorithm can construct the Bayesian network using alternative network structures and estimators for finding the conditional probability tables ([Chen and Pollino, 2012](#)). In a Bayesian network, conditional probability tables define the probability distribution of output values for every possible combination of input variables ([Aguilera et al., 2011](#); [Landuyt et al., 2013](#)). Unlike the use of expert elicitation or Bayesian network training (e.g., [Marcot et al. \(2006\)](#)), the machine learning approach fits the *structure of the model*, as well as the conditional probabilities, a process also called structural learning ([Fig. 1](#)). In this example, we evaluated 16 alternatives for parameterising the Bayesian network learning (see [Appendix 1](#)). We used 10 cross-fold validation ([Varma and Simon, 2006](#); [Wiens et al., 2008](#)), repeated 10 times with different seeds, for creating the random folds.

ARIES has recently incorporated the Weka machine learning algorithms into its modelling framework, with the aim of enabling use of DDM within the ES community (see [Villa et al. \(2014\)](#) for a description of the ARIES framework). In our second example, we used the ARIES implementation of Weka *BayesNet* to propagate site-based expert estimates of 'biodiversity value' and so build a map for the entire Sicilian region ([Li et al., 2011](#)). Here, biodiversity value does not refer to an economic value, but to a spatially explicit relative ranking. The original biodiversity value observations were the result of assessments made with multiple visits by flora, fauna and soil experts ([Fig. 2](#)). The same experts who had ranked high-value sites were asked to identify sites of low biodiversity value, with the constraint that the low value depended on natural factors and not on human intervention, as datasets combining high and

low value observations generally produce more accurate models ([Liu et al., 2016](#)). These data were originally interpolated using an inverse distance weighted technique to provide a map of biodiversity value to support policy- and decision-making ([Fig. 2a](#)), and our DDM attempts to improve on this map. The DDM process involved 20 repetitions, each using 75% of the data to train the model and 25% to validate it. Using ARIES, we instructed the machine learning algorithm to access explanatory variables, indicated by the same experts who provided the estimates used in training as the most likely predictors of biodiversity value (see [Appendix 2](#)). The data used by the machine learning process ([Appendix 2](#)) included distance to coastline and primary roads metric calculated using citizen science data from Open Street Map (<https://www.openstreetmap.org/>; [Haklay and Weber \(2008\)](#)). The trained model was then used to build a map of biodiversity value for the entire island, computing the distribution of biodiversity values for all locations not sampled by the experts. The machine learning algorithms used quantitative variables, discretised in 10 equal intervals, for both inputs and outputs ([Friedman and Goldszmidt, 1996](#); [Nojavan et al., 2017](#)). The resulting map was subsequently discussed and qualitatively validated by the same experts who collected the data, as well as quantitatively using a confusion matrix accuracy assessment.

3. Results

In the first example, the results for all configurations of the DDM created for firewood use in South Africa had a classification accuracy above 80% (see [Appendix 1](#)). The model predictions are statistically significant with a confidence level of 0.05 (two tailed) when compared to the ZeroR classifier (a baseline classifier that always predicts the majority class). Using ArcGIS v 10.5.1, we spatially mapped the outputs of the most accurate Bayesian network DDM ([Figs. 3, 4](#); [Appendix 3](#)). The confusion matrix for this model shows that 186 out of the 226 local municipalities were correctly classified (an overall classification accuracy of 82%), and, out of 56 municipalities classified in the upper quartile (Q4), 36 were correct predictions (64% recall [i.e. the percentage of the most important sites for firewood ES correctly identified], comparable with conventional modelling methods evaluated against independent data [[Table 3](#); [Willcock et al. \(submitted for publication\)](#)]; [Appendix 3](#)). The DDM also produces probabilistic outputs for the respective inputs ([Appendix 4](#)).

For biodiversity value in Sicily, 43% of the testing subsample was correctly classified into 1 of 10 biodiversity value categories, with a majority of the incorrectly classified results falling into immediately close numeric ranges ([Appendix 5](#)). During a workshop in June 2017, the same Sicilian experts that provided the training set (a team of five including an academic conservationist, an academic ornithologist, an academic botanist and an expert on agricultural biodiversity) qualitatively evaluated the output in non-sampled but well-known regions and deemed it a distinct improvement on previously computed biodiversity value assessments, built through conventional GIS overlapping and interpolation techniques; an assessment that was embraced by other participants from both local governmental and conservation institutions ([Fig. 2](#)). As the map reflects the human assessment of biodiversity value rather than objective measurements, the consensus of experts and practitioners was deemed equivalent to a satisfactory validation. The confusion matrix ([Appendix 5](#)) shows how the majority of misclassifications are between similar value categories. For example, 73% of test data were predicted within one class above or below their actual class, and 84% of test data were correctly classified within two classes above and below their actual class. A Spearman Rho test

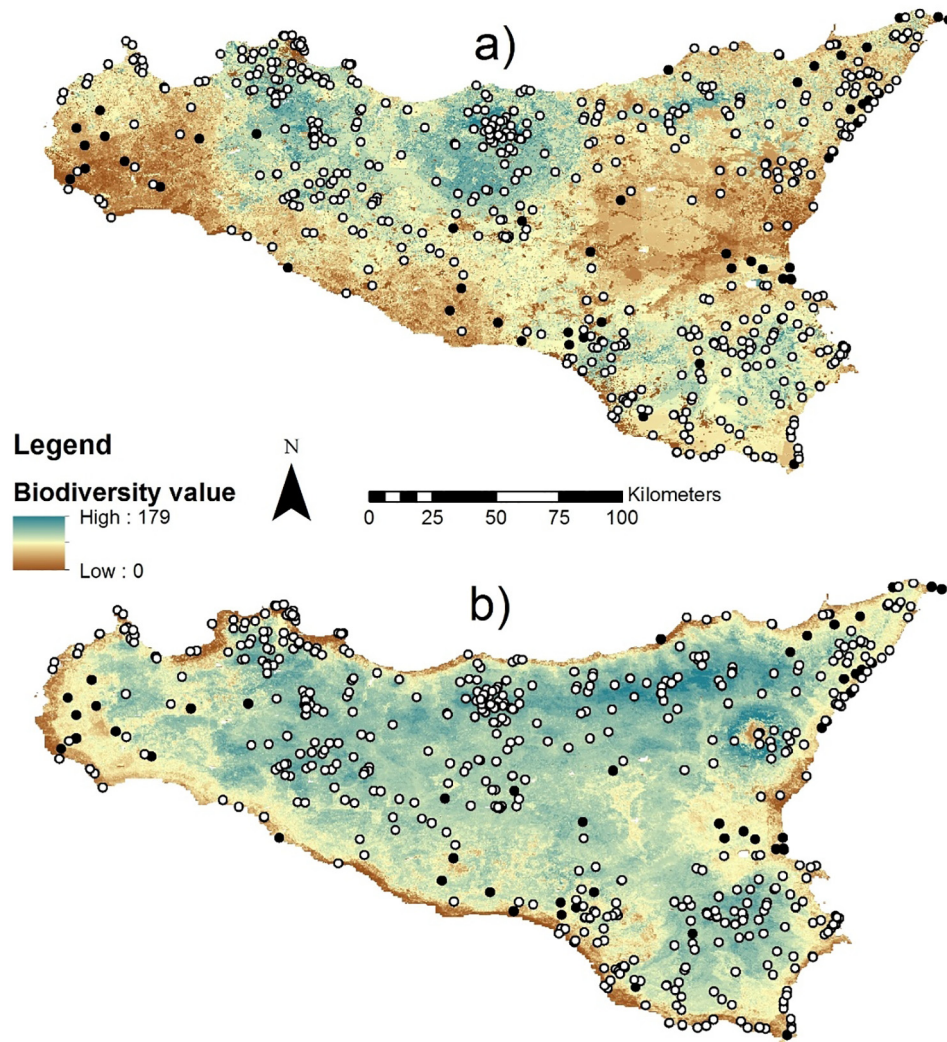


Fig. 2. The relative value of terrestrial biodiversity in Sicily estimated by a) inverse distance weighted interpolation of observed values and b) Bayesian networks using data-driven modelling. Both original (white) biodiversity value observations and the additional sites of low biodiversity value (black) are shown as points.

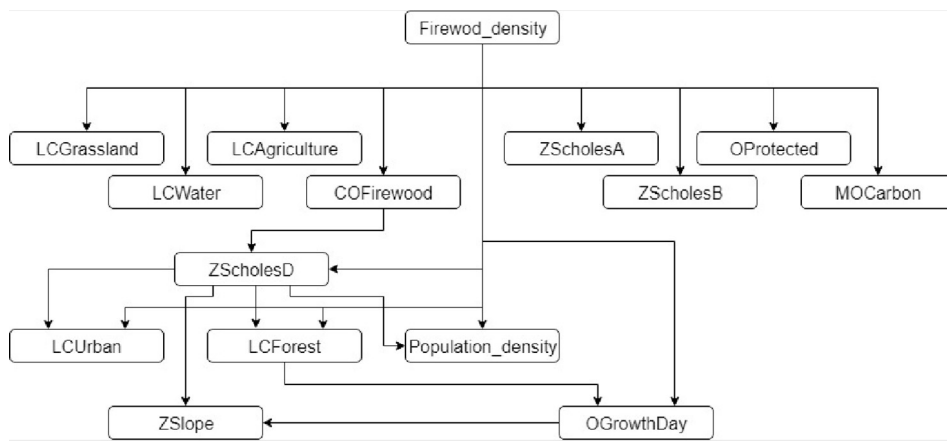


Fig. 3. Diagrammatic representation of the machine-learned Bayesian network model of firewood use in South Africa (see Table 2 for category codes). The structure of the model was informed by the machine learning algorithm with no predetermined restrictions.

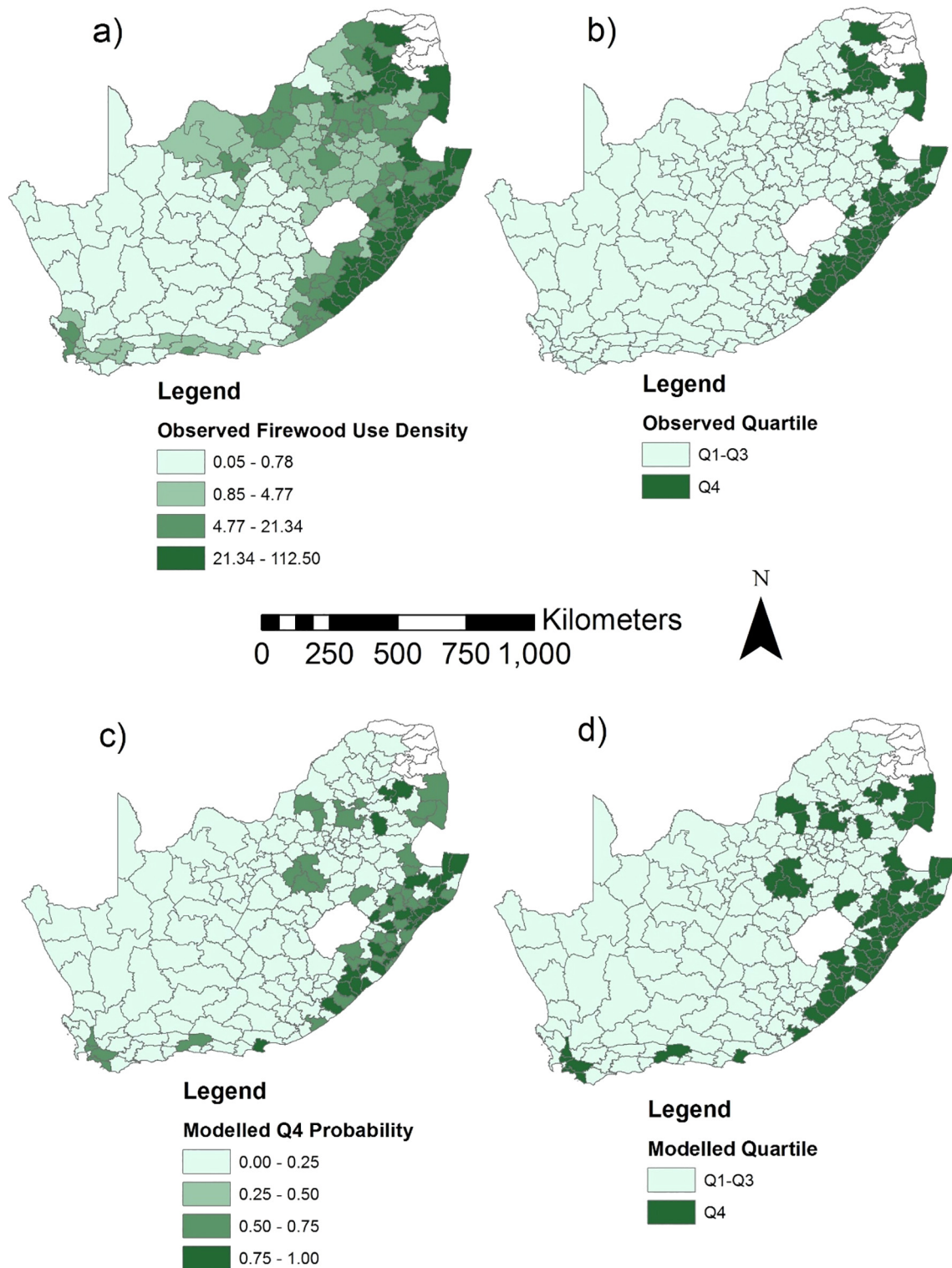


Fig. 4. Observed (a and b) and modelled (c and d) data on firewood use density within South Africa. The Weka *BayesNet* DDM process derives a probabilistic output (c) from the observed data (a). The modelled output can be categorised into quartiles (Q1–4, with Q4 being the upper quartile; d) and compared to the observed data within the same categories (b).

highlights the significant correlation between the ranked model and validation data categories ($Rho: 0.58$; $p\text{-value} < 0.001$). The root-mean-squared error of the model prediction was also computed and resulted in a value of 0.26 (Hyndman and Koehler, 2006).

4. Discussion

Lack of credibility, salience and legitimacy are the major reasons for the ‘implementation gap’ between ES research and its incorporation into policy- and decision-making (Clark et al.,

Table 3

Comparing recall of DDM outputs with conventional models when producing estimates of firewood use in South Africa. Outputs from conventional models of varying complexity were validated using independent data (see Willcock et al. (submitted for publication) for full model descriptions and model complexity analysis). DDM outputs were validated using k-fold cross validation (see Section 2).

Model	Model criteria	Recall for the upper quartile of firewood use (%)
Bayesian network within Weka (Frank et al., 2016; Hall et al., 2009)	Assignment threshold = 50%	64.3
	Assignment threshold = 75%	90.9
Conventional model A (Complexity score: 2; Willcock et al. (submitted for publication)) [*]	Gridcell size = 1 km	75.0
	Gridcell size = 10 km	73.2
Conventional model B (Complexity score: 4; Willcock et al. (submitted for publication)) [*]	Gridcell size = 1 km	75.0
	Gridcell size = 10 km	76.8
Conventional model C (Complexity score: 4; Willcock et al. (submitted for publication)) [*]	Gridcell size = 1 km	60.7
	Gridcell size = 10 km	60.7
Conventional model D (Complexity score: 36; Willcock et al. (submitted for publication)) [*]	Gridcell size = 55.6 km	76.8
Conventional model A (Complexity score: 31; Willcock et al. (submitted for publication)) [*]	Gridcell size = 5 km	53.6

^{*} Models have been anonymised as identification of the best specific model for a particular use is likely to be location specific and may shift as new models are developed (Willcock et al., submitted for publication).

2016; Olander et al., 2017; Wong et al., 2014). A lack of uncertainty information and the inability to run models in data-poor environments and/or under conditions where underlying processes are poorly understood may contribute to the implementation gap. However, DDM can help to address these current shortcomings in ES modelling. Here, we have demonstrated that DDM is feasible within ES science and is capable of providing estimates of uncertainty.

For our South African case study, the machine learning algorithms were able to produce a modelled output of comparable accuracy to conventional modelling methods when using the same input variables, despite our DDM using data at a much coarser (local municipality) scale (Table 3). Using the spatially attributed uncertainty (i.e., the probability of each local municipality being in Q4), decision-makers would be able to set their own level of acceptable uncertainty. In our example, since we have two categorical bins (i.e., Q1–3 and Q4), any local municipality with a modelled Q4 probability over 0.5 is assigned to the Q4 category. This assignment threshold can be varied; e.g., it is possible to state that municipalities where modelled Q4 probability is less than 0.25 or greater than 0.75 are likely to be grouped within Q1–3 and Q4 respectively, and to admit that we are less certain for the remaining municipalities. In our example, this would result in a 96% (135 out of 140) categorisation accuracy for Q1–3 and a 91% (30 out of 33) categorisation accuracy for Q4, with 53 local municipalities left uncategorised due to uncertainty.

Thus, using Bayesian networks and machine learning, we are able to convey to decision-makers not only which sites show the highest ES use or value, but also how confident we are in our estimate at each site (Aguilera et al., 2011; Chen and Pollino, 2012; Landuyt et al., 2013). This information allows decision-makers to 1) apply an assignment threshold of their choosing to the modelled output before making a policy- or management-decision, and 2) use their own judgement for potentially contentious decisions, where uncertainty is higher (Olander et al., 2017). For example, whilst it is perhaps obvious that sites where we are highly certain that there is high ES value should be appropriately managed, it is unclear which sites should be the next highest management priority. Given a limited budget, is a medium-ES value site with high certainty more or less worthy of management than a potentially high-value site with medium or low certainty? Decision-makers show both capacity and willingness to engage with the uncertainty information should these data be made available (McKenzie et al., 2014; Scholes et al., 2013; Willcock et al., 2016), even when results may indicate high levels of uncertainty. This is illustrated by a Sicilian case study, in which decision-makers, when advised of the relatively low overall classification accuracy (43%), accepted

it as predictions were close to their actual value (i.e. 73% of test data were predicted within one class above or below their actual class) and were viewed as an improvement on previous estimates (Fig. 2). Thus, providing estimates of uncertainty should become standard practice within the ES community (Hamel and Bryant, 2017).

There are both advantages and disadvantages to using machine learning algorithms for the ‘data mining’ step of DDM (Fayyad et al., 1996). As highlighted above, machine learning algorithms provide indications of uncertainty that could usefully support decision-making. However, similar uncertainty metrics can also be obtained using conventional modelling (i.e., via the confidence intervals surrounding regressions (Willcock et al., 2014) or Bayesian belief networks (Balbi et al., 2016)). Similar to conventional modelling, the performance of model algorithms substantially depends on the parameters, model structure and algorithm settings applied (Zhang and Wallace, 2015). For example, many machine learning algorithms require categorical data and so potentially an additional step of data processing whereby continuous data are discretised. In our South African case study, we divided firewood use data into five bins but acknowledge that the number of bins may affect model performance and the impact of this warrants further investigation (Friedman and Goldszmidt, 1996; Nojavan et al., 2017; Pradhan et al., 2017). However, a variety of machine learning algorithms are available (Table 1) and not all of them required discretised data (Jordan and Mitchell, 2015; Witten et al., 2016). Furthermore, for our firewood models, we used machine learning to create the model structure. Structural learning can yield better performing models (i.e., all our South African model configurations had a classification accuracy above 80%; Appendix 1) and may highlight relationships that have not yet been theorised (or have previously been discarded) (Gibert et al., 2008; Suominen and Toivanen, 2016). However, the obtained structures (Fig. 3) may not be causal and could confuse end-users (Schmidhuber, 2015). Thus, predefined network structures may be preferred for applications where causality is particularly important. Further generalisations useful for ES modellers considering machine learning algorithms include the following: 1) Multi-classification problems may have lower accuracy – as highlighted by comparing our South African (2 category output, 82% accuracy) and Sicilian (10 category output, 43% accuracy) examples – the more categories in the modelled output, the lower the apparent accuracy. Thus, the number of categories in the output should be considered when interpreting the model accuracy metric. For example, a random model with a two category output and a four category output will be accurate 50% and 25% of the time respectively. Thus, a machine-learned model with an accuracy of 40% is

poor if the output had two categories, but learned more (and so might be of more use) if a four category output was being considered; 2) Supervised learning can be used when drivers are known – for example, with no *a priori* assumptions, unsupervised learning could cluster beneficiaries into groups, but these may not match known beneficiary groups (i.e., livelihoods) and so might be difficult to interpret (Schmidhuber, 2015). Supervised learning can be used to align the outputs from machine learning algorithms with decision-maker specified beneficiary groups; 3) machine learning algorithms are best applied to the past and present, but not the future – Although machine learning algorithms can detect strong relationships, accurately describing past events and providing useful predictions where process-based understanding is lacking (Jean et al., 2016), the relationships identified may not be causally linked and so may not hold when extrapolating across space or time (Mullainathan and Spiess, 2017). Thus, where the process is well understood, DDM is unlikely to be more appropriate than conventional process-based models (Jordan and Mitchell, 2015). Understanding the caveats and limitations of machine learning algorithms is important before the algorithms are used for DDM.

A further critique of DDM is that it can appear as a ‘black box’ in which the machine learning processes are not clear to the user and so they could widen the implementation gap (Clark et al., 2016; Olander et al., 2017; Wong et al., 2014). However, we have demonstrated that utilisation of machine learning algorithms can be transparent and replicable. For example, Bayesian networks allow the links between data to be visualised (Fig. 3) (Aguilera et al., 2011; Chen and Pollino, 2012; Landuyt et al., 2013). The standalone Weka software is user friendly and requires minimal expertise, and ease of use has been further simplified within the ARIES software as DDM can be run merely by selecting a spatiotemporal modelling context and then using the ‘drag-drop’ function to start the machine learning process (Villa et al., 2014). Machine learning and machine reasoning (Bottou, 2014) are facilitated within the ARIES system through semantic data annotation, which makes data and models machine readable and allows for automated data selection and acquisition from cloud-hosted resources, as well as automated model building (Villa et al., 2017). To ensure that this complex process remains transparent, the Bayesian network is described using a provenance diagram (Fig. S2), characterising the DDM process, i.e., which data and models were selected by ARIES (Fig. 1). Furthermore, work has begun to enable the ARIES software to produce automated reports that describe the DDM process and modelling outputs in readily understandable language (see Appendix 2 for a preliminary automated report for the ARIES example used in this study). Advances such as this may enable decision-makers to run and interpret ES models with minimal support from scientists, potentially increasing ownership in the modelled results and closing the implementation gap (Olander et al., 2017).

The DDM process encourages scientists to use as much data as possible to generate the highest quality knowledge. Machine learning algorithms provide a tool by which ‘big data’ can be incorporated into ES assessments (Hampton et al., 2013; Lokers et al., 2016; Richards and Tunçer, 2018). For example, using the ARIES software, we demonstrated how Open Street Map data can be included in the machine learning process (Haklay and Weber, 2008). Whilst future research is needed to determine how much data is actually needed, it is clear that ES scientists must contribute to and make use of large datasets to participate in the information age (Hampton et al., 2013), particularly where data are standardised and made machine-readable (Villa et al., 2017). Using machine learning algorithms to interpret big data may help provide a wide range of ES information across the variety of temporal and spatial scales required by decision-makers (McKenzie et al., 2014; Scholes et al., 2013; Willcock et al.,

2016). There has been a recent call-to-arms within the ES modelling community to shift focus from models of biophysical supply towards understanding the beneficiaries of ES and quantifying their demand, access and utilisation of services, as well as the consequences for well-being (Bagstad et al., 2014; Poppy et al., 2014). Combining social science theory and data to explain the social-ecological processes of ES co-production, use and well-being consequences will likely result in substantial improvements to ES models (Bagstad et al., 2014; Díaz et al., 2015; Pascual et al., 2017; Suich et al., 2015; Willcock et al., submitted for publication). Such social science data are sometimes available at large scales (e.g., via national censuses) but, with some notable exceptions (e.g., Hamann et al. (2016, 2015)), are rarely used within ES models (Egoh et al., 2012; Martínez-Harms and Balvanera, 2012; Wong et al., 2014). The process of DDM guides researchers in how to incorporate of big data into ES models, scaling up results from sites to continents (Hampton et al., 2013; Lokers et al., 2016). DDM allows an interdisciplinary approach across a large scale and so may help guide global policy-making, e.g., within the Intergovernmental Science-Policy Platform for Biodiversity and Ecosystem Services (IPBES; www.ipbes.net).

In conclusion, DDM could be a useful tool to scale up ES models for greater policy- and decision-making relevance. DDM allows for the incorporation of big data, producing interdisciplinary models and holistic solutions to complex socio-ecological issues. It is crucial that the approach and results of machine learning algorithms are conveyed to the user to enhance transparency, including the uncertainty associated with the modelled results. In fact, we hope that the validation of ES models becomes standard practice with the ES community for both process-based and DDM. In the future, automation of the modelling processes may enable users to run ES models with minimal support from scientists, increasing ownership in the final output. Such automation should be accompanied by transparent provenance information and procedures for a computerised system to select context-appropriate data and models. Taken together, the advances described here could help to ensure ES research contributes to and inform ongoing policy processes, such as IPBES, as well as national-, subnational-, and local-scale decision making.

Contributions

SW, INA, JML, KJB, SB, BV & FV conceived the paper. SW, DAPH, JMB & INA carried out the South African case study. JML, SB, AM, CP, SS, GS & FV carried out the Sicilian case study. SW & INA wrote the manuscript, with comments and revisions from all other authors.

Acknowledgements

This work took place under the ‘WISER: Which Ecosystem Service Models Best Capture the Needs of the Rural Poor?’ project (NE/L001322/1), funded by the UK Ecosystem Services for Poverty Alleviation program (ESPA; www.espa.ac.uk). ESPA receives its funding from the UK Department for International Development, the Economic and Social Research Council and the Natural Environment Research Council. Any use of trade, firm, or product names is for descriptive purposes only and does not imply endorsement by the U.S. Government. We would like to thank the two anonymous reviewers whose comments improved the manuscript.

Appendix A. Supplementary data

Supplementary data associated with this article can be found, in the online version, at <https://doi.org/10.1016/j.ecoser.2018.04.004>.

References

- Aguilera, P.A., Fernandez, A., Fernandez, R., Rumi, R., Salmeron, A., 2011. Bayesian networks in environmental modelling. *Environ. Model. Softw.* 26, 1376–1388. <https://doi.org/10.1016/j.envsoft.2011.06.004>.
- Bagstad, K.J., Villa, F., Batker, D., Harrison-Cox, J., Voigt, B., Johnson, G.W., 2014. From theoretical to actual ecosystem services: mapping beneficiaries and spatial flows in ecosystem service assessments. *Ecol. Soc.* 19. <https://doi.org/10.5751/ES-06523-190264>. art64.
- Balbi, S., Villa, F., Mojtahed, V., Hegetschweiler, K.T., Giupponi, C., 2016. A spatial Bayesian network model to assess the benefits of early warning for urban flood risk to people. *Nat. Hazards Earth Syst. Sci.* 16, 1323–1337. <https://doi.org/10.5194/nhess-16-1323-2016>.
- Baveye, P.C., 2017. Quantification of ecosystem services: Beyond all the “guesstimates”, how do we get real data? *Ecosyst. Serv.* 24, 47–49. <https://doi.org/10.1016/j.ecoser.2017.02.006>.
- Baziliana, M., Rogner, H., Howells, M., Herrmann, S., Arent, D., Gielen, D., Steduto, P., Mueller, A., Komor, P., Tol, R.S.J., Yumkella, K.K., 2011. Considering the energy, water and food nexus: Towards an integrated modelling approach. *Energy Policy* 39, 7896–7906. <https://doi.org/10.1016/j.enpol.2011.09.039>.
- Blum, A.L., Langley, P., 1997. Selection of relevant features and examples in machine learning. *Artif. Intell.* 97, 245–271. [https://doi.org/10.1016/S0004-3702\(97\)00063-5](https://doi.org/10.1016/S0004-3702(97)00063-5).
- Bottou, L., 2014. From machine learning to machine reasoning. *Mach. Learn.* 94, 133–149. <https://doi.org/10.1007/s10994-013-5335-x>.
- Bullock, J.M., Dhanjal-Adams, K.L., Milne, A., Oliver, T.H., Todman, L.C., Whitmore, A. P., Pywell, R.F., 2017. Resilience and food security: rethinking an ecological concept. *J. Ecol.* 105, 880–884. <https://doi.org/10.1111/1365-2745.12791>.
- Chen, S.H., Pollino, C.A., 2012. Good practice in Bayesian network modelling. *Environ. Model. Softw.* 37, 134–145. <https://doi.org/10.1016/j.envsoft.2012.03.012>.
- Cheung, G.W., Rensvold, R.B., 2002. Evaluating goodness-of-fit indexes for testing measurement invariance. *Struct. Equ. Model. A Multidiscip. J.* 9, 233–255. https://doi.org/10.1207/S15328007SEM0902_5.
- Clark, J.S., 2003. Uncertainty and variability in demography and population growth: a hierarchical approach. *Ecology* 84, 1370–1381. [https://doi.org/10.1890/0012-9658\(2003\)084\[1370:UAVIDA\]2.0.CO;2](https://doi.org/10.1890/0012-9658(2003)084[1370:UAVIDA]2.0.CO;2).
- Clark, W.C., Tomich, T.P., van Noordwijk, M., Guston, D., Catacutan, D., Dickson, N. M., McNie, E., 2016. Boundary work for sustainable development: Natural resource management at the Consultative Group on International Agricultural Research (CGIAR). *Proc. Natl. Acad. Sci. USA* 113, 4615–4622. <https://doi.org/10.1073/pnas.0900231108>.
- Díaz, S., Demissew, S., Carabias, J., Joly, C., Lonsdale, M., Ash, N., Larigauderie, A., Adhikari, J.R., Arico, S., Báldi, A., Bartuska, A., Baste, I.A., Bilgin, A., Brondizio, E., Chan, K.M., Figueroa, V.E., Duraipapp, A., Fischer, M., Hill, R., Koetz, T., Leadley, P., Lyver, P., Mace, G.M., Martin-Lopez, B., Okumura, M., Pacheco, D., Pascual, U., Pérez, E.S., Reyers, B., Roth, E., Saito, O., Scholes, R.J., Sharma, N., Tallis, H., Thaman, R., Watson, R., Yahara, T., Hamid, Z.A., Akosim, C., Al-Hafedh, Y., Allahverdiyev, R., Amankwah, E., Asah, S.T., Asfaw, Z., Bartus, G., Brooks, L.A., Caillaux, J., Dalle, G., Darnaedi, D., Driver, A., Erpul, G., Escobar-Eyzaguirre, P., Failler, P., Fouda, A.M.M., Fu, B., Gundimeda, H., Hashimoto, S., Homer, F., Lavorel, S., Lichtenstein, G., Mala, W.A., Mandivenyi, W., Matczak, P., Mbizvo, C., Mehrdadi, M., Metzger, J.P., Mikissa, J.B., Moller, H., Mooney, H.A., Mumby, P., Nagendra, H., Neshover, C., Oteng-Yeboah, A.A., Pataki, G., Roué, M., Rubis, J., Schultz, M., Smith, P., Sumaila, R., Takeuchi, K., Thomas, S., Verma, M., Yeo-Chang, Y., Zlatanova, D., 2015. The IPBES Conceptual Framework – connecting nature and people. *Curr. Opin. Environ. Sustain.* 14, 1–16. <https://doi.org/10.1016/j.cosust.2014.11.002>.
- Egoh, B., Drakou, E.G., Dunbar, M.B., Maes, J., Willems, L., 2012. Indicators for mapping ecosystem services: a review. Report EUR 25456 EN. Luxembourg, Luxembourg.
- Fayyad, U., Piatetsky-Shapiro, G., Smyth, P., 1996. From data mining to knowledge discovery in databases. *AI Mag.* 17, 37. <https://doi.org/10.1609/AIMAG.V17I3.1230>.
- Frank, E., Hall, M.A., Witten, I.H., 2016. The WEKA Workbench. Online Appendix for “Data Mining: Practical Machine Learning Tools and Techniques. Morgan Kaufmann.
- Friedman, N., Goldszmidt, M., 1996. Discretizing continuous attributes while learning bayesian networks. *ICML*, 157–165.
- Galelli, S., Humphrey, G.B., Maier, H.R., Castelletti, A., Dandy, G.C., Gibbs, M.S., 2014. An evaluation framework for input variable selection algorithms for environmental data-driven models. *Environ. Model. Softw.* 62, 33–51. <https://doi.org/10.1016/j.envsoft.2014.08.015>.
- Gama, J., Medas, P., Castillo, G., Rodrigues, P., 2004. Learning with drift detection. In: *Brazilian Symposium on Artificial Intelligence*. Springer, Berlin, Heidelberg, pp. 286–295.
- GeoTerraImage, 2015. 2013–2014 South African National Land-Cover Dataset version 05.
- Ghahramani, Z., 2015. Probabilistic machine learning and artificial intelligence. *Nature* 521, 452–459.
- Gibert, K., Spate, J., Sánchez-Marré, M., Athanasiadis, I.N., Comas, J., 2008. Data mining for environmental systems. *Dev. Integr. Environ. Assess.* 3, 205–228. [https://doi.org/10.1016/S1574-101X\(08\)00612-1](https://doi.org/10.1016/S1574-101X(08)00612-1).
- Haklay, M., Weber, P., 2008. OpenStreetMap: User-generated street maps. *IEEE Pervasive Comput.* 7, 12–18. <https://doi.org/10.1109/MPRV.2008.80>.
- Hall, M., Frank, E., Holmes, G., Pfahringer, B., Reutemann, P., Witten, I.H., 2009. The WEKA data mining software: an update. *SIGKDD Explor.* 11, 10–18.
- Hamann, M., Biggs, R., Reyers, B., 2015. Mapping social-ecological systems: Identifying “green-loop” and “red-loop” dynamics based on characteristic bundles of ecosystem service use. *Glob. Environ. Chang.* 34, 218–226. <https://doi.org/10.1016/j.gloenvcha.2015.07.008>.
- Hamann, M., Biggs, R., Reyers, B., Pomeroy, R., Abunge, C., Galafassi, D., 2016. An exploration of human well-being bundles as identifiers of ecosystem service use patterns. *PLoS One* 11, e0163476. <https://doi.org/10.1371/journal.pone.0163476>.
- Hamel, P., Bryant, B.P., 2017. Uncertainty assessment in ecosystem services analyses: Seven challenges and practical responses. *Ecosyst. Serv.* 24, 1–15. <https://doi.org/10.1016/j.ecoser.2016.12.008>.
- Hampton, S.E., Strasser, C.A., Tewksbury, J.J., Gram, W.K., Budden, A.E., Batcheller, A. L., Duke, C.S., Porter, J.H., 2013. Big data and the future of ecology. *Front. Ecol. Environ.* 11, 156–162. <https://doi.org/10.1890/120103>.
- Hastie, T., Tibshirani, R., Friedman, J., 2009. Overview of Supervised Learning. Springer, New York, NY, pp. 9–41. https://doi.org/10.1007/978-0-387-84858-7_2.
- Hyndman, R.J., Koehler, A.B., 2006. Another look at measures of forecast accuracy. *Int. J. Forecast.* 22, 679–688. <https://doi.org/10.1016/j.jforecast.2006.03.001>.
- Isaac, N.J.B., van Strien, A.J., August, T.A., de Zeeuw, M.P., Roy, D.B., 2014. Statistics for citizen science: extracting signals of change from noisy ecological data. *Methods Ecol. Evol.* 5, 1052–1060. <https://doi.org/10.1111/2041-210X.12254>.
- Jean, N., Burke, M., Xie, M., Davis, W.M., Lobell, D.B., Ermon, S., 2016. Combining satellite imagery and machine learning to predict poverty. *Science* 353, 790–794. <https://doi.org/10.1126/science.127894>.
- Jordan, M.I., Mitchell, T.M., 2015. Machine learning: trends, perspectives, and prospects. *Science* 349, 255–260. <https://doi.org/10.1126/science.127894>.
- Landuyt, D., Broekx, S., D’hondt, R., Engelen, G., Aertsens, J., Goethals, P.L.M., 2013. A review of Bayesian belief networks in ecosystem service modelling. *Environ. Model. Softw.* 46, 1–11. <https://doi.org/10.1016/j.envsoft.2013.03.011>.
- Li, J., Heap, A.D., Potter, A., Daniell, J.J., 2011. Application of machine learning methods to spatial interpolation of environmental variables. *Environ. Model. Softw.* 26, 1647–1659. <https://doi.org/10.1016/j.envsoft.2011.07.004>.
- Liu, C., Newell, G., White, M., 2016. On the selection of thresholds for predicting species occurrence with presence-only data. *Ecol. Evol.* 6, 337–348. <https://doi.org/10.1002/ece3.1878>.
- Lokers, R., Knapen, R., Janssen, S., van Randen, Y., Jansen, J., 2016. Analysis of Big Data technologies for use in agro-environmental science. *Environ. Model. Softw.* 84, 494–504. <https://doi.org/10.1016/j.envsoft.2016.07.017>.
- Marcot, B.G., Steventon, J.D., Sutherland, G.D., Mccann, R.K., 2006. Guidelines for developing and updating Bayesian belief networks applied to ecological modeling and conservation. *Can. J. For. Res.* 36, 3063–3074. <https://doi.org/10.1139/X06-135>.
- Martínez-Harms, M.J., Balvanera, P., 2012. Methods for mapping ecosystem service supply: a review. *Int. J. Biodivers. Sci. Ecosyst. Serv. Manage.* 8, 17–25. <https://doi.org/10.1080/21513732.2012.663792>.
- McKenzie, E., Posner, S., Tillmann, P., Bernhard, J.R., Howard, K., Rosenthal, A., 2014. Understanding the use of ecosystem service knowledge in decision making: lessons from international experiences of spatial planning. *Environ. Plan. C Gov. Policy* 32, 320–340. <https://doi.org/10.1068/c12292j>.
- Mjolsness, E., DeCoste, D., 2001. Machine learning for science: state of the art and future prospects. *Science* 293, 2051–2055. <https://doi.org/10.1126/science.293.5537.2051>.
- Mouchet, M.A., Lamarque, P., Martin-Lopez, B., Crouzat, E., Gos, P., Byczek, C., Lavorel, S., 2014. An interdisciplinary methodological guide for quantifying associations between ecosystem services. *Glob. Environ. Chang.* 28, 298–308. <https://doi.org/10.1016/j.gloenvcha.2014.07.012>.
- Mullainathan, S., Spiess, J., 2017. Machine learning: an applied econometric approach. *J. Econ. Perspect.* 31, 87–106. <https://doi.org/10.1257/jep.31.2.87>.
- Nojavan, F., Qian, S.S., Stow, C.A., 2017. Comparative analysis of discretization methods in Bayesian networks. *Environ. Model. Softw.* 87, 64–71. <https://doi.org/10.1016/j.envsoft.2016.10.007>.
- Olander, L., Polasky, S., Kagan, J.S., Johnston, R.J., Waigner, L., Saah, D., Maguire, L., Boyd, J., Yoskowitz, D., 2017. So you want your research to be relevant? Building the bridge between ecosystem services research and practice. *Ecosyst. Serv.* 26, 170–182. <https://doi.org/10.1016/j.ecoser.2017.06.003>.
- Pascual, U., Balvanera, P., Díaz, S., Pataki, G., Roth, E., Stenseke, M., Watson, R.T., Başak Dessane, E., Islar, M., Kelemen, E., Maris, V., Quaa, M., Subramanian, S.M., Wittmer, H., Adlan, A., Ahn, S., Al-Hafedh, Y.S., Amankwah, E., Asah, S.T., Berry, P., Bilgin, A., Breslow, S.J., Bullock, C., Cáceres, D., Daly-Hassen, H., Figueroa, E., Golden, C.D., Gómez-Baggethun, E., González-Jiménez, D., Houdet, J., Keune, H., Kumar, R., Ma, K., May, P.H., Mead, A., O’Farrell, P., Pandit, R., Pengue, W., Pichis-Madruga, R., Popa, F., Preston, S., Pacheco-Balanza, D., Saarikoski, H., Strassburg, B.B., van den Belt, M., Verma, M., Wickson, F., Yagi, N., 2017. Valuing nature’s contributions to people: the IPBES approach. *Curr. Opin. Environ. Sustain.* 26–27, 7–16. <https://doi.org/10.1016/j.cosust.2016.12.006>.
- Poppy, G.M., Chiotha, S., Eigenbrod, F., Harvey, C.A., Honzák, M., Hudson, M.D., Jarvis, A., Madise, N.J., Schreckenberg, K., Shackleton, C.M., Villa, F., Dawson, T.P., 2014. Food security in a perfect storm: using the ecosystem services framework to increase understanding. *Philos. Trans. R. Soc. Lond. B Biol. Sci.*, 369.
- Pradhan, B., Seenii, M.I., Kalantar, B., 2017. Performance evaluation and sensitivity analysis of expert-based, statistical, machine learning, and hybrid models for producing landslide susceptibility maps. In: *Laser Scanning Applications in*

- Landslide Assessment. Springer International Publishing, Cham, pp. 193–232. https://doi.org/10.1007/978-3-319-55342-9_11.
- Richards, D.R., Tunçer, B., 2018. Using image recognition to automate assessment of cultural ecosystem services from social media photographs. *Ecosyst. Serv.* 31, 318–325. <https://doi.org/10.1016/j.ecoser.2017.09.004>.
- Russell, S., Norvig, P., 2003. *Artificial Intelligence: A Modern Approach*. Prentice Hall.
- Schmidhuber, J., 2015. Deep learning in neural networks: an overview. *Neural Networks* 61, 85–117. <https://doi.org/10.1016/j.neunet.2014.09.003>.
- Scholes, R., Reyers, B., Biggs, R., Spierenburg, M., Duriappah, A., 2013. Multi-scale and cross-scale assessments of social-ecological systems and their ecosystem services. *Curr. Opin. Environ. Sustain.* 5, 16–25. <https://doi.org/10.1016/j.cosust.2013.01.004>.
- Scholes, R.J., 1998. The South African 1: 250 000 maps of areas of homogeneous grazing potential.
- Smith, T.C., Frank, E., 2016. *Introducing Machine Learning Concepts with WEKA*. Humana Press, New York, NY, pp. 353–378. https://doi.org/10.1007/978-1-4939-3578-9_17.
- Suich, H., Howe, C., Mace, G., 2015. Ecosystem services and poverty alleviation: A review of the empirical links. *Ecosyst. Serv.* 12, 137–147. <https://doi.org/10.1016/j.ecoser.2015.02.005>.
- Suominen, A., Toivanen, H., 2016. Map of science with topic modeling: Comparison of unsupervised learning and human-assigned subject classification. *J. Assoc. Inf. Sci. Technol.* 67, 2464–2476. <https://doi.org/10.1002/asi.23596>.
- Tarca, A.L., Carey, V.J., Chen, X., Romero, R., Drăghici, S., 2007. Machine learning and its applications to biology. *PLoS Comput. Biol.* 3, e116. <https://doi.org/10.1371/journal.pcbi.0030116>.
- Uusitalo, L., 2007. Advantages and challenges of Bayesian networks in environmental modelling. *Ecol. Modell.* 203, 312–318. <https://doi.org/10.1016/j.ecolmodel.2006.11.033>.
- Varma, S., Simon, R., 2006. Bias in error estimation when using cross-validation for model selection. *BMC Bioinform.* 7, 91.
- Villa, F., Bagstad, K.J., Voigt, B., Johnson, G.W., Portela, R., Honzák, M., Batker, D., 2014. A methodology for adaptable and robust ecosystem services assessment. *PLoS One* 9, e91001. <https://doi.org/10.1371/journal.pone.0091001>.
- Villa, F., Balbi, S., Athanasiadis, I.N., Caracciolo, C., 2017. Semantics for interoperability of distributed data and models: Foundations for better-connected information. *F1000Research* 6, 686. <https://doi.org/10.12688/f1000research.11638.1>.
- Wiens, T.S., Dale, B.C., Boyce, M.S., Kershaw, G.P., 2008. Three way k-fold cross-validation of resource selection functions. *Ecol. Modell.* 212, 244–255. <https://doi.org/10.1016/j.ecolmodel.2007.10.005>.
- Willcock, S., Hooftman, D.A.P., Balbi, S., Blanchard, R., Dawson, T.P., O'Farrell, P.J., Hickler, T., Hudson, M.D., Lindeskog, M., Martinez-Lopez, J., Mulligan, M., Reyers, B., Schreckenber, K., Shackleton, C., Sitas, N., Villa, F., Watts, S.M., Eigenbrod, F., Bullock, J.M., submitted for publication. Continental scale validation of ecosystem service models.
- Willcock, S., Hooftman, D., Sitas, N., O'Farrell, P., Hudson, M.D., Reyers, B., Eigenbrod, F., Bullock, J.M., 2016. Do ecosystem service maps and models meet stakeholders' needs? A preliminary survey across sub-Saharan Africa. *Ecosyst. Serv.* 18, 110–117. <https://doi.org/10.1016/j.ecoser.2016.02.038>.
- Willcock, S., Phillips, O.L., Platts, P.J., Balmford, A., Burgess, N.D., Lovett, J.C., Ahrends, A., Bayliss, J., Daggart, N., Doody, K., Fanning, E., Green, J.M., Hall, J., Howell, K.L., Marchant, R., Marshall, A.R., Mbilinyi, B., Munishi, P.K., Owen, N., Swetnam, R.D., Topp-Jorgensen, E.J., Lewis, S.L., 2014. Quantifying and understanding carbon storage and sequestration within the Eastern Arc Mountains of Tanzania, a tropical biodiversity hotspot. *Carbon Balance Manage.* 9, 2. <https://doi.org/10.1186/1750-0680-9-2>.
- Witten, I.H., Frank, E., Hall, M.A., Pal, C.J., 2016. *Data Mining: Practical machine learning tools and techniques*. Morgan Kaufmann, San Francisco, USA.
- Wong, C.P., Jiang, B., Kinzig, A.P., Lee, K.N., Ouyang, Z., 2014. Linking ecosystem characteristics to final ecosystem services for public policy. *Ecol. Lett.* 18, 108–118. <https://doi.org/10.1111/ele.12389>.
- Wu, X., Zhu, X., Wu, G.-Q., Ding, W., 2014. Data mining with big data. *IEEE Trans. Knowl. Data Eng.* 26, 97–107. <https://doi.org/10.1109/TKDE.2013.109>.
- Yamashita, T., Yamashita, K., Kamimura, R., 2007. A stepwise AIC method for variable selection in linear regression. *Commun. Stat. Theory Methods* 36, 2395–2403. <https://doi.org/10.1080/03610920701215639>.
- Zhang, Y., Wallace, B., 2015. A Sensitivity Analysis of (and Practitioners' Guide to) Convolutional Neural Networks for Sentence Classification. *arXiv Prepr. arXiv1510.03820*.