

Additional file 1 - *Revisiting VERTIGO and VERTIGO-CI: Identifying confidentiality breaches and introducing a statistically sound, efficient alternative*

## Appendix A Notation

**Table A1** Glossary for general notation conventions

|                                                                                         |                                                       |                                                          |
|-----------------------------------------------------------------------------------------|-------------------------------------------------------|----------------------------------------------------------|
| Random variable in $\mathbb{R}$                                                         | $A$                                                   | Uppercase Non-italic                                     |
| Random vector in $\mathbb{R}^p$                                                         | $\mathbf{A}$                                          | Uppercase Non-italic bold                                |
| Scalar in $\mathbb{R}$                                                                  | $a$                                                   | Lowercase Italic                                         |
| Vector in $\mathbb{R}^p$                                                                | $\mathbf{a}$                                          | Lowercase Italic bold                                    |
| Vector in $\mathbb{R}^p$ with all components equal to 1                                 | $\mathbf{1}_p$                                        | -                                                        |
| Matrix in $\mathbb{R}^{n \times p}$                                                     | $\mathbf{A}$                                          | Uppercase Italic bold                                    |
| Identity matrix in $\mathbb{R}^{n \times n}$                                            | $\mathbf{I}_n$                                        | -                                                        |
| Gradient of $f(\boldsymbol{\theta})$ (column vector)                                    | $\nabla_{\boldsymbol{\theta}} f(\boldsymbol{\theta})$ | $\mathbf{H}$ for Hessian                                 |
| $\nabla_{\boldsymbol{\theta}} f(\boldsymbol{\theta}) _{\boldsymbol{\theta}=\mathbf{a}}$ | $\nabla_{\boldsymbol{\theta}} f(\mathbf{a})$          | $\mathbf{H}(\mathbf{a})$ for Hessian                     |
| $\max_{1 \leq j \leq p}  a_j $                                                          | $\ \mathbf{a}\ _{\infty}$                             | Infinite norm                                            |
| $\sum_{j=1}^p  a_j $                                                                    | $\ \mathbf{a}\ _1$                                    | $\ell_1$ -norm                                           |
| $\sqrt{\sum_{j=1}^p a_j^2}$                                                             | $\ \mathbf{a}\ _2$                                    | $\ell_2$ -norm                                           |
| Diagonal matrix with entries of $\mathbf{a}$ on diagonal                                | $\text{diag}(\mathbf{a})$                             | Dimension $p \times p$ for $\mathbf{a} \in \mathbb{R}^p$ |
| Quantity $\cdot$ at iteration $t$ (step count)                                          | $\cdot_{(t)}$                                         | Starts with $\cdot_{(0)}$                                |

**Table A2** Glossary for quantities that pertain to the regression settings

|                                                                    |                                                                                                                                                                                                                  |
|--------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Analytical dataset with sample size $n$                            | $\mathcal{D} = \{\dots\}_{i=1}^n$                                                                                                                                                                                |
| Covariate vector for $i$ th individual                             | $\mathbf{x}_i = [x_{i1}, \dots, x_{ip}]^{\top}$                                                                                                                                                                  |
| Covariate matrix in $\mathbb{R}^{n \times p}$                      | $\mathbf{X} = \begin{bmatrix} x_{11} & \dots & x_{1p} \\ \vdots & \ddots & \vdots \\ x_{n1} & \dots & x_{np} \end{bmatrix} = \begin{bmatrix} \mathbf{x}_1^{\top} \\ \vdots \\ \mathbf{x}_n^{\top} \end{bmatrix}$ |
| Gram matrix                                                        | $\mathbf{K} = [\mathbf{X} \mathbf{1}_n] [\mathbf{X} \mathbf{1}_n]^{\top}$                                                                                                                                        |
| True (unknown) parameters                                          | $\beta_{0*}, \boldsymbol{\beta}_*$                                                                                                                                                                               |
| Penalized estimate of the parameter (penalty parameter $\lambda$ ) | $\hat{\beta}_0^{\lambda}, \hat{\boldsymbol{\beta}}^{\lambda}$                                                                                                                                                    |
| Penalized estimate of the parameter (intercept not penalized)      | $\hat{\beta}_{0,\text{alt}}^{\lambda}, \hat{\boldsymbol{\beta}}_{\text{alt}}^{\lambda}$                                                                                                                          |
| Log-likelihood                                                     | $\ell_n(\beta_0, \boldsymbol{\beta}) = \frac{1}{n} \sum_{i=1}^n \log(\cdot)$                                                                                                                                     |
| Penalized log-likelihood                                           | $l_n^{\lambda}(\beta_0, \boldsymbol{\beta})$                                                                                                                                                                     |
| Penalized log-likelihood (intercept not penalized)                 | $l_{n,\text{alt}}^{\lambda}(\beta_0, \boldsymbol{\beta})$                                                                                                                                                        |

**Table A3** Glossary for quantities specific to the vertical setting

|                                                       |                                         |
|-------------------------------------------------------|-----------------------------------------|
| Number of covariate-nodes                             | $K$                                     |
| Number of covariates at covariate-node $k$            | $p^{(k)}$                               |
| Covariate matrix at covariate-node $k$                | $\mathbf{X}^{(k)}$                      |
| Dual parameter estimates                              | $\hat{\boldsymbol{\alpha}}^{\lambda}$   |
| Penalized estimate associated with covariate-node $k$ | $\hat{\boldsymbol{\beta}}^{\lambda(k)}$ |

## 047 Appendix B VERTIGO and VERTIGO-CI

### 048 Algorithms

051  
 052 Algorithm 3 provides the original VERTIGO algorithm [1], while algo-  
 053 rithm 4 provides the original VERTIGO-CI algorithm [2]. Some adjust-  
 054 ments have been made to the representation of the original algorithms  
 055 using the harmonized notation to facilitate the comparison. We note that  
 056 constrained optimization technique should be considered when numeri-  
 057 cally applying the Newton-Raphson steps because the components of the  
 058 parameter  $\alpha$  are constrained in  $(0, 1)$ .  
 059  
 060  
 061  
 062  
 063  
 064

---

#### 066 Algorithm 3 VERTIGO Original [1]

---

067 **Input** Local data in each node  $\mathbf{X}^{(k)}$ ,  $k = 1, \dots, K$ , shared response  $\mathbf{y}$ , parameter  $\lambda$ ,  
 068 initial solution  $\alpha_{(s)}$ ,  $s = 0$   
 069 **Output** Global solution  $\hat{\alpha}^\lambda$   
 070 **Procedure:**  
 071  
 072 1. Node  $k$ : compute the local matrix  $\mathcal{K}^{(k)}$  and send  $\mathcal{K}^{(k)}$  to CC.  
 073 2. CC: compute the global matrix  $\mathcal{K}$ .  
 074 3. Do:  
 075 (a) Node  $k$ : compute  $\mathbf{e}^{(k)}(\alpha_{(s)})$ , and send  $\mathbf{e}^{(k)}(\alpha_{(s)})$  to CC.  
 076 (b) CC: calculate  $\mathbf{e}_{(s)} = \sum_{k=1}^K \mathbf{e}^{(k)}(\alpha_{(s)})$  and  $\nabla_{\alpha} J^\lambda(\alpha_{(s)})$  (using Equation (10)).  
 077 (c) CC: compute the Hessian matrix  $\mathbf{H}(\alpha_{(s)})$  (using Equation (11)).  
 078 (d) CC: update  $\alpha_{(s)} \leftarrow \alpha_{(s)} - \mathbf{H}^{-1}(\alpha_{(s)}) \nabla_{\alpha} J^\lambda(\alpha_{(s)})$  and send update  $\alpha_{(s)}$  to  
 079 each node.  
 080  $s \leftarrow s + 1$ , until  $\alpha_{(s)}$  converges.  
 081 4.  $\hat{\alpha}^\lambda \leftarrow \alpha_{(s)}$ .  
 082  
 083  
 084  
 085  
 086

---

087  
 088  
 089  
 090  
 091  
 092

---

**Algorithm 4** VERTIGO-CI Original [2]

---

- Input** Local data in each node  $\mathbf{X}^{(k)}$ ,  $k = 1, \dots, K$ , shared response  $\mathbf{y}$ , parameter  $\lambda$ ,  
initial solution  $\boldsymbol{\alpha}_{(s)}$ ,  $s = 0$
- Output** Global solution  $\hat{\boldsymbol{\beta}}^\lambda$  and associated variances.
- Procedure:**
1. Node  $k$ : compute the local matrix  $\mathbf{K}^{(k)}$  and send  $\mathbf{K}^{(k)}$  to CC.
  2. CC: compute the global matrix  $\mathbf{K}$ .
  3. Do:
    - (a) Node  $k$ : compute  $\mathbf{e}^{(k)}(\boldsymbol{\alpha}_{(s)})$ , and send  $\mathbf{e}^{(k)}(\boldsymbol{\alpha}_{(s)})$  to CC.
    - (b) CC: calculate  $\mathbf{e}_{(s)} = \sum_{k=1}^K \mathbf{e}^{(k)}(\boldsymbol{\alpha}_{(s)})$  and  $\nabla_{\boldsymbol{\alpha}} J^\lambda(\boldsymbol{\alpha}_{(s)})$  (using Equation (10)).
    - (c) CC: compute the modified Hessian matrix  $\mathbf{H}(\boldsymbol{\alpha}_{(s)}) = \lambda^{-1} \text{diag}(\mathbf{y})\mathbf{K} \text{diag}(\mathbf{y}) + C\mathbf{I}_n$ .
    - (d) CC: update  $\boldsymbol{\alpha}_{(s)} \leftarrow \boldsymbol{\alpha}_{(s)} - \mathbf{H}^{-1}(\boldsymbol{\alpha}_{(s)})\nabla_{\boldsymbol{\alpha}} J^\lambda(\boldsymbol{\alpha}_{(s)})$  and send update  $\boldsymbol{\alpha}_{(s)}$  to each node.
  - $s \leftarrow s + 1$ , until  $\boldsymbol{\alpha}_{(s)}$  converges.
  4. CC:  $\hat{\boldsymbol{\alpha}}^\lambda \leftarrow \boldsymbol{\alpha}_{(s)}$ .
  5. CC: send  $\hat{\boldsymbol{\alpha}}^\lambda$  to each node.
  6. Node  $k$ ,  $k \leq K - 1$ : Calculate  $\hat{\boldsymbol{\beta}}^{\lambda(k)} = \lambda^{-1} \sum_{i=1}^n \hat{\alpha}_i^\lambda y_i \mathbf{x}_i^{(k)}$ , and send  $\hat{\boldsymbol{\beta}}^{\lambda(k)}$  to CC.
  7. Node  $K$ : Calculate  $\hat{\beta}_0^\lambda = \lambda^{-1} \sum_{i=1}^n \hat{\alpha}_i^\lambda y_i$  and  $\hat{\boldsymbol{\beta}}^{\lambda(K)} = \lambda^{-1} \sum_{i=1}^n \hat{\alpha}_i^\lambda y_i \mathbf{x}_i^{(K)}$ , and send  $\hat{\beta}_0^\lambda$  and  $\hat{\boldsymbol{\beta}}^{\lambda(K)}$  to CC.
  8. Client-to-client communication:
    - (a) Node  $k$ ,  $k \leq K - 1$ : compute  $\exp\{(\mathbf{x}_i^{(k)})^\top \hat{\boldsymbol{\beta}}^{\lambda(k)}\}$  for all  $i \in \{1, \dots, n\}$  and send result to node 1.
    - (b) Node  $K$ : compute and send  $\exp\{\hat{\beta}_0^\lambda\} \exp\{(\mathbf{x}_i^{(K)})^\top \hat{\boldsymbol{\beta}}^{\lambda(K)}\}$  for all  $i \in \{1, \dots, n\}$  and send result to node 1.
    - (c) Node 1: combine the results using  $\exp\{\hat{\beta}_0^\lambda + \mathbf{x}_i^\top \hat{\boldsymbol{\beta}}^\lambda\} = \exp\{\hat{\beta}_0^\lambda\} \prod_{k=1}^K \exp\{(\mathbf{x}_i^{(k)})^\top \hat{\boldsymbol{\beta}}^{\lambda(k)}\}$  for all  $i \in \{1, \dots, n\}$ , compute  $\mathbf{V}^\lambda$  (using Equation (12)) and send  $\mathbf{V}$  to all nodes  $k$ ,  $k \in \{2, \dots, K\}$ .
  9. Node  $k$ ,  $k \leq K - 1$ : calculate  $(\mathbf{X}^{(k)})^\top (\mathbf{V}^\lambda)^{1/2}$  and send this quantity to CC.
  10. Node  $K$ : calculate  $[\mathbf{X}^{(K)} \mathbf{1}_n]^\top (\mathbf{V}^\lambda)^{1/2}$  and send this quantity to CC.
  11. CC: build  $[\mathbf{X} \mathbf{1}_n]^\top \mathbf{V}^\lambda [\mathbf{X} \mathbf{1}_n]$  by block (see Algorithm 2 in [2] for details), and invert the result to obtain variance estimates.
-

## Appendix C Equivalence between matrix and expanded notation in dual optimization

To obtain the expression of  $J^\lambda(\boldsymbol{\alpha}) = \frac{1}{2\lambda} \boldsymbol{\alpha}^\top \text{diag}(\mathbf{y}) \mathcal{K} \text{diag}(\mathbf{y}) \boldsymbol{\alpha} + \sum_{i=1}^n L(\alpha_i)$ , it suffices to note that

$$\begin{aligned} \boldsymbol{\alpha}^\top \text{diag}(\mathbf{y}) \mathcal{K} \text{diag}(\mathbf{y}) \boldsymbol{\alpha} &= \begin{bmatrix} \alpha_1 y_1 & \cdots & \alpha_n y_n \end{bmatrix} \begin{bmatrix} \mathbf{x}_1^\top \mathbf{x}_1 + 1 & \cdots & \mathbf{x}_1^\top \mathbf{x}_n + 1 \\ \vdots & \ddots & \vdots \\ \mathbf{x}_n^\top \mathbf{x}_1 + 1 & \cdots & \mathbf{x}_n^\top \mathbf{x}_n + 1 \end{bmatrix} \begin{bmatrix} \alpha_1 y_1 \\ \vdots \\ \alpha_n y_n \end{bmatrix} \\ &= \begin{bmatrix} \sum_{i=1}^n \alpha_i y_i (\mathbf{x}_i^\top \mathbf{x}_1 + 1) & \cdots & \sum_{i=1}^n \alpha_i y_i (\mathbf{x}_i^\top \mathbf{x}_n + 1) \end{bmatrix} \begin{bmatrix} \alpha_1 y_1 \\ \vdots \\ \alpha_n y_n \end{bmatrix} \\ &= \sum_{j=1}^n \sum_{i=1}^n \alpha_i y_i (\mathbf{x}_i^\top \mathbf{x}_j + 1) \alpha_j y_j \\ &= \sum_{j=1}^n \sum_{i=1}^n \alpha_i \alpha_j y_i y_j (\mathbf{x}_i^\top \mathbf{x}_j + 1). \end{aligned}$$

## Appendix D Mathematical development for the dual problem of the ridge logistic regression maximum log-likelihood problem with no penalty over the intercept

**Proposition 1** Consider any  $\mathcal{D} = \{(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n)\}$ , with  $y_i \in \{-1, 1\}$ ,  $\mathbf{x}_i \in \mathbb{R}^p$  for  $i = 1, \dots, n$  and  $n \geq 2$ . Assume that  $\sum_{i=1}^n \mathbb{I}(y_i = 1) \in \{1, \dots, n-1\}$  (i.e., there exists at least one  $i$  such that  $y_i = 1$ , and at least one  $j$  such that  $y_j = -1$ ). Then, for any  $\lambda > 0$ ,

there exists a unique solution  $(\hat{\beta}_{0,\text{alt}}^\lambda, \hat{\beta}_{\text{alt}}^\lambda)$  to the optimization problem

$$\max_{\beta_0 \in \mathbb{R}, \beta = [\beta_1, \dots, \beta_p]^\top \in \mathbb{R}^p} l_{n,\text{alt}}^\lambda(\beta_0, \beta). \quad (\text{D1})$$

Furthermore,  $(\hat{\beta}_{0,\text{alt}}^\lambda, \hat{\beta}_{\text{alt}}^\lambda)$  satisfies

$$\hat{\beta}_{0,\text{alt}}^\lambda = \frac{1}{n} \sum_{i=1}^n (y_i \log\{(\hat{\alpha}_{i,\text{alt}}^\lambda)^{-1} - 1\} - \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda),$$

$$\hat{\beta}_{\text{alt}}^\lambda = \lambda^{-1} \mathbf{X}^\top \text{diag}(\mathbf{y}) \hat{\alpha}_{\text{alt}}^\lambda,$$

where  $\hat{\alpha}_{\text{alt}}^\lambda = [\hat{\alpha}_{1,\text{alt}}^\lambda, \dots, \hat{\alpha}_{n,\text{alt}}^\lambda]^\top$  is the unique solution to the following minimization problem:

$$\min_{\alpha \in (0,1)^n} J_{\text{alt}}^\lambda(\alpha) \quad \text{s.t.} \quad \sum_{i=1}^n y_i \alpha_i = 0.$$

*Proof of Proposition 1.* We begin by showing that the optimization problem in (D1) has a unique finite solution. To do this, we first compute that

$$\nabla_{\beta} l_{n,\text{alt}}^\lambda(\beta_0, \beta) = \sum_{i=1}^n \frac{y_i \mathbf{x}_i}{1 + \exp\{y_i(\beta_0 + \mathbf{x}_i^\top \beta)\}} - \lambda \beta,$$

and that

$$\frac{\partial}{\partial \beta_0} l_{n,\text{alt}}^\lambda(\beta_0, \beta) = \sum_{i=1}^n \frac{y_i}{1 + \exp\{y_i(\beta_0 + \mathbf{x}_i^\top \beta)\}}.$$

From this, we obtain that

$$\begin{aligned} & \nabla_{\beta, \beta_0}^2 l_{n,\text{alt}}^\lambda(\beta_0, \beta) \\ &= \sum_{i=1}^n \frac{-\exp\{y_i(\beta_0 + \mathbf{x}_i^\top \beta)\}}{(1 + \exp\{y_i(\beta_0 + \mathbf{x}_i^\top \beta)\})^2} \begin{bmatrix} \mathbf{x}_i \mathbf{x}_i^\top & \mathbf{x}_i \\ \mathbf{x}_i^\top & 1 \end{bmatrix} - \lambda \begin{bmatrix} \mathbf{I}_p & 0 \\ 0 & 0 \end{bmatrix}. \end{aligned}$$

Hence,  $l_{n,\text{alt}}^\lambda(\beta_0, \beta)$  is strictly concave in  $(\beta_0, \beta)$ . Therefore, if  $l_{n,\text{alt}}^\lambda(\beta_0, \beta)$  admits a maximizer, it must be unique. To conclude that the optimization problem in (D1) has a unique solution, it remains to show that  $l_{n,\text{alt}}^\lambda(\beta_0, \beta)$  indeed has a maximizer.

To establish this, it suffices to show that  $l_{n,\text{alt}}^\lambda(\beta_0, \beta) \rightarrow -\infty$  as  $\max\{|\beta_0|, \|\beta\|_\infty\} \rightarrow \infty$ , where we denote the infinite norm  $\|\mathbf{a}\|_\infty := \max_{1 \leq j \leq p} |a_j|$  for  $\mathbf{a} \in \mathbb{R}^p$ . Observe that the log-likelihood term is bounded above by 0. Therefore, if  $\|\beta\|_\infty \rightarrow \infty$ , then regardless of whether  $|\beta_0| \rightarrow \infty$  or  $|\beta_0| < \infty$ , the penalty term  $\lambda \beta^\top \beta$  dominates, and we have  $l_{n,\text{alt}}^\lambda(\beta_0, \beta) \rightarrow -\infty$ .

Now suppose  $\|\beta\|_\infty \not\rightarrow \infty$  but  $\max\{|\beta_0|, \|\beta\|_\infty\} \rightarrow \infty$ , which implies that  $|\beta_0| \rightarrow \infty$  while  $\|\beta\|_\infty$  remains bounded. Because there is at least one  $y_i = 1$  and at least one  $y_j = -1$ , there exists at least one term in the log-likelihood of the form  $\log\left(\left[1 + \exp\{-y_i(\beta_0 + \mathbf{x}_i^\top \beta)\}\right]^{-1}\right)$  that diverges to  $-\infty$  as  $|\beta_0| \rightarrow \infty$  (since  $\mathbf{x}_i^\top \beta$  remains bounded), while the others remain bounded from above by 0. Therefore,  $l_{n,\text{alt}}^\lambda(\beta_0, \beta) \rightarrow -\infty$  in this case as well.

The conclusion of the above discussion is that for any  $\lambda > 0$ , there exists a unique solution  $(\hat{\beta}_{0,\text{alt}}^\lambda, \hat{\beta}_{\text{alt}}^\lambda)$  to the optimization problem in (D1). This concludes the proof of the first part of the proposition. The rest of the proof is dedicated to show the second part of the proposition.

Since  $l_{n,\text{alt}}^\lambda(\beta_0, \beta)$  is strictly concave and two times continuously differentiable,  $(\hat{\beta}_{0,\text{alt}}^\lambda, \hat{\beta}_{\text{alt}}^\lambda)$  is a stationary point of  $l_{n,\text{alt}}^\lambda(\beta_0, \beta)$ , i.e.,

$$\begin{aligned} \nabla_{\beta} l_{n,\text{alt}}^\lambda(\hat{\beta}_{0,\text{alt}}^\lambda, \hat{\beta}_{\text{alt}}^\lambda) \\ = \sum_{i=1}^n \frac{y_i \mathbf{x}_i \exp\{-y_i(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)\}}{1 + \exp\{-y_i(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)\}} - \lambda \beta = 0, \end{aligned}$$

and

$$\begin{aligned} \frac{\partial}{\partial \beta_0} l_{n,\text{alt}}^\lambda(\hat{\beta}_{0,\text{alt}}^\lambda, \hat{\beta}_{\text{alt}}^\lambda) &= \sum_{i=1}^n \frac{y_i \exp\{-y_i(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)\}}{1 + \exp\{-y_i(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)\}} \\ &= \sum_{i=1}^n \left\{ \mathbb{I}(y_i = 1) - \frac{\exp(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)}{1 + \exp(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)} \right\} = 0, \end{aligned} \quad (\text{D2})$$

where, on the last line,  $\mathbb{I}(A)$  denotes the indicator function taking the value 1 if  $A$  is true, and 0 otherwise.

Using this result, we will next show that the search space in the optimization problem at (D1) can be narrowed to a compact set (see (D5) below).

From the first equation, we deduce that

$$\|\hat{\beta}_{\text{alt}}^\lambda\|_\infty \leq \lambda^{-1} \sum_{i=1}^n \|\mathbf{x}_i\|_\infty. \quad (\text{D3})$$

Rearranging the equation for  $\frac{\partial}{\partial \beta_0} l_{n,\text{alt}}^\lambda(\hat{\beta}_{0,\text{alt}}^\lambda, \hat{\beta}_{\text{alt}}^\lambda)$  gives

$$\begin{aligned} & \sum_{i:y_i=1} \left\{ 1 - \frac{\exp(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)}{1 + \exp(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)} \right\} \\ &= \sum_{i:y_i=-1} \frac{\exp(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)}{1 + \exp(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)}. \end{aligned} \quad (\text{D4})$$

From the latter equation, we will derive in turns a lower bound and an upper bound for  $\hat{\beta}_{0,\text{alt}}^\lambda$ . Starting with the upper bound, let  $n_1 = \#\{i : y_i = 1\}$  and  $n_{-1} = \#\{i : y_i = -1\}$ . Since we have assumed  $n_1 \geq 1$ , we have

$$\begin{aligned} & \sum_{i:y_i=1} \left\{ 1 - \frac{\exp(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)}{1 + \exp(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)} \right\} \\ & \geq 1 - \max_{i:y_i=1} \frac{\exp(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)}{1 + \exp(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)} \\ & = 1 - \max_{i:y_i=1} \frac{1}{1 + \exp\{-(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)\}} \\ & \geq 1 - \frac{1}{1 + \exp\{-\hat{\beta}_{0,\text{alt}}^\lambda - \|\hat{\beta}_{\text{alt}}^\lambda\|_\infty \max_{1 \leq i \leq n} \|\mathbf{x}_i\|_1\}}, \end{aligned}$$

where  $\|\mathbf{a}\|_1 := \sum_{j=1}^p |a_j|$  for  $\mathbf{a} \in \mathbb{R}^p$ . On the other hand, the term on the second line of (D4) satisfies

$$\begin{aligned} & \sum_{i:y_i=-1} \frac{\exp(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)}{1 + \exp(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)} \\ & \leq \frac{n_{-1}}{1 + \exp\{-\hat{\beta}_{0,\text{alt}}^\lambda - \|\hat{\beta}_{\text{alt}}^\lambda\|_\infty \max_{1 \leq i \leq n} \|\mathbf{x}_i\|_1\}}. \end{aligned}$$

Consequently,

$$\begin{aligned} & 1 - \frac{1}{1 + \exp\{-\hat{\beta}_{0,\text{alt}}^\lambda - \|\hat{\beta}_{\text{alt}}^\lambda\|_\infty \max_{1 \leq i \leq n} \|\mathbf{x}_i\|_1\}} \\ & \leq \frac{n_{-1}}{1 + \exp\{-\hat{\beta}_{0,\text{alt}}^\lambda - \|\hat{\beta}_{\text{alt}}^\lambda\|_\infty \max_{1 \leq i \leq n} \|\mathbf{x}_i\|_1\}}, \end{aligned}$$

which implies that

$$\begin{aligned} \hat{\beta}_{0,\text{alt}}^\lambda & \geq -\log(n_{-1}) - \|\hat{\beta}_{\text{alt}}^\lambda\|_\infty \max_{1 \leq i \leq n} \|\mathbf{x}_i\|_1 \\ & \geq -\log(n) - p \|\hat{\beta}_{\text{alt}}^\lambda\|_\infty \sum_{i=1}^n \|\mathbf{x}_i\|_\infty, \end{aligned}$$

where, to obtain the last line, we used the fact that  $n_{-1} \leq n$ , that for any  $\mathbf{a} \in \mathbb{R}^p$ ,  $\|\mathbf{a}\|_1 \leq p\|\mathbf{a}\|_\infty$ , and that  $\max_{1 \leq i \leq n} \|\mathbf{x}_i\|_\infty \leq \sum_{i=1}^n \|\mathbf{x}_i\|_\infty$ .

Since from similar arguments we can show that  $\widehat{\beta}_{0,\text{alt}}^\lambda \leq \log(n) + p \max_{1 \leq i \leq n} \|\mathbf{x}_i\|_\infty \|\widehat{\beta}_{\text{alt}}^\lambda\|_\infty$ , we conclude using (D3) that

$$|\widehat{\beta}_{0,\text{alt}}^\lambda| \leq \log(n) + \lambda^{-1} p \left\{ \sum_{i=1}^n \|\mathbf{x}_i\|_\infty \right\}^2.$$

To summarize, we have just shown that the solution  $(\widehat{\beta}_{0,\text{alt}}^\lambda, \widehat{\beta}_{\text{alt}}^\lambda)$  of (D1) lies in the compact set  $\Theta_{\mathbf{x},\lambda} = \Theta_{\mathbf{x},\lambda}^0 \times (\Theta_{\mathbf{x},\lambda}^1)^p$ , with

$$\Theta_{\mathbf{x},\lambda}^0 = \{\beta \in \mathbb{R} : |\beta| \leq \log(n) + \lambda^{-1} p \left\{ \sum_{i=1}^n \|\mathbf{x}_i\|_\infty \right\}^2\}$$

and

$$\Theta_{\mathbf{x},\lambda}^1 = \{\beta \in \mathbb{R} : |\beta| \leq \lambda^{-1} \sum_{i=1}^n \|\mathbf{x}_i\|_\infty\}.$$

Consequently, the search space  $\mathbb{R} \times \mathbb{R}^p$  in the optimization problem at (D1) can be restricted to  $\Theta_{\mathbf{x},\lambda}$ , and the optimization problem at (D1) is equivalent to

$$\max_{\beta_0 \in \Theta_{\mathbf{x},\lambda}^0, \beta \in (\Theta_{\mathbf{x},\lambda}^1)^p} l_{n,\text{alt}}^\lambda(\beta_0, \beta). \quad (\text{D5})$$

Now from the relationship

$$\begin{aligned} \min_{\alpha \in (0,1)} \alpha x + (1-\alpha) \log(1-\alpha) + \alpha \log(\alpha) \\ = \log\left(\frac{1}{1+e^{-x}}\right), \end{aligned}$$

since for any  $\beta_0 \in \Theta_{\mathbf{x},\lambda}^0$  and  $\beta \in (\Theta_{\mathbf{x},\lambda}^1)^p$ , we have

$$\begin{aligned} |\beta_0 + \beta^\top \mathbf{x}_j| &\leq \log(n) + \lambda^{-1} p \left\{ \sum_{i=1}^n \|\mathbf{x}_i\|_\infty \right\}^2 \\ &\quad + \lambda^{-1} p \left\{ \sum_{i=1}^n \|\mathbf{x}_i\|_\infty \right\} \max_{1 \leq j \leq n} \|\mathbf{x}_j\|_\infty \\ &\leq \log(n) + 2\lambda^{-1} p \left\{ \sum_{i=1}^n \|\mathbf{x}_i\|_\infty \right\}^2 \\ &< \log(n) + 2 + 2\lambda^{-1} p \left\{ \sum_{i=1}^n \|\mathbf{x}_i\|_\infty \right\}^2 := c_{n,\mathbf{x},\lambda}, \end{aligned}$$



where we used the fact that  $|\beta^\top \mathbf{x}_j| \leq p \|\beta\|_\infty \|\mathbf{x}_j\|_\infty \leq p \|\beta\|_\infty \sum_{i=1}^n \|\mathbf{x}_i\|_\infty$ , we conclude that

$$\begin{aligned} l_{n,\text{alt}}^\lambda(\beta_0, \beta) &= \min_{\alpha \in (\Theta_{\mathbf{x},\lambda}^2)^n} l_{n,\text{alt}}^\lambda(\beta_0, \beta, \alpha) \\ \text{with } l_{n,\text{alt}}^\lambda(\beta_0, \beta, \alpha) &= \sum_{i=1}^n \left\{ \alpha_i y_i (\beta_0 + \mathbf{x}_i^\top \beta) + L(\alpha_i) \right\} - \frac{\lambda}{2} \beta^\top \beta \\ \text{and } \Theta_{\mathbf{x},\lambda}^2 &= \left[ \frac{1}{1 + \exp(c_{n,\mathbf{x},\lambda})}, \frac{1}{1 + \exp(-c_{n,\mathbf{x},\lambda})} \right]. \end{aligned} \quad (\text{D6})$$

We recall  $L(\alpha_i) = (1 - \alpha_i) \log(1 - \alpha_i) + \alpha_i \log(\alpha_i)$ . Since the solution  $\alpha(\beta_0, \beta)$  of the optimization problem  $\min_{\alpha \in (\Theta_{\mathbf{x},\lambda}^2)^n} l_{n,\text{alt}}^\lambda(\beta_0, \beta, \alpha)$  satisfies

$$\alpha_i(\beta_0, \beta) = \frac{\exp(-y_i(\beta_0 + \mathbf{x}_i^\top \beta))}{1 + \exp(-y_i(\beta_0 + \mathbf{x}_i^\top \beta))} \text{ for all } i \in \{1, \dots, n\},$$

and as the optimal  $(\hat{\beta}_{0,\text{alt}}^\lambda, \hat{\beta}_{\text{alt}}^\lambda)$  are such that

$$\begin{aligned} 0 &= \sum_{i=1}^n y_i \frac{\exp(-y_i(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda))}{1 + \exp(-y_i(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda))} \\ &= \sum_{i=1}^n y_i \alpha_i(\hat{\beta}_{0,\text{alt}}^\lambda, \hat{\beta}_{\text{alt}}^\lambda), \end{aligned}$$

(see (D2)), then, the search space  $(\Theta_{\mathbf{x},\lambda}^2)^n$  can be further narrowed to

$$\Theta_{\mathbf{x},\lambda}^3 := (\Theta_{\mathbf{x},\lambda}^2)^n \cap \{\alpha : \sum_{i=1}^n \alpha_i y_i = 0\}. \quad (\text{D7})$$

Since for any  $\alpha \in \Theta_{\mathbf{x},\lambda}^3$  the function  $l_{n,\text{alt}}^\lambda(\beta_0, \beta, \alpha)$  satisfies

$$\begin{aligned} l_{n,\text{alt}}^\lambda(\beta_0, \beta, \alpha) &:= l_{n,\text{alt}}^\lambda(\beta, \alpha) \\ &= \sum_{i=1}^n \left\{ \alpha_i y_i \mathbf{x}_i^\top \beta + L(\alpha_i) \right\} - \frac{\lambda}{2} \beta^\top \beta, \end{aligned}$$

the optimization problem in (D5) can be expressed as

$$\max_{\beta_0 \in \Theta_{\mathbf{x},\lambda}^0, \beta \in (\Theta_{\mathbf{x},\lambda}^1)^p} l_{n,\text{alt}}^\lambda(\beta_0, \beta) = \max_{\beta \in (\Theta_{\mathbf{x},\lambda}^1)^p} \left( \min_{\alpha \in \Theta_{\mathbf{x},\lambda}^3} l_{n,\text{alt}}^\lambda(\beta, \alpha) \right). \quad (\text{D8})$$

Now for any fixed  $\beta$  the function  $\alpha \mapsto l_{n,\text{alt}}^\lambda(\beta, \alpha)$  is convex, and for any fixed  $\alpha$  the function  $\beta \mapsto l_{n,\text{alt}}^\lambda(\beta, \alpha)$  is concave. Since  $\Theta_{\mathbf{x},\lambda}^1$  and  $\Theta_{\mathbf{x},\lambda}^3$  are compact convex sets, we can apply Sion's

minimax theorem [3] and swap the max and the min in the above equation and conclude that

$$\max_{\beta \in (\Theta_{\mathbf{x}, \lambda}^1)^p} \left( \min_{\alpha \in \Theta_{\mathbf{x}, \lambda}^3} l_{n, \text{alt}}^\lambda(\beta, \alpha) \right) = \min_{\alpha \in \Theta_{\mathbf{x}, \lambda}^3} \left( \max_{\beta \in (\Theta_{\mathbf{x}, \lambda}^1)^p} l_{n, \text{alt}}^\lambda(\beta, \alpha) \right)$$

The inner problem can be solved exactly, by equating to 0 the gradient of  $l_{n, \text{alt}}^\lambda(\beta, \alpha)$  with respect to  $\beta$ . The solution  $\beta(\alpha)$  of this equation is given by the following equation:

$$\beta(\alpha) = \lambda^{-1} \sum_{i=1}^n \alpha_i y_i \mathbf{x}_i.$$

Plugging the expression of  $\beta(\alpha)$  into the definition of  $l_{n, \text{alt}}^\lambda(\beta, \alpha)$  yields  $l_{n, \text{alt}}^\lambda(\beta(\alpha), \alpha) = J_{\text{alt}}^\lambda(\alpha)$ , and therefore

$$\begin{aligned} \max_{\beta \in (\Theta_{\mathbf{x}, \lambda}^1)^p} \left( \min_{\alpha \in \Theta_{\mathbf{x}, \lambda}^3} l_{n, \text{alt}}^\lambda(\beta, \alpha) \right) &= \min_{\alpha \in \Theta_{\mathbf{x}, \lambda}^3} \left( J_{\text{alt}}^\lambda(\alpha) \right) \\ &= \min_{\alpha \in (\Theta_{\mathbf{x}, \lambda}^2)^n} \left( J_{\text{alt}}^\lambda(\alpha) \right) \quad \text{s.t.} \quad \sum_{i=1}^n y_i \alpha_i = 0. \end{aligned}$$

To obtain the last line, we used the definition of  $\Theta_{\mathbf{x}, \lambda}^3$  in (D7).

The optimization problem on the last line is strongly convex with a linear equality constraint. It therefore has a unique solution provided the feasible set  $\{\alpha \in (\Theta_{\mathbf{x}, \lambda}^2)^n : \sum_{i=1}^n \alpha_i y_i = 0\}$  is non empty. This is the case, since, as we have assumed  $n_1 \geq 1$  and  $n_{-1} \geq 1$  with  $n \geq 2$ , the point  $\alpha^f = (\alpha_1^f, \dots, \alpha_n^f)^\top$  with  $\alpha_i^f = 1/(2n_{-1})$  if  $y_i = -1$  and  $\alpha_i^f = 1/(2n_1)$  if  $y_i = 1$  is feasible. To see why this is the case, recalling the definition of  $\Theta_{\mathbf{x}, \lambda}^2$  in (D6), it suffices to note that the inequality  $1/(1 + \exp\{x\}) \leq \exp(-x)$  implies that

$$\begin{aligned} \frac{1}{2 \max(n_{-1}, n_1)} &\geq (2n)^{-1} \geq \exp\{-\log(n) - \log(2)\} \\ &\geq \exp\{-\log(n) - 2\} \geq \frac{1}{1 + \exp\{\log(n) + 2\}} \\ &\geq \frac{1}{1 + \exp(c_{n, \mathbf{x}, \lambda})}, \end{aligned}$$

and the inequality  $1/2 \leq (1 + \exp\{-2\})^{-1}$  implies that

$$\frac{1}{2 \min(n_{-1}, n_1)} \leq \frac{1}{2} \leq \frac{1}{1 + \exp(-2)} \leq \frac{1}{1 + \exp(-c_{n, \mathbf{x}, \lambda})}.$$

Combining these relationships shows that  $\alpha_i^f \in \Theta_{\mathbf{x},\lambda}^2$  for all  $1 \leq i \leq n$ . The proof of the proposition follows from the fact that, since the solution to the optimization problem

$$\min_{\boldsymbol{\alpha} \in (\Theta_{\mathbf{x},\lambda}^2)^n} \left( J_{\text{alt}}^\lambda(\boldsymbol{\alpha}) \right) \quad \text{s.t.} \quad \sum_{i=1}^n y_i \alpha_i = 0$$

is reached at a stationary point that cancels the Lagrangian, it can be equivalently expressed as

$$\min_{\boldsymbol{\alpha} \in (0,1)^n} \left( J_{\text{alt}}^\lambda(\boldsymbol{\alpha}) \right) \quad \text{s.t.} \quad \sum_{i=1}^n y_i \alpha_i = 0.$$

Finally, the expression  $\hat{\beta}_{0,\text{alt}}^\lambda = \frac{1}{n} \sum_{i=1}^n (y_i \log\{(\hat{\alpha}_{i,\text{alt}}^\lambda)^{-1} - 1\} - \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda)$  can directly be retrieved using the following:

$$\hat{\alpha}_{i,\text{alt}}^\lambda = \frac{1}{1 + \exp(y_i(\hat{\beta}_{0,\text{alt}}^\lambda + \mathbf{x}_i^\top \hat{\beta}_{\text{alt}}^\lambda))} \quad \text{for all } i \in \{1, \dots, n\}.$$

□

## Appendix E Alternative reverse-engineering procedure

An alternative procedure can be employed to construct the diagonal matrix  $\mathbf{V}^\lambda$  at the CC when executing VERTIGO-CI, which enables the CC to follow the steps outlined in the main text and subsequently reverse-engineer the feature data of all individuals. Recall the definition of the diagonal matrix  $\mathbf{V}^\lambda$ , with its entries specified in Equation (12). Also, recall the expression of  $\nabla_{\boldsymbol{\alpha}} J^\lambda(\boldsymbol{\alpha})$  (Equation (10)), that  $\hat{\boldsymbol{\alpha}}^\lambda$  satisfies  $\nabla_{\boldsymbol{\alpha}} J^\lambda(\hat{\boldsymbol{\alpha}}^\lambda) = 0$ , that  $\hat{\beta}^{\lambda(k)} = \lambda^{-1} \sum_{i=1}^n \hat{\alpha}_i^\lambda y_i \mathbf{x}_i^{(k)}$  and that  $\hat{\beta}_0^\lambda = \lambda^{-1} \sum_{i=1}^n \hat{\alpha}_i^\lambda y_i$ .

507 Using these expression, we derive that the  $j$ th component of  $\nabla_{\alpha} J^{\lambda}(\alpha)$   
 508  
 509 satisfies

$$\begin{aligned}
 510 & \\
 511 & \\
 512 \quad 0 &= [\nabla_{\alpha} J^{\lambda}(\hat{\alpha}^{\lambda})]_j \\
 513 & \\
 514 &= \sum_{k=1}^K e_j^{(k)}(\hat{\alpha}^{\lambda}) + \log \left( \frac{\hat{\alpha}_j^{\lambda}}{1 - \hat{\alpha}_j^{\lambda}} \right) \\
 515 & \\
 516 &= \sum_{k=1}^K (\lambda^{-1} \sum_{i=1}^n \hat{\alpha}_i^{\lambda} y_i y_j (\mathbf{x}_i^{(k)})^{\top} \mathbf{x}_j^{(k)}) + \lambda^{-1} \sum_{i=1}^n \hat{\alpha}_i^{\lambda} y_i y_j + \log \left( \frac{\hat{\alpha}_j^{\lambda}}{1 - \hat{\alpha}_j^{\lambda}} \right) \\
 517 & \\
 518 &= y_j \sum_{k=1}^K ((\mathbf{x}_j^{(k)})^{\top} \hat{\boldsymbol{\beta}}^{\lambda(k)}) + y_j \hat{\beta}_0^{\lambda} + \log \left( \frac{\hat{\alpha}_j^{\lambda}}{1 - \hat{\alpha}_j^{\lambda}} \right) \\
 519 & \\
 520 &= y_j \sum_{k=1}^K ((\mathbf{x}_j^{(k)})^{\top} \hat{\boldsymbol{\beta}}^{\lambda(k)}) + y_j \hat{\beta}_0^{\lambda} + \log \left( \frac{\hat{\alpha}_j^{\lambda}}{1 - \hat{\alpha}_j^{\lambda}} \right) \\
 521 & \\
 522 &= y_j (\mathbf{x}_j^{\top} \hat{\boldsymbol{\beta}}^{\lambda} + \hat{\beta}_0^{\lambda}) + \log \left( \frac{\hat{\alpha}_j^{\lambda}}{1 - \hat{\alpha}_j^{\lambda}} \right). \\
 523 & \\
 524 & \\
 525 & \\
 526 & \\
 527 & \\
 528 &
 \end{aligned}$$

529 By isolating  $\hat{\alpha}_j^{\lambda}$ , we obtain that

$$\begin{aligned}
 530 & \\
 531 & \\
 532 \quad \hat{\alpha}_j^{\lambda} &= [1 + \exp\{y_j(\mathbf{x}_j^{\top} \hat{\boldsymbol{\beta}}^{\lambda} + \hat{\beta}_0^{\lambda})\}]^{-1}. \\
 533 & \\
 534 & \\
 535 &
 \end{aligned}$$

536 Since  $1 - \hat{\alpha}_j^{\lambda} = [1 + \exp\{-y_j(\mathbf{x}_j^{\top} \hat{\boldsymbol{\beta}}^{\lambda} + \hat{\beta}_0^{\lambda})\}]^{-1}$ , and as  $y_j \in \{-1, 1\}$ , we  
 537  
 538 deduce from the latter equation that

$$\begin{aligned}
 539 & \\
 540 & \\
 541 & \\
 542 \quad \hat{\alpha}_j^{\lambda}(1 - \hat{\alpha}_j^{\lambda}) & \\
 543 &= [1 + \exp\{(\mathbf{x}_j^{\top} \hat{\boldsymbol{\beta}}^{\lambda} + \hat{\beta}_0^{\lambda})\}]^{-1} [1 + \exp\{-(\mathbf{x}_j^{\top} \hat{\boldsymbol{\beta}}^{\lambda} + \hat{\beta}_0^{\lambda})\}]^{-1} \\
 544 & \\
 545 &= V_{jj}^{\lambda}. \\
 546 & \\
 547 & \\
 548 & \\
 549 & \\
 550 & \\
 551 & \\
 552 &
 \end{aligned}$$

Since  $\hat{\alpha}_j^\lambda$  is computed at the CC, the latter can determine each diagonal entry  $\mathbf{V}_{jj}^\lambda$  for all  $j \in 1, \dots, n$  and proceed, as outlined in the main text, to reverse-engineer the individual feature data. Notably, this approach allows the CC to reconstruct the matrix  $\mathbf{V}^\lambda$  without requiring access to the response vector  $\mathbf{y}$ .

## References

- [1] Li Y, Jiang X, Wang S, Xiong H, Ohno-Machado L. VERTIcal Grid lOgistic regression (VERTIGO). Journal of the American Medical Informatics Association. 2016;23(3):570–579. <https://doi.org/10.1093/jamia/ocv146>.
- [2] Kim J, Li W, Bath T, Jiang X, Ohno-Machado L. VERTIcal Grid lOgistic regression with Confidence Intervals (VERTIGO-CI). AMIA Summits on Translational Science Proceedings. 2021;p. 355–363.
- [3] Sion M. On general minimax theorems. Pacific Journal of Mathematics. 1958;8(1):171–176.