
Reproducibility of Fast and Accurate Least-Mean-Squares Solvers

Anonymous Author(s)

Affiliation

Address

email

Abstract

Least-mean squares(LMS) solvers, one of the most elementary supervised learning algorithms, have been heavily used in solving various practical issues due to its high interpretability and nice mathematical closed-form solution. However, in many real world cases, computing the covariance matrix to produce such solution becomes impossible due to some reasons.

In this project, we investigated reproducibility of results from a paper submitted and accepted by 2019 Neural Information Processing Systems (NeurIPS), named *Fast and Accurate Least-Mean-Squares Solvers* [Maalouf et al., 2019], which proposes a novel method to derive a much smaller but maintainable covariance matrix without accuracy loss, based on a new time-efficient implementation of Caratheodory’s Theorem.

In our study, we first reproduce the tests’ results in the paper and examine the effect of the method. And our experiments on extension of the method reveals some property and limitations in dealing with new cases.

1 INTRODUCTION

1.1 Preliminary

Information explosion, firstly indicated and used by the Online Oxford English Dictionary in 1964 [1964], represents the rapid increase in the amount of published information and data. The problem of managing such an enormous amount of data turns to be a worldwide interest, also give birth to and inspire couples of interdisciplinary subjects. Thus, it is of huge significance to expand existing mathematical tools to apply into a big data issues, which will assist to solve a lot of practical problems in machine learning.

Least-Mean-Squares (LMS) solvers are the family of multiple optimization and regression models. containing Linear Regression, Principal Component Analysis, Lasso and Ridge Regression, Elastic Net and many more. LMS solvers have been extensively utilized in solving the classification and regression of practical big data problem, such as graph theory [Zhang and Rohe, 2018], spectral clustering [Peng et al., 2015] and many more.

Generally, there are two main approaches to a LMS solver. Like in a typical Linear Regression model:

$$y = w_0x_0 + w_1x_1 + \dots + w_nx_n = \sum_{m=0}^n w^T x$$

$$\text{LMSError}(w) = \frac{1}{n} \|y - W^T X\|^2$$

i. Normal equations (closed-form solution) $\hat{w} = (X^T X)^{-1} X^T y$

30 ii. Optimization algorithm (Gradient Descent, Stochastic Gradient Descent, Newton’s method,
31 etc.)

32 Compared with optimization algorithm like gradient descent, closed-form equation can provide an
33 "exact" solution. However, a conundrum has existed in this solution due to limit of computation
34 power in LMS solvers— matrix multiplication and inversion. Applying matrix multiplication, which
35 takes $O(m^2n)$ ($X \in \mathbb{R}^{m \times n}$), might result in a series of problems when the dataset is enormous.
36 (details in related work)

37 Here, we reproduced a fast caratheodory method with coresets and sketch approximation, which
38 is capable of reducing the size of dataset to a small subset. The covariance matrix from matrix
39 multiplication of new subset is equal to the original dataset. In this way, the closed-form solution is
40 able to deal with the big data issues, provide solid and reliable solutions. And we tested this method
41 with four LMS solvers — Linear/Ridge/Lasso Regression and Elastic Net.

42 1.2 Related Work

43 In the closed-form solution, the essential operation comes to the matrix multiplication.

44 **Issues in Covariance Matrix Computation** There are two mathematical methods to calculate this
45 covariance matrix: (1) Directly compute the outer product of each row in matrix X : $X^T X =$
46 $\sum_{i=0}^n a_i a_i^T$ or (2) factorization of matrix X by singular value decomposition or QR decomposition:
47 $X^T X = V D^2 V^T$. However, in method (1), each time we operate an addition in outer product,
48 numerical errors will be introduced and accumulated over time. The errors will increase to a
49 significant amount especially when 32-bit floating number data type is used in the algorithm, and
50 finally influence the whole accuracy of covariance matrix. A possible solution to this issue is
51 Cholesky decomposition [Pourahmadi, 2007], but it applies strict restriction on the data matrix
52 [Dostal et al., 2011], which makes it not useful in real cases. On the other hand, though method (2)
53 make us worry less about the issue of accuracy loss, we could not ignore both decomposition method
54 prolong the computation time by introducing additional matrix for factorization so that the hidden
55 constant behind $O(nd^2)$ would be pretty large.

56 **Caratheodory’s Theorem** first published on 1907, states that any point $a \in \mathbb{R}^d$ lies in the convex
57 hull of set P , can be written in convex combination of at most $d + 1$ points in P [Carathéodory,
58 1907]. This theorem indicates that if we have a dataset containing n data and d features (X is a $n \times d$
59 matrix and always $n \gg d$), it is possible to find a smaller $(d^2 + 1) \times d$ matrix S , whose covariance
60 matrix is the same as X . Since each row in $X^T X$ (n points in total) can be considered as a point in
61 $\mathbb{R}^{d^2}$, so $d^2 + 1$ points are required to represent these n points. In this way, both the numerical error
62 and time cost will significantly reduce. Unfortunately, it takes $O(nd^3)$ or $O(n^2 d^2)$ to calculate the
63 caratheodory’s set in LMS solvers [Binder et al., 2014], which is still not affordable when the dataset
64 is so huge.

65 **Coresets and Sketches for Approximation** A coreset is a reduced data set that can be used as proxy
66 to the full data set. Same algorithm runs on coreset is supposed to provide approximate result as on
67 full data set. A sketch is a compressed mapping of the full data set onto a data structure which is
68 easy to update or change data, also allow approximate result when certain queries are on the sketch
69 and full data set [Phillips, 2016]. So a matrix S is a coreset when its rows are a weighted subset of a
70 larger matrix X , or a sketch when each row of S can be written in a combination of rows in X .

71 An obvious advantage of coresets over sketches is that coresets is capable of preserving the sparsity of
72 the original large matrix, which will significantly reduce the computation time [Feldman et al., 2016].
73 Unfortunately, many new methods nowadays proposed to provide approximations to the covariance
74 matrix give rise to some multiplicative errors and eigenvalues of the reduced matrix could departs
75 evidently from the original matrix [Drineas et al., 2006]. The new method, based on Caratheodory’s
76 Theorem, discussed here addresses this problem and provide an accurate approximation via both
77 coresets over sketches that also maintains the ability to handling data streams.

78 1.3 Task Introduction

79 The paper provided an algorithm which:

- 80 *i.* Using a novel approach, which combine the coresets and sketches approximations together,
81 to discover the Caratheodory’s set only taking $O(\log n)$ calls to LMS solvers.
- 82 *ii.* The new matrix $S \in \mathbb{R}^{(d^2+1) \times d}$ based on the former step, is a weighted subsets of the input
83 data matrix $A \in \mathbb{R}^{n \times d}$. Plus, the covariance matrix of S and A are the same: $S^T S = A^T A$.
- 84 *iii.* Validating this S matrix with Linear/Ridge/Lasso Regression and Elastic Net.

85 Our task is to:

- 86 *i.* Examining the validity of this novel approach by reproducing the results of the experiments
87 mentioned in paper and comparing the values we obtain.
- 88 *ii.* Applying this method to new datasets and exploring the power of the method in dealing
89 with more complex problems.
- 90 *iii.* Extending and changing the parameters of the methods and situations it applies to, discussing
91 and testing the limitations and requirements for this method to be useful.
- 92 *iv.* In order to change the parameter α (regularization item) and obtain its effects on the
93 model, we need to re-implement partial model. (α is a pre-determined constant in default
94 implementation).
- 95 *v.* Re-implementing the part to create coresets and compare the performance with default
96 implementation from authors.

97 2 DATASET AND SETUP

98 2.1 Dataset for Test

99 All 3 sets of data for testing in the experiment are public available. (And they are the same in original
100 paper for reproducing test)

- 101 *i.* 3D Road Network (North Jutland, Denmark) [2015]. It contains 434874 instances with 4
102 types of information. We choose two — longitude and latitude to predict the height.
- 103 *ii.* Individual household electric power consumption Data Set [2017]. It contains 2075259
104 instances with 9 types of information. We choose two — global active power and global
105 reactive power to predict the voltage.
- 106 *iii.* House Sales in King County (USA) [2014-2015]. It contains 21614 instance with 21 types
107 of information. We choose eight — bedrooms, living area, lot size, floors, waterfront, above
108 area, basement area and built year to predict the house price.

109 The remaining datasets used are created by random number generation with the function provided
110 in package numpy. Since the focus here is examining the effect of time saving in LMS solver, it
111 wouldn’t be an issue to use synthetic data.

112 2.2 Data Preprocess

113 Errors, like NaN, are removed from the datasets. Data type casting is used to satisfy the input
114 requirements. For synthetic data, different sample sizes n and parameters’ values $\alpha \in \mathbb{A}$ and feature
115 size d is chosen in order to create more artificial data and investigate the power of this new method.

116 3 EXPERIMENT AND RESULTS

117 3.1 Environmental Setup

118 In this experiment, we first download code file named "Booster" provided by the authors of this paper
119 and find the datasets used in authors’ experiments on the Internet. The experiments with default
120 implementation is run on the 2018 MacBook Pro with processor - 2.3 GHz Intel Core i5 - and memory
121 - 16 GB 2133 MHz LPDDR3. And re-implementation of the same algorithm based on Caratheodory’s
122 Theorem is run on the Lenovo Thinkpad Yoga with Intel(R) Core(R) i7-4600U @ 2.10GHz processor
123 and 8GB DDR3 1600MHz memory. We all use CPython distribution.

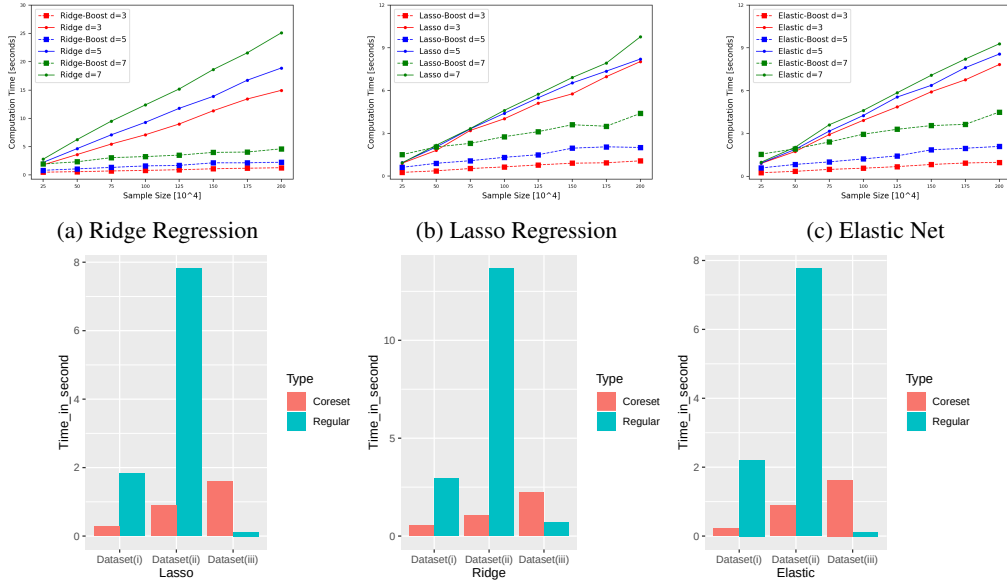
This limited the parameters' order of magnitudes we could test and the environment to examine effect of the method on streaming data is hard to set up. Whereas, it is enough for many available choices of the combination of parameters we could select to test, reproducing parts of its results and showing the power this novel approach bears in solving many issues.

3.2 Our Experiment

Our experiments are also based on four widely used LMS solvers Linear / Ridge / Elastic / Lasso-Regression.

3.2.1 Reproduction and Method Validity

We follow the steps from the paper, use existing datasets and construct new synthetic datasets with the same parameter. The result obtained from the LMS solvers with or without the boosting of this novel method is recorded. Also, departure in accuracy is collected in testifying the effect of the method on numerical stability. Then, we compare the result with the ones shown in the paper. Fig.1



(d) The runtime for three datasets with and without boosting when alpha is set to 100

Figure 1: Reproducibility tests on default method implementation

3.2.2 Hyperparameter and Extension

After carefully examining the content of the paper, it is not hard to imagine that there could be many cases different from the ones shown and used in the experiments displayed in the paper. In real world, situations might be more complex and so it is reasonable and intuitive to study the impact of different sample size n , feature size d and the number of alpha $|A|$ with the aim of justifying the method's flexibility and improvement.

The Fig.2 shows the method's influence when facing the cases with feature size d goes large with three LMS solver. Next, n is changed and we expect to find how large the n should be for this novel approach boosted regression model to out-compete the unvarnished one Fig.3.

3.2.3 Optimized Re-Implementation

After re-implementing algorithms detailed in the default paper, we have found some differences on performance between the re-implementation and the author's default implementation. Notably, we've found slight increase of performance, due to cleaner code. In Fig.4, we can notice a very slight

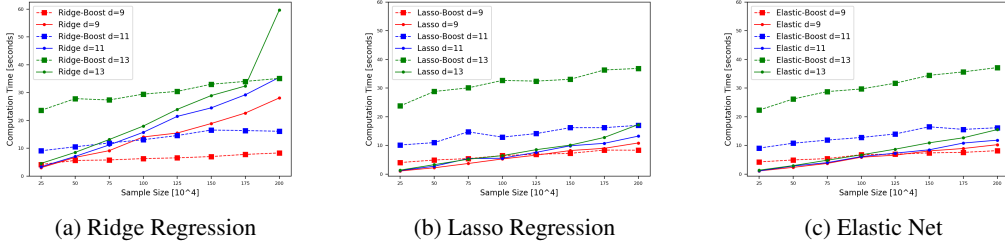


Figure 2: Time comparison of changing d on different regression models on default method implementation

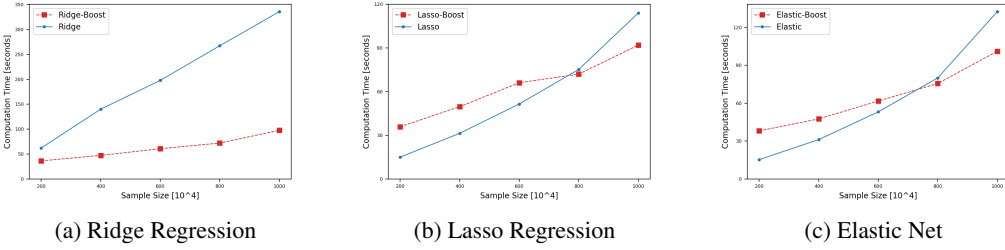


Figure 3: Time comparison of changing sample size n on different regression models on default method implementation

149 decrease in computation time, but due to hardware and small data samples, we also observe a increase
 150 in instability of the performance.

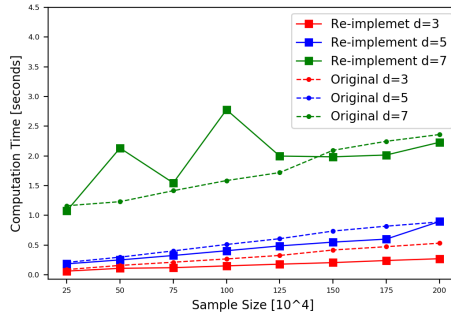


Figure 4: The effect of d on re-implement and original model

151 Additionally, our implementation allows us to tweak the value of k as describe in the algorithm
 152 formulation, whereas in the default implementation, k is fixed as $2d^2 + 2$. This let us explore the
 153 accuracy/speed trade-off by varying k as hinted in the paper. In Fig.5, we observe that computation
 154 time increase with the value of k . Note that values lower than $2d^2 + 2$ causes sometimes unexpected
 155 behaviour, such as non convergence of the algorithm. Thus we confirm that a value of $2d^2 + 2$ as
 156 theorized by the authors is indeed optimal.

157 3.3 Results

158 We successfully acquire the similar result displayed in the paper Fig.1 with the default implementation,
 159 which is a good indication for the time saving effect brought by this method. Also there is no accuracy
 160 loss, as a result of the merge-and-reduce property of this method. See Fig.6. The result is further
 161 testified by our re-implementation of the methods, which achieves even higher speed compared to the
 162 default implementation provided on the datasets without cross validation.

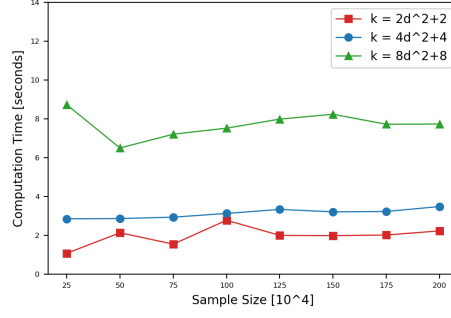


Figure 5: The effect of k on re-implement model. $d = 5$.

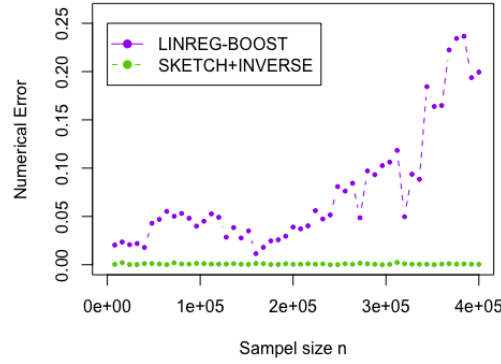


Figure 6: Numerical error on default method implementation

163 After finishing the experiments on new datasets new parameters, we discover that the boosting even
 164 exerts some negative influence as it is indicated in Fig.2 and Fig.3 suggests that in order for the
 165 boosting effect to be positive again, the demand of sample size increases dramatically.

166 4 DISCUSSION AND CONCLUSION

167 In this project, after we reproduced part of results shown in the paper, re-implemented the method
 168 codes, examined the robustness of this novel approach and investigated its limitations, we have many
 169 interesting discoveries.

170 The building up of coreset matrix with around d^2 rows within $O(nd^2)$ times does make acceleration
 171 possible for covariance matrix computation and there is no accuracy loss ignoring $\pm 10^{-16}$ Fig.6
 172 Applying coreset reveals its effectiveness even when the sample size is not large (Fig.1), but the
 173 effectiveness remains restricted when dealing with more complicated cases.

174 We originally intended to also experiment on the effect of regularization. α is the parameter that
 175 tunes the regularization on the ridge/lasso regression and elastic net models. The process to explore a
 176 well-performance α value is expected to optimize the LMS solvers models. However, we found that
 177 the original testing framework used by the authors selects alpha automatically based on results of
 178 cross-validation. We've attempted to create a new testing environment without success. We've found
 179 unexpected errors, caused by unclear details of the default implementation. We couldn't produce
 180 any interesting and consistent results, since the overall computation time increased and the models'
 181 accuracy decreased.

182 In real world, when adopting the linear regression or other LMS solvers, we could be confronted
 183 with cases where sample size is limited, where many features that might need to be preprocessed, or

both. And it is not rare for such cases to happen when we apply machine learning to problems in various fields such as advertising or bioinformatics [Saeys et al., 2007]. Cases where some features requires to be combined as a polynomial or simply where there is a need of slightly larger feature size to make accurate predictions would result in a rapid growth in the number of features and reduces this method's effectiveness, as it is indicated by Fig.2. Although the effect of a relatively large feature size can be offset by increasing the sample size Fig.3, the cost in collecting more samples or the circumstances under which we choose LMS solvers such as in problems like text classification [Sordo and Zeng, 2005] render this novel method impractical.

Notwithstanding, there exists many restrictions for this approach, it remains powerful when we have high demand for numerical accuracy on calculation results. Usually, when the feature size goes up, more computation time will be required. If the problem requires very low numerical error and computation time is not an concern, this approach would probably be our ideal choice [Xue et al., 2014]. In addition, there exists many dimension reduction techniques which can help, in addition of enough domain knowledge, reduce the feature size and fully utilise this approach. [Bharti and Singh, 2015].

5 Further Development

There are many potential improvements on this novel LMS booster method. For example, more efficient and complete implementation of the code is possible. Higher dimension data might introduce some other unexpected issues and must be considered carefully. And due to some limitations, the effect of this method with streamed or distributed data, or GPU acceleration requires more testing.

6 STATEMENT OF CONTRIBUTIONS

All members have made significant contributions to the project. The distribution of work for each member is described as follows:

Jiewen Liu : Model modification, experimentation, report writing.

Jianchen Zhao : Algorithm re-implementation, experimentation, report writing.

Zhenzhe Zhang : Publication discovering, data organization, report writing.

References

- Alaa Maalouf, Ibrahim Jubran, and Dan Feldman. Fast and Accurate Least-Mean-Squares Solvers. *arXiv e-prints*, art. arXiv:1906.04705, Jun 2019.
- information explodation. 1964. URL https://en.wikipedia.org/wiki/Information_explosion.
- Yilin Zhang and Karl Rohe. Understanding Regularized Spectral Clustering via Graph Conductance. *arXiv e-prints*, art. arXiv:1806.01468, Jun 2018.
- Xi Peng, Zhang Yi, and Huajin Tang. Robust subspace clustering via thresholding ridge regression, 2015. URL <https://www.aaai.org/ocs/index.php/AAAI/AAAI15/paper/view/9416>.
- Mohsen Pourahmadi. Cholesky Decompositions and Estimation of A Covariance Matrix: Orthogonality of Variance–Correlation Parameters. *Biometrika*, 94(4):1006–1013, 12 2007. ISSN 0006-3444. doi: 10.1093/biomet/asm073. URL <https://doi.org/10.1093/biomet/asm073>.
- Zdenek Dostal, Tomáš Kozubek, Alexandros Markopoulos, and Martin Mensík. Cholesky decomposition of a positive semidefinite matrix with known kernel. *Applied Mathematics and Computation*, 217:6067–6077, 03 2011. doi: 10.1016/j.amc.2010.12.069.
- C. Carathéodory. Über den variabilitätsbereich der koeffizienten von potenzreihen, die gegebene werte nicht annehmen. *Mathematische Annalen*, 64(1):95–115, Mar 1907. ISSN 1432-1807. doi: 10.1007/BF01449883.

228 Ilia Binder, Cristobal Rojas, and Michael Yampolsky. Computable carathéodory theory. *Ad-*
229 *vances in Mathematics*, 265:280 – 312, 2014. ISSN 0001-8708. doi: [https://doi.org/10.](https://doi.org/10.1016/j.aim.2014.07.039)
230 [1016/j.aim.2014.07.039](https://doi.org/10.1016/j.aim.2014.07.039). URL [http://www.sciencedirect.com/science/article/pii/](http://www.sciencedirect.com/science/article/pii/S0001870814002813)
231 [S0001870814002813](http://www.sciencedirect.com/science/article/pii/S0001870814002813).

232 Jeff M. Phillips. Coresets and Sketches. *arXiv e-prints*, art. arXiv:1601.00617, Jan 2016.

233 Dan Feldman, Mikhail Volkov, and Daniela Rus. Dimensionality reduction of massive sparse
234 datasets using coresets. pages 2766–2774, 2016. URL [http://papers.nips.cc/paper/](http://papers.nips.cc/paper/6596-dimensionality-reduction-of-massive-sparse-datasets-using-coresets.pdf)
235 [6596-dimensionality-reduction-of-massive-sparse-datasets-using-coresets.](http://papers.nips.cc/paper/6596-dimensionality-reduction-of-massive-sparse-datasets-using-coresets.pdf)
236 [pdf](http://papers.nips.cc/paper/6596-dimensionality-reduction-of-massive-sparse-datasets-using-coresets.pdf).

237 Petros Drineas, Michael W. Mahoney, and S. Muthukrishnan. Sampling algorithms for l2 regression
238 and applications. pages 1127–1136, 2006. URL [http://dl.acm.org/citation.cfm?id=](http://dl.acm.org/citation.cfm?id=1109557.1109682)
239 [1109557.1109682](http://dl.acm.org/citation.cfm?id=1109557.1109682).

240 Ryan A. Rossi and Nesreen K. Ahmed. The network data repository with interactive graph analytics
241 and visualization. In *AAAI*, 2015. URL <http://networkrepository.com>.

242 Dheeru Dua and Casey Graff. UCI machine learning repository, 2017. URL [http://archive.ics.](http://archive.ics.uci.edu/ml)
243 [uci.edu/ml](http://archive.ics.uci.edu/ml).

244 House sales in king county, usa, 2014-2015. URL [https://www.kaggle.com/harlfoxem/](https://www.kaggle.com/harlfoxem/housesalesprediction)
245 [housesalesprediction](https://www.kaggle.com/harlfoxem/housesalesprediction).

246 Yvan Saeys, Iñaki Inza, and Pedro Larrañaga. A review of feature selection techniques in bioinformat-
247 ics. *Bioinformatics*, 23(19):2507–2517, 08 2007. ISSN 1367-4803. doi: [10.1093/bioinformatics/](https://doi.org/10.1093/bioinformatics/btm344)
248 [btm344](https://doi.org/10.1093/bioinformatics/btm344). URL <https://doi.org/10.1093/bioinformatics/btm344>.

249 Margarita Sordo and Qing Zeng. On sample size and classification accuracy: A performance
250 comparison. In José Luís Oliveira, Víctor Maojo, Fernando Martín-Sánchez, and António Sousa
251 Pereira, editors, *Biological and Medical Data Analysis*, pages 193–201, Berlin, Heidelberg, 2005.
252 Springer Berlin Heidelberg. ISBN 978-3-540-31658-9.

253 J. Xue, J. Li, D. Yu, M. Seltzer, and Y. Gong. Singular value decomposition based low-footprint
254 speaker adaptation and personalization for deep neural network. In *2014 IEEE International*
255 *Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 6359–6363, May 2014.
256 doi: [10.1109/ICASSP.2014.6854828](https://doi.org/10.1109/ICASSP.2014.6854828).

257 Kusum Kumari Bharti and Pramod Kumar Singh. Hybrid dimension reduction by integrating feature
258 selection with feature extraction method for text clustering. *Expert Systems with Applications*, 42
259 (6):3105 – 3114, 2015. ISSN 0957-4174. doi: <https://doi.org/10.1016/j.eswa.2014.11.038>. URL
260 <http://www.sciencedirect.com/science/article/pii/S0957417414007301>.