

Double Double Descent: On Generalization Errors in Transfer Learning between Linear Regression Tasks

Yehuda Dar* and Richard G. Baraniuk†

Abstract. We study the *transfer learning* process between two linear regression problems. An important and timely special case is when the regressors are *overparameterized* and perfectly *interpolate* their training data. We examine a parameter transfer mechanism whereby a subset of the parameters of the target task solution are constrained to the values learned for a related source task. We analytically characterize the generalization error of the target task in terms of the salient factors in the transfer learning architecture, i.e., the number of examples available, the number of (free) parameters in each of the tasks, the number of parameters transferred from the source to target task, and the relation between the two tasks. Our non-asymptotic analysis shows that the generalization error of the target task follows a *two-dimensional double descent* trend (with respect to the number of free parameters in each of the tasks) that is controlled by the transfer learning factors. Our analysis points to specific cases where the transfer of parameters is beneficial as a substitute for extra overparameterization (i.e., additional free parameters in the target task). Specifically, we show that the usefulness of a transfer learning setting is fragile and depends on a delicate interplay among the set of transferred parameters, the relation between the tasks, and the true solution. We also demonstrate that overparameterized transfer learning is not necessarily more beneficial when the source task is closer or identical to the target task.

Key words. Overparameterized learning, linear regression, transfer learning, double descent.

AMS subject classifications. 62J05, 68Q32

1. Introduction. Transfer learning [21] is a prominent strategy to address a machine learning task of interest using information and parameters already learned and/or available for a related task. Such designs significantly aid training of overparameterized models like deep neural networks (e.g., [4, 24, 16]), which are inherently challenging due to the vast number of parameters compared to the number of training data examples. There are various ways to integrate the previously-learned information from the source task in the learning process of the target task; often this is done by taking subsets of parameters (e.g., layers in neural networks) learned for the source task and plugging them in the target task model as parameter subsets that can be set fixed, finely tuned, or serve as non-random initialization for a thorough learning process. Obviously, transfer learning is useful only if the source and target tasks are sufficiently related with respect to the transfer mechanism utilized, e.g., [23, 29, 13]. Moreover, finding a successful transfer learning setting for deep neural networks was shown in [22] to be a delicate engineering task. The importance of transfer learning in contemporary practice should motivate fundamental understanding of its main aspects via *analytical* frameworks that may consider linear structures (see, e.g., [14]).

In general, the impressive success of overparameterized architectures for supervised learning have raised fundamental questions on the classical role of the bias-variance tradeoff that guided the traditional designs towards seemingly-optimal underparameterized models [5]. Em-

*Electrical and Computer Engineering Department, Rice University (ydar@rice.edu).

†Electrical and Computer Engineering Department, Rice University (richb@rice.edu).

empirical studies from recent years [25, 9, 2] have demonstrated the phenomenon that overparameterized supervised learning corresponds to a generalization error curve with a *double descent* trend (with respect to the number of parameters in the learned model). This double descent shape means that the generalization error peaks when the learned model starts to interpolate the training data (i.e., to achieve zero training error), but then the error continuously decreases as the overparameterization increases, often arriving to a global minimum that outperforms the best underparameterized solution. This phenomenon has been studied theoretically from the linear regression perspective in an extensive series of papers, e.g., in [3, 11, 28, 17, 1, 19]. The next stage is to provide corresponding fundamental understanding to learning problems beyond a single fully-supervised regression problem (see, for example, the study in [7] on overparameterized linear subspace fitting in unsupervised and semi-supervised settings).

In this paper we study the fundamentals of the natural meeting point between *overparameterized models* and the *transfer learning* concept. Our analytical framework is based on the least squares solutions to two related linear regression problems: the first is a source task whose solution has been found independently, and the second is a target task that is addressed using the solution already available for the source task. Specifically, the target task is carried out while keeping a subset of its parameters fixed to values transferred from the source task solution. Accordingly, the target task includes three types of parameters: free to-be-optimized parameters, transferred parameters set fixed to values from the source task, and parameters fixed to zeros (which in our case correspond to the elimination of input features). The mixture of the parameter types defines the parameterization level (i.e., the relation between the number of free parameters and the number of examples given) and the transfer-learning level (i.e., the portion of transferred parameters among the solution layout).

We conduct a non-asymptotic statistical analysis of the generalization errors in this transfer learning structure where the minimum ℓ_2 -norm solutions are used when the source and/or target tasks are overparameterized. Clearly, since the source task is solved independently, its generalization error follows a regular (one-dimensional) double descent shape with respect to the number of examples and free parameters available in the source task. Hence, our main contribution and interest are in the characterization of the generalization error of the target task that is carried out using the transfer learning approach described above. We show that the generalization error of the target task follows a double descent trend that depends on the double descent shape of the source task and on the transfer learning factors such as the number of parameters transferred and the relation between the source and target tasks. We also examine the generalization error of the target task as a function of two quantities: the number of free parameters in the source task and the number of free parameters in the target task. This interpretation presents the generalization error of the target task as having a *two-dimensional double descent* trend.

We also show how the generalization error of the target task is affected by the *specific set* of transferred parameters and its delicate interplay with the forms of the true solution and the source-target task relation. By that, we provide an analytical theory to the fragile nature of successful transfer learning designs. We demonstrate that the practical approach of arbitrary selection of transferred parameters signifies the fragile nature of successful parameter transfer settings, especially when the true (unknown) parameters have a sparse form in the feature

space.

We characterize the settings where transferring a set of parameters is more beneficial (in the sense of improved generalization in the target task) than defining them as additional free parameters or zeroing them. We also prove that the transfer of parameters from an overparameterized solution of a source task is not necessarily optimal when the source task is closer or identical to the target task.

The majority of this paper is focused on the analytical and empirical study of the minimum ℓ_2 -norm solution in our transfer learning setting. Yet, we also empirically examine the utilization of the minimum ℓ_1 -norm solution and the ridge regression approach in our transfer learning framework. The minimum ℓ_1 -norm (interpolating) solution also induces generalization errors that follow a double descent shape, and can outperform the minimum ℓ_2 -norm solution especially if the true parameters are sparse. The ridge regression method includes a regularization term that prevents interpolation and, when properly tuned, eliminates the double descent peak. Ridge regression can outperform the interpolating solutions for a wide range of overparameterization levels. However, at the proximity of maximal overparameterization, the minimum ℓ_2 -norm solution performs comparably to the ridge regression. This demonstrates that interpolation at extreme overparameterization levels can substitute properly tuned regularization in our transfer learning setting.

1.1. Related Work. Despite the prevalence of transfer learning in contemporary practice, there are only few analytical theories for this topic. First, Lampinen and Ganguli [14] studied the optimization dynamics in transfer learning of multi-layer linear models, which is a different research objective than our focus on double descent phenomena. Interestingly, when we posted the first version of our work on arXiv in June 2020, there was no literature on double descent phenomena and interpolating solutions in transfer learning. Later on, Dhifallah and Lu [8] studied single-layer nonlinear models that are suitable for both regression and classification problems, where the source task is ridge regularized and therefore prevents interpolation and double descent phenomena. Gerace et al. [10] studied a binary classification problem that is addressed by transfer learning of the first layer in a two-layer model that includes nonlinearities. The learning settings in [10] include regularization on both the source and target task and therefore attenuate some of the double descent behavior. The analysis in [10] requires numerical optimizations and empirical estimation of covariance matrices as inputs for their asymptotic formulations.

Remotely from overparameterized learning but related to transfer learning theory, Obst et al. [20] studied fine-tuning (based on gradient descent) between linear regression problems in underparameterized settings.

To summarize, each of the existing transfer learning theories employs different analytical tools, assumptions, and accordingly presents a unique analytical perspective on the topic. We focus on transfer learning between linear regression problems for Gaussian data and without regularization. This allows us to consider the minimum ℓ_2 -norm solution and explicitly characterize the generalization error and its double descent behavior in a detailed, completely closed-form formulation that is unavailable elsewhere. This also contributes to our discussion on beneficial transfer learning and lets us to analytically characterize the optimal source task to transfer from.

1.2. Paper Organization. This paper is organized as follows. In Section 2 we define the transfer learning architecture examined in this paper. In Section 3 we analytically characterize the double descent phenomenon in our transfer learning setting. In Section 4 we analyze the conditions for beneficial transfer of parameters compared to the alternatives of free and zeroed parameters. In Section 5 we characterize the optimal source task to transfer from. Sections 3-5 are focused on the minimum ℓ_2 -norm solution in the overparameterized regime of our transfer learning setting; in Section 6 we examine transfer learning in conjunction with the minimum ℓ_1 -norm solution and the ridge regression method. Section 7 concludes the paper. The Supplementary Materials include all of the proofs and mathematical developments as well as additional details and results for the empirical part of the paper.

2. Transfer Learning between Linear Regression Tasks: Problem Definition.

2.1. Source Task: Data Model and Solution Form. We start with the *source* data model, where a d -dimensional Gaussian input vector $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$ is connected to a response value $v \in \mathbb{R}$ via the noisy linear model

$$(2.1) \quad v = \mathbf{z}^T \boldsymbol{\theta} + \xi,$$

where $\xi \sim \mathcal{N}(0, \sigma_\xi^2)$ is a Gaussian noise component independent of $\mathbf{z}$, $\sigma_\xi > 0$, and $\boldsymbol{\theta} \in \mathbb{R}^d$ is an unknown vector. The data user is unfamiliar with the distribution of $(\mathbf{z}, v)$, however gets a dataset of $\tilde{n}$ independent and identically distributed (i.i.d.) draws of $(\mathbf{z}, v)$ pairs denoted as $\tilde{\mathcal{D}} \triangleq \left\{ (\mathbf{z}^{(i)}, v^{(i)}) \right\}_{i=1}^{\tilde{n}}$. The $\tilde{n}$ data samples can be rearranged as $\mathbf{Z} \triangleq [\mathbf{z}^{(1)}, \dots, \mathbf{z}^{(\tilde{n})}]^T$ and $\mathbf{v} \triangleq [v^{(1)}, \dots, v^{(\tilde{n})}]^T$ that satisfy the relation $\mathbf{v} = \mathbf{Z}\boldsymbol{\theta} + \boldsymbol{\xi}$ where $\boldsymbol{\xi} \triangleq [\xi^{(1)}, \dots, \xi^{(\tilde{n})}]^T$ is an unknown noise vector that its i^{th} component $\xi^{(i)}$ participates in the relation $v^{(i)} = \mathbf{z}^{(i),T} \boldsymbol{\theta} + \xi^{(i)}$ underlying the i^{th} data sample.

The *source* task is defined for a new (out of sample) data pair $(\mathbf{z}^{(\text{test})}, v^{(\text{test})})$ drawn from the distribution induced by (2.1) independently of the $\tilde{n}$ examples in $\tilde{\mathcal{D}}$. For a given $\mathbf{z}^{(\text{test})}$, the source task is to estimate the response value $v^{(\text{test})}$ by the value $\hat{v}$ that minimizes the corresponding out-of-sample squared error (i.e., the generalization error of the *source* task)

$$(2.2) \quad \tilde{\mathcal{E}}_{\text{out}} \triangleq \mathbb{E} \left[\left(\hat{v} - v^{(\text{test})} \right)^2 \right] = \sigma_\xi^2 + \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}} - \boldsymbol{\theta} \right\|_2^2 \right]$$

where the second equality stems from the data model in (2.1) and the corresponding linear form of $\hat{v} = \mathbf{z}^{(\text{test},T)} \hat{\boldsymbol{\theta}}$ where $\hat{\boldsymbol{\theta}}$ estimates $\boldsymbol{\theta}$ based on $\tilde{\mathcal{D}}$.

To address the source task based on the $\tilde{n}$ examples, one should choose the number of free parameters in the estimate $\hat{\boldsymbol{\theta}} \in \mathbb{R}^d$. Consider a predetermined layout where $\tilde{p}$ out of the d components of $\hat{\boldsymbol{\theta}}$ are free to be optimized, whereas the remaining $d - \tilde{p}$ components are constrained to zero values. The coordinates of the free parameters are specified in the set $\mathcal{S} \triangleq \{s_1, \dots, s_{\tilde{p}}\}$ where $1 \leq s_1 < \dots < s_{\tilde{p}} \leq d$ and the complementary set $\mathcal{S}^c \triangleq \{1, \dots, d\} \setminus \mathcal{S}$ contains the coordinates constrained to be zero valued. We define the $|\mathcal{S}| \times d$ matrix $\mathbf{Q}_{\mathcal{S}}$ as the linear operator that extracts from a d -dimensional vector its $|\mathcal{S}|$ -dimensional subvector of components residing at the coordinates specified in $\mathcal{S}$. Specifically, the values of the (k, s_k) components ($k = 1, \dots, |\mathcal{S}|$) of $\mathbf{Q}_{\mathcal{S}}$ are ones and the other components of $\mathbf{Q}_{\mathcal{S}}$ are zeros. The

definition given here for $\mathbf{Q}_{\mathcal{S}}$ can be adapted also to other sets of coordinates (e.g., $\mathbf{Q}_{\mathcal{S}^c}$ for $\mathcal{S}^c$) as denoted by the subscript of $\mathbf{Q}$. We now turn to formulate the *source* task using the linear regression form of

$$(2.3) \quad \hat{\boldsymbol{\theta}} = \arg \min_{\mathbf{r} \in \mathbb{R}^d} \|\mathbf{v} - \mathbf{Z}\mathbf{r}\|_2^2 \quad \text{subject to} \quad \mathbf{Q}_{\mathcal{S}^c}\mathbf{r} = \mathbf{0}$$

that its minimum ℓ_2 -norm solution (see details in Appendix A.1) is

$$(2.4) \quad \hat{\boldsymbol{\theta}} = \mathbf{Q}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}}^+ \mathbf{v}$$

where $\mathbf{Z}_{\mathcal{S}}^+$ is the pseudoinverse of $\mathbf{Z}_{\mathcal{S}} \triangleq \mathbf{Z}\mathbf{Q}_{\mathcal{S}}^T$. Note that $\mathbf{Z}_{\mathcal{S}}$ is a $\tilde{n} \times \tilde{p}$ matrix that its i^{th} row is formed by the $\tilde{p}$ components of $\mathbf{z}^{(i)}$ specified by the coordinates in $\mathcal{S}$, namely, only $\tilde{p}$ out of the d features of the input data vectors are utilized. Hence, in the underparameterized case of $\tilde{p} \leq \tilde{n}$ the solution (2.4) almost surely reduces to the unique least squares form of

$$(2.5) \quad \hat{\boldsymbol{\theta}} = \mathbf{Q}_{\mathcal{S}}^T (\mathbf{Z}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}})^{-1} \mathbf{Z}_{\mathcal{S}}^T \mathbf{v}.$$

Moreover, in both (2.4)-(2.5), $\hat{\boldsymbol{\theta}}$ is a d -dimensional vector that may have nonzero values only in the $\tilde{p}$ coordinates specified in $\mathcal{S}$ (this can be easily observed by noting that for an arbitrary $\mathbf{w} \in \mathbb{R}^{|\mathcal{S}|}$, the vector $\mathbf{u} = \mathbf{Q}_{\mathcal{S}}^T \mathbf{w}$ is a d -dimensional vector that its components satisfy $u_{s_k} = w_k$ for $k = 1, \dots, |\mathcal{S}|$ and $u_j = 0$ for $j \notin \mathcal{S}$). While the specific optimization form in (2.3) was not explicit in previous studies of non-asymptotic settings, e.g., [5, 3], the solution in (2.4) coincides with theirs and, thus, the formulation of the generalization error of our *source* task (which is a linear regression problem that, by itself, does not have any transfer learning aspect) is available from [5, 3] and provided in Appendix A.2 in our notations for completeness of presentation.

2.2. Target Task: Data Model and Solution using Transfer Learning. A second data class, which is our main interest, is modeled by $(\mathbf{x}, y) \in \mathbb{R}^d \times \mathbb{R}$ that satisfy

$$(2.6) \quad y = \mathbf{x}^T \boldsymbol{\beta} + \epsilon$$

where $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$ is a Gaussian input vector including d features, $\epsilon \sim \mathcal{N}(0, \sigma_\epsilon^2)$ is a Gaussian noise component independent of $\mathbf{x}$, $\sigma_\epsilon > 0$, and $\boldsymbol{\beta} \in \mathbb{R}^d$ is an unknown vector related to the $\boldsymbol{\theta}$ from (2.1) via

$$(2.7) \quad \boldsymbol{\theta} = \mathbf{H}\boldsymbol{\beta} + \boldsymbol{\eta}$$

where $\mathbf{H} \in \mathbb{R}^{d \times d}$ is a deterministic matrix and $\boldsymbol{\eta} \sim \mathcal{N}(\mathbf{0}, \sigma_\eta^2 \mathbf{I}_d)$ is a Gaussian noise vector with $\sigma_\eta \geq 0$. Here $\boldsymbol{\eta}$, $\mathbf{x}$, ϵ , $\mathbf{z}$ and ξ are independent. The data user does not know the distribution of $(\mathbf{x}, y)$ but receives a small dataset of n i.i.d. draws of $(\mathbf{x}, y)$ pairs denoted as $\mathcal{D} \triangleq \left\{ (\mathbf{x}^{(i)}, y^{(i)}) \right\}_{i=1}^n$. The n data samples can be organized in a $n \times d$ matrix of input variables $\mathbf{X} \triangleq [\mathbf{x}^{(1)}, \dots, \mathbf{x}^{(n)}]^T$ and a $n \times 1$ vector of responses $\mathbf{y} \triangleq [y^{(1)}, \dots, y^{(n)}]^T$ that together satisfy the relation $\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\epsilon}$ where $\boldsymbol{\epsilon} \triangleq [\epsilon^{(1)}, \dots, \epsilon^{(n)}]^T$ is an unknown noise vector that its i^{th} component $\epsilon^{(i)}$ is involved in the connection $y^{(i)} = \mathbf{x}^{(i),T} \boldsymbol{\beta} + \epsilon^{(i)}$ underlying the i^{th} example pair.

The *target* task considers a new (out of sample) data pair $(\mathbf{x}^{(\text{test})}, y^{(\text{test})})$ drawn from the model in (2.6) independently of the training examples in $\mathcal{D}$. Given $\mathbf{x}^{(\text{test})}$, the goal is to establish an estimate $\hat{y}$ of the response value $y^{(\text{test})}$ such that the out-of-sample squared error, i.e., the generalization error of the *target* task,

$$(2.8) \quad \mathcal{E}_{\text{out}} \triangleq \mathbb{E} \left[\left(\hat{y} - y^{(\text{test})} \right)^2 \right] = \sigma_\epsilon^2 + \mathbb{E} \left[\left\| \hat{\boldsymbol{\beta}} - \boldsymbol{\beta} \right\|_2^2 \right]$$

is minimized, where $\hat{y} = \mathbf{x}^{(\text{test})T} \hat{\boldsymbol{\beta}}$, and the second equality stems from the data model in (2.6).

The target task is addressed via linear regression that seeks for an estimate $\hat{\boldsymbol{\beta}} \in \mathbb{R}^d$ with a layout including three disjoint sets of coordinates $\mathcal{F}, \mathcal{T}, \mathcal{Z}$ that satisfy $\mathcal{F} \cup \mathcal{T} \cup \mathcal{Z} = \{1, \dots, d\}$ and correspond to three types of parameters:

- p parameters are *free* to be optimized and their coordinates are specified in $\mathcal{F}$.
- t parameters are *transferred* from the co-located coordinates of the estimate $\hat{\boldsymbol{\theta}}$ already formed for the source task. Only the free parameters of the *source* task are relevant to be transferred to the target task and, therefore, $\mathcal{T} \subseteq \mathcal{S}$ and $t \in \{0, \dots, \tilde{p}\}$. The transferred parameters are taken as is from $\hat{\boldsymbol{\theta}}$ and set fixed in the corresponding coordinates of $\hat{\boldsymbol{\beta}}$, i.e., for $k \in \mathcal{T}$, $\hat{\beta}_k = \hat{\theta}_k$ where $\hat{\beta}_k$ and $\hat{\theta}_k$ are the k^{th} components of $\hat{\boldsymbol{\beta}}$ and $\hat{\boldsymbol{\theta}}$, respectively.
- ℓ parameters are set to *zeros*. Their coordinates are included in $\mathcal{Z}$ and effectively correspond to ignoring features in the same coordinates of the input vectors.

Clearly, the layout should satisfy $p+t+\ell = d$. Then, the constrained linear regression problem for the target task is formulated as

$$(2.9) \quad \begin{aligned} \hat{\boldsymbol{\beta}} &= \arg \min_{\mathbf{b} \in \mathbb{R}^d} \|\mathbf{y} - \mathbf{X}\mathbf{b}\|_2^2 \\ \text{subject to } \mathbf{Q}_{\mathcal{T}}\mathbf{b} &= \mathbf{Q}_{\mathcal{T}}\hat{\boldsymbol{\theta}} \\ \mathbf{Q}_{\mathcal{Z}}\mathbf{b} &= \mathbf{0} \end{aligned}$$

where $\mathbf{Q}_{\mathcal{T}}$ and $\mathbf{Q}_{\mathcal{Z}}$ are the linear operators extracting the subvectors corresponding to the coordinates in $\mathcal{T}$ and $\mathcal{Z}$, respectively, from d -dimensional vectors. Here $\hat{\boldsymbol{\theta}} \in \mathbb{R}^d$ is the *pre-computed* estimate for the source task and considered a constant vector for the purpose of the target task. The examined transfer learning structure includes a single computation of the source task (2.3), followed by a single computation of the target task (2.9) that produces the eventual estimate of interest $\hat{\boldsymbol{\beta}}$ using the given $\hat{\boldsymbol{\theta}}$. The minimum ℓ_2 -norm solution of the target task in (2.9) is (see details in Appendix A.3)

$$(2.10) \quad \hat{\boldsymbol{\beta}} = \mathbf{Q}_{\mathcal{F}}^T \mathbf{X}_{\mathcal{F}}^+ (\mathbf{y} - \mathbf{X}_{\mathcal{T}} \hat{\boldsymbol{\theta}}_{\mathcal{T}}) + \mathbf{Q}_{\mathcal{T}}^T \hat{\boldsymbol{\theta}}_{\mathcal{T}}$$

where $\hat{\boldsymbol{\theta}}_{\mathcal{T}} \triangleq \mathbf{Q}_{\mathcal{T}} \hat{\boldsymbol{\theta}}$, $\mathbf{X}_{\mathcal{T}} \triangleq \mathbf{X} \mathbf{Q}_{\mathcal{T}}^T$, and $\mathbf{X}_{\mathcal{F}}^+$ is the pseudoinverse of $\mathbf{X}_{\mathcal{F}} \triangleq \mathbf{X} \mathbf{Q}_{\mathcal{F}}^T$. In the underparameterized case of $p \leq n$ the solution (2.10) almost surely reduces to the unique least squares form of

$$(2.11) \quad \hat{\boldsymbol{\beta}} = \mathbf{Q}_{\mathcal{F}}^T (\mathbf{X}_{\mathcal{F}}^T \mathbf{X}_{\mathcal{F}})^{-1} \mathbf{X}_{\mathcal{F}}^T (\mathbf{y} - \mathbf{X}_{\mathcal{T}} \hat{\boldsymbol{\theta}}_{\mathcal{T}}) + \mathbf{Q}_{\mathcal{T}}^T \hat{\boldsymbol{\theta}}_{\mathcal{T}}.$$

Note that the desired layout is indeed implemented by the $\hat{\beta}$ forms in (2.10), (2.11): the components corresponding to $\mathcal{Z}$ are zeros, the components corresponding to $\mathcal{T}$ are taken as is from $\hat{\theta}$, and only the p coordinates specified in $\mathcal{F}$ are adjusted for the purpose of minimizing the in-sample error in the optimization cost of (2.9) while considering the transferred parameters. In this paper we study the generalization ability of overparameterized solutions (i.e., when $p > n$) to the *target* task formulated in (2.9): In Sections 3–5 we analyze the minimum ℓ_2 -norm solution (2.10) and in Section 6 we examine other solutions.

3. The Double Descent Phenomenon in Transfer Learning.

3.1. Analytical Characterization of the Double Descent Phenomenon. Consider an *overall layout of coordinate subsets* $\mathcal{L} \triangleq \{\mathcal{S}, \mathcal{F}, \mathcal{T}, \mathcal{Z}\}$. The following theorem formulates the generalization error of the target task that is solved using a set of t parameters that are transferred as is from their co-located coordinates (indicated by $\mathcal{T}$) at the source task solution.

Theorem 3.1. *Let $\mathcal{L} = \{\mathcal{S}, \mathcal{F}, \mathcal{T}, \mathcal{Z}\}$ be a deterministic (i.e., non-random) coordinate layout. Then, the out-of-sample error of the target task has the form of*

$$(3.1) \quad \mathcal{E}_{\text{out}}^{(\mathcal{L})} = \begin{cases} \frac{n-1}{n-p-1} \left(\|\beta_{\mathcal{Z}}\|_2^2 + \sigma_{\epsilon}^2 + \mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})} \right) & \text{for } p \leq n-2, \\ \infty & \text{for } n-1 \leq p \leq n+1, \\ \frac{p-1}{p-n-1} \left(\|\beta_{\mathcal{Z}}\|_2^2 + \sigma_{\epsilon}^2 + \mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})} \right) + \left(1 - \frac{n}{p}\right) \|\beta_{\mathcal{F}}\|_2^2 & \text{for } p \geq n+2, \end{cases}$$

where

$$(3.2) \quad \mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})} \triangleq \underbrace{\left\| \mathbb{E} \left[\hat{\theta}_{\mathcal{T}} \right] - \beta_{\mathcal{T}} \right\|_2^2}_{\text{transfer bias}} + \underbrace{\mathbb{E} \left[\left\| \hat{\theta}_{\mathcal{T}} - \mathbb{E} \left[\hat{\theta}_{\mathcal{T}} \right] \right\|_2^2 \right]}_{\text{transfer variance}}$$

is the error in the transferred coordinates in the target task solution.

The last theorem is proved using non-asymptotic properties of Wishart matrices (see Appendix B).

The error formulation in (3.1) expresses a double descent form in terms of p and n , similarly to the formulation given in [3] for a linear regression task without transfer learning. However, here (3.1) introduces a new term, $\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})}$, that encapsulates the transfer learning aspect of our setting.

The formulation of the transfer error term in (3.2) demonstrates a bias-variance decomposition where the bias in the transferred parameters is measured with respect to the true parameters of the *target* task, and the variance depends only on the source task. This bias-variance decomposition is affected by the coordinates of the transferred parameters (i.e., $\mathcal{T}$) and also indirectly (through $\hat{\theta}_{\mathcal{T}}$) by the coordinates of the free parameters in the source task (i.e., $\mathcal{S}$), as we will see next.

The following corollary explicitly formulates the transfer bias and variance terms and their detailed dependency on the *parameterization level* of the *source* task, i.e., the $\tilde{p}, \tilde{n}$ pair, and the set $\mathcal{S}$ of free parameters in the source task (see proof in Appendix B.3).

Corollary 3.2. *The transfer bias term from (3.2) can be written as*

$$(3.3) \quad \text{Bias}_{\mathcal{T}}^2 \triangleq \left\| \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}] - \boldsymbol{\beta}_{\mathcal{T}} \right\|_2^2 = \left\| \mathbf{Q}_{\mathcal{T}} (r\mathbf{H} - \mathbf{I}_d) \boldsymbol{\beta} \right\|_2^2 \quad \text{where } r \triangleq \begin{cases} 1 & \text{for } \tilde{p} \leq \tilde{n}, \\ \frac{\tilde{n}}{\tilde{p}} & \text{for } \tilde{p} > \tilde{n}. \end{cases}$$

The transfer variance term from (3.2) can be formulated as

$$(3.4) \quad \begin{aligned} \text{Var}_{\mathcal{T}, \mathcal{S}} &\triangleq \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{T}} - \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}] \right\|_2^2 \right] \\ &= \begin{cases} t \left(\sigma_{\eta}^2 + \frac{\zeta_{\mathcal{S}^c} + (d - \tilde{p})\sigma_{\eta}^2 + \sigma_{\xi}^2}{\tilde{n} - \tilde{p} - 1} \right) & \text{for } 1 \leq \tilde{p} \leq \tilde{n} - 2, \\ \infty & \text{for } \tilde{n} - 1 \leq \tilde{p} \leq \tilde{n} + 1, \\ \frac{\tilde{n}}{\tilde{p}} \left(\frac{(\tilde{p} - \tilde{n})t\zeta_{\mathcal{S} \setminus \mathcal{T}} + ((\tilde{p} - \tilde{n})t - 1 + \frac{\tilde{n}}{\tilde{p}})\zeta_{\mathcal{T}}}{\tilde{p}^2 - 1} + t \left(\sigma_{\eta}^2 + \frac{\zeta_{\mathcal{S}^c} + (d - \tilde{p})\sigma_{\eta}^2 + \sigma_{\xi}^2}{\tilde{p} - \tilde{n} - 1} \right) \right) & \text{for } \tilde{p} \geq \tilde{n} + 2, \end{cases} \end{aligned}$$

where $\zeta_{\mathcal{T}} \triangleq \left\| \mathbf{Q}_{\mathcal{T}} \mathbf{H} \boldsymbol{\beta} \right\|_2^2$, $\zeta_{\mathcal{S} \setminus \mathcal{T}} \triangleq \left\| \mathbf{Q}_{\mathcal{S} \setminus \mathcal{T}} \mathbf{H} \boldsymbol{\beta} \right\|_2^2$, and $\zeta_{\mathcal{S}^c} \triangleq \left\| \mathbf{Q}_{\mathcal{S}^c} \mathbf{H} \boldsymbol{\beta} \right\|_2^2$.

Note that the out-of-sample error formulation in (3.1) depends on the parameterization level of the *target* task (i.e., the p, n pair) and also on the parameterization level of the *source* task (i.e., the $\tilde{p}, \tilde{n}$ pair) via the transfer bias and variance terms that are formulated in (3.3)-(3.4). The formulations in Corollary 3.2 imply that the target error peaks not only around $p = n$, but also around $\tilde{p} = \tilde{n}$ when transfer learning is applied ($t > 0$). This induces the *double double descent* behavior that will be demonstrated in the next subsection. Later, in Sections 4-5 we will analyze the formulations in Theorem 3.1 and Corollary 3.2 in more detail and characterize the conditions for beneficial transfer learning.

3.2. On-Average Analysis of Arbitrarily Selected Parameters. In this subsection, we consider the *overall layout of coordinate subsets* $\mathcal{L} \triangleq \{\mathcal{S}, \mathcal{F}, \mathcal{T}, \mathcal{Z}\}$ as a random structure. This will let us to formulate the expected value (with respect to $\mathcal{L}$) of the generalization error of interest. The simplified setting in this subsection provides useful insights towards Section 3.3 where we return to analyze the transfer of a set of parameters which is induced by a single layout $\mathcal{L}$ (i.e., in Section 3.3 we will return to use Theorem 3.1 and Corollary 3.2 where there is no expectation over a random $\mathcal{L}$).

For given $d, \tilde{p}, p$, and t , we consider a uniform distribution of the coordinate layout $\mathcal{L}$ which is defined as follows.

Definition 3.3. *A coordinate subset layout $\mathcal{L} = \{\mathcal{S}, \mathcal{F}, \mathcal{T}, \mathcal{Z}\}$ that is $\{\tilde{p}, p, t\}$ -uniformly distributed, for $\tilde{p} \in \{1, \dots, d\}$ and $(p, t) \in \{0, \dots, d\} \times \{0, \dots, \tilde{p}\}$ such that $p + t \leq d$, satisfies: $\mathcal{S}$ is uniformly chosen at random from all the subsets of $\tilde{p}$ unique coordinates of $\{1, \dots, d\}$. Given $\mathcal{S}$, the target-task coordinate layout $\{\mathcal{F}, \mathcal{T}, \mathcal{Z}\}$ is uniformly chosen at random from all the layouts where $\mathcal{F}, \mathcal{T}, \mathcal{Z}$ are three disjoint sets of coordinates that satisfy $\mathcal{F} \cup \mathcal{T} \cup \mathcal{Z} = \{1, \dots, d\}$ such that $|\mathcal{F}| = p$, $|\mathcal{T}| = t$ and $\mathcal{T} \subseteq \mathcal{S}$, and $|\mathcal{Z}| = d - p - t$.*

Then, the following formulates the out-of-sample error of the target task under expectation with respect to a uniformly distributed $\mathcal{L}$ (more details are provided in Appendix B.4).

Corollary 3.4. *Let $\mathcal{L} = \{\mathcal{S}, \mathcal{F}, \mathcal{T}, \mathcal{Z}\}$ be a coordinate subset layout that is $\{\tilde{p}, p, t\}$ -uniformly distributed. Then, the expected out-of-sample error of the target task has the form of*

$$(3.5) \quad \mathbb{E}_{\mathcal{L}} [\mathcal{E}_{\text{out}}] = \begin{cases} \frac{n-1}{n-p-1} \left(\left(1 - \frac{p+t}{d}\right) \|\boldsymbol{\beta}\|_2^2 + \sigma_{\epsilon}^2 + \mathcal{E}_{\text{transfer}} \right) & \text{for } p \leq n-2, \\ \infty & \text{for } n-1 \leq p \leq n+1, \\ \frac{p-1}{p-n-1} \left(\left(1 - \frac{p+t}{d}\right) \|\boldsymbol{\beta}\|_2^2 + \sigma_{\epsilon}^2 + \mathcal{E}_{\text{transfer}} \right) + \frac{p-n}{d} \|\boldsymbol{\beta}\|_2^2 & \text{for } p \geq n+2, \end{cases}$$

where

$$(3.6) \quad \mathcal{E}_{\text{transfer}} \triangleq \mathbb{E}_{\mathcal{L}} [\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})}] = \mathbb{E}_{\mathcal{L}} [\text{Bias}_{\mathcal{T}}^2] + \mathbb{E}_{\mathcal{L}} [\text{Var}_{\mathcal{T}, \mathcal{S}}]$$

$$(3.7) \quad \mathbb{E}_{\mathcal{L}} [\text{Bias}_{\mathcal{T}}^2] = \frac{t}{d} \|(r\mathbf{H} - \mathbf{I}_d) \boldsymbol{\beta}\|_2^2 \quad \text{where } r \triangleq \begin{cases} 1 & \text{for } \tilde{p} \leq \tilde{n}, \\ \frac{\tilde{n}}{\tilde{p}} & \text{for } \tilde{p} > \tilde{n}. \end{cases}$$

$$(3.8) \quad \mathbb{E}_{\mathcal{L}} [\text{Var}_{\mathcal{T}, \mathcal{S}}] = \begin{cases} t \left(\sigma_{\eta}^2 + \frac{(1-\frac{\tilde{p}}{d})\zeta + (d-\tilde{p})\sigma_{\eta}^2 + \sigma_{\xi}^2}{\tilde{n}-\tilde{p}-1} \right) & \text{for } 1 \leq \tilde{p} \leq \tilde{n}-2, \\ \infty & \text{for } \tilde{n}-1 \leq \tilde{p} \leq \tilde{n}+1, \\ \frac{\tilde{n}}{\tilde{p}} t \left(\frac{(\tilde{p}-\tilde{n})\tilde{p} + \frac{\tilde{n}}{\tilde{p}} - 1}{d(\tilde{p}^2-1)} \zeta + \sigma_{\eta}^2 + \frac{(1-\frac{\tilde{p}}{d})\zeta + (d-\tilde{p})\sigma_{\eta}^2 + \sigma_{\xi}^2}{\tilde{p}-\tilde{n}-1} \right) & \text{for } \tilde{p} \geq \tilde{n}+2, \end{cases}$$

and $\zeta \triangleq \|\mathbf{H}\boldsymbol{\beta}\|_2^2$.

Figure 1 presents the curves of $\mathbb{E}_{\mathcal{L}} [\mathcal{E}_{\text{out}}]$ with respect to the number of free parameters p in the target task, whereas the source task has $\tilde{p} = d$ free parameters. In Fig. 1, the solid-line curves correspond to analytical values induced by Corollary 3.4, and the respective empirically computed values are denoted by circles (all the presented results are for $d = 120$, $n = 20$, $\tilde{n} = 50$, $\|\boldsymbol{\beta}\|_2^2 = d$, $\sigma_{\epsilon}^2 = 0.05 \cdot d$, $\sigma_{\xi}^2 = 0.025 \cdot d$. See additional details in Appendix C. The number of free parameters p is upper bounded by $d - t$ that gets smaller for a larger number of transferred parameters t (see, in Fig. 1, the earlier stopping of the curves when t is larger). Observe that the generalization error peaks at $p = n$ and, then, decreases as p grows in the overparameterized range of $p > n + 1$. We identify this behavior as a *double descent* phenomenon, but without the first descent in the underparameterized range (double descent curves without the first descent are common when the parameters are selected arbitrarily, for example, see the results in [3, 7]).

The error of a trivial solution in the form of the null estimate, i.e., the estimate $\hat{\boldsymbol{\beta}}$ is all zeros, is presented as black dotted horizontal lines in the figures. In various settings where the source and target tasks are sufficiently related and the number of parameters is sufficiently far from the peak of the double descent error curve, the examined transfer learning method (with $t > 0$ arbitrarily selected parameters) outperforms the null estimate.

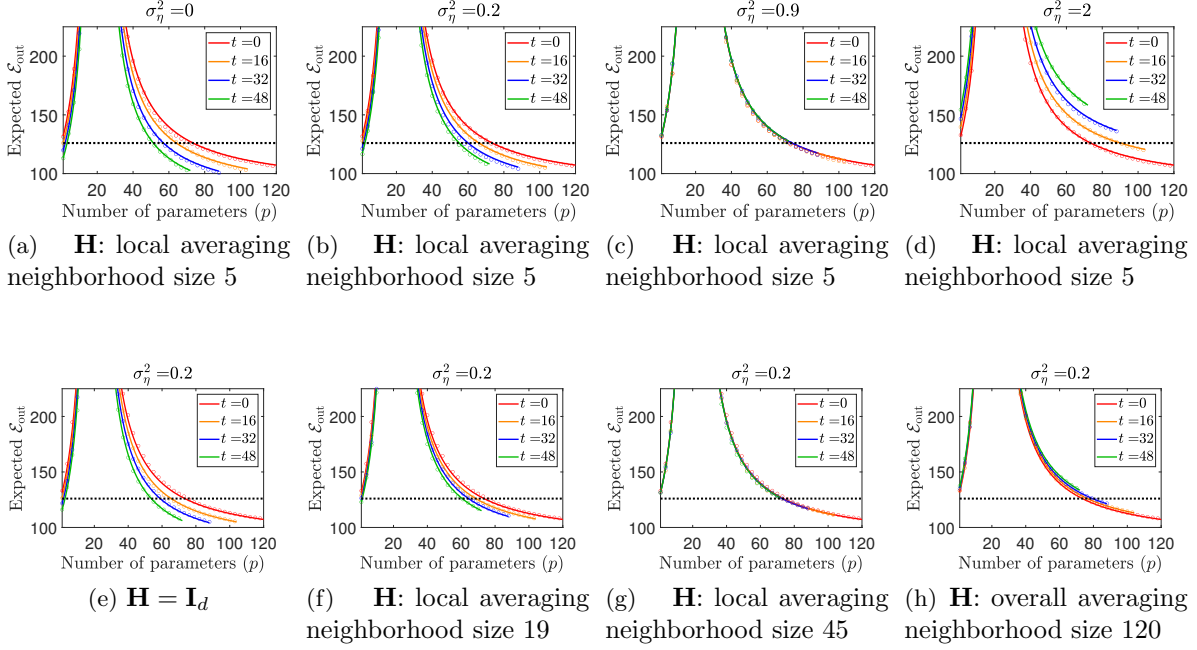


Figure 1: The expected generalization error of the target task, $\mathbb{E}_{\mathcal{L}}[\mathcal{E}_{\text{out}}]$, with respect to the number of free parameters (in the target task). The analytical values, computed using Corollary 3.4, are presented using solid-line curves, and the respective empirical results obtained from averaging over 250 experiments are denoted by circle markers. The horizontal dotted lines denote the error level of the null estimate. Each subfigure considers a different case of the source-target task relation (2.7) with a different pair of σ_η^2 and $\mathbf{H}$. The second row of subfigures corresponds to $\mathbf{H}$ operators that perform local averaging, each subfigure (e)-(h) is w.r.t. a different size of local averaging neighborhood. Each curve color refers to a different number of transferred parameters.

Each subfigure in Fig. 1 considers a different task relation with a different pair of noise level σ_η^2 and operator $\mathbf{H}$. The first row of subfigures in Fig. 1 emphasizes the effect of the noise variance σ_η^2 in the task relation model on the generalization errors in the target task. The second row of subfigures in Fig. 1 emphasizes the effect of the linear operator $\mathbf{H}$ in the task relation model on the generalization errors in the target task.

We can interpret the results in Figure 1 as examples for important cases of transfer learning settings. Figs. 1a,1e correspond to transfer learning between two *highly related tasks*, therefore, transfer learning is *beneficial* in the sense that for a given $p \notin \{n-1, n, n+1\}$ the error decreases as t increases (i.e., as more parameters are transferred instead of being omitted). Figs. 1b,1f correspond to transfer learning between two *moderately related tasks*, hence, transfer learning is still *beneficial*, but *less* than in the former case of highly related tasks. Figs. 1c,1g correspond to transfer learning between two *unrelated tasks* (although not extremely different), hence, transfer learning is *useless*, but not harmful (i.e., for a given p , the number of transferred

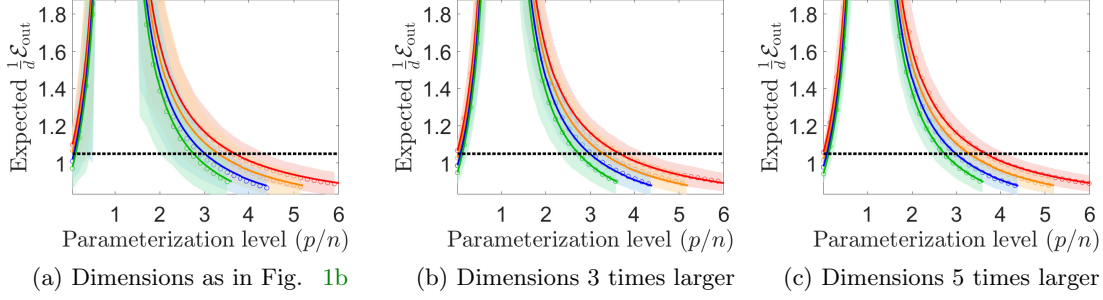


Figure 2: The concentration of the generalization error at three proportional scales of the same problem. The empirical standard deviations are denoted as shaded areas in colors corresponding to the on-average error curves (solid lines and markers denote the analytical and empirical evaluations of the expected error, respectively). Subfigure (a) corresponds to the setting of Fig. 1b. Subfigure (b) corresponds to a setting where all the dimensions and dimension-dependent quantities are 3 times their values in Fig. 1b (see a more detailed explanation in the text). Subfigure (c) corresponds to a setting where all dimensions and dimension-dependent quantities are 5 times their values in Fig. 1b. Lines, markers and areas in red correspond to $t = 0$ (no parameters are transferred); orange corresponds to transferring $t = 16 \times \frac{d}{120}$ parameters; blue corresponds to $t = 32 \times \frac{d}{120}$; green corresponds to $t = 48 \times \frac{d}{120}$. Note that $\frac{d}{120}$ equals to 1, 3, 5, in (a), (b), (c), respectively. The axes in this figure are normalized to be dimension-independent.

parameters t does not affect the out-of-sample error). Figs. 1d, 1h correspond to transfer learning between two *very different tasks* and, accordingly, transfer learning *degrades* the generalization performance (namely, for a given p , transferring more parameters increases the out-of-sample error).

Figure 2 demonstrates the empirical standard deviations of the generalization errors (denoted as shaded areas around the curves and markers that denote the average errors). The results show that the empirical errors are more concentrated around their (theoretical and empirical) expectations as the dimensions of the problem (e.g., input dimension, number of data examples and parameters) are proportionally increased. For this example, we consider the setting of Fig. 1b in three proportional scales of the dimensions and dimension-dependent quantities:

- Fig. 2a corresponds to the same dimensions as in Fig. 1b: $d = 120$, $n = 20$, $\tilde{n} = 50$, $t \in \{0, 16, 32, 48\}$, and local averaging with neighborhood size 5.
- Fig. 2b corresponds to dimensions that are proportionally increased 3 times: $d = 360$, $n = 60$, $\tilde{n} = 150$, $t \in \{0, 48, 96, 144\}$, and local averaging with neighborhood size 15.
- Fig. 2c corresponds to dimensions that are proportionally increased 5 times: $d = 600$, $n = 100$, $\tilde{n} = 250$, $t \in \{0, 80, 160, 240\}$, and local averaging with neighborhood size 25.

In all the settings in Figure 2, $\|\beta\|_2^2 = d$ and the noise levels are $\sigma_\epsilon^2 = 0.05 \times d$, $\sigma_\xi^2 = 0.025 \times d$, $\sigma_\eta^2 = 0.2$. More examples are provided in Fig. 12.

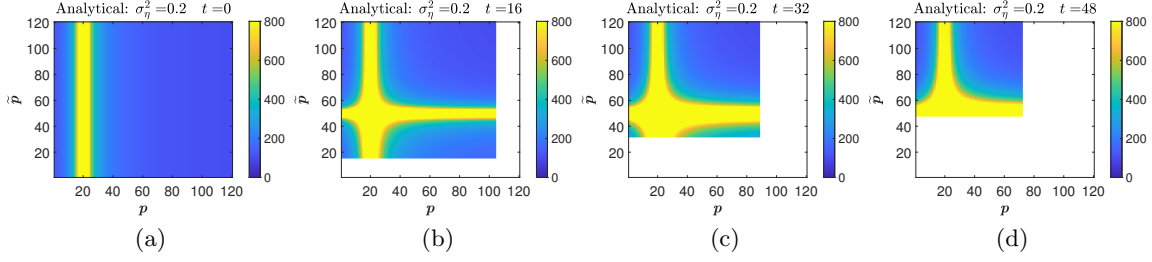


Figure 3: Analytical evaluation of the expected generalization error of the target task, $\mathbb{E}_{\mathcal{L}}[\mathcal{E}_{\text{out}}]$, with respect to the number of free parameters $\tilde{p}$ and p (in the source and target tasks, respectively). Each subfigure considers a different number of transferred parameters t . The white regions correspond to $(\tilde{p}, p)$ settings eliminated by the value of t in the specific subfigure. The yellow-colored areas correspond to values greater or equal to 800. All of the subfigures are for $\sigma_{\eta}^2 = 0.2$, $\mathbf{H}$ a local averaging operator with neighborhood of 5 samples, and β that has a piecewise-constant form (Fig. 15). See Fig. 10 for settings with additional values of σ_{η}^2 . See Fig. 11 for the corresponding empirical evaluation.

By considering the generalization error formula from Theorem 3.1 as a function of $\tilde{p}$ and p (i.e., the number of free parameters in the source and target tasks, respectively) we receive a *two-dimensional double descent* behavior as presented in Fig. 3 and its extended version Fig. 10 in Appendix C.1 that presents results for additional pairs of t and σ_{η}^2 . The results show a double descent trend along the p axis (with a peak at $p = n$) and also, when parameter transfer is applied (i.e., $t > 0$), a double descent trend along the $\tilde{p}$ axis (with a peak at $\tilde{p} = \tilde{n}$). Our solution structure implies that $\tilde{p} \in \{t, \dots, d\}$ and $p \in \{0, \dots, d - t\}$, hence, a larger number of transferred parameters t eliminates a larger portion of the underparameterized range of the source task and also eliminates a larger portion of the overparameterized range of the target task (see in Fig. 3 the white eliminated regions that grow with t). When t is high, the wide elimination of portions from the $(\tilde{p}, p)$ -plane hinders the complete form of the two-dimensional double descent phenomenon (see, e.g., Fig. 3d).

Conceptually, we can observe a *tradeoff between overparameterized learning and transfer learning* where parameters are transferred as is from their co-located coordinates of the source task solution: an increased transfer of parameters limits the level of overparameterization applicable in the target task and, in turn, this may limit the overall potential gains from the transfer learning. Yet, when the source task is *sufficiently related* to the target task (see, e.g., Figs. 1a, 1b), the parameter transfer can compensate (sometimes only partially) for an insufficient number of free parameters (in the target task). The last claim is also evident in Figs. 1a, 1b, 1e, 1f where, for $p > n + 1$, there is a range of generalization error values that is achievable by several settings of (p, t) pairs (i.e., specific error levels can be attained by curves of different colors in the same subfigure). E.g., in Fig. 1b the error achieved by $p = 112$ free parameters and no parameter transfer can be also achieved using $p = 70$ free parameters and $t = 48$ parameters transferred from the source task.

3.3. The Fragile Nature of Transfer Learning: Analysis of a Single Layout of Arbitrarily Selected Parameters. Section 3.2 considers an on-average error for a random coordinate layout. We now turn to discuss the generalization behavior of a *single* coordinate layout that was formed arbitrarily (i.e., without using any knowledge on the problem setting). Hence, we return to consider the generalization error $\mathcal{E}_{\text{out}}^{(\mathcal{L})}$ for a given layout $\mathcal{L}$ using Theorem 3.1 and Corollary 3.2 from Section 3.1.

Figure 4 shows the curves of $\mathcal{E}_{\text{out}}^{(\mathcal{L})}$ for specific coordinate layouts $\mathcal{L}$ that evolve with respect to the number of free parameters p in the target task (for more examples see Figures 13-14 in the Appendices). The excellent fit of the analytical results (that were computed using Theorem 3.1 and Corollary 3.2) to the empirical values (computed by averaging over 250 experiments with the same evolution of $\mathcal{L}$) is evident. The effect of the specific coordinate layout is clearly visible by the less-smooth curves (compared to the on-average results in Fig. 1). We examine two different cases for the true β : a linearly-increasing (Fig. 4a) and a sparse (Fig. 4e) layout of values, both have the same ℓ_2 norm. The difference in the true β forms yields error curves that significantly differ despite the use of the same sequential construction of the coordinate layouts with respect to p (e.g., compare Figs. 4f and 4b). The operator $\mathbf{H}$ in the task relation model greatly affects the generalization error curves as evident from comparing our results for different types of $\mathbf{H}$: an identity, local averaging (with neighborhood size 11), and discrete derivative operators (e.g., compare subfigures within the first row of Fig. 4). The results clearly show that the interplay among the structures of $\mathbf{H}$, β , and the coordinate layout significantly affects the generalization performance.

Our results also exhibit that an arbitrarily selected set $\mathcal{T}$ of t transferred parameters can be the best setting for a given set $\mathcal{F}$ of p free parameters but not necessarily for an extended set $\mathcal{F}' \supset \mathcal{F}$ of $p' > p$ free parameters. This is especially evident when the true parameters have a sparse form over an unknown support in the feature domain (i.e., true parameters with non-zero values are scarce and their coordinates are unknown). For example, see Fig. 4g where the green colored curve does not consistently maintain its relative vertical order with the red color curve at the overparameterized range of solutions; as a second example see Fig. 4h and observe that the blue and red curves do not maintain their vertical order in the overparameterized range. This exemplifies that, when arbitrary selection of parameters is employed due to unknown task relation, transfer learning settings can be fragile and hence finding a successful setting may require delicate, trial and error engineering. Therefore, our theory qualitatively explains similar practical aspects in deep neural networks (see, for example, [22]).

4. When is Transfer Learning Beneficial?. The formulation of the generalization error in the target task (Theorem 3.1 and Corollary 3.2) shows that the benefits from parameter transfer depend on various aspects of the learning setting. In this section we characterize the conditions for beneficial transfer from the viewpoint of the following question: Given a setting where $\mathcal{T}$ is the intended set of coordinates for parameter transfer, can avoiding this parameter transfer improve generalization?

4.1. Benefits in Transferred versus Zeroed Parameters. To accurately evaluate the difference in the out-of-sample error of the target task, consider the following definitions. First, recall the coordinate layout $\mathcal{L} = \{\mathcal{S}, \mathcal{F}, \mathcal{T}, \mathcal{Z}\}$ where $|\mathcal{S}| = \tilde{p}$, $|\mathcal{F}| = p$, $|\mathcal{T}| = t$. Second, we define a coordinate layout $\mathcal{L}' = \{\mathcal{S}, \mathcal{F}, \mathcal{T}', \mathcal{Z}'\}$ which is a modified version of $\mathcal{L}$ without trans-

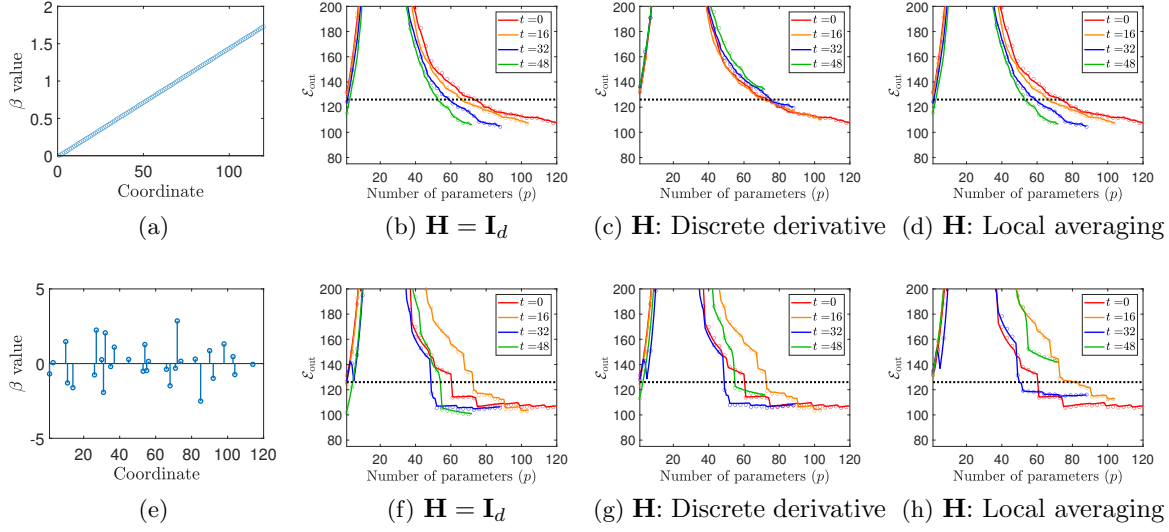


Figure 4: Analytical (solid lines) and empirical (circle markers) values of $\mathcal{E}_{\text{out}}^{(\mathcal{L})}$ for an arbitrary coordinate layout $\mathcal{L}$. In the first row of subfigures the true β has feature-domain values that follow a linear form, see (a). In the second row of subfigures the true β has a sparse form in the feature domain with non-zero values at 30 coordinates selected randomly out of the $d = 120$, see (e). In each row there are three error plots for different circulant forms of the operator $\mathbf{H}$ (here the local averaging is defined for a neighborhood of 11 samples). Here $\sigma_\eta^2 = 0.2$, $d = 120$, $n = 20$, $\tilde{n} = 50$, $\|\beta\|_2^2 = d$, $\sigma_\epsilon^2 = 0.05 \cdot d$, $\sigma_\xi^2 = 0.025 \cdot d$, and $\tilde{p} = d$ for all settings.

ferred parameters, specifically, $\mathcal{T}' = \emptyset$ and $\mathcal{Z}' = \mathcal{Z} \cup \mathcal{T}$. Namely, $\mathcal{L}'$ is obtained by zeroing all the parameters that are transferred in $\mathcal{L}$. Then, we define the following error difference due to transferring the parameters in $\mathcal{T}$ *instead of* zeroing them:

$$(4.1) \quad \Delta \mathcal{E}_{\text{TvsZ}}^{(\mathcal{L})} \triangleq \mathcal{E}_{\text{out}}^{(\mathcal{L})} - \mathcal{E}_{\text{out}}^{(\mathcal{L}')}$$

where $\mathcal{E}_{\text{out}}^{(\mathcal{L})}$ and $\mathcal{E}_{\text{out}}^{(\mathcal{L}')}$ are the out-of-sample errors in the target task for the coordinate layouts $\mathcal{L}$ and $\mathcal{L}'$, respectively. Using Theorem 3.1 we can write (4.1) as

$$(4.2) \quad \Delta \mathcal{E}_{\text{TvsZ}}^{(\mathcal{L})} = \left(\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})} - \|\beta_{\mathcal{T}}\|_2^2 \right) \times \begin{cases} \frac{n-1}{n-p-1} & \text{for } p \leq n-2, \\ \frac{p-1}{p-n-1} & \text{for } p \geq n+2, \end{cases}$$

where $\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})}$ was defined and formulated in Theorem 3.1 and Corollary 3.2. Note that $\Delta \mathcal{E}_{\text{TvsZ}}^{(\mathcal{L})}$ is undefined for $p \in \{n-1, n, n+1\}$. The definition in (4.1) implies that transferring the parameters in $\mathcal{T}$ is beneficial (over zeroing these coordinates in addition to the coordinates in $\mathcal{Z}$) if $\Delta \mathcal{E}_{\text{TvsZ}}^{(\mathcal{L})} < 0$. We define

$$(4.3) \quad \Delta \mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})} \triangleq \mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})} - \|\beta_{\mathcal{T}}\|_2^2.$$

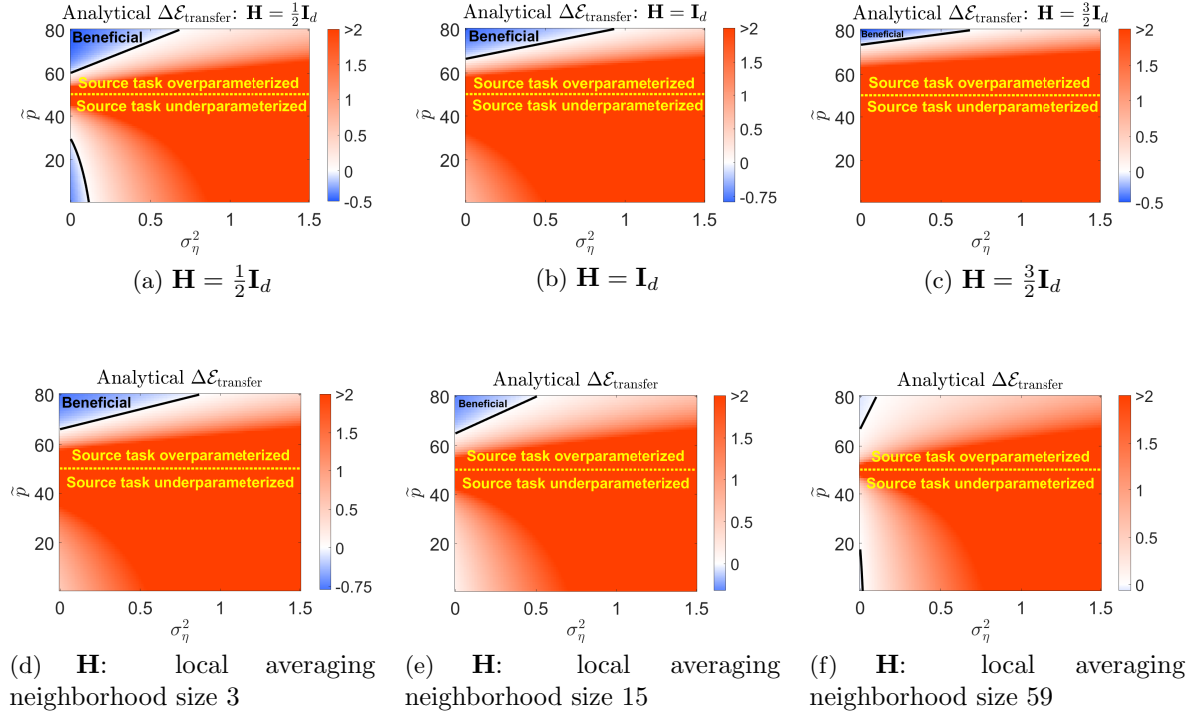


Figure 5: The analytical values of $\Delta\mathcal{E}_{\text{transfer}}$ (here normalized by t) as a function of $\tilde{p}$ and σ_η^2 . The positive and negative values of $\Delta\mathcal{E}_{\text{transfer}}$ appear in color scales of red and blue, respectively. The regions of negative values (appear in shades of blue) correspond to beneficial transfer of parameters (compared to zeroing them). The positive values were truncated at the value of 2 for the clarity of visualization. Each subfigure corresponds to a different task relation model induced by the definitions of $\mathbf{H}$ as: (a)-(c) different scalings of the identity matrix, (d)-(f) circulant matrices that correspond to local averaging operators with different neighborhood sizes. For all the subfigures, $d = 80$, $\tilde{n} = 50$, $\sigma_\xi^2 = 0.025 \cdot d$, $\|\beta\|_2^2 = d$. The components of β have a linear form (see Fig. 4a) in (a)-(c) and a piecewise-constant form (see Fig. 15) in (d)-(f). See corresponding empirical results in Figs. 16-17.

Then, according to (4.2), $\mathcal{E}_{\text{TVSZ}}^{(\mathcal{L})} < 0$ occurs when $p \notin \{n-1, n, n+1\}$ and $\Delta\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})} < 0$.

The examples in Fig. 5 show the values of $\Delta\mathcal{E}_{\text{transfer}} \triangleq \mathbb{E}_{\mathcal{L}} \left[\Delta\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})} \right]$ due to transferring an arbitrarily selected parameter. The error difference is presented with respect to the number $\tilde{p}$ of free parameters in the source task and the variance σ_η^2 of the noise in the task relation model. Each subfigure in Fig. 5 shows results for a different definition of $\mathbf{H}$. All the subfigures demonstrate that transferring an arbitrarily selected parameter is more beneficial when the solution to the source task is highly overparameterized; indeed, then the generalization error in the source task itself is lower due to its own double descent phenomena (e.g., recall the behavior of the error along the vertical axis in Fig. 3b-3d). As expected, a low level of noise in the task relation is also important for beneficial transfer.

The examples in Fig. 5 also show a somewhat unexpected behavior: parameter transfer

from an *underparameterized* source task can be more beneficial as the tasks are less related, e.g., compare the lower left corners in Figures 5a and 5b. To understand this observation mathematically one may recall the transfer bias and variance formulations in Corollary 3.2 and notice the following effects of $\mathbf{H}$, which in Figs. 5a-5b has the form of $\mathbf{H} = c\mathbf{I}_d$ where $c \in (0, 1)$. First, the transfer bias increases as $c \in (0, 1)$ gets smaller; also note that the transfer bias sums only over the $|\mathcal{T}| = t$ transferred coordinates. Second, the (underparameterized) transfer variance decreases as $c \in (0, 1)$ gets smaller; here note that the transfer variance is affected by $|\mathcal{S}^c| = d - \tilde{p}$ coordinates of $\mathbf{H} = c\mathbf{I}_d$ in addition to the number t of transferred parameters. In the experiment setting of Fig. 5a, the source task is underparameterized with a sufficiently large $d - \tilde{p}$ value such that the reduction in the transfer variance dominates the increase in the transfer bias, resulting in beneficial transfer.

More generally, the lesson here is that beneficial transfer learning can be achieved also in non-intuitive settings where the tasks are not necessarily highly related. In Section 5 we will provide additional insights into counter-intuitive beneficial cases.

4.2. Benefits in Transferred versus Free Parameters. Now we turn to discuss the question of when the transfer of the parameters in $\mathcal{T}$ is more beneficial than defining them as free for optimization.

Consider the coordinate layout $\mathcal{L} = \{\mathcal{S}, \mathcal{F}, \mathcal{T}, \mathcal{Z}\}$ where, as usual, $|\mathcal{S}| = \tilde{p}$, $|\mathcal{F}| = p$, $|\mathcal{T}| = t$. Here we also define a second coordinate layout $\mathcal{L}'' = \{\mathcal{S}, \mathcal{F}'', \mathcal{T}'', \mathcal{Z}\}$ which is a modified version of $\mathcal{L}$ without transferred parameters, specifically, $\mathcal{T}'' = \emptyset$ and $\mathcal{F}'' = \mathcal{F} \cup \mathcal{T}$. Hence, the consideration of $\mathcal{L}''$ reflects the case where all the transferred parameters in $\mathcal{L}$ are replaced by free parameters. Then, we define the following error difference term due to transferring the parameters in $\mathcal{T}$ *instead of allocating them as free parameters*:

$$(4.4) \quad \Delta\mathcal{E}_{\text{TvsF}}^{(\mathcal{L})} \triangleq \mathcal{E}_{\text{out}}^{(\mathcal{L})} - \mathcal{E}_{\text{out}}^{(\mathcal{L}'')}$$

where $\mathcal{E}_{\text{out}}^{(\mathcal{L})}$ and $\mathcal{E}_{\text{out}}^{(\mathcal{L}'')}$ are the out-of-sample errors in the target task (recall Theorem 3.1) for the coordinate layouts $\mathcal{L}$ and $\mathcal{L}''$, respectively. The definition in (4.4) implies that transferring the parameters in $\mathcal{T}$ is beneficial (over allocating these coordinates as t free parameters in addition to the p free parameters in $\mathcal{F}$) if $\Delta\mathcal{E}_{\text{TvsF}}^{(\mathcal{L})} < 0$.

We now turn to examine the effects of p and t on beneficial transfer in the sense of $\Delta\mathcal{E}_{\text{TvsF}}^{(\mathcal{L})} < 0$. For this, note that $\mathcal{L}''$ includes $p'' = p + t$ free parameters and hence is not necessarily in the same parameterization regime as $\mathcal{L}$ that includes p free parameters. The parameterization regimes of $\mathcal{L}$ and $\mathcal{L}''$ affect the behavior of the benefits from parameter transfer, as described by the following corollaries of Theorem 3.1 (these corollaries are simply proved by setting Eq. (3.1) in (4.4) and reorganizing the formulations).

Corollary 4.1. *For $p + t \leq n - 2$, namely, when both $\mathcal{L}$ and $\mathcal{L}''$ correspond to underparameterized settings, the error difference due to transferred versus free parameters is*

$$(4.5) \quad \Delta\mathcal{E}_{\text{TvsF}}^{(\mathcal{L})} = \frac{n-1}{n-p-1} \left(\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})} - \frac{t}{n-p-t-1} \left(\|\boldsymbol{\beta}_{\mathcal{Z}}\|_2^2 + \sigma_{\epsilon}^2 \right) \right).$$

Accordingly, the transfer of the $t > 0$ parameters in $\mathcal{T} \subseteq \mathcal{S}$ is more beneficial than allocating

them as free parameters (i.e., $\Delta\mathcal{E}_{\text{TvsF}}^{(\mathcal{T},\mathcal{S})} < 0$) if

$$(4.6) \quad \mathcal{E}_{\text{transfer}}^{(\mathcal{T},\mathcal{S})} < \frac{t}{n-p-t-1} \left(\|\beta_{\mathcal{Z}}\|_2^2 + \sigma_\epsilon^2 \right).$$

Corollary 4.2. For $p+t \geq n+2$, namely, when $\mathcal{L}''$ corresponds to an overparameterized setting, the error difference due to transferred versus free parameters is

$$(4.7) \quad \Delta\mathcal{E}_{\text{TvsF}}^{(\mathcal{L})} = \gamma_0 \mathcal{E}_{\text{transfer}}^{(\mathcal{T},\mathcal{S})} + \gamma_1 \|\beta_{\mathcal{F}}\|_2^2 + \gamma_2 \|\beta_{\mathcal{T}}\|_2^2 + \gamma_3 \left(\|\beta_{\mathcal{Z}}\|_2^2 + \sigma_\epsilon^2 \right)$$

where the coefficients $\gamma_0, \gamma_1, \gamma_2, \gamma_3$ depend on the parameterization level of $\mathcal{L}$:

- For an underparameterized $\mathcal{L}$ with $p \leq n-2$:
 $\gamma_0 = \frac{n-1}{n-p-1}$, $\gamma_1 = -\left(1 - \frac{n}{p+t}\right)$, $\gamma_2 = -\left(1 - \frac{n}{p+t}\right)$, $\gamma_3 = -\frac{n(n-1)-p(p+t-1)}{(p+t-n-1)(n-p-1)}$.
- For an overparameterized $\mathcal{L}$ with $p \geq n+2$:
 $\gamma_0 = \frac{p-1}{p-n-1}$, $\gamma_1 = -\frac{nt}{p(p+t)}$, $\gamma_2 = -\left(1 - \frac{n}{p+t}\right)$, $\gamma_3 = \frac{nt}{(p+t-n-1)(p-n-1)}$.

Accordingly, the transfer of the $t > 0$ parameters in $\mathcal{T} \subseteq \mathcal{S}$ is more beneficial than allocating them as free parameters (i.e., $\Delta\mathcal{E}_{\text{TvsF}}^{(\mathcal{L})} < 0$) if

$$(4.8) \quad \mathcal{E}_{\text{transfer}}^{(\mathcal{T},\mathcal{S})} < -\frac{\gamma_1}{\gamma_0} \|\beta_{\mathcal{F}}\|_2^2 - \frac{\gamma_2}{\gamma_0} \|\beta_{\mathcal{T}}\|_2^2 - \frac{\gamma_3}{\gamma_0} \left(\|\beta_{\mathcal{Z}}\|_2^2 + \sigma_\epsilon^2 \right).$$

Clearly, the formulation of $\Delta\mathcal{E}_{\text{TvsF}}^{(\mathcal{L})}$ is more intricate in the case of overparameterized $\mathcal{L}''$ in Corollary 4.2 than in the case of underparameterized $\mathcal{L}$ and $\mathcal{L}''$ in Corollary 4.1. Specifically, the set of free parameters $\mathcal{F}$ appears in (4.7)-(4.8) in addition to the sets $\mathcal{T}$ and $\mathcal{Z}$.

To intuitively understand the behavior of $\Delta\mathcal{E}_{\text{TvsF}}^{(\mathcal{L})}$, consider expectation over coordinates that are selected uniformly at random (recall Definition 3.3). Specifically, we define $\Delta\mathcal{E}_{\text{TvsF}} \triangleq \mathbb{E}_{\mathcal{L}} \left[\Delta\mathcal{E}_{\text{TvsF}}^{(\mathcal{L})} \right]$ while noting that $\mathcal{L}''$ is deterministically related to $\mathcal{L}$ and hence it is sufficient to consider the expectation over $\mathcal{L}$. Figure 6 shows the value of $\Delta\mathcal{E}_{\text{TvsF}}$ as a function of p and t for several operators $\mathbf{H}$ in the task relation. Figure 6 demonstrates the following typical behaviors:

- For underparameterized $\mathcal{L}$ and $\mathcal{L}''$ (i.e., $p+t \leq n-2$, which corresponds to the region to the left of the dotted black lines in the bottom subfigures in Fig. 6): Parameter transfer is more likely to be beneficial, and with larger gains, as $p+t$ increases towards $n-2$. Moreover, when (4.6) is satisfied, it is beneficial to have more transferred than free parameters (i.e., having $t > p$ for a fixed sum of $p+t$). The intuition here is that, when we are constrained to the underparameterized regime in both options (i.e., $\mathcal{L}$ and $\mathcal{L}''$), transferring parameters instead of having more free parameters can keep the error farther away from the peak of the double descent curve, which yields a beneficial transfer. Interestingly, due to the significant peak of the generalization error around the interpolation threshold, this behavior also occurs when the source and target tasks are quite different (see Fig. 6c).
- For underparameterized $\mathcal{L}$ and overparameterized $\mathcal{L}''$ (i.e., $p \leq n-2$ and $p+t \geq n+2$, which correspond to the region to the right of the dotted black lines in the bottom subfigures in Fig. 6): Parameter transfer is more likely to degrade as p increases

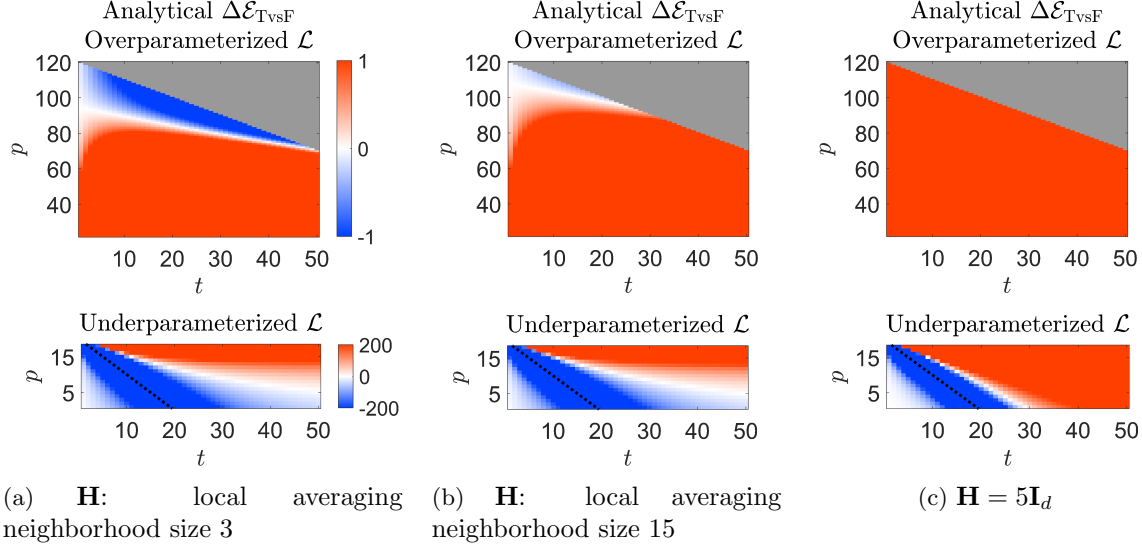


Figure 6: The analytical values of $\Delta\mathcal{E}_{\text{TvsF}}$ (namely, the expected error difference due to transfer of an arbitrarily-selected set of t parameters versus setting them as free parameters) as a function of t and p . The positive and negative values of $\Delta\mathcal{E}_{\text{TvsF}}$ appear in color scales of red and blue, respectively. The regions of negative values (appear in shades of blue) correspond to beneficial transfer of parameters (compared to defining them as free parameters). The gray regions correspond to $p+t > d$ where parameter transfer cannot be performed. For better visual clarity, the underparameterized and overparameterized cases of $\mathcal{L}$ are shown in different subfigures with significantly different range of values. Also, the positive values were truncated at the values of 1 and 200 in the overparameterized and underparameterized cases, respectively. Corresponding truncations are applied on the negative values at the values of -1 and -200. The dotted black line corresponds to $p+t=n$, which is the interpolation threshold of the auxiliary layout $\mathcal{L}''$. Each column of subfigures corresponds to a different task relation model induced by the definitions of $\mathbf{H}$. For all the subfigures, $d=120$, $\tilde{p}=d$, $\tilde{n}=50$, $\sigma_\xi^2=0.025 \cdot d$, $n=20$, $\sigma_\epsilon^2=0.05 \cdot d$, $\|\beta\|_2^2=d$ where β components have a piecewise-constant form (see Fig. 15). See corresponding empirical results in Fig. 19 in Appendix D.4.

towards $n-2$ while $p+t \geq n+2$. Moreover, increasing the number t of transferred parameters degrades the benefits from transfer learning. The intuition here is that, when the transfer option $\mathcal{L}$ is underparameterized and the no-transfer option $\mathcal{L}''$ is overparameterized, any increase in $p+t$ makes $\mathcal{L}''$ more overparameterized and hence farther away from the double descent peak (hence having an improved generalization ability).

- For *overparameterized* $\mathcal{L}$ and $\mathcal{L}''$ (i.e., $p \geq n+2$, which corresponds to the upper subfigures in Fig. 6): Here, the important behavior is that, for a given $\mathcal{T}$, parameter transfer improves as the number of free parameters p increases. Of course that this improvement becomes beneficial (at high overparameterization levels) only when the source task is sufficiently related to the target task (see, e.g., Figures 6a and 6b). The

intuition here is that, when both options (i.e., $\mathcal{L}$ and $\mathcal{L}''$) are in the overparameterized regime then the increase in p moves both $\mathcal{L}$ and $\mathcal{L}''$ away from the peak of the double descent. The transfer learning option $\mathcal{L}$ is closer to the double descent peak and hence typically improves more (due to the steeper slope at the corresponding part of the generalization error curve, e.g., see in Fig. 1) than the no-transfer case $\mathcal{L}''$.

5. The Optimal $\mathbf{H}$ in a Componentwise Task Relation. Theorem 3.1 shows that the task relation aspect is encapsulated in the term $\mathcal{E}_{\text{transfer}}^{(\mathcal{S}, \mathcal{T})}$. Consequently, we demonstrated in Section 4 that $\mathcal{E}_{\text{transfer}}^{(\mathcal{S}, \mathcal{T})}$ greatly affects the potential benefits from parameter transfer compared to both zero and free parameters. Specifically, $\Delta \mathcal{E}_{\text{TvsZ}}^{(\mathcal{L})}$ and $\Delta \mathcal{E}_{\text{TvsF}}^{(\mathcal{L})}$ both decrease as $\mathcal{E}_{\text{transfer}}^{(\mathcal{S}, \mathcal{T})}$ decreases. This motivates us to characterize the optimal task relation in the sense of minimum $\mathcal{E}_{\text{transfer}}^{(\mathcal{S}, \mathcal{T})}$ for a given coordinate layout $\mathcal{L}$. Namely, we characterize the best source task to transfer from when the other aspects of the parameter transfer are fixed.

Consider an operator $\mathbf{H}$ that has the diagonal form of

$$(5.1) \quad \mathbf{H} = \text{diag}\{\lambda_{\mathbf{H}}^{(1)}, \dots, \lambda_{\mathbf{H}}^{(d)}\}$$

where $\{\lambda_{\mathbf{H}}^{(j)}\}_{j=1}^d \in \mathbb{R}$. We refer to $\{\lambda_{\mathbf{H}}^{(j)}\}_{j=1}^d$ as the eigenvalues of $\mathbf{H}$ due to the fact that applying the same orthonormal rotation on the feature spaces of the source and target tasks can induce a non-diagonal $\mathbf{H}$ that its eigenvalues are the same $\{\lambda_{\mathbf{H}}^{(j)}\}_{j=1}^d$ as in (5.1). The diagonal form in (5.1) implies that the task relation in Eq. (2.7) reduces to the componentwise form of

$$(5.2) \quad \theta^{(j)} = \lambda_{\mathbf{H}}^{(j)} \beta^{(j)} + \eta^{(j)}, \quad j = 1, \dots, d$$

where $\theta^{(j)}$ and $\beta^{(j)}$ are the j^{th} components of the true parameter vectors of the source and target tasks, respectively, and $\eta^{(j)}$ is the j^{th} noise component in $\boldsymbol{\eta}$. Then, we can further develop the expressions from Corollary 3.2. First, the transfer bias term from (3.3) can be written as

$$(5.3) \quad \text{Bias}_{\mathcal{T}}^2 = \sum_{j \in \mathcal{T}} \left(r \lambda_{\mathbf{H}}^{(j)} - 1 \right)^2 \left(\beta^{(j)} \right)^2 \quad \text{where } r \triangleq \begin{cases} 1 & \text{for } \tilde{p} \leq \tilde{n}, \\ \frac{\tilde{n}}{\tilde{p}} & \text{for } \tilde{p} > \tilde{n}. \end{cases}$$

Second, the transfer variance term $\text{Var}_{\mathcal{T}, \mathcal{S}}$ from (3.4) can be also rewritten using

$$(5.4) \quad \zeta_{\mathcal{T}} = \sum_{j \in \mathcal{T}} \left(\lambda_{\mathbf{H}}^{(j)} \beta^{(j)} \right)^2, \quad \zeta_{\mathcal{S} \setminus \mathcal{T}} = \sum_{j \in \mathcal{S} \setminus \mathcal{T}} \left(\lambda_{\mathbf{H}}^{(j)} \beta^{(j)} \right)^2, \quad \zeta_{\mathcal{S}^c} = \sum_{j \in \mathcal{S}^c} \left(\lambda_{\mathbf{H}}^{(j)} \beta^{(j)} \right)^2.$$

Consequently, we get the following result (proof is provided in Appendix E.1).

Theorem 5.1. *Consider a task relation operator of the form $\mathbf{H} = \text{diag}\{\lambda_{\mathbf{H}}^{(1)}, \dots, \lambda_{\mathbf{H}}^{(d)}\}$. Also, $\tilde{p} \notin \{\tilde{n} - 1, \tilde{n}, \tilde{n} + 1\}$. Then, for given coordinate sets $\mathcal{S}$ (of size $\tilde{p}$) and $\mathcal{T} \subseteq \mathcal{S}$ (of size $t \leq \tilde{p}$), the transfer error term $\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})}$ attains its minimal value with respect to the eigenvalues of $\mathbf{H}$*

at

$$(5.5) \quad \text{for } j \in \mathcal{T}, \beta^{(j)} \neq 0 : \quad \lambda_{\mathbf{H}}^{(j)} = \begin{cases} 1 & \text{for } \tilde{p} \leq \tilde{n} - 2, \\ \frac{\tilde{p}^2 - 1}{\tilde{n}\tilde{p} - 1 + t(\tilde{p} - \tilde{n})} & \text{for } \tilde{p} \geq \tilde{n} + 2, \end{cases}$$

$$(5.6) \quad \text{for } j \in \mathcal{S} \setminus \mathcal{T}, \beta^{(j)} \neq 0 : \quad \lambda_{\mathbf{H}}^{(j)} = \begin{cases} \text{any value} & \text{for } \tilde{p} \leq \tilde{n} - 2, \\ 0 & \text{for } \tilde{p} \geq \tilde{n} + 2, \end{cases}$$

$$(5.7) \quad \text{for } j \in \mathcal{S}^c, \beta^{(j)} \neq 0 : \quad \lambda_{\mathbf{H}}^{(j)} = 0.$$

For $j \in \{1, \dots, d\}$ where $\beta^{(j)} = 0$, $\lambda_{\mathbf{H}}^{(j)}$ can have any value.

Let us interpret the meaning of Theorem 5.1 for the case of $\beta^{(j)} \neq 0$ for $j \in \{1, \dots, d\}$. The theorem shows that the linear operator in the optimal task relation has two different characterizations depending on whether the *source* task is under or over parameterized:

- For an *underparameterized source* task it is best that the true parameters in the transferred coordinates of the source task are the *same* (up to the additive noise terms from $\boldsymbol{\eta}$) as the corresponding true parameters of the target task (see (5.5) and recall (5.2)).
- For an *overparameterized source* task it is best that the true parameters in the transferred coordinates of the source task are *amplified* versions by a factor of $\frac{\tilde{p}^2 - 1}{\tilde{n}\tilde{p} - 1 + t(\tilde{p} - \tilde{n})}$ (see (5.5)) of the corresponding true parameters of the target task. This amplification intends to partially compensate for the increased transfer bias in the case of overparameterized source task (see the $\tilde{p} > \tilde{n}$ case in (5.3)). Indeed, the optimal amplification increases as the source task is more overparameterized (i.e., as p increases towards d). While this amplification reduces the transfer bias, it increases the transfer variance (see (3.4) and (5.4)) and hence the amplification is somewhat restrained. Also, the optimal amplification $\frac{\tilde{p}^2 - 1}{\tilde{n}\tilde{p} - 1 + t(\tilde{p} - \tilde{n})}$ decreases as the number t of transferred parameters increases. Specifically, when $t = \tilde{p}$ (i.e., all the free parameters of the source task are transferred to the target task) the optimal $\lambda_{\mathbf{H}}^{(j)}$ values for $j \in \mathcal{T}$ become 1 and there is no amplification.

Eq. (5.6) considers the coordinates of the free parameters of the source task that are *not* transferred (i.e., $\mathcal{S} \setminus \mathcal{T}$) and shows that the optimal true parameters of the source task in these coordinates are zeros in the overparameterized regime due to the dependency of the transfer variance on the true parameters of the source task at $\mathcal{S} \setminus \mathcal{T}$ (see (3.4) and (5.4)). In the underparameterized regime of the source task, there is no such dependency on $\boldsymbol{\theta}_{\mathcal{S} \setminus \mathcal{T}}$ and hence these true parameters may have any value without affecting the minimization of $\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})}$.

According to Eq. (5.7), in both the under and over parameterized cases, it is best that the true parameters of the source task are all zeros in the coordinates that do not participate in the source task solution (i.e., $\mathcal{S}^c$). This optimal form reduces the misspecification in the solution of the source task to originate only in the task relation noise $\boldsymbol{\eta}_{\mathcal{S}^c}$. Avoiding these misspecifications makes the estimate of each free source parameter closer to its true value (i.e., instead of also trying to compensate for omitted parameters that are informative) and, hence, this improves the relevance of the parameters $\boldsymbol{\theta}_{\mathcal{T}}$ that are transferred “as is” to the

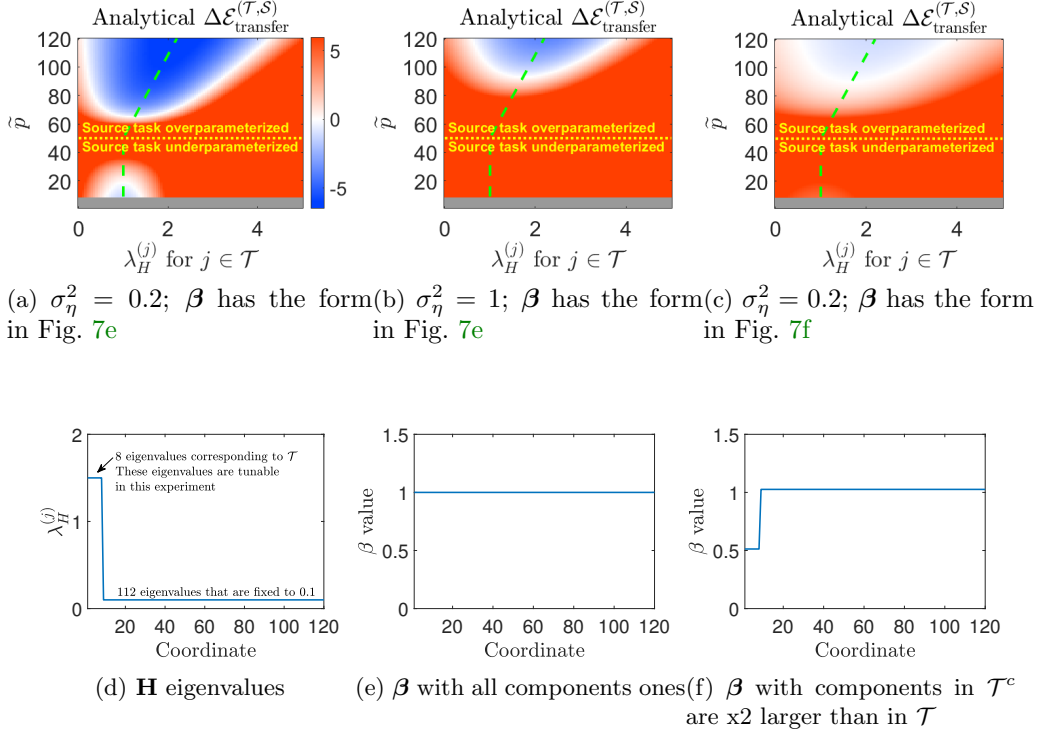


Figure 7: The analytical values of $\Delta\mathcal{E}_{\text{transfer}}^{(\mathcal{T},\mathcal{S})}$ as a function of $\tilde{p}$ and a value that determines the eigenvalues of $\mathbf{H}$ in $\mathcal{T}$. Here $\mathcal{T} = \{1, \dots, 8\}$. The operator $\mathbf{H}$ is diagonal with a main diagonal in the form described in subfigure (d). In subfigures (a)-(c), the positive and negative values of $\Delta\mathcal{E}_{\text{transfer}}^{(\mathcal{T},\mathcal{S})}$ appear in color scales of red and blue, respectively. The color scales are the same for all the subfigures in the first row. The regions of negative values (appear in shades of blue) correspond to beneficial transfer of parameters. The positive values were truncated at the value of 6 for the clarity of visualization. The gray regions correspond to $\tilde{p} < t$ where parameter transfer cannot be performed. The dashed green lines (in all subfigures) denote the optimal eigenvalues (corresponding to $\mathcal{T}$) as formulated in Theorem 5.1. Here, $d = 120$, $\tilde{n} = 50$, $\sigma_\xi^2 = 0.025 \cdot d$. Subfigure (e) shows the form of β in the experiments of subfigures (a), (b). Subfigure (f) shows the form of β in the experiments of subfigure (c). The corresponding empirical evaluations are provided in Fig. 20 in Appendix E.2.

target task from the co-located coordinates in the source task.

Theorem 5.1 describes the eigenvalues of $\mathbf{H}$ that minimize $\mathcal{E}_{\text{transfer}}^{(\mathcal{T},\mathcal{S})}$. However, $\mathcal{E}_{\text{transfer}}^{(\mathcal{T},\mathcal{S})}$ also depends on additional factors such as the noise level in the task relation and the true parameters β of the target task, which determine whether transferring the parameters in $\mathcal{T}$ is preferred, e.g., over zeroing them. Thus, evaluation of the sign of $\Delta\mathcal{E}_{\text{TVsZ}}^{(\mathcal{L})}$ is required for understanding whether the optimal eigenvalues of $\mathbf{H}$ induce beneficial transfer and, if so, how much can the eigenvalues deviate from their optimal values while still having a beneficial transfer.

In Fig. 7 we show the analytical values of $\Delta\mathcal{E}_{\text{transfer}}^{(\mathcal{T},\mathcal{S})}$ (recall the definition in (4.3)) as a function of $\tilde{p}$ and the eigenvalues of $\mathbf{H}$ that correspond to the transferred parameters (see Fig. 20 in Appendix E.2 for the corresponding empirical evaluations). Specifically, we consider a family of $\mathbf{H}$ operators that are diagonal in the feature domain and have the following step-shaped structure for their eigenvalues (see Fig. 7d): the first 8 eigenvalues are all equal to a value that varies among the $\mathbf{H}$ operators in this family; the next 112 eigenvalues are 0.1 for any $\mathbf{H}$ operator in this family. Here we consider transfer of the first eight parameters, i.e., $\mathcal{T} = \{1, \dots, 8\}$ that correspond to the eigenvalues with a tunable value in our evaluations (this is reflected by the horizontal axes in Figs. 7a-7c). The results in Fig. 7 demonstrate that the range of eigenvalues that induce beneficial transfer increases together with overparameterization (see the wider blue regions around the green dashed line that denotes the optimal eigenvalue from Theorem 5.1). Fig. 7a shows that benefits can be also obtained in underparameterized settings where the eigenvalues (in $\mathcal{T}$) are around 1, however, these eigenvalue regions are smaller and produce lower benefits than their overparameterized alternatives. The effect of the noise in the task model is also apparent by comparing Fig. 7a to Fig. 7b that shows results for the same settings but with a higher σ_η^2 . Specifically, the higher noise level in the task relation reduces the size of the beneficial regions and their gains in the overparameterized range, and leaves the underparameterized range without any beneficial regions. Next, compare Fig. 7a to Fig. 7c that shows results for the same settings but with β of the less favorable structure in Fig. 7f instead of the structure in Fig. 7e. This shows how the structure of the true parameters in β can also affect the parameter transfer performance.

6. Additional Linear Regression Methods. The previous sections presented analytical and empirical results on transfer learning between two least squares solutions that take their minimum ℓ_2 -norm forms (2.4), (2.10) in the overparameterized regime. In this section we examine the same transfer learning mechanism (i.e., transferring co-located parameters from the already-learned source model to the to-be-learned target model) for two more kinds of solutions to linear regression: least squares with the minimum ℓ_1 -norm form in the overparameterized regime, and ridge regression.

6.1. The Minimum ℓ_1 -Norm Interpolator. This setting also emerges from the optimization problems formulated in (2.3) and (2.9) for the source and target tasks, respectively. In the underparameterized regime, the source and target tasks almost surely have the unique least squares solutions as provided in (2.5) and (2.11).

For an overparameterized source task, i.e., $\tilde{p} > \tilde{n}$, the problem (2.3) has infinite interpolating solutions among which we choose here the minimum ℓ_1 -norm solution:

$$(6.1) \quad \begin{aligned} \hat{\boldsymbol{\theta}} &= \arg \min_{\mathbf{r} \in \mathbb{R}^d} \|\mathbf{r}\|_1 \\ \text{subject to } &\mathbf{Q}_{\mathcal{S}^c} \mathbf{r} = \mathbf{0} \\ &\mathbf{Z} \mathbf{r} = \mathbf{v}. \end{aligned}$$

This solution equals to $\hat{\boldsymbol{\theta}}$ where $\hat{\boldsymbol{\theta}}_{\mathcal{S}^c} = \mathbf{0}$ and $\hat{\boldsymbol{\theta}}_{\mathcal{S}} = \arg \min_{\mathbf{k} \in \mathbb{R}^{\tilde{p}}} \|\mathbf{k}\|_1$ subject to $\mathbf{Z}_{\mathcal{S}} \mathbf{k} = \mathbf{v}$, which is a basis pursuit problem [6] that can be addressed via linear programming methods.

For an overparameterized target task, i.e., $p > n$, the problem (2.9) has infinite interpolating solutions, from which we choose here the minimum ℓ_1 -norm solution

$$(6.2) \quad \begin{aligned} \hat{\beta} &= \arg \min_{\mathbf{b} \in \mathbb{R}^d} \|\mathbf{b}\|_1 \\ \text{subject to } \mathbf{Q}_{\mathcal{T}} \mathbf{b} &= \mathbf{Q}_{\mathcal{T}} \hat{\theta} \\ \mathbf{Q}_{\mathcal{Z}} \mathbf{b} &= \mathbf{0} \\ \mathbf{X} \mathbf{b} &= \mathbf{y} \end{aligned}$$

that can be formulated also as $\hat{\beta}$ where $\hat{\beta}_{\mathcal{Z}} = \mathbf{0}$, $\hat{\beta}_{\mathcal{T}} = \hat{\theta}_{\mathcal{T}}$ and

$$(6.3) \quad \begin{aligned} \hat{\beta}_{\mathcal{F}} &= \arg \min_{\mathbf{f} \in \mathbb{R}^p} \|\mathbf{f}\|_1 \\ \text{subject to } \mathbf{X}_{\mathcal{F}} \mathbf{f} &= \mathbf{y} - \mathbf{X}_{\mathcal{T}} \hat{\theta}_{\mathcal{T}}, \end{aligned}$$

which is a basis pursuit problem that can be solved via linear programming.

The empirical out-of-sample errors of the minimum ℓ_1 -norm transfer learning are shown in Figure 8 as blue curves where the shade of blue denotes the number of transferred parameters. The results demonstrate that the double descent behavior occurs also for the minimum ℓ_1 -norm solution. Double descent phenomena for minimum ℓ_1 -norm solutions to linear regression (without transfer learning) were studied in [19, 18]. The asymptotic analyses in [15, 27] show that a triple descent can be observed if the true parameters are sufficiently sparse. Here we do not clearly observe an additional error peak in the overparameterized regime, which could possibly be due to the misspecification strategy (i.e., zeroing parameters in predetermined coordinates) that we use for controlling the parameterization levels in our non-asymptotic setting. For β of a linear shape the minimum ℓ_2 -norm solution is better than the minimum ℓ_1 -norm solution in the entire overparameterized range (see Figs. 8a-8c). Nevertheless, for a sparse β the minimum ℓ_1 -norm solution can provide the best performance in the high overparameterization levels due to its sparsity-promoting nature (see Figs. 8d, 8g-8i).

6.2. Ridge Regression. Now we turn to formulate our transfer learning approach for ridge regression. So far we considered the optimization problems (2.3) and (2.9) that do not include explicit regularization terms in their cost functions. Hence, the ridge regression extension of (2.3) is

$$(6.4) \quad \hat{\theta} = \arg \min_{\mathbf{r} \in \mathbb{R}^d} \|\mathbf{v} - \mathbf{Z}\mathbf{r}\|_2^2 + \tilde{\alpha} \|\mathbf{r}\|_2^2 \quad \text{subject to } \mathbf{Q}_{\mathcal{S}^c} \mathbf{r} = \mathbf{0},$$

where $\tilde{\alpha} > 0$ determines the level of ridge regularization. The formulation in (6.4) is equivalent to $\hat{\theta}$ where $\hat{\theta}_{\mathcal{S}^c} = 0$ and

$$(6.5) \quad \hat{\theta}_{\mathcal{S}} = \arg \min_{\mathbf{k} \in \mathbb{R}^{\tilde{p}}} \|\mathbf{v} - \mathbf{Z}_{\mathcal{S}} \mathbf{k}\|_2^2 + \tilde{\alpha} \|\mathbf{k}\|_2^2.$$

The optimization in (6.5) has a standard ridge regression form and, thus, its closed-form solution is

$$(6.6) \quad \hat{\theta}_{\mathcal{S}} = (\mathbf{Z}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}} + \tilde{\alpha} \mathbf{I}_{\tilde{p}})^{-1} \mathbf{Z}_{\mathcal{S}}^T \mathbf{v}.$$

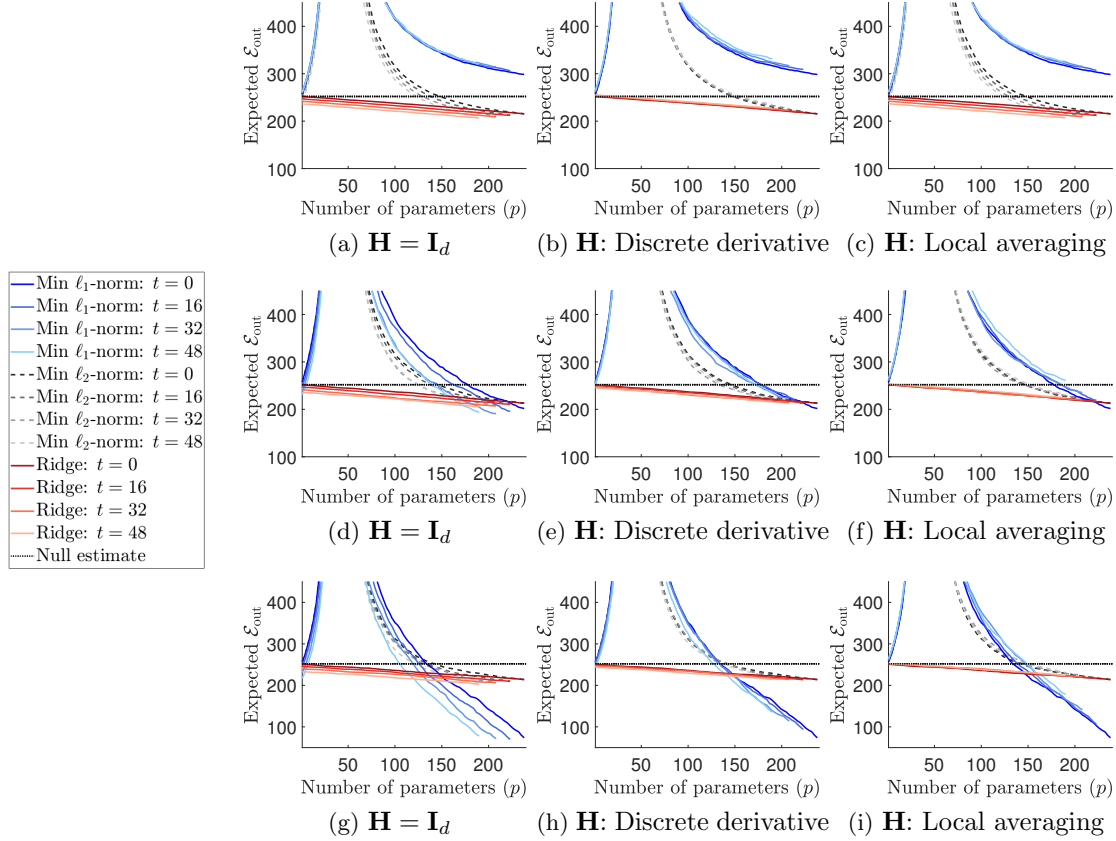


Figure 8: The expected generalization error of the target task, $\mathbb{E}_{\mathcal{L}}[\mathcal{E}_{\text{out}}]$, compared for minimum ℓ_2 -norm, minimum ℓ_1 -norm, and ridge solutions. In this figure all the errors are empirically evaluated. The errors of the minimum ℓ_1 -norm solution are shown in curves of blue shades. The errors of the ridge regression are shown in curves of red shades. The errors of the minimum ℓ_2 -norm solution are shown in dashed curves of gray shades. For each of the solution types the darkest shade corresponds to no transfer ($t = 0$), a lighter shade denotes more transferred parameters (larger t), and the lightest shade corresponds to $t = 48$. In the first row of subfigures the true β has values that follow a linear form. In the second row of subfigures the true β has a sparse form with non-zero values at 25% of the coordinates (selected randomly out of the $d = 240$). In the third row of subfigures the true β is sparse with only 5% non-zero values. Examples for linear and sparse β of a smaller dimension are provided in Figs. 4a, 4e. In each row there are three error plots for different circulant forms of the operator $\mathbf{H}$ (here the local averaging is defined for a neighborhood of 11 samples). Here $\sigma_\eta^2 = 0.2$, $d = 240$, $n = 40$, $\tilde{n} = 100$, $\|\beta\|_2^2 = d$, $\sigma_\epsilon^2 = 0.05 \cdot d$, $\sigma_\xi^2 = 0.025 \cdot d$, and $\tilde{p} = d$ for all settings.

In appendix F.1 we formulate the out-of-sample error of the ridge regression solution to the source task, and show that optimal tuning is provided by $\tilde{\alpha} = \frac{\tilde{p}(\sigma_\xi^2 + \|\theta_{S^c}\|_2^2)}{\|\theta_S\|_2^2}$. In our experiments

we assume that $\|\boldsymbol{\theta}_{\mathcal{S}}\|_2^2, \|\boldsymbol{\theta}_{\mathcal{S}^c}\|_2^2$ are unknown and only $\|\boldsymbol{\theta}\|_2^2, \sigma_\xi^2$ are known about the source task. Thus, we assume $\|\boldsymbol{\theta}_{\mathcal{S}}\|_2^2 \approx \frac{\tilde{p}}{d} \|\boldsymbol{\theta}\|_2^2, \|\boldsymbol{\theta}_{\mathcal{S}^c}\|_2^2 \approx \left(1 - \frac{\tilde{p}}{d}\right) \|\boldsymbol{\theta}\|_2^2$ that let us to approximate the optimal tuning using $\tilde{\alpha} = \frac{d\sigma_\xi^2}{\|\boldsymbol{\theta}\|_2^2} + d - \tilde{p}$.

Proceeding to the target task, the ridge regression extension of (2.9) is

$$(6.7) \quad \begin{aligned} \hat{\boldsymbol{\beta}} &= \arg \min_{\mathbf{b} \in \mathbb{R}^d} \|\mathbf{y} - \mathbf{X}\mathbf{b}\|_2^2 + \alpha \|\mathbf{b}\|_2^2 \\ \text{subject to } \mathbf{Q}_{\mathcal{T}}\mathbf{b} &= \mathbf{Q}_{\mathcal{T}}\hat{\boldsymbol{\theta}} \\ \mathbf{Q}_{\mathcal{Z}}\mathbf{b} &= \mathbf{0}, \end{aligned}$$

for $\alpha > 0$. The solution of (6.7) has the closed form of $\hat{\boldsymbol{\beta}}$ where $\hat{\boldsymbol{\beta}}_{\mathcal{Z}} = \mathbf{0}, \hat{\boldsymbol{\beta}}_{\mathcal{T}} = \hat{\boldsymbol{\theta}}_{\mathcal{T}}$, and

$$(6.8) \quad \hat{\boldsymbol{\beta}}_{\mathcal{F}} = (\mathbf{X}_{\mathcal{F}}^T \mathbf{X}_{\mathcal{F}} + \alpha \mathbf{I}_p)^{-1} \mathbf{X}_{\mathcal{F}}^T (\mathbf{y} - \mathbf{X}_{\mathcal{T}} \hat{\boldsymbol{\theta}}_{\mathcal{T}}).$$

In appendix F.2 we provide further details on (6.8), the corresponding out-of-sample error, and show that optimal tuning is given by $\alpha = \frac{p(\sigma_\epsilon^2 + \|\boldsymbol{\beta}_{\mathcal{Z}}\|_2^2 + \mathbb{E}[\|\boldsymbol{\beta}_{\mathcal{T}} - \hat{\boldsymbol{\theta}}_{\mathcal{T}}\|_2^2])}{\|\boldsymbol{\beta}_{\mathcal{F}}\|_2^2}$. In our experiments we approximate the optimal tuning by setting $\alpha = \frac{d\sigma_\epsilon^2}{\|\boldsymbol{\beta}\|_2^2} + d - p - t + t \frac{\|(\mathbf{I}_d - \mathbf{H})\boldsymbol{\beta}\|_2^2 + d\sigma_\eta^2}{\|\boldsymbol{\beta}\|_2^2}$ that assumes the knowledge of $\|\boldsymbol{\beta}\|_2^2, \|(\mathbf{I}_d - \mathbf{H})\boldsymbol{\beta}\|_2^2, \sigma_\eta, \sigma_\epsilon$. See Appendix F.2 for more details.

The empirical out-of-sample errors for the ridge regression transfer learning appear in Figure 8 as red curves where the shade of red denotes the number of transferred parameters. As expected, the ridge regularization (which is approximately optimally tuned) resolves the error peak and eliminates the double descent shape of the minimum norm interpolating solutions. In the case of suboptimally tuned ridge regularization the double descent peak may not be fully resolved (see Fig. 9). For $\boldsymbol{\beta}$ with a linear shape, the ridge regularization suits best among the examined methods at any parameterization level, except for the maximal overparameterization levels where the minimum ℓ_2 -norm solution performs comparably (see Figs. 8a-8c). However, for a sparse $\boldsymbol{\beta}$ the minimum ℓ_1 -norm outperforms the ridge solution for high overparameterization levels (see Figs. 8d-8f).

7. Conclusions. In this work we have established an analytical framework for the fundamental study of transfer learning in conjunction with overparameterized models. We used least squares solutions to linear regression problems for shedding clarifying light on the generalization performance induced for a target task addressed using parameters transferred from an already completed source task. We formulated the generalization error of the target task and presented its two-dimensional double descent shape as a function of the number of free parameters individually available in the source and target tasks. We characterized the conditions for a beneficial transfer of parameters and demonstrated its high sensitivity to the delicate interaction among crucial aspects such as the source-target task relation, the specific choice of transferred parameters, and the form of the true solution. We importantly showed that overparameterized transfer learning is not necessarily improved by using a source task which is closer or identical to the target task. Our focus was mainly on the analytical and empirical study of the minimum ℓ_2 -norm solution to overparameterized transfer learning. We

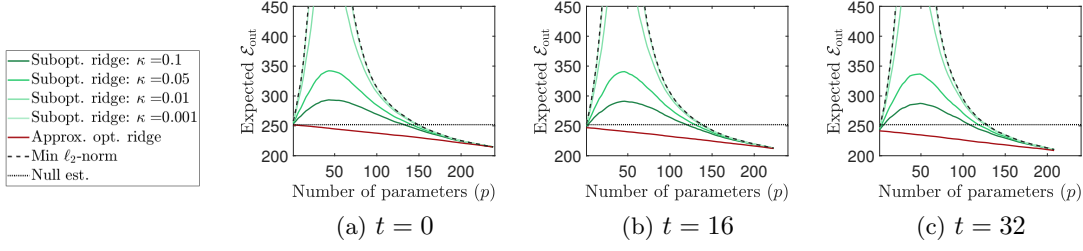


Figure 9: The expected generalization error of the target task, $\mathbb{E}_{\mathcal{L}}[\mathcal{E}_{\text{out}}]$, for transfer learning with suboptimally tuned ridge regularization. Each of the subfigures corresponds to a different number of transferred parameters t , and shows the error curves for four suboptimal ridge tunings where the parameters α and $\tilde{\alpha}$ are κ times their (approximately) optimal values. Here β has linear shape, $\mathbf{H}$ is a local averaging operator (over a neighborhood size 11), $\sigma_{\eta}^2 = 0.2$, $d = 240$, $n = 40$, $\tilde{n} = 100$, $\|\beta\|_2^2 = d$, $\sigma_{\epsilon}^2 = 0.05 \cdot d$, $\sigma_{\xi}^2 = 0.025 \cdot d$, and $\tilde{p} = d$.

also empirically examined the performance of the minimum ℓ_1 -norm solution and ridge regression in our transfer learning framework. We believe that our work opens a new research direction for the fundamental understanding of the generalization ability of transfer learning designs. Future work may study the theory and practice of additional transfer learning layouts such as fine tuning of the transferred parameters, inclusion of various regularization methods, well-specified and other settings where the task relation model is known (to some extent) and utilized in the actual learning process.

Acknowledgments. This work was supported by NSF grants CCF-1911094, IIS-1838177, and IIS-1730574; ONR grants N00014-18-12571, N00014-20-1-2534, and MURI N00014-20-1-2787; AFOSR grant FA9550-18-1-0478; and a Vannevar Bush Faculty Fellowship, ONR grant N00014-18-1-2047.

Appendices. The following appendices support the main paper as follows. Appendix A provides additional details on the mathematical developments leading to the formulations in Section 2 of the main paper. In Appendix B we present the proofs of Theorem 3.1 and Corollaries 3.2, 3.4 from Section 3, which formulate the generalization error of the target task. Appendix C provides additional empirical results and details for Section 3 of the main paper. In Appendices D, E we provide analytical proofs and empirical results for Sections 4, 5 of the main paper. In Appendix F we provide the mathematical developments for the ridge regression setting of our transfer learning problem from Section 6.2 of the main paper.

Appendix A. Mathematical Developments for Section 2.

A.1. The Estimate $\hat{\theta}$ in Eq. (2.4). Let us solve the optimization problem provided in (2.3). Using the relation $\mathbf{Q}_{\mathcal{S}}^T \mathbf{Q}_{\mathcal{S}} + \mathbf{Q}_{\mathcal{S}^c}^T \mathbf{Q}_{\mathcal{S}^c} = \mathbf{I}_d$ we can rewrite (2.3) as

$$(A.1) \quad \begin{aligned} \hat{\theta} &= \arg \min_{\mathbf{r} \in \mathbb{R}^d} \|\mathbf{v} - \mathbf{Z}_{\mathcal{S}} \mathbf{Q}_{\mathcal{S}} \mathbf{r} - \mathbf{Z}_{\mathcal{S}^c} \mathbf{Q}_{\mathcal{S}^c} \mathbf{r}\|_2^2 \\ &\text{subject to } \mathbf{Q}_{\mathcal{S}^c} \mathbf{r} = \mathbf{0} \end{aligned}$$

where $\mathbf{Z}_S \triangleq \mathbf{Z}\mathbf{Q}_S^T$ and $\mathbf{Z}_{S^c} \triangleq \mathbf{Z}\mathbf{Q}_{S^c}^T$. By setting the equality constraint in the optimization cost, the problem in (A.1) becomes

$$(A.2) \quad \begin{aligned} \hat{\boldsymbol{\theta}} &= \arg \min_{\mathbf{r} \in \mathbb{R}^d} \|\mathbf{v} - \mathbf{Z}_S \mathbf{Q}_S \mathbf{r}\|_2^2 \\ &\text{subject to } \mathbf{Q}_{S^c} \mathbf{r} = \mathbf{0}. \end{aligned}$$

Without the equality constraint, (A.2) is just an unconstrained least squares problem that its minimum ℓ_2 -norm solution is

$$(A.3) \quad \hat{\boldsymbol{\theta}} = \mathbf{Q}_S^T \mathbf{Z}_S^+ \mathbf{v}$$

where $\mathbf{Z}_S^+$ is the Moore-Penrose pseudoinverse of $\mathbf{Z}_S$. Note that $\hat{\boldsymbol{\theta}}$ in (A.3) satisfies the equality constraint in (A.1) and, therefore, (A.3) is also the solution for the constrained optimization problems in (A.1), (A.2), and (2.3).

A.2. The Double Descent Formulation for the Generalization Error of the Source Task. The generalization error of a single linear regression problem (that includes noise) in non-asymptotic settings is provided in [3] for a given coordinate subset (i.e., deterministic $\mathcal{S}$ in our terms). The result from [3] can be written in our notations as

$$(A.4) \quad \tilde{\mathcal{E}}_{\text{out}} = \begin{cases} \frac{\tilde{n}-1}{\tilde{n}-\tilde{p}-1} \left(\|\boldsymbol{\theta}_{S^c}\|_2^2 + \sigma_\xi^2 \right) & \text{for } \tilde{p} \leq \tilde{n} - 2, \\ \infty & \text{for } \tilde{n} - 1 \leq \tilde{p} \leq \tilde{n} + 1, \\ \frac{\tilde{p}-1}{\tilde{p}-\tilde{n}-1} \left(\|\boldsymbol{\theta}_{S^c}\|_2^2 + \sigma_\xi^2 \right) + \frac{\tilde{p}-\tilde{n}}{\tilde{p}} \|\boldsymbol{\theta}_S\|_2^2 & \text{for } \tilde{p} \geq \tilde{n} + 2. \end{cases}$$

In case that the coordinate subset $\mathcal{S}$ is uniformly chosen at random from all the subsets of $\tilde{p} \in \{1, \dots, d\}$ unique coordinates of $\{1, \dots, d\}$, then we get that $\mathbb{E}_{\mathcal{S}} [\|\boldsymbol{\theta}_S\|_2^2] = \frac{\tilde{p}}{d} \|\boldsymbol{\theta}\|_2^2$ and $\mathbb{E}_{\mathcal{S}} [\|\boldsymbol{\theta}_{S^c}\|_2^2] = \frac{d-\tilde{p}}{d} \|\boldsymbol{\theta}\|_2^2$. Accordingly, the expectation over $\mathcal{S}$ of the generalization error of the source task leads to the following result

$$(A.5) \quad \mathbb{E}_{\mathcal{S}} [\tilde{\mathcal{E}}_{\text{out}}] = \begin{cases} \frac{\tilde{n}-1}{\tilde{n}-\tilde{p}-1} \left(\left(1 - \frac{\tilde{p}}{d}\right) \|\boldsymbol{\theta}\|_2^2 + \sigma_\xi^2 \right) & \text{for } \tilde{p} \leq \tilde{n} - 2, \\ \infty & \text{for } \tilde{n} - 1 \leq \tilde{p} \leq \tilde{n} + 1, \\ \frac{\tilde{p}-1}{\tilde{p}-\tilde{n}-1} \left(\left(1 - \frac{\tilde{p}}{d}\right) \|\boldsymbol{\theta}\|_2^2 + \sigma_\xi^2 \right) + \frac{\tilde{p}-\tilde{n}}{d} \|\boldsymbol{\theta}\|_2^2 & \text{for } \tilde{p} \geq \tilde{n} + 2. \end{cases}$$

The formulation in (A.5) considers $\boldsymbol{\theta}$ as a deterministic vector. For the analysis of the target task, where the task relation model (2.7) is assumed to hold, it is also useful to formulate the expectation of the out-of-sample error of the source task with respect to both $\mathcal{S}$ and the noise vector $\boldsymbol{\eta}$ from the task relation model. This leads us to consider $\boldsymbol{\theta}$ as a random vector and to formulate the following expectation.

$$(A.6) \quad \mathbb{E}_{\mathcal{S}, \boldsymbol{\eta}} [\tilde{\mathcal{E}}_{\text{out}}] = \begin{cases} \frac{\tilde{n}-1}{\tilde{n}-\tilde{p}-1} \left(\left(1 - \frac{\tilde{p}}{d}\right) \kappa + \sigma_\xi^2 \right) & \text{for } \tilde{p} \leq \tilde{n} - 2, \\ \infty & \text{for } \tilde{n} - 1 \leq \tilde{p} \leq \tilde{n} + 1, \\ \frac{\tilde{p}-1}{\tilde{p}-\tilde{n}-1} \left(\left(1 - \frac{\tilde{p}}{d}\right) \kappa + \sigma_\xi^2 \right) + \frac{\tilde{p}-\tilde{n}}{d} \kappa & \text{for } \tilde{p} \geq \tilde{n} + 2. \end{cases}$$

where $\kappa \triangleq \mathbb{E}_{\boldsymbol{\eta}} [\|\boldsymbol{\theta}\|_2^2] = \|\mathbf{H}\boldsymbol{\beta}\|_2^2 + d\sigma_\eta^2$.

A.3. The Estimate $\hat{\beta}$ in Eq. (2.10). The optimization problem in (2.9), given for the target task, can be addressed using the relation $\mathbf{Q}_{\mathcal{F}}^T \mathbf{Q}_{\mathcal{F}} + \mathbf{Q}_{\mathcal{T}}^T \mathbf{Q}_{\mathcal{T}} + \mathbf{Q}_{\mathcal{Z}}^T \mathbf{Q}_{\mathcal{Z}} = \mathbf{I}_d$ and rewritten as

$$\begin{aligned} \hat{\beta} &= \arg \min_{\mathbf{b} \in \mathbb{R}^d} \|\mathbf{y} - \mathbf{X}_{\mathcal{F}} \mathbf{Q}_{\mathcal{F}} \mathbf{b} - \mathbf{X}_{\mathcal{T}} \mathbf{Q}_{\mathcal{T}} \mathbf{b} - \mathbf{X}_{\mathcal{Z}} \mathbf{Q}_{\mathcal{Z}} \mathbf{b}\|_2^2 \\ \text{subject to } &\mathbf{Q}_{\mathcal{T}} \mathbf{b} = \mathbf{Q}_{\mathcal{T}} \hat{\theta} \\ &\mathbf{Q}_{\mathcal{Z}} \mathbf{b} = \mathbf{0} \end{aligned} \quad (\text{A.7})$$

where $\mathbf{X}_{\mathcal{F}} \triangleq \mathbf{X} \mathbf{Q}_{\mathcal{F}}^T$, $\mathbf{X}_{\mathcal{T}} \triangleq \mathbf{X} \mathbf{Q}_{\mathcal{T}}^T$, and $\mathbf{X}_{\mathcal{Z}} \triangleq \mathbf{X} \mathbf{Q}_{\mathcal{Z}}^T$. By setting the equality constraints of (A.7) in its optimization cost, the problem (A.7) can be translated into the form of

$$\begin{aligned} \hat{\beta} &= \arg \min_{\mathbf{b} \in \mathbb{R}^d} \left\| \mathbf{y} - \mathbf{X}_{\mathcal{T}} \mathbf{Q}_{\mathcal{T}} \hat{\theta} - \mathbf{X}_{\mathcal{F}} \mathbf{Q}_{\mathcal{F}} \mathbf{b} \right\|_2^2 \\ \text{subject to } &\mathbf{Q}_{\mathcal{T}} \mathbf{b} = \mathbf{Q}_{\mathcal{T}} \hat{\theta} \\ &\mathbf{Q}_{\mathcal{Z}} \mathbf{b} = \mathbf{0}. \end{aligned} \quad (\text{A.8})$$

The last optimization is a restricted least squares problem that can be solved using the method of Lagrange multipliers to show that

$$\hat{\beta} = \mathbf{Q}_{\mathcal{F}}^T \mathbf{X}_{\mathcal{F}}^+ (\mathbf{y} - \mathbf{X}_{\mathcal{T}} \hat{\theta}_{\mathcal{T}}) + \mathbf{Q}_{\mathcal{T}}^T \hat{\theta}_{\mathcal{T}} \quad (\text{A.9})$$

where $\hat{\theta}_{\mathcal{T}} \triangleq \mathbf{Q}_{\mathcal{T}} \hat{\theta}$ and $\mathbf{X}_{\mathcal{F}}^+$ is the Moore-Penrose pseudoinverse of $\mathbf{X}_{\mathcal{F}}$.

Appendix B. Proofs for Section 3.

In this section we outline the proof of Theorem 3.1 for the generalization error of the target task in the setting where a specific coordinate subset layout $\mathcal{L}$ determines the transferred set of parameters. We start in Section B.1 by providing auxiliary results that use non-asymptotic properties of Gaussian and Wishart matrices. Then, in Section B.2 we prove Theorem 3.1, in Section B.3 we prove Corollary 3.2, and in Section B.4 we prove Corollary 3.4.

B.1. Auxiliary Results using Non-Asymptotic Properties of Gaussian and Wishart Matrices. The random matrix $\mathbf{X}_{\mathcal{F}} \triangleq \mathbf{X} \mathbf{Q}_{\mathcal{F}}^T$ is of size $n \times p$ and all its components are i.i.d. standard Gaussian variables. Then, almost surely,

$$\mathbb{E} [\mathbf{X}_{\mathcal{F}}^+ \mathbf{X}_{\mathcal{F}}] = \mathbf{I}_p \times \begin{cases} 1 & \text{for } p \leq n, \\ \frac{n}{p} & \text{for } p > n, \end{cases} \quad (\text{B.1})$$

where $\mathbf{X}_{\mathcal{F}}^+ \mathbf{X}_{\mathcal{F}}$ is the $p \times p$ projection operator onto the range of $\mathbf{X}_{\mathcal{F}}$. Accordingly, let $\mathbf{a} \in \mathbb{R}^p$ be a random vector independent of the matrix $\mathbf{X}_{\mathcal{F}}$ and, then,

$$\mathbb{E} [\|\mathbf{X}_{\mathcal{F}}^+ \mathbf{X}_{\mathcal{F}} \mathbf{a}\|_2^2] = \mathbb{E} [\|\mathbf{a}\|_2^2] \times \begin{cases} 1 & \text{for } p \leq n, \\ \frac{n}{p} & \text{for } p > n. \end{cases} \quad (\text{B.2})$$

The components of $\mathbf{X}_{\mathcal{F}}$ are i.i.d. standard Gaussian variables, hence $\mathbf{X}_{\mathcal{F}}^T \mathbf{X}_{\mathcal{F}} \sim \mathcal{W}_p(\mathbf{I}_p, n)$ is a $p \times p$ Wishart matrix with n degrees of freedom, and $\mathbf{X}_{\mathcal{F}} \mathbf{X}_{\mathcal{F}}^T \sim \mathcal{W}_n(\mathbf{I}_n, p)$ is a $n \times n$

Wishart matrix with p degrees of freedom. The pseudoinverse of the $n \times n$ Wishart matrix (almost surely) satisfies

$$(B.3) \quad \mathbb{E} \left[(\mathbf{X}_{\mathcal{F}} \mathbf{X}_{\mathcal{F}}^T)^+ \right] = \mathbb{E} \left[\mathbf{X}_{\mathcal{F}}^{+,T} \mathbf{X}_{\mathcal{F}}^+ \right] = \mathbf{I}_n \times \begin{cases} \frac{1}{n-p-1} \cdot \frac{p}{n} & \text{for } p \leq n-2, \\ \infty & \text{for } n-1 \leq p \leq n+1, \\ \frac{1}{p-n-1} & \text{for } p \geq n+2, \end{cases}$$

where the result for $p \geq n+2$ corresponds to the common case of inverse Wishart matrix with more degrees of freedom than its dimension, and the result for $p \leq n-2$ is based on constructions provided in the proof of Theorem 1.3 in [5].

Following (B.3), let $\mathbf{u} \in \mathbb{R}^n$ be a random vector independent of $\mathbf{X}_{\mathcal{F}}$. Then,

$$(B.4) \quad \mathbb{E} \left[\|\mathbf{X}_{\mathcal{F}}^+ \mathbf{u}\|_2^2 \right] = \frac{1}{n} \mathbb{E} \left[\|\mathbf{u}\|_2^2 \right] \times \begin{cases} \frac{p}{n-p-1} & \text{for } p \leq n-2, \\ \infty & \text{for } n-1 \leq p \leq n+1, \\ \frac{n}{p-n-1} & \text{for } p \geq n+2, \end{cases}$$

that specifically for $\mathbf{u} = \mathbf{X}_{\mathcal{F}^c} \boldsymbol{\beta}_{\mathcal{F}^c}$ becomes

$$(B.5) \quad \mathbb{E} \left[\|\mathbf{X}_{\mathcal{F}}^+ \mathbf{X}_{\mathcal{F}^c} \boldsymbol{\beta}_{\mathcal{F}^c}\|_2^2 \right] = \mathbb{E} \left[\|\boldsymbol{\beta}_{\mathcal{F}^c}\|_2^2 \right] \times \begin{cases} \frac{p}{n-p-1} & \text{for } p \leq n-2, \\ \infty & \text{for } n-1 \leq p \leq n+1, \\ \frac{n}{p-n-1} & \text{for } p \geq n+2. \end{cases}$$

The results in (B.1)-(B.5) are presented using notions of the *target* task, specifically, using the data matrix $\mathbf{X}$ and the coordinate subset $\mathcal{T}$. One can obtain the corresponding results for the *source* task by updating (B.1)-(B.5) by replacing $\mathbf{X}$, $\mathcal{T}$, n and p with $\mathbf{Z}$, $\mathcal{S}$, $\tilde{n}$ and $\tilde{p}$, respectively. For example, the result corresponding to (B.1) is

$$(B.6) \quad \mathbb{E} \left[\mathbf{Z}_{\mathcal{S}}^+ \mathbf{Z}_{\mathcal{S}} \right] = \mathbf{I}_{\tilde{p}} \times \begin{cases} 1 & \text{for } \tilde{p} \leq \tilde{n}, \\ \frac{\tilde{n}}{\tilde{p}} & \text{for } \tilde{p} > \tilde{n}, \end{cases}$$

where $\mathbf{Z}_{\mathcal{S}}^+ \mathbf{Z}_{\mathcal{S}}$ is the $\tilde{p} \times \tilde{p}$ projection operator onto the range of $\mathbf{Z}_{\mathcal{S}}$.

The next auxiliary results consider a coordinate subset layout $\mathcal{L} = \{\mathcal{S}, \mathcal{F}, \mathcal{T}, \mathcal{Z}\}$ which is specific, i.e., non random, and therefore the induced operators such as $\mathbf{Q}_{\mathcal{S}}$, $\mathbf{Q}_{\mathcal{F}}$, $\mathbf{Q}_{\mathcal{T}}$, $\mathbf{Q}_{\mathcal{Z}}$ are also fixed and do not have any random aspect. Recall that $\mathbf{Q}_{\mathcal{S}}^T \mathbf{Q}_{\mathcal{S}}$ is a $d \times d$ diagonal matrix with its j^{th} diagonal component equals 1 if $j \in \mathcal{S}$ and 0 otherwise. Similarly holds for the other coordinate subsets. Accordingly, here the norms of vector forms such as $\boldsymbol{\beta}_{\mathcal{T}} \triangleq \mathbf{Q}_{\mathcal{T}} \boldsymbol{\beta}$, $\boldsymbol{\beta}_{\mathcal{F}} \triangleq \mathbf{Q}_{\mathcal{F}} \boldsymbol{\beta}$, and $\boldsymbol{\beta}_{\mathcal{Z}} \triangleq \mathbf{Q}_{\mathcal{Z}} \boldsymbol{\beta}$, are directly referred to as $\|\boldsymbol{\beta}_{\mathcal{T}}\|_2^2$, $\|\boldsymbol{\beta}_{\mathcal{F}}\|_2^2$, $\|\boldsymbol{\beta}_{\mathcal{Z}}\|_2^2$, respectively.

Recall that $\mathcal{T} \subseteq \mathcal{S}$. Then, for a deterministic vector $\mathbf{w} \in \mathbb{R}^d$,

$$(B.7) \quad \mathbb{E} \left[\|\mathbf{Q}_{\mathcal{T}} \mathbf{Q}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}}^+ \mathbf{Z}_{\mathcal{S}} \mathbf{Q}_{\mathcal{S}} \mathbf{w}\|_2^2 \right] = \begin{cases} \|\mathbf{w}_{\mathcal{T}}\|_2^2 & \text{for } \tilde{p} \leq \tilde{n}, \\ \frac{\tilde{n}}{\tilde{p}(\tilde{p}+1)} \left(\left(\tilde{n} + \frac{\tilde{n}-1}{\tilde{p}-1} \right) \|\mathbf{w}_{\mathcal{T}}\|_2^2 + \left(1 - \frac{\tilde{n}-1}{\tilde{p}-1} \right) t \|\mathbf{w}_{\mathcal{S}}\|_2^2 \right) & \text{for } \tilde{p} > \tilde{n}. \end{cases}$$

$$(B.8) \quad \mathbb{E} \left[\left\| \mathbf{Q}_{\mathcal{T}} \mathbf{Q}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}}^+ \mathbf{Z}_{\mathcal{S}^c} \mathbf{Q}_{\mathcal{S}^c} \mathbf{w} \right\|_2^2 \right] = \frac{t}{\tilde{p}} \|\mathbf{w}_{\mathcal{S}^c}\|_2^2 \times \begin{cases} \frac{\tilde{p}}{\tilde{n}-\tilde{p}-1} & \text{for } \tilde{p} \leq \tilde{n} - 2, \\ \infty & \text{for } \tilde{n} - 1 \leq \tilde{p} \leq \tilde{n} + 1, \\ \frac{\tilde{n}}{\tilde{p}-\tilde{n}-1} & \text{for } \tilde{p} \geq \tilde{n} + 2. \end{cases}$$

For two deterministic vectors $\mathbf{w}, \mathbf{a} \in \mathbb{R}^d$,

$$(B.9) \quad \mathbb{E} \left[\langle \mathbf{Q}_{\mathcal{T}} \mathbf{a}, \mathbf{Q}_{\mathcal{T}} \mathbf{Q}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}}^+ \mathbf{Z}_{\mathcal{S}} \mathbf{Q}_{\mathcal{S}^c} \mathbf{w} \rangle \right] = \langle \mathbf{a}_{\mathcal{T}}, \mathbf{w}_{\mathcal{T}} \rangle \times \begin{cases} 1 & \text{for } \tilde{p} \leq \tilde{n}, \\ \frac{\tilde{n}}{\tilde{p}} & \text{for } \tilde{p} > \tilde{n}. \end{cases}$$

For a deterministic vector $\mathbf{r} \in \mathbb{R}^{\tilde{n}}$,

$$(B.10) \quad \mathbb{E} \left[\left\| \mathbf{Q}_{\mathcal{T}} \mathbf{Q}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}}^+ \mathbf{r} \right\|_2^2 \right] = \frac{t}{\tilde{n}\tilde{p}} \|\mathbf{r}\|_2^2 \times \begin{cases} \frac{\tilde{p}}{\tilde{n}-\tilde{p}-1} & \text{for } \tilde{p} \leq \tilde{n} - 2, \\ \infty & \text{for } \tilde{n} - 1 \leq \tilde{p} \leq \tilde{n} + 1, \\ \frac{\tilde{n}}{\tilde{p}-\tilde{n}-1} & \text{for } \tilde{p} \geq \tilde{n} + 2. \end{cases}$$

In our case we have the $\tilde{n} \times \tilde{p}$ matrix $\mathbf{Z}_{\mathcal{S}}$ that its components are i.i.d. standard Gaussian variables, thus, $\mathbf{Z}_{\mathcal{S}}$ can be decomposed into a form that involves an independent Haar-distributed matrix, i.e., a random orthonormal matrix that is uniformly distributed over the set of orthonormal matrices of the relevant size. This lets us to prove the results in (B.7)-(B.10) using some algebra and the non-asymptotic properties of random Haar-distributed matrices, see examples for such properties in Lemma 2.5 in [26] and also in Proposition 1.2 in [12].

B.2. Proof Outline of Theorem 3.1. The generalization error $\mathcal{E}_{\text{out}}$ of the target task was expressed in its basic form in Eq. (2.8) for a specific coordinate subset layout $\mathcal{L} = \{\mathcal{S}, \mathcal{F}, \mathcal{T}, \mathcal{Z}\}$. Please note that the expectations below do not include any expectation with respect to $\mathcal{L}$ or its components, which are non-random here.

We start with the relevant decomposition of the error expression, namely,

$$(B.11) \quad \begin{aligned} \mathcal{E}_{\text{out}} &= \sigma_{\epsilon}^2 + \mathbb{E} \left[\left\| \hat{\boldsymbol{\beta}} - \boldsymbol{\beta} \right\|_2^2 \right] \\ &= \sigma_{\epsilon}^2 + \|\boldsymbol{\beta}_{\mathcal{Z}}\|_2^2 + \mathbb{E} \left[\left\| \mathbf{X}_{\mathcal{F}}^+ (\mathbf{y} - \mathbf{X}_{\mathcal{T}} \hat{\boldsymbol{\theta}}_{\mathcal{T}}) - \boldsymbol{\beta}_{\mathcal{F}} \right\|_2^2 \right] + \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{T}} - \boldsymbol{\beta}_{\mathcal{T}} \right\|_2^2 \right]. \end{aligned}$$

Then, we use the expression for the estimate $\hat{\boldsymbol{\theta}}$ given in (2.4) and the relation $\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\epsilon}$, to

decompose the third term in (B.11) as follows

$$\begin{aligned}
& \mathbb{E} \left[\left\| \mathbf{X}_{\mathcal{F}}^+ (\mathbf{y} - \mathbf{X}_{\mathcal{T}} \hat{\boldsymbol{\theta}}_{\mathcal{T}}) - \boldsymbol{\beta}_{\mathcal{F}} \right\|_2^2 \right] \\
&= \mathbb{E} \left[\left\| \mathbf{X}_{\mathcal{F}}^+ (\mathbf{y} - \mathbf{X}_{\mathcal{T}} \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}]) - \boldsymbol{\beta}_{\mathcal{F}} \right\|_2^2 \right] + \mathbb{E} \left[\left\| \mathbf{X}_{\mathcal{F}}^+ \mathbf{X}_{\mathcal{T}} (\hat{\boldsymbol{\theta}}_{\mathcal{T}} - \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}]) \right\|_2^2 \right] \\
&= \mathbb{E} \left[\left\| \mathbf{X}_{\mathcal{F}}^+ (\mathbf{X}_{\mathcal{T}} \boldsymbol{\beta}_{\mathcal{T}} + \mathbf{X}_{\mathcal{T}^c} \boldsymbol{\beta}_{\mathcal{T}^c} + \boldsymbol{\epsilon} - \mathbf{X}_{\mathcal{T}} \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}]) - \boldsymbol{\beta}_{\mathcal{F}} \right\|_2^2 \right] + \mathbb{E} \left[\left\| \mathbf{X}_{\mathcal{F}}^+ \mathbf{X}_{\mathcal{T}} (\hat{\boldsymbol{\theta}}_{\mathcal{T}} - \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}]) \right\|_2^2 \right] \\
&= \mathbb{E} \left[\left\| \mathbf{X}_{\mathcal{F}}^+ (\mathbf{X}_{\mathcal{T}^c} \boldsymbol{\beta}_{\mathcal{T}^c} + \boldsymbol{\epsilon}) - \boldsymbol{\beta}_{\mathcal{F}} \right\|_2^2 \right] + \mathbb{E} \left[\left\| \mathbf{X}_{\mathcal{F}}^+ \mathbf{X}_{\mathcal{T}} (\boldsymbol{\beta}_{\mathcal{T}} - \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}]) \right\|_2^2 \right] \\
&\quad + \mathbb{E} \left[\left\| \mathbf{X}_{\mathcal{F}}^+ \mathbf{X}_{\mathcal{T}} (\hat{\boldsymbol{\theta}}_{\mathcal{T}} - \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}]) \right\|_2^2 \right]
\end{aligned}
\tag{B.12}$$

We further develop the last expression using the result in (B.4) and that $\mathbf{X}_{\mathcal{T}}^T \mathbf{X}_{\mathcal{T}} \sim \mathcal{W}_t(\mathbf{I}_t, n)$ is a Wishart matrix with mean $\mathbb{E} [\mathbf{X}_{\mathcal{T}}^T \mathbf{X}_{\mathcal{T}}] = n \mathbf{I}_t$, and get

$$\begin{aligned}
& \mathbb{E} \left[\left\| \mathbf{X}_{\mathcal{F}}^+ (\mathbf{y} - \mathbf{X}_{\mathcal{T}} \hat{\boldsymbol{\theta}}_{\mathcal{T}}) - \boldsymbol{\beta}_{\mathcal{F}} \right\|_2^2 \right] \\
&= \mathbb{E} \left[\left\| \mathbf{X}_{\mathcal{F}}^+ (\mathbf{X}_{\mathcal{T}^c} \boldsymbol{\beta}_{\mathcal{T}^c} + \boldsymbol{\epsilon}) - \boldsymbol{\beta}_{\mathcal{F}} \right\|_2^2 \right] \\
&\quad + \left(\left\| \boldsymbol{\beta}_{\mathcal{T}} - \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}] \right\|_2^2 + \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{T}} - \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}] \right\|_2^2 \right] \right) \times \begin{cases} \frac{p}{n-p-1} & \text{for } p \leq n-2, \\ \infty & \text{for } n-1 \leq p \leq n+1, \\ \frac{n}{p-n-1} & \text{for } p \geq n+2. \end{cases}
\end{aligned}
\tag{B.13}$$

We proceed to the fourth error term in (B.11), i.e., the error in the subvector induced by the specific $\mathcal{T}$ of interest, and develop its formulation as follows:

$$\begin{aligned}
& \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{T}} - \boldsymbol{\beta}_{\mathcal{T}} \right\|_2^2 \right] = \\
&= \left\| \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}] - \boldsymbol{\beta}_{\mathcal{T}} \right\|_2^2 + \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{T}} - \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}] \right\|_2^2 \right] + 2 \mathbb{E} \left[\left(\hat{\boldsymbol{\theta}}_{\mathcal{T}} - \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}] \right)^T \left(\mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}] - \boldsymbol{\beta}_{\mathcal{T}} \right) \right] \\
&= \left\| \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}] - \boldsymbol{\beta}_{\mathcal{T}} \right\|_2^2 + \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{T}} - \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}] \right\|_2^2 \right]
\end{aligned}
\tag{B.14}$$

Then, setting (B.13) and (B.14) in (B.11) gives

$$\begin{aligned} \mathcal{E}_{\text{out}} &= \sigma_\epsilon^2 + \|\beta_{\mathcal{Z}}\|_2^2 + \mathbb{E} \left[\|\mathbf{X}_{\mathcal{F}}^+ (\mathbf{X}_{\mathcal{T}^c} \beta_{\mathcal{T}^c} + \epsilon) - \beta_{\mathcal{F}}\|_2^2 \right] \\ &+ \left(\left\| \mathbb{E} [\hat{\theta}_{\mathcal{T}}] - \beta_{\mathcal{T}} \right\|_2^2 + \mathbb{E} \left[\left\| \hat{\theta}_{\mathcal{T}} - \mathbb{E} [\hat{\theta}_{\mathcal{T}}] \right\|_2^2 \right] \right) \left(1 + \begin{cases} \frac{p}{n-p-1} & \text{for } p \leq n-2, \\ \infty & \text{for } n-1 \leq p \leq n+1, \\ \frac{n}{p-n-1} & \text{for } p \geq n+2. \end{cases} \right). \end{aligned} \quad (\text{B.15})$$

Also,

$$\mathbb{E} \left[\|\mathbf{X}_{\mathcal{F}}^+ (\mathbf{X}_{\mathcal{T}^c} \beta_{\mathcal{T}^c} + \epsilon) - \beta_{\mathcal{F}}\|_2^2 \right] = \mathbb{E} \left[\|\mathbf{X}_{\mathcal{F}}^+ \mathbf{X}_{\mathcal{F}} \beta_{\mathcal{F}} - \beta_{\mathcal{F}}\|_2^2 \right] + \mathbb{E} \left[\|\mathbf{X}_{\mathcal{F}}^+ (\mathbf{X}_{\mathcal{Z}} \beta_{\mathcal{Z}} + \epsilon)\|_2^2 \right] \quad (\text{B.16})$$

where the first term can be developed using (B.2) into

$$\mathbb{E} \left[\|\mathbf{X}_{\mathcal{F}}^+ \mathbf{X}_{\mathcal{F}} \beta_{\mathcal{F}} - \beta_{\mathcal{F}}\|_2^2 \right] = \|\beta_{\mathcal{F}}\|_2^2 \times \begin{cases} 0 & \text{for } p \leq n, \\ 1 - \frac{n}{p} & \text{for } p > n, \end{cases} \quad (\text{B.17})$$

and the second term in (B.16) can be rewritten using the result in (B.4) and that $\mathbf{X}_{\mathcal{Z}}^T \mathbf{X}_{\mathcal{Z}} \sim \mathcal{W}_{d-p-t}(\mathbf{I}_{d-p-t}, n)$ is a Wishart matrix with mean $\mathbb{E} [\mathbf{X}_{\mathcal{Z}}^T \mathbf{X}_{\mathcal{Z}}] = n \mathbf{I}_{d-p-t}$:

$$\mathbb{E} \left[\|\mathbf{X}_{\mathcal{F}}^+ (\mathbf{X}_{\mathcal{Z}} \beta_{\mathcal{Z}} + \epsilon)\|_2^2 \right] = \left(\|\beta_{\mathcal{Z}}\|_2^2 + \sigma_\epsilon^2 \right) \times \begin{cases} \frac{p}{n-p-1} & \text{for } p \leq n-2, \\ \infty & \text{for } n-1 \leq p \leq n+1, \\ \frac{n}{p-n-1} & \text{for } p \geq n+2. \end{cases} \quad (\text{B.18})$$

Hence, (B.16)-(B.18) let us write (B.15) in the form which is provided in Theorem 3.1.

B.3. Proof of Corollary 3.2. To prove Corollary 3.2 we first formulate the expectation of the transferred parameters:

$$\begin{aligned} \mathbb{E} [\hat{\theta}_{\mathcal{T}}] &= \mathbb{E} [\mathbf{Q}_{\mathcal{T}} \mathbf{Q}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}}^+ \mathbf{Z} \mathbf{H} \beta] = \mathbb{E} [\mathbf{Q}_{\mathcal{T}} \mathbf{Q}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}}^+ \mathbf{Z}_{\mathcal{S}} \mathbf{Q}_{\mathcal{S}} \mathbf{H} \beta] \\ &= \mathbf{Q}_{\mathcal{T}} \mathbf{H} \beta \times \begin{cases} 1 & \text{for } \tilde{p} \leq \tilde{n}, \\ \frac{\tilde{n}}{\tilde{p}} & \text{for } \tilde{p} > \tilde{n}, \end{cases} \end{aligned} \quad (\text{B.19})$$

where the last equality stems from (B.6). Consequently, we use (B.19) to formulate the transfer bias term from Theorem 3.1 as

$$\text{Bias}_{\mathcal{T}}^2 = \left\| \mathbb{E} [\hat{\theta}_{\mathcal{T}}] - \beta_{\mathcal{T}} \right\|_2^2 = \|\mathbf{Q}_{\mathcal{T}} (r \mathbf{H} - \mathbf{I}_d) \beta\|_2^2 \quad \text{where } r \triangleq \begin{cases} 1 & \text{for } \tilde{p} \leq \tilde{n}, \\ \frac{\tilde{n}}{\tilde{p}} & \text{for } \tilde{p} > \tilde{n}, \end{cases} \quad (\text{B.20})$$

which corresponds to the formulation in Corollary 3.2.

Now we turn to prove the transfer variance formulation from Corollary 3.2. For a start, note that

$$\begin{aligned}
 \text{Var}_{\mathcal{T}, \mathcal{S}} &= \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{T}} - \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}] \right\|_2^2 \right] = \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{T}} \right\|_2^2 \right] - \left\| \mathbb{E} [\hat{\boldsymbol{\theta}}_{\mathcal{T}}] \right\|_2^2 \\
 (B.21) \quad &= \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{T}} \right\|_2^2 \right] - \|\mathbf{Q}_{\mathcal{T}} \mathbf{H} \boldsymbol{\beta}\|_2^2 \times \begin{cases} 1 & \text{for } \tilde{p} \leq \tilde{n}, \\ \left(\frac{\tilde{n}}{\tilde{p}}\right)^2 & \text{for } \tilde{p} > \tilde{n}. \end{cases}
 \end{aligned}$$

To further develop the last expression, we will now explicitly formulate $\mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{T}} \right\|_2^2 \right]$. Using the auxiliary notation $\boldsymbol{\beta}^{(\mathbf{H})} \triangleq \mathbf{H} \boldsymbol{\beta}$ we get

$$\begin{aligned}
 (B.22) \quad \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{T}} \right\|_2^2 \right] &= \mathbb{E} \left[\left\| \mathbf{Q}_{\mathcal{T}} \mathbf{Q}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}}^+ (\mathbf{Z} \mathbf{H} \boldsymbol{\beta} + \mathbf{Z} \boldsymbol{\eta} + \boldsymbol{\xi}) \right\|_2^2 \right] \\
 &= \mathbb{E} \left[\left\| \mathbf{Q}_{\mathcal{T}} \mathbf{Q}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}}^+ \mathbf{Z}_{\mathcal{S}} (\boldsymbol{\beta}_{\mathcal{S}}^{(\mathbf{H})} + \boldsymbol{\eta}_{\mathcal{S}}) \right\|_2^2 \right] \\
 &\quad + \mathbb{E} \left[\left\| \mathbf{Q}_{\mathcal{T}} \mathbf{Q}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}}^+ \mathbf{Z}_{\mathcal{S}^c} (\boldsymbol{\beta}_{\mathcal{S}^c}^{(\mathbf{H})} + \boldsymbol{\eta}_{\mathcal{S}^c}) \right\|_2^2 \right] + \mathbb{E} \left[\left\| \mathbf{Q}_{\mathcal{T}} \mathbf{Q}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}}^+ \boldsymbol{\xi} \right\|_2^2 \right]
 \end{aligned}$$

that using (B.7)-(B.8), (B.10) leads to

$$\begin{aligned}
 (B.23) \quad \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{T}} \right\|_2^2 \right] &= \\
 &= \left(\begin{cases} \left\| \boldsymbol{\beta}_{\mathcal{T}}^{(\mathbf{H})} \right\|_2^2 + t \sigma_{\eta}^2 & \text{for } \tilde{p} \leq \tilde{n}, \\ \frac{\tilde{n}}{\tilde{p}(\tilde{p}+1)} \left(\left(\tilde{n} + \frac{\tilde{n}-1}{\tilde{p}-1} \right) \left(\left\| \boldsymbol{\beta}_{\mathcal{T}}^{(\mathbf{H})} \right\|_2^2 + t \sigma_{\eta}^2 \right) + \left(1 - \frac{\tilde{n}-1}{\tilde{p}-1} \right) t \left(\left\| \boldsymbol{\beta}_{\mathcal{S}}^{(\mathbf{H})} \right\|_2^2 + \tilde{p} \sigma_{\eta}^2 \right) \right) & \text{for } \tilde{p} > \tilde{n}, \end{cases} \right) \\
 &\quad + \frac{t}{\tilde{p}} \left(\left\| \boldsymbol{\beta}_{\mathcal{S}^c}^{(\mathbf{H})} \right\|_2^2 + (d - \tilde{p}) \sigma_{\eta}^2 + \sigma_{\xi}^2 \right) \times \left(\begin{cases} \frac{\tilde{p}}{\tilde{n} - \tilde{p} - 1} & \text{for } \tilde{p} \leq \tilde{n} - 2, \\ \infty & \text{for } \tilde{n} - 1 \leq \tilde{p} \leq \tilde{n} + 1, \\ \frac{\tilde{n}}{\tilde{p} - \tilde{n} - 1} & \text{for } \tilde{p} \geq \tilde{n} + 2, \end{cases} \right) \\
 &= \begin{cases} \zeta_{\mathcal{T}} + t \sigma_{\eta}^2 + t \cdot \frac{\zeta_{\mathcal{S}^c} + (d - \tilde{p}) \sigma_{\eta}^2 + \sigma_{\xi}^2}{\tilde{n} - \tilde{p} - 1} & \text{for } \tilde{p} \leq \tilde{n} - 2, \\ \infty & \text{for } \tilde{n} - 1 \leq \tilde{p} \leq \tilde{n} + 1, \\ \frac{\tilde{n}}{\tilde{p}} \left(\frac{(\tilde{p} - \tilde{n}) t (\zeta_{\mathcal{S}} + \tilde{p} \sigma_{\eta}^2) + (\tilde{n} \tilde{p} - 1) (\zeta_{\mathcal{T}} + t \sigma_{\eta}^2)}{\tilde{p}^2 - 1} + t \cdot \frac{\zeta_{\mathcal{S}^c} + (d - \tilde{p}) \sigma_{\eta}^2 + \sigma_{\xi}^2}{\tilde{p} - \tilde{n} - 1} \right) & \text{for } \tilde{p} \geq \tilde{n} + 2. \end{cases}
 \end{aligned}$$

where $\zeta_{\mathcal{T}} \triangleq \|\mathbf{Q}_{\mathcal{T}} \mathbf{H} \boldsymbol{\beta}\|_2^2$, $\zeta_{\mathcal{S}^c} \triangleq \|\mathbf{Q}_{\mathcal{S}^c} \mathbf{H} \boldsymbol{\beta}\|_2^2$ and $\zeta_{\mathcal{S}} \triangleq \|\mathbf{Q}_{\mathcal{S}} \mathbf{H} \boldsymbol{\beta}\|_2^2$. Then, we set (B.23) in (B.21) and using some algebra obtain the transfer variance formulation in Corollary 3.2.

B.4. Proof of Corollary 3.4. Corollary 3.4 formulates the generalization error of the target task under the expectation over a coordinate layout $\mathcal{L}$ which is chosen uniformly at

random. Recall Definition 3.3 in the main text that characterizes a coordinate subset layout $\mathcal{L} = \{\mathcal{S}, \mathcal{F}, \mathcal{T}, \mathcal{Z}\}$ that is $\{\tilde{p}, p, t\}$ -uniformly distributed, for $\tilde{p} \in \{1, \dots, d\}$ and $(p, t) \in \{0, \dots, d\} \times \{0, \dots, \tilde{p}\}$ such that $p + t \leq d$. Here we provide several auxiliary results that are induced by this random structure and utilized in the proof of Corollary 3.4.

For $\mathcal{S}$ that is uniformly chosen at random from all the subsets of $\tilde{p}$ unique coordinates of $\{1, \dots, d\}$, we get that the mean of the projection operator $\mathbf{Q}_S^T \mathbf{Q}_S$ is

$$(B.24) \quad \mathbb{E}_{\mathcal{L}} [\mathbf{Q}_S^T \mathbf{Q}_S] = \mathbb{E}_{\mathcal{S}} [\mathbf{Q}_S^T \mathbf{Q}_S] = \frac{\binom{d-1}{\tilde{p}-1}}{\binom{d}{\tilde{p}}} \mathbf{I}_d = \frac{\tilde{p}}{d} \mathbf{I}_d$$

where we used the structure of $\mathbf{Q}_S^T \mathbf{Q}_S$ that is a $d \times d$ diagonal matrix with its j^{th} diagonal component equals 1 if $j \in \mathcal{S}$ and 0 otherwise.

Definition 3.3 also specifies that, given $\mathcal{S}$, the target-task coordinate layout $\{\mathcal{F}, \mathcal{T}, \mathcal{Z}\}$ is uniformly chosen at random from all the layouts where $\mathcal{F}$, $\mathcal{T}$, and $\mathcal{Z}$ are three disjoint sets of coordinates that satisfy $\mathcal{F} \cup \mathcal{T} \cup \mathcal{Z} = \{1, \dots, d\}$ such that $|\mathcal{F}| = p$, $|\mathcal{T}| = t$ and $\mathcal{T} \subseteq \mathcal{S}$, and $|\mathcal{Z}| = d - p - t$. Accordingly,

$$(B.25) \quad \begin{aligned} \mathbb{E}_{\mathcal{L}} [\mathbf{Q}_{\mathcal{T}}^T \mathbf{Q}_{\mathcal{T}}] &= \mathbb{E}_{\mathcal{S}} [\mathbb{E}_{\mathcal{L}|\mathcal{S}} [\mathbf{Q}_{\mathcal{T}}^T \mathbf{Q}_{\mathcal{T}}]] \\ &= \frac{\binom{\tilde{p}-1}{t-1}}{\binom{\tilde{p}}{t}} \mathbb{E}_{\mathcal{S}} [\mathbf{Q}_S^T \mathbf{Q}_S] \\ &= \frac{t}{\tilde{p}} \mathbb{E}_{\mathcal{S}} [\mathbf{Q}_S^T \mathbf{Q}_S] \\ &= \frac{t}{d} \mathbf{I}_d, \end{aligned}$$

and similarly

$$(B.26) \quad \mathbb{E}_{\mathcal{L}} [\mathbf{Q}_{\mathcal{F}}^T \mathbf{Q}_{\mathcal{F}}] = \frac{p}{d} \mathbf{I}_d,$$

$$(B.27) \quad \mathbb{E}_{\mathcal{L}} [\mathbf{Q}_{\mathcal{Z}}^T \mathbf{Q}_{\mathcal{Z}}] = \frac{d - p - t}{d} \mathbf{I}_d.$$

Another useful auxiliary result, based on the relation $\mathbf{Q}_S \mathbf{Q}_S^T = \mathbf{I}_{\tilde{p}}$ (carefully note the transpose appearance), is provided by

$$(B.28) \quad \begin{aligned} \mathbb{E}_{\mathcal{L}} [\mathbf{Q}_S \mathbf{Q}_{\mathcal{T}}^T \mathbf{Q}_{\mathcal{T}} \mathbf{Q}_S^T] &= \mathbb{E}_{\mathcal{S}} [\mathbf{Q}_S \mathbb{E}_{\mathcal{L}|\mathcal{S}} [\mathbf{Q}_{\mathcal{T}}^T \mathbf{Q}_{\mathcal{T}}] \mathbf{Q}_S^T] \\ &= \frac{t}{\tilde{p}} \mathbb{E}_{\mathcal{S}} [\mathbf{Q}_S \mathbf{Q}_S^T \mathbf{Q}_S \mathbf{Q}_S^T] \\ &= \frac{t}{\tilde{p}} \mathbf{I}_{\tilde{p}}. \end{aligned}$$

The results in (B.25)–(B.27) imply that

$$(B.29) \quad \mathbb{E}_{\mathcal{L}} [\|\beta_{\mathcal{T}}\|_2^2] = \beta^T \mathbb{E}_{\mathcal{L}} [\mathbf{Q}_{\mathcal{T}}^T \mathbf{Q}_{\mathcal{T}}] \beta = \frac{t}{d} \|\beta\|_2^2,$$

$$(B.30) \quad \mathbb{E}_{\mathcal{L}} [\|\beta_{\mathcal{F}}\|_2^2] = \beta^T \mathbb{E}_{\mathcal{L}} [\mathbf{Q}_{\mathcal{F}}^T \mathbf{Q}_{\mathcal{F}}] \beta = \frac{p}{d} \|\beta\|_2^2,$$

$$(B.31) \quad \mathbb{E}_{\mathcal{L}} [\|\beta_{\mathcal{Z}}\|_2^2] = \beta^T \mathbb{E}_{\mathcal{L}} [\mathbf{Q}_{\mathcal{Z}}^T \mathbf{Q}_{\mathcal{Z}}] \beta = \frac{d - p - t}{d} \|\beta\|_2^2,$$

where $\beta_{\mathcal{T}} \triangleq \mathbf{Q}_{\mathcal{T}}\beta$, $\beta_{\mathcal{F}} \triangleq \mathbf{Q}_{\mathcal{F}}\beta$, and $\beta_{\mathcal{Z}} \triangleq \mathbf{Q}_{\mathcal{Z}}\beta$. Note that the expressions in (B.29)-(B.31) hold also for d -dimensional deterministic vectors other than β , e.g., (B.29)-(B.31) hold for $\beta^{(\mathbf{H})} \triangleq \mathbf{H}\beta$.

Then, the auxiliary results in (B.24)-(B.31) can be utilized to formulate the expectation over $\mathcal{L}$ of the analytical results in Theorem 3.1 and, by that, proving Corollary 3.4.

Appendix C. Additional Results for Section 3.

C.1. Additional Results for the On-Average Analysis in Section 3.2. In Fig. 11 we present the empirically computed values of the out-of-sample squared error of the target task, $\mathbb{E}_{\mathcal{L}}[\mathcal{E}_{\text{out}}]$, with respect to the number of free parameters $\tilde{p}$ and p (in the source and target tasks, respectively). The empirical values in Fig. 11 (and also the values denoted by circle markers in Fig. 1 in Section 3.2) were obtained by averaging over 250 experiments where each experiment was carried out based on new realizations of the data matrices, noise components, and the sequential order of adding coordinates to subsets (such as $\mathcal{S}$) for the gradual increase of $\tilde{p}$ and p within each experiment. Note that the results in Fig. 4 do not include averaging over the sequential order of adding coordinates to subsets. Each single evaluation of the expectation of the squared error for an out-of-sample data pair $(\mathbf{x}^{(\text{test})}, y^{(\text{test})})$ was empirically carried out by averaging over a set of 1000 out-of-sample realizations of data pairs. Here $d = 120$, $\tilde{n} = 50$, $n = 20$, $\|\beta\|_2^2 = d$, $\sigma_{\epsilon}^2 = 0.05 \cdot d$, $\sigma_{\xi}^2 = 0.025 \cdot d$. The deterministic $\beta \in \mathbb{R}^d$ used in the experiments satisfies $\|\beta\|_2^2 = d$.

One can observe the excellent match between the empirical results in Fig. 11 and the analytical results provided in Fig. 10. This further establishes the formulations given in Corollary 3.4.

C.2. Additional Results for the Single-Layout Analysis in Section 3.3. The following results are for two different forms of the true solution β : the first is a form with linearly increasing values (Fig. 4a), the second is a form with sparse values where only 25% of coordinates have non-zero value (Fig. 4e). Note that both forms satisfy $\|\beta\|_2^2 = d$.

The three types of linear operator $\mathbf{H}$ in the evaluations in Section 3.3 are as follows. First, $\mathbf{H} = \mathbf{I}_d$ that is the identity operator. Second, is the circulant matrix $\mathbf{H}$ that corresponds to a shift-invariant local averaging operator that uniformly considers 11-coordinates neighborhood around the computed coordinate (note that in other parts of this paper we consider also averaging operators with neighborhood sizes other than 11). Third, is the circulant matrix $\mathbf{H}$ that corresponds to discrete derivative operator based on the convolution kernel $[-0.5, 0.5]$.

Figures 13-14 present the analytical and empirical values of the generalization error of the target task with respect to specific coordinate layouts $\mathcal{L}$ that evolve with respect to the value of p (this evolution of $\mathcal{L}$ is the same in each of the subfigures and it is not particularly designed to any of the combinations of the true β , $\mathbf{H}$, and σ_{η}^2). It is clear from Figures 13-14 that the increase in σ_{η}^2 , which by its definition corresponds to less related source and target tasks, reduces the benefits or even increases the harm due to transfer of parameters (one can observe that in Figs. 13-14 by comparing the error curves among subfigures in the same row).

The effect of $\mathbf{H}$ with respect to the true β is also evident. First, the identity operator $\mathbf{H} = \mathbf{I}_d$ does not reduce the relation between the source and target tasks and therefore does not degrade the parameter transfer performance by itself (i.e., for $\mathbf{H} = \mathbf{I}_d$, only the

additive noise level σ_η^2 can reduce the relation between the tasks). Second, when $\mathbf{H}$ is a local averaging operator it does not reduce the benefits from transfer learning (e.g., compare second to first row of subfigures in Figs. 13-14) in the case of linearly-increasing β shape (because local averaging does not affect a linear function, except to the few first and last coordinates where the periodic averaging is applied), in contrast, the local averaging operator significantly degrades the parameter transfer performance in the case of the sparse β form. Lastly, when $\mathbf{H}$ is a discrete derivative operator it renders transfer learning harmful in the case of linearly-increasing β shape (e.g., compare third to first row of subfigures in Fig. 13). In the case of the sparse β form the discrete derivative reduces the potential benefits of the parameter transfer but does not eliminate them completely in case these benefits exist for $\mathbf{H} = \mathbf{I}_d$ (e.g., compare third to first row of subfigures in Fig. 14).

Appendix D. Additional Empirical Results on Parameter Transfer Usefulness.

D.1. Details on the Empirical Evaluation of $\Delta\mathcal{E}_{\text{transfer}}$. The analytical formula for $\Delta\mathcal{E}_{\text{transfer}} \triangleq \mathbb{E}_{\mathcal{L}} \left[\Delta\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})} \right]$, which is based on (4.3) and Corollary 3.4, measures the (normalized) expected difference in the generalization error (of the target task) due to transferring parameters instead of setting them to zero. Accordingly, the empirical evaluation of $\Delta\mathcal{E}_{\text{transfer}}$ for a given $\tilde{p}$ can be computed by

$$(D.1) \quad \hat{\Delta}\mathcal{E}_{\text{transfer}} = \frac{1}{d-3} \sum_{p=1, \dots, n-2, n+2, \dots, d} \frac{\hat{\mathbb{E}}_{\mathcal{L}} \left\{ \mathcal{E}_{\text{out}}^{(\tilde{p}, p, t=m)} \right\} - \hat{\mathbb{E}}_{\mathcal{L}} \left\{ \mathcal{E}_{\text{out}}^{(\tilde{p}, p, t=0)} \right\}}{m \cdot \alpha(p)}$$

where

$$(D.2) \quad \alpha(p) \triangleq \begin{cases} 1 + \frac{p}{n-p-1} & \text{for } p \leq n-2, \\ 1 + \frac{n}{p-n-1} & \text{for } p \geq n+2 \end{cases}$$

is a normalization factor required for independence from p . The value measured in (D.1) is also normalized by the number of transferred parameters. Here $\hat{\mathbb{E}}_{\mathcal{L}} \left\{ \mathcal{E}_{\text{out}}^{(\tilde{p}, p, t=m)} \right\}$ is the out-of-sample error of the target task that is *empirically* computed for m transferred parameters, p free parameters in the target task, and $\tilde{p}$ free parameters in the source task. Correspondingly, $\hat{\mathbb{E}}_{\mathcal{L}} \left\{ \mathcal{E}_{\text{out}}^{(\tilde{p}, p, t=0)} \right\}$ is the empirically computed error induced by avoiding parameter transfer. Therefore, the formula in (D.1) empirically measures the average error difference for a single transferred parameter by averaging over the various settings induced by different values of p while $\tilde{p}$ is kept fixed. To obtain a good numerical accuracy with averaging over a moderate number of experiments we use the value $m = 5$.

Each empirical evaluation of $\hat{\mathbb{E}}_{\mathcal{L}} \left\{ \mathcal{E}_{\text{out}}^{(\tilde{p}, p, t)} \right\}$, for a specific set of values $\tilde{p}, p, t$ corresponds to averaging over 500 experiments where each experiment was conducted for new realizations of the data matrices, noise components, and the sequential order of adding coordinates to subsets. Each single evaluation of the expectation of the squared error for an out-of-sample data pair $(\mathbf{x}^{(\text{test})}, y^{(\text{test})})$ was empirically computed by averaging over 1000 out-of-sample realizations of data pairs.

D.2. Results for $\tilde{n} < d$. In Figures 16-17 we present analytical and empirical values of $\frac{1}{t}\Delta\mathcal{E}_{\text{transfer}}$ induced by settings where $\tilde{n} < d$ (specifically, $\tilde{n} = 50$ and $d = 80$), which naturally enable the corresponding overparameterized (i.e., $\tilde{n} < \tilde{p} < d$) and underparameterized (i.e., $\tilde{p} < \tilde{n} < d$) settings of the source task. In Fig. 16 we provide the analytical and empirical results for cases where $\mathbf{H}$ is local averaging and discrete derivative operators. In the main text only the analytical results were provided and here we show them again near their empirical counterparts that excellently match them (up to the resolution of the empirical settings).

In Figure 17 we provide additional results for cases where the operator $\mathbf{H}$ is a scaled identity matrix.

D.3. Results for $\tilde{n} > d$. Here we provide in Fig. 18 the analytical and empirical evaluations of $\Delta\mathcal{E}_{\text{transfer}}$ that correspond to settings where $\tilde{n} > d$, we specifically consider $\tilde{n} = 150$ and $d = 80$. Note that $\tilde{n} > d$ implies that, by the definition of $\tilde{p}$, the corresponding settings (of the source task) are underparameterized with $\tilde{p} \leq d < \tilde{n}$. Like in Fig. 17, the results in Fig. 18 show the excellent match between the analytical and empirical results.

D.4. Empirical Results for Benefits in Transferred versus Free Parameters. We provide in Figure 19 the empirical evaluation of $\Delta\mathcal{E}_{\text{TVsF}}$ that corresponds to the analytical evaluation in Figure 6 in the main text. The empirical evaluation was conducted by averaging over 750 experiments where, in each of them, $\Delta\mathcal{E}_{\text{TVsF}}$ was evaluated based on its definition in (4.4) and using a different realization of $\mathcal{L}$ (from the uniform distribution that we use) and its corresponding $\mathcal{L}''$.

Appendix E. The Optimal Componentwise Task Relation: Additional Analytical and Empirical Details.

E.1. Proof of Theorem 5.1. The proof outline for Theorem 5.1, which characterizes the optimal $\mathbf{H}$ in a componentwise task relation, is as follows. Recall that in this theorem we consider $\mathbf{H} = \text{diag}\{\lambda_{\mathbf{H}}^{(1)}, \dots, \lambda_{\mathbf{H}}^{(d)}\}$. The proof starts by setting the expressions from (5.3)-(5.4) in the formulations of the transfer bias and variance from Corollary 3.2, then we can also reformulate the expression of $\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})}$ from Theorem 3.1.

In the case of an underparameterized source task where $1 \leq \tilde{p} \leq \tilde{n} - 2$, we get (E.1)

$$\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})} = t\sigma_{\eta}^2 + \sum_{j \in \mathcal{T}} \left(\lambda_{\mathbf{H}}^{(j)} - 1 \right)^2 \left(\beta^{(j)} \right)^2 + \frac{t}{\tilde{n} - \tilde{p} - 1} \left((d - \tilde{p}) \sigma_{\eta}^2 + \sum_{j \in \mathcal{S}^c} \left(\lambda_{\mathbf{H}}^{(j)} \beta^{(j)} \right)^2 + \sigma_{\xi}^2 \right)$$

Recall that $\mathcal{T} \subseteq \mathcal{S}$ and hence $\mathcal{T} \cap \mathcal{S}^c = \emptyset$. Accordingly, the two sums in (E.1) include distinct sets of coordinates, which simplify the derivative of $\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})}$ with respect to $\lambda_{\mathbf{H}}^{(k)}$ for a particular k . Hence, in this case of an underparameterized source task and $\beta^{(k)} \neq 0$, the condition $\frac{\partial \mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})}}{\partial \lambda_{\mathbf{H}}^{(k)}} = 0$ is satisfied by $\lambda_{\mathbf{H}}^{(k)} = 1$ for $k \in \mathcal{T}$ and by $\lambda_{\mathbf{H}}^{(k)} = 0$ for $k \in \mathcal{S}^c$, which minimize $\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})}$ due to convexity. Note that the eigenvalues $\{\lambda_{\mathbf{H}}^{(k)}\}_{k \in \mathcal{S} \setminus \mathcal{T}}$ do not appear in (E.1) and hence they can have any value without affecting the minimization of $\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})}$ for $\tilde{p} \leq \tilde{n} - 2$.

In the case of an overparameterized source task where $\tilde{p} \geq \tilde{n} + 2$, we get

$$(E.2) \quad \begin{aligned} \mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})} = & \frac{\tilde{n}}{\tilde{p}} \left(t\sigma_\eta^2 + \frac{\tilde{p} - \tilde{n}}{\tilde{p}^2 - 1} t \sum_{j \in \mathcal{S} \setminus \mathcal{T}} \left(\lambda_{\mathbf{H}}^{(j)} \beta^{(j)} \right)^2 \right. \\ & + \sum_{j \in \mathcal{T}} \left(\frac{(\tilde{p} - \tilde{n})t + \tilde{n}\tilde{p} - 1}{\tilde{p}^2 - 1} \lambda_{\mathbf{H}}^{(j)} - 2 \right) \lambda_{\mathbf{H}}^{(j)} \left(\beta^{(j)} \right)^2 \\ & \left. + \frac{t}{\tilde{p} - \tilde{n} - 1} \left((d - \tilde{p}) \sigma_\eta^2 + \sum_{j \in \mathcal{S}^c} \left(\lambda_{\mathbf{H}}^{(j)} \beta^{(j)} \right)^2 + \sigma_\xi^2 \right) \right) \end{aligned}$$

where the three sums refer to disjoint sets of coordinates. Accordingly, in this case of an overparameterized source task and $\beta^{(k)} \neq 0$, the condition $\frac{\partial \mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})}}{\partial \lambda_{\mathbf{H}}^{(k)}} = 0$ is satisfied by $\lambda_{\mathbf{H}}^{(k)} = \frac{\tilde{p}^2 - 1}{\tilde{n}\tilde{p} - 1 + t(\tilde{p} - \tilde{n})}$ for $k \in \mathcal{T}$ and by $\lambda_{\mathbf{H}}^{(k)} = 0$ for $k \in \mathcal{T}^c$, which minimize $\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})}$ due to convexity.

In both cases (E.1) and (E.2), if $k \in \{1, \dots, d\}$ corresponds to $\beta^{(k)} = 0$ then $\lambda_{\mathbf{H}}^{(k)}$ may have any value without affecting the minimization of $\mathcal{E}_{\text{transfer}}^{(\mathcal{T}, \mathcal{S})}$. By this we complete the proof outline for Theorem 5.1.

E.2. Empirical Results for The Optimal $\mathbf{H}$ in a Componentwise Task Relation.

In Figure 20 we provide the empirical evaluations that correspond to the analytical results in Figure 7. The empirical values were computed by averaging over 500 experiments.

Appendix F. Additional Details and Results for Section 6.

F.1. Ridge Regression: Error Expression and Optimal Tuning for the Source Task.

The out-of-sample error of the ridge solution (6.5)-(6.6) to the source task is developed as follows. First, due to the layout of free parameters in $\hat{\boldsymbol{\theta}}$ we get that

$$(F.1) \quad \begin{aligned} \tilde{\mathcal{E}}_{\text{out}} &= \sigma_\xi^2 + \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}} - \boldsymbol{\theta} \right\|_2^2 \right] \\ &= \sigma_\xi^2 + \|\boldsymbol{\theta}_{\mathcal{S}^c}\|_2^2 + \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{S}} - \boldsymbol{\theta}_{\mathcal{S}} \right\|_2^2 \right]. \end{aligned}$$

Now we set the closed-form ridge solution from (6.6) to develop the third term in (F.1):

$$(F.2) \quad \begin{aligned} \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{S}} - \boldsymbol{\theta}_{\mathcal{S}} \right\|_2^2 \right] &= \mathbb{E} \left[\left\| (\mathbf{Z}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}} + \tilde{\alpha} \mathbf{I}_{\tilde{p}})^{-1} \mathbf{Z}_{\mathcal{S}}^T \mathbf{v} - \boldsymbol{\theta}_{\mathcal{S}} \right\|_2^2 \right] \\ &= \mathbb{E} \left[\left\| (\mathbf{Z}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}} + \tilde{\alpha} \mathbf{I}_{\tilde{p}})^{-1} \mathbf{Z}_{\mathcal{S}}^T (\mathbf{Z}_{\mathcal{S}^c} \boldsymbol{\theta}_{\mathcal{S}^c} + \boldsymbol{\xi}) \right\|_2^2 \right] + \mathbb{E} \left[\left\| ((\mathbf{Z}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}} + \tilde{\alpha} \mathbf{I}_{\tilde{p}})^{-1} \mathbf{Z}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}} - \mathbf{I}_{\tilde{p}}) \boldsymbol{\theta}_{\mathcal{S}} \right\|_2^2 \right]. \end{aligned}$$

Next, note that

$$(F.3) \quad \mathbb{E} \left[\left\| (\mathbf{Z}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}} + \tilde{\alpha} \mathbf{I}_{\tilde{p}})^{-1} \mathbf{Z}_{\mathcal{S}}^T (\mathbf{Z}_{\mathcal{S}^c} \boldsymbol{\theta}_{\mathcal{S}^c} + \boldsymbol{\xi}) \right\|_2^2 \right] = \left(\sigma_\xi^2 + \|\boldsymbol{\theta}_{\mathcal{S}^c}\|_2^2 \right) \mathbb{E} \left[\text{Tr} \left\{ (\mathbf{Z}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}} + \tilde{\alpha} \mathbf{I}_{\tilde{p}})^{-2} \mathbf{Z}_{\mathcal{S}}^T \mathbf{Z}_{\mathcal{S}} \right\} \right]$$

Also, we use the eigendecomposition $\mathbf{Z}_S^T \mathbf{Z}_S = \Phi_{\mathbf{Z}_S} \Lambda_{\mathbf{Z}_S} \Phi_{\mathbf{Z}_S}^T$ where $\Phi_{\mathbf{Z}_S}$ is a $\tilde{p} \times \tilde{p}$ orthonormal matrix and $\Lambda_{\mathbf{Z}_S} \triangleq \text{diag}\{\lambda_{\mathbf{Z}_S,1}, \dots, \lambda_{\mathbf{Z}_S,\tilde{p}}\}$ is the $\tilde{p} \times \tilde{p}$ diagonal matrix formed by the eigenvalues of $\mathbf{Z}_S^T \mathbf{Z}_S$. By using this eigendecomposition and (F.3), we develop (F.2) into the form of

$$\begin{aligned}
 \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_S - \boldsymbol{\theta}_S \right\|_2^2 \right] &= \\
 &= \left(\sigma_\xi^2 + \|\boldsymbol{\theta}_{S^c}\|_2^2 \right) \mathbb{E} \left[\text{Tr} \left\{ (\Lambda_{\mathbf{Z}_S} + \tilde{\alpha} \mathbf{I}_{\tilde{p}})^{-2} \Lambda_{\mathbf{Z}_S} \right\} \right] + \frac{\|\boldsymbol{\theta}_S\|_2^2}{\tilde{p}} \tilde{\alpha}^2 \mathbb{E} \left[\text{Tr} \left\{ (\Lambda_{\mathbf{Z}_S} + \tilde{\alpha} \mathbf{I}_{\tilde{p}})^{-2} \right\} \right] \\
 &= \mathbb{E} \left[\sum_{k=1}^{\tilde{p}} \frac{\left(\sigma_\xi^2 + \|\boldsymbol{\theta}_{S^c}\|_2^2 \right) \lambda_{\mathbf{Z}_S,k} + \frac{\|\boldsymbol{\theta}_S\|_2^2}{\tilde{p}} \tilde{\alpha}^2}{(\lambda_{\mathbf{Z}_S,k} + \tilde{\alpha})^2} \right].
 \end{aligned}
 \tag{F.4}$$

The optimal tuning is achieved by the $\tilde{\alpha}$ value that satisfies $\frac{\partial \tilde{\mathcal{E}}_{\text{out}}}{\partial \tilde{\alpha}} = 0$, which by using (F.1) and (F.4) is $\tilde{\alpha} = \frac{\tilde{p}(\sigma_\xi^2 + \|\boldsymbol{\theta}_{S^c}\|_2^2)}{\|\boldsymbol{\theta}_S\|_2^2}$. In the main text we explain how we approximate the optimal $\tilde{\alpha}$ in our experiments.

F.2. Ridge Regression: Error Expression and Optimal Tuning for the Target Task.

Let us start by explaining in more detail the ridge solution in (6.8). For this, note that the optimization constraints in (6.7) imply $\|\mathbf{y} - \mathbf{X}\mathbf{b}\|_2^2 = \left\| \mathbf{y} - \mathbf{X}_{\mathcal{F}}\mathbf{b}_{\mathcal{F}} - \mathbf{X}_{\mathcal{T}}\hat{\boldsymbol{\theta}}_{\mathcal{T}} \right\|_2^2$. Hence, the solution of (6.7) is equivalent to $\hat{\boldsymbol{\beta}}$ where $\hat{\boldsymbol{\beta}}_{\mathcal{Z}} = 0$, $\hat{\boldsymbol{\beta}}_{\mathcal{T}} = \hat{\boldsymbol{\theta}}_{\mathcal{T}}$, and

$$\hat{\boldsymbol{\beta}}_{\mathcal{F}} = \arg \min_{\mathbf{f} \in \mathbb{R}^p} \left\| \left(\mathbf{y} - \mathbf{X}_{\mathcal{T}}\hat{\boldsymbol{\theta}}_{\mathcal{T}} \right) - \mathbf{X}_{\mathcal{F}}\mathbf{f} \right\|_2^2 + \alpha \|\mathbf{f}\|_2^2.
 \tag{F.5}$$

The optimization in (F.5) has a standard ridge regression form and, accordingly, its closed-form solution is provided in (6.8).

Let us develop the out-of-sample error expression of the target task. Based on the layout of free, transferred, and zeroed parameters in $\hat{\boldsymbol{\beta}}$ we have

$$\begin{aligned}
 \mathcal{E}_{\text{out}} &= \sigma_\epsilon^2 + \mathbb{E} \left[\left\| \hat{\boldsymbol{\beta}} - \boldsymbol{\beta} \right\|_2^2 \right] \\
 &= \sigma_\epsilon^2 + \|\boldsymbol{\beta}_{\mathcal{Z}}\|_2^2 + \mathbb{E} \left[\left\| \hat{\boldsymbol{\theta}}_{\mathcal{T}} - \boldsymbol{\beta}_{\mathcal{T}} \right\|_2^2 \right] + \mathbb{E} \left[\left\| \hat{\boldsymbol{\beta}}_{\mathcal{F}} - \boldsymbol{\beta}_{\mathcal{F}} \right\|_2^2 \right].
 \end{aligned}
 \tag{F.6}$$

Then, similar to the proof given for the source task in Section F.1, we get that

$$\begin{aligned}
 & \mathbb{E} \left[\left\| \hat{\beta}_{\mathcal{F}} - \beta_{\mathcal{F}} \right\|_2^2 \right] = \\
 & = \left(\sigma_{\epsilon}^2 + \|\beta_{\mathcal{Z}}\|_2^2 + \mathbb{E} \left[\left\| \beta_{\mathcal{T}} - \hat{\theta}_{\mathcal{T}} \right\|_2^2 \right] \right) \mathbb{E} \left[\text{Tr} \left\{ (\Lambda_{\mathbf{X}_{\mathcal{F}}} + \alpha \mathbf{I}_p)^{-2} \Lambda_{\mathbf{X}_{\mathcal{F}}} \right\} \right] \\
 & \quad + \frac{\|\beta_{\mathcal{F}}\|_2^2}{p} \alpha^2 \mathbb{E} \left[\text{Tr} \left\{ (\Lambda_{\mathbf{X}_{\mathcal{F}}} + \alpha \mathbf{I}_p)^{-2} \right\} \right]. \\
 & = \mathbb{E} \left[\sum_{k=1}^p \frac{\left(\sigma_{\epsilon}^2 + \|\beta_{\mathcal{Z}}\|_2^2 + \mathbb{E} \left[\left\| \beta_{\mathcal{T}} - \hat{\theta}_{\mathcal{T}} \right\|_2^2 \right] \right) \lambda_{\mathbf{X}_{\mathcal{F}},k} + \frac{\|\beta_{\mathcal{F}}\|_2^2}{p} \alpha^2}{(\lambda_{\mathbf{X}_{\mathcal{F}},k} + \alpha)^2} \right],
 \end{aligned}
 \tag{F.7}$$

where we use the eigendecomposition $\mathbf{X}_{\mathcal{F}}^T \mathbf{X}_{\mathcal{F}} = \Phi_{\mathbf{X}_{\mathcal{F}}} \Lambda_{\mathbf{X}_{\mathcal{F}}} \Phi_{\mathbf{X}_{\mathcal{F}}}^T$ where $\Phi_{\mathbf{X}_{\mathcal{F}}}$ is a $p \times p$ orthonormal matrix and $\Lambda_{\mathbf{X}_{\mathcal{F}}} \triangleq \text{diag}\{\lambda_{\mathbf{X}_{\mathcal{F}},1}, \dots, \lambda_{\mathbf{X}_{\mathcal{F}},p}\}$ is the $p \times p$ diagonal matrix of the eigenvalues of $\mathbf{X}_{\mathcal{F}}^T \mathbf{X}_{\mathcal{F}}$.

Here, the optimal tuning is obtained by the α value that provides $\frac{\partial \mathcal{E}_{\text{out}}}{\partial \alpha} = 0$. Then, (F.6) and (F.7) imply that the optimal tuning is given by

$$\alpha = \frac{p \left(\sigma_{\epsilon}^2 + \|\beta_{\mathcal{Z}}\|_2^2 + \mathbb{E} \left[\left\| \beta_{\mathcal{T}} - \hat{\theta}_{\mathcal{T}} \right\|_2^2 \right] \right)}{\|\beta_{\mathcal{F}}\|_2^2}.
 \tag{F.8}$$

In our experiments we assume that only $\|\beta\|_2^2$, $\|\beta - \theta\|_2^2$, σ_{ϵ} , p , t , and d are known; thus, we use the approximations $\|\beta_{\mathcal{F}}\|_2^2 \approx \frac{p}{d} \|\beta\|_2^2$, $\|\beta_{\mathcal{Z}}\|_2^2 \approx \left(1 - \frac{p+t}{d}\right) \|\beta\|_2^2$, and

$$\mathbb{E} \left[\left\| \beta_{\mathcal{T}} - \hat{\theta}_{\mathcal{T}} \right\|_2^2 \right] \approx \frac{t}{d} \mathbb{E} \left[\|\beta - \theta\|_2^2 \right] = \frac{t}{d} \mathbb{E} \left[\|\beta - (\mathbf{H}\beta + \boldsymbol{\eta})\|_2^2 \right] = \frac{t}{d} \|(\mathbf{I}_d - \mathbf{H})\beta\|_2^2 + t\sigma_{\eta}^2,$$

to approximate (F.8) as $\alpha = \frac{d\sigma_{\epsilon}^2}{\|\beta\|_2^2} + d - p - t + t \frac{\|(\mathbf{I}_d - \mathbf{H})\beta\|_2^2 + d\sigma_{\eta}^2}{\|\beta\|_2^2}$ in our experiments.

REFERENCES

- [1] P. L. BARTLETT, P. M. LONG, G. LUGOSI, AND A. TSIGLER, *Benign overfitting in linear regression*, Proc. Natl. Acad. Sci. USA, 117 (2020), pp. 30063–30070.
- [2] M. BELKIN, D. HSU, S. MA, AND S. MANDAL, *Reconciling modern machine-learning practice and the classical bias–variance trade-off*, Proc. Natl. Acad. Sci. USA, 116 (2019), pp. 15849–15854.
- [3] M. BELKIN, D. HSU, AND J. XU, *Two models of double descent for weak features*, SIAM J. Math. Data Sci., 2 (2020), pp. 1167–1180.
- [4] Y. BENGIO, *Deep learning of representations for unsupervised and transfer learning*, in ICML Workshop on Unsupervised and Transfer Learning, 2012, pp. 17–36.
- [5] L. BREIMAN AND D. FREEDMAN, *How many variables should be entered in a regression equation?*, J. Amer. Statist. Assoc., 78 (1983), pp. 131–136.
- [6] S. S. CHEN, D. L. DONOHO, AND M. A. SAUNDERS, *Atomic decomposition by basis pursuit*, SIAM Rev., 43 (2001), pp. 129–159.

- [7] Y. DAR, P. MAYER, L. LUZI, AND R. G. BARANIUK, *Subspace fitting meets regression: The effects of supervision and orthonormality constraints on double descent of generalization errors*, in International Conference on Machine Learning (ICML), 2020.
- [8] O. DHIFALLAH AND Y. M. LU, *Phase transitions in transfer learning for high-dimensional perceptrons*, Entropy, 23 (2021).
- [9] M. GEIGER, A. JACOT, S. SPIGLER, F. GABRIEL, L. SAGUN, S. D’ASCOLI, G. BIROLI, C. HONGLER, AND M. WYART, *Scaling description of generalization with number of parameters in deep learning*, J. Stat. Mech., 2 (2020), p. 023401.
- [10] F. GERACE, L. SAGLIETTI, S. S. MANNELLI, A. SAXE, AND L. ZDEBOROVÁ, *Probing transfer learning with a model of synthetic correlated datasets*, Mach. Learn.: Sci. Technol., 3 (2022), p. 015030.
- [11] T. HASTIE, A. MONTANARI, S. ROSSET, AND R. J. TIBSHIRANI, *Surprises in high-dimensional ridgeless least squares interpolation*, Ann. Statist., 50 (2022), pp. 949 – 986.
- [12] F. HIAI AND D. PETZ, *Asymptotic freeness almost everywhere for random matrices*, Acta Sci. Math. Szeged, 66 (2000), pp. 801–826.
- [13] S. KORNBLITH, J. SHLENS, AND Q. V. LE, *Do better imagenet models transfer better?*, in IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2019, pp. 2661–2671.
- [14] A. K. LAMPINEN AND S. GANGULI, *An analytic theory of generalization dynamics and transfer learning in deep linear networks*, in International Conference on Learning Representations (ICLR), 2019.
- [15] Y. LI AND Y. WEI, *Minimum ℓ_1 -norm interpolators: Precise asymptotics and multiple descent*, preprint, arXiv:2110.09502, (2021).
- [16] M. LONG, H. ZHU, J. WANG, AND M. I. JORDAN, *Deep transfer learning with joint adaptation networks*, in International Conference on Machine Learning (ICML), 2017, pp. 2208–2217.
- [17] S. MEI AND A. MONTANARI, *The generalization error of random features regression: Precise asymptotics and the double descent curve*, Comm. Pure Appl. Math., 75 (2022), pp. 667–766.
- [18] P. P. MITRA, *Understanding overfitting peaks in generalization error: Analytical risk curves for ℓ_2 and ℓ_1 penalized interpolation*, preprint, arXiv:1906.03667, (2019).
- [19] V. MUTHUKUMAR, K. VODRAHALI, V. SUBRAMANIAN, AND A. SAHAI, *Harmless interpolation of noisy data in regression*, IEEE J. Sel. Areas Inf. Theory, (2020).
- [20] D. OBST, B. GHATTAS, J. CUGLIARI, G. OPPENHEIM, S. CLAUDEL, AND Y. GOUDE, *Transfer learning for linear regression: a statistical test of gain*, preprint, arXiv:2102.09504, (2021).
- [21] S. J. PAN AND Q. YANG, *A survey on transfer learning*, IEEE Trans. Knowl. Data Eng., 22 (2009), pp. 1345–1359.
- [22] M. RAGHU, C. ZHANG, J. KLEINBERG, AND S. BENGIO, *Transfusion: Understanding transfer learning for medical imaging*, in Advances in Neural Information Processing Systems (NeurIPS), 2019, pp. 3347–3357.
- [23] M. T. ROSENSTEIN, Z. MARX, L. P. KAEHLING, AND T. G. DIETTERICH, *To transfer or not to transfer*, in NIPS Workshop on Transfer Learning, 2005.
- [24] H.-C. SHIN, H. R. ROTH, M. GAO, L. LU, Z. XU, I. NOGUES, J. YAO, D. MOLLURA, AND R. M. SUMMERS, *Deep convolutional neural networks for computer-aided detection: CNN architectures, dataset characteristics and transfer learning*, IEEE Trans. Med. Imag., 35 (2016), pp. 1285–1298.
- [25] S. SPIGLER, M. GEIGER, S. D’ASCOLI, L. SAGUN, G. BIROLI, AND M. WYART, *A jamming transition from under-to over-parametrization affects loss landscape and generalization*, J. Phys. A, 52 (2019), p. 474001.
- [26] A. M. TULINO AND S. VERDÚ, *Random matrix theory and wireless communications*, Foundations and Trends in Communications and Information Theory, (2004), pp. 1–182.
- [27] G. WANG, K. DONHAUSER, AND F. YANG, *Tight bounds for minimum ℓ_1 -norm interpolation of noisy data*, in International Conference on Artificial Intelligence and Statistics (AISTATS), 2022, pp. 10572–10602.
- [28] J. XU AND D. J. HSU, *On the number of variables to use in principal component regression*, in Advances in Neural Information Processing Systems (NeurIPS), 2019, pp. 5095–5104.
- [29] A. R. ZAMIR, A. SAX, W. SHEN, L. J. GUIBAS, J. MALIK, AND S. SAVARESE, *Taskonomy: Disentangling task transfer learning*, in IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2018, pp. 3712–3722.

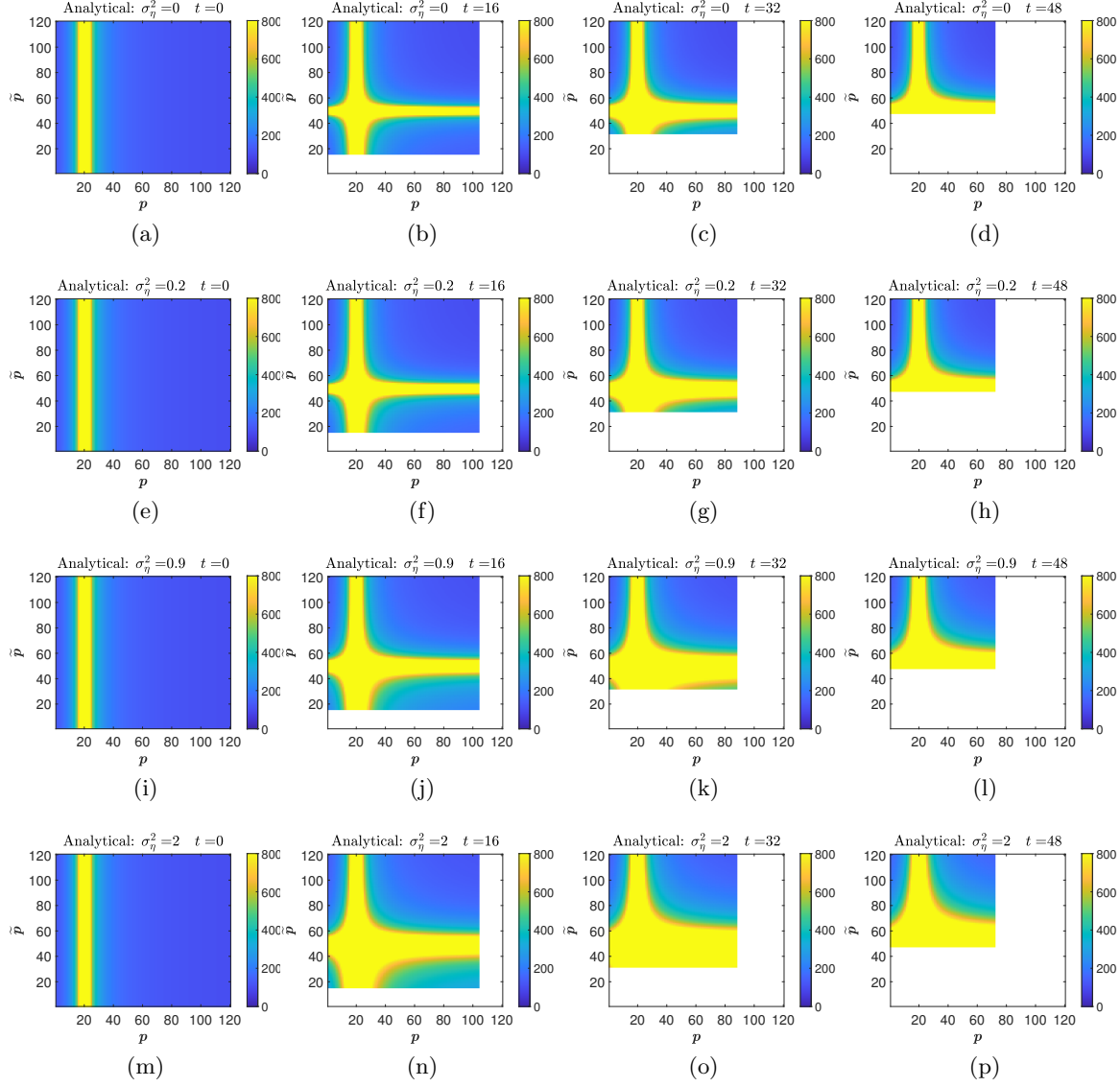


Figure 10: **Analytical** evaluation of the expected out-of-sample squared error of the target task, $\mathbb{E}_{\mathcal{L}}[\mathcal{E}_{\text{out}}]$, with respect to the number of free parameters $\tilde{p}$ and p (in the source and target tasks, respectively). Each row of subfigures considers a different case of the relation (2.7) between the source and target tasks in the form of a different noise variance σ_{η}^2 whereas $\mathbf{H}$ is a local averaging operator with neighborhood size 5 for all. Each column of subfigures considers a different number of transferred parameters t . Here $d = 120$, $\tilde{n} = 50$, $n = 20$, $\|\beta\|_2^2 = d$, $\sigma_{\epsilon}^2 = 0.05 \cdot d$, $\sigma_{\xi}^2 = 0.025 \cdot d$. The white regions correspond to $(\tilde{p}, p)$ settings eliminated by the value of t in the specific subfigure. The yellow-colored areas correspond to values greater or equal to 800. See Fig. 11 for the corresponding empirical evaluation.

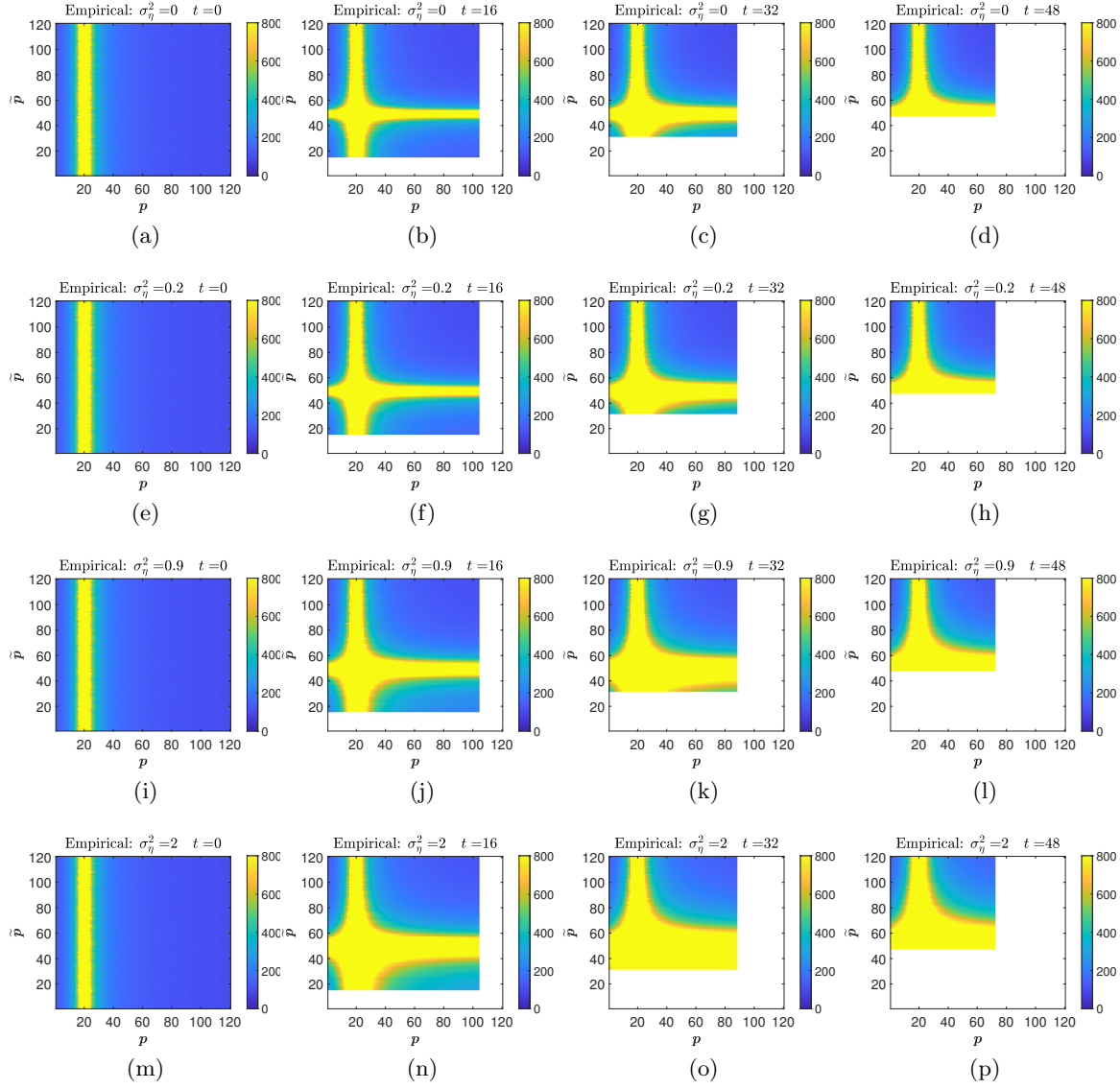


Figure 11: **Empirical** evaluation of the expected out-of-sample squared error of the target task, $\mathbb{E}_{\mathcal{L}}[\mathcal{E}_{\text{out}}]$, with respect to the number of free parameters $\tilde{p}$ and p (in the source and target tasks, respectively). The presented values obtained by averaging over 250 experiments. The figures here have settings as in the corresponding figures in the analytical evaluation in Fig. 10.

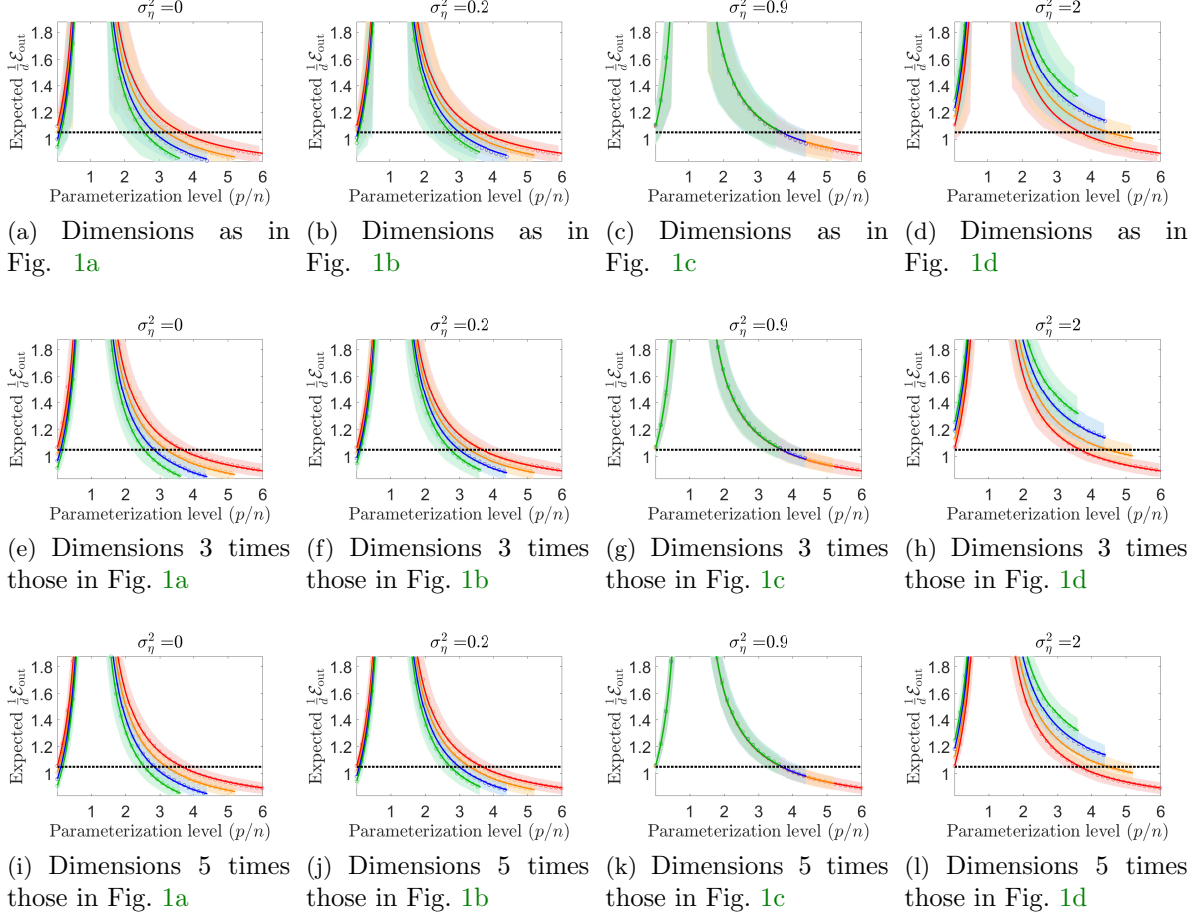


Figure 12: Demonstration of error concentration. Each row of subfigures corresponds to the settings of Figures 1a-1d but with a different proportional increase of dimensions and dimension-dependent quantities (see explanation in Section 3.2 in the main text). The empirical standard deviations are denoted as shaded areas in colors corresponding to the on-average error curves (solid lines and markers denote the analytical and empirical evaluations of the expected error, respectively). Lines, markers and areas in red correspond to $t = 0$ (no parameters are transferred); orange corresponds to transferring $t = 16 \times \frac{d}{120}$ parameters; blue corresponds to $t = 32 \times \frac{d}{120}$; green corresponds to $t = 48 \times \frac{d}{120}$. Note that $\frac{d}{120}$ equals to 1, 3, 5, in the first, second, third rows of subfigures, respectively. The axes in this figure are normalized to be dimension-independent.

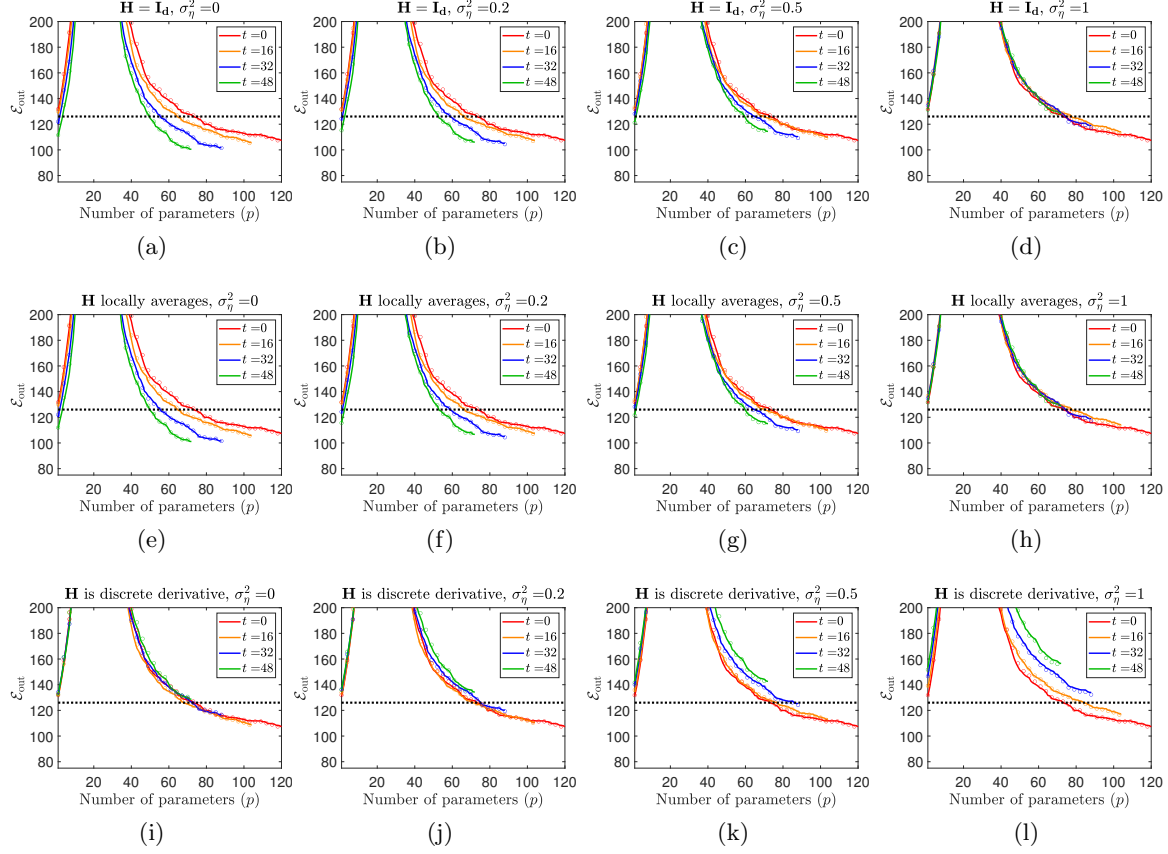


Figure 13: Analytical (solid lines) and empirical (circle markers) values of $\mathcal{E}_{\text{out}}^{(\mathcal{L})}$ for specific, non-random coordinate layouts. **The true solution β has linearly-increasing values.** All subfigures use the same sequential evolution of $\mathcal{L}$ with p . Each subfigure considers a different case of the relation (2.7) between the source and target tasks: each column of subfigures has a different σ_η^2 value, and each row of subfigures corresponds to a different linear operator $\mathbf{H}$. The analytical values, computed using Theorem 3.1, are presented using solid-line curves, and the respective empirical results obtained from averaging over 250 experiments are denoted by circle markers. Each curve color refers to a different number of transferred parameters.

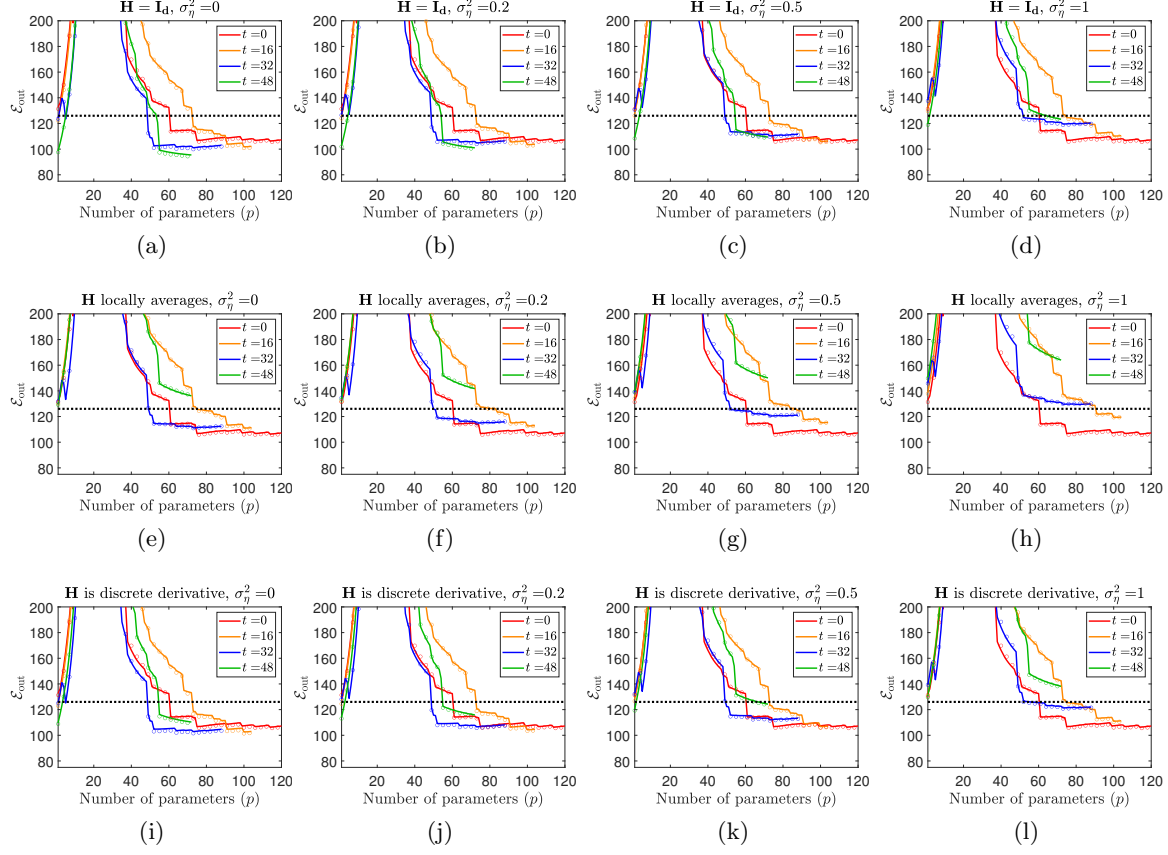


Figure 14: Analytical (solid lines) and empirical (circle markers) values of $\mathcal{E}_{\text{out}}^{(\mathcal{L})}$ for specific, non-random coordinate layouts. **The true solution β has a sparse form of values.** All subfigures use the same sequential evolution of $\mathcal{L}$ with p . Each subfigure considers a different case of the relation (2.7) between the source and target tasks: each column of subfigures has a different σ_η^2 value, and each row of subfigures corresponds to a different linear operator $\mathbf{H}$. The analytical values, computed using Theorem 3.1, are presented using solid-line curves, and the respective empirical results obtained from averaging over 250 experiments are denoted by circle markers. Each curve color refers to a different number of transferred parameters.

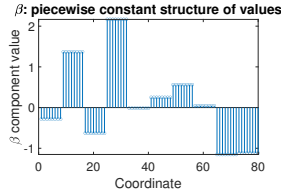


Figure 15: The piecewise-constant structure of β that was used in part of the experiments.

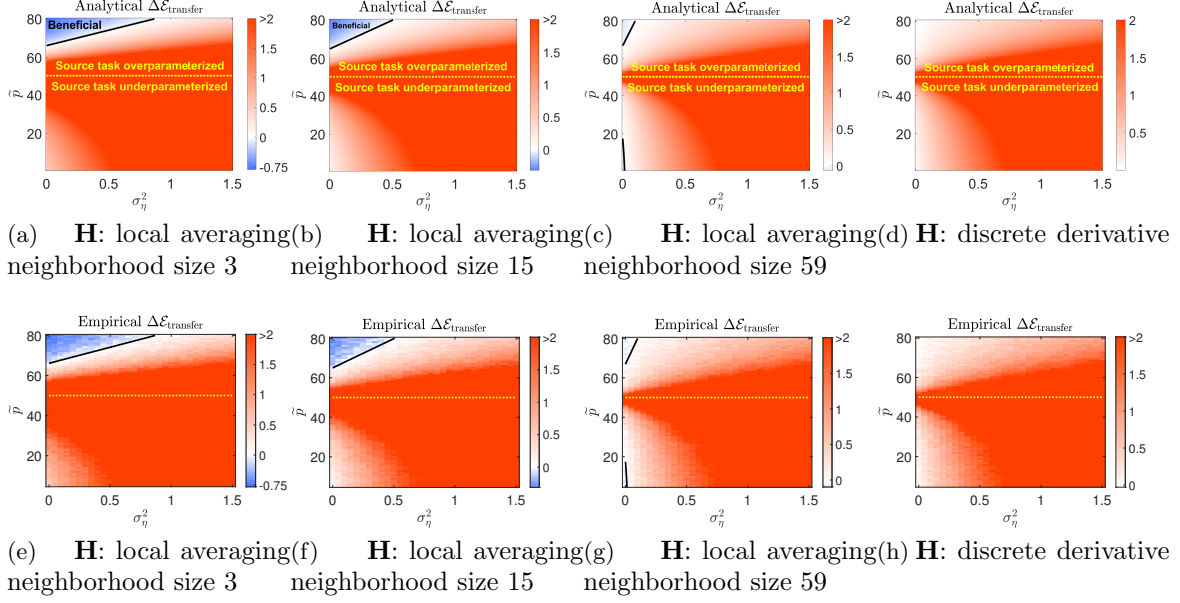


Figure 16: The analytical and empirical values of $\Delta\mathcal{E}_{\text{transfer}}$ defined in Corollary 3.4 (here, normalized by t , namely, the expected error difference due to transfer of a parameter from the source to target task) as a function of $\tilde{p}$ and σ_η^2 . The positive and negative values of $\frac{1}{t}\Delta\mathcal{E}_{\text{transfer}}$ appear in color scales of red and blue, respectively. The regions of negative values (appear in shades of blue) correspond to beneficial transfer of parameters. The positive values were truncated in the value of 2 for the clarity of visualization. Each subfigure corresponds to a different task relation model induced by the definitions of $\mathbf{H}$ as: (a)-(c),(e)-(g) local averaging operators with different neighborhood sizes, (d),(h) discrete derivative. For all the subfigures, $d = 80$, $\tilde{n} = 50$, $\|\beta\|_2^2 = d$, $\sigma_\xi^2 = 0.025 \cdot d$. Here, all the subfigures correspond to a β vector with a piecewise-constant form (see Fig. 15).

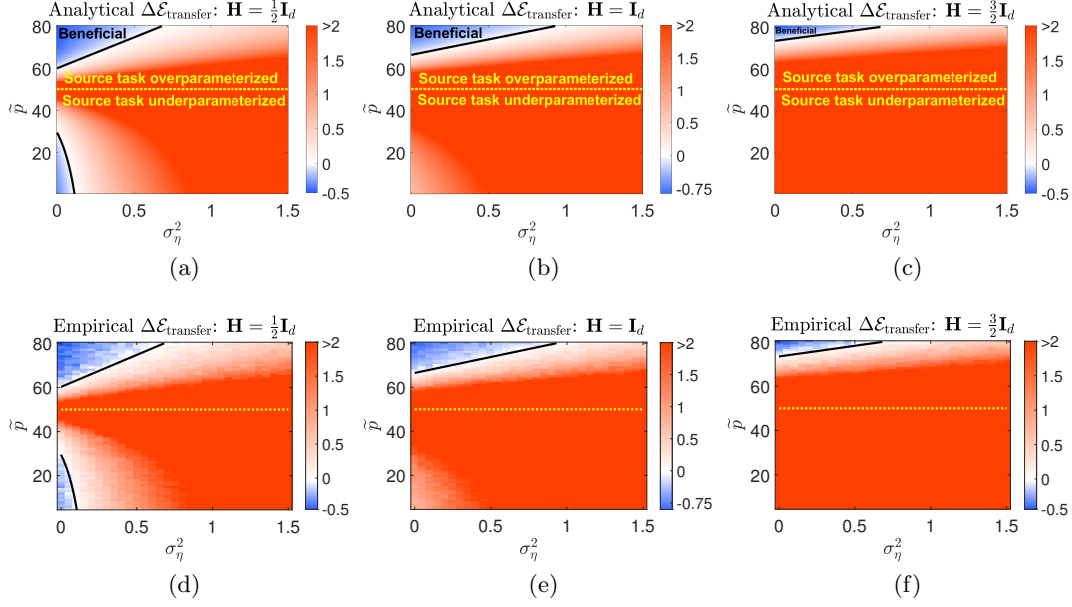


Figure 17: The analytical and empirical values of $\Delta \mathcal{E}_{\text{transfer}}$ defined in Corollary 3.4 (here, normalized by t , namely, the expected error difference due to transfer of a parameter from the source to target task) as a function of $\tilde{p}$ and σ_η^2 . The positive and negative values of $\frac{1}{t}\Delta \mathcal{E}_{\text{transfer}}$ appear in color scales of red and blue, respectively. The regions of negative values (appear in shades of blue) correspond to beneficial transfer of parameters. The positive values were truncated in the value of 2 for the clarity of visualization. Each subfigure corresponds to a different task relation model induced by the definitions of $\mathbf{H}$ as $\mathbf{H} = \frac{1}{2}\mathbf{I}_d$, $\mathbf{H} = \mathbf{I}_d$, and $\mathbf{H} = \frac{3}{2}\mathbf{I}_d$. For all the subfigures, $d = 80$, $\tilde{n} = 50$, $\|\boldsymbol{\beta}\|_2^2 = d$, $\sigma_\xi^2 = 0.025 \cdot d$. Here, all the subfigures correspond to a $\boldsymbol{\beta}$ vector with a linear form (see Fig. 4a).

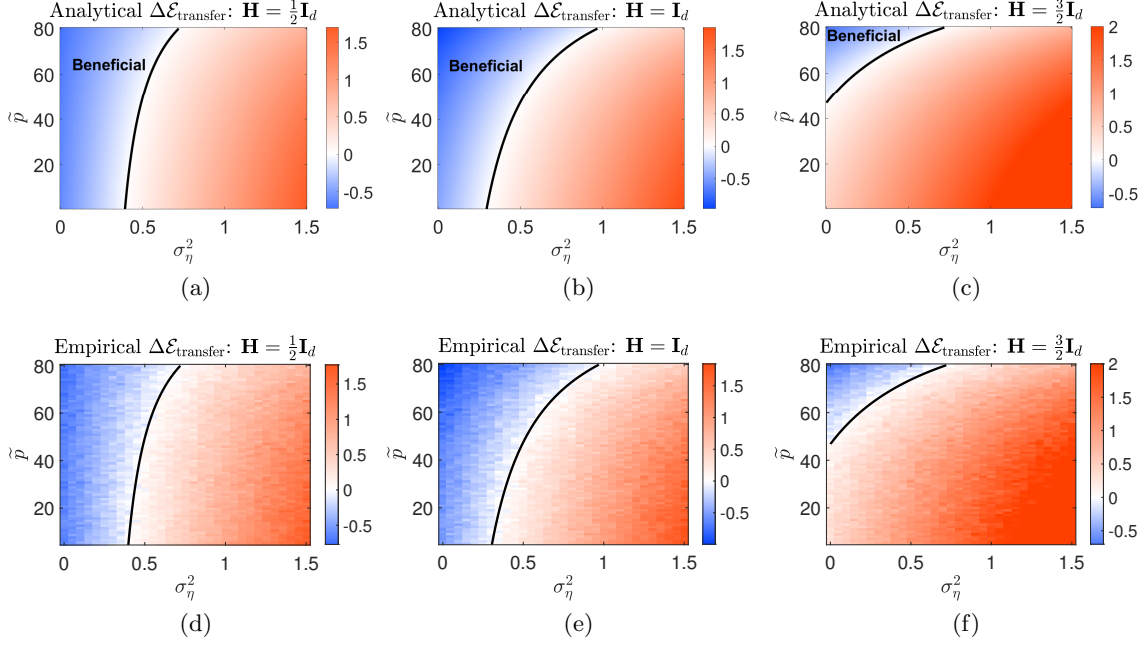


Figure 18: The analytical (top row of subfigures) and empirical (bottom row of subfigures) values of $\Delta \mathcal{E}_{\text{transfer}}$ defined in Corollary 3.4 (here, normalized by t , namely, the expected error difference due to transfer of a parameter from the source to target task) as a function of $\tilde{p}$ and σ_η^2 . The positive and negative values of $\frac{1}{t} \Delta \mathcal{E}_{\text{transfer}}$ appear in color scales of red and blue, respectively. The regions of negative values (appear in shades of blue) correspond to beneficial transfer of parameters. The positive values were truncated in the value of 2 for the clarity of visualization. Each column of subfigures correspond to a different task relation model induced by the definitions of $\mathbf{H}$ as $\mathbf{H} = \frac{1}{2} \mathbf{I}_d$, $\mathbf{H} = \mathbf{I}_d$, and $\mathbf{H} = \frac{3}{2} \mathbf{I}_d$. For all the subfigures, $d = 80$, $\tilde{n} = 150$, $\|\boldsymbol{\beta}\|_2^2 = d$, $\sigma_\xi^2 = 0.025 \cdot d$. Here, all the subfigures correspond to a $\boldsymbol{\beta}$ vector with a linear form (see Fig. 4a). Note that the results in this figure are for $\tilde{n} > d$.

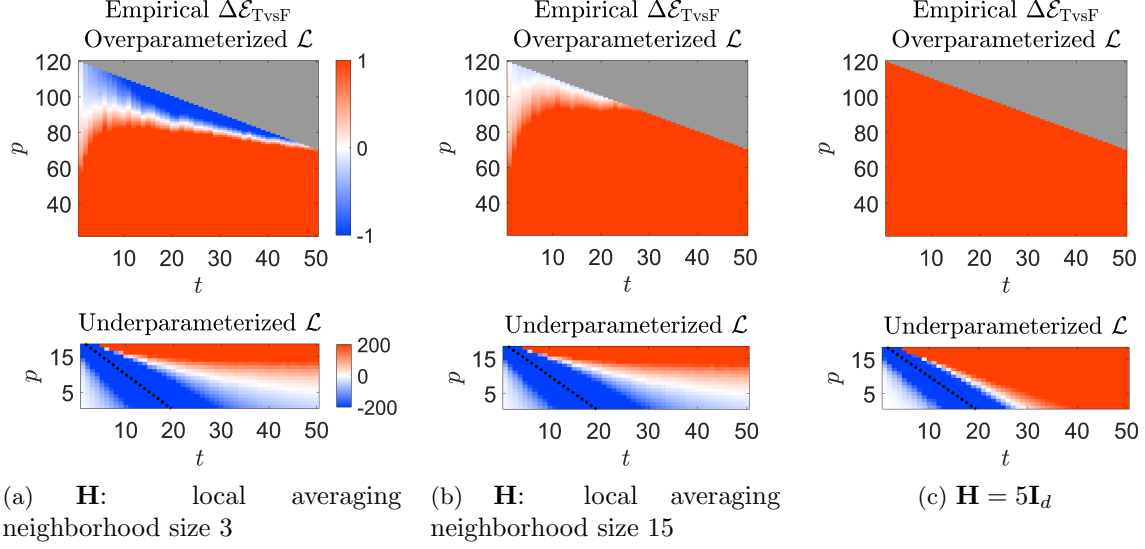


Figure 19: The empirical values of $\Delta\mathcal{E}_{\text{TvsF}}$ (namely, the average error difference due to transfer of an arbitrarily-selected set of t parameters versus setting them as free parameters) as a function of t and p . The settings and visualization of results are as in the corresponding analytical results in Fig. 6.

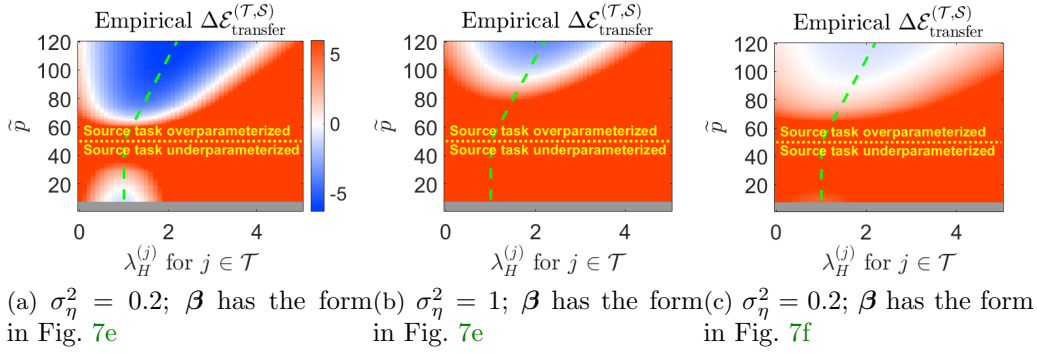


Figure 20: The empirical evaluation of $\Delta\mathcal{E}_{\text{transfer}}^{(\mathcal{T},\mathcal{S})}$ as a function of $\tilde{p}$ and a value that determines the eigenvalues of $\mathbf{H}$ in $\mathcal{T}$. The settings and visualization of results are as in the corresponding analytical results in Fig. 7.