

Calibration of Machine Learning Models in glmnet

Walter K. Kremers, Mayo Clinic, Rochester MN

12 September 2026

The Package

The “An Overview of glmnet” vignette shows how to run the main package function `nested.glmnet()` and how to summarize model performances. If one identifies a well performing model according to the metrics in this summary, e.g. concordance, correlation, deviance ratio, linear calibration coefficients, one may want to do further evaluation in terms of calibration. The strongest calibration and validation will involve calibration with new, independent datasets. Frequently one will not have immediate access to such new data sets, or one may want first to do an internal validation before subjecting a model to an external validation. Here we consider an internal validation approach using cross validation or bootstrap re-sampling, similar to how we numerically assessed model performance.

An Example Analysis

To explore calibration we first consider the `nested.glmnet()` call from the “An Overview of glmnet” vignette which fit machine learning models to survival data with `family="cox"`, i.e.

```
set.seed(465783346)
nested.cox.fit = nested.glmnet(xs, NULL, yt, event, family="cox",
                              dolasso=1, dostep=1, steps_n=40, folds_n=10, track=1)
```

How we generated the data for this function call, i.e. the `xs` matrix as well as the `yt` and `event` vectors, is too described in the “An Overview of glmnet” vignette.

Linear Calibration

Using either `print()` or `summary()` on the output object `nested.cox.fit` one gets, amongst other information, summaries for the linear calibration slopes and intercepts as in

```
summary( nested.cox.fit )

## Sample information including number of records, events, number of columns in
## design (predictor, X) matrix, and df (rank) of design matrix:
##           family              n          nevent
##           cox              1000           698
##           xs.columns        xs.df      null.dev/nevent
##           100              94           12.43
## null.m2LogLik/nevent  sat.m2LogLik/nevent
```

```
##                12.43                0
##
## For LASSO, and Stepwise regression tuned by df and p, average (Ave) model
## performance measures from the 10-fold (NESTED) Cross Validation are given together
## with naive summaries calculated using all data without cross validation
##
##                Ave DevRat Ave Slope Ave Concordance Ave Non Zero
## lasso                0.2412    1.1202            0.8717        44.0
## lassoR               0.2425    1.0316            0.8718        17.5
## lassoR0              0.2400    0.9693            0.8728        16.1
## elastic net          0.2454    1.0698            0.8734        22.1
## ridge                0.2316    1.1232            0.8666        99.0
## elastic net G0        0.2422    0.9715            0.8740        18.1
## elastic net G1        0.2412    1.1202            0.8717        44.0
##
##                Naive DevRat Naive Concordance Non Zero
## lasso                0.1696            0.8794         45
## lassoR               0.1710            0.8791         20
## lassoR0              0.1663            0.8759         14
## elastic net          0.1710            0.8791         20
## ridge                0.1718            0.8822         99
## elastic net G0        0.1663            0.8759         14
## elastic net G1        0.1696            0.8794         45
##
##                Ave DevRat Ave Slope Ave Concordance Ave Non Zero
## Stepwise df tuned      0.2496    0.9941            0.8743        14.7
## Stepwise p tuned       0.2522    1.0010            0.8752        14.8
## Full regression        0.2251    0.0000            0.8650        93.0
##
##                Naive DevRat Naive Concordance Non Zero
## Stepwise df tuned      0.1700            0.8778         14
## Stepwise p tuned       0.1711            0.8785         15
## Full regression        0.1797            0.8813         93
```

Here we see that for many of the models the linear calibration slope term is near 1, the ideal for perfect calibration. For the Cox model any intercept term can be absorbed into the baseline survival function and there is no pertinent intercept term for calibration.

At first glance the lower deviance ratios for the naive models than for the average deviance ratio for the hold out models looks wrong. The deviance and deviance ratio for the Cox model cannot be calculated simply by adding the deviance from separate sample observations like for the standard gaussian or binomial models. Inspecting the deviance formula one sees that observations associated with events have a smaller impact on deviance when they occur at times when there are more sample elements at risk. In this way the deviance ratio tends to be smaller for larger samples from the same population and the smaller this difference between the average and naive deviance ratios is not surprising.

A First Visual

An initial calibration consideration was made in the overview vignette by regressing observed outcomes on the predicted from the final model based upon the relaxed lasso. This regression was made using splines, in particular the `pspline()` function from within a `coxph()` call, as in

```
# Get predicted from CV relaxed lasso model embedded in nested CV output object
xb.hat = predict( object=nested.cox.fit , xs_new=xs, lam=NULL, gam=NULL, comment=FALSE)
# Fit a spline to xb.hat using coxph, and plot
```

```
#library( survival )                ## load survival package for Cox model fits
fit1 = coxph(Surv(yt, event) ~ pspline(xb.hat, df=3))
```

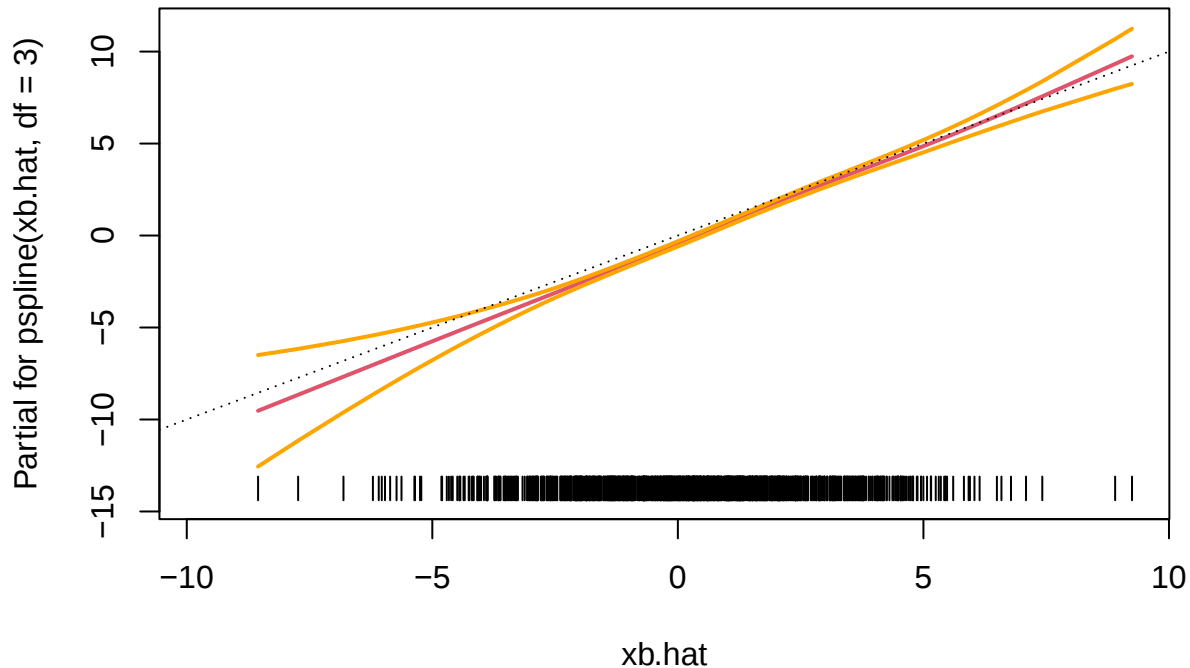
```
summary(fit1)
```

```
## Call:
## coxph(formula = Surv(yt, event) ~ pspline(xb.hat, df = 3))
##
##      n= 1000, number of events= 698
##
##              coef se(coef) se2      Chisq  DF    p
## pspline(xb.hat, df = 3),  1.073 0.03333  0.03333 1035.61 1.00 3.3e-227
## pspline(xb.hat, df = 3),              2.17 2.03  3.4e-01
##
##              exp(coef) exp(-coef) lower .95 upper .95
## ps(xb.hat)3  1.068e+01  9.361e-02 1.642e+00 6.948e+01
## ps(xb.hat)4  1.141e+02  8.764e-03 4.178e+00 3.116e+03
## ps(xb.hat)5  1.217e+03  8.215e-04 1.710e+01 8.665e+04
## ps(xb.hat)6  1.244e+04  8.042e-05 1.183e+02 1.307e+06
## ps(xb.hat)7  1.324e+05  7.554e-06 1.269e+03 1.381e+07
## ps(xb.hat)8  1.713e+06  5.837e-07 1.661e+04 1.767e+08
## ps(xb.hat)9  1.442e+07  6.933e-08 1.380e+05 1.507e+09
## ps(xb.hat)10 1.718e+08  5.819e-09 1.575e+06 1.875e+10
## ps(xb.hat)11 2.481e+09  4.031e-10 1.910e+07 3.221e+11
## ps(xb.hat)12 3.677e+10  2.719e-11 1.773e+08 7.627e+12
##
## Iterations: 4 outer, 14 Newton-Raphson
##      Theta= 0.7945991
## Degrees of freedom for terms= 3
## Concordance= 0.879 (se = 0.005 )
## Likelihood ratio test= 1492 on 3.03 df,  p=<2e-16
```

followed by plotting with

```
termplot(fit1,term=1,se=TRUE, rug=TRUE, lwd.term=2, lwd.se=2, lty.se=1)
abline(a=0,b=1,lty=3)
title ("Naive calibration curve for relaxed lasso model")
```

Naive calibration curve for relaxed lasso model



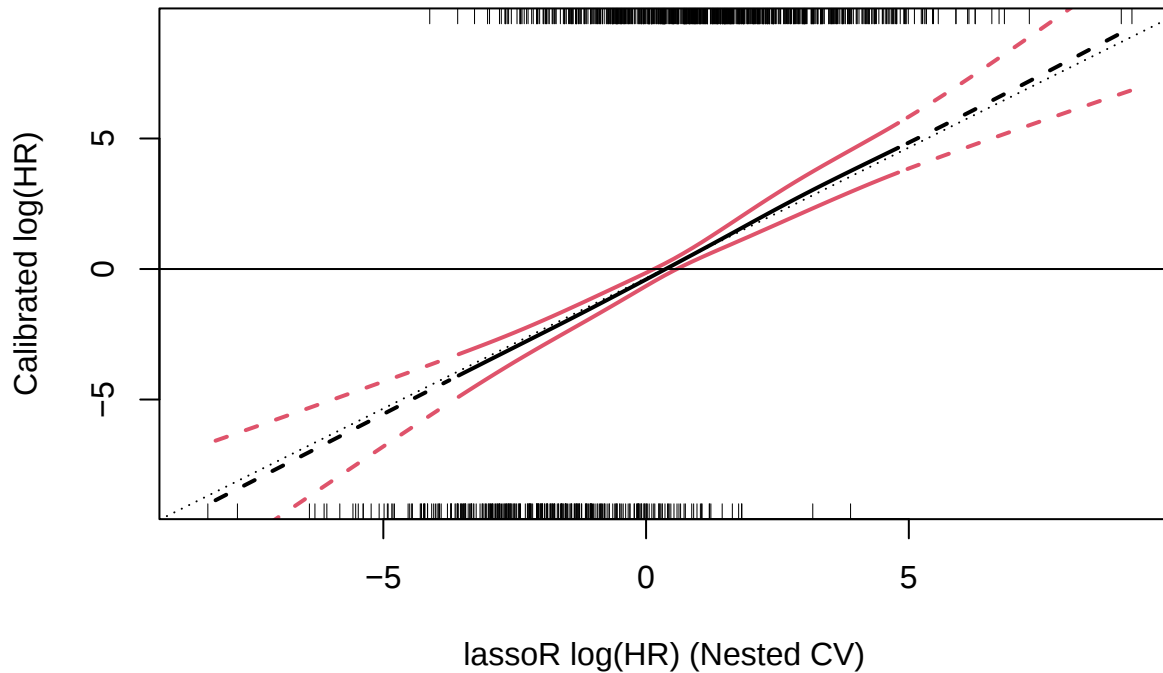
The spline fits may help to understand potential nonlinearities in the model. Here we see, a calibration line which is not far from linear. Still, as noted in the “An Overview of glmnet” vignette, because the same data are used for model evaluation as well as model derivation, it is hard to put much confidence in such a calibration plot which may suggest a better fit than can be expected for new data.

Calibration Using Spline Fits and Resampling

For each of the models fit, `nested.glmnet()` saves the $X \cdot \beta$'s from the model fit using all data. The `nested.glmnet()` function also calculates the $X \cdot \beta$'s for the hold out data for each partitioning, i.e. each hold out fold of the validation (outer) loop of nested cross validation. In this manner there are multiple model fits and hold out subsets, e.g. k from the k -fold nested CV, for calibration based upon independent observations. Here for each of these subsets we calibrate by regressing outcome on the $X \cdot \beta$'s. While each of these calibrations will individually have limited information, when combined (averaged) following the principles of cross validation or bootstrap sampling, they will collectively provide a more meaningful evaluation. When using the bootstrap option the function also stores the $X \cdot \beta$'s for both the out-of-bag sample elements not selected by the sample with replacement of each bootstrap sample of the validation loop, as well as the in-bag sample elements. A calibration plot can be drawn using the `calplot()` function as in

```
calplot(nested.cox.fit, wbeta=3)
```

```
## Range of X*Beta for calibration:
## -8.339908 9.246226
## Range of calibrated naive confidence intervals:
## -11.62559 11.66253
```

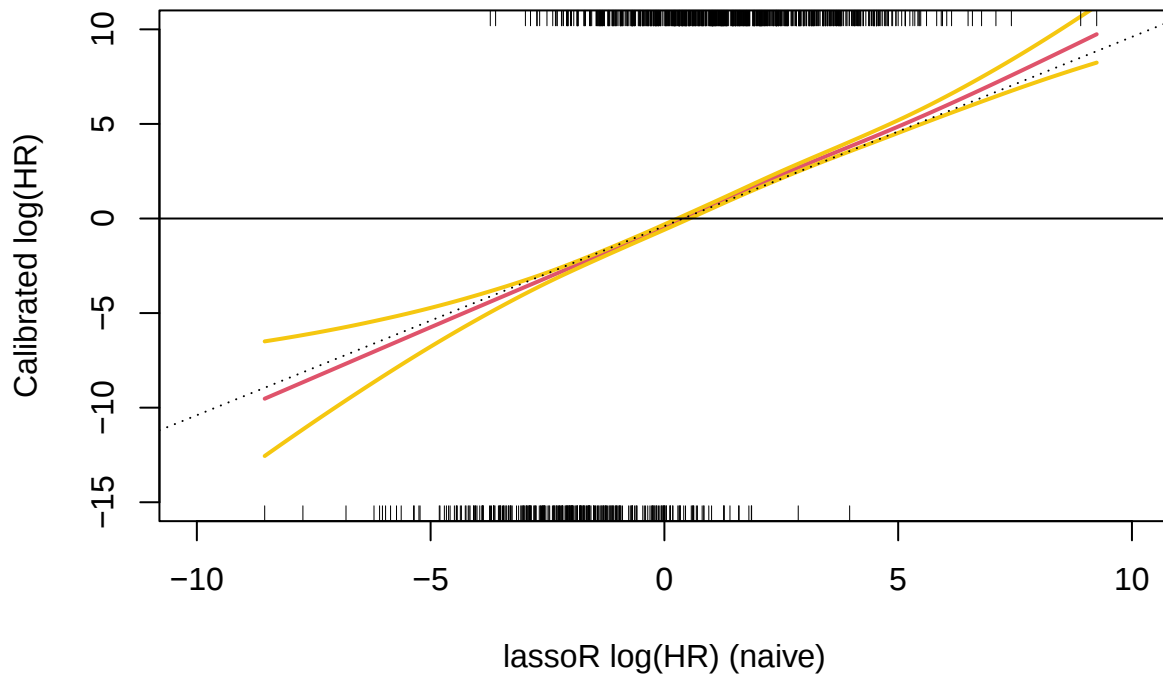


Here we see a smooth, nearly linear predicted log hazard ratio as a function of the model $X \cdot \text{Beta}$ from the relaxed lasso model. The bounding lines in red depict the average $\pm qt(0.975, k-1)$ standard errors (SE) calculated using the Nadeau-Bengio adjustment (usual $SE \cdot \sqrt{1+k/(k-1)}$) (Nadeau et al.) to assist in assessing meaningfulness in any deviation from the ideal identity line, and non linearities. In these curves the central region with solid lines denotes the region within the range of all the calibration spline fits, i.e. spline fits from all the different leave-out folds of the CV overlap without extrapolation. The dashed lines depict areas out of range for at least one of the leave out folds. Because spline fits can be rather uncertain when extrapolating beyond the data range, one should be more cautious when drawing conclusions in the dashed regions of these plots. Here we see somewhat wider confidence bounds about the overall calibration curve than in the naive calibration figure.

Bengio et al. showed that the standard deviations from cross validation might not be accurate for estimation of the actual standard errors. Bates et al. showed that the naive cross validation (CV) estimates and standard errors may be biased, and should thus be viewed with caution. In particular the CV estimates may more closely estimate the expected performance measures over multiple samples than the performance of the model based upon all observed data, and usage the naive CV standard errors, when constructing confidence intervals might be associated with non-coverage three times the nominal non-coverage. This is discussed further in the vignette “An Overview of glmnet”.

A naive calibration curve similar to that shown above (the first figure) can also be easily gotten using the `calplot()` function when specifying to not use the resample for construction as in

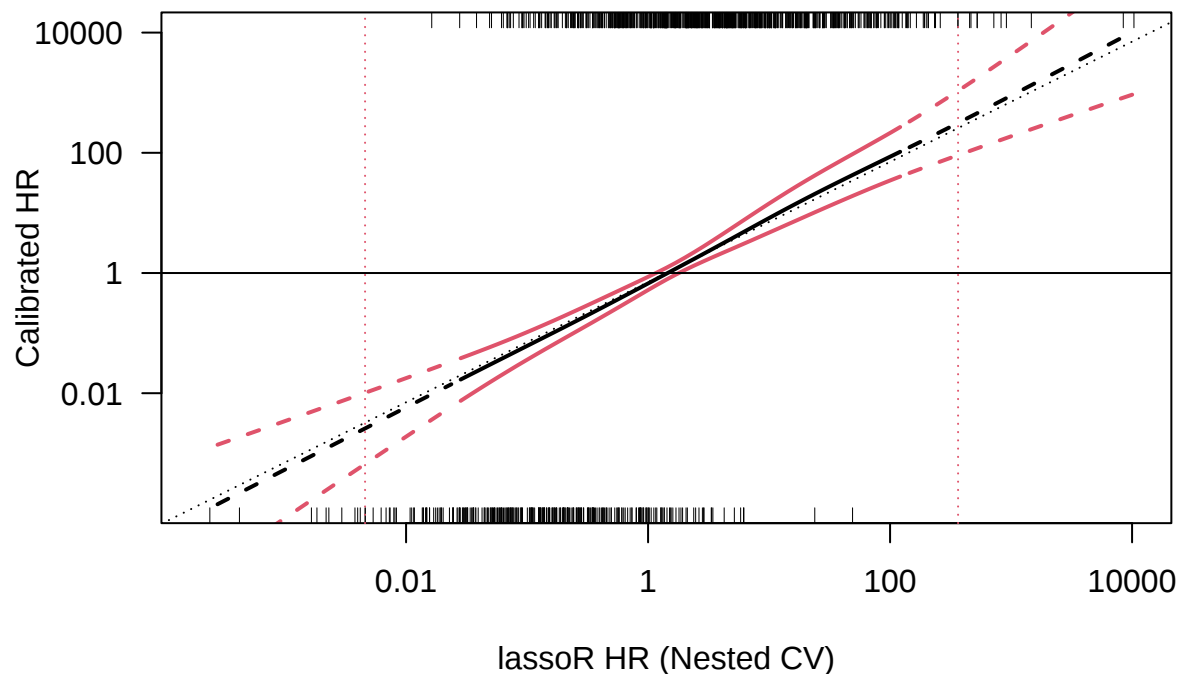
```
calplot(nested.cox.fit, wbeta=3, resample=0, xlim=c(-10,10), ylim=c(-15,10),
        col.term=2, col.se=7)
```



In the `calplot()` generated figures (so far, for the Cox and binomial models) we see two rugs, one below and one above the plotted region. When using cross validation to evaluate model performance the rug below depicts the model X^*Beta 's which are not associated with an event and the rug above depicts X^*Beta 's which are associated with events. When cross validation is used to evaluate model performance the X^*Beta are taken from the estimate for the leave out data. When bootstrap is used the X^*Beta are either the averages of the Out-Of-Bag or In-Bag X^*Beta 's depending on the users input to `calplot()`. When there are lots of data points it can be hard to read these rugs. One can use the `vref` option in `calplot()` to draw two vertical lines where the first separates the smaller `vref%` of the X^*Beta 's from the rest, and a second which separates the larger `vref%` of the data. To depict the hazard ratios (HR) instead of the X^*Beta for the Cox model one can use the option `plotr`, where one assigns a numerical value for the product between tick marks, e.g. `exp(1)` or `10`. Combining these two options we have the example

```
calplot(nested.cox.fit, wbeta=3, vref=1, plothr=100)
```

```
## Range of X*Beta for calibration:
## -8.339908 9.246226
## Range of calibrated naive confidence intervals:
## -11.62559 11.66253
```

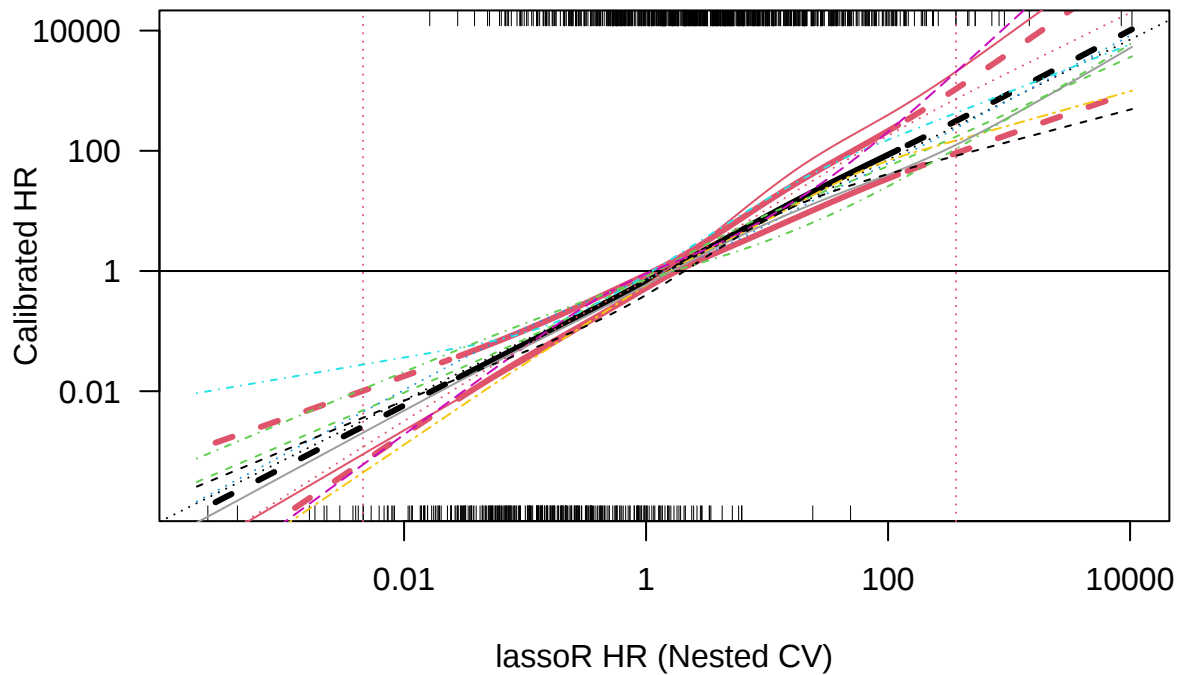


The user can also use different colors for the lines with the options `col.term`, `col.se`. One can also specify `xlim` and `ylim` in case a few data points cause an excessive amount of white space or odd aspect ratio in the plots.

To view the calibration plots for the individual leave out cross validation folds, one may specify `plotfold = 2`. In that this generates many figures, we omit in this vignette actually producing plots using this option specification, and instead assign `plotfold=1` which overlays the individual calibration curves, albeit without the $\pm qt(0.975, k-1)$ SE limits for the individual CV folds. The overall calibration (average of the individual CV fold calibrations) and overall $\pm qt(0.975, k-1)$ SE limits though are maintained.

```
calplot(nested.cox.fit, 3, plotfold=1, vref=1, plothr=100)
```

```
## Range of X*Beta for calibration:
## -8.339908 9.246226
## Range of calibrated naive confidence intervals:
## -11.62559 11.66253
```



As we see from the above calls the first term in the `calplot()` function call is an output object from a `nested.glmnet()` call. The second term, `wbeta`, specifies “which beta” or model is to be used for deriving the model $X \cdot \text{Beta}$ ’s. Here, as we see in the figure x-axis label, the 3 determines the relaxed lasso model. Instead of making a hard to remember key the user can leave this term unspecified and a key will be directed to the R console. The actual numbers for the different models will depend on which models are fit and so this key is dynamic.

```
calplot(nested.cox.fit)
```

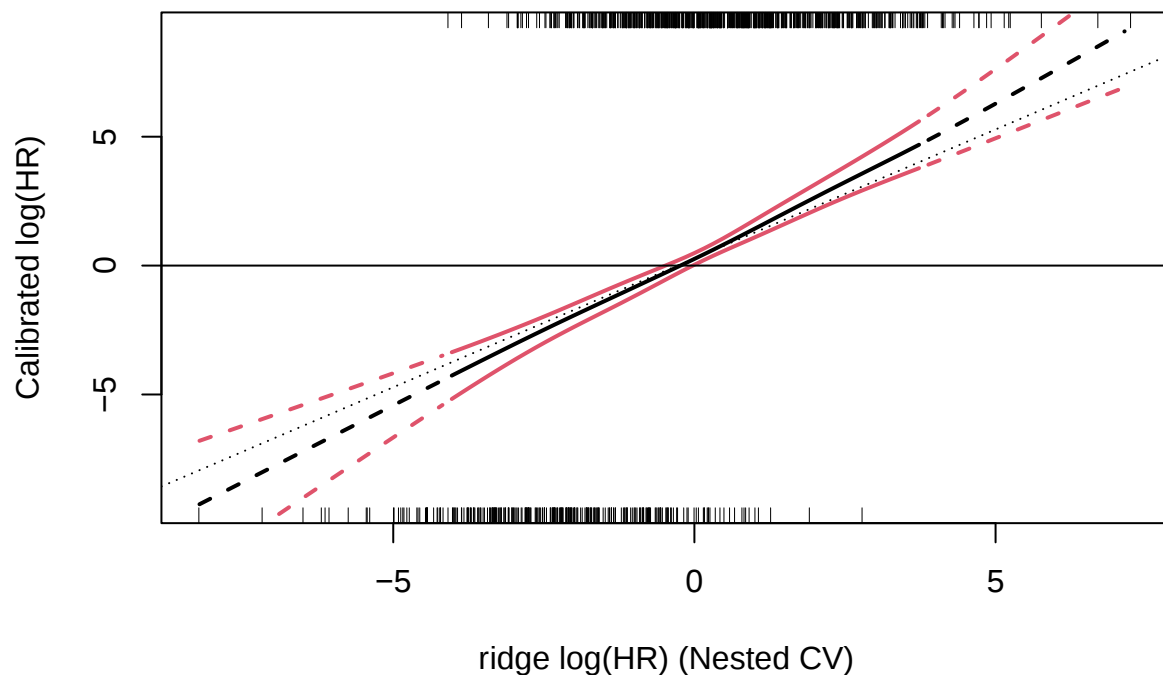
```
## specify num for wbeta =
##           Var wbeta
## null      0.000000    1
## lasso      5.142054    2
## lassoR     6.046005    3
## lassoR0    6.842611    4
## elastic   5.821383    5
## ridge     4.633838    6
## elastic G0 6.928897    7
## elastic G1 5.142054    8
## lasso cal  7.225475    9
## lassoR cal 6.883119   10
## lassoR0 cal 6.842685   11
## elastic cal 7.061203   12
## ridge cal  7.923128   13
## elastic G0 cal 6.928922 14
## elastic G1 cal 7.225475 15
```

```
## step.df      7.071502    16
## step.p       7.104980    17
```

From this key we read off the numbers corresponding to the respective models. The variance in $X\beta$ for the “null” model is 0 as the intercept for the Cox model is arbitrarily assigned the value of 0 for each resample model fit. From this key we see we can produce a calibration plot for the ridge regression model by setting $w\beta = 6$ ($w\beta$ for “which beta”), as in

```
calplot(nested.cox.fit, 6)
```

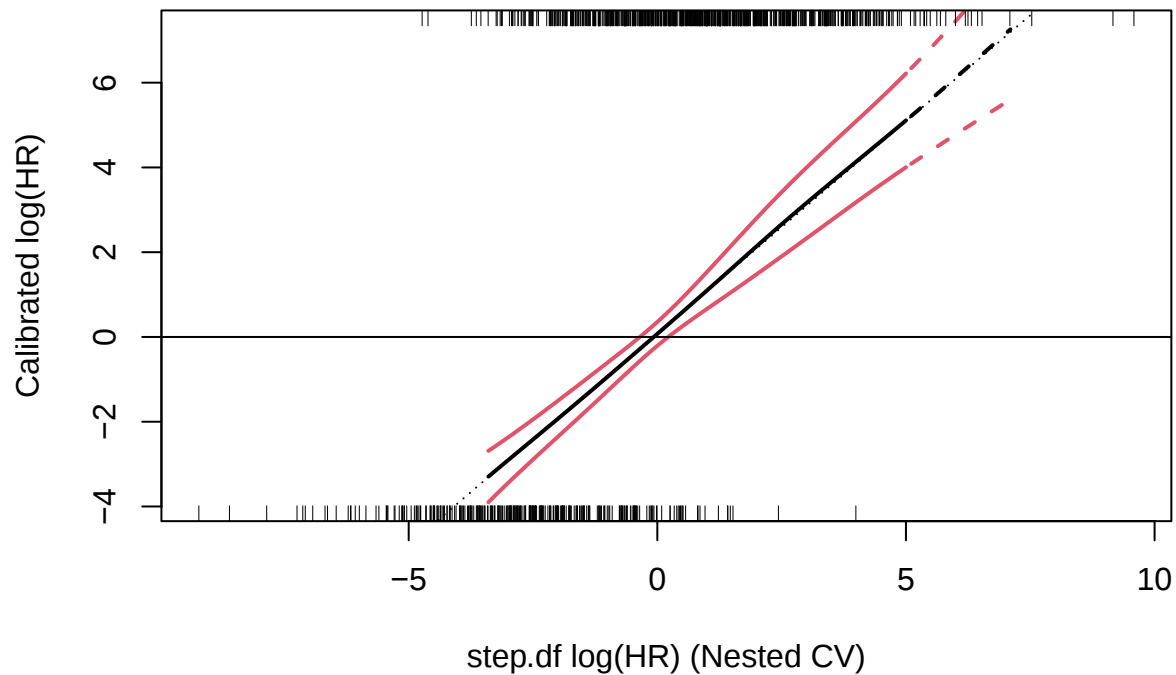
```
## Range of X*Beta for calibration:
## -8.227856 7.241285
## Range of calibrated naive confidence intervals:
## -11.71907 11.52762
```



Here we see the calibration fit is slightly more separated from the identity line than the calibration curve for the relaxed lasso model. Inspecting the calibration curve for a step wise regression model, where the number of terms included in the model is informed by cross validation

```
calplot(nested.cox.fit, 16)
```

```
## Range of X*Beta for calibration:
## -9.220216 9.585343
## Range of calibrated naive confidence intervals:
## -3.899754 8.91537
```



we see that the response deviates some form the identity line but not significantly so.

To obtain the numerical values used to construct these calibration plots one may specify `plot=0` (or `plot=2` to plot and obtain the numerical data) in list format as in

```
tmp = calplot(nested.cox.fit, 3, plot=0)
```

```
## Range of X*Beta for calibration:
## -8.339908 9.246226
## Range of calibrated naive confidence intervals:
## -11.62559 11.66253
```

```
str(tmp)
```

```
## List of 5
## $ estimates      : num [1:101, 1:5] -8.55 -8.37 -8.19 -8.02 -7.84 ...
## .. attr(*, "dimnames")=List of 2
## .. ..$ : chr [1:101] "1" "2" "3" "4" ...
## .. ..$ : chr [1:5] "plotxb" "est" "se" "lower" ...
## $ est.resample   : num [1:10, 1:101] -10.74 -8.08 -8.83 -4.68 -12.07 ...
## $ se.resample    : num [1:10, 1:101] 7.18 3.69 3.94 3.46 5.6 ...
## $ lower.resample : num [1:10, 1:101] -25.1 -15.5 -16.7 -11.6 -23.3 ...
## $ upper.resample : num [1:10, 1:101] 3.612 -0.709 -0.948 2.245 -0.87 ...
```

These data may be further processed by the user.

A Binomial Model

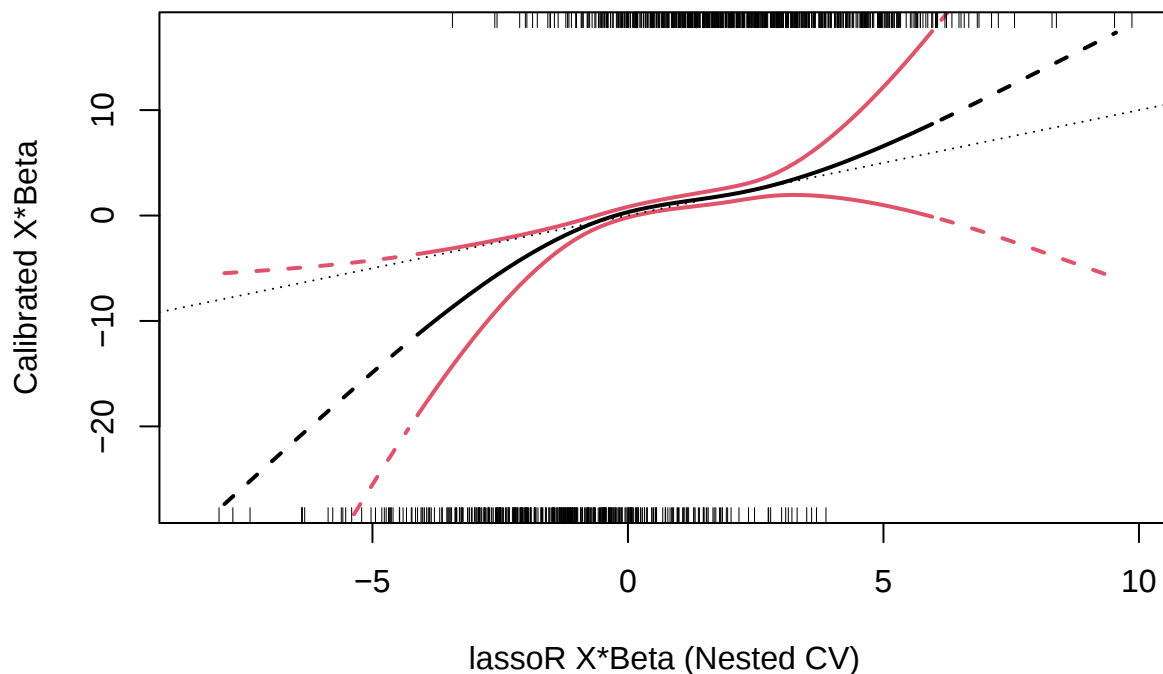
For `nested.glmnet()` analyses with family = “binomial” with the call

```
yb = simdata$yb
nested.bin.fit = nested.glmnet(xs=NULL,yb=NULL,family="binomial",
  dolasso=1, doxgb=list(nrounds=250), dorf=1, doorf=1, doann=1,
  folds_n=10, seed=219301029, track=1)
```

an example calibration plot is

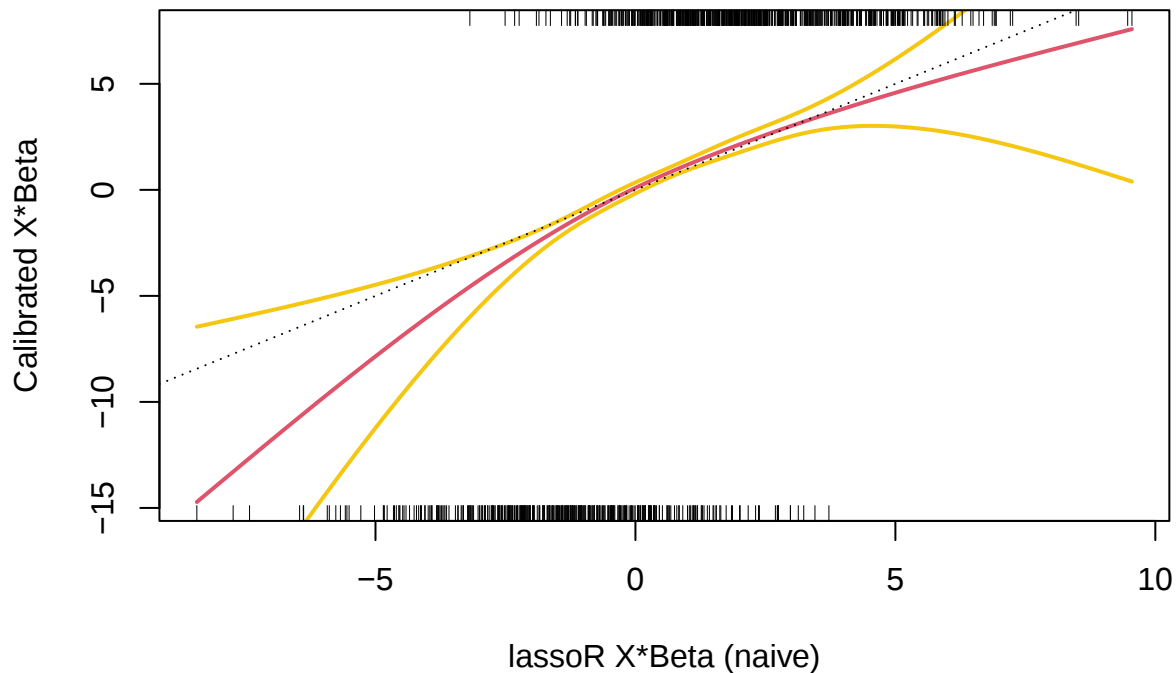
```
calplot(nested.bin.fit, 3, plotfold=0)
```

```
## Range of X*Beta for calibration:
## -8.002414 9.862854
## Range of calibrated naive confidence intervals:
## -53.8154 40.62347
```



Since these data were generated with probabilities $\exp(X*\text{Beta})/(1+\exp(X*\text{Beta}))$ we want that the model would calibrate linearly. The wide variability for the larger $X*\text{Beta}$, positively or negatively, is likely due to the lack of stability in estimates for probabilities near 0 and 1. Note that the confidence limits include the identity line over the entire range of the $X*\text{Beta}$, just barely, meaning there is not strong evidence supporting the nonlinearities of the calibration curve. Examination of the naive calibration plot which uses the same data for model derivation and calibration

```
calplot(nested.bin.fit, 3, resample=0, df=3, col.term=2, col.se=7)
```



shows a self consistency better than the above plots. While this may calibrate slightly better it too is not as linear as we might expect.

A Binomial Model Calibrated Using Bootstrap

Considering this we next evaluate model performances using bootstrap, that is fit models based upon random samples from the original sample (with replacement) of same sample size as the original sample. Then we fit calibration curves for the out-of-bag sample units for each bootstrap sample fitted model, that is the elements of the original sample that are not selected by the bootstrap sample. The number of bootstrap samples for calculation is specified using the *bootstrap* option in the *nested.glmnetr()* call,

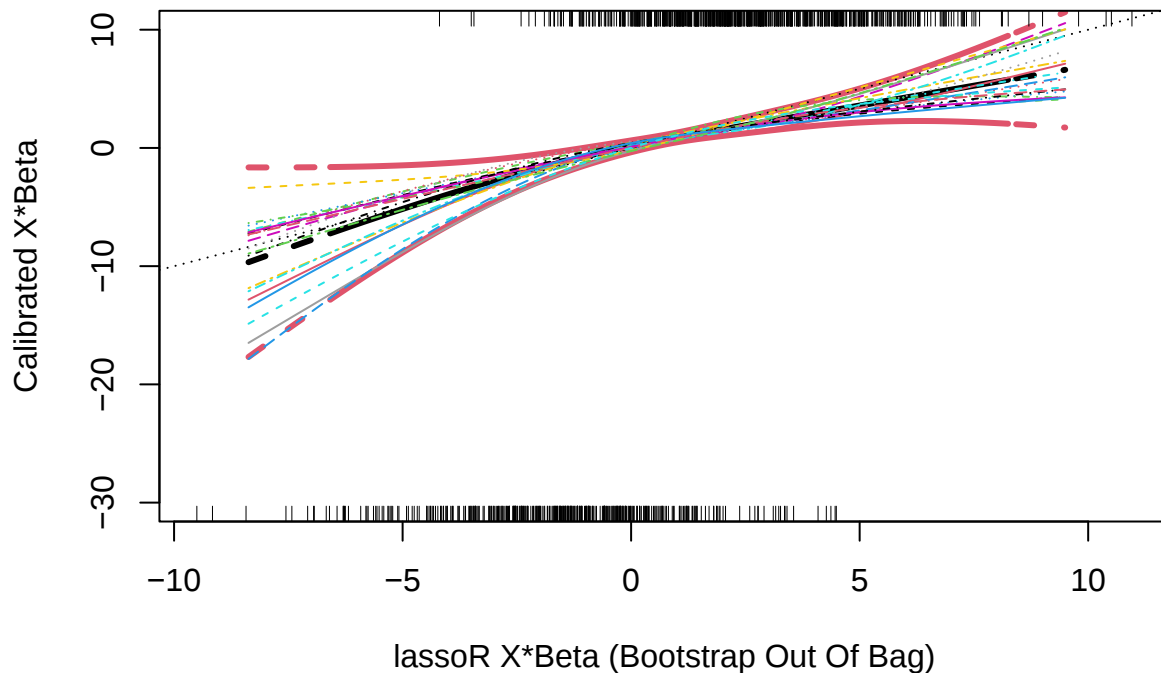
```
yb = simdata$yb
nested.bin.boot.fit = nested.glmnetr(xs, NULL, yb, NULL, family="binomial",
  dolasso=1, dorf=1,
  folds_n=10, seed=219301029, track=1, bootstrap=20)
```

Bootstrap Out-Of-Bag (OOB) Calibrations

The Out-Of-Bag (OOB) calibration plots can then be constructed using the *calplot()* function as before,

```
calplot(nested.bin.boot.fit, wbeta=3, plotfold=1, ylim=c(-30,10))
```

```
## Range of X*Beta for calibration:
## -9.501871 10.95765
## Range of calibrated naive confidence intervals:
## -17.68289 11.49037
```



As for the nested CV, the solid center black line is the average of the calibration lines from the multiple bootstrap resamples. The upper and lower thick red lines are the average $\pm$ $qt(0.975, k-1)$ (the 0.975 quantile of the t-distribution with degrees of freedom the number of folds in the cross validation, or bootstrap replications) standard deviations from the bootstrap resample calibrations. Here we see a calibration plot that seems more on target than we did for the nested cross-validation plots. This may be due to the larger sample size for the hold out datasets in the individual validations, circa 37% of the original sample as opposed to 10% for 10-fold cross-validation. For the cross validation one could consider 3-fold cross validation to increase the hold-out sample sizes, repeat the whole process many times using different seeds (see the `glmnet` package vignettes) and then average, but we have not done this here. For the bootstrap one can also easily run a larger number of re-samples to improve estimate stability. One might also consider a method similar to the bootstrap which selects unique observations in the resamples and then evaluate on the hold out data for each resample. This can be done using the `unique` option in `nested.glmnetr()`, e.g. setting `unique=0.63` to randomly select 63% of the sample for model fitting leaving 37% for model evaluation in each resample.

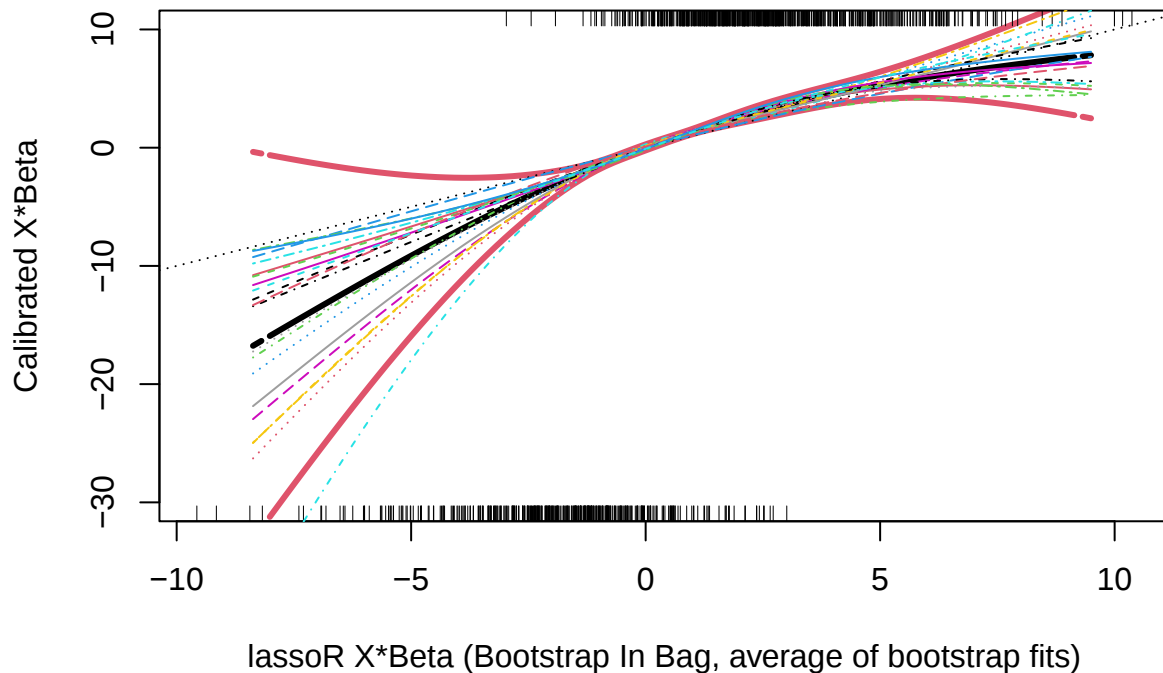
Bootstrap In-Bag Calibrations

In earlier works Austin et al. as well as Riley et al. fit models based upon bootstrap samples and then fit calibration curves based upon the in-bag data points and model $X \cdot \beta$ (predicted). This mimics the usual

procedure for bootstrap analysis. This can be done using the `calplot()` function by setting the *oob* option to 0,

```
calplot(nested.bin.boot.fit, wbeta=3, oob=0, plotfold=1, ylim=c(-30,10))
```

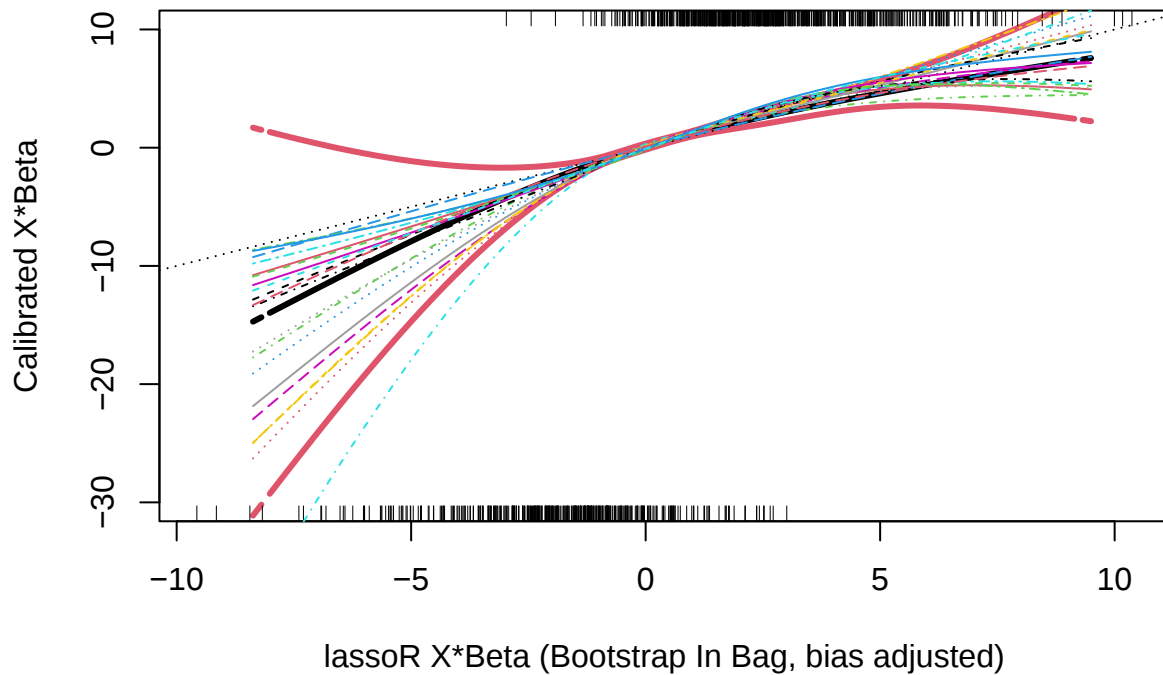
```
## Range of X*Beta for calibration:
## -9.56948 10.36934
## Range of calibrated naive confidence intervals:
## -33.16328 13.15909
```



Here the thin lines are taken from in-bag calibration plots from the individual bootstrap resamples. The center thick line is the average of the calibration plots from the individual bootstrap resamples, and the upper and lower thick red lines are the average $\pm$ $qt(0.975, k-1)$ standard deviations from the bootstrap resample calibrations. When using the bootstrap one common approach for estimation is to take the estimate from the full sample analysis and describe confidence intervals in terms of this estimate and the variability in the in-bag resample estimates. This may be gotten using the `bootci=1` option as in

```
calplot(nested.bin.boot.fit, wbeta=3, oob=0, bootci=1, plotfold=1, ylim=c(-30,10))
```

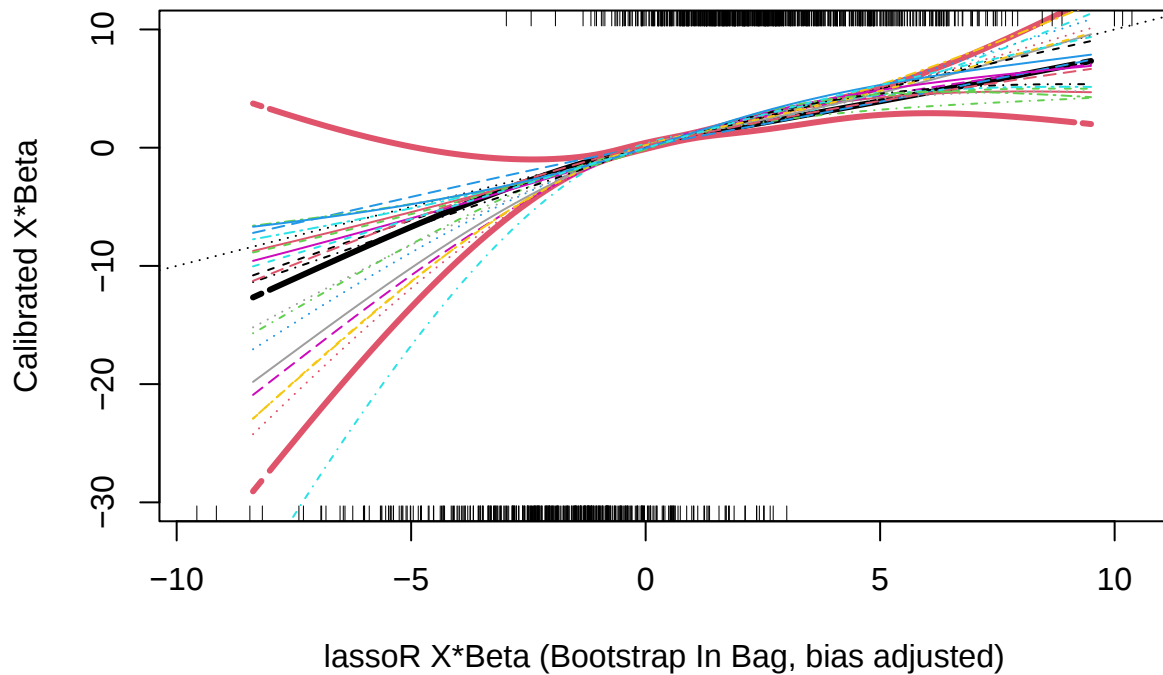
```
## Range of X*Beta for calibration:
## -9.56948 10.36934
## Range of calibrated naive confidence intervals:
## -31.11647 12.91578
```



where we see a shift and slight change in shapes of the curves. The bootstrap estimate can also be biased in the full data model by $(\text{Ave}(\text{XBeta}_i) - \text{XBeta}_{\text{full}})$. We can make this bias adjustment using the `biasbc=1` (bc for bias adjustment) as with

```
calplot(nested.bin.boot.fit,wbeta=3,oob=0,bootci=1,bootbc=1,plotfold=1,ylim=c(-30,10))
```

```
## Range of X*Beta for calibration:
## -9.56948 10.36934
## Range of calibrated naive confidence intervals:
## -29.06965 12.67247
```



This figure again differs slightly from the above figures. The in-bag bootstrap calibrations here seem different from the out-of-bag calibrations.

Bootstrap Calibration for a Random Forest

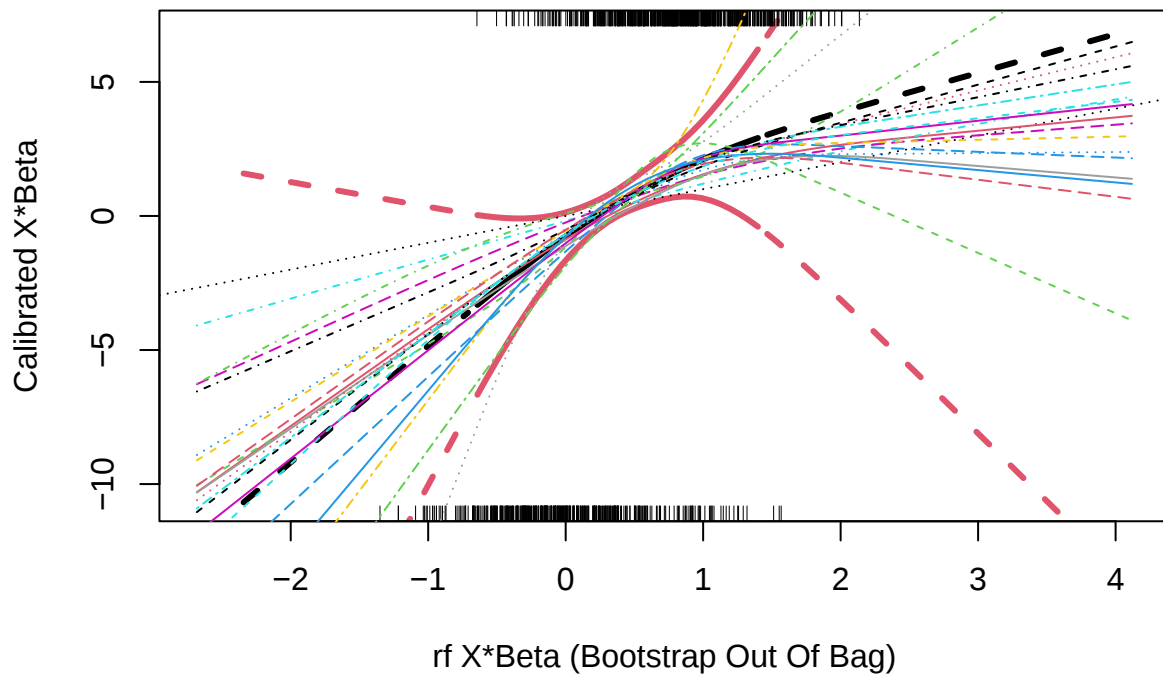
These bootstrap plots above were for the relaxed lasso model fit. We now examine similar bootstrap calibration plots but for random forest models which have greater potential for overfit.

Bootstrap Out-Of-Bag (OOB) Calibrations for a Random Forest

For the out-of-bag calibration bootstrap plots

```
calplot(nested.bin.boot.fit, wbeta=9, plotfold=1)
```

```
## Range of X*Beta for calibration:  
## -1.350711 2.136162  
## Range of calibrated naive confidence intervals:  
## -26.23269 27.67918
```



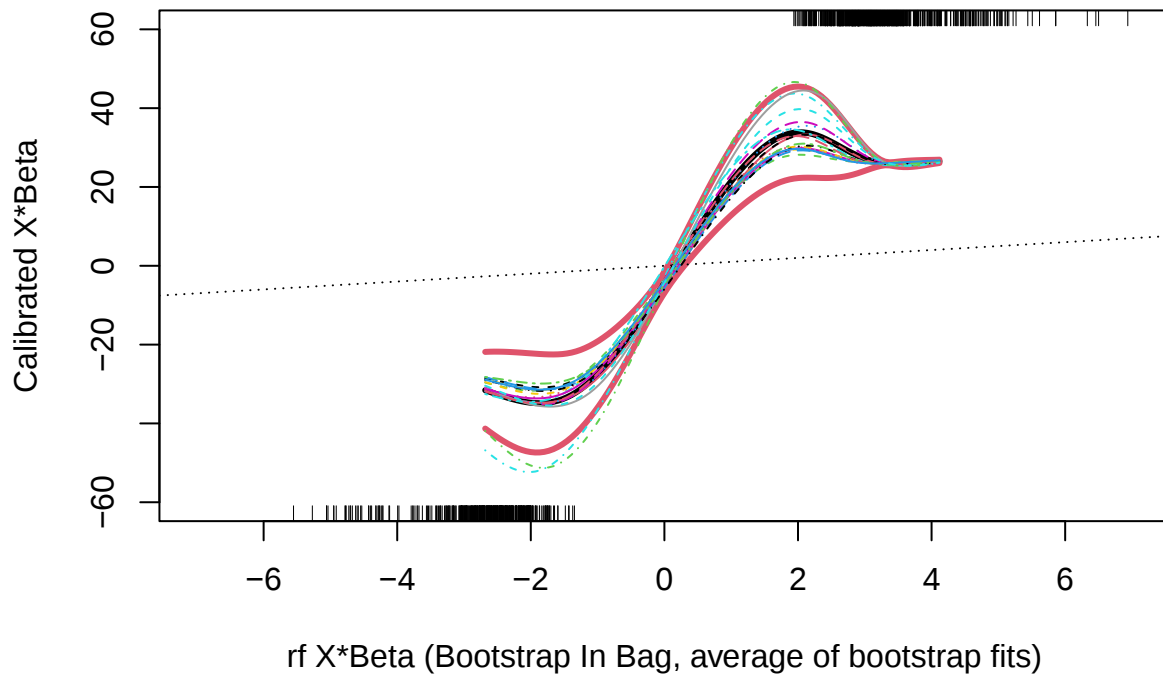
we see that these are highly variable yet consistent with the ideal calibration curve, the identity line.

Bootstrap In-Bag Calibration for a Random Forest

The variability in the in-bag calibration curves from in-bag bootstrap resamples and for the random forest model is depicted in the graph

```
calplot(nested.bin.boot.fit, wbeta=9, oob=0, plotfold=1, df=6,
        ylim=c(-60,60), xlim=c(-7,7))
```

```
## Range of X*Beta for calibration:
## -5.552771 6.93748
## Range of calibrated naive confidence intervals:
## -47.37791 45.48823
```



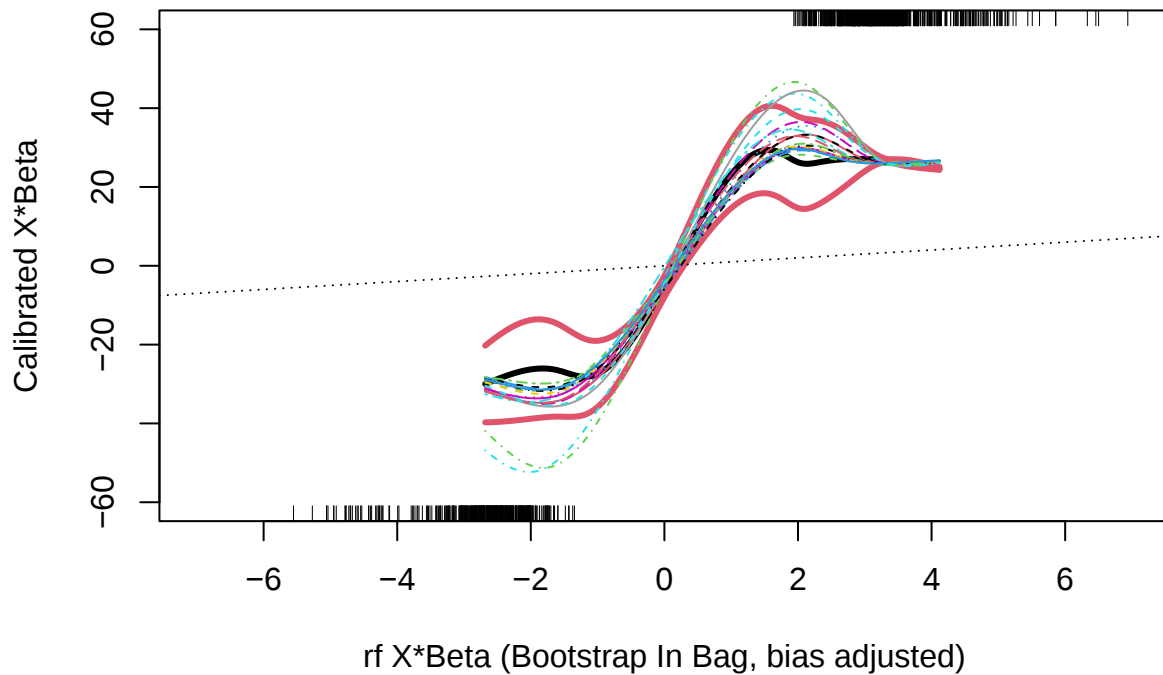
Note, here we used the `df=6` option for the degree of freedom in spline fitting to get a set of more reasonable fits.

Bootstrap In-Bag Calibration for a Random Forest using a bootstrap approach

As described above a common way to use the bootstrap to construct confidence intervals is to take the estimate based upon all data and use the bootstrap in-bag sample estimates to inform about the variability of the full sample variability. This is depicted in the graph

```
calplot(nested.bin.boot.fit, wbeta=9, oob=0, bootci=1, plotfold=1, df=6,
        ylim=c(-60,60), xlim=c(-7,7))
```

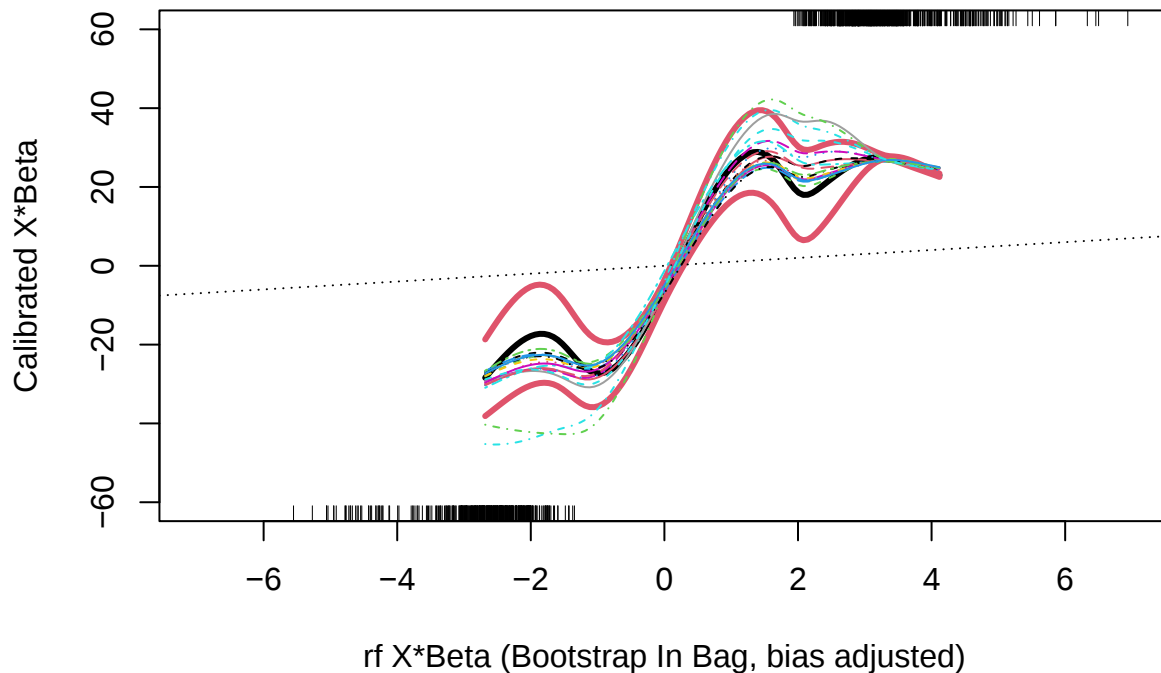
```
## Range of X*Beta for calibration:
## -5.552771 6.93748
## Range of calibrated naive confidence intervals:
## -39.72065 40.61569
```



Here we see a strong deviation from the identity line. The tick marks for the rugs (based upon average X^*Bata 's over the many bootstrap fits) also show that the random forest may largely be separating events from non-events for the logistic model framework, suggesting overfit. The bias corrected bootstrap calibration curves are depicted by

```
calplot(nested.bin.boot.fit, wbeta=9, oob=0, bootci=1, bootbc=1, plotfold=1, df=6,
        ylim=c(-60,60), xlim=c(-7,7))
```

```
## Range of X*Beta for calibration:
## -5.552771 6.93748
## Range of calibrated naive confidence intervals:
## -38.13846 39.48236
```



have a rather unusual and unexpected deviation from the ideal calibration curve of the identity line. This too may be due to overfit.

While the two in-bag calibration curve sets strongly depart from the ideal of the identity line, the out-of-bag curves while possibly highly variable were consistent as a group with the ideal of the identity line. The bootstrap out-of-bag calibration curves have a clearer basis and depict more reasonable findings. In a sense this is not unexpected. When fitting machine learning models one typically evaluates model fit based upon the out-of-bag data points and not the in-bag data points. A difficulty here could be that the calibration curves, calculated at particular values for $X \cdot \beta$, do not describe a well defined parameter in terms of the distribution function. Potentially as well the fitting process for machine learning models might not lend it self to the assumptions typically used to support bootstrap inferences. In particular as we see here when fitting random forest models, with the multiple repeat observations in the bootstrap resamples the model refits are overly optimistic, more strongly overfit and give more extreme $X \cdot \beta$ as evidenced in these figures. We present the in-bag and biased adjusted bootstrap calibration curves not to promote their usage but to 1) show the variability in model fit and 2) how they may give poor inferences for a simple dataset. The user may want to generate other datasets of different known form and see how the bootstrap performs for their data.

A Normal (Gaussian Errors) Model

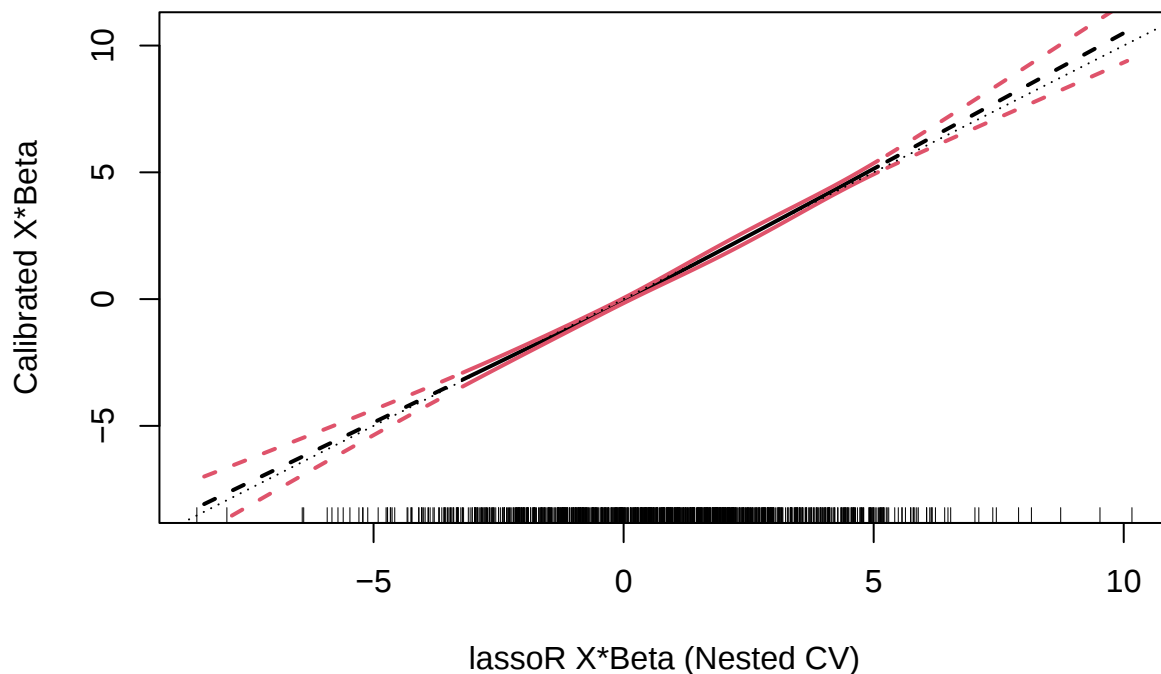
First calculating the numerical summaries of prediction performance

```
nested.gau.fit = nested.glmnetr(xs,NULL,y_,NULL,family="gaussian",
                                dolasso=1, doxgb=list(nrounds=250), dorf=1, doorf=1, doann=list(bestof=10),
                                folds_n=10, seed=219301029, track=1)
```

and then plotting

```
calplot(nested.gau.fit, wbeta=3)
```

```
## Range of X*Beta for calibration:
## -8.533496 10.1643
## Range of calibrated naive confidence intervals:
## -9.169553 11.74125
```



we see a small not significant deviation from the ideal calibration line, the identity function.

Perspective

The summary and calibration plot functions used here do not address all needed for model validation and calibration but do allow a meaningful and un (or minimally) biased summary of model fits. The original outcome variable and X*betas are stored as vectors or matrices in the output with names `y_`, `xbetas` (full

model) and `xbetas.cv` (cross-validation) and `xbetas.boot.oob` and `betas.boot.inb` (bootstrap out-of-bag and in-bag) allowing the user to further inspect model fits and to perform other means of calibration. For cross-validation analyses the fold information is contained in the output object in the vector `foldid`. For bootstrap analyses the first column of `xbetas.boot.oob` and `betas.boot.inb` describe the replication of the bootstrap sampling process and the second columns the sequential index for the data point in the input dataset.

References

- Austin PC, Steyerberg EW. Graphical assessment of internal and external calibration of logistic regression models by using loess smoothers. *Stat Med.* 2014; 33(3):517-35. doi: 10.1002/sim.5941 .
- Bates S, Hastie T, Tibshirani R, “Cross-validation: what does it estimate and how well does it do it?”, 2022, <https://arxiv.org/abs/2104.00673> .
- Bengio Y & Grandvalet Y, “No Unbiased Estimator of the Variance of K-Fold Cross-Validation”, *Journal of Machine Learning Research* 5, 2004, 1089–1105, <https://www.jmlr.org/papers/volume5/grandvalet04a/grandvalet04a.pdf> .
- Nadeau C, Bengio Y, “Inference for the Generalization Error”, *Machine Learning*, 52, 239–281 (2003).
- Riley RD, Pate A, Dhiman P, Archer L, Martin GP, Collins GS. Clinical prediction models and the multiverse of madness. *BMC Med.* 2023; 21(1):502. doi: 10.1186/s12916-023-03212-y .