

A Coefficient of Determination for GLMs

by Dabao Zhang

Rearranged by Jae Ho, Chang
Spring, 2018

OLS

R^2_{OLS} ; A Coefficient of Determination for OLS

Consider a linear model ;

$$\mathbf{y}_{n \times 1} = \mathbf{X}_{n \times (p+1)} \boldsymbol{\beta}_{(p+1) \times 1} + \boldsymbol{\epsilon}_{n \times 1} , \quad \boldsymbol{\epsilon} \sim N_n(\mathbf{0}, \sigma^2 \mathbf{I}_n)$$

Then we have

$$SSE = \sum_{i=1}^n (y_i - \hat{y}_i)^2, \quad SST = \sum_{i=1}^n (y_i - \bar{y}_i)^2$$

$$R^2_{OLS} \stackrel{def}{=} 1 - \frac{SSE}{SST}$$

OLS

- ① R_{OLS}^2 measures **the goodness of fit**
- ② However, for GLMs?
- ③ Some GLMs with well-defined likelihood

Specified Likelihood

For a logistic regression, perfectly fitted values result in $l(\mathbf{y}, \hat{\mu}(\mathbf{X})) = 0$. That is, samples can be perfectly separated by a linear function.

$$\max(R_{LR}^2) = 1 - \exp \left\{ \frac{2}{n} l(\mathbf{y}, \hat{\mu}(\mathbf{1}_n)) \right\}$$

For example, with balanced case-control data (Bernoulli-distributed with $\pi = 0.5$), $\max(R_{LR}^2) = 0.75$. That is, R_{LR}^2 is bounded from above by $l(\mathbf{y}, \hat{\mu}(\mathbf{1}_n))$ and will never attain value 1.

This is inconsistent with the existing concept of R^2 .

Specified Likelihood

Nagelkerke(1991) ; Suggested the corrected one.

$$R_N^2 \stackrel{def}{=} \frac{R_{LR}^2}{1 - \exp\left\{\frac{2}{n}l(\mathbf{y}, \hat{\mu}(\mathbf{1}_n))\right\}} = \frac{R_{LR}^2}{\max(R_{LR}^2)} \in [0, 1]$$

But still inconsistent with the classical definition of coefficient of determination.

Cameron and Windmeijer(1997) ; Use the KL divergence to quantify the uncertainty remaining in the response after accounting for predictors.

Limits

All the aforementioned generalized coefficients of determination are given through **completely specified likelihood function**.

However, these are not applicable for more general GLMs like quasi-models which specify only the **mean and variance functions**.

Bartlett's Identities

The first and second Bartlett results ensures that under suitable conditions (see Leibniz integral rule), for a density function dependent on θ , $f_\theta(\cdot)$,

$$E_\theta \left[\frac{\partial}{\partial \theta} \log(f_\theta(y)) \right] = 0$$

$$Var_\theta \left[\frac{\partial}{\partial \theta} \log(f_\theta(y)) \right] + E_\theta \left[\frac{\partial^2}{\partial \theta^2} \log(f_\theta(y)) \right] = 0$$

Bartlett's Identities

Expected value of $\mathbf{Y}$: Taking the first derivative with respect to θ of the log of the density in the exponential family form described above, we have

$$\frac{\partial}{\partial \theta} \log(f(y, \theta, \phi)) = \frac{\partial}{\partial \theta} \left[\frac{y\theta - b(\theta)}{\phi} - c(y, \phi) \right] = \frac{y - b'(\theta)}{\phi}$$

. Then taking the expected value and setting it equal to zero leads to,

$$E_{\theta} \left[\frac{y - b'(\theta)}{\phi} \right] = \frac{E_{\theta}[y] - b'(\theta)}{\phi} = 0$$

$$E_{\theta}[y] = b'(\theta)$$

Bartlett's Identities

Variance of Y: To compute the variance we use the second Bartlett identity,

$$\text{Var}_\theta \left[\frac{\partial}{\partial \theta} \left(\frac{y\theta - b(\theta)}{\phi} - c(y, \phi) \right) \right] + \text{E}_\theta \left[\frac{\partial^2}{\partial \theta^2} \left(\frac{y\theta - b(\theta)}{\phi} - c(y, \phi) \right) \right]$$

is 0. Then,

$$\text{Var}_\theta \left[\frac{y - b'(\theta)}{\phi} \right] + \text{E}_\theta \left[\frac{-b''(\theta)}{\phi} \right] = 0, \quad \text{Var}_\theta [y] = b''(\theta)\phi$$

Bartlett's Identities

We have now a relationship between μ and θ , namely $\mu = b'(\theta)$ and $\theta = b'^{-1}(\mu)$, which allows for a relationship between μ and the variance,

$V(\theta) = b''(\theta)$ = the part of the variance that depends on θ

$$V(\mu) = b''(b'^{-1}(\mu))$$

or,

$$V(\mu) = (b'' \circ b'^{-1})(\mu)$$

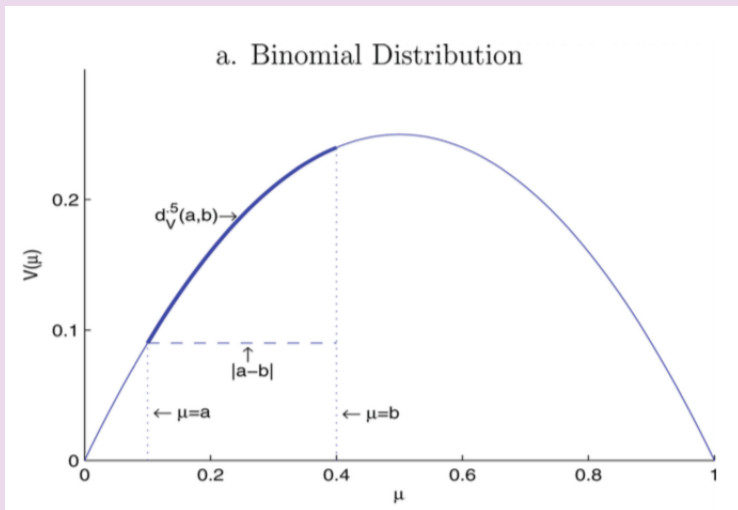
Variance Function - Bernoulli case

This give us

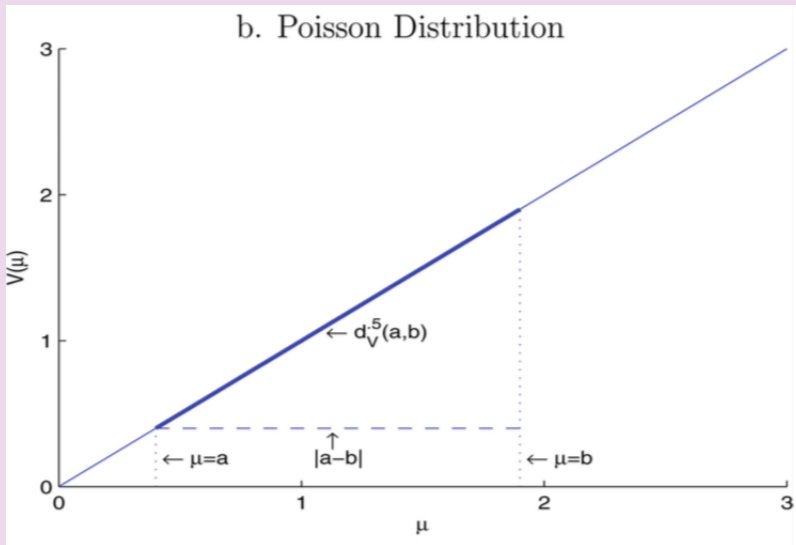
$$V(\mu) = \mu(1 - \mu) = \mu - \mu^2.$$

In this case, dispersion parameter $\phi = 1$.

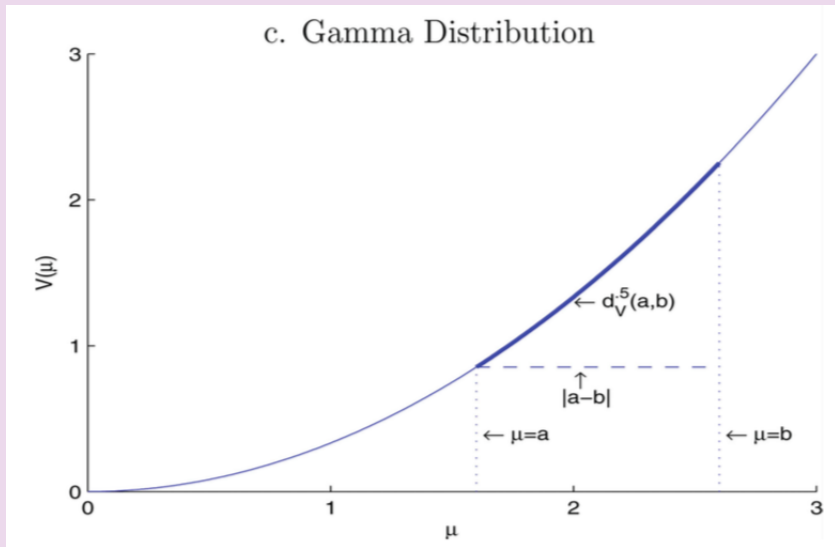
Variation change along the variance function



Variation change along the variance function



Variation change along the variance function



Consistency of R_V^2 and its Extension

$V'(\cdot)$ is **constant** for normal and Poisson distributions.
That is, $R_V^2 = R^2$ for OLS with normal distribution and log-linear model with Poisson distribution.

Similar to the coefficient of determination, the coefficient of **partial determination** is well-defined for linear models.

Measures the proportion of variation in the response variable not explained by a set of predictors that can be explained by an additional set of predictors.

Partial Determination

For example, considering two sets of predictors X_1 and X_2 in a linear regression model, we have

$$R^2(X_2|X_1) = 1 - \frac{SSE(X_1, X_2)}{SSE(X_1)} = \frac{R^2(X_1, X_2) - R^2(X_1)}{1 - R^2(X_1)}$$

($\therefore$) Recall that $1 - R^2(X_1) = \frac{SSE(X_1)}{SST}$.

measuring the proportion of remaining variation in the response, when including X_1 , explained by X_2 .

Partial Determination

With our definition of R_V^2 , we can easily extend it to a coefficient of partial determination for more general models.

$$R_V^2(X_2|X_1) = \frac{R_V^2(X_1, X_2) - R_V^2(X_1)}{1 - R_V^2(X_1)}$$

Binomial Model

When $\beta = 0$, we have $\mu_1 = \mu_2$.

Thus corresponding R^2 's are supposed to report 0 (or very close to 0).

On the other hand, $\mu_1 \rightarrow 0$, $\mu_2 \rightarrow 1$ as $\beta \rightarrow \infty$. Which leads to the single-population based samples.

Therefore, it is not surprising to observe that, when including the true predictor X_1 , R_N^2 , R_{KL}^2 , and R_V^2 approach 1.

See Fig.2.

Candidates

$$R_{LR}^2 = 1 - \exp \left\{ \frac{2}{n} l(\mathbf{y}, \hat{\mu}(\mathbf{1}_n)) - \frac{2}{n} l(\mathbf{y}, \hat{\mu}(\mathbf{X})) \right\}$$

$$R_N^2 = \frac{R_{LR}^2}{\max(R_{LR}^2)}$$

$$R_{KL}^2 = 1 - \frac{\hat{K}L(\beta, \mathbf{X})}{\hat{K}L(\beta, \mathbf{1}_n)}$$

$$R_V^2 = 1 - \frac{\sum_{i=1}^n d_V\{y_i, \hat{\mu}_i(\mathbf{X})\}}{\sum_{i=1}^n d_V\{y_i, \hat{\mu}_i(\mathbf{1}_n)\}}.$$

Poisson & Gamma Model

When $\beta \rightarrow \infty$, $|\mu_1 - \mu_2|$ reaches the maximum in poisson model but is bounded by 50 in the gamma model.

Therefore, it is not surprising to observe that, when including the true predictor X_1, R_{KL}^2 and R_V^2 approach one in the Poisson model, but are barely bounded away from one in the gamma model.

Why R_N^2, R_{LR}^2 falsely claim high R^2 ?

We consider the case $\beta = 5$.

Denote X_{21} as the subset of X_2 with $E(X_1) = \mu_1$ and X_{22} as the subset of X_2 with $E(X_1) = \mu_2$.

With the total sample size at 50, we have $\bar{X}_{21} - \bar{X}_{22} \sim N(0, 0.08)$.

A large value of $\bar{X}_{21} - \bar{X}_{22}$ implies falsely correlated X_1 and X_2 although X_1 and X_2 are truly independent.

Plot of calculated R^2 versus $\bar{X}_{21} - \bar{X}_{22}$ in fig.5.

It shows the robustness of different coefficients of determination when only a falsely correlated X_2 is included.

Why R_N^2 , R_{LR}^2 falsely claim high R^2 ?

R_{LR}^2 and R_N^2 ; **Have tendency to severely overstate the variation proportion explained by the poisson or gamma model.**

On the other hand, both R_{KL}^2 and R_V^2 are more robust to such false correlation.

X_1 vs (X_1, X_2)

$R^2(X_1, X_2)$ is similar to $R^2(X_1)$

. Though the former models always have slightly higher values.

We defined $R_{V,adj}^2$ to adjust for increasing numbers of predictors. Since the K-L divergences used to define R_{KL}^2 are indeed deviances, we can also address the issue of different degrees of freedom in the deviances by defining an adjusted version of R_{KL}^2 as follows;

$$R_{KL,adj}^2 = 1 - \frac{\hat{KL}(\mathbf{y}, \hat{\mu}(\mathbf{X})) / (n - p)}{\hat{KL}(\mathbf{y}, \hat{\mu}(\mathbf{1}_n)) / (n - 1)}$$

X_1 vs (X_1, X_2)

As shown in Figure 6, both $R_{V,adj}^2$ and $R_{KL,adj}^2$ have lowered values when irrelevant predictors X_2 and X_3 are added to the model, with X_3 independently simulated from a standard normal distribution.

Real Data Analysis

We illustrate the different definitions of R^2 by applying them to the data from a study of nesting horseshoe crabs included in Agresti (1996).

(C ; colors), (SC ; spine conditions), (CW ; carapace widths), (W ; weights) of 173 female crabs, each with a male crab attached to her in her nest.

This study intended to investigate whether these factors affect the number of satellites, that is, any other males riding near a female crab.

Real Data Analysis

We consider binomial models ;

$\mathbf{y} = \textit{Whether a female crab had any satellites}$

and Poisson models ;

$\mathbf{y} = \# \textit{ of satellites a female crab had}$

As demonstrated in the previous section, both R_{LR}^2 and R_N^2 indeed severely overstate the variation proportion explained by the Poisson models.

Real Data Analysis

Quasi-binomial models and quasi-Poisson models can be fitted for the above cases, with R_V^2 and partial R_V^2 calculated.

$\hat{\phi} = 1.0266$ for the binomial full model, and $\hat{\phi} = 3.2354$ for the Poisson full model.

Use quasi-Poisson models to allow overdispersion(that is, use of likelihood-based R^2 is inappropriate)!

For Poisson or quasi-Poisson models we have same regression coefficients. This leads to the same R_V^2 and partial R_V^2 .

Conclusions

R^2 ; Measures goodness of fit , also provides a measure of predictability.

R^2 can be used to choose the optimal set of predictors **when the model size**, that is, the number of predictors, **is fixed**.

The *adj.* versions can be used to compare models including different numbers of predictors. For this reason, R_{adj}^2 can be also used to help **model selection, tuning parameter selection**, etc.

Our extension $R_{V,adj}^2$ makes all these possible when any statistical model with a well-defined variance function, such as quasi model.