



Preprint-Series: Department of Mathematics - Applied Mathematics

Nyström type subsampling analyzed as a regularized projection

Galyna Kriukova, Sergiy Pereverzyev Jr., Pavlo Tkackenko



APPLIEDMATHEMATICS

Technikerstraße 13 - 6020 Innsbruck - Austria
Tel.: +43 512 507 53803 Fax: +43 512 507 53898
<https://applied-math.uibk.ac.at>

i
1



Our last assumption describes the regularity of $f^\dagger$ in terms of source condition concept that is fairly standard in the regularization theory [8]. In the context of the learning theory this concept has been introduced in [2]. Within this concept, we assume that $f^\dagger$ admits the representation

$$f^\dagger = \varphi(\mathbf{C})v^\dagger, v^\dagger \in \mathcal{H}_K, \|v^\dagger\|_K \leq R, \quad (6)$$

where the function φ is operator monotone on $[0, d]$, $d > \|\mathbf{C}\|_{\mathcal{H}_K \rightarrow \mathcal{H}_K}$, and such that $\varphi(0) = 0$ and φ^2 is a concave function.

As it has been shown in [9] an important implication of operator monotonicity is that there is a number d_φ depending only on φ such that for any self-adjoint operators C, C_1 with spectra in $[0, d]$ it holds

$$\|\varphi(C) - \varphi(C_1)\|_{\mathcal{H}_K \rightarrow \mathcal{H}_K} \leq d_\varphi \varphi(\|C - C_1\|_{\mathcal{H}_K \rightarrow \mathcal{H}_K}). \quad (7)$$

Moreover, as a consequence of the concavity of φ^2 we have (see Proposition 2 [9])

$$\|(\mathbf{I} - \mathbf{P}_{\mathbf{z}^\nu})\varphi(\mathbf{C})\|_{\mathcal{H}_K \rightarrow \mathcal{H}_K} \leq \varphi(\|\mathbf{C}^{1/2}(\mathbf{I} - \mathbf{P}_{\mathbf{z}^\nu})\|_{\mathcal{H}_K \rightarrow \mathcal{H}_K}^2). \quad (8)$$

Note that our assumption (6) generalizes Assumption 4 of [11], where only the case of operator monotone functions $\varphi(t) = t^s$, $0 < s \leq \frac{1}{2}$, has been studied.

In the sequel we extensively use the following bounds (see, e.g., [2]) that hold under the above assumptions with probability at least $1 - \delta$ and quantify the perturbation effect of random sampling:

$$\|\mathbf{C} - \mathbf{S}_{\mathbf{z}}^* \mathbf{S}_{\mathbf{z}}\|_{\mathcal{H}_K \rightarrow \mathcal{H}_K} \leq d_{1,\delta} |\mathbf{z}|^{-\frac{1}{2}}, \quad (9)$$

$$\|\mathbf{S}_{\mathbf{z}}^* \mathbf{S}_{\mathbf{z}} f - \mathbf{S}_{\mathbf{z}}^* \mathbf{Y}\|_K \leq d_{2,\delta} |\mathbf{z}|^{-\frac{1}{2}}, \quad (10)$$

where $d_{1,\delta}$ and $d_{2,\delta}$ are of order $\mathcal{O}(\log \frac{1}{\delta})$ and depend only on K and ρ .

The following capacity independent learning rates have been proven in [2] for KRR (1)

Theorem 1 ([2]). *Consider a sampling space $Z = X \times [-D, D]$, where the input space $X \subset \mathbb{R}^d$ is closed. Consider also a bounded and continuous kernel K defined on X . If minimizer $f^\dagger$ of the expected risk $\mathcal{E}(f)$ over $\mathcal{H}_K$ meets the assumption (6), then for $\alpha = \alpha_{\mathbf{z}} = \Theta^{-1}(|\mathbf{z}|^{-1/2})$, $\Theta(t) = \varphi(t)t$, we have with probability at least $1 - \delta$ that*

$$\|f^\dagger - f_{\mathbf{z}}^{\alpha_{\mathbf{z}}}\|_\rho = \mathcal{O}\left(\varphi(\Theta^{-1}(|\mathbf{z}|^{-1/2}))\sqrt{\Theta^{-1}(|\mathbf{z}|^{-1/2})\log \frac{1}{\delta}}\right). \quad (11)$$

Note that for $\varphi(t) = t^s$ the above theorem gives us the learning rate $\mathcal{O}\left(|\mathbf{z}|^{-\frac{s+\frac{1}{2}}{2(s+1)}}\right)$ that matches the result obtained in seminal paper by Smale and Zhou [14]. Moreover, for $\varphi(t) = t^s$ the rate (11) can be thought of as the limit case of the capacity dependent learning rate $\mathcal{O}\left(|\mathbf{z}|^{-\frac{(s+\frac{1}{2})\mu}{2s\mu+\mu+1}}\right)$ obtained in [3] under the assumptions that the eigenvalues λ_i of the covariance operator $\mathbf{C}$ have polynomial decay $\lambda_i \asymp i^{-\mu}$ with $\mu > 1$.

Now we are going to prove that the same learning rate (11) can be achieved in Nyström type subsampling (2) if the approximation power of $\mathbf{P}_{\mathbf{z}^\nu}$ is high enough.

Theorem 2. *Assume the conditions of Theorem 1, and let (5) be satisfied. If the size $m = |\mathbf{z}^\nu|$ of a subsampling $\mathbf{z}^\nu$ is chosen such that*

$$\Delta_m \leq \sqrt{\Theta_{1/2}^{-1}(|\mathbf{z}|^{-1/2})}, \Theta_{1/2}(t) = \varphi(t)\sqrt{t},$$

then with probability at least $1 - \delta$ we have

$$\|f^\dagger - f_{\mathbf{z}, \mathbf{z}^\nu}^{\alpha_{\mathbf{z}}}\|_\rho = \mathcal{O}\left(\varphi\left(\Theta^{-1}(|\mathbf{z}|^{-1/2})\right) \sqrt{\Theta^{-1}(|\mathbf{z}|^{-1/2}) \log^{\beta_2} \frac{1}{\delta}}\right), \quad (12)$$

where $\beta_2 = \max\{1, \beta_1\}$, and β_1 is the same as in (5).

Before proving this statement, we first comment on the computational complexity of Nyström approximation (2) with a subsampling size $|\mathbf{z}^\nu|$ chosen according to Theorem 2.

In view of the assumption (5) it is clear that the condition of the theorem can be satisfied with

$$|\mathbf{z}^\nu| \asymp [\Theta_{1/2}^{-1}(|\mathbf{z}|^{-1/2})]^{-\frac{1}{2\beta}}.$$

Let the assumption (6) be satisfied with

$$\varphi(t) = o(t^{\frac{1-\beta}{2\beta}}) \text{ as } t \rightarrow 0, \quad (13)$$

i.e. $\Theta_{1/2}(t) = o(t^{1/2\beta})$. Then

$$|\mathbf{z}|^{-\beta} = o(\Theta_{1/2}^{-1}(|\mathbf{z}|^{-1/2})) = o(|\mathbf{z}^\nu|^{-2\beta}),$$

which means that $|\mathbf{z}^\nu|^2 = o(|\mathbf{z}|)$ as $|\mathbf{z}| \rightarrow \infty$.

On the other hand, the computational complexity of (2) is of order $\mathcal{O}(|\mathbf{z}||\mathbf{z}^\nu|^2)$ (see, e.g. [11]), and under the condition (13) it is subquadratic, because $|\mathbf{z}||\mathbf{z}^\nu|^2 = o(|\mathbf{z}|^2)$.

Thus, under the conditions of Theorem 2 Nyström subsampling has the same learning rate as the one guaranteed by Theorem 1 for KRR based on the whole sample $\mathbf{z}$. Moreover, Theorem 2 allows an estimation of a regularity range, such as (13), for which the above mentioned learning rate can be achieved with subquadratic complexity. Note, that the condition (13) is automatically satisfied with $\beta \geq 1$, for example.

Proof of Theorem 2. It is known (see, e.g. [9]) that the following inequality holds true for functions φ mentioned in the assumption (6)

$$\sup_t |(1 - (\alpha + t)^{-1}t)\varphi(t)t^q| \leq h_{\varphi,q}\varphi(\alpha)\alpha^q, q \in [0, 1/2], \quad (14)$$

where $h_{\varphi,q}$ depends only on φ and q .

Note also that, by very definition, $\Theta_{1/2}(|\mathbf{z}|^{-1/2}) > \Theta(|\mathbf{z}|^{-1/2})$, and therefore

$$\Delta_m^2 = \Theta_{1/2}^{-1}(|\mathbf{z}|^{-1/2}) < \Theta^{-1}(|\mathbf{z}|^{-1/2}) = \alpha_{\mathbf{z}}. \quad (15)$$

Moreover, without loss of generality we can assume that $|\mathbf{z}|$ is so large that

$$\varphi(\max\{d_{1,\delta}, d_{2,\delta}\}|\mathbf{z}|^{-1/2}) < [\max\{d_{1,\delta}, d_{2,\delta}\}], \quad (16)$$

where $d_{1,\delta}, d_{2,\delta}$ are the numbers appearing in (9), (10). This is not a real restriction, because the left-hand side of (16) tends to zero as $|\mathbf{z}| \rightarrow \infty$. A direct implication of (16) is that with probability at least $1 - \delta$

$$\alpha_{\mathbf{z}} = \Theta^{-1}(|\mathbf{z}|^{-1/2}) > \max\{\|\mathbf{C} - \mathbf{C}_{\mathbf{z}}\|_{\mathcal{H}_K \rightarrow \mathcal{H}_K}, \|\mathbf{C}_{\mathbf{z}}f^\dagger - \mathbf{S}_{\mathbf{z}}^*\mathbf{Y}\|_K\}. \quad (17)$$

Consider the decomposition

$$f^\dagger - f_{\mathbf{z},\mathbf{z}^\nu}^{\alpha_{\mathbf{z}}} = \sigma_1 + \sigma_2 + \sigma_3, \quad (18)$$

where

$$\begin{aligned} \sigma_1 &= f^\dagger - \mathbf{P}_{\mathbf{z}^\nu}f^\dagger, \\ \sigma_2 &= \mathbf{P}_{\mathbf{z}^\nu}f^\dagger - (\alpha_{\mathbf{z}}\mathbf{I} + \mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu})^{-1}\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu}f^\dagger, \\ \sigma_3 &= (\alpha_{\mathbf{z}}\mathbf{I} + \mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu})^{-1}(\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu}f^\dagger - \mathbf{P}_{\mathbf{z}^\nu}\mathbf{S}_{\mathbf{z}}^*\mathbf{Y}), \end{aligned}$$

and we use notation $\mathbf{C}_{\mathbf{z}} = \mathbf{S}_{\mathbf{z}}^* \mathbf{S}_{\mathbf{z}}$.

Now we are going to bound each term of (18). From (4)–(6) and (8) we have

$$\begin{aligned} \|\sigma_1\|_\rho &= \|\mathbf{C}^{1/2}(\mathbf{I} - \mathbf{P}_{\mathbf{z}^\nu})\varphi(\mathbf{C})v^\dagger\|_\mathbf{K} \\ &\leq R\|\mathbf{C}^{1/2}(\mathbf{I} - \mathbf{P}_{\mathbf{z}^\nu})\|_{\mathcal{H}_\mathbf{K} \rightarrow \mathcal{H}_\mathbf{K}}\|(\mathbf{I} - \mathbf{P}_{\mathbf{z}^\nu})\varphi(\mathbf{C})\|_{\mathcal{H}_\mathbf{K} \rightarrow \mathcal{H}_\mathbf{K}} \\ &\leq R\Delta_m\varphi(\Delta_m^2) = R\Theta_{1/2}(\Delta_m^2) \\ &\leq R\Theta_{1/2}(\Theta_{1/2}^{-1}(|\mathbf{z}|^{-1/2})) = R|\mathbf{z}|^{-1/2} \end{aligned} \quad (19)$$

To prove (12) we also need to bound σ_2, σ_3 in the norms $\|\cdot\|_\mathbf{K}$ and $\|\cdot\|_\rho$. We start with the decomposition

$$\sigma_2 = \sigma_{2,1} + \sigma_{2,2}, \quad (20)$$

where

$$\begin{aligned} \sigma_{2,1} &= (\mathbf{I} - (\alpha_{\mathbf{z}}\mathbf{I} + \mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu})^{-1}\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu})\varphi(\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu})v^\dagger, \\ \sigma_{2,2} &= (\mathbf{I} - (\alpha_{\mathbf{z}}\mathbf{I} + \mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu})^{-1}\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu})\sigma_{2,2,1}, \\ \sigma_{2,2,1} &= (\mathbf{P}_{\mathbf{z}^\nu}\varphi(\mathbf{C}) - \mathbf{P}_{\mathbf{z}^\nu}\varphi(\mathbf{C})\mathbf{P}_{\mathbf{z}^\nu} + \mathbf{P}_{\mathbf{z}^\nu}\varphi(\mathbf{C})\mathbf{P}_{\mathbf{z}^\nu} \\ &\quad - \varphi(\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}\mathbf{P}_{\mathbf{z}^\nu}) + \varphi(\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}\mathbf{P}_{\mathbf{z}^\nu}) - \varphi(\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu}))v^\dagger. \end{aligned}$$

From (14) it follows that

$$\|\sigma_{2,1}\|_\mathbf{K} \leq R \sup_t |(1 - (\alpha_{\mathbf{z}} + t)^{-1}t)\varphi(t)| \leq Rh_{\varphi,0}\varphi(\alpha_{\mathbf{z}})$$

Moreover,

$$\begin{aligned} \|\sigma_{2,1}\|_\rho &= \|\mathbf{C}^{1/2}\sigma_{2,1}\|_\mathbf{K} \\ &\leq \|\mathbf{C}_{\mathbf{z}}^{1/2}\mathbf{P}_{\mathbf{z}^\nu}\sigma_{2,1}\|_\mathbf{K} + \|(\mathbf{C}^{1/2} - \mathbf{C}_{\mathbf{z}}^{1/2})\mathbf{P}_{\mathbf{z}^\nu}\sigma_{2,1}\|_\mathbf{K}, \end{aligned}$$

and

$$\begin{aligned} \|\mathbf{C}_{\mathbf{z}}^{1/2}\mathbf{P}_{\mathbf{z}^\nu}\sigma_{2,1}\|_\mathbf{K} &\leq \|(\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu})^{1/2}\sigma_{2,1}\|_\mathbf{K} \\ &\leq R \sup_t |(1 - (\alpha_{\mathbf{z}} + t)^{-1}t)t^{1/2}\varphi(t)| \leq Rh_{\varphi,\frac{1}{2}}\alpha_{\mathbf{z}}^{1/2}\varphi(\alpha_{\mathbf{z}}). \end{aligned}$$

Keeping in mind that $\psi(t) = \sqrt{t}$ is an operator monotone function, from (7), (15) and (17), we have

$$\|(\mathbf{C}^{1/2} - \mathbf{C}_{\mathbf{z}}^{1/2})\mathbf{P}_{\mathbf{z}^\nu}\sigma_{2,1}\|_\mathbf{K} \leq d_{1/2}\|\mathbf{C} - \mathbf{C}_{\mathbf{z}}\|_{\mathcal{H}_\mathbf{K} \rightarrow \mathcal{H}_\mathbf{K}}^{1/2}\|\sigma_{2,1}\|_\mathbf{K} \leq d_{1/2}Rh_{\varphi,0}\alpha_{\mathbf{z}}^{1/2}\varphi(\alpha_{\mathbf{z}}).$$

All together this gives us the bound

$$\|\sigma_{2,1}\|_\rho = \mathcal{O}(\varphi(\alpha_{\mathbf{z}})\alpha_{\mathbf{z}}^{1/2}) = \mathcal{O}\left(\varphi(\Theta^{-1}(|\mathbf{z}|^{-1/2}))\sqrt{\Theta^{-1}(|\mathbf{z}|^{-1/2})}\right).$$

To estimate $\|\sigma_{2,2}\|_\rho$ we need to bound $\|\sigma_{2,2,1}\|_\kappa$. For this end, we use the following known estimate (see Proposition 3 [9])

$$\|\mathbf{P}_{\mathbf{z}^\nu}\varphi(\mathbf{C})\mathbf{P}_{\mathbf{z}^\nu} - \varphi(\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}\mathbf{P}_{\mathbf{z}^\nu})\|_{\mathcal{H}_\kappa \rightarrow \mathcal{H}_\kappa} \leq \bar{d}_\varphi \varphi(\|\mathbf{C}^{1/2}(\mathbf{I} - \mathbf{P}_{\mathbf{z}^\nu})\|_{\mathcal{H}_\kappa \rightarrow \mathcal{H}_\kappa}^2).$$

Moreover, (7), (8) and (15), (17) give us

$$\|\varphi(\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}\mathbf{P}_{\mathbf{z}^\nu}) - \varphi(\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu})\|_{\mathcal{H}_\kappa \rightarrow \mathcal{H}_\kappa} \leq d_\varphi \varphi(\|\mathbf{C} - \mathbf{C}_{\mathbf{z}}\|_{\mathcal{H}_\kappa \rightarrow \mathcal{H}_\kappa}) \leq d_\varphi \varphi(\alpha_{\mathbf{z}}),$$

and

$$\begin{aligned} \|\mathbf{P}_{\mathbf{z}^\nu}\varphi(\mathbf{C}) - \mathbf{P}_{\mathbf{z}^\nu}\varphi(\mathbf{C})\mathbf{P}_{\mathbf{z}^\nu}\|_{\mathcal{H}_\kappa \rightarrow \mathcal{H}_\kappa} &\leq \|\varphi(\mathbf{C})(\mathbf{I} - \mathbf{P}_{\mathbf{z}^\nu})\|_{\mathcal{H}_\kappa \rightarrow \mathcal{H}_\kappa} \\ &= \|(\mathbf{I} - \mathbf{P}_{\mathbf{z}^\nu})\varphi(\mathbf{C})\|_{\mathcal{H}_\kappa \rightarrow \mathcal{H}_\kappa} \leq \varphi(\|\mathbf{C}^{1/2}(\mathbf{I} - \mathbf{P}_{\mathbf{z}^\nu})\|_{\mathcal{H}_\kappa \rightarrow \mathcal{H}_\kappa}^2) \leq \varphi(\alpha_{\mathbf{z}}). \end{aligned}$$

Therefore, $\|\sigma_{2,2,1}\|_\kappa \leq R(\bar{d}_\varphi + d_\varphi + 1)\varphi(\alpha_{\mathbf{z}})$, and

$$\|\sigma_{2,2}\|_\kappa \leq \|\sigma_{2,2,1}\|_\kappa \sup_t |1 - \frac{t}{\alpha_{\mathbf{z}} + t}| \leq \|\sigma_{2,2,1}\|_\kappa = \mathcal{O}(\varphi(\alpha_{\mathbf{z}})).$$

Then, using the same argument as for $\|\sigma_{2,2,1}\|_\rho$ we obtain

$$\begin{aligned} \|\sigma_{2,2}\|_\rho &= \mathcal{O}\left(\varphi(\Theta^{-1}(|\mathbf{z}|^{-1/2}))\sqrt{\Theta^{-1}(|\mathbf{z}|^{-1/2})}\right), \text{ and} \\ \|\sigma_2\|_\rho &= \mathcal{O}\left(\varphi(\Theta^{-1}(|\mathbf{z}|^{-1/2}))\sqrt{\Theta^{-1}(|\mathbf{z}|^{-1/2})}\right). \end{aligned}$$

Finally, we need to estimate $\|\sigma_3\|_\rho$. Observe that

$$\begin{aligned} \|\sigma_3\|_\rho &\leq \sup_t |(\alpha_{\mathbf{z}} + t)^{-1}| \|\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu}f^\dagger - \mathbf{P}_{\mathbf{z}^\nu}\mathbf{S}_{\mathbf{z}}^*\mathbf{Y}\|_\kappa \\ &\leq \frac{1}{\alpha_{\mathbf{z}}} (\|\mathbf{P}_{\mathbf{z}^\nu}(\mathbf{C}_{\mathbf{z}}f^\dagger - \mathbf{S}_{\mathbf{z}}^*\mathbf{Y})\|_\kappa + \|\mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}f^\dagger - \mathbf{P}_{\mathbf{z}^\nu}\mathbf{C}_{\mathbf{z}}\mathbf{P}_{\mathbf{z}^\nu}f^\dagger\|_\kappa) \end{aligned}$$

Then using (8)–(10) we obtain

$$\|\mathbf{P}_{\mathbf{z}^\nu}(\mathbf{C}_{\mathbf{z}}f^\dagger - \mathbf{S}_{\mathbf{z}}^*\mathbf{Y})\|_\kappa \leq d_{2,\delta}|\mathbf{z}|^{-1/2},$$

$$\begin{aligned} \|P_{\mathbf{z}^\nu} \mathbf{C}_{\mathbf{z}} f^\dagger - P_{\mathbf{z}^\nu} \mathbf{C}_{\mathbf{z}} P_{\mathbf{z}^\nu} f^\dagger\|_{\mathbf{K}} &\leq \|P_{\mathbf{z}^\nu} (\mathbf{C}_{\mathbf{z}} - \mathbf{C}) f^\dagger\|_{\mathbf{K}} + \|P_{\mathbf{z}^\nu} \mathbf{C} f^\dagger - P_{\mathbf{z}^\nu} \mathbf{C} P_{\mathbf{z}^\nu} f^\dagger\|_{\mathbf{K}} \\ &+ \|P_{\mathbf{z}^\nu} (\mathbf{C} - \mathbf{C}_{\mathbf{z}}) P_{\mathbf{z}^\nu} f^\dagger\|_{\mathbf{K}} \leq 2d_{1,\delta} \|f^\dagger\|_{\mathbf{K}} |\mathbf{z}|^{-1/2} + \|\mathbf{C}(\mathbf{I} - P_{\mathbf{z}^\nu})(\mathbf{I} - P_{\mathbf{z}^\nu})\varphi(\mathbf{C})v^\dagger\|_{\mathbf{K}} \\ &\leq d_{3,\delta}(|\mathbf{z}|^{-1/2} + \Delta_m \varphi(\Delta_m^2)) \leq d_{3,\delta}(|\mathbf{z}|^{-1/2} + \Theta_{1/2}(\Theta_{1/2}^{-1}(|\mathbf{z}|^{-1/2}))) = 2d_{3,\delta}|\mathbf{z}|^{-1/2}, \end{aligned}$$

that allows us to write

$$\begin{aligned} \|\sigma_3\|_{\mathbf{K}} &= \mathcal{O}(\alpha_{\mathbf{z}}^{-1}|\mathbf{z}|^{-1/2}) = \mathcal{O}(\alpha_{\mathbf{z}}^{-1}\Theta(\Theta^{-1}(|\mathbf{z}|^{-1/2}))) \\ &= \mathcal{O}([\Theta^{-1}(|\mathbf{z}|^{-1/2})]^{-1}\varphi(\Theta^{-1}(|\mathbf{z}|^{-1/2}))\Theta^{-1}(|\mathbf{z}|^{-1/2})) = \mathcal{O}(\varphi(\Theta^{-1}(|\mathbf{z}|^{-1/2}))). \end{aligned}$$

Using again the same argument as for $\|\sigma_{2,1}\|_\rho$ we obtain

$$\|\sigma_3\|_\rho = \mathcal{O}\left(\varphi(\Theta^{-1}(|\mathbf{z}|^{-1/2}))\sqrt{\Theta^{-1}(|\mathbf{z}|^{-1/2})}\right).$$

Summing up the above bounds for $\|\sigma_i\|$, $i = 1, 2, 3$, we prove the statement of the theorem. $\square$

3 Dealing with uncertainty in the sampling size $|\mathbf{z}^\nu|$

Theorem 2 contains a recipe for choosing the subsampling size $|\mathbf{z}^\nu|$ depending on the regularity of the target function and on the approximation power of the corresponding projection method. Both of them, especially the first, may not be exactly given in the form described above. Then several subsampling sizes $|\mathbf{z}^{\nu_1}|, |\mathbf{z}^{\nu_2}|, \dots, |\mathbf{z}^{\nu_l}|$ may be tried in Nyström method, provided that $|\mathbf{z}^{\nu_i}| = o(|\mathbf{z}|^{1/2})$, $i = 1, 2, \dots, l$. Of course, the number l of possible size candidates should not be too large to allow a calculation of all corresponding approximants $f_{\mathbf{z}, \mathbf{z}^{\nu_1}}^\alpha, f_{\mathbf{z}, \mathbf{z}^{\nu_2}}^\alpha, \dots, f_{\mathbf{z}, \mathbf{z}^{\nu_l}}^\alpha$ with a subquadratic complexity. Nevertheless, the question appears of how to select a good approximant among the calculated ones, or how to use all of them. This question is similar to the one in the regularization theory, where some strategy for aggregating all calculated regularized approximants has been discussed recently [4]. In [7] the strategy [4] has been adjusted in the context of learning and presented in several versions.

According to the simplest version, the intention is to approximate the vector $c^* = (c_1^*, c_2^*, \dots, c_l^*) \in \mathbb{R}^l$ solving the following minimization problem

$$\left\| f^\dagger - \sum_{i=1}^l c_i f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \right\|_\rho \rightarrow \min. \quad (21)$$

Recall that $\|\cdot\|_\rho$ is the norm of the Hilbert space $L_2(X, \rho_X)$. Therefore, (21) is equivalent to the matrix problem

$$Gc = g^\dagger, \quad (22)$$

where G and $g^\dagger$ are respectively a Gram matrix and a vector of inner products $\langle \cdot, \cdot \rangle_\rho$ in $L_2(X, \rho_X)$, i.e.

$$G = \left(\left\langle f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha, f_{\mathbf{z}, \mathbf{z}^{\nu_j}}^\alpha \right\rangle_\rho \right)_{i,j=1}^l, \quad g^\dagger = \left(\left\langle f^\dagger, f_{\mathbf{z}, \mathbf{z}^{\nu_j}}^\alpha \right\rangle_\rho \right)_{j=1}^l \quad (23)$$

Note that neither Gram matrix G nor the vector $g^\dagger$ is accessible, since the target function $f^\dagger$ is unknown and the marginal probability distribution ρ_X , which is involved in the definition of $\langle \cdot, \cdot \rangle_\rho$, is not assumed to be given.

On the other hand, $f^\dagger, f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha, i = 1, 2, \dots, l$, belong to the space $\mathcal{H}_K$. That is assumed to be continuously embedded into $L_2(X, \rho_X)$. Then, for example,

$$\begin{aligned} \langle f^\dagger, f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \rangle_\rho &= \langle J_K f^\dagger, J_K f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \rangle_K = \langle C f^\dagger, f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \rangle_K \\ &= \langle (C - C_{\mathbf{z}}) f^\dagger, f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \rangle_K + \langle C_{\mathbf{z}} f^\dagger - S_{\mathbf{z}}^* \mathbf{Y}, f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \rangle_K + \langle S_{\mathbf{z}}^* \mathbf{Y}, f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \rangle_K \end{aligned} \quad (24)$$

In view of (9) the first term of the last equality (24) can be estimated as follows:

$$\begin{aligned} \left| \langle (C - C_{\mathbf{z}}) f^\dagger, f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \rangle_K \right| &\leq \|C - C_{\mathbf{z}}\|_{\mathcal{H}_K \rightarrow \mathcal{H}_K} \cdot \|f^\dagger\|_K \cdot \|f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha\|_K \\ &\leq d_{1,\delta} |\mathbf{z}|^{-1/2} \|f^\dagger\|_K \cdot \|f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha\|_K \end{aligned} \quad (25)$$

Moreover, the norm $\|f^\dagger\|_K$ does not depend on $|\mathbf{z}|, |\mathbf{z}^{\nu_i}|$, and the norm $\|f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha\|_K$ can be controlled. So, with a high probability it holds

$$\left| \langle (C - C_{\mathbf{z}}) f^\dagger, f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \rangle_K \right| = \mathcal{O}(|\mathbf{z}|^{-1/2}). \quad (26)$$

In the same way, with the use of (10) we have

$$\left| \langle C f^\dagger - S_{\mathbf{z}}^* \mathbf{Y}, f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \rangle_K \right| = \mathcal{O}(|\mathbf{z}|^{-1/2}). \quad (27)$$

As to the third term of the last equality (24), it can be directly calculated from the training data since

$$\langle S_{\mathbf{z}}^* \mathbf{Y}, f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \rangle_K = \langle \mathbf{Y}, S_{\mathbf{z}} f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \rangle_{\mathbb{R}^{|\mathbf{z}|}} = |\mathbf{z}|^{-1} \sum_{k=1}^{|\mathbf{z}|} y_k f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha(x_k) \quad (28)$$

Therefore, from (24)–(28) we have with high probability

$$\langle f^\dagger, f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \rangle_\rho = |\mathbf{z}|^{-1} \sum_{k=1}^{|\mathbf{z}|} y_k f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha(x_k) + \mathcal{O}(|\mathbf{z}|^{-1/2}), i = 1, 2, \dots, l. \quad (29)$$

Similar reasoning gives us the relations

$$\langle f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha, f_{\mathbf{z}, \mathbf{z}^{\nu_j}}^\alpha \rangle_\rho = |\mathbf{z}|^{-1} \sum_{k=1}^{|\mathbf{z}|} f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha(x_k) f_{\mathbf{z}, \mathbf{z}^{\nu_j}}^\alpha(x_k) + \mathcal{O}(|\mathbf{z}|^{-1/2}), i, j = 1, 2, \dots, l. \quad (30)$$

In view of (29), (30) the matrix

$$\tilde{G} = \left(|\mathbf{z}|^{-1} \sum_{k=1}^{|\mathbf{z}|} f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha(x_k) f_{\mathbf{z}, \mathbf{z}^{\nu_j}}^\alpha(x_k) \right)_{i,j=1}^l$$

and the vector

$$\tilde{g} = \left(|\mathbf{z}|^{-1} \sum_{k=1}^{|\mathbf{z}|} y_k f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha(x_k) \right)_{i=1}^l$$

can be considered as approximations of G and $g^\dagger$ respectively. Moreover, with probability at least $1 - \delta$

$$\|G - \tilde{G}\|_{\mathbb{R}^l} = \mathcal{O}\left(|\mathbf{z}|^{-1/2} \log \frac{1}{\delta}\right), \quad \|g^\dagger - \tilde{g}\|_{\mathbb{R}^l} = \mathcal{O}\left(|\mathbf{z}|^{-1/2} \log \frac{1}{\delta}\right).$$

With the matrix $\tilde{G}$ in hand one can easily test whether or not $\tilde{G}^{-1}$ exists. For sufficiently large $|\mathbf{z}|$ in case of positive test result a standard perturbation argument (see, e.g. [7] for details) implies the invertibility of G^{-1} , the existence of the vectors $c^* = G^{-1}g^\dagger$, $\tilde{c} = \tilde{G}^{-1}\tilde{g}$ and the bound

$$\|c^* - \tilde{c}\|_{\mathbb{R}^l} = \mathcal{O}\left(|\mathbf{z}|^{-1/2} \log \frac{1}{\delta}\right)$$

that holds with probability at least $1 - \delta$.

Consider now the function

$$f_{\mathbf{z}}^* = \sum_{i=1}^l c_i^* f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha,$$

that solves (21), and its approximation

$$\tilde{f}_{\mathbf{z}}^* = \sum_{i=1}^l \tilde{c}_i^* f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha,$$

where $\tilde{c}_i$, $i = 1, 2, \dots, l$, are the components of the vector $\tilde{c} = \tilde{G}^{-1}\tilde{g}$. Since $f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha$, $i = 1, 2, \dots, l$, are up to our choice, their norms can be controlled such that

$$\left\| f_{\mathbf{z}}^* - \tilde{f}_{\mathbf{z}} \right\|_\rho \leq l \max_i \left\| f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \right\|_\rho \|c^* - \tilde{c}\|_{\mathbb{R}^l} = \mathcal{O} \left(|\mathbf{z}|^{-1/2} \log \frac{1}{\delta} \right).$$

This gives us the following statement

Theorem 3. *Assume that $\tilde{G}$ is invertible and consider $\tilde{f}_{\mathbf{z}} = \sum_{i=1}^l \tilde{c}_i f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha$, $\tilde{c} = (\tilde{c}_i)_{i=1}^l = \tilde{G}^{-1}\tilde{g}$. Then under the conditions of Theorem 2 for sufficiently large $|\mathbf{z}|$ we have with probability at least $1 - \delta$*

$$\left\| f^\dagger - \tilde{f}_{\mathbf{z}} \right\|_\rho = \min_{c_i} \left\| f^\dagger - \sum_{i=1}^l c_i f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha \right\|_\rho + \mathcal{O} \left(|\mathbf{z}|^{-1/2} \log \frac{1}{\delta} \right),$$

where a coefficient implicit in $\mathcal{O}$ -symbol may depend on the cardinality l of the family $\{f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha\}$ and on the distribution ρ , but does not depend on $|\mathbf{z}|$ and δ .

Note that in Theorem 3 the term $\mathcal{O} \left(|\mathbf{z}|^{-1/2} \log \frac{1}{\delta} \right)$ is negligible because, as we know from [3], $|\mathbf{z}|^{-1/2}$ is of higher order than the best guaranteed accuracy of a reconstruction of the target function $f^\dagger \in \mathcal{H}_{\mathbf{K}}$ in $L_2(X, \rho_X)$ from a training set $\mathbf{z}$.

Thus, Theorem 3 tells us that the effectively constructed linear combination of the candidates $f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha$, $i = 1, 2, \dots, l$, is almost as accurate as the best linear aggregator of them.

In the next section we present some numerical experiments illustrating the performance of the aggregator $\tilde{f}_{\mathbf{z}}$.

4 Numerical experiments

For our first experiment we simulate data in the same way as in [17], where another strategy for learning with big data called divide and conquer algorithm

or distributed learning has been studied. Following that paper, we simulate training data sets $\mathbf{z} = \{(x_i, y_i)\}_{i=1}^{|\mathbf{z}|}$, $|\mathbf{z}| \in \{2^8, 2^9, \dots, 2^{13}\}$ from the regression model $y_i = f^\dagger(x_i) + \xi_i$, $i = 1, 2, \dots, |\mathbf{z}|$, where $f^\dagger(x) = \min\{x, 1-x\}$, the random samples x_i are uniformly distributed over $[0, 1]$, and the noise random variables ξ_i are normally distributed with zero mean and variance $\sigma^2 = 1/5$. This simulated problem can be seen as a supervised learning with $X = [0, 1]$ and $\rho_X = \text{Uni}[0, 1]$.

As in [17], all kernel ridge regression estimators appearing in this experiment are constructed in $\mathcal{H}_K$ with $K(x, x') = 1 + \min\{x, x'\}$ and $\alpha = |\mathbf{z}|^{-2/3}$.

We perform plain Nyström subsampling and construct estimators $f_{\mathbf{z}, \mathbf{z}^{\nu_1}}^\alpha$, $f_{\mathbf{z}, \mathbf{z}^{\nu_2}}^\alpha$ with $|\mathbf{z}^{\nu_1}| = \lfloor |\mathbf{z}|^{4/10} \rfloor$ and $|\mathbf{z}^{\nu_2}| = \lfloor |\mathbf{z}|^{3/10} \rfloor$, such that the computational complexity of their construction is of order $o(|\mathbf{z}|^2)$, i.e. subquadratic. Then, as has been discussed in Theorem 3, we construct the aggregator $\tilde{f}_{\mathbf{z}} = \tilde{c}_1 f_{\mathbf{z}, \mathbf{z}^{\nu_1}}^\alpha + \tilde{c}_2 f_{\mathbf{z}, \mathbf{z}^{\nu_2}}^\alpha$.

The accuracy of $\tilde{f}_{\mathbf{z}}$ is compared with the one of divide and conquer algorithm [17]. That algorithm is based on splitting a large training set $\mathbf{z}$ into p much smaller equal-sized subsets $\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_p$, $|\mathbf{z}_i| = \lfloor |\mathbf{z}|/p \rfloor$, $i = 1, 2, \dots, p$; then, each data set $\mathbf{z}_i$ is used as a training set for constructing the minimizer $f_{\mathbf{z}_i}^\alpha$ of (1), where $\mathbf{z}$ is substituted for $\mathbf{z}_i$; finally, the approximations $f_{\mathbf{z}_i}^\alpha$, $i = 1, 2, \dots, p$, are aggregated linearly with equal coefficients (averaged) into

$$f_{\mathbf{z}, p}^\alpha = p^{-1} \sum_{i=1}^p f_{\mathbf{z}_i}^\alpha.$$

In our experiment we compare the errors $\|f^\dagger - \tilde{f}_{\mathbf{z}}\|$, $\|f^\dagger - f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha\|$, $i = 1, 2$, and $\|f^\dagger - f_{\mathbf{z}, p}^\alpha\|$. As in [17] we consider $p = 1, 4, 16, 64$, and execute each simulation 20 times to obtain average values of the errors. In Figure 1 we plot these values versus the total number of samples $|\mathbf{z}|$, where the values corresponding to $\|f^\dagger - \tilde{f}_{\mathbf{z}}\|$, $\|f^\dagger - f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha\|$, and $\|f^\dagger - f_{\mathbf{z}, p}^\alpha\|$ are respectively depicted by dotted, dashed and solid lines.

Figure 1 shows that in the considered case the aggregated approximation $\tilde{f}_{\mathbf{z}}$ outperforms all others, including the baseline KRR-solution $f_{\mathbf{z}, 1}^\alpha$ constructed by the full sample $\mathbf{z}$. It is also interesting to note, that the Nyström approximation $f_{\mathbf{z}, \mathbf{z}^{\nu_2}}^\alpha$, $|\mathbf{z}^{\nu_2}| = \lfloor |\mathbf{z}|^{3/10} \rfloor$, performs poorly, but the aggregated approximation $\tilde{f}_{\mathbf{z}}$ automatically uses the best of available options.

In our second experiment we follow the paper [11], where the dataset `pumadyn32nh` and `cpuSmall` have been used for an empirical study of the

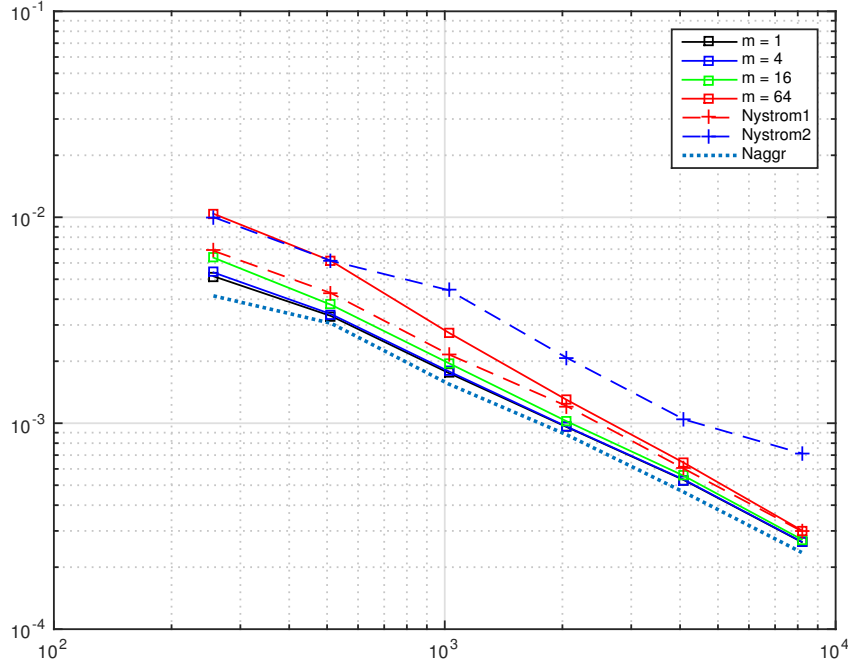


Figure 1: The mean square error between $f^\dagger$ and the averaged estimate $f_{\mathbf{z},p}^\alpha$ for $p = 1, 4, 16, 64$ (solid), Nyström solutions $f_{\mathbf{z},\mathbf{z}^{\nu_1}}^\alpha$ ($|\mathbf{z}^{\nu_1}| = \lfloor |\mathbf{z}|^{4/10} \rfloor$) and $f_{\mathbf{z},\mathbf{z}^{\nu_2}}^\alpha$, $|\mathbf{z}^{\nu_2}| = \lfloor |\mathbf{z}|^{3/10} \rfloor$ (dashed) and aggregated solution $\tilde{f}_{\mathbf{z}}$ (dotted)

Nyström subsampling method. These datasets have been splitted in training and test sets and Gaussian kernels $K(x, x') = \exp(-\|x - x'\|^2 / 2\sigma^2)$ have been used in construction of $f_{\mathbf{z},\mathbf{z}^\nu}^\alpha$. Moreover, 20% of the training points have been hold out for tuning such parameters as σ and α , and the performance of the selected models has been reported on the test sets.

In [11] the performance has been measured in particular by comparing the root-mean-square-errors (RMSE) of the approximations $f_{\mathbf{z},\mathbf{z}^{\nu_1}}^\alpha$, $f_{\mathbf{z},\mathbf{z}^{\nu_2}}^\alpha$ with large $|\mathbf{z}^{\nu_1}|$ and small $|\mathbf{z}^{\nu_2}|$.

It turns out that in the case of `cpuSmall` the effectiveness of the Nyström subsampling is not so high, since comparable values of RMSE of $f_{\mathbf{z},\mathbf{z}^{\nu_1}}^\alpha$, $f_{\mathbf{z},\mathbf{z}^{\nu_2}}^\alpha$ have been observed when both $|\mathbf{z}^{\nu_1}|$, $|\mathbf{z}^{\nu_2}|$, as well as $|\mathbf{z}|$, are of order of 10^3 .

At the same time, in the case of `pumadyn32nh` the same RMSE of 0.033

has been observed for $f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha$, $i = 1, 2$, with $|\mathbf{z}^{\nu_1}| = 1000$ and $|\mathbf{z}^{\nu_2}| = 62$.

Such different performances may hardly be explained by different capacities of the used hypothesis spaces $\mathcal{H}_K$, because in both considered cases they are generated by Gaussian kernels, and, moreover, the dimension of the input space X for `cpuSmall` is smaller than in case of `pumadyn32nh`.

In our Theorem 2 one may find a plausible explanation of the above mentioned behaviour of Nyström approximations. Namely, that is because of the regularities of the target functions corresponding to `pumadyn32nh` and `cpuSmall` are described by source condition (6) with functions φ tending to zero with essentially different rates. This is an example of how Theorem 2 can be used for interpreting empirical results and explaining limitations of the Nyström approach.

Now we use `pumadyn32nh` dataset for illustrating the performance of the aggregators $\tilde{f}_{\mathbf{z}}$. As in [11] we construct the approximants $f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha$, $i = 1, 2, 3$, in $\mathcal{H}_K$ generated by the Gaussian kernel of width $\sigma = 2.66$, and we use $\alpha = 10^{-7}$, $|\mathbf{z}| = 4096$, $|\mathbf{z}^{\nu_1}| = 200$, $|\mathbf{z}^{\nu_2}| = 60$, $|\mathbf{z}^{\nu_3}| = 20$. Table 1 reports the performance of $f_{\mathbf{z}, \mathbf{z}^{\nu_i}}^\alpha$, $i = 1, 2, 3$, and $\tilde{f}_{\mathbf{z}}$.

Approximant	RMSE
$f_{\mathbf{z}, \mathbf{z}^{\nu_1}}^\alpha$	0.03381
$f_{\mathbf{z}, \mathbf{z}^{\nu_2}}^\alpha$	0.03325
$f_{\mathbf{z}, \mathbf{z}^{\nu_3}}^\alpha$	0.03442
Aggregator $\tilde{f}_{\mathbf{z}}$	0.03325

Table 1: Performance of Nyström approximants and their aggregator on a testing set of 4096 data points from `cpuSmall`

As can be seen from Table 1, the aggregation approach described in Section 3 again automatically uses the best of the available options and can be recommended as a reliable strategy to be implemented together with the Nyström subsampling when dealing with uncertainty in the subsampling size.

Acknowledgements

The authors affiliated with Johann Radon Institute for Computational and Applied Mathematics (RICAM) gratefully acknowledge the support of the the Austrian Science Fund (FWF), grant P25424.

References

- [1] Bach F 2013 Sharp analysis of low-rank kernel matrix approximations *JMLR W&CP* **30** 185–209
- [2] Bauer F, Pereverzev S and Rosasco L 2007 On regularization algorithms in learning theory *Journal of Complexity* **23** 52–72 ISSN 0885-064X
- [3] Caponnetto A and De Vito E 2007 Optimal rates for the regularized least-squares algorithm *Found. Comput. Math.* **7** 331–368
- [4] Chen J, Pereverzyev Jr S and Xu Y 2015 Aggregation of regularized solutions from multiple observation models *Inverse Problems* **31** 075005
- [5] De Vito E, Rosasco L, Caponnetto A, De Giovannini U and Odone F 2005 Learning from examples as an inverse problem *J. Mach. Learn. Res.* **6** 883–904
- [6] Kimeldorf G S and Wahba G 1970 A correspondence between bayesian estimation on stochastic processes and smoothing by splines *The Annals of Mathematical Statistics* **41** 495–502 ISSN 00034851
- [7] Kriukova G, Panasiuk O, Pereverzyev S V and Tkachenko P 2016 A linear functional strategy for regularized ranking *Neural Networks* **73** 26–35 ISSN 0893-6080
- [8] Mathé P and Hofmann B 2008 How general are general source conditions? *Inverse Problems* **24** 015009
- [9] Mathé P and Pereverzev S V 2003 Discretization strategy for linear ill-posed problems in variable Hilbert scales *Inverse Problems* **19** 1263–1277
- [10] Plato R and Vainikko G 1990 On the regularization of projection methods for solving ill-posed problems *Numerische Mathematik* **57** 63–79
- [11] Rudi A, Camoriano R and Rosasco L 2015 Less is more: Nyström computational regularization *Advances in Neural Information Processing Systems 28* ed Cortes C, Lawrence N, Lee D, Sugiyama M, Garnett R and Garnett R (Curran Associates, Inc.) pp 1648–1656 also arXiv:1507.04717 [stat.ML]

- [12] Schölkopf B and Smola A J 2001 *Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond* (MIT Press) ISBN 0262194759
- [13] Shawe-Taylor J and Cristianini N 2004 *Kernel Methods for Pattern Analysis* (Cambridge University Press)
- [14] Smale S and Zhou D X 2007 Learning theory estimates via integral operators and their approximations *Constructive Approximation* **26** 153–172 ISSN 0176-4276
- [15] Smola A J and Schölkopf B 2000 Sparse greedy matrix approximation for machine learning *Proceedings of the Seventeenth International Conference on Machine Learning (ICML 2000), Stanford University, Stanford, CA, USA, June 29 - July 2, 2000* pp 911–918
- [16] Williams C and Seeger M 2001 Using the Nyström method to speed up kernel machines *Proceedings of the 14th Annual Conference on Neural Information Processing Systems* pp 682–688
- [17] Zhang Y, Duchi J C and Wainwright M J 2013 Divide and conquer kernel ridge regression: A distributed algorithm with minimax optimal rates *arXiv:1305.5029 [math.ST]* Also *JMLR W&CP* 30: 592–617