



Universiteit
Leiden
The Netherlands

Machine learning algorithms performed no better than regression models for prognostication in traumatic brain injury

Gravesteijn, B.Y.; Nieboer, D.; Ercole, A.; Lingsma, H.F.; Nelson, D.; Calster, B. van; ... ; CENTER-TBI Collaborators

Citation

Gravesteijn, B. Y., Nieboer, D., Ercole, A., Lingsma, H. F., Nelson, D., Calster, B. van, & Steyerberg, E. W. (2020). Machine learning algorithms performed no better than regression models for prognostication in traumatic brain injury. *Journal Of Clinical Epidemiology*, 122, 95-107. doi:10.1016/j.jclinepi.2020.03.005

Version: Publisher's Version
License: [Creative Commons CC BY 4.0 license](#)
Downloaded from: <https://hdl.handle.net/1887/3627143>

Note: To cite this publication please use the final published version (if applicable).

ORIGINAL ARTICLE

Machine learning algorithms performed no better than regression models for prognostication in traumatic brain injury

Benjamin Y. Gravesteijn^{a,*}, Daan Nieboer^b, Ari Ercole^c, Hester F. Lingsma^b, David Nelson^d, Ben van Calster^{e,f}, Ewout W. Steyerberg^{b,f}, the CENTER-TBI collaborators

^aDepartments of Public Health, Erasmus MC – University Medical Centre Rotterdam, Postbus 2040, 3000 CA, Rotterdam, the Netherlands

^bDepartments of Public Health, Erasmus MC – University Medical Centre Rotterdam, Rotterdam, the Netherlands

^cDivision of Anaesthesia, University of Cambridge, Cambridge, United Kingdom

^dDepartment of Physiology and Pharmacology, Section of Perioperative Medicine and Intensive Care, Karolinska Institutet, Stockholm, Sweden

^eDepartment of Development and Regeneration, KU Leuven, Belgium

^fDepartment of Biomedical Data Sciences, Leiden University Medical Centre, Leiden, the Netherlands

Accepted 9 March 2020; Published online 20 March 2020

Abstract

Objective: We aimed to explore the added value of common machine learning (ML) algorithms for prediction of outcome for moderate and severe traumatic brain injury.

Study Design and Setting: We performed logistic regression (LR), lasso regression, and ridge regression with key baseline predictors in the IMPACT-II database (15 studies, $n = 11,022$). ML algorithms included support vector machines, random forests, gradient boosting machines, and artificial neural networks and were trained using the same predictors. To assess generalizability of predictions, we performed internal, internal-external, and external validation on the recent CENTER-TBI study (patients with Glasgow Coma Scale < 13 , $n = 1,554$). Both calibration (calibration slope/intercept) and discrimination (area under the curve) was quantified.

Results: In the IMPACT-II database, 3,332/11,022 (30%) died and 5,233(48%) had unfavorable outcome (Glasgow Outcome Scale less than 4). In the CENTER-TBI study, 348/1,554(29%) died and 651(54%) had unfavorable outcome. Discrimination and calibration varied widely between the studies and less so between the studied algorithms. The mean area under the curve was 0.82 for mortality and 0.77 for unfavorable outcomes in the CENTER-TBI study.

Conclusion: ML algorithms may not outperform traditional regression approaches in a low-dimensional setting for outcome prediction after moderate or severe traumatic brain injury. Similar to regression-based prediction models, ML algorithms should be rigorously validated to ensure applicability to new populations. © 2020 The Authors. Published by Elsevier Inc. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

Keywords: Machine learning; Prognosis; Traumatic brain injury; Prediction; Data science; Cohort study

Funding: Data used in preparation of this manuscript were obtained in the context of CENTER-TBI, a large collaborative project with the support of the European Union 7th Framework program (EC grant 602150). Additional funding was obtained from the Hannelore Kohl Stiftung (Germany), the OneMind (USA) and the Integra LifeSciences Corporation (USA). The funder had no role in the study design, enrollment, collection of data, writing, or publication decisions.

Ethics approval and consent to participate: The authors declare that all participants signed informed consent to be included in the study. Ethical approval was obtained for each recruiting sites.

Conflict of interest: The authors declare to have no competing interests.

Consent for publication: The authors declare that approval for publication was obtained.

Availability of data and material: As an EU-funded project, CENTER-TBI is an open-access database. Access can be obtained as collaborator, after declaring to adhere to the CENTER-TBI data use agreement. For more information, see <https://www.center-tbi.eu/data>.

Trial registration: ClinicalTrials.gov Identifier: NCT02210221.

Take home message: Flexible machine learning algorithms may not perform better than traditional regression approaches in a low-dimensional setting for outcome prediction after moderate or severe TBI. Similar to regression-based prediction models, ML algorithms should be rigorously validated to ensure applicability to new populations.

* Corresponding author: Tel.: +316-83448055; fax: +31-10-7044724.

E-mail address: b.gravesteijn@erasmusmc.nl (B.Y. Gravesteijn).

What is new?

Key findings

- Considering discrimination and calibration and overall performance, no clear difference was seen in performance between machine learning (ML) algorithms or regression-based models.
- More variability in performance (both discrimination and calibration) was seen between study populations than between algorithms.

What this adds to what was known?

- A recent systematic review showed that studies that suggested superior performance of ML methods are more prone to bias. However, these studies mainly focused on comparing discriminative performance of these models. Our study also focused on performance in terms of calibration and generalizability.

What is the implication and what should change now?

- Using novel ML algorithms will likely not improve outcome prediction. Instead, prediction research should focus including predictors with substantial incremental prognostic value.
- Prediction models, based on both ML algorithms and regression-based methods, need continuous validation and updating to ensure applicability to new populations.

1. Introduction

Traumatic brain injury (TBI) is a common disease, with a significant societal burden [1]: TBI is estimated to be responsible for around 300 hospital admissions and 12 deaths per 100,000 persons per year in Europe [2]. TBI is a heterogeneous disease in terms of phenotype and prognosis [3]. Therefore, prognostic models, which predict outcome for a patient, given a particular combination of baseline characteristics, are important: they may give us insight in mechanisms of disease that lead to poor outcome and allow for risk-based stratification of patients for logistic, research, and clinical reasons.

A large number of prediction models have been developed to predict outcome for patients with TBI, mostly using traditional regression techniques [4]. However, these models have not yet been widely implemented in clinical practice. In recent years, more flexible machine learning (ML) algorithms have enjoyed enthusiasm as potentially promising techniques to improve outcome prognostication [5]. Frequently used methods are support vector machines

(SVMs) [6], deep neural networks (NNs) [7], random forests (RFs) [8], and gradient boosting machine (GBM) [9]. Some of these algorithms have been used to develop prediction models on small data sets (<200 events) [10–12]. Because ML algorithms are more prone to overfitting [13], it remains unclear what the impact on prognostication is of these novel techniques.

Although the incremental value of flexible ML methods has been previously assessed, these comparisons were potentially subject to bias [14]. The incremental value of ML algorithms is potentially overrated because studies up to this point mainly focused on the ability of the methods to discriminate between patients with good and poor outcomes [15–19]. Performance of prediction models is however commonly measured across at least two dimensions: calibration and discrimination [20,21]. Calibration refers to the agreement of predicted probabilities of a model and observed outcomes (e.g., “if the risk of death is x%, do x% of the patients with this prediction actually die?”). Poor calibration of prediction models may lead to harmful decision-making when applying these models [22–24].

One of the more thoroughly validated prediction models with good performance exists in the field of TBI: the IMPACT model [25]. This model comprises baseline clinical characteristics, presence of secondary insults, imaging findings, and laboratory characteristics. Using the variables of this model, the present study aims to fairly assess the potential incremental value of flexible ML methods beyond classical regression approaches.

2. Methods

This study was reported to conform with the TRIPOD guidelines [23].

2.1. Study population

We included 15 studies from the IMPACT-II database. These include four observational studies and eleven randomized controlled trials on moderate to severe TBI (Glasgow Coma Scale [GCS] ≤ 12), which were conducted between 1984 and 2004 [26]. Furthermore, we validated models in the patients with moderate to severe TBI (GCS ≤ 12) from the CENTER-TBI core study. This is a recent prospective study, which included patients from 2014 to 2018 [27]. Data for the CENTER-TBI study have been collected through the Quesgen e-CRF (Quesgen Systems Inc, Burlingame, CA, USA), hosted on the INCF platform, and extracted via the INCF Neurobot tool (INCF, Sweden). Version 1.0 of the CENTER-TBI data was used for this analysis.

2.2. Model specification

The outcomes which were predicted were 6 months mortality and unfavorable outcome (Glasgow Outcome Scale < 3, or Glasgow Outcome Scale–Extended <5).

The predictors included in the models were 11 predictors of the IMPACT laboratory model [25]. Continuous variables were included as continuous variables in the model (no categorization). An overview of the included variables, and their specifications, is shown in Table 1. The baseline GCS score was defined as the last GCS in the emergency department (“poststabilization”). If this score was missing, the nearest GCS at an earlier moment was used. In total, eleven predictors were included, representing 19 parameters (or degrees of freedom [df]). In the case of mortality, 3,491 events (or 184 events per parameter) were on average present in our database for each training. The variables were normalized or one-hot encoded because this is standard practice for training algorithms which use gradient descent optimization.

2.3. Regression techniques

The regression techniques which were compared with the ML algorithms included standard logistic regression, but also penalized regression: lasso and ridge regressions [28]. These algorithms were developed to improve the performance of logistic regression models by shrinking the coefficients during estimation [29,30]. The objective is to obtain models that are less prone to making too extreme predictions (overfitting). The *glmnet* function from the *glmnet* package was used (alpha = 0 for ridge, and alpha = 1 for lasso). No nonlinear or interaction terms were included in the regression models.

Table 1. Model specification: 11 predictors, with 19 degrees of freedom

Variable in the model	Characteristics
Age	Continuous
Motor GCS score	Categorical, 1–6
w	Categorical, 3 levels: • Both reactive • One reactive • Two reactive
CT class	Categorical, 5 levels: • No visible pathology • Diffuse injury • Diffuse injury with swelling • Diffuse injury with shift • Mass
Traumatic subarachnoid hemorrhage	Binary
Epidural hematoma	Binary
Hypoxia	Binary
Hypotension	Binary
Glucose, first measured	Continuous
Sodium, first measured	Continuous
Hemoglobin, first measured	Continuous

Abbreviations: GCS, Glasgow Coma Scale; CT, computed tomography.

2.4. Machine learning algorithms

All analyses were performed using R (R Core Team (2013). R: A language and environment for statistical computing. R Foundation for Statistical Computing, Vienna, Austria). The script can be found on https://github.com/bgravesteyn/ML_baseline_pred_code.

The flexible ML algorithms that were compared with logistic regression were SVM, NN, RF, and GBM. All these algorithms have so-called “hyperparameters,” which need to be optimized for the algorithms to work optimally. To select the optimal hyperparameters, the framework of the *caret* package was used. The best combination of hyperparameters of the algorithms was chosen based on the highest log likelihood. The average log likelihood over 10 repetitions of tenfold cross-validation was used to select the optimal parameters (Fig. 1). For a detailed description of what algorithms were used and what hyperparameters were considered, see Appendix B.

The included flexible ML methods, just like regression, do not allow for missing values. Unlike regression, however, they are not readily compatible with multiple imputation: not every algorithm uses weights as core operators. Moreover, for the algorithms that use weights, there is no implementation of pooling these weights over multiple data sets using Rubin’s rules [31]. Therefore, multiple imputation using the *mice* package was performed [32], but only one imputed data set was used to train the models. The outcome and all predictors were included in the imputation model. To check for stability of results, a sensitivity analysis was performed with a different imputed data set.

2.5. Cross-validation

The models were validated using three different strategies. First, they were cross-validated per study: the algorithms were trained on all but one study, and calibration and discrimination were assessed by applying the models to the study not used at model development. This procedure has been referred to as ‘Internal-external cross-validation’ [33,34]. For an overview of the analytical steps of internal-external cross-validation, see Fig. 1. Second, internal validation was performed in the IMPACT-II database using 10 times 10-fold random cross-validation. For this method, the data were randomly divided by deciles. The model was developed on 9/10 and validated on 1/10 of the data. This process was repeated until all patients were used once as validation sample. Finally, a fully external validation was performed, with training of the models in the IMPACT-II database and validating in CENTER-TBI.

The performance was assessed in three domains. First, calibration was examined graphically and quantified using a calibration slope and the calibration intercept: the calibration test proposed by Cox [35]. Second, discrimination was quantified using the *c*-statistic, also known as the area under the receiver operating curve. The confidence intervals of

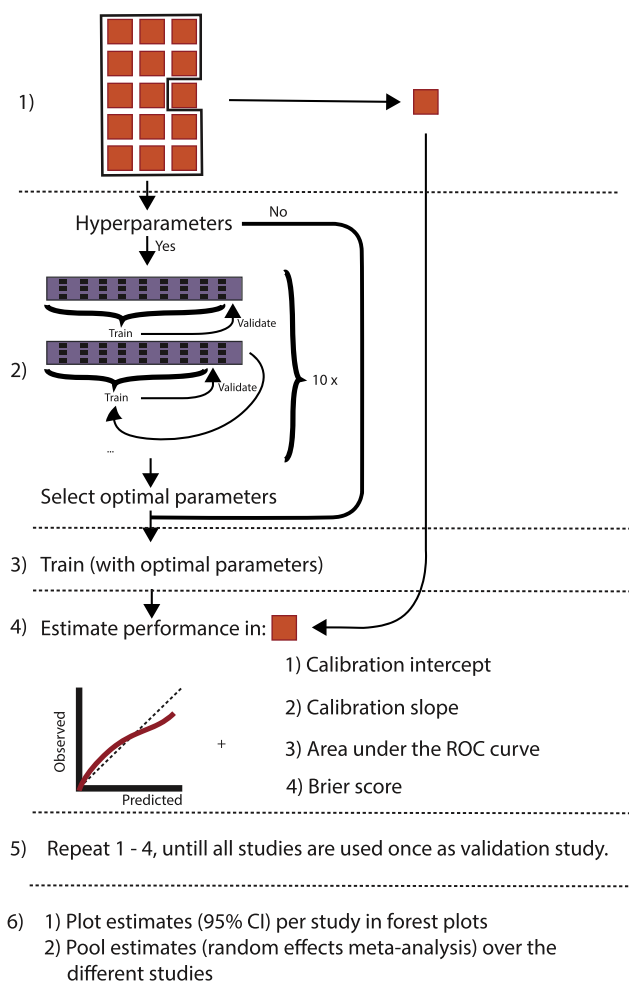


Fig. 1. Overview of the experimental setup. Step 1 is selecting a study as a validation study. Step 2 is selecting the optimal hyperparameters through 10 times 10-fold cross-validation. If the algorithm did not require hyperparameters, this step was skipped. Step 3 is the training of the final model with optimal hyperparameters on the full training data. The model of step 3 was validated in step 4 with the study that was left out of the training set. Step 5 is repeating step 1–4 until all studies are used once as validation study. Finally step 6 is the presentation of the results, and pooling the results over the different studies.

the c-statistic were obtained using the DeLong et al. method [36], using the *ci.auc* function from the *pROC* package. Third, as a measure of overall performance, the Brier score was calculated [37]. More extensive descriptions of these metrics can be found in Appendix B.

The estimates and 95% confidence intervals were plotted in forest plots, to visually inspect the variation. To obtain estimates per model and outcome, the estimates (and standard errors) in every validation were pooled using a random effects meta-analysis, using the DerSimonian and Laird estimator for τ^2 [38]. Because the CENTER-TBI database is a recent study, unlike the IMPACT-II studies, the estimates obtained from validating in this study were presented separately.

To compare whether observed variation of the performance measures can be attributed to differences in performance across study population or type of model used, we

used mixed effects linear regression. This was performed in the internal-external validation framework. The performance measure was used as dependent variable, and two random intercepts were included in the model: one for what algorithm was used and one for what study the models were validated in. These random intercepts were assumed to follow a normal distribution with mean 0 and variance τ^2 . The percentage variation in performance attributable to in which study the model was validated was calculated by dividing the τ^2 of study by the total variance (the sum of the variance of the random intercepts of study and algorithm, and the residuals): $\tau^2_{\text{study}} / (\tau^2_{\text{study}} + \tau^2_{\text{algorithm}} + \tau^2_{\text{residuals}})$. Similarly, the percentage variation in performance attributable to what algorithm was trained was calculated.

3. Results

3.1. Patient characteristics

The baseline characteristics differed substantially between the IMPACT-II and the CENTER-TBI data. In the IMPACT-II database, patients were younger (35 vs. 47.4 years), had less traumatic subarachnoid hemorrhages (4,016 [45%] vs. 759 [74%]), and presented less often with a motor GCS of one (1,565 [16%] vs. 615 [45%]). However, the patients showed similar Glasgow Outcome Scale in the two studies: In the IMPACT-II database, 3,332 (30%) died and 5,233 (48%) had an unfavorable outcome, and in the CENTER-TBI study, 348 (29%) died and 651 (54%) had unfavorable outcome (Table 2). For an overview of the patient characteristics per study in IMPACT-II and CENTER-TBI, see Table A1.

3.2. Discrimination

At internal-external validation, the difference between maximum and minimum c-statistic of the algorithms was only 0.02 for mortality and unfavorable outcome. The discriminatory performance of the implementation of RF was suboptimal: the median and IQR of c-statistic of the RF were 0.79 (0.77–0.82) for mortality (the overall average was 0.81) and 0.79 (0.76–0.81) for unfavorable outcome (the overall average was 0.80). The discriminative performances varied substantially per study (Fig. A2 and Table 3). At internal validation in IMPACT-II, a similar pattern was seen, but the c-statistics were somewhat higher. For example, the GBM showed a c-statistic of 0.81 (0.79–0.83) at internal-external validation and 0.83 (0.82–0.84) at internal validation. When performing external validation in CENTER-TBI, this pattern was also seen: The RF showed a median and 95% CI for the c-statistic of 0.81 (0.78–0.84) for mortality (overall average was 0.82) and 0.76 (0.74–0.79) for unfavorable outcome (overall average was 0.77). Similar results were observed over a different imputed set, see Table A5.

Table 2. Baseline characteristics of the CENTER-TBI and IMPACT-II databases

Characteristic	IMPACT-II	CENTER-TBI	Missing data, total %
N	11,022	1,375	
Age (median [IQR])	31 [22, 46]	48 [28, 65]	0.0
Hypoxia (%)	1,707 (22)	217 (16.8)	26.3
Hypotension (%)	1,518 (17.2)	205 (15.9)	18.3
Marshall CT class (%)			40.6
1	379 (5.9)	81 (8.3)	
2	2,281 (36)	428 (43.9)	
3	1,259 (20)	86 (8.8)	
4	248 (3.9)	19 (2.0)	
5	2,223 (35)	360 (37.0)	
Traumatic subarachnoid hemorrhage (%)	4,016 (44.6)	759 (73.6)	19.1
Epidural hematoma (%)	1,275 (13.4)	172 (16.7)	14.8
Glucose (median mmol/L (SD))	8.84 (3.46)	8.18 (2.95)	44.5
Hemoglobin (mean g/dL (SD))	12.46 (2.42)	7.96 (2.36)	52.2
GCS motor (%)			7.4
1	1,565 (15.5)	615 (44.7)	
2	1,285 (12.7)	77 (5.6)	
3	1,362 (13.5)	80 (5.8)	
4	2,438 (24.1)	136 (9.9)	
5	2,791 (27.6)	357 (26.0)	
6	658 (6.5)	110 (8.0)	
Pupil (%)			12.8
Both reactive	6,292 (66.3)	973 (73.7)	
One reactive	1,192 (12.6)	110 (8.3)	
None reactive	2,010 (21.2)	238 (18.0)	
Glasgow outcome scale (%)			1.4
2	3,322 (30.1)	348 (29.0)	
3	1,911 (17.3)	303 (25.2)	
4	2,262 (20.5)	246 (20.5)	
5	3,527 (32.0)	303 (25.2)	

Abbreviations: CT, computed tomography; GCS, Glasgow Coma Scale; SD, standard deviation; IQR, interquartile range.

3.3. Calibration

At internal-external validation, the average calibration intercepts across the algorithms did not vary substantially: the range of calibration intercepts was -0.08 to -0.02 for mortality, and for unfavorable outcome, the calibration intercepts were 0.02 (Fig. 2B and Table A2). The range of calibration slopes was larger: 0.85 – 1.05 for mortality and 0.89 – 1.06 for unfavorable outcome (Fig. 2C and Table A3). The RF made too extreme predictions, with a median (95% CI) calibration slope of 0.85 (0.77 – 0.93) for mortality, whereas the overall mean was 0.97 , and 0.89 (0.82 – 0.96) for unfavorable outcome, whereas the overall mean was 0.99 . At internal validation in IMPACT-II, calibration slopes and intercepts were similar. In external validation in CENTER-TBI, the RF had again a too low calibration slope (0.88 , 95% CI: 0.77 – 0.99 for mortality).

The calibration intercept for mortality was generally low in CENTER-TBI: the overall mean was -0.58 , indicating

that the 6-month mortality was lower than expected in CENTER-TBI.

3.4. Overall predictive ability

The Brier score was very similar at internal-external validation, internal, and external validation for both outcomes (Table A4). The brier score was somewhat higher at external validation but consistent for all methods (e.g., 0.19 vs. 0.18 for logistic regression to predict unfavorable outcome).

3.5. Explained heterogeneity

At internal-external validation, variation in c-statistic, calibration intercept, and Brier score was mainly attributable to the study in which the algorithm was validated (Table 4): for mortality, the variation in c-statistic was 97% attributable to the study in which the algorithm was

Table 3. Results for discriminative performance of all algorithms, in all three validation strategies: internal-external (per-study CV), internal (10-fold CV), and external (CENTER-TBI) validation

Algorithm	Outcome	Internal-external	Internal	External
Logistic regression	Mortality	0.81 (0.79–0.84)	0.82 (0.81–0.83)	0.82 (0.79–0.84)
Support vector machine		0.81 (0.78–0.83)	0.82 (0.82–0.83)	0.81 (0.79–0.84)
Random forest		0.79 (0.77–0.82)	0.79 (0.78–0.81)	0.81 (0.78–0.84)
Neural network		0.81 (0.79–0.84)	0.82 (0.81–0.83)	0.82 (0.79–0.84)
Gradient boosting machine		0.81 (0.79–0.84)	0.83 (0.82–0.84)	0.83 (0.81–0.86)
Lasso regression		0.81 (0.79–0.84)	0.82 (0.82–0.83)	0.82 (0.79–0.84)
Ridge regression		0.81 (0.79–0.84)	0.82 (0.82–0.83)	0.82 (0.79–0.84)
Logistic regression	Unfavorable outcome	0.81 (0.79–0.83)	0.82 (0.81–0.82)	0.77 (0.75–0.80)
Support vector machine		0.80 (0.79–0.82)	0.81 (0.81–0.82)	0.78 (0.75–0.80)
Random forest		0.79 (0.76–0.81)	0.79 (0.78–0.80)	0.76 (0.74–0.79)
neural network		0.80 (0.79–0.82)	0.81 (0.81–0.82)	0.78 (0.76–0.80)
Gradient boosting machine		0.80 (0.78–0.82)	0.81 (0.80–0.82)	0.78 (0.76–0.80)
Lasso regression		0.81 (0.79–0.83)	0.81 (0.80–0.82)	0.77 (0.75–0.80)
Ridge regression		0.81 (0.79–0.83)	0.81 (0.80–0.82)	0.77 (0.75–0.80)

Abbreviation: CV, cross-validation.

Estimates and 95% CI are shown.

validated (vs. 2.0% to what algorithm was used), whereas the variation in calibration intercept was 98% attributable to the study in which the algorithm was validated (vs. 0.3% to what algorithm was used); and variation in Brier

score was 96% attributable to the study in which the algorithm was validated (vs. 2.0% to what algorithm was used). Variation in calibration slope was slightly more attributable to what algorithm was used, compared with the other

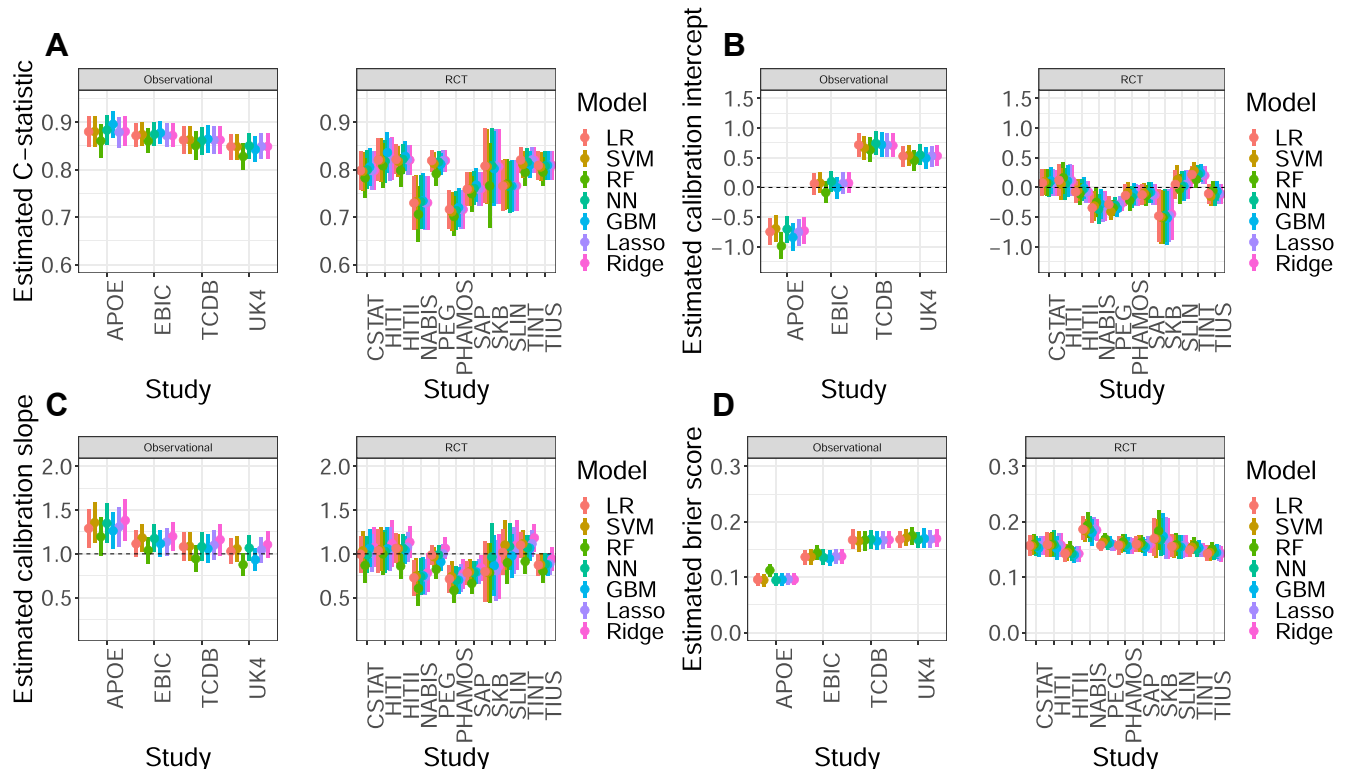
**Fig. 2.** Results of the internal-external cross-validation for mortality. Panel A shows the results of the c-statistic/area under the ROC curve, panel B shows the calibration intercept, panel C shows the calibration slope, and panel D shows the Brier score. The validation results are displayed per study (left: observational, right: randomized controlled trials), and per algorithm. Abbreviations: LR, logistic regression; SVM, support vector machine; RF, random forest; NN, neural network; GBM, gradient boosting machine; ROC, receiver operating curve.

Table 4. Percentage of variation in performance attributable to what study the algorithms were validated in

Outcome	C-statistic	Calibration intercept	Calibration slope	Brier score
Mortality				
Algorithm	2.0	0.3	11	2.0
Study	97	98	86	96
Unfavorable				
Algorithm	2.9	0.0	12	2.5
Study	96	99	85	97

An example is shown in the supplemental material (Fig. 1).

metrics (Fig. A1). For mortality, the variation in calibration slope was 11% attributable to the algorithm used and 86% attributable to the study in which the algorithm was validated. This was mostly caused by the low calibration slope of the RF algorithm. This algorithm displayed the worst calibration slope, as indicated in Fig. 2C. For unfavorable outcome, the results were similar.

3.6. Nonadditivity and nonlinearity

To explore whether nonadditive and nonlinear effects were frequently appropriate to assume in our data, we performed a post hoc analysis. Per study, logistic regression models allowing for nonadditivity and nonlinearity were tested with likelihood ratio tests (omnibus tests) to the model which did not allow for relaxation of those assumptions [20]. It was observed that the model predicting mortality had a better fit when nonlinearity was allowed for in 7 (44%) studies. Less often, the assumption of nonadditivity improved the model fit (Table A6).

4. Discussion

This study aimed to compare flexible ML algorithms to more traditional logistic regression in contemporary patient data. We trained the algorithms to obtain a model with both high discrimination and good calibration. This was achieved by optimizing the log-likelihood for both regression and ML algorithms. All models and algorithms were developed and validated in large data sets, including the recent prospective cohort study CENTER-TBI [27]. Performance was assessed in terms of both discrimination and calibration, which are both important characteristics to be assessed in algorithm validation [22,24,39]. Similar performance of most methods was found across a large number of studies from different time periods.

The algorithm that relatively underperformed was the RF: the discrimination was somewhat lower, but it clearly underperformed in terms of calibration. In particular, the RF showed a calibration slope that was far below one. This indicates overfitting, a problem often arising in small data sets [37]. According to theoretical arguments, the RF algorithm

should not overfit [40]. The discrepancy between the theory and the empirical evidence of our study should be explored further. There could be a role for the selection of hyperparameters, in particular the number of random variables at the split, and the fraction of observations in the training sample [41]. Because the RF shows signs of overfitting, even in large data sets, the discriminative performance should be interpreted with caution: due to optimism, the discrimination in new data sets can be lower [21]. As a contrast, this method was one of the better performing methods in other studies [15,42], which however did not assess calibration. Because calibration is a crucial step before implementation of a prediction model in clinical practice [20,39,43], our study encourages the use of other modeling techniques compared with RFs for outcome prediction.

The variation in observed performance was more explained by the cohorts where the algorithms were validated than by which algorithms were used. This implies that prediction models need continuous updating and validation because their performance is often worse in new cohorts [44]. This is a limitation which needs to be addressed, to effectively use these models in clinical practice [45]. This finding does raise concerns about the validity of individual patient data meta-analysis in the context of prediction modeling.

A recent systematic review compared flexible ML methods to traditional statistical techniques in relatively small data sets (median sample size was 1,250) and did not find incremental value [14]. This was perhaps to be expected because modern ML methods are known to be data hungry compared with classical statistical techniques [13,46]. However, due to the increased sharing of data, international collaborations, and the availability of data from electronic health records and other data sets with routinely collected data, data sets are becoming increasingly large [47–49]. Our study shows that in this situation, flexible ML methods are not improving outcome prognostication as well.

A limitation of our study is that we only used a linear kernel function of SVM. Other kernels could have increased the performance of the algorithm because the performance of the algorithm is substantially dependent on its hyperparameters [41]. Unfortunately, the computation time increased drastically when this kernel was implemented (the expected running time for one series of cross-validation was 21 days). Because the first six iterations did not show substantial increase in discriminative performance, we decided to use the linear basis function instead.

Second, we only considered a relatively small number of predictors (11 predictors, with 19 df). The reason for not including more predictors is that there were no other common data elements between all databases. This potentially limits the performance of ML techniques because it has been suggested that flexible ML techniques perform better than traditional regression techniques when a large number of predictors are being considered, that is, high-dimensional data [50,51]. A reason for such presumed superiority is the flexibility of these algorithms, enabling

them to capture complex nonlinear and interaction effects. It should be noted that regression-based techniques can also be extended by nonlinear and interaction effects [20]. Given that ML algorithms did not outperform regression, these effects are not likely to be essential in the field of outcome prediction in patients with TBI. Our study was not able to fully use the potential benefit of multidimensional data because of a phenomenon that is expected in big data research: larger volumes of data for better models may come at the price of less detailed or lower quality data.

We do believe that although we could perhaps not use the full potential performance of ML algorithms, our comparison is just as relevant. Published ML-based prediction algorithms often include a large number of predictors, sometimes with the suggestion to result in high discriminative performance [52,53]. We note that external validation of these high-dimensional prediction algorithms is challenging because availability of predictors may differ from one setting to the other. For prediction with genomics data, this may be feasible if sufficient standardization and harmonization was performed [54]. However, clinical variables often have different definitions, notations, or units, which complicate the validation procedure with a large number (say $n > 50$) of predictors. External validation remains an essential step before implementing prediction algorithms in clinical practice. To train and validate high-dimensional data, a sophisticated IT environment is necessary [55]. Therefore, we believe that the low-dimensional setting, such as our study, might be more relevant for clinical practice, also for the near future. Powerful predictions for outcome after TBI can apparently be made with linear effects which are captured with simple algorithms.

Finally, this study should be replicated in other fields than TBI to ensure the generalizability of our findings, again from a largely neutral perspective [54]. Preferably, a wide range of studies should be used, representing different settings in terms of study design (randomized controlled trials vs. observational), geography (different countries), types of centers (level I trauma centers vs. other), and so forth. Most studies that compared algorithms used only one or a limited number of study populations [15–19]. Because the performance heavily relies on the study population, comparing the methods in multiple populations is recommended.

5. Conclusion

In a low-dimensional setting, flexible ML algorithms do not perform better than more traditional regression models in outcome prediction after moderate or severe TBI. This is potentially explained by the most important prognostic effects acting as independent, linear effects. Predictive performance is more dependent on the population in which the model is applied than the type of algorithm used. This finding has strong implications: continuous validation and updating of prediction models is necessary to ensure

applicability to new populations of both ML algorithms and regression-based models. To improve prognostication for TBI, future studies should extend current prognostic models with new predictors (biomarkers, imaging, genomics) with strong incremental value, for the reliable identification of patients with poor vs. good prognosis.

Acknowledgments

B.Y.G., Daan N., B.v.C., and E.W.S. contributed to conceptualization; B.Y.G. and Daan N. contributed to data curation and formal analysis; A.E., H.F.L. and E.W.S. contributed to funding acquisition; H.F.L. and E.W.S. contributed to investigation and project administration; B.Y.G., D.N., and E.W.S. contributed to methodology; H.F.L. contributed to resources; Daan N., H.F.L., and E.W.S. contributed to software; Daan N. and H.F.L. contributed to supervision; B.Y.G. contributed to validation, visualization and writing—original draft; All the authors contributed to writing—review & editing.

The CENTER-TBI participants and investigators:

Cecilia Åkerlund¹, Krisztina Amrein², Nada Andelic³, Lasse Andreassen⁴, Audny Anke⁵, Anna Antoni⁶, Gérard Audibert⁷, Philippe Azouvi⁸, Maria Luisa Azzolini⁹, Ronald Bartels¹⁰, Pál Barzó¹¹, Romuald Beauvais¹², Ronny Beer¹³, Bo-Michael Bellander¹⁴, Antonio Belli¹⁵, Habib Benali¹⁶, Maurizio Berardino¹⁷, Luigi Beretta⁹, Morten Blaabjerg¹⁸, Peter Bragge¹⁹, Alexandra Brazinova²⁰, Vibeke Brinck²¹, Joanne Brooker²², Camilla Brorsson²³, Andras Buki²⁴, Monika Bullinger²⁵, Manuel Cabeleira²⁶, Alessio Caccioppola²⁷, Emiliana Calappi²⁷, Maria Rosa Calvi⁹, Peter Cameron²⁸, Guillermo Carbayo Lozano²⁹, Marco Carbonara²⁷, Giorgio Chevallard³⁰, Arturo Chieregato³⁰, Giuseppe Citerio^{31, 32}, Maryse Cnossen³³, Mark Coburn³⁴, Jonathan Coles³⁵, D. Jamie Cooper³⁶, Marta Correia³⁷, Amra Čović³⁸, Nicola Curry³⁹, Endre Czeiter²⁴, Marek Czosnyka²⁶, Claire Dahyot-Fizelier⁴⁰, Helen Dawes⁴¹, Véronique De Keyser⁴², Vincent Degos¹⁶, Francesco Della Corte⁴³, Hugo den Boogert¹⁰, Bart Depreitere⁴⁴, Đula Đilvesi⁴⁵, Abhishek Dixit⁴⁶, Emma Donoghue²², Jens Dreier⁴⁷, Guy-Loup Dulière⁴⁸, Ari Ercole⁴⁶, Patrick Esser⁴¹, Erzsébet Ezer⁴⁹, Martin Fabricius⁵⁰, Valery L. Feigin⁵¹, Kelly Foks⁵², Shirin Frisvold⁵³, Alex Furmanov⁵⁴, Pablo Gagliardo⁵⁵, Damien Galanaud¹⁶, Dashiell Gantner²⁸, Guoyi Gao⁵⁶, Pradeep George⁵⁷, Alexandre Ghuysen⁵⁸, Lelde Giga⁵⁹, Ben Glocker⁶⁰, Jagoš Golubovic⁴⁵, Pedro A. Gomez⁶¹, Johannes Gratz⁶², Benjamin Gravesteijn³³, Francesca Grossi⁴³, Russell L. Gruen⁶³, Deepak Gupta⁶⁴, Juanita A. Haagsma³³, Iain Haitsma⁶⁵, Raimund Helbok¹³, Eirik Helseth⁶⁶, Lindsay Horton⁶⁷, Jilske Huijben³³, Peter J. Hutchinson⁶⁸, Bram Jacobs⁶⁹, Stefan Jankowski⁷⁰, Mike Jarrett²¹, Ji-yao Jiang⁵⁶, Kelly Jones⁵¹, Mladen Karan⁴⁷, Angelos G. Kolias⁶⁸, Erwin Kompanje⁷¹, Daniel Kondziella⁵⁰, Evgenios Koraropoulos⁴⁶, Lars-Owe Koskinen⁷², Noémi Kovács⁷³, Alfonso

Lagares⁶¹, Linda Lanyon⁵⁷, Steven Laureys⁷⁴, Fiona Lecky⁷⁵, Rolf Lefering⁷⁶, Valerie Legrand⁷⁷, Aurelie Lejeune⁷⁸, Leon Levi⁷⁹, Roger Lightfoot⁸⁰, Hester Lingsma³³, Andrew I.R. Maas⁴², Ana M. Castaño-León⁶¹, Marc Maegele⁸¹, Marek Majdan²⁰, Alex Manara⁸², Geoffrey Manley⁸³, Costanza Martino⁸⁴, Hugues Maréchal⁴⁸, Julia Mattern⁸⁵, Catherine McMahon⁸⁶, Béla Melegh⁸⁷, David Menon⁴⁶, Tomas Menovsky⁴², Davide Mulazzi²⁷, Visakh Muraleedharan⁵⁷, Lynnette Murray²⁸, Nandesh Nair⁴², Ancuta Negru⁸⁸, David Nelson¹, Virginia Newcombe⁴⁶, Daan Nieboer⁵³, Quentin Noirhomme⁷⁴, József Nyirádi², Otesile Olubukola⁷⁵, Matej Oresic⁸⁹, Fabrizio Ortolano²⁷, Aarno Palotie^{90, 91, 92}, Paul M. Parizel⁹³, Jean-François Payen⁹⁴, Natascha Perera¹², Vincent Perlberg¹⁶, Paolo Persona⁹⁵, Wilco Peul⁹⁶, Anna Piippo-Karjalainen⁹⁷, Matti Pirinen⁹⁰, Horia Ples⁸⁸, Suzanne Polinder³³, Inigo Pomposo²⁹, Jussi P. Posti⁹⁸, Louis Puybasset⁹⁹, Andreea Radoi¹⁰⁰, Arminas Ragauskas¹⁰¹, Rahul Raj⁹⁷, Malinka Rambadagalla¹⁰², Ruben Real³⁸, Jonathan Rhodes¹⁰³, Sylvia Richardson¹⁰⁴, Sophie Richter⁴⁶, Samuli Ripatti⁹⁰, Saulius Rocka¹⁰¹, Cecilie Roe¹⁰⁵, Olav Roise^{106, 140}, Jonathan Rosand¹⁰⁷, Jeffrey V. Rosenfeld¹⁰⁸, Christina Rosenlund¹⁰⁹, Guy Rosenthal⁵⁴, Rolf Rossaint³⁴, Sandra Rossi⁹⁵, Daniel Rueckert⁶⁰, Martin Rusnák¹¹⁰, Juan Sahuquillo¹⁰⁰, Oliver Sakowitz^{85, 111}, Renan Sanchez-Porras¹¹¹, Janos Sandor¹¹², Nadine Schäfer⁷⁶, Silke Schmidt¹¹³, Herbert Schoechl¹¹⁴, Guus Schoonman¹¹⁵, Rico Frederik Schou¹¹⁶, Elisabeth Schwendenwein⁶, Charlie Sewalt³³, Toril Skandsen^{117, 118}, Peter Smielewski²⁶, Abayomi Sorinola¹¹⁹, Emmanuel Stamatakis⁴⁶, Simon Stanworth³⁹, Ana Kowark³⁴, Robert Stevens¹²⁰, William Stewart¹²¹, Ewout W. Steyerberg^{33, 122}, Nino Stocchetti¹²³, Nina Sundström¹²⁴, Anneliese Synnot^{22, 125}, Riikka Takala¹²⁶, Viktória Tamás¹¹⁹, Tomas Tamosuitis¹²⁷, Mark Steven Taylor²⁰, Braden Te Ao⁵¹, Olli Tenovuo⁹⁸, Alice Theadom⁵¹, Matt Thomas⁸², Dick Tibboel¹²⁸, Marjolein Timmers⁷¹, Christos Tolia¹²⁹, Tony Tzani²⁸, Cristina Maria Tudora⁸⁸, Peter Vajkoczy¹³⁰, Shirley Vallance²⁸, Egils Valeinis⁵⁹, Zoltán Vámos⁴⁹, Gregory Van der Steen⁴², Joukje van der Naalt⁶⁹, Jeroen T.J.M. van Dijk⁹⁶, Thomas A. van Essen⁹⁶, Wim Van Hecke¹³¹, Caroline van Heugten¹³², Dominique Van Praag¹³³, Thijs Vande Vyvere¹³¹, Audrey Vanhauudenhuysse^{16, 74}, Roel P. J. van Wijk⁹⁷, Alessia Vargiolu³², Emmanuel Vega⁷⁹, Kimberley Velt³³, Jan Verheyden¹³¹, Paul M. Vespa¹³⁴, Anne Vik^{117, 135}, Rimantas Vilcinis¹²⁷, Victor Volovici⁶⁵, Nicole von Steinbüchel³⁸, Daphne Voormolen³³, Petar Vulekovic⁴⁵, Kevin K.W. Wang¹³⁶, Eveline Wiegers³³, Guy Williams⁴⁶, Lindsay Wilson⁶⁷, Stefan Winzeck⁴⁶, Stefan Wolf¹³⁷, Zhihui Yang¹³⁶, Peter Ylén¹³⁸, Alexander Younsi⁸⁵, Frederik A. Zeiler^{46, 139}, Veronika Zelinkova²⁰, Agate Ziverte⁵⁹, Tommaso Zoerle²⁷

¹ Department of Physiology and Pharmacology, Section of Perioperative Medicine and Intensive Care, Karolinska Institutet, Stockholm, Sweden.

² János Szentágothai Research Centre, University of Pécs, Pécs, Hungary.

³ Division of Surgery and Clinical Neuroscience, Department of Physical Medicine and Rehabilitation, Oslo University Hospital and University of Oslo, Oslo, Norway.

⁴ Department of Neurosurgery, University Hospital Northern Norway, Tromsø, Norway.

⁵ Department of Physical Medicine and Rehabilitation, University Hospital Northern Norway, Tromsø, Norway.

⁶ Trauma Surgery, Medical University Vienna, Vienna, Austria.

⁷ Department of Anesthesiology & Intensive Care, University Hospital Nancy, Nancy, France.

⁸ Raymond Poincaré hospital, Assistance Publique—Hopitaux de Paris, Paris, France.

⁹ Department of Anesthesiology & Intensive Care, S Raffaele University Hospital, Milan, Italy.

¹⁰ Department of Neurosurgery, Radboud University Medical Center, Nijmegen, The Netherlands.

¹¹ Department of Neurosurgery, University of Szeged, Szeged, Hungary.

¹² International Projects Management, ARTTIC, München, Germany.

¹³ Department of Neurology, Neurological Intensive Care Unit, Medical University of Innsbruck, Innsbruck, Austria.

¹⁴ Department of Neurosurgery & Anesthesia & intensive care medicine, Karolinska University Hospital, Stockholm, Sweden.

¹⁵ NIHR Surgical Reconstruction and Microbiology Research Centre, Birmingham, UK.

¹⁶ Anesthésie-Réanimation, Assistance Publique—Hopitaux de Paris, Paris, France.

¹⁷ Department of Anesthesia & ICU, AOU Città della Salute e della Scienza di Torino—Orthopedic and Trauma Center, Torino, Italy.

¹⁸ Department of Neurology, Odense University Hospital, Odense, Denmark.

¹⁹ BehaviourWorks Australia, Monash Sustainability Institute, Monash University, Victoria, Australia.

²⁰ Department of Public Health, Faculty of Health Sciences and Social Work, Trnava University, Trnava, Slovakia.

²¹ Quesgen Systems Inc., Burlingame, California, USA.

²² Australian & New Zealand Intensive Care Research Centre, Department of Epidemiology and Preventive Medicine, School of Public Health and Preventive Medicine, Monash University, Melbourne, Australia.

²³ Department of Surgery and Perioperative Science, Umeå University, Umeå, Sweden.

²⁴ Department of Neurosurgery, Medical School, University of Pécs, Hungary and Neurotrauma Research Group, János Szentágothai Research Centre, University of Pécs, Hungary.

²⁵ Department of Medical Psychology, Universitätsklinikum Hamburg-Eppendorf, Hamburg, Germany.

²⁶ Brain Physics Lab, Division of Neurosurgery, Dept of Clinical Neurosciences, University of Cambridge, Addenbrooke's Hospital, Cambridge, UK.

²⁷ Neuro ICU, Fondazione IRCCS Cà Granda Ospedale Maggiore Policlinico, Milan, Italy.

- ²⁸ ANZIC Research Centre, Monash University, Department of Epidemiology and Preventive Medicine, Melbourne, Victoria, Australia.
- ²⁹ Department of Neurosurgery, Hospital of Cruces, Bilbao, Spain.
- ³⁰ NeuroIntensive Care, Niguarda Hospital, Milan, Italy.
- ³¹ School of Medicine and Surgery, Università Milano Bicocca, Milano, Italy.
- ³² NeuroIntensive Care, ASST di Monza, Monza, Italy.
- ³³ Department of Public Health, Erasmus Medical Center-University Medical Center, Rotterdam, The Netherlands.
- ³⁴ Department of Anaesthesiology, University Hospital of Aachen, Aachen, Germany.
- ³⁵ Department of Anesthesia & Neurointensive Care, Cambridge University Hospital NHS Foundation Trust, Cambridge, UK.
- ³⁶ School of Public Health & PM, Monash University and The Alfred Hospital, Melbourne, Victoria, Australia.
- ³⁷ Radiology/MRI department, MRC Cognition and Brain Sciences Unit, Cambridge, UK.
- ³⁸ Institute of Medical Psychology and Medical Sociology, Universitätsmedizin Göttingen, Göttingen, Germany.
- ³⁹ Oxford University Hospitals NHS Trust, Oxford, UK.
- ⁴⁰ Intensive Care Unit, CHU Poitiers, Poitiers, France.
- ⁴¹ Movement Science Group, Faculty of Health and Life Sciences, Oxford Brookes University, Oxford, UK.
- ⁴² Department of Neurosurgery, Antwerp University Hospital and University of Antwerp, Edegem, Belgium.
- ⁴³ Department of Anesthesia & Intensive Care, Maggiore Della Carità Hospital, Novara, Italy.
- ⁴⁴ Department of Neurosurgery, University Hospitals Leuven, Leuven, Belgium.
- ⁴⁵ Department of Neurosurgery, Clinical centre of Vojvodina, Faculty of Medicine, University of Novi Sad, Novi Sad, Serbia.
- ⁴⁶ Division of Anaesthesia, University of Cambridge, Addenbrooke's Hospital, Cambridge, UK.
- ⁴⁷ Center for Stroke Research Berlin, Charité—Universitätsmedizin Berlin, corporate member of Freie Universität Berlin, Humboldt-Universität zu Berlin, and Berlin Institute of Health, Berlin, Germany.
- ⁴⁸ Intensive Care Unit, CHR Citadelle, Liège, Belgium.
- ⁴⁹ Department of Anaesthesiology and Intensive Therapy, University of Pécs, Pécs, Hungary.
- ⁵⁰ Departments of Neurology, Clinical Neurophysiology and Neuroanesthesiology, Region Hovedstaden Rigshospitalet, Copenhagen, Denmark.
- ⁵¹ National Institute for Stroke and Applied Neurosciences, Faculty of Health and Environmental Studies, Auckland University of Technology, Auckland, New Zealand.
- ⁵² Department of Neurology, Erasmus MC, Rotterdam, the Netherlands.
- ⁵³ Department of Anesthesiology and Intensive care, University Hospital Northern Norway, Tromsø, Norway.
- ⁵⁴ Department of Neurosurgery, Hadassah-hebrew University Medical center, Jerusalem, Israel.
- ⁵⁵ Fundación Instituto Valenciano de Neurorehabilitación (FIVAN), Valencia, Spain.
- ⁵⁶ Department of Neurosurgery, Shanghai Renji hospital, Shanghai Jiaotong University/school of medicine, Shanghai, China.
- ⁵⁷ Karolinska Institutet, INCF International Neuroinformatics Coordinating Facility, Stockholm, Sweden.
- ⁵⁸ Emergency Department, CHU, Liège, Belgium.
- ⁵⁹ Neurosurgery clinic, Pauls Stradins Clinical University Hospital, Riga, Latvia.
- ⁶⁰ Department of Computing, Imperial College London, London, UK.
- ⁶¹ Department of Neurosurgery, Hospital Universitario 12 de Octubre, Madrid, Spain.
- ⁶² Department of Anesthesia, Critical Care and Pain Medicine, Medical University of Vienna, Austria.
- ⁶³ College of Health and Medicine, Australian National University, Canberra, Australia.
- ⁶⁴ Department of Neurosurgery, Neurosciences Centre & JPN Apex trauma centre, All India Institute of Medical Sciences, New Delhi-110029, India.
- ⁶⁵ Department of Neurosurgery, Erasmus MC, Rotterdam, the Netherlands.
- ⁶⁶ Department of Neurosurgery, Oslo University Hospital, Oslo, Norway.
- ⁶⁷ Division of Psychology, University of Stirling, Stirling, UK.
- ⁶⁸ Division of Neurosurgery, Department of Clinical Neurosciences, Addenbrooke's Hospital & University of Cambridge, Cambridge, UK.
- ⁶⁹ Department of Neurology, University of Groningen, University Medical Center Groningen, Groningen, Netherlands.
- ⁷⁰ Neurointensive Care, Sheffield Teaching Hospitals NHS Foundation Trust, Sheffield, UK.
- ⁷¹ Department of Intensive Care and Department of Ethics and Philosophy of Medicine, Erasmus Medical Center, Rotterdam, The Netherlands.
- ⁷² Department of Clinical Neuroscience, Neurosurgery, Umeå University, Umeå, Sweden.
- ⁷³ Hungarian Brain Research Program - Grant No. KTIA_13_NAP-A-II/8, University of Pécs, Pécs, Hungary.
- ⁷⁴ Cyclotron Research Center, University of Liège, Liège, Belgium.
- ⁷⁵ Emergency Medicine Research in Sheffield, Health Services Research Section, School of Health and Related Research (ScHARR), University of Sheffield, Sheffield, UK.
- ⁷⁶ Institute of Research in Operative Medicine (IFOM), Witten/Herdecke University, Cologne, Germany.
- ⁷⁷ VP Global Project Management CNS, ICON, Paris, France.
- ⁷⁸ Department of Anesthesiology-Intensive Care, Lille University Hospital, Lille, France.

- ⁷⁹ Department of Neurosurgery, Rambam Medical Center, Haifa, Israel.
- ⁸⁰ Department of Anesthesiology & Intensive Care, University Hospitals Southampton NHS Trust, Southampton, UK.
- ⁸¹ Cologne-Merheim Medical Center (CMMC), Department of Traumatology, Orthopedic Surgery and Sportmedicine, Witten/Herdecke University, Cologne, Germany.
- ⁸² Intensive Care Unit, Southmead Hospital, Bristol, Bristol, UK.
- ⁸³ Department of Neurological Surgery, University of California, San Francisco, California, USA.
- ⁸⁴ Department of Anesthesia & Intensive Care, M. Bufalini Hospital, Cesena, Italy.
- ⁸⁵ Department of Neurosurgery, University Hospital Heidelberg, Heidelberg, Germany.
- ⁸⁶ Department of Neurosurgery, The Walton centre NHS Foundation Trust, Liverpool, UK.
- ⁸⁷ Department of Medical Genetics, University of Pécs, Pécs, Hungary.
- ⁸⁸ Department of Neurosurgery, Emergency County Hospital Timisoara, Timisoara, Romania.
- ⁸⁹ School of Medical Sciences, Örebro University, Örebro, Sweden.
- ⁹⁰ Institute for Molecular Medicine Finland, University of Helsinki, Helsinki, Finland.
- ⁹¹ Analytic and Translational Genetics Unit, Department of Medicine; Psychiatric & Neurodevelopmental Genetics Unit, Department of Psychiatry; Department of Neurology, Massachusetts General Hospital, Boston, MA, USA.
- ⁹² Program in Medical and Population Genetics; The Stanley Center for Psychiatric Research, The Broad Institute of MIT and Harvard, Cambridge, MA, USA.
- ⁹³ Department of Radiology, Antwerp University Hospital and University of Antwerp, Edegem, Belgium.
- ⁹⁴ Department of Anesthesiology & Intensive Care, University Hospital of Grenoble, Grenoble, France.
- ⁹⁵ Department of Anesthesia & Intensive Care, Azienda Ospedaliera Università di Padova, Padova, Italy.
- ⁹⁶ Dept. of Neurosurgery, Leiden University Medical Center, Leiden, The Netherlands and Dept. of Neurosurgery, Medical Center Haaglanden, The Hague, The Netherlands.
- ⁹⁷ Department of Neurosurgery, Helsinki University Central Hospital.
- ⁹⁸ Division of Clinical Neurosciences, Department of Neurosurgery and Turku Brain Injury Centre, Turku University Hospital and University of Turku, Turku, Finland.
- ⁹⁹ Department of Anesthesiology and Critical Care, Pitié-Salpêtrière Teaching Hospital, Assistance Publique, Hôpitaux de Paris and University Pierre et Marie Curie, Paris, France.
- ¹⁰⁰ Neurotraumatology and Neurosurgery Research Unit (UNINN), Vall d'Hebron Research Institute, Barcelona, Spain.
- ¹⁰¹ Department of Neurosurgery, Kaunas University of technology and Vilnius University, Vilnius, Lithuania.
- ¹⁰² Department of Neurosurgery, Rezekne Hospital, Latvia.
- ¹⁰³ Department of Anaesthesia, Critical Care & Pain Medicine NHS Lothian & University of Edinburgh, Edinburgh, UK.
- ¹⁰⁴ Director, MRC Biostatistics Unit, Cambridge Institute of Public Health, Cambridge, UK.
- ¹⁰⁵ Department of Physical Medicine and Rehabilitation, Oslo University Hospital/University of Oslo, Oslo, Norway.
- ¹⁰⁶ Division of Orthopedics, Oslo University Hospital.
- ¹⁰⁷ Broad Institute, Cambridge MA Harvard Medical School, Boston MA, Massachusetts General Hospital, Boston MA, USA.
- ¹⁰⁸ National Trauma Research Institute, The Alfred Hospital, Monash University, Melbourne, Victoria, Australia.
- ¹⁰⁹ Department of Neurosurgery, Odense University Hospital, Odense, Denmark.
- ¹¹⁰ International Neurotrauma Research Organisation, Vienna, Austria.
- ¹¹¹ Klinik für Neurochirurgie, Klinikum Ludwigsburg, Ludwigsburg, Germany.
- ¹¹² Division of Biostatistics and Epidemiology, Department of Preventive Medicine, University of Debrecen, Debrecen, Hungary.
- ¹¹³ Department Health and Prevention, University Greifswald, Greifswald, Germany.
- ¹¹⁴ Department of Anaesthesiology and Intensive Care, AUVA Trauma Hospital, Salzburg, Austria.
- ¹¹⁵ Department of Neurology, Elisabeth-TweeSteden Ziekenhuis, Tilburg, the Netherlands.
- ¹¹⁶ Department of Neuroanesthesia and Neurointensive Care, Odense University Hospital, Odense, Denmark.
- ¹¹⁷ Department of Neuromedicine and Movement Science, Norwegian University of Science and Technology, NTNU, Trondheim, Norway.
- ¹¹⁸ Department of Physical Medicine and Rehabilitation, St. Olavs Hospital, Trondheim University Hospital, Trondheim, Norway.
- ¹¹⁹ Department of Neurosurgery, University of Pécs, Pécs, Hungary.
- ¹²⁰ Division of Neuroscience Critical Care, John Hopkins University School of Medicine, Baltimore, USA.
- ¹²¹ Department of Neuropathology, Queen Elizabeth University Hospital and University of Glasgow, Glasgow, UK.
- ¹²² Dept. of Department of Biomedical Data Sciences, Leiden University Medical Center, Leiden, The Netherlands.
- ¹²³ Department of Pathophysiology and Transplantation, Milan University, and Neuroscience ICU, Fondazione IRCCS Cà Granda Ospedale Maggiore Policlinico, Milano, Italy.
- ¹²⁴ Department of Radiation Sciences, Biomedical Engineering, Umeå University, Umeå, Sweden.
- ¹²⁵ Cochrane Consumers and Communication Review Group, Centre for Health Communication and Participation, School of Psychology and Public Health, La Trobe University, Melbourne, Australia.

¹²⁶ Perioperative Services, Intensive Care Medicine and Pain Management, Turku University Hospital and University of Turku, Turku, Finland.

¹²⁷ Department of Neurosurgery, Kaunas University of Health Sciences, Kaunas, Lithuania.

¹²⁸ Intensive Care and Department of Pediatric Surgery, Erasmus Medical Center, Sophia Children's Hospital, Rotterdam, The Netherlands.

¹²⁹ Department of Neurosurgery, Kings college London, London, UK.

¹³⁰ Neurologie, Neurochirurgie und Psychiatrie, Charité—Universitätsmedizin Berlin, Berlin, Germany.

¹³¹ icoMetrix NV, Leuven, Belgium.

¹³² Movement Science Group, Faculty of Health and Life Sciences, Oxford Brookes University, Oxford, UK.

¹³³ Psychology Department, Antwerp University Hospital, Edegem, Belgium.

¹³⁴ Director of Neurocritical Care, University of California, Los Angeles, USA.

¹³⁵ Department of Neurosurgery, St.Olavs Hospital, Trondheim University Hospital, Trondheim, Norway.

¹³⁶ Department of Emergency Medicine, University of Florida, Gainesville, Florida, USA.

¹³⁷ Department of Neurosurgery, Charité—Universitätsmedizin Berlin, corporate member of Freie Universität Berlin, Humboldt-Universität zu Berlin, and Berlin Institute of Health, Berlin, Germany.

¹³⁸ VTT Technical Research Centre, Tampere, Finland.

¹³⁹ Section of Neurosurgery, Department of Surgery, Rady Faculty of Health Sciences, University of Manitoba, Winnipeg, MB, Canada.

¹⁴⁰ Institute of Clinical Medicine, Faculty of Medicine, University of Oslo.

Supplementary data

Supplementary data to this article can be found online at <https://doi.org/10.1016/j.jclinepi.2020.03.005>.

References

- [1] Maas AIR, Menon DK, Adelson PD, Andelic N, Bell MJ, Belli A, et al. Traumatic brain injury: integrated approaches to improve prevention, clinical care, and research. *Lancet Neurol* 2017;16:987.
- [2] Majdan M, Plancikova D, Brazinova A, Rusnak M, Nieboer D, Feigin V, et al. Epidemiology of traumatic brain injuries in Europe: a cross-sectional analysis. *Lancet Public Heal* 2016;1:e76–83.
- [3] Saatman KE, Duhaime A, Bullock R, Maas AI, Valadka A, Manley GT, et al. Classification of traumatic brain injury for targeted therapies. *J Neurotrauma* 2008;25:719–38.
- [4] Lingsma HF, Roozenbeek B, Steyerberg EW, Murray GD, Maas AI. Early prognosis in traumatic brain injury: from prophecies to predictions. *Lancet Neurol* 2010;9:543–54.
- [5] Liu NT, Salinas J. Machine learning for predicting outcomes in trauma. *Shock* 2017;48:504–10.
- [6] Burges CJC. A tutorial on support vector machines for pattern recognition. *Data Mining Knowledge Discovery* 1998;2:121–67.
- [7] Jain AK, Mao J, Mohiuddin KM. Artificial neural networks: a tutorial. *Computer (Long Beach Calif)* 1996;29:31–44.
- [8] Afanador NL, Smolinska A, Tran TN, Blanchet L. Unsupervised random forest: a tutorial with case studies. *J Chemom* 2016;30:232–41.
- [9] Natekin A, Knoll A. Gradient boosting machines, a tutorial. *Front Neurorobot* 2013;7:21.
- [10] Rau C-S, Kuo P-J, Chien P-C, Huang C-Y, Hsieh H-Y, Hsieh C-H. Mortality prediction in patients with isolated moderate and severe traumatic brain injury using machine learning models. *PLoS One* 2018;13:e0207192.
- [11] Matsuo K, Aihara H, Nakai T, Morishita A, Tohma Y, Kohmura E. Machine learning to predict in-hospital Morbidity and mortality after traumatic brain injury. *J Neurotrauma* 2018;37:202–10.
- [12] Feng J, Wang Y, Peng J, Sun M, Zeng J, Jiang H. Comparison between logistic regression and machine learning algorithms on survival prediction of traumatic brain injuries. *J Crit Care* 2019;54:110–6.
- [13] van der Ploeg T, Austin PC, Steyerberg EW. Modern modelling techniques are data hungry: a simulation study for predicting dichotomous endpoints. *BMC Med Res Methodol* 2014;14:137.
- [14] Christodoulou E, Jie MA, Collins GS, Steyerberg EW, Verbakel JY, van Calster B. A systematic review shows no performance benefit of machine learning over logistic regression for clinical prediction models. *J Clin Epidemiol* 2019;110:12–22.
- [15] van Os HJA, Ramos LA, Hilbert A, van Leeuwen M, van Walderveen MAA, Kruij ND, et al. Predicting outcome of endovascular treatment for acute ischemic stroke: potential value of machine learning algorithms. *Front Neurol* 2018;9:784.
- [16] Churpek MM, Yuen TC, Winslow C, Meltzer DO, Kattan MW, Edelson DP. Multicenter comparison of machine learning methods and conventional regression for predicting clinical deterioration on the wards. *Crit Care Med* 2016;44:368–74.
- [17] Lee H-C, Yoon S, Yang S-M, Kim W, Ryu H-G, Jung C-W, et al. Prediction of acute kidney injury after liver transplantation: machine learning approaches vs. logistic regression model. *J Clin Med* 2018;7:428.
- [18] Bisaso KR, Karungi SA, Kiragga A, Mukonzo JK, Castelnovo B. A comparative study of logistic regression based machine learning techniques for prediction of early virological suppression in antiretroviral initiating HIV patients. *BMC Med Inform Decis Mak* 2018;18:77.
- [19] Decruyenaere A, Decruyenaere P, Peeters P, Vermassen F, Dhaene T, Couckuyt I. Prediction of delayed graft function after kidney transplantation: comparison between logistic regression and machine learning methods. *BMC Med Inform Decis Mak* 2015;15:83.
- [20] Harrell FE. *Regression Modeling Strategies*. New York, NY: Springer New York; 2001.
- [21] Steyerberg EW. *Clinical Prediction Models*. New York, NY: Springer New York; 2009.
- [22] Van Calster B, Nieboer D, Vergouwe Y, De Cock B, Pencina MJ, Steyerberg EW. A calibration hierarchy for risk models was defined: from utopia to empirical data. *J Clin Epidemiol* 2016;74:167–76.
- [23] Collins GS, Reitsma JB, Altman DG, Moons KGM. Transparent reporting of a multivariable prediction model for individual prognosis or diagnosis (TRIPOD): the TRIPOD Statement. *Ann Intern Med* 2015;162:55.
- [24] Moons KGM, Kengne AP, Grobbee DE, Royston P, Vergouwe Y, Altman DG, et al. Risk prediction models: II. External validation, model updating, and impact assessment. *Heart* 2012;98:691–8.
- [25] Steyerberg EW, Mushkudiani N, Perel P, Butcher I, Lu J, McHugh GS, et al. Predicting outcome after traumatic brain injury: development and international validation of prognostic scores based on admission characteristics. *PLoS Med* 2008;5:e165.
- [26] Marmarou A, Lu J, Butcher I, McHugh GS, Mushkudiani NA, Murray GD, et al. IMPACT database of traumatic brain injury: design and description. *J Neurotrauma* 2007;24:239–50.

- [27] Maas AIR, Menon DK, Steyerberg EW, Citerio G, Lecky F, Manley GT, et al. Collaborative European neurotrauma effectiveness research in traumatic brain injury (CENTER-TBI): a prospective longitudinal observational study. *Neurosurgery* 2015;76: 67–80.
- [28] Steyerberg EW. Clinical prediction models: a practical approach to development, validation and updating. New York: Springer; 2009.
- [29] Tibshirani R. Regression shrinkage and selection via the lasso. *J R Stat Soc Ser B* 1996;58:267–88.
- [30] FIRTH D. Bias reduction of maximum likelihood estimates. *Biometrika* 1993;80:27–38.
- [31] Rubin DB. Multiple imputation for nonresponse in surveys. Hoboken, New Jersey: John Wiley & Sons; 2004.
- [32] Buuren S van. Flexible imputation of missing data. Cleveland, Ohio: CRC Press; 2018.
- [33] Royston P, Parmar MKB, Sylvester R. Construction and validation of a prognostic model across several studies, with an application in superficial bladder cancer. *Stat Med* 2004;23:907–26.
- [34] Steyerberg EW, Harrell FE. Prediction models need appropriate internal, internal-external, and external validation HHS Public Access. *J Clin Epidemiol* 2016;69:245–7.
- [35] Cox D. Two further applications of a model for binary regression. *Biometrika* 1958;45:562–5.
- [36] DeLong ER, DeLong DM, Clarke-Pearson DL. Comparing the areas under two or more correlated receiver operating characteristic curves: a nonparametric approach. *Biometrics* 1988;44:837–45.
- [37] Steyerberg EW, Vickers AJ, Cook NR, Gerds T, Gonen M, Obuchowski N, et al. Assessing the performance of prediction models: a framework for traditional and novel measures. *Epidemiology* 2010;21:128–38.
- [38] DerSimonian R, Laird N. Meta-analysis in clinical trials. *Control Clin Trials* 1986;7:177–88.
- [39] Steyerberg EW, Vergouwe Y. Towards better clinical prediction models: Seven steps for development and an ABCD for validation. *Eur Heart J* 2014;35:1925–31.
- [40] Breiman L. Random Forests. *Mach Learn* 2001;45:5–32.
- [41] Probst P, Boulesteix A-L, Bischl B. Tunability: importance of hyperparameters of machine learning algorithms. *J Mach Learn Res* 2019; 20:1–32.
- [42] Sakr S, Elshawi R, Ahmed AM, Qureshi WT, Brawner CA, Keteyian SJ, et al. Comparison of machine learning techniques to predict all-cause mortality using fitness data: the Henry ford exercise testing (FIT) project. *BMC Med Inform Decis Mak* 2017;17:174.
- [43] König IR, Malley JD, Weimar C, Diener H-C, Ziegler A, German Stroke Study Collaboration. Practical experiences on the necessity of external validation. *Stat Med* 2007;26:5499–511.
- [44] Thelin EP, Nelson DW, Vehvilä Inen J, Nyström H, Kivisaari R, Siironen J, et al. Evaluation of novel computerized tomography scoring systems in human traumatic brain injury: an observational, multicenter study. *PLoS Med* 2017;14:e1002368.
- [45] Ngiam KY, Khor IW. Big data and machine learning algorithms for health-care delivery. *Lancet Oncol* 2019;20:e262–73.
- [46] van der Ploeg T, Nieboer D, Steyerberg EW. Modern modeling techniques had limited external validity in predicting mortality from traumatic brain injury. *J Clin Epidemiol* 2016;78:83–9.
- [47] Poldrack RA, Gorgolewski KJ. Making big data open: data sharing in neuroimaging. *Nat Neurosci* 2014;17:1510–7.
- [48] Neurology TL. The changing landscape of traumatic brain injury research. *Lancet Neurol* 2012;11:651.
- [49] Charles D, Gabriel M, Searcy T. Adoption of electronic health record systems among U.S. non-federal acute care hospitals: 2008–2014. *ONC Data Brief* 2016;35:1–9.
- [50] Rajkomar A, Oren E, Chen K, Dai AM, Hajaj N, Hardt M, et al. Scalable and accurate deep learning with electronic health records. *NPJ Digit Med* 2018;1:18.
- [51] Beam AL, Kohane IS. Big data and machine learning in health care. *JAMA* 2018;319:1317.
- [52] Delahanty RJ, Kaufman D, Jones SS. Development and evaluation of an automated machine learning algorithm for in-hospital mortality risk adjustment among critical care patients. *Crit Care Med* 2018; 46:e481–8.
- [53] Desautels T, Das R, Calvert J, Trivedi M, Summers C, Wales DJ, et al. Prediction of early unplanned intensive care unit readmission in a UK tertiary care hospital: a cross-sectional machine learning approach. *BMJ Open* 2017;7:e017199.
- [54] Klau S, Jurinovic V, Hornung R, Herold T, Boulesteix A-L. Priority-Lasso: a simple hierarchical approach to the prediction of clinical outcome using multi-omics data. *BMC Bioinformatics* 2018;19:322.
- [55] Reps JM, Schuemie MJ, Suchard MA, Ryan PB, Rijnbeek PR. Design and implementation of a standardized framework to generate and evaluate patient-level prediction models using observational healthcare data. *J Am Med Inform Assoc* 2018;25:969–75.