

Appendix

A. Proofs of the technical results in Section 2

A.1. Proof of Proposition 1

Proposition 1 Given a fixed training set $\mathcal{S}$, let $\boldsymbol{\mu} = [\mu_q]_{q \in [Q]}$ be the Lagrangian multipliers for the constraints $\{\frac{1}{|V_q|} \sum_{j \in V_q} (y_j - h_{\mathbf{w}}(\mathbf{x}_j))^2 \leq \delta + \xi_q\}_{q \in [Q]}$ in the optimization problem (2) and $F(\mathbf{w}, \boldsymbol{\mu}, \mathcal{S})$ be defined as follows:

$$F(\mathbf{w}, \boldsymbol{\mu}, \mathcal{S}) = \sum_{i \in \mathcal{S}} [\lambda \|\mathbf{w}\|^2 + (y_i - h_{\mathbf{w}}(\mathbf{x}_i))^2] + \sum_{q \in [Q]} \mu_q \left[\frac{\sum_{j \in V_q} (y_j - h_{\mathbf{w}}(\mathbf{x}_j))^2}{|V_q|} - \delta \right] \quad (16)$$

Then, for the fixed set $\mathcal{S}$, the dual of the optimization problem (2) for estimating $\mathbf{w}$ and $\{\xi_q\}$ is given by,

$$\max_{\mathbf{0} \leq \boldsymbol{\mu} \leq C\mathbf{1}} \min_{\mathbf{w}} F(\mathbf{w}, \boldsymbol{\mu}, \mathcal{S}) \quad (17)$$

Proof The dual problem of our data selection problem (2) is given as:

$$\max_{\boldsymbol{\mu} \geq \mathbf{0}, \boldsymbol{\nu}} \min_{\mathbf{w}, \{\xi_q\}_{q \in [Q]}} \sum_{i \in \mathcal{S}} [\lambda \|\mathbf{w}\|^2 + (y_i - h_{\mathbf{w}}(\mathbf{x}_i))^2] + C \sum_{q \in V_q} \xi_q + \sum_{q \in [Q]} \mu_q \left[\frac{\sum_{j \in V_q} (y_j - h_{\mathbf{w}}(\mathbf{x}_j))^2}{|V_q|} - \delta - \xi_q \right] - \nu_q \xi_q$$

Differentiating with respect to $\boldsymbol{\xi}$, we get $\boldsymbol{\mu} + \boldsymbol{\nu} = C\mathbf{1}$, which proves the Proposition (giving us the constraint $\mathbf{0} \leq \boldsymbol{\mu} \leq C\mathbf{1}$). ■

A.2. Proof of Proposition 3

Proposition 3 Both the variants of the data selection problems (4) and (8) are NP-Hard.

Proof Consider our data selection problem as follows:

$$\begin{aligned} & \min_{\mathcal{S} \subset \mathcal{D}, \mathbf{w}, \{\xi_q\}_{q \in [Q]}} \sum_{i \in \mathcal{S}} [\lambda \|\mathbf{w}\|^2 + (y_i - h_{\mathbf{w}}(\mathbf{x}_i))^2] + C \sum_{q \in V_q} \xi_q, \\ & \text{such that, } \frac{\sum_{j \in V_q} (y_j - h_{\mathbf{w}}(\mathbf{x}_j))^2}{|V_q|} \leq \delta + \xi_q \quad \forall q \in [Q], \\ & \xi_q \geq 0 \quad \forall q \in [Q] \text{ and, } |\mathcal{S}| = k \end{aligned} \quad (18)$$

We make $C = 0$ and $h_{\mathbf{w}}(\mathbf{x}) = \mathbf{w}^\top \mathbf{x}$. Then the problem becomes equivalent to the robust regression problem (Bhatia et al., 2017), i.e.,

$$\min_{\mathcal{S} \subset \mathcal{D}, \mathbf{w}} \sum_{i \in \mathcal{S}} [\lambda \|\mathbf{w}\|^2 + (y_i - \mathbf{w}^\top \mathbf{x})^2], \quad \text{such that, } |\mathcal{S}| = k, \quad (19)$$

which is known to be NP-hard. ■

B. Techninal results on Section 3 and their proofs

B.1. Proof of Proposition 5

Proposition 5 For any model h_w , $f(S)$ is monotone, i.e., $f(S \cup a) - f(S) \geq 0$ for all $S \subset \mathcal{D}$ and $a \in \mathcal{D} \setminus S$.

Proof We note that

$$f(S \cup a) - f(S) = F(w^*(\mu^*(S \cup a), S \cup a), \mu^*(S \cup a), S \cup a) - F(w^*(\mu^*(S), S), \mu^*(S), S) \quad (20)$$

$$= \underbrace{F(w^*(\mu^*(S \cup a), S \cup a), \mu^*(S \cup a), S \cup a) - F(w^*(\mu^*(S), S \cup a), \mu^*(S), S \cup a)}_{\geq 0} + F(w^*(\mu^*(S), S \cup a), \mu^*(S), S \cup a) - F(w^*(\mu^*(S), S), \mu^*(S), S) \quad (21)$$

$$\stackrel{(i)}{\geq} F(w^*(\mu^*(S), S \cup a), \mu^*(S), S \cup a) - F(w^*(\mu^*(S), S), \mu^*(S), S) \quad (22)$$

$$= F(w^*(\mu^*(S), S \cup a), \mu^*(S), S \cup a) - F(w^*(\mu^*(S), S \cup a), \mu^*(S), S) + \underbrace{F(w^*(\mu^*(S), S \cup a), \mu^*(S), S) - F(w^*(\mu^*(S), S), \mu^*(S), S)}_{\geq 0} \quad (23)$$

$$\stackrel{(ii)}{\geq} F(w^*(\mu^*(S), S \cup a), \mu^*(S), S \cup a) - F(w^*(\mu^*(S), S \cup a), \mu^*(S), S) = \sum_{i \in S \cup a} [\lambda \|w^*(\mu^*(S), S \cup a)\|^2 + (y_i - h_{w^*(\mu^*(S), S \cup a)}(x_i))^2] + \sum_{q \in [Q]} \mu_q^*(S) \left[\frac{\sum_{j \in V_q} (y_j - h_{w^*(\mu^*(S), S \cup a)}(x_j))^2}{|V_q|} - \delta \right] \quad (24)$$

$$- \sum_{i \in S} [\lambda \|w^*(\mu^*(S), S \cup a)\|^2 + (y_i - h_{w^*(\mu^*(S), S \cup a)}(x_i))^2] - \sum_{q \in [Q]} \mu_q^*(S) \left[\frac{\sum_{j \in V_q} (y_j - h_{w^*(\mu^*(S), S \cup a)}(x_j))^2}{|V_q|} - \delta \right] \quad (25)$$

$$= \lambda \|w^*(\mu^*(S), S \cup a)\|^2 + (y_a - h_{w^*(\mu^*(S), S \cup a)}(x_a))^2 \quad (26)$$

Here, inequality (i) is due to the fact that: $\mu^*(S \cup a) = \arg\max_{0 \leq \mu \leq C} F(w^*(\mu, S \cup a), \mu, S \cup a)$; and, inequality (ii) is due to the fact that: $w^*(\mu^*(S), S) = \arg\min_w F(w, \mu^*(S), S)$. ■

B.2. Proof of Theorem 6

Theorem 6 Assume that $|y| \leq y_{\max}$; $h_w(x) = 0$ if $w = 0$, i.e., $h_w(x)$ has no bias term; h_w is H -Lipschitz, i.e., $|h_w(x)| \leq H \|w\|$; the eigenvalues of the Hessian matrix of $(y - h_w(x))^2$ have a finite upper bound, i.e., $\text{Eigenvalue}(\nabla_w^2 (y - h_w(x))^2) \leq 2\chi_{\max}^2$; and, define $\ell^* = \min_{a \in \mathcal{D}} \min_w \chi_{\max}^2 \cdot \|w\|^2 + (y_a - h_w(x_a))^2 > 0$. Then, for $\lambda \geq \max \{ \chi_{\max}^2, 32(1 + CQ)^2 y_{\max}^2 H^2 / \ell^* \}$, $f(S)$ is a α -submodular set function, where

$$\alpha \geq \hat{\alpha}_f = 1 - \frac{32(1 + CQ)^2 y_{\max}^2 H^2}{\lambda \ell^*}, \quad (27)$$

Proof We assume that: $S \subset \mathcal{T}$. Hence, $|\mathcal{T}| > 0$. Let us define: $\ell_a(w) = \lambda \|w\|^2 + (y_a - h_w(x_a))^2$, $\bar{w} = \arg\min_w \ell_a(w)$. Finally, we denote $\ell^* = \min_{a \in \mathcal{D}} \min_w \chi_{\max}^2 \|w\|^2 + (y_a - h_w(x_a))^2$. Next, we have that:

$$\begin{aligned} \frac{f(S \cup a) - f(S)}{f(\mathcal{T} \cup a) - f(\mathcal{T})} &\geq \frac{\ell_a(w^*(\mu^*(S), S \cup a))}{\ell_a(w^*(\mu^*(\mathcal{T} \cup a), \mathcal{T}))} \quad (\text{Due to Lemma 12}) \\ &\geq \frac{\ell_a(\bar{w})}{\ell_a(w^*(\mu^*(\mathcal{T} \cup a), \mathcal{T}))} \quad (\text{Since } \ell_a(w^*(\mu^*(S), S \cup a)) \geq \ell_a(\bar{w})) \end{aligned}$$

$$\begin{aligned}
 &\stackrel{(i)}{\geq} \frac{\ell_a(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T}), \mathcal{T} \cup a)) - (\lambda + \chi_{\max}^2) \|\bar{\mathbf{w}} - \mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T})\|^2}{\ell_a(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T}))} \\
 &\geq 1 - \frac{(\lambda + \chi_{\max}^2) \|\bar{\mathbf{w}} - \mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T})\|^2}{\ell_a(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T}))} \\
 &\geq 1 - \frac{(\lambda + \chi_{\max}^2)}{\ell_a(\bar{\mathbf{w}})} \|\bar{\mathbf{w}} - \mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T})\|^2 \quad (\text{Since } \ell_a(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S} \cup a)) \geq \ell_a(\bar{\mathbf{w}})) \\
 &\stackrel{(ii)}{\geq} 1 - \frac{32\lambda(1+CQ)^2 y_{\max}^2 H^2}{\ell^* \lambda^2} \\
 &= 1 - \frac{32(1+CQ)^2 y_{\max}^2 H^2}{\lambda \ell^*}, \tag{28}
 \end{aligned}$$

Inequality (i) is due to the following:

$$\ell_a(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T})) \tag{29}$$

$$= \ell_a(\bar{\mathbf{w}}) + \nabla \ell_a(\bar{\mathbf{w}})^\top (\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T}) - \bar{\mathbf{w}}) + (\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T}) - \bar{\mathbf{w}})^\top \nabla^2 \ell_a(\bar{\mathbf{w}}) (\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T}) - \bar{\mathbf{w}}) \tag{30}$$

$$\leq \ell_a(\bar{\mathbf{w}}) + \frac{\max\{\text{eig}(\nabla^2 \ell_a)\}}{2} \|\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T}) - \bar{\mathbf{w}}\|^2 \quad (\nabla \ell_a(\bar{\mathbf{w}}) = 0) \tag{31}$$

$$\leq \ell_a(\bar{\mathbf{w}}) + (\lambda + \chi_{\max}^2) \|\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T}) - \bar{\mathbf{w}}\|^2; \tag{32}$$

and inequality (ii) follows from

$$\begin{aligned}
 (1) \quad &\ell_a(\bar{\mathbf{w}}) = \lambda \|\bar{\mathbf{w}}\|^2 + (y_a - h_{\bar{\mathbf{w}}}(\mathbf{x}_a))^2 \geq \chi_{\max}^2 \|\bar{\mathbf{w}}\|^2 + (y_a - h_{\bar{\mathbf{w}}}(\mathbf{x}_a))^2 \geq \min_{\mathbf{w}} \chi_{\max}^2 \|\mathbf{w}\|^2 + (y_a - h_{\mathbf{w}}(\mathbf{x}_a))^2 = \ell^*, \\
 (2) \quad &\|\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T}) - \bar{\mathbf{w}}\| \leq 2w_{\max} = \frac{4(1+CQ)y_{\max}H}{\lambda} \quad (\text{Due to Lemma 13}), \\
 (3) \quad &\lambda \geq \chi_{\max}^2. \tag{33}
 \end{aligned}$$

■

B.3. Proof of Proposition 7

Proposition 7 Given $0 < y_{\min} \leq |y| \leq y_{\max}$, $h_{\mathbf{w}}(\mathbf{x}) = \mathbf{w}^\top \mathbf{x}$, $\|\mathbf{x}\| \leq x_{\max}$, we set the regularizing coefficient as $\lambda \geq \max\{x_{\max}^2, 16(1+CQ)^2 y_{\max}^2 x_{\max}^2 / y_{\min}^2\}$. Then $f(\mathcal{S})$ is a α -submodular set function, where

$$\alpha \geq \hat{\alpha}_f = 1 - \frac{16(1+CQ)^2 y_{\max}^2 x_{\max}^2}{\lambda y_{\min}^2}. \tag{34}$$

Proof The proof exactly follows the previous proof, except in the highlighted part. We assume that: $\mathcal{S} \subset \mathcal{T}$. Hence, $|\mathcal{T}| > 0$ and define $\ell_a(\mathbf{w}) = \lambda \|\mathbf{w}\|^2 + (y_a - h_{\mathbf{w}}(\mathbf{x}_a))^2$, $\bar{\mathbf{w}} = \text{argmin}_{\mathbf{w}} \ell_a(\mathbf{w})$; $\ell^* = \min_{a \in \mathcal{D}} \min_{\mathbf{w}} \chi_{\max}^2 \|\mathbf{w}\|^2 + (y_a - h_{\mathbf{w}}(\mathbf{x}_a))^2$. Then, we have that:

$$\begin{aligned}
 &\frac{f(\mathcal{S} \cup a) - f(\mathcal{S})}{f(\mathcal{T} \cup a) - f(\mathcal{T})} \geq \frac{\ell_a(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S} \cup a))}{\ell_a(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T}))} \\
 &\geq \frac{\ell_a(\bar{\mathbf{w}})}{\ell_a(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T}))} \\
 &\geq \frac{\ell_a(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T}), \mathcal{T} \cup a)) - (\lambda + \chi_{\max}^2) \|\bar{\mathbf{w}} - \mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T})\|^2}{\ell_a(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T}))} \\
 &\geq 1 - \frac{(\lambda + \chi_{\max}^2) \|\bar{\mathbf{w}} - \mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T})\|^2}{\ell_a(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T}))} \\
 &\geq 1 - \frac{(\lambda + \chi_{\max}^2)}{\ell_a(\bar{\mathbf{w}})} \|\bar{\mathbf{w}} - \mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T})\|^2 \\
 &\geq 1 - \frac{8\lambda(1+CQ)^2 y_{\max}^2 x_{\max}^2}{\ell^* \lambda^2}
 \end{aligned}$$

$$\begin{aligned}
 &= 1 - \frac{8(1+CQ)^2 y_{\max}^2 x_{\max}^2}{\lambda \ell^*} \\
 &\geq 1 - \frac{16(1+CQ)^2 y_{\max}^2 x_{\max}^2}{\lambda y_{\min}^2}
 \end{aligned} \tag{35}$$

where the highlighted part is due to second part of Lemma 13 which gives:

$$\|\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{T} \cup a), \mathcal{T}) - \bar{\mathbf{w}}\| \leq 2w_{\max} = \frac{2(1+CQ)y_{\max}x_{\max}}{\lambda}; \tag{36}$$

and, Claim 1 which shows that $\ell^* = \frac{\lambda y_{\min}^2}{\lambda + x_{\max}^2} \geq y_{\min}^2/2$. ■

B.4. Proof of Proposition 8

Proposition 8 *Given the assumptions stated in Theorem 6. the generalized curvature $k_f(\mathcal{S})$ for any set $\mathcal{S}$ satisfies $\kappa_f(\mathcal{S}) \leq \hat{\kappa}_f = 1 - \frac{\ell^*}{(CQ+1)y_{\max}^2}$.*

Proof Let us define: $\ell_a(\mathbf{w}) = \lambda \|\mathbf{w}\|^2 + (y_a - h_{\mathbf{w}}(\mathbf{x}_a))^2$, $\bar{\mathbf{w}} = \operatorname{argmin}_{\mathbf{w}} \ell_a(\mathbf{w})$. Finally, we denote $\ell^* = \min_{a \in \mathcal{D}} \min_{\mathbf{w}} \chi_{\max}^2 \|\mathbf{w}\|^2 + (y_a - h_{\mathbf{w}}(\mathbf{x}_a))^2$. By definition, we have $1 - \kappa_f(\mathcal{S}) = \min_{a \in \mathcal{D}} \frac{f(a|\mathcal{S} \setminus a)}{f(a|\emptyset)}$. We show that, from Lemma 12, we have that:

$$f(a|\mathcal{S} \setminus a) \geq \lambda \|\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \setminus a), \mathcal{S})\|^2 + (y_a - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \setminus a), \mathcal{S})}(\mathbf{x}_a))^2 \geq \ell_a(\bar{\mathbf{w}}) \geq \ell^* \tag{37}$$

Next, we note that:

$$f(a|\emptyset) = f(a) - f(\emptyset) \tag{38}$$

$$\begin{aligned}
 &= \lambda \|\mathbf{w}^*(\boldsymbol{\mu}^*(a), a)\|^2 + (y_a - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(a), a)}(\mathbf{x}_a))^2 \\
 &\quad + \sum_{q \in [Q]} \mu_q^*(a) \sum_{j \in V_q} \left[\frac{(y_j - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(a), a)}(\mathbf{x}_j))^2}{|V_q|} - \delta \right] - \sum_{q \in [Q]} \mu_q^*(\emptyset) \sum_{j \in V_q} \left[\frac{(y_j - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(\emptyset), \emptyset)}(\mathbf{x}_j))^2}{|V_q|} - \delta \right] \\
 &\stackrel{(i)}{\leq} \lambda \|\mathbf{w}^*(\boldsymbol{\mu}^*(a), a)\|^2 + (y_a - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(a), a)}(\mathbf{x}_a))^2 \\
 &\quad + \sum_{q \in [Q]} \mu_q^*(a) \sum_{j \in V_q} \left[\frac{(y_j - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(a), a)}(\mathbf{x}_j))^2}{|V_q|} - \delta \right] - \sum_{q \in [Q]} \mu_q^*(a) \sum_{j \in V_q} \left[\frac{(y_j - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(\emptyset), \emptyset)}(\mathbf{x}_j))^2}{|V_q|} - \delta \right] \\
 &= \lambda \|\mathbf{w}^*(\boldsymbol{\mu}^*(a), a)\|^2 + (y_a - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(a), a)}(\mathbf{x}_a))^2 \\
 &\quad + \sum_{q \in [Q]} \mu_q^*(a) \sum_{j \in V_q} \left[\frac{(y_j - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(a), a)}(\mathbf{x}_j))^2}{|V_q|} \right] - \sum_{q \in [Q]} \mu_q^*(a) \sum_{j \in V_q} \left[\frac{(y_j - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(\emptyset), \emptyset)}(\mathbf{x}_j))^2}{|V_q|} \right] \\
 &\leq \lambda \|\mathbf{w}^*(\boldsymbol{\mu}^*(a), a)\|^2 + (y_a - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(a), a)}(\mathbf{x}_a))^2 \\
 &\quad + \sum_{q \in [Q]} \mu_q^*(a) \sum_{j \in V_q} \left[\frac{(y_j - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(a), a)}(\mathbf{x}_j))^2}{|V_q|} \right]
 \end{aligned} \tag{39}$$

$$\stackrel{(ii)}{\leq} (CQ+1)y_{\max}^2 \tag{40}$$

Here, (i) is because $\boldsymbol{\mu}^*(\emptyset) = \operatorname{argmax}_{\boldsymbol{\mu}} \sum_{q \in [Q]} \mu_q \sum_{j \in V_q} \left[\frac{(y_j - h_{\mathbf{w}^*(\boldsymbol{\mu}, \emptyset)}(\mathbf{x}_j))^2}{|V_q|} - \delta \right]$, (ii) is obtained by putting $\mathbf{w} = \mathbf{0}$ in Eq. (39) which is now at the minimum, i.e.,

$$\mathbf{w}^*(\boldsymbol{\mu}^*(a), a) = \operatorname{argmin}_{\mathbf{w}} \lambda \|\mathbf{w}\|^2 + (y_a - h_{\mathbf{w}}(\mathbf{x}_a))^2 + \sum_{q \in [Q]} \mu_q^*(a) \sum_{j \in V_q} \left[\frac{(y_j - h_{\mathbf{w}}(\mathbf{x}_j))^2}{|V_q|} \right] \tag{41}$$

Hence, Eqs. (37) and (40) show that, $\kappa_f(\mathcal{S}) \leq 1 - \frac{\ell^*}{(CQ+1)y_{\max}^2}$. ■

B.5. Auxiliary Lemmas

Lemma 12 If $f(\cdot)$ defined in Eq. (7), we have that

$$f(\mathcal{S} \cup a) - f(\mathcal{S}) \geq \lambda \|\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S} \cup a)\|^2 + (y_a - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S} \cup a)}(\mathbf{x}_a))^2. \quad (42)$$

and,

$$f(\mathcal{S} \cup a) - f(\mathcal{S}) \leq \lambda \|\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})\|^2 + (y_a - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})}(\mathbf{x}_a))^2. \quad (43)$$

Proof The proof of the lower bound of the marginal gain

$$f(\mathcal{S} \cup a) - f(\mathcal{S}) \geq \lambda \|\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S} \cup a)\|^2 + (y_a - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S} \cup a)}(\mathbf{x}_a))^2. \quad (44)$$

follows from the proof of Proposition 5.

Next we prove that

$$f(\mathcal{S} \cup a) - f(\mathcal{S}) \leq \lambda \|\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})\|^2 + (y_a - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})}(\mathbf{x}_a))^2. \quad (45)$$

To show this, we prove that:

$$\begin{aligned} f(\mathcal{S} \cup a) - f(\mathcal{S}) &= F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S} \cup a), \boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S} \cup a) - F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S}), \boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S}) \\ &= F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S} \cup a), \boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S} \cup a) - F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S}), \boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S}) \end{aligned} \quad (46)$$

$$+ \underbrace{F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S}), \boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S}) - F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S}), \boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S})}_{\leq 0} \quad (47)$$

$$\stackrel{(i)}{\leq} F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S} \cup a), \boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S} \cup a) - F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S}), \boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S}) \quad (48)$$

$$= \underbrace{F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S} \cup a), \boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S} \cup a) - F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S}), \boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S} \cup a)}_{\leq 0}$$

$$+ F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S}), \boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S} \cup a) - F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S}), \boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})$$

$$\stackrel{(ii)}{\leq} F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S}), \boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S} \cup a) - F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S}), \boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})$$

$$= \sum_{i \in \mathcal{S} \cup a} [\lambda \|\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})\|^2 + (y_i - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})}(\mathbf{x}_i))^2]$$

$$+ \sum_{q \in [Q]} \mu_q^*(\mathcal{S} \cup a) \left[\frac{\sum_{j \in V_q} (y_j - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})}(\mathbf{x}_j))^2}{|V_q|} - \delta \right] \quad (49)$$

$$- \sum_{i \in \mathcal{S}} [\lambda \|\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})\|^2 + (y_i - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})}(\mathbf{x}_i))^2]$$

$$- \sum_{q \in [Q]} \mu_q^*(\mathcal{S} \cup a) \left[\frac{\sum_{j \in V_q} (y_j - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})}(\mathbf{x}_j))^2}{|V_q|} - \delta \right] \quad (50)$$

$$= \lambda \|\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})\|^2 + (y_a - h_{\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S})}(\mathbf{x}_a))^2. \quad (51)$$

Here (i) is due to the fact that,

$$\boldsymbol{\mu}^*(\mathcal{S}) = \operatorname{argmax}_{\boldsymbol{\mu}} F(\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S}), \boldsymbol{\mu}, \mathcal{S}) \quad (52)$$

and (ii) is due to the fact that:

$$\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S} \cup a) = \operatorname{argmin}_{\mathbf{w}} F(\mathbf{w}, \boldsymbol{\mu}^*(\mathcal{S} \cup a), \mathcal{S} \cup a) \quad (53)$$

■

Lemma 13 Given that $\mathcal{S} \neq \emptyset$; $h_{\mathbf{w}}(\mathbf{x}) = 0$ for $\mathbf{w} = \mathbf{0}$; and, $h_{\mathbf{w}}(\mathbf{x})$ is H -Lipschitz, i.e., $|h_{\mathbf{w}}(\mathbf{x})| \leq H \|\mathbf{w}\|$. Then, we have $\|\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})\| \leq \frac{2(1+CQ)y_{\max}H}{\lambda}$. Moreover, if $h_{\mathbf{w}}(\mathbf{x}) = \mathbf{w}^\top \mathbf{x}$, we have that $\|\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})\| \leq \frac{(1+CQ)y_{\max}x_{\max}}{\lambda}$. Note that, for the linear model, we are able to exploit the structure of the model much better and therefore the bound is tighter.

Proof First we define $\nabla h_0(\mathbf{x}) = \nabla_{\mathbf{w}} h_{\mathbf{w}}(\mathbf{x})|_{\mathbf{w}=\mathbf{0}}$.

$$\begin{aligned}
 & F(\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S}), \boldsymbol{\mu}, \mathcal{S}) \\
 &= \lambda \|\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})\|^2 |\mathcal{S}| + \sum_{i \in \mathcal{S}} (y_i - h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}(\mathbf{x}_i))^2 + \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} [(y_j - h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}(\mathbf{x}_j))^2 - \delta] \\
 &= \lambda \|\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})\|^2 |\mathcal{S}| + \sum_{i \in \mathcal{S}} y_i^2 + \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} y_j^2 - \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} \delta \\
 &\quad - 2 \sum_{i \in \mathcal{S}} y_i h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}(\mathbf{x}_i) - 2 \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} y_j h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}(\mathbf{x}_j) + \underbrace{\sum_{i \in \mathcal{S}} h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}^2(\mathbf{x}_i) + \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}^2(\mathbf{x}_j)}_{\geq 0} \\
 &\stackrel{(i)}{\geq} \lambda \|\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})\|^2 |\mathcal{S}| + \underbrace{\sum_{i \in \mathcal{S}} y_i^2 + \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} y_j^2 - \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} \delta}_{F(\mathbf{0}, \boldsymbol{\mu}, \mathcal{S})} - 2 \sum_{i \in \mathcal{S}} y_i h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}(\mathbf{x}_i) - 2 \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} y_j h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}(\mathbf{x}_j). \tag{54}
 \end{aligned}$$

Here (i) is due to $\sum_{i \in \mathcal{S}} h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}^2(\mathbf{x}_i) + \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}^2(\mathbf{x}_j) \geq 0$. Now since $F(\mathbf{0}, \boldsymbol{\mu}, \mathcal{S}) = \sum_{i \in \mathcal{S}} y_i^2 + \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} y_j^2 - \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} \delta$, Eq. (54) gives us:

$$\begin{aligned}
 F(\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S}), \boldsymbol{\mu}, \mathcal{S}) - F(\mathbf{0}, \boldsymbol{\mu}, \mathcal{S}) &\geq \lambda \|\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})\|^2 |\mathcal{S}| \\
 &\quad - 2 \sum_{i \in \mathcal{S}} y_i h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}(\mathbf{x}_i) - 2 \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} y_j h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}(\mathbf{x}_j) \tag{55}
 \end{aligned}$$

Now since $\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S}) = \operatorname{argmin}_{\mathbf{w}} F(\mathbf{w}, \boldsymbol{\mu}, \mathcal{S})$, we have that $F(\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S}), \boldsymbol{\mu}, \mathcal{S}) \leq F(\mathbf{0}, \boldsymbol{\mu}, \mathcal{S})$. Then, Eq. (55) implies that

$$\begin{aligned}
 & \lambda \|\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})\|^2 |\mathcal{S}| - 2 \sum_{i \in \mathcal{S}} y_i h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}(\mathbf{x}_i) - 2 \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} y_j h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}(\mathbf{x}_j) \leq 0 \\
 & \stackrel{(i)}{\implies} \lambda \|\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})\|^2 |\mathcal{S}| \leq 2|\mathcal{S}| y_{\max} H \|\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})\| + 2 \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} y_{\max} H \|\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})\| \\
 & \leq 2(|\mathcal{S}| + CQ) y_{\max} H \|\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})\| \\
 & \implies \|\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})\| \leq \frac{2(|\mathcal{S}| + CQ) y_{\max} H}{\lambda |\mathcal{S}|} \leq \frac{2(1 + CQ) y_{\max} H}{\lambda} \tag{56}
 \end{aligned}$$

Here (i) is due to H -Lipschitzness of $h_{\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})}(\mathbf{x})$. For linear model, we have $H = x_{\max}$. However, we use the structure of the model to obtain a better bound. More specifically, for linear model, we have:

$$\begin{aligned}
 \mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S}) &= \left(\lambda |\mathcal{S}| \mathbb{I} + \sum_{i \in \mathcal{S}} \mathbf{x}_i \mathbf{x}_i^\top + \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} \mathbf{x}_j \mathbf{x}_j^\top \right)^{-1} \left(\sum_{i \in \mathcal{S}} y_i \mathbf{x}_i + \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} y_j \mathbf{x}_j \right) \\
 \implies \|\mathbf{w}^*(\boldsymbol{\mu}, \mathcal{S})\| &\leq \frac{(|\mathcal{S}| + CQ) x_{\max} y_{\max}}{\operatorname{Eig}_{\min} \left(\lambda |\mathcal{S}| \mathbb{I} + \sum_{i \in \mathcal{S}} \mathbf{x}_i \mathbf{x}_i^\top + \sum_{q \in [Q]} \frac{\mu_q}{|V_q|} \sum_{j \in V_q} \mathbf{x}_j \mathbf{x}_j^\top \right)} \\
 &\leq \frac{(|\mathcal{S}| + CQ) x_{\max} y_{\max}}{\lambda |\mathcal{S}|} \\
 &\leq \frac{(1 + CQ) x_{\max} y_{\max}}{\lambda}. \tag{57}
 \end{aligned}$$

■

Claim 1 $\min_{\mathbf{w}} [\lambda \|\mathbf{w}\|^2 + (y_a - \mathbf{w}^\top \mathbf{x}_a)^2] = \frac{\lambda y_a^2}{\lambda + \|\mathbf{x}_a\|^2}$

Proof We note that:

$$\bar{\mathbf{w}} = y_a(\lambda + \mathbf{x}_a \mathbf{x}_a^\top)^{-1} \mathbf{x}_a \quad (58)$$

Hence, we have that:

$$\lambda \|\bar{\mathbf{w}}\|^2 + (y_a - \bar{\mathbf{w}}^\top \mathbf{x}_a)^2 = y_a^2 - 2y_a \bar{\mathbf{w}}^\top \mathbf{x}_a + \bar{\mathbf{w}}^\top (\lambda \mathbb{I} + \mathbf{x}_a \mathbf{x}_a^\top) \bar{\mathbf{w}} \quad (59)$$

$$= y_a^2 - y_a \bar{\mathbf{w}}^\top \mathbf{x}_a \quad (60)$$

$$= y_a^2 - y_a^2 \mathbf{x}_a^\top (\lambda + \mathbf{x}_a \mathbf{x}_a^\top)^{-1} \mathbf{x}_a \quad (61)$$

$$= y_a^2 - y_a^2 \mathbf{x}_a^\top \left[\frac{1}{\lambda} - \frac{\mathbf{x}_a \mathbf{x}_a^\top / \lambda^2}{1 + \mathbf{x}_a^\top \mathbf{x}_a / \lambda} \right] \mathbf{x}_a \quad (\text{Due to Sherman Morrison formula}) \quad (62)$$

$$\begin{aligned} &= \frac{\lambda y_a^2}{\lambda + \|\mathbf{x}_a\|^2} \\ &\geq \frac{\lambda y_{\min}^2}{\lambda + x_{\max}^2} \end{aligned} \quad (63)$$

■

C. Proofs of the technical results in Section 4

C.1. Proof of Lemma 9

Lemma 9 Given a fixed set $\hat{\mathcal{S}}$ and an α -submodular function $f(\mathcal{S})$, let the modular function $m_{\hat{\mathcal{S}}}^f[\mathcal{S}]$ be defined as follows:

$$\begin{aligned} m_{\hat{\mathcal{S}}}^f[\mathcal{S}] = & f(\hat{\mathcal{S}}) - \sum_{i \in \hat{\mathcal{S}}} \alpha f(i|\hat{\mathcal{S}} \setminus \{i\}) \\ & + \sum_{i \in \hat{\mathcal{S}} \cap \mathcal{S}} \alpha f(i|\hat{\mathcal{S}} \setminus \{i\}) + \sum_{i \in \mathcal{S} \setminus \hat{\mathcal{S}}} \frac{f(i|\emptyset)}{\alpha}. \end{aligned} \quad (64)$$

Then, $f(\mathcal{S}) \leq m_{\hat{\mathcal{S}}}^f[\mathcal{S}]$ for all $\mathcal{S} \subseteq \mathcal{D}$.

Proof Recall that f is α -submodular with coefficient $\hat{\alpha}_f$ if $f(a|\mathcal{S}) \geq \hat{\alpha}_f f(a|\mathcal{T})$, $a \notin \mathcal{T}$, $\mathcal{S} \subseteq \mathcal{T}$. Given this, the following inequalities follow directly from:

$$\hat{\alpha}_f [f(\mathcal{S}) - f(\hat{\mathcal{S}} \cap \mathcal{S})] \leq \sum_{i \in \mathcal{S} \setminus \hat{\mathcal{S}}} f(i|\emptyset) \quad (65)$$

and similarly,

$$[f(\hat{\mathcal{S}}) - f(\hat{\mathcal{S}} \cap \mathcal{S})] \geq \hat{\alpha}_f \sum_{i \in \hat{\mathcal{S}} \setminus \mathcal{S}} f(i|\hat{\mathcal{S}} \setminus i) \quad (66)$$

The inequalities above hold by considering a chain of sets from $\hat{\mathcal{S}} \cap \mathcal{S}$ to either $\hat{\mathcal{S}}$ or $\mathcal{S}$ and applying the weak-submodularity definition by considering sets $\mathcal{S}$ and $\mathcal{T}$ appropriately. We then multiply -1 to inequality (66), multiply $1/\hat{\alpha}_f$ to equation (65) and add both of them together. We then achieve:

$$f(\mathcal{S}) \leq f(\hat{\mathcal{S}}) - \hat{\alpha}_f \sum_{i \in \hat{\mathcal{S}} \setminus \mathcal{S}} f(i|\hat{\mathcal{S}} \setminus i) + \frac{1}{\hat{\alpha}_f} \sum_{i \in \mathcal{S} \setminus \hat{\mathcal{S}}} f(i|\emptyset) \quad (67)$$

Rearranging this, we get the expression for the Lemma. ■

C.2. Proof of Theorem 10

Theorem 10 If the training algorithm in Algorithm 1 (lines 3, 6, 8) provides perfect estimates of the model parameters, it obtains a set $\hat{\mathcal{S}}$ which satisfies:

$$f(\hat{\mathcal{S}}) \leq \frac{k}{\hat{\alpha}_f(1 + (k-1)(1 - \hat{\kappa}_f)\hat{\alpha}_f)} f(\mathcal{S}^*) \quad (68)$$

where $\hat{\alpha}_f$ and $\hat{\kappa}_f$ are as stated in Theorem 6 and Proposition 8 respectively.

Proof From the definition of α -submodularity, note that $\hat{\alpha}_f f(\mathcal{S}) \leq \sum_{i \in \mathcal{S}} f(i)$. Next, we can obtain the following inequality for any $k \in \mathcal{S}$ using weak submodularity:

$$f(\mathcal{S}) - f(k) \geq \hat{\alpha}_f \sum_{j \in \mathcal{S} \setminus k} (f(j|\mathcal{S} \setminus j)) \quad (69)$$

We can add this up for all $k \in \mathcal{S}$ and obtain:

$$\begin{aligned} |\mathcal{S}|f(\mathcal{S}) - \sum_{k \in \mathcal{S}} f(k) & \geq \hat{\alpha}_f \sum_{k \in \mathcal{S}} \sum_{j \in \mathcal{S} \setminus k} (f(j|\mathcal{S} \setminus j)) \\ & \geq \hat{\alpha}_f (|\mathcal{S}| - 1) \sum_{k \in \mathcal{S}} f(k|\mathcal{S} \setminus k) \end{aligned} \quad (70)$$

Finally, from the definition of curvature, note that $f(k|\mathcal{S} \setminus k) \leq (1 - \hat{\kappa}_f)f(k)$. Combining all this together, we obtain:

$$|\mathcal{S}|f(\mathcal{S}) \geq (1 + \hat{\alpha}_f(1 - \hat{\kappa}_f)(|\mathcal{S}| - 1)) \sum_{j \in \mathcal{S}} f(j) \quad (71)$$

which implies:

$$\sum_{j \in \mathcal{S}} f(j) \leq \frac{|\mathcal{S}|}{1 + \hat{\alpha}_f(1 - \hat{\kappa}_f)(|\mathcal{S}| - 1)} f(\mathcal{S}) \quad (72)$$

Combining this with the fact that $\hat{\alpha}_f f(\mathcal{S}) \leq \sum_{i \in \mathcal{S}} f(i)$, we obtain that:

$$f(\mathcal{S}) \leq \frac{1}{\hat{\alpha}_f} \sum_{i \in \mathcal{S}} f(i) \leq \frac{|\mathcal{S}|}{\hat{\alpha}_f(1 + \hat{\alpha}_f(1 - \hat{\kappa}_f)(|\mathcal{S}| - 1))} f(\mathcal{S}) \quad (73)$$

The approximation guarantee then follows from some simple observations. In particular, given an approximation

$$m^f(\mathcal{S}) = \frac{1}{\hat{\alpha}_f} \sum_{i \in \mathcal{S}} f(i) \quad (74)$$

which satisfies $f(\mathcal{S}) \leq m^f(\mathcal{S}) \leq \beta_f f(\mathcal{S})$, we claim that optimizing m^f essentially gives a β_f approximation factor. To prove this, let $\mathcal{S}^*$ be the optimal subset, and $\hat{\mathcal{S}}$ be the subset obtained after optimizing m^f . The following chain of inequalities holds:

$$f(\hat{\mathcal{S}}) \leq m^f(\hat{\mathcal{S}}) \leq m^f(\mathcal{S}^*) \leq \beta_f f(\mathcal{S}^*) \quad (75)$$

This shows that $\hat{\mathcal{S}}$ is a β_f approximation of $\mathcal{S}^*$. Finally, note that this is just the first iteration of SELCON, and with subsequent iterations, SELCON is guaranteed to reduce the objective value (see Appendix C.4). ■

C.3. Proof of Theorem 11

Theorem 11 *If the training algorithm (lines 3, 6, 8) in Algorithm 1 provides imperfect estimates, so that $\|F(\hat{\mathbf{w}}, \hat{\boldsymbol{\mu}}, \mathcal{S}) - F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S}), \boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S})\| \leq \epsilon$ for any $\mathcal{S}$, then Algorithm 1 obtains a set $\hat{\mathcal{S}}$ that satisfies:*

$$f(\hat{\mathcal{S}}) \leq \left(\frac{k}{\hat{\alpha}_f(1 + (k-1)(1 - \hat{\kappa}_f)\hat{\alpha}_f)} + \frac{2k\epsilon}{\ell} \right) f(\mathcal{S}^*),$$

where $\ell = \min_{a \in \mathcal{D}} \min_{\mathbf{w}} \lambda \|\mathbf{w}\|^2 + (y_i - h_{\mathbf{w}}(\mathbf{x}_i))^2$, $\hat{\alpha}_f$ and $\hat{\kappa}_f$ are obtained in Theorem 6 and Proposition 8, respectively.

Proof Define:

$$\beta_f = \frac{k}{(1 + (k-1)(1 - \hat{\kappa}_f)\hat{\alpha}_f)} \quad (76)$$

and also define, $\hat{f}(\mathcal{S}) = F(\hat{\mathbf{w}}, \hat{\boldsymbol{\mu}}, \mathcal{S})$ and $f(\mathcal{S}) = F(\mathbf{w}^*(\boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S}), \boldsymbol{\mu}^*(\mathcal{S}), \mathcal{S})$. Note that instead of having access to f , the algorithm has access to $\hat{f}$ which satisfies:

$$|f(\mathcal{S}) - \hat{f}(\mathcal{S})| \leq \epsilon, \forall \mathcal{S} \quad (77)$$

Let us assume that $\hat{f}$ is always smaller compared to f , i.e. in other words,

$$f(\mathcal{S}) \leq \hat{f}(\mathcal{S}) \leq f(\mathcal{S}) + \epsilon \quad (78)$$

Combining this with the fact that:

$$f(\mathcal{S}) \leq \frac{1}{\hat{\alpha}_f} \sum_{j \in \mathcal{S}} f(j) \leq \frac{\beta_f}{\hat{\alpha}_f} f(\mathcal{S}) \quad (79)$$

we obtain the following chain of inequalities:

$$f(\mathcal{S}) \leq \frac{1}{\hat{\alpha}_f} \sum_{j \in \mathcal{S}} f(j) \leq \frac{1}{\hat{\alpha}_f} \sum_{j \in \mathcal{S}} [\hat{f}(j)] \leq \frac{1}{\hat{\alpha}_f} \sum_{j \in \mathcal{S}} [f(j) + \epsilon] \leq \frac{\beta_f}{\hat{\alpha}_f} f(\mathcal{S}) + \frac{k\epsilon}{\hat{\alpha}_f} \quad (80)$$

where $|\mathcal{S}| = k$. Finally, we get the approximation factor by dividing by a lower bound of $l = \min_{\mathcal{S}: |\mathcal{S}|=k} f(\mathcal{S})$ which can be obtained via a very similar proof technique to the weak submodularity and curvature results. Hence we get the final approximation factor as $\frac{\beta_f}{\hat{\alpha}_f} + \frac{k\epsilon}{l\hat{\alpha}_f}$.

We end by pointing out that we can get a similar result even if we do not assume that $\hat{f}$ is always smaller compared to f and in fact, assume the more general condition:

$$f(\mathcal{S}) - \epsilon \leq \hat{f}(\mathcal{S}) \leq f(\mathcal{S}) + \epsilon \quad (81)$$

The only difference is we have an additional factor of 2 in the additive bound. In particular, we get the following chain of inequalities:

$$f(\mathcal{S}) \leq \frac{1}{\hat{\alpha}_f} \sum_{j \in \mathcal{S}} f(j) \leq \frac{1}{\hat{\alpha}_f} \sum_{j \in \mathcal{S}} [\hat{f}(j) + \epsilon] \leq \frac{1}{\hat{\alpha}_f} \sum_{j \in \mathcal{S}} [f(j) + 2\epsilon] \leq \frac{\beta_f}{\hat{\alpha}_f} f(\mathcal{S}) + \frac{2k\epsilon}{\hat{\alpha}_f} \quad (82)$$

The chain of inequalities holds because $f(j) \leq \hat{f}(j) + \epsilon$ and $\hat{f}(j) \leq f(j) + \epsilon$. ■

C.4. Convergence property

We begin this section by showing that SELCON is guaranteed to reduce the objective value at every iteration as long as we obtain perfect solutions from the training algorithm (lines 3, 6, 8 in Algorithm 1).

Lemma 14 SELCON (Algorithm 1) is guaranteed to reduce the objective value of f at every iteration as long as we obtain perfect solutions from the training sub-routine.

Proof SELCON essentially uses modular upper bounds m^f of f at every iteration. Denote $\mathcal{S}_l$ as the set obtained in the l th iteration and let $\mathcal{S}_{l+1}$ be the one from the $l + 1$ th iteration. Then the following chain of inequalities hold:

$$f(\mathcal{S}_{l+1}) \leq m^f(\mathcal{S}_{l+1}) \leq m^f(\mathcal{S}_l) = f(\mathcal{S}_l) \quad (83)$$

The first inequality holds because m^f is a modular upper bound, the second inequality holds because $\mathcal{S}_{l+1}$ is the solution of minimizing m^f (and hence $m^f(\mathcal{S}_{l+1})$ is lower in value compared to $m^f(\mathcal{S}_l)$). The last equality holds because m^f is a modular upper bound which is tight at $\mathcal{S}_l$ and hence $m^f(\mathcal{S}_l) = f(\mathcal{S}_l)$. This shows that $f(\mathcal{S}_{l+1}) \leq f(\mathcal{S}_l)$. ■

We end this section by pointing out that this chain of inequalities does not hold if we get inexact or approximation solutions to the training sub-routine. In practice, we observe that the objective value of f still reduces even though we obtain only inexact solutions since the inexact solutions are often close to the true solutions of the training step.

D. Additional details about experimental setup

D.1. Dataset details

- **Cadata:** California housing dataset is obtained from the LIBSVM package ⁶. This spatial dataset contains 20,640 observations on housing prices with 9 economic covariates. As described in (Pace & Barry, 1997), here $\mathbf{x}$ are information about households in a block, say median age, median income, total rooms/population, bedrooms/population, population/households, households etc. and y is median price in median housing prices by all California census blocks. It has $\text{dimension}(\mathbf{x}) = 8$.
- **Law:** This refers to the dataset on Law School Admissions Council’s National Longitudinal Bar Passage Study (Wightman, 1998). Here $\mathbf{x}$ is information about a law student, including information on gender, race, family income, age, etc. and y indicates GPA normalised to $[0, 1]$. We use race as a protected attribute for the fairness experiments. It has $\text{dimension}(\mathbf{x}) = 10$.
- **NYSE-High:** This dataset is obtained from the New York stock exchange (NYSE) ⁷ dataset as follows. Given the set $\{s_i\}$ with s_i corresponding to the the highest stock price of the i^{th} day, we define $s_{k+1} = \sum_{i \in [100]} w_i s_{k+1-i}$. Here $y_k = s_{k+1}$ and $\mathbf{x}_k = [s_k, s_{k-1}, \dots, s_{k-99}]$. This dataset has $\text{dimension}(\mathbf{x}) = 100$.
- **NYSE-close:** This dataset is obtained from the New York stock exchange (NYSE) ⁸ dataset as follows. Given the set $\{s_i\}$ with s_i corresponding to the the closing stock price of the i^{th} day, we define $s_{k+1} = \sum_{i \in [100]} w_i s_{k+1-i}$. Here $y_k = s_{k+1}$ and $\mathbf{x}_k = [s_k, s_{k-1}, \dots, s_{k-99}]$. This dataset has $\text{dimension}(\mathbf{x}) = 100$.

D.2. Implementation details

Our models. We use two models— a simple linear regression model and a two layer neural network that consists of a linear layer of 5 hidden nodes and a ReLU activation unit. In all our experiments, we use a learning rate of 0.01. We choose the value of δ as the 30% of the mean validation error obtained using Full-selection.

Implementation of CRAIG. CRAIG (Mirzasoleiman et al., 2020) requires computing a $\mathcal{D} \times \mathcal{D}$ matrix with similarity measure for each pair of points in the training set. For the larger datasets, i.e., NYSE-close and NYSE-high, such a computation requires a large amount of memory. Hence, we use a stochastic version where we randomly select R points and build $R \times R$ matrix and select $\frac{kR}{\mathcal{D}}$ each time and repeat the process $\frac{\mathcal{D}}{R}$ times. We use $R = 50000$. Note that, for other datasets, since $|\mathcal{D}| < 50000$ the stochastic version is same as the original version.

CRAIG requires us to select the subset only once, since features will not change even as the training proceeds. However, since CRAIG is an adaptive method, for the non-linear setting, we need to run CRAIG every epoch. Despite using the stochastic version, we found CRAIG to be very slow in the non-linear setting and therefore we don’t report it.

Implementation of GLISTER. GLISTER (Killamsetty et al., 2021b), an another adaptive subset selection method where we select a new subset every 35th epoch to help make a fair comparison against SELCON. We update the model parameters after every selection step.

Machine configuration. We performed our experiments on a computer system with Ubuntu 16.04.6 LTS, an i-7 with 6 cores CPU and a total RAM of 125 GBs. The system had a single GeForce GTX 1080 GPU which was employed in our experiments.

E. Additional experiments

E.1. Discussion on adding offsets to the response variable y

The approximation ratio of SELCON is $f(\hat{\mathcal{S}})/f(\mathcal{S}^*) \leq \frac{k}{\hat{\alpha}_f(1 + (k-1)(1 - \hat{\kappa}_f)\hat{\alpha}_f)}$ when the training method is accurate.

A trite calculation shows that this quantity is $O(y_{\max}^4/y_{\min}^4)$. If $y_{\max}/y_{\min}$ is very high, the approximation ratio is affected. Such a problem can be easily overcome by adding an offset to y and then augmenting the feature $\mathbf{x}$ with an additional term 1— which incorporates the effect of the added offset. We summarize the effect of this offset on the approximation ratio (for different datasets) in Figure 5 which shows that adding an offset improves the approximation factor. Note that in the case of

⁶<https://www.csie.ntu.edu.tw/~cjlin/libsvmtools/datasets/regression.html>

⁷<https://www.kaggle.com/dgawlik/nyse>

⁸<https://www.kaggle.com/dgawlik/nyse>

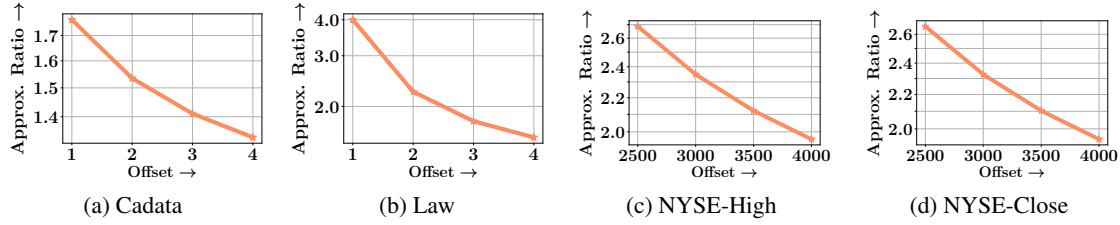


Figure 5. Variation in the approximation ratio with respect to the offset added to the response variables y .

Cadata, y indicates the house price; whereas in the case of Law, y indicates student GPA. Therefore, the approximation factor of these datasets is reasonable even without adding an offset. Whereas, for NYSE-High and NYSE-Close, the approximation factor is somewhat poorer at lower values of the offset (not shown in the plot).

E.2. Significance Tests

RANDOM-SELECTION	0.000089				
RANDOM-WITH-CONSTRAINTS	0.00012	0.50159			
CRAIG	0.040043	0.00014	0.00078		
GLISTER	0.001713	0.601212	0.88129	0.00803	
SELCON-WITHOUT-CONSTRAINTS	0.0001	0.00014	0.00059	0.0001	0.0001
SELCON					

Table 6. Pairwise significance p-values using Wilcoxon signed rank test.

In Table 6, we show the p-values of two-tailed Wilcoxon signed-rank test (Wilcoxon, 1992) performed on every possible pair of data selection strategies to determine whether there is a significant statistical difference between the strategies in each pair, across all datasets. Our null hypothesis is that there is no difference between each pair of data selection strategies. From the results, it is evident that SELCON significantly outperforms other baselines at $p < 0.01$.