



Differential Privacy for Eye-Tracking Data

Ao Liu
Rensselaer Polytechnic Institute
Troy, New York
liua6@rpi.edu

Lirong Xia
Rensselaer Polytechnic Institute
Troy, New York
xial@cs.rpi.edu

Andrew Duchowski
Clemson University
Clemson, South Carolina
duchowski@clemson.edu

Reynold Bailey
Rochester Institute of Technology
Rochester, New York
rjb@cs.rit.edu

Kenneth Holmqvist
University of Regensburg
Regensburg, Germany
Kenneth.Holmqvist@psychologie.uni-regensburg.de

Eakta Jain
University of Florida
Gainesville, Florida
ejain@cise.ufl.edu

ABSTRACT

As large eye-tracking datasets are created, data privacy is a pressing concern for the eye-tracking community. De-identifying data does not guarantee privacy because multiple datasets can be linked for inferences. A common belief is that aggregating individuals' data into composite representations such as heatmaps protects the individual. However, we analytically examine the privacy of (noise-free) heatmaps and show that they do not guarantee privacy. We further propose two noise mechanisms that guarantee privacy and analyze their privacy-utility tradeoff. Analysis reveals that our Gaussian noise mechanism is an elegant solution to preserve privacy for heatmaps. Our results have implications for interdisciplinary research to create differentially private mechanisms for eye tracking.

CCS CONCEPTS

• **Security and privacy** → **Human and societal aspects of security and privacy; Privacy protections.**

KEYWORDS

Eye-tracking, Differential Privacy, Privacy-Utility Tradeoff, Heatmaps

ACM Reference Format:

Ao Liu, Lirong Xia, Andrew Duchowski, Reynold Bailey, Kenneth Holmqvist, and Eakta Jain. 2019. Differential Privacy for Eye-Tracking Data. In *2019 Symposium on Eye Tracking Research and Applications (ETRA '19)*, June 25–28, 2019, Denver, CO, USA. ACM, New York, NY, USA, 10 pages. <https://doi.org/10.1145/3314111.3319823>

1 INTRODUCTION

With advances in mobile and ubiquitous eye tracking, there is ample opportunity to collect eye tracking data at scale. A user's gaze encodes valuable information including attention, intent, emotional state, cognitive ability, and health. This information can be used to gain insight into human behavior (e.g. in marketing and user

experience design), create computational models (e.g. for smart environments and vehicles), and enable interventions (e.g. health and education). When combined with physiological sensing and contextual data, this information facilitates the modeling and prediction of human behavior and decision making. As users become increasingly conscious about what their data reveals about them, there is mounting pressure on policymakers and corporations to introduce robust privacy regulations and processes [gdp 2018; Graham 2018]. The eye tracking community must actively pursue research about privacy for broad public acceptance of this technology.

Data privacy for eye tracking has been raised as a concern in the community [Khamis et al. 2018; Ling et al. 2014]. At a recent Dagstuhl seminar on ubiquitous gaze sensing and interaction¹, privacy considerations were highlighted in a number of papers in the proceedings [Chuang et al. 2018]. Privacy as a general term has a wide range of meanings and different levels of importance for different users. Privacy can obviously be preserved by distorting or randomizing the answers to queries, however doing so renders the information in the dataset useless.

To maintain privacy while preserving the utility of the information, we propose to apply the concept of *differential privacy* (DP) which has been developed by theoretical computer scientists and applied to database applications over the past decade [Dwork 2011]. Differential privacy can be summarized as follows:

Privacy is maintained if an individual's records cannot be accurately identified, even in the worst case when all other data has been exposed by adversaries.

Our technical contributions are: (1) We introduce the notion of differential privacy for eye tracking data. (2) We formally examine the privacy of aggregating eye tracking data as heatmaps and show that aggregating into heatmaps does not guarantee privacy from a DP perspective. (3) We propose two mechanisms to improve the privacy of aggregated gaze data. (4) We analyze the privacy-utility trade-off of these mechanisms from a DP-point of view.

From a practical perspective, the notion of differential privacy is both achievable and theoretically verifiable. Though the proofs may be mathematically sophisticated, the implementation is straightforward, and can be integrated into the eye tracking data collection pipeline. Figure 1 illustrates how this may be achieved. Privacy is guaranteed for the worst case when an adversary has already gained access to the data of all other individuals in a dataset (by

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.
ETRA '19, June 25–28, 2019, Denver, CO, USA

© 2019 Copyright held by the owner/author(s). Publication rights licensed to ACM.
ACM ISBN 978-1-4503-6709-7/19/06...\$15.00
<https://doi.org/10.1145/3314111.3319823>

¹<https://www.dagstuhl.de/18252>

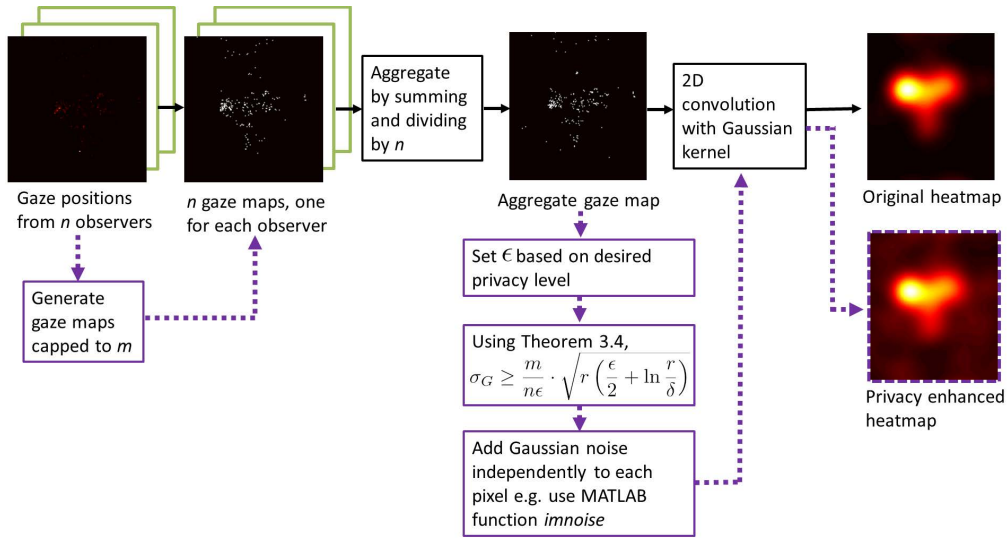


Figure 1: Workflow for researchers and practitioners to create the desired strength of privacy level. The solid lines illustrate the standard workflow for generating an aggregate static heatmap from eye tracking data. The dotted lines show how to implement a privacy protocol with small modifications to this workflow. The hotspots on the privacy enhanced heatmap are visually in the same locations as the original heatmap. The supplementary materials show several examples of privacy enhanced heatmaps for the same noise level.

hacking them for example). Even in this case, the adversary will still not be able to accurately infer data records of the individual. In applying the general definition of differential privacy to eye tracking, we acknowledge that individual users, service providers, and policy makers may have different positions on what level of privacy versus utility is desirable. Our work provides a theoretically grounded analysis of privacy preserving mechanisms to empower these stakeholders to make such decisions.

Implications. Table 1 presents some of the threats that may be posed if an adversary was to access eye tracking data with no privacy protocol in place. Specifically, we elaborate three scenarios

where eye tracking data is collected with good intentions, but if hacked, could have consequences for the individuals concerned.

Scenario 1: A hospital or doctor's office collects eye tracking data as part of patients' general examination. A research grant enables a team to use this data to build a machine learning model that can predict whether someone has a certain neurological disorder. A hacker gains unauthorized access to this database and is able to identify specific individuals with the disorder. The hacker then sells or publicly releases the identity of these individuals, negatively impacting their employment opportunities, inflating their health insurance costs, and elevating their social and emotional anxiety.

Table 1: In most cases, eye tracking data is released with the stimuli. This table illustrates the threats posed by releasing this data if no privacy protocol is in place.

Type of data	Example of intended use	What adversary can access in worst case	What adversary can do now	Does DP apply?
Raw eye movements	Foveated rendering	Raw eye movements	Neurological diagnoses (see Scenario 1)	yes, future work
Aggregated data without temporal information (static heatmaps)	Marketing, UX design, education	Individual's heatmap	Behavioral diagnoses (see Scenario 2)	yes, this paper
Aggregated data with temporal information (dynamic heatmaps)	Training models for autonomous vehicles	Individual's heatmap	Establish driver's liability (see Scenario 3)	yes, future work
Areas of Interest (AOI) analysis	Expert vs novice analysis	Individual's AOI visit order	Autism spectrum diagnoses	yes, future work

Scenario 2: A parent signs a consent form allowing her child to be eye tracked in a classroom. The consent form says that this data is for a research project to understand and characterize learning disabilities and build interventions. The anonymized dataset will be released as part of an NIH big data initiative. If an adversary manages to access an individual child's data and analyze it for markers of dyslexia (for example), they may sell the information to a marketing company that will contact the parent with unsolicited advertising for therapies.

Scenario 3: A publicly funded research team is using eye tracking to study awareness and fatigue of commercial truck drivers. The eye movement data along with the scene being viewed is streamed to a remote server for later analysis. A driver in the study was involved in an accident that resulted in a fatality. Although drivers were told their data would be de-identified, a private investigator, hired by the family of the deceased, was able to extract his/her data record from the database, revealing evidence that (s)he was at fault in the accident.

In scenarios such as these, research teams may reassure participants that raw data will not be released, or that individual data will be de-identified or aggregated (often in the form of heatmaps), providing the impression that privacy is preserved.

2 BACKGROUND

2.0.1 The problem with de-identification. The first “solution” that occurs to many of us is to simply anonymize, or de-identify the dataset. This operation refers to removing personal identifiers such as the name of the participant from the dataset. The problem with this approach is that it is not future-proof; as newer datasets are released, multiple datasets can be linked, and the identity of a participant can then be inferred [Holland and Komogortsev 2011; Komogortsev et al. 2010; Nissim et al. 2017; Ohm 2009].

2.0.2 The problem with running queries. A second “solution” would be to not release the dataset as is, rather allow the analyst to query the dataset. The dataset would not allow queries on individual items, but only on large numbers of items. In other words, a query such as “Where did the student with the lowest grade look?” would be disallowed. But then, the analyst can run queries such as “Where did the students who did not have the lowest grade look?”, and “Where did all the students look?”, and use these queries to infer the disallowed query. This “solution” is not able to guarantee privacy in the worst case, for example, if the adversary hacks the data of $n - 1$ out of n persons in the dataset. Then (s)he can easily infer the n th person's data by querying the average or sum of the dataset.

These issues are well known in database research. One widely accepted formal definition of privacy that has emerged from this extensive research is as follows: an individual's privacy is preserved if the inferences that are made from the dataset do not indicate in any significant way whether this individual is part of the dataset or not. This notion is called *differential privacy*.

2.0.3 Differential privacy. Differential privacy as a concept was conceived through insights by theoretical computer scientists aiming to formalize the notion of privacy that was practically achievable as well as theoretically verifiable [Dwork 2011]. A survey of differential privacy in different fields is presented by Dwork [2008]. Relevant to eye tracking are the works that have applied differential privacy definitions to machine learning [Abadi et al. 2016; Ji et al.

2014] and time-series analysis [Fan and Xiong 2014; Rastogi and Nath 2010]. From a societal impact perspective, the eye tracking industry has as much to gain from these ideas.

2.0.4 Mathematical definition of differential privacy. Formally, given datasets D and D' that differ in at most one entry, let $\mathcal{M}$ denote a randomized mechanism that outputs a query of a database with some probability. Then, let $\mathcal{S}$ denote a subset of query outcomes (called an “event”). Then, we say the mechanism $\mathcal{M}$ is ϵ -differentially private (or ϵ -DP in short) if for any $\mathcal{S}$, D and D' ,

$$\Pr[\mathcal{M}(D) \in \mathcal{S}] \leq e^\epsilon \Pr[\mathcal{M}(D') \in \mathcal{S}], \quad (1)$$

In the above inequality, the probability comes from the randomness of mechanism $\mathcal{M}$. Such randomness is necessary as we will see in Section 3.3. We note that this is a worst-case analysis that offers a strong guarantee of privacy, because the inequality must hold for all $\mathcal{S}$, and all neighboring datasets D and D' .

Another more applicable notion of differential privacy is (ϵ, δ) -differential privacy, which is a generalization of ϵ -DP. Using the notation above, we say the mechanism $\mathcal{M}$ is (ϵ, δ) -differentially private (or (ϵ, δ) -DP in short) if for any $\mathcal{S}$, D and D' (D and D' differs at most one entry),

$$\Pr[\mathcal{M}(D) \in \mathcal{S}] \leq e^\epsilon \Pr[\mathcal{M}(D') \in \mathcal{S}] + \delta,$$

Typically it is believed that $\delta = \Omega\left(\frac{1}{n}\right)$ means poor privacy [Dwork et al. 2014] because it allows some individuals' data to be fully recovered, where n is the input size. We note that a mechanism can be (ϵ, δ) -DP for multiple combinations of (ϵ, δ) . As a rule of thumb, smaller ϵ 's and δ 's means better privacy, though we must point out that directly comparing different numerical values is not informative, e.g. (0.1, 0.1) and (1, 0) are not comparable.

2.0.5 Toy example. As part of a general wellness dataset D , the heights of five people are collected. The mean value as the average height of the population is released. Here, $\mathcal{S}$ is the set of outputting average height. In this example, an adversary obtains the heights of four of these five persons through hacking. In this way, the adversary has a dataset D' that contains all persons except the fifth. The adversary computes the average height of the dataset D' and finds that it is much lower than the average height of the dataset D . The adversary thus infers that the fifth person must be very tall.² In other words, even though the fifth person was not known by the adversary, and the dataset D was not released (only the average height was released), the fifth person is also compromised because his or her height can be reverse engineered by the adversary. Now, we introduce a mechanism $\mathcal{M}$ that perturbs the average height of the dataset D by a random amount before releasing it. If the level of perturbation is high enough, the adversary will not be able to even infer whether the fifth person is tall or not. Thus the mechanism $\mathcal{M}$ protects the privacy of the fifth person. Of course, if we add too much perturbation (or, output totally at random), the utility of the dataset will be affected because the output average height contains little information and does not reflect the average height of the population. This is the privacy-utility tradeoff (see Section 4).

²The adversary can also compute the exact height of the fifth person.

2.0.6 Privacy in eye tracking. For much of the past two decades, the focus of eye tracking research has been on making eye tracking ubiquitous, and on discovering the breadth of inferences that can be made from this data, especially in the contexts of health [Leigh and Zee 2015] and education [Jarodzka et al. 2017]. Privacy has not been a high priority because of the benefits of identifying pathology and designing personalized interventions. The relevance of privacy to eye tracking data was eloquently discussed by Liebling and Preibusch [2014]. Ling et al. [2014] and Khamis et al. [2018] have also highlighted the need for eye-tracking data. Privacy considerations have been raised both for streaming data, as well as pre-recorded datasets. Despite growing awareness and concern, few solutions have been proposed. Our work provides a technical solution for the privacy of individuals.

2.0.7 Why heatmaps as the first for privacy analysis. Besides *scanpaths*, the *heatmap* is a popular method of visualizing eye movement data [Duchowski 2018]. Heatmaps, or attentional landscapes as introduced by Pomplun et al. [1996] and popularized by Wooding [2002], are used to represent aggregate fixations [Duchowski et al. 2012]. Other similar approaches involve gaze represented as height maps [Elias et al. 1984; van Gisbergen et al. 2007] or Gaussian Mixture Models [Mital et al. 2011]. Heatmaps are generated by accumulating exponentially decaying intensity $I(i, j)$ at pixel coordinates (i, j) relative to a fixation at coordinates (x, y) ,

$$I(i, j) = \exp\left(-((x - i)^2 + (y - j)^2)/(2\sigma^2)\right)$$

where the exponential decay is modeled by the Gaussian point spread function. A GPU-based implementation [Duchowski et al. 2012] is available for real-time visualization. Though heatmaps are very popular as a visualization, AOI analyses and temporal data analysis is key to eye-tracking research. We have focused on static heatmaps as a proof of concept for the applicability of differential privacy (DP) to eye tracking data. Insights from this work will inform future research on privacy in eye tracking.

3 ANALYZING DIFFERENTIAL PRIVACY OF THE PROPOSED PRIVACY-PRESERVING MECHANISMS

In this section, we analyze the differential privacy of four natural random mechanisms. We show two of these mechanisms cannot preserve privacy under the notion of DP. For the other two mechanisms, we provide theoretically guaranteed lower bounds on the noise required for any user-defined privacy level. Because a heatmap is created from aggregation of gaze maps, and because this is a reversible (convolution) process, the privacy of a heatmap is equivalent to that of the aggregated gaze map on which it is based.

3.1 Notations

We use n to denote the number of observers in the database and r to denote the total number of pixels in the gaze maps. For example, an image of resolution 800×600 corresponds to $r = 4.8 \times 10^5$. We introduce an integer $m > 1$ to cap every observer's gaze map. For example, if an observer looked at one pixel more than m times, we only count m in his/her gaze map.³ In Section 4.2, we will discuss

³Think of this as if the gaze map *saturated*.

the privacy-utility trade off and provide an algorithm for finding the “optimal cap”. Let G_i denote the i -th observer's personal gaze map (after applying cap). The aggregated gaze map of all n observers in the database is denoted by $G = \frac{1}{n} \sum_{i=1}^n G_i$. Here, we normalize G by the number of observers in order to compare the noise-level under different setups. To simplify notations, we use $\mathcal{G} = (G_1, \dots, G_n)$ to denote the collection of all observers' gaze maps. Similarly, we use $\mathcal{G}_{-i} = (G_1, \dots, G_{i-1}, G_{i+1}, \dots, G_n)$ to denote the collection of all observers' personal gaze maps except the i -th observer. Then, we will define several gaze-map-aggregation mechanisms as follows:

- $\mathcal{M}_{\text{noise-free}}$: Directly output the aggregated gaze map. Formally, $\mathcal{M}_{\text{noise-free}}(G_1, \dots, G_n) = G = \frac{1}{n} \sum_{i=1}^n G_i$.
- $\mathcal{M}_{\text{rs1}(c)}$: Randomly select cn gaze maps from dataset (without replacement) and calculate aggregated gaze map accordingly. Formally, assuming the selected gaze maps are $G_{j_1}, \dots, G_{j_{cn}}$, $\mathcal{M}_{\text{rs1}(c)}(G_1, \dots, G_n) = G = \frac{1}{cn} \sum_{k=1}^{cn} G_{j_k}$.
- $\mathcal{M}_{\text{rs2}(c)}$: Similar with $\mathcal{M}_{\text{rs1}(c)}$, the only difference is the sampling process is with replacement.
- $\mathcal{M}_{\text{Gaussian}(\sigma_N)}$: Adding Gaussian noise with standard deviation (noise-level) σ_N to all pixels independently. Formally, $\mathcal{M}_{\text{Gaussian}(\sigma_N)}(G_1, \dots, G_n) = G + \epsilon_{\sigma_N}$, where ϵ_{σ_N} is a r -dimensional Gaussian noise term with zero mean and standard deviation σ_N (all dimensions are mutually independent).
- $\mathcal{M}_{\text{Laplacian}(\sigma_L)}$: Similar with $\mathcal{M}_{\text{Gaussian}(\sigma_N)}$, the only difference is Laplacian noise with noise level σ_L is added instead of Gaussian noise.

In short, $\mathcal{M}_{\text{rs1}(c)}$ and $\mathcal{M}_{\text{rs2}(c)}$ inject sampling noise to the output while $\mathcal{M}_{\text{Gaussian}(\sigma_N)}$ and $\mathcal{M}_{\text{Laplacian}(\sigma_L)}$ inject additive noise.

3.2 Defining eye-tracking differential privacy

We start with re-phrasing the definition of (ϵ, δ) -differential privacy to eye tracking data. In the following discussion, we assume that the aggregated gaze map G (or its noisy version) has been publicly released⁴. The goal of our research is to protect observers' personal gaze maps $G_1, \dots, G_n$ by adding appropriate noise to the aggregated gaze map. Using the notation in Section 3.1, we assume that $\mathcal{G}_{-i}$, all gaze maps other than G_i , are known by the adversary. For any set $\mathcal{S}$ of outputting gaze maps, (ϵ, δ) -differential privacy is formally defined as follows.

Definition 3.1 ((ϵ, δ) -DP). For any set of event $\mathcal{S}$, any collection of gaze maps $\mathcal{G}_{-i}$ known by the adversary, we say a mechanism $\mathcal{M}$ is (ϵ, δ) -differentially private if and only if

$$\Pr[\mathcal{M}(G_i^*, \mathcal{G}_{-i}) \in \mathcal{S} \mid \mathcal{G}_{-i}] \leq e^\epsilon \Pr[\mathcal{M}(G_i^{**}, \mathcal{G}_{-i}) \in \mathcal{S} \mid \mathcal{G}_{-i}] + \delta, \quad (2)$$

where G_i^* and G_i^{**} are any gaze maps of the i -th observer.

According to differential privacy literatures [Dwork et al. 2014], there is no hard threshold between good and poor privacy. For the purpose of illustration, we define the following “privacy levels” in the remainder of this paper:

- Poor privacy: $\delta = \Omega(1/n)$.
- Okay privacy: $\epsilon = 3$ and $\delta = n^{-3/2}$.
- Good privacy: $\epsilon = 1$ and $\delta = n^{-3/2}$.

⁴Because DP focuses on worst case scenarios, the adversary also knows all other observers individual gazemaps.

Note “okay privacy” and “good privacy” are two examples we used for implementation. Practitioners can set their values of ϵ and δ according to their requirements (smaller ϵ and δ means better privacy). Note again $\delta = \Omega(1/n)$ is widely acknowledged as poor privacy [Dwork et al. 2014].

3.3 There is no free privacy

We first use $\mathcal{M}_{\text{noise-free}}$ (poor privacy) as an example to connect intuition and the definition of DP. Intuitively, if the adversary has the noiseless aggregated gaze map G and all other observers’ gaze maps $\mathcal{G}_{-i}$, he/she can perfectly recover G_i by calculating $nG - \sum_{j \neq i} G_j = \left(\sum_{j=1}^n G_j\right) - \sum_{j \neq i} G_j = G_i$.

Using Definition 3.1 and letting $G_i^* = G_i \neq G_i^{**}$ and $\mathcal{S} = \{G\}$,

$$\Pr[\mathcal{M}(G_i^*, \mathcal{G}_{-i}) \in \mathcal{S} \mid \mathcal{G}_{-i}] = 1 \text{ and } \Pr[\mathcal{M}(G_i^{**}, \mathcal{G}_{-i}) \in \mathcal{S} \mid \mathcal{G}_{-i}] = 0,$$

because G will not be a possible output if $G_i \neq G_i^*$. Thus, we know δ can’t be less than 1 to make Inequality 2 hold. Considering $\delta = \Omega(1)$ corresponds to poor privacy, we know $\mathcal{M}_{\text{noise-free}}$ has poor privacy in the language of (ϵ, δ) -DP defined in Definition 3.1.

3.4 Random selection gives poor privacy

In Section 3.1, we proposed two versions of random selection mechanisms. The first version ($\mathcal{M}_{\text{rs1}}$) randomly selects cn observers without replacement while the second version ($\mathcal{M}_{\text{rs2}}$) selects cn with replacement.

THEOREM 3.2 (WITHOUT REPLACEMENT). *Mechanism $\mathcal{M}_{\text{rs1}}$ has poor privacy.*

PROOF. We prove $\mathcal{M}_{\text{rs1}}$ ’s privacy by considering the following case: assuming resolution $r = 1^5$, all observers other than the i -th did not look at the only pixel, we have,

$$\begin{aligned} \Pr \left[\mathcal{M}_{\text{rs1}}(G_1, \dots, G_n) = \frac{1}{cn} \mid G_i = \mathbf{1}, \mathcal{G}_{-i} = \mathbf{0} \right] &= c \text{ and} \\ \Pr \left[\mathcal{M}_{\text{rs1}}(G_1, \dots, G_n) = \frac{1}{cn} \mid G_i = \mathbf{0}, \mathcal{G}_{-i} = \mathbf{0} \right] &= 0, \end{aligned}$$

where $\mathcal{G}_{-i} = \mathbf{0}$ means all elements in collection $\mathcal{G}_{-i}$ equals to 0. Thus, we know δ can’t be less than c to make (2) hold. Then, Theorem 3.2 follows because $c = \Omega(1/n)$ ($cn = \Omega(1)$ is the number of observers selected). $\square$

THEOREM 3.3 (WITH REPLACEMENT). *Mechanism $\mathcal{M}_{\text{rs2}}$ has poor privacy.*

Proof of Theorem 3.3 (see Appendix A in Supplementary materials) is similar to the proof of Theorem 3.2.

3.5 Achieving good privacy with random noise

In this section, we show that adding Gaussian or Laplacian noise can give good privacy if the noise level satisfies certain conditions based on user-defined privacy levels.

3.5.1 Gaussian Noise. Gaussian noise is widely used noise in many optical systems. In $\mathcal{M}_{\text{Gaussian}(\sigma_N)}$, we add Gaussian noise with standard deviation σ_N independently to all pixels of the aggregated gaze map. The probability density p_N of outputting aggregated gaze map $G^{(N)}$ is

$$p_N \left(\mathcal{M}_{\text{Gaussian}(\sigma_N)}(G_1, \dots, G_n) = G^{(N)} \right) = \frac{1}{(2\pi\sigma_N)^{r/2}} \cdot \exp \left(-\frac{\|G^{(N)} - G\|_2^2}{2\sigma_N^2} \right), \quad (3)$$

which is a r dimensional Gaussian distribution such that all dimensions are independent. Note all ℓ_2 norm in main paper and appendix represent Frobenius norm of matrices. For simplification, we use $p_N(G^{(N)})$ to represent $p_N(\mathcal{M}_{\text{Gaussian}(\sigma_N)}(G_1, \dots, G_n) = G^{(N)})$ when without ambiguity. The next Theorem shows announcing $G^{(N)}$ ($\mathcal{M}_{\text{Gaussian}(\sigma_N)}$ ’s output) will not give much information to adversary if the noise-level is as required (for any (ϵ, δ) , we can always find noise level σ_N to guarantee (ϵ, δ) -DP).

THEOREM 3.4 (GAUSSIAN NOISE). *For any noise level $\sigma_N \geq \frac{m}{n\epsilon}$, $\sqrt{r \left(\frac{\epsilon}{2} + \ln \frac{r}{\delta} \right)} \cdot \mathcal{M}_{\text{Gaussian}(\sigma_N)}$ is (ϵ, δ) -differentially private.*

Theorem 3.4 basically says we can always find a noise level σ_N to meet any user-defined privacy level (any ϵ and δ).

PROOF. Let G_i^* and G_i^{**} to denote any two possible gaze maps of the i -th observer. To simplify notation, we use $G_{-i} = \frac{1}{n-1} \sum_{j \neq i} G_j$ to denote the aggregated gaze map from observers other than the i -th. If the i -th observer’s gaze map is G_i^* , the probability density of the outputting $p_N(G^{(N)} \mid G_i = G_i^*)$ is

$$p_N(G^{(N)} \mid G_i = G_i^*) = \frac{1}{(2\pi\sigma_N)^{r/2}} \exp \left(-\frac{1}{2\sigma_N^2} \left\| \frac{G_i^*}{n} + \frac{n-1}{n} G_{-i} - G^{(N)} \right\|_2^2 \right),$$

Similarly, if the i -th observer’s gaze map is G_i^{**} , we have,

$$p_N(G^{(N)} \mid G_i = G_i^{**}) = \frac{1}{(2\pi\sigma_N)^{r/2}} \exp \left(-\frac{1}{2\sigma_N^2} \left\| \frac{G_i^{**}}{n} + \frac{n-1}{n} G_{-i} - G^{(N)} \right\|_2^2 \right),$$

For any G_i^* , G_i^{**} and G_{-i} , we have,

$$\begin{aligned} & \frac{p_N(G^{(N)} \mid G_i = G_i^{**})}{p_N(G^{(N)} \mid G_i = G_i^*)} \\ &= \exp \left(\frac{1}{2\sigma_N^2} \cdot \left(\left\| \frac{G_i^*}{n} + \frac{n-1}{n} G_{-i} - G^{(N)} \right\|_2^2 - \left\| \frac{G_i^{**}}{n} + \frac{n-1}{n} G_{-i} - G^{(N)} \right\|_2^2 \right) \right) \\ &\leq \exp \left(\frac{2 \left\| \frac{G_i^*}{n} + \frac{n-1}{n} G_{-i} - G^{(N)} \right\|_2 \cdot \left\| G_i^{**} - G_i^* \right\|_2 + \left\| G_i^{**} - G_i^* \right\|_2^2}{2\sigma_N^2} \right). \end{aligned}$$

Letting $\mu = \frac{G_i^*}{n} + \frac{n-1}{n} G_{-i}$ and considering $\|G_i^{**} - G_i^*\|_2 \leq \frac{m\sqrt{r}}{n}$, we have,

$$\frac{p_N(G^{(N)} \mid G_i = G_i^{**})}{p_N(G^{(N)} \mid G_i = G_i^*)} \leq \exp \left(\frac{\frac{2m\sqrt{r}}{n} \|G^{(N)} - \mu\|_2 + \frac{m^2 r}{n^2}}{2\sigma_N^2} \right).$$

Thus, for any $G^{(N)}$ such that $\|G^{(N)} - \mu\|_2 \leq \frac{n}{m\sqrt{r}} \epsilon \sigma^2 - \frac{m\sqrt{r}}{2n}$, the ϵ requirement of DP is always met. Then, we bound the tail

⁵This case also holds for $r > 1$ because the first pixel already leaked information.

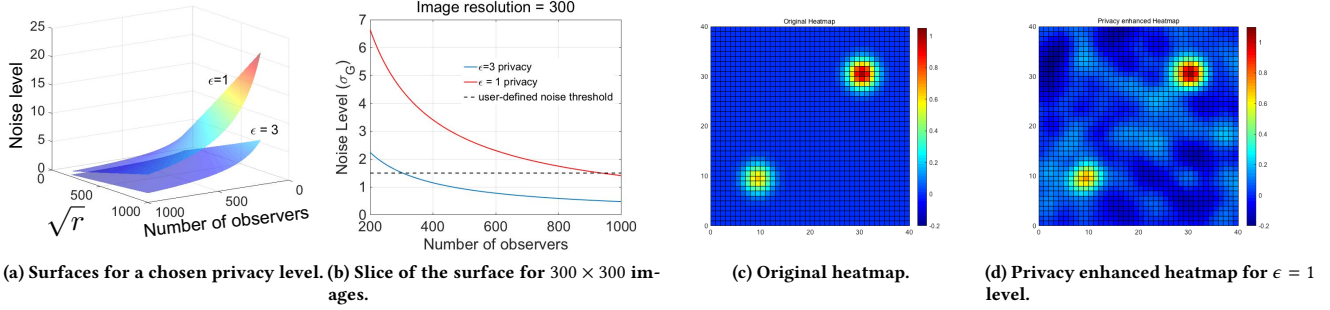


Figure 2: We examine the privacy-utility tradeoff for selected values of σ_N for a simulated heatmap. The greater the noise level we choose to add, the stronger is the privacy guarantee. The relevant stakeholders decide what level of noise is acceptable for a given application. For example, in Figure 2d, the hotspots are still clear, and a UX designer may find this acceptable for the purpose of getting feedback on the design of a website.

probability for all cases where ϵ 's requirement is not met.

$$\begin{aligned} & \Pr \left[\left\| G^{(N)} - \mu \right\|_2 > \frac{n}{m\sqrt{r}} \epsilon \sigma_N^2 - \frac{m\sqrt{r}}{2n} \right] \\ & \leq \sum_{j=1}^r \Pr \left[\left| G_j^{(N)} - \mu_j \right| > \frac{n}{m\sqrt{r}} \epsilon \sigma_N^2 - \frac{m\sqrt{r}}{2n} \right] \\ & \leq r \cdot \exp \left(-\frac{n^2 \epsilon^2 \sigma_N^2}{2m^2 r} + \frac{\epsilon}{4} \right). \end{aligned}$$

When $\sigma_N \geq \frac{m}{n\epsilon} \cdot \sqrt{r \left(\frac{\epsilon}{2} + \ln \frac{r}{\delta} \right)}$, we have,

$$\Pr \left[\left\| G^{(N)} - \mu \right\|_2 > \frac{n}{m\sqrt{r}} \epsilon \sigma_N^2 - \frac{m\sqrt{r}}{2n} \right] \leq r \cdot \exp \left(-\frac{n^2 \epsilon^2 \sigma_N^2}{2m^2 r} + \frac{\epsilon}{4} \right) \leq \delta$$

Then, Theorem 3.4 follows by the definition of (ϵ, δ) -DP. $\square$

3.5.2 Laplacian Noise. Laplacian noise is the most widely used in many differential privacy problems. However, we will show Laplacian noise is not as suitable as Gaussian noise for protecting eye tracking data. The next Theorem shows $G^{(L)}$ will not give much information to the adversary if the noise-level is as required.

THEOREM 3.5 (LAPLACIAN NOISE). *Using the notations above, for any $\sigma_L \geq \frac{\sqrt{2} \cdot m r}{\epsilon n}$, $\mathcal{M}_{\text{Laplacian}(\sigma_L)}$ is $(\epsilon, 0)$ -differentially private.*

Proof of Theorem 3.5 (see Appendix A.2 in supplementary material) is very similar with Theorem 3.4. However, the required noise level, $\sigma_L \geq \frac{\sqrt{2} \cdot m r}{\epsilon n}$, normally is much higher than the requirement of Gaussian noise, $\sigma_N \geq \frac{m}{n\epsilon} \cdot \sqrt{r \left(\frac{\epsilon}{2} + \ln \frac{r}{\delta} \right)}$. One can see the Laplacian mechanism requires one more $\sqrt{r}$ term on noise level, which normally corresponds to $\sim 10^2$ times higher noise level.

4 PRIVACY-UTILITY TRADEOFF

According to Theorem 3.4 and Theorem 3.5, we know better privacy (smaller ϵ and δ) usually requires higher noise level. In this section, we will conduct experiments to show how Gaussian and Laplacian noise influence the utility, i.e., the corresponding heatmap.

4.1 Noise level vs. information loss

In Figure 2(a), we show a three-dimensional plot where the x and y axes are $\sqrt{r}$ and n respectively. The reader may revisit notations in Section 3.1. On the vertical z-axis, we plot σ_N , specifically based on the formula given by Theorem 3.4. The upper surface shows σ_N for good privacy ($\epsilon = 1$ and $\delta = n^{-3/2}$). The lower surface shows for okay privacy ($\epsilon = 3$ and $\delta = n^{-3/2}$). Any value of σ_N above this surface will provide okay privacy, and any value above the upper surface will provide good privacy.

In Figure 2(a), as the image resolution increases, a larger number of observers is needed in the dataset to maintain the guarantee of good privacy. If there is a small number of observers, good privacy can be achieved by downsampling the image. In Figure 2(b) we show a slice of this surface at $\sqrt{r} = 300$. The dotted lines show an example noise level that we could have set based on what we find acceptable for utility. This is of course user-defined, and will vary depending on the application. The graphs illustrate that at a selected noise level, e.g., $\sigma_N = 1.5$, we can achieve good privacy for a 300×300 image if we have of the order of $n = 900$ observers. For a dataset that has $n = 300$ observers, we can tell the participants that we can achieve Okay privacy. We show two simulated heatmaps in Figure 2(c) and (d). The location of the hotspots is unchanged for all practical purposes in the noisy but private heatmap.

We quantify the privacy-utility tradeoff in Figure 3. 100 noisy heatmaps are generated using the workflow in Figure 1. Real-world 1050×1680 gaze maps from five observers looking at a portrait of a woman are used here.⁶ The original heatmap is shown in Figure 1 to the right. For the purpose of the noisy heatmap, we assume the number of observers in dataset is 50,000⁷ (the noise is added according to $n = 50,000$ and Theorem 3.4 and Theorem 3.5). We simulate this large number of observers by replicating each of the five real observers 10000 times.

In the supplementary materials, we show the original heatmap overlaid on the stimulus image in high resolution (original.png). We

⁶Gaze data from dataset of Raiturkar et al. [2018], stimulus image from Farid and Bravo [2012] and Mader et al. [2017].

⁷If the number of observer is much smaller than 50,000, the practitioner could either down-sample gaze maps or sacrifice privacy (setting larger ϵ and δ) to get an acceptable noise level.

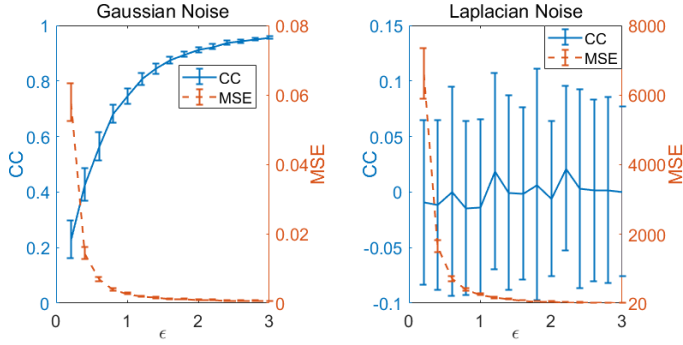


Figure 3: Similarity between the privacy enhanced heatmap and original heatmap when ϵ is varied. The smaller the value of ϵ the stronger is the privacy guarantee from the DP perspective. This graph illustrates the privacy-utility tradeoff: as ϵ is made smaller, the mean squared error increases and the cross-correlation decreases. The Laplacian mechanism results in lesser similarity than the Gaussian mechanism.

also show examples of privacy enhanced heatmaps for this original heatmap at the $\epsilon = 1.5$ privacy level (privacynhanced.mpg). For this image resolution, $\sigma_G = 0.0986$ based on Theorem 3.5.

We numerically analyzed *correlation coefficient* (CC) and *mean square error* (MSE) of noisy heatmaps under different privacy levels (different values of ϵ while fixing $\delta = n^{-3/2}$). The cap $m = 1$ is decided according to Algorithm 1 (see Section 4.2 for details). 100 noisy heatmaps are generated under each setting. The average CC and MSE of those generated noisy heatmaps are plotted in Figure 3. Error bars in Figure 3 represent the standard deviations.

It can be seen from the Figure 3 that Laplacian mechanism results in much more information loss than Gaussian mechanism to achieve same level of privacy under our setting. For both Gaussian and Laplacian mechanisms, one can see that better privacy (smaller ϵ) usually means more information loss in the outputting heatmap.

We note that these graphs are based on real data of only five observers on one stimulus image. This graph is an example of how a practitioner may visualize the privacy-utility tradeoff in any given application domain. In practice, stakeholders would use our proposed workflow on their dataset to prepare such visualizations for different settings of the internal parameters (m, ϵ, δ) to help them evaluate the privacy-utility tradeoff. We note also that Theorem 3.4 is specific to aggregate heatmaps. For any other mechanism, the appropriate theorems would need to be worked out and the workflow modified to be consistent with the problem definition. We also point out that while mean squared error and cross-correlation are readily computed, they do not fully reflect the information lost or retained when noise is added. As an illustration, in Figure 2, the hotspots in the privacy enhanced heatmap are still clear, and a UX designer may find that the heatmap acceptable for their use case even though the MSE and CC metrics suggest otherwise.

4.2 Computing the optimal “cap”

In order to achieve better privacy with less information loss, we set a cap on the maximum number of times an observer’s gaze position

falls on a pixel. This cap was denoted by m in Section 3. Here we discuss the information loss on different settings on m .

When m is larger, higher noise level is required to get the same privacy (both the upper bound for σ_N and σ_L are proportional to m). However, larger m also corresponds to less information loss on every observer’s gaze map. In other words, there is tradeoff between variance (noise) and bias (cap) on cap. Let $G^{(N, m, \sigma^*)}$ to denote the gaze map outputted by Gaussian mechanism with cap m and noise level $\sigma_N = m\sigma^*$. Thus, $G^{(N, \infty, 0)}$ denotes the original aggregated gaze map and $G^{(N, m, 0)}$ denotes the aggregated gaze map with cap m without adding any Gaussian noise. Algorithm 1 provides an implementable way to choose the best value of m to optimize *mean square error* (MSE). In the next theorem, we will

Algorithm 1: Utility optimization algorithm on choosing m

- 1 **Input:** Individual gaze maps $G_1, \dots, G_n$ and noise factor σ^* .
 - 2 **Initialization:** Calculate $G^{(N, \infty, 0)}$ and the maximum number of times one observer look at one pixel: $g_{\max} = \max_{i,j} [G_i(j)]$.
 - 3 **for** $m = 1, \dots, g_{\max}$ **do**
 - 4 Calculate $\mathbb{E} [\text{MSE} (G^{(N, m, \sigma^*)})]$ according to Theorem 4.1.
 - 5 **end**
 - 6 **Output:** m value with the smallest expected MSE.
-

analytically analyze the expectation of MSE.

THEOREM 4.1. *The expected MSE of Gaussian mechanism with cap m and noise level $\sigma_N = m\sigma^*$ is*

$$\mathbb{E} [\text{MSE} (G^{(N, m, \sigma^*)})] = m^2 \sigma^{*2} + \frac{1}{r} \sum_{j=1}^r (G^{(N, m, 0)}(j) - G^{(N, \infty, 0)}(j))^2.$$

PROOF. By the definition of MSE, we have,

$$\mathbb{E} [\text{MSE} (G^{(N, m, \sigma^*)})] = \frac{1}{r} \sum_{j=1}^r \mathbb{E} [(G^{(N, m, \sigma^*)}(j) - G^{(N, \infty, 0)}(j))^2]$$

Using the notations defined above, the expected square error on j -th pixel is

$$\begin{aligned} & \mathbb{E} [(G^{(N, m, \sigma^*)}(j) - G^{(N, \infty, 0)}(j))^2] \\ &= \mathbb{E} [(G^{(N, m, \sigma^*)}(j))^2] - 2G^{(N, \infty, 0)}(j) \cdot \mathbb{E} [G^{(N, m, \sigma^*)}(j)] + (G^{(N, \infty, 0)}(j))^2. \end{aligned} \quad (4)$$

because $G_j^{(N, m, \sigma^*)} \sim \mathcal{N}(G_j^{(N, m, 0)}, m^2 \sigma^{*2})$, we have,

$$\begin{aligned} \mathbb{E} [G^{(N, m, \sigma^*)}(j)] &= G^{(N, m, 0)}(j) \quad \text{and} \\ \mathbb{E} [(G^{(N, m, \sigma^*)}(j))^2] &= \sigma^{*2} m^2 + (G^{(N, m, 0)}(j))^2. \end{aligned} \quad (5)$$

Theorem 4.1 follows by combining (4) and (5). $\square$

Then, we analyze the complexity of Algorithm 1 in the next theorem, which says Algorithm 1 is with linear-time complexity.

THEOREM 4.2. *The complexity of Algorithm 1 is $O(g_{\max} \cdot nr)$.*

PROOF. Rewriting the expected MSE in Theorem 4.1, we have,

$$\mathbb{E}[\text{MSE}] = m^2 \sigma^{*2} + \frac{1}{r} \left\| G^{(N, m, 0)} - G^{(N, \infty, 0)} \right\|_2^2,$$

where the ℓ_2 norm still represents Frobenius norm. Since $G^{(N, m, 0)}$ and $G^{(N, \infty, 0)}$ are r -dimensional vectors, the complexity of computing expected MSE for a given m and $G^{(N, m, 0)}$ will be $O(m)$.

Then, we evaluate the complexity of calculating the capped noise-free aggregated gaze map $G^{(N, m, 0)}$. Since we are adding cap to each observer's individual gaze map, we can add cap to every pixel of all observers. Thus, one can see there are nr pixels from n observers in total. Considering the for loop in Algorithm 1 runs $g_{\max}$ times, Theorem 4.2 follows. $\square$

5 IMPLICATIONS

5.0.1 Datasets are growing. In contrast to the previous research paradigm where datasets were collected, archived, and then released, there is a growing trend to crowd-source data collection, via mobile apps for example, so that new data is continually being added to the dataset. With the methods presented, the new data is safe as long as the publicly available dataset is put through the Gaussian noise mechanism. Another way that eye tracking datasets might seek to preserve a user's privacy is by releasing their eye movements, but not *what* they were looking at. With the methods we present, releasing the stimulus image/video that observers look at is safe because even in the worst case an adversary will not be able to guess at what a particular individual looked at.

5.0.2 Why can the generic theorem of differential privacy not be applied to eye tracking? Unlike classical databases, every observer in eye tracking database contributes much richer information (i.e., millions of pixels) than individuals in classical databases. However, the generic theorems in differential privacy do not focus on high-dimensional data. Simply applying union bounds will result in very loose privacy bounds and unacceptable noise levels.

5.0.3 Why are we adding noise when the field is spending so much time and effort removing it? There has been much research in eye tracking to improve the accuracy of eye tracking to maximize the utility and applicability of eye tracking devices for diverse use cases. This work has been directed at sources of noise that are inherent to the process, such as sensor and measurement noise. However, as eye tracking becomes ubiquitous, there is a cost for the individual user whose data is being recorded and for the organizations who are safeguarding and distributing this data. This cost is the sacrifice of privacy of the individual. We do not argue for reversing the technological push towards reliable, accurate eye tracking. Rather, we argue that our objective as a community must expand to include privacy in addition to utility. For those situations where privacy is deemed to be worth protecting, we introduce flexible mechanisms to do so. Noise is added to data in the aggregate, not to any individual's data. Further, the noise function is fully understood, and its parameters are set based on the desired privacy-utility tradeoff. Unlike measurement noise, whose source may not be fully understood, we add noise in a controlled and measured way to achieve a specific objective.

5.0.4 Why should the research community care? This research requires an interdisciplinary approach. The eye tracking community cannot just "leave it to the privacy researchers" because the theoretical guarantees that form the basis of this framework are highly dependent on the particular mechanisms that the data goes through (the functional forms in the equations, the particular thresholds, etc.). These mechanisms have to be developed collaboratively to preserve the utility of the output for eye tracking applications.

5.0.5 Why should the industry care? The push towards ubiquitous eye tracking is being driven by large investments by major industry players. While their applications are highly data-dependent, their customers are increasingly data-sensitive. This paper proposes the first of a class of solutions which pair theoretical analysis from a DP-perspective with a practically implementable workflow for developers. This work opens the door for a responsible industry that can inform their users that while they may eye track the users at very high accuracy and resolution to enable foveated rendering (for example), they would put this data through mechanism A or B before releasing it to the app developers.

6 CONCLUSIONS AND FUTURE DIRECTIONS

We have proposed to apply the notion of differential privacy toward the analysis of privacy of eye-tracking data. We have analyzed the privacy guarantees provided by the mechanisms of random selection, and additive noise (Gaussian and Laplacian noise). The main takeaway from this paper is that adding Gaussian noise will guarantee differential privacy; the noise level should be appropriately selected based on the application. Our focus is on static heatmaps as a sandbox to understand how the definitions of differential privacy apply to eye tracking data. In this sense, this paper is a proof of concept. Eye tracking data is fundamentally temporal in nature, and the privacy loss if an adversary could access saccade velocities and dynamic attention allocation would be much greater than static heatmaps. Future work would systematically consider all the different ways in which eye tracking data is analyzed and stored.

We have considered two noise models (Gaussian and Laplacian noise). Follow up work might consider the privacy-utility trade-off for different noise models like pink noise. For temporal data such as raw eye movements, it may even be relevant to understand which noise models are more realistic. In other words, if the user's virtual avatar was driven by privacy enhanced eye tracking data, it should still appear realistic and natural.

The mechanisms and analyses presented here apply to real-valued data that can be aligned to a grid and capped to a maximum value without loss of utility. Though our focus has been on eye tracking heatmaps, there are other data that fall in this category, for example, gestures on a touchscreen, or readings from a force plate. It would also be interesting to generalize these mechanisms and analyses to other physiological data such as heart rate, galvanic skin response, and even gestures or gait. These data are conceptually similar to eye tracking data in that they carry signatures of the individual's identity and markers of their health and well-being. Furthermore, in physiological domains many data and analyses are temporal in nature. It would be interesting and important to define and analyze differential privacy for temporal data.

ACKNOWLEDGMENTS

This material is based upon work supported by the National Science Foundation under Grant No. 1566481. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the National Science Foundation.

REFERENCES

- (2018). Eu general data protection regulation. <https://eugdpr.org>.
- Abadi, M., Chu, A., Goodfellow, I., McMahan, H. B., Mironov, I., Talwar, K., and Zhang, L. (2016). Deep learning with differential privacy. In *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security*, pages 308–318. ACM.
- Chuang, L., Duchowski, A., Qvarfordt, P., and Weiskopf, D., editors (2018). *Ubiquitous Gaze Sensing and Interaction*, volume 8 of *Dagstuhl Reports*, Schloss Dagstuhl–Leibniz-Zentrum für Informatik, Germany. Dagstuhl Publishing.
- Duchowski, A. T. (2018). Gaze-based interaction: A 30 year retrospective. *Computers & Graphics*, pages –. Special Section on Serious Games and Virtual Environments.
- Duchowski, A. T., Price, M. M., Meyer, M., and Orero, P. (2012). Aggregate gaze visualization with real-time heatmaps. In *Proceedings of the Symposium on Eye Tracking Research and Applications*, pages 13–20. ACM.
- Dwork, C. (2008). Differential privacy: A survey of results. In *International Conference on Theory and Applications of Models of Computation*, pages 1–19. Springer.
- Dwork, C. (2011). A firm foundation for private data analysis. *Communications of ACM*, 54(1):86–95.
- Dwork, C., Roth, A., et al. (2014). The algorithmic foundations of differential privacy. *Foundations and Trends® in Theoretical Computer Science*, 9(3–4):211–407.
- Elias, G., Sherwin, G., and Wise, J. (1984). Eye movements while viewing NTSC format television. Technical report, SMPTE Psychophysics Subcommittee.
- Fan, L. and Xiong, L. (2014). An adaptive approach to real-time aggregate monitoring with differential privacy. *IEEE Transactions on knowledge and data engineering*, 26(9):2094–2106.
- Farid, H. and Bravo, M. J. (2012). Perceptual discrimination of computer generated and photographic faces. *Digital Investigation*, 8(3–4):226–235.
- Graham, J. (2018). Is apple really better about privacy? here’s what we found out. <https://www.usatoday.com/story/tech/talkingtech/2018/04/17/apple-make-simpler-download-your-privacy-data-year/521786002/>.
- Holland, C. and Komogortsev, O. V. (2011). Biometric identification via eye movement scanpaths in reading. In *Biometrics (IJCB), 2011 International Joint Conference on*, pages 1–8. IEEE.
- Jarodzka, H., Holmqvist, K., and Gruber, H. (2017). Eye tracking in educational science: Theoretical frameworks and research agendas. *Journal of Eye Movement Research*, 10(1).
- Ji, Z., Lipton, Z. C., and Elkan, C. (2014). Differential privacy and machine learning: a survey and review. *arXiv preprint arXiv:1412.7584*.
- Khamis, M., Alt, F., and Bulling, A. (2018). The past, present, and future of gaze-enabled handheld mobile devices: Survey and lessons learned. In *Proceedings of the 20th International Conference on Human-Computer Interaction with Mobile Devices and Services (MobileHCI)*, pages 38:1–38:17.
- Komogortsev, O. V., Jayarathna, S., Aragon, C. R., and Mahmoud, M. (2010). Biometric identification via an oculomotor plant mathematical model. In *Proceedings of the 2010 Symposium on Eye-Tracking Research & Applications*, pages 57–60. ACM.
- Leigh, R. J. and Zee, D. S. (2015). *The neurology of eye movements*. Oxford University Press, USA.
- Liebling, D. J. and Preibusch, S. (2014). Privacy considerations for a pervasive eye tracking world. In *Proceedings of the 2014 ACM International Joint Conference on Pervasive and Ubiquitous Computing: Adjunct Publication*, pages 1169–1177. ACM.
- Ling, R., Mardanbeigi, D., and Hansen, D. W. (2014). Synergies between head-mounted displays and headmounted eye tracking: The trajectory of development and its social consequences. In *Living inside social mobile information*, pages 131–156, 2251 Arbor Blvd. Dayton, OH 45439, USA.
- Mader, B., Banks, M. S., and Farid, H. (2017). Identifying computer-generated portraits: The importance of training and incentives. *Perception*, 46(9):1062–1076.
- Mital, P. K., Smith, T. J., Hill, R. L., and Henderson, J. M. (2011). Clustering of Gaze During Dynamic Scene Viewing is Predicted by Motion. *Cognitive Computation*, 3:5–24.
- Nissim, K., Steinke, T., Wood, A., Altman, M., Bembek, A., Bun, M., Gaboardi, M., O’Brien, D. R., and Vadhan, S. (2017). Differential privacy: A primer for a non-technical audience. *Working Group Privacy Tools Sharing Res. Data, Harvard Univ., Boston, MA, USA, Tech. Rep. TR-2017-03*.
- Ohm, P. (2009). Broken promises of privacy: Responding to the surprising failure of anonymization. *Ucla L. Rev.*, 57:1701.
- Pomplun, M., Ritter, H., and Velichkovsky, B. (1996). Disambiguating complex visual information: Towards communication of personal views of a scene. *Perception*, 25(8):931–948.
- Raiturkar, P., Farid, H., and Jain, E. (2018). Identifying computer-generated portraits: an eye tracking study. Technical report, University of Florida.
- Rastogi, V. and Nath, S. (2010). Differentially private aggregation of distributed time-series with transformation and encryption. In *Proceedings of the 2010 ACM SIGMOD International Conference on Management of data*, pages 735–746. ACM.
- van Gisbergen, M. S., van der Most, J., and Aelen, P. (2007). Visual Attention to Online Search Engine Results. Technical report, De Vos & Jansen in cooperation with Checkit. URL: http://www.iprospect.nl/wp-content/themes/iprospect/pdf/checkit/eyetracking_research.pdf (last accessed Dec. 2011).
- Wooding, D. S. (2002). Fixation maps: quantifying eye-movement traces. In *Proceedings of the Symposium on Eye tracking research & Applications (ETRA)*, pages 31–36. ACM.

A MISSING PROOFS FOR THEOREMS

A.1 Proof and discussion for Theorem 3.3

Theorem 3.3 says $\mathcal{M}_{rs2}$ has poor privacy.

PROOF. Considering the same case as the proof of Theorem 3.2, we have,

$$\begin{aligned} & \Pr \left[\mathcal{M}_{rs2}(G_1, \dots, G_n) \geq \frac{1}{cn} \mid G_i = \mathbf{0}, \mathcal{G}_{-i} = 0 \right] = 0 \quad \text{and} \\ & \Pr \left[\mathcal{M}_{rs2}(G_1, \dots, G_n) \geq \frac{1}{cn} \mid G_i = \mathbf{1}, \mathcal{G}_{-i} = 0 \right] \\ &= 1 - \left(1 - \frac{1}{n} \right)^{cn} \approx 1 - e^{-c} = \Theta(c) \end{aligned}$$

Thus, we know δ can't be less than c to make Inequality 2 hold and Theorem 3.3 follows. $\square$

A.2 Proof for Theorem 3.5

PROOF. The probability density function p_L of output $G^{(L)} = (G^{(L)}(1), \dots, G^{(L)}(r))$ is

$$\begin{aligned} & p_L \left(\mathcal{M}_{\text{Laplacian}(\sigma_L)}(G_1, \dots, G_n) = G^{(L)} \right) \\ &= \frac{1}{(\sqrt{2}\sigma_L)^r} \cdot \exp \left(-\frac{\sqrt{2}}{\sigma_L} \left\| G - G^{(L)} \right\|_1 \right). \end{aligned}$$

To simplify notations, we use $p_L \left(G^{(L)} \right)$ to represent $p_L \left(\mathcal{M}_{\text{Laplacian}(\sigma_L)}(G_1, \dots, G_n) = G^{(L)} \right)$ when without ambiguity. Let G_i^* and G_i^{**} to denote any two possible gaze maps of the i -th

observer. If the i -th observer's gaze map is G_i^* , the probability density function of the outputting $p_L(G^{(L)} \mid G_i = G_i^*)$ is

$$p_L(G^{(L)} \mid G_i = G_i^*) = \frac{1}{(\sqrt{2}\sigma_L)^r} \exp \left(-\frac{\sqrt{2}}{\sigma_L} \left\| \frac{G_i^*}{n} + \frac{n-1}{n} G_{-i} - G^{(L)} \right\|_1 \right),$$

where we abused notation to let $G_{-i} = \frac{1}{n-1} \sum_{j \neq i} G_j$, which is the aggregated gaze map except G_i . Similarly, if the i -th observer's gaze map is G_i^{**} , we have,

$$p_L(G^{(L)} \mid G_i = G_i^{**}) = \frac{1}{(\sqrt{2}\sigma_L)^r} \exp \left(-\frac{\sqrt{2}}{\sigma_L} \left\| \frac{G_i^{**}}{n} + \frac{n-1}{n} G_{-i} - G^{(L)} \right\|_1 \right).$$

For any G_i^* , G_i^{**} and G_{-i} , we have,

$$\begin{aligned} & \frac{p_L(G^{(L)} \mid G_i = G_i^{**})}{p_L(G^{(L)} \mid G_i = G_i^*)} \\ &= \exp \left(\frac{\sqrt{2}}{\sigma_L} \cdot \left(\left\| \frac{G_i^*}{n} + \frac{n-1}{n} G_{-i} - G^{(L)} \right\|_1 - \left\| \frac{G_i^{**}}{n} + \frac{n-1}{n} G_{-i} - G^{(L)} \right\|_1 \right) \right) \\ &\leq \exp \left(\frac{\sqrt{2} \cdot \|G_i^{**} - G_i^*\|_1}{\sigma_L n} \right) \leq \exp \left(\frac{\sqrt{2} \cdot mr}{\sigma_L n} \right). \end{aligned}$$

Since the probability is the integral of PDF, the above upper bound for PDF ratio is also an upper bound for probability ratio. Thus, for any possible output set S , we have,

$$\begin{aligned} & \Pr \left[\mathcal{M}_{\text{Laplacian}(\sigma_L)}(G_i^*, \mathcal{G}_{-i}) \in S \mid \mathcal{G}_{-i} \right] \\ &\leq \exp \left(\frac{\sqrt{2} \cdot mr}{\sigma_L n} \right) \Pr \left[\mathcal{M}_{\text{Laplacian}(\sigma_L)}(G_i^{**}, \mathcal{G}_{-i}) \in S \mid \mathcal{G}_{-i} \right]. \end{aligned}$$

and Theorem 3.5 follows by applying Definition 3.1. $\square$