

Optimal learning rates for Kernel Conjugate Gradient regression

Gilles Blanchard Nicole Krämer

Weierstrass Institute for

Applied Analysis and Stochastics

Mohrenstr. 39, 10117 Berlin, Germany

`{gilles.blanchard;nicole.kraemer}@wias-berlin.de`

September 30, 2010

Abstract

We prove rates of convergence in the statistical sense for kernel-based least squares regression using a conjugate gradient algorithm, where regularization against overfitting is obtained by early stopping. This method is directly related to Kernel Partial Least Squares, a regression method that combines supervised dimensionality reduction with least squares projection. The rates depend on two key quantities: first, on the regularity of the target regression function and second, on the intrinsic dimensionality of the data mapped into the kernel space. Lower bounds on attainable rates depending on these two quantities were established in earlier literature, and we obtain upper bounds for the considered method that match these lower bounds (up to a log factor) if the true regression function belongs to the reproducing kernel Hilbert space. If this assumption is not fulfilled, we obtain similar convergence rates provided additional unlabeled data are available. The order of the learning rates match state-of-the-art results that were recently obtained for least squares support vector machines and for linear regularization operators.

1 Introduction

The contribution of this paper is the learning theoretical analysis of kernel-based least squares regression in combination with conjugate gradient techniques. The goal is to estimate a regression function f^* based on random noisy observations. We have an i.i.d. sample of n observations $(X_i, Y_i) \in \mathcal{X} \times \mathbb{R}$ from an unknown distribution $P(X, Y)$ that follows the model

$$Y = f^*(X) + \varepsilon,$$

where ε is a noise variable whose distribution can possibly depend on X , but satisfies $\mathbb{E}[\varepsilon|X] = 0$. We assume that the true regression function f^* belongs to the space $\mathcal{L}_2(P_X)$ of square-integrable functions. Following the kernelization principle, we implicitly map the data into a reproducing kernel Hilbert space $\mathcal{H}$ with a kernel k . We denote by $K_n = \frac{1}{n}(k(X_i, X_j)) \in \mathbb{R}^{n \times n}$ the normalized kernel matrix and by $\mathbf{Y} = (Y_1, \dots, Y_n)^\top \in \mathbb{R}^n$ the n -vector of response observations. The task is to find coefficients α such that the function defined by the normalized kernel expansion

$$f_\alpha(X) = \frac{1}{n} \sum_{i=1}^n \alpha_i k(X_i, X)$$

is an adequate estimator of the true regression function f^* . The closeness of the estimator f_α to the target f^* is measured via the $\mathcal{L}_2(P_X)$ distance,

$$\begin{aligned} \|f_\alpha - f^*\|_2^2 &= \mathbb{E}_{X \sim P_X} [(f_\alpha(X) - f^*(X))^2] \\ &= \mathbb{E}_{XY} [(f_\alpha(X) - Y)^2] - \mathbb{E}_{XY} [(f^*(X) - Y)^2], \end{aligned}$$

The last equality recalls that this criterion is the same as the excess generalization error for the squared error loss $\ell(f, x, y) = (f(x) - y)^2$.

In empirical risk minimization, we use the training data empirical distribution as a proxy for the generating distribution, and minimize the *training* squared error. This gives rise to the linear equation

$$K_n \alpha = \mathbf{Y} \quad \text{with } \alpha \in \mathbb{R}^n. \quad (1)$$

Assuming K_n invertible, the solution of the above equation is given by $\alpha = K_n^{-1} \mathbf{Y}$, which yields a function in $\mathcal{H}$ interpolating perfectly the training data but having poor generalization error. It is well-known that to avoid overfitting, some form of regularization is needed. There is a considerable variety of possible approaches (see e.g. [12] for an overview). Perhaps the most well-known one is

$$\alpha = (K_n + \lambda I)^{-1} \mathbf{Y}, \quad (2)$$

known alternatively as kernel ridge regression, least squares support vector machine, or Tikhonov's regularization. A powerful generalization of this is to consider

$$\alpha = F_\lambda(K_n) \mathbf{Y}, \quad (3)$$

where $F_\lambda : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ is a fixed function depending on a parameter λ . The notation $F_\lambda(K_n)$ is to be interpreted as F_λ applied to each eigenvalue of K_n in its eigen decomposition. Intuitively, F_λ should be a "regularized" version of the inverse function $F(x) = x^{-1}$. This type of regularization, which we refer to as linear regularization methods, is directly inspired from the theory of inverse problems. Popular examples include as particular cases kernel Ridge Regression, Principal components regression and L_2 -Boosting. Their application in

a learning context has been studied extensively [2, 3, 6, 7, 14]. Results obtained in this framework will serve as a comparison yardstick in the sequel.

In this paper, we study conjugate gradient (CG) techniques in combination with early stopping for the regularization of the kernel based learning problem (1). The principle of CG techniques is to restrict the learning problem onto a nested set of data-dependent subspaces, the so-called Krylov subspaces, defined as

$$\mathcal{K}_m(\mathbf{Y}, K_n) = \text{span} \{ \mathbf{Y}, K_n \mathbf{Y}, \dots, K_n^{m-1} \mathbf{Y} \}. \quad (4)$$

Denote by $\langle \cdot, \cdot \rangle$ the usual euclidean scalar product on $\mathbb{R}^n$ rescaled by the factor n^{-1} . We define the K_n -norm as $\|\alpha\|_{K_n}^2 = \langle \alpha, K_n \alpha \rangle$. The CG solution after m iterations is formally defined as

$$\alpha_m = \arg \min_{\alpha \in \mathcal{K}_m(\mathbf{Y}, K_n)} \|\mathbf{Y} - K_n \alpha\|_{K_n}; \quad (5)$$

and the number m of CG iterations is the model parameter. To simplify notation we define $f_m := f_{\alpha_m}$. In the learning context considered here, regularization corresponds to early stopping. Conjugate gradients have the appealing property that the optimization criterion (5) can be computed by a simple iterative algorithm that constructs basis vectors $d_1, \dots, d_m$ of $\mathcal{K}_m(\mathbf{Y}, K_n)$ by using only *forward multiplication* of vectors by the matrix K_n . Algorithm 1 displays the computation of the CG kernel coefficients α_m defined by (5).

Algorithm 1 Kernel Conjugate Gradient regression

Input kernel matrix K_n , response vector $\mathbf{Y}$, maximum number of iterations m

Initialization: $\alpha_0 = \mathbf{0}_n$; $r_1 = \mathbf{Y}$; $d_1 = \mathbf{Y}$

for $i = 1, \dots, m$ **do**

$d_i = d_i / \|K_n d_i\|_{K_n}$ (normalization of the basis vector)

$\gamma_i = \langle \mathbf{Y}, K_n d_i \rangle_{K_n}$ (step size)

$\alpha_i = \alpha_{i-1} + \gamma_i d_i$ (update)

$r_{i+1} = r_i - K_n \gamma_i d_i$ (residuals)

$d_{i+1} = r_{i+1} - \langle K_n d_i, r_{i+1} \rangle_{K_n} / \|r_{i+1}\|_{K_n}^2$ (new basis vector)

end for

Return: CG kernel coefficients α_m , CG function $f_m = \sum_{i=1}^n \alpha_{i,m} k(X_i, \cdot)$

The CG approach is also inspired by the theory of inverse problems, but it is not covered by the framework of linear operators defined in (3): As we restrict the learning problem onto the Krylov space $\mathcal{K}_m(\mathbf{Y}, K_n)$, the CG coefficients α_m are of the form $\alpha_m = q_m(K_n) \mathbf{Y}$ with q_m a polynomial of degree $\leq m - 1$. However, the polynomial q_m is not fixed but depends on $\mathbf{Y}$ as well, making the CG method nonlinear in the sense that the coefficients α_m depend on $\mathbf{Y}$ in a nonlinear fashion.

We remark that in machine learning, conjugate gradient techniques are often used as fast solvers for operator equations, e.g. to obtain the solution for the regularized equation

(2). We stress that in this paper, we study conjugate gradients as a *regularization approach* for kernel based learning, where the regularity is ensured via early stopping. This approach is not new. As mentioned in the abstract, the algorithm that we study is closely related to Kernel Partial Least Squares [21]. The latter method also restricts the learning problem onto the Krylov subspace $\mathcal{K}_m(\mathbf{Y}, K_n)$, but it minimizes the euclidean distance $\|\mathbf{Y} - K_n\alpha\|$ instead of the distance $\|\mathbf{Y} - K_n\alpha\|_{K_n}$ defined above¹. Kernel Partial Least Squares has shown competitive performance in benchmark experiences (see e.g [21, 22]). Moreover, a similar conjugate gradient approach for non-definite kernels has been proposed and empirically evaluated by Ong et al [19]. The focus of the current paper is therefore not to stress the usefulness of CG methods in practical applications (and we refer to the above mentioned references) but to examine its theoretical convergence properties. In particular, we establish the existence of early stopping rules that lead to optimal convergence rates. We summarize our main results in the next session.

2 Main results

For the presentation of our convergence results, we require suitable assumptions on the learning problem. We first assume that the kernel space $\mathcal{H}$ is separable and that the kernel function is measurable. (This assumption is satisfied for all practical situations that we know of.) Furthermore, for all results, we make the (relatively standard) assumption that the kernel is *bounded*:

$$k(x, x) \leq \kappa \text{ for all } x \in \mathcal{X}. \quad (6)$$

We consider – depending on the result – one of the following assumptions on the *noise*:

(Bounded) (Bounded Y): $|Y| \leq M$ almost surely.

(Bernstein) (Bernstein condition): $\mathbb{E}[\varepsilon^p | X] \leq (1/2)p!M^p$ almost surely, for all integers $p \geq 2$.

The second assumption is weaker than the first. In particular, the first assumption implies that not only the noise, but also the target function f^* is bounded in supremum norm, while the second assumption does not put any additional restriction on the target function.

The *regularity* of the target function f^* is measured in terms of a *source condition* as follows. The kernel integral operator is given by

$$K : \mathcal{L}_2(P_X) \rightarrow \mathcal{L}_2(P_X), g \mapsto \int k(., x)g(x)dP(x).$$

The **source condition** for the parameter $r \geq 0$ is defined by:

$$\mathbf{SC}(r) : f^* = K^r u \quad \text{with} \quad \|u\| \leq \kappa^{-r} \rho.$$

¹This is generalized to a CG- l algorithm ($l \in \mathbb{N}_{\geq 0}$) by replacing the K_n -norm in (5) with the norm defined by K_n^l . Corresponding fast iterative algorithms to compute the solution exist for all l (see e.g. [13]).

It is a known fact (see, e.g., [10]) that if $r \geq 1/2$, then f^* coincides almost surely with a function belonging to $\mathcal{H}_k$. Therefore, we refer to $r \geq 1/2$ as the “inner case” and to $r < 1/2$ as the “outer case”.

The regularity of the kernel operator K with respect to the marginal distribution P_X is measured in terms of the **intrinsic dimensionality** parameter, defined by the condition

$$\mathbf{ID}(s) : \text{tr}(K(K + \lambda I)^{-1}) \leq D^2(\kappa^{-1}\lambda)^{-s} \text{ for all } \lambda \in (0, 1].$$

It is known that the best attainable rates of convergence, as a function of the number of examples n , is determined by the parameters r and s . It was shown in [11] that the minimax learning rate given these two parameters is lower bounded by $\mathcal{O}(n^{-2r/(2r+s)})$.

We now expose our main results in different situations. In all the cases considered, the early stopping rule takes the form of a so-called **discrepancy stopping rule**: For some sequence of thresholds Λ_m to be specified (and possibly depending on the data), define the (data-dependent) stopping iteration $\hat{m}$ as the first iteration m for which

$$\left\| (f_m(X_1), \dots, f_m(X_n))^{\top} - \mathbf{Y} \right\|_{K_n} < \Lambda_m. \quad (7)$$

(Only in the first result below, the threshold Λ_m actually depends on the iteration m and on the data.)

2.1 Inner case without knowledge on intrinsic dimension

The inner case corresponds to $r \geq 1/2$, i.e. the target function f^* lies in $\mathcal{H}$ almost surely. For some constants $\tau > 1$ and $1 > \gamma > 0$, we consider the discrepancy stopping rule with the threshold sequence

$$\Lambda_m = 4\tau \sqrt{\frac{\kappa \log(2\gamma^{-1})}{n}} \left(\sqrt{\kappa} \|\alpha_m\|_{K_n} + M \sqrt{\log(2\gamma^{-1})} \right). \quad (8)$$

For technical reasons, we consider a slight variation of the rule in that we stop at step $\hat{m} - 1$ instead of $\hat{m}$ if $q_{\hat{m}}(0) \geq 4\kappa \sqrt{\log(2\gamma^{-1})/n}$, where q_m is the iteration polynomial such that $\alpha_m = q_m(K_n)\mathbf{Y}$. Denote $\tilde{m}$ the resulting stopping step. We obtain the following result.

Theorem 2.1. *Suppose that Y is bounded (**Bounded**), and that the source condition **SC**(r) holds for $r \geq 1/2$. With probability $1 - 2\gamma$, the estimator $f_{\tilde{m}}$ obtained by the (modified) discrepancy stopping rule (8) satisfies*

$$\|f_{\tilde{m}} - f^*\|_2^2 \leq c(r, \tau)(M + \rho)^2 \left(\frac{\log^2 \gamma^{-1}}{n} \right)^{\frac{2r}{2r+1}}.$$

We present the proof in Section 4.

2.2 Optimal rates in inner case

We now introduce a stopping rule yielding order-optimal convergence rates as a function of the two parameters r and s in the “inner” case ($r \geq 1/2$, which is equivalent to saying that the target function belongs to $\mathcal{H}$ almost surely). For some constant $\tau' > 3/2$ and $1 > \gamma > 0$, we consider the discrepancy stopping rule with the fixed threshold

$$\Lambda_m \equiv \Lambda = \tau' M \sqrt{\kappa} \left(\frac{4D}{\sqrt{n}} \log \frac{6}{\gamma} \right)^{\frac{2r+1}{2r+s}}. \quad (9)$$

for which we obtain the following:

Theorem 2.2. *Suppose that the noise fulfills the Bernstein assumption (**Bernstein**), that the source condition **SC**(r) holds for $r \geq 1/2$, and that **ID**(s) holds. With probability $1 - 3\gamma$, the estimator $f_{\hat{m}}$ obtained by the discrepancy stopping rule (9) satisfies*

$$\|f_{\hat{m}} - f^*\|_2^2 \leq c(r, \tau')(M + \rho)^2 \left(\frac{16D^2}{n} \log^2 \frac{6}{\gamma} \right)^{\frac{2r}{2r+s}}.$$

Due to space limitations, the proof is presented in the supplementary material.

2.3 Optimal rates in outer case, given additional unlabeled data

We now turn to the “outer” case. In this case, we make the additional assumption that *unlabeled* data is available. Assume that we have $\tilde{n}$ i.i.d. observations $X_1, \dots, X_{\tilde{n}}$, out of which only the first n are labeled. We define a new response vector $\tilde{\mathbf{Y}} = \frac{\tilde{n}}{n} (Y_1, \dots, Y_n, 0, \dots, 0) \in \mathbb{R}^{\tilde{n}}$ and run the CG algorithm 1 on $X_1, \dots, X_{\tilde{n}}$ and $\tilde{\mathbf{Y}}$. We use the same threshold (9) as in the previous section for the stopping rule, except that the factor M is replaced by $\max(M, \rho)$.

Theorem 2.3. *Suppose assumptions (**Bounded**), **SC**(r) and **ID**(s), with $r + s \geq \frac{1}{2}$. Assume unlabeled data is available with*

$$\frac{\tilde{n}}{n} \geq \left(\frac{16D^2}{n} \log^2 \frac{6}{\gamma} \right)^{-\frac{(1-2r)_+}{2r+s}}.$$

Then with probability $1 - 3\gamma$, the estimator $f_{\hat{m}}$ obtained by the discrepancy stopping rule defined above satisfies

$$\|f_{\hat{m}} - f^*\|_2^2 \leq c(r, \tau')(M + \rho)^2 \left(\frac{16D^2}{n} \log^2 \frac{6}{\gamma} \right)^{\frac{2r}{2r+s}}.$$

A sketch of the proof can be found in the supplementary material.

3 Discussion and comparison to other results

For the inner case – i.e. $f^* \in \mathcal{H}$ almost surely – we provide two different consistent stopping criteria. The first one (Section 2.1) is oblivious to the intrinsic dimension parameter s , and the obtained bound corresponds to the “worst case” with respect to this parameter (that is, $s = 1$). However, an interesting feature of stopping rule (8) is that the rule itself does not depend on the a priori knowledge of the regularity parameter r , while the achieved learning rate does (and with the optimal dependence in r when $s = 1$). Hence, Theorem 2.1 implies that the obtained rule is automatically *adaptive* with respect to the regularity of the target function. This contrasts with the results obtained in [2] for linear regularization schemes of the form (3), (also in the case $s = 1$) for which the choice of the regularization parameter λ leading to optimal learning rates required the knowledge of r beforehand.

When taking into account also the intrinsic dimensionality parameter s , Theorem 2.2 provides the order-optimal convergence rate in the inner case (up to a log factor). A noticeable difference to Theorem 2.1 however, is that the stopping rule is no longer adaptive, that is, it depends on the *a priori* knowledge of parameters r and s . We observe that previously obtained results for linear regularization schemes of the form (2) in [7] and of the form (3) in [6], also rely on the a priori knowledge of r and s to determine the appropriate regularization parameter λ .

The outer case – when the target function does not lie in the reproducing Kernel Hilbert space $\mathcal{H}$ – is more challenging and to some extent less well understood. The fact that additional assumptions are made is not a particular artefact of CG methods, but also appears in the studies of other regularization techniques. Here we follow the semi-supervised approach that is proposed in e.g. [6] (to study linear regularization of the form (3)) and assume that we have sufficient additional unlabeled data in order to ensure learning rates that are optimal as a function of the number of *labeled* data. We remark that other forms of additional requirements can be found in the recent literature in order to reach optimal rates. For regularized M-estimation schemes studied in [23], availability of unlabeled data is not required, but a condition is imposed of the form $\|f\|_\infty \leq C \|f\|_{\mathcal{H}}^p \|f\|_2^{1-p}$ for all $f \in \mathcal{H}$ and some $p \in (0, 1]$. In [15], assumptions on the supremum norm of the eigenfunctions of the kernel integral operator are made (see [23] for an in-depth discussion on this type of assumptions).

Finally, as explained in the introduction, the term ‘conjugate gradients’ comprises a class of methods that approximate the solution of linear equations on Krylov subspaces. In the context of learning, our approach is most closely linked to Partial Least Squares (PLS) [24] and its kernel extension [21]. While PLS has proven to be successful in a wide range of applications and is considered one of the standard approaches in chemometrics, there are only few studies of its theoretical properties. In [9, 16], consistency properties are provided for linear PLS under the assumption that the target function f^* depends on a finite known number of orthogonal latent components. These findings were recently extended to the nonlinear case and without the assumption of a latent components model

[4], but all results come without optimal rates of convergence. For the slightly different CG approach studied by Ong et al [19], bounds on the difference between the empirical risks of the CG approximation and of the target function are derived in [18], but no bounds on the generalization error were derived.

4 Proofs

Convergence rates for regularization methods of the type (2) or (3) have been studied by casting kernel learning methods into the framework of *inverse problems* (see [11]). We use this framework for the present results as well, and recapitulate here some important facts.

We first define the *empirical evaluation operator* T_n as follows:

$$T_n : \quad g \in \mathcal{H} \mapsto T_n g := (g(X_1), \dots, g(X_n))^\top \in \mathbb{R}^n$$

and the *empirical integral operator* T_n^* as:

$$T_n^* : u = (u_1, \dots, u_n) \in \mathbb{R}^n \mapsto T_n^* u := \frac{1}{n} \sum_{i=1}^n u_i k(X_i, \cdot) \in \mathcal{H}.$$

Using the reproducing property of the kernel, it can be readily checked that T_n and T_n^* are adjoint operators, i.e. they satisfy $\langle T_n^* u, g \rangle_{\mathcal{H}} = \langle u, T_n g \rangle$, for all $u \in \mathbb{R}^n, g \in \mathcal{H}$. Furthermore, $K_n = T_n T_n^*$, and therefore $\|\alpha\|_{K_n} = \|f_\alpha\|_{\mathcal{H}}$. Based on these facts, equation (5) can be rewritten as

$$f_m = \arg \min_{f \in \mathcal{K}_m(T_n^* \mathbf{Y}, S_n)} \|T_n^* Y - S_n f\|_{\mathcal{H}}, \quad (10)$$

where $S_n = T_n^* T_n$ is a self-adjoint operator of $\mathcal{H}$, called empirical covariance operator. This definition corresponds to that of the “usual” conjugate gradient algorithm formally applied to the so-called normal equation (in $\mathcal{H}$)

$$S_n f_\alpha = T_n^* \mathbf{Y},$$

which is obtained from (1) by left multiplication by T_n^* . The advantage of this reformulation is that it can be interpreted as a “perturbation” of a *population, noiseless* version (of the equation and of the algorithm), wherein Y is replaced by the target function f^* and the empirical operator T_n^*, T_n are respectively replaced by their population analogues, the kernel integral operator

$$T^* : g \in L_2(P_X) \mapsto T^* g := \int k(\cdot, x) g(x) dP_X(x) = \mathbb{E} [k(X, \cdot) g(X)] \in \mathcal{H},$$

and the change-of-space operator

$$T : \quad g \in \mathcal{H} \mapsto g \in \mathcal{L}_2(P_X).$$

The latter maps a function to itself but between two Hilbert spaces which differ with respect to their geometry – the inner product of $\mathcal{H}$ being defined by the kernel function k , while the inner product of $\mathcal{L}_2(P_X)$ depends on the data generating distribution (this operator is well defined: since the kernel is bounded, all functions in $\mathcal{H}$ are bounded and therefore square integrable under any distribution P_X).

The following results, taken from [2] (Propositions 21 and 22) quantify more precisely that the empirical covariance operator $S_n = T_n^* T_n$ and the empirical integral operator applied to the data, $T_n^* \mathbf{Y}$, are close to the population covariance operator $S = T^* T$ and to the kernel integral operator applied to the noiseless target function, $T^* f^*$ respectively.

Proposition 4.1. *Provided that condition (6) is true, the following holds:*

$$\mathbb{P} \left[\|S_n - S\|_{HS} \leq \frac{4\kappa}{\sqrt{n}} \sqrt{\log \frac{2}{\gamma}} \right] \geq 1 - \gamma, \quad (11)$$

where $\|\cdot\|_{HS}$ denotes the Hilbert-Schmidt norm. If the representation $f^* = T f_{\mathcal{H}}^*$ holds, and under assumption **(Bernstein)**, we have the following:

$$\mathbb{P} \left[\|T_n^* \mathbf{Y} - S f_{\mathcal{H}}^*\| \leq \frac{4M\sqrt{\kappa}}{\sqrt{n}} \log \frac{2}{\gamma} \right] \geq 1 - \gamma. \quad (12)$$

We note that $f^* = T f_{\mathcal{H}}^*$ implies that the target function f^* coincides with a function $f_{\mathcal{H}}^*$ belonging to $\mathcal{H}$ (remember that T is just the change-of-space operator). Hence, the second result (12) is valid for the case with $r \geq 1/2$, but it is not true in general for $r < 1/2$.

4.1 Nemirovskii's result on conjugate gradient regularization rates

We recall a sharp result due to Nemirovskii [17] establishing convergence rates for conjugate gradient methods in a deterministic context. We present the result in an abstract context, then show how, combined with the previous section, it leads to a proof of Theorem 2.1. Consider the linear equation

$$Az^* = b,$$

where A is a bounded linear operator over a Hilbert space $\mathcal{H}$. Assume that the above equation has a solution and denote z^* its minimal norm solution; assume further that a self-adjoint operator $\bar{A}$, and an element $\bar{b} \in \mathcal{H}$ are known such that

$$\|A - \bar{A}\| \leq \delta; \quad \|b - \bar{b}\| \leq \varepsilon, \quad (13)$$

(with δ and ε known positive numbers). Consider the CG algorithm based on the noisy operator $\bar{A}$ and data $\bar{b}$, giving the output at step m

$$z_m = \text{Arg Min}_{z \in \mathcal{K}_m(\bar{A}, \bar{b})} \|\bar{A}z - \bar{b}\|^2. \quad (14)$$

The *discrepancy principle* stopping rule is defined as follows. Consider a fixed constant $\tau > 1$ and define

$$\bar{m} = \min \{m \geq 0 : \|\bar{A}z_m - \bar{b}\| < \tau(\delta \|z_m\| + \varepsilon)\}.$$

We output the solution obtained at step $\max(0, \bar{m} - 1)$. Consider a minor variation of this rule:

$$\hat{m} = \begin{cases} \bar{m} & \text{if } q_{\bar{m}}(0) < \eta\delta^{-1} \\ \max(0, \bar{m} - 1) & \text{otherwise,} \end{cases}$$

where $q_{\bar{m}}$ is the degree $m - 1$ polynomial such that $z_{\bar{m}} = q_{\bar{m}}(\bar{A})\bar{b}$, and η is an arbitrary positive constant such that $\eta < 1/\tau$. Nemirovskii established the following theorem:

Theorem 4.2. *Assume that (a) $\max(\|A\|, \|\bar{A}\|) \leq L$; and that (b) $z^* = A^\mu u^*$ with $\|u^*\| \leq R$ for some $\mu > 0$. Then for any $\theta \in [0, 1]$, provided that $\hat{m} < \infty$ it holds that*

$$\|A^\theta(z_{\hat{m}} - z^*)\|^2 \leq c(\mu, \tau, \eta) R^{\frac{2(1-\theta)}{1+\mu}} (\varepsilon + \delta RL^\mu)^{2(\theta+\mu)/(1+\mu)}.$$

4.2 Proof of Theorem 2.1

We apply Nemirovskii's result in our setting (assuming $r \geq \frac{1}{2}$): By identifying the approximate operator and data as $\bar{A} = S_n$ and $\bar{b} = T_n^* \mathbf{Y}$, we see that the CG algorithm considered by Nemirovskii (14) is exactly (10), more precisely with the identification $z_m = f_m$.

For the population version, we identify $A = S$, and $z^* = f_{\mathcal{H}}^*$ (remember that provided $r \geq \frac{1}{2}$ in the source condition, then there exists $f_{\mathcal{H}}^* \in \mathcal{H}$ such that $f^* = T f_{\mathcal{H}}^*$).

Condition (a) of Nemirovskii's theorem 4.2 is satisfied with $L = \kappa$ by the boundedness of the kernel. Condition (b) is satisfied with $\mu = r - 1/2 \geq 0$ and $R = \kappa^{-r} \rho$, as implied by the source condition **SC**(r). Finally, the concentration result 4.1 ensures that the approximation conditions (13) are satisfied with probability $1 - 2\gamma$, more precisely with $\delta = \frac{4\kappa}{\sqrt{n}} \sqrt{\log \frac{2}{\gamma}}$ and $\varepsilon = \frac{4M\sqrt{\kappa}}{\sqrt{n}} \log \frac{2}{\gamma}$. (Here we replaced γ in (11) and (12) by $\gamma/2$, so that the two conditions are satisfied simultaneously, by the union bound). The operator norm is upper bounded by the Hilbert-Schmidt norm, so that the deviation inequality for the operators is actually stronger than what is needed.

We consider the discrepancy principle stopping rule associated to these parameters, the choice $\eta = 1/(2\tau)$, and $\theta = \frac{1}{2}$, thus obtaining the result, since

$$\left\| A^{\frac{1}{2}}(z_{\hat{m}} - z^*) \right\|^2 = \left\| S^{\frac{1}{2}}(f_{\hat{m}} - f_{\mathcal{H}}^*) \right\|_{\mathcal{H}}^2 = \|f_{\hat{m}} - f_{\mathcal{H}}^*\|_2^2.$$

4.3 Notes on the proof of Theorems 2.2 and 2.3

The above proof shows that an application of Nemirovskii’s fundamental result for CG regularization of inverse problems under deterministic noise (on the data and the operator) allows us to obtain our first result. One key ingredient is the concentration property 4.1 which allows to bound deviations in a quasi-deterministic manner.

To prove the sharper results of Theorems 2.2 and 2.3, such a direct approach does not work unfortunately, and a complete rework and extension of the proof is necessary. The proof of Theorem 2.2 is presented in the supplementary material to the paper. In a nutshell, the concentration result 4.1 is too coarse to prove the optimal rates of convergence taking into account the intrinsic dimension parameter. Instead of that result, we have to consider the deviations from the mean in a “warped” norm, i.e. of the form $\left\| (S + \lambda I)^{-\frac{1}{2}} (T_n^* \mathbf{Y} - T^* f^*) \right\|$ for the data, and $\left\| (S + \lambda I)^{-\frac{1}{2}} (S_n - S) \right\|_{HS}$ for the operator (with an appropriate choice of $\lambda > 0$) respectively. Deviations of this form were introduced and used in [6, 7] to obtain sharp rates in the framework of Tikhonov’s regularization (2) and of the more general linear regularization schemes of the form (3). Bounds on deviations of this form can be obtained via a Bernstein-type concentration inequality for Hilbert-space valued random variables.

On the one hand, the results concerning linear regularization schemes of the form (3) do not apply to the nonlinear CG regularization. On the other hand, Nemirovskii’s result does not apply to deviations controlled in the warped norm. Moreover, the “outer” case introduces additional technical difficulties. Therefore, the proofs for Theorems 2.2 and 2.3, while still following the overall fundamental structure and ideas introduced by Nemirovskii, are significantly different in that context. As mentioned above, we present the complete proof of Theorem 2.2 in the supplementary material and a sketch of the proof of Theorem 2.3.

5 Conclusion

In this work, we derived early stopping rules for kernel Conjugate Gradient regression that provide optimal learning rates to the true target function. Depending on the situation that we study, the rates are adaptive with respect to the regularity of the target function in some cases. The proofs of our results rely most importantly on ideas introduced by Nemirovskii [17] and further developed by Hanke [13] for CG methods in the deterministic case, and moreover on ideas inspired by [6, 7].

Certainly, in practice, as for a large majority of learning algorithms, cross-validation remains the standard approach for model selection. The motivation of this work is however mainly theoretical, and our overall goal is to show that from the learning theoretical point of view, CG regularization stands on equal footing with other well-studied regularization methods such as kernel ridge regression or more general linear regularization methods (which

includes between many others L_2 boosting). We also note that theoretically well-grounded model selection rules can generally help cross-validation in practice by providing a well-calibrated parametrization of regularizer functions, or, as is the case here, of thresholds used in the stopping rule.

One crucial property used in the proofs is that the proposed CG regularization schemes can be conveniently cast in the reproducing kernel Hilbert space $\mathcal{H}$ as displayed in e.g (10). This reformulation is not possible for Kernel Partial Least Squares: It is also a CG type method, but uses the standard Euclidean norm instead of the K_n -norm used here. This point is the main technical justification on why we focus on (5) rather than kernel PLS. Obtaining optimal convergence rates also valid for Kernel PLS is an important future direction and should build on the present work.

Another important direction for future efforts is the derivation of stopping rules that do not depend on the confidence parameter γ . Currently, this dependence prevents us to go from convergence in high probability to convergence in expectation, which would be desirable. Perhaps more importantly, it would be of interest to find a stopping rule that is adaptive to both parameters r (target function regularity) and s (intrinsic dimension parameter) without their a priori knowledge. We recall that our first stopping rule is adaptive to r but at the price of being worst-case in s . In the literature on linear regularization methods, the optimal choice of regularization parameter is also non-adaptive, be it when considering optimal rates with respect to r only [2] or to both r and s [6]. An approach to alleviate this problem is to use a hold-out sample for model selection; this was studied theoretically in [8] for linear regularization methods (see also [5] for an account of the properties of hold-out in a general setup). We strongly believe that the hold-out method will yield theoretically founded adaptive model selection for CG as well. However, hold-out is typically regarded as inelegant in that it requires to throw away part of the data for estimation. It would be of more interest to study model selection methods that are based on using the whole data in the estimation phase. The application of Lepskii's method is a possible step towards this direction.

References

- [1] R. Bathia. *Matrix Analysis*, volume 169 of *Graduate texts in mathematics*. Springer, 1997.
- [2] F. Bauer, S. Pereverzev, and L. Rosasco. On Regularization Algorithms in Learning Theory. *Journal of Complexity*, 23:52–72, 2007.
- [3] N. Bissantz, T. Hohage, A. Munk, and F. Ruymgaart. Convergence Rates of General Regularization Methods for Statistical Inverse Problems and Applications. *SIAM Journal on Numerical Analysis*, 45(6):2610–2636, 2007.

- [4] G. Blanchard and N. Krämer. Kernel Partial Least Squares is Universally Consistent. *Proceedings of the 13th International Conference on Artificial Intelligence and Statistics, JMLR Workshop & Conference Proceedings*, 9:57–64, 2010.
- [5] G. Blanchard and P. Massart. Discussion of V. Koltchinskii’s 2004 IMS Medallion Lecture paper, ”Local Rademacher complexities and oracle inequalities in risk minimization”. *Annals of Statistics*, 34(6):2664–2671, 2006.
- [6] A. Caponnetto. Optimal Rates for Regularization Operators in Learning Theory. Technical Report CBCL Paper 264/ CSAIL-TR 2006-062, Massachusetts Institute of Technology, 2006.
- [7] A. Caponnetto and E. De Vito. Optimal Rates for Regularized Least-squares Algorithm. *Foundations of Computational Mathematics*, 7(3):331–368, 2007.
- [8] A. Caponnetto and Y. Yao. Cross-validation based Adaptation for Regularization Operators in Learning Theory. *Analysis and Applications*, 8(2):161–183, 2010.
- [9] H. Chun and S. Keles. Sparse Partial Least Squares for Simultaneous Dimension Reduction and Variable Selection. *Journal of the Royal Statistical Society B*, 72(1):3–25, 2010.
- [10] F. Cucker and S. Smale. On the Mathematical Foundations of Learning. *Bulletin of the American Mathematical Society*, 39(1):1–50, 2002.
- [11] E. De Vito, L. Rosasco, A. Caponnetto, U. De Giovannini, and F. Odone. Learning from Examples as an Inverse Problem. *Journal of Machine Learning Research*, 6(1):883, 2006.
- [12] L. Györfi, M. Kohler, A. Krzyżak, and H. Walk. *A Distribution-Free Theory of Non-parametric Regression*. Springer, 2002.
- [13] M. Hanke. *Conjugate Gradient Type Methods for Linear Ill-posed Problems*. Pitman Research Notes in Mathematics Series, 327, 1995.
- [14] L. Lo Gerfo, L. Rosasco, E. Odone, F. and De Vito, and A. Verri. Spectral Algorithms for Supervised Learning. *Neural Computation*, 20:1873–1897, 2008.
- [15] S. Mendelson and J. Neeman. Regularization in Kernel Learning. *The Annals of Statistics*, 38(1):526–565, 2010.
- [16] P. Naik and C.L. Tsai. Partial Least Squares Estimator for Single-index Models. *Journal of the Royal Statistical Society B*, 62(4):763–771, 2000.

- [17] A. S. Nemirovskii. The Regularizing Properties of the Adjoint Gradient Method in Ill-posed Problems. *USSR Computational Mathematics and Mathematical Physics*, 26(2):7–16, 1986.
- [18] C. S. Ong. Kernels: Regularization and Optimization. *Doctoral dissertation, Australian National University*, 2005.
- [19] C. S. Ong, X. Mary, S. Canu, and A. J. Smola. Learning with Non-positive Kernels. In *Proceedings of the 21st International Conference on Machine Learning*, pages 639 – 646, 2004.
- [20] I. F. Pinelis and A. I. Sakhanenko. Remarks on inequalities for probabilities of large deviations. *Theory Probab. Appl.*, 1(30):143–148, 1985.
- [21] R. Rosipal and L.J. Trejo. Kernel Partial Least Squares Regression in Reproducing Kernel Hilbert Spaces. *Journal of Machine Learning Research*, 2:97–123, 2001.
- [22] R. Rosipal, L.J. Trejo, and B. Matthews. Kernel PLS-SVC for Linear and Nonlinear Classification. In *Proceedings of the Twentieth International Conference on Machine Learning*, pages 640–647, Washington, DC, 2003.
- [23] I. Steinwart, D. Hush, and C. Scovel. Optimal Rates for Regularized Least Squares Regression. In *Proceedings of the 22nd Annual Conference on Learning Theory*, pages 79–93, 2009.
- [24] S. Wold, H. Ruhe, H. Wold, and W.J. Dunn III. The Collinearity Problem in Linear Regression. The Partial Least Squares (PLS) Approach to Generalized Inverses. *SIAM Journal of Scientific and Statistical Computations*, 5:735–743, 1984.
- [25] V. Yurinski. *Sums and Gaussian vectors*, volume 1617 of *Lecture notes in mathematics*. Springer, 1995.

A Supplementary Material

A.1 Notation

We follow the notation used in the main part, in particular the operators T_n, T_n^*, T, T^* defined in Section 4.1, and we recall that $S_n := T_n^* T_n$; $S = T^* T$; $K_n = T_n T_n^*$; and $K = T T^*$.

We denote by $(\xi_i)_{i \geq 1}$ the possibly finite sequence in $[0, \kappa]$ of nonzero eigenvalues of S and K , and by $(\xi_{j,n})_{1 \leq j \leq n}$ the n -sequence of eigenvalues of S_n and K_n respectively (in each case in decreasing order and with multiplicity). Finally, $(F_u)_{u \geq 0}$ denotes the spectral family of the operator S_n , i.e. F_u is the orthogonal projector on the subspace of $\mathcal{H}$ spanned by eigenvectors of S_n corresponding to eigenvalues strictly less than u .

It is useful to consider the spectral integral representation: If $(e_{i,n})_{1 \leq i \leq n}$ denotes the orthogonal eigensystem of S_n associated to the non-zero eigenvalues $(\lambda_{i,n})_{1 \leq i \leq n}$, for any integrable function h on $[0, \kappa]$, we set

$$\int_0^\kappa h(u) d \|F_{u,n} T_n^* \mathbf{Y}\|^2 := \langle T_n^* \mathbf{Y}, h(S_n) T_n^* \mathbf{Y} \rangle = \sum_{i=1}^n h(\lambda_{i,n}) \langle T_n^* \mathbf{Y}, e_{i,n} \rangle^2.$$

By its definition, the output of the m -th iteration of the CG algorithm can be put under the form $f_m = q_m(S_n) T_n^* \mathbf{Y}$, where $q_m \in \mathcal{P}_{m-1}$, the set of real polynomials of degree less than $m-1$. A crucial role is played by the *residual polynomial*

$$p_m(x) = 1 - x q_m(x) \in \mathcal{P}_m^0,$$

where $\mathcal{P}_m^0$ is the set of real polynomials of degree less than m and having constant term equal to 1. In particular $T_n^* \mathbf{Y} - S_n f_m = p_m(S_n) T_n^* \mathbf{Y}$. Furthermore, the definition of the CG algorithm implies that the sequence $(p_m)_{m \geq 0}$ are orthogonal polynomials for the scalar product $[\cdot, \cdot]_{(1)}$, where for $i \geq 0$ we define

$$[p, q]_{(i)} := \langle p(S_n) T_n^* \mathbf{Y}, S_n^i q(S_n) T_n^* \mathbf{Y} \rangle = \int_0^\kappa p(u) q(u) u^i d \|F_{u,n} T_n^* \mathbf{Y}\|^2.$$

This can be shown as follows: p_m is the orthogonal projection, of the origin onto the affine space $\mathcal{P}_m^0 = 1 + x \mathcal{P}_{m-1}$ with the scalar product $[\cdot, \cdot]_{(0)}$, where $x \mathcal{P}_{m-1}$ denotes (with some abuse of notation) the set of polynomials of degree less than m with constant coefficient equal to zero. Thus $0 = [p_m, xq]_{(0)} = [p_m, q]_{(1)}$ for any $q \in \mathcal{P}_{m-1}$. From the theory of orthogonal polynomials, it results that for any $m \leq m_{final} := \#\{i : 1 \leq i \leq n, \xi_{i,n} \langle T_n^* \mathbf{Y}, e_{i,n} \rangle \neq 0\}$, the polynomial p_m has exactly m distinct roots belonging to $[0, \kappa]$, which we denote by $(x_{k,m})_{1 \leq k \leq m}$ (in increasing order). Finally, we use the notation $c(a, b)$ to denote a function depending on the stated parameters only, and whose exact value can change from line to line.

A.2 Preparation of the proof

We follow the general architecture of Nemirovskii's proof to establish rates. We recall that since we assume $r \geq 1/2$, the representation $f^* = T f_{\mathcal{H}}^*$ holds. The main difference to Nemirovskii's original result is that (similar to the approach of [6, 7]) we use deviation bounds in a “warped” norm rather than in the standard norm. More precisely, we consider the following type of assumptions:

$$\mathbf{B1}(\lambda) \quad \left\| (S + \lambda I)^{-\frac{1}{2}} (T_n^* \mathbf{Y} - S_n f_{\mathcal{H}}^*) \right\| \leq \delta(\lambda),$$

B2(λ) $\|(S + \lambda I)(S_n + \lambda I)^{-1}\| \leq \Upsilon^2$, with $\Upsilon \geq 1$
 (this implies in particular $\|(S + \lambda I)^{\frac{1}{2}}(S_n + \lambda I)^{-\frac{1}{2}}\| \leq \Upsilon$ via (34) below),

B3 $\|S - S_n\| \leq \kappa \Delta$.

In the rest of this section we set $\mu = r - 1/2$. Under the source condition assumption **SC**(r), for $r \geq \frac{1}{2}$ the representation $f^* = K^r u$ can be rewritten

$$f^* = (TT^*)^r u = T(T^*T)^{r-\frac{1}{2}}(T^*T)^{-\frac{1}{2}}T^*u = TS^\mu(T^*T)^{-\frac{1}{2}}T^*u,$$

by identification we therefore have the source condition for $f_{\mathcal{H}}$ given by $f_{\mathcal{H}} = S^\mu w$ with $w = (T^*T)^{-\frac{1}{2}}T^*u$, and $\|w\|_{\mathcal{H}} \leq \|u\|$, since $(T^*T)^{-\frac{1}{2}}T^*$ is a restricted isometry from $L_2(P_X)$ into $\mathcal{H}$.

We define the shortcut notation

$$Z_\mu(\lambda) = \begin{cases} \lambda^\mu & \text{for } \mu \leq 1, \\ \kappa^\mu \Delta & \text{for } \mu > 1. \end{cases} \quad (15)$$

We start with preliminary technical lemmas, before turning to the proof of Theorem 2.2.

Lemma A.1. *For any $\lambda > 0$, if assumptions **SC**(r), **B1**(λ), **B2**(λ) and **B3** hold, then for any iteration step $1 \leq m \leq m_{\text{final}}$*

$$\begin{aligned} \|T_n^*(T_n f_m - \mathbf{Y})\| &\leq c(\mu) \Upsilon^2 \left(|p'_m(0)|^{-(\mu+1)} + Z_\mu(\lambda) |p'_m(0)|^{-1} \right) \kappa^{-\mu-\frac{1}{2}} \rho \\ &\quad + \left(|p'_m(0)|^{-\frac{1}{2}} + \lambda^{\frac{1}{2}} \right) \Upsilon \delta(\lambda). \end{aligned} \quad (16)$$

Proof. Recall that $(x_{k,m})_{1 \leq k \leq m}$ denote the m roots of the polynomial p_m ; define further the function φ_m on the interval $[0, x_{1,m}]$ as

$$\varphi_m(x) = p_m(x) \left(\frac{x_{1,m}}{x_{1,m} - x} \right)^{\frac{1}{2}}$$

Following the idea introduced by Nemirovski, it can be shown that

$$\begin{aligned} \|T_n^*(T_n f_m - \mathbf{Y})\| &= \|p_m(S_n)T_n^*\mathbf{Y}\| \\ &\leq \|F_{x_{1,m}}\varphi_m(S_n)T_n^*\mathbf{Y}\| \\ &\leq \|F_{x_{1,m}}\varphi_m(S_n)S_n f_{\mathcal{H}}^*\| + \|F_{x_{1,m}}\varphi_m(S_n)(T_n^*\mathbf{Y} - S_n f_{\mathcal{H}}^*)\| := (I) + (II). \end{aligned}$$

Above, the first inequality (lemma 3.7. in [13]) is the crucial point, and relies fundamentally on the fact that (p_m) is an orthogonal polynomial sequence.

We start with controlling the second term:

$$\begin{aligned}
(II) &= \|F_{x_{1,m}}\varphi_m(S_n)(T_n^*\mathbf{Y} - S_nf_{\mathcal{H}}^*)\| \\
&= \|F_{x_{1,m}}\varphi_m(S_n)(S + \lambda I)^{\frac{1}{2}}(S + \lambda I)^{-\frac{1}{2}}(T_n^*\mathbf{Y} - S_nf_{\mathcal{H}}^*)\| \\
&\leq \|F_{x_{1,m}}\varphi_m(S_n)(S_n + \lambda I)^{\frac{1}{2}}\| \Upsilon\delta(\lambda) \\
&\leq \left(\sup_{x \in [0, x_{1,m}]} x^{\frac{1}{2}}\varphi_m(x) + \lambda^{\frac{1}{2}} \sup_{x \in [0, x_{1,m}]} \varphi_m(x) \right) \Upsilon\delta(\lambda) \\
&\leq \left(|p'_m(0)|^{-\frac{1}{2}} + \lambda^{\frac{1}{2}} \right) \Upsilon\delta(\lambda),
\end{aligned}$$

where the last line used the inequality (see (3.10) in [13])

$$\sup_{x \in [0, x_{1,m}]} x^\nu \varphi_m^2(x) \leq \nu^\nu |p'_m(0)|^{-\nu}, \quad (17)$$

for any $\nu \geq 0$ (using the convention $0^0 = 1$), which we applied above for $\nu = 0, 1$. For the first term, we use assumption **SC**(r); first consider the case $\mu > 1$:

$$\begin{aligned}
(I) &= \|F_{x_{1,m}}\varphi_m(S_n)S_nf_{\mathcal{H}}^*\| = \|F_{x_{1,m}}\varphi_m(S_n)S_nS^\mu w\| \\
&\leq (\|F_{x_{1,m}}\varphi_m(S_n)S_n^{\mu+1}\| + \|F_{x_{1,m}}\varphi_m(S_n)S_n\| \|S^\mu - S_n^\mu\|) \kappa^{-\mu-\frac{1}{2}}\rho \\
&\leq c(\mu) \left(|p'_m(0)|^{-(\mu+1)} + \kappa^\mu \Delta |p'_m(0)|^{-1} \right) \kappa^{-\mu-\frac{1}{2}}\rho,
\end{aligned}$$

where we applied (17) with $\nu = 2(\mu + 1)$, $\nu = 2$ and (33). $\square$

For the case $\mu \leq 1$, using (34) and arguments similar to the previous case:

$$\begin{aligned}
(I) &= \|F_{x_{1,m}}\varphi_m(S_n)S_nf_{\mathcal{H}}^*\| \\
&= \|F_{x_{1,m}}\varphi_m(S_n)S_nS^\mu w\| \\
&\leq \|F_{x_{1,m}}\varphi_m(S_n)S_n(S_n + \lambda I)^\mu\| \|(S_n + \lambda I)^{-\mu}(S + \lambda I)^\mu\| \|(S + \lambda I)^{-\mu}S^\mu\| \kappa^{-\mu-\frac{1}{2}}\rho \\
&\leq c(\mu)\Upsilon^2 \left(|p'_m(0)|^{-(\mu+1)} + \lambda^\mu |p'_m(0)|^{-1} \right) \kappa^{-\mu-\frac{1}{2}}\rho.
\end{aligned}$$

Lemma A.2. *For any $\lambda > 0$, if assumptions **SC**(r), **B1**(λ), **B2**(λ) and **B3** hold, then for any iteration step $1 \leq m \leq m_{\text{final}}$, for any $\varepsilon \in (0, x_{1,m})$:*

$$\begin{aligned}
\|T(f_m - f_{\mathcal{H}}^*)\| &\leq \Upsilon \left(3(1 + \lambda(|p'_m(0)| + \varepsilon^{-1})) \Upsilon\delta(\lambda) + c(\mu)\Upsilon^2 \left(\varepsilon^{\frac{1}{2}} + \lambda^{\frac{1}{2}} \right) (\varepsilon^\mu + Z_\mu(\lambda)) \kappa^{-\mu-\frac{1}{2}}\rho \right. \\
&\quad \left. + \sqrt{2} \left(1 + \frac{\lambda^{\frac{1}{2}}}{\varepsilon^{\frac{1}{2}}} \right) \varepsilon^{-\frac{1}{2}} \|T_n^*(T_nf_m - \mathbf{Y})\| \right)
\end{aligned}$$

If $m = 0$, the above inequality is valid for any $\varepsilon > 0$.

Proof. Set $\bar{f}_m = q_m(S_n)S_nf_{\mathcal{H}}^*$. This is the element in $\mathcal{H}$ that we obtain by applying the m th-iteration CG polynomial q_m to the *noiseless* data. We have

$$\begin{aligned}\|T(f_m - f_{\mathcal{H}}^*)\| &= \left\| S^{\frac{1}{2}}(f_m - f_{\mathcal{H}}^*) \right\| \leq \Upsilon \left\| (S_n + \lambda I)^{\frac{1}{2}}(f_m - f_{\mathcal{H}}^*) \right\| \\ &\leq \Upsilon \left(\left\| F_{\varepsilon}(S_n + \lambda I)^{\frac{1}{2}}(f_m - \bar{f}_m) \right\| + \left\| F_{\varepsilon}(S_n + \lambda I)^{\frac{1}{2}}(\bar{f}_m - f_{\mathcal{H}}^*) \right\| \right. \\ &\quad \left. + \left\| F_{\varepsilon}^{\perp}(S_n + \lambda I)^{\frac{1}{2}}(f_m - f_{\mathcal{H}}^*) \right\| \right) \\ &:= \Upsilon((I) + (II) + (III)),\end{aligned}$$

where we denote $F_{\varepsilon}^{\perp} := (I - F_{\varepsilon})$. *First summand:*

$$\begin{aligned}(I) &= \left\| F_{\varepsilon}(S_n + \lambda I)^{\frac{1}{2}}(f_m - \bar{f}_m) \right\| = \left\| F_{\varepsilon}(S_n + \lambda I)^{\frac{1}{2}}q_m(S_n)(S_n + \lambda I)^{\frac{1}{2}}(S_n + \lambda I)^{-\frac{1}{2}}(T_n^* \mathbf{Y} - S_nf_{\mathcal{H}}^*) \right\| \\ &\leq \Upsilon \left\| F_{\varepsilon}(S_n + \lambda I)^{\frac{1}{2}}q_m(S_n)(S_n + \lambda I)^{\frac{1}{2}} \right\| \delta(\lambda) \\ &\leq \Upsilon \delta(\lambda) \left(\sup_{x \in [0, \varepsilon]} x q_m(x) + \lambda \sup_{x \in [0, \varepsilon]} q_m(x) \right) \\ &\leq \Upsilon \delta(\lambda) (1 + \lambda |p'_m(0)|) .\end{aligned}$$

The last inequality is obtained by the following argument: if $m \geq 1$, since $\varepsilon \leq x_{1,m}$, p_m is convex in $[0, \varepsilon]$, we have

$$q_m(x) = \frac{1 - p_m(x)}{x} \leq |p'_m(0)| \quad \text{for } x \in [0, \varepsilon];$$

and also $xq_m(x) = 1 - p_m(x) \leq 1$ for $x \in [0, \varepsilon]$. If $m = 0$, we have $p_0 \equiv 1$ and $q_m \equiv 0$, so that the above is also trivially satisfied for any x .

Second summand: first subcase, $\mu > 1$, using (33), and the fact that $|p_m|(x) \leq 1$ for $x \in [0, \varepsilon]$:

$$\begin{aligned}(II) &= \left\| F_{\varepsilon}(S_n + \lambda I)^{\frac{1}{2}}(\bar{f}_m - f_{\mathcal{H}}^*) \right\| \\ &= \left\| F_{\varepsilon}(S_n + \lambda I)^{\frac{1}{2}}p_m(S_n)S_n^{\mu}w \right\| \\ &\leq \left(\left\| F_{\varepsilon}(S_n + \lambda I)^{\frac{1}{2}}p_m(S_n)S_n^{\mu} \right\| + \left\| F_{\varepsilon}(S_n + \lambda I)^{\frac{1}{2}}p_m(S_n) \right\| c(\mu)\kappa^{\mu}\Delta \right) \kappa^{-\mu-\frac{1}{2}}\rho \\ &\leq \left(\varepsilon^{\mu+\frac{1}{2}} + \lambda^{\frac{1}{2}}\varepsilon^{\mu} + c(\mu)\kappa^{\mu} \left(\varepsilon^{\frac{1}{2}} + \lambda^{\frac{1}{2}} \right) \Delta \right) \kappa^{-\mu-\frac{1}{2}}\rho \\ &\leq c(\mu) \left(\varepsilon^{\frac{1}{2}} + \lambda^{\frac{1}{2}} \right) (\varepsilon^{\mu} + \kappa^{\mu}\Delta) \kappa^{-\mu-\frac{1}{2}}\rho .\end{aligned}$$

Bounding the second summand: second subcase, $\mu \leq 1$:

$$(II) = \left\| F_\varepsilon(S_n + \lambda I)^{\frac{1}{2}} p_m(S_n) S^\mu w \right\| \leq \left\| F_\varepsilon(S_n + \lambda I)^{\mu + \frac{1}{2}} p_m(S_n) \right\| \Upsilon^2 \kappa^{-\mu - \frac{1}{2}} \rho \\ \leq c(\mu)(\varepsilon + \lambda)^{\mu + \frac{1}{2}} \Upsilon^2 \kappa^{-\mu - \frac{1}{2}} \rho.$$

Third summand:

$$(III) = \left\| F_\varepsilon^\perp(S_n + \lambda I)^{\frac{1}{2}}(f_m - f_{\mathcal{H}}^*) \right\| \leq \left\| F_\varepsilon^\perp S_n^{\frac{1}{2}}(f_m - f_{\mathcal{H}}^*) \right\| + \lambda^{\frac{1}{2}} \left\| F_\varepsilon^\perp(f_m - f_{\mathcal{H}}^*) \right\| \\ \leq \left(\frac{(\varepsilon + \lambda)^{\frac{1}{2}}}{\varepsilon^{\frac{1}{2}}} + \lambda^{\frac{1}{2}} \frac{(\varepsilon + \lambda)^{\frac{1}{2}}}{\varepsilon} \right) \left\| F_\varepsilon^\perp(S_n + \lambda I)^{-\frac{1}{2}} S_n(f_m - f_{\mathcal{H}}^*) \right\| \\ \leq \left(1 + \frac{\lambda^{\frac{1}{2}}}{\varepsilon^{\frac{1}{2}}} \right) \left(1 + \frac{\lambda}{\varepsilon} \right)^{\frac{1}{2}} \left\| F_\varepsilon^\perp(S_n + \lambda I)^{-\frac{1}{2}} S_n(f_m - f_{\mathcal{H}}^*) \right\| \\ \leq \left(1 + \frac{\lambda^{\frac{1}{2}}}{\varepsilon^{\frac{1}{2}}} \right) \left(1 + \frac{\lambda}{\varepsilon} \right)^{\frac{1}{2}} \left(\left\| F_\varepsilon^\perp(S_n + \lambda I)^{-\frac{1}{2}} T_n^*(T_n f_m - \mathbf{Y}) \right\| + \left\| (S_n + \lambda I)^{-\frac{1}{2}}(T_n^* \mathbf{Y} - S_n f_{\mathcal{H}}^*) \right\| \right) \\ \leq \left(1 + \frac{\lambda^{\frac{1}{2}}}{\varepsilon^{\frac{1}{2}}} \right) \varepsilon^{-\frac{1}{2}} \|T_n^*(T_n f_m - \mathbf{Y})\| + \sqrt{2} \Upsilon \left(1 + \frac{\lambda}{\varepsilon} \right) \left\| (S_n + \lambda I)^{-\frac{1}{2}}(T_n^* \mathbf{Y} - S_n f_{\mathcal{H}}^*) \right\| \\ \leq \left(1 + \frac{\lambda^{\frac{1}{2}}}{\varepsilon^{\frac{1}{2}}} \right) \varepsilon^{-\frac{1}{2}} \|T_n^*(T_n f_m - \mathbf{Y})\| + \sqrt{2} \Upsilon \left(1 + \frac{\lambda}{\varepsilon} \right) \delta(\lambda).$$

□

We now consider the sequence of polynomials that are orthogonal with respect to the scalar product $[\cdot, \cdot]_{(2)}$, which we denote by $p_m^{(2)}$, and its roots by $x_m^{(2)}$.

Lemma A.3. *For any $\lambda > 0$, if assumptions **SC**(r), **B1**(λ), **B2**(λ) and **B3** hold, then for any iteration step $1 \leq m \leq m_{\text{final}}$, for any $\varepsilon \in (0, x_{1,m-1})$:*

$$[p_{m-1}, p_{m-1}]_{(0)}^{\frac{1}{2}} = \|p_{m-1}(S_n) T_n^* \mathbf{Y}\| \\ \leq \Upsilon(\varepsilon + \lambda)^{\frac{1}{2}} \delta(\lambda) + c(\mu) \Upsilon^2 \varepsilon (\varepsilon^\mu + Z_\mu(\lambda)) \kappa^{-\mu - \frac{1}{2}} \rho + \varepsilon^{-\frac{1}{2}} \left[p_{m-1}^{(2)}, p_{m-1}^{(2)} \right]_{(1)}^{\frac{1}{2}}. \quad (18)$$

Proof. By the optimality property defining our CG algorithm,

$$\|p_{m-1}(S_n) T_n^* \mathbf{Y}\| \leq \|p_{m-1}^{(2)}(S_n) T_n^* \mathbf{Y}\| \leq \|F_\varepsilon p_{m-1}^{(2)}(S_n) T_n^* \mathbf{Y}\| + \|F_\varepsilon^\perp p_{m-1}^{(2)}(S_n) T_n^* \mathbf{Y}\| \\ \leq \|F_\varepsilon T_n^* \mathbf{Y}\| + \varepsilon^{-\frac{1}{2}} \|p_{m-1}^{(2)}(S_n) S_n^{\frac{1}{2}} T_n^* \mathbf{Y}\| = \|F_\varepsilon T_n^* \mathbf{Y}\| + \varepsilon^{-\frac{1}{2}} \left[p_{m-1}^{(2)}, p_{m-1}^{(2)} \right]_{(1)}^{\frac{1}{2}}$$

For the last inequality, we have used the fact that $|p_{m-1}^{(2)}|(x) \leq 1$ for $x \in [0, x_{m-1}^{(2)}]$, along with the assumption $0 < \varepsilon < x_{1,m-1} \leq x_{1,m-1}^{(2)}$; the latter inequality is due to interlacing properties of the roots of orthogonal polynomials for $[\cdot, \cdot]_{(i)}$ and $[\cdot, \cdot]_{(i+1)}$ (see [13], Cor 2.7). We now bound

$$\begin{aligned} \|F_\varepsilon T_n^* \mathbf{Y}\| &\leq \|F_\varepsilon (T_n^* \mathbf{Y} - S_n f_{\mathcal{H}}^*)\| + \|F_\varepsilon S_n S^\mu w\| \\ &\leq \left\| F_\varepsilon (S_n + \lambda I)^{\frac{1}{2}} \right\| \left\| (S_n + \lambda I)^{-\frac{1}{2}} (T_n^* \mathbf{Y} - S_n f_{\mathcal{H}}^*) \right\| + \|F_\varepsilon S_n S^\mu w\| \\ &\leq \Upsilon(\varepsilon + \lambda)^{\frac{1}{2}} \delta(\lambda) + \|F_\varepsilon S_n S^\mu w\| ; \end{aligned}$$

for the second term, we divide as usual into two cases: for $\mu > 1$:

$$\|F_\varepsilon S_n S^\mu w\| \leq \|F_\varepsilon S_n^{\mu+1} w\| + \|F_\varepsilon S_n (S_n^\mu - S^\mu) w\| \leq \varepsilon c(\mu) (\varepsilon^\mu + \kappa^\mu \Delta) \kappa^{-\mu-\frac{1}{2}} \rho ,$$

and for $\mu \leq 1$:

$$\|F_\varepsilon S_n S^\mu w\| \leq \|F_\varepsilon S_n (S_n + \lambda I)^\mu\| \Upsilon^2 \kappa^{-\mu-\frac{1}{2}} \rho \leq \varepsilon (\varepsilon^\mu + \lambda^\mu) \Upsilon^2 \kappa^{-\mu-\frac{1}{2}} \rho .$$

□

A.3 Proof of Theorem 2.2

We fix

$$\lambda_* = \left((4D/\sqrt{n}) \log(6/\gamma) \right)^{\frac{2}{2\mu+s+1}} \kappa . \quad (19)$$

and assume n is big enough to ensure $\lambda_* \leq \kappa$. Furthermore we denote $\tilde{\lambda}_* = \kappa^{-1} \lambda_*$ (this normalization was introduced in [6]).

We rewrite equivalently the discrepancy stopping rule as follows: for some fixed $\tau > 0$,

$$\hat{m} := \min \left\{ 0 \leq m : \|T_n^* (T_n f_m - \mathbf{Y})\| \leq (2 + \tau) \lambda_*^{\frac{1}{2}} \delta(\lambda_*) \right\} , \quad (20)$$

where

$$\delta(\lambda_*) := \frac{3}{4} M \tilde{\lambda}_*^{\mu+\frac{1}{2}} . \quad (21)$$

(Observe that the above $\tau > 0$ is deduced from the constant $\tau' > 3/2$ considered in the main part of the paper via $\tau = \frac{4}{3}(\tau' - \frac{3}{2})$.)

We first check **B1**(λ_*), **B2**(λ_*) and **B3** are satisfied simultaneously with large probability, using for this concentration results which are recalled in Section A.5. Concerning

$\mathbf{B1}(\lambda_*)$, inequality (31) ensures that with probability $1 - \gamma$, we have

$$\begin{aligned}
\left\| (S + \lambda_* I)^{-\frac{1}{2}} (T_n^* \mathbf{Y} - S_n f_{\mathcal{H}}^*) \right\| &\leq 2M \left(\sqrt{\frac{\mathcal{N}(\lambda_*)}{n}} + \frac{2\sqrt{\kappa}}{\sqrt{\lambda_* n}} \right) \log \frac{6}{\gamma} \\
&\leq \frac{2M}{\sqrt{n}} D \tilde{\lambda}_*^{-\frac{s}{2}} \left(1 + \frac{1}{2D^2} \left(\frac{4D}{\sqrt{n}} \log \frac{6}{\gamma} \right) \tilde{\lambda}_*^{\frac{s-1}{2}} \right) \log \frac{6}{\gamma} \\
&\leq \frac{M}{2} \tilde{\lambda}_*^{\mu+\frac{1}{2}} \left(1 + \frac{1}{2D^2} \tilde{\lambda}_*^{\mu+s} \right) \\
&\leq \frac{3}{4} M \tilde{\lambda}_*^{\mu+\frac{1}{2}} = \delta(\lambda_*), \tag{22}
\end{aligned}$$

where we have used $\mathbf{SC}(r)$, (19) and the assumptions $D \geq 1$ and $\tilde{\lambda}_* \leq 1$. We now turn to $\mathbf{B2}(\lambda_*)$. Inequality (32) along with a repetition of the above reasoning yields that with probability $1 - \gamma$:

$$\left\| (S + \lambda_* I)^{-\frac{1}{2}} (S_n - S) \right\|_{HS} \leq \frac{\sqrt{\kappa}}{M} \delta(\lambda_*),$$

so that

$$\left\| (S + \lambda_* I)^{-\frac{1}{2}} (S_n - S) (S + \lambda_* I)^{-\frac{1}{2}} \right\| \leq \frac{\sqrt{\kappa}}{M} \lambda_*^{-\frac{1}{2}} \delta(\lambda_*).$$

Observe that

$$\frac{\sqrt{\kappa}}{M} \lambda_*^{-\frac{1}{2}} \delta(\lambda_*) = \frac{3}{4} \tilde{\lambda}_*^{\mu} \leq \frac{3}{4}, \tag{23}$$

so that with Lemma A.5, we obtain that $\mathbf{B2}(\lambda_*)$ is satisfied with $\Upsilon := 2$ (with probability $1 - \gamma$). Finally, equation (11) in the main paper implies that $\mathbf{B3}$ is also satisfied with probability $1 - \gamma$, with

$$\Delta := \frac{2}{\sqrt{n}} \log \frac{1}{\gamma}. \tag{24}$$

To conclude, by the union bound, the event that $\mathbf{B1}(\lambda_*)$, $\mathbf{B2}(\lambda_*)$ and $\mathbf{B3}$ satisfied simultaneously has probability larger than $1 - 3\gamma$, and we assume for the rest of the proof that we are on this event.

We will assume $\hat{m} \geq 1$ for the remainder of the proof and postpone to the end the (simpler) case $\hat{m} = 0$.

First step: upper bound on $|p'_{\hat{m}-1}(0)|$. By definition of the stopping rule we have $\|T_n^*(T_n f_{\hat{m}-1} - \mathbf{Y})\| > (2 + \tau) \lambda_*^{\frac{1}{2}} \delta(\lambda_*)$. Now applying this together with the upper bound of Lemma A.1 we get

$$\begin{aligned}
\tau \lambda_*^{\frac{1}{2}} \delta(\lambda_*) &\leq c(\mu) \left(|p'_{\hat{m}-1}(0)|^{-(\mu+1)} + Z_{\mu}(\lambda_*) |p'_{\hat{m}-1}(0)|^{-1} \right) \kappa^{-\mu-\frac{1}{2}} \rho + 2 |p'_{\hat{m}-1}(0)|^{-\frac{1}{2}} \delta(\lambda_*) \\
&\leq 3 \max \left(2 |p'_{\hat{m}-1}(0)|^{-\frac{1}{2}} \delta(\lambda_*), c(\mu) \rho \kappa^{-\mu-\frac{1}{2}} |p'_{\hat{m}-1}(0)|^{-(\mu+1)}, c(\mu) \rho \kappa^{-\mu-\frac{1}{2}} Z_{\mu}(\lambda_*) |p'_{\hat{m}-1}(0)|^{-1} \right).
\end{aligned}$$

We examine in succession the possibility that the maximum in the above expression is attained for each of the terms which comprise it. If the first term attains the maximum, this implies $|p'_{\widehat{m}-1}(0)| \leq (9/\tau^2)\lambda_*^{-1}$. If the second term attains the maximum, this entails

$$c(\mu)\rho\kappa^{-\mu-\frac{1}{2}}|p'_{\widehat{m}-1}(0)|^{-(\mu+1)} \geq \tau\lambda_*^{\frac{1}{2}}\delta(\lambda_*),$$

which using (21) yields:

$$|p'_{\widehat{m}-1}(0)| \leq c(\mu, \tau) \left(\frac{\rho}{M}\right)^{\frac{1}{\mu+1}} \lambda_*^{-1}.$$

Finally, if the third term attains the maximum, we have

$$c(\mu)\rho Z_\mu(\lambda_*)\kappa^{-\mu-\frac{1}{2}}|p'_{\widehat{m}-1}(0)|^{-1} \geq \tau\lambda_*^{\frac{1}{2}}\delta(\lambda_*),$$

which using (21) yields:

$$|p'_{\widehat{m}-1}(0)| \leq c(\mu, \tau) \frac{\rho}{M} \lambda_*^{-\mu-1} Z_\mu(\lambda_*).$$

We now establish the inequality

$$Z_\mu(\lambda_*)\lambda_*^{-\mu} \leq 1. \quad (25)$$

The inequality is trivial if $\mu \leq 1$ given the definition of $Z_\mu(\lambda_*)$ in (15). If $\mu > 1$ holds, from the definition (24), it holds that $\Delta \leq \frac{1}{2}\widetilde{\lambda}_*^{\frac{2\mu+s+1}{2}}$, hence

$$Z_\mu(\lambda_*)\lambda_*^{-\mu} = \Delta\widetilde{\lambda}_*^{-\mu} \leq \frac{1}{2}\widetilde{\lambda}_*^{\frac{s+1}{2}} \leq \frac{1}{2}.$$

Gathering all three cases, we obtain that it always holds that

$$|p'_{\widehat{m}-1}(0)| \leq c(\mu, \tau) \max\left(\frac{\rho}{M}, 1\right) \lambda_*^{-1}. \quad (26)$$

Second step: upper bound on $|p'_{\widehat{m}}(0)|$. For this we use the result of the first step and relate $|p'_{\widehat{m}-1}(0)|$ to $|p'_{\widehat{m}}(0)|$. It is a property of orthogonal polynomials (see Hanke, Corollary 2.6) that for any $m \geq 1$

$$p_{m-1}'(0) - p_m'(0) = \frac{[p_{m-1}, p_{m-1}]_{(0)} - [p_m, p_m]_{(0)}}{\left[p_{m-1}^{(2)}, p_{m-1}^{(2)}\right]_{(1)}} \leq \frac{[p_{m-1}, p_{m-1}]_{(0)}}{\left[p_{m-1}^{(2)}, p_{m-1}^{(2)}\right]_{(1)}}. \quad (27)$$

To upper bound the above quantity, we apply Lemma A.3 with the choice $\lambda = \lambda_*$ and

$$\varepsilon = \varepsilon_* := a(\mu, \tau) \min\left(\frac{M}{\rho}, 1\right) \lambda_*,$$

where $0 < a(\mu, \tau) \leq 1$ should be chosen small enough in order to satisfy some constraints to be specified below. The first constraint is the requirement $\varepsilon_* \in (0, x_{1,m-1})$ in order to apply Lemma A.3. For this, it can be seen from (26) that $a(\mu, \tau)$ can be chosen small enough to ensure

$$\varepsilon_* \leq |p'_{m-1}(0)|^{-1} \leq x_{1,m-1},$$

the last inequality is an easy consequence of the fact that p_{m-1} has exactly $(m-1)$ positive real roots and $p_{m-1}(0) = 1$. We now turn to upper bound the following quantity appearing on the RHS of (18):

$$\begin{aligned} \Upsilon(\varepsilon_* + \lambda_*)^{\frac{1}{2}} \delta(\lambda_*) + c(\mu) \Upsilon^2 \varepsilon_* (\varepsilon_*^\mu + Z_\mu(\lambda_*)) \kappa^{-\mu-\frac{1}{2}} \rho \\ \leq 2(a(\mu, \tau) + 1) \lambda_*^{\frac{1}{2}} \delta(\lambda_*) + c(\mu) a(\mu, \tau) \min(\rho, M) \lambda_* \tilde{\lambda}_*^\mu \kappa^{-\frac{1}{2}} \rho \\ \leq (c(\mu) a(\mu, \tau) + 2) \lambda_*^{\frac{1}{2}} \delta(\lambda_*), \end{aligned} \quad (28)$$

where we have used the definition (21) for $\delta(\lambda_*)$ and inequality $Z_\mu(\lambda_*) \leq \lambda_*^\mu$, see (25). Now, we chose $a(\mu, \tau)$ so that the factor in the last display satisfies $c(\mu) a(\mu, \tau) \leq \frac{\tau}{2}$. Remember that the definition of the stopping rule entails

$$[p_{m-1}, p_{m-1}]_{(0)}^{\frac{1}{2}} = \|T_n^*(T_n f_{\widehat{m}-1} - \mathbf{Y})\| > (2 + \tau) \lambda_*^{\frac{1}{2}} \delta(\lambda_*) > (2 + \tau) \lambda_*^{\frac{1}{2}} \delta(\lambda_*), \quad (29)$$

Now combining (18), (29) and (28), we obtain

$$\left(1 - \frac{\tau + \frac{1}{2}}{\tau + 2}\right) [p_{m-1}, p_{m-1}]_{(0)}^{\frac{1}{2}} \leq \varepsilon_*^{-\frac{1}{2}} \left[p_{m-1}^{(2)}, p_{m-1}^{(2)}\right]_{(1)}^{\frac{1}{2}};$$

using this inequality in relation with (27) and (26), we obtain

$$|p'_{\widehat{m}}(0)| \leq |p'_{\widehat{m}-1}(0)| + c(\tau) \varepsilon_*^{-1} \leq c(\mu, \tau) \max\left(\frac{\rho}{M}, 1\right) \lambda_*^{-1}.$$

Final step. We apply Lemma A.2 (with $\lambda = \lambda_*$ and $\varepsilon = \varepsilon_*$), together with the bound on $|p'_{\widehat{m}}(0)|$ just obtained, and the inequality (by definition of the stopping rule)

$$\|T_n^*(T_n f_{\widehat{m}} - \mathbf{Y})\| \leq (2 + \tau) \lambda_*^{\frac{1}{2}} \delta(\lambda_*),$$

obtaining, using again (25):

$$\begin{aligned} \|f_{\widehat{m}} - f^*\|_2 &= \|T(f_{\widehat{m}} - f_{\mathcal{H}}^*)\| \\ &\leq c(\mu, \tau) \left(\delta(\lambda_*) \max\left(\frac{\rho}{M}, 1\right) + \min(\rho, M) \tilde{\lambda}_*^{\mu+\frac{1}{2}} \right) \leq c(\mu, \tau) (M + \rho) \tilde{\lambda}_*^{\mu+\frac{1}{2}}. \end{aligned}$$

If $\widehat{m} = 0$, we can apply directly Lemma A.2 as above without requiring the two previous steps, since in this case $p'_0(0) = 0$, so that we obtain the same final bound.

A.4 Sketch of the proof of Theorem 2.3

For the proof of Theorem 2.3, the condition **B1**(λ) is replaced by

$$\mathbf{B1}'(\lambda) \quad \left\| (S + \lambda I)^{-\frac{1}{2}} (T_n^* \mathbf{Y} - T^* f^*) \right\| \leq \delta(\lambda).$$

We check that **B1'**(λ_*), **B2**(λ_*) and **B3** are satisfied in the setting of Theorem 2.3. To check **B1'**(λ_*), we use (30) instead of (31). Since the easily checked relation $T_n^* \mathbf{Y} = T_{\tilde{n}}^* \tilde{\mathbf{Y}}$ holds, the upper bound obtained here has the same form as for Theorem 2.2, therefore we can use the same value $\delta(\lambda^*)$ for condition **B1'**(λ_*) as in the previous section, given by (22). Notice however that we must now use the condition $\mu + s = r + s - \frac{1}{2} \geq 0$ to ensure that the chain of inequalities leading to (22) is valid.

For condition **B2**(λ_*), we can apply the deviation inequality (32) but with n replaced by $\tilde{n}$, since we make use of all the unlabeled data. Using the fact that $\frac{n}{\tilde{n}} \leq \tilde{\lambda}_*^{-(1-2r)_+}$ and some elementary algebra leads to **B2**(λ_*) being satisfied with $\Upsilon := 2$.

Finally condition **B3** is satisfied with Δ given by (24) with n replaced by $\tilde{n}$.

Once these conditions are established, intermediate results similar in structure to Lemmas A.1, A.2 and A.3 can be derived, but where **B1**(λ) is replaced by **B1'**(λ). The details are omitted here.

A.5 More technical lemmas

In this section we collect some technical lemmas which underpin the main results. These are taken from previous sources and are recalled here for completeness. The main statistical tool is the following deviation inequality:

Lemma A.4. *Let λ be a positive number. Under assumption **(Bounded)**, the following holds:*

$$\mathbb{P} \left[\left\| (S + \lambda I)^{-\frac{1}{2}} (T_n^* \mathbf{Y} - T^* f^*) \right\| \leq 2M \left(\sqrt{\frac{\mathcal{N}(\lambda)}{n}} + \frac{2\sqrt{\kappa}}{\sqrt{\lambda n}} \right) \log \frac{6}{\gamma} \right] \geq 1 - \gamma. \quad (30)$$

If the representation $f^ = T f_{\mathcal{H}}^*$ holds and under assumption **(Bernstein)**, we have the following:*

$$\mathbb{P} \left[\left\| (S + \lambda I)^{-\frac{1}{2}} (T_n^* \mathbf{Y} - S_n f_{\mathcal{H}}^*) \right\| \leq 2M \left(\sqrt{\frac{\mathcal{N}(\lambda)}{n}} + \frac{2\sqrt{\kappa}}{\sqrt{\lambda n}} \right) \log \frac{6}{\gamma} \right] \geq 1 - \gamma. \quad (31)$$

Finally, the following holds:

$$\mathbb{P} \left[\left\| (S + \lambda I)^{-\frac{1}{2}} (S_n - S) \right\|_{HS} \leq 2\sqrt{\kappa} \left(\sqrt{\frac{\mathcal{N}(\lambda)}{n}} + \frac{2\sqrt{\kappa}}{\sqrt{\lambda n}} \right) \log \frac{6}{\gamma} \right] \geq 1 - \gamma, \quad (32)$$

where we recall that $\|\cdot\|_{HS}$ denotes the Hilbert-Schmidt norm.

The proof can be found in [7], and is based on a Bernstein-type inequality for random variables taking values in a Hilbert space, as established in [20, 25].

Inequality (32) can be fruitfully combined with the following:

Lemma A.5. *Assume there exists $\eta > 0$ such that the following inequality holds:*

$$\left\| (S + \lambda)^{-\frac{1}{2}} (S_n - S) (S + \lambda)^{-\frac{1}{2}} \right\| < 1 - \eta,$$

then

$$\left\| (S + \lambda)^{\frac{1}{2}} (S_n + \lambda)^{-\frac{1}{2}} \right\| \leq \frac{1}{\sqrt{\eta}}.$$

Proof. First we have

$$\left\| (S + \lambda)^{\frac{1}{2}} (S_n + \lambda)^{-\frac{1}{2}} \right\| = \left\| (S + \lambda)^{\frac{1}{2}} (S_n + \lambda)^{-1} (S + \lambda)^{\frac{1}{2}} \right\|^{\frac{1}{2}};$$

then simple algebraic manipulation shows

$$(S + \lambda)^{\frac{1}{2}} (S_n + \lambda)^{-1} (S + \lambda)^{\frac{1}{2}} = \left(I - (S + \lambda)^{\frac{1}{2}} (S - S_n)^{-1} (S + \lambda)^{\frac{1}{2}} \right)^{-1}.$$

Finally, using the inequality $\|(I - A)^{-1}\| = \|\sum_{k \geq 0} A^k\| \leq (1 - \|A\|)^{-1}$ for $\|A\| < 1$ yields the conclusion. $\square$

We make use of the following operator inequalities:

Lemma A.6. *Let A, B be two positive, self-adjoint operators with $\max(\|A\|, \|B\|) \leq C$. Then for any $r \geq 0$, putting $\zeta = (r - 1)_+$, the following inequality holds:*

$$\|A^r - B^r\| \leq (\zeta + 1)C^\zeta \|A - B\|^{r-\zeta}. \quad (33)$$

Proof. Follows from the fact that the power function $x \mapsto x^r$ is operator monotone for $r \leq 1$ and Lipschitz with constant rC^{r-1} on $[0, C]$ if $r > 1$. $\square$

Lemma A.7 ([1], Theorem IX.2.1-2). *Let A, B be to self-adjoint, positive operators. Then for any $s \in [0, 1]$:*

$$\|A^s B^s\| \leq \|AB\|^s. \quad (34)$$

Note: this result is stated for positive matrices in [1], but it is easy to check that the proof applies as well to positive operators on a Hilbert space.