



Computer Science and Artificial Intelligence Laboratory

Technical Report

MIT-CSAIL-TR-2006-062
CBCL-264

September 10, 2006

Optimal Rates for Regularization Operators in Learning Theory

Andrea Caponnetto

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE 10 SEP 2006		2. REPORT TYPE		3. DATES COVERED 00-09-2006 to 00-09-2006	
4. TITLE AND SUBTITLE Optimal Rates for Regularization Operators in Learning Theory				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Massachusetts Institute of Technology, Computer Science and Artificial Intelligence Laboratory (CSAIL), Cambridge, MA, 02139				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES The original document contains color images.					
14. ABSTRACT					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT	18. NUMBER OF PAGES 18	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

Optimal rates for regularization operators in learning theory

Andrea Caponnetto ^{a b c}

^a *C.B.C.L., McGovern Institute, Massachusetts Institute of Technology, Bldg. 46-5155, , 77 Massachusetts Avenue, Cambridge, MA 02139*

^b *D.I.S.I., Università di Genova, Via Dodecaneso 35, 16146 Genova, Italy*

^c *Department of Computer Science, University of Chicago, 1100 East 58th Street, Chicago, IL 60637*

Abstract

We develop some new error bounds for learning algorithms induced by regularization methods in the regression setting. The “hardness” of the problem is characterized in terms of the parameters r and s , the first related to the “complexity” of the target function, the second connected to the *effective dimension* of the marginal probability measure over the input space. We show, extending previous results, that by a suitable choice of the regularization parameter as a function of the number of the available examples, it is possible to attain the optimal minimax rates of convergence for the expected squared loss of the estimators, over the family of priors fulfilling the constraint $r + s \geq \frac{1}{2}$. The setting considers both labelled and unlabelled examples, the latter being crucial for the optimality results on the priors in the range $r < \frac{1}{2}$.

This report describes research done at the Center for Biological & Computational Learning, which is in the McGovern Institute for Brain Research at MIT, as well as in the Dept. of Brain & Cognitive Sciences, and which is affiliated with the Computer Sciences & Artificial Intelligence Laboratory (CSAIL).

This research was sponsored by grants from: Office of Naval Research (DARPA) Contract No. MDA972-04-1-0037, Office of Naval Research (DARPA) Contract No. N00014-02-1-0915, National Science Foundation (ITR/SYS) Contract No. IIS-0112991, National Science Foundation (ITR) Contract No. IIS-0209289, National Science Foundation-NIH (CRCNS) Contract No. EIA-0218693, National Science Foundation-NIH (CRCNS) Contract No. EIA-0218506, and National Institutes of Health (Conte) Contract No. 1 P20 MH66239-01A1.

Additional support was provided by: Central Research Institute of Electric Power Industry (CRIEPI), Daimler-Chrysler AG, Compaq/Digital Equipment Corporation, Eastman Kodak Company, Honda R&D Co., Ltd., Industrial Technology Research Institute (ITRI), Komatsu Ltd., Eugene McDermott Foundation, Merrill-Lynch, NEC Fund, Oxygen, Siemens Corporate Research, Inc., Sony, Sumitomo Metal Industries, and Toyota Motor Corporation.

This work was also supported by the NSF grant 0325113, by the FIRB Project ASTAA and the IST Programme of the European Community, under the PASCAL Network of Excellence, IST-2002-506778.

1. INTRODUCTION

We consider the setting of semi-supervised statistical learning. We assume $Y \subset [-M, M]$ and the supervised part of the training set equal to

$$\mathbf{z} = (z_1, \dots, z_m),$$

with $z_i = (x_i, y_i)$ drawn i.i.d. according to the probability measure ρ over $Z = X \times Y$. Moreover consider the unsupervised part of the training set $(x_{m+1}^u, \dots, x_{\tilde{m}}^u)$, with x_i^u drawn i.i.d. according to the marginal probability measure over X , ρ_X . For sake of brevity we will also introduce the complete training set

$$\tilde{\mathbf{z}} = (\tilde{z}_1, \dots, \tilde{z}_{\tilde{m}}),$$

with $\tilde{z}_i = (\tilde{x}_i, \tilde{y}_i)$, where we introduced the compact notations $\tilde{x}_i$ and $\tilde{y}_i$, defined by

$$\tilde{x}_i = \begin{cases} x_i & \text{if } 1 \leq i \leq m, \\ x_i^u & \text{if } m < i \leq \tilde{m}, \end{cases}$$

and

$$\tilde{y}_i = \begin{cases} \frac{\tilde{m}}{m} y_i & \text{if } 1 \leq i \leq m, \\ 0 & \text{if } m < i \leq \tilde{m}. \end{cases}$$

It is clear that, in the supervised setting, the semi-supervised part of the training set is missing, whence $\tilde{m} = m$ and $\tilde{\mathbf{z}} = \mathbf{z}$.

In the following we will study the generalization properties of a class of estimators $f_{\tilde{\mathbf{z}}, \lambda}$ belonging to the *hypothesis space* $\mathcal{H}$: the RKHS of functions on X induced by the bounded Mercer kernel K (in the following $\kappa = \sup_{x \in X} K(x, x)$). The learning algorithms that we consider, have the general form

$$(1) \quad f_{\tilde{\mathbf{z}}, \lambda} = G_\lambda(T_{\tilde{\mathbf{x}}}) g_{\mathbf{z}},$$

where $T_{\tilde{\mathbf{x}}} \in \mathcal{L}(\mathcal{H})$ is given by,

$$T_{\tilde{\mathbf{x}}} f = \frac{1}{\tilde{m}} \sum_{i=1}^{\tilde{m}} K_{\tilde{x}_i} \langle K_{\tilde{x}_i}, f \rangle_{\mathcal{H}},$$

$g_{\mathbf{z}} \in \mathcal{H}$ is given by,

$$g_{\mathbf{z}} = \frac{1}{\tilde{m}} \sum_{i=1}^{\tilde{m}} K_{\tilde{x}_i} \tilde{y}_i = \frac{1}{m} \sum_{i=1}^m K_{x_i} y_i,$$

and the *regularization parameter* λ lays in the range $(0, \kappa]$. We will often used the shortcut notation $\dot{\lambda} = \frac{\lambda}{\kappa}$.

The functions $G_\lambda : [0, \kappa] \rightarrow \mathbb{R}$, which select the *regularization method*, will be characterized in terms of the constants A and B_r in $[0, +\infty]$, defined as follows

$$(2) \quad A = \sup_{\lambda \in (0, \kappa]} \sup_{\sigma \in [0, \kappa]} |(\sigma + \lambda) G_\lambda(\sigma)|$$

$$(3) \quad B_r = \sup_{t \in [0, r]} \sup_{\lambda \in (0, \kappa]} \sup_{\sigma \in [0, \kappa]} |1 - G_\lambda(\sigma) \sigma| \sigma^t \lambda^{-t}, \quad r > 0.$$

Finiteness of A and B_r (with r over a suitable range) are standard in the literature of ill-posed inverse problems (see for reference [12]). Regularization methods have been recently studied in the context of learning theory in [13, 9, 8, 10, 1].

The main results of the paper, Theorems 1 and 2, describe the convergence rates of $f_{\tilde{\mathbf{z}}, \lambda}$ to the *target function* $f_{\mathcal{H}}$. Here, the target function is the “best” function which can be arbitrarily well approximated by elements of our hypothesis space $\mathcal{H}$. More formally, $f_{\mathcal{H}}$ is the projection of the regression function $f_\rho(x) = \int_Y y d\rho_x(y)$ onto the closure of $\mathcal{H}$ in $\mathcal{L}^2(X, \rho_X)$.

The convergence rates in Theorems 1 and 2, will be described in terms of the constants C_r and D_s in $[0, +\infty]$ characterizing the probability measure ρ . These constants can be

described in terms of the integral operator $L_K : \mathcal{L}^2(X, \rho_X) \rightarrow \mathcal{L}^2(X, \rho_X)$ of kernel K . Note that the same integral operator is denoted by T , when seen as a bounded operator from $\mathcal{H}$ to $\mathcal{H}$.

The constants C_r characterize the conditional distributions $\rho_{|x}$ through $f_{\mathcal{H}}$, they are defined as follows

$$(4) \quad C_r = \begin{cases} \kappa^r \|L_K^{-r} f_{\mathcal{H}}\|_{\rho} & \text{if } f_{\mathcal{H}} \in \text{Im } L_K^r \\ +\infty & \text{if } f_{\mathcal{H}} \notin \text{Im } L_K^r \end{cases}, \quad r > 0.$$

Finiteness of C_r is a common *source condition* in the inverse problems literature (see [12] for reference). This type of condition has been introduced in the statistical learning literature in [7, 18, 3, 17, 4].

The constants D_s characterize the marginal distribution ρ_X through the *effective dimension* $\mathcal{N}(\lambda) = \text{Tr}[T(T + \lambda)^{-1}]$, they are defined as follows

$$(5) \quad D_s = 1 \vee \sup_{\lambda \in (0,1]} \sqrt{\mathcal{N}(\lambda) \lambda^s}, \quad s \in (0, 1].$$

Finiteness of D_s was implicitly assumed in [3, 4].

The paper is organized as follows. In Section 2 we focus on the RLS estimators $f_{\mathbf{z}, \lambda}^{\text{ls}}$, defined by the optimization problem

$$f_{\mathbf{z}, \lambda}^{\text{ls}} = \underset{f \in \mathcal{H}}{\text{argmin}} \frac{1}{\tilde{m}} \sum_{i=1}^{\tilde{m}} (f(\tilde{x}_i) - \tilde{y}_i)^2 + \lambda \|f\|_K^2,$$

and corresponding to the choice $G_{\lambda}(\sigma) = (\sigma + \lambda)^{-1}$ (see for example [5, 7, 18]). The main result of this Section, Theorem 1, extends the convergence analysis performed in [3, 4] from the range $r \geq \frac{1}{2}$ to arbitrary $r > 0$ and $s \geq \frac{1}{2} - r$. Corollary 1 gives optimal s -independent rates for $r > 0$.

The analysis of the RLS algorithm is a useful preliminary step for the study of general regularization methods, which is performed in Section 3. The aim of this Section is develop a s -dependent analysis in the case $r > 0$ for general regularization methods G_{λ} . In Theorem 2 we extend the results given in Theorem 1 to general regularization methods. In fact, in Theorem 2 we obtain optimal minimax rates of convergence (see [3, 4]) for the involved problems, under the assumption that $r + s \geq \frac{1}{2}$. Finally, Corollary 2 extends Corollary 1 to general G_{λ} .

In Sections 4 and 5 we give the proofs of the results stated in the previous Sections.

2. RISK BOUNDS FOR RLS.

We state our main result concerning the convergence of $f_{\mathbf{z}, \lambda}^{\text{ls}}$ to $f_{\mathcal{H}}$. The function $|x|_+$, appearing in the text of Theorem 1, is the “positive part” of x , that is $\frac{x+|x|}{2}$.

Theorem 1. *Let r and s be two reals in the interval $(0, 1]$, fulfilling the constraint $r + s \geq \frac{1}{2}$.*

Furthermore, let m and λ satisfy the constraints $\lambda \leq \|T\|$ and

$$(6) \quad \dot{\lambda} = \left(\frac{4D_s \log \frac{6}{\delta}}{\sqrt{m}} \right)^{\frac{2}{2r+s}},$$

for $\delta \in (0, 1)$. Finally, assume $\tilde{m} \geq m \dot{\lambda}^{-|1-2r|_+}$. Then, with probability greater than $1 - \delta$, it holds

$$\|f_{\mathbf{z}, \lambda}^{\text{ls}} - f_{\mathcal{H}}\|_{\rho} \leq 4(M + C_r) \left(\frac{4D_s \log \frac{6}{\delta}}{\sqrt{m}} \right)^{\frac{2r}{2r+s}}.$$

Some comments are in order.

First, while eq. (6) expresses λ in terms of m and δ , it is straightforward verifying that the condition $\lambda \leq \|T\|$ is satisfied for

$$\sqrt{m} \geq 4D_s \left(\frac{\kappa}{\|T\|} \right)^{r+\frac{s}{2}} \log \frac{6}{\delta}.$$

Second, the asymptotic rate of convergence $O\left(m^{-\frac{r}{2r+s}}\right)$ of $\|f_{\mathbf{z},\lambda}^{\text{ls}} - f_{\mathcal{H}}\|_{\rho}$, is optimal in the minimax sense of [11, 4]. Indeed, in Th.2 of [4], it was showed that this asymptotic order is optimal over the class of probability measures ρ , such that $f_{\mathcal{H}} \in \text{Im } L_K^r$, and the eigenvalues of T , λ_i , have asymptotic order $O\left(i^{-\frac{1}{s}}\right)$. In fact, the condition on $f_{\mathcal{H}}$ implies the finiteness of C_r and the condition on the spectrum of T implies the finiteness of D_s (see Prop.3 in [4]).

Upper bounds of the type given in [17] or [3] (and stated in [6, 4], under a weaker noise condition, and in the more general framework of vector-valued functions) can be obtained as a corollary of Theorem 1, considering the case $r \geq \frac{1}{2}$.

However, the advantage of using extra unlabelled data, is evident when $r < \frac{1}{2}$. In this case, the unlabelled examples (enforcing the assumption $\tilde{m} \geq m\dot{\lambda}^{2r-1}$) allow (if $s \geq \frac{1}{2} - r$) again the rate of convergence $O\left(m^{-\frac{r}{2r+s}}\right)$, over classes of measures ρ defined in terms of finiteness of the constants C_r and D_s . It is not known to the author whether the same rate of convergence can be achieved by the RLS estimator, for $s < \frac{1}{2} - r$.

A simple corollary of Theorem 1, encompassing all the values of r in $(0, 1]$, can be obtained observing that $D_1 = 1$, for every kernel K and marginal distribution ρ_X (see Prop. 2).

Corollary 1. *Let $\tilde{m} \geq m\dot{\lambda}^{-|1-2r|_+}$ hold with r in the interval $(0, 1]$. If λ satisfies the constraints $\lambda \leq \|T\|$ and*

$$\dot{\lambda} = \left(\frac{4 \log \frac{6}{\delta}}{\sqrt{m}} \right)^{\frac{2}{2r+1}},$$

for $\delta \in (0, 1)$, then, with probability greater than $1 - \delta$, it holds

$$\|f_{\mathbf{z},\lambda}^{\text{ls}} - f_{\mathcal{H}}\|_{\rho} \leq 4(M + C_r) \left(\frac{4 \log \frac{6}{\delta}}{\sqrt{m}} \right)^{\frac{2r}{2r+1}}.$$

3. RISK BOUNDS FOR GENERAL REGULARIZATION METHODS.

In this Section we state a result which generalizes Theorem 1 from RLS to general regularization algorithms of type described by equation (1). In this general framework we need $(\dot{\lambda}^{-|2-2r-s|_+} - 1)m$ unlabelled examples in order to get minimax optimal rates, slightly more than the $(\dot{\lambda}^{-|1-2r|_+} - 1)m$ required in Theorem 1 for the RLS estimator. We adopt the same notations and definitions introduced in the previous section.

Theorem 2. *Let $r > 0$ and $s \in (0, 1]$ fulfill the constraint $r + s \geq \frac{1}{2}$. Furthermore, let m and λ satisfy the constraints $\lambda \leq \|T\|$ and*

$$(7) \quad \dot{\lambda} = \left(\frac{4D_s \log \frac{6}{\delta}}{\sqrt{m}} \right)^{\frac{2}{2r+s}},$$

for $\delta \in (0, \frac{1}{3})$. Finally, assume $\tilde{m} \geq 4 \vee m\dot{\lambda}^{-|2-2r-s|_+}$. Then, with probability greater than $1 - 3\delta$, it holds

$$\|f_{\mathbf{z},\lambda} - f_{\mathcal{H}}\|_{\rho} \leq E_r \left(\frac{4D_s \log \frac{6}{\delta}}{\sqrt{m}} \right)^{\frac{2r}{2r+s}},$$

where

$$(8) \quad E_r = C_r (30A + 2(3 + r)B_r + 1) + 9MA.$$

The proof of the above Theorem is postponed to Section 5.

For the particular case $G_\lambda(\sigma) = (\sigma + \lambda)^{-1}$, $f_{\mathbf{z},\lambda} = f_{\mathbf{z},\lambda}^{\text{ls}}$ and the result above can be compared with Theorem 1. In this case, it is easy to verify that $A = 1$, $B_r \leq 1$ for $r \in [0, 1]$ and $C_r = +\infty$ for $r > 1$. The maximal value of r for which $C_r < +\infty$ is usually denoted as the *qualification* of the regularization method.

For a description of the properties of common regularization methods, in the inverse problems literature we refer to [12]. In the context of learning theory a review of these techniques can be found in [10] and [1]. In particular in [10] some convergence results of algorithms induced by Lipschitz continuous G_λ can be found.

A simple corollary of Theorem 2 which generalizes Corollary 1 to arbitrary regularization methods, can be obtained observing that $D_1 = 1$, for every kernel K and marginal distribution ρ_X (see Prop. 2).

Corollary 2. *Let $\tilde{m} \geq 4 \vee m\dot{\lambda}^{-|1-2r|+}$ hold with $r > 0$. If λ satisfies the constraints $\lambda \leq \|T\|$ and*

$$\dot{\lambda} = \left(\frac{4 \log \frac{6}{\delta}}{\sqrt{m}} \right)^{\frac{2}{2r+1}},$$

for some $\delta \in (0, \frac{1}{3})$, then, with probability greater than $1 - 3\delta$, it holds

$$\left\| f_{\mathbf{z},\lambda}^{\text{ls}} - f_{\mathcal{H}} \right\|_{\rho} \leq E_r \left(\frac{4 \log \frac{6}{\delta}}{\sqrt{m}} \right)^{\frac{2r}{2r+1}},$$

with E_r defined by eq. (8).

4. PROOF OF THEOREM 1

In this section we give the proof of Theorem 1. First we need some preliminary propositions.

Proposition 1. *Assume $\lambda \leq \|T\|$ and*

$$(9) \quad \lambda \tilde{m} \geq 16\kappa \mathcal{N}(\lambda) \log^2 \frac{6}{\delta},$$

for some $\delta \in (0, 1)$. Then, with probability greater than $1 - \delta$, it holds

$$\left\| (T + \lambda)^{\frac{1}{2}} (f_{\mathbf{z},\lambda}^{\text{ls}} - f_{\lambda}^{\text{ls}}) \right\|_{\mathcal{H}} \leq 8 \left(M + \sqrt{\kappa \frac{m}{\tilde{m}}} \|f_{\lambda}^{\text{ls}}\|_{\mathcal{H}} \right) \left(\frac{2}{m} \sqrt{\frac{\kappa}{\lambda}} + \sqrt{\frac{\mathcal{N}(\lambda)}{m}} \right) \log \frac{6}{\delta},$$

where

$$f_{\lambda}^{\text{ls}} = (T + \lambda)^{-1} L_K f_{\mathcal{H}}.$$

Proof. Assuming

$$(10) \quad S_1 := \left\| (T + \lambda)^{-\frac{1}{2}} (T - T_{\mathbf{x}}) (T + \lambda)^{-\frac{1}{2}} \right\|_{\text{HS}} < 1,$$

by simple algebraic computations we obtain

$$\begin{aligned}
f_{\mathbf{z},\lambda}^{\text{ls}} - f_{\lambda}^{\text{ls}} &= (T_{\mathbf{z}} + \lambda)^{-1} g_{\mathbf{z}} - (T + \lambda)^{-1} g \\
&= (T_{\mathbf{z}} + \lambda)^{-1} \{ (g_{\mathbf{z}} - g) + (T - T_{\mathbf{z}})(T + \lambda)^{-1} g \} \\
&= (T_{\mathbf{z}} + \lambda)^{-1} (T + \lambda)^{\frac{1}{2}} \left\{ (T + \lambda)^{-\frac{1}{2}} (g_{\mathbf{z}} - g) + (T + \lambda)^{-\frac{1}{2}} (T - T_{\mathbf{z}})(T + \lambda)^{-1} g \right\} \\
&= (T + \lambda)^{-\frac{1}{2}} \left\{ \text{Id} - (T + \lambda)^{-\frac{1}{2}} (T - T_{\mathbf{z}})(T + \lambda)^{-\frac{1}{2}} \right\}^{-1} \\
&\quad \left\{ (T + \lambda)^{-\frac{1}{2}} (g_{\mathbf{z}} - g) + (T + \lambda)^{-\frac{1}{2}} (T - T_{\mathbf{z}}) f_{\lambda} \right\}.
\end{aligned}$$

Therefore we get

$$\left\| (T + \lambda)^{\frac{1}{2}} (f_{\mathbf{z},\lambda}^{\text{ls}} - f_{\lambda}^{\text{ls}}) \right\|_{\mathcal{H}} \leq \frac{S_2 + S_3}{1 - S_1},$$

where

$$\begin{aligned}
S_2 &:= \left\| (T + \lambda)^{-\frac{1}{2}} (g_{\mathbf{z}} - g) \right\|_{\mathcal{H}}, \\
S_3 &:= \left\| (T + \lambda)^{-\frac{1}{2}} (T - T_{\mathbf{z}}) f_{\lambda} \right\|_{\mathcal{H}}.
\end{aligned}$$

Now we want to estimate the quantities S_1 , S_2 and S_3 using Prop. 4. In fact, choosing the correct vector-valued random variables ξ_1 , ξ_2 and ξ_3 , the following common representation holds,

$$S_h = \left\| \frac{1}{m_h} \sum_{i=1}^{m_h} \xi_h(\omega_i) - \mathbb{E}[\xi_h] \right\|, \quad h = 1, 2, 3.$$

Indeed, in order to let the equality above hold, $\xi_1 : X \rightarrow \mathcal{L}_{\text{HS}}(\mathcal{H})$ is defined by

$$\xi_1(x)[\cdot] = (T + \lambda)^{-\frac{1}{2}} K_x \langle K_x, \cdot \rangle_{\mathcal{H}} (T + \lambda)^{-\frac{1}{2}},$$

and $m_1 = \tilde{m}$.

Moreover, $\xi_2 : Z \rightarrow \mathcal{H}$ is defined by

$$\xi_2(x, y) = (T + \lambda)^{-\frac{1}{2}} K_x y,$$

with $m_2 = m$.

And finally, $\xi_3 : X \rightarrow \mathcal{H}$ is defined by

$$\xi_3(x) = (T + \lambda)^{-\frac{1}{2}} K_x f_{\lambda}^{\text{ls}}(x),$$

with $m_3 = \tilde{m}$.

Hence, applying three times Prop. 4, we can write

$$\mathbb{P} \left[S_h \leq 2 \left(\frac{H_h}{m_h} + \frac{\sigma_h}{\sqrt{m_h}} \right) \log \frac{6}{\delta} \right] \geq 1 - \frac{\delta}{3}, \quad h = 1, 2, 3,$$

where, as it can be straightforwardly verified, the constants H_h and σ_h are given by the expressions

$$\begin{aligned}
H_1 &= 2 \frac{\kappa}{\lambda}, & \sigma_1^2 &= \frac{\kappa}{\lambda} \mathcal{N}(\lambda), \\
H_2 &= 2M \sqrt{\frac{\kappa}{\lambda}}, & \sigma_2^2 &= M^2 \mathcal{N}(\lambda), \\
H_3 &= 2 \|f_{\lambda}^{\text{ls}}\|_{\mathcal{H}} \frac{\kappa}{\sqrt{\lambda}}, & \sigma_3^2 &= \kappa \|f_{\lambda}^{\text{ls}}\|_{\mathcal{H}} \mathcal{N}(\lambda).
\end{aligned}$$

Now, recalling the assumptions on λ and $\tilde{m}$, with probability greater than $1 - \delta/3$, we get

$$\begin{aligned} S_1 &\leq 2 \left(\frac{2\kappa}{\tilde{m}\lambda} + \sqrt{\frac{\mathcal{N}(\lambda)\kappa}{\tilde{m}\lambda}} \right) \log \frac{6}{\delta} \\ \left(\mathcal{N}(\lambda) \geq \frac{\|T\|}{\|T\| + \lambda} \geq \frac{1}{2} \right) &\leq 4 \frac{\kappa \mathcal{N}(\lambda) \log^2 \frac{6}{\delta}}{\lambda \tilde{m}} + \sqrt{\frac{\kappa \mathcal{N}(\lambda) \log^2 \frac{6}{\delta}}{\lambda \tilde{m}}} \\ (\text{eq. (9)}) &\leq \frac{1}{4} + \frac{1}{2} = \frac{3}{4}. \end{aligned}$$

Hence, since $\tilde{m} \geq m$, with probability greater than $1 - \delta$,

$$\begin{aligned} \left\| (T + \lambda)^{\frac{1}{2}} (f_{\mathbf{z},\lambda}^{\text{ls}} - f_{\lambda}^{\text{ls}}) \right\|_{\mathcal{H}} &\leq 4(S_2 + S_3) \\ &\leq 8 \left(M + \sqrt{\kappa \frac{m}{\tilde{m}}} \|f_{\lambda}^{\text{ls}}\|_{\mathcal{H}} \right) \left(\frac{2}{m} \sqrt{\frac{\kappa}{\lambda}} + \sqrt{\frac{\mathcal{N}(\lambda)}{m}} \right) \log \frac{6}{\delta}. \end{aligned}$$

□

Proposition 2. *For every probability measure ρ_X and $\lambda > 0$, it holds*

$$\|T\| \leq \kappa,$$

and

$$\lambda \mathcal{N}(\lambda) \leq \kappa.$$

Proof. First, observe that

$$\text{Tr}[T] = \int_X \text{Tr}[K_x \langle K_x, \cdot \rangle_{\mathcal{H}}] d\rho_X(x) = \int_X K(x, x) d\rho_X(x) \leq \sup_{x \in X} K(x, x) \leq \kappa.$$

Therefore, since T is a positive self-adjoint operator, the first inequality follows observing that

$$\|T\| \leq \text{Tr}[T] \leq \kappa.$$

The second inequality can be proved observing that, since $\psi_{\lambda}(\sigma^2) := \frac{\lambda \sigma^2}{\sigma^2 + \lambda} \leq \sigma^2$, it holds

$$\lambda \mathcal{N}(\lambda) = \text{Tr}[\psi_{\lambda}(T)] \leq \text{Tr}[T] \leq \kappa.$$

□

Proposition 3. *Let $f_{\mathcal{H}} \in \text{Im } L_K^r$ for some $r > 0$. Then, the following estimates hold,*

$$\begin{aligned} \|f_{\lambda}^{\text{ls}} - f_{\mathcal{H}}\|_{\rho} &\leq \lambda^r \|L_K^{-r} f_{\mathcal{H}}\|_{\rho}, \quad \text{if } r \leq 1 \\ \|f_{\lambda}^{\text{ls}}\|_{\mathcal{H}} &\leq \begin{cases} \lambda^{-\frac{1}{2}+r} \|L_K^{-r} f_{\mathcal{H}}\|_{\rho} & \text{if } r \leq \frac{1}{2}, \\ \kappa^{-\frac{1}{2}+r} \|L_K^{-r} f_{\mathcal{H}}\|_{\rho} & \text{if } r > \frac{1}{2}. \end{cases} \end{aligned}$$

Proof. The first estimate is standard in the theory of inverse problems, see, for example, [14, 12] or [18].

Regarding the second estimate, if $r \leq \frac{1}{2}$, since T is positive, we can write,

$$\begin{aligned} \|f_\lambda^{\text{ls}}\|_{\mathcal{H}} &\leq \|(T + \lambda)^{-1} L_K f_{\mathcal{H}}\|_{\mathcal{H}} \\ &\leq \left\| (T + \lambda)^{-\frac{1}{2}+r} (T(T + \lambda)^{-1})^{\frac{1}{2}+r} L_K^{\frac{1}{2}-r} f_{\mathcal{H}} \right\|_{\mathcal{H}} \\ &\leq \|(T + \lambda)^{-1}\|^{\frac{1}{2}-r} \|L_K^{-r} f_{\mathcal{H}}\|_{\rho} \leq \lambda^{-\frac{1}{2}+r} \|L_K^{-r} f_{\mathcal{H}}\|_{\rho}. \end{aligned}$$

On the contrary, if $r > \frac{1}{2}$, since by Prop. 2 $\|T\| \leq \kappa$, we obtain,

$$\begin{aligned} \|f_\lambda^{\text{ls}}\|_{\mathcal{H}} &\leq \|(T + \lambda)^{-1} L_K f_{\mathcal{H}}\|_{\mathcal{H}} \\ &\leq \left\| T^{r-\frac{1}{2}} T(T + \lambda)^{-1} L_K^{\frac{1}{2}-r} f_{\mathcal{H}} \right\|_{\mathcal{H}} \\ &\leq \|T\|^{r-\frac{1}{2}} \|L_K^{-r} f_{\mathcal{H}}\|_{\rho} \leq \kappa^{r-\frac{1}{2}} \|L_K^{-r} f_{\mathcal{H}}\|_{\rho}. \end{aligned}$$

□

We also need the following probabilistic inequality based on a result of [16], see also Th. 3.3.4 of [19]. We report it without proof.

Proposition 4. *Let $(\Omega, \mathcal{F}, P)$ be a probability space and ξ be a random variable on Ω taking value in a real separable Hilbert space $\mathcal{K}$. Assume that there are two positive constants H and σ such that*

$$\begin{aligned} \|\xi(\omega)\|_{\mathcal{K}} &\leq \frac{H}{2} \quad \text{a.s,} \\ \mathbb{E}[\|\xi\|_{\mathcal{K}}^2] &\leq \sigma^2, \end{aligned}$$

then, for all $m \in \mathbb{N}$ and $0 < \delta < 1$,

$$(11) \quad \mathbb{P}_{(\omega_1, \dots, \omega_m) \sim P^m} \left[\left\| \frac{1}{m} \sum_{i=1}^m \xi(\omega_i) - \mathbb{E}[\xi] \right\|_{\mathcal{K}} \leq 2 \left(\frac{H}{m} + \frac{\sigma}{\sqrt{m}} \right) \log \frac{2}{\delta} \right] \geq 1 - \delta.$$

We are finally ready to prove Theorem 1.

Proof of Theorem 1. The Theorem is a corollary of Prop. 1. We proceed by steps.

First. Observe that, by Prop. 2, it holds

$$\dot{\lambda} \leq \frac{\|T\|}{\kappa} \leq 1.$$

Second. Condition (9) holds. In fact, since $\dot{\lambda} \leq 1$ and by the assumption $\tilde{m} \geq m \dot{\lambda}^{-|1-2r|_+}$, we get,

$$\dot{\lambda} \tilde{m} \geq \dot{\lambda}^{-|1-2r|_++1} m \geq \dot{\lambda}^{2r} m.$$

Moreover, by eq. (6) and definition (5), we find

$$\dot{\lambda}^{2r} m = 16 D_s^2 \dot{\lambda}^{-s} \log^2 \frac{6}{\delta} \geq 16 \mathcal{N}(\lambda) \log^2 \frac{6}{\delta}.$$

Third. Since $\dot{\lambda} \leq 1$, recalling definition (4) and Prop. 3, for every r in $(0, 1]$, we can write,

$$\begin{aligned} \|f_\lambda^{\text{ls}} - f_{\mathcal{H}}\|_{\rho} &\leq \dot{\lambda}^r C_r, \\ \kappa \|f_\lambda^{\text{ls}}\|_{\mathcal{H}}^2 &\leq \dot{\lambda}^{-|1-2r|_+} C_r^2. \end{aligned}$$

Therefore we can apply Prop. 1, and using the two estimates above, the assumption $\tilde{m} \geq m \dot{\lambda}^{-|1-2r|_+}$ and the definition of D_s , to obtain the following bound,

$$\begin{aligned}
\|f_{\mathbf{z},\lambda}^{\text{ls}} - f_{\mathcal{H}}\|_{\rho} &\leq \|f_{\mathbf{z},\lambda}^{\text{ls}} - f_{\lambda}^{\text{ls}}\|_{\rho} + \|f_{\lambda} - f_{\mathcal{H}}\|_{\rho} \\
\left(\|f\|_{\rho} = \|\sqrt{T}f\|_{\mathcal{H}}\right) &\leq \|(T + \lambda)^{\frac{1}{2}}(f_{\mathbf{z},\lambda}^{\text{ls}} - f_{\lambda}^{\text{ls}})\|_{\mathcal{H}} + \|f_{\lambda} - f_{\mathcal{H}}\|_{\rho} \\
&\leq 8(M + C_r) \frac{1}{\sqrt{m}} \left(\frac{2}{\sqrt{m\lambda}} + \frac{D_s}{\sqrt{\lambda^s}} \right) \log \frac{6}{\delta} + \dot{\lambda}^r C_r \\
(\text{eq. (6)}) &= 2(M + C_r) \dot{\lambda}^r \left(1 + \frac{\dot{\lambda}^{r+s-\frac{1}{2}}}{2D_s^2 \log \frac{6}{\delta}} \right) + \dot{\lambda}^r C_r \\
\left(r + s \geq \frac{1}{2}\right) &\leq (3(M + C_r) + C_r) \dot{\lambda}^r \leq 4(M + C_r) \dot{\lambda}^r.
\end{aligned}$$

Substituting the expression (6) for $\dot{\lambda}$ in the inequality above, concludes the proof. $\square$

5. PROOF OF THEOREM 2

In this section we give the proof of Theorem 2. It is based on Proposition 1 which establishes an upper bound on the sample error for the RLS algorithm in terms of the constants C_r and D_s . When need some preliminary results. Proposition 5 shows properties of the truncated functions f_{λ}^{tr} , defined by equation (12), analogous to those given in Proposition 3 for the functions f_{λ}^{ls} .

Proposition 5. *Let $f_{\mathcal{H}} \in \text{Im } L_K^r$ for some $r > 0$. For any $\lambda > 0$ let the truncated function f_{λ}^{tr} be defined by*

$$(12) \quad f_{\lambda}^{\text{tr}} = P_{\lambda} f_{\mathcal{H}}$$

where P_{λ} is the orthogonal projector in $\mathcal{L}^2(X, \rho_X)$ defined by

$$(13) \quad P_{\lambda} = \Theta_{\lambda}(L_K),$$

with

$$(14) \quad \Theta_{\lambda}(\sigma) = \begin{cases} 1 & \text{if } \sigma \geq \lambda, \\ 0 & \text{if } \sigma < \lambda. \end{cases}$$

Then, the following estimates hold,

$$\begin{aligned}
\|f_{\lambda}^{\text{tr}} - f_{\mathcal{H}}\|_{\rho} &\leq \lambda^r \|L_K^{-r} f_{\mathcal{H}}\|_{\rho}, \\
\|f_{\lambda}^{\text{tr}}\|_{\mathcal{H}} &\leq \begin{cases} \lambda^{-\frac{1}{2}+r} \|L_K^{-r} f_{\mathcal{H}}\|_{\rho} & \text{if } r \leq \frac{1}{2}, \\ \kappa^{-\frac{1}{2}+r} \|L_K^{-r} f_{\mathcal{H}}\|_{\rho} & \text{if } r > \frac{1}{2}. \end{cases}
\end{aligned}$$

Proof. The first estimate follows simply observing that

$$\|f_{\lambda}^{\text{tr}} - f_{\mathcal{H}}\|_{\rho} = \|P_{\lambda}^{\perp} f_{\mathcal{H}}\|_{\rho} = \|P_{\lambda}^{\perp} L_K^r\| \|L_K^{-r} f_{\mathcal{H}}\|_{\rho} \leq \lambda^r \|L_K^{-r} f_{\mathcal{H}}\|_{\rho},$$

where we introduced the orthogonal projector $P_{\lambda}^{\perp} = \text{Id} - P_{\lambda}$.

Now let us consider the second estimate. Firstly observe that, since the compact operators L_K and T have a common eigensystem of functions on X , then P_{λ} can also be seen as an orthogonal projector in $\mathcal{H}$, and $f_{\lambda}^{\text{tr}} \in \mathcal{H}$. Hence we can write,

$$\begin{aligned}
\|f_{\lambda}^{\text{tr}}\|_{\mathcal{H}} &= \|P_{\lambda} f_{\mathcal{H}}\|_{\mathcal{H}} \leq \|L_K^{-\frac{1}{2}} P_{\lambda} f_{\mathcal{H}}\|_{\rho} \\
&\leq \|L_K^{-\frac{1}{2}+r} \Theta_{\lambda}(L_K)\| \|L_K^{-r} f_{\mathcal{H}}\|_{\rho}.
\end{aligned}$$

The proof is concluded observing that by Prop. 2, $\|L_K\| = \|T\| \leq \kappa$, and that, for every $\sigma \in [0, \kappa]$, it holds

$$\sigma^{-\frac{1}{2}+r} \Theta_\lambda(\sigma) \leq \begin{cases} \lambda^{-\frac{1}{2}+r} & \text{if } r \leq \frac{1}{2}, \\ \kappa^{-\frac{1}{2}+r} & \text{if } r > \frac{1}{2}. \end{cases}$$

□

Proposition 6 below estimates one of the terms appearing in the proof of Theorem 2 for any $r > 0$. The case $r \geq \frac{1}{2}$ had already been analyzed in the proof of Theorem 7 in [1].

Proposition 6. *Let $r > 0$ and define*

$$(15) \quad \gamma = \lambda^{-1} \|T - T_{\bar{x}}\|.$$

Then, if $\lambda \in (0, \kappa]$, it holds

$$\left\| \sqrt{T} (G_\lambda(T_{\bar{x}}) T_{\bar{x}} - \text{Id}) f_\lambda^{\text{tr}} \right\|_{\mathcal{H}} \leq B_r C_r (1 + \sqrt{\gamma}) (2 + r\gamma \lambda^{\frac{3}{2}-r} + \gamma^\eta) \lambda^r,$$

where

$$\eta = |r - \frac{1}{2}| - \lfloor |r - \frac{1}{2}| \rfloor.$$

Proof. The two inequalities (16) and (17) will be useful in the proof. The first follows from Theorem 1 in [15],

$$(16) \quad \|T^\alpha - T_{\bar{x}}^\alpha\| \leq \|T - T_{\bar{x}}\|^\alpha, \quad \alpha \in [0, 1]$$

where we adopted the convention $0^0 = 1$. The second is a corollary of Theorem 8.1 in [2]

$$(17) \quad \|T^p - T_{\bar{x}}^p\| \leq p\kappa^{p-1} \|T - T_{\bar{x}}\|, \quad p \in \mathbb{N}.$$

We also need to introduce the orthogonal projector in $\mathcal{H}$, $P_{\bar{x}, \lambda}$, defined by

$$P_{\bar{x}, \lambda} = \Theta_\lambda(T_{\bar{x}}),$$

with Θ_λ defined in (14).

We analyze the cases $r \leq \frac{1}{2}$ and $r \geq \frac{1}{2}$ separately.

Case $r \leq \frac{1}{2}$: In the three steps below we subsequently estimate the norms of the three terms of the expansion

$$(18) \quad \begin{aligned} \sqrt{T} (G_\lambda(T_{\bar{x}}) T_{\bar{x}} - \text{Id}) f_\lambda^{\text{tr}} &= \sqrt{T} P_{\bar{x}, \lambda}^\perp r_\lambda(T_{\bar{x}}) f_\lambda^{\text{tr}} \\ &\quad + P_{\bar{x}, \lambda} r_\lambda(T_{\bar{x}}) T_{\bar{x}}^{\frac{1}{2}} f_\lambda^{\text{tr}} \\ &\quad + (\sqrt{T} - \sqrt{T_{\bar{x}}}) P_{\bar{x}, \lambda} r_\lambda(T_{\bar{x}}) f_\lambda^{\text{tr}}, \end{aligned}$$

where $P_{\bar{x}, \lambda}^\perp = \text{Id} - P_{\bar{x}, \lambda}$ and $r_\lambda(\sigma) = \sigma G_\lambda(\sigma) - 1$.

Step 1: Observe that

$$\begin{aligned} \left\| \sqrt{T} P_{\bar{x}, \lambda}^\perp \right\|^2 &= \left\| P_{\bar{x}, \lambda}^\perp T P_{\bar{x}, \lambda}^\perp \right\| \leq \sup_{\phi \in \text{Im } P_{\bar{x}, \lambda}^\perp} \frac{(\phi, T\phi)_{\mathcal{H}}}{\|\phi\|_{\mathcal{H}}^2} \\ &\leq \sup_{\phi \in \text{Im } P_{\bar{x}, \lambda}^\perp} \frac{(\phi, T_{\bar{x}}\phi)_{\mathcal{H}}}{\|\phi\|_{\mathcal{H}}^2} + \sup_{\phi \in \mathcal{H}} \frac{(\phi, (T - T_{\bar{x}})\phi)_{\mathcal{H}}}{\|\phi\|_{\mathcal{H}}^2} \\ &\leq \lambda + \|T_{\bar{x}} - T\| = \lambda(1 + \gamma). \end{aligned}$$

Therefore, from definitions (2) and (4) and Proposition 5, it follows

$$\begin{aligned} \left\| \sqrt{T} P_{\bar{x}, \lambda}^\perp r_\lambda(T_{\bar{x}}) f_\lambda^{\text{tr}} \right\|_{\mathcal{H}} &\leq \left\| \sqrt{T} P_{\bar{x}, \lambda}^\perp \right\| \|r_\lambda(T_{\bar{x}})\| \|f_\lambda^{\text{tr}}\|_{\mathcal{H}} \\ &\leq B_r C_r \sqrt{1 + \gamma} \lambda^r. \end{aligned}$$

Step 2: Observe that, from inequality (16), definition (4) and Proposition 5

$$\begin{aligned}
 \left\| T_{\bar{x}}^{\frac{1}{2}-r} f_{\lambda}^{\text{tr}} \right\|_{\mathcal{H}} &\leq \left\| \Theta_{\lambda}(T) \right\| \left\| T^{\frac{1}{2}-r} f_{\mathcal{H}} \right\|_{\mathcal{H}} + \left\| T^{\frac{1}{2}-r} - T_{\bar{x}}^{\frac{1}{2}-r} \right\| \left\| f_{\lambda}^{\text{tr}} \right\|_{\mathcal{H}} \\
 (19) \qquad \qquad \qquad &\leq \kappa^{-r} C_r (1 + \gamma^{\frac{1}{2}-r}).
 \end{aligned}$$

Therefore from definition (3), it follows

$$\begin{aligned}
 \left\| P_{\bar{x},\lambda} r_{\lambda}(T_{\bar{x}}) T_{\bar{x}}^{\frac{1}{2}} f_{\lambda}^{\text{tr}} \right\|_{\mathcal{H}} &\leq \left\| P_{\bar{x},\lambda} \right\| \left\| r_{\lambda}(T_{\bar{x}}) T_{\bar{x}}^r \right\| \left\| T_{\bar{x}}^{\frac{1}{2}-r} f_{\lambda}^{\text{tr}} \right\|_{\mathcal{H}} \\
 &\leq B_r C_r (1 + \gamma^{\frac{1}{2}-r}) \dot{\lambda}^r.
 \end{aligned}$$

Step 3: Recalling the definition of $P_{\bar{x},\lambda}$, and applying again inequality (19) and Proposition 5, we get

$$\begin{aligned}
 &\left\| (\sqrt{T} - \sqrt{T_{\bar{x}}}) P_{\bar{x},\lambda} r_{\lambda}(T_{\bar{x}}) f_{\lambda}^{\text{tr}} \right\|_{\mathcal{H}} \\
 &\leq \left\| (\sqrt{T} - \sqrt{T_{\bar{x}}}) T_{\bar{x}}^{-\frac{1}{2}+r} P_{\bar{x},\lambda} r_{\lambda}(T_{\bar{x}}) T_{\bar{x}}^{\frac{1}{2}-r} f_{\lambda}^{\text{tr}} \right\|_{\mathcal{H}} \\
 &\leq \left\| \sqrt{T} - \sqrt{T_{\bar{x}}} \right\| \left\| T_{\bar{x}}^{-\frac{1}{2}+r} P_{\bar{x},\lambda} \right\| \left\| r_{\lambda}(T_{\bar{x}}) \right\| \left\| T_{\bar{x}}^{\frac{1}{2}-r} f_{\lambda}^{\text{tr}} \right\|_{\mathcal{H}} \\
 &\leq B_r C_r \gamma^{\frac{1}{2}} (1 + \gamma^{\frac{1}{2}-r}) \dot{\lambda}^r.
 \end{aligned}$$

Since we assumed $0 < r \leq \frac{1}{2}$, and therefore $\eta = \frac{1}{2} - r$, the three estimates above prove the statement of the Theorem in this case.

Case $r \geq \frac{1}{2}$: Consider the expansion

$$\begin{aligned}
 (G_{\lambda}(T_{\bar{x}}) T_{\bar{x}} - \text{Id}) f_{\lambda}^{\text{tr}} &\leq r_{\lambda}(T_{\bar{x}}) T^{r-\frac{1}{2}} v \\
 &\leq r_{\lambda}(T_{\bar{x}}) T_{\bar{x}}^{r-\frac{1}{2}} v + r_{\lambda}(T_{\bar{x}}) \left(T^{r-\frac{1}{2}} - T_{\bar{x}}^{r-\frac{1}{2}} \right) v \\
 &\leq r_{\lambda}(T_{\bar{x}}) T_{\bar{x}}^{r-\frac{1}{2}} v + r_{\lambda}(T_{\bar{x}}) T_{\bar{x}}^p \left(T^{r-\frac{1}{2}-p} - T_{\bar{x}}^{r-\frac{1}{2}-p} \right) v \\
 &\quad + r_{\lambda}(T_{\bar{x}}) (T^p - T_{\bar{x}}^p) T^{r-\frac{1}{2}-p} v
 \end{aligned}$$

where $v = P_{\lambda} T^{\frac{1}{2}-r} f_{\mathcal{H}}$, $r_{\lambda}(\sigma) = \sigma G_{\lambda}(\sigma) - 1$ and $p = \lfloor r - \frac{1}{2} \rfloor$.

Now, for any $\beta \in [0, \frac{1}{2}]$, from the expansion above using inequalities (16) and (17), and definition (3), we get

$$\begin{aligned}
 &\left\| T_{\bar{x}}^{\beta} (G_{\lambda}(T_{\bar{x}}) T_{\bar{x}} - \text{Id}) f_{\lambda}^{\text{tr}} \right\|_{\mathcal{H}} \leq \left\| r_{\lambda}(T_{\bar{x}}) T_{\bar{x}}^{r-\frac{1}{2}+\beta} \right\| \|v\|_{\mathcal{H}} \\
 (20) \qquad \qquad \qquad &+ \left\| r_{\lambda}(T_{\bar{x}}) T_{\bar{x}}^{p+\beta} \right\| \left\| T^{r-\frac{1}{2}-p} - T_{\bar{x}}^{r-\frac{1}{2}-p} \right\| \|v\|_{\mathcal{H}} \\
 &+ \left\| r_{\lambda}(T_{\bar{x}}) T_{\bar{x}}^{\beta} \right\| \|T^p - T_{\bar{x}}^p\| \left\| T^{r-\frac{1}{2}-p} \right\| \|v\|_{\mathcal{H}} \\
 &\leq B_r C_r \kappa^{-\frac{1}{2}+\beta} \left(\dot{\lambda}^{r-\frac{1}{2}+\beta} (1 + \gamma^{r-\frac{1}{2}-p}) + p \dot{\lambda}^{1+\beta} \gamma \right) \\
 &\leq B_r C_r \kappa^{-\frac{1}{2}+\beta} \left(\dot{\lambda}^{-\frac{1}{2}+\beta} (1 + \gamma^{\eta}) + r \gamma \dot{\lambda}^{1+\beta-r} \right) \dot{\lambda}^r.
 \end{aligned}$$

Finally, from the expansion

$$\begin{aligned}
 \left\| \sqrt{T} (G_{\lambda}(T_{\bar{x}}) T_{\bar{x}} - \text{Id}) f_{\lambda}^{\text{tr}} \right\|_{\mathcal{H}} &\leq \left\| \sqrt{T} - \sqrt{T_{\bar{x}}} \right\| \left\| r_{\lambda}(T_{\bar{x}}) f_{\lambda}^{\text{tr}} \right\|_{\mathcal{H}} \\
 &\quad + \left\| \sqrt{T_{\bar{x}}} r_{\lambda}(T_{\bar{x}}) f_{\lambda}^{\text{tr}} \right\|_{\mathcal{H}},
 \end{aligned}$$

using (16) and inequality (20) with $\beta = 0$ and $\beta = \frac{1}{2}$, we get the claimed result also in this case. $\square$

We need an additional preliminary result.

Proposition 7. *Let the operator Ω_λ be defined by*

$$(21) \quad \Omega_\lambda = \sqrt{T} G_\lambda(T_{\tilde{\mathbf{x}}}) (T_{\tilde{\mathbf{x}}} + \lambda) (T + \lambda)^{-\frac{1}{2}}.$$

Then, if $\lambda \in (0, \kappa]$, it holds

$$\|\Omega_\lambda\| \leq (1 + 2\sqrt{\gamma}) A,$$

with γ defined in eq. (15).

Proof. First consider the expansion

$$\Omega_\lambda = \left(\sqrt{T} - \sqrt{T_{\tilde{\mathbf{x}}}} \right) \Delta_\lambda (T + \lambda)^{-\frac{1}{2}} - \Delta_\lambda \left(\sqrt{T} - \sqrt{T_{\tilde{\mathbf{x}}}} \right) (T + \lambda)^{-\frac{1}{2}} + \Delta_\lambda \sqrt{T} (T + \lambda)^{-\frac{1}{2}},$$

where we introduced the operator

$$\Delta_\lambda = G_\lambda(T_{\tilde{\mathbf{x}}}) (T_{\tilde{\mathbf{x}}} + \lambda).$$

By condition (2), it follows $\|\Delta_\lambda\| \leq A$. Moreover, from inequality (16)

$$(22) \quad \left\| \sqrt{T} - \sqrt{T_{\tilde{\mathbf{x}}}} \right\| \leq \sqrt{\|T - T_{\tilde{\mathbf{x}}}\|}.$$

From the previous observations we easily get

$$\begin{aligned} \|\Omega_\lambda\| &\leq 2 \|\Delta_\lambda\| \left\| \sqrt{T} - \sqrt{T_{\tilde{\mathbf{x}}}} \right\| \left\| (T + \lambda)^{-\frac{1}{2}} \right\| + \|\Delta_\lambda\| \left\| \sqrt{T} (T + \lambda)^{-\frac{1}{2}} \right\| \\ &\leq A (1 + 2\sqrt{\gamma}), \end{aligned}$$

the claimed result. $\square$

We are now ready to show the proof of Theorem 2.

Proof of Theorem 2. We consider the expansion

$$\begin{aligned} \sqrt{T}(f_{\tilde{\mathbf{z}}, \lambda} - f_{\mathcal{H}}) &= \sqrt{T} (G_\lambda(T_{\tilde{\mathbf{x}}}) g_{\mathbf{z}} - f_{\lambda}^{\text{tr}}) + \sqrt{T}(f_{\lambda}^{\text{tr}} - f_{\mathcal{H}}) \\ &= \Omega_\lambda (T + \lambda)^{\frac{1}{2}} (f_{\tilde{\mathbf{z}}, \lambda}^{\text{ls}} - f_{\tilde{\mathbf{z}}', \lambda}^{\text{ls}}) + \sqrt{T} (G_\lambda(T_{\tilde{\mathbf{x}}}) T_{\tilde{\mathbf{x}}} - \text{Id}) f_{\lambda}^{\text{tr}} + \sqrt{T}(f_{\lambda}^{\text{tr}} - f_{\mathcal{H}}) \\ &= \Omega_\lambda \left((T + \lambda)^{\frac{1}{2}} (f_{\tilde{\mathbf{z}}, \lambda}^{\text{ls}} - f_{\lambda}^{\text{ls}}) + (T + \lambda)^{\frac{1}{2}} (f_{\lambda}^{\text{ls}} - \tilde{f}_{\lambda}^{\text{ls}}) + (T + \lambda)^{\frac{1}{2}} (\tilde{f}_{\lambda}^{\text{ls}} - f_{\tilde{\mathbf{z}}', \lambda}^{\text{ls}}) \right) \\ &\quad + \sqrt{T} (G_\lambda(T_{\tilde{\mathbf{x}}}) T_{\tilde{\mathbf{x}}} - \text{Id}) f_{\lambda}^{\text{tr}} + \sqrt{T}(f_{\lambda}^{\text{tr}} - f_{\mathcal{H}}) \end{aligned}$$

where the operator Ω_λ is defined by equation (21), the ideal RLS estimators are $f_{\lambda}^{\text{ls}} = (T + \lambda)^{-1} T f_{\mathcal{H}}$ and $\tilde{f}_{\lambda}^{\text{ls}} = (T + \lambda)^{-1} T f_{\lambda}^{\text{tr}}$, and $f_{\tilde{\mathbf{z}}', \lambda}^{\text{ls}} = (T_{\tilde{\mathbf{x}}} + \lambda)^{-1} T_{\tilde{\mathbf{x}}} f_{\lambda}^{\text{tr}}$ is the RLS estimator constructed by the training set

$$\tilde{\mathbf{z}}' = ((\tilde{x}_1, f_{\lambda}^{\text{tr}}(\tilde{x}_1)) \dots, (\tilde{x}_{\tilde{m}}, f_{\lambda}^{\text{tr}}(\tilde{x}_{\tilde{m}}))).$$

Hence we get the following decomposition,

$$(23) \quad \|f_{\tilde{\mathbf{z}}, \lambda} - f_{\mathcal{H}}\|_{\rho} \leq D \left(S^{\text{ls}} + R + \tilde{S}^{\text{ls}} \right) + P + P^{\text{tr}},$$

with

$$\begin{aligned}
S^{\text{ls}} &= \left\| (T + \lambda)^{\frac{1}{2}} (f_{\tilde{\mathbf{z}}, \lambda}^{\text{ls}} - f_{\lambda}^{\text{ls}}) \right\|_{\mathcal{H}}, \\
\bar{S}^{\text{ls}} &= \left\| (T + \lambda)^{\frac{1}{2}} (f_{\tilde{\mathbf{z}}', \lambda}^{\text{ls}} - \bar{f}_{\lambda}^{\text{ls}}) \right\|_{\mathcal{H}}, \\
D &= \|\Omega_{\lambda}\|, \\
P &= \left\| \sqrt{T} (G_{\lambda}(T_{\tilde{\mathbf{x}}}) T_{\tilde{\mathbf{x}}} - \text{Id}) f_{\lambda}^{\text{tr}} \right\|_{\mathcal{H}}, \\
P^{\text{tr}} &= \|f_{\lambda}^{\text{tr}} - f_{\mathcal{H}}\|_{\rho}, \\
R &= \left\| (T + \lambda)^{\frac{1}{2}} (\bar{f}_{\lambda}^{\text{ls}} - f_{\lambda}^{\text{ls}}) \right\|_{\mathcal{H}}.
\end{aligned}$$

Terms S^{ls} and $\bar{S}^{\text{ls}}$ will be estimated by Proposition 1, term D by Proposition 7, term P by Proposition 6 and finally terms P^{tr} and R by Proposition 5.

Let us begin with the estimates of S^{ls} and $\bar{S}^{\text{ls}}$. First observe that, by the same reasoning in the proof of Theorem 1, the assumptions of the Theorem imply inequality (9) in the text of Proposition 1.

Regarding the estimate of S^{ls} . Applying Proposition 1 and reasoning as in the proof of Theorem 1 (recall that by assumption $\tilde{m} \geq m\dot{\lambda}^{-|2-2r-s|_+} \geq m\dot{\lambda}^{-|1-2r|_+}$ and from Proposition 3, $\sqrt{\kappa} \|f_{\lambda}^{\text{ls}}\|_{\mathcal{H}} \leq C_r \dot{\lambda}^{-|\frac{1}{2}-r|_+}$), we get that with probability greater than $1 - \delta$

$$\begin{aligned}
(24) \quad S^{\text{ls}} &\leq 8 \left(M + \sqrt{\frac{m}{\tilde{m}}} C_r \dot{\lambda}^{-|\frac{1}{2}-r|_+} \right) \left(\frac{2}{m} \sqrt{\frac{\kappa}{\lambda}} + \sqrt{\frac{\mathcal{N}(\lambda)}{m}} \right) \log \frac{6}{\delta} \\
&\leq 8(M + C_r) \frac{1}{\sqrt{m}} \left(\frac{2}{\sqrt{m\lambda}} + \frac{D_s}{\sqrt{\lambda^s}} \right) \log \frac{6}{\delta} \\
&\stackrel{(eq. (7))}{=} 2(M + C_r) \dot{\lambda}^r \left(1 + \frac{\dot{\lambda}^{r+s-\frac{1}{2}}}{2D_s^2 \log \frac{6}{\delta}} \right) \\
&\stackrel{\left(r+s \geq \frac{1}{2}\right)}{\leq} 3(M + C_r) \dot{\lambda}^r
\end{aligned}$$

The term $\bar{S}^{\text{ls}}$ can be estimated observing that $\tilde{\mathbf{z}}'$ is a training set of $\tilde{m}$ supervised samples drawn i.i.d. from the probability measure ρ' with marginal ρ_X and conditional $\rho'_{|x}(y) = \delta(y - f_{\lambda}^{\text{tr}}(x))$. Therefore the regression function induced by ρ' is $f_{\rho'} = f_{\lambda}^{\text{tr}}$, and the support of ρ' is included in $[-M', M'] \times X$, with $M' = \sup_{x \in X} f_{\rho'}(x) \leq \sqrt{\kappa} \|f_{\lambda}^{\text{tr}}\|_{\mathcal{H}}$. Again applying Proposition 1 and reasoning as in the proof of Theorem 1, we obtain that

with probability greater than $1 - \delta$ it holds

$$\begin{aligned}
(25) \quad \bar{S}^{\text{ls}} &\leq 8 \left(M' + \sqrt{\kappa} \left\| \bar{f}_\lambda^{\text{ls}} \right\|_{\mathcal{H}} \right) \left(\frac{2}{\tilde{m}} \sqrt{\frac{\kappa}{\lambda}} + \sqrt{\frac{\mathcal{N}(\lambda)}{\tilde{m}}} \right) \log \frac{6}{\delta} \\
&\leq 16 \sqrt{\kappa} \left\| f_\lambda^{\text{tr}} \right\|_{\mathcal{H}} \left(\frac{2}{\tilde{m}} \sqrt{\frac{\kappa}{\lambda}} + \sqrt{\frac{\mathcal{N}(\lambda)}{\tilde{m}}} \right) \log \frac{6}{\delta} \\
(Prop.5) \quad &\leq 16 \sqrt{\frac{m}{\tilde{m}}} C_r \dot{\lambda}^{-|\frac{1}{2}-r|_+} \left(\frac{2}{m} \sqrt{\frac{\kappa}{\lambda}} + \sqrt{\frac{\mathcal{N}(\lambda)}{m}} \right) \log \frac{6}{\delta} \\
&\leq 16 C_r \frac{1}{\sqrt{m}} \left(\frac{2}{\sqrt{m \dot{\lambda}}} + \frac{D_s}{\sqrt{\dot{\lambda}^s}} \right) \log \frac{6}{\delta} \\
(eq. (7)) \quad &= 4 C_r \dot{\lambda}^r \left(1 + \frac{\dot{\lambda}^{r+s-\frac{1}{2}}}{2 D_s^2 \log \frac{6}{\delta}} \right) \\
\left(r + s \geq \frac{1}{2} \right) &\leq 6 C_r \dot{\lambda}^r
\end{aligned}$$

In order to get an upper bound for D and P , we have first to estimate the quantity γ (see definition (15)) appearing in the Propositions 6 and 7. Our estimate for γ follows from Proposition 4 applied to the random variable $\xi : X \rightarrow \mathcal{L}_{\text{HS}}(\mathcal{H})$ defined by

$$\xi(x)[\cdot] = \lambda^{-1} K_x \langle K_x, \cdot \rangle_{\mathcal{H}}.$$

We can set $H = \frac{2\kappa}{\lambda}$ and $\sigma = \frac{H}{2}$, and obtain that with probability greater than $1 - \delta$

$$\begin{aligned}
\gamma &\leq \lambda^{-1} \|T - T_{\tilde{\mathbf{x}}}\|_{\text{HS}} \leq \frac{2}{\lambda} \left(\frac{2\kappa}{\tilde{m}} + \frac{\kappa}{\sqrt{\tilde{m}}} \right) \log \frac{2}{\delta} \leq 4 \frac{1}{\lambda \sqrt{\tilde{m}}} \log \frac{2}{\delta} \\
&\leq 4 \frac{\dot{\lambda}^{|1-r-\frac{s}{2}|_+ - 1}}{\sqrt{m}} \log \frac{2}{\delta} \leq \dot{\lambda}^{|1-r-\frac{s}{2}|_+ - (1-r-\frac{s}{2})} \leq \dot{\lambda}^{|r+\frac{s}{2}-1|_+} \leq 1,
\end{aligned}$$

where we used the assumption $\tilde{m} \geq 4 \vee m \dot{\lambda}^{-|2-2r-s|_+}$ and the expression for $\dot{\lambda}$ in the text of the Theorem.

Hence, since $\gamma \leq 1$, from Proposition 7 we get

$$(26) \quad D \leq 3A,$$

and from Proposition 6

$$\begin{aligned}
(27) \quad P &\leq 2B_r C_r (3 + r \gamma \dot{\lambda}^{\frac{3}{2}-r}) \dot{\lambda}^r \\
&\leq 2B_r C_r (3 + r \dot{\lambda}^{|r+\frac{s}{2}-1|_+ + \frac{3}{2}-r}) \dot{\lambda}^r \\
&\leq 2B_r C_r (3 + r \dot{\lambda}^{\frac{s+1}{2}}) \dot{\lambda}^r \leq 2B_r C_r (3 + r) \dot{\lambda}^r.
\end{aligned}$$

Regarding terms P^{tr} and R . From Proposition 5 we get

$$(28) \quad P^{\text{tr}} \leq C_r \dot{\lambda}^r,$$

and hence,

$$\begin{aligned}
(29) \quad R &= \left\| (T + \lambda)^{-\frac{1}{2}} T (\bar{f}_\lambda^{\text{ls}} - f_\lambda^{\text{ls}}) \right\|_{\mathcal{H}} \\
&\leq \left\| \sqrt{T} (\bar{f}_\lambda^{\text{ls}} - f_\lambda^{\text{ls}}) \right\|_{\mathcal{H}} \leq P^{\text{tr}} \leq C_r \dot{\lambda}^r.
\end{aligned}$$

The proof is completed by plugging inequalities (24), (25), (26), (27), (28) and (29) in (23) and recalling the expression for $\dot{\lambda}$. $\square$

ACKNOWLEDGEMENTS

The author wish to thank E. De Vito, T. Poggio, L. Rosasco, S. Smale, A. Verri and Y. Yao for useful discussions and suggestions.

REFERENCES

- [1] F. Bauer, S. Pereverzev, and L. Rosasco. On regularization algorithms in learning theory. preprint, 2005.
- [2] M. S. Birman and M. Solomyak. Double operators integrals in hilbert scales. *Integr. Equ. Oper. Theory*, pages 131–168, 2003.
- [3] A. Caponnetto and E. De Vito. Fast rates for regularized least-squares algorithm. Technical report, Massachusetts Institute of Technology, Cambridge, MA, April 2005. CBCL Paper#248/AI Memo#2005-013.
- [4] A. Caponnetto and E. De Vito. Optimal rates for regularized least-squares algorithm. 2005. *to appear in Foundations of Computational Mathematics*.
- [5] F. Cucker and S. Smale. On the mathematical foundations of learning. *Bull. Amer. Math. Soc. (N.S.)*, 39(1):1–49 (electronic), 2002.
- [6] E. De Vito and A. Caponnetto. Risk bounds for regularized least-squares algorithm with operator-valued kernels. Technical report, Massachusetts Institute of Technology, Cambridge, MA, May 2005. CBCL Paper #249/AI Memo #2005-015.
- [7] E. De Vito, A. Caponnetto, and L. Rosasco. Model selection for regularized least-squares algorithm in learning theory. *Foundation of Computational Mathematics*, 5(1):59–85, February 2005.
- [8] E. De Vito, L. Rosasco, and A. Caponnetto. Discretization error analysis for tikhonov regularization. to appear in *Analysis and Applications*, 2005.
- [9] 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*, 6:883–904, 2005.
- [10] E. De Vito, L. Rosasco, and A. Verri. Spectral methods for regularization in learning theory. preprint, 2005.
- [11] R. DeVore, G. Kerkycharian, D. Picard, and V. Temlyakov. Mathematical methods for supervised learning. Technical report, Industrial Mathematics Institute, University of South Carolina, 2004.
- [12] H. W. Engl, M. Hanke, and A. Neubauer. *Regularization of inverse problems*, volume 375 of *Mathematics and its Applications*. Kluwer Academic Publishers Group, Dordrecht, 1996.
- [13] T. Evgeniou, M. Pontil, and T. Poggio. Regularization networks and support vector machines. *Adv. Comp. Math.*, 13:1–50, 2000.
- [14] C. W. Groetsch. *The theory of Tikhonov regularization for Fredholm equations of the first kind*, volume 105 of *Research Notes in Mathematics*. Pitman (Advanced Publishing Program), Boston, MA, 1984.
- [15] P. Mathé and S. V. Pereverzev. Moduli of continuity for operator valued functions. *Numer. Funct. Anal. Optim.*, 23(5-6):623–631, 2002.
- [16] I. F. Pinelis and A. I. Sakhanenko. Remarks on inequalities for probabilities of large deviations. *Theory Probab. Appl.*, 30(1):143–148, 1985.
- [17] S. Smale and D. Zhou. Learning theory estimates via integral operators and their approximations. preprint, 2005.
- [18] S. Smale and D.X. Zhou. Shannon sampling II : Connections to learning theory. *Appl. Comput. Harmonic Anal.* to appear.
- [19] V. Yurinsky. *Sums and Gaussian vectors*, volume 1617 of *Lecture Notes in Mathematics*. Springer-Verlag, Berlin, 1995.

