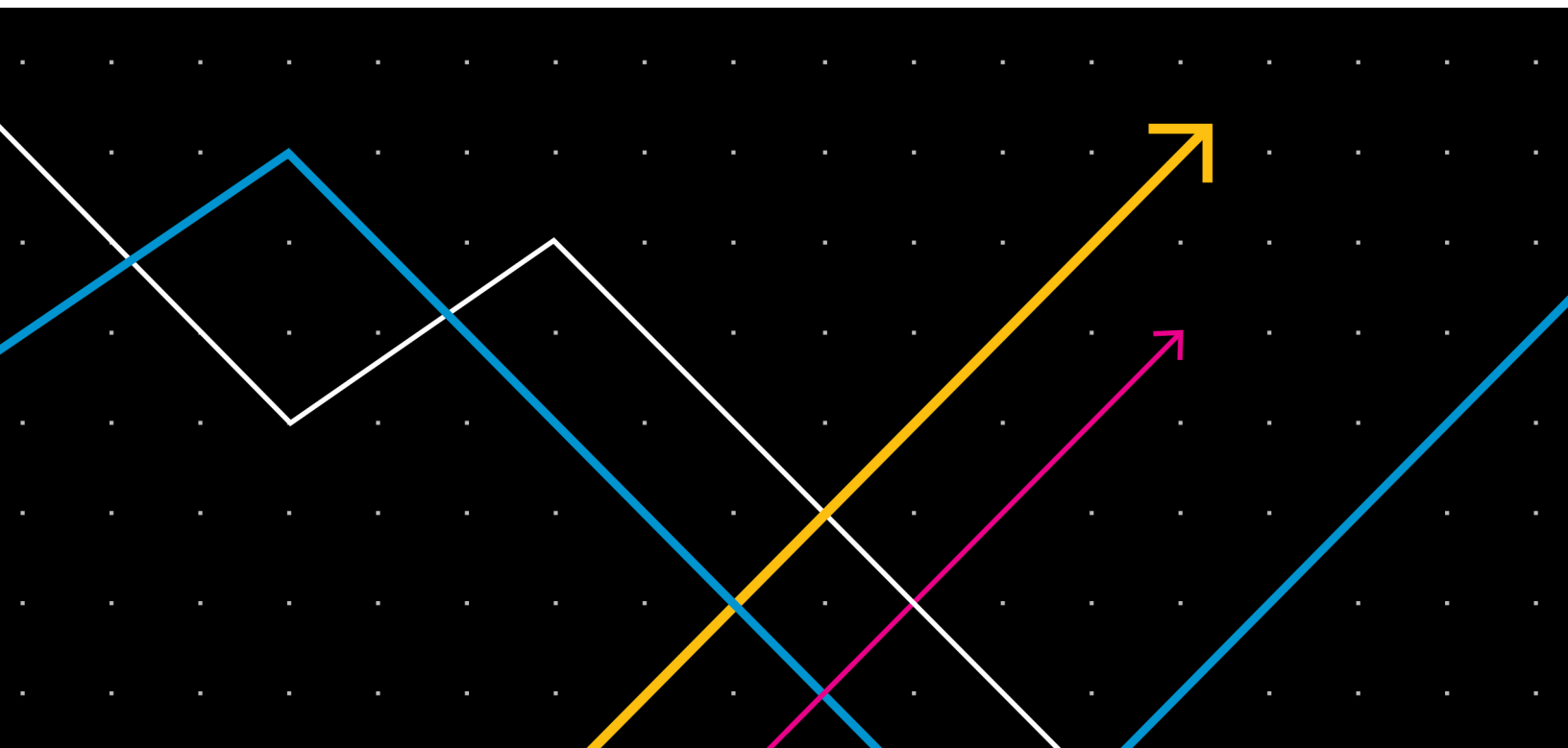




RESEARCH REPORT

Decennial Disclosure

An Explainer on Formal Privacy and the TopDown Algorithm

Claire McKay Bowen

Aaron R. Williams

Madeline Pickens

September 2022



ABOUT THE URBAN INSTITUTE

The nonprofit Urban Institute is a leading research organization dedicated to developing evidence-based insights that improve people's lives and strengthen communities. For 50 years, Urban has been the trusted source for rigorous analysis of complex social and economic issues; strategic advice to policymakers, philanthropists, and practitioners; and new, promising ideas that expand opportunities for all. Our work inspires effective decisions that advance fairness and enhance the well-being of people and places.

Contents

Acknowledgments	iv
Decennial Disclosure	1
Introduction to the 2020 Census and Data Privacy	1
Data Privacy Definitions and Terminology	4
Data Privacy Methodology Workflow	6
Introduction to Formal Privacy	10
Formal Privacy	10
Differential Privacy and Other Formally Private Definitions	11
Privacy-Loss Budget	13
Global Sensitivity	16
Gaussian Mechanism	17
Models of Differential Privacy	19
Introduction to 2020 Disclosure Avoidance System	21
Privacy and Utility Measures	21
Statistical Disclosure Control Method	22
Takeaways and Ongoing Challenges	26
Notes	29
References	30
About the Authors	31
Statement of Independence	32

Acknowledgments

This report was funded by the Tableau Foundation. We are grateful to them and to all our funders, who make it possible for Urban to advance its mission.

The views expressed are those of the authors and should not be attributed to the Urban Institute, its trustees, or its funders. Funders do not determine research findings or the insights and recommendations of Urban experts. Further information on the Urban Institute's funding principles is available at urban.org/fundingprinciples.

The authors thank the following individuals who generously provided invaluable feedback that greatly improved this explainer:

- Constance Citro, Senior Scholar, Committee on National Statistics at the National Academies of Sciences, Engineering, and Medicine
- Ron Prevost, Research Professor, Massive Data Institute, McCourt School of Public Policy at Georgetown University
- Leslie Reynolds, Research Support Specialist, Program on Applied Demographics, Cornell Jeb E. Brooks School of Public Policy
- Joseph Salvo, Fellow, Social and Decision Analytics Division at the University of Virginia Biocomplexity Institute
- Meghan Stuessy, Analyst, Government Organization and Management at Congressional Research Service
- David Van Riper, Director of Spatial Analysis, Institute for Social Research and Data Innovation at the University of Minnesota
- Jan Vink, Extension Associate, Program on Applied Demographics, Cornell Jeb E. Brooks School of Public Policy
- Izzy Youngs, Research Specialist, Massive Data Institute, McCourt School of Public Policy at Georgetown University

Decennial Disclosure

Although collecting more and better data can provide great benefits to society, such as furthering medical research or targeting investments to those most in need, data privacy concerns surface from those charged with protecting data when that information can be de-anonymized and used maliciously.

For example, the US Census Bureau conducted a simulated attack on the 2010 Decennial Census and discovered they could reidentify about one-sixth of the US population using publicly available data (such as name, sex, and age) from external sources, like public social media profiles (Leclerc 2019). This type of attack on the 2020 Decennial Census has the potential to be even more disclosive because of the detailed information collected, such as more race and ethnicity categories, that could lead to more individuals being identified with great specificity. The reconstruction attack results and the more detailed information available in the decennial census motivated the Census Bureau to update their Disclosure Avoidance System (DAS) from traditional statistical disclosure control methods to a formally private method—the TopDown Algorithm—for the 2020 Decennial Census.

However, this drastic change in how data privacy and confidentiality was defined for the 2020 DAS caused significant friction between the US Census Bureau and census data users. For instance, leaders from states, counties, cities, and towns rely on census data for school planning, budgeting, social program provisions, redistricting, revenue sharing, and a multitude of other statutory requirements. These data users want more accurate data at granular geographic areas and fear that the updated DAS will lead to incorrect public policy decisions.

This explainer aims to help readers better understand what formal privacy is and how the TopDown Algorithm works. The explainer is also a continuation of “Personal Privacy and the Public Good: Balancing Data Privacy and Data Utility” (Bowen 2021) and we encourage readers to read that report first.

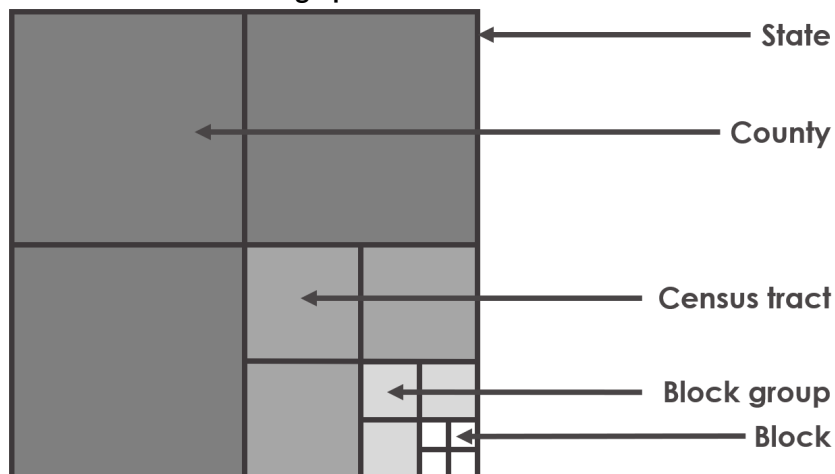
Introduction to the 2020 Census and Data Privacy

The decennial census data products affect how the United States apportion the 435 seats for the United States House of Representatives, redistrict voting lines, plan for natural disasters, and conduct many other purposes. Therefore, the Census Bureau’s mission is “...to count everyone once, only once,

and in the right place.” With this goal in mind, the US Census Bureau collects information on every person and household at various geographic levels for the United States (figure 1).

FIGURE 1

US Census Bureau’s Geographic Levels



Source: Authors' illustration.

Because the US Census Bureau collects such detailed information about individuals, the 1929 Census Act requires the Census Bureau to alter decennial census data with privacy-preserving methods.¹ Specifically, this act enforces that individuals and businesses cannot be identified in publicly released data. Since then, several laws have required the Census Bureau to protect census data products. The most cited law is Title 13 of the US Code, which protects individual-level data. A discussion on the US Census Bureau's history of privacy protection and the interpretation of Title 13 is beyond the scope of this explainer. Interested readers should see work by Hotz and Salvo (2022).

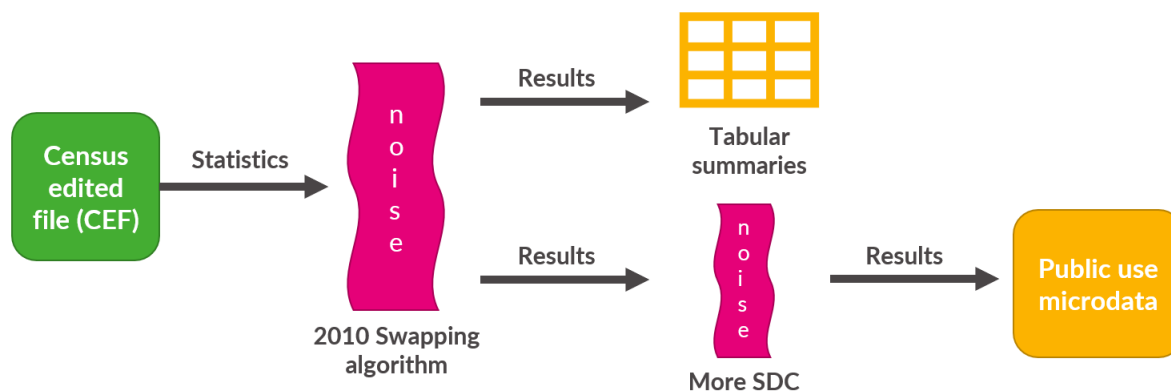
Note that the Census Bureau uses other important geographic levels not shown in figure 1, such as places, minor civil divisions, and American Indian and Alaska Native areas. We do not highlight these other areas because the US Census Bureau focuses on the geographic levels shown in figure 1 when protecting the data.

In addition to the legal requirements, some people might not be ethically comfortable with data users knowing certain characteristics of a group or area, such as where many people of certain racial groups live (e.g., Asian Americans, considering the legacy of internment camps during World War II and the racial prejudice and discrimination that recently accompanied the COVID-19 pandemic). On the other hand, data users, such as Asian American advocacy groups, might want access to such data

to provide targeted services like financial support for Asian-owned businesses that struggled during the pandemic. This is another example of the tension between data privacy and data utility.

The Census Bureau refers to the overall methodology to protect a census data product as the DAS. The last time US Census Bureau updated the decennial DAS was for the 1990 Census, by applying data swapping (figure 2 provides a summary of the 2010 DAS process). The Census Bureau periodically updates the DAS because the technological landscape is constantly evolving. For instance, modern smart phones have more computational power than the average desktop computer had in 2010.

FIGURE 2
2010 Disclosure Avoidance System Framework



Source: Authors' illustration.

To reassess if the US Census Bureau needed to update the DAS, they conducted a database reconstruction attack. In other words, this type of attack evaluates whether too many independent statistics are published based on confidential data to recreate the underlying confidential data with little or no error. The Census Bureau tested this by

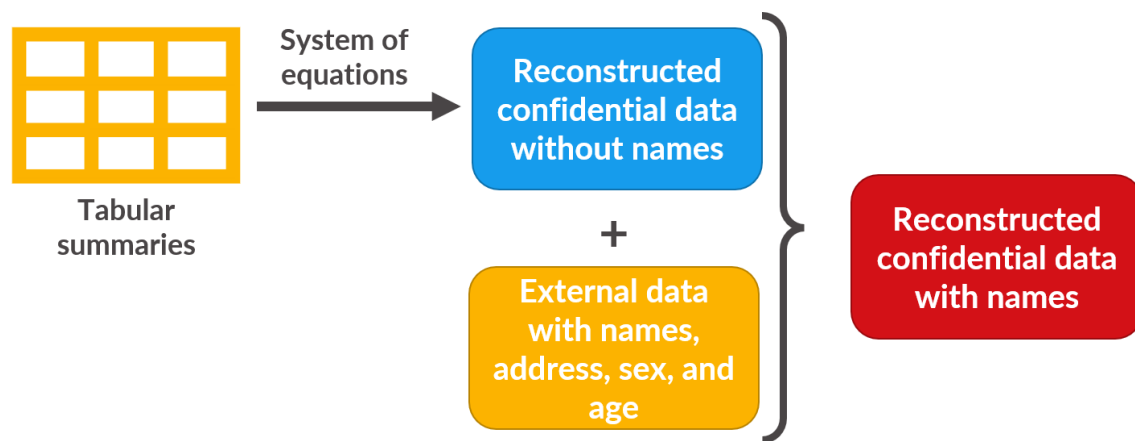
1. recreating the individual level 2010 Census (i.e., age, sex, race, and Hispanic or Non-Hispanic ethnicity for every individual in each census block) from nine summary tables, and then
2. uniquely identifying approximately one in six records using publicly available data, such as what could be found on social media profiles (Leclerc 2019). This rate is higher for smaller groups, such as underrepresented racial groups in rural areas.

Figure 3 illustrates a high-level explanation of how the US Census Bureau executed the reconstruction attack. For more detailed information about the reconstruction attack, see “The Census

Bureau's Simulated Reconstruction-Abetted Re-identification Attack on the 2010 Census" webinar materials.²

Although the rate of reidentification from the 2010 Census is troubling, a potential data attacker could not confirm whether (1) a match was correct or (2) the reconstructed data were correct before the match without access to the Census Edited File, the confidential data that have been edited for mistakes. Also, the Census Bureau has received criticism for their reconstruction attack. Ruggles and Van Riper (2021) claim that the US Census Bureau did not test whether identifying individuals through their reconstruction attack is more effective than a random guessing. Consider an analogy of clinical trials, where the experiment must have a control group to confirm whether people get better or not after a treatment. The authors describe the US Census Bureau's reconstruction attack as using just a treatment group without a control group for comparison. Some people in the treatment group would get better regardless of whether they received a treatment, and some people could be identified regardless of whether they were included in the reconstruction attack.

FIGURE 3
2010 Census Reconstruction Attack



Source: Authors' illustration.

Data Privacy Definitions and Terminology

The debate over the 2010 Census reidentification attack raises the question of what a realistic data privacy threat is. If you asked this question to a dozen different people, you would likely receive a dozen different responses. This is because data privacy is a broad topic that includes data security, encryption, access to data, and more.

In the context of the census data products, our explainer focuses on applying data privacy and confidentiality methods that provide privacy-preserving access to sensitive data. Although this area of data privacy is very important, especially within the federal statistical system, a smaller share of people know about it. Therefore, we need to cover the many definitions and terminologies that are widely used in the data privacy and confidentiality field before discussing how the Census Bureau implemented the 2020 DAS. We outline several definitions and terminology to keep discussions consistent and avoid confusion, because the data privacy and confidentiality field often has conflicting terms, or several terms are used to represent the same concept. We will also refer to “data privacy and confidentiality” as “data privacy” to be concise, but as stated, outside of this context, the phrase “data privacy” has many meanings.

DATA PRIVACY AND CONFIDENTIALITY TERMS

Although data privacy and data confidentiality are certainly related, they are different, and both play a role in limiting statistical disclosure risk.

Data privacy: the ability “...to determine what information about ourselves we will share with others” (Fellegi 1972).

Data confidentiality: “the agreement, explicit or implicit, between data subject and data collector regarding the extent to which access by others to personal information is allowed” (Fienberg and Jin 2018).

Statistical disclosure control or limitation: statistical approaches to ensure data confidentiality as a means of maintaining privacy.

As we learned in Bowen (2021), there is a necessary balance between data privacy and data utility (or usefulness). This tension is often referred to in the data privacy literature as the “privacy-utility trade-off.”

Data utility, quality, accuracy, or usefulness: how practically useful or accurate to the data are for research and analysis purposes.

Original data: the uncleaned, unprotected version of the data, such as the raw census microdata, which are never publicly released.

Confidential data: the cleaned version (meaning edited for inaccuracies or inconsistencies) of the data; often referred to as the gold standard or actual data for analysis. For example, the Census Edited File that is the final confidential data for the 2020 Census. This dataset is never publicly released but may be made available to others who are sworn to protect confidentiality and who are provided access in a secure environment, such as a Federal Statistical Research Data Center.

Public data: the publicly released version of the confidential data, such as the US Census Bureau’s public tables and datasets.

DATA PRIVACY AND CONFIDENTIALITY COMMUNITY

The data privacy community or ecosystem encompasses a wide range of stakeholders:

Data users: individuals who consume the data, such as analysts, researchers, planners, and decisionmakers.

Data privacy experts or researchers: individuals who specialize in developing data privacy and confidentiality methods.

Data curators, maintainers, or stewards: individuals who own the data and are responsible for its safekeeping.

Data intruders, attackers, or adversaries: individuals who try to gather sensitive information from the confidential data.

Data Privacy Methodology Workflow

Given the importance of the 2020 Census and other data products, how does the US Census Bureau and other data curators provide data users information from these confidential data? Generally, data users obtain the information in two ways:

1. Direct access to the confidential data if they are trusted users (e.g., obtaining Special Sworn Status to use the Federal Statistical Research Data Centers).
2. Access to public data or statistics, such as public microdata and summary tables, that are produced by data curators and modified to protect confidentiality.

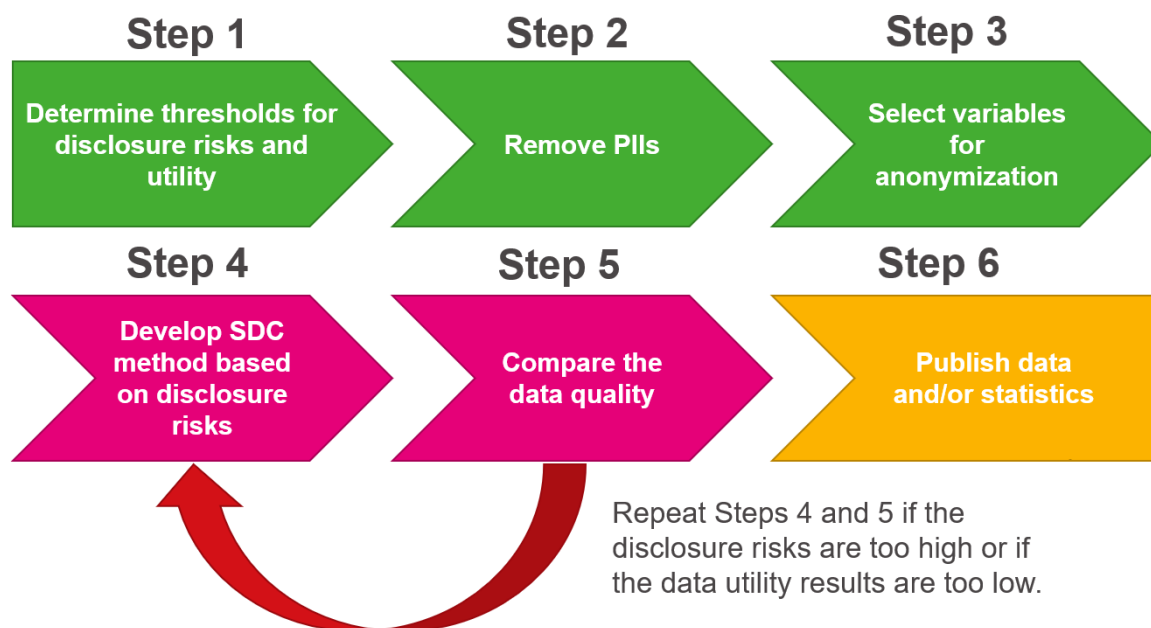
The latter is how most data users gain access to information from confidential data and is the focus of this explainer. To create public data or statistics, data curators rely on statistical disclosure control (SDC) methods to preserve data confidentiality. The process of releasing this information publicly often involves the steps shown in figure 4.

We see that step 1 requires the data curator to determine the acceptable thresholds of disclosure risk and utility. For the disclosure risks, the thresholds are frequently determined by law, such as Title 13 of the US Code³ to “provide strong protection for the information [that the Census] collect[s] from individuals and businesses.” In this example, the Data Stewardship Executive Policy Committee “... serves as the focal point for decision-making and communication on policy issues related to privacy, security, confidentiality and administrative records”⁴ for the Census Bureau, including the interpretation of Title 13. Similar groups exist within various other federal agencies that make these decisions.

For the data utility, data curators often (and should) consult data users and establish data quality metrics based on how the data users will analyze the data. The 2020 Census data, for example, are used to determine boundaries of legislative districts. Thus, the Census Bureau produced and published several relevant metrics to ensure data quality.

FIGURE 4

Data Privacy Workflow



Source: Authors' illustration.

Note: PII = personally identifiable information.

In step 2, the data curator must remove any personally identifiable information that is unnecessary for the public data or statistics release, such as names or Social Security numbers. In step 3, the data curator must then identify what features of the data should be altered with SDC methods. Sometimes the data curator must decide which features of the data are high, medium, and low priority for preserving the information. This helps the data curator and the privacy researcher determine to what extent certain parts of the data should be altered to help balance the privacy-utility trade-off.

Steps 4 and 5 are the hardest parts of the workflow. We can imagine in the extreme case, if data users want full data utility, then the data curator would release the confidential data unaltered. On the other hand, to achieve only privacy, the data curator would never release the confidential data. In this example, we see how data privacy and data utility naturally oppose one another. This is why steps 4

and 5 become an iterative process in developing an SDC method, and making one part of the data more useful reduces the data privacy guarantee (and vice versa).

For step 4, the privacy researcher must carefully determine how much to alter, change, or sanitize the confidential information using a particular SDC method. Within the data privacy field, the terminology defining each step of the SDC process can be inconsistent.

We break down SDC methods into three steps (but note that some SDC methods do not have the last step).

1. **Preprocessing:** prioritizing which statistics or information to preserve (i.e., could be considered step 3 in the workflow).
2. **Privacy:** applying a sanitizer to the desired statistic or information (i.e., altering the statistic).
3. **Postprocessing:** ensuring the results of the statistic or information are consistent with realistic constraints (e.g., population counts should not be negative).

Note that in the privacy step, sanitizer is used with a lowercase “s”; some SDC methods use capitalized “Sanitizer” as part of their formal name.

The privacy step requires the privacy expert to know the type of disclosure risk to protect the confidential data against. Traditionally, there are generally three types of disclosure risk:

1. **Identity disclosure risk** occurs if the data intruder associates a known individual with a public data record (e.g., a record linkage attack or when a data adversary combines one or more external data sources to identify individuals in the public data).
2. **Attribute disclosure risk** occurs if the data intruder determines new characteristics (or attributes) of an individual based on the information available through public data or statistics (e.g., if a dataset shows that all people age 50 or older in a city are on Medicaid, then the data adversary knows that any person in that city above age 50 is on Medicaid).
3. **Inferential disclosure risk** occurs if the data intruder predicts the value of some characteristic from an individual more accurately with the public data or statistic than would otherwise have been possible (e.g., if a public homeownership dataset reports a high correlation between the purchase price of a home and family income, a data adversary could infer another person’s income based on purchase price listed on Redfin or Zillow).

Note that some federal statistical agencies are not concerned about inferential disclosure risk for two reasons. First, one of the main reasons for releasing public data is to allow data users to infer and identify relationships among various attributes. If an agency considered inferential disclosure risk, then few datasets and statistics would be released. Second, inferential disclosure risk is predicting aggregated attributes instead of individual, which means the data intruder would poorly predict individual values. However, some other federal statistical agencies assess inferential disclosure risk when there are high statistical relationships between certain attributes and an adversary can create an extremely accurate model (Federal Committee on Statistical Methodology 2005).

After developing an SDC method that protects against certain types of disclosure risks, the data curator and privacy research must assess data utility. Broadly, there are two ways to measure it:

1. **General utility or global utility:** measures the univariate and multivariate distributional similarity between the confidential data and the public data (e.g., sample means, sample variances, and the variance-covariance matrix).
2. **Specific utility or outcome-specific utility:** measures the similarity of results for a specific analysis (or analyses) of the confidential and public data (e.g., comparing the coefficients in regression models).

Some in the data privacy community argue that data utility does not necessarily mean accuracy and are actively exploring other measures that best convey data quality, consistency, and accuracy.

Once the data curator and privacy expert find the right balance between privacy and utility, they may proceed to publishing the data or statistics in step 6. However, the data curators and privacy researchers *should* consult the data user community to determine what kind of published data and/or statistics to release and ensure that information are fit for use, which is why achieving the balance is so difficult. Here, we list examples of possible data products that a data curator could release after applying SDC methods, roughly from most to least detailed:

- microdata (e.g., public use microdata series or PUMS)
- summary tables (e.g., American Community Survey tables)
- summary statistics (e.g., multiple statistics on income in a state)
- single statistics (e.g., maximum age in a county)

Curators could release one of these products after applying an SDC method, or they could release them “on demand” to answer different questions using the data. Questions asked of the data are

referred to in computer science terminology as *queries*, which are statistics. We will therefore refer to them as statistics throughout the explainer to avoid confusion. Note that when reading more technical data privacy papers, these questions are more commonly referred to as queries.

Introduction to Formal Privacy

We now better understand the challenges the US Census Bureau faces when creating the DAS to protect against a privacy threat. In particular, the Census Bureau, as the data curator, must make assumptions or judgement calls on how a data intruder would obtain sensitive information from public data or statistics. They must ask themselves the following questions: How much disclosure risk is too much, and what type? When evaluating disclosure risk, what assumptions can be made about how the data intruder will approach the data? What about the resources the intruder has access to? Do these assumptions hold in the context of the specific, real-world application?

These questions and many others motivated the creation of a concept known as formal privacy, which provides a mathematical bound on the disclosure risk for any statistic applied to the confidential data. Although methods developed within the formal privacy framework are considered SDC methods, data privacy researchers often separate formal privacy from other SDC methods. We will refer to the SDC methods and disclosure risk measures *not* developed under formal privacy as *traditional SDC methods* and *traditional disclosure risk definitions*.

In this part of the explainer, we will cover a high-level overview of formal privacy, differential privacy, and differentially private mechanisms. This summary will involve some mathematical intuition and present some mathematical equations to prepare the reader for the next section on how the 2020 DAS works. For readers interested in a more technical review of similar content, see work by Bowen and Garfinkel (2021).

Formal Privacy

We begin with what makes a privacy definition formally private. Although the privacy community has not fully agreed on a common definition, formal privacy is defined by the Census Bureau⁵ as a subset of SDC methods that give “formal and quantifiable guarantees on inference disclosure risk and known algorithmic mechanisms for releasing data that satisfy these guarantees.”

Traits of formally private mechanisms include the following:

- Ability to quantify and adjust the privacy-utility trade-off, typically through parameters.
- Ability to rigorously and mathematically prove the maximum privacy-loss that can result from the release of information (Bowen and Garfinkel 2021).
- Formal privacy definitions also allow one to “compose” multiple statistics. In other words, a data curator can compute the total privacy-loss from multiple individual information releases (Bowen and Garfinkel 2021).

Simply put, the main difference between traditional SDC methods and formally private methods is the ability to account for each piece of information being “leaked” from the confidential data. We can think of traditional SDC methods as akin to someone charging a limitless credit card; formally private methods are akin to someone charging to a debit card with a set budget. In both scenarios, there is a running bill, but only one requires constantly checking the balance. We can easily imagine that not tracking that bill is the equivalent of releasing too many statistics with enough accuracy, which could compromise the confidential data (Bowen and Garfinkel 2021). Although in both traditional and formal privacy settings data curators must limit the type and number of questions asked of the data, they are faced with “tracking the bill” under a formal privacy framework.

Differential Privacy and Other Formally Private Definitions

We now understand the key differences between formally private definitions and traditional disclosure risk definitions. But what are some formally private definitions? The most well-known formal privacy definition is differential privacy (DP), first introduced by Dwork and colleagues (2006). We emphasize that DP is a strict mathematical definition that a method must satisfy (or meet the mathematical conditions) to be considered differentially private, not a statement or description of the data itself.

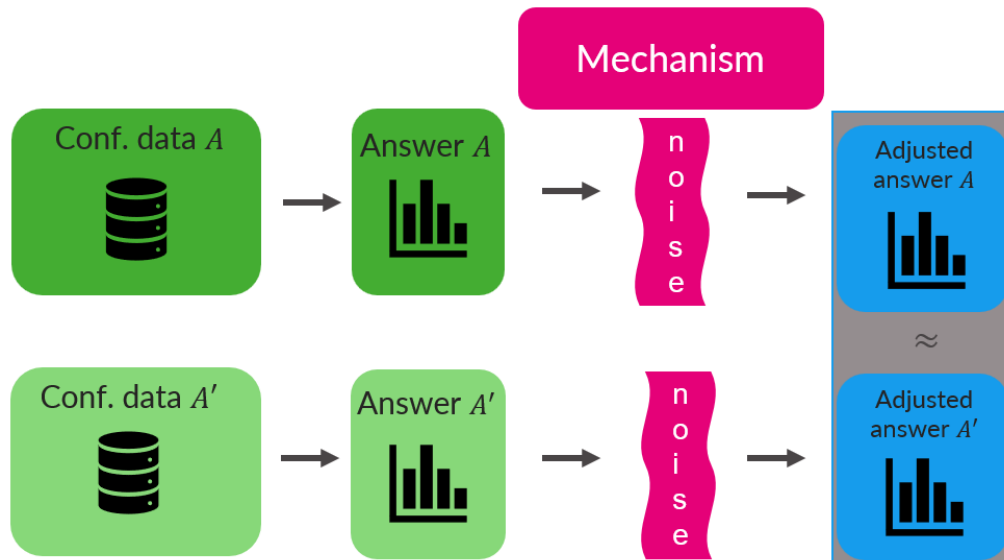
Simply put, DP does not make assumptions about how a data intruder will attack the data and the amount of external information or computing power an actor has access to, now or in the future.⁶ Instead, DP assumes the worst-case scenario to provide a strong privacy guarantee:

- The data intruder has information on every observation except one
- The intruder has unlimited computational power
- The missing observation is the most extreme possible observation (or an extreme outlier) that could alter the statistic

Mathematically, DP states that the log of the ratio of the probability that any individual observation was or was not in the data that generated the output is bounded by the value of ϵ , where $\epsilon > 0$. This means that if we use a privacy-loss budget of 1, then that ratio converts to $e^1 \approx 2.72$ and represents the bound on the probability that the above assumptions fail because we are releasing information. Informally, DP guarantees the output of a differentially private mechanism will be roughly the same whether the individual observation is in the data or not (figure 5).

FIGURE 5

Visual Representation of Differential Privacy



Source: Authors' illustration.

Privacy researchers have also developed other formally private definitions and consider these alternative definitions as relaxations of the DP definition because they “ease up” on the strong privacy guarantee that DP provides (i.e., the worst-case scenario listed earlier). We briefly cover the two most popular ones at a high level with some math to explain the relationships among the three formal privacy definitions. Note that privacy experts often refer to the original definition of DP as **pure-DP** or ϵ -DP given the many DP relaxations.

A popular DP relaxation is **approximate-DP** or (ϵ, δ) -DP, which has similar levels of privacy guarantee as ϵ -DP, but with a small probability (i.e., $\delta \in [0,1]$) that the DP ratio does not hold (Dwork et al. 2006; Dwork and Roth 2014). In other words, if $\delta = 10^{-3}$, then there will be a 0.001 percent chance that a (ϵ, δ) -DP method will release the confidential value.

Dwork and Rothblum (2016) created **concentrated DP** with the purpose of reducing the privacy loss over multiple computations. Bun and Steinke (2016) later improved the definition and called it **zero-concentrated-DP** or ρ -zCDP. The authors also proved that if a method satisfies ρ -zCDP, then it satisfies (ϵ, δ) -DP, where $\epsilon = \rho + 2\sqrt{\rho \log(1/\delta)}$ for any $\delta > 0$.

There is also a direct relationship between ρ -zCDP and ϵ -DP, where if a method satisfies ϵ -DP, then the method satisfies ρ -zCDP, where $\rho = 1/2 \epsilon^2$. These conversions will be particularly relevant for the section on the 2020 Decennial Census. The US Census Bureau initially announced they would be using pure-DP but changed to approximate-DP and zero-concentrated-DP to reduce the amount of noise to the 2020 Census data products. Unfortunately, this shift in formally private definitions created confusion among the census data user community. We discuss this communication issue further in our takeaways section.

Privacy-Loss Budget

In contrast to the traditional disclosure risk types, how does DP and the other DP relaxations compute or account for the privacy-loss or disclosure risk when releasing information? These definitions use the concept of a privacy-loss budget, typically represented mathematically as ϵ . Although there are two other privacy parameters (δ and ρ), we will focus on ϵ for simplicity and ease of conceptual explanations until the next section on the 2020 Decennial Census. The privacy-loss budget bounds the disclosure risk associated with releasing data or statistics. It can be thought of as a knob that adjusts the trade-off between data privacy and utility. Some things to keep in mind about the privacy-loss budget are as follows:

- The data curator must decide the privacy-loss budget (i.e., the total amount of ϵ) before the release of any data or statistic. Like a real budget, when privacy-loss budget is exhausted, no more information from the confidential data is released.
- A larger value of ϵ increases the maximum disclosure risk (i.e., the upper bound of the disclosure risk) associated with a given release of information. Simply put,
 - » larger ϵ = less noise potentially added to a statistic = more accuracy, but less privacy, and
 - » smaller ϵ = more noise potentially added to a statistic = less accuracy, but more privacy.

For a visual representation, figure 6 shows the image becoming clearer or more accurate as ϵ increases.

FIGURE 6

Illustration of Increased ϵ Results in a Clearer Image



As ϵ increases, the image becomes clearer (more accuracy).

Source: Authors' illustration. Original image is *Flowers in a Vase* by Philip van Kouwenbergh and is in the public domain.

Earlier in the explainer, we stated that the trade-off between data privacy and utility could be explained in the extreme case of releasing the confidential data (i.e., maximum utility) or not releasing the confidential data (i.e., maximum privacy). In the DP framework, we can explain the scenario with ϵ . When $\epsilon \rightarrow \infty$, we obtain perfect utility, but no privacy. When $\epsilon \rightarrow 0$, we obtain perfect privacy, but no utility. In other words, as with traditional SDC methods, the privacy-loss budget cannot eliminate all risk. When the data curator adjusts the privacy-loss budget, they are adjusting the strength of the privacy guarantee provided by DP.

Additionally, the data curator must also determine how to distribute the privacy-loss budget over the many possible public datasets and statistics. For instance, we can imagine the privacy-loss budget as a set monthly budget for household expenses (e.g., housing, groceries, utilities, and transportation). Some people might want to equally allocate their funds to each expense, whereas others might think that groceries should cost more than transportation, but not more than housing. Likewise, some data curators might prioritize releasing multiple statistics, while others might allocate the full privacy budget to allow the release of microdata. In other words, data curators must consider how they will allocate the privacy-loss budget for each individual release of information while maintaining the overall budget for the system.

Given that data curators could distribute their overall privacy-loss budget across several public datasets or statistics in many ways, some would want guidance on allocating the privacy-loss budget. Although Dwork and colleagues (2006) proposed DP over 16 years ago, setting an appropriate privacy-loss budget is still an open question. All members of the data privacy community should be involved in this discussion, but many advise that the choice is ultimately up to public policymakers.

However, although policymakers are the most equipped to understand the consequences of the privacy-loss, they are likely the least equipped to understand what ϵ means.

For instance, public policymakers would probably not know that DP defines ϵ as logarithmic (e.g., $\epsilon = 1, 2$, and 3 becomes approximately 2.7 , 7.34 , and 20.01 , respectively) and an inequality (i.e., ϵ represents the upper bound for disclosure risk, which means the actual disclosure risk could be lower).

What do privacy experts suggest? Early privacy research considered ϵ being 1 or 2 as the maximum privacy-loss budget for releasing public data or statistics. However, much larger values are appearing in more recent applications. For example, the US Census Bureau applied their new differentially private method, called the TopDown Algorithm, to the 1940 and 2010 Census and published the resulting data as a demonstration. Data users could then compare the demonstration data against the unaltered 1940 Census data⁷ and the original 2010 Census data release. Table 1 shows the values that the Census Bureau used for the demonstration data throughout multiple releases.

TABLE 1

Demonstration File ϵ and Ratio Values

The ϵ values that the US Census Bureau used for the demonstration persons file

ϵ	Ratio
0.25	1.28
0.50	1.65
0.75	2.12
1.00	2.72
2.00	7.39
4.00	54.60
4.50	90.02
6.00	403.43
8.00	2,981.96
10.30	29,733.62

Source: Authors.

In June 2021, the Census Bureau⁸ committed to ϵ of 17.14 for the persons file and 2.47 for the housing file, which are part of the redistricting data. Although the difference between the two numbers is 14.67 , the privacy-loss budget for the persons file converts to $27,784,809$ (i.e., $e^{17.14}$), whereas the housing file privacy-loss budget becomes 11.82 (i.e., $e^{2.47}$). In other words, the persons file has a privacy-loss budget that is roughly 2.35 million times larger than the housing file.

What about other real-world applications? Rogers and colleagues (2020) compared several industry applications of DP. One of the largest values of ϵ is 769 (a monthly budget), which converts to a value rapidly approaching infinity. In recent years, there have been more differentially private applications that provide further context on setting an appropriate privacy-loss budget.

At this point, you're likely still confused about what all these privacy budget values mean for public policy decisions. The simple answer is the community still doesn't know, other than that higher privacy-loss budget values mean more information is leaked from the confidential data. We need more applications and conversations with data users and privacy researchers to best inform policymaking decisions on the best balance between data privacy and utility. We discuss this further at the end of this report.

Global Sensitivity

In addition to the privacy-loss budget, most differentially private methods rely on the concept called global sensitivity, which describes how resistant the differentially private sanitizer is to the presence of outliers (Bowen and Garfinkel 2021). We can think of the global sensitivity as another value that helps determine how much noise is needed to protect the released data or statistic, because some information is more sensitive than other information to outliers.

TABLE 2

Demonstration File ϵ and Ratio Values

A fictitious socioeconomic dataset with participants' names, age, and wealth, along with Elon Musk's information

Person	Age	Wealth
Alex	28	\$51,489
Andrea	26	\$36,072
Bob	62	\$85,356
Beth Ann	58	\$77,226
Daniel	17	\$623
Donna	34	\$41,543
Edward	45	\$115,879
Elizabeth	53	\$99,253
Elon	51	\$263.6 billion
Nikola	86	\$0

Source: Authors' hypothetical and "The Real-Time Billionaires List," *Forbes*, accessed July 15, 2022, <https://www.forbes.com/real-time-billionaires/#2cfd93953d78>.

We borrow the example explained in Bowen (2021) to help explain this concept along with table 5 to provide example values. Imagine the data we want to protect contains socioeconomic information and the statistic we want answered is, “What is the median wealth of a group of individuals?” Under DP, we must consider the change of the most extreme possible record that could exist in any given data that has demographic and financial information. For our example, that person is Elon Musk, who was the wealthiest person in the world as of the publication of this explainer (table 2). If Musk is present or absent in the data, the median wealth should not change too much. For instance, from the values in table 2, the median wealth is \$64,357.50 with Musk in the data. The median wealth becomes \$51,489 without Musk. This means we can provide a more accurate answer by applying less alterations to the median income statistic, because it is less sensitive to (or more robust against) the extreme outlier, Musk. Now consider the question, “What is the average wealth of a group of individuals?” Unlike the previous statistic, the answer would significantly change if Musk were present or absent from the data. From our values in table 2, the average wealth would be \$26,360,050,744 with Musk and \$56,382.33 without Musk. To protect the extreme case, a differentially private method would need to provide a significantly less accurate answer by altering the statistic more.

Gaussian Mechanism

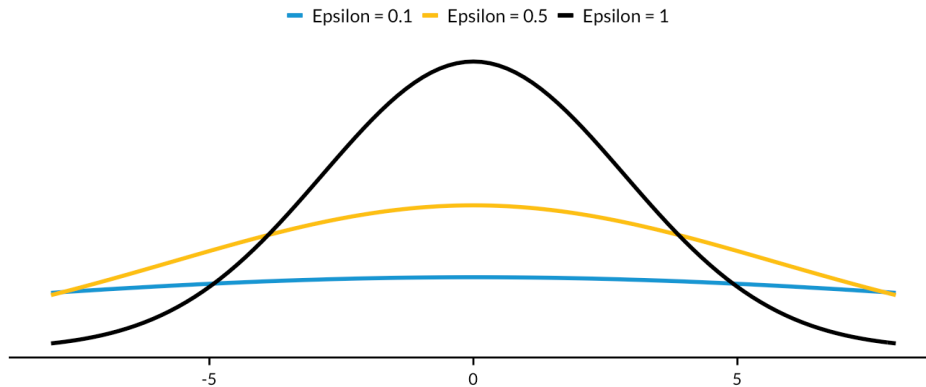
So what is an example of a differentially private sanitizer (or a differentially private method that alters a statistic) that uses the privacy-loss budget and global sensitivity? We will walk through how to apply a popular differentially private sanitizer, called the Gaussian mechanism. This sanitizer adds noise to a statistic by drawing values from a Gaussian distribution (i.e., a normal or bell-curve distribution). This distribution is centered at zero and its variability (i.e., how wide or narrow the distribution is) changes based on the privacy-loss parameters, ϵ and δ , and the global sensitivity of the target statistics. Having the distribution centered at zero means there is a higher probability of adding very little or no noise to the confidential data statistics, which is ideal for data utility.

How do the privacy parameters and the global sensitivity of a statistic affect the noise variability? Suppose the statistic we want to release is a count, which has a global sensitivity of 1. If we want to have a higher probability of adding very little noise to our count statistic (more accuracy), then we want to increase the privacy parameter values (ϵ and/or δ). If we want to add more noise to our count statistic (more privacy), then we want to decrease the privacy parameter values. In figure 7, we show how the variability of the Gaussian distribution increases (i.e., the curve flattens out more) when we decrease ϵ . This translates to having a higher probability of adding more noise to our count statistic. In

our figure, we don't change δ , but changing δ will also affect the variability of the Gaussian distribution similarly.

FIGURE 7

Gaussian Distribution with Different Values of ϵ with Global Sensitivity = 1 and $\delta = 10^{-7}$

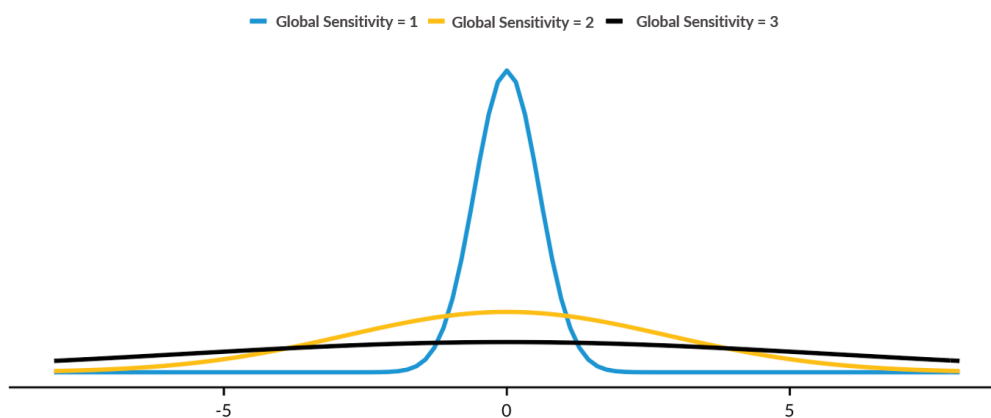


Source: Authors' illustration.

Suppose now we have three different statistics that have global sensitivities of 1, 2, and 3, respectively. We also want to equally allocate the same amount of privacy-loss budget to each statistic (e.g., $\epsilon = 1$ and $\delta = 10^{-7}$). When a statistic has a higher global sensitivity (i.e., less robust to outliers), we will need to add more noise to protect that statistic for the same privacy-loss budget. Figure 8 illustrates how the variability of the Gaussian distribution increases when the global sensitivity of a statistic is large for a set privacy-loss budget.

FIGURE 8

Gaussian Distribution with Different Values of Global Sensitivity with $\epsilon = 1$ and $\delta = 10^{-7}$



Source: Authors' illustration.

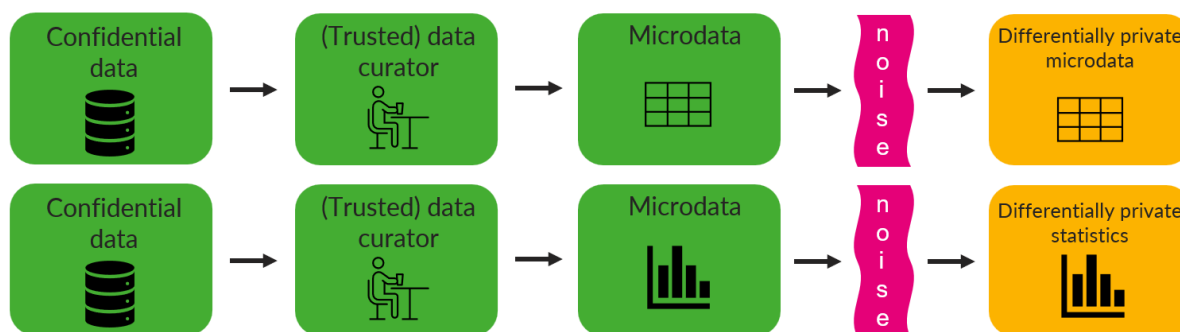
Note that for the 2020 Decennial Census, the Census Bureau used a similar sanitizer called the discrete Gaussian distribution, which adds discrete values to a statistic instead of continuous values to ensure the noise added to count statistics resulted in integer values.

Models of Differential Privacy

Now that we understand how a formally private sanitizer works, how do privacy researchers implement them in practice? Over the years, roughly two models or frameworks for applying formally or differentially private sanitizers and methods have arisen. Bowen and Garfinkel (2021) present two models (trusted curator and local), but we will discuss differentially private synthetic microdata as a third model here because it encompasses a large part of the literature.

FIGURE 9

Trusted Curator Model Illustration



Source: Authors' illustration.

TRUSTED CURATOR MODEL

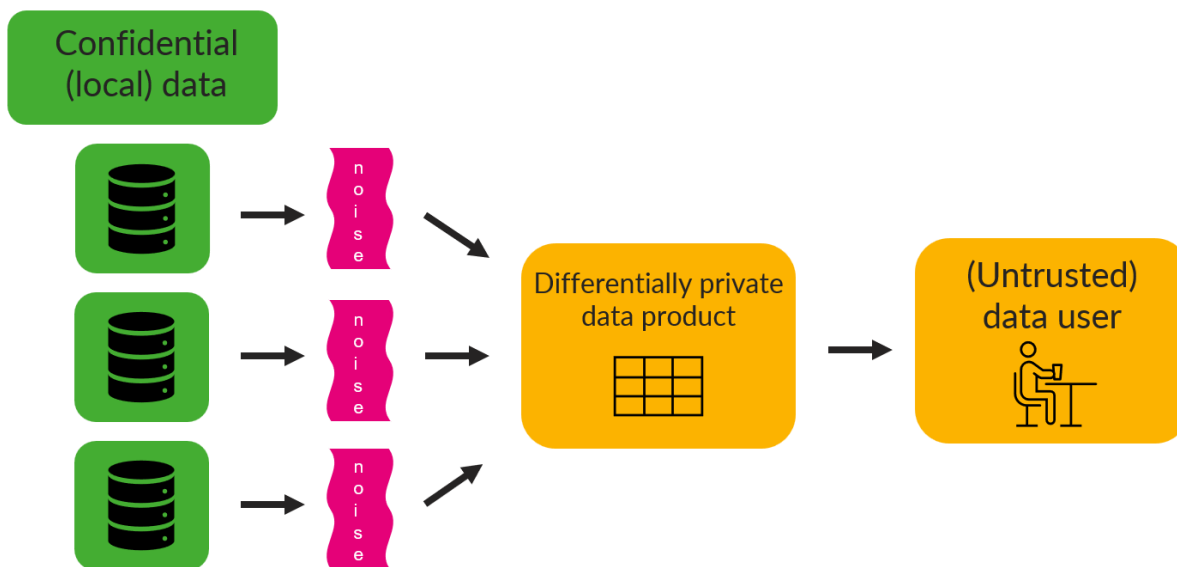
In the trusted curator model, a centralized data curator receives confidential data, creates the data products, applies the differentially private method, and releases the results. This means that if the data curator has a set privacy-loss budget, then the curator must stop releasing information when the budget is reached. For example, Uber created a differentially private system that allowed their analysts within the company to evaluate customer experience through targeted requests without seeing confidential individual trip or rider details (Johnson, Near, and Song 2018). In this situation, one part of Uber is the data curator, and the other is the data user. Figure 9 shows how the trusted curator model works for generating differentially private microdata and statistics.

DIFFERENTIALLY PRIVATE SYNTHETIC MICRODATA

Differentially private synthetic microdata is DP applied to a statistical model of the confidential data (shown in the top part of figure 9). The privacy-protected model is then used to create individual records for release, like synthetic data generation. Although this model is considered a type of the trusted curator model, it is one of the most common applications of DP. The differentially private synthetic microdata is popular because once the microdata is generated it can be distributed or repeatedly analyzed without adding to overall privacy loss. However, creating accurate differentially private synthetic microdata is very difficult, particularly for data with more than a few columns.

FIGURE 10

Local Differential Privacy Model Illustration



Source: Authors' illustration.

LOCAL DIFFERENTIAL PRIVACY

Local DP allows the participant to add DP locally to their own data before sending that information to the curator (figure 10). Essentially, the local model sanitizes the data upon collection. This framework trusts no one, not even the data curator. The general idea is that rather than a global or overall privacy-loss budget being applied to the entire confidential data, each data participant or data collection point receives its own privacy budget. However, this model substantially adds more noise to locally noised microdata than to data products created by a trusted curator (Bowen and Garfinkel 2021).

Introduction to 2020 Disclosure Avoidance System

Although the formal privacy framework avoids the ad hoc nature of traditional SDC methods, we learned some of the difficulties in implementing formally private methods. For instance, a data curator must answer new questions: Which formally private definition should be used? What is an appropriate value for the privacy-loss budget? What are the utility measures to ensure data quality and usability? Once a model of formal privacy is identified, how is a formally private method created? And what communication materials are needed to explain the formally private method to data users, such as public policymakers?

In this section, we walk through how the US Census Bureau tackles these questions when developing the 2020 DAS.

Privacy and Utility Measures

After the 2010 reconstruction attack, the US Census Bureau started to explore various formally private definitions. The Census Bureau began with pure-DP and eventually settled on zero-concentrated differential privacy (ρ -zCDP, a DP relaxation), which can be converted to (ϵ, δ) -DP. The US Census Bureau changed to ρ -zCDP because of how multiple Gaussian distributions compose. At a high level, the Census Bureau can reduce the overall amount of noise added to the data because multiple Gaussian distributions create a Gaussian distribution. Also, adding noise from a Gaussian distribution does not satisfy pure-DP.

In terms of communication, the US Census Bureau continued to report the privacy parameter values as ϵ and δ . The Census Bureau likely made this decision to avoid confusion after first using pure-DP for many of the demonstration data (Abowd and colleagues 2022; US Census Bureau 2021).

When selecting the privacy-loss budget, the US Census Bureau tested out several values (table 1) against various utility metrics. However, there are thousands of data use applications, ranging from allocating congressional seats to determining the number of restaurant permits to issue. The Census Bureau therefore implemented several utility metrics, such as, the following:

- General utility
 - » mean absolute error
 - » mean numeric error
 - » root mean squared error

- » mean absolute percent error
- » coefficient of variation
- » total absolute error of shares
- Outcome specific
 - » decisions on redistricting voting lines or school districts
 - » total absolute error of shares metric by county within each state as a share of that state, by incorporated place as a share of that state, and by minor civil divisions as a share of that state

A full list of the metrics can be found in the “Revised Data Metrics for 2020 Disclosure Avoidance” document.⁹

The US Census Bureau used these metrics and more when determining the privacy-loss budget. During testing, they decided to split the redistricting file into two parts, a persons file and a housing units file. For the persons file, the Census Bureau used $\delta = 10^{-10}$ and $\rho = 2.56$. These values convert to $\epsilon = 17.14$ (i.e., $\rho + 2\sqrt{\rho \log(1/\delta)} = \epsilon$). For the housing units file, the US Census Bureau used $\delta = 10^{-10}$ and $\rho = 0.07$, which converts to $\epsilon = 2.47$. We will cover how ρ is allocated to each target statistic in the next subsection.

Statistical Disclosure Control Method

Similar to traditional SDC methods, we can break down formally private methods into the preprocessing, privacy, and postprocessing steps. Here, we will focus on the persons file of the 2020 DAS, but general steps apply to the housing units file as well. For more technical details of the 2020 DAS, see work by Abowd and colleagues (2022).

PREPROCESSING STEP

The Census Bureau first had to calculate the crosstabulation (or marginal counts) of all variables for each geographic level (from state to census blocks) from the confidential data or Census Edited File. The US Census Bureau (2021) lists the statistics of interest as follows:

- 1 total count
- 63 race, 2 ethnicity (Hispanic or Latino/Not Hispanic or Latino), 2 voting age (under 18 years/18 years and older)

- 3 institutional versus noninstitutional group quarter types
- 1 residential and 7 possible group quarter types for a total of 8 (e.g., dorms and prisons)
- 126 possible combinations of race and ethnicity
- 126 possible combinations of race and voting age
- 4 possible combinations of ethnicity and voting age
- 252 possible combinations of race, ethnicity, and voting age
- 2,016 possible combinations of race, ethnicity, and voting age at each residential and group quarter type

For example, one of the possible statistics is the number of Asian Americans alone who are under 18 in a residential housing unit at the census tract level.

PRIVACY STEP

Essentially, the US Census Bureau applies the Gaussian mechanism to all the possible combinations listed earlier unless that combination has no observations (i.e., treat as a structural zero) at each geographic level. However, the confusing part is how the Census Bureau allocates the privacy-loss budget for the different statistics at each geographic level.

TABLE 2

Privacy-Loss Budget: People

The ρ allocation that the US Census Bureau used for the person file of the 2020 Census

Geographic level	Rho allocation
United States	$104/4,099 \approx 0.025$
State	$1,440/4,099 \approx 0.351$
County	$447/4,099 \approx 0.109$
Tract	$687/4,099 \approx 0.168$
Optimized block group	$1,256/4,099 \approx 0.306$
Block	$165/4,099 \approx 0.040$

Source: US Census Bureau (2021).

We recreate the tables from the US Census Bureau (2021) to help explain how the US Census Bureau allocated the privacy-loss budget (ρ in this case). First, the Census Bureau allocated a proportion of the total or global ρ to geographic levels. Table 2 lists the ρ allocations, which sum to 4,099/4,099.

Mathematically, this becomes

$$\begin{aligned}\rho_{total} &= \rho_{US} + \rho_{State} + \rho_{County} + \rho_{Tract} + \rho_{Block\ Group} + \rho_{Block} \\ &= \frac{104}{4,099} \times 2.56 + \frac{1,440}{4,099} \times 2.56 + \frac{447}{4,099} \times 2.56 + \frac{687}{4,099} \times 2.56 + \frac{1,256}{4,099} \times 2.56 + \frac{165}{4,099} \times 2.56 = 2.56\end{aligned}$$

After allocating ρ to each geographic level, US Census Bureau divided that geographic level's privacy-loss budget further for each of the race, ethnicity, and voting age statistics. For the sake of simplicity, we focus on the “Block” column in the “Per Query Privacy-Loss Budget: People” in US Census Bureau (2021) and recreate this information in this report in table 3. From table 2, we know that $\rho_{Block} = 165/4,099 \times 2.56 \approx 0.10$. This means $\rho_{Block} \approx 0.10$ is allocated proportionally to the 11 statistics in table 3.

What does that mean for a specific statistic? Suppose we wanted to sanitize the block total population. Table 3 lists that $5/4,097$ of the 0.10 privacy-loss budget is allocated to that statistic, which comes out to roughly 0.0001. Because the Census Bureau used a discrete Gaussian distribution, the spread parameter (this is not the variance) is $\sigma^2 = 1/\rho = 9939$. This means if we made 100 draws from this distribution, then about 90 of them will be between -164 and 164. In other words, 90 percent of the time, the potential noise added to the block population statistic will be between -164 and 164.

TABLE 3

Per Query Privacy-Loss Budget: People

The ρ allocation that the US Census Bureau used for the personal file of the 2020 Census

Query	Rho allocation
TOTAL	$5/4,097 \approx 0.001$
CENRACE	$9/4,097 \approx 0.002$
HISPANIC	$5/4,097 \approx 0.001$
VOTINGAGE	$5/4,097 \approx 0.001$
HHINSTLEVELS ¹⁰	$5/4,097 \approx 0.001$
HHGQ	$5/4,097 \approx 0.001$
HISPANIC $\times$ CENRACE	$21/4,097 \approx 0.005$
VOTINGAGE $\times$ CENRACE	$21/4,097 \approx 0.005$
VOTINGAGE $\times$ HISPANIC	$5/4,097 \approx 0.001$
VOTINGAGE $\times$ HISPANIC $\times$ CENRACE	$71/4,097 \approx 0.017$
HHGQ $\times$ VOTINGAGE $\times$ HISPANIC $\times$ CENRACE	$3,945/4,097 \approx 0.963$

Source: US Census Bureau (2021).

However, note that this potential noise is for changing one statistic and the redistricting file has thousands of statistics. Further, the 2020 DAS truncates or reduces the potential noise from all these statistics further using mathematical properties when composing multiple Gaussian distributions. Simply put, the -164 and 164 range is reduced to a much smaller one. How much has yet to be reported. For more technical details, see work by Abowd and colleagues (2022).

POSTPROCESSING STEP

After adding noise to each statistic, the TopDown Algorithm is the procedure that enforces the invariant statistics (i.e., no change to the statistics) and constraints (e.g., the population counts in all counties in the state should equal the state population) listed below. Note that many refer to the 2020 DAS as the TopDown Algorithm even though the TopDown Algorithm only encompasses the postprocessing step.

- Invariant statistics:
 - » Total population in each state, the District of Columbia, and Puerto Rico
 - » Total number of housing units within each census block
 - » Number of group quarter facilities by type within each census block
- Constraints:
 - » Counts must be integers
 - » Sums of rows and column margins must sum to the total populations
 - » Counts must be consistent within tables, across tables, and across geographies
 - » If there are zero housing units and zero group quarters at a geographic level, then no people may be assigned to that geography
 - » Number of people in a group quarter is equal to or greater than 1
 - » Number of people in a housing unit or group quarter is less than or equal to 99,999
 - » Geographic areas cannot have everyone under the age of 18 except areas with certain group quarter populations (e.g., juvenile detention centers)
 - » Census Edited File constraints (such as, if person two is the “natural child” of person one, then person two cannot be older than person one).

Technical details on how the TopDown Algorithm optimizes these invariants and constraints can be found in work by Abowd and colleagues (2022). Also, note that the full list of edit specifications for each Census Edited File is not public record.

Takeaways and Ongoing Challenges

In this explainer, we recapped the basics of data privacy, learned about formal privacy, and understood how the US Census Bureau implemented the TopDown Algorithm for the 2020 Census. We learned the following:

- Unlike traditional SDC methods, formally private methods quantify and bound the disclosure risk associated with releasing information from the confidential data.
- Formal privacy definitions use the idea of a privacy-loss budget that adjust the amount of maximum disclosure risk (the upper bound of the disclosure risk) associated with releasing information from the confidential data.
 - » larger ϵ = less noise potentially added to a statistic = more accuracy, but less privacy
 - » smaller ϵ = more noise potentially added to a statistic = less accuracy, but more privacy
- How public policymakers set the privacy-loss budget is still an open question.
- Census data users must now answer the question, “How good is good enough?” and provide new utility measures and use cases to the US Census Bureau.

The final two takeaways leave us with three major challenges. First, we do not have clear interpretations of the worst-case privacy-loss for the privacy parameters ϵ , δ , or ρ . As shown in table 1, when ϵ is one or two, the ratio of probabilities is around 2.7 and 7.4, respectively. Those familiar with the DP literature can interpret these ratios because of familiarity. However, when $\epsilon = 17.14$, the ratio of probabilities is 27,784,809—a value far larger than what was typically in the literature before more real-world applications. Furthermore, there is a small probability of 10^{-10} that the ratio does not hold. Most people cannot interpret this privacy budget, including privacy experts.

If privacy researchers cannot interpret the budget, then we are left wondering, “How can policymakers make informed decisions about trade-offs between utility and privacy?” One option is they cannot make an informed decision and select parameters without understanding the bound. The other option is they use ad hoc and post hoc measures of data privacy to interpret the results of the chosen privacy parameters. This latter option results in decisions based on assumptions similar to the traditional SDC methods. In other words, without better privacy-loss parameter interpretations, we revert the formally private methods to the traditional SDC methods.

The second challenge is we need *even more* formally private use cases. Although we have more use cases, part of the reason we do not know reasonable values for various privacy parameters (e.g.,

ϵ , δ , and ρ) is because most formally private research is still largely theoretical. More privacy researchers need to implement formally private methods on real-world applications to fully understand the privacy-utility trade-off under several conditions. For example, privacy experts should explore more small, practical differentially private applications rather than highly complicated, theoretical scenarios to better discern some of the data challenges and how we should address them. The same idea applies to other SDC methods, such as synthetic data, where we do not have enough use cases.

The last challenge is improving and creating data privacy and confidentiality communication and education materials. For example, suppose someone told you that they had data that contained records of individuals, including demographics such as their age, their sex, and their race along with financial information. They want to explore applying machine learning methods to gain unique insights into the data. What resources would you recommend? Now, suppose this person, with the same data, asked you how to apply data privacy and confidentiality methods. Besides this document, would you have any idea what resources to recommend?

Snoke and Bowen (2021) posed these scenarios and stated that “a significantly higher percentage of readers probably will have answers to the questions posed in the first hypothetical scenario than to those in the second, which raises the question of why. Statisticians often use public microdata or tables, or access sensitive data through restricted data centers or agreements. Yet, few develop and implement data privacy and confidentiality methods that enable that access.”

These hypothetical scenarios highlight the lack of well-written and well-designed computational resources. For the latter, not having readily available computational tools will hinder the accessibility for data users to implement SDC methods. They might not have the proper computing environment to run these methods or the technical background (expert knowledge and/or programming skills) to hand-code the methods. Moreover, hand-coding is more prone to errors and might be less efficient.

Very few people in general have the technical knowledge and the coding ability to implement SDC methods. Some propose that we need to teach the next generation of data privacy researchers. However, most higher-education institutions do not offer data privacy courses. If these courses are taught, professors usually teach them at the graduate level in computer science departments, which is not representative of those who depend on and contribute to the field.

What can we do to improve this situation? Hu and Bowen propose the following:¹¹

- Incorporate data privacy and confidentiality into undergraduate curriculum and go beyond the basic introduction by, for example, applying appropriate methods to real data and evaluating their effectiveness.
- Focus on how to translate theory to applications and deployment rather than on theory alone.
- Advocate for more funding for applied research and deployment (i.e., computational tools and educational resources) instead of only on new method development.

These measures alone will not address all the communication and education needs for data privacy and confidentiality, but they are a promising start.

Notes

- ¹ “A Monograph on the Confidentiality and Privacy in the U.S. Census,” accessed August 4, 2022, <https://www.census.gov/history/pdf/ConfidentialityMonograph.pdf>
- ² “The Census Bureau’s Simulated Reconstruction-Abetted Re-identification Attack on the 2010 Census,” US Census Bureau, accessed July 18, 2022, <https://www.census.gov/data/academy/webinars/2021/disclosure-avoidance-series/simulated-reconstruction-abetted-re-identification-attack-on-the-2010-census.html>
- ³ “Title 13, U.S. Code | History,” US Census Bureau, accessed July 15, 2022, https://www.census.gov/history/www/reference/privacy_confidentiality/title_13_us_code.html.
- ⁴ “Data Stewardship Executive Policy Committee,” accessed August 4, 2022, https://www.census.gov/about/policies/privacy/data_stewardship/dsep_committee.html
- ⁵ “Consistency of data products and formally private methods for the 2020 census,” US Census Bureau, accessed July 15, 2022, <https://irp.fas.org/agency/dod/jason/census-privacy.pdf>
- ⁶ Joseph Near, David Darais, and Kaitlin Boeckl, “Differential Privacy for Privacy-Preserving Data Analysis: An Introduction to our Blog Series,” Cybersecurity Insights (National Institute of Standard and Technology blog), July 27, 2020, <https://www.nist.gov/blogs/cybersecurity-insights/differential-privacy-privacy-preserving-data-analysis-introduction-our>.
- ⁷ Title 44 allows census records to be publicly available (unaltered) after 72 years.
- ⁸ “Census Bureau Sets Key Parameters to Protect Privacy in 2020 Census Results,” news release, US Census Bureau, June 9, 2021, www.census.gov/newsroom/press-releases/2021/2020-census-key-parameters.html
- ⁹ “Revised Data Metrics for 2020 Disclosure Avoidance,” US Census Bureau, September 17, 2020, <https://www2.census.gov/programs-surveys/decennial/2020/program-management/data-product-planning/disclosure-avoidance-system/2020-09-17-data-metrics-overview.pdf>.
- ¹⁰ 2010 Demonstration Privacy-Protected Microdata Files: Production Settings,” August 12, 2021, https://www2.census.gov/programs-surveys/decennial/2020/program-management/data-product-planning/2010-demonstration-data-products/01-Redistricting_File--PL_94-171/2021-06-08_ppmf_Production_Settings/2021-06-08-ppmf-factsheet-production-release.pdf
- ¹¹ Jingchen Hu and Claire McKay Bowen, “Prescribing Privacy: Human and Computational Resource Limitations and What Statisticians and Data Scientists Can Do,” *Amstat News*, September 1, 2022, <https://magazine.amstat.org/blog/2022/09/01/prescribing-privacy/>.

References

- Abowd, John M., Robert Ashmead, Ryan Cumings-Menon, Simson Garfinkel, Micah Heineck, Christine Heiss, Robert Johns, et al. 2022. "The 2020 Census Disclosure Avoidance System TopDown Algorithm." arXiv preprint 2204.08986. Ithaca, NY: Cornell University arXiv.
- Bowen, Claire M. 2021. "Personal Privacy and the Public Good: Balancing Data Privacy and Data Utility." Washington, DC: Urban Institute.
- Bowen, Claire M., and Simson Garfinkel. 2021. "Philosophy of Differential Privacy." *Notices of the American Mathematical Society* 68: 1727–39.
- Bun, Mark, and Thomas Steinke. 2016. "Concentrated Differential Privacy: Simplifications, Extensions, and Lower Bounds." In *Proceedings of the 14th International Theory of Cryptography Conference, Part I*, edited by Martin Hirt and Adam Smith, 635–58. Berlin: Springer.
- Dwork, Cynthia, Frank McSherry, Kobbi Nissim, and Adam Smith. 2006. "Calibrating Noise to Sensitivity in Private Data Analysis." In *Proceedings of the Third Theory of Cryptography Conference*, edited by Shai Halevi and Tal Rabin, 265–84. Berlin: Springer.
- Dwork, Cynthia, and Aaron Roth. 2014. "The Algorithmic Foundations of Differential Privacy." *Foundations and Trends in Theoretical Computer Science* 9 (3–4): 211–407.
- Dwork, Cynthia, and Guy N. Rothblum. 2016. "Concentrated Differential Privacy." arXiv preprint 1603.01887. Ithaca, NY: Cornell University arXiv.
- Federal Committee on Statistical Methodology. 2005. "Statistical Policy Working Paper 22: Report on Statistical Disclosure Limitation Methodology (2nd version)." Washington, DC: Office of Management and Budget.
- Fellegi, Ivan P. 1972. "On the Question of Statistical Confidentiality." *Journal of the American Statistical Association* 67 (337): 7–18.
- Fienberg, Stephen E., and Jiashun Jin. 2018. "Statistical Disclosure Limitation for Data Access." In *Encyclopedia of Database Systems* (2nd ed.), edited by Ling Liu and M. Tamer Özsu. Berlin: Springer.
- Hotz, V. Joseph, and Joseph Salvo. 2022. "A Chronicle of the Application of Differential Privacy to the 2020 Census." *Harvard Data Science Review* (special issue 2).
- Johnson, Noah, Joseph P. Near, and Dawn Song, D. 2018. "Towards Practical Differential Privacy for SQL Queries." *Proceedings of the VLDB Endowment* 11 (5): 526–39.
- Leclerc, Philip. 2019. "Results from a Consolidated Database Reconstruction and Intruder Re-identification Attack on the 2010 Decennial Census." Slide deck presented at workshop "Challenges and New Approaches for Protecting Privacy in Federal Statistical Programs," Washington, DC, June 6–7.
- Rogers, Ryan, Subbu Subramaniam, Sean Peng, David Durfee, Seunghyun Lee, Santosh Kumar Kancha, Shraddha Sahay, and Parvez Ahammad. 2020. "LinkedIn's Audience Engagements API: A Privacy Preserving Data Analytics System at Scale." arXiv preprint 2002.05839. Ithaca, NY: Cornell University arXiv.
- Ruggles, Steven, and David Van Riper. 2022. "The Role of Chance in the Census Bureau Database Reconstruction Experiment." *Population Research and Policy Review* 41 (3): 781–88.
- Snoke, Joshua, and Claire McKay Bowen. 2020. "How Statisticians Should Grapple with Privacy in a Changing Data Landscape." *Chance* 33 (4): 6–13.
- US Census Bureau. 2021. "Disclosure Avoidance for the 2020 Census: An Introduction." Washington, DC: US Government Publishing Office.

About the Authors

Claire McKay Bowen is a principal research associate in the Center on Labor, Human Services, and Population and leads the Statistical Methods Group at the Urban Institute. Her research focuses on developing and assessing the quality of differentially private data synthesis methods and science communication. In 2021, the Committee of Presidents of Statistical Societies identified her as an emerging leader in statistics for her technical contributions and leadership to statistics and the field of data privacy and confidentiality. She is also a member of the Census Scientific Advisory Committee, an advisory board member of the Future of Privacy Forums, and an adjunct professor at Stonehill College.

Bowen holds a BS in mathematics and physics from Idaho State University and an MS and PhD in statistics from the University of Notre Dame. After completing her PhD, she worked at Los Alamos National Laboratory, where she investigated cosmic ray effects on supercomputers.

Aaron R. Williams is a senior data scientist in the Income and Benefits Policy Center at the Urban Institute, where he works on retirement policy, microsimulation models, data privacy, and data imputation methods. He has worked on Urban's Dynamic Simulation of Income (DYNASIM) microsimulation model, the Social Security Administration's Modeling Income in the Near Term (MINT) microsimulation model, and the Tax Policy Center's synthesis of individual tax records. He holds a BS in economics from Virginia Commonwealth University, a BA in music from Virginia Commonwealth University, and an MS in mathematics and statistics from Georgetown University where he is currently an adjunct professor in the McCourt School of Public Policy.

Madeline Pickens is a data scientist at the Urban Institute. She works with Urban's Technology and Data Science team to support research and analysis relating to data privacy. Before joining Urban, she worked as a data analyst at the Universal Service Administrative Company on improving broadband access. She also worked as a graduate fellow in the University of Virginia's Social and Decision Analytics Division.

Madeline holds a bachelor's degree in economics from the University of Arizona, where she was a Flinn Scholar, and a master's degree in data science for public policy from Georgetown University, where she was a Whittington Scholar.

STATEMENT OF INDEPENDENCE

The Urban Institute strives to meet the highest standards of integrity and quality in its research and analyses and in the evidence-based policy recommendations offered by its researchers and experts. We believe that operating consistent with the values of independence, rigor, and transparency is essential to maintaining those standards. As an organization, the Urban Institute does not take positions on issues, but it does empower and support its experts in sharing their own evidence-based views and policy recommendations that have been shaped by scholarship. Funders do not determine our research findings or the insights and recommendations of our experts. Urban scholars and experts are expected to be objective and follow the evidence wherever it may lead.



500 L'Enfant Plaza SW
Washington, DC 20024

www.urban.org