



Universiteit
Leiden
The Netherlands

The many faces of online learning

Hoeven, D. van der

Citation

Hoeven, D. van der. (2021, March 4). *The many faces of online learning*. Retrieved from <https://hdl.handle.net/1887/3147345>

Version: Publisher's Version

License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/3147345>

Note: To cite this publication please use the final published version (if applicable).

Cover Page



Universiteit Leiden



The handle <https://hdl.handle.net/1887/3147345> holds various files of this Leiden University dissertation.

Author: Hoeven, D. van der

Title: The many faces of online learning

Issue Date: 2021-03-4

User-Specified Local Differential Privacy in Unconstrained Adaptive Online Learning

This chapter is based on Van der Hoeven, D. (2019). User-specified local differential privacy in unconstrained adaptive online learning. In *Advances in Neural Information Processing Systems 32*, pages 14103–14112.

Abstract

Local differential privacy is a strong notion of privacy in which the provider of the data guarantees privacy by perturbing the data with random noise. In the standard application of local differential privacy the distribution of the noise is constant and known by the learner. In this chapter we generalize this approach by allowing the provider of the data to choose the distribution of the noise without disclosing any parameters of the distribution to the learner, under the constraint that the distribution is symmetrical. We consider this problem in the unconstrained Online Convex Optimization setting with noisy feedback. In this setting the learner receives the subgradient of a loss function, perturbed by noise, and aims to achieve sublinear regret with respect to some competitor, without constraints on the norm of the competitor. We derive the first algorithms that have adaptive regret bounds in this setting, i.e. our algorithms adapt to the unknown competitor norm, unknown noise, and unknown sum of the norms of the subgradients, matching state of the art bounds in all cases.

3.1 Introduction

In learning, a natural tension exists between learners and the providers of data. The learner aims to make optimal use of the data, perhaps even at the cost of the privacy of the providers. To nevertheless ensure sufficient privacy the provider can add random noise to the data that he sends to the learner. This idea is called ϵ -local differential privacy (Wasserman and Zhou, 2010; Duchi et al., 2014) and the standard implementation has constant ϵ for all providers. However, not all providers care equivalently about their privacy (Song et al., 2015). Some providers may wish to aid the learner in making optimal use of their data, while other providers value their privacy over helping the learner. For instance, celebrities might care more for their privacy than others because they want to preserve the privacy they have left. To complicate things further, the providers of the data may not wish to reveal how much they care about their privacy, because when privacy levels differ between providers these privacy levels become privacy sensitive themselves. Furthermore, not all parts of the data are equally privacy sensitive. For example, tweets are already publicly available, but browsing history may contain sensitive information that should be kept private. To capture these varying privacy constraints we allow each provider to choose how much noise is added for each dimension of the data.

In this chapter, we consider these problems in the Online Convex Optimization (OCO) setting (Hazan et al., 2016) with local differential privacy guarantees. The OCO framework is a popular and successful framework to design and analyse many algorithms used to train machine learning models. The OCO setting proceeds in rounds $t = 1, \dots, T$. In a given round t the learner is to provide a prediction $\mathbf{w}_t \in \mathbb{R}^d$. An adversary then chooses a convex loss function ℓ_t and sends a subgradient $\mathbf{g}_t \in \partial \ell_t(\mathbf{w}_t)$ to the learner. We work with an unconstrained domain for $\mathbf{w}$, which has recently grown in popularity (see McMahan and Orabona (2014); Foster et al. (2015); Orabona and Pál (2016); Foster et al. (2017); Cutkosky and Boahen (2017); Kotłowski (2017); Cutkosky and Orabona (2018); Foster et al. (2018b); Jun and Orabona (2019)). We aim to develop online learning methods that make the best use of data providers who wish to help the learner while at the same time guaranteeing the desired level of privacy for providers that care about their privacy, without knowing how much each provider cares for their privacy.

We consider the local differential privacy model with varying levels of privacy unknown to the learner. Differential privacy (Dwork and Roth, 2014) is a privacy model that is used in many recent machine-learning applications. The local differential privacy model is a variant of differential privacy in which the learner can only access the data of the provider via noisy estimates (Wasserman and Zhou, 2010;

Duchi et al., 2014). The local differential privacy model with varying levels of privacy appeared before in Song et al. (2015), but with known levels of noise and only two levels of noise.

Learning in our setting is modelled by the OCO framework with noisy estimates of the subgradient (see also Jun and Orabona (2019)). To ensure local differential privacy the provider adds zero-mean noise $\xi_t \in \mathbb{R}^d$ to the subgradient g_t . The learner then receives the perturbed subgradient $\tilde{g}_t = g_t + \xi_t$. We allow each ξ_t to follow a different distribution each round to satisfy different privacy guarantees. In the standard OCO framework the goal of the learner is to minimize the *regret* with respect to some parameter $u \in \mathbb{R}^d$:

$$\mathcal{R}_T(u) = \sum_{t=1}^T (\ell_t(w_t) - \ell_t(u)).$$

However, since the learner receives perturbed subgradients we consider the expected regret $\mathbb{E}[\mathcal{R}(u)]$, where the expectation is over the randomness in w_t due to the noisy subgradients. The setting will be formally introduced in Section 3.2. Because $\tilde{g}_t \in \mathbb{R}^d$, standard algorithms for unconstrained domains do not work since they require bounded $\tilde{g}_t$. Initial work in this setting by Jun and Orabona (2019) was motivated by a lower bound of Cutkosky and Boahen (2017), which shows that one can suffer an exponential penalty when both the domain and subgradients are unbounded. They replace the boundedness assumption on $\tilde{g}_t$ by a boundedness assumption on $\mathbb{E}[\tilde{g}_t]$ and an assumption on the tails of the noise distribution. Jun and Orabona (2019) achieved expected regret guarantees of $O(\|u\| \sqrt{(G^2 + \sigma^2)T \ln(1 + \|u\|T)})$, where σ^2 is a uniform upper bound on $\mathbb{E}[\|\xi_t\|_*^2]$, G^2 is a uniform upper bound on $\|g_t\|_*^2$, and $\|\cdot\|$ and $\|\cdot\|_*$ are dual norms. This bound is useful when the distribution of the noise is constant and known and an adversary selects g_t . We derive an algorithm that satisfies

$$\mathbb{E}[\mathcal{R}_T(u)] = O\left(\|u\| \sqrt{\left(G^2T + \sum_{t=1}^T \sigma_t^2\right) \ln(1 + \|u\|T)}\right), \quad (3.1.1)$$

where $\sigma_t^2 = \mathbb{E}[\|\xi_t\|_*^2]$. This bound can be smaller in cases where only a few σ_t are large but most are small, for example when only few providers have privacy requirements. In fact, we will prove something stronger than (3.1.1):

$$\mathbb{E}[\mathcal{R}_T(u)] = O\left(\mathbb{E}[\|u\|] \sqrt{\sum_{t=1}^T \|\tilde{g}_t\|_*^2 \ln(1 + \|u\|T)}\right), \quad (3.1.2)$$

which implies (3.1.1) via Jensen's inequality and $\mathbb{E}[\|\tilde{\mathbf{g}}_t\|_*^2] \leq 3\mathbb{E}[\|\xi_t\|_*^2] + 3\mathbb{E}[\|\mathbf{g}_t\|_*^2]$. This bound was motivated by work in the noiseless setting, where $O(\|\mathbf{u}\| \sqrt{\sum_{t=1}^T \|\mathbf{g}_t\|_*^2 \ln(1 + \|\mathbf{u}\|T)})$ bounds are possible (Cutkosky and Orabona, 2018). With these type of bounds, when the sum of the squared norms of the subgradients is small the regret is also small. To achieve (3.1.2) we require two assumptions: bounded $\|\mathbf{g}_t\|_*$ and zero-mean symmetrical noise ξ_t . The assumption on $\mathbf{g}_t$ is common in standard OCO. The symmetrical noise assumption is satisfied for common mechanisms to ensure local differential privacy. The dependence on $\mathbb{E}[\|\xi_t\|_*^2]$ and $\mathbb{E}[\|\mathbf{g}_t\|_*^2]$ is unimprovable, which is shown by the lower bound for this setting by Jun and Orabona (2019).

The algorithms in this chapter are built using the recently developed wealth-regret duality approach (McMahan and Streeter, 2012). We provide two algorithms. The first achieves the bound in (3.1.2). The second algorithm satisfies (3.1.2) for each dimension separately. This second algorithm can exploit sparse privacy structures, which combined with sparse subgradients yields low expected regret bounds.

Contributions We extend the known results in several directions. Many common local differential privacy applications use symmetric additive noise (Laplace mechanism, normal mechanism). We use the symmetry of the noise to adapt to unknown levels of privacy and achieve adaptive expected regret bounds. We also adapt to dimension specific privacy requirements, again without requiring knowledge of the structure of the noise other than symmetry in each dimension. Our algorithms interpolate between no noise and maximum noise, matching state of the art bounds in both cases. This can reduce the cost of privacy in some cases, outlined in Section 3.4. Our work partially answers two problems left open by Jun and Orabona (2019). The first question asks whether or not data-dependent bounds are possible in the noisy OCO setting, which we answer affirmatively. The second question is how to adapt to different levels of noise without using extra parameters compared to the noiseless setting, which we do for symmetric noise.

Related work There has been significant work on unconstrained and adaptive methods in OCO with noiseless subgradients $\mathbf{g}_t$ (Foster et al., 2015; Orabona and Pál, 2016; Foster et al., 2017; Cutkosky and Boahen, 2017; Kotłowski, 2017; Cutkosky and Orabona, 2018; Foster et al., 2018b). However, these results do not extend to the setting with noisy unbounded subgradients $\tilde{\mathbf{g}}_t$, which is possible with our work. For bounded domains regret bounds of $O(D \sqrt{\sum_{t=1}^T \|\tilde{\mathbf{g}}_t\|_*^2})$ are possible without knowledge of the noise (Duchi et al., 2011; Orabona and Pál, 2018), where D is an upper bound on $\|\mathbf{u}\|$. However, these bounds do not adapt to unknown $\|\mathbf{u}\|$,

which may be costly for large D but small $\|u\|$. We provide an algorithm that both scales with $\|u\|$ instead of D and does not require knowledge of the noise.

There is a body of literature in the differential privacy setting with online feedback (Jain et al., 2012; Jain and Thakurta, 2014; Thakurta and Smith, 2013; Agarwal and Singh, 2017; Abernethy et al., 2019). In this chapter we consider *local* differential privacy (Wasserman and Zhou, 2010; Duchi et al., 2014), which is a stronger notion of privacy than differential privacy. Duchi et al. (2014) provide an algorithm with constant local differential privacy that learns by using SGD. (Song et al., 2015) derive how to use knowledge of several levels of local differential privacy for SGD, but only with two different levels of noise. Jun and Orabona (2019) consider local privacy with an unbounded domain and constant noise. With knowledge of the noise it is possible to extend the results of Jun and Orabona (2019) to achieve (3.1.1), but not (3.1.2).

Outline In Section 3.2 we introduce our problem formally and introduce the key techniques. In Section 3.3 we derive a one-dimensional algorithm that achieves our goals, which we use in a black-box reduction in Section 3.3.1 and we apply it coordinate-wise in Section 3.3.2. Section 3.4 contains two scenarios in which our new algorithm achieves improvements compared to current algorithms. Finally, in Section 3.5 we present our conclusions.

3.2 Problem Formulation and Preliminaries

In this Section we describe our notation, introduce the version of local differential privacy we use, briefly introduce the OCO setting with noisy subgradients, and provide some background to the reward-regret duality paradigm.

Notation. A random variable x is called symmetric if the density function ρ of the random variable $z = x - \mathbb{E}[x]$ satisfies $\rho(z) = \rho(-z)$. The inner product between vectors $g \in \mathbb{R}^d$ and $w \in \mathbb{R}^d$ is denoted by $\langle w, g \rangle$. The Fenchel conjugate of a convex function F , F^* is defined as $F^*(w) = \sup_g \langle w, g \rangle - F(g)$. $\|\cdot\|$ denotes a norm and $\|g\|_* = \sup_{w: \|w\| \leq 1} \langle w, g \rangle$ denotes the dual norm. $g_{t,j}$ indicates the j^{th} component of vector g_t .

3.2.1 User-Specified Local Differential Privacy

In the local differential privacy setting each datum is kept private from the learner. The standard definition of local privacy requires a randomiser R that perturbs g_t with random noise ξ_t , where $\xi_1, \dots, \xi_T$ are independently distributed (Wasserman

and Zhou, 2010; Kasiviswanathan et al., 2011; Duchi et al., 2014). The amount of perturbation is controlled by ϵ , where smaller ϵ means more privacy. We allow the provider to specify his desired level of privacy, so in a given round t we have ϵ_t -local differential privacy.

Definition 1. [Duchi et al. (2014)] Let $A = (X_1, \dots, X_T)$ be a sensitive dataset where each $X_t \in A$ corresponds to data about individual t . A randomiser R which outputs a disguised version of $S = (U_1, \dots, U_T)$ of A is said to provide ϵ -local differential privacy to individual t , if for all $x, x' \in A$ and for all $S \subseteq S$,

$$\Pr(U_t \in S | X_t = x) \leq \exp(\epsilon) \Pr(U_t \in S | X_t = x').$$

In this chapter we make use of randomisers of the form $R_t(\mathbf{g}_t) = \mathbf{g}_t + \boldsymbol{\xi}_t$, where $\boldsymbol{\xi}_t$ is generated by a zero-mean symmetrical distribution ρ_t . A common choice for ρ_t is $\rho_t(\mathbf{z}) \propto \exp(-\frac{\epsilon_t}{2} \|\mathbf{z}\|)$ (Song et al., 2015). This randomiser is ϵ_t -local differentially private for $\|\mathbf{g}_t\| \leq 1$ (Song et al., 2015, Theorem 1). We use a small variation of this randomiser, which we call the local Laplace randomiser: $\rho_t(\mathbf{z}) \propto \exp(-\sum_{j=1}^d \frac{\tau_{t,j}}{2} |z_j|)$, where $\sum_{j=1}^d \tau_{t,j} = \epsilon_t$, $\tau_{t,j} \geq 0$. The following result shows that the local Laplace randomiser preserves ϵ_t -local differential privacy.

Lemma 4. Suppose $|g_{t,j}| \leq 1$, then the local Laplace randomiser is ϵ_t -local differentially private, where $\epsilon_t = \sum_{j=1}^d \tau_{t,j}$.

The proof follows from applying Theorem 1 of Song et al. (2015) to each dimension and summing the $\tau_{t,j}$. For completeness the proof is provided in Section 3.6. This randomiser is the Laplace randomiser (Dwork and Roth, 2014) applied to each dimension with a possibly different ϵ per dimension. The local Laplace randomiser gives the user more control over the details of the privacy guarantees: with the local Laplace randomiser each dimension j is $\tau_{t,j}$ -local differentially private. This can also lead to lower regret in some cases, of which we give an example in Section 3.4.

3.2.2 Online Convex Optimization with Noisy Subgradients

The analysis of many efficient online learning tools has been influenced by the Online Convex Optimization framework. As mentioned in the introduction, the OCO setting with noisy subgradients proceeds in rounds $t = 1, \dots, T$. In each round t

1. The learner sends $\mathbf{w}_t \in \mathbb{R}^d$ to the provider of the t^{th} subgradient.
2. The provider samples $\boldsymbol{\xi}_t$ from zero-mean and symmetrical ρ_t and computes subgradient $\mathbf{g}_t \in \partial \ell_t(\mathbf{w}_t)$, where $\|\mathbf{g}_t\|_* \leq G$.

3. The provider sends $\tilde{\mathbf{g}}_t = \mathbf{g}_t + \boldsymbol{\xi}_t \in \mathbb{R}^d$ to the learner.

This protocol is a slight adaptation of the protocol of Duchi et al. (2014), where we allow a different ρ_t in each round t instead of using a constant ρ . In each round the provider only sends $\tilde{\mathbf{g}}_t$ to the learner. The learner has no information about ρ_t other than that ρ_t is symmetrical and zero-mean. Also note that ρ_t is allowed to change with each round t , complicating things even further. Since the feedback the learner receives is random we are interested in the expected regret. To bound the expected regret we upper bound the losses by their tangents:

$$\mathbb{E}[\mathcal{R}_T(\mathbf{u})] \leq \mathbb{E}\left[\sum_{t=1}^T \langle \mathbf{w}_t - \mathbf{u}, \mathbf{g}_t \rangle\right] = \mathbb{E}\left[\sum_{t=1}^T \langle \mathbf{w}_t - \mathbf{u}, \tilde{\mathbf{g}}_t \rangle\right], \quad (3.2.1)$$

where the equality holds because of the law of total expectation. The analysis focusses on bounding the r.h.s of (3.2.1), which is a standard approach in OCO. In the following we introduce a recently popularized method to control the regret when $\mathbf{w}_t$ and $\mathbf{u}$ are unbounded.

3.2.3 Reward Regret Duality

In this Section we introduce the main technical workhorse in this chapter: the reward regret duality (McMahan and Orabona, 2014, Theorem 1). Informally, for noiseless $\mathbf{g}_t$, suppose we are able to guarantee $-\sum_{t=1}^T \langle \mathbf{w}_t, \mathbf{g}_t \rangle \geq F_T(-\sum_{t=1}^T \mathbf{g}_t) - c_T$ for a convex F_T and $c_T \in \mathbb{R}$. We will refer to F_T as the potential function. Here, $-\sum_{t=1}^T \langle \mathbf{w}_t, \mathbf{g}_t \rangle$ is seen as the reward. By Fenchel's inequality we have $F_T(-\sum_{t=1}^T \mathbf{g}_t) \geq -F_T^*(\mathbf{u}) - \sum_{t=1}^T \langle \mathbf{u}, \mathbf{g}_t \rangle$, which gives us a bound on the regret after using that $-\sum_{t=1}^T \langle \mathbf{w}_t, \mathbf{g}_t \rangle \geq F_T(-\sum_{t=1}^T \mathbf{g}_t) - c_T$ and reordering the terms. For noisy $\tilde{\mathbf{g}}_t$, the formal result is found in the following lemma (see also Theorem 3 of Jun and Orabona (2019)).

Lemma 5. *If $-\mathbb{E}[\sum_{t=1}^T \langle \mathbf{w}_t, \mathbf{g}_t \rangle] \geq \mathbb{E}[F_T(-\sum_{t=1}^T \tilde{\mathbf{g}}_t) - c_T]$ for some convex function F_T and $c_T \in \mathbb{R}$, then $\mathbb{E}[\mathcal{R}_T(\mathbf{u})] \leq \mathbb{E}[c_T] + F_T^*(\mathbf{u})$.*

Proof. From the definition of Fenchel conjugates we have $\mathbb{E}[F_T(-\sum_{t=1}^T \tilde{\mathbf{g}}_t)] \geq \mathbb{E}[-F_T^*(\mathbf{u}) - \sum_{t=1}^T \langle \mathbf{u}, \tilde{\mathbf{g}}_t \rangle] = -F_T^*(\mathbf{u}) - \sum_{t=1}^T \langle \mathbf{u}, \tilde{\mathbf{g}}_t \rangle$. Using $-\mathbb{E}[\sum_{t=1}^T \langle \mathbf{w}_t, \mathbf{g}_t \rangle] \geq \mathbb{E}[F_T(-\sum_{t=1}^T \tilde{\mathbf{g}}_t) - c_T]$ and reordering the terms completes the proof. $\square$

The difficulty lies in finding a suitable F_T and c_T . For example, we could use gradient descent with learning rate η to find $F_T(-\sum_{t=1}^T \tilde{\mathbf{g}}_t) = \frac{\eta}{2} \|\sum_{t=1}^T \tilde{\mathbf{g}}_t\|_2^2$ and $c_T = \sum_{t=1}^T \frac{\eta}{2} \|\tilde{\mathbf{g}}_t\|_2^2$. However, it would be impossible to tune η optimally

due to the dependence on the unknown $\mathbf{u}$ in $F_T^*(\mathbf{u}) = \frac{1}{2\eta}\|\mathbf{u}\|_2^2$. For noiseless subgradients $\mathbf{g}_t$ (Cutkosky and Orabona, 2018) provide a route to find a suitable F_T , with a constant c_T . Jun and Orabona (2019) extend this idea to noisy subgradients $\tilde{\mathbf{g}}_t$: one needs to find an F_t , F_{t-1} , and $\mathbf{w}_t$ that satisfy $F_{t-1}(\mathbf{x}) - \langle \mathbf{w}_t, \mathbf{g}_t \rangle \geq \mathbb{E}_{\tilde{\mathbf{g}}_t}[F_t(\mathbf{x} - \tilde{\mathbf{g}}_t)]$. By assuming that $-\mathbb{E}[\sum_{s=1}^t \langle \mathbf{w}_s, \mathbf{g}_s \rangle] \geq \mathbb{E}[F_t(-\sum_{s=1}^t \tilde{\mathbf{g}}_s)]$ holds one can show that if F_t and F_{t-1} satisfy $F_{t-1}(\mathbf{x}) - \langle \mathbf{w}_t, \mathbf{g}_t \rangle \geq \mathbb{E}_{\tilde{\mathbf{g}}_t}[F_t(\mathbf{x} - \tilde{\mathbf{g}}_t)]$, then $-\mathbb{E}[\sum_{t=1}^T \langle \mathbf{w}_t, \mathbf{g}_t \rangle] \geq \mathbb{E}[F_T(-\sum_{t=1}^T \tilde{\mathbf{g}}_t)]$ holds by induction. The result is given in the following lemma, of which the proof can be found in Section 3.6.

Lemma 6. *Suppose that $F_{t-1}(\mathbf{x}) - \langle \mathbf{w}_t, \mathbf{g}_t \rangle \geq \mathbb{E}_{\tilde{\mathbf{g}}_t}[F_t(\mathbf{x} - \tilde{\mathbf{g}}_t)]$ holds for all t , then*

$$-\mathbb{E}[\sum_{t=1}^T \langle \mathbf{w}_t, \mathbf{g}_t \rangle] \geq \mathbb{E}[F_T(-\sum_{t=1}^T \tilde{\mathbf{g}}_t)].$$

3.3 One-Dimensional Private Adaptive Potential Function

Algorithm 1 Local Differentially Private Adaptive Potential Function

Input: G such that $|\mathbb{E}[\tilde{g}_t]| \leq G$ and prior P on $v \in [-\frac{1}{5G}, \frac{1}{5G}]$.

- 1: **for** $t = 1, \dots, T$ **do**
 - 2: Play $w_t = \mathbb{E}_{v \sim P}[v \exp(-\sum_{s=1}^{t-1} (v\tilde{g}_s + (v\tilde{g}_s)^2))]$.
 - 3: Receive symmetric $\tilde{g}_t \in \mathbb{R}$.
 - 4: **end for**
-

In this Section we derive a suitable potential function for a one-dimensional problem. In the remainder of this chapter we use this one-dimensional potential to derive new algorithms. To derive our one-dimensional potential function we rely on a property of symmetric random variables with bounded means. The following Lemma is key deriving our potential function F_T .

Lemma 7. *Suppose $\mathbf{x}$ is a symmetrical random variable with $|\mathbb{E}[\langle \mathbf{v}, \mathbf{x} \rangle]| \leq \frac{1}{5}$ for some $\mathbf{v}$. Then $\mathbb{E}[\exp(\langle \mathbf{v}, \mathbf{x} \rangle - \langle \mathbf{v}, \mathbf{x} \rangle^2)] \leq 1 + \mathbb{E}[\langle \mathbf{v}, \mathbf{x} \rangle]$.*

The proof of Lemma 7 can be found in Section 3.7. We can now use Lemma 7 to derive a one-dimensional potential function. Suppose $\tilde{g}_t \in \mathbb{R}$ is a symmetrical random variable with $|\mathbb{E}[\tilde{g}_t]| \leq G$. Then $v\tilde{g}_t$ with $v \in [-\frac{1}{5G}, \frac{1}{5G}]$ satisfies the assumptions in Lemma 7. Multiplying the lower bound of Lemma 7 for $1 - \mathbb{E}[v\tilde{g}_t]$,

3.3. One-Dimensional Private Adaptive Potential Function

for $t = 1, \dots, T$, yields a potential function via Lemma 6. The potential we find is

$$F_t(-\sum_{s=1}^t \tilde{g}_s) = \mathbb{E}_{v \sim P}[\exp(-\sum_{s=1}^t (v\tilde{g}_s + (v\tilde{g}_s)^2)) - 1], \quad (3.3.1)$$

where P is an (improper) prior on $v \in [-\frac{1}{5G}, \frac{1}{5G}]$, the first expectation is over $\tilde{g}_1, \dots, \tilde{g}_t$, and $F_0(0) = 0$. This kind of potential function has been used before by Chernov and Vovk (2010); Koolen and Van Erven (2015); Jun and Orabona (2019). The novelty in this particular potential function is that it allows for the incorporation of the symmetrical noise in the analysis. The $\sum_{s=1}^t (v\tilde{g}_s)^2$ term is unique to our potential function and allows us to derive adaptive regret bounds for unconstrained u . Note that the $c_T = 1$ term has moved inside the definition of F_T . While this does not influence the analysis for proper priors it does influence the analysis for improper priors. The corresponding prediction strategy is given by

$$w_t = \mathbb{E}_{v \sim P}[v \exp(-\sum_{s=1}^{t-1} (v\tilde{g}_s + (v\tilde{g}_s)^2))]. \quad (3.3.2)$$

Algorithm 1 summarizes the strategy. Note that Algorithm 1 does not require any extra parameters compared to the setting with noiseless subgradients.

The following result shows that F_T defined by (3.3.1) and w_t defined by (3.3.2) satisfy our assumptions.

Lemma 8. *Suppose $\tilde{g}_t$ is a symmetrical random variable with $|\mathbb{E}[\tilde{g}_t]| \leq G$. Then F_t defined by (3.3.1) and w_t defined by (3.3.2) satisfy $\mathbb{E}_{\tilde{g}_t}[F_t(-\sum_{s=1}^t \tilde{g}_s)] \leq F_{t-1}(-\sum_{s=1}^{t-1} \tilde{g}_s) - w_t \mathbb{E}[\tilde{g}_t]$.*

The proof follows from an application of Lemma 7 and can be found in Section 3.7. We consider two types of priors. The first type are proper priors that are of the form:

$$\frac{dP(v)}{dv} = \frac{\nu(v) \exp(-bv^2)}{Z}, \quad (3.3.3)$$

Where $b \geq 0$, $\nu : [-\frac{1}{5G}, \frac{1}{5G}] \mapsto \mathbb{R}_+$, and $Z = \int_{-\frac{1}{5G}}^{\frac{1}{5G}} \nu(v) e^{-bv^2} dv$ is a normalizing constant. This captures several priors used in literature, including the conjugate prior $\frac{dP}{dv} = \frac{\exp(-bv^2)}{Z}$ (Koolen and Van Erven, 2015), a variant of the CV prior $\frac{dP}{dv} = \frac{1}{Z|v| \ln(|v|)^2}$ (for $G > \frac{1}{5}$), (Chernov and Vovk, 2010; Koolen and Van Erven, 2015), and the uniform prior on $[-\frac{1}{5G}, \frac{1}{5G}]$ (Jun and Orabona, 2019).

The second type of prior is an improper prior: $\frac{dP}{dv} = \frac{1}{|v|}$. A variant of this prior was previously used by (Koolen and Van Erven, 2015). For all priors we derive a

regret bound by computing an upper bound on the convex conjugate of F_T, F_T^* . For conciseness we only present the regret bound for the conjugate prior in the main text. In Section 3.8 we present the analysis of the regret of the improper prior, for which a slightly different analysis is required compared to the proper priors. The analysis for all priors can be seen as performing a Laplace approximation of the integral over v to show that the prior places sufficient mass in a neighbourhood of the optimal v .

Abbreviating $B_t = b + \sum_{s=1}^{t-1} \tilde{g}_s^2$, $L_t = -\sum_{s=1}^{t-1} \tilde{g}_s$, and $C = \frac{1}{5G}$, the predictions (3.3.2) with the conjugate prior are given by:

$$w_t = \frac{\sqrt{b}L_t \exp\left(\frac{(L_t+2CB_t)^2}{4B_t}\right) \left(\operatorname{erf}\left(\frac{L_t-2CB_t}{2\sqrt{B_t}}\right) - \operatorname{erf}\left(\frac{L_t+2CB_t}{2\sqrt{B_t}}\right)\right)}{\operatorname{erf}(C\sqrt{b}) \exp(C(L_t + CB_t)) 4B_t^{\frac{3}{2}}} + \frac{2\sqrt{B_t}(\exp(2CL_t) - 1)}{\operatorname{erf}(C\sqrt{b}) \exp(C(L_t + CB_t)) 4B_t^{\frac{3}{2}}}. \quad (3.3.4)$$

These w_t can be computed efficiently, but see Koolen and Van Erven (2015) for numerically stable evaluation. With the conjugate prior we find the following result:

Theorem 10. Suppose $\tilde{g}_t$ is a symmetrical random variable with $|\mathbb{E}[\tilde{g}_t]| \leq G$ for all t . The predictions (3.3.4) satisfy:

$$\mathbb{E}[\mathcal{R}_T(u)] \leq 1 + |u| \max \left\{ 11G \left(\ln(|u|11G) - 1 + \ln \left(\frac{\sqrt{5}G\sqrt{\pi}}{4\sqrt{b}} \right) \right), \right. \\ \left. \mathbb{E} \left[\sqrt{8 \left(b + \sum_{t=1}^T \tilde{g}_t^2 \right) \ln(16|u|^2 \left(b + \sum_{t=1}^T \tilde{g}_t^2 \right)^{\frac{3}{2}} \frac{\sqrt{\pi}}{\sqrt{b}} + 1)} \right] \right\}.$$

The proof of Theorem 10 can be found in Section 3.7.1 and follows from computing the Fenchel conjugate of the potential function. For noisy subgradients this is the first bound that is adaptive to the sum of the squares of the noisy subgradients. Compared to the expected regret bound for the improper prior (see Theorem 12 in Section 3.8) this bound has worse constants. However, with the conjugate prior all non-constant terms scale with $|u|$, which is not the case with the improper prior. For all proper priors of the form (3.3.3) a similar regret bound can be computed. This can be seen from Lemma 11 in Section 3.7.1, which shows that the convex conjugate of the potential function for these priors is $O(\mathbb{E}[|u| \sqrt{\sum_{t=1}^T \tilde{g}_t^2 \ln(|u|(T+1))}])$.

Algorithm 2 Black-Box Reduction

Input: G such that $\|\mathbb{E}[\tilde{\mathbf{g}}_t]\|_* \leq G$ and Algorithm $\mathcal{A}_{\mathcal{Z}}$ with domain $\mathcal{Z} = \{\mathbf{z} : \|\mathbf{z}\| \leq 1\}$

- 1: **for** $t = 1, \dots, T$ **do**
- 2: Get $\mathbf{z}_t \in \mathcal{Z}$ from $\mathcal{A}_{\mathcal{Z}}$
- 3: Get $v_t \in \mathbb{R}$ from Algorithm 1
- 4: Play $\mathbf{w}_t = v_t \mathbf{z}_t$, receive symmetrical $\tilde{\mathbf{g}}_t$ such that $\|\mathbb{E}[\tilde{\mathbf{g}}_t]\|_* \leq G$
- 5: Send $\tilde{\mathbf{g}}_t$ to $\mathcal{A}_{\mathcal{Z}}$
- 6: Send $\langle \mathbf{z}_t, \tilde{\mathbf{g}}_t \rangle$ to Algorithm 1
- 7: **end for**

3.3.1 Black-Box Reductions

In this Section we use our potential function in a black-box reduction: we take a constrained noisy OCO algorithm $\mathcal{A}_{\mathcal{Z}}$ and turn it into an unconstrained algorithm using our potential function. The same reduction is used by Cutkosky and Orabona (2018) and Jun and Orabona (2019). The algorithm can be found in Figure 2. The potential function and the OCO algorithm each have their task: the potential function is to learn the norm of $\mathbf{u}$ and the constrained OCO algorithm is to learn the direction of $\mathbf{u}$. In each round t we play $\mathbf{w}_t = v_t \mathbf{z}_t$, where $\mathbf{z}_t \in \mathcal{Z}$, $\mathcal{Z} = \{\mathbf{z} : \|\mathbf{z}\| \leq 1\}$, is the prediction of the OCO algorithm and v_t is the prediction of Algorithm 1. We feed $\tilde{\mathbf{g}}_t$ as feedback to $\mathcal{A}_{\mathcal{Z}}$ and $\langle \mathbf{z}_t, \tilde{\mathbf{g}}_t \rangle$ as feedback to Algorithm 1. Since $\tilde{\mathbf{g}}_t$ is a symmetrical random variable and $\mathbb{E}[\langle \mathbf{z}_t, \tilde{\mathbf{g}}_t \rangle] \leq G$, $\langle \mathbf{z}_t, \tilde{\mathbf{g}}_t \rangle$ satisfies the assumptions in Lemma 7. This allows us to control the regret for learning the norm of $\mathbf{u}$ using Theorem 10.

As outlined by Cutkosky and Orabona (2018) the expected regret of Algorithm 2 decomposes into two parts. The first part of the regret is for learning the norm of $\mathbf{u}$, and is controlled by Algorithm 1. The second part of the regret for learning the direction of $\mathbf{u}$ and is controlled by $\mathcal{A}_{\mathcal{Z}}$. The proof is given by Cutkosky and Orabona (2018), but for completeness we provide the proof in Section 3.7.2.

Lemma 9. Suppose $\tilde{\mathbf{g}}_t$ is a symmetrical random variable with $\|\mathbb{E}[\tilde{\mathbf{g}}_t]\|_* \leq G$ for all t . Let $\mathcal{R}_T^{\mathcal{V}}(\|\mathbf{u}\|) = \mathbb{E}[\sum_{t=1}^T (v_t - \|\mathbf{u}\|) \langle \mathbf{z}_t, \tilde{\mathbf{g}}_t \rangle]$ be the regret for learning $\|\mathbf{u}\|$ by Algorithm 1 and let $\mathcal{R}_T^{\mathcal{Z}}(\frac{\mathbf{u}}{\|\mathbf{u}\|}) = \mathbb{E}[\sum_{t=1}^T \langle \mathbf{z}_t - \frac{\mathbf{u}}{\|\mathbf{u}\|}, \tilde{\mathbf{g}}_t \rangle]$ be the regret for learning $\frac{\mathbf{u}}{\|\mathbf{u}\|}$ by $\mathcal{A}_{\mathcal{Z}}$. Then Algorithm 2 satisfies $\mathbb{E}[\mathcal{R}_T(\mathbf{u})] = \mathcal{R}_T^{\mathcal{V}}(\|\mathbf{u}\|) + \|\mathbf{u}\| \mathcal{R}_T^{\mathcal{Z}}(\frac{\mathbf{u}}{\|\mathbf{u}\|})$.

Orabona and Pál (2018) show that Mirror Descent with learning rates $\eta_t = (\sqrt{\sum_{s=1}^t \|\tilde{\mathbf{g}}_s\|_*^2})^{-1}$ yields $\mathcal{R}_T^{\mathcal{Z}}(\frac{\mathbf{u}}{\|\mathbf{u}\|}) = O(\mathbb{E}[\sqrt{\sum_{t=1}^T \|\tilde{\mathbf{g}}_t\|_*^2}])$. Since Algorithm 1

satisfies $\mathcal{R}_T^{\mathcal{V}}(\|\mathbf{u}\|) = O(\mathbb{E}[\|\mathbf{u}\| \sqrt{\sum_{t=1}^T \|\tilde{\mathbf{g}}_t\|_*^2 \ln(\|\mathbf{u}\| \sum_{t=1}^T \|\tilde{\mathbf{g}}_t\|_*^2 + 1)}])$ the total regret of Algorithm 2 is

$$\mathbb{E}[\mathcal{R}_T(\mathbf{u})] = O \left(\|\mathbf{u}\| \mathbb{E} \left[\sqrt{\sum_{t=1}^T \|\tilde{\mathbf{g}}_t\|_*^2 \ln(\|\mathbf{u}\| \sum_{t=1}^T \|\tilde{\mathbf{g}}_t\|_*^2 + 1)} \right] \right). \quad (3.3.5)$$

This bound matches state of the art bounds for noiseless subgradients and is never worse than the bound of Jun and Orabona (2019) for noisy subgradients, but can be substantially better.

3.3.2 Private Unconstrained Adaptive Sparse Gradient Descent

Algorithm 3 Private Unconstrained Adaptive Sparse Gradient Descent

Input: G such that $|\mathbb{E}[\tilde{g}_{t,j}]|_* \leq G$.

- 1: **for** $t = 1, \dots, T$ **do**
 - 2: Play $\mathbf{w}_t$
 - 3: **for** $j = 1, \dots, d$ **do**
 - 4: Receive symmetrical $\tilde{g}_{t,j}$ such that $|\tilde{g}_{t,j}| \leq G$
 - 5: Send $\tilde{g}_{t,j}$ to the j -th instance of Algorithm 1
 - 6: Receive $v_{t+1,j} \in \mathbb{R}$ from the j -th instance of Algorithm 1 with the conjugate prior
 - 7: Set $\mathbf{w}_{t+1,j} = v_{t+1,j}$
 - 8: **end for**
 - 9: **end for**
-

In this Section we propose a noisy unconstrained OCO algorithm that can exploit sparse subgradients. The algorithm is summarized in Algorithm 3. Algorithm 3 runs a copy of Algorithm 1 with the conjugate prior coordinate-wise. A similar strategy is used by Orabona and Tommasi (2017). This strategy can exploit sparse privacy structures, which, combined with sparse subgradients, may yield low regret (see Section 3.4). Its expected regret bound is given below. The proof follows from applying Theorem 10 per dimension.

Theorem 11. Suppose $\tilde{g}_{t,j}$ is a symmetric random variable with $|\mathbb{E}[\tilde{g}_{t,j}]| \leq G$ for

all t and j . Then the expected regret of Algorithm 3 satisfies

$$\mathbb{E}[\mathcal{R}_T(u)] \leq d + \sum_{j=1}^d |u_j| \max \left\{ 11G \left(\ln(|u_j|11G) - 1 + \ln \left(\frac{\sqrt{5}G\sqrt{\pi}}{4\sqrt{b_j}} \right) \right), \right. \\ \left. \mathbb{E} \left[\sqrt{8 \left(b_j + \sum_{t=1}^T \tilde{g}_{t,j}^2 \right) \ln(16|u_j|^2 \left(b_j + \sum_{t=1}^T \tilde{g}_{t,j}^2 \right)^{\frac{3}{2}} \frac{\sqrt{\pi}}{\sqrt{b_j}} + 1)} \right] \right\}.$$

3.4 Motivating Examples

In this Section we present two scenarios in which our algorithms provide better expected regret guarantees than standard algorithms. The first scenario concerns a case where many providers do not care for their privacy (so they do not perturb the subgradients) and few providers care substantially for their privacy. Suppose that the providers who care for their privacy are $\lceil \ln(T) \rceil$ of the total number of providers T . Suppose that $\|g_t\|_2^2 \leq 1$ and that the providers who care for their privacy use $\rho(z) \propto \exp(-\frac{\epsilon}{2}\|z\|_2)$, then $\mathbb{E}[\|\xi_t\|_2^2] \leq 4 + 4\frac{d^2+d}{\epsilon^2}$ (Song et al., 2015, Theorem 1). Using Algorithm 2, Jensen's inequality, and the fact that the square root is subadditive we see from (3.3.5) that the expected regret is upper bounded by $O(\|u\|_2 \sqrt{\sum_{t=1}^T \|g_t\|_2^2} \ln(1 + \|u\|_2 T) + \|u\|_2 \frac{d}{\epsilon} \ln(\|u\|_2 T + T))$ instead of $O(\|u\|_2 \frac{d}{\epsilon} \sqrt{T \ln(1 + \|u\|_2 T)})$ had we used the maximum privacy guarantee for all providers instead of letting the providers choose their desired level of privacy.

In the second scenario the providers use the local Laplace randomiser. Suppose that g_t is sparse. A standard algorithm that has good performance for sparse g_t is AdaGrad (Duchi et al., 2011). AdaGrad achieves $O(\mathbb{E}[D \sum_{j=1}^d \sqrt{\sum_{t=1}^T \tilde{g}_{t,j}^2}])$ expected regret, where $\max_j |u_j| \leq D$, and D has to be guessed prior to running AdaGrad. Using Jensen's inequality and the fact that the square root is subadditive the expected regret can be upper bounded by $O(D \sum_{j=1}^d (\sqrt{3 \sum_{t=1}^T g_{t,j}^2} + \sqrt{\sum_{t=1}^T 3 \mathbb{E}[\xi_{t,j}^2]}))$. Algorithm 3 achieves $O(\sum_{j=1}^d |u_j| (\sqrt{3 \sum_{t=1}^T g_{t,j}^2} \ln(|u_j|T + 1) + \sqrt{3 \sum_{t=1}^T \mathbb{E}[\xi_{t,j}^2]} \ln(|u_j|T + 1)))$ regret, which can be significantly smaller than the bound of AdaGrad if D is much larger than all u_j or if u is sparse. Furthermore, since we allow the provider of the data to choose $\tau_{t,j}$, the parameter of the Laplace randomiser for dimension j , ξ_t can be sparse as well. While this does not give local differential privacy guarantees for all attributes it does give local differential privacy guarantees for attributes with $\tau_j < \infty$.

3.5 Conclusions

In this chapter, we extended the local differential privacy framework in unconstrained Online Convex Optimization by allowing the provider of the data to choose their privacy guarantees. Standard algorithms do not yield satisfactory regret bounds in this setting, either due to dependence on the unknown parameters of the noise or due to dependence on bounded subgradients. Hence, we proposed two new algorithms that match state of the art regret algorithms in both the noisy and noiseless setting, without requiring knowledge of the noise other than symmetry. Our algorithms do not require parameters other than a bound on the norm of the expectation of the subgradients, which allows the privacy requirements of all providers to be private itself. The new algorithms are a step towards practically useful algorithms with local differential privacy guarantees that have sound theoretical guarantees. Furthermore, our algorithms are the first adaptive unconstrained algorithms in the noisy OCO setting without requiring extra parameters compared to the standard OCO setting, solving two problems left open by Jun and Orabona (2019).

3.6 Details from Section 3.2

Proof. (of Lemma 4) Evaluating and rewriting Definition 1 gives

$$\begin{aligned} \prod_{j=1}^d \frac{\exp(-\frac{\tau_{t,j}}{2} |\tilde{g}_{t,j} - g_{t,j'}|)}{\exp(-\frac{\tau_{t,j}}{2} |\tilde{g}_{t,j} - g_{t,j'}|)} &\leq \prod_{j=1}^d \exp(\frac{\tau_{t,j}}{2} (|g_{t,j}| + |g_{t,j'}|)) \\ &\leq \prod_{j=1}^d \exp(\tau_{t,j}) = \exp(\epsilon_t), \end{aligned}$$

where the first inequality follows from applying the triangle inequality for each j and the second inequality follows from the assumption that $|g_{t,j}| \leq 1$. $\square$

Proof. (of Lemma 6) We will prove the result by induction. In a given round t assume that $-\mathbb{E}[\sum_{s=1}^t \langle \mathbf{w}_s, \mathbf{g}_s \rangle] \geq \mathbb{E}[F_t(-\sum_{s=1}^t \tilde{\mathbf{g}}_s)]$ holds. Now,

$$\begin{aligned} -\mathbb{E}[\sum_{s=1}^{t+1} \langle \mathbf{w}_s, \mathbf{g}_s \rangle] &= \mathbb{E}[-\langle \mathbf{w}_{t+1}, \mathbf{g}_{t+1} \rangle - \sum_{s=1}^t \langle \mathbf{w}_s, \mathbf{g}_s \rangle] \\ &\geq \mathbb{E}[F_t(-\sum_{s=1}^t \tilde{\mathbf{g}}_s) - \langle \mathbf{w}_{t+1}, \mathbf{g}_{t+1} \rangle] \\ &\geq \mathbb{E}[F_{t+1}(-\sum_{s=1}^{t+1} \tilde{\mathbf{g}}_s)], \end{aligned}$$

where the first inequality comes from the inductive hypothesis and the second inequality is by the assumption that $F_{t-1}(\mathbf{x}) - \langle \mathbf{w}_t, \mathbf{g}_t \rangle \geq \mathbb{E}_{\tilde{\mathbf{g}}_t}[F_t(\mathbf{x} - \tilde{\mathbf{g}}_t)]$ for all t . Now, by induction $-\mathbb{E}[\sum_{t=1}^T \langle \mathbf{w}_t, \mathbf{g}_t \rangle] \geq \mathbb{E}[F_T(-\sum_{t=1}^T \tilde{\mathbf{g}}_t)]$. $\square$

3.7 Details from Section 3.3

Proof. (of Lemma 7) We start by rewriting the l.h.s.:

$$\begin{aligned} & \mathbb{E}[\exp(\langle \mathbf{v}, \mathbf{x} \rangle - \langle \mathbf{v}, \mathbf{x} \rangle^2)] \\ &= \mathbb{E}[\exp(y\langle \mathbf{v}, \mathbf{z} \rangle - \langle \mathbf{v}, \mathbf{z} \rangle^2)] \exp(\mathbb{E}[\langle \mathbf{v}, \mathbf{x} \rangle] - \mathbb{E}[\langle \mathbf{v}, \mathbf{x} \rangle]^2). \end{aligned}$$

where $\mathbf{z} = \mathbf{x} - \mathbb{E}[\mathbf{x}]$ and $y = 1 - 2\mathbb{E}[\langle \mathbf{v}, \mathbf{x} \rangle]$. $\mathbf{z}$ is a random variable with mean $\mathbf{0}$ and $|y| \leq 1.4$ due to the restrictions on $\mathbb{E}[\langle \mathbf{v}, \mathbf{x} \rangle]$. By Lemma 10, $\mathbb{E}[\exp(y\langle \mathbf{v}, \mathbf{z} \rangle - \langle \mathbf{v}, \mathbf{z} \rangle^2)] \leq 1$. It remains to show that $\exp(\mathbb{E}[\langle \mathbf{v}, \mathbf{x} \rangle] - \mathbb{E}[\langle \mathbf{v}, \mathbf{x} \rangle]^2) \leq 1 + \mathbb{E}[\langle \mathbf{v}, \mathbf{x} \rangle]$, which holds for $\mathbb{E}[\langle \mathbf{v}, \mathbf{x} \rangle] \geq -\frac{1}{2}$ (Cesa-Bianchi and Lugosi, 2006, Lemma 2.4). $\square$

Lemma 10. *Let $\mathbf{z} \in \mathbb{R}^d$ be a zero-mean symmetrical random variable. Then for $|y| \leq 1.4$ and arbitrary $\mathbf{v} \in \mathbb{R}^d$*

$$\mathbb{E}[\exp(y\langle \mathbf{v}, \mathbf{z} \rangle - \langle \mathbf{v}, \mathbf{z} \rangle^2)] \leq 1.$$

Proof. Due to symmetry of $\mathbf{z}$ we can write

$$\begin{aligned} & \mathbb{E}[\exp(y\langle \mathbf{v}, \mathbf{z} \rangle - \langle \mathbf{v}, \mathbf{z} \rangle^2)] \\ &= \mathbb{E}\left[\frac{1}{2} \exp(-y\langle \mathbf{v}, \mathbf{z} \rangle - \langle \mathbf{v}, \mathbf{z} \rangle^2) + \frac{1}{2} \exp(y\langle \mathbf{v}, \mathbf{z} \rangle - \langle \mathbf{v}, \mathbf{z} \rangle^2)\right]. \end{aligned}$$

We continue by showing that the expression inside the expectation is smaller than 1:

$$\begin{aligned} & \frac{1}{2} \exp(-y\langle \mathbf{v}, \mathbf{z} \rangle - \langle \mathbf{v}, \mathbf{z} \rangle^2) + \frac{1}{2} \exp(y\langle \mathbf{v}, \mathbf{z} \rangle - \langle \mathbf{v}, \mathbf{z} \rangle^2) \leq 1 \\ & \ln(\cosh(y\langle \mathbf{v}, \mathbf{z} \rangle)) - \langle \mathbf{v}, \mathbf{z} \rangle^2 \leq 0. \end{aligned}$$

which holds because for $|y| \leq 1.4$ $f(x) = \ln(\cosh(yx)) - x^2$ is concave and maximized at $x = 0$, which gives $f(0) = 0$. $\square$

Proof. (of Lemma 8) Let $\ell_t(v) = v\tilde{g}_t + (v\tilde{g}_t)^2$

$$\begin{aligned}\mathbb{E}_{\tilde{g}_t}[F_t(-\sum_{s=1}^t \tilde{g}_s)] &= \mathbb{E}_v[\mathbb{E}_{\tilde{g}_t}[\exp(-\ell_t(v) - \sum_{s=1}^{t-1} \ell_t(v)) - 1]] \\ &\leq \mathbb{E}_v[(1 - v \mathbb{E}[\tilde{g}_t]) \exp(-\sum_{s=1}^{t-1} \ell_t(v)) - 1] \\ &= F_{t-1}(-\sum_{s=1}^{t-1} \tilde{g}_s) - w_t \mathbb{E}[\tilde{g}_t]\end{aligned}$$

where the first equality is due to Tonelli's theorem and the inequality is due to Lemma 7, which applies due to the restrictions on v and $\mathbb{E}[\tilde{g}_t]$. Since $F_0(x) = 0$ the proof is complete. $\square$

3.7.1 Regret Analysis for Proper Priors

Proof. (of Theorem 10). By Lemma 5, Lemma 6, and Lemma 8 we only have to compute the convex conjugate of the potential function. We do the analysis for $-\sum_{t=1}^T \tilde{g}_t \geq 0$. The analysis for $-\sum_{t=1}^T \tilde{g}_t \leq 0$ is analogous. We have $-\sum_{t=1}^T w_t \tilde{g}_t \geq F_T(-\sum_{t=1}^T \tilde{g}_t) \geq -1$. Suppose $\sum_{t=1}^T \tilde{g}_t \leq \sqrt{2(\sum_{t=1}^T \tilde{g}_t^2 + b)}$, then $\mathbb{E}[\mathcal{R}_T(u)] = \mathbb{E}[\sum_{t=1}^T w_t \tilde{g}_t - u\tilde{g}_t] \leq \mathbb{E}[\sum_{t=1}^T |u| |\sum_{t=1}^T \tilde{g}_t|] + 1 \leq |u| \mathbb{E}[\sqrt{2(\sum_{t=1}^T \tilde{g}_t^2 + b)}] + 1$, which implies the result.

Now, suppose $\sum_{t=1}^T \tilde{g}_t \geq \sqrt{2(\sum_{t=1}^T \tilde{g}_t^2 + b)}$. For the conjugate prior $\nu([\eta, \mu]) = \eta - \mu$ and $Z \leq \frac{\sqrt{\pi}}{\sqrt{b}}$. In the case where $-\sum_{t=1}^T \tilde{g}_t \leq \frac{2}{5G}(\sum_{t=1}^T \tilde{g}_t^2 + b)$ set $\mu = \frac{-\sum_{t=1}^T \tilde{g}_t}{2(\sum_{t=1}^T \tilde{g}_t^2 + b)}$. Using Lemma 11 we obtain:

$$\begin{aligned}F_T^*(u) &\leq \sqrt{8|u|^2 \left(\sum_{t=1}^T \tilde{g}_t^2 + b\right) \ln(16|u|^2 \left(\sum_{t=1}^T \tilde{g}_t^2 + b\right) \sqrt{\pi} \frac{\sqrt{\sum_{t=1}^T \tilde{g}_t^2 + b}}{\sqrt{b}} + 1)} + 1.\end{aligned}\tag{3.7.1}$$

In the case where $-\sum_{t=1}^T \tilde{g}_t \geq \frac{2}{5G}(\sum_{t=1}^T \tilde{g}_t^2 + b)$ set $\eta = \frac{5-\sqrt{5}}{50G}$ and $\mu = \frac{1}{2}$ to obtain:

$$F_T^*(u) \leq 11G|u|(\ln(|u|11G) - 1 + \ln\left(\frac{\sqrt{5}G\sqrt{\pi}}{4\sqrt{b}}\right)) + 1.\tag{3.7.2}$$

Combining the expectations of (3.7.1) and (3.7.2) completes the proof. $\square$

3.7. Details from Section 3.3

Lemma 11. Suppose $L > \sqrt{2(V+b)}$. Let $F_T(L) = \mathbb{E}_{v \sim P}[\exp(vL - v^2V) - 1]$ with P as in (3.3.3). If $L \leq \frac{2}{5G}(V+b)$ then

$$F_T^*(u) \leq \sqrt{8|u|^2(V+b) \ln(16|u|^2(V+b)S_t([\eta_1, \mu_1]) + 1)} + 1,$$

where $S_t([\eta, \mu]) = \frac{Z}{\nu([\eta, \mu])}$, $\eta_1 = \frac{L}{2(V+b)} - \frac{1}{\sqrt{2(V+b)}}$, $|\mu_1| \in [\eta_1, \frac{1}{5G}]$ such that $\mu_1 \leq \frac{L}{2(V+b)}$, and $\nu([\eta, \mu]) = \int_{\eta}^{\mu} \nu(v)dv$. If $L \geq \frac{2}{5G}(V+b)$ then

$$F_T^*(u) \leq \frac{|u|}{\eta - \eta^2 \frac{5}{2}G} \left(\ln \left(\frac{|u|}{\eta_2 - \eta_2^2 \frac{5}{2}G} \right) - 1 + \ln(S_T([\eta_2, \mu_2])) \right) + 1,$$

where $[\eta_2, \mu_2] \subseteq [-\frac{1}{5G}, \frac{1}{5G}]$ such that $\mu_2 \leq \frac{L}{2(V+b)}$.

Proof. The initial part of the analysis is parallel to the analysis of Theorem 3 by Koolen and Van Erven (2015). Denote by $B = V + b$. For $v \leq \hat{\eta} = \frac{L}{2B}$, $vL - v^2B$ is non-decreasing in v . Therefore, for $[\eta, \mu] \subseteq [-\frac{1}{5G}, \frac{1}{5G}]$ such that $\mu \leq \hat{\eta}$:

$$\begin{aligned} F_T(-\sum_{t=1}^T x_t) &= \frac{1}{Z} \int_{-\frac{1}{5G}}^{\frac{1}{5G}} \nu(v) \exp(vL - v^2B) dv - 1 \\ &\geq \frac{1}{Z} \nu([\eta, \mu]) \exp(\eta L - \eta^2 B) - 1, \end{aligned}$$

where $\nu([\eta, \mu]) = \int_{\eta}^{\mu} \nu(v)dv$. First suppose that $\hat{\eta} \leq \frac{1}{5G}$. Take $\eta = \hat{\eta} - \frac{1}{\sqrt{2B}}$, which yields

$$F_T(L) \geq \frac{\nu([\eta, \mu])}{Z} \exp\left(\frac{L^2}{4B} - \frac{1}{2}\right) - 1 = g(m(L)) - 1$$

where $g(x) = \exp(x - \frac{1}{2} - \ln(\frac{Z}{\nu([\eta, \mu])}))$ and $m(x) = \frac{x^2}{4B}$. By Hiriart-Urruty (2006, Theorem 2) we have

$$\begin{aligned} F_T^*(u) &\leq (g(m(u)))^* = \inf_{\gamma \geq 0} g^*(\gamma) + \gamma m^*\left(\frac{u}{\gamma}\right) \\ &= \inf_{\gamma \geq 0} \gamma \ln(\gamma) + \gamma \left(\ln\left(\frac{Z}{\nu([\eta, \mu])}\right) - \frac{1}{2} \right) + \frac{1}{\gamma} 4|u|^2 B + 1. \end{aligned} \tag{3.7.3}$$

Denote by $S = \ln(\frac{Z}{\nu([\eta, \mu])})$ and $H = 4|u|^2 B$. Setting the derivative to 0 we find that $\hat{\gamma} = \sqrt{\frac{2H}{W(2H \exp(S^2 + \frac{1}{2}))}}$ minimizes (3.7.3), where W is the Lambert function.

Plugging $\hat{\gamma}$ in (3.7.3) gives

$$F_T^*(u) \leq \frac{H(2W(2H \exp(S + \frac{1}{2})) - 1)}{\sqrt{2H(W(2H \exp(S + \frac{1}{2})))}} + 1 \leq \sqrt{2H(W(2H \exp(S + \frac{1}{2})))} + 1.$$

Using $W(x) \leq \ln(x + 1)$ (Orabona and Pál, 2016, Lemma 17) we obtain

$$F_T^*(u) \leq \sqrt{2H \ln(2H \exp(S + \frac{1}{2})) + 1} \leq \sqrt{8|u|^2 B \ln(16|u|^2 B \exp(S) + 1) + 1}.$$

Now suppose that $\hat{\eta} > \frac{1}{5G}$, which is equivalent to $\frac{5}{2}GL > B$. Then

$$F_T(L) \geq \frac{\nu([\eta, \mu])}{Z} \exp((\eta - \eta^2 \frac{5}{2}G)L) - 1.$$

The convex conjugate of this lower bound is well known and is an upper bound on F_T^* :

$$F_T^*(u) \leq \frac{|u|}{\eta - \eta^2 \frac{5}{2}G} \left(\ln \left(\frac{|u|}{\eta - \eta^2 \frac{5}{2}G} \right) - 1 + \ln \left(\frac{Z}{\nu([\eta, \mu])} \right) \right) + 1,$$

which concludes the proof. $\square$

3.7.2 Details From section 3.3.1

Proof. (of Lemma 9) We have

$$\begin{aligned} \mathbb{E}[\mathcal{R}_u(\mathbf{u})] &= \mathbb{E} \left[\sum_{t=1}^T \langle \mathbf{w}_t - \mathbf{u}, \tilde{\mathbf{g}}_t \rangle \right] \\ &= \mathbb{E} \left[\sum_{t=1}^T \langle \mathbf{z}_t, \tilde{\mathbf{g}}_t \rangle (v_t - \|\mathbf{u}\|) \right] + \|\mathbf{u}\| \mathbb{E} \left[\sum_{t=1}^T \langle \mathbf{z}_t - \frac{\mathbf{u}}{\|\mathbf{u}\|}, \tilde{\mathbf{g}}_t \rangle \right] \\ &= \mathcal{R}_T^{\mathcal{V}}(\|\mathbf{u}\|) + \|\mathbf{u}\| \mathcal{R}_T^{\mathcal{Z}} \left(\frac{\mathbf{u}}{\|\mathbf{u}\|} \right) \end{aligned}$$

$\square$

3.8 Regret Analysis for the Improper Prior

Abbreviating $B_t = \sum_{s=1}^{t-1} \tilde{g}_s^2$, $L_t = -\sum_{s=1}^{t-1} \tilde{g}_s$, and $C = \frac{1}{5G}$, the predictions (3.3.2) with the improper prior are given by:

$$\frac{\sqrt{\pi} \exp(\frac{L^2}{4B}) \left(2 \operatorname{erf} \left(\frac{L}{2\sqrt{B}} \right) - \operatorname{erf} \left(\frac{L+2CB}{2\sqrt{B}} \right) - \operatorname{erf} \left(\frac{L-2CB}{2\sqrt{B}} \right) \right)}{2\sqrt{B}}. \quad (3.8.1)$$

With the predictions in (3.8.1) we can show the following result.

Theorem 12. *Suppose $\tilde{g}_t$ is a symmetrical random variable with $|\mathbb{E}[\tilde{g}_t]| \leq G$ for all t . The the expected regret of algorithm 1 with the improper prior $\frac{dP}{dv} = \frac{1}{|v|}$ satisfies*

$$\begin{aligned} \mathbb{E}[\mathcal{R}_T(u)] \leq \max \left\{ |u| \mathbb{E} \left[\sqrt{8 \sum_{t=1}^T \tilde{g}_t^2} \left(\sqrt{\ln(8|u|^2 \sum_{t=1}^T \tilde{g}_t^2 + 1)} + 1 \right) \right], \right. \\ |u| 11G(\ln(|u| 11G \ln(2)) - 1) + \ln(2), \\ \left. |u| \mathbb{E} \left[\sqrt{2 \sum_{t=1}^T \tilde{g}_t^2} + 1 + \mathbb{E} \left[\ln \left(1 + 2 \sqrt{2 \sum_{t=1}^T \tilde{g}_t^2} \right) \right] \right] \right\}. \end{aligned} \quad (3.8.2)$$

Proof. By Lemma 5, Lemma 6, and Lemma 8 we only have to compute the convex conjugate of the potential function. The initial part of the analysis is parallel to the analysis of Theorem 4 by Koolen and Van Erven (2015). Denote by $L = -\sum_{t=1}^T \tilde{g}_t$ and by $V = \sum_{t=1}^T \tilde{g}_t^2$. We do the analysis for $L \geq 0$. The analysis for $L \leq 0$ is analogous. We start by considering the case where $L \leq \sqrt{2V}$. We have

$$\begin{aligned} F_T(L) &\geq \int_0^\epsilon \frac{1}{v} (\exp(-vL - v^2V) - 1) + \int_\epsilon^{\frac{1}{5G}} \frac{1}{v} (\exp(-vL - v^2V) - 1) \\ &\geq -\epsilon L - \epsilon^2 V + \ln(5G\epsilon), \end{aligned}$$

where we used $\exp(x) \geq 1 + x$. Choosing $\epsilon = \frac{1}{5G+2\sqrt{2V}}$ gives $-\mathbb{E}[\sum_{t=1}^T w_t \tilde{g}_t] \geq \mathbb{E}[F_T(L)] \geq -1 - \mathbb{E}[\ln(1 + 2\sqrt{2V})]$. Now, $\mathbb{E}[\mathcal{R}_T(u)] = \mathbb{E}[\sum_{t=1}^T w_t \tilde{g}_t - u \tilde{g}_t] \leq \mathbb{E}[\sum_{t=1}^T |u| |\tilde{g}_t|] + 1 + \mathbb{E}[\ln(1 + 2\sqrt{2V})] \leq |u| \mathbb{E}[\sqrt{2V}] + 1 + \mathbb{E}[\ln(1 + 2\sqrt{2V})]$.

Now consider the case where $L > \sqrt{2V}$. For $v \leq \hat{\eta} = \frac{L}{2V}$, $vL - v^2V$ is non-decreasing in v . Therefore, for $[\eta, \mu] \subseteq [0, \frac{1}{5G}]$ such that $\mu \leq \hat{\eta}$, we have:

$$\begin{aligned} F_T(L) &= \int_{-\frac{1}{5G}}^{\frac{1}{5G}} \frac{1}{|v|} (\exp(vL - v^2V) - 1) dv \\ &\geq (\exp(\eta L - \eta^2 V) - 1) \int_\eta^\mu \frac{1}{v} dv - \int_\mu^{\frac{1}{5G}} \frac{1}{v} dv \\ &= (\exp(\eta L - \eta^2 V) - 1) \ln\left(\frac{\mu}{\eta}\right) + \ln(5G\mu). \end{aligned}$$

First, suppose that $\hat{\eta} \leq \frac{1}{5G}$. Set $\mu = \hat{\eta}$ and $\eta = \hat{\eta} - \frac{1}{\sqrt{2V}}$ and use $L \geq 2\sqrt{V}$ to obtain

$$\begin{aligned} F_T(L) &\geq \exp\left(\frac{L^2}{4V} - \frac{1}{2}\right) \ln\left(\frac{1}{1 - \frac{\sqrt{2V}}{L}}\right) + \ln\left(\frac{L}{V}\right) \\ &\geq \exp\left(\frac{L^2}{4V} - \frac{1}{2}\right) \ln\left(\frac{1}{1 - \frac{\sqrt{2V}}{L}}\right) - \frac{1}{2} \ln\left(\frac{V}{4}\right) \\ &\geq \exp\left(\frac{1}{2} \left(\frac{L}{\sqrt{2V}} - 1\right)^2\right) - 1, \end{aligned}$$

where the last inequality follows by using $\exp\left(\frac{1}{2}(x^2 - 1)\right) \geq \exp\left(\frac{1}{2}(x - 1)^2\right) x$, $-1 \geq -\frac{L}{\sqrt{2V}}$, and $-\ln(1 - x) \geq x$. Write $\exp\left(\frac{1}{2} \left(\frac{L}{\sqrt{2V}} - 1\right)^2\right) - 1 = g(m(x))$, where $g(x) = \exp(x) - 1$ and $m(x) = \left(\frac{x}{\sqrt{2V}} - 1\right)^2$. By Hiriart-Urruty (2006, Theorem 2) we have

$$\begin{aligned} F_T^*(u) &\leq (g(m(u)))^* = \inf_{\gamma \geq 0} g^*(\gamma) + \gamma m^*\left(\frac{u}{\gamma}\right) \\ &= \inf_{\gamma \geq 0} \gamma \ln(\gamma) - \gamma + \frac{1}{\gamma} 4|u|^2 V + 2|u| \sqrt{2V}. \end{aligned} \quad (3.8.3)$$

Setting the derivative to 0 we find that $\hat{\gamma} = \exp\left(\frac{1}{2} W(8|u|^2 V)\right)$ minimizes (3.8.3), where W is the Lambert function. Plugging $\hat{\gamma}$ in (3.8.3) gives

$$F_T^*(u) \leq |u| \sqrt{8V W(8|u|^2 V)} - \hat{\gamma} + 2|u| \sqrt{2V}.$$

Using $W(x) \leq \ln(x + 1)$ (Orabona and Pál, 2016, Lemma 17) and dropping the negative term we obtain

$$F_T^*(u) \leq |u| \sqrt{8V} \left(\sqrt{\ln(8|u|^2 V + 1)} + 1 \right).$$

Now suppose that $\hat{\eta} > \frac{1}{5G}$. Using that $\frac{5G}{2} L \geq V$, choosing $\mu = \frac{1}{5G}$, and $\eta = \frac{5 - \sqrt{5}}{50G}$ we obtain

$$\begin{aligned} F_T(L) &\geq \left(\exp\left(\frac{2(\sqrt{5} - 1)}{25G}\right) L - 1 \right) \ln\left(\frac{1}{1 - \frac{1}{\sqrt{5}}}\right) \\ &\geq \left(\exp\left(\frac{1}{11G}\right) L - 1 \right) \ln(2). \end{aligned} \quad (3.8.4)$$

The convex conjugate of the last expression in (3.8.4) is well known and given by

$$F_T^*(u) \leq |u| 11G (\ln(|u| 11G \ln(2)) - 1) + \ln(2).$$

Combining the above completes the proof. $\square$