

Harvard Journal of Law & Technology
Volume 31, Number 2 Spring 2018

**BRIDGING THE GAP BETWEEN COMPUTER SCIENCE
AND LEGAL APPROACHES TO PRIVACY**

*Kobbi Nissim, Aaron Bembenek, Alexandra Wood, Mark Bun, Marco
Gaboardi, Urs Gasser, David R. O'Brien, Thomas Steinke, & Salil
Vadhan**

TABLE OF CONTENTS

I. INTRODUCTION	689
II. BACKGROUND.....	694
A. <i>The Setting: Privacy in Statistical Computation</i>	695
1. What Is a Computation?	696
2. Privacy as a Property of Computation.....	697
3. Privacy Risks in Computation	699
B. <i>An Introduction to the Computer Science Approach to Defining Privacy</i>	702
C. <i>An Introduction to Legal Approaches to Privacy</i>	706
D. <i>Related Research to Bridge Between Legal and Computer Science Approaches to Privacy</i>	711
III. INTRODUCTION TO TWO PRIVACY CONCEPTS: DIFFERENTIAL PRIVACY AND FERPA	713

* Kobbi Nissim is a McDevitt Chair in Computer Science at Georgetown University. Aaron Bembenek is a PhD student in computer science at Harvard University. Alexandra Wood is a Fellow at the Berkman Klein Center for Internet & Society at Harvard University. Mark Bun is a postdoctoral researcher in the Theory of Computation Group at Princeton University. Marco Gaboardi is an Assistant Professor in the Computer Science and Engineering department at the State University of New York at Buffalo. Urs Gasser is the Executive Director of the Berkman Klein Center for Internet & Society at Harvard University and a Professor of Practice at Harvard Law School. David R. O'Brien is a Senior Researcher at the Berkman Klein Center for Internet & Society at Harvard University. Thomas Steinke is a postdoctoral researcher at IMB Research – Almaden. Salil Vadhan is the Vicky Joseph Professor of Computer Science and Applied Mathematics at Harvard University.

This Article is the product of a working group of the *Privacy Tools for Sharing Research Data* project at Harvard University. Kobbi Nissim led the working group discussions that informed the development of the approach articulated in this Article. Aaron Bembenek, Kobbi Nissim, and Alexandra Wood were the lead authors of the manuscript. Working group members Mark Bun, Marco Gaboardi, Urs Gasser, Kobbi Nissim, David O'Brien, Thomas Steinke, Salil Vadhan, and Alexandra Wood contributed to this work's conception.

The authors wish to thank Joseph Calandrino, Jonathan Frankle, Simson Garfinkel, Ann Kristin Glenster, Deborah Hurley, Georgios Kellaris, Paul Ohm, Leah Plunkett, Michel Raymond, Or Sheffet, Olga Slobodyanyuk, and the members of the *Privacy Tools for Sharing Research Data* project for their helpful comments. A preliminary version of this work was presented at the 9th Annual Privacy Law Scholars Conference (PLSC 2016), and the authors thank the participants for contributing thoughtful feedback. This material is based upon research supported by the National Science Foundation under Grant No. 1237235, grants from the Alfred P. Sloan Foundation, and a Simons Investigator grant to Salil Vadhan.

A. <i>Differential Privacy</i>	714
1. The Privacy Definition and Its Guarantee	714
2. Differential Privacy in the Real World.....	717
B. <i>The Family Educational Rights and Privacy Act of 1974</i> (FERPA)	718
1. The Applicability of FERPA's Requirements to Formal Privacy Models Such as Differential Privacy	720
2. The Distinction Between Directory and Non-Directory Information	724
3. The Definition of Personally Identifiable Information	725
C. <i>Gaps Between Differential Privacy and FERPA</i>	730
D. <i>Value in Bridging These Gaps</i>	733
IV. EXTRACTING A FORMAL PRIVACY DEFINITION FROM FERPA	734
A. <i>A Conservative Approach to Modeling</i>	737
B. <i>Modeling Components of a FERPA Privacy Game</i>	739
1. Modeling FERPA's Implicit Adversary	739
2. Modeling the Adversary's Knowledge.....	741
3. Modeling the Adversary's Capabilities and Incentives.....	743
4. Modeling Student Information	746
5. Modeling a Successful Attack	746
C. <i>Towards a FERPA Privacy Game</i>	749
1. Accounting for Ambiguity in Student Information	750
2. Accounting for the Adversary's Baseline Success.....	752
D. <i>The Game and Definition</i>	753
1. Mechanics	754
2. Privacy Definition.....	755
3. Privacy Loss Parameter	757
E. <i>Applying the Privacy Definition</i>	758
F. <i>Proving that Differential Privacy Satisfies the</i> <i>Requirements of FERPA</i>	759
G. <i>Extending the Model</i>	761
V. DISCUSSION	763

<i>A. Introducing a Formal Legal-Technical Approach to Privacy</i>	763
<i>B. Policy Implications and Practical Benefits of this New Approach</i>	765
1. Benefits for Privacy Practitioners.....	766
2. Benefits for Privacy Scholars	767
<i>C. Applying this Approach to Other Laws and Technologies</i>	769
<i>D. Setting the Privacy Parameter</i>	771
VI. CONCLUSION.....	771
APPENDIX I. PROVING THAT DIFFERENTIAL PRIVACY SATISFIES FERPA	772
APPENDIX II. EXTENSIONS OF THE PRIVACY GAME	777
<i>A. Untargeted Adversaries</i>	777
1. Mechanics	778
2. Discussion of the Untargeted Scenario.....	779
<i>B. Multiple Releases</i>	779

I. INTRODUCTION

The analysis and release of statistical data about individuals and groups of individuals carries inherent privacy risks,¹ and these risks have been conceptualized in different ways within the fields of law and computer science.² For instance, many information privacy laws adopt notions of privacy risk that are sector- or context-specific, such as in the case of laws that protect from disclosure *certain* types of information contained within health, educational, or financial records.³ In addition, many privacy laws refer to specific techniques, such as de-identification, that are designed to address a subset of possible attacks on privacy.⁴ In doing so, many legal standards for privacy protection

1. See generally Cynthia Dwork et al., *Robust Traceability from Trace Amounts*, 56 PROC. IEEE ANN. SYMP. ON FOUND. COMPUTER SCI. 650 (2015).

2. For a discussion of the differences between legal and computer science definitions of privacy, see Felix T. Wu, *Defining Privacy and Utility in Data Sets*, 84 U. COLO. L. REV. 1117, 1124–25 (2013).

3. See, e.g., Health Insurance Portability and Accountability Act Privacy Rule, 45 C.F.R. pt. 160, §§ 164.102–106, 500–534 (2017) (protecting “individually identifiable health information”); Family Educational Rights and Privacy, 34 C.F.R. pt. 99 (2017) (protecting non-directory “personally identifiable information contained in education records”); Gramm-Leach-Bliley Act, 15 U.S.C. §§ 6801–09 (2012) (protecting “personally identifiable financial information provided by a consumer to a financial institution; resulting from any transaction with the consumer or any service performed for the consumer; or otherwise obtained by the financial institution”).

4. See OFFICE FOR CIVIL RIGHTS, DEP’T OF HEALTH & HUMAN SERVS., GUIDANCE REGARDING METHODS FOR DE-IDENTIFICATION OF PROTECTED HEALTH INFORMATION IN ACCORDANCE WITH THE HEALTH INSURANCE PORTABILITY AND ACCOUNTABILITY ACT

rely on individual organizations to make case-by-case determinations regarding concepts such as the identifiability of the types of information they hold.⁵ These regulatory approaches are intended to be flexible, allowing organizations to (1) implement a variety of specific privacy measures that are appropriate given their varying institutional policies and needs, (2) adapt to evolving best practices, and (3) address a range of privacy-related harms. However, in the absence of clear thresholds and detailed guidance on making case-specific determinations, flexibility in the interpretation and application of such standards also creates uncertainty for practitioners and often results in ad hoc, heuristic processes. This uncertainty may pose a barrier to the adoption of new technologies that depend on unambiguous privacy requirements. It can also lead organizations to implement measures that fall short of protecting against the full range of data privacy risks.

Emerging concepts from computer science provide formal mathematical models for quantifying and mitigating privacy risks that differ significantly from traditional approaches to privacy such as de-identification. This creates conceptual challenges for the interpretation and application of existing legal standards, many of which implicitly or explicitly adopt concepts based on a de-identification approach to privacy. An example of a formal privacy model is *differential privacy*, which provides a provable guarantee of privacy against a wide range of potential attacks, including types of attacks currently unknown or unforeseen.⁶ New technologies relying on differential privacy are now making significant strides towards practical implementation. Several first-generation real-world implementations of differential privacy have been deployed by organizations such as Google, Apple, Uber, and the U.S. Census Bureau, and researchers in industry and academia are

(HIPAA) PRIVACY RULE 6–7 (2012), https://www.hhs.gov/sites/default/files/ocr/privacy/hipaa/understanding/coveredentities/De-identification/hhs_deid_guidance.pdf [https://perma.cc/76QG-5EW4].

5. See, e.g., Family Educational Rights and Privacy, 73 Fed. Reg. 74,805, 74,853 (Dec. 9, 2008) (stating that “determining whether a particular set of methods for de-identifying data and limiting disclosure risk is adequate [under FERPA] cannot be made without examining the underlying data sets, other data that have been released,” and other contextual factors).

6. Differential privacy was introduced by Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam Smith. Cynthia Dwork et al., *Calibrating Noise to Sensitivity in Private Data Analysis*, 3 PROC. THEORY CRYPTOGRAPHY CONF. 265 (2006). For a general introduction to differential privacy, see generally Cynthia Dwork, *A Firm Foundation for Private Data Analysis*, 54 COMM. ACM 86 (2011); Ori Heffetz & Katrina Ligett, *Privacy and Data-Based Research*, 28 J. ECON. PERSP. 75 (2014); Erica Klarreich, *Privacy by the Numbers: A New Approach to Safeguarding Data*, QUANTA MAG. (Dec. 10, 2012), <https://www.quantamagazine.org/20121210-privacy-by-the-numbers-a-new-approach-to-safeguarding-data> [https://perma.cc/FUQ7-BHC5]; Kobbi Nissim et al., *Differential Privacy: A Primer for a Non-technical Audience*, 21 VAND. J. ENT. & TECH. L. (forthcoming 2018).

currently building and testing additional tools for differentially private analysis.⁷

Demonstrating that these new technological tools are sufficient to satisfy relevant legal requirements will be key to their adoption and deployment for use with sensitive personal information. However, making such an argument is challenging due to significant gaps between legal and mathematical approaches to defining privacy.⁸ For instance, legal standards for privacy protection, and the definitions they employ, often vary according to industry sector, jurisdiction, institution, types of information involved, and other contextual factors.⁹ Variations between laws create challenges for the implementation of technological tools for privacy protection that are designed to be generally applicable. Legal approaches often, implicitly or explicitly, focus on a limited scope of attacks, such as re-identification by matching a named individual to a record in a database through linkage to information from public records databases. Adopting this conceptualization of privacy risks, many legal standards turn on the presence of personally identifiable information in a data release. Personally identifiable information is defined differently in various settings,¹⁰ involves substantial ambiguity as many definitions are context-dependent and evolving in response to technological developments,¹¹ and does not have a clear analog in mathematical definitions of privacy.¹²

In addition, standards and implementation guidance for privacy regulations such as the Family Educational Rights and Privacy Act of

7. See Ctr. for Econ. Studies, U.S. Census Bureau, *OnTheMap*, U.S. CENSUS BUREAU, <https://onthemap.ces.census.gov> (last visited May 5, 2018); Úlfar Erlingsson, Vasily Pihur & Aleksandra Korolova, *RAPPOR: Randomized Aggregatable Privacy-Preserving Ordinal Response*, 21 ACM CONF. ON COMPUTER & COMM. SECURITY 1054, 1054–55 (2014); Andrew Eland, *Tackling Urban Mobility with Technology*, GOOGLE EUR. BLOG (Nov. 15, 2015), <https://googlepolicyeurope.blogspot.com/2015/11/tackling-urban-mobility-with-technology.html> [<https://perma.cc/BDX6-QJR9>]; *Privacy Integrated Queries (PIQ)*, MICROSOFT (June 22, 2009), <https://research.microsoft.com/en-us/projects/pinq> [<https://perma.cc/ZQR8-UMK6>]; *Putting Differential Privacy to Work*, UNIV. OF PA., <http://privacy.cis.upenn.edu/index.html> [<https://perma.cc/C229-GK9Z>]; *Airavat: Improved Security for MapReduce*, UNIV. OF TEX. AT AUSTIN, <https://z.cs.utexas.edu/users/osa/airavat> [<https://perma.cc/8DKZ-WNZC>]; Prashanth Mohan et al., *GUPT: Privacy Preserving Data Analysis Made Easy*, 2012 ACM SIGMOD/PODS CONF. 349, 349–50; *PSI (Ψ): a Private Data Sharing Interface*, DATAVERSE, <https://beta.dataverse.org/custom/DifferentialPrivacyPrototype> (last visited May 5, 2018).

8. For an extended discussion of this argument, see *infra* Section III.C.

9. See Paul M. Schwartz & Daniel J. Solove, *The PII Problem: Privacy and a New Concept of Personally Identifiable Information*, 86 N.Y.U. L. REV. 1814, 1827–28 (2011).

10. See *id.*

11. See *id.* at 1818 (observing that “the line between PII and non-PII is not fixed but rather depends upon changing technological developments” and that “the ability to distinguish PII from non-PII is frequently contextual”).

12. See Arvind Narayanan & Vitaly Shmatikov, *Myths and Fallacies of “Personally Identifiable Information”*, 53 COMM. ACM 24, 26 (2010).

1974 (“FERPA”),¹³ the Health Insurance Portability and Accountability Act (“HIPAA”) Privacy Rule,¹⁴ and the Privacy Act of 1974,¹⁵ often emphasize techniques for protecting information released at the individual level but provide little guidance for releasing aggregate data, where the latter setting is particularly relevant to formal privacy models.¹⁶ In addition, there is limited or no available guidance for cases in which personal information may be partly leaked, or in which it may be possible to make inferences about individuals’ personal information with less than absolute certainty. Mathematical models of privacy must make determinations within difficult gray areas, where the boundaries of cognizable legal harms may not be fully defined. For these reasons, current regulatory terminology and concepts do not provide clear guidance for the implementation of formal privacy models such as differential privacy.

The main contribution of this Article is an argument that bridges the gap between the privacy requirements of FERPA, a federal law that protects the privacy of education records in the United States, and differential privacy, a formal mathematical model of privacy. Two arguments are made along the way. The first is a legal argument establishing that FERPA’s requirements for privacy protection are relevant to analyses computed with differential privacy. The second is a technical argument demonstrating that differential privacy satisfies FERPA’s requirements for privacy protection. To address the inherent ambiguity reflected in different interpretations of FERPA, the analysis takes a conservative, worst-case approach and extracts a mathematical requirement that is robust to differences in interpretation. The resulting argument thereby demonstrates that differential privacy satisfies a large class of reasonable interpretations of the FERPA privacy standard.

While FERPA and differential privacy are used to illustrate an application of the approach, the argument in this Article is an example of a more general method that may be developed over time and applied, with potential modifications, to bridge the gap between technologies other than differential privacy and privacy laws other than FERPA. Its degree of rigor enables the formulation of strong arguments about the privacy requirements of statutes and regulations. Further, with the level of generalization afforded by this conservative approach to modeling,

13. Family Educational Rights and Privacy Act of 1974, 20 U.S.C. § 1232g (2012).

14. Health Insurance Portability and Accountability Act Privacy Rule, 45 C.F.R. pt. 160, §§ 164.102–106, 500–534 (2017).

15. Privacy Act of 1974, 5 U.S.C. § 552a (2012).

16. Although documents have been produced to describe various techniques that are available for protecting privacy in aggregate data releases, *see, e.g.*, FED. COMM. ON STATISTICAL METHODOLOGY, REPORT ON STATISTICAL DISCLOSURE LIMITATION METHODOLOGY 99–103 (2005) (commonly referred to as “Statistical Policy Working Paper 22”), such documents are not designed to help practitioners determine when the use of certain techniques is sufficient to meet regulatory requirements.

differences between sector- and institution-specific standards are blurred, making significant portions of the argument broadly applicable. Future research exploring potential modifications to the argument and the extent to which it can address different interpretations of various legal standards could inform understandings of the significance of variations between different legal standards and, similarly, with respect to different technological standards of privacy.¹⁷

Arguments that are rigorous from both a legal and technical standpoint can also help support the future adoption of emerging privacy-preserving technologies. Such arguments could be used to assure actors who release data with the protections afforded by a particular privacy technology that they are doing so in compliance with the law. These arguments may also be used to better inform data subjects of the privacy protection to which they are legally entitled and to demonstrate that their personal information will be handled in accordance with that guarantee of privacy protection. More generally, the process of formalizing privacy statutes and regulations can help scholars and practitioners better understand legal requirements and identify aspects of law and policy that are ambiguous or insufficient to address real-world privacy risks. In turn, it can serve as a foundation for future extensions to address other problems in information privacy, an area that is both highly technical and highly regulated and therefore is well-suited to combined legal-technical solutions.¹⁸

Note that, while this Article demonstrates that differential privacy is sufficient to satisfy the requirements for privacy protection set forth by FERPA, it does not make the claim that differential privacy is the only technological standard that could be shown to be sufficient under the law. Indeed, future extensions analyzing other technologies may support sufficiency arguments with respect to various legal and policy requirements for privacy protection. Differential privacy is just one of a large collection of privacy interventions that may be considered as an appropriate component of an organization's data management program. This Article is intended to complement other resources providing practical guidance on incorporating data sharing models, including those that rely on formal privacy frameworks like differential privacy, within a specific institutional setting.¹⁹

17. For a further discussion of anticipated differences with respect to modeling legal standards beyond FERPA, see *infra* Part V.

18. For other research formalizing legal requirements for privacy protection using mathematical approaches, see, for example, Omar Chowdhury et al., *Privacy Promises That Can Be Kept: A Policy Analysis Method with Application to the HIPAA Privacy Rule*, 26 PROC. INT'L CONF. COMPUTER AIDED VERIFICATION 131, 134–39 (2014); Henry DeYoung et al., *Experiences in the Logical Specification of the HIPAA and GLBA Privacy Laws*, 9 PROC. ACM WORKSHOP ON PRIVACY ELECTRONIC SOC'Y 73, 76–80 (2010).

19. See Micah Altman et al., *Towards a Modern Approach to Privacy-Aware Government Data Releases*, 30 BERKELEY TECH. L.J. 1967, 1975–2009 (2015) (providing an extended

The following sections establish a rigorous approach to modeling FERPA and bridging between differential privacy and FERPA's privacy requirements. Part II describes the setting in which the privacy issues relevant to this discussion arise and provides a high-level overview of approaches from computer science and law that have emerged to address these issues. Part III provides an introduction to the two privacy concepts at the heart of the analysis, differential privacy and FERPA. It also discusses the applicability of FERPA's requirements to differentially private computations and articulates the gaps between differential privacy and FERPA that create challenges for implementing differential privacy in practice. The later sections present a novel argument for formally proving that differential privacy satisfies FERPA's privacy requirements. Part IV describes the process of extracting a formal mathematical requirement of privacy protection under FERPA. It demonstrates how to construct a model of the attacker that is implicitly recognized by FERPA, based on the definitions found in the FERPA regulations and informed by relevant administrative guidance. Part V concludes with a discussion exploring how formal argumentation can help enable real-world implementation of emerging formal privacy technologies and the development of more robust privacy regulations.

For reference, the Article includes two Appendices with supplementary arguments and models. Appendix I provides a sketch of the mathematical proof to demonstrate that differential privacy satisfies the mathematical definition of privacy extracted from FERPA. Appendix II presents two possible extensions of the model described in Part IV.

II. BACKGROUND

Privacy is conceptualized differently across a range of legal contexts, from surveillance to criminal procedure to public records releases to research ethics to medical decision making.²⁰ Because privacy law is so broad in its reach and privacy measures can be designed to address a wide range of harms, it is important to define the scope of the analysis

discussion of the array of legal, technical, and procedural privacy controls that are available and how to determine which are suitably aligned with the privacy risks and intended uses in a particular context); SIMSON GARFINKEL, U.S. DEP'T OF COMMERCE, DE-IDENTIFYING GOVERNMENT DATASETS 35–57 (2016), https://csrc.nist.gov/CSRC/media/Publications/sp/800-188/draft/documents/sp800_188_draft2.pdf [<https://perma.cc/HU4N-469S>].

20. See, e.g., Daniel J. Solove, *A Taxonomy of Privacy*, 154 U. PA. L. REV. 477, 484–90 (2006) (grouping privacy problems into a wide range of categories, including those listed above); Bert-Jaap Koops et al., *A Typology of Privacy*, 38 U. PA. J. INT'L L. 483, 566–68 (2017) (surveying constitutional protections and privacy scholarship from nine North American and European countries, resulting in a classification of privacy into types); ALAN F. WESTIN, *PRIVACY AND FREEDOM* 31 (1967) (identifying four states of privacy: solitude, intimacy, anonymity, and reserve).

in this Article. Specifically, this Article focuses on privacy in the context of the statistical analysis of data or the release of statistics derived from personal data.

A. The Setting: Privacy in Statistical Computation

Many government agencies, commercial entities, and research organizations collect, process, and analyze data about individuals and groups of individuals. They also frequently release such data or statistics based on the data. Various data protection laws and policies are in place, restricting the degree to which data containing personal information can be disclosed, including the formats in which the data can be published.

Federal and state statistical agencies such as the Census Bureau, the Bureau of Labor Statistics, and the National Center for Education Statistics release large quantities of statistics about individuals, households, and establishments, upon which policy and business investment decisions are based. For instance, the National Center for Education Statistics collects data from U.S. schools and colleges, using surveys and administrative records, and releases statistics derived from these data to the public.²¹ Statistics such as total enrollments at public and private schools by grade level, standardized test scores by state and student demographics, and high school graduation rates, among other figures, are used to shape education policy and school practices. The release of such statistics by federal statistical agencies is subject to strict requirements to protect the privacy of the individuals in the data.²²

Companies such as Google and Facebook also collect personal information, which they use to provide services to individual users and third-party advertisers. For instance, Facebook enables advertisers to target ads on the Facebook platform based on the locations, demographics, interests, and behaviors of their target audiences, and provides tailored estimates of the number of recipients of an ad given different parameters.²³ The Federal Trade Commission regulates the

21. See THOMAS D. SNYDER & SALLY A. DILLOW, NAT'L CTR. FOR EDUC. STATISTICS, U.S. DEP'T OF EDUC., DIGEST OF EDUCATION STATISTICS: 2013 (2015), <https://nces.ed.gov/pubs2015/2015011.pdf> (last visited May 5, 2018).

22. See 44 U.S.C. § 3501 note (2012) (Confidential Information Protection and Statistical Efficiency Act) (prohibiting the release of statistics in "identifiable form" and establishing specific confidentiality procedures to be followed by employees and contractors when collecting, processing, analyzing, and releasing data).

23. See *Facebook for Business: How to Target Facebook Ads*, FACEBOOK (Apr. 27, 2016), <https://www.facebook.com/business/a/online-sales/ad-targeting-details> [https://perma.cc/XX7A-TZKV]. These audience-reach statistics are evidently rounded, in part, to protect the privacy of Facebook users. See Andrew Chin & Anne Klinefelter, *Differential Privacy as a Response to the Reidentification Threat: The Facebook Advertiser Case Study*, 90 N. CAROLINA L. REV. 1417, 1441–43 (2012).

commercial collection, use, and release of such data.²⁴ Federal and state privacy laws also govern certain commercial data activities.²⁵

In addition, researchers and their institutions release statistics to the public and make their data available to other researchers for secondary analysis. These research activities are subject to oversight by an institutional review board in accordance with the Federal Policy for the Protection of Human Subjects.²⁶ Researchers governed by these regulations are required to obtain consent from participants and to employ data privacy safeguards when collecting, storing, and sharing data about individuals.²⁷

The following discussion explores how statistical computations can be performed by government agencies, commercial entities, and researchers while providing strong privacy protection to individuals in the data.

1. What Is a Computation?

A *computation* (alternatively referred to in the literature as an algorithm, mechanism, or analysis) is a mechanizable procedure for producing an output given some input data, as illustrated in Figure 1.²⁸ This general definition does not restrict the nature of the relationship between the input data and the output. For instance, a computation could output its input without any transformation, or it could even ignore its input and produce an output that is independent of the input. Some computations are *deterministic* and others are *randomized*. A deterministic computation will always produce the same output given the same input, while a randomized computation does not provide such a guarantee. A computation that returns the mean age of participants in a dataset is an example of a deterministic computation. However, a similar computation that estimates the mean age in the dataset by sampling at random a small subset of the records in a database and returning the mean age for

24. FED. TRADE COMM'N, PROTECTING CONSUMER PRIVACY IN AN ERA OF RAPID CHANGE (2012).

25. See, e.g., Children's Online Privacy Protection Act, 15 U.S.C. §§ 6501–06 (2012) (requiring certain web site operators to provide notice and obtain consent when collecting personal information from children under 13 years of age); CAL. CIV. CODE § 1798.82 (2017) (requiring businesses to disclose any data breach to California residents whose unencrypted personal information was acquired by an unauthorized person).

26. See 45 C.F.R. § 46 (2017).

27. See *id.* § 46.111.

28. This figure is adapted from Nissim et al., *supra* note 6, at 5.

that sample is a randomized computation. As another example, a computation that outputs the mean age with the addition of some random noise is a randomized computation.²⁹

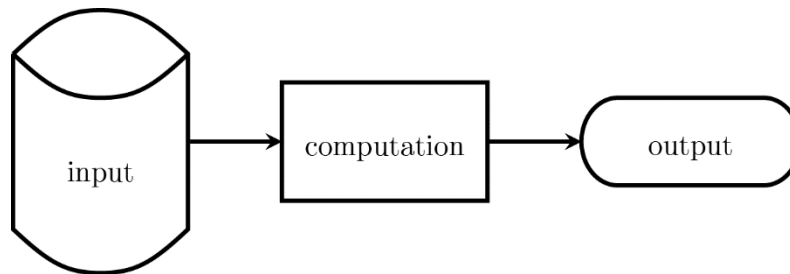


Figure 1: A computation produces an output given some input data.

Public releases of statistics derived from personal information, such as the statistical information many government agencies publish, result from applying certain computations on the personal information collected. The U.S. Census Bureau, for example, collects information about individuals and households, including sensitive personal information, in order to release statistics about the larger population.³⁰ Statistics such as the median household income of a city or region are computed after collecting household-level income information from a sample population in that area.³¹ Similarly, a school district, seeking to study and report on the performance of students in its schools, applies a set of computations to the raw data it has collected about individual students to generate more general statistics for publication. Releasing statistics about personal information can benefit society by enabling research and informing policy decisions. However, the computations used to create the statistics must provide sufficient privacy protection for the individuals in the data. The next subsection discusses how it is possible to release statistics based on personal information while protecting individual privacy.

2. Privacy as a Property of Computation

This Article considers privacy to be a property of computations. More specifically, a computation is said to be “privacy-preserving” if

29. Randomness may be introduced into computations for several reasons. Foremost among them are efficiency and the need to create uncertainty for an adversary (for example, when cryptographic keys are selected or for achieving privacy protection).

30. U.S. CENSUS BUREAU, AMERICAN COMMUNITY SURVEY: INFORMATION GUIDE (2017), https://www.census.gov/content/dam/Census/programs-surveys/acs/about/ACS_Information_Guide.pdf [<https://perma.cc/Q9V7-PWAV>].

31. See *id.* at 10.

the relationship between input and output of the computation must satisfy the requirements of a particular definition of privacy. In other words, it is the computation itself that is private rather than a particular output of the computation that is private.

To see why distinguishing between a purportedly private or non-private output fails to capture a reasonable notion of privacy, consider a policy requiring all statistics to be coarsened to the nearest ten (e.g., 0–9, 10–19, 20–29, etc.), under the assumption that such coarsening would hide the information of individuals. Suppose a school releases data for the fall semester showing that between twenty and twenty-nine of its students have a disability — an output that may seem innocuous in terms of risk to individual privacy. However, suppose that in the spring, the school releases updated data showing that between thirty and thirty-nine of its students have a disability, another output that may seem innocuous in isolation. However, these data in combination reveal that at least one student with a disability joined the school between the fall and spring semesters. Although both *outputs* seem innocuous, reasoning about how they *depend on the input data* — a dependency created by the computation of the statistics — reveals sensitive information.

Computer scientists seek to reason about the properties of computations and treating privacy as a computational property fits naturally in this worldview. This approach has successful precedents in established areas of computer science such as cryptography.³² This approach — defining privacy as a property of a computation — is also applicable from a legal perspective. While privacy laws and policies might not explicitly refer to computation, they often attempt to define privacy implicitly as a property that certain computations possess. For instance, consider the HIPAA Privacy Rule’s safe harbor method of de-identification, which is satisfied when certain pieces of personal information deemed to be identifying have been removed from a dataset prior to release.³³ This provision is in effect specifying a computation that produces an output in which the identifying information has been redacted. Such a computation is considered to provide sufficient de-identification in accordance with the law. Similarly, regulatory requirements or related guidance prescribing minimum cell counts for aggregate data tables produced using personal information,³⁴ or

32. For instance, a cryptographic computation might be considered to provide privacy protection if any hypothetical adversary given an encrypted message can do no better at guessing a property of that message than another hypothetical adversary that is not given the encrypted message. See Shafi Goldwasser & Silvio Micali, *Probabilistic Encryption & How to Play Mental Poker Keeping Secret All Partial Information*, 14 ACM SYMP. ON THEORY COMPUTING 365, 365–66 (1982).

33. See 45 C.F.R. § 164.514 (2017).

34. For example, in accordance with the Elementary and Secondary Education Act of 1965, states must define minimum cell counts for the publication of student achievement results

recommending transforming data using k -anonymization techniques prior to release,³⁵ effectively treat certain computations as providing privacy protection.³⁶

A benefit of viewing privacy as a property of computation is that computations are formal mathematical objects and as such are amenable to rigorous mathematical arguments. This Article likewise treats privacy as a property of computations in order to make rigorous arguments about the privacy requirements of laws and regulations. In turn, these arguments can assure actors who release data that they are acting in accordance with the law. They can also provide data subjects with a better understanding of the privacy protection to which they are legally entitled and an assurance that their personal information will be handled in accordance with that guarantee of privacy protection. Additionally, the process itself of formalizing legal requirements as computational objects can lead to a better understanding of those requirements and help identify aspects of the law that are ambiguous or potentially insufficient to provide adequate privacy protection.

3. Privacy Risks in Computation

Several approaches have been developed and widely implemented to limit the disclosure of personal information when sharing statistical data about individuals. Traditional approaches include obtaining consent from data subjects, entering into data use agreements restricting the use and re-disclosure of data, and applying various techniques for de-identifying data prior to release.³⁷ Statistics about individuals or

“[b]ased on sound statistical methodology” that “[y]ield[s] statistically reliable information for each purpose for which disaggregated data are used” and does not “reveal personally identifiable information about an individual student.” 34 C.F.R. § 200.17 (2017).

35. “A release provides k -anonymity protection if the information for each person contained in the release cannot be distinguished from at least $k - 1$ individuals whose information also appears in the release.” Latanya Sweeney, *k-Anonymity: A Model for Protecting Privacy*, 10 INT’L J. ON UNCERTAINTY, FUZZINESS & KNOWLEDGE-BASED SYSTEMS 557, 557 (2002). Guidance from the Department of Health and Human Services covers the application of k -anonymity as one approach to protecting health records subject to the HIPAA Privacy Rule. See OFFICE FOR CIVIL RIGHTS, DEP’T OF HEALTH & HUMAN SERVS., *supra* note 4, at 21 (providing guidance on applying k -anonymity and noting that “[t]he value for k should be set at a level that is appropriate to mitigate risk of identification by the anticipated recipient of the data set,” while declining to “designate a universal value for k ”).

36. Note that by specifying properties of the output, such requirements restrict the computation, but do not necessarily restrict the informational relationship between input and output. These examples are included here only to support the claim that conceptualizing the problem in terms of restrictions on computation is compatible with some existing legal and regulatory approaches to privacy protection.

37. See, e.g., DEP’T OF EDUC., DATA-SHARING TOOL KIT FOR COMMUNITIES: HOW TO LEVERAGE COMMUNITY RELATIONSHIPS WHILE PROTECTING STUDENT PRIVACY 1–6 (2016), <https://www2.ed.gov/programs/promiseneighborhoods/datasharingtool.pdf> [<https://perma.cc/LR8S-YXQJ>].

groups of individuals are generally made available after de-identification techniques have transformed the data by removing, generalizing, aggregating, and adding noise to pieces of information determined to be identifiable. At the core of this approach is the concept of *personally identifiable information* (“PII”), which is based in a belief that privacy risks lurk in tables of individual-level information, and that protecting privacy depends on the removal of information deemed to be identifying, such as names, addresses, and Social Security numbers.³⁸

Privacy in the release of statistical data about individuals is under increasing scrutiny. Numerous data privacy breaches have demonstrated that risks can be discovered even in releases of data that have been redacted of PII. For example, in the late 1990s, Latanya Sweeney famously demonstrated that the medical record of Massachusetts Governor William Weld could be identified in a release of data on state employee hospital visits that had been stripped of the names and addresses of patients.³⁹ Using Governor Weld’s date of birth, ZIP code, and gender, which could be found in public records, Sweeney was able to locate his record in the dataset that had been released, as it was the only record that matched all three attributes. Indeed, well over 50% of the U.S. population can be uniquely identified using these three pieces of information.⁴⁰

Repeated demonstrations across many types of data have confirmed that this type of privacy breach is not merely anecdotal but is in fact widespread. For instance, it has been shown that individuals can be identified in releases of Netflix viewing records and AOL search query histories, despite efforts to remove identifying information from the

38. Many privacy laws explicitly or implicitly endorse the practice of removing personal information considered to be identifying prior to release. See, e.g., Family Educational Rights and Privacy Act of 1974, 20 U.S.C. § 1232g (2012); Family Educational Rights and Privacy, 34 C.F.R. pt. 99 (2017) (permitting educational agencies and institutions to release, without the consent of students or their parents, information from education records that have been de-identified through the removal of “personally identifiable information”). For an extended discussion of various legal approaches to de-identification, see Section II.C below. Such approaches also appear in a wide range of guidance materials on privacy and confidentiality in data management. See, e.g., FED. COMM. ON STATISTICAL METHODOLOGY, *supra* note 16, at 103 (recommending the removal of identifiers but also acknowledging that “[a]fter direct identifiers have been removed, a file may still remain identifiable, if sufficient data are left on the file with which to match with information from an external source that also contains names or other direct identifiers”).

39. See *Recommendations to Identify and Combat Privacy Problems in the Commonwealth: Hearing on H.R. 351 Before the H. Select Comm. on Information Security*, 189th Sess. (Pa. 2005) (statement of Latanya Sweeney, Associate Professor, Carnegie Mellon University), <https://dataprivacylab.org/dataprivacy/talks/Flick-05-10.html> [<https://perma.cc/L6C9-AZZA>]; see also Latanya Sweeney, *Weaving Technology and Policy Together to Maintain Confidentiality*, 25 J.L. MED. & ETHICS 98, 102–07 (1997) (describing Sweeney’s methods in greater detail).

40. See Latanya Sweeney, *Simple Demographics Often Identify People Uniquely*, DATA PRIVACY LAB 27–28 (2000), <https://dataprivacylab.org/projects/identifiability/paper1.pdf> [<https://perma.cc/SMS6-RGUC>].

data prior to release.⁴¹ Other research has demonstrated that just four records of an individual's location at a point in time can be sufficient to identify 95% of individuals in mobile phone data and 90% of individuals in credit card purchase data.⁴² Narayanan and Shmatikov generalize these and other privacy failures, stating that "[a]ny information that distinguishes one person from another can be used for re-identifying anonymous data."⁴³

Other privacy attacks have exposed vulnerabilities in releases of aggregate data and revealed inference risks that are distinct from the risk of re-identification.⁴⁴ More specifically, many successful attacks on privacy have focused not on discovering the identities of individuals but rather on learning or inferring sensitive details about them. For example, researchers have discovered privacy risks in databases containing information about mixtures of genomic DNA from hundreds of people.⁴⁵ Although the data were believed to be sufficiently aggregated so as to pose little risk to the privacy of individuals, it was shown that an individual's participation in a study could be confirmed using the aggregated data, thereby revealing that the individual suffers from the medical condition that was the focus of the research study.

Privacy risks have also been uncovered in online recommendation systems used by web sites such as Amazon, Netflix, and Last.fm, which employ algorithms that recommend similar products to users based on an analysis of data generated by the behavior of millions of users.⁴⁶ These attacks have shown that recommendation systems can leak information about the transactions made by individuals. In addition to such

41. See Arvind Narayanan & Vitaly Shmatikov, *Robust De-anonymization of Large Sparse Datasets*, 29 PROC. IEEE SYMP. ON SECURITY & PRIVACY 111, 118–23 (2008); Michael Barbaro & Tom Zeller Jr., *A Face Is Exposed for AOL Searcher No. 4417749*, N.Y. TIMES (Aug. 9, 2006), <https://www.nytimes.com/2006/08/09/technology/09aol.html> (last visited May 5, 2018).

42. See Yves-Alexandre de Montjoye et al., *Unique in the Shopping Mall: On the Reidentifiability of Credit Card Metadata*, 347 SCIENCE 536, 537 (2015) [hereinafter *Credit Card Metadata*]; Yves-Alexandre de Montjoye et al., *Unique in the Crowd: The Privacy Bounds of Human Mobility*, 3 SCI. REPS. 1376, 1377–78 (2013) [hereinafter *Human Mobility*].

43. Narayanan & Shmatikov, *supra* note 12, at 26.

44. For a recent survey of different classes of privacy attacks, including both re-identification attacks and tracing attacks (i.e., attacks that aim to determine whether information about a target individual is in a database), see Cynthia Dwork et al., *Exposed! A Survey of Attacks on Private Data*, 4 ANN. REV. STAT. & ITS APPLICATION 61, 65–77 (2017).

45. See Nils Homer et al., *Resolving Individuals Contributing Trace Amounts of DNA to Highly Complex Mixtures Using High-Density SNP Genotyping Microarrays*, 4 PLOS GENETICS, at 6 (2008), <http://journals.plos.org/plosgenetics/article/file?id=10.1371/journal.pgen.1000167&type=printable> [<https://perma.cc/XF2F-8YZV>]. This attack involves using published genomic statistics and some information about the underlying population to determine whether an individual's genomic sequence was used in the study. This attack has been strengthened and generalized in several works. See, e.g., Cynthia Dwork et al., *Robust Traceability from Trace Amounts*, 56 PROC. IEEE SYMP. ON FOUND. COMPUTER SCI. 650, 651–52 (2015).

46. See Joseph A. Calandrino et al., *"You Might Also Like:" Privacy Risks of Collaborative Filtering*, 32 PROC. IEEE SYMP. ON SECURITY & PRIVACY 231, 232 (2011).

attacks that have been demonstrated on real-world data, theoretical work has resulted in other, more general privacy attacks on aggregate data releases in a variety of settings.⁴⁷ These findings, taken together, demonstrate that privacy risks may be present when releasing not just individual-level records but also aggregate statistics, and that privacy risks include not only re-identification risks but other inference risks as well.

Moreover, techniques for inferring information about individuals in a data release are rapidly advancing and exposing vulnerabilities in many commonly used measures for protecting privacy. Although traditional approaches to privacy may have been sufficient at the time they were developed, they are becoming increasingly ill-suited for protecting information in the Internet age.⁴⁸ It is likely that privacy risks will continue to grow and evolve over time, enabled by rapid advances in analytical capabilities and the increasing availability of personal data from different sources, and motivated by the tremendous value of sensitive personal data. These realizations have led computer scientists to seek new approaches to data privacy that are robust against a wide range of attacks, including types of attacks unforeseen at the time they were developed.

B. An Introduction to the Computer Science Approach to Defining Privacy

In the field of computer science, privacy is often formalized as a game, or a thought experiment, in which an adversary attempts to exploit a computation to learn protected information. A system is considered to provide privacy protection if it can be demonstrated, via a mathematical proof, that no adversary can win the game with a probability that is “too high.”⁴⁹ Every aspect of a game framework must be carefully and formally defined, including any constraints on the adversary, the mechanics of the game, what it means for the adversary to win the game, and with what probability it is acceptable for the adversary to win the game. This formalization allows one to prove that a given system is private under the explicit assumptions of the model.

47. For example, Dinur and Nissim showed that publishing estimates to randomly issued statistical queries can lead to a very accurate reconstruction of the information in a statistical database, and hence result in a massive leakage of individual information. Irit Dinur & Kobbi Nissim, *Revealing Information While Preserving Privacy*, 22 *PROC. ACM PODS* 202, 206–08 (2003).

48. It is important to also note that, in addition to vulnerability to privacy attacks, traditional de-identification techniques may be unsuitable in practice for other reasons, such as the impact their use has on data quality. See, e.g., Jon P. Daries et al., *Privacy, Anonymity, and Big Data in the Social Sciences*, 57 *COMM. ACM* 56, 63 (2014).

49. See, e.g., JONATHAN KATZ & YEHUDA LINDELL, *INTRODUCTION TO MODERN CRYPTOGRAPHY* 43–44 (Douglas R. Stinson ed., 2d ed. 2014).

These formal privacy games have the following components: an adversary, a computation, and game mechanics. Each of these components is discussed in turn below.

An *adversary* attempts to exploit the system to learn private information. The adversary is not defined in terms of the specific techniques he might use to exploit the system, but rather by the computational power, access to the system, and background knowledge he can leverage. As a consequence, the adversary does not represent a uniquely specified attack, but rather a whole class of attacks, including attacks that have not been conceived by the system designer. This means that a system cannot be tested for its privacy, as testing its resilience to known attacks would not rule out its vulnerability to other attacks. Rather, privacy needs to be proved mathematically. Proving that the system provides privacy protection against such an adversary makes a strong claim: the system is private no matter the particular attack or attacker, provided that the model's assumptions are not violated.

A *computation* takes as input a dataset (potentially containing private information) and produces some output. For instance, one could envision a computation that takes as input a spreadsheet of students' grades and returns the average grade (or an approximation of it). Unlike the adversary, the computation needs to be uniquely specified and, furthermore, known to the adversary.⁵⁰ The game framework is used to prove that a given computation (or a class of computations) preserves privacy under the assumptions of the game. For example, one might prove that in a particular game framework a computation that reports the average grade does so in such a way as to maintain the individual privacy of the students in the original dataset.

The *game mechanics* act as an intermediary between the adversary and the private data. The adversary does not have direct access to non-public data, and instead receives information via the game mechanics. The game mechanics enforce a specific protocol and determine the ways in which the adversary can interact with the computation. In addition, the game mechanics define when the adversary is deemed to have won the game and when a system is deemed to provide a sufficient level of privacy. For instance, the game mechanics might specify the following protocol: the previously described computation is used to calculate the average grade on a test, and then the computation result is released to the adversary. The adversary responds with a guess of the

50. This follows a design principle widely accepted in cryptography and is known as Kerckhoffs' principle: a cryptographic system should maintain security even if all details of its design and implementation are known to the attacker. See AUGUSTE KERCKHOFFS, *LA CRYPTOGRAPHIE MILITAIRE* 8 (1883).

grade of a particular student and is considered to have won the game if his guess is within a certain range of that student's true grade.⁵¹

A privacy game provides a precise definition of privacy. A computation satisfies the definition of privacy if no adversary can win the game (with the computation) "too often." Note that it does not require that an adversary never win the game. Even if intuitive, such a requirement would not be achievable because there is always some probability that an adversary will win the game by sheer chance, even without seeing any outcome of the computation. Thus, the standard is that no adversary should be able to win the game with a probability significantly greater than some baseline probability, generally taken to be the probability of winning without access to the computation outcome.

To illustrate, consider a game in which an adversary's goal is to guess a person's gender. It is possible for the adversary to guess the person's gender correctly with a probability of 50% even without obtaining any information about that person. This means that the baseline for comparison should be a probability of (at least) 50% for guessing a target individual's gender, as it would not be reasonable to establish a requirement that all adversaries must guess correctly with a success rate that is lower than this probability. Additionally, adversaries are typically allowed to win with a probability that is slightly higher than the baseline value because any system that provides utility necessarily leaks at least a tiny bit of information.⁵² How much the adversary is allowed to win beyond the baseline probability while the computation is still considered to satisfy the privacy definition is often quantified as a parameter, and this parameter can be tuned to provide less or more privacy protection.⁵³ The cryptographic standard for encryption is to

51. Note that the game mechanics do not necessarily correspond to any specific "real-world" entity. A privacy game is a thought experiment and in general there is not a direct correspondence between the components of a privacy game and the parties in a real-world privacy scenario.

52. See FED. COMM. ON STATISTICAL METHODOLOGY, *supra* note 16, at 3 (2005). Intuitively, a release "provides utility" only if it reveals some information that was not previously known about the group of individuals whose information is in the dataset. For a concrete example illustrating why absolute disclosure prevention is impossible in any system that provides utility, consider an illustration by Dwork and Naor. Rephrasing their example, consider a privacy attacker who has access to the auxiliary information: "Terry Gross is two inches shorter than the average American woman," but possesses no other knowledge about the height of American women (hence, the auxiliary information is initially not helpful to the attacker). Given access to a system that allows estimating the average height of American women, the attacker can estimate Terry Gross' height. Informally, the attacker's auxiliary information contained sensitive information (Terry Gross' height), but this information was "encrypted" and hence useless to the attacker. The decryption key (the average height of American women) could be learned by invoking the system's computation. See Cynthia Dwork & Moni Naor, *On the Difficulties of Disclosure Prevention in Statistical Databases or the Case for Differential Privacy*, 2 J. PRIVACY & CONFIDENTIALITY 93, 103 (2010).

53. For further discussion of the privacy parameter, see *infra* Part IV.

only allow adversaries a negligible advantage over the 50% probability, say a 50.00000000001% chance of winning the game.⁵⁴

For a concrete example of a game mechanics, consider the following privacy game from cryptography. Alice wants to send a private message to Bob, but she is worried that a third party, Eve, might be eavesdropping on their communications. Therefore, Alice decides to encrypt the message before sending it to Bob. She wants to be confident that the computation she uses to encrypt the message will ensure that Eve cannot learn much about the content of the original message (i.e., the plaintext) from seeing the encrypted version of the message (i.e., the ciphertext). To gain confidence in the security of the system, Alice can formalize her privacy desiderata as a game and then use an encryption computation that is proven to meet the game's definition of privacy. Here is one possible description of the mechanics for the game, also represented visually in Figure 2.⁵⁵

- (1) An adversary Eve chooses two distinct plaintext messages and passes them to the game mechanics. Intuitively, she is asserting that she can distinguish between the encrypted messages and hence the encryption is not secure.
- (2) The game mechanics toss a fair coin to choose between the two messages with equal probability, encrypt the chosen message (denoted "plaintext" in Figure 2) with a computation C , and give the resulting ciphertext to the adversary.
- (3) The adversary wins if she is able to guess from seeing the ciphertext which of the two original messages was encrypted.

54. The difference of 0.00000000001% between the baseline of 50% and the set threshold of 50.00000000001% determines the adversary's cost-benefit tradeoff. For example, if Eve's goal is to differentiate between the two messages, "Attack at dawn" and "Attack at sunset," then she would have to accumulate a number of encrypted messages that is inversely proportional to this difference (i.e., in the order of magnitude of 10^{12} messages). The difference that is deemed to be acceptable may also affect the efficiency of the cryptographic scheme. For instance, it can affect Alice's computational costs. Exactly how much higher the adversary is permitted to go above the baseline probability is often captured as a parameter that can be tuned to different privacy (and computational cost) levels. For instance, if Alice feels that the message she is sending is not especially sensitive, she might decide that it is acceptable for an adversary to win the game with a probability of 50.000001%.

55. This game is modeled after the notion of indistinguishability of ciphertexts introduced by Goldwasser & Micali. Goldwasser & Micali, *supra* note 32, at 367–70.

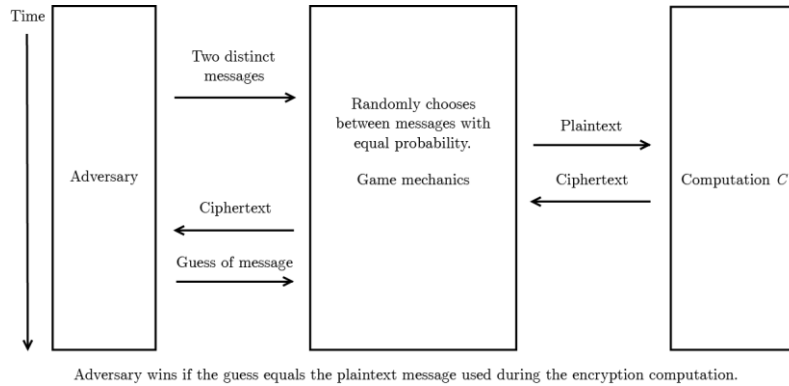


Figure 2: Example cryptographic privacy game.

Notice that an adversary who ignores the ciphertext she is given in Step 2 of the game and simply outputs one of the messages she selected in Step 1 already has a 50% chance of winning. This means that any reasonable adversary would have a winning probability of at least 50%, so the game mechanics may declare an adversary to have won this game if the adversary can decipher the message correctly 50.00000000001% of the time. Now that her privacy desiderata have been formalized, Alice can use any encryption computation that is mathematically proven to meet this definition of privacy to send an encrypted message to Bob with confidence that Eve cannot learn much from eavesdropping on their communication.

Although privacy laws are not written with an explicit game-based definition of privacy, this Article argues that it is possible to extract a suitable privacy game from a law, its legislative history, and administrative guidance. Furthermore, the privacy game that is extracted can be used to establish that particular computations meet the privacy requirements of the law. In Part IV below, the Article extracts such a privacy game from FERPA. Based on this game, the Article sketches in Appendix I a proof showing that all differentially private computations provide sufficient privacy protection to satisfy the privacy requirements of FERPA.

C. An Introduction to Legal Approaches to Privacy

Information privacy laws around the world vary substantially with respect to their scope of coverage and the protections they afford.⁵⁶ This discussion focuses on a subset of information privacy laws relevant to

56. For a broad survey and classification of privacy laws across many jurisdictions, see generally Koops et al., *supra* note 20.

privacy in statistical computation, namely those laws that restrict the release of statistical information about individuals or groups of individuals, whether released as raw data, de-identified data, or statistical summaries.

The applicability of such laws typically turns on the definition of “personal information” or a similar term.⁵⁷ Typically, if the information to be released falls within a particular law’s definition of personal information, then the law applies and restricts the disclosure of the information.⁵⁸ Conversely, if it does not meet the particular law’s definition of personal information, then the information is often afforded minimal or no protection under that law.⁵⁹ In addition, some privacy laws expressly exclude a category of de-identified information. In some cases, de-identified information can even be released publicly without further restriction on use or redistribution.⁶⁰

Definitions of personal information vary considerably across information privacy laws. The inconsistency between definitions and the reliance on ambiguous and narrow interpretations of these definitions are widely cited as weaknesses of the legal framework for privacy protection.⁶¹ It is beyond the scope of this Article to detail all legal approaches to privacy and definitions of personal information currently in place around the world. Instead, the following discussion provides an overview of selected approaches in order to illustrate a range of different approaches and definitions, and some of the challenges that have arisen in developing, interpreting, and complying with various regulatory definitions and standards for privacy protection.

Privacy law in the United States takes a sectoral approach. Many privacy laws are in place at the federal and state level, and each law is drawn rather narrowly to protect certain types of information in particular contexts, from video rental records to medical records.⁶² Despite

57. For an overview of various definitions of personal information, see Schwartz & Solove, *supra* note 9, at 1828–35.

58. See, e.g., Children’s Online Privacy Protection Act, 15 U.S.C. § 6502(a)(1) (2012).

59. See Schwartz & Solove, *supra* note 9, at 1816 (finding that many laws “share the same basic assumption — that in the absence of PII, there is no privacy harm”).

60. See, e.g., HIPAA Privacy Rule, 45 C.F.R. § 164.502(d)(2) (2017).

61. See Schwartz & Solove, *supra* note 9, at 1893 (“[T]here is no uniform definition of PII in information privacy law. Moreover, the definitions that do exist are unsatisfactory.”).

62. For example, the Video Privacy Protection Act protects the privacy of individuals in records of video sales and rentals, defining “personally identifiable information” as “information which identifies a person as having requested or obtained specific video materials or services from a video tape service provider.” 18 U.S.C. § 2710(a)(3) (2013). In contrast, the California Confidentiality of Medical Information Act defines “medical information” as “individually identifiable health information . . . regarding a patient’s medical history, mental or physical condition, or treatment” which “includes . . . any element of personal identifying information sufficient to allow identification of the individual, such as . . . name, address, electronic mail address, telephone number, or social security number, or other information that . . . in combination with other publicly available information, reveals the individual’s identity.” CAL. CIV. CODE § 56.05 (West 2014). For a discussion of the evolution and nature of the U.S.

their context-specific differences, many U.S. laws share some variation of a definition of protected information that depends on whether the information either does or does not identify a person.⁶³

In contrast to laws that protect all information that identifies a person, some laws aim to provide a more bright-line standard, by setting forth an exhaustive list of the types of information that the law protects. An example is the Massachusetts data security regulation, which defines “personal information” as a Massachusetts resident’s name in combination with his or her Social Security number, driver’s license number, or financial account number.⁶⁴ As another example, the HIPAA Privacy Rule provides a safe harbor, by which data can be shared widely once all information from a list of eighteen categories of information have been removed.⁶⁵

Beyond a small subset of such laws that attempt to take such a bright-line approach, privacy laws in the U.S. generally employ standards that require case-by-case determinations that rely on some degree of interpretation. These determinations are complicated by advances in analytical capabilities, the increased availability of data about individuals, and advances in the scientific understanding of privacy. These developments have created substantial uncertainty in determining whether a specific piece of information does in fact identify a person in practice.⁶⁶

In combination with limited guidance on interpreting and applying regulatory standards for privacy, these factors have led individual actors who manage personal data to incorporate a wide range of different standards and practices for privacy protection.⁶⁷ Recognizing the need for case-specific determinations, the HIPAA Privacy Rule, for example, provides an alternative approach to de-identification that allows data to be shared pursuant to (1) an expert’s determination that “generally accepted statistical and scientific principles and methods for rendering information not individually identifiable” have been applied and

sectoral approach to privacy, see generally Paul M. Schwartz, *Preemption and Privacy*, 118 *YALE L.J.* 902 (2009).

63. See Schwartz & Solove, *supra* note 9, at 1829–30.

64. See 201 C.M.R. § 17.02 (West 2018).

65. See 45 C.F.R. § 164.514 (2017). Note, however, that HIPAA’s safe harbor standard creates ambiguity by requiring that the entity releasing the data “not have actual knowledge that the information could be used alone or in combination with other information to identify an individual who is a subject of the information.” *Id.*

66. See Narayanan & Shmatikov, *supra* note 12, at 26 (discussing successful re-identification of AOL search queries, movie viewing histories, social network data, and location information and concluding that the distinction between identifying information and non-identifying information is meaningless as “any attribute can be identifying in combination with others”).

67. See, e.g., Benjamin C.M. Fung et al., *Privacy-Preserving Data Publishing: A Survey of Recent Developments*, 42 *ACM COMPUTING SURVEYS*, June 2010, at 42 (2010).

(2) provision of documentation that “the risk is very small that the information could be used, alone or in combination with other reasonably available information, by an anticipated recipient to identify an individual who is a subject of the information.”⁶⁸ Practitioners frequently comment on the ambiguity of these standards and the lack of clarity in interpreting definitions such as PII.⁶⁹

In a 2012 survey of commentary on the U.S. legal framework for privacy protection, the Federal Trade Commission (“FTC”) concluded that “the traditional distinction between PII and non-PII has blurred” and that it is appropriate to more comprehensively examine data to determine the data’s privacy implications.⁷⁰ In response, the FTC, which has authority to bring enforcement actions against companies that engage in unfair and deceptive trade practices including practices related to data privacy and security, takes a different approach. The FTC has developed a privacy framework that applies to commercial entities that collect or use consumer data that “can be reasonably linked to a specific consumer, computer, or other device.”⁷¹ FTC guidance has set forth a three-part test for determining whether data are “reasonably linkable” under this standard. To demonstrate that data are not reasonably linkable to an individual identity, a company must (1) take reasonable measures to ensure the data are de-identified, (2) publicly commit not to try to re-identify the data, and (3) contractually prohibit downstream recipients from attempting to re-identify the data.⁷² The first prong is satisfied when there is “a reasonable level of justified confidence that the data cannot reasonably be used to infer information about, or otherwise be linked to, a particular consumer, computer, or other device.”⁷³ Noting that it will follow the flexible standard that it follows in data security cases,⁷⁴ the FTC clarifies that “what qualifies as a reasonable level of justified confidence depends upon the particular circumstances, including the available methods and technologies,” as well as “the nature of the data at issue and the purposes for which it will be used.”⁷⁵ The FTC notes that various technical approaches can be used to satisfy this standard, and that it “encourages companies and researchers to continue innovating in the development and evaluation of new and better approaches to deidentification. FTC staff will continue to monitor and

68. 45 C.F.R. § 164.514(b). The Department of Health & Human Services has declined to provide specific instructions for carrying out an expert determination. OFFICE FOR CIVIL RIGHTS, DEP’T OF HEALTH & HUMAN SERVS., *supra* note 4, at 10.

69. For example, the 2008 rulemaking to update FERPA acknowledged the confusion expressed by commentators regarding the potential applicability of the law’s definition of PII. See Family Educational Rights and Privacy, 73 Fed. Reg. 74,806, 74,830–31 (Dec. 9, 2008).

70. FED. TRADE COMM’N, *supra* note 24, at 2.

71. *Id.* at 15.

72. *See id.* at 21.

73. *Id.*

74. *See id.* at 21 n.108.

75. *Id.* at 21.

assess the state of the art in de-identification.”⁷⁶ As such, the FTC’s approach is likely to evolve over time in response to the development of new technologies for privacy protection.

It is noteworthy that privacy law in the European Union relies on a definition of personal data that is broader than corresponding definitions in the United States. The new EU General Data Protection Regulation (“GDPR”) defines “personal data” as “any information relating to a data subject,” where a data subject is defined as a person who “can be identified, directly or indirectly, by means reasonably likely to be used.”⁷⁷ The GDPR also distinguishes between “pseudonymous” information, which “could be attributed to a natural person by the use of additional information,” and “anonymous” information, which is “namely information which does not relate to an identified or identifiable natural person or to personal data rendered anonymous in such a manner that the data subject is not or no longer identifiable.”⁷⁸ The GDPR does not apply to the processing of anonymous information, while pseudonymous information “should be considered to be information on an identifiable natural person.”⁷⁹

Recital 26 of the GDPR aims to clarify this standard, in articulating that “[t]o determine whether a natural person is identifiable, account should be taken of all the means reasonably likely to be used, such as singling out, either by the controller or by another person to identify the natural person directly or indirectly.”⁸⁰ In turn, when determining whether means are reasonably likely to be used for identification, “account should be taken of all objective factors, such as the costs of and the amount of time required for identification, taking into consideration the available technology at the time of the processing and technological developments.”⁸¹

The wide variation among legal definitions of personal information and how they have been interpreted, the gaps created by the narrowness of their scope (particularly within the U.S. framework), ambiguity regarding the context-specific applicability along the boundaries, and the dynamic nature of the definitions and their interpretation in response to technological developments over time, create challenges for demonstrating that the use of a privacy-preserving technology is sufficient to satisfy legal requirements for privacy protection. In the language of the selected laws introduced in this Section, this Article aims to overcome

76. *Id.*

77. Council Regulation (EU) 2016/679, art. 4, of the European Parliament and of the Council of 27 April 2016 on the Protection of Natural Persons with Regard to the Processing of Personal Data and on the Free Movement of Such Data, and repealing Directive 95/46/EC (General Data Protection Regulation), 2016 O.J. (L 119) 1, 33 [hereinafter GDPR].

78. *Id.* at 5.

79. *Id.*

80. *Id.*

81. *Id.*

these challenges by proposing an approach that can potentially be used to formalize legal definitions such as “information which identifies a person;”⁸² “a reasonable level of justified confidence that the data cannot reasonably be used to infer information about, or otherwise be linked to, a particular consumer, computer, or other device;”⁸³ and “a person who can be identified, directly or indirectly, by means reasonably likely to be used”⁸⁴ so that a privacy technology’s compliance with the regulations can be rigorously demonstrated. The proposed approach could also be used by regulators and advisory bodies in future assessments regarding whether an emerging privacy technology satisfies regulatory requirements.

D. Related Research to Bridge Between Legal and Computer Science Approaches to Privacy

Numerous research directions have previously sought to address the gaps between legal and technical approaches to privacy protection. Most closely related to this Article is work by Haney et al., who modeled the legal requirements for data collected by the U.S. Census Bureau.⁸⁵ Their research proposes formal privacy definitions to match the requirements of Title 13 of the U.S. Code, which protects information furnished to the Census Bureau by individuals and establishments.⁸⁶ In order to specify a formal model of Title 13’s privacy requirements, Haney et al. introduce a technical definition of privacy that is a variant of differential privacy and that is informed by an understanding of the Census Bureau Disclosure Review Board’s historical interpretations of Title 13. The privacy definition they present goes beyond protecting the privacy of individuals, which is inherent to the differential privacy definition, to also protect quantitative establishment information from being inferred with a level of accuracy that is “too high.”

While similar, the analysis in this Article differs from that explored by Haney et al. in significant ways. First, the model of FERPA in this Article aims to capture a large class of interpretations of the law in order to address potential differences in interpretation, whereas the definition presented by Haney et al. adopts a singular understanding of the Census Bureau’s Disclosure Review Board’s interpretation of Title 13. In addition to being a narrower model of legal requirements, it is one that relies on internal analysis by the Census Bureau Disclosure Review

82. Video Privacy Protection Act, 18 U.S.C. § 2710(a)(3) (2013).

83. FED. TRADE COMM’N, *supra* note 24, at 21.

84. GDPR, *supra* note 77, at 33.

85. Samuel Haney et al., *Utility Cost of Formal Privacy for Releasing National Employer-Employee Statistics*, 2017 ACM SIGMOD/PODS CONF. 1339, 1342–46.

86. *See id.* at 1342.

Board, rather than drawing directly from the text of statutes, regulations, and publicly available policy. Whereas Haney et al. describe specific computations that meet the definition they introduce, the technical analysis in this Article shows that a rich class of computations (i.e., differentially private computations) meet the privacy requirements of the model. Finally, this Article aims to introduce concepts from the technical literature on privacy and to describe an approach to combining legal and technical analysis in a way that is accessible to a legal audience.

This Article also relates to a line of work that encodes legal requirements for privacy protection using formal logic in order to verify the compliance of technical systems.⁸⁷ A prominent line of research uses Nissenbaum's framework of contextual integrity,⁸⁸ which models privacy norms in the flow of information between agents in a system.⁸⁹ This framework is used to extract norms from regulatory requirements for privacy protection that can be encoded using the language of formal logic.⁹⁰ This approach, which this Article will refer to as the *formal logic model*, has been used to formalize large portions of laws such as the Gramm-Leach-Bliley Act and the Health Insurance Portability and Accountability Act.⁹¹

The analysis in this Article differs from the formal logic model in several substantive ways. First, this Article relies on a different type of formalization, i.e., a privacy game instead of logic formulae. Second, the formal logic model typically avoids directly addressing ambiguity in the law by using under-specified predicates and making the assumption that their exact unambiguous meaning can be defined exogenously.⁹² The model does not specify how to determine whether these predicates actually hold.⁹³ In contrast, this Article aims to address am-

87. See Paul N. Otto & Annie I. Antón, *Addressing Legal Requirements in Requirements Engineering*, 15 PROC. IEEE INT'L REQUIREMENTS ENGINEERING CONF. 5, 7–11 (2007) (highlighting a survey of logic-based approaches, as well as other approaches).

88. See, e.g., Adam Barth et al., *Privacy and Contextual Integrity: Framework and Applications*, 27 PROC. IEEE SYMP. ON SECURITY & PRIVACY 184, 187–189 (2006).

89. HELEN NISSENBAUM, *PRIVACY IN CONTEXT: TECHNOLOGY, POLICY, AND THE INTEGRITY OF SOCIAL LIFE* (2010).

90. More specifically, the requirements are encoded in a first-order logic extended with primitive operators for reasoning about time. See, e.g., Barth et al., *supra* note 88, at 187–89.

91. See DeYoung et al., *supra* note 18, at 76–80.

92. For instance, a model might want to restrict the transmission of a message *m* about an individual *q* if *m* contains the attribute *t* and *t* is considered non-public information. In doing so, it might use predicates such as **contains**(*m*, *q*, *t*) and *t* ∈ **npi**. Predicates are functions that evaluate to either True or False. The predicates given as examples here are drawn from Barth et al., *supra* note 88, at 191.

93. The model does not specify when a message *m* contains an attribute *t* about individual *q* and when attribute *t* is non-public information. This is an intentional modeling decision as is stated explicitly in the seminal paper on this approach: “Much of the consternation about

biguity directly by modeling the requirements of the law very conservatively so as to capture a wide-range of reasonable interpretations. Third, the formal logic model uses predicates, which are functions that evaluate to either true or false (but to no other “intermediate” value). For example, the logic model may use a predicate that evaluates to true when a message contains information about a specific attribute of an individual (e.g., whether Jonas failed the standardized math exam). In contrast, this Article’s proposed model attempts to also account for cases in which an adversary can use a message to infer partial information about an attribute.⁹⁴ Fourth, the formal logic model supports reasoning about inference through explicit rules,⁹⁵ which only accounts for types of inferences anticipated by the authors of the model. In contrast, this Article’s proposed model does not limit the type of inferences that agents can make. Finally, formal logic models might specify that a doctor may share a specific patient’s private medical information with that patient or that researchers may publish aggregate statistics based on private medical information. However, such models do not enable expressing restrictions on the communication of aggregate statistics such as “the average salary of bank managers can be released only if it does not identify a particular individual’s salary.”⁹⁶ The latter is precisely the type of restriction this Article aims to capture through the proposed model; it focuses on modeling the degree to which an adversary should be prevented from inferring private information about an individual from an aggregate statistic.

III. INTRODUCTION TO TWO PRIVACY CONCEPTS: DIFFERENTIAL PRIVACY AND FERPA

This Article demonstrates an approach for bridging technical and regulatory privacy concepts. In order to illustrate its use, this Article focuses on two specific real-world privacy concepts — one technical,

GLBA [Gramm-Leach-Bliley Act] revolves around the complex definition of which companies are affiliates and what precisely constitutes non-public personal information. Our formalization of these norms sidesteps these issues by taking the role *affiliate* and the attribute *npi* to be defined exogenously.” *Id.* at 192 (citation omitted).

94. Technically, this inference results in a change in the probability that the adversary correctly guesses the value of the attribute, without necessarily empowering the adversary observer to guess correctly with certainty.

95. These rules are in the form (T, t) , where T is a set of attributes and t is a specific attribute: if an agent knows the value of every attribute in T about some individual q , then the agent also knows the value of attribute t about q . For instance, the rule $(\{\text{weight, height}\}, \text{BMI})$ specifies that if an agent knows the weight and height of some individual, then the agent also knows the body mass index of that individual. See Barth et al., *supra* note 88, at 185.

96. *Id.* at 184.

differential privacy, and the other regulatory, FERPA. This Section introduces the two concepts and sets forth the definitions that will form the basis of the analysis that will follow in later sections.

This Article relies on differential privacy because its rich and developed theory can serve as an initial subject of an examination of how formal notions of privacy can be compared to regulatory standards. In addition, demonstrating that differential privacy is in accordance with relevant laws may be essential for some practical uses of differential privacy.

A. Differential Privacy

Differential privacy is a nascent privacy concept that has emerged in the theoretical computer science literature in response to evidence of the weaknesses of commonly used techniques for privacy protection.⁹⁷ Differential privacy, first presented in 2006, is the result of ongoing research to develop a privacy technology that provides robust protection even against unforeseen attacks. Differential privacy by itself is not a single technological solution but a definition — sometimes referred to as a standard — that states a concrete requirement. Technological solutions are said to satisfy differential privacy if they adhere to the definition. As a strong, quantitative notion of privacy, differential privacy is provably resilient to a very large class of potential data misuses. Differential privacy therefore represents a solution that moves beyond traditional approaches to privacy — which require modification as new vulnerabilities are discovered.

1. The Privacy Definition and Its Guarantee

This Section offers an intuitive view of the privacy guarantee provided by differentially private computations.⁹⁸ Consider a hypothetical individual John who may participate in a study exploring the relationship between socioeconomic status and medical outcomes. All participants in the study must complete a questionnaire encompassing topics such as their location, their health, and their finances. John is aware that re-identification attacks have been performed on de-identified data. He is concerned that, should he participate in this study, some sensitive information about him, such as his HIV status or annual income, might be revealed by a future analysis based in part on his responses to the questionnaire. If leaked, this personal information could embarrass

97. For a discussion of differential privacy in this context, see Nissim et al., *supra* note 6. For the literature on differential privacy, see sources cited *supra* note 6.

98. For the mathematical definition, see Dwork et al., *supra* note 6. For a less technical presentation, see Nissim et al., *supra* note 6.

him, lead to a change in his life insurance premium, or affect the outcome of a future bank loan application.

If a computation of the data from this study is differentially private, then John is guaranteed that the outcome of the analysis will not disclose anything that is specific to him even though his information is used in the analysis. To understand what this means, consider a thought experiment, which is referred to as John's *privacy-ideal scenario* and is illustrated in Figure 3. John's privacy-ideal scenario is one in which his personal information is omitted but the information of all other individuals is provided as input as usual. Because John's information is omitted, the outcome of the computation cannot depend on John's specific information.

Differential privacy aims to provide John with privacy protection in the real-world scenario that approximates his privacy-ideal scenario. Hence, what can be learned about John from a differentially private computation is essentially limited to what could be learned about him from everyone else's data *without him being included in the computation*. Crucially, this very same guarantee is made not only with respect to John, but also with respect to every other individual contributing his or her information to the analysis.

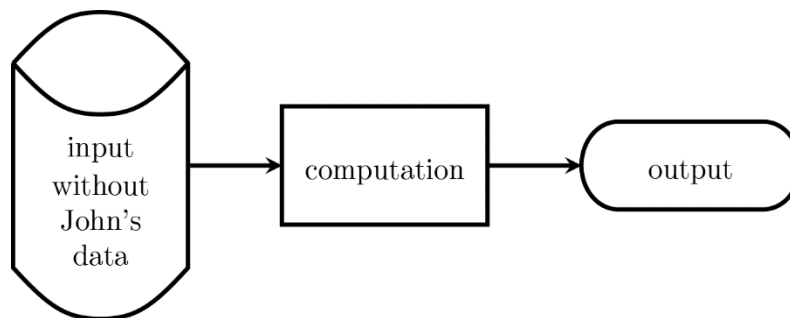


Figure 3: John's privacy-ideal scenario.

A parameter quantifies and limits the extent of the deviation between the privacy-ideal and real-world scenarios. As shown in Figure 4 below, this parameter is usually denoted by the Greek letter ϵ (epsilon) and is referred to as the *privacy parameter*, or, more accurately, the *privacy loss parameter*. The parameter ϵ measures the effect of each individual's information on the output of the analysis. It can also be viewed as a measure of the additional privacy risk an individual could incur beyond the risk incurred in the privacy-ideal scenario.⁹⁹

99. For a more detailed, technical discussion of how the privacy parameter ϵ controls risk, see Nissim et al., *supra* note 6, at 12.

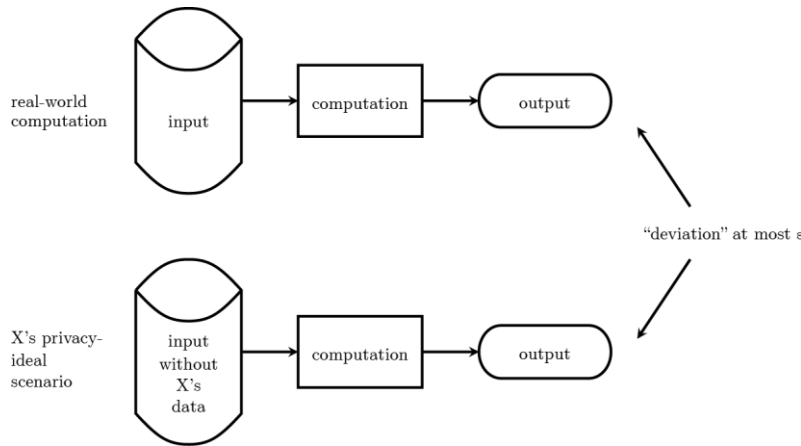


Figure 4: Differential privacy.

Note that in Figure 4 John is replaced with a prototypical individual X to emphasize that the differential privacy guarantee is made simultaneously to *all* individuals in the sample. The maximum deviation between the privacy-ideal scenario and real-world computation should hold simultaneously for each individual X whose information is included in the input.

It is important to note that differential privacy does not guarantee that an observer will not be able to learn anything about John from the outcome of the survey. Consider an observer, Alice, who possesses prior knowledge of some information about John, such as that he regularly consumes a lot of red wine. If the study reports a correlation between drinking red wine and the occurrence of a certain type of liver disease, Alice might conclude that John has a heightened risk of liver disease. However, Alice would be able to draw the conclusion that he has a heightened liver disease risk just as other red wine drinkers do even if information about John is not used in the study. In other words, this risk is present in both John's privacy-ideal scenario and his real-world scenario.

John may be adversely affected by the discovery of the results of a differentially private computation — for example, if sales of red wine were made illegal as a result of the discovery. The guarantee is that such harm is not due to the presence of John's data in particular.

2. Differential Privacy in the Real World

Despite being a relatively new concept, differential privacy has already found use in several real-world applications, and more applications are currently under development. The U.S. Census Bureau makes available an online interface for exploring the commuting patterns of workers across the United States, using confidential data collected through the Longitudinal Employer-Household Dynamics program over the period of 2002–2014.¹⁰⁰ Users of this interface interact with synthetic datasets that have been carefully generated from confidential agency records. The computations used to synthesize the data provide formal privacy guarantees and meet a variant of differential privacy.¹⁰¹

Several businesses are also experimenting with differentially private applications. For example, a system developed and implemented by Google employs differentially private computations to collect information from users of the company's Chrome web browser, in order to gather statistics used to monitor how unwanted software hijacks the browser settings of their users.¹⁰² This application allows analysts at Google to study trends present in the extensive Chrome user base, with strong guarantees that not much can be learned that is specific to any individual user.¹⁰³ Other companies, such as Apple, have also deployed implementations of differential privacy, as interest in differential privacy continues to grow in the private sector.¹⁰⁴

The academic community is also in the process of developing practical platforms for performing differentially private analyses. The Putting Differential Privacy to Work project at the University of Pennsylvania strives to build a general system that enables the average programmer to employ differentially private computations in a range of applications.¹⁰⁵ As part of the Privacy Tools for Sharing Research Data project at Harvard University,¹⁰⁶ differential privacy will be integrated with TwoRavens,¹⁰⁷ a browser-based software interface for exploring

100. See Ctr. for Econ. Studies, U.S. Census Bureau, *supra* note 7.

101. See Ashwin Machanavajjhala et al., *Privacy: Theory Meets Practice on the Map*, 24 PROC. IEEE INT'L CONF. ON DATA ENGINEERING 277, 283–85 (2008).

102. See Úlfar Erlingsson, *Learning Statistics with Privacy, Aided by the Flip of a Coin*, GOOGLE RES. BLOG (Oct. 30, 2014), <https://googleresearch.blogspot.com/2014/10/learning-statistics-with-privacy-aided.html> [<https://perma.cc/DJF6-G4MR>]; Erlingsson et al., *supra* note 7.

103. Another example implementation of differential privacy is Apple's use of differential privacy in iOS 10. See Eland, *supra* note 7.

104. Andy Greenberg, *Apple's 'Differential Privacy' Is About Collecting Your Data—But Not Your Data*, WIRED (June 13, 2016, 7:20 PM), <https://www.wired.com/2016/06/apples-differential-privacy-collecting-data/> [<https://perma.cc/RSB4-ZBDM>].

105. See *Putting Differential Privacy to Work Project*, *supra* note 7.

106. See *Harvard University Privacy Tools Project*, HARV. U., <https://privacytools.seas.harvard.edu> (last visited May 5, 2018).

107. See Inst. for Quantitative Soc. Sci., *TwoRavens*, GITHUB, <https://github.com/tworavens/tworavens> (last visited May 5, 2018).

and analyzing data that are hosted on the Dataverse repository platform,¹⁰⁸ which is currently used by institutions throughout the world. This will allow researchers to see the results of statistical analyses created by differentially private computations on sensitive datasets, including datasets that cannot otherwise be shared widely due to privacy concerns.

B. The Family Educational Rights and Privacy Act of 1974 (FERPA)

FERPA is a U.S. federal law requiring the protection of personal information contained in education records.¹⁰⁹ Education records are defined as records that directly relate to a student¹¹⁰ and are maintained by an educational agency or institution that receives funding under a program administered by the U.S. Department of Education.¹¹¹ Entities covered by FERPA include elementary and secondary schools, school districts, colleges and universities, state educational agencies, and other institutions providing educational services or directing institutions that do.¹¹²

FERPA provides parents with certain rights with respect to personal information contained in their child's education records, rights that transfer to an eligible student upon turning 18.¹¹³ These rights include the right to inspect, request amendment to, and consent to the disclosure of such personal information.¹¹⁴ FERPA distinguishes between two types of personal information contained in education records: directory information and non-directory PII.¹¹⁵ Generally, a parent or eligible student must provide written consent before an educational agency or institution can disclose non-directory PII from a student's education record.¹¹⁶ Information that a school has designated as

108. See Inst. for Quantitative Soc. Sci., *Dataverse*, GITHUB, <https://github.com/IQSS/dataverse> (last visited May 5, 2018).

109. Family Educational Rights and Privacy Act of 1974, 20 U.S.C. § 1232g (2012); Family Educational Rights and Privacy, 34 C.F.R. pt. 99 (2017).

110. See 20 U.S.C. § 1232g(a)(4).

111. See *id.* § 1232g(a)(1)(A).

112. See 34 C.F.R. § 99.1(a).

113. See *id.* § 99.3.

114. See *id.* §§ 99.10, 99.20, 99.30.

115. See *infra* Part III.

116. See 34 C.F.R. § 99.30(a). The term *disclosure* is defined broadly, meaning “to permit access to or the release, transfer, or other communication of personally identifiable information contained in education records by any means, including oral, written, or electronic means, to any party except the party identified as the party that provided or created the record.” *Id.* § 99.3. FERPA provides several exceptions to the written consent requirement. See, e.g., *id.* § 99.31(a)(1)(i)(A) (excepting from this requirement disclosures to school officials with a legitimate educational interest in the information). However, these exceptions are not a focus of the analysis in this Article because its goal is to analyze FERPA's requirements for protecting non-directory PII when publishing statistics based on such information. This analysis of the definition of non-directory PII is not affected by exceptions to FERPA allowing the full disclosure of such information to certain parties for specified purposes.

directory information can be disclosed without the consent of parents or eligible students, as long as they were provided with notice and an opportunity to opt out of the disclosure of directory information beforehand.

FERPA also permits the disclosure of de-identified information without consent “after the removal of all personally identifiable information provided that the educational agency or institution or other party has made a reasonable determination that a student’s identity is not personally identifiable, whether through single or multiple releases, and taking into account other reasonably available information.”¹¹⁷ By authorizing the disclosure of de-identified information without restriction, this provision enables the widespread publication and use of statistics on educational programs. According to the Department of Education guidance, this provision is intended to strike “an appropriate balance that facilitates the release of appropriate information for school accountability and educational research purposes while preserving the statutory privacy protections in FERPA.”¹¹⁸

The discussion below introduces the definitions of directory information, PII, and de-identified information as set forth by the FERPA regulations. These definitions form the basis of the formal model of FERPA’s privacy requirements in Part IV. The analysis also refers to guidance from the Department of Education explaining and interpreting the definitions found in the regulations.¹¹⁹ It is important to note that the Department of Education’s interpretations of the regulations may or may not qualify for controlling weight, depending on a number of factors.¹²⁰ However, this Article’s model of FERPA is based on the regulatory text itself. Interpretations from Department of Education

117. *Id.* § 99.31(b)(1).

118. Family Educational Rights and Privacy, 73 Fed. Reg. 74,806, 74,834 (Dec. 9, 2008).

119. In addition to the regulatory text, this analysis refers to guidance appearing in the preamble to the 2008 final rule to update the FERPA regulations, which provides an extended discussion of the definition of PII in justification of the latest revision to the definition. *See* Family Educational Rights and Privacy, 73 Fed. Reg. at 74,806–55. Note that, while the Department of Education subsequently promulgated rules in 2011, these rules do not amend the definitions of “personally identifiable information” or “de-identified information” set forth under FERPA, nor does the 2011 rulemaking provide guidance on interpreting these concepts. Therefore, these regulations are not pertinent to the analysis in this Article.

120. As a general rule, courts defer to an agency’s interpretation of its own regulations unless the interpretation is “plainly erroneous or inconsistent with the regulation.” *Bowles v. Seminole Rock & Sand Co.*, 325 U.S. 410, 414 (1945); *accord* *Auer v. Robbins*, 519 U.S. 452, 461 (1997). However, this rule is limited in several ways. *See* *Christopher v. SmithKline Beecham Corp.*, 567 U.S. 142, 155 (2012). In such cases, the agency’s interpretation of its regulations is persuasive rather than controlling. *See* *Skidmore v. Swift & Co.*, 323 U.S. 134, 140 (1944). Guidance that appears in the preamble to a final rule, such as the Department of Education’s guidance in the preamble to the 2008 final rule interpreting definitions from the amended FERPA regulations that is referenced in this Article, arguably merits greater weight than the many other sources of agency guidance. *See* Kevin M. Stack, *Preambles as Guidance*, 84 GEO. WASH. L. REV. 1252, 1272–77 (2016).

guidance appear in the discussion only for the purposes of demonstrating that the conservative modeling choices capture these interpretations among others falling within a broad class of reasonable interpretations of FERPA.

1. The Applicability of FERPA's Requirements to Formal Privacy Models Such as Differential Privacy

FERPA's protections likely apply to the release of statistics, and in particular to releases produced by methods that satisfy formal privacy models such as differential privacy. Therefore, it is important to understand exactly how information privacy laws such as FERPA would govern the use of new technologies based on formal privacy models, in anticipation of the practical implementation of these technologies. This Section points to sources that highlight the ways in which FERPA arguably applies to the release of differentially private statistics. This discussion anticipates and sets the stage for later sections that aim to interpret this language more formally. Although it is plausible that some legal scholars would view the release of aggregate statistics or synthetic data as falling outside the scope of FERPA, this analysis takes the more conservative view that FERPA does apply to the release of aggregate statistics in order to ensure the analysis is valid despite possible ambiguity regarding FERPA's scope of applicability.

Generally, FERPA governs the disclosure of non-directory PII about students in education records maintained by educational agencies and institutions. Here, "disclosure" is defined broadly, meaning "to permit access to or the release, transfer, or other communication of personally identifiable information contained in education records by any means, including oral, written, or electronic means, to any party except the party identified as the party that provided or created the record."¹²¹ PII is also defined broadly to include "information that, alone or in combination, is linked or linkable to a specific student that would allow a reasonable person in the school community, who does not have personal knowledge of the relevant circumstances, to identify the student with reasonable certainty."¹²² Because, as explained in Section II.A.3, it has been demonstrated that releases of aggregate statistics can leak some information about individuals, these definitions, which cover any communication of any information linkable to a specific student with reasonable certainty, are arguably written broadly enough to encompass privacy risks associated with statistical releases.

The preamble to the 2008 final rule updating the FERPA regulations also supports the conclusion that FERPA applies to releases of

121. 34 C.F.R. § 99.3.

122. *Id.* For further discussion of the definition of PII, see *infra* Part III.

statistical data that adhere to formal privacy models like differential privacy. The Department of Education's rationale in revising FERPA's definition of PII explicitly addresses privacy issues in releases of statistics. In the preamble, the Department refers to the capability of "re-identifying statistical information or redacted records."¹²³ By recognizing the privacy risks associated with both statistical information and redacted records, it instructs educational agencies and institutions to consider the privacy risks in both aggregate and individual-level data releases.¹²⁴ The preamble provides specific examples to illustrate some of the privacy risks associated with the release of statistical information.¹²⁵ The Department of Education notes, for example, that "a school may not release statistics on penalties imposed on students for cheating on a test where the local media have published identifiable information about the only student . . . who received that penalty," explaining that "statistical information or redacted record is now personally identifiable to the student or students because of the local publicity."¹²⁶ It also explains how the publication of a series of tables about the same set of students, with the data broken down in different ways, can, in combination, reveal PII about individual students.¹²⁷ In addition, the guidance notes that educational institutions are prohibited from reporting that 100% of students achieved specified performance levels, as a measure to prevent the leakage of PII.¹²⁸ These references from the preamble to the 2008 final rule are evidence that the Department of Education recognizes privacy risks associated with the release of aggregate statistics.

Educational agencies and institutions share statistics from education records with other agencies, researchers, and the public, for the purposes of research and evaluation. Indeed, such institutions must by law release certain education statistics in order to be eligible for some federal education funding,¹²⁹ unless such statistics would reveal PII as defined by FERPA.¹³⁰ For instance, educational agencies and institutions are prohibited from releasing tables containing statistics on groups of individuals falling below some minimum size.¹³¹ It is important to

123. Family Educational Rights and Privacy, 73 Fed. Reg. at 74,831.

124. *See id.* at 74,835.

125. *See id.*

126. *Id.* at 74,832.

127. *See id.* at 74,835.

128. *See id.* at 74,835–36.

129. *See* 20 U.S.C. § 6311(h) (2012).

130. Family Educational Rights and Privacy, 34 C.F.R. § 200.17(b)(1) (2017).

131. NAT'L CTR. FOR EDUC. STATISTICS, U.S. DEP'T OF EDUC., STATISTICAL METHODS FOR PROTECTING PERSONALLY IDENTIFIABLE INFORMATION IN AGGREGATE REPORTING 1 (2011), <https://nces.ed.gov/pubs2011/2011603.pdf> [<https://perma.cc/G5J9-XEPZ>] (noting that "[i]ndividual states have adopted minimum group size reporting rules, with the minimum number of students ranging from 5 to 30 and a modal category of 10 (used by 39 states in the most recent results available on state websites in late winter of 2010).").

note, however, that the Department of Education does not consider specific procedures, such as adherence to minimum cell sizes, to be sufficient to meet FERPA's privacy requirements in all cases.¹³²

Additionally, the Department of Education's National Center for Education Statistics has released implementation guidance devoted to helping such institutions apply privacy safeguards in accordance with FERPA when releasing aggregate statistics.¹³³ This guidance aims to clarify the goal of FERPA in the context of aggregate data releases:

Protecting student privacy means publishing data only in a manner that does not reveal individual students' personally identifiable information, either directly or in combination with other available information. Another way of putting this is that the goal is to publish summary results that do not allow someone to learn information about a specific student.¹³⁴

This guidance further "demonstrates how disclosures occur even in summary statistics,"¹³⁵ discusses how common approaches to privacy may fall short of FERPA's standard for privacy protection,¹³⁶ and provides some best practices for applying disclosure limitation techniques in releases of aggregate data.¹³⁷ This practical guidance is further evidence that the agency recognizes some leakages of information about individuals from aggregate data releases to be FERPA violations.

Not only does the Department of Education require educational agencies and institutions to address privacy risks in the release of aggregate statistics, but the scientific literature on privacy also suggests this is a prudent approach. Numerous attacks have demonstrated that it is often possible to link particular individuals to information about them in aggregate data releases.¹³⁸ Moreover, the potential leakage of PII through releases of aggregate statistics is a concern that is anticipated

132. "[I]t is not possible to prescribe or identify a single method to minimize the risk of disclosing personally identifiable information . . . that will apply in every circumstance, including determining whether defining a minimum cell size is an appropriate means . . . and, if so, selection of an appropriate number." Family Educational Rights and Privacy, 73 Fed. Reg. at 74,835.

133. See NAT'L CTR. FOR EDUC. STATISTICS, U.S. DEP'T OF EDUC., *supra* note 131, at 27–29.

134. *Id.* at 4. Note that one could simply adopt this interpretation and argue that the goal of FERPA's privacy requirements is "not [to] allow someone to learn information about a specific student," *id.*, as this language arguably directly implies the differential privacy definition. However, this Article does not adopt this singular interpretation and instead presents an argument that captures a significantly wider class of potential interpretations of FERPA's requirements, thereby strengthening the argument.

135. *Id.*

136. See *id.* at 7–13.

137. See *id.* at 14–26.

138. See *supra* Part II.

to evolve and grow over time, as analytical capabilities advance and the availability of large quantities of personal information from various sources increases. It is likely that the approaches identified in current agency guidance will in the future be shown not to provide adequate privacy protection, while also significantly limiting the usefulness of the data released, requiring future updates to the guidance.

For instance, FERPA defines de-identified information to mean that “the educational agency or institution or other party has made a reasonable determination that a student’s identity is not personally identifiable, whether through single or multiple releases, and taking into account other reasonably available information.”¹³⁹ In interpreting this language, the preamble to the 2008 final rule explicitly recognizes that privacy risk from data disclosures is cumulative, and accordingly requires educational agencies and institutions to take into account the accumulated privacy risk from multiple disclosures of information:

The existing professional literature makes clear that public directories and previously released information, including local publicity and even information that has been de-identified, is sometimes linked or linkable to an otherwise de-identified record or dataset and renders the information personally identifiable. The regulations properly require parties that release information from education records to address these situations.¹⁴⁰

However, the agency does not provide guidance on addressing cumulative privacy risks from successive disclosures. Rather, it notes that “[i]n the future [it] will provide further information on how to monitor and limit disclosure of personally identifiable information in successive statistical data releases.”¹⁴¹ Indeed, research from the computer science literature demonstrates that it is very difficult to account for the cumulative privacy risk from multiple statistical releases.¹⁴² This suggests that Department of Education guidance is likely to evolve over time in response to new understandings of privacy risk, particularly with respect to the risk that accumulates from multiple data releases. It also lends support to the claim that use of differential privacy is sufficient to satisfy FERPA’s requirements, as differential privacy provides provable guarantees with respect to the cumulative risk from successive data

139. Family Educational Rights and Privacy, 34 C.F.R. § 99.31(b)(1) (2017).

140. Family Educational Rights and Privacy, 73 Fed. Reg. 74,806, 74,831 (Dec. 9, 2008).

141. *Id.* at 74,835.

142. See Srivatsava Ranjit Ganta, Shiva Prasad Kasiviswanathan & Adam Smith, *Composition Attacks and Auxiliary Information in Data Privacy*, 14 PROC. ACM SIGKDD INT’L CONF. ON KNOWLEDGE DISCOVERY & DATA MINING 265, 266 (2008).

releases and, to the authors' knowledge, is the only known approach to privacy that provides such a guarantee.

2. The Distinction Between Directory and Non-Directory Information

FERPA distinguishes between directory information and non-directory PII. Directory information is defined as "information contained in an education record of a student that would not generally be considered harmful or an invasion of privacy if disclosed."¹⁴³ Each educational agency or institution produces a list of categories of information it designates as directory information. The regulations provide some categories of information for illustration of the types of information an educational agency or institution may designate as directory information, such as a student's name, address, telephone number, e-mail address, photograph, and date and place of birth.¹⁴⁴ Note that many of these examples of directory information overlap with FERPA's definition of PII. However, the law does not require consent for the disclosure of directory information — even directory information constituting PII — as long as the relevant educational agency or institution has provided parents and eligible students with public notice of the types of information it has designated as directory information as well as an opportunity to opt out certain information being designated as disclosable directory information.¹⁴⁵ An educational agency or institution may also disclose directory information about former students without reissuing the notice and opportunity to opt out.¹⁴⁶

In contrast, educational agencies and institutions must take steps to protect non-directory PII from release. This category of information can only be disclosed without consent under certain exceptions.¹⁴⁷ In addition, information from education records that would otherwise be considered non-directory PII can be released to the public, without consent, if it has been rendered de-identified, meaning "the educational agency or institution or other party has made a reasonable determination that a student's identity is not personally identifiable, whether through single or multiple releases, and taking into account other reasonably available information."¹⁴⁸

When releasing data, educational agencies and institutions must assess the privacy-related risks in light of publicly available information,

143. 34 C.F.R. § 99.3.

144. *See id.*

145. *See id.* § 99.37(a).

146. *See id.* § 99.37(b).

147. *See, e.g., id.* § 99.31(a)(6)(i)(B) (providing an exception for the sharing of data for the purposes of administering student aid programs).

148. *Id.* § 99.31(b)(1).

including directory information.¹⁴⁹ Because the scope of directory information varies from school to school, knowing what information may be available to a potential adversary is uncertain, creating a challenge for an educational agency or institution assessing the privacy risks associated with a planned data release. Part IV below proposes an approach to formally modeling directory information and privacy risks in the release of education data more generally, which addresses the ambiguity created by differences in what may be classified as directory information across different educational agencies and institutions, currently and in the future. Specifically, the model addresses this ambiguity by making an assumption that a potential privacy attacker is given unrestricted access to information that does not require consent for release, including all potential directory information.

The definition of non-directory PII and how it has been interpreted by the Department of Education serves as the basis for the formal model of FERPA's requirements constructed in this Article. This requires a careful examination of the definition of PII set forth by the regulations and its interpretation in agency guidance.

3. The Definition of Personally Identifiable Information

FERPA defines PII by way of a non-exhaustive list of categories of information included within the definition. The definition includes, but is not limited to:

- (a) The student's name; (b) The name of the student's parent or other family members; (c) The address of the student or student's family; (d) A personal identifier, such as the student's social security number, student number, or biometric record; (e) Other indirect identifiers, such as the student's date of birth, place of birth, and mother's maiden name; (f) Other information that, alone or in combination, is linked or linkable to a specific student that would allow a reasonable person in the school community, who does not have personal knowledge of the relevant circumstances, to identify the student with reasonable certainty; or (g) Information requested by a person who

149. The Department of Education observes that "the risk of re-identification may be greater for student data than other information because of the regular publication of student directories, commercial databases, and de-identified but detailed educational reports" Family Educational Rights and Privacy, 73 Fed. Reg. 74,806, 74,834 (Dec. 9, 2008).

the educational agency or institution reasonably believes knows the identity of the student to whom the education record relates.¹⁵⁰

This analysis focuses in particular on paragraph (f) of the above definition, taking it as the basis of the formal model of FERPA's privacy requirements presented in Part IV. Agency guidance from the preamble to the 2008 final rule and other documents provide interpretations of this definition.¹⁵¹ The definition, by referring to "a reasonable person in the school community"¹⁵² makes use of an objective reasonableness standard. By relying on this standard, the Department of Education states that it intended to "provide the maximum privacy protection for students" because "a reasonable person in the school community is . . . presumed to have at least the knowledge of a reasonable person in the local community, the region or State, the United States, and the world in general."¹⁵³ The agency also notes that the standard was not intended to refer to the "technological or scientific skill level of a person who would be capable of re-identifying statistical information or redacted records."¹⁵⁴ Rather, it refers to the knowledge a reasonable person might have, for example, "based on local publicity, communications, and other ordinary conditions."¹⁵⁵ The preamble also explains that FERPA's reasonableness standard is not to be interpreted as subjective or based on the motives or capabilities of a potential attacker.¹⁵⁶

In 2008, the Department of Education updated the definition of PII that appears in the FERPA regulations. Previously, it had included "information that would make a student's identity easily traceable" in lieu of clauses (e)–(g) found in the current definition.¹⁵⁷ The Department of Education explained that it removed the "easily traceable" language from the definition "because it lacked specificity and clarity" and "suggested that a fairly low standard applied in protecting education records, i.e., that information was considered personally identifiable only if it was easy to identify the student."¹⁵⁸

In providing guidance on interpreting the definitions of PII and de-identified information, the Department of Education acknowledged that

150. 34 C.F.R. § 99.3.

151. See Family Educational Rights and Privacy, 73 Fed. Reg. at 74,831; NAT'L CTR. FOR EDUC. STATISTICS, DEP'T OF EDUC., *supra* note 131, at 3.

152. 34 C.F.R. § 99.3. The preamble includes some examples of members of the school community, such as "students, teachers, administrators, parents, coaches, volunteers," and others at the local school. Family Educational Rights and Privacy, 73 Fed. Reg. at 74,832.

153. Family Educational Rights and Privacy, 73 Fed. Reg. at 74,832.

154. *Id.* at 74,831.

155. *Id.* at 74,832.

156. See *id.* at 74,831.

157. *Id.* at 74,829.

158. *Id.* at 74,831.

FERPA's privacy standard requires case-by-case determinations. It noted that the removal of "nominal or direct identifiers," such as names and Social Security numbers, "does not necessarily avoid the release of personally identifiable information."¹⁵⁹ Furthermore, the removal of other information such as address, date and place of birth, race, ethnicity, and gender, may not be sufficient to prevent one from "indirectly identify[ing] someone depending on the combination of factors and level of detail released."¹⁶⁰ However, the preamble declines "to list all the possible indirect identifiers and ways in which information might indirectly identify a student" (1) because "[i]t is not possible" and (2) "[i]n order to provide maximum flexibility to educational agencies and institutions."¹⁶¹ In this way, FERPA's privacy requirements, unlike the HIPAA Privacy Rule, "do not attempt to provide a 'safe harbor' by listing all the information that may be removed in order to satisfy the de-identification requirements."¹⁶² The preamble also emphasizes that de-identified information that is released could be linked or linkable to "public directories and previously released information, including local publicity and even information that has been de-identified," rendering it personally identifiable.¹⁶³ Ultimately, "determining whether a particular set of methods for de-identifying data and limiting disclosure risk is adequate cannot be made without examining the underlying data sets, other data that have been released, publicly available directories, and other data that are linked or linkable to the information in question."¹⁶⁴ The Department of Education therefore declines to "provide examples of rules and policies that necessarily meet the de-identification requirements," opting instead to leave the entity releasing the data "responsible for conducting its own analysis and identifying the best methods to protect the confidentiality of information from education records it chooses to release."¹⁶⁵

Examples provided in the preamble contribute to a lack of clarity around applying FERPA's privacy requirements. For instance, it is not clear how an educational agency or institution should differentiate between special knowledge and information known to a reasonable person in the school community. The agency provides the following example to illustrate how the language "personal knowledge of the relevant circumstances," found in paragraph (f) of the definition of PII, is to be interpreted:

159. *Id.*

160. *Id.*

161. *Id.* at 74,833.

162. *Id.*

163. *Id.* at 74,831.

164. *Id.* at 74,835.

165. *Id.*

[I]f it is generally known in the school community that a particular student is HIV-positive, or that there is an HIV-positive student in the school, then the school could not reveal that the only HIV-positive student in the school was suspended. However, if it is not generally known or obvious that there is an HIV-positive student in school, then the same information could be released, even though someone with special knowledge of the student's status as HIV-positive would be able to identify the student and learn that he or she had been suspended.¹⁶⁶

This example seems counterintuitive because it does not address whether members of the school community might know, or might in the future learn, that a particular student was suspended. Possession of such knowledge would enable one to learn that this student is HIV-positive. Enabling this type of disclosure through a release of information — highly sensitive information such as a student's HIV status, no less — is likely not what the agency intended. In another example, the agency notes that a student's nickname, initials, or personal characteristics are not personally identifiable “if teachers and other individuals in the school community generally would not be able to identify a student” using this information.¹⁶⁷ This seems to imply a weak privacy standard, as a student's initials, nickname, or personal characteristics are likely to be uniquely identifying in many cases, regardless of whether such characteristics are considered to be generally known within the community.

Ambiguity in interpreting this definition is also reflected in a discrepancy between the preamble and an interpretation of the regulation developed by the Privacy Technical Assistance Center (“PTAC”), a Department of Education contractor that develops guidance on complying with FERPA's requirements. In guidance interpreting the standard used to evaluate disclosure risk when releasing information from education records, PTAC advises that “[s]chool officials, including teachers, administrators, coaches, and volunteers, are not considered in making the reasonable person determination since they are presumed to have inside knowledge of the relevant circumstances and of the identity of the students.”¹⁶⁸ This interpretation from PTAC appears to be substantially weaker than the regulatory text and directly contradicts the language of the 2008 final rule, which interprets the regulations to

166. *Id.* at 74,832.

167. *Id.* at 74,831.

168. Privacy Tech. Assistance Ctr., *Frequently Asked Questions—Disclosure Avoidance*, DEP'T EDUC. (May 2013), https://studentprivacy.ed.gov/sites/default/files/resource_document/file/FAQs_disclosure_avoidance.pdf [<https://perma.cc/W56R-RMN8>].

protect non-directory PII from disclosure if it can be identified by this same constituency of the school community. Such individuals would seem to be reasonable persons in the school community based on the plain meaning of this language. In the preamble, the Department of Education provides the following example:

[I]t might be well known among students, teachers, administrators, parents, coaches, volunteers, or others at the local high school that a student was caught bringing a gun to class last month but generally unknown in the town where the school is located. In these circumstances, a school district may not disclose that a high school student was suspended for bringing a gun to class last month, even though a reasonable person in the community where the school is located would not be able to identify the student, because a reasonable person in the high school would be able to identify the student.¹⁶⁹

As discussed above, a court would likely give more weight to the agency's interpretation in the preamble to the final rule than to an agency contractor's subsequent interpretation. Nevertheless, this example illustrates the potential for alternative interpretations of FERPA's definition of PII.

Numerous commentators have likewise expressed uncertainty regarding interpretations of the privacy requirements of FERPA. For example, the preamble to the 2008 final rule refers to many public comments seeking clarification on the de-identification standard. Comments refer to the standard as being "too vague and overly broad," as the definition of PII "could be logically extended to cover almost any information about a student."¹⁷⁰ Other commenters question whether the standard provides privacy protection as strong as the agency intends, in light of concerns about the difficulty of de-identifying data effectively.¹⁷¹

In Part IV, this Article aims to overcome ambiguities in the regulatory requirements by modeling them formally, based on conservative, "worst-case" assumptions. This approach can be used to demonstrate that a privacy technology satisfies a large family of reasonable interpretations of FERPA's requirements. This formal model of FERPA is based on the regulatory definitions and, to a lesser extent, agency interpretations of these definitions. In particular, the model of the FERPA

169. Family Educational Rights and Privacy, 73 Fed. Reg. at 74,832.

170. *Id.* at 74,830.

171. *See id.* at 74,833–34.

standard aims to honor the Department of Education's intent to "provide the maximum privacy protection for students,"¹⁷² by providing protection against a strong adversary.¹⁷³ The model outlined below does not require a determination of the subjective judgment or inquiries into an attacker's motives and declines to presuppose the attacker's goal. In this way, the model is able to consider attackers with different goals, different pieces of outside knowledge about students, and different ways of using the information, consistent with the flexible, case-by-case approach taken by the FERPA regulations.

C. Gaps Between Differential Privacy and FERPA

The emergence of formal privacy models such as differential privacy represents a shift in the conceptualization of data privacy risks and ways to mitigate such risks. The FERPA regulations and guidance from the Department of Education were drafted, for the most part, prior to the development and practical implementation of formal privacy models. They are largely based on traditional approaches to privacy through, for example, their emphasis on the importance of removing PII from data prior to release. The regulatory approach therefore differs from the approach relied upon by formal privacy models, creating challenges for translating between the two notions. Several key differences between the two privacy concepts are outlined below in order to illustrate these gaps. Note, however, that this discussion is limited, in that it is an illustration of the challenges of applying interpretations of FERPA to implementations of differential privacy. It is not intended as a more general assessment or critique of the level of privacy protection provided by FERPA.

Overall scope of privacy protection. FERPA does not apply across the board to protect all types of data in all settings. Instead, it is designed to protect certain types of information in specific contexts. For instance, FERPA applies only to educational agencies and institutions and protects only certain types of information from education records known as non-directory PII. In addition, it primarily appears to address releases of information that could be used in record linkage attacks, or the linkage of a named individual to a record in a release of data, that leverage publicly available data. Differential privacy, in contrast, is designed to be broadly applicable, providing formal bounds on the leakage of *any* information about an individual, not just an individual's identity or non-directory PII. Given its comparatively narrow scope,

172. *Id.* at 74,832.

173. Such an adversary (1) is not constrained by a limited technological skill level; (2) possesses knowledge about students that a reasonable person in the school community may have; and (3) has motives that are unknown. *See id.* at 74,831–32.

FERPA's applicability to a more general definition of privacy like differential privacy is arguably unclear.

Range of attacks. While guidance from the Department of Education expresses its intent for the FERPA regulations to provide strong privacy protection that addresses a wide range of privacy risks, in effect the regulations seem to address a narrower category of attacks. By permitting the release of de-identified information, these provisions seem primarily aimed at addressing record linkage attacks that could, using certain data known to be available, enable the linkage of a named individual with a record in a released set of data. In contrast, differential privacy is a quantitative guarantee of privacy that is provably resilient to a very large class of potential data misuses. Its guarantee holds no matter what computational techniques or resources the adversary brings to bear. In this way, differential privacy protects against inference attacks and attacks unforeseen at the time the privacy-preserving technique is applied. Because the FERPA regulations and implementation guidance focus on certain types of known privacy attacks, it is not clear how the regulations apply to formal privacy models that provide more general protection, including protection against inference attacks and attacks currently unknown.

Scope of private information. FERPA's privacy requirements focus on PII. They draw a sharp binary distinction between non-directory PII and de-identified information, protecting the former but placing the latter outside of the scope of the regulations altogether. The regulatory definition of PII provides a non-exhaustive list of the categories of information that should be protected from release. Accordingly, a common practice historically has been for educational agencies and institutions to withhold or redact certain pieces of information, such as names, Social Security numbers, and addresses, when disclosing information from education records.¹⁷⁴ The literature recognizes, however, that privacy risks are not limited to certain categories of information; indeed, information not typically used for identification purposes can often be used to identify individuals in de-identified data.¹⁷⁵ For example, a small number of data points about an individual's characteristics,

174. See, e.g., Family Policy Compliance Office, Dep't of Educ., Letter to Miami University re: Disclosure of Information Making Student's Identity Easily Traceable (Oct. 19, 2004), <https://www2.ed.gov/policy/gen/guid/fpco/ferpa/library/unofmiami.html> [https://perma.cc/MHF9-2CAF] (noting that "the redaction of [certain] items of information (student's name, social security number, student ID number, and the exact date and time of the incident) may generally be sufficient to remove all 'personally identifiable information' under FERPA," but concluding that such redaction may be insufficient in light of previous releases of information about the same individuals).

175. Narayanan and Shmatikov make an even bolder statement: "[a]ny information that distinguishes one person from another can be used for re-identifying anonymous data." Narayanan & Shmatikov, *supra* note 12, at 26.

behavior, or relationships can be sufficient to identify an individual.¹⁷⁶ Moreover, a wide range of inferences of personal information about an individual, not just an individual's identity, can be made based on a release of de-identified information. Differential privacy takes this broader conception of private information into account by putting formal bounds on any leakage of any information about an individual from a system. How FERPA's binary conceptions of non-directory PII and de-identification apply to a formal privacy model, which bounds the incremental leakage of any information specific to an individual, is uncertain.

Form of data release. By relying on terminology such as PII and de-identification, the FERPA regulations seems to be written with microdata, or individual-level data, as their primary use case. To a lesser extent, guidance on interpreting FERPA refers to protecting information in releases of statistical tables, and it is limited to the specification of minimum cell sizes and related approaches from the traditional statistical disclosure limitation literature.¹⁷⁷ In addition, by referring explicitly to de-identification and permitting the disclosure of information from which categories of non-directory PII have been removed, FERPA appears to endorse heuristic de-identification techniques, such as redaction of pieces of information deemed to be direct or indirect identifiers. These references to de-identification, approaches to risk in microdata releases, and traditional disclosure limitation techniques are difficult to generalize to other types of techniques. For techniques that rely on formal privacy models like differential privacy, which address risk in non-microdata releases, it is not clear how FERPA's privacy requirements should be applied.

Scope of guidance. The FERPA regulations and implementation guidance focus on practices for protecting information in cases where the risks to individual privacy are clear. Consider, for example, FERPA's definition of PII, which includes a non-exhaustive list of identifiers such as names, addresses, Social Security numbers, dates of birth, places of birth, and mother's maiden names.¹⁷⁸ When de-identifying data prior to release, an educational agency or institution must take steps to suppress fields containing these categories of information. It is more difficult to determine how to protect other information also falling within this definition, such as "information that, alone or in combination, is linked or linkable to a specific student that would allow a

176. See *Credit Card Metadata*, *supra* note 42, at 537; *Human Mobility*, *supra* note 42, at 1376.

177. See Family Educational Rights and Privacy, 73 Fed. Reg. at 74,835 (directing educational agencies and institutions to consult the Federal Committee on Statistical Methodology's Statistical Policy Working Paper 22 for guidance on applying methods for protecting information in a data release).

178. Family Educational Rights and Privacy, 34 C.F.R. § 99.3 (2017).

reasonable person in the school community, who does not have personal knowledge of the relevant circumstances, to identify the student with reasonable certainty.”¹⁷⁹ Indeed, the Department of Education acknowledges this, stating that “[i]t is not possible, however, to list all the possible indirect identifiers and ways in which information might indirectly identify a student.”¹⁸⁰ Educational agencies and institutions are provided little guidance on the steps necessary for protecting privacy in accordance with FERPA beyond redacting certain categories of obvious identifiers. The guidance is particularly unclear regarding how to interpret this language with respect to statistical computations, which do not contain direct or indirect identifiers. Where guidance refers to aggregate data, it flags the disclosure risks associated with the release of statistics derived from as few as one or two individuals.¹⁸¹ Making a determination whether the steps that have been taken to protect privacy are sufficient, once these clear categories of concern have been addressed, is left as a flexible, case-specific analysis by design.¹⁸² Because FERPA does not set forth a general privacy goal — one that can be demonstrably satisfied with certainty for any statistical computation — it is difficult to determine when a privacy-preserving technology provides protection that is sufficient to satisfy the regulatory requirements.

On first impression, it seems likely that the conceptual differences between FERPA’s privacy requirements and differential privacy are too great to be able to make an argument that differentially private tools can be used to satisfy FERPA. However, in the Sections that follow, such an argument is provided. Moreover, other regulations such as the HIPAA Privacy Rule share some of the characteristics of FERPA identified in this Section, suggesting that parts of the analysis presented in this Article could be applied to make similar arguments with respect to other information privacy laws.

D. Value in Bridging These Gaps

Legal and computer science concepts of privacy are evolving side by side, and it is becoming increasingly important to understand how they can work together to provide strong privacy protection in practice. The evolution of concepts and understandings of privacy in the field of computer science can benefit from an understanding of normative legal

179. *Id.*

180. Family Educational Rights and Privacy, 73 Fed. Reg. at 74,833.

181. *Id.* at 74,834.

182. *Id.* at 74,835 (“[W]e are unable to provide examples of rules and policies that necessarily meet the de-identification requirements The releasing party is responsible for conducting its own analysis and identifying the best methods to protect the confidentiality of information from education records it chooses to release.”).

and ethical privacy desiderata. Similarly, legal privacy scholarship can be informed and influenced by concepts originating in computer science thinking and by the understanding of what privacy desiderata can be technically defined and achieved. This Article argues, however, that in order to promote a mutual influence between the fields it is necessary to overcome the substantial gaps between the two approaches. Furthermore, in light of the vast bodies of theory and scholarship underlying the concepts, an approach that is rigorous from both a legal and a technical standpoint is required.

Bridging the gap between technical and regulatory approaches to privacy can help support efforts to bring formal privacy models such as differential privacy to practice. Uncertainty about compliance with strict regulatory requirements for privacy protection may act as barriers to adoption and use of emerging techniques for analyzing sensitive information in the real world. If data holders, owners, or custodians can be assured that the use of formal privacy models will satisfy their legal obligations, they will be more likely to begin using such tools to make new data sources available for research and commercial use.

At the same time, this interdisciplinary approach is also important for the future of robust privacy regulation. Because information privacy is in large part a highly technical issue, it will be critical for policymakers to understand the privacy technologies being developed and their guarantees. This understanding is needed in order to approve and guide the use of formal privacy models as a means of satisfying regulatory requirements. In addition, a clear understanding of the principles underlying formal approaches to privacy protection and the ways in which they differ from the concepts underlying existing regulatory definitions, and technical approaches outlined in current agency guidance, help illustrate the weaknesses in the regulatory framework and point to a new path forward. Taken together, these insights can be used to help pave the way for the adoption of robust privacy practices in the future.

IV. EXTRACTING A FORMAL PRIVACY DEFINITION FROM FERPA

The goal of this Section is to extract a mathematical model of FERPA's privacy requirements. Using such a model, it is possible to prove that privacy technologies that adhere to a model such as differential privacy meet the requirements of the law. An example of such a proof is presented later in this Article.

The approach used in this Article is inspired by the game-based privacy definitions found in the field of computer science. As discussed in Section II.B *supra*, a game-based approach defines privacy via a hypothetical game in which an adversary attempts to learn private information based on the output of a computation performed on private data.

If it can be shown that the adversary cannot win the game “too much,” the computation is considered to protect privacy.

To demonstrate that a given privacy technology meets the privacy requirements of FERPA, the first step is to define a game that faithfully encompasses the privacy desiderata that were envisioned by the Department of Education when drafting the regulations. Doing so requires carefully defining the capabilities of the adversary and the mechanics of the game in ways that capture the privacy threats that FERPA was designed to protect against. Because FERPA is written in the language of regulation and not as a formal mathematical definition, it is open to different, context-dependent interpretations by design. To deal with the inherent ambiguity in the language, the model presented in this Article conservatively accounts for a large class of such interpretations. As described in detail below, this requires designing games that give the adversary what might be considered to be an unrealistic advantage. However, if a system is proven to preserve privacy under extremely conservative assumptions, it follows that the system also preserves privacy in more realistic scenarios.

First, consider a simple, informal view of a privacy game based on FERPA’s requirements. Imagine a game in which a school classifies the student information it maintains into two distinct categories as defined by FERPA: (1) directory information, which can be made publicly available in accordance with FERPA, and therefore this modeling assumes it is available to the attacker, and (2) non-directory PII, which FERPA protects from disclosure, so the modeling does not assume it is available to the attacker. In the game, a statistical computation is performed by the game mechanics over this information and the result is shown to the adversary. The adversary wins the game if she successfully guesses a sensitive attribute of a student, i.e., links a named student to non-directory PII about that student. Figure 5 represents this game visually.¹⁸³

183. The description here omits important details that will be introduced in the following sections. Note that the mathematical formalization of a game abstracts the adversary as an arbitrary computation. This is a reasonable abstraction considering that to optimize its success in winning the game, an adversary needs to perform a sequence of inferences based on all the information it has gathered on students including, in particular, the adversary’s *a priori* knowledge, the directory information, and the computation result. Note, however, that by modeling the adversary as an arbitrary computation the model does not make any assumptions regarding this sequence of inferences, and in particular, the adversary does not have to adhere to any known attack strategy. This abstraction is necessary for making the mathematical arguments concrete and precise.

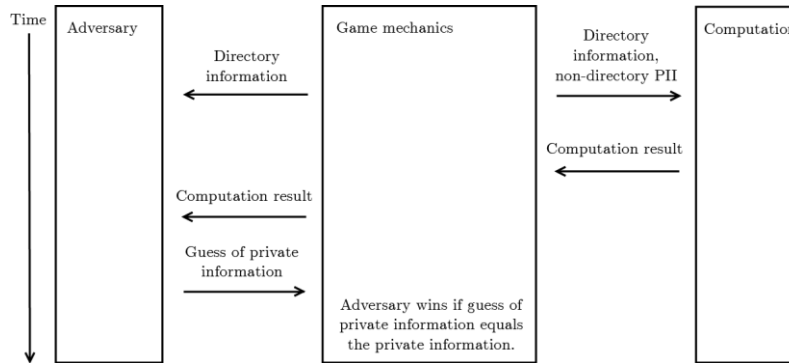


Figure 5: Simplified FERPA game.

To make this game framework more concrete, consider the following hypothetical scenario. The National Center for Education Statistics has decided to release publicly a dataset containing information from education records obtained from schools across the United States. To protect the students' privacy, statisticians have de-identified the dataset using a technique such as a k -anonymization algorithm.¹⁸⁴ After the de-identified dataset has been published online, a data broker attempts to glean non-directory PII about students from the k -anonymized data with the goal of selling the re-identified information to marketers. If the data broker successfully extracts a student's non-directory PII from the release, then the privacy of that student has been violated and the data broker's attack has been successful. In the language of a privacy game, the data broker is playing the role of the adversary and the k -anonymization algorithm is the computation that is intended to provide privacy protection. The data broker wins the game by performing a sequence of inferences, at the end of which the data broker successfully guesses a student's non-directory PII.

Note, however, that a privacy breach does not necessarily occur any time an adversary wins the privacy game. For instance, consider a game in which gender is treated as a private attribute of a student. If challenged to guess the gender of an arbitrary student, an adversary can be expected to win the game with a probability of 0.5. Moreover, if the student's name, for example, "Susan," is publicly available, the adversary could correctly guess the student's gender with probability close to 1. While this example might seem contrived, it is important to acknowledge that, even in richer domains, the adversary always has some likelihood of correctly guessing private information about a student and thus winning the privacy game even without seeing the results

184. For a reminder of the definition of k -anonymization, see *supra* note 35.

of computations performed on that student's private information. A privacy breach has not occurred in such scenarios because the adversary's successful guess was not informed by a leak of private information. Section IV.C.2 discusses how to account for the adversary's baseline probability of success, that is, the adversary's probability of success prior to seeing the results of the computation.

The subsections that follow present an approach to formalizing a game-based privacy definition for FERPA. Justifications for the assumptions made in developing this definition, and an outline of the game-based formalization itself, are also provided. The resulting model captures a broad scope of potential interpretations of the privacy risks that FERPA is intended to protect against. In this way, a computation that can be proven to satisfy the resulting model can be used to analyze private data or release statistics with a high degree of confidence that the computation adheres to the privacy requirements of FERPA.

A. A Conservative Approach to Modeling

Although statutory and regulatory definitions are generally more precise than language used in everyday interactions, they are nevertheless ambiguous. On the one hand, this ambiguity is an advantageous feature, since it builds flexibility into legal standards and leaves room for interpretation, value judgments, and adaptability to new scenarios as practices and social norms inevitably evolve over time. On the other hand, technological solutions that rely on formal mathematical models require the specification of exact, unambiguous definitions. This presents a challenge for the practical implementation of privacy technologies. This Section explores how a legal standard such as the privacy protection required by FERPA can be translated into precise mathematical concepts against which a technological tool can be evaluated.

Consider, for instance, the fact that FERPA explicitly classifies a student's Social Security number as non-directory PII and bars it from release.¹⁸⁵ It would clearly be unacceptable for a school to publicly release all nine digits of a student's Social Security number. It would also clearly be acceptable to release zero digits of the number, as such a release would not leak any information about the Social Security number. However, it is a fair question to ask whether it would be acceptable to release three, four, or five digits of a Social Security number. In other words, it is not necessarily clear at what point between disclosing zero and nine digits does acceptable data release become a prohibited data release.¹⁸⁶

185. Family Educational Rights and Privacy, 34 C.F.R. § 99.3 (2017).

186. Note that while this example may appear to be a purely hypothetical concern, regulators do in fact grapple with this type of problem. For instance, guidance from the Department of Health and Human Services on de-identifying data in accordance with the HIPAA Privacy

One way to approach this problem is to analyze the text of a statute or regulation and decide on a reasonable interpretation as applied to a given privacy technology. However, an interpretation that seems reasonable to one person might not seem reasonable to another. There is a concern that, if the model of FERPA's privacy requirements were based on a particular interpretation of the law, this interpretation could be disputed by other legal experts. In addition, future applications of the law may lead to new interpretations that challenge the interpretation adopted by the model.

To overcome these issues, the model proposed in this Article aims to err on the side of a conservative interpretation of the regulation's privacy requirements. That is, wherever there is a choice between different interpretations of FERPA's requirements, the more restrictive interpretation is selected. For instance, it is not clear from the regulatory requirements how skilled an adversary the educational institutions must withstand.¹⁸⁷ Therefore, the model assumes that the adversary is well-resourced and capable of carrying out a sophisticated privacy attack. This approach effectively strengthens the claim that a given technical approach to privacy protection satisfies a particular legal requirement of privacy protection. If a system can be proved to be secure in strongly antagonistic circumstances, including unrealistically antagonistic circumstances, it is certainly secure given assumptions that are more realistic.

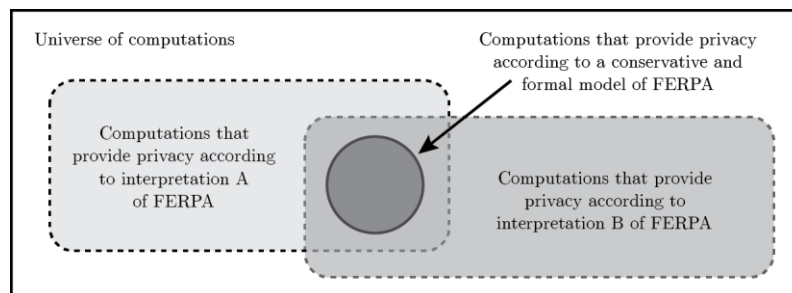


Figure 6: A conservative approach to modeling.

Rule's safe harbor standard for de-identification states that in general "parts or derivatives" of identifiers cannot be released. OFFICE FOR CIVIL RIGHTS, DEP'T OF HEALTH & HUMAN SERVS, *supra* note 4, at 25. However, the standard permits the first three digits of a ZIP code to be released without violating patients' privacy as long as the populations of all ZIP codes that begin with those three digits sum to over 20,000 individuals. *See id.* at 7–8.

187. Note, for example, that the definition of PII is not limited by the technological skill level of a potential adversary. *See* Family Educational Rights and Privacy, 73 Fed. Reg. at 74,831.

Figure 6 provides a visual representation of this desideratum for the model. The figure recognizes that there could be many possible interpretations of the privacy requirements of FERPA. Each possible interpretation can be understood as specifying a set of computations that satisfy the privacy requirements of the law, as well as a set of computations that it considers as failing to provide sufficient privacy protection. In addition, if such an interpretation is not specified with precise, mathematical language, then there are most likely also computations that fall into a gray area and are neither unambiguously endorsed nor rejected by that interpretation. A conservative approach to modeling the law attempts to identify those computations that fall unambiguously within the intersection of all (or, at least, a large class of) reasonable interpretations of the law. That is, if a conservative model considers a certain computation to provide sufficient privacy protection, then all these reasonable interpretations of the law would also consider it to provide sufficient privacy protection.

B. Modeling Components of a FERPA Privacy Game

The following sections describe this Article's model, which aims to capture the privacy desiderata of the Department of Education in promulgating the FERPA regulations. While conservative decisions were made at many points in the modeling process, there are places in the model where the decisions were not fully conservative and these are described below. In future work, the authors plan to extend the analysis to reflect on these decisions.

1. Modeling FERPA's Implicit Adversary

FERPA does not specify an explicit model of the adversary its protections are intended to withstand. For instance, neither the statutory nor regulatory text specifies the capabilities of a hypothetical attacker that PII must be protected from.¹⁸⁸ Despite the lack of an explicit description of the adversary envisioned, the Department of Education provided some details relevant to determining the types of attacks and attackers that were considered when drafting the regulations. In particular, FERPA's definition of PII describes what or whom the law is designed to protect against. This definition and how it has been interpreted in agency guidance provides useful details that can serve as a basis for modeling the implicit adversary the Department had in mind when formulating FERPA's requirements.

188. For instance, the Department of Education declined to interpret the reference to a "reasonable person in the school community" within FERPA's definition of PII as restricting the technological or skill level of a potential adversary. *Id.*

As discussed in detail above in Section III.B, FERPA prohibits the disclosure of non-directory PII from education records, except with the consent of the student or her parent or guardian, or in accordance with one of the limited exceptions set forth by FERPA.¹⁸⁹ One might naturally ask how the agency originally envisioned an improper disclosure. In investigating this question, of particular interest is the case in which a school or educational agency releases information pursuant to the provision of FERPA permitting the release of de-identified data.

To qualify under this exception, the released data must not contain non-directory PII. As discussed in Section III.B.3, FERPA defines PII to include direct and indirect identifiers, as well as other information that “would allow a reasonable person in the school community, who does not have personal knowledge of the relevant circumstances, to identify the student with reasonable certainty.”¹⁹⁰

By the inclusion of the quoted language, the agency emphasized concerns that a member of the school community might be able to learn non-directory PII about a student from a disclosure of de-identified data and established this as a category of privacy breach to protect against. As described above in Section III.B.1, this type of privacy breach is relevant to the use of formal privacy models. Therefore, the “reasonable person in the school community, who does not have personal knowledge of the relevant circumstances”¹⁹¹ is taken to be the implicit adversary embedded within FERPA’s requirements for the purposes of the model.

Next, consider who is a “reasonable person in the school community” and the knowledge such an individual is expected to have. The preamble to the 2008 final rule updating the FERPA regulations provides some limited guidance on these questions. As noted in Section III.B, a “reasonable person” is described in the preamble as a “hypothetical, rational, prudent, average individual.”¹⁹² In addition to enjoying the insider information about students that comes from being a member of the “school community,” this individual is “also presumed to have at least the knowledge of a reasonable person in the local community, the region or State, the United States, and the world in general.”¹⁹³ Moreover, the agency expressed an intent for this standard to

189. See 34 C.F.R. § 99.30.

190. 34 C.F.R. § 99.3. FERPA’s definition of PII also includes “[i]nformation requested by a person who the educational agency or institution reasonably believes knows the identity of the student to whom the education record relates.” *Id.* The privacy model as presented does not extensively address this part of the definition. However, the model guarantees that, no matter the adversary’s *a priori* knowledge about a student, the adversary’s ability to guess private information about that student does not improve much by seeing the result of a computation that meets the given definition of privacy.

191. *Id.*

192. Family Educational Rights and Privacy, 73 Fed. Reg. at 74,832.

193. *Id.*

provide “the maximum privacy protection for students.”¹⁹⁴ At the same time, since the adversary is assumed not to have “personal knowledge of the relevant circumstances,” the model should assume that the adversary also has some uncertainty about the student information. In other words, the regulations appear to recognize that it is not necessarily a privacy breach if some private information is identified in a data release. Rather, it amounts to a privacy breach only if the adversary had some uncertainty about that information before the data were made public.

2. Modeling the Adversary’s Knowledge

The regulatory language suggests that a “reasonable person in the school community” brings with her some knowledge about student information, and this knowledge should be modeled in the privacy game. How the model represents the adversary’s knowledge will affect the adversary’s probabilities of correctly guessing the private information of a student both with and without access to the result a computation performed on that private information.¹⁹⁵

A key choice made in the model is to represent the adversary’s knowledge about students as probability distributions over student attributes. Probability distributions provide a sufficiently rich language to model the type of knowledge that the regulations anticipate an adversary could have about students in protected education records. For instance, distributions can describe statistics that the adversary knows about the entire student body (e.g., demographic information), as well as beliefs that the adversary might hold about a particular student. For each student, the adversary is presumed to have some *a priori* beliefs about that student, represented as a probability distribution. That is, in the model, each student is associated with a probability distribution that represents the adversary’s beliefs about the non-directory PII of that student.¹⁹⁶ The following examples demonstrate the versatility of probability distributions for modeling the adversary’s knowledge.

First, probability distributions can be used to describe an adversary’s general knowledge about the demographics of a school district. Suppose there is a school district with two schools, Abington and

194. *Id.*

195. Note that, while having more knowledge makes it easier for the adversary to win the privacy game, having more knowledge also increases the baseline probability of success without access to the computation. While the former makes violating privacy easier, the latter makes violating privacy harder. It follows that an adversary having more knowledge does not necessarily make protecting privacy harder.

196. The model assumes that student attributes are drawn independently from the distributions; that is, the fact that one student has certain attributes does not affect the probability of another student having particular attributes. This is a limitation of the model, which will be discussed later in this Section and in Part V.

Bragdon. The percentage of low-income students at Abington is 95%, while only 40% of the students at Bragdon fall into this category. Furthermore, say that 35% of the students in Abington and 15% of the students in Bragdon scored below proficient on a statewide math exam. Without knowing anything else about her, if the adversary knows that Alice attends Abington, he might believe that there is a 95% chance that Alice is from a low-income family and a 35% chance that she scored below proficient on the exam. On the other hand, knowing that Grace attends Bragdon might lead the adversary to believe that there is a 40% chance that she is from a low-income family and a 15% chance that she scored below proficient on the exam. This general knowledge is modeled with two distributions: a student sampled randomly from the first distribution has a 95% chance of coming from a low-income family and a 35% chance of scoring below proficient on the assessment, and a student sampled randomly from the second distribution will fall into each of these categories with a likelihood of 40% and 15%, respectively.

Furthermore, distributions can be tailored to reflect more complex beliefs about trends across demographics. For instance, the two distributions described above could be reworked to reflect the adversary's belief that a student from a low-income family has a greater chance of scoring below proficient on the examination than a student who is not from a low-income family. Consider the example distributions given in Table 1. The table shows that, if one were to sample randomly from the distribution describing a student from Abington, the sampled student scored at least proficient on the exam and is from a low-income family with a probability of 0.601. The probability that the sampled student scored below proficient and is not from a low-income family is only 0.001. Taken together, the distributions in this table still reflect the facts that 95% of the students at Abington are low-income and 40% of the students at Bragdon are low-income and that 35% of students at Abington scored below proficient on the exam and 15% at Bragdon scored at this level. However, the distributions now also reflect the fact that, if a student from either school is from a low-income family, that student scored below proficient on the exam with a likelihood of about 35%. Similarly, if a student is not from a low-income family, the likelihood that the student scored below proficient is only about 2%.

Table 1: Distributions describing students at Abington and Bragdon.
Numbers are given as probabilities.

	Distribution describing student from Abington		Distribution describing student from Bragdon	
	Low-income	Not low-income	Low-income	Not low-income
Proficient or higher	60.1%	4.9%	26.0%	59.0%
Below proficient	34.9%	0.1%	14.0%	1.0%

Second, in addition to describing beliefs based on demographic information, distributions can also reflect an adversary's more specific beliefs about individuals. This can be used to model the case in which the adversary has a substantial amount of outside knowledge about a particular student. For instance, say an adversary is aware that Ashley's mother is a professor of mathematics. Because of her family environment, the adversary might believe that Ashley is more proficient at math than the average student and so has a greater chance of having passed the state math exam. The adversary's beliefs are modeled by associating with Ashley a distribution that describes a student that has a higher than average probability of having passed the math exam.

One important limitation of the current model is that the characteristics of a given student are assumed to be independent of the attributes of all other students. In other words, correlations between students are not modeled, even though one might see such correlations in the real world. For example, there may be correlations between siblings. If one sibling comes from a low-income family, then the other siblings should also have this characteristic. The best the current model can do is to give each sibling an equal probability of coming from a low-income family. There might also be correlations among larger groups of students. For example, it might be known that Alfred was the top achiever in his class on the state standardized math exam. Alfred's grade is not independent of the grades of the other students in his class. His grade is at least as high as the exam grades of the other students. While the current model does not account for correlations between students, Section IV.G discusses some ways to address this limitation in future work.

3. Modeling the Adversary's Capabilities and Incentives

As discussed in Section III.B.3, the preamble to the 2008 final rule interprets the regulatory language with respect to the capabilities of a

potential adversary. Specifically, the Department of Education explained that the reasonable person standard “was not intended to describe the technological or scientific skill level of a person who would be capable of re-identifying statistical information or redacted records.”¹⁹⁷ Accordingly, the model does not assume anything about the skill level of the adversary. The model also makes no assumptions about the analyses the adversary can perform nor about the computational resources available to her.¹⁹⁸ Furthermore, the model makes no assumptions about the motivation of the adversary. By not constraining the resources available to the adversary or the adversary’s motive, the model conservatively accounts for the types of adversaries contemplated by many reasonable interpretations of FERPA.

The model is also greatly expanded by treating the adversary as representative of a whole class of potential attackers and attacks, rather than being restricted to only a specific attacker or attack. In fact, the model accounts even for attacks that have yet to be conceived. This approach is consistent with the commitment to adhering to a conservative interpretation of FERPA, and it is also well supported by modern cryptographic theory, which emphasizes the design of cryptographic systems that are secure against attacks that were unknown at the time the cryptographic system was designed.¹⁹⁹ By not making assumptions about the capabilities of the adversary in the modeling — beyond adhering to a general computational model — the model is ensured to be robust not only to currently known privacy attacks, but also to attacks developed in the future, as is consistent with a modern cryptographic approach.²⁰⁰

This approach is also consistent with the preamble, which recognizes that the capabilities of attackers are increasing over time. For instance, the Department of Education recommends limiting the extent of

197. Family Educational Rights and Privacy, 73 Fed. Reg. at 74,831.

198. Note, however, that the adversary only tries to learn private information from analyzing a data release rather than by stealing a hard drive containing sensitive data or hacking into the system to inspect its memory or actively manipulate it. This modeling is consistent with the privacy threats addressed by the regulations in the context of releasing de-identified information.

199. Before modern cryptography’s emphasis on providing robust protection against known and unknown attacks, cryptographic techniques typically only addressed known privacy threats. The failure of these cryptographic methods to provide privacy against new attacks motivated the development of new cryptographic techniques, which in turn were broken by even more sophisticated attacks. To break this cycle, modern cryptographers strive to develop methods that are guaranteed to be invulnerable to any feasible attack under a general computational model.

200. In particular, the adversary in the model may possess greater computational power and perform more sophisticated computations than those performed by the system designed to protect privacy or those conceived by the system designer. This is justified, in part, considering that once a system is put to use and the results of its computations are made public, an attacker may have virtually unlimited time to exploit it, applying in the process technological advances not known at the time of the system’s design and deployment.

information publicly released as directory information, with the justification that “since the enactment of FERPA in 1974, the risk of re-identification from such information has grown as a result of new technologies and methods.”²⁰¹ Furthermore, in explaining why they left ambiguity in the definition of PII, the regulators emphasize that holders of education data must consider potential adversaries who may attempt to learn private student information through many possible attack vectors.²⁰² The regulators contrast this approach with the HIPAA Privacy Rule, which effectively assumes that an adversary will only attempt re-identification attacks and base these attacks on a small set of a person’s attributes.²⁰³ Additionally, while traditional re-identification attacks against microdata are emphasized, FERPA does not contemplate only this type of attack, and implementation guidance also addresses exploits aimed at learning private student information from aggregate data.²⁰⁴ Because the Department of Education has not explained the full scope of the types of attacks covered by FERPA, a conservative approach ensures the model accounts for *any* attack that can be perpetuated by an adversary.

In contrast with the approach taken in this Article, some frameworks for measuring the privacy risk associated with a data release explicitly require making assumptions about the motives and capabilities of a potential adversary. For instance, some experts suggest that estimates of the motives and capabilities of a potential adversary should be used as an input when calculating the re-identification risk of a data release.²⁰⁵ If it is believed that any potential adversary will have limited resources to apply towards a privacy attack or little incentive to attempt an attack, the risk of re-identification is assumed to be small, so the data can presumably be safely released with fewer protections, or with administrative controls in place.²⁰⁶ Such an approach has clear benefits to utility as data can be released with little or no modifications, and it may be justifiable in specific settings, such as where the payoff of a successful attack on privacy can be demonstrated to be significantly lower than the cost of the attack. However, the reliance on assumptions regarding the adversary may result in future susceptibility to attacks as incentives and analytic capabilities change over time. The model in this Article takes the more conservative stance that the adversary is fully capable of any privacy attack and fully incentivized to attack. This is necessary

201. Family Educational Rights and Privacy, 73 Fed. Reg. at 74,834.

202. *See id.* at 74,833 (“It is not possible, however, to list all the possible indirect identifiers and ways in which information might indirectly identify a student.”).

203. *See id.*

204. *See, e.g.*, NAT’L CTR. FOR EDUC. STATISTICS, U.S. DEP’T OF EDUC., *supra* note 131, at 5.

205. *See* KHALED EL EMAM & LUK ARBUCKLE, ANONYMIZING HEALTH DATA 23 (Andy Oram & Allyson MacDonald eds., 2013).

206. *See id.*

in order to capture a broad class of possible interpretations of FERPA, in light of the Department of Education's explanation that FERPA's privacy requirements are intended to provide "the maximum privacy protection for students,"²⁰⁷ and its recognition that "the risk of re-identification may be greater for student data than other information because of the regular publication of student directories, commercial databases, and de-identified but detailed educational reports by States and researchers that can be manipulated with increasing ease by computer technology."²⁰⁸

4. Modeling Student Information

Like FERPA, the model distinguishes between directory information and non-directory PII. Because directory information can be disclosed in accordance with FERPA, the model assumes that the adversary has access to it.²⁰⁹ Non-directory PII is private information generally not available to the adversary, although the adversary might hold some *a priori* beliefs about it. As explained in Section IV.B.2, these beliefs are modeled via probability distributions over student attributes. Therefore, in the model each student is associated with a record consisting of two pieces of information: a set of concrete values that constitutes the student's directory information, and a probability distribution describing the adversary's beliefs about that student's private attributes. Section IV.C.1 discusses the modeling of ambiguity concerning which student attributes are considered directory information and which attributes are considered non-directory PII.

5. Modeling a Successful Attack

FERPA prohibits the non-consensual disclosure of non-directory PII. According to the regulations, disclosure "means to permit access to or the release, transfer, or other communication of personally identifiable information contained in education records by any means, including oral, written, or electronic means, to any party except the party identified as the party that provided or created the record."²¹⁰ Therefore, outside of a disclosure exception, a data release by an educational agency or institution should not communicate any non-directory PII to

207. Family Educational Rights and Privacy, 73 Fed. Reg. at 74,832.

208. *Id.* at 74,834.

209. Note that some schools make their directory information available only to certain members of the school community, such as teachers and parents, not the general public. However, consistent with a conservative approach, the model assumes it is available to the adversary.

210. Family Educational Rights and Privacy, 34 C.F.R. § 99.3 (2017).

a recipient or user of the data it discloses. In the model, this is formalized by saying that the adversary wins the privacy game if she correctly guesses non-directory PII of a particular student after seeing the output of a computation performed on the private data.

More precisely, the adversary wins the game if she successfully guesses a function of the private information of a student. In other words, the model recognizes that often learning (or inferring) something about private information could be a breach of privacy, even if the private information itself is not learned directly or completely. For instance, consider a scenario in which the private information of a particular student is his test score, which happens to be a 54 (graded on a scale from 0 to 100). If the adversary learns from a data release that the student failed the test, the adversary is considered to have won the privacy game, even if the adversary has not learned the student's exact numerical score. This is a more conservative approach than only considering the adversary to have won the game if the adversary guesses the exact private information.

The FERPA regulations require a data release to preserve uncertainty about private student information. In addition, many states prohibit reporting that 100% of students achieved certain performance levels. This recommendation is presumably based on the assumption that there is likely a reasonable person in the school community who does not know the academic performance of every student in the school. Reporting that 100% of students achieved a certain performance level removes any uncertainty from that individual's mind about each student's individual performance; therefore, non-directory PII has been communicated.

For a more nuanced example, consider the release of a table that contains information about student test scores, such as Table 2. The table distinguishes between native English speakers and English language learners and gives the number of students from each category who received scores at various levels, including below basic, basic, proficient, and advanced. Suppose the table reports that one English language learner achieved a proficient score on the exam and the other nine English language learners scored at either a below basic or basic level. This would constitute a disclosure of non-directory PII because the sole English language learner who achieved a proficient score, Monifa, could learn from the publication of this table that her peer English language learners, such as her friend Farid, scored at one of the lower levels. Prior to the release of the table, Monifa presumably did not know the scores of her classmates. Thus, if the "relevant circumstances" are taken to be the performance of her peers, she did not have "personal knowledge of the relevant circumstances."²¹¹ However, after

211. *Id.*

the release, she can infer with certainty that each of the other language learners, like Farid, performed below a proficient level on the test.²¹²

Table 2: A release of non-directory PII prohibited by FERPA.

	Test score levels			
	Below basic	Basic	Proficient	Advanced
Native English speaker	4	6	11	5
English language learner	5	4	1	0

Guidance on de-identifying data protected by FERPA also suggests that very low or high percentages should be masked when reported.²¹³ For example, consider a statistic reporting what percentage of a class of forty-one students scored at the proficient level on an exam, where only one student, Siobhan, scored at the proficient level. Guidance from the National Center for Education Statistics recommends reporting that less than 5% scored at a proficient level, instead of the true percentage, which is around 2.5%.²¹⁴ If the true percentage were reported, Siobhan would learn that all her peers failed to achieve this level, since she knows that she accounts for the entire 2.5% of proficient students. However, if the “masked” percentage is reported, Siobhan would no longer be able to reach this conclusion with certainty, since another student could have scored at her level (i.e., 5% of forty-one is slightly more than two students out of forty-one.). As there is only a 2.5% probability that a given student besides Siobhan scored at a proficient level, Siobhan might be able to form a strong belief about the performance of each of her peers; however, she has not learned with certainty the performance level of any particular classmate based on this publication. Consequently, the data release is considered to preserve uncertainty about private student information.²¹⁵

To capture this fully, the model needs to preserve uncertainty about non-directory PII in a data release. The model does so by requiring a stronger property. It says that a system preserves privacy only if gaining

212. This example is adapted from NAT’L CTR. FOR EDUC. STATISTICS, U.S. DEP’T OF EDUC., *supra* note 131, at 5.

213. *See id.* at 14.

214. *See id.* at 23.

215. *See id.*

access to the system does not change by very much an adversary's ability to guess correctly the private information of any student. In the model, an adversary holds some *a priori* beliefs about the non-directory PII of each student. Based on these beliefs alone, the adversary can make a guess about the private information of a student. Thus, a system preserves privacy if the adversary's chances of guessing the private information correctly after seeing a data release are about the same as his chances of guessing correctly based only on his *a priori* beliefs. Specifically, the model allows the adversary's probability of success to change by a multiplicative factor that is close to one.²¹⁶ For instance, it might require that the adversary's probability of success only change by a multiplicative factor of at most 1.1. In this case, if some adversary is able to correctly guess that a particular student failed an exam with a probability of 0.6 based on his *a priori* beliefs about that student, he would be able to successfully guess that the same student failed the exam with a probability of no more than $0.6 \times 1.1 = 0.66$ after seeing the output of a privacy-preserving computation. In other words, his chance of successfully guessing the performance of that student has changed, but the change is bounded.

This approach also preserves uncertainty when the initial uncertainty is small. If the adversary can initially guess that a student failed with a probability of success of 0.99, the model guarantees that the output of a privacy-preserving computation would not increase his chances of guessing correctly beyond a probability of success of 0.991. Intuitively, this is because the probability of success of the "complementary adversary" — who is interested in the probability that the student did not fail the exam — cannot change significantly.²¹⁷

C. Towards a FERPA Privacy Game

Section IV.D below describes a privacy game and a corresponding privacy definition that fits the model extracted based on FERPA's requirements for protecting privacy in releases of education records. There are two topics to cover before introducing the game. First, Section IV.B.4 described how student information is classified as directory

216. More formally, this multiplicative factor is captured by the parameter ϵ , which is typically a small constant. The adversary's belief may change by a factor of $e^{\pm\epsilon}$. For a small epsilon, $e^\epsilon \approx 1 + \epsilon$ and $e^{-\epsilon} \approx 1 - \epsilon$. The epsilon parameter can be tuned to provide different levels of privacy protection.

217. Given the same information, the "complementary adversary" makes the opposite guess of the original adversary. The "complementary adversary" has an initial chance of $1 - 0.99 = 0.01$ of successfully guessing that the student did not fail the exam. This probability of success can only change by a multiplicative factor between 0.9 and 1.1, which means that after observing the computation output, the "complementary adversary" will have a probability of success of at least $0.01 \times 0.9 = 0.009$. Returning to the original adversary, her probability of success is therefore at most $1 - 0.009 = 0.991$.

information or non-directory PII but left some ambiguity as to what comprises student information. Formalizing the requirements in a game would require addressing this ambiguity, and this is discussed in Section IV.C.1. Second, Section IV.B.5 modeled an adversary winning the game as correctly linking a student's identity with a function of that student's sensitive information. Note, however, that there is always some non-zero probability that an adversary can win the game, even without receiving the output of the computation. Section IV.C.2 defines this success probability as a baseline, and a successful attack is one where an adversary wins the FERPA privacy game with probability that is significantly higher than the baseline.

1. Accounting for Ambiguity in Student Information

Before the privacy game can be run, the model must define what constitutes directory information and non-directory PII. In the game, the adversary will have direct access to the directory information and will also have some indirect knowledge of the private information. The computation will use both types of information to produce some output.

There is a degree of ambiguity in the regulatory definitions of directory information and non-directory PII. The regulations provide a non-exhaustive list of examples of directory information,²¹⁸ but each educational agency or institution is granted discretion in determining what to designate and release as directory information. Seeking generality, the model does not make any assumptions about the content of directory information. Instead, it allows the adversary to decide what information is published as directory information. While this may seem like an unintuitive and unrealistic modeling choice, as it is unlikely that any adversary would have this ability in the real world, this modeling decision helps establish a privacy requirement that will be robust both to a wide range of potential interpretations of FERPA and to different choices made within specific institutional settings in accordance with FERPA. This modeling also exhibits conceptual clearness that is instrumental to understanding intricate concepts like privacy. To recap this modeling decision, it is effectively allowing the adversary to choose directory information that is the worst-case scenario for the privacy protection of the system. If the system is proven to provide privacy protection even in this worst-case scenario, then one can be confident that it preserves privacy no matter what the directory information actually is.²¹⁹

218. See Family Educational Rights and Privacy, 34 C.F.R. § 99.3 (2017).

219. By allowing the adversary to choose the directory information, the model enables the adversary to win the game more easily, while at the same time raising the bar for what qualifies as a privacy breach. See *supra* note 195.

Similarly, there is uncertainty regarding exactly what information constitutes non-directory PII. From the definition of PII, it is clear that the system must protect information that “is linked or linkable to a specific student.”²²⁰ However, the model should not make any assumptions about the capabilities of the adversary or the methods he might use to identify a student in released records that have been stripped of non-directory PII. Indeed, the guidance on interpreting FERPA instructs against making any assumptions of this nature.²²¹ It is impossible to say what information falls into these categories. Furthermore, the model should not make any assumptions about the auxiliary knowledge that the adversary may have about students. Accordingly, as with the directory information, the model allows the adversary to have a say about the non-directory PII of the student or students whom he is targeting. More precisely, for each student in the dataset, it allows the adversary to choose the distribution over student attributes that models that student’s non-directory PII. For example, the adversary could choose the student information presented in Table 3. The directory information consists of the name, age, and ZIP code of three students. The adversary chooses concrete values for these attributes. The private attributes are whether a student has a learning disability and whether a student has been suspended in the last year. In the model, the adversary assigns a probability distribution that describes the likelihood of each combination of these attributes being true.

Table 3: Example student information chosen by the adversary.

Directory information			Probability distribution over private information			
Name	Age	ZIP code	Disability & suspended	Disability & not suspended	No disability & suspended	No disability & not suspended
Robi McCabe	18	00034	0.3	0.3	0.2	0.2
Launo Cooney	17	00035	0.2	0.2	0.3	0.3
Samuel Strudwick	17	00034	0.1	0.2	0.3	0.4

220. 34 C.F.R. § 99.3.

221. See Family Educational Rights and Privacy, 73 Fed. Reg. 74,806, 74,831 (Dec. 9, 2008).

Table 3 explicitly breaks down the distribution for each student. Consider the third row in the table. As directory information, the adversary chooses that this student's name is Samuel Strudwick, that he is seventeen years-old, and that he lives in ZIP code 00034. The adversary assigns to Samuel a distribution describing how likely he is to have a disability and how likely he is to have been suspended. This distribution gives an exact probability that each combination of the two binary attributes is true. For instance, according to this distribution chosen by the adversary, there is a 10% likelihood that Samuel has a learning disability and has been suspended in the last year, and a 20% chance of Samuel having a learning disability but not having been suspended. This distribution reflects the adversary's *a priori* beliefs about Samuel's private information, which will be used to set the baseline for the adversary's success in Section IV.C.2. As will be discussed in Section IV.D, during the game, Samuel's actual private information will be sampled from this distribution.

2. Accounting for the Adversary's Baseline Success

To understand how the model accounts for the adversary's baseline success, recall the cryptographic privacy game from Section II.B, in which one of two plaintext messages is encrypted. The adversary is given the resulting ciphertext and must then identify which of the two original plaintext messages is behind the ciphertext. Since each message was encrypted with a probability of 0.5, the adversary can win the game with a probability of 0.5 without even examining the ciphertext, e.g., by guessing randomly between the two messages. This is the adversary's baseline for success. Even if a "perfect" cryptographic computation were used to encrypt the message, the adversary could still be expected to win the game 50% of the time.

Similarly, it is unreasonable to expect that the adversary will never win the FERPA privacy game proposed in this Article, as there is always the possibility that the adversary guesses correctly by chance. For example, consider that Siobhan attends a school where 50% of the students scored below proficient on the state reading exam. Without any other knowledge, the adversary might guess that Siobhan scored below proficient on the test and have a 50% chance of being correct, provided that each student scored below proficient with an equal probability.

A system can still be considered to preserve privacy even if the adversary wins the game with some likelihood (e.g., due purely to chance). To see why this is the case, consider a computation C that, on every input, outputs the result 0. Because the outcome of C does not depend on the input, it provides perfect privacy protection. Suppose that the adversary knows that 80% of the students in a school have a learning disability. If C is performed on the private student data, the

adversary receives only the output 0, which provides her with no useful information for learning private information about the students, i.e., whether a particular student has a learning disability. However, if the adversary guesses that a given student has a learning disability, she will win the game with a probability of 0.8. The fact that the adversary wins with a high probability should not be considered to be in contradiction to C providing perfect privacy protection, as the adversary could win the game with the same probability without access to the outcome of C .

To account for the adversary's baseline chance of success, a computation is considered to preserve privacy if the probability of the adversary winning the game when she has access to the system is not much greater than the probability of her successfully guessing sensitive information without having access to the system. This approach closely mirrors the approach taken in Section III.A. By this standard, the 0-outputting computation C from the previous example is considered to provide perfect privacy protection, since the adversary's probability of winning the game has not improved at all by having access to the output of this computation. Since there is no informational relationship between the input dataset and C 's output, C necessarily cannot leak information about the input, and so seeing C 's output cannot possibly enable the adversary to win the game with a higher probability than before seeing the output. Of course, C provides no utility; it is a completely useless computation. Useful computations will involve some meaningful relationship between the input dataset and the computation output, and so one cannot expect them to provide the same level of privacy protection as C (i.e., perfect privacy). Nonetheless, if the adversary's probability of winning the game after seeing the output of one of these computations does not improve very much relative to her probability of success before seeing the output, then intuitively the computation must not have leaked very much private student information, and so that computation is considered to preserve privacy.

D. The Game and Definition

This Section describes a game representing a scenario in which the adversary is committed to learning non-directory PII about a particular student. For instance, consider the case of a local journalist trying to learn the private information of a school's star basketball player from a data release about school disciplinary actions. The journalist only cares about learning the information of this one student.²²²

222. This scenario is related to subsection (g) of FERPA's definition of PII, which includes "[i]nformation requested by a person who the educational agency or institution reasonably believes knows the identity of the student to whom the education record relates." 34 C.F.R. § 99.3. In both cases, the contemplated adversary is trying to learn information about

In this model of a targeted attack, shown in Figure 7, the adversary commits to attacking a specific student *before* seeing the output of a computation performed on private student data, and then attempts to guess private information about only that student. This scenario is more restricted than other possible scenarios, such as an untargeted attack, in which the adversary chooses which student to attack *after* seeing the output of the computation. This discussion focuses on the targeted scenario because it is conceptually easier to understand (especially with respect to the proof of differential privacy's sufficiency) and still captures a large number of reasonable interpretations of FERPA's privacy requirements. Appendix II sketches models of more general scenarios like the untargeted attack scenario; differential privacy can also be shown to provide privacy in these cases. While the targeted attack scenario is itself not the most general scenario, the model for it presented here is conservative. Taking for granted that the adversary commits to a victim student before seeing the computation output, the model otherwise errs on the side of giving the adversary more power.

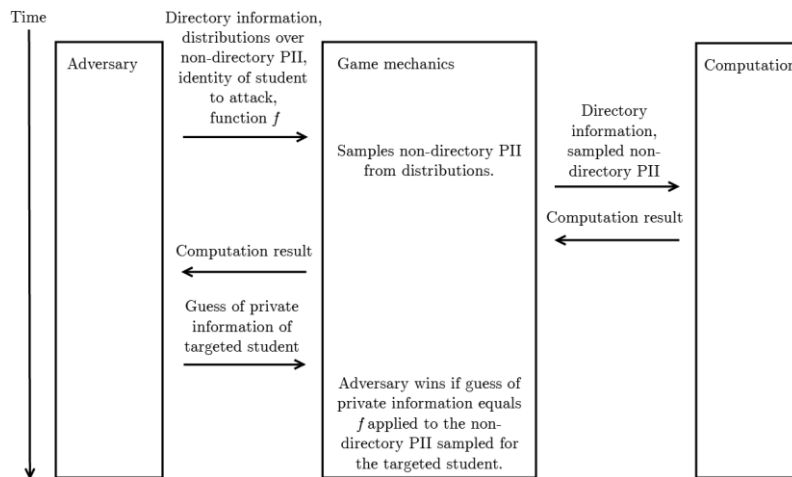


Figure 7: The real-world scenario.

1. Mechanics

In line with a conservative approach to modeling, the adversary is allowed to choose a directory of public student information. The adversary assigns to each student a probability distribution that describes the

one particular student. However, in the model the adversary does not request education records. Instead, the adversary observes the result of a statistical computation performed on education records, e.g., attempts to learn information about the star basketball player from a de-identified data release. *See supra* note 190.

adversary's *a priori* beliefs about the non-directory PII for that student and also chooses a function f whose domain is private student information. Intuitively, f represents the aspect of private student information that the adversary is interested in learning. For example, consider a scenario in which the private information is a student's test score. The adversary might want to learn whether a student, Bill, passed or failed an exam but might not care about the exact score that Bill earned. If this is the case, the adversary can choose a function that takes as input Bill's numerical score and outputs "passed" if Bill earned a passing score and which outputs "failed" otherwise.

The adversary provides the directory, the probability distributions over student information, and the function f to the game mechanics. Additionally, the adversary chooses a particular student to attack (i.e., the adversary commits to trying to learn the non-directory PII of that student, and that student alone), and provides the identity of the chosen student to the game mechanics. The game mechanics instantiate a database of non-directory student information by making a random draw from each of the probability distributions chosen by the adversary. This database, along with the directory information, is given to some computation C , resulting in an output statistic. Note that C is not supplied with the identity of the student that the adversary is attacking or the function that the adversary has chosen.

The game mechanics pass the result of C to the adversary. Based on this result, the adversary reports a guess about some aspect of the student's non-directory PII. The game mechanics declare that the adversary has won if the guess matches the result of applying the function f to the distinguished student's non-directory PII. Otherwise, the game mechanics declare that the adversary has lost. To continue the example from above, the adversary will guess either "passed" or "failed," and will win if his guess matches the result of applying the function f to Bill's test score.²²³

2. Privacy Definition

Informally, this Article considers a computation to preserve privacy in a targeted scenario if any adversary's chance of successfully winning the game is about the same as that of an adversary who does not have access to the output of the computation. This notion is formalized by comparing the game described above, called the *real-world scenario*, to another game, called the *ideal-world scenario*, as depicted in Figure 8. In an ideal-world scenario, the adversary chooses the same

223. The game mechanics' declaration of whether the adversary won or lost is only for purposes of the privacy definition (as a proxy for measuring the adversary's success rate) and does not correspond to an actual "real world" declaration. In fact, such a declaration would itself constitute a privacy breach (e.g., if the adversary has multiple chances to guess).

distributions over student information, directory information, student to target, and function f over student information that the real-world adversary chose. In both scenarios, the game mechanics instantiate a database of student information by sampling randomly from the probability distributions chosen by the adversary.

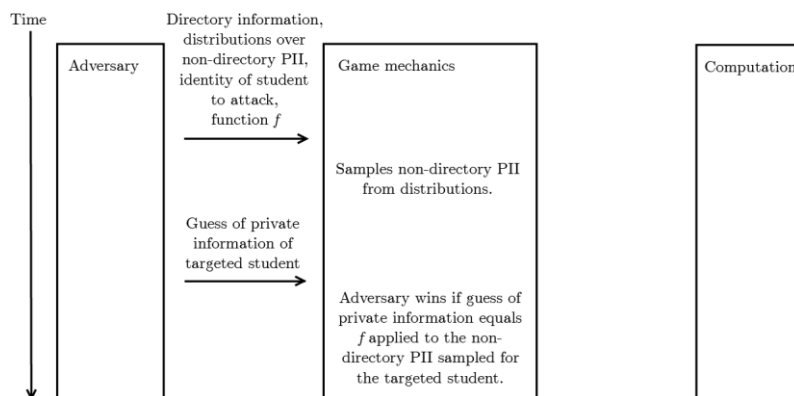


Figure 8: Ideal-world scenario. Note that the computation is left only for visual comparison with Figure 7.

However, unlike in the real world, in the ideal world no computation is performed on this database, and no information flows from the game mechanics to the adversary. Hence, the adversary clearly does not receive any private information about the targeted student during the course of the game and can only guess based on his *a priori* beliefs about the student. Thus, as the name implies, the ideal-world scenario provides perfect privacy protection to the targeted student.²²⁴

It is important to note that, technically speaking, the ideal-world adversary and the real-world adversary represent distinct classes of adversaries. A real-world adversary makes a guess about the private student information based on seeing the computation result, whereas an ideal-world adversary makes a guess without receiving any information. A computation preserves privacy if for every real-world adversary playing against that computation, there is an ideal-world adversary with the same *a priori* knowledge who wins against that computation with nearly the same chance of success.²²⁵ Note that the only difference

224. Note that the adversary can still win the game although no private information is leaked. See discussion *supra* Section IV.C.2.

225. Note that, in both real- and ideal-world scenarios, the “actual” student attributes are drawn from distributions supplied by the adversary. This means that, in both scenarios, the adversary necessarily has correct prior beliefs about the private information of the students. To see why this is a reasonable design choice, note that it would be too restrictive to consider an adversary with wrong beliefs, as that would imply that the computation output would not

between the information available to these adversaries is that the real-world adversary gets to see the computation output and the ideal-world adversary does not. As argued above, the ideal-world adversary learns absolutely nothing specific to the targeted student during the course of the game. If he is nearly as successful at guessing the private information of the student as the real-world adversary (who does see the computation result) then it must be that the computation result does not reveal significant information that is specific to the student to the real-world adversary. Thus, if for every real-world adversary there exists an ideal-world adversary with nearly the same level of success, one can be confident that the computation does not leak significant private information and hence preserves privacy in the targeted setting.

3. Privacy Loss Parameter

The definition above states that every real-world adversary playing against a certain computation should win with *nearly* the same level of success as the corresponding ideal-world adversary playing against that same computation. The so-called privacy loss parameter captures exactly how much better the real-world adversary may perform than the ideal-world adversary if the computation is to be considered privacy-preserving. This parameter is a small multiplicative factor denoted ε . If the ideal-world adversary has success with probability p , the real-world adversary is only allowed to have success up to probability $p \times (1 + \varepsilon)$. For instance, consider ε set to 0.01. If an ideal-world adversary's probability of success is 0.5, then the real-world adversary's probability of success is only allowed to be $0.5 \times (1 + 0.01) = 0.505$. Thus, while the ideal-world adversary has exactly a 50% chance of winning the game in this example, the real-world adversary can be expected to win the game slightly more than 50% of the time.²²⁶

allow the real-world scenario adversary to learn that her prior beliefs were wrong. As a concrete example, consider an adversary whose (incorrect) prior belief is that each of the students taking the exam had a 99% chance of failing, while in reality each had a 99% chance of passing the exam, and in fact 97% of students passed. In addition, say the school wants to release the passing rate via some private computation. If a computation is not considered private when it improves this clueless adversary's chance of guessing whether or not a student passed the exam, then the school will not be able to release any useful statistic.

226. It is important that this parameter is a multiplicative factor, and not additive. In the above example, an additive parameter allows the adversary to win with a probability of $0.5 + \varepsilon = 0.51$. While this might seem acceptable, using an additive factor is problematic when the ideal-world adversary's chance of winning is very small. For instance, say that the ideal-world adversary is guessing about some property that only one out of ten thousand students has. The ideal-world adversary has a chance of success of $p = 0.0001$. If ε were an additive parameter, then the real-world adversary is allowed to win with probability $p + \varepsilon = 0.0101$. This would mean the real-world adversary is over one hundred times as successful as the ideal-world adversary: the ideal-world adversary only has a one in ten thousand chance of winning the game, but the real-world adversary wins the game roughly one in one hundred times. This is highly unlikely to be an acceptable level of privacy leakage. If, on the other hand, ε is a

The fact that this definition is parameterized should be considered a strength because it allows the privacy guarantees of the definition to be tuned to an acceptable level. Such parameters are typically set after a societal-technical process during which the mathematical-theoretical understandings are matched with practical real-life experience, and then reviewed and adjusted periodically. Examples include the 5% confidence often used as an acceptable level of uncertainty in statistics. Where included in a parameterized privacy definition, these parameters enable the definition to be tuned to meet ethical or legal standards of privacy protection.²²⁷

E. Applying the Privacy Definition

This Section presents a few examples of the privacy definition in action. Keep in mind that if a computation meets the privacy definition, then a real-world adversary is only marginally more successful at guessing private student information than its corresponding ideal-world adversary — who does not have access to the computation output at all. Relatedly, the computation output will not enable any adversary to guess private student information with certainty, assuming, of course, that the adversary could not have guessed the information with certainty before seeing the computation result.

These strong guarantees automatically rule out a wide collection of attacks. To see how, consider the example of a linkage attack, in which an adversary breaches privacy by combining the computation output with an external database. Since the ideal-world adversary does not see the computation output, it cannot perform any sort of linkage attack. Because the real-world adversary can only be slightly more successful than the ideal-world adversary, whatever linkage attack the real-world adversary might attempt to perform when it sees the computation output cannot help the adversary guess private student information very much.

The use of computations that meet the definition of privacy from Section IV.D.2 would also have prevented the privacy breaches demonstrated in examples in previous sections. In Section IV.B.5, a student named Monifa learned private information about her peers from the release of a table showing the number of students who scored at the below

multiplicative factor, any computation meeting the given privacy definition guarantees that the real-world adversary would not win the game with a probability of more than $p \times (1 + \epsilon) = 0.000101$, which is only slightly more than the ideal-world adversary's chance of success.

227. In future work, the authors plan to explore paradigms based on legal standards and best practices to identify a range of values for an initial setting of the privacy loss parameter that would be considered reasonable in terms of satisfying the requirements of FERPA and other regulations. See discussion *infra* Section IV.G.

basic, basic, proficient, and advanced levels on a standardized test, separated by whether the students were native English speakers or English language learners. Since Table 2 shows that only one English language learner performed at a proficient level, and Monifa knows that she achieved this score, she can infer that all her peer English language learners scored below proficient. Since she presumably did not know this before the release of the table, the release has violated the privacy of the other students.

Envision an adversary who has similar knowledge to Monifa. This adversary has no uncertainty about the private attributes of Monifa but has some uncertainty about the private attributes of the other students. Further, suppose this adversary accurately knows the *a priori* probability that each of Monifa's peer English language learners failed the exam. A computation that meets the given definition of privacy comes with the guarantee that this adversary cannot use the computation output to guess with certainty the private information of any of the students about whom he or she had initial uncertainty. For instance, the adversary would be able to guess the private information of Monifa with certainty, which makes sense, since the adversary is designed to model her. However, the adversary would not be able to guess whether Monifa's friend Farid passed the exam with much more success than the adversary's *a priori* beliefs would allow.

What this means is that even if Monifa had accurate *a priori* beliefs about the private information of her peers, if the table were released via a computation that meets the given definition of privacy, she could not learn anything from the table that would enable her to guess the academic performance of one of her peers with certainty, unless she were able to do so before the release of the table.²²⁸

F. Proving that Differential Privacy Satisfies the Requirements of FERPA

In the modeling process detailed above, language from the regulations and implementation guidance were used to construct a formal privacy game that captures a very conservative interpretation of the privacy desiderata of FERPA (i.e., where there are ambiguities in the regulatory language, the model errs towards stronger privacy requirements). The model focuses on a targeted attack scenario, in which a

228. Note that if Monifa had very inaccurate prior beliefs about the likelihood of her peers passing the exam, the computation output might radically change her probability of successfully guessing whether or not a peer passed the exam. The model allows this for the reasons described in Section IV.B.5, *supra*. Regardless, Monifa would never be able to guess the performance of her peers with certainty, no matter her initial beliefs (unless she were certain before seeing the computation output).

hypothetical adversary seeks to recover non-directory PII about a particular student. During the game, the adversary selects the directory information and distributions describing private attributes of the students that will be available as resources to be leveraged in the attack. The adversary wins the game if, after receiving the result of a computation performed on the student data, she is able to successfully guess a function of the non-directory PII sampled for the targeted student. This game provides a precise, mathematical definition for determining whether a statistical computation meets a conservative model of the privacy requirements of FERPA.

Every differentially private computation provably meets this definition and, since the requirements of this definition are likely stricter than those of FERPA, thus satisfies the privacy requirements of FERPA.

This Section contains a brief overview of the mathematical proof.²²⁹ Consider a computation that is being evaluated for its adherence to the definition of privacy extracted from FERPA. To show that it does, this Section gives a proof using the real-world and ideal-world games presented in Section IV.D above, which were based on conservative interpretations of FERPA's privacy requirements. Recall that in the ideal-world game, no computation is performed on the data and no information flows from the game mechanics to the adversary, while, in the real-world game, the computation is performed and the output is shared with the adversary. The goal of the proof is to demonstrate that, for every real-world adversary who plays the game against the computation, there exists an ideal-world adversary with the same *a priori* knowledge who is nearly as successful at winning the game. If this is true, one concludes that the real-world adversary also cannot win the game significantly better than in the case where it does not get to see the computation result.

The proof follows a paradigm known as a *hybrid argument*,²³⁰ which involves a progression of hybrid games.²³¹ The proof's progression begins with the ideal-world game and ends with the real-world game, with intermediate "hybrid" games that share features with both the ideal- and real-world games. The argument demonstrates that there is little to no difference in privacy loss between every consecutive step in the progression, and hence little privacy loss between the ideal- and real-world games.

This proof demonstrates that using differentially private computations to analyze educational records is consistent with a very conservative interpretation of FERPA. Therefore, organizations that wish to

229. For a technical sketch of the proof, see *infra* Appendix I.

230. This type of argument originated in Shafi Goldwasser & Silvio Micali, *Probabilistic Encryption*, 28 J. COMPUTER & SYS. SCI. 270 (1984).

231. For a pictorial description of a hybrid argument, see *infra* Figure 9.

perform analyses on non-directory PII using these tools, or to allow others to do so using their data, can do so with high confidence that the risk of thereby being found in violation of FERPA is very low.

G. Extending the Model

The model discussed in this Article accounts only for adversaries who target a particular student in the data and see only the result of a single computation performed on the student data. This Section introduces two possible extensions to this model, including extensions for untargeted attack scenarios and scenarios involving multiple data releases. A more detailed, technical discussion of these extensions is provided in Appendix II.

In an untargeted attack scenario, the adversary does not commit to attacking a particular student, but rather decides which student to attack after seeing the computation result. This attack scenario captures an adversary such as a data broker, who wants to extract from a data release information about *any* student in the dataset but does not initially have any particular student in mind. The multiple-release scenario reflects the fact that it is common for multiple statistics to be published about the same group of students, and there is the possibility that the releases in aggregate fail to preserve the students' privacy, even if each release in itself is privacy-preserving. Clearly, these extensions of the model reflect more general privacy protection than the targeted, single-release scenario focused on in this Article. In other words, every computation that preserves privacy in the untargeted scenario or the multiple-release scenario also preserves privacy in the targeted, single-release scenario, but the converse is not true.

This Article describes informally why this is so in the case of the targeted and untargeted scenarios. Although it *allows* adversaries to choose which student to attack after seeing the computation output, the untargeted scenario also accounts for adversaries who commit at the outset to attacking a particular student. Since this is exactly the class of adversaries that the targeted model contemplates, any computation that preserves privacy in the untargeted scenario also preserves privacy in the targeted scenario.

On the other hand, not every computation that preserves privacy in the targeted scenario preserves privacy in the untargeted scenario. For example, consider a computation that leaks information about a randomly chosen student. This computation might still meet the requirements of the targeted model, since the probability that the computation leaks information about the same student that the adversary chooses to attack might be sufficiently small. However, if the adversary can decide which student to attack after seeing the computation output, the adversary can take advantage of the leaked information and win the game

with an unacceptably high probability by choosing to attack the student whose private information the computation leaked.

By introducing the extensions to the untargeted attack and multiple release scenarios it is not meant to imply that the single release model formalized in Section IV.D is insufficient or incorrect. On the contrary, the legal analysis and technical arguments presented throughout Part IV suggest that the single release model is likely much stronger than many interpretations of what is needed to comply with FERPA. These extensions are presented for two reasons. First, these extensions do correspond to what may be privacy concerns in real-world uses of educational data, so it makes sense to consider them when examining privacy-preserving technologies to be used with such data. Second, it can be shown that differentially private computations preserve privacy in either extension to the model, making the argument that the use of differential privacy complies with the standard presented in FERPA more robust.

Further technical research is needed to extend the model to address situations in which student attributes are dependent and to select an appropriate value for the privacy parameter based on a legal standard of privacy protection. As noted in Section IV.B.2, the model treats the attributes of a student as being independent of the attributes of any other student even though this does not fully capture reality. For instance, if a student has a learning disability, it is more likely that his identical twin also has that disability. The analysis in Part IV could be modified to protect privacy when there is dependence between members of small groups of students. In addition to correlations between a few individuals such as family members, there could be global correlations among the students not accounted for by the model. In general, it is not possible to provide non-trivial inferential privacy guarantees when the adversary has arbitrary auxiliary information about correlations concerning members of a dataset.²³² However, guarantees can be given when the correlations are limited.²³³ Therefore, one direction for addressing this issue in the model is to alter the game mechanics to allow the adversary to construct limited types of correlations between students. In this case, one would still be able to give privacy guarantees when students are not independent, provided that the correlations between the students fit the ones allowed in the model. The authors intend to explore the suitability of this direction in future work.

In addition, both the privacy definition developed throughout Part IV and differential privacy are parameterized, i.e., they both refer to a privacy loss parameter ϵ that controls the level of privacy required (and,

232. See Dwork & Naor, *supra* note 52, at 95.

233. See Arpita Ghosh & Robert Kleinberg, *Inferential Privacy Guarantees for Differentially Private Mechanisms*, ARXIV 3, 7–8 (May 23, 2017), <https://arxiv.org/pdf/1603.01508.pdf> [<https://perma.cc/75S7-7MVY>].

as a result, the accuracy of the computations that can be performed). Intuitively, ϵ corresponds to the requirement that risk to privacy should be kept “very small,” and the question of setting ϵ (discussed in Section IV.D.3) corresponds to quantifying “very small.” Setting ϵ is likely to be a dynamic process, in which the value is adjusted over time, depending on the understanding of the various factors affected by the real-world use of differential privacy. It is important, however, to begin this process, and the question therefore reduces to developing and implementing a methodology for the initial setting of the privacy loss parameter. In future work, the authors intend to explore a methodology whereby an initial value for ϵ is selected to match the differential privacy guarantees to the intent underlying current best practices — regardless of whether policymakers’ expectations were actually met by a given standard and how it has been interpreted in practice.

V. DISCUSSION

On first impression, the gaps between the standards for privacy protection found in statutes and case law, on one hand, and formal mathematical models of privacy, on the other, may seem vast and insurmountable. Despite these challenges, this Article demonstrates a case where it is indeed possible to bridge between these diverging privacy concepts using arguments that are rigorous from both a legal and a mathematical standpoint. The main contribution of this Article is therefore its approach to formulating a legal-technical argument demonstrating that a privacy technology provides protection that is sufficient to satisfy regulatory requirements. This Section discusses the significance, as well as the policy implications, of this approach.

A. Introducing a Formal Legal-Technical Approach to Privacy

This Article’s primary contribution is a rigorously supported claim that a particular technical definition of privacy (i.e., differential privacy) satisfies the requirements of a particular legal standard of privacy (i.e., FERPA). The argument consists of two components. The first part of the argument provides support for the claim that FERPA’s privacy standard is relevant to performing differentially private analyses of educational records. This part is supported by a legal analysis of the text of the FERPA regulations as well as agency guidance on interpreting FERPA that pertains to releases of aggregate statistics. It also builds on results from the technical literature showing that the release of aggregate statistics can leak information about individuals.

The second part of the argument extracts a formal mathematical model from FERPA. This analysis identifies and characterizes an ad-

versary, formulates a privacy definition, and constructs a proof demonstrating that differential privacy satisfies the mathematical definition extracted from the regulation. Based on FERPA's definition of PII, this Article constructs a detailed model of an adversary, or a privacy attacker trying to learn private information from the system. This analysis is rooted in FERPA's conceptualization of a potential attacker as a "reasonable person in the school community, who does not have personal knowledge of the relevant circumstances."²³⁴ Using this language as a reference point, the model specifies the type of prior knowledge an adversary may have. The model also includes a specification of the information provided to the attacker from the privacy system, namely, the information designated by the school as directory information and the result of a particular statistical analysis. Finally, it defines the attacker's goal in terms of a violation of FERPA, i.e., disclosing non-directory PII about a student.

Next, following a well-established paradigm from the field of cryptography, this Article extracts a game-based definition of privacy based on FERPA's requirements. This privacy game is a thought experiment that aims to capture what it means for a statistical computation to preserve privacy under (a conservative interpretation of) FERPA. More specifically, the game aims to model what it means for non-directory PII to be protected from leakage. The game is formalized as an interaction between an adversary and a system that provides a statistical analysis, and this interaction is mediated by the mechanics of the game. To model the variety of settings in which a privacy-preserving mechanism may be challenged, the adversary is allowed to choose the directory information that is available for each student, the prior beliefs the adversary possesses about the private information for each student, and the particular student to target in the attack. These features of the model are formalized within the game by having the adversary supply each of these pieces of information to the game mechanics. The game mechanics generate hypothetical educational records that reflect the directory information supplied by the adversary, as well as the adversary's beliefs about the students' non-directory PII. Then, the game mechanics feed these records into the statistical analysis and supply the adversary with the result of the analysis. Finally, the game mechanics monitor whether the adversary is successful in guessing the non-directory PII of the targeted student using the results of the analysis she receives.

Key features of the approach are its mathematical formalism and its conservative design choices. These aspects of the approach help to construct a proof that differential privacy satisfies the definition of privacy extracted from FERPA. This approach provides a very strong basis for arguing that differentially private tools can be used to protect

234. Family Educational Rights and Privacy, 34 C.F.R. § 99.3 (2017).

student privacy in accordance with FERPA when releasing education statistics. Moreover, it provides a rationale for asserting that the use of differential privacy would be sufficient to satisfy a wide range of potential interpretations of FERPA's privacy requirements, including very strict interpretations of the regulations. Extensions to the privacy definition formulated in this Article are provided in Appendix II, which demonstrates the robustness of the argument that the use of differential privacy is sufficient to satisfy FERPA with respect to an even wider range of interpretations of the regulations.

B. Policy Implications and Practical Benefits of this New Approach

The analysis in this Article embraces a scientific understanding of privacy. In particular, it recognizes the importance of relying on formal, mathematical definitions of privacy. Adopting formal definitions is critical to ensuring that the privacy technologies being developed today will provide strong privacy protection, withstand future kinds of attacks, and remain robust over the long term despite the increasingly wide availability of big data and growing sophistication of privacy attacks. Bringing mathematical formalism to the law holds promise for addressing ongoing challenges in information privacy law.

This Article advocates the application of a scientific understanding of privacy and mathematical formalism within the practice of information privacy law, as elements of a new regulatory regime that will provide strong, long-term protection of information privacy. When evaluating whether a technology provides sufficient privacy protection in accordance with a legal standard, experts should justify their determinations with a rigorous analysis. In particular, experts' determinations should be based on a formal argument that is well-supported from both legal and technical perspectives, and their analysis should include a detailed description of a formal mathematical model extracted from the regulation.

In this way, a legal-technical approach to privacy can facilitate the real-world implementation of new privacy technologies that satisfy strong, formal definitions of privacy protection. A challenge to bringing formal definitions to practice in real-world settings is ensuring that the definitions that emerge in the mathematical study of privacy are consistent with the normative privacy expectations set in the law. Privacy regulations often rely on concepts and approaches that are fundamentally different from those underlying formal privacy models, and their requirements are inherently ambiguous. For instance, many information privacy laws have historically adopted a framing of privacy

risks that is based on traditional and mostly heuristic approaches to privacy protection like de-identification.²³⁵ This potentially creates uncertainty and risk for practitioners who would seek to use tools that do not rely on traditional techniques such as de-identification that are seemingly better supported by the law. However, an approach to formally modeling legal requirements that is conservative and consistent with a broad range of legal interpretations can provide a high level of assurance that the use of a privacy technology is sufficient under a legal standard.

Arguments such as the one presented in this Article can help lower the barrier for adoption of technologies that move beyond these traditional conceptions of privacy, including privacy technologies that adhere to formal privacy models. In particular, given a formalization of a regulation's privacy desiderata as a mathematical definition, privacy researchers can make and substantiate claims that certain computations satisfy the requirements of the formalization and hence comply with the regulation. This approach can be used to overcome ambiguities in legal standards for privacy protection and be used to design implementations of new privacy technologies that can be shown to satisfy a legal standard of privacy protection with more certainty than is afforded by alternative approaches. In turn, applications of this approach can support the wider use of formal privacy models, which provide strong guarantees of privacy protection for individuals whose information is being collected, analyzed, and shared using technologies relying on such models.

1. Benefits for Privacy Practitioners

This approach has potential applications and benefits for government agencies, corporations, research institutions, regulators, data subjects, and the public. Organizations such as research universities and corporations currently manage large amounts of personal data that hold tremendous research potential. Many of these organizations are reluctant to share data that may contain sensitive information about individuals due to recent high-profile privacy breaches and the specter of legal liability. Many government agencies, most notably statistical agencies, are interested in adopting public-facing tools for differentially private analysis. However, before sharing sensitive data with the public, agencies typically must demonstrate that the data release meets relevant regulatory requirements and satisfies various institutional policies, such as those related to any applicable internal disclosure limitation review. The proposed methodology could be used in this case to demonstrate that an agency's use of a formal privacy model satisfies its obligations to protect the privacy of data subjects pursuant to applicable laws such

235. See Altman et al., *supra* note 16, at 2068.

as the Privacy Act of 1974²³⁶ or the Confidential Information Protection and Statistical Efficiency Act.²³⁷

Similarly, corporations and research institutions that collect, analyze, and share statistics about individuals may seek to use differentially private tools or other tools based on formal privacy models, but wish to do so only with assurance that doing so will be in accordance with regulatory requirements for privacy protection. Formal modeling, especially when done conservatively, could enable actors within each of these sectors to begin using and sharing data with assurance that the risk of an administrative enforcement action is low. For example, a formal legal-technical approach could be applied towards satisfaction of the expert determination method of de-identifying protected health information in accordance with the HIPAA Privacy Rule. An expert applying a privacy-preserving technology could use this formal approach both to “determine[] that the risk is very small that the information could be used, alone or in combination with other reasonably available information, by an anticipated recipient to identify an individual who is a subject of the information” and to “[d]ocument[] the methods and results of the analysis that justify such determination.”²³⁸ This approach can be used to provide documentation of compliance with other privacy laws as well, providing a high degree of confidence in justifying — both *ex ante* and *ex post* — that an organization’s practices with respect to privacy-preserving analysis and publishing meet the requirements of the law.

In addition, regulatory agencies such as the Department of Education and corresponding state agencies often develop specific guidance on implementing privacy safeguards in accordance with regulatory requirements. Agencies could employ a legal-technical approach to privacy in the future to evaluate and approve the use of new technologies, based on a rigorous determination regarding whether they satisfy existing regulatory and statutory requirements for privacy protection. This approach could also create benefits for data subjects, and the public at large, facilitating the use of tools that provide stronger privacy protection than data subjects have been afforded in many historical cases, as well as enable new and wider uses of data that carry benefits to society.

2. Benefits for Privacy Scholars

Precise mathematical modeling can also enhance scholars’ understanding of the privacy protection afforded by various laws and technologies. Consider a technology that offers privacy guarantees that are

236. Privacy Act of 1974, 5 U.S.C. § 552a (2012).

237. 44 U.S.C. § 3501 note (2012) (Confidential Information Protection and Statistical Efficiency Act).

238. *See* Security and Privacy, 45 C.F.R. § 164.514(b)(1) (2017).

weaker than those provided by differential privacy (but potentially outperforming differential privacy in other respects, such as utility). An expert seeking to demonstrate that this technology satisfies FERPA may choose to extract and justify a weaker model of FERPA than the one extracted in this Article. This can enable a comparison of the various models — and, consequently, multiple privacy definitions — that have been developed based on FERPA. Such a comparison makes it possible to understand, in precise terms, the privacy guarantees offered by each model, as well as the assumptions relied upon by each model. Because the privacy guarantees and assumptions of each model are made explicit, policymakers, scholars, and the public can better understand the tradeoffs of each model. This provides a basis for scholarly and public policy debates regarding the privacy guarantees that should be required from normative and technical perspectives and the types of assumptions that are reasonable in a privacy analysis. Feedback from the normative debates can, in turn, help inform practitioners constructing models for evaluating privacy technologies and provide more clarity to those developing the technologies themselves.

Comparing the models extracted from different regulations may inform scholars' understanding of the ways in which regulations differ, including evaluating the relative strength of protection provided by different privacy laws. The large number of privacy regulations that are applicable based on jurisdiction and industry sector make regulatory compliance complex.²³⁹ The high level of abstraction offered by the analytical approach in this Article has the potential advantage of simplifying the analysis of privacy regulation and its applicability to particular privacy technologies. This approach can serve as a new lens for examining and comparing privacy regulations. For instance, it could be used to identify essential features that regulations ought to share and be used to reform such regulations in the future.

Although this Article does not advocate any particular regulatory changes, it argues that evaluations of new technologies should be done rigorously so as to ensure their design and implementation adhere to the legal standards of privacy protection. This Article highlights a number of gaps between the current legal framework for privacy protection and recent advances in the scientific understanding of privacy. Updating information privacy laws based on a modern scientific approach to privacy could bring greater certainty and stronger privacy measures to practice. Using mathematical formalism, reforms can be made to specify requirements that are meaningful and accurate and yet permissive enough to promote the development and use of innovative privacy technologies. Furthermore, mathematical formalism can serve as a tool for guiding regulatory decisions based on rigorous conceptions of privacy.

239. See Schwartz & Solove, *supra* note 9, at 1827–28.

To begin to close the gap between privacy law and technology, this Article recommends that future regulations aim to define the objective of a privacy standard, rather than providing a list of permissible privacy protection technologies or certain identifiers that, if redacted, presumably render the data de-identified.²⁴⁰ This shift in regulatory approach would enable practitioners and legal scholars to assess whether new technologies meet the goals provided by the regulations. This would, in turn, help ensure that new privacy technologies can be brought to practice with greater certainty that they satisfy legal standards for privacy protection.

C. Applying this Approach to other Laws and Technologies

This Article demonstrates that it is possible to construct a rigorous argument that a privacy technology satisfies a legal standard of privacy protection, using differential privacy and FERPA for illustration. However, in doing so, it does not argue that using differential privacy is *necessary* to satisfy FERPA's requirements for privacy protection, nor that privacy-preserving technologies other than differential privacy are insufficient to satisfy FERPA. It also does not claim that the model presented in Part IV is the only possible model that can be extracted from FERPA. Rather, it is likely there are several possible models that can be extracted from FERPA to support the use of various privacy technologies.

In future work, the authors anticipate extending the technical results presented in this Article to a general methodology that can be applied to privacy models other than differential privacy and regulations apart from FERPA. The computational focus of this analysis works at a level of abstraction in which many of the differences between particular regulations are likely to become less important. Achieving this level of abstraction, however, will require deepening our understanding of how different regulations lend themselves to these kinds of analyses. It will also require extending our "toolkit," i.e., the collection of argument paradigms that can be used in making a claim that a differential privacy, or another formal privacy model, satisfies a regulatory standard of privacy protection.

Establishing a general methodology will require examining regulations that are different from FERPA in terms of the type of analysis

240. For further discussion and examples illustrating how regulatory requirements could be stated in terms of a general privacy goal, see Letter from Salil Vadhan et al. to Dep't of Health & Human Servs., Office of the Sec'y and Food & Drug Admin. (Oct. 26, 2011), <https://privacytools.seas.harvard.edu/files/commonruleanprm.pdf> [<https://perma.cc/Z4FP-RZ9U>]; Letter from Alexandra Wood et al. to Jerry Menikoff, M.D., J.D., Dep't of Health & Human Servs. (Jan. 6, 2016), <https://www.regulations.gov/document?D=HHS-OPHS-2015-0008-2015> (last visited May 5, 2018).

that would be required. Two examples of regulations with privacy requirements that differ in significant ways from FERPA's include the HIPAA Privacy Rule and Title 13 of the U.S. Code. This brief overview discusses some of the ways in which the privacy standards in these laws differ from that set forth by FERPA. These differences will likely require future modifications to the argument made in this Article in order to achieve a generally applicable methodology.

The HIPAA Privacy Rule presents a unique set of challenges for modeling its requirements formally. The Privacy Rule sets forth two alternative methods of de-identification: a detailed safe harbor method and a method relying on expert determination. While FERPA includes a description of an adversary, its knowledge, and its goal as a "reasonable person in the school community, who does not have personal knowledge of the relevant circumstances, to identify the student with reasonable certainty,"²⁴¹ neither the regulatory language nor relevant guidance for HIPAA includes such an explicit description of the envisioned adversary. Furthermore, HIPAA's expert determination method delegates the responsibility to confirm that the "risk [of identification] is very small" to a qualified statistician,²⁴² but neither the text of HIPAA nor guidance from the Department of Health & Human Services provides the form of reasoning an expert must perform to argue that the risk is small. For instance, where the guidance describes "principles for considering the identification risk of health information," it notes that such "principles should serve as a starting point for reasoning and are not meant to serve as a definitive list."²⁴³ The provision of general principles meant to serve as a starting point rather than specific criteria to be met creates challenges for modeling the regulatory requirements with precision.

Other challenges arise when modeling privacy standards that lack detail, such as the confidentiality provisions of Title 13 of the U.S. Code, the statute which establishes the authority of the U.S. Census Bureau and mandates the protection of data furnished by respondents to its censuses and surveys. The model extracted from FERPA is based on detailed language from the regulations and is informed by a large body of guidance from the Department of Education. In comparison, Title 13's privacy standard is succinct, providing only that the Census Bureau is prohibited from "mak[ing] any publication whereby the data furnished by any particular establishment or individual under this title can be identified."²⁴⁴ Rather than delineating specific privacy requirements in regulations or guidance, the Census Bureau delegates much of the interpretation of the confidentiality requirements of Title 13 to the

241. Family Educational Rights and Privacy, 34 C.F.R. § 99.3 (2017).

242. 45 C.F.R. § 164.514(b).

243. OFFICE FOR CIVIL RIGHTS, DEP'T OF HEALTH & HUMAN SERVS., *supra* note 4, at 13.

244. 13 U.S.C. § 9(a)(2) (2012).

technical experts of its internal Disclosure Review Board. Modeling the privacy requirements of Title 13 would likely require heavy reliance on sources beyond the text of the statute, such as policies published by the Census Bureau and prior determinations made its Disclosure Review Board.

In future work, the authors intend to model additional laws such as the HIPAA Privacy Rule and Title 13, as well as other privacy technologies and to advocate further application of this approach by other scholars and practitioners. The results of this research are anticipated to point to alternative modeling techniques and, ultimately, inform a more general methodology for formally modeling privacy regulations.

D. Setting the Privacy Parameter

As noted in Part IV, the privacy definition extracted from FERPA is parameterized by a numeric value denoted ϵ . Any real-world implementation of a privacy technology relying on this definition from FERPA would require the appropriate value of ϵ to be determined. An understanding of the parameter ϵ can be informed, in part, by the definition of differential privacy, which is also parameterized by a value often denoted ϵ . As illustrated by the proof presented in Appendix I demonstrating that differential privacy satisfies the requirements of the definition extracted from FERPA, the parameter ϵ is used similarly in both definitions. As it relates to the differential privacy definition, the parameter ϵ intuitively measures the increase in risk to an individual whose information is used in an analysis. In order to provide strong privacy protection, the recommendation would be to keep ϵ “small.” However, a smaller ϵ implies less accurate computations. Setting this parameter therefore involves reaching a compromise between data privacy and accuracy. Because this tradeoff has both technical and normative components, choosing a value of ϵ should also be based on both technical and normative considerations.

In future work, the authors plan to explore methodologies for setting ϵ . A setting of ϵ controls both the level of privacy risk and quality of analysis. Arguments with respect to setting ϵ in real-world settings will likely rely both on a quantitative understanding of this tradeoff and on legal-normative arguments for balancing these concerns. In particular, relevant legal and normative concerns may include policymakers’ expectations in terms of this tradeoff between privacy protection and information quality, as well as the expectations of data subjects.

VI. CONCLUSION

As new privacy technologies advance from the research stage to practical implementation and use, it is important to verify that these

technologies adhere to normative legal and ethical standards of privacy protection. Being able to demonstrate that a privacy technology satisfies a normative standard will be critical to enabling practitioners to adopt such tools with confidence that they have satisfied their obligations under the law and their responsibilities to the subjects of the data.

However, making a rigorous, substantive claim that the protection offered by a privacy technology meets a normative standard requires the development of a common language as well as formal argumentation. These features are essential to bridging the gaps between divergent normative and technical approaches to defining and reasoning about privacy. Towards this end, this Article details a combined legal-technical argument for connecting the privacy definition found in a legal standard, FERPA, and that utilized by a privacy technology, differential privacy. This argument is grounded in a legal analysis of FERPA and uses mathematical formalism to address uncertainty in interpreting the legal standard.

Instrumenting the law with a modern scientific understanding of privacy, together with precise mathematical language and models, could help guide the development of modern conceptions of privacy in the law. This approach could also facilitate the development and implementation of new privacy technologies that demonstrably adhere to legal requirements for privacy protection. Introduction of combined legal-technical tools for privacy analysis such as these may also help lawmakers articulate privacy desiderata that more closely match today's information regime, in which a large number of information sources can be accessed and analyzed in ways that lead to breaches of individual privacy. Such techniques can provide robustness with respect to unknown future privacy attacks, providing long-term, future-proof privacy protection in line with technical, legal, and ethical requirements for protecting the privacy of individuals when handling data about them.

APPENDIX I. PROVING THAT DIFFERENTIAL PRIVACY SATISFIES FERPA

Part IV describes the process of extracting a formal model of FERPA's privacy requirements, as well as a mathematical definition of a sufficient level of privacy protection provided by a computation over non-public student data in this model. The model is envisioned as a privacy game, in which a hypothetical adversary tries to beat a computation by using the output of the computation to successfully guess the private information of a specific student whose records were part of the input data fed to the computation.

Section IV.C.2 explains why it would be unreasonable to expect an adversary to lose every time, even when playing against a computation

that provides “perfect privacy,” since the adversary always has some possibility of winning without having access to the output of the computation output — if only by random chance. Instead, the model considers a computation to preserve privacy if no adversary can win the game against it “too often,” the intuition being that no adversary should be substantially more successful at guessing student information when she has access to the output of the computation than some adversary would have been without having access to the output. In this discussion, the term “real-world scenario” describes the game in which an adversary plays against a computation while having access to the computation output and the term “ideal-world scenario” describes the game in which an adversary plays against a computation *without* having access to the computation output. These scenarios were illustrated in Figure 7 and Figure 8, respectively.

This modeling exercise provided a precise, mathematical definition of privacy. Proving that a statistical computation meets the formal definition of privacy extracted from FERPA shows that the computation provides a sufficient level of privacy protection to satisfy the privacy requirements of FERPA. This Appendix shows that every differentially private computation meets this definition of privacy. That is, for every real-world adversary playing the game against a differentially private computation, there exists an ideal-world adversary — who has the same *a priori* knowledge but does not see the computation’s output — that wins the game with nearly the same probability. The proof derives a sequence of games, enumerated H_0 through H_4 . These games are referred to as “hybrid” because they represent steps between the ideal-world (H_0 , identical to the game in Figure 8) and real-world (H_4 , identical to the game in Figure 7) scenarios. The claim is that the adversary in the real world does only a little better than an adversary in the ideal world. To substantiate this claim, the proof considers any pair of consecutive hybrid games, denoted H_i and H_{i+1} , and demonstrates that, for any adversary participating in the hybrid game H_{i+1} , there exists an adversary (with the same *a priori* knowledge) for the hybrid game H_i with an identical or almost equal winning probability. Intuitively, this means that the difference in privacy loss between two consecutive games is null or very small, and hence it is also the case that the accumulated difference between the privacy loss in the ideal-world game, H_0 , and the real-world game, H_4 , is small. Given that there is no privacy loss in the ideal-world game, one concludes that the privacy loss in the real-world game is sufficiently small.

Some mathematical notation is used in the diagrams and in the discussion below. First, $\bar{d} = (d_1, \dots, d_n)$ denotes a list of directory information, and each d_i denotes the directory information for student i . Similarly, $\bar{P} = (P_1, \dots, P_n)$ denotes the list of distributions over private

student information, and P_i denotes the distribution that describes student i 's private information. Next, each p_i represents the "actual" value for student i 's private information that is sampled from P_i . In addition, s denotes the targeted student, and f is a function over private student information, representing the adversary's goal. DB represents a database of student information: the i th row in DB consists of d_i and p_i ; that is, that i th record consists of the directory information for student i and student i 's private information. Finally, c denotes the computation output and g_s is the adversary's guess about the private information of student s . H_0 describes the ideal world given in Figure 8, and H_4 describes the real world given in Figure 7.

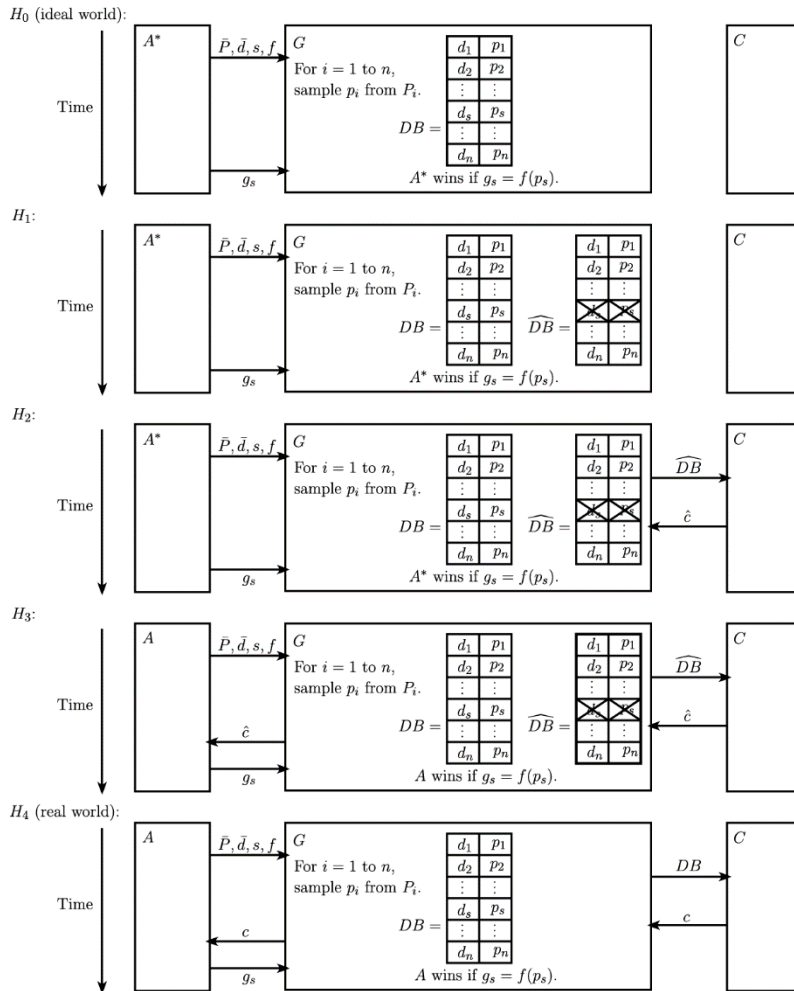


Figure 9: Progression of hybrid games.

The argument proceeds by showing that there is little to no privacy loss between the games appearing consecutively in this proof. Each of the hybrid games is reviewed in turn.

Hybrid H_0 (the ideal-world game):

The starting point of the argument is the ideal-world game. In this scenario, the game mechanics generate a database with the student information but perform no computation on it and pass no information on to the adversary. The adversary makes a guess about the private information of the targeted student using only his *a priori* knowledge about the student.

Hybrids H_0 and H_1 :

The first step is a modification of the ideal-world hybrid H_0 . In the ideal-world game, the game mechanics generate the database of all students; this database is referred to as DB . In H_1 , the game mechanics also generate a second, related, database $\widehat{DB}$ that is identical to the original DB except that the entry corresponding to the attacked student, s , is removed.

In comparing how an adversary A^* fares in the two games, it is important to note that no information flows from the game mechanics to the adversary in either the ideal-world game or H_1 . Therefore, for any adversary A^* the winning probability is identical in both games.

The only difference between H_0 (the ideal world) and H_1 occurs internally within the game mechanics. In H_0 , the game mechanics create a database DB of student information from the directory information $\bar{d}$ and private student information p_i sampled from the distributions over student information $\bar{P} = (P_1, \dots, P_n)$. The game mechanics in H_1 also take these steps. The game mechanics then create a new database $\widehat{DB}$ that is identical to DB , except that the record of student s has been removed.

All entries of $\widehat{DB}$ are chosen independently from the value p_s of student s . These entries include $d_{s'}, p_{s'}$ for any student s' different than s . Because all these values are independent of p_s , revealing $\widehat{DB}$ fully or in part would not yield any information about p_s .

Hybrids H_1 and H_2 :

In Hybrid H_2 , the game mechanics apply the mechanism on the database $\widehat{DB}$ and receive the result of the computation $\hat{c}$. However, this result is not transferred to the adversary, and hence the adversary's view in H_2 is identical to the adversary's view in H_1 . It follows that for any adversary the winning probability is identical in both games.

Hybrids H_2 and H_3 :

H_3 differs from H_2 in that the adversary A in H_3 sees the result of C computed on $\widehat{DB}$ before making his guess while the adversary A^* in H_2 does not get any information from the game mechanics. However, by the note above, the outcome of the computation $\hat{c} = C(\widehat{DB})$ does

not carry any information about p_s and hence does not help A obtain any advantage over A^* .²⁴⁵

Hybrids H_3 and H_4 :

The real-world game differs from H_3 in that, in the real-world game, the game mechanics do not create a database from which the information of student s is removed. Instead, the game mechanics invoke C on the database DB containing the actual data for student s , and give this result to the adversary. As DB contains the student's private information p_s , there is a risk that the output of the computation, which depends on DB , conveys information about p_s , and therefore a risk to the privacy of student s . The fact that C is a differentially private computation helps show that this is not the case. Recall from Section III.A that if a computation is differentially private, then its outputs when invoked on two neighboring databases (i.e., databases that differ on one entry) will be similar. This implies that, for any adversary A , the winning probability in the real-world game is similar to the winning probability in game H_3 .

Hybrid H_4 (the real-world game):

In this scenario, the computation is performed on the actual student data and the adversary receives the computation output. The adversary can use any knowledge gained by seeing the computation output to more accurately guess the private information of the targeted student. Yet, seeing the output of the computation performed on the actual student data does not give the adversary much of an advantage as compared to game H_3 if the computation is differentially private.

In sum, the proof demonstrates that no privacy is lost transitioning from H_0 to H_3 , and that there is only a small loss in privacy stepping from H_3 to H_4 . Consequently, the privacy loss in the real world (H_4) is not significantly higher than that in the ideal world (H_0). Since there is no privacy loss in the ideal-world game, it follows that there is no significant privacy loss also in the real-world game. Note that the argument holds for any differentially private computation. Thus, every differentially private computation meets the privacy definition given in Section IV.D.2 and fulfills the privacy requirements of FERPA, up to the limitations of the model.

245. Making this argument formally is a bit subtler. Recall the need to demonstrate that, for any adversary A participating in H_3 , there is an adversary A^* participating in H_2 that has the same winning probability. This is done by defining an adversary A^* that sends the same initial information as A does to the game mechanics. This A^* samples database $\widehat{DB}$ consisting of the directory information of all students except s and fake entries for these students ($\tilde{p}_i$ sampled from P_i for each student i). (Note that database $\widehat{DB}$ is not the same database as the $\widehat{DB}$ described above, as $\widehat{DB}$ is sampled by A^* while $\widehat{DB}$ is sampled by the game mechanics.) Then, A^* applies C on $\widehat{DB}$ to obtain $\tilde{c} = C(\widehat{DB})$, which it passes to A . A^* takes the guess made by A and uses it as its own guess. Because neither $\hat{c}$ (on which the A in H_3 bases its guess) nor $\tilde{c}$ (on which the A in H_2 bases its guess) contains any information specific to student s , the winning probability in both games is the same.

APPENDIX II. EXTENSIONS OF THE PRIVACY GAME

Section IV.G explained how the model presented in Section IV.D can be extended to account for even stronger adversaries. The focus was on adversaries who can choose which student to target *after* seeing the computation output and who have access to the output of *multiple* computations performed on the student data. This Appendix describes these extensions in further detail and sketch out how they might be modeled.

The privacy definitions for these extensions are more mathematically involved than the definition considered in Section IV.D and Appendix I. Nevertheless, these definitions can be formalized, and an argument analogous to that presented in Appendix I can be made that differential privacy also provides the required protection with respect to these extensions.

A. Untargeted Adversaries

Recall that Figure 7 above depicts a *targeted* game. In this game, the adversary has a particular student s in mind, and the adversary's goal is to improve on guessing this specific student's private information. For instance, the adversary may examine publications of educational data from a Washington, D.C. school with the purpose of learning sensitive information about the President's daughter's grades. This *targeted* aspect of the adversary is manifested in the game by the adversary's declaration of the target student s and the target function f . The privacy definition presented in Section IV.D.2 refers to targeted adversaries, and a computation C that satisfies the requirements of the definition provides protection against such adversaries.

The game in Figure 7 can be modified to also consider *untargeted* adversaries, i.e., adversaries that are interested in learning the non-directory PII of *any* student in the dataset. An example of such a scenario could be a data broker mining a data release for private information about any student it can. The data broker's goal is to glean non-publicly available information from the data release that it can sell to third parties such as marketers. Although the data broker wants to learn information specific to individual students, the data broker is not committed to targeting a particular student. Instead, it is hoping to learn private information specific to any student in the dataset. Hence, the game and definition presented in Section IV.D — which would require the data broker to commit upfront to targeting a single student — do not capture this scenario.

1. Mechanics

The untargeted game is shown graphically in Figure 10 below. As before, the adversary chooses a directory of public student information and associates with each student in that directory a probability distribution that describes the adversary's beliefs about the non-directory PII of that student. Unlike the previous game, at this point the adversary does not commit to trying to guess the non-directory PII of any particular student. Instead, the adversary passes just the directory and the probability distributions to the game mechanics.

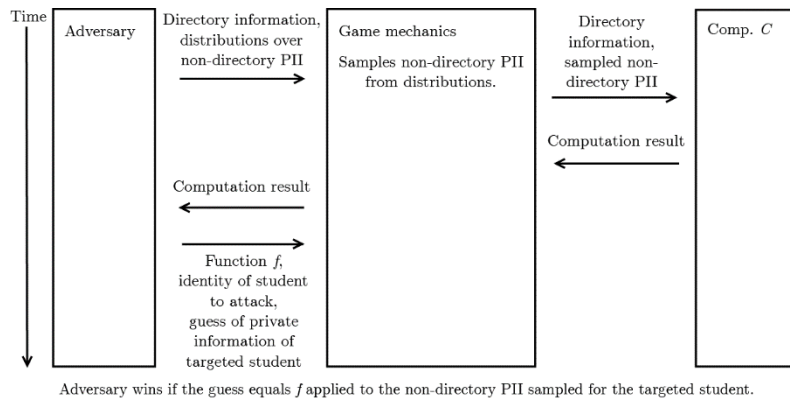


Figure 10: Untargeted scenario. Note that the adversary does not choose a student to attack and a function f until after seeing the computation result.

The game mechanics build a database of non-directory student information by sampling randomly from each of the probability distributions given by the adversary and pass this database, with the directory, to a computation C , which outputs some result.

The game mechanics then give this result to the adversary. After seeing this result, the adversary chooses a student to try to attack, and makes a guess about an aspect of that student's non-directory PII. The adversary gives the identity of this student, a function f that it tries to guess, and the guess to the game mechanics. The game mechanics then declare that the adversary has won if the adversary's guess matches the result of applying the function f to that chosen student's non-directory PII. Otherwise the adversary is considered to have lost the game.²⁴⁶

246. For a more detailed explanation, see *supra* note 223.

2. Discussion of the Untargeted Scenario

The privacy definition for the untargeted scenario is subtler and more mathematically involved than the definition given in Section IV.D.2 for the targeted scenario. The detailed definition and proof is omitted from this Article.

It is important to note, however, that while protection against untargeted adversaries implies protection against targeted adversaries, it is not the case that protection against targeted adversaries necessarily implies protection against untargeted adversaries. To see why this is the case, consider a hypothetical privacy protection mechanism that chooses a random student among a district's 10,000 students and publishes this student's entire record. A targeted adversary would only have a chance of one in 10,000 to learn new information about the targeted student. An untargeted adversary, on the other hand, would always be able to learn new information about some student.²⁴⁷

B. Multiple Releases

The second extension is designed to account for scenarios in which an adversary has access to the outputs of multiple computations on student data. The formalization above currently only accounts for a single release. That is, it assumes that the adversary only has access to the output of a single computation performed on the private data. However, it is also important for a robust model to consider the multiple release scenarios. In fact, the FERPA regulations require that, before de-identified data can be released, a determination must be made that "a student's identity is not personally identifiable, whether through single or multiple releases."²⁴⁸

One can model this requirement with a game in which the adversary sees the output of multiple computations performed on the student data before guessing the private student information. Such a game is given in Figure 11. Like the targeted game presented in Section IV.D, the game starts with the adversary supplying the directory information, distributions over private student information, the identify of a student to attack, and a function over student information f . The game mechanics create a database of student information from the directory information and samples drawn from the private attribute distributions. Unlike the games presented above, here the game mechanics invoke multiple computations (enumerated C_1 through C_m) on the database.

247. Note that the weak privacy protection considered in this Section is given only as an example. The authors do not recommend employing such weak protection in practice.

248. Family Educational Rights and Privacy, 34 C.F.R. § 99.31(b)(1) (2017) (emphasis added).

After each computation, the game mechanics pass the computation result to the adversary. After all computations have been completed, the adversary makes a guess about the private information of the targeted student and wins if the guess equals f applied to the student's database entry.

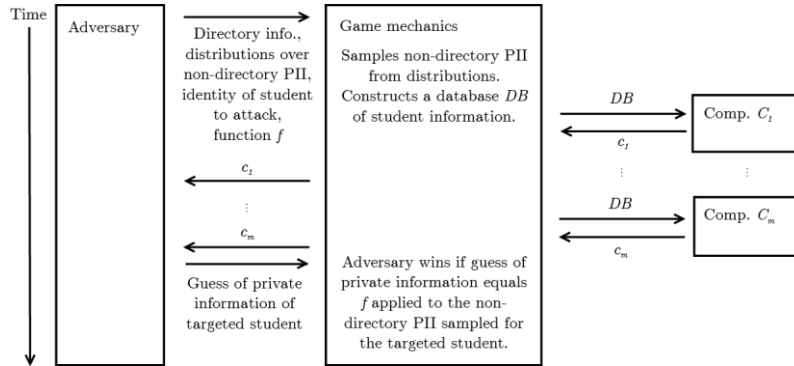


Figure 11: Multiple release scenario.

This formulation of the game reflects a targeted scenario, where the adversary commits ahead of time to a student to attack. One could alternatively formulate the game for an untargeted scenario with multiple releases, where the adversary only decides which student to attack after he has seen the output of all the computations. Either scenario could be further modified by allowing the adversary to adaptively choose the computations to be performed. This game proceeds as follows. The adversary first chooses an initial computation C_1 . Thereafter, after seeing each c_i (the result of computation C_i), the adversary chooses C_{i+1} , the next computation to be performed on the student data.