



University of Antwerp
| Faculty of Science

Cellwise outliers

From identification to regression

Jakob Raymaekers

Dept. of Mathematics, University of Antwerp, Belgium

Peter J. Rousseeuw

Dept. of Mathematics, KU Leuven, Belgium

DSSV-26 Conference, June 2026

Types of outliers

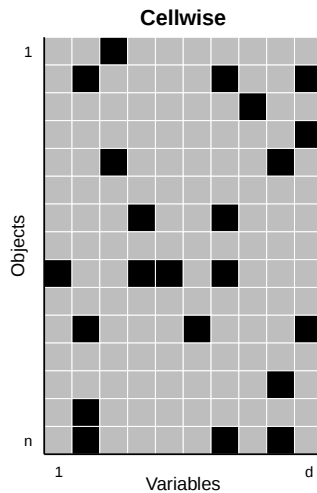
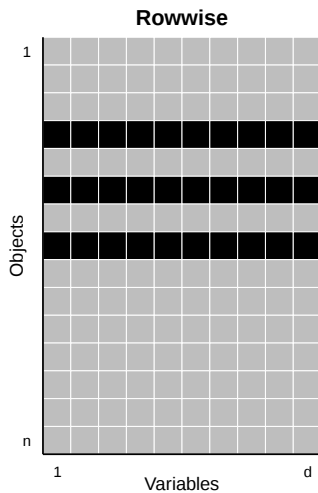
We want to analyze a data matrix $\mathbf{X}$ with n rows (cases) and $d > 1$ columns (variables). But the data may contain outliers.

The usual **rowwise contamination model** of Tukey (1960), Huber (1964),... assumes that some **rows** x_i have been replaced by arbitrary rows.

Rowwise robust methods downweight outlying rows and require at least 50% of clean rows.

The **cellwise contamination model** of Alqallaf, Van Aelst, Yohai and Zamar (2009 AOS) assumes that some **cells** x_{ij} have been replaced.

Downweighting/dropping an entire row loses a lot of information.



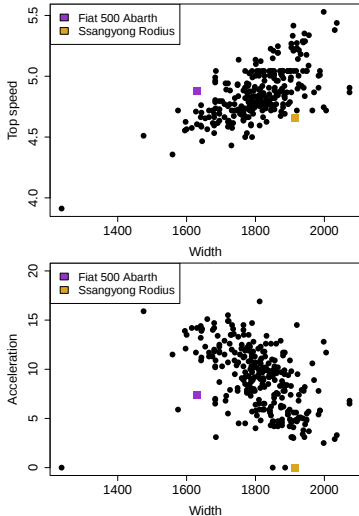
Detecting cellwise outliers

To **detect** outlying cells:

- By column only: based on the usual robust z-scores, or with the univariate Gervini-Yohai filter (2002 AOS)
- Modified Stahel-Donoho outlyingness (Van Aelst et al., 2012 CSDA)
- Bivariate filters, e.g. Leung et al (2017)
- Detecting Deviating Data Cells (Rousseeuw and Van den Bossche, Technometrics 2018).

The majority of approaches do not quantify cellwise outlyingness by the **residual from a robust fit**.

The challenge of identifying cellwise outliers



DetectDeviatingCells

	31.5	1968	177	280	8.2	143	61	1480	4701	1826	1477
Audi A4	31.5	1968	177	280	8.2	143	61	1480	4701	1826	1477
BMW i3	33.8	947	170	184	7.9	93	470	1315	3999	1775	1578
Chevrolet Cruze	16.7	1991	125	221	10.3	122	51	1427	4597	1788	1477
Corvette C6	68.5	6102	437	424	4.3	186	21	1460	4435	1844	1246
Fiat 500 Abarth	17.7	1368	160	169	7.4	131	43	1075	3660	1630	1490
Ford Kuga	26.8	1997	140	250	10.6	118	53	NA	4524	1838	1702
Honda CR-V	26.1	2199	150	258	9.7	118	50	1653	4570	1820	1685
Land Rover Defender	28.2	2198	122	265	14.7	90	25	2120	4785	1790	1790
Mazda CX-5	24.7	2191	150	280	9.4	122	54	1605	4555	1840	1710
Mercedes-Benz G	82.9	2987	211	398	9.1	108	25	2500	4662	1760	1951
Mini Coupe	19.8	1598	184	192	6.9	143	48	NA	3734	1683	1378
Peugeot 107	9.3	998	68	70	14.2	100	65	219	3430	1630	1465
Porsche Boxster	38.2	2706	265	206	5.8	164	34	1310	4374	1801	1282
Renault Clio	14.3	1461	90	162	12	112	88	1071	4062	1731	1448
Ssangyong Rodius	18	1998	155	265	0	105	38	2059	5125	1915	1845
Suzuki Jimmy	12	1328	84	81	14.1	87	39	1158	3665	1645	1705
Volkswagen Golf	20.1	1598	105	184	10.7	119	74	1295	4255	1799	1452

Price

Displacement

BHP

Torque

Acceleration

TopSpeed

MPG

Weight

Length

Width

Height

variables

Cellwise deviation with respect to a fit

CellHandler (Raymaekers and Rousseeuw 2021, JDSSV) was our first answer to the challenge of detecting outlying cells based on estimates $\hat{\mu}$ and $\hat{\Sigma}$.

The main idea is to find a sparse δ which shifts a few cells of an observation z to obtain $z - \delta$ which fits better with the model suggested by $\hat{\mu}$ and $\hat{\Sigma}$.

It is based on the identity

$$\text{MD}^2(z - \delta, \mu, \Sigma) = \|\tilde{Y} - \tilde{X}\delta\|_2^2,$$

with $\tilde{Y} := \Sigma^{-1/2}(z - \mu)$ and $\tilde{X} := \Sigma^{-1/2}$ with coefficient vector δ .

Finding such a sparse δ can now be done through solving a penalized regression!

Unfortunately, CellHandler does not integrate naturally into a covariance estimator with attractive properties, but it did give us direction.

Likelihood for missing data

Rather than shifting cells, we instead **select** cells. For incomplete data, denote the missingness pattern of the $n \times d$ data matrix $\mathbf{X}$ by the $n \times d$ matrix $\mathbf{W}$ with entries w_{ij} that are 0 for missing x_{ij} and 1 otherwise.

The **observed likelihood** of row i is defined as:

$$f(\mathbf{x}_i^{(\mathbf{w}_i)}, \boldsymbol{\mu}^{(\mathbf{w}_i)}, \boldsymbol{\Sigma}^{(\mathbf{w}_i)}) = \frac{1}{(2\pi)^{d(\mathbf{w}_i)/2} |\boldsymbol{\Sigma}^{(\mathbf{w}_i)}|^{1/2}} e^{-\text{MD}^2(\mathbf{x}_i, \mathbf{w}_i, \boldsymbol{\mu}, \boldsymbol{\Sigma})/2}$$

in which $\text{MD}(\mathbf{x}_i, \mathbf{w}_i, \boldsymbol{\mu}, \boldsymbol{\Sigma})$ is the **partial Mahalanobis distance**. Here

- $\mathbf{x}_i^{(\mathbf{w}_i)}$ is the vector with only the entries for which $w_{ij} = 1$
- $\boldsymbol{\Sigma}^{(\mathbf{w}_i)}$ is the submatrix of $\boldsymbol{\Sigma}$ for the variables j with $w_{ij} = 1$.
- $d(\mathbf{w}_i)$ is the dimension of $\mathbf{x}_i^{(\mathbf{w}_i)}$, i.e. the number of non-missing entries of $\mathbf{x}_i$.

Likelihood for missing data

The MLE of $(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ for the observed likelihood thus minimizes

$$\sum_{i=1}^n L(\mathbf{x}_i, \mathbf{w}_i, \boldsymbol{\mu}, \boldsymbol{\Sigma}) = \sum_{i=1}^n \left(\ln |\boldsymbol{\Sigma}^{(\mathbf{w}_i)}| + d(\mathbf{w}_i) \ln(2\pi) + \text{MD}^2(\mathbf{x}_i, \mathbf{w}_i, \boldsymbol{\mu}, \boldsymbol{\Sigma}) \right) .$$

This estimate is typically computed by the EM algorithm (Dempster et al., 1977).

When constructing a cellwise MCD, the matrix $\mathbf{W}$ is unknown and must be estimated!

Cellwise MCD

The **cellwise MCD** (**cellMCD**) finds $(\boldsymbol{\mu}, \boldsymbol{\Sigma}, \boldsymbol{W})$ that minimize the objective function

$$\sum_{i=1}^n L(\mathbf{x}_i, \mathbf{w}_i, \boldsymbol{\mu}, \boldsymbol{\Sigma}) + \sum_{j=1}^d q_j \|\mathbf{1}_d - \mathbf{W}_{.j}\|_0$$

under the constraints $\|\mathbf{W}_{.j}\|_0 \geq h$ for all $j = 1, \dots, d$ and $\lambda_d(\boldsymbol{\Sigma}) \geq a > 0$.

Here $\|\mathbf{1}_d - \mathbf{W}_{.j}\|_0$ is the number of zero weights in column j of $\boldsymbol{W}$, i.e. the number of cells in column j of $\mathbf{X}$ flagged as outlying.

From covariance to linear regression

Consider a standard linear regression model

$$\mathbf{y} = \alpha + \mathbf{X}\boldsymbol{\beta} + \text{noise}$$

Now both $\mathbf{y}$ and $\mathbf{X}$ can contain cellwise outliers.

We aim to:

- rely as little as possible on assumptions of ellipticity
- handle *test* data with potential outliers
- treat in-sample and out-of-sample in a uniform way

Simply cellMCD?

It may be tempting to simply treat $\mathbf{Z} = [\mathbf{y} ; \mathbf{X}]$ as a $(d + 1)$ -dimensional data matrix and apply cellMCD, i.e. find $(\boldsymbol{\mu}_Z, \boldsymbol{\Sigma}_Z, \mathbf{W})$ that minimize:

$$\sum_{i=1}^n L(\mathbf{z}_i, \mathbf{w}_i, \boldsymbol{\mu}_Z, \boldsymbol{\Sigma}_Z) + \sum_{j=1}^{d+1} q_j \|\mathbf{1}_d - \mathbf{w}_{\cdot j}\|_0,$$

and then recover $\hat{\alpha}, \hat{\beta}$ from $\hat{\boldsymbol{\Sigma}}_Z$. This is similar in spirit to the 3SGS estimator. Unfortunately:

- this strongly relies on ellipticity of $\mathbf{Z}$
- $\mathbf{w}_{ij}$ now depends on all $\mathbf{w}_{ik}$, $k \neq j$ hence the information in the response is used to flag outliers in the predictors. This cannot be replicated out-of-sample.

To address these shortcomings, we take two measures: **symmetrization** and a **two-step procedure**.

Symmetrization

For the **population** model $Y = \beta^\top \mathbf{X} + E$, and an independent copy $(\mathbf{X}', E', Y')$ of $(\mathbf{X}, E, Y)$, we have $Y - Y' = \beta^\top (\mathbf{X} - \mathbf{X}') + E - E'$. We thus obtain

$$\tilde{Y} = \beta^\top \tilde{\mathbf{X}} + \tilde{E},$$

with $\tilde{Y} := Y - Y'$, $\tilde{\mathbf{X}} := \mathbf{X} - \mathbf{X}'$ and $\tilde{E} := E - E'$.

By construction, we have $\mathbb{E}[\tilde{Y}] = \mathbb{E}[\tilde{E}] = 0$, $\mathbb{E}[\tilde{\mathbf{X}}] = \mathbf{0}$ and $\mathbb{E}[\tilde{\mathbf{X}} \tilde{\mathbf{X}}^\top] = 2 \mathbb{E}[\mathbf{X} \mathbf{X}^\top]$.

For **finite samples**:

We symmetrize $\mathbf{y} = \{y_1, \dots, y_n\}$ by $\text{Sym}(\mathbf{y}) = \{y_\ell - y_i; i \neq \ell\}$.

We also symmetrize $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$ to $\{\mathbf{x}_\ell - \mathbf{x}_i; i \neq \ell\}$.

To speed up computations, we can work with a subset of the pairwise differences.

cellTS: A two-step procedure

In *cellTS*, we execute 2 steps:

Step 1: Obtain a cleaned data matrix of regressors $\hat{X}$.

Step 2: Obtain robust estimates of α and β .

cellLTS: A two-step procedure

Step 1: Obtain a cleaned data matrix of regressors $\hat{\mathbf{X}}$.

- (a) Obtain $\hat{\Sigma}_{\mathbf{X}}$ through cellMCD with fixed center on $\text{Sym}(\mathbf{X})$ and dividing by 2.
- (b) Obtain $\hat{\mu}_{\mathbf{X}}$ by a robust univariate location estimate on each column of $\mathbf{X}$.
- (c) Obtain $\mathbf{W}_{\mathbf{X}}$ by iterations that optimize the cellMCD objective with respect to $\mathbf{W}_{\mathbf{X}}$, while keeping $\hat{\mu}_{\mathbf{X}}$ and $\hat{\Sigma}_{\mathbf{X}}$ fixed. The starting value is based on a marginal detection rule.
- (d) Obtain $\hat{\mathbf{X}}$ by

$$\hat{x}_{ij} = \begin{cases} x_{ij} & \text{if } W_{ij} = 1 \\ \hat{\mu}_j + \hat{\Sigma}_{j,(\mathbf{w}_{i\cdot})} \left(\hat{\Sigma}^{(\mathbf{w}_{i\cdot})} \right)^{-1} (\mathbf{x}_i^{(\mathbf{w}_{i\cdot})} - \hat{\mu}^{(\mathbf{w}_{i\cdot})}) & \text{if } W_{ij} = 0 . \end{cases}$$

Step 2: Obtain robust estimates of α and β .

cellTS: A two-step procedure

Step 1: Obtain a cleaned data matrix of regressors $\hat{\mathbf{X}}$.

Step 2: Obtain robust estimates of α and β .

(e) Symmetrize $\mathbf{Z} = [\mathbf{y} ; \hat{\mathbf{X}}]$ and standardize the result with the scales $(\hat{\sigma}_y, \sqrt{\text{diag}(\hat{\Sigma}_{\mathbf{X}})})$ to obtain $\tilde{\mathbf{Z}} = [\tilde{\mathbf{y}} ; \tilde{\mathbf{X}}]$.

(f) Both $\tilde{\mathbf{X}}$ and $\tilde{\mathbf{y}}$ have $\tilde{n} := n(n-1)$ rows, and we will consider $\tilde{h}$ -subsets $\tilde{H}$ of $\{1, \dots, \tilde{n}\}$ with $\tilde{h} = h(h-1)$ elements. The cellLTS estimator of β is now

$$\hat{\beta} = \underset{\beta(\tilde{H})}{\operatorname{argmin}} \left(\sum_{\ell \in \tilde{H}} r^2(\tilde{y}_{\ell}, \tilde{\mathbf{x}}_{\ell}, \beta) + \lambda \|\beta\|_2^2 \right)$$

with residuals $r(\tilde{y}_{\ell}, \tilde{\mathbf{x}}_{\ell}, \beta) := \tilde{y}_{\ell} - \tilde{\mathbf{x}}_{\ell}^{\top} \beta$.

(g) Estimate α by the univariate MCD of the pseudo residuals $\mathbf{r}^* = \mathbf{y} - \hat{\mathbf{X}} \hat{\beta}$.

(h) Transform the regression coefficients to undo the standardization in step (e).

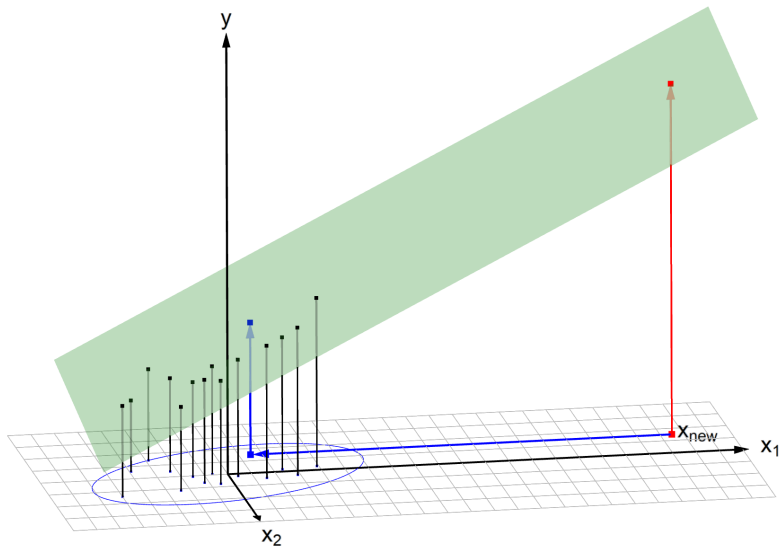
Out-of-sample predictions

We want to obtain the prediction at a new $\mathbf{x}$, which may itself contain missing values and/or outlying cells! Therefore we cannot just use $\hat{\alpha} + \mathbf{x}^\top \hat{\beta}$. Instead:

1. Obtain an initial $\mathbf{w}^{(0)}$ through the same marginal detection rule as in (c).
2. Obtain $\mathbf{w}$ by cellmcd iterations with fixed $\hat{\mu}_{\mathbf{x}}$ and $\hat{\Sigma}_{\mathbf{x}}$ as in (c)
3. Obtain $\hat{\mathbf{x}}$ as in (d)
4. Calculate $\hat{y} = \hat{\alpha} + \hat{\mathbf{x}}^\top \hat{\beta}$

Note: the imputation in a new point $\mathbf{x}$ starts from the same marginal detection rule that we used for $\mathbf{W}_{\mathbf{x}}$, we treat in-sample prediction and out-of-sample prediction consistently.

Out-of-sample predictions



Breakdown value

Let $\mathbf{Z} = [\mathbf{y} ; \mathbf{X}]$ as before and $\boldsymbol{\gamma} = (\alpha, \beta_1, \dots, \beta_d)^\top$ the combined parameters. Denote with $\mathbf{Z}^{\text{con}}$ a contaminated matrix obtained by replacing at most m cells in each column of $\mathbf{Z}$ by arbitrary values.

The *cellwise finite-sample breakdown value* of the regression estimator $\hat{\boldsymbol{\gamma}}$ at the dataset $\mathbf{Z}$ as

$$\varepsilon_n^*(\hat{\boldsymbol{\gamma}}, \mathbf{Z}) := \min \left\{ \frac{m}{n} : \sup_{\mathbf{Z}^{\text{con}}} \|\hat{\boldsymbol{\gamma}}(\mathbf{Z}^{\text{con}}) - \hat{\boldsymbol{\gamma}}(\mathbf{Z})\| = \infty \right\}$$

If the dataset $\mathbf{X}$ is in general position and $h \geq [(n + d + 1)/2]$, then the breakdown value of cellLTS with $\tilde{h} = h(h - 1)$ is

$$\varepsilon_n^*(\hat{\boldsymbol{\gamma}}, \mathbf{Z}) = (n - h + 1)/n.$$

Example: US Cancer dataset

Data on cancer rates in the period 2010-2016 in US counties, compiled from the U.S. Census Bureau and the National Cancer Institute.

The dataset contains 3047 rows, corresponding to the counties, and 33 columns.

The response variable is the death rate due to cancer (in cases per 100 000).

For illustrative purposes, we build a regression model using five covariates:

feature	explanation
incidencerate	Incidence rate of cancer (per 100 000).
medincome	Median income in the county (in thousands of dollars).
medianage	Median age in the county.
pcths18_24	Percentage of residents aged 18-24 whose highest attained education is a high school diploma.
pctemployed16_over	Percentage of people aged 16 and over who are employed.

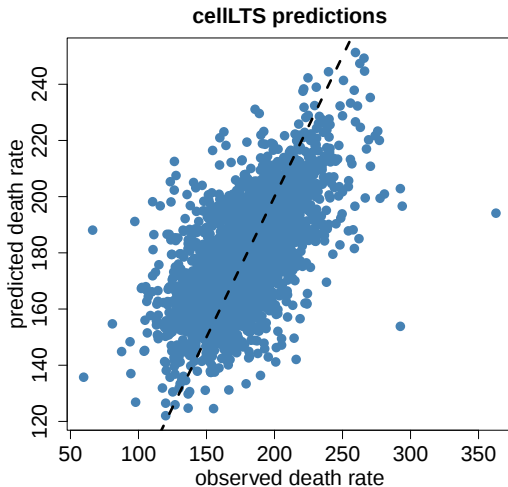
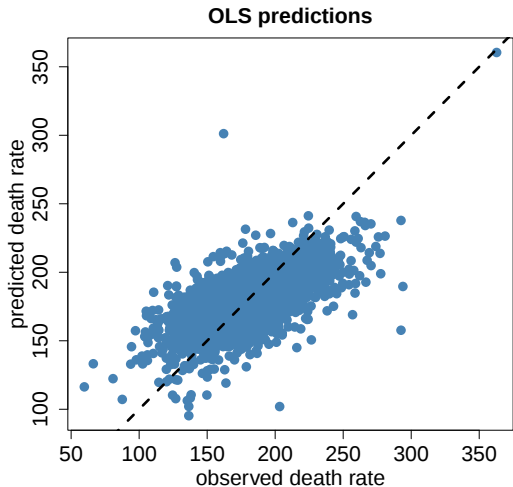
OLS vs. cellLTS on US Cancer data

	OLS	cellLTS	OLS of cleaned
intercept	122.50	157.16	147.27
incidencerate	0.23	0.24	0.23
medincome	-0.66	-0.75	-0.65
medianage	-0.01	-0.73	-0.61
pcths18_24	0.47	0.62	0.55
pctemployed16_over	-0.56	-0.83	-0.62

The cleaned OLS fit is obtained by removing the gross errors in the medianage variable (values larger than 400).

OLS vs. cellITS on US Cancer data

Predicted vs. observed death rate for OLS (left) and cellITS (right).



Cellmap of US Cancer data

Cellmap of the 10 counties with the highest sum of absolute standardized cellwise residuals:

cellmap of the most outlying counties

cases	Iosco County, Michigan	193	406	37	624	32	40
	Union County, Florida	363	1207	40	40	45	NA
	Brooks County, Georgia	159	469	33	497	36	45
	Alamance County, North Carolina	179	469	41	470	26	58
	Tangipahoa Parish, Louisiana	207	490	40	413	25	55
	Lake County, Oregon	139	397	40	580	54	46
	Williamsburg city, Virginia	162	1014	47	25	10	44
	Hall County, Texas	223	416	33	536	42	48
	Pittsylvania County, Virginia	177	399	44	546	41	54
	Knox County, Missouri	218	432	38	523	28	54
		death rate	incidence rate	median income	median age	pcths18_24	pctemployed16_over
		variables					

Conclusions

Implementation:

The software is in the R package *cellWise* on CRAN.

References:

R&R (2026). Least trimmed squares regression with missing values and cellwise outliers. *arXiv preprint* arXiv:2603.04632.

R&R (2026). Challenges of cellwise outliers. *Econometrics and Statistics*, 38, 6-25.

R&R (2024). The cellwise minimum covariance determinant estimator. *Journal of the American Statistical Association* 119.548, 2610-2621.

R&R (2021). Handling Cellwise Outliers by Sparse Regression and Robust Covariance. *Journal of Data Science, Statistics, and Visualisation*, 1.3.

Thank you!