

Guía para la **anonimización de datos estructurados**

Departamento Administrativo Nacional (DANE)

Archivo General

ARCHIVO GENERAL DE LA NACIÓN JORGE PALACIOS PRECIADO

Establecimiento público adscrito al
Ministerio de Cultura

Laura Sánchez Alvarado

Directora General (E) del Archivo General
de la Nación Jorge Palacios Preciado

Colaboradores:

Adela Díaz Acuña

Subdirectora de Transformación Digital e
Innovación Archivística

Carlos Enrique Rojas Núñez

Profesional Especializado de la
Subdirección de Transformación Digital e
Innovación Archivística

DANE

Julieth Alejandra Solano Villa

Directora Técnica de la Dirección DIRPEN

Andrea Milena Roncancio Sánchez

Coordinadora GIT Prospectiva y Análisis de
Datos

Profesionales:

Alexander González Coca

Johana Catherine Avila Alvarado

Gildardo Andrés Vargas Acuña

Bryan Felipe Lugo

German Leonidas Orjuela

Francisco Javier Lara Carrillo

Daniel Felipe Ortiz

Cristian Alejandro Ruge

Mónica Andrea Quiroga Rivera

Contenido

Introducción

1	Antecedentes nacionales e internacionales frente al uso de la anonimización	8
1.1.	Experiencia del DANE en los procesos de anonimización.	
1.2.	Experiencia del Archivo General de la Nación (AGN) en procesos de anonimización	
1.3.	Antecedentes en Colombia	
1.4.	Contexto internacional	
1.5.	Metodologías aplicadas en la anonimización de los censos	
1.5.1.	Eurostat	
1.5.2.	CEPAL	
1.5.3.	Australia	
1.5.4.	México	
1.5.5.	Estados Unidos	
1.5.6.	Técnicas de anonimización aplicadas a Censos – revisión de literatura	
1.6.	Implementación de instrumentos, software, aplicativos o sistemas que permitan la anonimización de datos	

2	Objetivo y alcance de la guía	20
3	Marco conceptual para la anonimización	21
4	Planeación del proceso de anonimización	26
	4.1 Revisión normativa sobre protección de datos e identificación de necesidades de información	
	4.1.1 Revisión de restricciones de publicación de la información	
	4.1.2 Revisión de las necesidades de los usuarios de la información	
5	Ejecución del proceso de anonimización	33
6	Recomendaciones	45
7	Bibliografía	98

Lista de tablas

Tabla 1.

Experiencias internacionales en el uso de software, aplicativos o sistemas para realizar procesos de anonimización

Tabla 2.

Clasificación de variables de base de datos

Tabla 3.

Medidas descriptivas principales para variables cuantitativas

Tabla 4.

Distribución de frecuencias para una variable con dos categorías

Tabla 5.

Clasificación por tipo de variable de la EAC en el 2016

Tabla 6.

Medidas descriptivas de algunas variables cuantitativas de la EAC en el 2016

Tabla 7.

Distribución de frecuencias de la variable Organización Jurídica de la EAC

Tabla 8.

Tabla resumen de la clasificación de las variables por su tipo de sensibilidad

Tabla 9.

Resumen de unidades de observación riesgosas en la anonimización teniendo en cuenta 3 riesgos

Tabla 10.

Clasificación de las variables por tipo de sensibilidad de COL20

Tabla 11.

Medidas descriptivas de las variables cuantitativas de COL20

Tabla 12.

Distribución de frecuencias del RH en COL20

Tabla 13.

Distribución de frecuencias del nivel de escolaridad en COL20

Tabla 14.

Distribución de frecuencias del grupo étnico en COL20

Tabla 15.

Unidades de observación riesgosas para los cinco riesgos más frecuentes para la anonimización de COL20

Tabla 16.

Técnicas basadas en la no perturbación de datos según el tipo de variable.

Tabla 17.

Técnicas basadas en la perturbación de datos según el tipo de variable

Tabla 18.

Planteamiento de técnicas de anonimización para cada uno de los riesgos identificados

Tabla 19.

Criterios para analizar la viabilidad de la anonimización de la base de datos

Tabla 21.

Porcentajes de unidades de observación por números de riesgos

Tabla 22.

Porcentajes de unidades de observación para cada riesgo priorizado

Tabla 23.

Unidades de observación riesgosas en Amazonas. Datos originales

Tabla 24.

Unidades de observación riesgosas en Amazonas. Datos anonimizados

Tabla 27.

Variaciones de los promedios de las variables numéricas a nivel departamental

Tabla 28.

Reidentificación de unidades de observación riesgosas

Lista de ilustraciones

Ilustración 1.

Esquema general del proceso de anonimización

Ilustración 2.

Flujograma reevaluación de riesgos de identificación

Introducción

Los Sistemas Estadísticos Nacionales (SEN) tienen como propósito proveer información estadística oficial, relevante, comprensiva, confiable y objetiva a los diferentes usuarios, siendo la información un elemento fundamental para la toma de decisiones. Este propósito implica importantes retos para los productores de información estadística, especialmente al observar una mayor demanda de información desagregada, un aumento en el uso de microdatos y una mayor necesidad de explotar los datos disponibles en los sistemas estadísticos.

El SEN colombiano establece dentro de sus objetivos promover entre sus miembros el acceso y el uso de microdatos para la producción y la difusión de estadísticas oficiales a nivel nacional y territorial que requiere el país de manera organizada y sistemática (Art. 5. Ley Estadística 2335 de 2023). De igual forma, el Código Nacional de Buenas Prácticas Estadísticas del SEN, en su principio 10 sobre accesibilidad de la información, incentiva a los miembros del SEN para que implementen prácticas que permitan el acceso de las estadísticas y los microdatos asociados a todo tipo de usuarios con el máximo detalle posible y en diferentes formatos y medios que faciliten la consulta, la visualización y el uso (DANE, 2017: 10). Además, el principio 11 del Código que trata sobre la confidencialidad, pretende incentivar en los productores de información estadística, así como la utilización de técnicas para la anonimización de los microdatos para garantizar la protección de la identificación o la localización geográfica de las fuentes empleadas en el proceso estadístico (DANE; 2017: 11).

Bajo estas premisas y adoptando la Ley Estatutaria 1266 de 2008 que establece el Habeas Data y regula el manejo de la información contenida en bases de datos personales, financieras, crediticias, comerciales, de servicios y provenientes de terceros países, además de establecer disposicio-

nes sobre la recolección, el tratamiento y la circulación de datos personales en el país (Congreso de Colombia, 2008) el Departamento Administrativo Nacional de Estadística (DANE), en su rol de coordinador del SEN, presenta la Guía para la anonimización de datos estructurados, cuyo propósito se centra en orientar a los integrantes del SEN sobre el proceso de anonimización de bases de datos que provienen de registros administrativos, operaciones estadísticas y fuentes secundarias.

Este documento se construye a partir de la experiencia del DANE y el Archivo General de la Nación (AGN), al implementar procesos propios de anonimización en distintas bases de datos y tomando como referencia experiencias nacionales e internacionales de oficinas de estadística, ministerios, secretarías, entre otros. Se espera que los integrantes del SEN puedan identificar en este documento las buenas prácticas, las herramientas y los instrumentos cuando se implementen procesos de anonimización para la producción y la publicación de sus propias estadísticas. La guía también presenta a lo largo de las etapas, la ejemplificación del proceso mediante el uso de una base de datos simulada, siendo este un insumo para que las entidades del SEN interesadas puedan seguir paso a paso el proceso y comparar los resultados que se presentan aquí. La guía se divide en seis partes: la primera parte se presentan antecedentes nacionales e internacionales sobre los usos de la anonimización; en la segunda parte se relaciona el objetivo y el alcance del documento; en la tercera se presenta el marco conceptual para el proceso de anonimización; en la cuarta parte se relaciona la planeación del proceso de anonimización; en la quinta parte se expone el proceso de anonimización, y finalmente, en la sexta parte se presentan algunas recomendaciones finales que deben ser tenidas en cuenta al momento de aplicar el proceso de anonimización.



1. Antecedentes nacionales e internacionales frente al uso de la anonimización

Son varias las experiencias internacionales que permiten observar la implementación de la anonimización, dado que su principal propósito es incrementar la desagregación de la información, así como mantener sus niveles de confidencialidad y generar un mayor aprovechamiento estadístico de la misma. Esta sección presentará algunas experiencias de países como Reino Unido, Estados Unidos y Países Bajos y los avances que se han tenido en el contexto nacional. Asimismo, se exponen experiencias y metodologías aplicadas a la anonimización de censos y herramientas o softwares utilizados en la aplicación de los procesos de anonimización.

1.1. Experiencia del DANE en los procesos de anonimización

El DANE ha demostrado un compromiso constante con la transparencia y la promoción del acceso a la información estadística en Colombia. Desde 2011 la entidad ha experimentado un proceso de transformación en la forma en que pone a disposición sus bases de datos, incluyendo encuestas de hogares, calidad de vida y encuestas estructurales como industria, comercio y servicios.

Anteriormente, el DANE operaba bajo la figura de convenios interinstitucionales es-

tablecidos con universidades y entidades del orden nacional para el suministro de sus bases de datos. Además, comercializaba algunas de sus operaciones estadísticas, permitiendo a investigadores y personas interesadas acceder a dichos datos previo pago, para ello, la entidad contaba con una dirección de mercadeo encargada de estas actividades. Sin embargo, con el objetivo de fomentar una mayor transparencia y apertura de datos, así como la de asegurar el acceso público a la información generada por la entidad, el DANE ha adoptado una nueva estrategia, en lugar de mantener un enfoque basado en la comercialización de datos, puso a disposición del público en general sus bases de datos a través de su catálogo central, el Archivo Nacional de Datos (ANDA)¹; esta transformación ha implicado cambios en la estructura organizativa al reemplazar la dirección de mercadeo por una dirección enfocada en potencializar la cultura estadística y promover el acceso abierto a la información.

Esta evolución ha permitido al DANE ampliar el acceso a sus operaciones estadísticas y facilitar su uso por parte de entidades públicas, entidades privadas, investigadores y el público en general. Igualmente, ha fortalecido la transparencia en la publicación de estadísticas del país y ha fomentado un mayor aprovechamiento de la información por parte de diferentes actores, impulsando así el desarrollo de la investigación y el análisis estadístico en Colombia.

¹ Disponible en: <https://www.dane.gov.co/index.php/servicios-al-ciudadano/tramites/transparencia-y-acceso-a-la-informacion-publica/informacion-de-interes?highlight=WyJhbmRhIl0=>.

De igual manera, la entidad cuenta con experiencia en los procesos de anonimización de bases de datos para garantizar la reserva estadística de las fuentes de datos y, al mismo tiempo, facilitar el máximo aprovechamiento de la información. En la actualidad, posee una amplia gama de operaciones estadísticas certificadas que incluyen la disponibilidad de microdatos y estas operaciones engloban principalmente unidades de observación como personas, empresas y establecimientos comerciales.

Este escenario presenta un desafío pues se debe garantizar la privacidad y la confidencialidad de sus fuentes y la información suministrada. Este hecho se refleja en el esfuerzo continuo de la entidad por promover el uso de técnicas y buenas prácticas en la producción de cada una de las operaciones estadísticas realizadas.

En la última década, el DANE ha realizado diversas actividades en el ámbito de la anonimización de bases de datos, lo que ha impulsado la exploración de nuevas soluciones. Estas actividades incluyen:

- La publicación de microdatos anonimizados que implica el uso de técnicas de anonimización que permitan a los usuarios replicar la mayoría de los indicadores generados con las bases de datos.
- La implementación de procesos de anonimización de bases de datos tradicionalmente no publicadas en forma completa. Sin embargo, el DANE se encuentra en el proceso de publicar microdatos del Censo Nacional Agropecuario (CNA) 2014, del Censo de Edificaciones (CEED) 2017-2023 I Trimestres, los Censos de Habitantes de Calle (CHC) 2019, 2020, 2021 y el Censo de Población y Vivienda 2018.
- La exploración de técnicas avanzadas de perturbación en los procesos de anonimización de datos, abarcando métodos de mayor complejidad, como la generación de datos sintéticos y la aplicación de técnicas como SWAP y PRAM, entre otras.
- La exploración y la implementación de tecnologías de privacidad en el análisis de datos sensibles, como la utilización de enclaves seguros, entornos de ejecución protegidos por hardware que permiten realizar cálculos en datos sensibles sin exponerlos a procesos no autorizados. Además, se ha considerado la aplicación de la privacidad diferencial, una técnica que agrega ruido aleatorio a los datos para preservar la privacidad de los individuos representados en ellos.
- El DANE ha observado la creación del Laboratorio de Datos de las Naciones Unidas, con su programa piloto enfocado en mejorar la privacidad en el intercambio internacional de datos. Este programa, conocido como el UN PET Lab, ha establecido colaboraciones con oficinas nacionales de estadística para abordar los desafíos de privacidad en el análisis de datos sensibles. Estas iniciativas internacionales refuerzan la necesidad de encontrar tecnologías y algoritmos que promuevan la privacidad en el contexto de la preservación de datos sensibles.
- La gestión de la privacidad en un entorno de cambios tecnológicos acelerados por las tecnologías de la información y la comunicación, así como al aprovechamiento del big data. En este contexto, es crucial abordar la vinculación de las fuentes de datos tradicionales con los accesos públicos de diferentes entidades y la información proporcionada por los individuos, con la creciente divulgación de datos a través de redes sociales y entidades privadas, existe la posibilidad de que la información suministrada sea vinculable, lo que plantea un desafío para garantizar la privacidad.
- El desarrollo de algoritmos eficientes para la anonimización de bases de datos de gran volumen. El DANE se enfrenta al desafío de implementar paquetes especializados en la anonimización de datos que sean capaces de manejar eficientemente bases de datos de gran volumen. Actualmente, los programas como SDC-micro, diseñados para tratar bases de

datos de tamaño muestral, presentan limitaciones cuando se supera el límite de 30.000 registros, lo que ocasiona problemas computacionales significativos. Esta limitación restringe su uso en fuentes de datos más grandes, como registros administrativos y censos.

- El aprovechamiento de fuentes secundarias como imágenes satelitales, datos de redes sociales como X y Facebook. Además, se vislumbra la posibilidad de utilizar otras fuentes audiovisuales, como audio, video y fotografías, para las que se ha ampliado recientemente la disponibilidad de algoritmos de análisis. Esto plantea la necesidad de considerar el tratamiento adecuado que se debe dar a estas fuentes de datos no convencionales. Para abordar este reto, el DANE debe establecer marcos éticos sólidos que guíen el uso responsable de estas fuentes de datos, tomando en cuenta los términos y las condiciones que estas redes establecen para su uso. Asimismo, se requiere desarrollar técnicas de anonimización adaptadas a estas nuevas fuentes, teniendo en cuenta su naturaleza y características específicas. Es esencial garantizar la confidencialidad y la privacidad de los individuos representados en estos datos y aplicando rigurosos procesos de anonimización que preserven su identidad.

La experiencia acumulada por el DANE en los procesos de anonimización, ha demostrado su capacidad para proporcionar datos de calidad, preservando la privacidad de los individuos. Este conocimiento adquirido puede ser aprovechado para la construcción de una guía de anonimización de datos más actualizada, que promueva la privacidad, la confidencialidad y el uso ético de la información estadística. La entidad sigue siendo líder en la aplicación de buenas prácticas en la anonimización de datos en Colombia y su experiencia puede contribuir significativamente a mejorar y actualizar las técnicas y los enfoques utilizados en este campo.

1.2. Experiencia del Archivo General de la Nación (AGN) en procesos de anonimización

El AGN coordina el Sistema Nacional de Archivos y desarrolla procesos de organización y fortalecimiento de los archivos de la administración pública a nivel nacional y territorial. Además, verifica el cumplimiento de la normativa archivística vigente por medio de la inspección, la vigilancia y el control de los archivos de todas las entidades del Estado y personas naturales y jurídicas de derecho público y privado, de acuerdo con la normatividad vigente.

En este contexto, el AGN en articulación con el Ministerio de Tecnologías de la Información y las Comunicaciones, la Superintendencia de Industria y Comercio, el Departamento Administrativo de la Función Pública, el Departamento Nacional de Planeación y el DANE desarrollaron una *Guía de anonimización de datos estructurado*² cuya finalidad es precisar conceptos y proporcionar los lineamientos metodológicos para realizar procesos de anonimización de datos personales e información producida o gestionada por entidades.

El AGN, además, aplica los procesos de anonimización en los siguientes trámites:

Emisión de Conceptos Técnicos: en cumplimiento de las funciones dadas al AGN por las Leyes 80 de 1989 y 527 de 2000, este responde las consultas técnicas sobre temas archivísticos que son formuladas por ciudadanos, entidades públicas y empresas privadas.

Dado que dichos conceptos sirven para resolver inquietudes de otros ciudadanos u organizaciones, se decidió publicarlos en la página web de la entidad. Para ello se impartió una directriz interna en la que se solicita a los profesionales encargados de la redacción de estos documentos omitir en

² Disponible en https://www.archivogeneral.gov.co/sites/default/files/Estructura_Web/5_/_Consulte/Recursos/Publicaciones/2021_06_25_Guia_de_Anonimizacion.pdf

los mismos cualquier información que pueda identificar al solicitante, ya sea persona natural o jurídica; para la remisión de la respuesta se elabora un oficio al cual se le anexa el correspondiente concepto.

Certificación de tiempo de servicio a Exfuncionarios del DAS: el AGN recibió la documentación de Historias Laborales y Nóminas del extinto DAS, por lo cual es responsable de atender las solicitudes de Certificaciones de Tiempos de Servicio y de Aportes pagados por pensión, provenientes de los exfuncionarios, de los Fondos de Pensiones y de la Rama Judicial, las cuales deben ir acompañadas de los soportes documentales, en especial de la nómina; para este último proceso, se realiza la digitalización de los folios correspondientes en formato PDF a cuyas imágenes digitales se les realiza un enmascaramiento de la información sensible en color negro para suprimir toda información que corresponda al solicitante, con lo cual se evita que datos personales de otros funcionarios quede expuesta.

1.3. Antecedentes en Colombia

En Colombia, el DANE ha sido pionero en el desarrollo de los procesos de anonimización de datos estructurados. En 2014 se publicaron los Lineamientos para la anonimización de microdatos que identificaban tres grandes etapas en el proceso: preanonimización, anonimización de microdatos de uso interno y anonimización (DANE, 2014: 11).

El Ministerio de Salud, basado los lineamientos del DANE, generó sus propios lineamientos para anonimizar microdatos a través del documento Lineamientos para la Anonimización de Datos del Sistema Nacional de Estudios y Encuestas Poblacionales para la Salud (Ministerio de Salud, s.f.). En este ejercicio, el Ministerio incluyó las etapas previstas por el DANE y agregó con mayor detalle, las técnicas de anonimización que se usarían en sus bases de información; específicamente presenta el uso de los métodos de perturbación o los métodos de reducción en la información (Ministerio de Salud, s.f.:12-13).

Adicional a los lineamientos y las orientaciones en materia de anonimización, Colombia ha gestionado la implementación de un marco legal en esta materia, partiendo del derecho a la intimidad personal, familiar y el derecho a conocer, actualizar y rectificar las informaciones que se hayan recogido en las bases de datos y en archivos de entidades públicas y privadas, que aparece consignado en el artículo 15 de la Constitución Política de 1991.

En términos de la confidencialidad, enmarcados en las directrices para el SEN, es importante tener en cuenta la Ley Estatutaria 1266 de 2008 que establece el Habeas Data y regula el manejo de la información contenida en bases de datos personales, financieras, crediticias, comerciales, de servicios y provenientes de terceros países, además de establecer disposiciones sobre la recolección, el tratamiento y la circulación de datos personales en el país (Congreso de Colombia, 2008).

Junto con estas leyes se cuenta con la Ley 1581 de 2012, que trata sobre la protección de datos personales, reglamentada parcialmente por el decreto 1377 de 2013, en el que se señala que el acceso a los datos se debe restringir y la información debe estar sujeta a tratamiento por parte del responsable, como lo indica en su artículo 4, manteniendo los principios de acceso y circulación restringida, de seguridad y de confidencialidad. Para el DANE, particularmente, la Ley 79 de 1993 establece la reserva estadística y la confidencialidad de las fuentes cuando se realizan procesos de recolección de información a través de censos o encuestas.

Respecto a la transparencia y el acceso de la información pública, la Ley 1712 de 2014, regula el derecho de acceso a la información pública, haciendo énfasis en el establecimiento de una política de datos abiertos por parte de las entidades públicas, y Decreto 2573 de 2014 del Ministerio de Tecnología de la información y las Comunicaciones (Min-TIC), donde se establecen los lineamientos generales de la Estrategia de Gobierno en línea, indicando los principios y los fundamentos a tener en cuenta en las entidades públicas destacando entre ellos, la excelen-

cia en el servicio ciudadano, la apertura y la reutilización de datos públicos, la estandarización, la innovación, entre otros. Igualmente, se establecen cuatro componentes que facilitarán la masificación de la oferta y la demanda en gobierno en línea, resaltando entre estos, la seguridad y la privacidad de la información, siendo un componente transversal.

Estas normas han tenido su complemento con la reglamentación del SEN, a través de la Ley 1753 de 2015, del Decreto 1743 de 2016 y la Ley Estadística 2335 de 2023. En esta última, se establece el fortalecimiento y el aprovechamiento amplio e intensivo de los registros administrativos, así como el intercambio de información entre los integrantes del SEN, como fuente para la producción de estadísticas oficiales y el mejoramiento de la calidad y la coherencia en las cifras (Art. 11. Ley Estadística Nacional 2335 de 2023).

Finalmente, en 2017 con la actualización del Código Nacional de Buenas Prácticas del SEN, se planteó la implementación de prácticas en materia de acceso y confidencialidad de la información por parte del SEN, incentivando un mayor acceso y uso de la información de los productores de estadísticas en Colombia, así como la difusión de información anonimizada para garantizar la confidencialidad de la misma.

Antecedentes de anonimización de datos no estructurados en el Centro Nacional de Memoria Histórica en Colombia

El Centro Nacional de Memoria Histórica de Colombia (CNMH), desarrolló la *Guía para la Anonimización de Datos e Información No Estructurada*³, cuyo objetivo es brindar los lineamientos técnicos y la orientación metodológica para garantizar cualquier información producida, gestionada o recolectada

por entidades públicas o privadas que contenga datos personales o información de identificación; se enmarca en las premisas de protección de derechos, transparencia y datos abiertos, acceso e interoperabilidad, eficiencia administrativa y reportes de información.

La CNMH utiliza como definición de dato no estructurado la aportada por el Consejo Nacional de Política Económica y Social en el CONPES 2018: *“información cuya organización y presentación no está guiada por ningún modelo o esquema. En esta categoría se incluye, por ejemplo, textos, audios, contenidos de redes sociales, etc”*.

En este contexto, las técnicas de anonimización varían dependiendo de la información que se utilice, textos, archivos de audio, archivos de video, formatos de imagen u otros formatos multimedia. A continuación, se relacionan las principales técnicas por tipo de información.

Técnicas de anonimización para documentos físicos

Los documentos en físico deben cumplir con organización archivística y en términos de proceso de digitalización se debe tener en cuenta la conservación física, las condiciones ambientales, las condiciones operacionales, la seguridad de la información y la preservación en el largo plazo. En el caso de documentos originales que posean valor histórico, estos no podrán ser destruidos aun cuando hayan sido reproducidos o almacenados en cualquier medio. Debe controlarse el acceso a los documentos y la autorización de acceso debe ser normalizada de acuerdo con los lineamientos de la entidad.

Los documentos físicos pueden necesitar anonimizarse de forma parcial o total. En este sentido, para retirar los documentos de un expediente debe indicarse el lugar de almacenamiento por medio de un testigo (formato en papel) y en el caso de tratarse

³ Disponible en <https://centrodememoriahistorica.gov.co/wp-content/uploads/2022/08/GUIA-DE-ANONIMIZACION.pdf>

de varios documentos, se puede incluir un listado al comienzo del expediente con los documentos que se retiran.

Para el proceso de anonimización se hace una copia del documento original y de la copia se retira la información que se requiera. Posteriormente se hace una copia al documento resultante y es esta la que se dará a conocer en forma de fotocopia o archivo digitalizado. Se sugiere marcar las copias 1 y 2 para identificarlas a partir de la original. El documento original debe conservarse garantizando su integridad.

Técnicas de anonimización para documentos textuales en formatos digitales

Para su anonimización se recomienda usar una metodología análoga a la descrita anteriormente, es decir, se debe preservar el documento original y paralelamente guardar una copia que contenga el borrado de las partes restringidas que será publicada como copia anonimizada. Para este proceso se pueden utilizar archivos electrónicos como editores y programas de diseño gráfico.

Técnicas de anonimización para archivos de tipo audiovisual

El lenguaje audiovisual contiene diversos tipos de información sensible de identificación, para ello, se pueden utilizar técnicas de desidentificación en contenido multimedia, cuyo objetivo es ocultar o eliminar identificadores personales, o sustituirlos por identificadores sustitutos en contenido multimedia para que se evite la divulgación y el uso de datos para fines no relacionados con el propósito para el que la información fue obtenida inicialmente.

- **Análisis de audio:** es el proceso de comprimir los datos y empaquetarlos en un formato de audio. Luego Audio Analytics realiza la extracción de significado e información de señales de audio para su análisis. El audio se puede presentar por medio de representación del sonido y archivos de sonido sin formato. Existen tres formatos de audio principales: formato de audio sin comprimir, for-

mato de audio comprimido sin pérdida y formato de audio comprimido Lossy.

El análisis de audio se puede aplicar a servicios de vigilancia, detección de amenazas, sistema de tele-monitoreo, sistema de red móvil, etc.

- **Análisis de video:** aunque los videos representan el 80% de los datos no estructurados cada día se generan y almacenan millones de pixeles de información, tanto en plataformas, como Youtube, como en cámaras de vigilancia y por personas independientes. Es así como en las dimensiones del análisis, el tamaño del video al ser mayor utiliza la red y el servidor en tiempo de procesamiento y las conexiones de bajo ancho de banda crean mayor tráfico en la red ya que los videos circulan lentamente.

El análisis de video se puede aplicar en la identificación en accidentes de tránsito, en la policía de tráfico, en negocios, en seguridad, en análisis para inteligencia empresarial, en análisis de objetivos y escenarios, en direction analytics, en eliminar la ecuación humana a través de la automatización, etc.

Los formatos de imagen son más difíciles de limpiar y la selección de técnicas de desidentificación se aplican según el tipo de información. Por ejemplo, varios archivos multimedia como los formatos de imágenes pueden ser accesibles para ciertas solicitudes en las que el algoritmo de los solicitantes busca metadatos de archivos de imágenes y produce resultados basados en texto para facilitar el borrado de imágenes sensibles.

Debido a las características particulares de la producción audiovisual y al gran número de marcadores de información personal que se puede encontrar en dicho material, el proceso de anonimización implica una actividad de entendimiento del proceso comunicativo que caracteriza el lenguaje natural. El proceso de desidentificación implica reconocer los insumos sonoros, visuales, espaciales, temporales, comportamentales, etc.

Actualmente no se han identificado herramientas informáticas que automáticamente realicen el proceso de anonimización total de los datos. Sin embargo, las nuevas herramientas de inteligencia artificial por medio de reconocimiento facial, el análisis de sonido, el análisis de texto, etc., han permitido un avance en términos de aplicación metodológica para este proceso.

1.4. Contexto internacional

La mayor parte de las experiencias internacionales sobre la anonimización de información estadística parten de las indicaciones o recomendaciones propias de los países, establecidas en sus sistemas nacionales de estadística. Este proceso se ha observado también como un posible mecanismo para dar seguimiento a la legislación internacional en materia de protección, privacidad y confidencialidad de la información.

En Reino Unido, por ejemplo, la Oficina del Comisionado de Información desarrolló el Código de Buenas Prácticas para la Anonimización (Oficina del Comisionado de información, 2012), en el cual se presentan los antecedentes jurídicos de la protección de datos de ese país y explica los beneficios de la anonimización y el por qué se debe hacer, todo enmarcado en los principios que deben guiar dicho proceso. También se señalan los riesgos que se pueden presentar al cruzar información en bases que ya se encuentran anonimizadas, generando posibles identificaciones de usuarios y dado que se debe realizar el control de los datos en el sistema estadístico en este país.

Otro esfuerzo internacional destacable en la documentación de procesos de anonimización para orientar a las entidades que producen estadística en los Sistemas Estadísticos Nacionales de los distintos países, es el realizado por la Comisión Económica para Europa de las Naciones Unidas (UNECE, por sus siglas en inglés), quien publicó el *Manual sobre el control de divulgación estadística*, en el que se proveen lineamientos técnicos para el control de la revelación de la información, se describen los métodos aplicables para la protección de la privacidad de

la información y se explica detalladamente el programa de anonimización ARGUS, una iniciativa que viene liderando UNECE para implementar mecanismos de anonimización en los productores de estadística. ARGUS es un programa interactivo y libre, cuyas funcionalidades permiten identificar los datos de una base (metadatos), seleccionar y calcular las tablas de frecuencia, establecer la base de los métodos de anonimización y aplicar las diferentes técnicas a las variables relevantes. El programa es compatible con otros programas estadísticos (Morales, 2017:10).

Igualmente, la Red Internacional de Encuestas de Hogares (IHSN por sus siglas en inglés), en 2014 publicó una introducción sobre los controles a tener en cuenta en la divulgación estadística de información para uso de los integrantes del SEN y donde explicaba los métodos para difundir la información de datos confidenciales, teniendo en presentaba los distintos riesgos a los que se enfrentan los productores de información estadística, así como posibles mecanismos para valorar este tipo de riesgos y métodos de anonimización (IHSN, 2014).

Por otra parte, la anonimización permite atender de una mejor manera las demandas y las necesidades de información de distintos usuarios. Dentro de los sectores que se han caracterizado por generar este tipo de requerimientos, respecto a mayores desagregaciones de la información, especialmente en países como Estados Unidos, han sido investigadores, centros de investigación, universidades y el sector público.

En Estados Unidos se ha observado una fuerte relación entre la información y la investigación, siendo este último uno de principales focos para el desarrollo de procesos de anonimización. La Oficina de Censos ha liderado el estudio de técnicas de anonimización e integración de información, así como la opción de acceso a microdatos no anonimizados con propósitos académicos y de investigación. Algunas de las técnicas implementadas se basan, por ejemplo, en eliminar las variables de identificación directa, utilizar técnicas como umbrales geográficos o categóricos de las variables, re-

dondeo, infusión de ruido, recodificación, intercambio de registros basado en rangos o en proximidad (Morales, 2017: 9).

Dentro las bases anonimizadas publicadas por esta institución estadounidense se destacan las de censos demográficos y de encuestas económicas y, para el caso del Censo Agropecuario, se disponen al público algunas tablas agregadas por Estado y Condado. Todo esto bajo el marco de la normatividad nacional e internacional de confidencialidad de la información de los usuarios.

La Oficina de Estadísticas de Países Bajos (CBS, por sus siglas en holandés), ha trabajado conjuntamente con la UNECE (por sus siglas en inglés) para el avance de la investigación sobre procesos de anonimización de la información, destacándose de esta cooperación, la implementación de la iniciativa *Control de Divulgación Estadística* que es un estudio exhaustivo sobre la importancia de la confidencialidad y que busca, mediante la implementación de proyectos, generar mecanismos que permitan garantizarla.

1.5. Metodologías aplicadas en la anonimización de los censos económicos en la práctica internacional

1.5.1. Eurostat

EuroStat⁴ estableció en 2008 una metodología para la anonimización de los datos de la Encuesta sobre la Estructura de los Ingresos, en ella se establecen un conjunto de criterios detallados y reglas generales aplicadas a la anonimización de los microdatos, además de dicha metodología, algunos países establecieron diversos parámetros para ampliar la anonimización específica de algunos datos individuales.

Los institutos nacionales de estadística de cada país donde se aplica la Encuesta sobre la Estructura de los Ingresos se encar-

gan de seleccionar la muestra, preparar los cuestionarios, aplicar la encuesta, sintetizar los resultados y enviarlos a Eurostat, quien posteriormente se encargará de procesar los datos. En tanto la anonimización de los datos, Eurostat en 2008 desarrolló una metodología para la anonimización de los datos de la encuesta donde se estableció un conjunto de criterios detallados y reglas generales aplicadas a la anonimización de los microdatos. Además, de dicha metodología, algunos países adoptaron diversos parámetros para ampliar la anonimización específica de algunos datos individuales.

Las reglas generales aplicadas en la anonimización de los microdatos de la Encuesta sobre la Estructura de los Ingresos son:

- La recodificación de los cuasi-identificadores categóricos NACE, NUTS y SIZE de las empresas. La codificación resulta en mezcla de secciones, subsecciones y divisiones NACE y niveles NUTS 0 o 1 según el Estado miembro.
- Realizar una recodificación global en la variable edad para restringir sus valores a 6 intervalos (14-19, 20-29, 30-39, 40-49, 50-59, 60+).
- Eliminar la ciudadanía y el identificador de la empresa.
- Suprimir el factor de extrapolación para unidades locales.
- Brindar protección adicional a los empleados mediante la microagregación de clasificación individual sin restricciones para las variables métricas (días de ausencia e ingresos), con ello se busca ocultar Información relativa a las personas y estas variables, siempre y cuando sean grupos de al menos tres empleados.

⁴ Disponible en <https://ec.europa.eu/eurostat/web/microdata/structure-of-earnings-survey>

1.5.2. Comisión Económica para América Latina y el Caribe (CEPAL)

La CEPAL realizó un estudio denominado “Control de divulgación estadística para las tablas censales del Caribe, una propuesta para ampliar la disponibilidad de datos censales desagregados”⁵. Este estudio revisa el problema del control de divulgación estadística para las tablas censales y las mejores prácticas internacionales, centrándose particularmente en el uso de métodos perturbativos; lleva a cabo análisis comparativos y pruebas del método de perturbación celular y métodos de redondeo aleatorio; recomienda que estos métodos se pongan a disposición de las oficinas de estadística a través del software Recuperación de Datos para Áreas pequeñas por Microcomputador (REDATAM).

Práctica actual en el Caribe

Al producir cuadros censales como los publicados en los informes de censos nacionales, las oficinas de estadística del Caribe limitan o evitan la publicación de recuentos pequeños que conduzcan a la divulgación, restringen la publicación de tablas para áreas geográficas y grupos pequeños o se recodifican las variables. Este tipo de supresión o rediseño de tablas se considera una forma de control de divulgación estadística, pero si se usa solo no logra un equilibrio eficiente entre el riesgo de divulgación y la utilidad de los datos. Cuando se utiliza REDATAM para difundir los datos del censo, el riesgo de divulgación se gestiona a través del diseño de la aplicación en línea. Los aspectos relevantes del diseño incluyen: las variables que están disponibles en la aplicación, la forma en que se codifican esas variables, la medida en que los usuarios pueden tabular diferentes variables entre sí (o usar variables para filtrar las consultas de la base de datos), el nivel de desagregación geográfica que se encuentra disponible, y si se le brinda al usuario la funcionalidad de utilizar el lenguaje de programación de RE-

DATAM para especificar consultas a la base de datos.

Métodos pre-tabulares

El método de perturbación pre-tabular más utilizado es el intercambio de registros que se ha utilizado en censos en Estados Unidos (censos de 1990, 2000 y 2010), Reino Unido (censos de 2001 y 2011), Suecia (2011), Austria (2011) y Bélgica (2011). La esencia del método de intercambio de registros (o intercambio de datos) es que se intercambia una pequeña proporción de hogares que son similares en términos de algunas características sociales y demográficas básicas, es decir, sus códigos geográficos se intercambian como si los hogares intercambiaran físicamente lugares. La rutina de intercambio se puede refinar para favorecer a los hogares con características raras que corren un mayor riesgo de identificación o divulgación de atributos (esto se conoce como intercambio de registros específicos). La distancia geográfica entre los hogares intercambiados también puede variar según el riesgo de divulgación, dentro de un cierto límite debido a que un hogar no se intercambiaría con otro en una parte completamente diferente del condado.

Métodos post-tabulares

La alternativa a la perturbación de los microdatos del censo es aplicar la perturbación después de la tabulación, es decir, perturbar los datos en las tablas finales del censo. De esta forma, el control de la divulgación se convierte efectivamente en un “complemento” del proceso de tabulación. Hay tres enfoques post-tabulares principales que han empleado las oficinas de estadística: ajustes de celdas pequeñas; redondeo aleatorio, y el método de perturbación de celdas de la Oficina Australiana de Estadísticas (ABS).

⁵ Disponible en https://repositorio.cepal.org/bitstream/handle/11362/46628/4/S2000895_en.pdf

1.5.3. Australia

La Oficina de Estadísticas de Australia (ABS) se compromete a garantizar que se implementen medidas de seguridad para proteger la información sensible recopilada del censo. La seguridad de la información incluye la conversión de nombres a números anónimos y solo se permite que una pequeña cantidad de funcionarios ABS usen estos códigos, mientras se aplican estrictos protocolos de seguridad⁶. La ABS contrató a expertos en criptografía de la Universidad de Melbourne para investigar diferentes métodos de codificación de nombres para su uso en proyectos de integración de datos, haciendo énfasis en el Censo de 2016 para mejorar el valor de los datos del censo, combinándolos con datos de diferentes fuentes (p.ej., encuestas o registros administrativos), obteniendo mejor comprensión del panorama social y económico para ayudar en la toma de decisiones gubernamentales en áreas como salud, educación, infraestructura y economía.

1.5.4. México

El Instituto Nacional de Estadística y Geografía (INEGI) realiza desde 1930 censos económicos y el último censo económico se llevó a cabo en 2019 y su objetivo era recopilar información estadística básica (correspondiente a 2018) de todos los establecimientos productores de bienes y servicios, con el fin de producir indicadores económicos del país a gran nivel de detalle geográfico, sectorial y temático. La información recopilada durante los censos económicos realizados en 2019 se sometió al sistema de confidencialidad llamado “supresión parcial con datos complementarios”⁷ con el fin de garantizar la confidencialidad de los datos suministrados por los informantes, cumpliendo con lo dispuesto en los artículos 37 y 38 de la Ley del Sistema Nacional de Información Estadística y Geográfica; este sistema consiste en suprimir los datos de

las variables económicas de los renglones con problemas de confidencialidad, donde permanecen el número de unidades económicas y para compensar se agregan indicadores relacionados con las variables que se suprimieron, por lo tanto, se muestra la totalidad de renglones existentes de cada cuadro (tengan o no problemas de confidencialidad).

1.5.5. Estados Unidos

La Oficina de Censo de Estados Unidos salvaguarda la información con los controles de divulgación estadística, disfrazando los datos originales con métodos de intercambio de datos en el censo económico para establecer estos símbolos válidos y tener un control de información. Como parte de las buenas prácticas esta oficina ha implementado desde 1970 métodos de proyecciones de privacidad del censo, lo que ha mejorado los procesos con el avance de la tecnología.

1.5.6. Técnicas de anonimización aplicadas a censos – revisión de literatura

Dick et al (2023) *Confidence-ranked reconstruction of census microdata from published statistics*⁸ (Reconstrucción clasificada por confianza de microdatos censales a partir de estadísticas publicadas): consiste en procesos de reconstrucción de datos a partir de información contenida en estadísticas agregadas publicadas para reconstruir filas completas de un conjunto privado de datos por medio de la aplicación de técnicas aleatorias de optimización no convexa. Esta técnica se aplicó a datos del censo de Estados Unidos de 2010 y en conjunto de datos de la Encuesta de la Comunidad Americana.

Fioretto et al (2021) *Differential privacy of hierarchical Census data. An optimization approach* (Privacidad diferencial de los datos censales jerárquicos. Un enfoque de optimización): implica el desarrollo de un mecanismo que optimiza el ruido introducido

⁶ Disponible en [https://www.ausstats.abs.gov.au/ausstats/subscriber.nsf/0/844269E83C4B6666CA25823C00178BBB/\\$File/1351055162_2018.pdf](https://www.ausstats.abs.gov.au/ausstats/subscriber.nsf/0/844269E83C4B6666CA25823C00178BBB/$File/1351055162_2018.pdf)

⁷ https://www.inegi.org.mx/contenidos/productos/prod_serv/contenidos/espanol/bvinegi/productos/nueva_estruc/702825196530.pdf

⁸ Disponible en: <https://www.pnas.org/doi/10.1073/pnas.2218605120>

para asegurar la privacidad diferencial que garantiza la coherencia entre diferentes niveles jerárquicos (nacional, estatal, condado). El mecanismo se basa en programación dinámica que aprovecha la naturaleza jerárquica del problema y la estructura de la función objetivo.

Parra-Arnau, et al (2020) *Differentially private data publishing via cross-moment microaggregation* (Publicación de datos diferencialmente privados mediante microagregación entre momentos): la microagregación de momentos cruzados reemplaza los registros originales en un grupo por estadísticas adicionales lo que reduce la sensibilidad de los datos y aporta mayor privacidad y utilidad que la microagregación normal.

1.6. Implementación de instrumentos, software, aplicativos o sistemas que permitan la anonimización de datos

En esta sección se presentan las experiencias de los referentes internacionales frente al uso o la aplicación de herramientas, instrumentos, softwares, aplicativos o sistemas para realizar procesos de anonimización, sus aplicaciones, actividades desarrolladas e implicaciones de uso.

Tabla 1. Experiencias internacionales en el uso de software, aplicativos o sistemas para realizar procesos de anonimización

Referente	Software, aplicativo o herramienta
EuroStat	Cuenta con un software que permite realizar anonimización, mediante el manual de estadística empresarial europeo, hace énfasis en softwares como T- ARGUS y Sdc-Table que permiten la anonimización de información o el control de divulgación estadística. Son herramientas que realizan sus procedimientos sobre dos tipos de salidas estadísticas, como los datos confidenciales presentados en tablas y datos confidenciales en archivos de microdatos.
CEPAL	La Cepal desarrolló el software REDATAM que es un sistema computacional de carácter social e interactivo que facilita el procesamiento, el análisis y la diseminación web de la información de censos, encuestas, registros administrativos, indicadores nacionales/ regionales y otras fuentes de datos. Asimismo, es una herramienta que permite el procesamiento en línea de la información censal y estadística producida por las instituciones gubernamentales, divulgando dicha información de manera segura y sin ningún costo ya que el usuario interactúa a través de páginas predefinidas seleccionando tabulados e indicadores para luego enviar una solicitud al servidor y posteriormente este devuelve el resultado en forma de tabulado, gráfico o mapa estático; la solicitud del usuario puede tomar formatos como frecuencias, cruces de variables, promedios, conteos o listas por área geográfica.

Referente	Software, aplicativo o herramienta
España	El ayuntamiento de Barcelona, anonimizó los datos personales que tienen a su disposición para proteger la información sensible de los ciudadanos con un software llamado Nymiz. Este es un software que detecta los datos personales en archivos no estructurados y también en datos estructurados, y anonimiza o seudonimiza de forma reversible o irreversible los datos de acuerdo con las necesidades del tratamiento de la información. Es un programa de pago y su costo depende del volumen de datos a procesar.
Canadá	Statistics Canada desarrolló el software de control de divulgación automatizado G- Confind con el fin de proporcionar el nivel adecuado de protección para celdas confidenciales y minimizar la pérdida de información. G-Confind es un sistema generalizado que evita la divulgación de información confidencial en datos tabulares por medio del método de supresión de celdas, en el que se identifican y suprimen las celdas sensibles y complementarias a proteger, además, puede usarse para auditar patrones de supresión de celdas y encontrar agregados confidenciales y datos tabulares redondos.
Países Bajos	Para el control de divulgación estadística Estadísticas de Países Bajos fue pionero en el desarrollo del software μ -ARGUS, el cual es un programa interactivo flexible que guía a los usuarios en el proceso de protección de los datos y que está basado en el enfoque de gestión de riesgos. Este software incorpora algunos métodos sofisticados tradicionales para la anonimización de microdatos como p. ej. métodos de enmascaramiento perturbativo y métodos de enmascaramiento no perturbativo y adicionalmente incorpora otros como micro agregación, PRAM, redondeo, codificación superior e inferior, intercambio de rango y adición de ruido.
Estados Unidos	Desde el Instituto Nacional de Estándares y Tecnología tienen el programa de ingeniería de privacidad y en él se describen 15 herramientas de desidentificación y el software ARX. Estas técnicas son aplicadas a un conjunto de datos para prevenir o limitar los riesgos de la privacidad de los individuos, los grupos protegidos y los establecimientos; estas técnicas pueden ser introducción de ruido como la privacidad diferencial, el enmascaramiento de datos y la creación de conjuntos de datos sintéticos que se basan en modelos que preservan la privacidad.

Fuente: DANE, a partir de información de referentes internacionales.

2. Objetivo y alcance de la guía



Alcance

La *Guía de anonimización de datos estructurados* tiene como objetivo orientar a las entidades integrantes del SEN sobre la planeación y la ejecución del proceso de anonimización de datos estructurados, que permitan el acceso y el aprovechamiento de los datos teniendo en cuenta el índice de información clasificada y reservada de la Ley de transparencia⁹, controlando el riesgo de identificación de la fuente; no obstante, la guía no contiene lineamientos obligatorios para la aplicación del proceso de anonimización.

De esta manera, esta guía puede ser consultada por:

- Instituciones generadoras de datos estructurados.
- Instituciones y personas que usan datos estructurados.
- Entidades del SEN y personas que deseen salvaguardar la información privada de sus unidades de observación.

Objetivos específicos

- Establecer los conceptos necesarios para comprender los procesos de anonimización de datos estructurados.

- Definir la estrategia de planeación que permita la implementación del proceso de anonimización.
- Presentar los aspectos metodológicos que debe considerar el usuario en la implementación de técnicas de anonimización.
- Orientar a las entidades del SEN sobre las técnicas de anonimización aplicables de acuerdo con las características de los conjuntos de datos que requieran el proceso de anonimización en el marco de la normatividad de protección de datos personales.

Tanto la conceptualización metodológica de la guía, sus lineamientos como su aplicación están orientados a la anonimización de datos estructurados y está dirigida a las entidades del SEN cuya misión es gestionar, almacenar, administrar, procesar, producir, custodiar y publicar información; el proceso consta de seis etapas: revisiones previas; análisis de riesgos; selección de técnicas; análisis de viabilidad; aplicación de técnicas, y evaluación de resultados.

⁹ Disponible en: <https://www.archivogeneral.gov.co/transparencia/datos-abiertos/indice-informacion-clasificada-reservada#:~:text=El%20Índice%20de%20Información%20Clasificada,calificada%20como%20clasificada%20o%20reservada.>



3. Marco conceptual para la anonimización

Se define la **anonimización de microdatos** como un proceso técnico que consiste en:

“(…) transformar los datos individuales de las unidades de observación, de tal modo que no sea posible identificar sujetos o características individuales de la fuente de información, preservando así las propiedades estadísticas en los resultados”
(Decreto 1743 de 2016: Art. 2.2.3.1.1).

La finalidad de la anonimización es impedir que, a partir de una información o de una combinación de informaciones, se logren identificar sujetos individuales en un archivo de microdatos (Morales, 2017: 5).

El concepto de **microdato** es fundamental en la concepción del proceso de anonimización, pues corresponde a:

“(…) los datos sobre las características asociadas a las unidades de observación que se encuentran consolidadas en una base de datos”
(Artículo 5. Ley Estadística 2335 de 2023).

A su vez, el Sistema de consulta de conceptos estandarizados del DANE define que una **base de datos** hace referencia a:

“... un conjunto o colección de datos interrelacionados entre sí, que se utilizan para la obtención de información de acuerdo con el contexto de los mismos y que son almacenados sistemáticamente para su posterior uso”
(Sistema de consultas de conceptos estandarizados, 2018).

Según lo anterior, el proceso de anonimización se aplica a aquellos datos que por su naturaleza son sensibles al público debido a la posibilidad de que sea violada su confidencialidad.

El proceso de anonimización puede aplicarse tanto a los microdatos obtenidos de **operaciones estadísticas como a los de los registros administrativos** que posee una entidad.

Respecto al primer tipo de datos, el Decreto 1743 de 2016 ha definido operación estadística como la “aplicación de un proceso estadístico sobre un objeto de estudio que conduce a la producción de información estadística” (Decreto 1743 de 2016: Art. 2.2.3.1.1).

El **proceso estadístico** se entiende como un:

“Conjunto sistemático de actividades encaminadas a la producción de estadísticas, entre las cuales están comprendidas: la detección de necesidades de información, el diseño, la construcción, la recolección; el procesamiento, el análisis, la difusión y la evaluación” (Artículo 5. Ley Estadística 2335 de 2023).

Por tanto, el proceso de anonimización que se aplique a las operaciones estadísticas debe contar con una metodología y una documentación a lo largo de las fases de producción para garantizar la calidad de la información estadística a generar.

La información que se obtiene del proceso estadístico a nivel de los datos de

las unidades de observación (hogares, personas, empresas) sirve como insumo para que los usuarios puedan aprovechar estadísticamente la información de dichas operaciones, así como para generar estudios o investigaciones propias que permitan mejorar la toma de decisiones, además de otros usos que consideren pertinentes los responsables de la información a anonimizar.

Es así como la anonimización fomenta en las entidades del SEN la transparencia de la información y priorizar, en este caso, la preservación de la confidencialidad.

Frente a los **registros administrativos**, estos se entienden como un:

“Conjunto de datos que contiene la información recogida y conservada por entidades y organizaciones en el cumplimiento de sus funciones o competencias misionales u objetos sociales. De igual forma, se consideran registros administrativos las bases de datos con identificadores únicos asociados a números de identificación personal números de identificación tributaria u otros, los datos geográficos que permitan identificar o ubicar espacialmente los datos, así como los listados de unidades y transacciones administrados por los integrantes del SEN” (Artículo 5. Ley Estadística 2335 de 2023).

Para esta fuente en particular, las entidades del SEN pueden contar con procesos de diagnósticos que permiten implementar prácticas para mejorar la captura de la información, mediante la implementación de indicadores de cali-

dad sobre el registro de la entidad¹⁰. Los resultados de los diagnósticos de los registros administrativos permitirán al mismo tiempo, desarrollar mejores procesos de anonimización y obtener mejoras en la publicación de su información estadística.

Los **datos estructurados** se entienden como:

“(...) datos que tienen un modelo de datos y formato predefinido y que se ajustan a una forma de tablas, de registros o filas con campos de significados fijos y relaciones o enlaces entre las tablas” (Departamento Administrativo Nacional de Estadística (DANE), Acuerdo equipo de trabajo DIRPEN).

Los datos estructurados siguen una estructura predeterminada, por ejemplo: bases de datos SQL, archivos de Excel, resultados de formularios web, registros administrativos, etc.

Los **datos no estructurados** se entienden como:

“(...) datos que no tienen un modelo de datos predefinido y no están organizados de manera predefinida o su estructura no se ajusta perfectamente a una tabla de datos relacional” (Departamento Administrativo Nacional de Estadística (DANE)).

Los datos no estructurados pueden presentarse en formato de texto, imagen, sonido, video u otros formatos, por ejemplo: correos electrónicos, archivos PDF, publicaciones en redes sociales, imágenes digitales, archivos de audio, archivos de video.

A continuación, se relacionan una serie de conceptos que permitirán al usuario un mayor entendimiento de la información:

Datos abiertos se entiende como:

“(...) aquellos datos primarios o sin procesar, que se encuentran en formatos estándar e interoperables que facilitan su acceso

¹⁰ Para mayor información sobre los procesos de diagnóstico de los registros administrativos, se puede consultar la Metodología de diagnóstico de los Registros administrativos para su Aprovechamiento estadístico del DANE. Disponible en: <http://www.dane.gov.co/files/sen/registros-administrativos/Metodologia-de-Diagnostico.pdf>

y reutilización, los cuales están bajo la custodia de las entidades públicas o privadas que cumplen con funciones públicas y que son puestos a disposición de cualquier ciudadano, de forma libre y sin restricciones, con el fin de que terceros puedan reutilizarlos y crear servicios derivados de los mismos” (Ley 1712 de 2014: Ley de Transparencia y del derecho de acceso a la información pública nacional).

Dato semiprivado es:

“(…) aquel que además de ser de interés para el titular, puede ser de interés para cierto sector o grupo de personas. Ejemplo: dirección de residencia y teléfono” (Registraduría Nacional del Estado Civil).

Dato personal se define como:

“Cualquier información vinculada o que pueda asociarse a una o varias personas naturales determinadas o determinables” (Ley 1581 de 2012, Congreso de Colombia).

Se entiende por **dato sensible**:

“(…) a aquellos que afectan la intimidad del titular o cuyo uso indebido puede generar su discriminación, tales como aquellos que revelen el origen racial o étnico, la orientación política, las convicciones religiosas o filosóficas, la pertenencia a sindicatos, organizaciones sociales, de derechos humanos o que promueva intereses de cualquier partido político o que garanticen los derechos y garantías de partidos políticos de oposición así como los datos relativos a la salud, a la vida sexual y los datos biométricos” (Ley 1581 de 2012, Congreso de Colombia).

Dato público se define como:

“(…) aquel que puede ser consultado por cualquier persona de manera directa, sin el consentimiento del titular. Ejemplo: número de identificación, apellidos, lugar y fecha de expedición del documento” (Registraduría Nacional del Estado Civil).

Se entiende **operación estadística** como:

“(…) la aplicación del conjunto de procesos y actividades que comprende la identificación de necesidades, diseño, construcción, recolección o acopio, procesamiento, análisis, difusión y evaluación, la cual conduce a la producción de información estadística sobre un tema de interés nacional o territorial” (Artículo 5. Ley Estadística 2335 de 2023).



4. Planeación del proceso de anonimización

Este capítulo tiene la finalidad de orientar sobre como delimitar el alcance de la anonimización que se pretende realizar, así como, establecer los recursos a utilizar en el proceso. Por ello se estableció la revisión normativa sobre protección de datos e identificación de necesidades de información y que implicó una exhaustiva revisión de restricciones normativas que puedan afectar la publicación de datos a anonimizar, considerando leyes, acuerdos de confidencialidad y normativas pertinentes.

Posteriormente, se realizó una revisión de las demandas de información de los usuarios, analizando solicitudes, derechos de petición y otras necesidades. Esta revisión permite clasificar y cuantificar las demandas e identificar las variables y los periodos más solicitados. Con base en estos resultados, el equipo de trabajo ejecuta el proceso de anonimización, tomando en cuenta requisitos previos como la disponibilidad de la base de datos, un diccionario de datos, una infraestructura tecnológica adecuada y mecanismos de seguridad para proteger la información.

4.1. Revisión normativa sobre protección de datos e identificación de necesidades de información

Este subproceso busca realizar una revisión normativa que pueda afectar la publicación de la información sujeta a anonimizar; asimismo, se sugiere identificar las necesidades de información que presentan los usuarios sobre la base de datos. El subproceso se compone de dos pasos:

1. Revisión de restricciones de publicación de la información.
2. Revisión de las necesidades de los usuarios de la información.

4.1.1. Revisión de restricciones de publicación de la información

El equipo de trabajo debe realizar una revisión de la normatividad que puede afectar la publicación de la información sujeta a ser anonimizada. En este caso es importante que la entidad verifique las normas y las cláusulas de confidencialidad que tengan

alcance en la información estadística que se espera dejar disponible para el SEN.

Para iniciar esta actividad, el equipo de trabajo tendrá que realizar la revisión de leyes, decretos, resoluciones, convenios institucionales, acuerdos de confidencialidad de la información, estatutos y toda la normatividad que fundamenta el origen del registro administrativo o de la operación estadística de la entidad. El resultado de este análisis determinará si existe alguna norma que impida la publicación de la información de la base de datos a anonimizar.

Un insumo a tener en cuenta en esta actividad para el equipo de trabajo es el alcance de las normas asociadas a la confidencialidad de la información que tiene impacto sobre las entidades del SEN. Por ejemplo, el principio de confidencialidad descrito en la Ley Estatutaria 1266 de 2008 o el artículo 2.2.3.3.5 del Decreto 1743 de 2016, donde se establecen indicaciones sobre la no exposición de la identificación y la ubicación de las unidades de observación al momento de publicar información (Recuadro 1).

Recuadro 1. Normas referentes a la confidencialidad de la información

Ley Estatutaria 1266 de 2008: el principio de la confidencialidad descrito en la ley indica que “Todas las personas naturales o jurídicas que intervengan en la administración de datos personales que no tengan la naturaleza de públicos están obligadas en todo tiempo a garantizar la reserva de la información”.

El artículo 2.2.3.3.5 del decreto 1743 de 2016 establece que “las entidades que conforman el SEN ..., deberán guardar la confidencialidad de los datos que permi-

tan la identificación y/o localización espacial de las fuentes, cuando estos fueren recolectados exclusivamente para la producción de las estadísticas oficiales y para fines estadísticos”.

Ley Estadística 2335 de 2023, Capítulo V establece que quienes producen estadísticas oficiales deberán proteger los datos confidenciales de tal forma que, en la publicación de resultados estadísticos, la unidad estadística no pueda ser identificada, directa ni indirectamente, además, deberán tomar todas las medidas regulatorias, administrativas, técnicas y organizativas necesarias para evitar el acceso de personas no autorizadas a los datos individuales.

Circular 13 del 29 de abril de 2020, Oficina Asesora Jurídica del DANE - FONDANE, “CRITERIOS JURÍDICOS DE CLASIFICACIÓN DE LA INFORMACIÓN DEL DANE - FONDANE” que consolida de manera clara y precisa cómo se debe clasificar la información teniendo como presupuesto lo dispuesto en la normatividad vigente respecto a información pública, información reservada y clasificada y protección de datos personales, así como los lineamientos de la Guía No. 5 para la Gestión y la Clasificación de Activos de Información de Seguridad y Privacidad de la Información del Ministerio de Tecnologías de la Información y Telecomunicaciones en Colombia.

Después de la revisión normativa y de tener en cuenta los principios de confidencialidad, el equipo de trabajo describirá los hallazgos de la correspondiente revisión. Esta descripción es necesario incluirla en el informe final de anonimización y puede relacionar aspectos como:

- **Normatividad identificada** que impida la publicación de la información, como decretos, leyes, artículos, entre otros.
- **Normas que respalden la publicación de la información** a los diferentes usuarios y demás entidades del SEN.
- **Observaciones y sugerencias** a tener en cuenta al momento de realizar la publicación de la información, para evitar sanciones legales y judiciales a la entidad del SEN.
- **Personal y fecha en la que se hizo la revisión de la normatividad** para presentar una trazabilidad del trabajo sobre la base de datos a anonimizar.

4.1.2. Revisión de las necesidades de los usuarios de la información

Este paso tiene como propósito realizar una caracterización de las necesidades de información de los usuarios sobre la base de datos a anonimizar. En este caso, el equipo de trabajo de la entidad del SEN debe tener en cuenta las demandas o los requerimientos realizados por parte de los usuarios.

Para el caso en que la entidad del SEN cuente con demandas de información por distintos usuarios, se recomienda elaborar un listado que contenga los requerimientos solicitados, las fechas

de solicitud, las variables requeridas, la frecuencia con la que se hacen las solicitudes y las respuestas dadas a dichos requerimientos.

Se recomienda al equipo de trabajo realizar una revisión del histórico de las solicitudes, mayor a un año y menor a dos años, teniendo en cuenta el volumen de las solicitudes recibidas, la recurrencia y la frecuencia que tiene cada solicitud. Algunas de las solicitudes que podrá revisar el equipo de trabajo se presentan a continuación:

- Solicitudes de información recibidas por distintos usuarios: ciudadanos, academia, empresas, entidades públicas del orden nacional o territorial, organismos internacionales, entre otros.
- Derechos de petición radicados en la entidad durante el último año.
- Necesidades de información propias del área temática de la entidad.

Con la información recolectada, el equipo de trabajo podrá clasificar y cuantificar las demandas de información, teniendo en cuenta:

- Tipo de usuarios.
- Tipo de solicitudes.
- Variables solicitadas.
- Nivel de desagregación requerido de la información.
- Periodos de la información (anual, semestral, trimestral, etc.).
- Frecuencia de las solicitudes
- Objetivo, finalidad, uso de la información requerida.

El Recuadro 2 presenta un ejemplo de clasificación de este tipo de solicitudes de información.

Recuadro 2. Ejemplo de clasificación de solicitudes recibidas.

Referente	Descripción	N° de solicitudes
Tipo de usuario	Entidades públicas	4
	Entidades privadas	3
	Investigadores	20
	Instituciones académicas	15
	Otros	8
Tipo de solicitudes	Derechos de petición	1
	Acceso a información	30
	Actualización de datos	19
Clasificación de variables	Variables de clasificación	15
	Variables de ubicación	10
	Variables temáticas	25
Niveles de desagregación	Nacional	3
	Departamental	10
	Municipal	30
	Temática	5
	Otros	2
Periodos de información solicitados	Anual	16
	Mensual	30
	Trimestral	20
	Diaria	23

Referente	Descripción	N° de solicitudes
Frecuencia de la solicitud	Mensual	120
	Diaria	26
Objetivo, finalidad, uso	Investigaciones o estudios	36
	Académica	12
	Política pública	59

Fuente: DANE/DIRPEN.

Basados en la clasificación y la cuantificación de las solicitudes, el equipo de trabajo podrá identificar las variables, los niveles de desagregación, los periodos de la información de la base de datos que tienen mayor demanda lo que le dará información para definir la base de datos y desagregaciones que satisficiera de la mejor forma las necesidades de los usuarios y que igualmente no expongan la identificación de las unidades de observación.

Una ventaja de publicar las bases anonimizadas, para la entidad del SEN, es la disminución de las cargas operativas para responder solicitudes de información que pueden ser repetitivas o recurrentes. Teniendo en cuenta la revisión de los diferentes recursos para encontrar un uso potencial, el equipo de trabajo de la entidad del SEN tomará la decisión del tipo de información que puede suministrar en la base de datos anonimizada para responder a las necesidades de los usuarios.

Después de la revisión de las necesidades de los usuarios de la información, el equipo de trabajo, mediante una bitácora, indicará los resultados encontrados:

- Periodo de revisión del equipo de trabajo de las solicitudes, las peticiones y los requerimientos de información.
- Listado de variables más solicitadas y los niveles de desagregación más demandados.
- Listado de los usuarios que con mayor frecuencia hacen solicitudes de información.
- Los periodos o las bases de datos con determinados cortes que más son solicitados.
- Listado de los diferentes usos identificados para los cuales son radicadas la mayoría de las solicitudes y los requerimientos.

El proceso de anonimización debe ser ejecutado por un equipo de trabajo que tenga acceso y conocimiento de la base de datos a anonimizar. Es importante que este equipo de trabajo cuente con la capacidad de:

- **Conocer temáticamente el contenido de la base de datos a anonimizar:** por ejemplo, si la base es insumo de una operación estadística de un área temática particular es importante que el equipo de trabajo se encuentre conformado por una persona o varias que conozcan el fenómeno registrado en la base de datos. La vinculación de este personal especializado tiene como fin apoyar la toma de decisiones de la anonimización a desarrollar, principalmente en las etapas de riesgos de identificación (Etapa II) y de análisis de viabilidad del proceso (Etapa IV).
- **Manejar herramientas que permitan el análisis exploratorio de datos:** en este caso se necesitará que el equipo cuente con miembros con habilidades para el manejo de paquetes estadísticos como *Python*, *R*, *SAS*, *SPSS*, *Stata*, entre otros. Igualmente, se recomienda que conozcan técnicas estadísticas que permitan apoyar la etapa de análisis exploratorio de la base de datos (subproceso de la Etapa I) y la etapa de aplicación de las técnicas de anonimización (Etapa V).

Cuando la entidad conforme estratégicamente su equipo de trabajo, este deberá tener en cuenta que existen requerimientos iniciales que permiten la ejecución satisfactoria del proceso de anonimización. Estos son:

- **Disponer de una base de datos:** la base definida por la entidad y que será difundida para el SEN. Esta puede contener una tabla o varias tablas relacionadas entre sí¹¹. La entidad debe saber si la base de datos es re-

sultado de una operación estadística o es la resultante de un registro administrativo.

- **Contar con el diccionario de datos de la base a anonimizar¹²:** este documento debe especificar claramente las propiedades básicas de las variables contenidas en la base de datos. Algunas de estas propiedades son: nombre, longitud, obligatoriedad de respuesta, descripción, reglas de validación, entre otros, así como la relación entre ellas.
- **Disponer de una infraestructura tecnológica:** la infraestructura definida por la entidad debe permitir el manejo estadístico de datos; en este caso, debe tener en cuenta paquetes estadísticos, equipos de cómputo que permitan el manejo de datos y en general tecnología que se encuentre acorde con el volumen de información a anonimizar. Cuando las bases de datos son de grandes dimensiones, el equipo de trabajo debe tener en cuenta que algunos paquetes de software no son compatibles, por lo cual requerirá revisiones sobre programas estadísticos y el alcance de estos sobre grandes cantidades de información.
- **Definir mecanismos de seguridad sobre la base de datos a anonimizar:** la entidad debe prever condiciones mínimas para salvaguardar la información, así como definir protocolos de acceso a la información por parte del equipo de trabajo que participará en el proceso de anonimización. Por ejemplo, acuerdos de confidencialidad del personal involucrado en el proceso (equipo de trabajo) debidamente firmados, usos de permisos y contraseñas para el uso de la información, entre otros.

¹¹ Modelo Entidad-Relación de la base de datos.

¹² Diccionario ejemplo dispuesto, disponible en: https://www.sen.gov.co/files/sen/lineamientos/Guía_Diccionario_de_Datos.xlsx.

Recuadro 3. Verificación especial de la base de datos de acuerdo a su origen

La entidad del SEN debe verificar si la base de datos a anonimizar es la resultante de alguna operación estadística o de un registro administrativo.

Caso 1: Resultado de operaciones estadísticas

La entidad tendrá que verificar que la base de datos para la anonimización es la resultante de la finalización de la Fase de Ejecución del Proceso Estadístico. Esto significa que la consistencia de los datos recolectados debe haber sido validada, según: i) las técnicas definidas por la entidad responsable; ii) las variables creadas (o calculadas) a partir de la información recolectada, y iii) que se encuentren en la base de datos. En caso de valores faltantes, tendrá que verificar que los métodos de imputación hayan sido aplicados.

Caso 2: Base de datos de registros administrativos

La entidad debe verificar que la base de datos sea la versión más actual. En este caso, se recomienda que el periodo de tiempo del registro administrativo que se anonimizará se defina con base en la periodicidad de consolidación de la información. Además, se recomienda revisar la consistencia y la calidad de la base de datos teniendo en cuenta la sección Revisión de la con-

sistencia de la base de datos de la Metodología de diagnóstico de registros administrativos.

Ejemplo:

El Sistema de Identificación de Potenciales Beneficiarios de Programas Sociales (SISBEN) tiene una periodicidad de recolección diaria y de consolidación de la información en el aplicativo SISBENNET mensual. En este caso, se recomendaría que SISBEN realice un corte para anonimizar la base de datos de un periodo, teniendo como fecha de cierre la última consolidación de la información.

Fuente: elaboración propia.

Después de que el equipo de trabajo confirme la disponibilidad de los requerimientos previos, empezará con la primera etapa del proceso de anonimización.

Insumos:

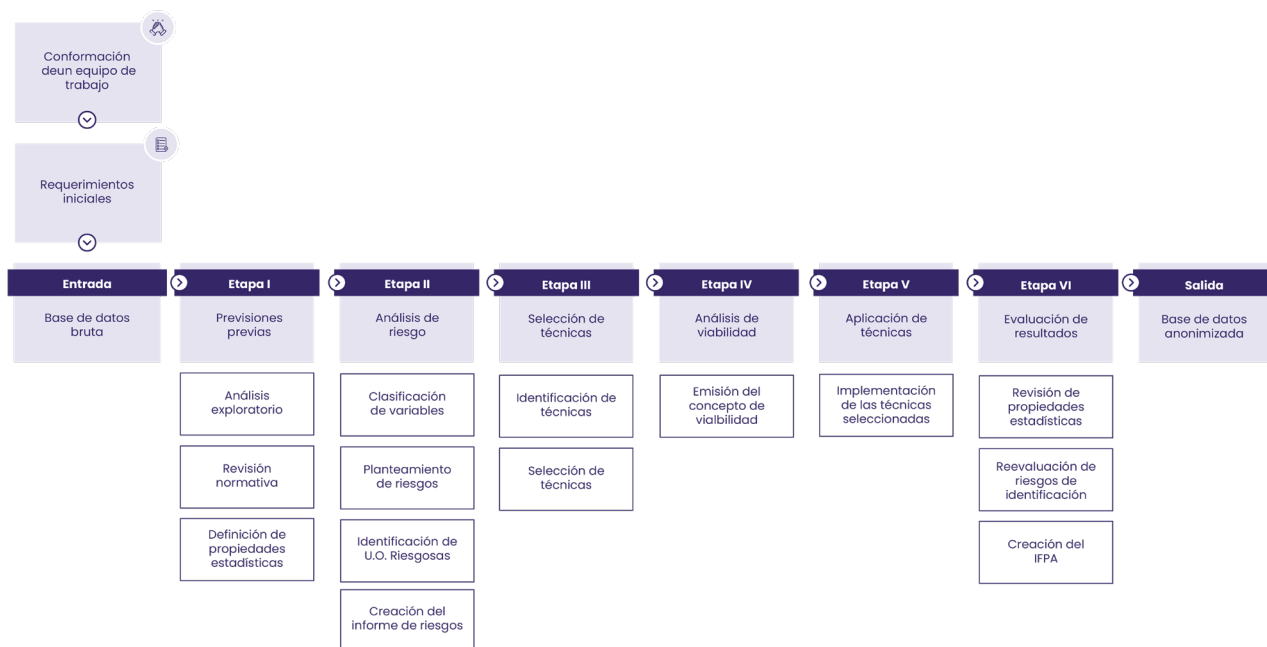
- Equipo de trabajo definido.
- Base de datos.
- Dirección de datos completo
- Infraestructura definida

5. Ejecución del proceso de anonimización

Este capítulo contiene las etapas en que se desarrollan el proceso de anonimización en datos estructurados. El proceso de anonimización de una base de datos se encuentra compuesto por las siguientes etapas: i) Revisiones previas; ii) Análisis de riesgos de identificación de

las fuentes de información; iii) Identificación y selección de técnicas de anonimización; iv) Análisis de viabilidad del proceso; v) Aplicación de técnicas de anonimización, y vi) Evaluación de resultados del proceso. Estas fases se presentan en la Ilustración 1:

Ilustración 1. Esquema general del proceso de anonimización



Fuente: DANE- DIRPEN.

Estas etapas presentan distintos subprocesos y actividades a desarrollar para realizar la anonimización de la base de datos deseada.

Debido a que es un proceso de anonimización, los contenidos de las fases se encuentran orientados a presentar entradas y salidas, así como recomendaciones sobre pasos previos que permitirán preparar a la entidad al iniciar el proceso de anonimización.

Igualmente, a lo largo de las etapas se realizan indicaciones y recomendaciones de documentación que permitan conocer la trazabilidad del proceso de anonimización dentro de la entidad. Al finalizar la etapa 6 del proceso, podrá encontrarse un modelo de Informe Final del Proceso de Anonimización (IFPA), el cual recoge todas las recomendaciones de documentación cuando se desarrolla la anonimización de la base de datos.

Etapas I. Análisis exploratorio de los datos

En esta etapa se deben caracterizar los datos a anonimizar, teniendo en cuenta aspectos procedimentales y temáticos. Está compuesto de cinco pasos:

A. Caracterización de la base de datos

Para iniciar con la etapa, el equipo de trabajo caracterizará cada una de las variables contenidas en la base de datos teniendo en cuenta si son cuantitativas (continuas o discretas) o categóricas (4).

Recuadro 4. Clasificación de las variables por su tipo

Las variables se pueden clasificar por el tipo de valores que toman:

- **Variable continua:** es una característica medida en las unidades de observación que puede tomar valores dentro de un intervalo específico (infinito) (Sistema de Consulta de Conceptos del DANE, 2018).

Por ejemplo: la altura medida en centímetros, el peso medido en kilogramos, la temperatura medida en grados centígrados, entre otros.

- **Variable discreta:** es una característica medida en las unidades de observación donde su conjunto de posibles valores no corresponde a un intervalo ya que presenta interrupciones (Sistema de Consulta de Conceptos del DANE, 2018).

Por ejemplo: número de hijos, número de intentos en un experimento, número de personas de un hogar, entre otros.

- **Variable categórica:** es una característica medida en las unidades de observación que asigna una de varias categorías cualitativas (Sistema de Consulta de Conceptos del DANE, 2018).

Por ejemplo: una variable dicotómica, que asigna el valor de 1 si el individuo cumple cierta característica y 0 si no la cumple, o estado civil, que asigna a cada individuo alguna de las categorías: soltero, casado, separado, divorciado o unión libre.

Para desarrollar este subproceso, el equipo de trabajo, con ayuda del diccionario de la base de datos (o de toda la documentación disponible), describirá las dimensiones de la base en términos del número de variables, el número de registros y la distribución de todas las variables con respecto a su tipo (cuantitativa o categórica).

Otra clasificación que el equipo de trabajo debe tener en cuenta para las variables de la base de datos es el tipo de información que contiene de la unidad

de observación. Estas se clasifican en variables de identificación, de ubicación y temáticas. Por ejemplo: el número de identificación de una persona es una variable de identificación; el municipio; la localidad y el departamento de una entidad, son variables de ubicación, y el ingreso promedio mensual, el nivel de escolaridad y el número de hijos, son variables temáticas.

La caracterización de la base de datos se puede hacer teniendo en cuenta como ejemplo la Tabla 2.

Tabla 2. Clasificación de variables de base de datos

Variables/tipo de variable	Cuantitativas	Categóricas	Total
Identificación	Escriba en este espacio el número de variables cuantitativas de identificación		
Ubicación		Escriba en este espacio el número de variables categóricas de ubicación.	
Temáticas			Escriba en este espacio el número de variables temáticas.
Total	Escriba en este espacio el número de variables cuantitativas.	Escriba en este espacio el número de variables categóricas.	Escriba en este espacio el total de variables de la base de datos.

Fuente: DANE- DIRPEN.

B. Cálculo de medidas descriptivas de las variables cuantitativas

Después de que el equipo realice la clasificación de las variables de la base de

datos, procederá a calcular las medidas descriptivas estadísticas para cada una de las variables cuantitativas. Un ejemplo de estas se presenta a continuación:

Tabla 3. Medidas descriptivas principales para variables cuantitativas

Variable cuantitativa	Media	Varianza	Cuartil 1	Cuartil 2	Cuartil 3
Escriba en este espacio el nombre de la variable cuantitativa.			Escriba en este espacio el primer cuartil de los datos de la variable.		

Fuente: DANE- DIRPEN.

Estas medidas servirán como insumo para el análisis de riesgos (Etapa II), el análisis de viabilidad del proceso de anonimización (Etapa IV) y la evaluación de resultados (Etapa VI). Este tipo de medidas se conocen como propiedades globales de las variables y estarán sujetas a verificación por parte del equipo de trabajo para examinar el resultado del proceso de anonimización.

C. Cálculo de frecuencias de las variables categóricas

Las distribuciones de frecuencias para las variables categóricas hacen parte de las medidas para verificar el proceso de anonimización de la base de datos. Estas distribuciones pueden realizarse teniendo en cuenta la siguiente tabla:

Tabla 4. Distribución de frecuencias para una variable con dos categorías

Variable categórica	Número de unidades de observación que cumplen la categoría	Porcentaje de registros que cumplen la categoría
Categoría 1	Escriba en este espacio el número de unidades de observación que cumplen con la categoría 1.	Escriba en este espacio el porcentaje de unidades de observación que cumplen con la categoría 1.

Variable categórica	Número de unidades de observación que cumplen la categoría	Porcentaje de registros que cumplen la categoría
Categoría 2	Escriba en este espacio el número de unidades de observación que cumplen con la categoría 2.	Escriba en este espacio el porcentaje de unidades de observación que cumplen con la categoría 2.
Total	Escriba en este espacio el número total de unidades de observación que cumplen con las categorías 1 y 2.	Escriba en este espacio el porcentaje de unidades de observación que cumplen con la categoría 2.

Fuente: DANE- DIRPEN.

D. Revisión temática del contenido de la base de datos

El equipo de trabajo, después de analizar cada una de las variables, debe realizar una revisión temática de la base de datos, teniendo en cuenta la documentación de la operación estadística o el registro administrativo. Para ello se sugiere responder las siguientes preguntas:

- ¿Cuál es el objetivo de la operación estadística o del registro administrativo?
- ¿Qué cambios metodológicos ha presentado la operación estadística (o registro administrativo) en los últimos tres años?
- ¿Qué tipo de estándares, clasificaciones o nomenclaturas siguen las variables de la base de datos obtenida por la operación estadística o por el registro administrativo?
- ¿Se están usando apropiadamente los estándares, las clasificaciones o las nomenclaturas?

- ¿Qué documentos existen acerca de la operación estadística (o registro administrativo)? ¿Son de fácil acceso?
- ¿Existen otras operaciones estadísticas (o registros administrativos) relacionadas con la temática de la base de datos?

Esta revisión temática servirá como insumo en el planteamiento de los riesgos de identificación de las unidades de observación (Etapa II) y en el análisis de viabilidad del proceso (Etapa IV). Con esto, finaliza el subproceso de Análisis exploratorio de la base de datos.

E. Definición de las propiedades estadísticas a conservar en la base de datos

En este subproceso el equipo de trabajo deberá establecer las propiedades estadísticas que se deben mantener en la base de datos anonimizada, en relación con la base de datos sin anonimizar. Algunas de estas propiedades estadísticas son:

- **Mantener tendencias en las variables a través del tiempo:** esta propiedad hace referencia a que las

variables conserven un comportamiento en determinados periodos de tiempo.

Por ejemplo, si la base de datos de una operación estadística de temática económica contiene la variable ingreso de los hogares colombianos y esta variable ha presentado un comportamiento creciente en el primer trimestre de 2016, al publicar la base de datos anonimizada el equipo de trabajo desea garantizar que esta tendencia se conserve.

- **Mantener propiedades globales de las variables:** el equipo de trabajo debe definir cuáles de las medidas estadísticas descritas en el análisis exploratorio de datos (sección 5.1), para las variables categóricas y cuantitativas, se deben mantener sin variación y para qué niveles de desagregación geográfica o temática. Asimismo, debe decidir cuáles de las propiedades globales pueden presentar alguna variación significativa y hasta qué porcentaje de variación es permitido en la base de datos anonimizada.

Por ejemplo, un equipo decidió que la propiedad global que desea mantener es el promedio de la variable “Ingreso por hogar”. Además, aceptará el proceso de anonimización, solamente si el promedio de la variable en la base de datos anonimizada difiere del promedio en la base de datos sin anonimizar en menos del 1%.

- **Mantener cifras por niveles de desagregación geográfica o temática:** el equipo de trabajo debe definir cuáles medidas estadísticas se deben conservar sin variación en los niveles de desagregación geográfica o temática, para garantizar a los usuarios análisis de estadísticas más sectorizados.

Por ejemplo, un equipo de trabajo decidió mantener para la variable grupos étnicos los totales de cada categoría a nivel departamental; esta propiedad permite caracterizar la población étnica

en cada departamento y con esta información los usuarios pueden realizar análisis estadístico por regiones.

- **Mantener correlaciones entre variables:** esta propiedad busca conservar las posibles relaciones lineales o no lineales que se puedan presentar entre las variables. El equipo de trabajo definirá si mantiene los coeficientes de correlación entre dos o más variables (cuantitativas o categóricas) en la base de datos anonimizada, con el fin de no distorsionar los resultados finales.

Por ejemplo, un equipo de trabajo decidió mantener las correlaciones entre estrato y avalúo comercial de un predio, esto para conservar la integridad de la información ya que son variables importantes para analizar.

Finalmente, en este subproceso el equipo de trabajo definirá las propiedades estadísticas que deberá tener la base de datos anonimizada, dado que son insumo para evaluar el proceso de anonimización. Además, describirá las propiedades estadísticas a conservar en la base de datos, teniendo en cuenta:

- Descripción de las propiedades estadísticas elegidas por el grupo de trabajo para conservar en la base de datos anonimizada.
- Listado de las variables con la respectiva propiedad estadística a conservar en la base de datos anonimizada.
- Nivel de desagregación geográfica o temática en el que se desea conservar las propiedades estadísticas.
- Porcentajes de variación permitidos por variables y niveles de desagregación geográfica o temática para las propiedades globales en la base de datos anonimizada.

Para visualizar estos elementos de los subprocesos de la Etapa I, se presenta a continuación un ejemplo.

Ejemplo de la Etapa I. Análisis exploratorio de la base de datos

La base de datos anonimizada de la Encuesta Anual de Comercio (EAC) de

2016¹³, cuenta con 64 variables y 10.242 unidades de observación, en este caso empresas. Las variables se pueden caracterizar de la siguiente forma:

Tabla 5. Clasificación por tipo de variable de la EAC en el 2016

Variables/tipo de variable	Cuantitativas	Categóricas	Total
IDENTIFICACIÓN	0	2	2
UBICACIÓN	0	0	0
TEMÁTICAS	59	3	62
TOTAL	59	5	64

Fuente: DANE - EAC, Cálculos DANE - DIRPEN.

¹³ Disponible en el Archivo Nacional de Datos.

Además, las medidas descriptivas para cinco variables cuantitativas temáticas de la operación estadística son:

Tabla 6. Medidas descriptivas de algunas variables cuantitativas de la EAC en el 2016

Variable	Media	Varianza	Cuartil 1	Cuartil 2	Cuartil 3
TOTAL SUELDOS*	973.902	2.76E+13	137.877	274.157	629.193
TOTAL PRESTACIONES*	491.284	9.30E+12	58.905	118.604	280.252
VALOR AGREGADO*	3.220.106	2.40E+14	360.691	834.109	2.058.826
VENTAS CAUSADAS*	24.022.327	1.89E+16	2.789.874	6.335.314	14.654.517
PERSONAL REMUNERADO**	52.5	114575.7999	12	19	38

*Cifras en millones de pesos

**Cifra en número de personas

Fuente: DANE- EAC, Cálculos DANE – DIRPEN.

Además, se presenta la distribución de frecuencias de la variable categórica Organización Jurídica:

Tabla 7. Distribución de Frecuencias de la variable Organización Jurídica de la EAC

Organización jurídica	Número de registros por categoría	% de registros por categoría
Sociedad en comandita simple	111	1,08%
Sociedad en comandita por acciones	37	0,36%
Sociedad limitada	1,676	16,36%
Sociedad anónima	1,567	15,3%
Sucursal de sociedad extranjera	30	0,29%

Organización jurídica	Número de registros por categoría	% de registros por categoría
Empresa unipersonal	63	0,62%
Persona natural	1,946	19%

Organización jurídica	Número de registros por categoría	% de registros por categoría
Organizaciones de economía solidaria	82	0,8%
Entidades sin ánimo de lucro	50	0,49%
Sociedad por acciones simplificada	4,652	45,42%
Otro	28	0,27%
Total	10,242	100%

Fuente: DANE- EAC, Cálculos DANE - DIRPEN

Finalmente, la revisión temática del contenido de la base de datos permite concluir que:

- El objetivo de la EAC es conocer la estructura y el comportamiento económico del sector comercio a nivel nacional y por grupo de actividad comercial, de manera que permita el análisis de la evolución del sector y de la conformación de agregados económicos.
- En la página web del DANE¹⁴ se encuentran disponibles todos los metadatos y los microdatos de la EAC. Los

documentos contienen información sobre:

- Recolección de los datos.
- Procesamiento de la información.
- Políticas de acceso a los microdatos.
- Diccionario de datos.
- Descripción de cada una de las variables (incluidos los estándares utilizados).
- Referentes internacionales.
- Cuestionario.
- Metodología y ficha metodológica.

¹⁴ Disponible en: <https://www.dane.gov.co/index.php/estadisticas-por-tema/comercio-interno/encuesta-anual-de-comercio-eac>

Revisión de restricciones de publicación de la información

- Ley 2ª de 1962, que permite el levantamiento de encuestas nacionales y en especial, las de industria, comercio y servicios.
- El Decreto 1633 de 1960 en su artículo 74 establece que “todas las personas naturales o jurídicas, domiciliadas en el territorio nacional y los empleados públicos, en todos sus niveles, están obligados a suministrar información al DANE, dentro de los plazos que al efecto se señalen, los datos que este requiera para el cumplimiento de sus finalidades”. Además, establece que los datos suministrados a la entidad tienen un carácter estrictamente reservado y no podrán darse a conocer al público ni a las entidades oficiales, sino únicamente en resúmenes numéricos.
- La Ley 0079 de octubre 20 de 1993, que regula la realización de los censos y las encuestas, decreta que las personas naturales o jurídicas, de cualquier orden o naturaleza, domiciliadas o residentes en el territorio nacional, están obligadas a suministrar información al DANE. Asimismo, especifica que la entidad podrá imponer multas a quienes incumplan esta disposición.
- El Principio de confidencialidad descrito en la Ley estatutaria 1266 de 2008 en su artículo 4 indica: “Todas las personas naturales o jurídicas que intervengan en la administración de datos personales que no tengan la naturaleza de públicos están obligadas en todo tiempo a garantizar la reserva de la información. Inclusive después de finalizada su relación con alguna de las labores que comprende la administración de datos”.
- El artículo 2.2.3.3.5 del Decreto 1743 de 2016 indica: “(...) guardar la confidencialidad de los datos que permitan la identificación y/o localización espacial de las fuentes, cuando estos fueren recolectados exclusivamente para la producción de las estadísticas oficiales y para fines estadísticos”.
- Ley 79 de 1993 Artículo 5: “Los datos suministrados al DANE, en desarrollo de censos y encuestas, no podrán darse a conocer al público ni a entidades u organismos oficiales, ni a las autoridades públicas, sino únicamente en resúmenes numéricos, que no hagan posible deducir de ellos información alguna de carácter individual que pudiera utilizarse para fines comerciales, de tributación fiscal, de investigación judicial o cualquier otro diferente del propiamente estadístico”.
- Ley Estadística 2335 de 2023 establece que quienes producen estadísticas oficiales deberán proteger los datos confidenciales de tal forma que, en la publicación de resultados estadísticos, la unidad estadística no pueda ser identificada, directa ni indirectamente. Además, deberán tomar todas las medidas regulatorias, administrativas, técnicas y organizativas necesarias para evitar el acceso de personas no autorizadas a los datos individuales.
- Es conveniente tener en cuenta además las excepciones descritas en la Ley 1712 de 2014: la Ley de Transparencia y del Derecho de Acceso a la Información Pública Nacional, la cual menciona en su título III las excepciones de acceso a la información donde se destacan:

Artículo 18. Información exceptuada por daño de derechos a personas naturales o jurídicas: es toda aquella información pública clasificada, cuyo acceso podrá ser rechazado o denegado de manera motivada y por escrito, siempre que el acceso pudiere causar un daño a los siguientes derechos (...).

Artículo 19. Información exceptuada por daño a los intereses públicos: es toda aquella información pública reservada, cuyo acceso podrá ser rechazado o denegado de manera motivada y por escrito en las siguientes circunstancias, siempre que dicho acceso estuviere expresamente prohibido por una norma legal o constitucional (...).

Artículo 20. Índice de Información clasificada y reservada: los sujetos obligados deberán mantener un índice actualizado de los actos, documentos e informaciones calificados como clasificados o reservados, de conformidad a esta ley. El índice incluirá sus denominaciones, la motivación y la individualización del acto en que conste tal calificación.

Artículo 21. Divulgación parcial y otras reglas: en aquellas circunstancias en que la totalidad de la información contenida en un documento no esté protegida por una excepción contenida en la presente ley, debe hacerse una versión pública que mantenga la reserva únicamente de la parte indispensable. (...).

Artículo 22. Excepciones temporales: la reserva de las informaciones amparadas por el artículo 19 no deberá extenderse por un período mayor a quince (15) años

Acorde con la revisión de esta normatividad, no se evidencia alguna norma que impida publicar la información siempre y cuando se proteja la identificación de las unidades de observación.

Revisión de las necesidades de los usuarios de la información

El grupo de trabajo, después de revisar las solicitudes de información en la entidad durante los últimos dos años, evidenció que:

- La mayoría de las solicitudes de información que llegan al DANE provienen de gremios, asociaciones, investigadores, académicos, centros de investigación, universidades.
- Las variables ingreso por ventas, gastos de personal (sueldos y prestaciones), costos de la mercancía vendida, gastos de operación, personal ocupado, inventarios y movimientos de activos fijos, son las que mayor demanda de información presentan.
- La información correspondiente al 2016 se encuentra de forma recurrente entre los periodos de tiempo solicitado, a través de las diferentes solicitudes de información que ha recibido la entidad.
- La información que los diferentes usuarios manifiestan en sus respectivas solicitudes corresponde a la información desagregada a nivel nacional.

Definición de las propiedades estadísticas a conservar en la base de datos

En la Encuesta Anual de Comercio (EAC) el equipo de trabajo decide que las propiedades estadísticas que la base de datos anonimizada debe conservar son:

- Para las variables ventas y personal ocupado **mantener la participación % directo de las divisiones CIIU Rev. 4. A.C. del sector comercio.**
- Para las variables número de empresas, ventas, costo de la mercancía, producción bruta, consumo intermedio, valor agregado, remuneración **mantener los totales a nivel nacional.**
- La información correspondiente al subsector vehículos automotores, motocicletas, sus partes, piezas y accesorios debe **mantener los totales a nivel nacional.**

Con la base de datos anonimizada los usuarios puedan replicar los diferentes cuadros de salida que son publicados en el boletín técnico¹⁵ de la EAC para 2016.

Etapas II: Identificación y medición de los riesgos presentes en datos

Esta etapa tiene como finalidad presentar los pasos a seguir para reconocer los riesgos de identificación de las fuentes en la base de datos que se pretende anonimizar, así como, proporcionar un procedimiento que permita medir los riesgos identificados. (Ejemplos basados en el censo de población + Unidades de Observación y Registros únicos que conlleven a singularizar).

Esta etapa se proyecta a desarrollar en términos de identificación y medición de riesgos en datos estructurados. En esta etapa el equipo de trabajo se planteará todos los posibles escenarios de riesgo de identificación de las unidades de observación de la base de datos.

Un escenario de riesgo de identificación de información confidencial es aquel en el cual existe una posibilidad de que, mediante la combinación de variables de la base de datos, se puedan identificar características de las unidades de observación que deben ser protegidas por el tipo de información que contiene.

La etapa de análisis de riesgos se compone de los siguientes subprocesos:

- Clasificación de variables por su nivel de sensibilidad.
- Planteamiento de riesgos de la base de datos
- Identificación de unidades de observación riesgosas.
- Creación del informe de riesgos.

Clasificación de variables por su nivel de sensibilidad

La etapa de análisis de riesgos inicia con la clasificación de las variables de la base de datos por su nivel de sensibilidad. El equipo debe realizar la revisión de los riesgos teniendo en cuenta el contexto de la base de datos y las personas encargadas del procesamiento de la base de datos.

¹⁵ https://www.dane.gov.co/files/investigaciones/boletines/eac/bol_eac_2016.pdf

Se debe tener en cuenta que una **variable se considera sensible** cuando es:

“aquella que puede afectar la intimidad, la honra, la reputación, la tranquilidad personal o la fama de las personas” (Ley 1266 de 2008).

En este punto los expertos temáticos que componen el equipo de trabajo juegan un rol preponderante.

Las variables pueden clasificarse en:

- **Identificadores directos:** las variables denominadas como identificadores directos son todas aquellas variables que contienen información sensible de identificación o ubicación de las unidades de observación. Estas variables pueden coincidir con las variables denominadas en el análisis exploratorio de la base de datos (Etapa I) como variables de identificación o de ubicación. Un ejemplo de este tipo de variables, es el número de cédula de una persona, NIT de una empresa, dirección de una entidad, entre otros.
- **Pseudoidentificadores:** las variables denominadas como pseudoidentificadores son todas aquellas que, combinadas con otras variables, conllevan a la identificación de las unidades de observación. Estas variables pueden coincidir comúnmente con las variables denominadas temáticas o de ubicación en el análisis exploratorio de la base de datos (Etapa I). Un ejemplo de este tipo de variables es la combinación del nivel de escolaridad con el ingreso promedio mensual en cierto municipio. En este caso, son tres variables consideradas como pseudoidentificadores que permiten la identificación de algunas unidades de observación.
- **No confidenciales:** las variables denominadas como no confidenciales son todas aquellas que no permiten la identificación de las unidades de observación de la base de datos, ni siquiera cuando son combinadas con pseudoidentificadores. Algunos ejemplos de este tipo de variables son: habilidades de conducción de un vehículo de una persona, gusto por el deporte, gustos culturales, entre otros.

El equipo de trabajo puede resumir la clasificación de acuerdo con la siguiente tabla:

Tabla 8. Tabla resumen de la clasificación de las variables por su tipo de sensibilidad

Tipo de variable / tipo de sensibilidad	Identificadores directos	Pseudoidentificadores	No confidenciales	Total
CUANTITATIVAS	Escriba en este espacio el número de variables cuantitativas clasificadas como identificadores directos.		Escriba en este espacio el número total de variables cuantitativas clasificadas como no confidenciales.	
CATEGÓRICAS		Escriba en este espacio el número total de variables categóricas clasificadas como pseudoidentificadores		
TOTAL				Escriba en este espacio el número total de variables.

Fuente: DANE - DIRPEN.

Planteamiento de riesgos de la base de datos

Después de la clasificación de todas las variables de la base de datos de acuerdo con su nivel de sensibilidad, el equipo de trabajo deberá establecer los riesgos de identificación de la base de datos que será anonimizada. Los riesgos que defina el equipo de trabajo servirán como insumo en la selección de técnicas de anonimización (Sección 5.3.).

En general, los riesgos se pueden entender como todas las posibles combinaciones de las variables (entre identificadores directos y pseudoidentificadores) y sus niveles de desagregación (geográfica o temática), que pueden aumentar la probabilidad de que una o varias unidades de observación sean identificadas por los usuarios de la información.

Es necesario que al realizar combinaciones entre variables, se identifiquen unidades de observación que deben ser protegidas dentro de sus niveles de desagregación geográfica o temática.

Por ejemplo:

- Si en la base de datos de un hospital de Cartagena, se tiene la información de los pacientes atendidos en 2016, al combinar la información de la edad, tipo de enfermedad, medicamentos recibidos y dirección, es posible identificar a las mujeres de más de 70 años que han tenido cáncer de seno para el barrio la Chinita; esta combinación de las variables permite reconocer a la unidad de observación.
- Al tener la información de los ingresos por hogar de los colombianos, al combinar las variables, ingresos anuales, edad, sexo, corregimiento

o vereda, y nivel escolar, podemos identificar que hay un hombre de 75 años con nivel educativo profesional en el corregimiento de Puerto Salazar, con ingresos anuales de \$ 35.601.804. Con esta combinación podría identificarse que se trata del único Médico del corregimiento.

- Al tener la información de ventas por catálogo de una empresa de arte, al combinar las variables departamento, sexo, ventas reportadas, artículos, se puede identificar que en Yopal hay un hombre y dos mujeres que poseen las ventas más altas de la región, por lo que se puede inferir que se traten de los líderes premium del departamento.

Una vez definidos los riesgos de la base de datos, el equipo de trabajo los organizará en un listado y posteriormente los priorizará teniendo en cuenta la frecuencia en que pueden ser identificadas las unidades de observación. Algunas recomendaciones sobre cómo definir los niveles de riesgos de identificación son:

- Verificar los identificadores que definitivamente deben ser suprimidos de la base de datos por su nivel de sensibilidad.
- Verificar los identificadores directos que podrían ser agrupados con el propósito de minimizar el riesgo de identificación
- Listar todas las posibles combinaciones de las variables que representen un riesgo de identificación de las unidades de observación.
- Analizar los niveles mínimos de desagregación de los identificadores directos o pseudoidentificadores que puedan ser más riesgosos para de-

terminar la identificación de las unidades de observación. Usualmente estos niveles de desagregación hacen referencia a desagregaciones geográficas o temáticas.

- Revisar las medidas descriptivas estadísticas calculadas para las variables cuantitativas en el *Análisis exploratorio de la base de datos* (Etapa I).
- Revisar la distribución de frecuencias de las variables categóricas calculadas en el *Análisis exploratorio de la base de datos* (Etapa I) y utilizarlas en la identificación de categorías con frecuencias considerablemente bajas. Usualmente este tipo de categorías se asocian como riesgosas porque permiten fácilmente identificar a las unidades de observación.
- Considerar la recodificación de las categorías de las variables que presentan frecuencias considerablemente bajas.
- Verificar las bases de datos de entidades externas encontradas en la revisión temática realizada en el *Análisis exploratorio de la base de datos* (subproceso de la Etapa I). El equipo deberá identificar si todas las variables contenidas en esas bases podrían convertirse en pseudoidentificadores respecto a la base de datos original sujeta a anonimizar. Esta revisión permitirá establecer qué variables de las bases de datos externas, combinadas con las variables de la base de datos a anonimizar, aumentan el riesgo de identificación de las unidades de observación.

Cuando el equipo de trabajo obtenga el listado priorizado de los riesgos de identificación de las unidades de observación de acuerdo con la frecuencia en la

que ocurran estos riesgos, procederá a evaluar temáticamente la base de datos con el equipo de trabajo. En este caso, el equipo revisará si es necesario considerar todos los riesgos identificados, si se pueden considerar sólo algunos, o si se pueden construir nuevos riesgos a partir de la primera ronda de riesgos planteados.

Después de que el equipo de trabajo haya evaluado todos los posibles riesgos de identificación y haya decidido finalmente cuáles son los definitivos (o más probables), procederá a identificar qué unidades de observación se considerarán riesgosas bajo los escenarios de riesgos definidos en este subproceso.

Identificación de unidades de observación riesgosas

Con el listado definitivo de riesgos de identificación priorizados, el equipo de trabajo procederá a identificar con mayor nivel de detalle las unidades de observación que son riesgosas bajo todos los riesgos planteados en el subproceso anterior.

Las unidades de observación riesgosas son aquellas que cumplen con al menos una de las condiciones planteadas por el equipo de trabajo para ser susceptibles a identificación. Una unidad de observación puede ser riesgosa por sólo un riesgo o por todos los riesgos planteados por el equipo de trabajo.

La identificación de las unidades riesgosas, por su parte, puede ser una actividad para realizar por parte del equipo del procesamiento de la base de datos dada su experticia y capacidades técnicas en el manejo de este tipo de archivos de información.

A continuación, la Tabla 9 presenta un resumen sobre la identificación de las unidades de observación que resultan riesgosas.

Tabla 9. Resumen de unidades de observación riesgosas en la anonimización teniendo en cuenta tres riesgos

Riesgo	Variables involucradas	¿Cuándo se considera una unidad de observación riesgosa?	Número de unidades de observación riesgosas	Porcentaje de unidades de observación riesgosas
RIESGO 1	Escriba en este espacio qué variables están involucradas en este riesgo.			Escriba en este espacio el porcentaje de unidades de observación riesgosas por el riesgo 1 con respecto al total de unidades de observación.
RIESGO 2		En este espacio explique qué condiciones debe cumplir una unidad de observación para considerarse riesgosa.		
RIESGO 3			Escriba en este espacio cuántas unidades de observación se consideran riesgosas bajo el riesgo 3.	

Riesgo	Variables involucradas	¿Cuándo se considera una unidad de observación riesgosa?	Número de unidades de observación riesgosas	Porcentaje de unidades de observación riesgosas
TOTAL	Escriba en este espacio el número de variables involucradas en el análisis de riesgos			Escriba en este espacio el porcentaje de unidades de observación riesgosas con respecto al total de unidades de observación.

Fuente: DANE - DIRPEN.

Creación del informe de riesgos

Finalmente, el equipo de trabajo creará un informe de riesgos que será utilizado en la identificación y la aplicación de técnicas de anonimización (Etapa III), el cual describirá cómo se clasifican las variables según su tipo de sensibilidad, los criterios utilizados para la definición de riesgos de identificación y las unidades de observación que son riesgosas a la hora de publicar la base de datos.

El planteamiento de los riesgos de identificación de las unidades de observación debe estar debidamente documentado. De manera que en caso de cambios en el equipo de trabajo responsable de la anonimización se pueda asegurar una trazabilidad del proceso.

Se recomienda que el informe contenga al menos la siguiente información:

- Criterios y aspectos considerados en la definición de los riesgos.
- Listado de riesgos definitivos priorizados.

- La tabla resumen de unidades de observación riesgosas obtenida en el tercer subproceso de esta etapa (Tabla 9).
- Fecha de emisión del informe.

Producto Etapa II:

- Clasificación de variables por su tipo de sensibilidad.
- Planteamiento de riesgos de identificación.
- Identificación de las unidades de observaciones riesgosas.
- Informe de riesgos de identificación.

Para visualizar estos elementos de los subprocesos de la Etapa I, se presenta a continuación un ejemplo.

Ejemplo de la Etapa II: Análisis de riesgos

Para ejemplificar la etapa de análisis de riesgo, se utilizará una base de datos simulada a la cual se llamará **COL20**¹⁶, la cual contiene información de 32 variables de identificación, ubicación y socioeconómicas, medidas en 496 personas que se encuentran presentes en

345 municipios a nivel nacional. La descripción de cada una de las variables de **COL20** se encuentra disponible en el Anexo C.

Para el análisis de riesgos, con base en la experiencia del DANE, se iniciará con la clasificación de las variables por su tipo de sensibilidad. Las 32 variables de **COL20**, pueden caracterizarse así:

Tabla 10. Clasificación de las variables por tipo de sensibilidad de COL20

Tipo de variable / tipo de sensibilidad	Identificadores directos	Pseudoidentificadores	No confidenciales	Total
CUANTITATIVAS	0	8	0	8
CATEGÓRICAS	6	14	3	23
TOTAL	6	22	3	31

Fuente: DANE - DIRPEN.

De la anterior tabla, se puede concluir que 28 variables son sensibles, entre identificadores directos y pseudoidentificadores, donde 20 de las variables son categóricas. Algunas variables contenidas en **COL20** que son identificadores directos son nombres, apellidos, dirección, número de identificación, y respecto a los pseudoidentificadores, se tienen variables como RH, grupo étnico, ingresos anuales, número de bienes raíces, entre otras.

Al revisar la distribución de las variables, se observa que la mayor parte de estas son categóricas, recordando que para este tipo de variables las unidades de observación cuentan con una valoración cuando cumplen una característica particular. Es el caso de la variable Grupo étnico que toma valores *afrocolombiano*, *gitano*, *indígena* o *ninguno*. Teniendo en cuenta estas características de las variables categóricas, es posible utilizar como una medida aproximada para la definición de riesgos de identificación, las distribuciones de frecuencias.

¹⁶ https://www.dane.gov.co/files/investigaciones/boletines/eac/bol_eac_2016.pdf

Posteriormente, siguiendo las recomendaciones propuestas en esta guía, se elaboró el siguiente listado de riesgos de la base de datos **COL20**:

- Los identificadores directos, como número de identificación, tipo de identificación, nombre, apellidos, fecha de nacimiento y dirección, deben ser suprimidos definitivamente de la base de datos porque representa una identificación inmediata de las unidades de observación.
- Las siguientes combinaciones entre pseudoidentificadores se consideran riesgosas porque el cruce entre ellas podría, eventualmente, permitir la identificación de una unidad de observación de la base de datos:
 - Ingresos anuales (o mensuales) y nivel de escolaridad a nivel municipal y departamental.
 - RH y edad a nivel municipal y departamental.
 - Ocupación, nivel de escolaridad, ingresos anuales (o mensuales) a nivel municipal.
 - Grupo étnico a nivel municipal y departamental.
 - Número de habitaciones de la vivienda, materiales de los pisos del hogar e ingresos anuales (o mensuales) a nivel municipal y departamental.
- Número de bienes raíces con la variable “tiene vehículo” (marca y modelo) a nivel municipal y departamental.
- Número de viajes fuera del país, ocupación e ingresos anuales (o mensuales).
- Valores frecuentes de la variable asistencia a eventos culturales (o deportivos), edad e ingresos a nivel municipal y departamental.
- La desagregación a nivel municipal es altamente riesgosa para la identificación de las unidades de observación con respecto a algunas variables temáticas como es el caso de la variable *Ingresos anuales*, porque se podría tener con certeza que en el municipio El Castillo en el departamento del Meta se encuentran 2 personas con ingresos anuales superiores a 10 millones de pesos.
- A partir de las medidas descriptivas estadísticas calculadas en la Etapa I, se identificaron las variables cuantitativas más sensibles. A continuación, se presentan los resultados de este ejercicio para **COL20**:

Tabla 11. Medidas descriptivas de las variables cuantitativas de COL20

Variable	Media	Varianza	Cuartil 1	Cuartil 2	Cuartil 3
Ingresos anuales*	\$ 52.174.008	2.15679E+15	\$ 13.188.150	\$ 40.048.606	\$ 89.000.899
Ingresos mensuales*	\$ 4,347,834	1.49777E+13	\$ 1.099.012	\$ 3.337.384	\$ 7.416.742
Número de hijos nacidos** vivos	1,97	2,84	0	2	3
Número de bienes raíces**	2,38	3,25	1	2	4
Número de viajes fuera del país**	3,47	6,74	1	3	6

*Unidades: millones de pesos

**Números enteros

Fuente: DANE- DIRPEN.

- Con el propósito de identificar las variables categóricas más sensibles, se buscan aquellas categorías con menor participación a nivel nacional. En este ejemplo, se utilizaron las distribuciones de frecuencia calculadas en la sección 5.2 (Tabla 5), de las variables categóricas que el equipo consideró eran las más sensibles, como por ejemplo: grupo étnico, nivel de escolaridad y RH.

Es importante recordar que las distribuciones considerablemente bajas presentan altos niveles de riesgo. Por ejemplo, se identificaron las

variables de RH; nivel de escolaridad y grupo étnico como variables que permiten fácilmente la identificación de las unidades de observación.

- Estas categorías deben revisarse cuidadosamente, ya que las unidades de observación que pertenezcan a esta categoría pueden ser riesgosas a nivel municipal.

A continuación, se presentan las frecuencias para las variables que tienen un mayor nivel de sensibilidad en la base de dato **COL20**:

Tabla 12. Distribución de frecuencias del RH en COL20

RH	Número de registros	Porcentaje de registros
O+	244	49.2%
A+	212	42.7%
B+	5	1.0%
AB+	3	0.6%
O-	23	4.6%
A-	6	1.2%
B-	1	0.2%
AB-	2	0.4%
Total	496	100%



En este caso, las categorías de RH B-, AB-, AB+ y A-, se consideran como las más sensibles, dada su baja frecuencia de aparición en las unidades de observación.

Fuente: DANE - DIRPEN.

Tabla 13. Distribución de frecuencias del nivel de escolaridad en COL20

RH	Número de registros	Porcentaje de registros
Primaria	54	10,9%
Secundaria	26	5,2%
Básica Media	70	14,1%
Técnico	30	6%
Tecnólogo	88	17,7%
Profesional	130	26,2%
Posgrado	98	19,8%
Total	496	100%



Para este segundo caso, se tendrían las categorías de secundaria y técnico como los niveles de escolaridad más sensibles dada su baja frecuencia en las unidades de observación.

Fuente: DANE - DIRPEN.

Tabla 14. Distribución de frecuencias del grupo étnico en COL20

Grupo étnico	Número de registros	Porcentaje de registros
Afrocolombiano	20	4.0%
Indígena	12	2.4%
Rrom	4	0.8%
Ninguno	460	92.7%
Total	496	100%



En este último caso, las categorías Rrom, indígena y afrocolombiano son los grupos étnicos más sensibles dada su baja frecuencia en las unidades de observación.

Fuente: DANE - DIRPEN.

Además, se presenta la tabla de identificación de unidades de observación riesgosas, la cual sirve como insumo para la

identificación de las técnicas de anonimización a utilizar y definir la viabilidad del proceso de anonimización:

Tabla 15. Unidades de observación riesgosas para los cinco riesgos más frecuentes para la anonimización de COL20

Riesgo	VARIABLES INVOLUCRADAS	¿Cuándo se considera una unidad de observación riesgosa?	Número de unidades de observación riesgosas	Porcentaje de unidades de observación riesgosas
1	Ingresos mensuales y departamento	Las tres personas con el ingreso anual más alto en cada departamento.	99	20%
2	Grupo étnico, departamento	Las personas pertenecientes a un grupo étnico particular a nivel departamental.	36	7,3%
3	Número de habitaciones de la vivienda, departamento	Todas las personas que vivan en una vivienda con un número de habitaciones por encima del promedio departamental.	235	47,4%
4	Número de viajes fuera del país y departamento	Todas las personas que hayan viajado fuera del país más veces que el promedio departamental.	225	45,4%

Riesgo	Variables involucradas	¿Cuándo se considera una unidad de observación riesgosa?	Número de unidades de observación riesgosas	Porcentaje de unidades de observación riesgosas
5	Nivel de escolaridad y departamento	Todas las personas con posgrado en aquellos departamentos con menos de cuatro personas a ese nivel de escolaridad.	31	6,3%

Fuente: DANE - DIRPEN.

La anterior tabla evidencia que, de los cinco riesgos más probables seleccionados, más del 40% de las unidades de observación de la base de datos tienen riesgo de ser identificadas por las variables *número de habitaciones de la vivienda* (47,4%) y *número de viajes fuera del país* (45,4%) (Riesgos 3 y 4). De la misma forma, 99 unidades de observación tienen

riesgo de ser identificadas por sus ingresos mensuales (Riesgo 1).

Teniendo en cuenta que estas tres variables son cuantitativas, es posible identificar las técnicas de anonimización más idóneas para minimizar el riesgo de identificación de aquellas unidades de observación.

Recuadro 5. Ejemplo Comisión de la Verdad

La Comisión de la Verdad de Colombia, establecida en el marco del Acuerdo de Paz firmado en 2016, tiene como objetivo principal esclarecer los hechos y responsabilidades en las violaciones de derechos humanos durante el conflicto armado interno. En línea con su mandato, la Comisión emprendió la tarea de recopilar y analizar información estadística sobre diversos tipos de victimización.

Este organismo llevó a cabo un proceso estadístico integral, utilizando la unión de 112 bases de datos relacionadas con diferentes hechos victimizantes en el conflicto armado. Posteriormente, tras la consolidación de la información, se generaron 100 réplicas de imputación múltiple específicamente para los componentes de homicidios, desapariciones forzadas, secuestros y reclutamiento ilícito. Estas bases de datos contenían registros con variables potencialmente identificadoras como edad, sexo, etnia y ubicación. Su publicación permitiría estudios sobre patrones de violencia, pero también conllevaría riesgos de reidentificación de las víctimas.

Por esta razón, el DANE realizó un riguroso análisis del proceso de anonimización aplicado por el equipo investigador sobre las 100 réplicas imputadas.

Análisis de riesgo

La metodología consistió en detectar registros únicos en la base de datos según distintas combinaciones de variables como llaves pseudoidentificadoras. Como resultado, entre más inhabitual la combinación (ej. una mujer indígena de 92 años secuestrada en un municipio de 500 habitantes) mayor es la probabilidad de reidentificación. Se probaron dos escenarios: (1) departamento y fecha del evento y (2) departamento, fecha del evento y municipio.

En el escenario 1 se encontraron 2.661 registros únicos de alto riesgo (1,1% del total), mientras que en el escenario 2 fueron 35.003 registros (4,4%).

Para evaluar la posibilidad real de reidentificación se realizaron ejercicios de rompimiento cruzando las bases de datos con registros externos como el Registro de Población y el Registro de Víctimas. Usando las mismas variables en ambas bases se buscaron coincidencias exactas, encontrando tasas de entre 1,1% y 11,1%.

** Los ejercicios de rompimiento (cruces de la información con bases externas), pueden realizarse para determinar si existen más unidades riesgosas, adicionales a las halladas en la base de datos a anonimizar.

Etapla III. Selección y aplicación de las técnicas de anonimización

Esta etapa presenta las posibles técnicas de anonimización que se pueden utilizar, en función a las particularidades de las variables contenidas en la base de datos y el uso que se pretenda dar a la base anonimizada.

Técnicas de anonimización en datos estructurados

En esta etapa el equipo de trabajo conocerá e identificará las técnicas de anonimización más comunes para variables cuantitativas y categóricas. Además, seleccionará una o más técnicas para aplicar a cada uno de los riesgos planteados en la etapa anterior.

La etapa de identificación y selección de técnicas de anonimización se compone de dos subprocesos así:

- **Identificación de técnicas de anonimización más comunes y adecuadas.**
- **Selección de técnicas de anonimización.**

Identificación de técnicas de anonimización más comunes

En este subproceso se presentan de manera general las técnicas de anonimización más comunes por tipo de varia-

ble, que permiten minimizar el riesgo de identificación de las unidades de observación.

Tenga en cuenta que las técnicas de anonimización se dividen en:

- **Técnicas basadas en la no perturbación de datos:** estas técnicas utilizan supresiones parciales, reducción o recodificación de la información para minimizar el riesgo de identificación de las unidades de observación. Este tipo de técnicas son comúnmente utilizadas para evitar que los datos atípicos sean de fácil identificación.
- **Técnicas basadas en la perturbación de datos:** estas técnicas se refieren a procedimientos que implican la modificación sistemática de datos (a veces en pequeñas cantidades aleatorias), de manera tal que las cifras no sean lo suficientemente precisas como para revelar información sobre casos individuales. Pueden incluirse nuevos datos, suprimir y modificar los existentes beneficiando la confidencialidad estadística¹⁷.

A continuación, se describen brevemente las técnicas y su implementación teniendo en cuenta el tipo de variables que pueden estar contenidas en las bases de datos. En este caso para aquellas técnicas basadas en no perturbación.

¹⁷ González M., 2017, p. 29.

Tabla 15. Unidades de observación riesgosas para los cinco riesgos más frecuentes para la anonimización de COL20

Técnicas	Descripción	Tipo de variable	Ejemplos variables (base col20)	Referencia bibliográfica
ELIMINACIÓN DE VARIABLES	Esta técnica suprime toda la información de una variable. Se usa cuando la variable contiene información de identificación directa de la unidad de observación.	En variables categóricas	CÉDULA; TIPO DE IDENTIFICACION; NOMBRE; APELLIDOS; DIRECCIÓN; BARRIO; MUNICIPIO; FECHA DE NACIMIENTO; RH Estas variables son eliminadas porque debido a la información contenida es posible identificar directamente a las unidades de observación. Además, algunas de ellas contienen información sensible, por lo tanto, permitirían reconocer las unidades de observación.	Hundepool et al. (2010).

Técnicas	Descripción	Tipo de variable	Ejemplos variables (base col20)	Referencia bibliográfica
RECODIFICACIÓN GLOBAL	Esta técnica consiste en combinar diversas categorías de las variables categóricas en una más general que tenga mayor frecuencia y menor información. En el caso de las variables continuas, consiste en agrupar por medio de intervalos para mantener la utilidad de los datos. Esta técnica es recomendable cuando se desean proteger unidades de observación con riesgo de identificación a partir de las variables pseudoidentificadoras.	En variables continuas o categóricas	GRUPO ÉTNICO; NÚMERO DE HABITACIONES DE LA CASA; NÚMERO DE VIAJES REALIZADOS FUERA DEL PAÍS; NIVEL DE ESCOLARIDAD Estas variables son recodificadas para combinar las categorías de las variables y poder tener en cada nueva categoría una mayor frecuencia de unidades de observaciones.	Hundepool et al., (2012). Templ et al., IHSN Working Paper No. 007 (2014).

Técnicas	Descripción	Tipo de variable	Ejemplos variables (base col20)	Referencia bibliográfica
CODIFICACIÓN SUPERIOR E INFERIOR	Esta técnica consiste en proteger la identificación de las unidades de observación que presentan los valores más altos o más bajos de cada variable. Se utiliza cuando se presentan valores máximos y mínimos en el nivel de desagregación geográfico o temático que son de fácil identificación.	En variables continuas o categóricas.	TÉCNICA NO UTILIZADA EN LA ANONIMIZACIÓN DE LA BASE DE DATOS EJEMPLO COL20	Hundepool et al. (2012).
SUPRESIÓN LOCAL	Esta técnica consiste en reemplazar los valores de una o más variables de las unidades de observación identificadas como riesgosas por valores faltantes. Esta técnica se usa cuando la combinación entre las variables pseudoidentificadoras permita la identificación de las unidades de observación.	En variables categóricas.	TÉCNICA NO UTILIZADA EN LA ANONIMIZACIÓN DE LA BASE DE DATOS EJEMPLO COL20	Hundepool y De Wolf (2012). Templ et al., IHSN Working Paper No. 007 (2014).

Fuente: DANE - DIRPEN.

Respecto a las técnicas basadas en perturbación, se tienen las siguientes recomendaciones:

Tabla 17. Técnicas basadas en la perturbación de datos según el tipo de variable

Técnicas	Descripción	Tipo de variable	Ejemplos variables	Referencia bibliográfica
			(base col20)	
MICROAGREGACIÓN	Esta técnica consiste en reemplazar los valores de algunas unidades de observación, por el valor promedio calculado sobre ellas.	En variables continuas.	EDAD;INGRESOS ANUALES;INGRESOS MENSUALES;NÚMERO DE HIJOS; NACIDOS VIVOS;CUANTAS PERSONAS COMPONEN EL HOGAR; NUMERO DE HABITACIONES DE LA CASA;NUMERO DE BIENES RAICES;NUMERO DE VIAJES FUERA DEL PAÍS Estas variables son microagregadas porque se busca proteger la identificación de las unidades de observaciones con los ingresos anuales más altos por departamento.	Hundepool et al., (2012) Templ et al., IHSN Working Paper No. 007 (2014).

Fuente: DANE - DIRPEN.

Técnicas	Descripción	Tipo de variable	Ejemplos variables	Referencia bibliográfica
			(base col20)	
REDONDEO	Esta técnica consiste en sustituir los valores de las unidades de observación en aquellas variables que tienen decimales por valores redondeados (cero decimales). Comúnmente, se usa después de aplicar microagregación y cuando la información de las variables técnicamente se debe expresar en unidades enteras y no decimales. Por ejemplo, número de hijos.	En variables continuas.	EDAD;NÚMERO DE HIJOS NACIDOS VIVOS;NÚMERO DE HABITACIONES DE LA CASA;NÚMERO DE BIENES RAÍCES;NÚMERO DE VIAJES FUERA DEL PAÍS. Estas variables son redondeadas para mantener la información de las variables en unidades enteras después de su microagregación.	Hundepool et al. (2012).

Técnicas	Descripción	Tipo de variable	Ejemplos variables	Referencia bibliográfica
			(base col20)	
INTERCAMBIO DE DATOS	Esta técnica consiste en intercambiar la información de las unidades de observación identificadas con riesgo, con la información de las unidades de observación que no tienen riesgo de identificación. Este intercambio de datos se realiza de manera aleatoria entre pares de observaciones (con riesgo de identificación y sin riesgo).	En variables continuas o categóricas	TÉCNICA NO UTILIZADA EN LA ANONIMIZACIÓN DE LA BASE DE DATOS EJEMPLO COL20	Hun-depool et al., (2010), p. 58.

Tabla 18. Planteamiento de técnicas de anonimización para cada uno de los riesgos identificados

Riesgo	Variables involucradas	¿Cuándo se considera una unidad de observación riesgosa?	Número de unidades de observación riesgosas	Porcentaje de unidades de observación riesgosas	Técnica(s) de anonimización a utilizar
RIESGO 1	Escriba en este espacio qué variables están involucradas en este riesgo.			Escriba en este espacio el porcentaje de unidades de observación riesgosas por el riesgo 1 con respecto al total de unidades de observación.	Escriba en este espacio la(s) técnica(s) de anonimización que minimizarán la ocurrencia del riesgo 1.
RIESGO 2		En este espacio explique qué condiciones debe cumplir una unidad de observación para considerarse riesgosa.			
RIESGO 3			Escriba en este espacio cuántas unidades de observación se consideran riesgosas bajo el riesgo 3.		

Riesgo	Variables involucradas	¿Cuándo se considera una unidad de observación riesgosa?	Número de unidades de observación riesgosas	Porcentaje de unidades de observación riesgosas	Técnica(s) de anonimización a utilizar
TOTAL	Escriba en este espacio el número de variables involucradas en el análisis de riesgos.			Escriba en este espacio el porcentaje de unidades de observación riesgosas con respecto al total de unidades de observación.	Escriba en este espacio todas las técnicas de anonimización que utilizará para minimizar la ocurrencia de todos los riesgos.

Fuente: DANE- DIRPEN.

Es importante destacar que el equipo de trabajo deberá realizar la selección de las técnicas de anonimización, teniendo en cuenta el tipo de variables involucradas en cada uno de los riesgos y de las propiedades estadísticas que se desean conservar en la base de datos.

Las técnicas de anonimización más comunes fueron presentadas en el anterior subproceso, con una referencia bibliográfica que se puede consultar para profundizar en su aplicación.

La Tabla 18 (Planteamiento de técnicas de anonimización para cada uno de los riesgos identificados) sirve como insumo para la toma de decisiones en el análisis de viabilidad que se realizará en la siguiente etapa.

Productos de la Etapa III:

- Características de las técnicas de anonimización.
- Técnicas de anonimización para cada uno de los riesgos identificados.

publicada a los usuarios. Si esta utilidad es baja y no permite la réplica de las cifras publicadas por la entidad, se considerará que el proceso de anonimización no es viable.

Cuando el equipo identifique que el proceso de anonimización se ve afectado por el primer criterio, definitivamente no es viable. Sin embargo, cuando se presentan los criterios 2 y 3, existe una forma alternativa de anonimizar la base de datos:

- El equipo definirá propiedades estadísticas más flexibles para que la base de datos anonimizada pueda conservar y garantizar la utilidad de la información para los usuarios. La flexibilidad de las propiedades estadísticas puede darse al cambiar el nivel de desagregación geográfica o temática, modificándolas a categorías más generales, al aumentar la posible variación en las propiedades globales o al eliminar variables sensibles que son de alto riesgo de identificación.

Por ejemplo: un equipo deseaba publicar la variable “Ingreso promedio por hogar” a nivel departamental. Sin embargo, con las técnicas existentes, evidenció que aún se podían identificar algunas unidades de observación. Por lo tanto, decidió modificar la desagregación de nivel departamental a nivel regional, que, aunque no conserva la utilidad que se esperaba, sigue siendo información relevante para el usuario y además permite la réplica de las cifras publicadas por la entidad.

Finalmente, si el equipo logra proponer nuevas propiedades estadísticas que permiten que la base de datos anonimizada siga siendo útil para los usuarios, el proceso se considera viable.

A modo de resumen frente a los tres criterios, el equipo de trabajo podría diseñar un marco para establecer el riesgo máximo tolerable para anonimizar la base de datos, teniendo en cuenta el siguiente esquema.

Tabla 19. Criterios para analizar la viabilidad de la anonimización de la base de datos

Criterio	Nivel de afectación	Decisión
1. Revisión temática y normativa	Alta	No es viable la anonimización de la base de datos.
	Todas las variables presentan riesgos.	
2. Técnicas de anonimización	Medio	Flexibilizar las propiedades estadísticas de la base de datos (Etapa III. Selección de las técnicas de anonimización).
	Se podrían identificar algunas unidades de observación.	
3. Nivel de utilidad de la información	Medio	
	Se podrían identificar algunas unidades de observación.	

Fuente: DANE - DIRPEN.

Producto Etapa IV:

- Informe y concepto de viabilidad del proceso de anonimización.

En términos conceptuales, el análisis de viabilidad que realiza el equipo de trabajo en esta etapa, puede verse de forma detallada en el Recuadro 7.

Etapa V. Aplicación de técnicas de anonimización

En esta etapa el grupo de trabajo implementará las técnicas de anonimización asociadas a los riesgos de identificación seleccionadas en la Etapa II, obteniendo así una primera versión de la base de datos anonimizada que será examinada cuidadosamente en la siguiente etapa. Para la aplicación de las técnicas el equipo debe tener en cuenta los siguientes pasos:

- Clasificación de las técnicas asociadas a cada riesgo identificado, tal y como se evidencia en la Tabla 18 (Planteamiento de técnicas de anonimización para cada uno de los riesgos identificados).
- Seleccionar el software que permita implementar las técnicas de anonimización elegidas con el fin de planear de manera eficiente el proceso de anonimización. Dentro de los paquetes de software a tener en cuenta están μ -Argus (es un software gratuito y de código abierto que proporciona técnicas de sustitución, supresión y perturbación de datos), desarrollado por el Instituto de Estadística de Países bajos, T-Argus (es un software gratuito de código abierto que proporciona una técnica de anonimización basada en la teoría de la

información) cuando se tienen datos agregados o tablas, SAS, también se ha publicado un paquete de anonimización para R, entre otros.

- Es importante que el equipo de trabajo crea rutinas (algoritmos) que ayuden a implementar las técnicas para que disminuya el riesgo de identificación de las unidades de observación. Se recomienda que en las rutinas el equipo de trabajo siga la siguiente estructura:
- **Cargue de base de datos a anonimizar.**
- **Tipos de riesgos identificados y la explicación de estos.**
- **Consolidación de riesgos identificados.**
- **Técnica de anonimización a aplicar.**
- **Verificación de riesgos.**
- **Exportación de base de datos anonimizada.**

A continuación, se presentan una parte de la rutina utilizada por el DANE en el ejercicio simulado para la anonimización de la Encuesta Anual de Comercio (EAC), utilizando el paquete estadístico SAS:

Recuadro 8. Ejercicio simulado para la anonimización de la EAC

Se presenta la forma en que fue cargada la base de datos a anonimizar:

```
libname EAC "\\BASES_ANONIMIZADAS\EAC"; run;
data base11;
set work.'2011_if_0000'n;
run;
```

Posteriormente, se presenta la forma en que identificó los valores máximos de

la variable Ventas, que corresponde al riesgo 1 previsto en la base de la EAC:

```
PROC SQL; /*RIESGO 1 MAXIMOS A NIVEL NACIONAL*/
CREATE TABLE MAXIMOS
AS SELECT *, MAX(VENTA) AS MAX_VENTA
FROM BASE11;
QUIT;
DATA MAXIMOS_VENTAS;
SET MAXIMOS;
IF VENTA = MAX_VENTA THEN ID_RIESGO1 = 1;
RUN;
```

Se relaciona la forma en que se considera el número de unidades de observa-

ción riesgosas, de acuerdo con los diferentes riesgos planteados en la EAC:

```
PROC FREQ DATA=EAC.EAC_RIESGO2016 ORDER=INTERNAL;
TABLES Riesgo1 / SCORES=TABLE;
TABLES Riesgo2 / SCORES=TABLE;
TABLES Riesgo3 / SCORES=TABLE;
TABLES Riesgo4 / SCORES=TABLE;
TABLES RiesgoTotal / SCORES=TABLE;
RUN;
```

Finalmente, se presenta la programación para aplicar la técnica de microagregación cuyo objetivo es eliminar, uno

de los riesgos de identificación planteado en la base de datos de la EAC.

```
/*GENERAR UN IDENTIFICADOR*/
data ranqueo (where =(ranqueo in (1,2,3)));
set orden;
retain ranqueo 0;
if first.llave then do ranqueo=0;
end;
ranqueo=ranqueo+1;
by llave;
run;

/*APLICA TÉCNICA: PROMEDIAR LOS DATOS MAS ALTOS
DE TODAS LAS VARIABLES POR LLAVE*/
proc sql;
create table Metodo_1 as select
CIIU4,
intio, mean (intio) as promedio_intio,
/*VARIABLES, MEAN(VARIABLE) AS PROMEDIO_VARIABLE,*/
from ranqueo
group by llave; quit;
```

Producto Etapa V:

- Base de datos anonimizada.
- Rutina del proceso de anonimización.

Ejemplo de la Etapa V: Aplicación de técnicas de anonimización

Retomando la base **COL20**, esta cuenta con seis variables que son identificadores directos, 22 variables que son pseudoidentificadores y tres variables no confidenciales.

Antes de identificar y aplicar las técnicas de anonimización, se define que, en

este ejercicio simulado, el objetivo del proceso de anonimización es mantener los promedios de las variables cuantitativas (edad, ingresos anuales, ingresos mensuales, número de hijos nacidos vivos, cuantas personas componen el hogar, número de bienes raíces) con una variación inferior al 5% a nivel departamental.

El equipo de trabajo lista los riesgos y acorde con el tipo de variable, indica la técnica de anonimización a utilizar, como se evidencia en la Tabla 20.

