

Statistical tutorials

Towards better clinical prediction models: seven steps for development and an ABCD for validation

Ewout W. Steyerberg* and Yvonne Vergouwe

Department of Public Health, Erasmus MC, University Medical Center Rotterdam, PO Box 2040, 3000 CA Rotterdam, The Netherlands

Received 7 October 2013; revised 22 April 2014; accepted 30 April 2014; online publish-ahead-of-print 4 June 2014

Clinical prediction models provide risk estimates for the presence of disease (diagnosis) or an event in the future course of disease (prognosis) for individual patients. Although publications that present and evaluate such models are becoming more frequent, the methodology is often suboptimal. We propose that seven steps should be considered in developing prediction models: (i) consideration of the research question and initial data inspection; (ii) coding of predictors; (iii) model specification; (iv) model estimation; (v) evaluation of model performance; (vi) internal validation; and (vii) model presentation. The validity of a prediction model is ideally assessed in fully independent data, where we propose four key measures to evaluate model performance: calibration-in-the-large, or the model intercept (A); calibration slope (B); discrimination, with a concordance statistic (C); and clinical usefulness, with decision-curve analysis (D). As an application, we develop and validate prediction models for 30-day mortality in patients with an acute myocardial infarction. This illustrates the usefulness of the proposed framework to strengthen the methodological rigour and quality for prediction models in cardiovascular research.

Keywords

Prediction model • Non-linearity • Missing values • Shrinkage • Calibration • Discrimination • Clinical usefulness

Introduction

Prediction of the presence of disease (diagnosis) or an event in the future course of disease (prognosis) becomes more and more important in the current era of personalized medicine. Improvements in imaging, biomarkers, and 'omics' research lead to many new predictors for diagnosis and prognosis.

Clinical prediction models may combine multiple predictors to provide insight into the relative effects of predictors in the model. For example, we may be interested in the independent prognostic value of inflammatory markers such as C-reactive protein for the clinical course and outcome of an acute coronary syndrome.¹ Clinical prediction models may also provide absolute risk estimates for individual patients.^{2,3} We focus here on the second role of such models, which are commonly developed with regression analysis techniques. Logistic regression analysis is most commonly used for the prediction of binary events, such as 30-day mortality. Cox regression is most common for time-to-event data, such as long-term mortality. We focus on prediction models for binary events, and indicate differences with time-to-event data where most relevant.

We note that rigorous development and validation of prediction models is important. Despite a substantial body of methodological literature and published guidance on how to perform prediction research, most models suffer from methodological shortcomings, or are at least reported poorly.^{4–7} We propose a systematic approach to model development and validation, illustrated with prediction of 30-day mortality in patients suffering from an acute myocardial infarction. We compare the performance of a simple model (including age only) to a more complex model (including age and other key predictors), which are developed either in a small or a large data set.

Case study: prediction of 30-day mortality in acute myocardial infarction

As an illustrative case study, we consider a cohort of 40 830 patients suffering from an acute myocardial infarction who were enrolled in the GUSTO-I trial. The data from this trial have been used for many analyses, including the development of a prediction model for 30-day mortality⁸ and various methodological studies.^{9–15}

* Corresponding author. Tel: +31 107038470, Fax: +31 104089455, Email: e.steyerberg@erasmusmc.nl

Published on behalf of the European Society of Cardiology. All rights reserved. © The Author 2014. For permissions please email: journals.permissions@oup.com.

We consider the development of prediction models in patients enrolled in the USA ($n = 23\,034$, 1565 deaths, and a small subset with $n = 259$, 20 deaths). We validate these models in patients enrolled outside of the USA ($n = 17\,796$, 1286 deaths). The first model only includes age as a continuous, linear term in a logistic regression analysis, while a slightly more complex model includes age, Killip class, systolic blood pressure, and heart rate. Programming code for the analyses is available at www.clinicalpredictionmodels.org with the R software (R Foundation for Statistical Computing, Vienna, Austria, www.r-project.org). The original GUSTO-I model was based on 40 830 patients and included 14 predictors.⁸

Seven steps to model development

We propose seven logically distinct steps in the development of prediction models with regression analysis. These steps are addressed below, with more detail provided elsewhere.¹⁶ A glossary is provided which summarizes definitions and characteristics of terms relevant to prediction model development and validation (Appendix).

Step 1: Problem definition and data inspection

An important preliminary step is to carefully consider the prediction problem. Questions to be addressed include:

- (i) What is the precise research question? In prediction research, we often are both interested in insight in what factors predict the endpoint, and in the pragmatic issue of providing estimates of the absolute risk for the endpoint, based on a combination of factors.²² Indeed, in the development of the GUSTO-I model, the title mentions 'Predictors of ...', suggesting a focus on insight in which prognostic factors predict 30-day mortality. The analysis, however, also describes the development of a multivariable prediction model, where a combination of factors predicts absolute risk with a logistic regression formula.⁸ For insight in the importance of predictors, effects are usually expressed on a relative scale, e.g. as an odds ratio (OR). For example, older age is associated with higher 30-day mortality, reflected in an OR of 2.3 per 10 years in GUSTO-I, with 95% confidence interval 2.2–2.4. In contrast, risk predictions are expressed as probabilities on an absolute scale between 0 and 100%. For example, predicted risks for 30-day mortality are 2 and 20% for 50 and 80-year-old patients, respectively. Uncertainty can be indicated with 95% confidence intervals. This is relevant for relative effects, but may confuse rather than help patients and clinicians when provided with absolute risk predictions.²³
- (ii) What is already known about the predictors? Literature review and close interactions between statisticians and clinical researchers are important to incorporate subject matter knowledge in the modelling process. The full GUSTO-I data set was of exceptional size, such that many modelling decision could reliably be guided by findings in the data. With smaller data sets, drawn from GUSTO-I, simulations showed that developed prediction models are unstable and produce too extreme predictions.¹¹
- (iii) How were patients selected? Patient data used for model development are commonly collected for a different purpose. The

GUSTO-I trial was designed to study the therapeutic effects of streptokinase and tissue plasminogen activator. The inclusion criteria for the trial were relatively broad, which implies that the GUSTO-I patients may be reasonably representative for the population of patients with an acute MI. We recognize that representativeness is usually a concern when RCT data are used for prognostic research, but this may be outbalanced by the superior quality of the data. For prediction of risk in current medical practice, the GUSTO-I data are likely outdated, since accrual in the trial was over 20 years ago.

- (iv) Another issue is how we should deal with any treatment effects in prognostic analyses. The treatment effect may be of specific interest, and adjustment for baseline prognostic factors has several advantages in the estimation of a treatment effect that is applicable to individual patients.^{24,25} If the focus is on absolute risk prediction, treatment effects have often been ignored, also since they are usually relatively small compared with the effects of prognostic factors. In our analyses of the GUSTO-I trial, we study prognostic effects without consideration of the treatment.

For diagnostic prediction models, we note that these should be developed in subjects suspected of having a particular condition.³ The selection should mimic the clinical setting, e.g. we may consider patients with chest pain suspected of obstructive coronary artery disease for a diagnostic prediction model that estimates the presence of obstructive disease.²⁶

- (v) Were the predictors reliably and completely measured? Many data sets are incomplete with respect to the values for some potential predictors. By default, patients with any missing value are excluded from statistical analyses (complete case analysis or available case analysis). This is inefficient since available information of other predictors is lost. To solve this problem, we may fill in the best estimates for the missing values, exploiting the correlation between variables in the data set (both predictor, endpoint, and auxiliary variables such as calendar time and place).¹⁸ Various statistical approaches are available to perform such imputation of missing values. Multiple imputation is a procedure to fill in missing values multiple times (typically at least five times) to appropriately address the randomness of the estimation procedure. We recognize that imputation should be performed carefully, but is usually preferable to a complete case analysis.¹⁸

In GUSTO-I, the collection of baseline data was prospective, since it was an RCT, and the definition of all predictors was carefully documented in the trial protocol. This makes us confident regarding the quality of the data. It also limited the occurrence of missing values, which were filled in with a single imputation procedure considering correlations between predictor variables.⁸

- (vi) Is the endpoint of interest? Hard endpoints are usually preferred, such as 30-day mortality in GUSTO-I, which was also the primary endpoint in the trial. Another important issue is the frequency of the endpoint, which determines the effective sample size (rather than the total number of patients). The GUSTO-I data set had 2851 deaths, which allows for reliable prediction modelling, where a common rule of thumb is to require at least 10 events per variable (EPV).^{2,12}

Step 2: Coding of predictors

Categorical and continuous predictor variables can be coded in different ways. Categories with infrequent occurrence may for instance be collapsed with others. In GUSTO-I, location of the infarction might well be coded as anterior vs. other, rather than as anterior, inferior, other, since a location other than anterior or inferior was infrequent (3% of the patients), and did not lead to a better model fit ($P = 0.30$, likelihood ratio test for improving the logistic regression model).

Continuous predictors can often be modelled as a linear association in a regression model, at least as a starting point.¹⁷ The interpretation of the relative effect of a predictor requires attention for the scale of measurement. For example, the importance of age is easily interpreted when scaled per decade (OR: 2.3, more than a doubling in odds per decade) than per year (OR: 1.09, 9% higher odds per year older). Linear terms are obviously not appropriate for predictors with non-linear associations, such as a J-shape or U-shape. Such shapes can efficiently be modelled using restricted cubic splines² or fractional polynomials.¹⁷ These functions provide flexibility but use only few extra regression coefficients. We emphasize that continuous predictors should not be dichotomized (categorization as below vs. above a certain cut-off) in the model development phase, since valuable information is lost.²⁷ In a later phase, we may search for a user-friendly format of the prediction model and categorize some predictors (e.g. in three, four, or five categories), if the loss of predictive information is limited.²⁸

Step 3: Model specification

Various strategies may be followed to choose predictors for inclusion in the prediction model. Stepwise selection methods are widely used to reduce a set of candidate predictors, but have many disadvantages. In particular, when the numbers of events are low, the selection is instable, the estimated regression coefficients are too extreme, and the performance of the selected model is overestimated.^{2,10,16,29} In the small sample of 259 patients, age and systolic blood pressure were statistically significant predictors, but Killip class ($P = 0.31$) and heart rate ($P = 0.92$) were not. In the total GUSTO-I data set, the sample size was large enough to rely on statistical testing for identification of predictors. All 14 predictors selected for the GUSTO-I model had P -values < 0.01 .⁸

A related issue is how researchers should deal with assumptions in regression models, such as that the effect of one predictor is the same irrespective of the value of other predictors. This additivity assumption can be relaxed by including statistical interaction terms.² In the full GUSTO-I data set, one such term was included in the prediction model (age * Killip class).⁸ It appeared that the prognostic effect of higher Killip class was less strong for older than for younger patients. For time-to-event data, the Cox model assumes proportionality of effects, which can be tested with interaction terms including time.²

We note that iterative cycles of testing of the importance of predictors, assumptions in the model, and adaptation may lead to a model that fits the data well. But such a model may provide predictions that do not generalize to new subjects outside the data under study ('overfitting').^{2,16} A simple, robust model may not fit the data perfectly, but should be preferred to an overly fine-tuned model for the specific data under study. Similarly, predictor selection

should usually consider clinical knowledge and previous studies rather than solely rely on statistical selection methods.^{2,16}

Step 4: Model estimation

Once a model is specified, regression coefficients need to be estimated. For logistic and Cox regression models, we commonly estimate coefficients with maximum likelihood (ML) methods. Some modern techniques have been developed that aim to limit overfitting of a model to the available data, such as statistical shrinkage techniques,³⁰ penalized ML estimation,³¹ and the least absolute shrinkage and selection operator (LASSO).^{11,32} For the model estimated in 259 patients, we found a shrinkage factor of 0.82. The regression coefficients should be multiplied by that value to provide more reliable predictions for new patients.

Step 5: Model performance

For a proposed model, researchers need to determine the quality. Several performance measures are commonly used, including measures for model calibration and discrimination. We discuss these and measures for clinical usefulness with model validation.

Step 6: Model validity

It is important to separate internal and external validity. Internal validity of a prediction model refers to the validity of claims for the underlying population that the data originated from ('reproducibility').¹⁹ Internal validation may especially address the stability of the selection of predictors, and the quality of predictions. Using a random sample for model development, and the remaining patients for validation ('split sample validation') is a common, but suboptimal form of internal validation.¹³ Better methods are cross-validation and bootstrap resampling, where samples are drawn with replacement from the development sample.² In the small sample of 259 patients, bootstrapping indicated that the discriminative ability was expected to decrease from 0.82 to 0.78 in new patients. Such internal validation should always be attempted when developing a prediction model (Tables 1 and 2).

External validity refers to generalizability of claims to 'plausibly related' populations.¹⁹ External validation is commonly considered a stronger test for prediction models than internal validation, since it addresses transportability rather than reproducibility. External validity may be evaluated by studying patients who were more recently treated (temporal validation), from other hospitals (geographic validation), or treated in fully different settings (strong external validation).^{19,33}

Step 7: Model presentation

As a final step we propose to consider is the presentation of a prediction model, such that it best addresses the clinical needs. Regression formulas can be used, such as the formula published with the GUSTO-I model.⁸ Many paper-based alternatives are available for easy applicability of a model, including score charts and nomograms.² A recent trend is to present prediction models as web-based calculators, or as apps for mobile phones and tablets. In the future, prediction models may well be embedded in decision aids and in electronic patient records to support clinical decision-making.

Table 1 Illustration of seven steps for developing a prediction model with the GUSTO-I data ($n = 40\,830$)⁸

Step	Specific issues	GUSTO-I model
1. Problem definition and data inspection	Aim: predictors or predictions? Selection, predictor definitions, and completeness, endpoint definition	Aim is both insights in which predictors are important and to provide individualized risk predictions Prospective data collection in a randomized trial with a hard endpoint (30-day mortality). Missing predictor values were imputed
2. Coding of predictors	Continuous predictors Combining categorical predictors	Extensive checks of transformations for continuous predictors Categories kept separate, e.g. for location of infarction
3. Model specification	Selection of main effects? Assessment of assumptions?	Stepwise selection; appropriate because of very large sample size Additivity checked with interaction terms; one included (age $\times$ Killip class)
4. Model estimation	Shrinkage included?	Not necessary
5. Model performance	Appropriate measures used?	Calibration and discrimination, but no indicators of clinical usefulness
6. Model validation	Internal validation including model specification and estimation?	Bootstrap and 10-fold cross-validation
7. Model presentation	Format appropriate for audience	No; complex formula in appendix

Table 2 Overview of four measures (ABCD) for model performance

Aspect	Measure	Visualization	Characteristics
Calibration	A: alpha Calibration-in-the-large	Calibration plot	Intercept in plot; compares mean observed with mean predicted
	B: beta Calibration slope		Regression slope in plot; related to shrinkage of regression coefficients
Discrimination	C-statistic	ROC curve	Interpretation as the probability of correct classification for a pair of subjects with and without the endpoint
Clinical usefulness	Decision-curve analysis NB	Decision curve	Net true-positive classification rate by using a model over a range of thresholds

An ABCD for model validation

Whatever the method used to develop a model, one could argue that validity is all that matters.⁷ We propose four key measures in the assessment of the validation of prediction models, related to calibration, discrimination, and clinical usefulness.^{20,34–36} The measures are illustrated by studying the external validity of the models developed in 259 or 23 034 patients enrolled in GUSTO-I in the USA and tested in 17 796 GUSTO-I patients from outside the US.

Alpha: calibration-in-the-large

Calibration refers to the agreement between observed endpoints and predictions.²⁰ For example, if we predict a 5% risk that a patient will die within 30 days, the observed proportion should be ~ 5 deaths per 100 with such a prediction. Calibration can be well assessed graphically, in a plot with predictions on the x-axis and the observed endpoint on the y-axis (Figure 1). The observed values on the y-axis are 0 or 1 (dead/alive), while the predictions on the x-axis range between 0 and 100%. To visualize the agreement between the observed and predicted values, smoothing techniques can well be used to depict the association.³⁷ We can also plot the observed proportions of death for groups of patients with similar predicted risk, for instance, by deciles of predictions. Considering such groups with their deviations from the ideal line makes the plot a graphical illustration of the often used Hosmer–Lemeshow goodness-of-fit test. We do not recommend

this test for assessment of calibration. It does not indicate the direction of any miscalibration and only provides a P -value for differences between observed and predicted endpoints per group of patients (commonly deciles).²⁰ Such grouping is arbitrary and imprecise, and P -values depend on the combination of the extent of miscalibration and sample size. Rather, we emphasize the older recalibration idea as proposed by Cox in 1958.³⁸ Perfect predictions should be on the ideal line, described with an intercept alpha ('A') of 0 and slope beta ('B') of 1. The log odds of predictions are used as the predictor of the 0/1 outcome, or the log (hazard) for time-to-event outcomes.^{38,39} Imperfect calibration can be characterized by deviations from these ideal values.

The intercept A relates to calibration-in-the-large, which compares the mean of all predicted risks with the mean observed risk. The parameter hence indicates the extent that predictions are systematically too low or too high. At model development, observed incidence and mean predicted risk are equal for regression models, and assessment of calibration-in-the-large makes no sense. At external validation, calibration-in-the-large problems have often been found, for example, for the Framingham model in multiple ethnic groups,⁴⁰ or the Framingham variant developed by NICE for the UK.³³ If we test the prediction model developed in 23 034 US patients outside of the USA, we note a slightly higher mortality ($A = 0.07$, $P = 0.38$; equivalent to an OR of non-US vs. US of $\exp(0.07) = 1.07$, right panel of Figure 1). Note that no intercept is

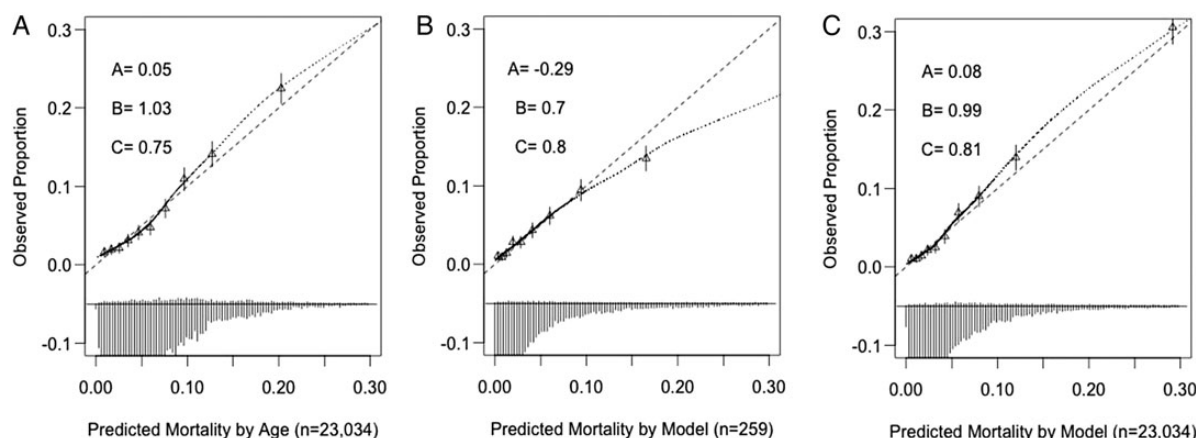


Figure 1 Validation plots for clinical prediction models applied in 17 796 patients enrolled in GUSTO-I outside the USA. The models contained the predictors age (left panel), or age plus Killip class, blood pressure, and heart rate (right panels, with $n = 259$ or $n = 23\,034$ US patients for model development). (A) calibration-in-the-large calculated as the logistic regression model intercept given that the calibration slope equals 1; (B) calibration slope in a logistic regression model with the linear predictor as the sole predictor; (C) c-statistic indicating discriminative ability. Triangles represent deciles of subjects grouped by similar predicted risk. The distribution of subjects is indicated with spikes at the bottom of the graph, stratified by endpoint (deaths above the x-axis, survivors below the x-axis).

calculated if a Cox model is used at external validation.³⁶ We can compare the mean of the predicted risks for a certain time point, e.g. 5 or 10 years, with the estimate from a Kaplan–Meier curve. Other types of models, such as Weibull regression models, can also be used to assess calibration-in-the-large for survival models.⁴¹

Beta: calibration slope

The calibration slope B is often smaller than 1 if a model was developed in a relatively small data set. Such a finding reflects that predictions were too extreme: low prediction too low, and high predictions too high. For the model based on 259 patients, the slope was 0.70 among the 17 796 non-US patients, in line with what was expected at internal validation (shrinkage factor 0.82). For the model based on 23 034 patients, B was close to one at internal validation by cross-validation or bootstrapping,⁸ as well as at external validation in non-US patients (slope 0.99, $P = 0.55$ for comparison with slope = 1, Figure 1).

Concordance statistic: discrimination

Discrimination refers to the ability of the model to distinguish a patient with the endpoint (dead) from a patient without (alive). A better discriminating model has more spread between the predictions than a poorly discriminating model.³⁵ Indeed, a model that predicts for all subjects the same predicted risk equal to the incidence shows perfect calibration, but is useless since it does not discriminate between patients. A validation plot showing a wide spread between predictions as a histogram, or between deciles of predicted risk, indicates a good discriminating model (Figure 1). Discriminative ability is commonly quantified with a concordance (c) statistic. For a binary endpoint, c is identical to the area under the receiver operating characteristic (ROC) curve, which plots the sensitivity (true-positive rate) against $1 - \text{specificity}$ (false-positive rate) for consecutive cut-offs for the predicted risk. For a time-to-event endpoint, such as survival, the calculation of c may be affected by the amount of

incomplete follow-up (censoring). We note that the c statistic is insensitive to systematic errors in calibration, and considers the rather artificial situation of classification in a pair of patients with and without the endpoint.

In GUSTO-I, the c-statistic indicates the probability that among two patients, one dying before 30 days, and one surviving, the patient bound to die will have a higher predicted risk than the surviving patient. The more complex prediction models had c-statistics over 0.8 [0.813 (0.802–0.824) and 0.812 (0.800–0.824) at internal and external validation, respectively]. This performance was much better than a model which only included age (0.75 at external validation, left panel in Figure 1). With development in only 259 patients, the apparent c-statistic was 0.82, but 0.78 at internal validation, and 0.80 (0.79–0.82) at external validation. This illustrates that the availability of a smaller data set decreases model performance at external validation.

Decision-curve analysis

Calibration and discrimination are important aspects of a prediction model, and consider the full range of predicted risks. However, these aspects do not assess clinical usefulness, i.e. the ability to make better decisions with a model than without.²⁰ If a prediction model aims to guide treatment decisions, a cut-off is required to classify patients as either low risk (no treatment) or high risk (treatment is indicated). The cut-off is a decision threshold. At the threshold, the likelihood of benefit, e.g. reduced mortality as a result of thrombolytic therapy, exactly balances the likelihood of harm, e.g. bleeding risk and financial costs. A threshold value of, for example, 2% indicates that death of a non-treated patient is $98 : 2 = 49$ times worse than the complications of a bleeding incident and costs of an unnecessarily treated patient. It is usually difficult to define a threshold since empirical evidence for the relative weight of benefits and harms is often lacking. Furthermore, some patients may be prepared to take higher risk for a possible benefit than others. It is therefore advised to consider a range of thresholds when quantifying the clinical usefulness of a prediction model.²¹

Once a threshold has been applied to classify patients as low vs. high risk, sensitivity and specificity are often used as measures for usefulness. Finding an optimal balance between these is again possible with consideration of harms and benefits of treatment, in combination with the incidence of the endpoint. The sum of sensitivity and specificity can only be used as a naive summary indicator of usefulness, since such a sum ignores the relative weight of true positives (considered in sensitivity) and false positives (considered in $1 - \text{specificity}$).⁴²

Recently proposed and more appropriate summary measures include the net benefit (NB).²¹ This measure is consistent with the use of an optimal decision threshold to classify patients.⁴³ The relative weight of harms and benefits is used to define the threshold, and is used to calculate a weighted sum of true- minus false-positive classifications.²¹

For GUSTO-I, we first note that the 7% overall 30-day mortality implies a maximum NB of 7%. This is obtained if we use a threshold of 0%. All patients are then candidates for tPA treatment since we assume no harm of treatment. If we appreciate that treatment involves some harm, the optimal decision threshold is $>0\%$. We find that treatment decision-making on the basis of risk predictions from a model would give a slightly higher NB than treating everyone. For example at a threshold of 2%, we have 1225 true-positive classifications (candidates for treatment, and died), but also 11 192 false-positive classifications (candidates for treatment, but survived) among the 17 796 non-US patients. The NB is calculated as $(1225 - 2/98 \times 11\,192) / 17\,796 = 5.6\%$ (Figure 2). This is only slightly better than treating all, where $\text{NB} = (1286 - 2/98 \times 16\,510) / 17\,796 = 5.3\%$, since 1286 died and 16 510 survived overall. The difference is 0.3%. This implies that for every 1000 patients where we apply the prediction model, three extra true positives are identified without increasing the false-positive rate. The NB was higher for a higher threshold, such as 5% (19 per 1000 extra net true positives). Based on age only, the NB was virtually absent for a 2% decision threshold (0.04%, Figure 2), and in-between for a model based on only 259 patients (0.2%, Figure 2). Figure 2 illustrates that using more prognostic information than age increases the clinical usefulness of a model, and that a larger sample size leads to a better performing model. Documentation and software for decision-curve analysis is publicly available (www.decisioncurveanalysis.org), considering both binary and time-to-event endpoints.

Concluding remarks

We discussed seven steps to reliably develop a prediction model, and four measures that are important at model validation. There are many details to consider for each development and validation step, which are discussed in the methodological literature. Involvement of statistical experts is usually required to well develop or validate a prediction model. We provided an illustration on predicting 30-day mortality in patients with an acute myocardial infarction. The exceptional size of the US part of the GUSTO-I trial implied that overfitting was not relevant. Overfitting is a key problem in many prediction models. This is either because the number of events is small, as illustrated with the substudy in 259 patients, or because many potential predictors are studied, such as in genetic or other 'omics' analyses.

At model validation, calibration, discrimination, and clinical usefulness should be considered. A validation graph such as Figure 1

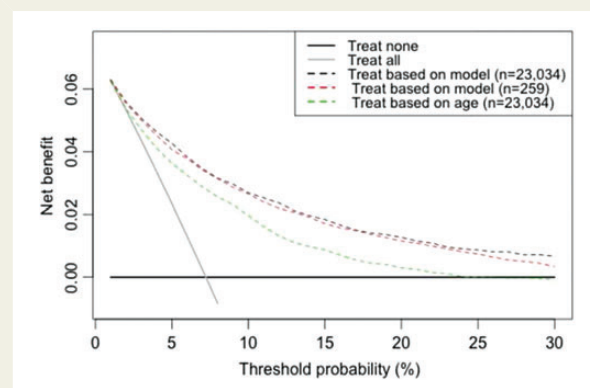


Figure 2 Decision curves for the prediction models applied in 17 796 patients enrolled in GUSTO-I outside the USA. Solid line: Assume no patients are treated, net benefit is zero (no true-positive and no false-positive classifications); Grey line: assume all patients are treated; Dotted lines: patients are treated if predictions exceed a threshold, with 30-day mortality risk predictions based on age only, or a prediction model with age, Killip class, blood pressure, and heart rate, developed in $n = 259$ or $n = 23\,034$ US patients. The graph gives the expected net benefit per patient relative to no treatment in any patient ('Treat none'). The threshold defines the weight w for false-positive (treat while patient survived) vs. true-positive (treat a patient who died) classifications. For example, a threshold of 2% implies that FP classifications are valued at $2/98$ of true-positive classifications, and w is $0.02 / (1 - 0.02) = 0.0204$. The clinical usefulness of a prediction model can then be summarized as: $\text{NB} = (\text{TP} - w \text{FP}) / N$, where N is the total number of patients.²¹

(or a variant with a 'calibration belt'⁴⁴) is an important summary tool. We furthermore recognize that a key question is nowadays how we can quantify an increase in predictive ability by new markers. For markers, we may consider changes in the A, B, C, and D measures related to calibration, discrimination, and clinical usefulness. Miscalibration of predictions is common at external validation (A different from 0, $B < 1$), but we would expect this to be similar when adding a marker to a model. Some performance measures, such as the net reclassification improvement (NRI) require well-calibrated predictions to be meaningful.^{45,46} Discrimination (C) only shows modest increases for most currently available markers in cardiology.⁴⁷ Some researchers have blamed the c -statistic as being insensitive. One might argue that we should merely accept that markers with statistically significant effects, but with a modest relative effect size, will not impact tremendously on identifying those with or without the endpoint.⁴⁸ Cost-effectiveness analyses should eventually guide us on the question whether an increase in performance is important enough to measure an additional marker in clinical practice.⁴⁹ Measures for clinical usefulness such as the NB (which may be shown in decision curves, D) are easy to calculate, increasingly used, and give a first impression of effectiveness in terms of potentially better patient outcomes.⁴³ The NRI is very similar in behaviour to measures for discrimination.⁵⁰ We hence did not discuss this quite popular measure in detail here. A fierce debate is ongoing on the relevance of the NRI in the evaluation of the predictive value of markers.^{42,46,51}

We see a role for our proposed framework to support methodological researchers who develop and validate prediction models. More importantly, clinical researchers may use the framework to systematically and critically assess a publication where a prediction model is developed or validated. We anticipate that following the framework, admittedly with room for refinements, will strengthen the methodological rigour and quality of prediction models in cardiovascular research.

Funding

E.W.S. was supported by a U award (1U01-NS086294-01, value of personalized risk information). Y.V. was supported by a VIDI grant (91711383, prognostic research with clustered data).

Conflict of interest: none declared.

Appendix

Glossary of terms relevant to prediction model development and validation

Term	Definition	Characteristics
Continuous predictors		
Restricted cubic splines	Smooth functions with linear tails	One to three extra regression coefficients are sufficient to adequately model most non-linear shapes ²
Fractional polynomials	A combination of 1 or 2 polynomial transformations such as $x^2 + x^{-1}$, with powers from the set $(-3, -2, -1, -0.5, 0, 0.5, 1, 2, 3)$	An automatic search can find transformations to adequately model non-linear shapes, including reciprocal, logarithm, square root, square, and cubic transformations ¹⁷
Missing values		
Complete case analysis	Analysis that considers patients with complete information on all predictors and the endpoint available	Performed by default in statistical packages, but inefficient ¹⁸
Multiple imputation	Fill in missing values multiple times to allow for analyses of completed data sets	Statistically efficient, but relies on assumptions of the missing value generating mechanism, and a correct imputation model ¹⁸
Reliable estimation		
EPV	The number of patients with the event of interest divided by the number of predictor variables considered	Indicates effective sample size, which is lower than the total number of patients. The count of predictor variables should include all regression coefficients considered for all candidate predictors (not only the predictors that are finally selected for a prediction model) ²
Stepwise selection	Procedure to eliminate candidate predictors that are not relevant to making predictions	Reduces the set of predictors. Many variants of stepwise selection are possible, all with disadvantages especially in small data sets ¹⁶
Shrinkage	Reduce regression coefficients towards zero, such that less extreme predictions are made	Improves predictions from models, especially in small data sets. Specific methods include penalized estimation, such as the LASSO ¹¹
Validation		
Internal validation	Assesses the validity of claims for the underlying population where the data originated from ('reproducibility')	Common methods are cross-validation and bootstrap resampling ²
External validation	Assesses the validity of claims for 'plausibly related' populations ('generalizability, or 'transportability')	Study patients who were more recently treated (temporal validation), from other hospitals (geographic validation), or treated in fully different settings (strong external validation) ¹⁹
Evaluation of predictions		
Calibration	Agreement between observed and predicted risks	Calibration is usually near perfect at model development, and especially of interest at external validation ¹⁶
Discrimination	Ability to distinguish a patient with the endpoint from a patient without	Discrimination is a key aspect of model performance, but the concordance statistic ('C') refers to the artificial situation of a considering a pair of patients, one with and one without the endpoint ²⁰
Evaluation of classifications		
Decision threshold	Cut-off for a predicted risk to define a low- and a high-risk group	The optimal threshold is defined by the balance between the harm of a false-positive classification and the benefit of a true-positive classification ²¹
Sensitivity (specificity)	Probability of true-positive (true-negative) classification among those with (without) the endpoint	Traditional measures to quantify classification performance, conditional on knowing the endpoint ²⁰
NB	A weighted sum of true-positive and false-positive classifications	Novel measure to quantify classification performance, taking a decision-analytic perspective ²¹

References

1. Van de Werf F, Bax J, Betriu A, Blomstrom-Lundqvist C, Crea F, Falk V, Filippatos G, Fox K, Huber K, Kastrati A, Rosengren A, Steg PG, Tubaro M, Verheugt F, Weidinger F, Weis M. Guidelines ESCCfP. Management of acute myocardial infarction in patients presenting with persistent ST-segment elevation: the Task Force on the Management of ST-Segment Elevation Acute Myocardial Infarction of the European Society of Cardiology. *Eur Heart J* 2008;**29**:2909–2945.
2. Harrell FE. *Regression Modeling Strategies: With Applications to Linear Models, Logistic Regression, and Survival Analysis*. New York: Springer, 2001.
3. Moons KG, Royston P, Vergouwe Y, Grobbee DE, Altman DG. Prognosis and prognostic research: what, why, and how? *BMJ* 2009;**338**:b375.
4. Mushkudiani NA, Hukkelhoven CW, Hernandez AV, Murray GD, Choi SC, Maas AI, Steyerberg EW. A systematic review finds methodological improvements necessary for prognostic models in determining traumatic brain injury outcomes. *J Clin Epidemiol* 2008;**61**:331–343.
5. Mallett S, Royston P, Dutton S, Waters R, Altman DG. Reporting methods in studies developing prognostic models in cancer: a review. *BMC Med* 2010;**8**:20.
6. Bouwmeester W, Zuithoff NP, Mallett S, Geerlings MI, Vergouwe Y, Steyerberg EW, Altman DG, Moons KG. Reporting and methods in clinical prediction research: a systematic review. *PLoS Med* 2012;**9**:1–12.
7. Collins GS, de Groot JA, Dutton S, Omar O, Shanyinde M, Tajar A, Voysey M, Wharton R, Yu LM, Moons KG, Altman DG. External validation of multivariable prediction models: a systematic review of methodological conduct and reporting. *BMC Med Res Methodol* 2014;**14**:40.
8. Lee KL, Woodlief LH, Topol EJ, Weaver WD, Betriu A, Col J, Simoons M, Aylward P, Van de Werf F, Califf RM. Predictors of 30-day mortality in the era of reperfusion for acute myocardial infarction. Results from an international trial of 41,021 patients. GUSTO-I Investigators. *Circulation* 1995;**91**:1659–1668.
9. Ennis M, Hinton G, Naylor D, Revow M, Tibshirani R. A comparison of statistical learning methods on the GUSTO database. *Stat Med* 1998;**17**:2501–2508.
10. Steyerberg EW, Eijkemans MJ, Habbema JD. Stepwise selection in small data sets: a simulation study of bias in logistic regression analysis. *J Clin Epidemiol* 1999;**52**:935–942.
11. Steyerberg EW, Eijkemans MJ, Harrell FE Jr, Habbema JD. Prognostic modelling with logistic regression analysis: a comparison of selection and estimation methods in small data sets. *Stat Med* 2000;**19**:1059–1079.
12. Steyerberg EW, Eijkemans MJ, Harrell FE Jr, Habbema JD. Prognostic modeling with logistic regression analysis: in search of a sensible strategy in small data sets. *Med Decis Making* 2001;**21**:45–56.
13. Steyerberg EW, Harrell FE Jr, Borsboom GJ, Eijkemans MJ, Vergouwe Y, Habbema JD. Internal validation of predictive models: efficiency of some procedures for logistic regression analysis. *J Clin Epidemiol* 2001;**54**:774–781.
14. Steyerberg EW, Borsboom GJ, van Houwelingen HC, Eijkemans MJ, Habbema JD. Validation and updating of predictive logistic regression models: a study on sample size and shrinkage. *Stat Med* 2004;**23**:2567–2586.
15. Steyerberg EW, Eijkemans MJ, Boersma E, Habbema JD. Equally valid models gave divergent predictions for mortality in acute myocardial infarction patients in a comparison of logistic regression models. *J Clin Epidemiol* 2005;**58**:383–390.
16. Steyerberg EW. *Clinical Prediction Models: A Practical Approach to Development, Validation, and Updating*. New York: Springer, 2009.
17. Royston P, Sauerbrei W. *Multivariable Model-building: A Pragmatic Approach to Regression Analysis Based on Fractional Polynomials for Modelling Continuous Variables*. Chichester, England, Hoboken, NJ: John Wiley, 2008.
18. Altman DG, Bland JM. Missing data. *BMJ* 2007;**334**:424.
19. Justice AC, Covinsky KE, Berlin JA. Assessing the generalizability of prognostic information. *Ann Intern Med* 1999;**130**:515–524.
20. Steyerberg EW, Vickers AJ, Cook NR, Gerds T, Gonen M, Obuchowski N, Pencina MJ, Kattan MW. Assessing the performance of prediction models: a framework for traditional and novel measures. *Epidemiology* 2010;**21**:128–138.
21. Vickers AJ, Elkin EB. Decision curve analysis: a novel method for evaluating prediction models. *Med Decis Making* 2006;**26**:565–574.
22. Shmueli G. To explain or to predict? *Statist Sci* 2010;**25**:289–310.
23. Kattan MW. Doc, what are my chances? A conversation about prognostic uncertainty. *Eur Urol* 2011;**59**:224.
24. Hauck WW, Anderson S, Marcus SM. Should we adjust for covariates in nonlinear regression analyses of randomized trials? *Control Clin Trials* 1998;**19**:249–256.
25. Steyerberg EW, Bossuyt PM, Lee KL. Clinical trials in acute myocardial infarction: should we adjust for baseline characteristics? *Am Heart J* 2000;**139**:745–751.
26. Genders TS, Steyerberg EW, Alkadhi H, Leschka S, Desbiolles L, Nieman K, Galema TW, Meijboom WB, Mollet NR, de Feyter PJ, Cademartini F, Maffei E, Dewey M, Zimmermann E, Laule M, Pugliese F, Barbaggio R, Sinitsyn V, Bogaert J, Goetschalckx K, Schoepf UJ, Rowe GW, Schuijff JD, Bax JJ, de Graaf FR, Knuti J, Kajander S, van Mieghem CA, Meijjs MF, Cramer MJ, Gopalan D, Feuchtnner G, Friedrich G, Krestin GP, Hunink MG, Consortium CAD. A clinical prediction rule for the diagnosis of coronary artery disease: validation, updating, and extension. *Eur Heart J* 2011;**32**:1316–1330.
27. Royston P, Altman DG, Sauerbrei W. Dichotomizing continuous predictors in multiple regression: a bad idea. *Stat Med* 2006;**25**:127–141.
28. Boersma E, Poldermans D, Bax JJ, Steyerberg EW, Thomson IR, Banga JD, van De Ven LL, van Urk H, Roelandt JR, Group DS. Predictors of cardiac events after major vascular surgery: role of clinical characteristics, dobutamine echocardiography, and beta-blocker therapy. *JAMA* 2001;**285**:1865–1873.
29. Derksen S, Keselman H. Backward, forward and stepwise automated subset selection algorithms: frequency of obtaining authentic and noise variables. *Br J Math Stat Psychol* 1992;**45**:265–282.
30. van Houwelingen JC, Le Cessie S. Predictive value of statistical models. *Stat Med* 1990;**9**:1303–1325.
31. Moons KG, Donders AR, Steyerberg EW, Harrell FE. Penalized maximum likelihood estimation to directly adjust diagnostic and prognostic prediction models for over-optimism: a clinical example. *J Clin Epidemiol* 2004;**57**:1262–1270.
32. Tibshirani R. Regression and shrinkage via the Lasso. *J R Stat Soc Ser B* 1996;**58**:267–288.
33. Collins GS, Altman DG. Predicting the 10 year risk of cardiovascular disease in the United Kingdom: independent and external validation of an updated version of QRISK2. *BMJ* 2012;**344**:e4181.
34. Vergouwe Y, Steyerberg EW, Eijkemans MJ, Habbema JD. Validity of prognostic models: when is a model clinically useful? *Semin Urol Oncol* 2002;**20**:96–107.
35. Royston P, Altman DG. Visualizing and assessing discrimination in the logistic regression model. *Stat Med* 2010;**29**:2508–2520.
36. Royston P, Altman DG. External validation of a Cox prognostic model: principles and methods. *BMC Med Res Methodol* 2013;**13**:33.
37. Austin PC, Steyerberg EW. Graphical assessment of internal and external calibration of logistic regression models by using loess smoothers. *Stat Med* 2014;**33**:517–535.
38. Cox DR. Two further applications of a model for binary regression. *Biometrika* 1958;**45**:562–565.
39. Miller ME, Langefeld CD, Tierney WM, Hui SL, McDonald CJ. Validation of probabilistic predictions. *Med Decis Making* 1993;**13**:49–58.
40. D'Agostino RB Sr, Grundy S, Sullivan LM, Wilson P. Validation of the Framingham coronary heart disease prediction scores: results of a multiple ethnic groups investigation. *JAMA* 2001;**286**:180–187.
41. van Houwelingen HC, Thorogood J. Construction, validation and updating of a prognostic model for kidney graft survival. *Stat Med* 1995;**14**:1999–2008.
42. Greenland S. The need for reorientation toward cost-effective prediction. *Stat Med* 2008;**27**:199–206.
43. Localio AR, Goodman S. Beyond the usual prediction accuracy metrics: reporting results for clinical decision making. *Ann Intern Med* 2012;**157**:294–295.
44. Nattino G, Finazzi S, Bertolini G. A new calibration test and a reappraisal of the calibration belt for the assessment of prediction models based on dichotomous outcomes. *Stat Med* 2014; in press. PMID 24497413.
45. Pencina MJ, D'Agostino RB Sr, Steyerberg EW. Extensions of net reclassification improvement calculations to measure usefulness of new biomarkers. *Stat Med* 2011;**30**:11–21.
46. Leening MJ, Vedder MM, Witteman JC, Pencina MJ, Steyerberg EW. Net reclassification improvement: computation, interpretation, and controversies: a literature review and clinician's guide. *Ann Intern Med* 2014;**160**:122–131.
47. Tzoulaki I, Liberopoulos G, Ioannidis JP. Assessment of claims of improved prediction beyond the Framingham risk score. *Jama* 2009;**302**:2345–2352.
48. Pepe MS, Janes H, Longton G, Leisenring W, Newcomb P. Limitations of the odds ratio in gauging the performance of a diagnostic, prognostic, or screening marker. *Am J Epidemiol* 2004;**159**:882–890.
49. Hlatky MA, Greenland P, Arnett DK, Ballantyne CM, Criqui MH, Elkind MS, Go AS, Harrell FE Jr, Hong Y, Howard BV, Howard VJ, Hsue PY, Kramer CM, McConnell JP, Normand SL, O'Donnell CJ, Smith SC Jr, Wilson PW. Criteria for evaluation of novel markers of cardiovascular risk: a scientific statement from the American Heart Association. *Circulation* 2009;**119**:2408–2416.
50. Van Calster B, Vickers AJ, Pencina MJ, Baker SG, Timmerman D, Steyerberg EW. Evaluation of markers and risk prediction models: overview of relationships between NRI and decision-analytic measures. *Med Decis Making* 2013;**33**:490–501.
51. Vickers AJ, Pepe M. Does the net reclassification improvement help us evaluate models and markers? *Ann Intern Med* 2014;**160**:136–137.