

Delete-2 Jackknife-After-Bootstrap in Regression

Ufuk Beyaztas^{a,b} and Aylin Alin^{a,*†}

In this study, we propose a delete-2 jackknife-after-bootstrap method to refine the cut-offs for the well-known diagnostic measure Cook's distance when the data have multiple influential data points with masking and swamping effects. The performance of the proposed method is compared with one of the most recent approaches in the literature through real world examples as well as a designed simulation study. Copyright © 2014 John Wiley & Sons, Ltd.

Keywords: Cook's distance; influential observation; jackknife-after-bootstrap; masking; swamping

1. Introduction

The detection and evaluation of influential observations and outliers are among critical aspects of any data analysis. There are lots of measures proposed to flag such observations. See Chatterjee and Hadi¹ for an excellent overview. Among them, single-case deletion diagnostics are very popular because they are easy to apply as well as being already available in most statistical software packages. The idea behind these measures is to compare a feature of the model fit obtained from the full data set with the one obtained from a reduced set not including a certain point. Observations whose removal causes major changes in the analysis are termed influential with respect to the feature of the model fit under consideration. In this paper, we only concentrate on Cook's distance because of its poor performance with respect to masking and swamping effects. The points with masking effect mask the influence of the other points while the ones with swamping effect cause good data points to be incorrectly diagnosed as influential. We may have only one data point masking the effect of another or multiple-point interactions, which can only be detected through deletion of multiple points that may cause a swamping effect. The masking problem is a topic of considerable recent research interest (See Martin *et al.*² and the references therein).

The assessment of whether the change in the model resulting from the inclusion/exclusion of each respective data point is major or not is usually based on cut-off values obtained from asymptotic approximations. However, the asymptotic approximations used for these diagnostic measures suffer both because the null distributions of these quantities are very complex and the approximations tend to be poor when sample sizes are small (See Martin and Roberts³ and Beyaztas and Alin⁴ for numerical evidence). As the computing power has exploded over the intervening decades, computer-intensive methods such as Efron's⁵ bootstrap have also been widely used in data analysis. In the present context, bootstrap method is used to estimate both the distribution of the diagnostic measure and the cut-off values. By bootstrap distribution, we mean nonparametric bootstrap distribution obtained by taking random samples from original sample with replacement, which makes the distribution sensitive to unusual observations. What is needed is a sampling distribution free from any effect of these observations. Hence, Efron's⁶ jackknife-after-bootstrap (JaB) technique has been proposed by Martin and Roberts³ and Beyaztas and Alin⁴ as a method for detecting influence in regression models through refining the cut-off values for the common influence diagnostics used in linear regression models. Later, Beyaztas and Alin⁷ proposed sufficient JaB, which reduces the computational burden by 30%.

The rationale of the JaB method is to generate the null sampling distribution for the measure under the null hypothesis that i th data point is not influential. The JaB sampling distribution is obtained from bootstrap samples where none of the values equal to the i th data point. Hence, the cut-off values are free of the influence of this particular data point. This method itself is based on single-case deletion, and it may not be successful under masking or/and swamping effects when it is used on any diagnostics. We propose Delete-2 JaB (D-2 JaB) method to form sampling distribution for Cook's distance. Through this approach, the sampling distribution as well as the cut-off values for Cook's distance will be free of the effect of any influential pair of points.

The linear regression model considered in our study is $Y = X\beta + \varepsilon$, where Y is an $n \times 1$ column vector of responses, X is a $n \times p$ ($p = k + 1$, where k is the number of covariates, and the model typically includes an intercept term) fixed full-rank design matrix, β is an $p \times 1$ vector

^aDepartment of Statistics, Dokuz Eylül University, Izmir, Turkey

^bDepartment of Statistics, Istanbul Medeniyet University, Istanbul, Turkey

*Correspondence to: Aylin Alin, Department of Statistics, Dokuz Eylül University, Izmir, Turkey.

†E-mail: aylin.alin@deu.edu.tr

of unknown parameters including intercept β_0 , and ε is an $n \times 1$ unobservable error vector having normal distribution with mean zero and variance σ^2 . Deleting a d ($d \geq 2$) group of observations requires $C_{n,d} = n!/[d!(n-d)!]$ different reduced samples with too much computational operation and time. Hence, we only focused on the $d=2$ (D-2) case, which is enough for our proposed method to flag masked observations.

Martin *et al.*² proposed a new approach to the use of single-case deletion diagnostics that involves applying these diagnostics to D-2 and D-3 jackknife replicates of the data. In the succeeding text is the single-case deleted Cook's distance formula for i th observation applied to the D-2 jackknife replicates when both k th and s th data points for $k=1, \dots, n-1$, $s=k+1, \dots, n$, $k \neq s$ are removed.

$$\hat{\theta}_{(i)}^{(k,s)} = \frac{\left(\hat{\beta}^{(k,s)} - \hat{\beta}_{(i)}^{(k,s)}\right)^T (X^T X) \left(\hat{\beta}^{(k,s)} - \hat{\beta}_{(i)}^{(k,s)}\right)}{p \hat{\sigma}^2} \quad (1)$$

$\hat{\beta}^{(k,s)}$ and $\hat{\beta}_{(i)}^{(k,s)}$, respectively, are the least squares estimates for the regression coefficient obtained from the data sets where (k,s) th and (k,s,i) th data points are missing. $\hat{\sigma}^2$ is the maximum likelihood estimate for the error variance. The cut-off point for the Cook's distance is chosen as 1 or the median value of corresponding F distribution. In our study, the observations are flagged as influential if its corresponding Cook's distance value is greater than 1. However, this approximation is ineffective to detect influential points under the presence of the points causing masking or/and swamping effect.

The following section details our D-2 JaB approach, and Section 3 includes numerical examples where our method is compared with the approach of Martin *et al.*²

2. Delete-2 jackknife-after-bootstrap method

Following Lemma is a generalization of the Efron's⁶ idea.

Lemma

A sample of size n from $z_1, z_2, \dots, z_{k-1}, z_{k+1}, \dots, z_{s-1}, z_{s+1}, \dots, z_n$ drawn at random with replacement has the same distribution as a bootstrap sample from $z_1, z_2, \dots, z_n$ in which none of the bootstrap values equal to z_k and z_s for $k=1, \dots, n-1$, $s=k+1, \dots, n$, $k \neq s$.

The proof of lemma follows directly Efron's⁶ Lemma 1. Our method requires about e^2 times the number of resamples for a regular bootstrap because the probability of a 'regular' resample not containing a pair of (k, s) th data points is, roughly, e^{-2} . For example, for any data set, $1000 \times e^2 \approx 7400$ resamples are required to obtain 1000 resamples without k th and s th data points. Then, these 1000 resamples are used to estimate the D-2 JaB sampling distribution of Cook's distance measure and the influence cut-offs. The rationale behind this approach is to generate a 'null' bootstrap distribution of Cook's distance under the hypothesis that the joint influence of k th and s th observations on the other observations is insignificant. Because these data points are not present in any of the resamples from which the D-2 JaB distribution is generated, the sampling distribution is free from their influence. The percentiles of this sampling distribution are then used as cut-offs to gauge how extreme the observed influence measure associated with that data point are in the context of that influence-free null distribution. The algorithm of D-2 JaB can be described briefly as follows:

Step 1. Fit the appropriate model to the original data set, and calculate the single-case deleted Cook's distance measure for i th data point, $\hat{\theta}_{(i)}$ for $i=1, 2, \dots, n$, for full data set as

$$\hat{\theta}_{(i)} = \frac{\left(\hat{\beta} - \hat{\beta}_{(i)}\right)^T (X^T X) \left(\hat{\beta} - \hat{\beta}_{(i)}\right)}{p \hat{\sigma}^2}$$

where $\hat{\beta}$ and $\hat{\beta}_{(i)}$ are the least squares estimates obtained from full set and reduced set without i th point, respectively.

Step 2. Draw B resamples with replacement from the original data set.

Step 3. Obtain a subset of the resamples within these not containing a specific pair of (k,s) th data points, say $(1,2)$.

Step 4. Calculate $\hat{\theta}_{(i)}^*$, $i=1, \dots, n$ as defined in Step 1 for each of these resamples, and collect these values into a single vector, say $\hat{\theta}^*$.

Step 5. Repeat Steps 3 and 4 for each pair of (k, s) for $k=1, \dots, n-1$, $s=k+1, \dots, n$, and $k \neq s$.

Step 6. Determine the suitable quantiles (say 2.5% and 97.5%) of this generated bootstrap distribution including approximately $nC_{n,2}B/e^2$ values.

Step 7. Compare these quantiles with each of the original $\hat{\theta}_{(i)}$ for $i=1, 2, \dots, n$ calculated in Step 1 to check if the points are influential or not.

For calculating Cook's distance based on the approach of Martin *et al.*,² we construct the jackknife sample by removing the (k, s) th ($k \neq s$) data points from the original data set and obtain the usual Cook's distance based on single-case deletion technique for remaining $n-2$ observations as in Equation (1). For this reduced data set, a point is flagged as influential if its Cook's distance value exceeds the cut-off value of 1.

Martin *et al.*² considered two different percentages to determine if any point is influential or not. The first one given in Equation (2) is calculated to see how unusual a data point appears to be.

$$dr = \frac{\#d}{C_{n-1,2}} \quad (2)$$

$\#d$ is the number of times that a certain data point is detected in $C_{n-1,2} = (n-1)!/[2!(n-3)!]$ different reduced data sets without a pair of (k, s) th data points. According to Martin *et al.*,² this overall percentage will typically be large for clearly influential, unmasked data points, but may remain small for an unusual point masked by another as the masking point will be present in many of the jackknife replicates as well. The second percentage is calculated to reveal points that are considered to be masked by another point. Martin *et al.*² considered $(n-2)$ jackknife replicates that do not contain the potentially masking data point in question so that they could use this as a measure of how unusual each data point is when the potentially masking point is excluded. If a low dr becomes significantly larger among replicates without a certain data point, this is considered to be an evidence of single-point masking.

3. Numerical study

We studied the performance of our proposed method with several data sets, but because of limited space, we can only present the results of designed simulation study and four well-known data sets, which are used in many papers in this context. All calculations have been performed under the assumption that the linear regression model is the correct model, and all the models satisfy the homoscedasticity assumption. For all of the real data examples, 8000 resamples were created from the original data set for each pair to have roughly 1000 resamples without the pair. The calculations were obtained using R 2.15 (see Appendix 1 for the R codes). The values in brackets in the Tables are the detection rates (dr) of flagged influential observations calculated as in Equation (2).

3.1. Real data examples

3.1.1. Life cycle savings data (Belsley *et al.*⁸, p. 41). The life cycle savings data for 50 countries are explained by four explanatory variables. According to Belsley *et al.*,⁸ the set includes influential observations determined by various diagnostic measures described in that monograph. Noted influential points correspond to countries such as Japan (point 23), Zambia (46), Libya (49), Canada (6), Chile (7), South Rhodesia (37), and the USA (44).

As visible from the normal quantile plot (Figure 1), some points seem unusual, but no point extends beyond the dotted boundaries suggesting that the normality assumption remains tenable. According to the regression influence plot presented in Figure 2, points 23, 46, 49, and 7 appear influential. In this plot, the areas of the circles represent the observations proportional to Cook's distance. It seems that there is no masking effect between the individual data points because there are no interlocking circles in this figure. On the other hand, discarding two observations at a time causes swamping effect for this data set. Points 47 and 48 are mistakenly flagged by D-2 jackknife Cook's distance while Cook's distance with the refined cut-offs by D-2 JaB method successfully flags three of the influential data points (Table I).

3.1.2. The Hertzsprung-Russell diagram of star cluster data (Rousseeuw and Leroy⁹, p. 27). For this data set, there is one explanatory variable to explain the logarithm of the light intensity of a star. Normality assumption is satisfied (Figure 3). According to Figure 4,

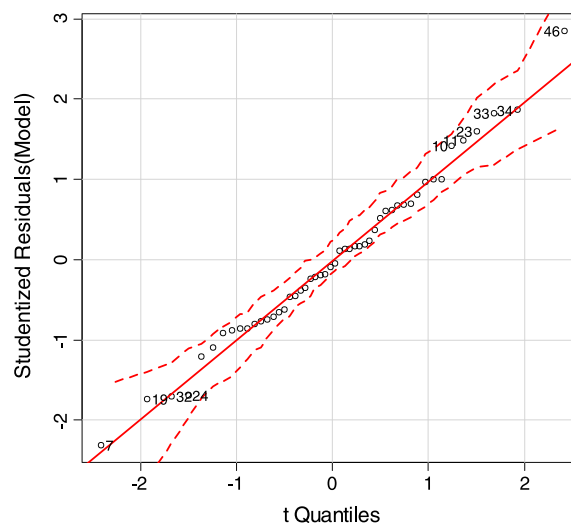


Figure 1. Normal quantile plot for life cycle savings data.

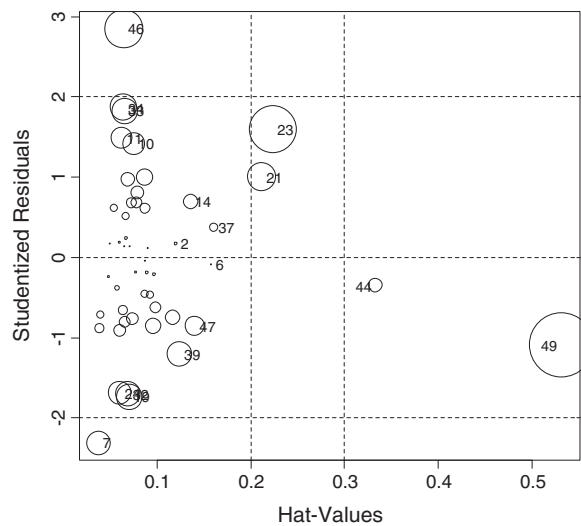


Figure 2. Influence plot for life cycle savings data.

Table I. Regression influence diagnostics for life cycle savings data, $n = 50$, $p = 5$		
Method		
D-2 jackknife Cook's distance	Cut-off	1.000
	Influential points	47 ($dr = 0.00387$) 48 ($dr = 0.00129$)
D-2 JaB Cook's distance	Cut-off	0.077
	Influential points	23, 46, 49

JaB: jackknife-after-bootstrap.

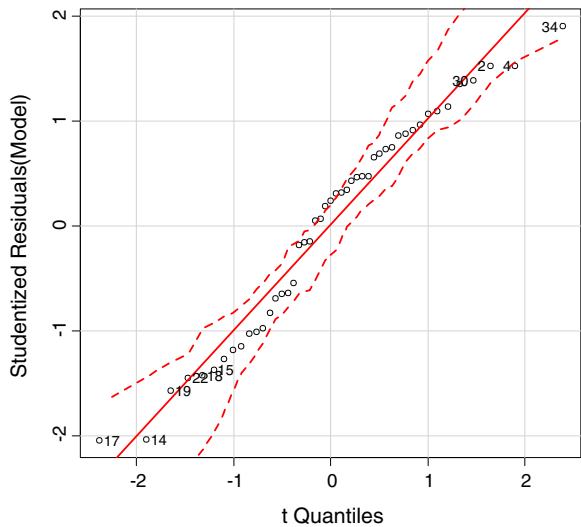


Figure 3. Normal quantile plot for Hertzprung-Russell diagram of star cluster data.

the data exhibit a pronounced masking effect for points 11, 20, 30, and 34, a problem well examined by Rousseeuw and Leroy⁹ (p. 230). Our results are presented in Table II. Cook's distance with cut-offs from D-2 JaB method detects points 14, 20, 30, and 34. However, its opponent not only misses these points but also flags some good ones as influential, such as point 13, because of the swamping effect.

3.1.3. Soil evaporation data. This data set given by Freund¹⁰ is also available in the *Teaching Demos* R package. There are 46 observations and 10 explanatory variables. According to Figures 5 and 6, points 2 and 33 distort the normality, and points 2, 31,

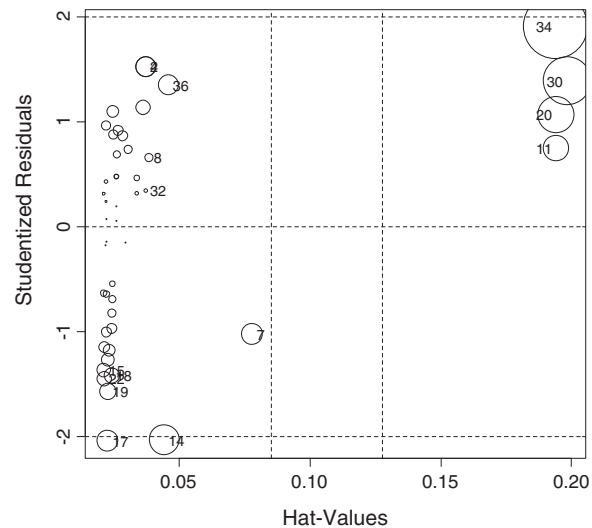


Figure 4. Influence plot for Hertzsprung-Russell diagram of star cluster data.

Table II. Regression influence diagnostics for Hertzsprung-Russell diagram of star cluster data, $n = 47$, $p = 2$		
Method		
D-2 jackknife Cook's distance	Cut-off	1.000
	Influential points	10 ($dr = 0.0058$)
		11 ($dr = 0.0395$)
		12 ($dr = 0.0761$)
		13 ($dr = 0.4216$)
D-2 JaB Cook's distance		14 ($dr = 0.9736$)
	Cut-off	0.085
	Influential points	14, 20, 30, 34

JaB: jackknife-after-bootstrap.

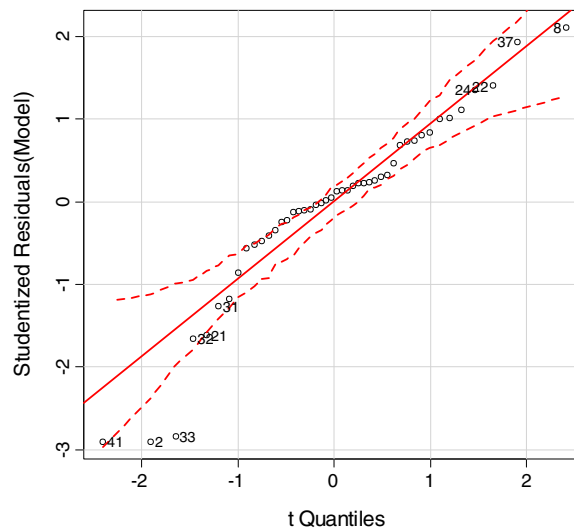


Figure 5. Normal quantile plot for soil evaporation data.

32, and 41 have potential to be influential. This data set also suffers from a masking problem for points 2 and 41 masking the impact of the other data points. Our results are presented in Table III. Masking seriously affects the D-2 jackknife Cook's distance so that it is unable to detect problematic data points. Nevertheless, proposed method is successful in detecting all of the suspicious points, overcoming the serious masking problem and not causing the swamping effect.

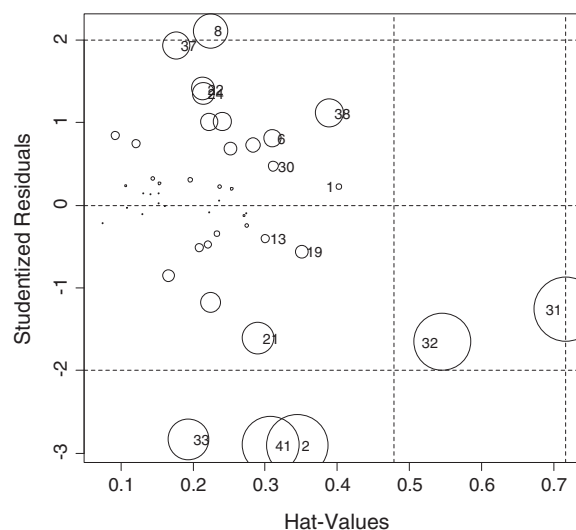


Figure 6. Influence plot for soil evaporation data.

Table III. Regression influence diagnostics for the soil evaporation data, $n = 46$, $p = 11$		
Method		
D-2 jackknife Cook's distance	Cut-off	1.000
	Influential points	30 ($dr = 0.0015$)
		32 ($dr = 0.0015$)
D-2 JaB Cook's distance	Cut-off	0.221
	Influential points	2, 31, 32, 41

JaB: jackknife-after-bootstrap.

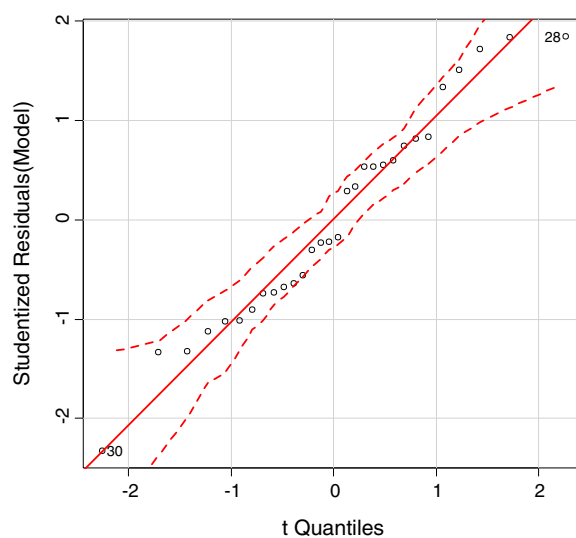


Figure 7. Normal quantile plot for health club data.

3.1.4. *Health club data* (Chatterjee and Hadi¹¹). The Health Club data consists of health records of 30 employees who are regular members of a given company's health club. For this data set, the normally distributed dependent variable (Figure 7) is modeled on four explanatory variables. Based on Figure 8, points 28 and 30 seem to be influential, which agree with the results of Chatterjee and Hadi¹¹ who suggest that point 28 is influential based on influence curve, point 30 is outlying in the residuals, and point 23 is outlying in the explanatory variables. Our results presented in Table IV show that unlike our proposed method, D-2 jackknife Cook's distance is unsuccessful to identify those points.

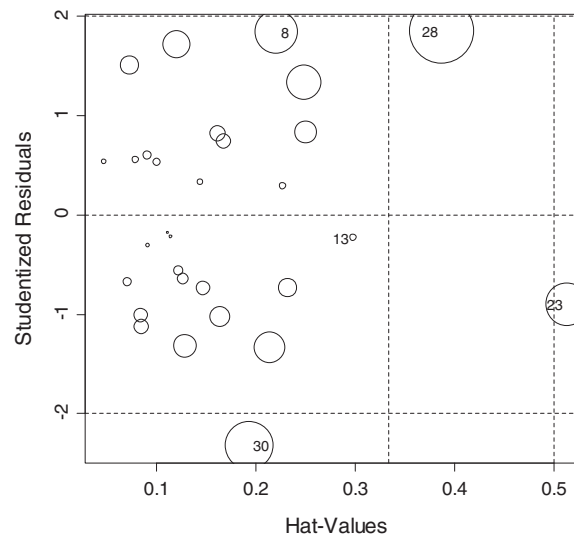


Figure 8. Influence plot for health club data.

Table IV. Regression influence diagnostics for the health club data, $n = 30$, $p = 5$		
Method		
D-2 jackknife Cook's distance	Cut-off	1.000
	Influential points	None
D-2 JaB Cook's distance	Cut-off	0.205
	Influential points	28, 30

JaB: jackknife-after-bootstrap.

Table V. Simulation results, $n = 20$, $p = 2$ with two influential observation for all distribution of errors				
Method		Distribution of errors		
		Normal	$t(3)$	Log-normal
D-2 jackknife Cook's distance	Cut-off	1.000	1.000	1.000
	Detection proportion for point 1	0.193	0.193	0.185
	Detection proportion for point 2	0.001	0.001	0.001
D-2 JaB Cook's distance	Cut-off	0.473	0.453	0.339
	Detection proportion for point 1	1.000	1.000	0.906
	Detection proportion for point 2	1.000	1.000	0.906

JaB: jackknife-after-bootstrap.

3.2. Simulation study

Simulation study is based on the design of Beyaztas *et al.*¹² where considered cases are $(n, p) = (20, 2)$ (small sample), $(n, p) = (50, 5)$ (large sample). The 'true' regression models are assumed as $Y = 1 + 2X_1 + 4X_2 + 3X_3 + 2X_4 + \varepsilon$ for the large sample cases and $Y = 1 + 2X + \varepsilon$ for the small sample cases, respectively. For each simulated case, X are generated as independent and identically distributed $N(2, 1)$ variates, and ε is generated with one of the following three error distributions: Normal ($N(0, 0.5625)$), $t_{(3)}$ (heavy-tailed), and centered lognormal ($1.5[\exp\{N(0, 0.5625)\} - \exp(\frac{1}{2})]$, skewed). For the small sample, the first two observations are replaced by the points $(x, y) = (5, 2)$. For the large sample, on the other hand, the points $\{(10, 10), (10, 10), (0, 0)\}$ take place of the first three data points. $M = 500$ simulations are runned for each scenario with the nominal significance level of $\alpha = 0.05$. The proportion of inserted influential data points detected in all 500 simulations is presented in Tables V and VI for the small and large samples, respectively, where the significant success of the proposed method under considered scenarios is obvious.

Table VI. Simulation results, $n = 50$, $p = 5$ with three influential observation for all distribution of errors				
Method		Distribution of errors		
		Normal	$t_{(3)}$	Log-normal
D-2 jackknife Cook's distance				
	Cut-off	1.000	1.000	1.000
	Detection proportion for point 1	0.532	0.532	0.471
	Detection proportion for point 2	0.448	0.396	0.445
	Detection proportion for point 3	0.024	0.027	0.011
D-2 JaB Cook's distance				
	Cut-off	0.154	0.137	0.154
	Detection proportion for point 1	0.996	0.996	0.992
	Detection proportion for point 2	0.988	0.996	1.000
	Detection proportion for point 3	0.964	0.996	1.000

JaB: jackknife-after-bootstrap.

4. Conclusion and discussion

In this study, we propose D-2 JaB method to refine the cut-offs for single case deleted Cook's distance under masking problem. The results for all the data sets and the simulation study show that with the refined cut-offs, which are robust against masking and swamping effects, Cook's distance successfully flags influential observations. The computing time gets longer as we increase the units to be deleted (d). However, for $d=2$ case, which is enough for the proposed method to detect influential observations, computing time is around 4.5 min. Finally, it should be noted that all the calculations in this paper are performed under the assumption that the linear model form is correct. Otherwise, as pointed out by a referee, the points flagged might not really be influential under the model that allows curvature.

We only worked on refining the cut-offs for Cook's distance because of its sensitivity to masking and swamping effects. However, the proposed method can also be used on other influence measures to improve their performance as well as being combined with other methods such as cluster-based bounded influence regression proposed by Lawrence *et al.*¹³

Acknowledgements

The authors are grateful to Tufan Alin for his suggestions on the algorithm having significant impact on the computation time and two anonymous referees for their constructive comments, which improved the previous version of the manuscript. The authors also thank the R Development Core Team (2009), for using R coding in the simulation study.

References

1. Chatterjee S, Hadi AS. Influential observations, high leverage points, and outliers in linear regression. *Statistical Science* 1986; **1**:379–416.
2. Martin MA, Roberts S, Zheng L. Delete-2 and delete-3 jackknife procedures for unmasking in regression. *Australian & New Zealand Journal of Statistics* 2010; **52**:45–60.
3. Martin MA, Roberts S. Jackknife-after-bootstrap regression influence diagnostics. *Journal of Nonparametric Statistics* 2010; **22**:257–269.
4. Beyaztas U, Alin A. Jackknife-after-bootstrap method for detection of influential observations in linear regression models. *Communications in Statistics – Simulation and Computation* 2013; **42**:1256–1267.
5. Efron B. Bootstrap methods: another look at the jackknife. *The Annals of Statistics* 1979; **7**:1–26.
6. Efron B. Jackknife-after-bootstrap standard errors and influence functions. *Journal of Royal Statistical Society* 1992; **54**:83–127.
7. Beyaztas U, Alin A. Sufficient jackknife-after-bootstrap method for detection of influential observations in linear regression models. *Statistical Papers* 2013. DOI: 10.1007/s00362-013-0548-4
8. Belsley DA, Kuh E, Welsch RE. Regression diagnostics. Wiley: New York, 1980.
9. Rousseeuw PJ, Leroy AM. Robust regression and outlier detection. Wiley: New York, 1987.
10. Freund RJ. Multicollinearity etc., some "new" examples. *Proceedings of the Statistical Computing Section* 1979; **4**:111–112.
11. Chatterjee S, Hadi AS. Sensitivity Analysis in Linear Regression. Wiley: USA, 1988.
12. Beyaztas U, Alin A, Martin MA. Robust BCa-JaB method as a diagnostic tool for linear regression models. *Journal of Applied Statistics* 2014; **41**:1593–1610.
13. Lawrence DE, Birch JB, Chen Y. Cluster-based bounded influence regression. *Quality and Reliability Engineering International* 2013. DOI: 10.1002/qre.1480

Appendix

R codes of Delete-2 Jackknife-after-Bootstrap method for Cook's distance

For Soil Evaporation data

library(TeachingDemos)

data(evap)

attach(evap)

y <- Evap

n <- length(y)

B <- 8000

alpha <- 0.05

X <- cbind(1, as.matrix(evap[, 4:13]))

Y <- matrix(y, ncol = 1)

index <- matrix(c(1:n), ncol = 1)

Design.data <- cbind(X, Y, index)

index.original <- c(index)

B.cap <- solve(crossprod(X)) %*% crossprod(X, Y)

P <- X %*% solve(crossprod(X)) %*% t(X)

Y.cap <- P %*% Y

e <- Y - Y.cap

dX <- nrow(X) - ncol(X)

var.cap <- crossprod(e) / (dX)

ei <- as.vector(Y - X %*% B.cap)

pi <- diag(P)

var.cap.i <- (((dX) * var.cap) / (dX - 1)) -

(ei^2 / ((dX - 1) * (1 - pi)))

ti <- ei / sqrt(var.cap * (1 - pi))

Ci <- (pi * ti^2) / (ncol(X) * (1 - pi))

boot.vector <- vector("list",)

boot.index <- vector("list",)

for (boot in 1:B) {

data <- Design.data[sample(n, n, replace = TRUE),]

dataX <- data[, 1:ncol(X)]

dataY <- data[, ncol(X) + 1]

index.bootstrap <- data[, ncol(X) + 2]

index.simulation <- c(index.bootstrap)

s.c.dX <- solve(crossprod(dataX))

B.cap.simulation <- s.c.dX %*% crossprod(dataX, dataY)

P.simulation <- dataX %*% s.c.dX %*% t(dataX)

Y.cap.simulation <- P.simulation %*% dataY

e.simulation <- dataY - Y.cap.simulation

dX.simulation <- nrow(dataX) - ncol(dataX)

var.cap.simulation <- crossprod(e.simulation) / (dX.simulation)

ei.simulation <- as.vector(dataY - dataX %*% B.cap.simulation)

pi.simulation <- diag(P.simulation)

var.cap.i.simulation <- (((dX.simulation) * var.cap.simulation) / (dX.simulation - 1)) - (ei.simulation^2 / ((dX.simulation - 1) * (1 - pi.simulation)))

ti.simulation <- ei.simulation / sqrt(var.cap.simulation * (1 - pi.simulation))

Ci.simulation <- (pi.simulation * ti.simulation^2) / (ncol(dataX) * (1 - pi.simulation))

boot.vector[[boot]] <- Ci.simulation

boot.index[[boot]] <- index.simulation

}

finalresult <- vector("list",)

result <- vector("list",)

for (j in 1:n-1) {

l = j + 1

for (k in l:n) {

for (i in 1:B) {

result[[i]] <- list(outCi.simulation = as.numeric(boot.vector[[i]]),

infl.u.obs = !(index.original[j] %in% as.numeric(boot.index[[i]])) & !(index.original[k] %in% as.numeric(boot.index[[i]])))

}

}

}

```
i.obs]]><- sapply(result, function(x) {x$infl.u.obs})
no.i.obs<- which(i.obs == TRUE)
data.no.i.obs<- result[no.i.obs]
D.2.JaB.data<- sapply(data.no.i.obs, function(x) {x$out.Ci.simulation})
finalresult<- c(finalresult, list(D.2.JaB.data))
}
data.JaB<- c(unlist(finalresult))
percentile.JaB<- quantile(data.JaB, 1-alpha)
which(Ci > percentile.JaB)
```

Authors' biographies

Ufuk Beyaztas is a PhD candidate in the Department of Statistics at Dokuz Eylul University, Turkey. He received both his BS and MS in Statistics from Dokuz Eylul University, Turkey. His research interest includes regression analysis and re-sampling methods.

Aylin Alin is an associate professor in the Department of Statistics at Dokuz Eylul University, Turkey, and is a visiting associate professor in the Department of Statistics and Probability at Michigan State University, USA. She received her MS in Applied Statistics from Michigan State University and her PhD in Statistics from Dokuz Eylul University. Her research interests are categorical data analysis, regression analysis, and re-sampling methods.