



SAFER • HEALTHIER • PEOPLE™

National Health and Nutrition Examination Survey: Sample Design, 2011–2014



U.S. DEPARTMENT OF HEALTH AND HUMAN SERVICES
Centers for Disease Control and Prevention
National Center for Health Statistics

Copyright information

All material appearing in this report is in the public domain and may be reproduced or copied without permission; citation as to source, however, is appreciated.

Suggested citation

Johnson CL, Dohrmann SM, Burt VL, Mohadjer LK. National Health and Nutrition Examination Survey: Sample design, 2011–2014. National Center for Health Statistics. Vital Health Stat 2(162). 2014.

Library of Congress Cataloging-in-Publication Data 614.4'273—dc23

For sale by the U.S. Government Printing Office
Superintendent of Documents
Mail Stop: SSOP
Washington, DC 20402–9328
Printed on acid-free paper.

Vital and Health Statistics

Series 2, Number 162

National Health and Nutrition Examination Survey: Sample Design, 2011–2014

Data Evaluation and Methods Research

U.S. DEPARTMENT OF HEALTH AND HUMAN SERVICES
Centers for Disease Control and Prevention
National Center for Health Statistics

Hyattsville, Maryland
March 2014
DHHS Publication No. 2014-1362

National Center for Health Statistics

Charles J. Rothwell, M.S., M.B.A., *Director*

Jennifer H. Madans, Ph.D., *Associate Director for Science*

Division of Health and Nutrition Examination Surveys

Kathryn S. Porter, M.D., M.S., *Director*

Contents

Acknowledgments	iv
Abstract	1
Introduction	1
Design Specifications	2
Survey Objectives	2
Domain and Precision Considerations	2
Operational Requirements	4
Sample Design	4
Sampling Rates	5
Sample Selection Methods	7
Special Samples	15
References	16
Appendix I. Glossary	17
Appendix II. Supporting Tables and Figure	20

Text Tables

A. Selected sample design parameters: Health and Nutrition Examination Surveys, 1971–2014	3
B. Sampling subdomains classified by race and Hispanic origin, income, sex, and age: National Health and Nutrition Examination Survey, 2011–2014	4
C. Values of A_k used in calculating primary sampling unit measures of size	7
D. State grouping for primary sampling unit stratification	9
E. Characteristics of major strata formed for selection of primary sampling units: National Health and Nutrition Examination Survey, 2011–2014	9
F. Release group distribution: National Health and Nutrition Examination Survey, 2011–2014	15

Appendix Figure

Major strata formed for selection of 2011–2014 primary sampling units: National Health and Nutrition Examination Survey, 2011–2014	25
--	----

Appendix Tables

I. Derivation of expected screening requirements: National Health and Nutrition Examination Survey, 2011–2014	20
II. Final sampling rates and base weights: National Health and Nutrition Examination Survey, 2011–2014	22
III. Description of subsamples: National Health and Nutrition Examination Survey, 2011–2014	24

Acknowledgments

The authors gratefully acknowledge the assistance of Michele Chiappa, Lester Curtin, Lisa Mirel, and Ryne Paulose-Ram in the preparation of this report.

Background

Data collection for the National Health and Nutrition Examination Survey (NHANES) consists of a household screener, an interview, and a physical examination. The screener primarily determines whether any household members are eligible for the interview and examination. Eligibility is established using preset selection probabilities for the desired demographic subdomains. After an eligible sample person is selected, the interview collects person-level demographic, health, and nutrition information, as well as information about the household. The examination includes physical measurements, tests such as hearing and dental examinations, and the collection of blood and urine specimens for laboratory testing.

Objectives

This report provides some background on the NHANES program, beginning with the first survey cycle in the 1970s and highlighting significant changes since its inception. The report then describes the broad design specifications for the 2011–2014 survey cycle, including survey objectives, domain and precision specifications, and operational requirements unique to NHANES. The report also describes details of the survey design, including the calculation of sampling rates and sample selection methods. Documentation of survey content, data collection procedures, estimation methods, and methods to assess nonsampling errors are reported elsewhere.

Keywords: weighting • sampling rates • NHANES

National Health and Nutrition Examination Survey: Sample Design, 2011–2014

by Clifford L. Johnson, M.S.P.H., National Center for Health Statistics; Sylvia M. Dohrmann, M.S., Westat; Vicki L. Burt, Sc.M., R.N., National Center for Health Statistics; and Leyla K. Mohadjer, Ph.D., Westat

Introduction

The National Health and Nutrition Examination Survey (NHANES) is one of a series of health-related programs conducted by the Centers for Disease Control and Prevention's (CDC) National Center for Health Statistics (NCHS). A unique feature of this survey is the collection of health examination data for a nationally representative sample of the resident, civilian noninstitutionalized U.S. population. The survey consists of questionnaires administered in the home, followed by a standardized health examination in specially equipped mobile examination centers (MECs).

NHANES provides information on the noninstitutionalized civilian resident population. It excludes all persons in supervised care or custody in institutional settings, all active-duty military personnel, active-duty family members living overseas, and any other U.S. citizens residing outside of the 50 states and District of Columbia. Noninstitutional group quarters are included in the sample; see the Glossary (Appendix I) for more details on institutional and noninstitutional group quarters.

NHANES I, the first cycle of NHANES, was conducted during 1971–1975 (1–5); two other cycles, and a cycle focusing on the Hispanic population (HHANES) (6), were conducted between 1976 and 1994 (7–8). NHANES was fielded again in 1999 (9) and, in the tradition of the

previous national surveys, it continues to provide information on the health and nutritional status of the U.S. population. This information is used to estimate the prevalence of various diseases and conditions and to provide information for many public health functions (e.g., reference data, nutrition and health monitoring, and prevention initiatives).

Differences in the sample sizes and designs for all cycles of NHANES and for HHANES should be considered when comparisons are made across the various surveys. For example, NHANES I, NHANES II, and HHANES did not include persons aged 75 and over, and NHANES I and NHANES II did not oversample Hispanic or black persons. The sample for NHANES 1999–2006 included a supplemental sample of pregnant women (9). NHANES 1999–2006 also included an oversample of Mexican-American persons.

In the 2007–2010 sample, the oversampling of the Mexican-American population was replaced by an oversample of the entire Hispanic population (10). In addition, the oversample of adolescents and the supplemental sample of pregnant women in the 1999–2006 survey were discontinued.

The primary change in the NHANES 2011–2014 design was the addition of an oversample of Asian persons. The design also featured a revised stratification scheme at the primary sampling unit (PSU) level, which included a representative sample for California. The sample design

parameters for the seven HANES surveys are compared in [Table A](#).

Data collection for NHANES consists of a household screener, an interview, and an examination, including selected objective measures of health status. The primary objective of the household screener is to determine whether any household members are eligible for the interview and examination. The interview collects person-level demographic, health, and nutrition information, as well as information about the household. The examination includes physical measurements such as blood pressure, a dental examination, and the collection of blood and urine specimens for laboratory testing.

The sample design used since 1999 allows the production of aggregate-level national estimates from NHANES each year. Because NHANES can cover only a small number of PSUs each year, parameter estimates for single-year data are relatively unstable (large variance estimates). In addition, releasing only 1 year of data increases the possibility of disclosing a sample person's identity. These two factors resulted in the decision to publicly release data in 2-year cycles. Annual estimates may only be made through the NCHS Research Data Center and should be produced only for the nation as a whole, for race and Hispanic origin subdomains, or for very broad sex-age subdomains within race and Hispanic origins, due to limited sample sizes and larger variances of point estimates. To improve the statistical reliability and stability of estimates, using combinations of 2-year cycles is advisable. Combining data from 2-year cycles is particularly appropriate for rare events, estimates pertaining to detailed demographic subdomains, and measures that may have considerable geographic variation.

This report describes the broad design specifications for the 2011–2014 survey, including survey objectives, domain and precision considerations, and operational requirements. In addition, the report describes survey design details, including the calculation of sampling rates and sample selection methods.

More information about the 2011–2014 estimation procedures, the creation of weights for the entire sample and subsamples, and appropriate variance estimation methods to be used when analyzing NHANES data can be found in the forthcoming “National Health and Nutrition Examination Survey: Estimation Procedures, 2011–2014.” Information about the 2007–2010 estimation procedures can be found in the “National Health and Nutrition Examination Survey: Estimation Procedures, 2007–2010” report (11). Other documentation of the survey content, data collection procedures, and methods to assess nonsampling errors is reported elsewhere; visit <http://www.cdc.gov/nchs/nhanes.htm> for more information.

Design Specifications

Survey Objectives

A primary objective of NHANES 2011–2014 is to produce a broad range of descriptive health and nutrition statistics for sex, race and Hispanic origin, and age subdomains of the U.S. population. These data can then be used to measure and monitor the health and nutritional status of the civilian noninstitutionalized population. The analytic goals of NHANES are to:

- Estimate the number and percentage of persons in the U.S. population and in designated subgroups with selected diseases and risk factors.
- Monitor trends in the prevalence, awareness, treatment, and control of selected diseases.
- Monitor trends in risk behaviors and environmental exposures.
- Study the relationship among diet, nutrition, and health.
- Explore emerging public health issues and new technologies.
- Provide baseline health characteristics that can be linked to mortality data from the National Death Index or other administrative records (e.g., enrollment and claims data from the Centers for Medicare & Medicaid Services).

Domain and Precision Considerations

The set of domains for which specified reliability was desired in NHANES 2011–2014 consisted of sex-age groups for non-Hispanic black persons, non-Hispanic non-black Asian persons, Hispanic persons, and low- and non-low-income groups for the remainder of the U.S. population. [Table B](#) provides the set of sampling domains in NHANES 2011–2014.

To meet target specifications for the domains, race and Hispanic origin designations were made exclusive. The Hispanic category includes all Hispanic persons regardless of any other self-identified race. The non-Hispanic black category includes all non-Hispanic persons who self-identify as black regardless of any other self-identified race. The non-Hispanic non-black Asian category (hereafter referred to as Asian) includes all non-Hispanic persons who do not self-identify as black and have origins in any of the original peoples of the Far East, Southeast Asia, or the Indian subcontinent, including, for example, Cambodia, China, India, Japan, Korea, Malaysia, Pakistan, Philippine Islands, Thailand, and Vietnam. All other persons not falling into the categories above are considered to be in the “white and other” category.

To increase the precision of estimates for certain subdomains, oversampling was carried out for Hispanic, non-Hispanic black, Asian, and white and other persons at or below 130% of the federal poverty level, and for white and other persons aged 80 and over. Although data are released in 2-year cycles, the accumulation of at least 4 years of data is required to obtain an acceptable level of reliability for the domains given in [Table B](#). Thus, to create estimates for 2 years (or any annual estimates), collapsing of some sampling domains is necessary to produce adequate sample sizes for analysis.

The NHANES 2011–2014 sample was designed to produce data that would meet two conditions:

1. An estimated prevalence statistic of approximately 10% in a sex-age

Table A. Selected sample design parameters: Health and Nutrition Examination Surveys, 1971–2014

Characteristic	NHANES I 1971–1974	NHANES II 1976–1980	Hispanic HANES 1982–1984	NHANES III 1988–1994	NHANES 1999–2006	NHANES 2007–2010	NHANES 2011–2014
Age of noninstitutionalized civilian target population	1–74 years	6 months–74 years	6 months–74 years	2 months and over	All ages from birth	All ages from birth	All ages from birth
Geographic areas	United States (excluding Alaska and Hawaii)	United States	Southwest for Mexican-American population; NY, NJ, CT for Puerto Rican population; Dade County, FL, for Cuban population	United States	United States	United States	United States
Average number of sample persons per eligible household.	1	1	^{1,2} 2.71	^{1,2} 2.03	² 2.02	2	2
Number of study locations.	100	64	17 in Southwest; 9 in NY, NJ, CT; 4 in Dade County, FL	89	117	60	60
Domains for oversampling.	Low income (at or below 100% of federal poverty level); Children aged 1–5 years; Women aged 20–44 years; Persons aged 65–74 years	Low income (at or below 100% of federal poverty level); Children aged 6 months–5 years; Persons aged 60–74 years	Southwest; NY, NJ, CT; Dade County, FL; Persons aged 6 months–19 years and 45–74 years	Predesignated: 52 subdomains of sex-age groups for black, Mexican-American, and other persons Oversampled: Mexican-American persons, black persons, young children (under age 1 year), and adults aged 60 and over	Predesignated: 76 subdomains ² of sex-age groups for black persons and Mexican-American persons, and income-sex-age groups for other persons Oversampled: Mexican-American persons, black persons, low-income white and other persons (at or below 130% of federal poverty level), adolescents aged 12–19, and adults aged 70 and over. A supplemental sample included pregnant women	Predesignated: 72 subdomains ³ of sex-age groups for non-Hispanic black persons and Hispanic persons, and income-sex-age groups for other persons Oversampled: Hispanic persons, non-Hispanic black persons, low-income non-Hispanic white and other persons (at or below 130% of federal poverty level), and non-Hispanic white and other persons aged 80 and over	Predesignated: 87 subdomains [†] of sex-age groups for non-Hispanic black persons, non-Hispanic non-black Asian persons, and Hispanic persons, and income-sex-age groups for other persons Oversampled: Hispanic persons, non-Hispanic black persons, non-Hispanic non-black Asian persons, low-income non-Hispanic non-black non-Asian white and other persons (at or below 130% of federal poverty level), and adults aged 80 and over
Number of selected persons	28,043	27,801	15,924	39,695	50,939	26,215	[†] 27,631
Number of interviewed persons.	27,753	25,286	13,689	33,994	41,474	20,686	[†] 20,491
Number of examined persons.	20,749	20,322	11,653	30,818	39,352	20,015	[†] 19,644

[†] Age-sex domains were the same as defined for 2007–2010 with the addition of 15 domains for Asian.

[‡] Values include actual numbers for 2011–2012 and targeted numbers for 2013–2014.

¹ Average number of sampled persons per eligible family.

² In 1999, 53 subdomains were predesignated for black, Mexican-American, and other persons.

³ Compared with previous age-sex domains, age domains 12–15 and 16–19 were combined, and age-minority domain 40–59 was split into 10-year age domains (40–49 and 50–59).

NOTE: NHANES is National Health and Nutrition Examination Survey; Hispanic HANES is Hispanic Health and Nutrition Examination Survey.

Table B. Sampling subdomains classified by race and Hispanic origin, income, sex, and age: National Health and Nutrition Examination Survey, 2011–2014

Sex	Non-Hispanic black	Non-Hispanic non-black Asian	Hispanic	Non-Hispanic white and other	
				Non-low income	Low income ¹
Age (years)					
Males and females. . .	Under 1	Under 1	Under 1	Under 1	Under 1
	1–2	1–2	1–2	1–2	1–2
	3–5	3–5	3–5	3–5	3–5
Males	6–11	6–11	6–11	6–11	6–11
	12–19	12–19	12–19	12–19	12–19
	20–39	20–39	20–39	20–29	20–29
	...	...	...	30–39	30–39
	40–49	40–49	40–49	40–49	40–49
	50–59	50–59	50–59	50–59	50–59
	60 and over	60 and over	60 and over	60–69	60–69
	...	...	...	70–79	70–79
	...	...	...	80 and over	80 and over
Females	6–11	6–11	6–11	6–11	6–11
	12–19	12–19	12–19	12–19	12–19
	20–39	20–39	20–39	20–29	20–29
	...	...	...	30–39	30–39
	40–49	40–49	40–49	40–49	40–49
	50–59	50–59	50–59	50–59	50–59
	60 and over	60 and over	60 and over	60–69	60–69
	...	...	...	70–79	70–79
	...	...	...	80 and over	80 and over

... Category not applicable.
¹At or below 130% of the federal poverty level.

domain should have a relative standard error of 30% or less.

2. Estimated (absolute) differences between domains of at least 10% should be detectable with a Type I error rate (α) of 0.05 or less, and a Type II error rate (β) of 0.10 or less.

To satisfy the first condition, a sample size of about 150 examined persons is needed. This assumes a design effect of 1.5 resulting from the variability in sampling rates across density strata necessary to accommodate oversampling. The sample size needed to satisfy the second condition is about 420 examined persons.

These general sample-size considerations used in the sample design provide guidance on whether the data can meet analytic objectives. For example, for a very small demographic group, combining 4 years of NHANES data for a specific variable and a specific analysis may be necessary. However, the sample design effects for each measured NHANES variable, and for specific demographic subdomains, may be quite different from the assumed

general design effect of 1.5. The issues of precision and statistical power should be addressed for each specific analysis.

Operational Requirements

A unique feature of NHANES is the collection of physical examination data for each respondent in the sample. To standardize their administration, these examinations are carried out in MECs. Three separate MECs are in service at any given time. Following a carefully designed schedule, two MECs are in operation at study locations while the third is either traveling or being prepared for operation at a new location.

To maintain a cost-efficient workload within each location while considering the time and the cost involved in moving a MEC between study locations, the maximum number of NHANES study locations in each annual sample is 15. Based on this parameter, the number of sampled participants selected in each study location should be between 300 and 600, with an average of approximately 450, to yield approximately 333 examined persons in each of the 15

locations for that year.

Previous experience with other NHANES surveys indicates that response rates increase when a larger sample of persons is selected within households. One of the factors thought to be responsible for increased household response rates was that each person was given remuneration for his or her time and participation. Therefore, another important factor considered in the final design was to maximize response rates and reduce screening costs by selecting as large an average number of sampled participants per household as possible. Another factor affecting response rates is the amount of travel necessary for respondents to visit a MEC. The PSUs for NHANES are typically defined as individual counties, rather than combinations of counties as in other area surveys, to increase the likelihood of achieving high response rates.

Sample Design

The NHANES sample represents the total noninstitutionalized civilian U.S. population residing in the 50 states and District of Columbia. As with previous NHANES samples, a four-stage sample design was used in NHANES 2011–2014. The first stage consisted of selecting PSUs from a frame of all U.S. counties. The first-stage PSUs were mostly counties; in a few cases, adjacent counties were combined to keep PSUs above a certain minimum size. NHANES PSUs were selected with probabilities proportionate to a measure of size (PPS).

The second stage of selection for the NHANES 2011–2014 sample included a sample of area segments, comprising census blocks or combinations of blocks. However, because these samples were based on 2000 census data, the measure of size (MOS) used for sampling was updated as necessary for PSUs experiencing large growth since 2000.

The sample was designed to produce approximately equal sample sizes per PSU. Noncertainty PSUs have 24 segments. PSUs selected with

certainty (with a probability of one) may have more or fewer than 24 segments, to ensure appropriate representation in the sample. Additionally, some large certainty PSUs were treated as multiple study locations with varying numbers of segments in each location to, again, ensure appropriate representation of the PSU. The segments were also selected with PPS. The MOS of the segments, when combined with the subsampling rates used within the segments, provided approximately equal numbers of sampled participants per segment.

The third stage of sample selection consisted of dwelling units (DUs), including noninstitutional group quarters such as dormitories. In a given PSU, following the selection of segments, a listing of all DUs in the sampled segments was prepared, and a subsample of these was designated for screening to identify potential sampled participants. The subsampling rates were set up to produce a national, approximately equal probability sample of households. The screening rate was designed to produce the desired number of sampled participants for the most difficult race-Hispanic origin-sex-age-income domain (i.e., the domain sampled at the highest rate).

The fourth stage of sample selection consisted of persons within occupied DUs, or households. All eligible members within a household were listed, and a subsample of individuals was selected based on sex, age, race and Hispanic origin, and income. The subsampling rates and designation of potential sampled participants within screened households were arranged to provide approximately self-weighting samples for each subdomain and to maximize the average number of sampled participants per sample household.

Expected annual sample sizes at the design stage are:

● Study locations	15
● Segments	360
● DUs to be screened	13,529
● Households to be screened	11,500
● Sampled persons	6,888
● Examined persons	5,000

Sampling Rates

The rates required for sampling race-Hispanic origin-sex-age-income domains are the driving force in all stages of sampling for NHANES.

Calculation of screening amounts and sampling rates to achieve self-weighting domain samples

NHANES is a multistage, national area probability survey with fixed sample-size targets for sampling domains defined by race and Hispanic origin, sex, age, and low-income status. Thus, the first step in determining the MOS to be used for sampling at each stage is to calculate the sampling rate for each domain. The sampling rate for a domain depends on the target examination sample size, the expected examination response rate, and the estimated population size. These sampling rates determine the amount of screening that will be required.

To calculate sampling rates, it is necessary to set expectations for response rates. Because this was the first design to include an oversample of the Asian population, it was determined that the estimates should be conservative, especially for those domains. The projected response rate for a sampling domain was set equal to the minimum of the unweighted achieved response rate for that domain for periods 2007–2008, 2006–2008, and 2003–2008, except in four Asian domains (males aged 40–49 and females aged 6–11, 12–19, and 40–49) where the minimum response rate was based on a sample size of less than 30. In three of those domains, the response rate from the time span with the next minimum value was used. In the fourth domain (Asian females aged 6–11), response rates were inconsistent with those of the other Asian domains and based on small sample sizes, so the projected response rate was set equal to that for Asian males aged 6–11. The response rates used in the calculations ranged from 50% to 91%, with the lowest response rates assigned to the most challenging sampling domains, such as persons aged 80 and over.

Several data sources were used to obtain national estimates of the noninstitutionalized civilian population by race and Hispanic origin for the NHANES 2011–2014 sample. At the time of PSU selection, population projections were not available for certain U.S. subpopulations such as noninstitutional civilian residents; multiple data sources were needed to create these estimates. Sex-age distributions for Hispanic, non-Hispanic black, and Asian persons, as well as for the entire U.S. population, were from the U.S. Census Bureau's projections of the total resident U.S. population. The proportion of the total resident population that is civilian and living outside of institutions in each sex-age domain was calculated from the census population estimates for November 2008. Race and ethnicity data from the American Community Survey (ACS) 2007 Public Use Microdata Sample File were used to make adjustments to the Asian population totals. Because the census population projections and population estimates do not include totals for non-Hispanic non-black Asian persons alone or in combination (AOIC), the totals for non-Hispanic Asian AOIC were adjusted by the ACS estimates for the proportion of this group that is not black. National poverty estimates for white and other persons (non-black, non-Hispanic, non-Asian) were derived from the Census Bureau's March 2008 Current Population Survey (CPS).

The information in [Table I](#) (Appendix II) was used to determine the overall sample size required to meet each domain target in NHANES 2011–2014. The domains requiring the most screening are low-income white and other persons aged 80 and over, and the Asian domains for age groups under 12, 60 and over, and males aged 12–19. Examination targets for these domains were set equal to the maximum attainable for a sample of approximately 11,500 screened households, the level set for NHANES 2007–2010.

A within-domain self-weighting sample of approximately 13,529 DUs per year, resulting in 11,500 screened households, was used for NHANES 2011–2014 (assuming that 15% of the

DUs were either vacant or included in the sample in error, such as a structure thought to be a residence that was actually a business). The total expected screening for the 4-year sample is 46,000 households. However, to ensure that the targeted number of sampled persons could be reached in the event of unexpected occurrences during the field period, additional households were selected and held in reserve. The following is a derivation of the original screening rate for the full 4-year sample:

- Screening sample size for 4 years: 46,000 households
- Eighty percent add-on for reserve: 82,800 households
- Projected total U.S. households in 2011–2014: 118,410,614 households
- Maximum sampling rate: $82,800 / 118,410,614 \approx 1/1,430$.

Final sampling rates

Table II (Appendix II) shows the sampling rates used for the selection of PSUs in NHANES 2011–2014 for each of the sampling domains. The sampling rates given in Table I (Appendix II) were designed to provide an 80% reserve sample, as well as a provision for expected nonresponse in each subdomain.

Sampling rates were calculated using the approach described in the preceding section, “Calculation of screening amounts and sampling rates to achieve self-weighting domain samples.” All screened persons in the subdomain having that maximum rate were to be retained in the sample. The screened persons in other subdomains were to be subsampled to bring the sampling rates for those subdomains down to the desired levels. The subsampling rates were designed to minimize the variability in sampling rates among strata while still achieving the desired precision.

Departures from self-weighting sample within domains

Calculating the sampling rates required several assumptions related to population size and response rates. To the extent that these assumptions were

not met, the actual screening required to reach the target sample sizes differed from the expected screening. As stated in the previous section on calculation of screening amounts and sampling rates, several data sources were used to develop the 2011–2014 population projections used in the sampling-rate calculations, including the Census Bureau’s projections of the resident population by age, sex, and race and Hispanic origin; its 2008 postcensal estimates of the resident and noninstitutionalized civilian population; the ACS 2007 data on race and ethnicity; and the March 2008 CPS national poverty estimates for white and other persons. Estimates of occupied housing units by race and Hispanic origin from the 2007 ACS also were used to project total occupied housing units for 2011–2014, and the national vacancy rate (15%) was from the 2005–2007 ACS. The population projections and resulting expected screening requirement numbers depended on the assumption that these proportions continued to hold in the years of data collection.

Finally, as noted in the previous section on calculation of screening amounts and sampling rates, the expected examination response rates were set equal to achieved examination response rates by domain for earlier years of NHANES, and response rates for Mexican-American persons were used to estimate those for Hispanic persons. Screening requirements also varied from expectations depending on how much these earlier response rates differed from the actual experience in 2011–2014.

Sampling rate modifications for NHANES 2013–2014 sample

Introduction of the Asian oversample resulted in a decrease in the number of examinations in the Hispanic domains, which are among the higher-responding domains in NHANES. Additionally, because approximately 15% of the targeted examination sample was in Asian domains, which required a larger screening effort than anticipated,

this design shift made the overall target number of examinations more difficult to achieve than in previous surveys.

In 2012, it appeared that the response rate patterns in this new design were such that the targeted 10,000 examinations would not be met for the 2011–2012 sample. Anticipating similar response rate patterns in the 2013–2014 sample, NCHS reduced the amount of screening required for the hardest-to-fill non-Asian domain. The result was a drop in the expected number of Asian examinations in the 4-year sample, from 15% to 13%.

Because reducing the amount of screening reduced the number of examinations overall, targeted examinations removed from the Asian domains were distributed to Hispanic and non-low-income white and other domains, as these were expected to have the largest deficits at the end of the 4 years. These targeted examinations were distributed inversely proportional to the numerator of the sampling rate, so that easier-to-fill domains were given more examinations than the other domains. No additional examinations were targeted for Hispanic infants (birth–11 months), because all enumerated infants in this domain are already included in the sample.

To attain the target of 20,000 examinations over the 4 years, an additional target of 150 examinations was added in the annual sample for 2013–2014. A similar increase was used for the 2009 and 2010 samples to compensate for the deficit experienced in 2007–2008. Because the estimated deficit for the 4-year sample was primarily in the Hispanic and white and other domains, the additional 150 examinations in each annual sample were distributed to these groups only. Again, the targets were distributed inversely proportional to the sampling rate (i.e., the higher the sampling rate, the smaller the number of additional examinations targeted). The revised sampling rates are included in Table II (Appendix II).

Sample Selection Methods

Stratification and selection of PSUs

The operational requirements for NHANES are such that the amount of travel necessary for a sampled participant to visit a MEC should be minimized to increase the likelihood of their visiting the MEC, leading to high response rates overall. As a result, individual counties were chosen as the PSUs for NHANES. However, some counties have such small populations that their probabilities of selection would be lower than what is required to attain the sampling rates for some of the domains. If selected for the sample, they would introduce considerable variability into the weights. Consequently, these small counties were combined with one or more adjacent counties to form more efficient sampling units. For the same reason, independent cities in Virginia were combined with nearby counties.

The frame for NHANES 2011–2014 included all counties in the entire country. From the approximately 3,100 counties and county equivalents in the United States, 2,846 PSUs were formed (most of which consisted of individual counties), a sample of 60 study locations was selected, and 15 of these locations per year were randomly allocated to each of the years.

Calculation of PSU MOS

The NHANES sample was designed to yield a self-weighting sample for each sampling domain while producing an efficient workload in each study location. PSUs were selected with probabilities proportionate to an MOS. The selection probability of a PSU determines the maximum rate at which persons residing in that particular PSU can be selected.

The expression used to define the PSU MOS is similar to that used in previous years. The MOS of PSU h , denoted by M_h , is a weighted average of estimated populations by race and Hispanic origin, calculated as

$$M_h = \sum_k A_k C_{hk}$$

and

$$A_k = \sum_l r_{kl} \frac{C_{.kl}^*}{C_{.k}^*}$$

where

k = Race-Hispanic origin-income subdomain

l = Sex-age subdomain

C_{hk} = Most recent population estimate for race-Hispanic origin-income subdomain k in PSU h (see below)

r_{kl} = Sampling rate of persons in the (k,l) -th race-Hispanic origin-income-sex-age subdomain

$C_{.kl}^*$ = Most recent projection of the 2008 total population count for race-Hispanic origin-income-sex-age subdomain (k,l)

$C_{.k}^*$ = Most recent projection of the 2008 total population count for race-Hispanic origin-income subdomain k

Because single counties, rather than larger areas made up of groups of counties, are optimal as NHANES PSUs, M_h was first calculated with h representing a single county.

At the time of PSU selection for the 2011–2014 PSUs, the most recent county-level estimates of population by race and Hispanic origin were from the 2008 Census Bureau population estimates. To obtain the estimates for C_{hk} , these were adjusted by census 2000 county-level estimates of the proportion of the population not living in institutional group quarters in 2000, and estimates for white and other persons were split into low income and non-low income based on the census 2000 county-level estimates of the proportion of non-Hispanic white persons with a 1999 income below the federal poverty level.

The factors A_k (Table C) are the weights used to assign the relative contribution from each race-Hispanic origin group in the computation of MOS.

Minimum MOS

The selection probability of a PSU determines the maximum rate at which persons residing in that particular PSU can be selected for NHANES while retaining the self-weighting nature of the sample. If the MOS of a PSU is too small, the required sampling rates for some subdomains cannot be achieved. Consequently, special weighting procedures would be required for such PSUs, and the resulting variability in weights would increase sampling errors. To ensure that all required sampling rates could be achieved, counties with a very small MOS were combined with other adjacent counties.

The condition that determines the minimum MOS of a PSU is

$$P_h \geq \hat{r} \text{ for all } h$$

where

P_h = Probability of selecting PSU h

$\hat{r}$ = Maximum sampling rate among the sampling domains

For certainty PSUs, this condition always holds, because $\hat{r} \leq 1$ and $P_h = 1$. For noncertainty PSUs, the probability of selecting PSU h is

$$P_h = c_{NC} \frac{M_h}{\sum_{h \in NC} M_h}$$

where

NC = Set of noncertainty counties

C_{NC} = Number of noncertainty PSUs to be selected

M_h = MOS for PSU h

Table C. Values of A_k used in calculating primary sampling unit measures of size

Race and Hispanic origin	A_k
Hispanic	0.000211
Non-Hispanic black.	0.000286
Non-Hispanic non-black Asian	0.000507
Non-Hispanic, low income, white and other.	0.000235
Non-Hispanic, non-low income, white and other	0.000071

Thus, the condition that determines the minimum MOS is equivalent to

$$M_h \geq \hat{r} \frac{\Sigma_{heNC} M_h}{c_{NC}}$$

For each county, it was necessary to check whether the MOS of the county met the minimum MOS condition. Because the righthand side is a constant, the first step in this check was to compute this product. The number of noncertainty locations, c_{NC} , was 52.

The second term on the righthand side of the expression was found to be

$$\frac{\Sigma_{heNC} M_h}{c_{NC}} = 761.8922$$

Based on this minimum MOS criterion, 328 counties were found to have an MOS that was too small. These counties were combined with neighboring noncertainty counties. The neighboring counties had to be adjacent, and the maximum distance between any two points in the combined-county area had to be less than 125 miles. Unless no alternatives met the aforementioned criteria, counties combined were also from the same state.

A total of 123 counties were combined into PSUs that either consisted of three or more counties or had a maximum distance greater than 125 miles. To avoid the complexity of working with more than two county administrations, and to reduce the listing and interviewing cost associated with traveling, the maps of these counties were reviewed and some were manually recombined.

After the necessary county combinations were made, the PSU MOS, M_h , was recalculated with h representing the combined counties as a single PSU.

Selection of certainty PSUs

Some counties had an MOS large enough that they were selected with certainty, and a few of these were selected multiple times. These certainty PSUs were removed from the county frame prior to noncertainty PSU selection.

A PSU was identified as a certainty if its weighted MOS exceeded 75% of the initial sampling interval; that is,

PSU h was included in the sample with certainty if

$$M_h > 0.75 \frac{\sum_{h=1}^H M_h}{60}$$

where H is the number of PSUs on the entire sampling frame.

Some certainty PSUs were so large that they warranted more than one study location; otherwise, weighting factors would have to be applied to ensure appropriate representation, and these weighting factors would reduce the efficiency of estimates. The number of study locations allocated to each certainty PSU was obtained by comparing the weighted MOS, M_h , for the PSU to the initial PSU sampling interval, $(1/60) \sum_{h=1}^H M_h$.

A total of 8 study locations in the full NHANES 2011–2014 out of the 60-location sample were assigned to certainty PSUs. These locations were in six counties; one county contained multiple study locations.

Stratification

The stratification scheme for NHANES 2011–2014 PSUs was developed with the primary goal of efficiency for the 4-year sample, and secondary goals of efficiency for 2-year and annual samples. For the 4-year sample, 13 major strata were defined based on state groupings and the urban-rural population distribution of the PSUs. Fifty-two minor strata were defined based on the demographics of the PSUs. Each major stratum included four minor strata, and one PSU was selected from each of these final strata.

The 4-year sample had a one-PSU-per-minor stratum design; each annual sample had a one-PSU-per-major stratum design. Two-year samples have a two-PSU-per-major stratum design, with each PSU coming from different minor strata. These major strata were also used as the strata for variance estimation. However, because certainty PSUs are not selected within strata, variance strata for these PSUs are formed based on the size of the PSU relative to the other PSUs. As a result, some of these variance strata may have up to three PSUs for variance estimation depending on the number and size of

the certainty locations that year. That is, all multiyear samples contained only one PSU per sampled minor stratum rather than multiple PSUs from the same stratum.

The NHANES 2011–2014 PSU sample selection employed a stratification scheme different from NHANES 2007–2010 in that, instead of using the census region and metropolitan status of the PSUs, the 13 major strata for the noncertainty PSUs were based on state groups that were relatively homogeneous in their derived health ranking (Figure). The state-level health-related variables considered were death rate (12), infant mortality rate (13), percentage of adults with high blood pressure (14), percentage of adults overweight or obese, percentage of adults with poor nutrition, and percentage of adults who smoke (15); all were available from the CDC. States were grouped based on the results from cluster and factor analyses. For most of the states, results using the two methods were consistent or close, and group assignment was straightforward. For a small number of states with inconsistent results, group membership was assigned by judgment based on reviewing individual health indicators.

Each of the 51 state-level units (one for each of the 50 states and the District of Columbia) was assigned to state groups, with the first group consisting of states with the highest health levels and the last group including those with the lowest health levels. Initially, four groups were formed. California, originally in the first group, was treated as a separate single group because of interest in subnational, multiyear estimates for a few larger states. With California as a separate group, a total of five final state groups were formed (Table D).

Thirteen major strata were defined by sorting the noncertainty PSUs by various geographical and urban-rural measures of the PSUs within each of the five state groups. Within each group, noncertainty PSUs were sorted by census region and percentage of population living in rural areas, using a serpentine sort method to place the noncertainty PSUs in alternating order with respect to each variable.

Table D. State grouping for primary sampling unit stratification

State group	State abbreviation
1	CT, HI, IA, MA, MN, ND, NH, NJ, NY, RI, UT, VT, WA
2	CA
3	AK, AZ, CO, FL, ID, IL, KS, ME, MT, NE, NM, OR, SD, VA, WI, WY
4	DE, IN, MD, MI, OH, PA, TX
5	AL, AR, DC, GA, KY, LA, MO, MS, NC, NV, OK, SC, TN, WV

After sorting, the desired number of major strata for each state group was formed by cumulating the MOS as close as possible to the desired mean size for that state group. Once the major strata were established, four minor strata were created within each major stratum. The variables used in forming the minor strata were minority concentration, percentage of rural population, and percentage of white and other population with income below the federal poverty level.

Two methods for forming the minor strata were compared:

- Nested stratification approach—This method (16) searched for an optimal stratification to minimize a weighted combination of the equal stratum size measure and the Euclidean distance measure. The Euclidean distance measure was constructed as a weighted combination of the demographic variables.
- Tree search algorithm—In defining minor strata, the algorithm (17)

followed a series of steps to sequentially split the group of PSUs within each major stratum to minimize the variance of the subpopulation (Hispanic persons, black persons, and Asian persons, among others) with the largest relative variance.

Based on a comparison of the two methods of minor strata formation, the minor strata formed by the two methods were combined to create “hybrid” minor strata as the final strata for noncertainty PSU selection. That is, instead of choosing one method as the final stratification scheme for all minor strata, simulation results were compared for each individual major stratum, and the method that had an overall better performance in controlling the variation in the relative difference between the expected and the targeted number of sample persons for the minority groups was chosen for that major stratum.

Table E presents characteristics of the PSUs within each major stratum.

Selection of noncertainty PSUs

To improve the geographic distribution and diversity of the sample and to limit the burden imposed on any particular PSU over time, the NHANES 2011–2014 noncertainty PSUs were selected so as to minimize the number of PSUs that were in both the 2007–2010 and 2011–2014 samples. This was achieved by using Ohlsson’s method, which is applicable to PPS samples with one unit selected per stratum and can accommodate changes in design across samples (18). The method uses permanent random numbers and exponential sampling to minimize overlap between two or more samples.

The first step in implementing Ohlsson’s method was to retrospectively assign a number between 0 and 1, called a permanent random number (PRN), to each noncertainty county on the 2007–2010 sampling frame. The PRN value depended on whether the county had been selected for the previous sample. PRNs were then transferred to the corresponding counties on the 2011–2014 frame. Any noncertainty counties in the 2007–2010 frame that had been selected with certainty in the previous sample were randomly assigned a new PRN from the uniform distribution; see Ohlsson (18) for a complete description of retrospectively assigning PRNs.

To minimize overlap between the two samples, the PRNs were shifted by

Table E. Characteristics of major strata formed for selection of primary sampling units: National Health and Nutrition Examination Survey, 2011–2014

State group	Major strata	Number of PSUs ¹	Number of minor strata	Census region	Mean percent rural population ²
1	11	23	4	Northeast	1.7
	12	313	4	Northeast, Midwest	28.0
	13 (one-half)	66	2	West	15.3
2	21	14	4	West	3.2
	22 (one-half)	40	2	West	17.1
3	31	34	4	South	4.6
	32	378	4	Midwest, West	23.4
	33	353	4	Northeast, Midwest, South	29.8
4	41	30	4	South	5.4
	42	42	4	Northeast, Midwest	6.8
	43	499	4	Northeast, Midwest, South	45.3
5	51	35	4	Midwest, South, West	3.8
	52	192	4	Midwest, South, West	26.7
	53	821	4	Midwest, South, West	67.8

¹PSU is primary sampling unit.

²Weighted by 2008 civilian noninstitutional population estimates of noncertainty PSUs.

one-half; that is, the goal was to minimize the overlap only between two successive samples, 2007–2010 and 2011–2014. To minimize overlap between more than two samples, a PRN adjustment of less than one-half would be used. PRNs were then transformed as

$$\xi_h = -\frac{\log(1 - X_h)}{P_h}$$

where

ξ_h = Transformed PRN for county h

X_h = Shifted PRN for county h

$$P_h = \frac{M_h}{\sum_{v \in H_h} M_v} = \frac{\text{MOS of county } h}{\text{Total MOS of stratum } H \text{ containing county } h}$$

Next, the PSU in each stratum with the minimum transformed random number was selected for the NHANES 2011–2014 sample.

The combinations of counties not meeting the minimum MOS criterion into PSUs for 2011–2014 were not identical to those in 2007–2010, due to differences in county MOS as a result of population changes and the addition of Asian sampling domains. To minimize overlap between the two samples, consistent sampling units were desirable. Therefore, counties or county equivalents were used as the sampling unit in Ohlsson's algorithm rather than PSUs. For 2011–2014, a PSU was selected if any counties in the PSU were selected. The selection probability for the PSU was the sum of the selection probabilities for the individual counties in the PSU.

The resulting sample consisted of one PSU from each stratum, with selection probabilities, P_h , proportional to the PSU MOS. In addition, as a result of the shifting of the PRN, the conditional probability of selecting a PSU for the 2011–2014 sample, given that it was selected for the previous sample, was less than if the samples had been selected independently.

Allocation of PSUs to time period

To have nationally representative annual samples, which is a design requirement of NHANES, study

locations had to be assigned to years randomly. One way to achieve nationally representative annual samples would be to select an independent sample of PSUs each year. This approach would lead to substantial overlap in PSUs each year, and the overlap could lead to increased clustering of the sample, resulting in less precise estimates. In addition, each annual sample can have only 5,000 examined persons in 15 study locations. Therefore, a 4-year sample in NHANES was selected from a nested structure of major and minor strata. This stratification scheme was developed so that annual and multiyear samples were distributed evenly in terms of geography and certain population characteristics.

The certainty PSUs were first sorted according to their MOS, then assigned in a manner that appropriately reflected their relative size. One PSU was large enough to be selected with certainty in 3 years of the 4-year sample and contained three study locations within the 60-location sample. In 1999, this PSU was divided into three study locations along tract boundaries: the northeastern, southern, and northwestern areas of the county. These locations have been and will continue to be fielded in this same order so that 1 of the 4 years is randomly selected to not contain one of these locations. These three pseudo-PSUs were then combined with the remaining five PSUs that were certainties in the 4-year sample, and the full set of eight study locations were paired off with each pair assigned to a year for geographic distribution across the 2-year samples.

Noncertainty PSUs were sorted by major and minor strata. The PSUs were then paired so that the first and third minor strata within a major stratum were in one pair, and the other two minor strata were in another pair. Then the year was assigned by a random process that provided each year with one PSU from each major stratum. Additionally, the adjacent minor strata 1 and 2, and 3 and 4, were never placed in the same 2-year sample. This method of allocation ensured that the noncertainty PSUs in the same major stratum were evenly distributed across the 2-year samples.

Targeted number of sampled persons in each PSU

The initial target number of examined persons per location was 333, based on the assumption of a total of 5,000 examined persons per year in 15 study locations. Once the sample of locations was selected, the examination targets were adjusted. For certainty locations, the initial target was adjusted by the relative size of the location to obtain the final target number of examined persons. This was calculated as the MOS allocated to the study location divided by the initial sampling interval, $(1/60) \sum_{h=1}^H M_h$. The relative size of certainty locations ranged from 0.76 to 1.03.

For all noncertainty locations, the initial examination target was adjusted by the relative contribution of the location's stratum to the total noncertainty MOS. The final target number of identified sampled participants for a given study location was derived by inflating the desired number of examined persons to account for the predicted combined screener, interview, and examination response rate for the study location.

NHANES response rates (combined screener, interview, and examination) for each location in an annual sample were predicted using a linear regression based on the actual response rates and location-level characteristics of prior study locations. Prediction based on previous experience has proven more accurate than simply applying a single response rate across all study locations.

Each year, the model was refitted with the most recent data available at the time of the prediction. A relatively large number of geographic, demographic, and economic variables from the Census Bureau were assembled and brought into a linear regression model as potential independent variables, with the study location response rate as the dependent variable. A stepwise regression was used. The final model for each year was decided based on a combination of the regression correlation coefficient and a statistic that adjusts for the total number of variables included in the model. The model was applied to the values of the

selected variables for the current year's study locations to predict their response rates.

After these predicted response rates were reviewed with senior project staff, some of the rates were adjusted based on past experience. Once the response rate predictions were final, they were applied to the target number of examinations to predict the number of identified sampled participants that would be required in each study location to achieve the targets. This became the initial target number of identified sampled participants.

Selection of segments

The second stage of the design involved sampling segments within each PSU. Segments were selected in a continuous process about 5–6 months prior to the start of the field period for the study location.

The usual practice in area samples is to list all DUs in sampled segments and apply a prespecified sampling rate to the listed DUs. This approach gives all DUs the desired probabilities of selection. For example, if the sampling rate is 50%, then one-half of the DUs listed in the segments will be included in the sample. If the number of DUs has tripled due to new construction (i.e., housing units built since the most recent decennial census), the same sampling rate will produce three times as many interviews and examinations as the number originally expected. Such dramatic changes in the segment size are expected when the data collection period is several years after the most recent decennial census for which data files are available.

If the segment contains much new construction, the segment MOS may be inaccurate. As a result, either a larger-than-expected sample must be drawn from that segment or a weighting factor must be applied to all sampled participants from that segment. Because highly variable sample sizes are not operationally feasible for NHANES, subsampling within PSUs would be necessary to attain equal sample sizes across PSUs. However, this would require the application of a weighting factor, which would reduce the

efficiency of the sample.

To update a sampling frame when the sample is to be selected with respect to an MOS but a reliable estimate of the MOS is not available, double (or two-phase) sampling can be used. This was executed in NHANES, by first having the home office and field staff determine the number of DUs in the larger-than-needed first-phase sample of segments through a combination of digital and windshield canvassing. Then an updated MOS was calculated, reflecting the ratio of the actual number of DUs to the expected number of DUs. The final sample of segments was selected by subsampling from the first-phase segments using the updated MOS.

Stratification within PSUs

The procedures for selecting the segment sample involve implicit stratification by minority density and geography. To keep combined blocks within a single block group, the stratification is based on characteristics of the block group in which segments are located. Within the geographical strata, implicit stratification is created by sorting the area segments by minority density, tract number, census-designated place, block group within tract, and segment number within block group, and selecting a systematic sample with PPS.

MOS of segments

The segment MOS calculation is similar to that for the PSU MOS. The data used to calculate the MOS at the segment level are the most recent decennial census data available at the time of segment selection. Because the decennial census data are not made available until 1 year after the census-taking, and segment selection begins 5–6 months prior to the start of the field period for the study location, segments in the first 12 locations of the 2011–2014 sample were selected using 2000 census data. A two-phase sampling procedure, in which the MOS is updated for growth at the second stage, was followed in any study locations that experienced significant growth between the 2000 decennial census and the segment selection. The more recent

2010 census data were used to form segments in the other 48 study locations.

Prior research on intraclass correlations and unit costs indicates that an average of 14 examined sampled participants per segment was reasonably close to an optimum for most statistics in NHANES. As indicated earlier, operational requirements make it necessary to have a fairly constant number of examined sampled participants per study location, usually about 333. This implies having 24 segments per PSU.

Because segments consist of census blocks or groups of blocks, the segment MOS is a sum of MOSs calculated at the block level. For the first phase, let $M_{hb(1)}$ denote the MOS of block b in PSU h , so that

$$M_{hb(1)} = \sum_{k^*} A_{k^*} C_{hbk^*}$$

$$A_{k^*} = \sum_l r_{k^*l} \frac{C_{.k^*l}^*}{C_{.k^*}^*}$$

where

h = PSU

b = Block

k^* = Race-Hispanic origin subdomain (non-Hispanic black, Hispanic, and white-and-other income levels combined)

C_{hbk^*} = 2000 population of race-Hispanic origin k^* in block b in PSU h

r_{k^*l} = Sampling rate of persons in the (k^*, l) -th race-Hispanic origin-sex-age subdomain

$C_{.k^*l}^*$ = Most recent projection of the total population count for race-Hispanic origin-sex-age subdomain (k^*, l)

$C_{.k^*}^*$ = Most recent projection of the total population count for race-Hispanic origin subdomain k^*

For the non-Hispanic black, non-Hispanic non-black Asian, and Hispanic subdomains, the factor A_{k^*} is the same as A_k used in the PSU sampling MOS calculation described in the previous "Calculation of PSU MOS" section. Because income level is

not available at the block level, the value of A_{k*} for all non-Hispanic white and other persons was calculated as a weighted average of the A_k values used in the PSU sampling for low-income and non-low-income white and other persons. For the 2011–2012 sample, A_{k*} for white and other persons was calculated as 0.000091. With the sampling rate change for 2013–2014, the value of A_{k*} increased to 0.000103.

The MOS for the first-phase segments $j^{[1]}$ are the sums of the MOS of the block(s) comprising each segment. These MOSs are denoted by $M_{hj^{[1]}(1)}$.

The MOS used for the second-phase segment selection is the segment growth rate. Based on the DU counts obtained for the first-phase sample segments, the growth rate is estimated by computing the ratio of the actual number of DUs (determined by home-office and field staff through a combination of digital and windshield canvassing) to the expected number of DUs in the segment (based on the 2000 decennial census data).

Let $U'_{hj^{[1]}}$ denote the number of DUs found by NHANES staff in the field for the first-phase segment $j^{[1]}$ in PSU h . The growth ($g_{hj^{[1]}}$) of the first-phase segment $j^{[1]}$ is estimated as:

$$g_{hj^{[1]}} = \frac{U'_{hj^{[1]}}}{U_{hj^{[1]}}^{[0]}}$$

where $U_{hj^{[1]}}^{[0]}$ is the number of DUs in segment $j^{[1]}$ according to the 2000 decennial census. If significant growth was not experienced in the location between the census-taking and segment selection, $g_{hj^{[1]}} = 1$.

Thus, the second-phase MOS for segment $j^{[1]}$ selected in the first phase of sampling in PSU h is equal to

$$M_{hj^{[1]}(2)} = g_{hj^{[1]}}$$

Number of segments and probability of selection

As discussed in the previous “Targeted number of sample persons in each PSU” section, the person sample sizes for some study locations selected with certainty were adjusted to account for their size relative to the other

selected locations, to minimize the effects of intraclass correlation. The number of segments selected in the certainty locations was also adjusted from 24 to account for the relative size of the location. As a result, some study locations selected with certainty contained as few as 17 segments or as many as 24 segments in the second phase of segment sampling, denoted as $n_{h(2)}$. To achieve proper within-segment sampling rates in study locations in which the segment sample was selected in two phases, the first-phase sample must be larger than the ultimate sample. In NHANES 2011–2014, 100 segments were selected initially in the study locations requiring the two-phase segment selection procedure.

For each study location, the conditional probability of selection of first-phase segment $j^{[1]}$ is

$$P_{hj^{[1]}(1)} = \min \left[\frac{n_{h(1)} M_{hj^{[1]}(1)}}{\sum_{j^{[1]}=1}^{N_{h(1)}} M_{hj^{[1]}(1)}}, 1 \right]$$

where

$n_{h(1)}$ = Total number of first-phase segments to be selected in the h th PSU

$N_{h(1)}$ = Total number of segments in first-phase segment frame in the h th PSU

$M_{hj^{[1]}(1)}$ = First-phase MOS of segment $j^{[1]}$ in the h th PSU

Given the first-phase segments, the conditional selection probability of second-phase segment $j^{[2]}$ is

$$P_{hj^{[2]}(2)} = \min \left[\frac{n_{h(2)} M_{hj^{[2]}(2)}}{\sum_{j^{[2]}=1}^{n_{h(1)}} M_{hj^{[2]}(2)}}, 1 \right]$$

where

$n_{h(2)}$ = Total number of second-phase segments to be selected in the h th PSU

$M_{hj^{[2]}(2)}$ = Second-phase MOS of segment $j^{[2]}$ in the h th PSU

The actual probability of selection of a segment depends on the MOS of the segment and the probability of selection of the location from which the segment is selected. So the overall probability of selection of a second-phase segment $j^{[2]}$ is

$$P_{hj^{[2]}(2)} P_{hj^{[1]}(1)} P_h = \frac{n_{h(2)} n_{h(1)}}{n_h(2) n_h(1)} \left(\frac{M_{hj^{[2]}(2)}}{\sum_{j^{[2]}=1}^{n_{h(1)}} M_{hj^{[2]}(2)}} \right) \left(\frac{M_{hj^{[1]}(1)}}{\sum_{j^{[1]}=1}^{N_{h(1)}} M_{hj^{[1]}(1)}} \right) P_h \quad [1]$$

Note that in study locations that do not require the two-phase procedure, $n_{h(2)} = n_{h(1)} = n_h$. Moreover, in study locations where no second phase is needed, $M_{hj^{[2]}(2)} = 1$. Substituting these into the second-phase selection probability above results in:

$$P_{hj^{[2]}(2)} = \min \left(\frac{n_{h^*}(1)}{\sum_{j=1}^{n_h} (1)}, 1 \right) = \min \left(\frac{n_h}{n_h}, 1 \right) = 1$$

which indicates that the segment probability of selection within one-phase locations is simply the first-phase probability of selection.

Minimum MOS of segments

One of the goals of the sample design is to create equal probabilities of selection for each domain within a study location. This enables the selection of a nearly within-domain self-weighting sample and facilitates the selection of persons. To create equal probabilities, the within-segment sampling rate for a domain in study locations selected without certainty should be

$$r_{hkl} = \frac{r_{kl}}{P_h P_{hj^{[1]}(1)} P_{hj^{[2]}(2)}} \quad [2]$$

For locations selected with certainty, $P_h = 1$, so the within-segment sampling rate should be

$$r_{hkl} = \frac{r_{kl}}{P_{hj^{[1]}(1)} P_{hj^{[2]}(2)}}$$

The within-segment rates must be less than or equal to 1. The most severe constraint is for domains with the highest value of r_{kl} . These maximum sampling rates are known as $\hat{r}$, that is,

$$\hat{r} = \max_{k,l} \{r_{kl}\}$$

Therefore,

$$\max_{k,l} \{r_{hkl}\} = \frac{\hat{r}}{P_h P_{hj^{[1]}(1)} P_{hj^{[2]}(2)}} \leq 1 \quad [3]$$

Replacing the denominator in equation [3] with its equivalent in equation [1], the condition in equation [3] becomes

$$\frac{\hat{r}}{n_{h(2)}n_{h(1)}\left(\frac{M_{hj^{2}}}{\sum_{j^{[2]=1}^{n_{h(1)}}M_{hj^{2}}}\right)\left(\frac{M_{hj^{1}}}{\sum_{j^{[1]=1}^{N_{h(1)}}M_{hj^{1}}}\right)}P_h \leq 1$$

which is equivalent to

$$M_{hj^{1}} \geq \left(\frac{\hat{r}}{P_h \frac{n_{h(2)}}{\sum_{j^{[1]=1}^{N_{h(1)}}M_{hj^{1}}}}\right) \cdot \left(\frac{\sum_{j^{[2]=1}^{n_{h(1)}}M_{hj^{2}}}{n_{h(1)}M_{hj^{2}}}\right) \quad [4]$$

Consequently, the first-phase minimum MOS is a product of two factors, (a) the first factor calculated for the study location based on known information, and (b) the second factor based on the second-phase MOS, which was not known at the time of selection of the first-phase segments.

For study locations that do not require the two-phase process, the second factor reduces to 1:

$$\frac{\sum_{j^{[2]=1}^{n_{h(1)}}M_{hj^{2}}}{n_{h(1)}M_{hj^{2}}} = \frac{\sum_{j^{[2]=1}^{n_{h(1)}}1}{n_{h*1}} = \frac{n_h}{n_h} = 1$$

and equation [4] reduces to

$$M_{hj^{1}} \geq \frac{\hat{r}}{P_h \frac{n_h}{\sum_{j^{[1]=1}^{N_{h(1)}}M_{hj^{1}}}} \quad [5]$$

which is the minimum MOS for segments.

For study locations where the full two-phase process was implemented, $M_{hj^{2}}$ was not known when the first-phase segment was selected. In this case, the second factor must be considered as

$$\frac{\sum_{j=1}^{n_{h(1)}}M_{hj^{2}}}{n_{h(1)}M_{hj^{2}}} = \frac{ave(M_{hj^{2}})}{M_{hj^{2}}}$$

This factor would inflate the minimum MOS to account for expected growth in the segment due to new construction. Because the actual values

are not known, an inflation factor constant across all segments was used. Based on empirical research, this inflation factor was set at 1.25. In implementing the sample selection, the minimum MOS was made 80% greater than needed.

Within each PSU, the blocks reported on the block-level census files were sorted by minority density, tract, census-designated place, block group, and block number. Blocks with MOS below the minimum were combined with succeeding blocks until the desired measure was achieved. To the extent possible, the combinations were kept to the same block group. When the combinations came to the end of a block group without reaching the minimum, earlier blocks within the same block group were added. When necessary, blocks were combined across block groups within the same tract to form segments; however, collapsing blocks across tracts was not permitted. Consequently, the combinations consisted of blocks in close geographic proximity, and, in most cases, they were adjacent blocks. As a result of this method of combination, some large blocks that could have been segments by themselves were combined with small blocks.

At the second phase of segment selection, the constraint in equation [4] is equivalent to

$$M_{hj^{2}} \geq \frac{\hat{r}}{P_h P_{hj^{1}} \frac{n_{h(2)}}{\sum_{j^{[2]=1}^{n_{h(1)}}M_{hj^{2}}}} \quad [6]$$

The righthand side of equation [6] is the minimum MOS for the second-phase segment selection. Any first-phase segments, $j^{[1]}$, with MOS less than the minimum second-phase MOS are combined with adjacent segments to form the second-phase segments, $j^{[2]}$, prior to selection.

After second-phase selection, any $j^{[2]}$ segments formed as a combination of first-phase segments to achieve the second-phase minimum MOS were disaggregated into their first-phase components for operational reasons. The within-PSU probability of selection was

equal for the constituent segments. After completing the segment selection, the selected segments are denoted by j (without the superscript) to simplify the notation.

Controlling sample size per PSU

Screening and interviewing begin approximately 3 weeks before the first examinations in a location. This ensures having enough identified and interviewed sampled participants to fill available examination sessions. Once the MEC team arrives at a location (after conducting examinations in a previous location only days before), examinations for interviewed sampled participants begin. Examinations continue for approximately 5 weeks. After the last examination day, the field staff has limited time to travel to the next study location.

This strict time schedule for examining the sampled participants in each study location requires advance establishment of a fixed screening and examination workload in each location (see “Operational Requirements” in the previous “Design Specifications” section). However, as with any survey, it is not possible to predict the exact number of screened households that will supply the desired number of sampled participants and completed examinations. This is further aggravated by variations in response rates from location to location.

A fixed number of sampled participants is expected in the locations selected without certainty as a result of the constant sampling rate defined for each domain across all study locations r_{kl} . Within the study location, the sampling rate used for domain (k,l) is

$$P_{hj} \frac{\hat{r}}{P_h P_{hj}} \frac{r_{kl}}{\hat{r}} = \frac{r_{kl}}{P_h}$$

It can be shown that $\frac{r_{kl}}{P_h}$ is constant across the locations selected without certainty, to the extent that the population distribution is approximately the same as that in the decennial census data. Therefore, the number of sampled participants is approximately constant across these locations.

Because the number of segments per location is constant (i.e., 24) for all but the certainty PSUs, the variation in quotas per location is also reflected in segment sample sizes. In addition, changes in population distribution since the most recent census are likely to be greater among segments than among locations. The average segment size is thus expected to vary more than the average location size, but even this variation will generally be within a moderate range.

The approximate equality that exists in participant-level sample sizes per location and segment does not occur in the screening sample. The amount of screening in a location is partially based on what proportion of the location population lives in high-minority areas, so the amount of screening per segment will vary considerably. Consequently, a procedure that can produce samples either somewhat larger or somewhat smaller than those arising from application of the self-weighting sampling rates must be used; see the following “Selection of sample persons” section for more information on this procedure.

Selection of DUs and persons

The third stage of sample selection consisted of DUs, including certain types of group quarters. All DUs in the sample segments were listed, and a subsample of DUs was designated for screening to identify potential sample persons for interviews and examinations. The subsampling rates were designed to produce a national, approximately equal probability sample of DUs in most of the 50 states. Within each geographical stratum, an approximately equal probability sample of DUs across all PSUs existed.

Within-segment sampling rates

Within segments, DUs were selected with equal probability at a rate equal to the maximum within-segment sampling rate required to attain the subdomain sampling rates. That is, the sampling rate used to select DUs within segment j in PSU h is

$$\frac{\max_{k,l} \{r_{kl}\}}{P_h P_{hj}}$$

Sampled participants were selected within DUs using the ratio of the subdomain sampling rate to the maximum subdomain sampling rate. Thus, the overall selection probability for a person in race-Hispanic origin-sex-age-income subdomain (k,l) is

$$\Pr\{\text{select PSU } h\} \cdot \Pr\{\text{select segment } hj \mid \text{select PSU } h\}$$

$$\Pr\{\text{select a DU in segment } hj \mid \text{select segment } hj\}$$

$$\Pr\{\text{domain } (k,l) \text{ flagged for selection in DU} \mid \text{DU in segment } hj \text{ selected}\}$$

$$= P_h \cdot P_{hj} \cdot \frac{\max_{k,l} \{r_{kl}\}}{P_h \cdot P_{hj}} \cdot \frac{r_{kl}}{\max_{k,l} \{r_{kl}\}} = r_{kl}$$

and it can easily be shown that these probabilities yield approximately equal sample sizes for each PSU.

Selection of sample persons

The fourth stage of sample selection consisted of selecting sample persons. After the DU sample was released to the field, each DU was screened to determine whether it was occupied, vacant, or for seasonal use only. Only occupied DUs, or households, were eligible. Once the sampled households were identified, a sample of persons to be interviewed and examined from each household was selected. All eligible members within a household were listed and a subsample of persons was selected based on sex, age, race and Hispanic origin, and income. Sampled participants were selected at rates established to ensure that the target sample sizes by subdomain were achieved, and the average number of sampled participants per household was maximized.

Considerable subsampling was needed to reduce the screening sample of households to the desired number of sample participants. If independent random or systematic selections had been made for the subdomains, in most cases, only one person in a household would have been selected, and the average sample size per household

would have been quite low—not much above one.

A method of subsampling was used to maximize the number of sampled participants per household. (Conversely, this method minimizes the number of households containing sampled participants.) The effect of within-household clustering is not a large concern for NHANES because most analyses are done within subdomains, and within-household clustering at the subdomain level is generally small.

The method begins with the designated screening sample from which persons are to be subsampled. The persons are classified into Q subdomains with sampling rates $r_1, r_2, \dots, r_Q$. The subdomains are ordered by subsampling rate so that $r_q \leq r_{q+1}$. Note that the screening rates are set so that $r_Q = 1$; that is, the screening rate is equal to the maximum subsampling rate.

The set of households designated for screening is partitioned into L unequally sized random subsets, such that the sizes of the subsets are proportionate to $r_1, r_{(2)} - r_{(1)}, r_{(3)} - r_{(2)}, \dots, r_{(q+1)} - r_{(q)}, r_Q - r_{(Q-1)}$. The sum of these proportions is equal to $r_Q = 1$, so that each screened household is assigned to exactly one of the sets.

This sampling procedure was implemented using a set of sampling flags that designate for each DU which domains were eligible for sampling. The interviewers were not required to carry out any subsampling operation. They were, instead, instructed by the automated system (based on the set of domain flags provided for each household) which persons to include as sampled participants. Note that because the sampling domain flags were prepared in advance of the screening, they were based on the expected distribution of the screened sample by race and Hispanic origin, sex, age, and income, rather than on the distribution actually achieved. Thus, this procedure was expected to produce small deviations in the sample from the desired number in each domain. Such deviations are inevitable when subsampling rates must be established before the screening is completed.

The subsampling is then carried out as follows:

- In the first random subset of households (corresponding to $100 * r_1\%$ of all screened households), all persons in the household are designated as sampled participants.
- In the second random subset of households corresponding to $100 * (r_{(2)} - r_{(1)})\%$ of all screened households], all persons in the household are sampled participants except those in subdomain 1; therefore, those persons in subdomain 1 are selected only in the first random subset, with probability r_1 .
- In the third random subset of households, all persons in the household are sampled participants except those in subdomains 1 and 2. Thus, those in subdomain 2 are selected only in the first two random subsets, with probability $r_1 + (r_{(2)} - r_{(1)}) = r_{(2)}$.
- This procedure is continued in this manner through the Q th random subset, for which only persons in subdomain Q are sampled participants.

Instead of unrestricted randomization, a pseudorandom procedure was used to guarantee that all sampled DUs within each sequence of 1,000 consecutive DUs were assigned different random numbers (because the random number assigned determined the set of domains to be selected). To start, a random number between 0.000 and 0.999 was created, and then an additional random number on the interval (0,1) was created at the maximum precision and appended to the initial number. For example, if the first random number is 0.345 and the second is 0.93826485..., the result is 0.34593826485.

The resulting number was then assigned to the first DU, with a separate initial random number used in each study location. A three-digit prime number, 0.419, was then used as a skip interval and added successively to the original random number with precision to the third decimal place, and a new

Table F. Release group distribution: National Health and Nutrition Examination Survey, 2011–2014

Release group	Percentage for 180% sample
A.	42.0
B.	8.0
C.	7.0
D.	7.0
E.	6.0
F.	6.0
G.	5.0
H.	4.0
I.	4.0
J.	3.0
K.	3.0
L.	2.0
M.	2.0
N.	1.0

maximum precision number was generated and appended to the result to obtain the random number for the next case. The random number was then used, in the manner described above, to determine the sampling domain flags assigned to each case.

Initially, a screening sample was drawn for each study location using sampling rates larger than those required to attain the target sample sizes in each domain. Each study location's screening sample was then divided into release groups. Each group was a systematic subsample of the screening sample, with the screening sample sequenced by segment number and a temporary, geographically based sequence number prior to subsampling. Thus, each release group contained cases from all segments, except as limited by release group and segment size. The reserve sample selected was 80% larger than required. [Table F](#) gives the expected distribution of the sample of DUs across release groups.

In most study locations, the first, and largest, release group (i.e., group A) was issued to the interviewers initially. The yield from this group was monitored and used to project estimates of the total yield of sampled participants expected from this group. Based on these figures, additional groups (or portions of groups) were released as needed. The sample was monitored on a daily basis to determine whether additional release groups were required.

Special Samples

Examination session subsamples

NHANES has two examination session subsamples: the morning subsample, and the afternoon or evening subsample. Sampled participants selected for the morning sessions were instructed to fast overnight; those selected for the afternoon or evening sessions were also instructed to fast, but for a shorter period of time. Data that are sensitive to fasting times should be analyzed separately for these two groups.

Because it is generally more convenient for household members to come to the MEC at the same time (which is believed to favorably affect response rates), the examination session subsample assignment was made at the household level. The assignment was based on the household identifier (ID). If the household ID was an even number, the household was assigned to the morning subsample; if it was an odd number, the household was assigned to the afternoon or evening subsample. The examination session subsample was assigned immediately after DUs were selected.

Although the examination session subsamples were designed to be approximately one-half samples, some deviations resulted. Additionally, sampled participants did not always report to the assigned examination

session. For example, some sampled participants assigned to be examined in a morning session may have been unable to report to the MEC at that time; in such cases, they were permitted to schedule afternoon or evening examinations.

Examination and laboratory subsamples

The examination component of NHANES consisted of physical, dental, and laboratory tests to assess various aspects of health. For some of these components, subsampling was required to reduce respondent burden and facilitate the scheduling and completion of examinations.

Sampled participants were assigned to laboratory subsamples by first using an algorithm to randomly divide them into 12 groups; combinations of these groups were predetermined to create the various subsamples. Subsamples are most often mutually exclusive. In rare cases, subsamples overlap with one another, but not completely; for example, the persons who are part of a 1/3 environmental subsample may also be found in the 1/2 fasting subsample. To combine the 1/3 environmental subsample with the 1/2 fasting subsample, new weights would need to be created by the researcher because they had not been created by NCHS. Sample sizes may get quite small when combining these subsamples, resulting in unstable and unreliable statistical estimates.

After subsample assignment, weighting factors were attached to each sampled participant record, as appropriate, to reflect this stage of subsampling. [Table III](#) (Appendix II) provides the specifications for the components requiring subsampling.

As stated previously, more information about the 2011–2014 estimation procedures, the creation of weights for the entire sample and subsamples, and appropriate variance estimation methods to be used when analyzing NHANES data can be found in the forthcoming “National Health and Nutrition Examination Survey: Estimation Procedures, 2011–2014.”

References

1. Miller HW. Plan and operation of the Health and Nutrition Examination Survey, United States, 1971–1973. National Center for Health Statistics. Vital Health Stat 1(10a). 1973.
2. National Center for Health Statistics. Plan and initial program of the Health Examination Survey. Washington, DC. Vital Health Stat 1(4). 1965.
3. National Center for Health Statistics. Plan and operation of a health examination survey of U.S. youths 12–17 years of age. Rockville, MD. Vital Health Stat 1(8). 1974.
4. National Center for Health Statistics. Plan and operation of the Health and Nutrition Examination Survey, United States, 1971–1973. Hyattsville, MD. Vital Health Stat 1(10b). 1977.
5. National Center for Health Statistics. Plan, operation, and response results of a program of children’s examinations. Washington, DC. Vital Health Stat 1(5). 1967.
6. McDowell A, Engel A, Massey JT, Maurer K. Plan and operation of the second National Health and Nutrition Examination Survey, 1976–80. National Center for Health Statistics. Vital Health Stat 1(15). 1981.
7. National Center for Health Statistics. Plan and operation of the third National Health and Nutrition Examination Survey, 1988–1994. Hyattsville, MD. Vital Health Stat 1(32). 1994.
8. Maurer KR. Plan and operation of the Hispanic Health and Nutrition Examination Survey, 1982–84. National Center for Health Statistics. Vital Health Stat 1(19). 1985.
9. Curtin LR, Mohadjer L, Dohrmann S, et al. The National Health and Nutrition Examination Survey: Sample design, 1999–2006. National Center for Health Statistics. Vital Health Stat 2(155). 2012.
10. Curtin LR, Mohadjer LK, Dohrmann SM, et al. National Health and Nutrition Examination Survey: Sample design, 2007–2010. National Center for Health Statistics. Vital Health Stat 2(160). 2013.
11. Mirel LB, Mohadjer LK, Dohrmann SM, et al. National Health and Nutrition Examination Survey: Estimation procedures, 2007–2010. National Center for Health Statistics. Vital Health Stat 2(159). 2013.
12. National Center for Health Statistics. Vital statistics. Annual mortality files, 2003–2005.
13. National Center for Health Statistics. Vital statistics. Linked birth/infant death data set, 2002–2004.
14. CDC. Behavioral risk factor surveillance system, 2001.
15. CDC. Behavioral risk factor surveillance system, 2002.
16. Krenzke T, Haung W. Revisiting nested stratification of primary sampling units. FCSM. 2009.
17. Benedetti R, Espa G, Lafratta G. A tree-based approach to forming strata in multipurpose business surveys. *Survey Methodology* 34(2):195–203. 2008.
18. Ohlsson E. Methods for PPS size one sample coordination. Stockholm, Sweden: Institute for Actuarial Mathematics and Mathematical Statistics, Stockholm University;(194). 1996.

Appendix I. Glossary

Domain—A demographic group of analytic interest (analytic domain). Analytic domains may also be sampling domains if a sample design is created to meet goals for specific demographic groups. For NHANES, sampling domains are defined by race and Hispanic origin, income, age, and sex. See also *Sampling domain*.

Domain flags—See *Sampling domain flags*.

Double sampling—A general term for a method used in a number of statistical applications, such as stratification and regression, or ratio estimation. One of the applications of double sampling is to update a sampling frame when the sample is to be selected with respect to a measure of size (MOS), but a reliable estimate of that MOS is not available. For NHANES, double (or two-phase) sampling was used in second-stage units (SSUs, or segments) late in a decade when population counts from the U.S. Census Bureau—used in calculating MOS—were old and potentially no longer representative of the study location. In the NHANES study locations for which an accurate MOS is not available, a larger-than-needed sample of segments was selected in the first phase. After field staff determined the number of dwelling units (DUs) in the first-phase sample of segments, an updated MOS that reflected the ratio of the actual number of DUs to the expected number of DUs was calculated. The final sample of segments was selected by subsampling from the first-phase segments using the updated MOS.

Dwelling unit (DU), housing unit—The house, apartment, mobile home or trailer, group of rooms, or single room occupied as separate living quarters (see *Group quarters*) or, if vacant, intended for occupancy as separate living quarters. Separate living quarters are those in which the occupants live separately from other individuals in the building and which have direct access from outside the building or through a common hall. In

this report, the term generally means those DUs that are eligible for the survey (i.e., excluding institutional group quarters), or that could become eligible (e.g., vacant at the time of sampling but which could be occupied once screening begins).

Group quarters—A place where people live or stay that is normally owned or managed by an entity or organization providing housing or services for the residents. These services may include custodial or medical care as well as other types of assistance, and residency is commonly restricted to those receiving these services. People living in group quarters usually are not related to each other. Group quarters include such places as college residence halls, residential treatment centers, skilled nursing facilities, group homes, military barracks, correctional facilities, workers' dormitories, and facilities for people experiencing homelessness. These are generally grouped into two categories: institutional group quarters and noninstitutional group quarters.

Institutional group quarters—Group quarters providing formally authorized supervised care or custody in institutional settings, such as correctional facilities, nursing and skilled nursing facilities, inpatient hospice facilities, mental health or psychiatric hospitals, and group homes and residential treatment centers for juveniles. Institutional group quarters are not included in the NHANES sample.

Noninstitutional group quarters—Group quarters that do not provide formally authorized supervised care or custody in institutional settings. These include college or university housing, group homes and residential treatment facilities for adults, workers' group living quarters and Job Corps centers, and religious group quarters. Noninstitutional group quarters are included in the NHANES sample.

Household—The person or group of persons living in an occupied dwelling unit.

Institutional group quarters—See listing under *Group quarters*.

Low income—Beginning in 2000, NHANES split the sampling domains for white and other persons based on their income status into low income and non-low income. Low-income persons are those at or below 130% of the poverty level. The poverty threshold used in this determination was based on the most recent poverty guidelines published by the U.S. Department of Health and Human Services (HHS); these thresholds are updated annually by the U.S. Census Bureau.

Maximum sampling rate—The largest probability of selection assigned to a demographic group within a survey design. This value within certain strata and demographic groups was used in determining the sample size and other sampling parameters in NHANES.

Measure of size (MOS)—A value assigned to every sampling unit in a sample selection, usually a count of units associated with the elements to be selected. For NHANES, the MOS is actually a weighted average of estimates of population counts for the race-Hispanic origin-income groups of interest.

National Center for Health Statistics (NCHS)—The nation's principal health statistics agency, which designs, develops, and maintains a number of systems that produce data related to demographic and health concerns. These include data on registered births and deaths collected through the National Vital Statistics System, National Health Interview Survey or NHIS, National Health and Nutrition Examination Survey or NHANES, National Health Care Surveys, and National Survey of Family Growth or NSFG, among others. NCHS is one of 13 centers within the Centers for Disease Control and Prevention, which is part of HHS.

Noninstitutional group quarters—See listing under *Group quarters*.

Noninstitutionalized civilian population—Includes all people living in households, excluding institutional group quarters and those persons on active duty with the military. This is the target population for NHANES.

Primary sampling unit (PSU)—The first-stage selection unit in a multistage area probability sample. In NHANES, PSUs are counties or groups of counties in the United States. Some PSUs have such a large MOS that they are selected into the survey with a probability of one. These are referred to as PSUs selected with certainty, or “certainty PSUs”; all other PSUs are selected without certainty and known as “noncertainty PSUs.”

Public-use file—An electronic data set containing respondent records from a survey with a subset of variables collected in the survey that have been reviewed by analysts within NCHS to assure that the respondents’ identities are protected. NCHS disseminates this file to encourage widespread use of the survey data.

Probability proportionate to size (PPS) sampling—In this method, the probability of selecting any unit varies with the size of the unit, giving larger units a greater probability of selection and smaller units a lower probability. NHANES uses PPS sampling in the selection of PSUs and SSUs.

Race and Hispanic origin—The term used in this report as it was used in NHANES sample selection, covering four groups: Hispanic persons, non-Hispanic black persons, non-Hispanic Asian persons, and a fourth group consisting of all others.

Release group—A systematic subsample of a study location’s screening sample, with the screening sample sequenced by segment number and a temporary, geographically based sequence number. Each release group contained cases from all segments, except as limited by release group and segment size. In most study locations, the largest release group (i.e., group A) was released to the interviewers first. The yield from this group was monitored and used to project estimates of the total yield of sample persons

expected from this group. Based on these figures, additional groups (or portions of groups) were released as needed. The sample was monitored daily to determine whether additional release groups were required.

Respondent—A person selected into a sample who agrees to participate in all aspects of a survey. In NHANES, persons agreeing to complete the in-home interviews are “interview respondents.” Persons agreeing to complete the in-home interviews and an examination at a mobile examination center (MEC) are “MEC respondents.”

Response rate—The number of survey respondents divided by the number of persons selected into the sample. Response rates in this report are MEC response rates, calculated as the number of people receiving examinations in the MEC divided by the total number of people sampled.

Restricted-use file—An electronic data set of survey respondent records containing some information that may, if released to the public, risk disclosing individual survey respondents. The data are available only through the NCHS Research Data Center. These special data sets are for (a) data items collected for an odd number of calendar years (1, 3, or 5 years); (b) data sets with data geographically linked to other contextual data files (often supplied by the data user); (c) data items determined to be too sensitive or detailed to be released to the public due to confidentiality restrictions; and (d) surplus sera projects where past biological samples have been stored and subsequently used based on a formal proposal submitted as a special study; these could be on either the full sample or a special subsample.

Sampling domain—NHANES 2011–2014 includes 87 sampling domains, and Table A in this report contains the specific sampling domains for those years; see *Domain*.

Sampling domain flags—Strings of zeroes and ones attached to each sampled DU in the computer-assisted personal interview system. Each race-Hispanic origin-income group comprised one string, with each digit of the string representing one of the specific age-sex sampling domains. If the digit corresponding to an age-sex

domain in a race-Hispanic origin string contained a 1, then all persons in that DU with matching demographic characteristics were included in the sample. The zeroes and ones in each string were set based on the sampling rates.

Sampling rate—The rate at which a unit is selected from a sampling frame. For NHANES, the rates required for sampling persons in the race-Hispanic origin-sex-age-income domains were designed to achieve the designated number of MEC examinations in each of those domains. The sampling rates are the driving force in all stages of sampling.

Sample weight—For each NHANES respondent, the sample weight is the estimated number of persons in the target population that he or she represents. For example, if a man in the sample represents 12,000 men in his race-Hispanic origin-income-age category, then his sample weight is 12,000. The NHANES sample weights were adjusted for different sampling rates (of the race-Hispanic origin-income-age-sex groups), different response rates, and different coverage rates among persons in the sample, so that accurate national estimates can be made from the sample. Because it is the product of all of these adjustments, it is sometimes called the “final sample weight.”

Screener—An interview (usually short) containing a set of questions asked of a household member to determine whether the household contains anyone eligible for the survey. In NHANES, the screener, or screening interview, consisted of a household roster collecting the income level of the household and the race and Hispanic origin, age, and sex of all members. In NHANES, only persons aged 18 and over can answer the screener.

Screening—The process of conducting, or attempting to conduct, the screening interview in the dwelling units contained in the groups released. Occupied dwelling units (households) are “screened” through the screening interview. Other units can also be screened; the process for these units is verification that they are either vacant or not DUs. See also *Screener*.

Screening sample—The sample of DUs selected for a study location.

Secondary sampling unit (SSU)—The second-stage selection unit in a multistage area probability sample. For NHANES, these are typically referred to as “segments.”

Segment—A group of housing units located near each another, all of which were considered for selection into the sample. For NHANES, segments consist of a census block or groups of blocks. For NHANES, the selection of segments comprises the second stage of sampling. Within each segment, a sample of DUs was selected.

Self-weighting sample—A sample for which each elementary unit in the population has the same, nonzero chance of selection into the sample; that is, they are selected with the same constant probability. Higher-stage sampling units may be selected with differing probabilities, but such differences in selection probabilities at various stages cancel out. NHANES is a self-weighting sample of persons within each sampling domain.

Simple random sample—A sample in which all members of the population are selected directly and have an equal chance to be selected for the sample. The NHANES sample is not a simple random sample. The NHANES sample was stratified, was selected in stages, and employed unequal chances of selection for the respondents by race and Hispanic origin, income, age, and sex. Such designs are referred to as “complex” and require special software to estimate the variance of statistics computed from a sample with a complex design.

Study location—The set of segments within a PSU that were fielded together with all MEC examinations conducted at the same physical location. The distinction between a PSU and a study location is necessary because some large certainty PSUs were divided into multiple study locations and fielded at different times.

Strata, stratification—The partitioning of a population of sampling units into mutually exclusive categories (strata). Typically, stratification is used to increase the precision of survey estimates for subpopulations important

to the survey’s objectives. For the selection of PSUs fielded in 2011–2014, PSUs were stratified based on derived health-based state group.

Target population—The population to be described by estimates from the survey. In NHANES, the target population was the resident civilian noninstitutionalized population of the United States, which excluded all persons in supervised care or custody in institutional settings, all active-duty military personnel, active-duty family members living overseas, and any other persons residing outside the 50 states and District of Columbia.

Two-phase segment selection—See *Double sampling*.

Variance—A measure of the dispersion of a set of numbers. In this report, the variance is specifically the sample variance, which is a measure of the variation of a statistic, such as a proportion or a mean, calculated as a function of the sampling design and the population parameter being estimated. Many common statistical software packages compute “population variances” by default, which may underestimate the sampling variance because they do not incorporate any effects of having taken a sample compared with collecting data from every person in the full population. Estimating the variance in NHANES requires special software, as discussed in this report.

Weight—See *Sample weight*.

Appendix II. Supporting Tables and Figure

Table I. Derivation of expected screening requirements: National Health and Nutrition Examination Survey, 2011–2014

Race and Hispanic origin-income-sex-age sampling domain ¹	Projected population	Target number of examinations for 1 year	Projected number of households screened to have one examined person	Projected number of households screened to attain target number of examinations over 4 years in self-weighting area sample
Non-Hispanic black				
Male and female:				
Under age 1	674,174	55	209	45,920
1–2.	1,330,381	92	102	37,627
3–5.	1,942,242	92	72	26,613
Male:				
6–11	1,918,151	92	73	26,756
12–19	2,529,339	92	57	20,818
20–39	5,280,450	98	31	11,979
40–49	2,321,704	49	73	14,338
50–59	2,232,042	49	80	15,658
60 and over	2,150,364	98	82	32,290
Female:				
6–11	1,867,886	92	77	28,365
12–19	2,543,160	92	56	20,477
20–39	6,110,951	98	24	9,552
40–49	2,818,348	49	59	11,649
50–59	2,706,781	49	71	13,878
60 and over	3,173,942	98	59	22,985
Hispanic				
Male and female:				
Under age 1	1,089,092	92	123	45,104
1–2.	2,141,008	92	63	23,252
3–5.	3,139,421	92	47	17,318
Male:				
6–11	3,132,004	92	46	17,051
12–19	3,897,529	92	37	13,507
20–39	8,407,598	92	20	7,457
40–49	3,562,104	46	46	8,547
50–59	2,391,547	46	65	12,047
60 and over	2,159,484	93	86	31,872
Female:				
6–11	3,035,955	92	47	17,424
12–19	3,701,758	92	39	14,239
20–39	7,923,139	92	18	6,799
40–49	3,333,519	46	46	8,450
50–59	2,464,093	46	65	11,875
60 and over	2,700,912	93	69	25,849
Non-Hispanic non-black Asian				
Male and female:				
Under age 1	233,012	17	658	44,719
1–2.	471,178	35	328	45,933
3–5.	706,885	45	255	45,982
Male:				
6–11	697,062	51	225	45,832
12–19	877,572	58	199	46,132
20–39	2,417,584	76	66	20,114
40–49	1,255,497	39	184	28,652
50–59	944,805	39	201	31,422
60 and over	1,057,490	61	189	46,105
Female:				
6–11	684,332	51	228	46,445
12–19	895,244	58	184	42,693
20–39	2,703,143	76	71	21,437
40–49	1,375,901	39	172	26,851
50–59	1,105,732	39	185	28,803
60 and over	1,375,374	72	160	46,170

See footnotes at end of table.

Table I. Derivation of expected screening requirements: National Health and Nutrition Examination Survey, 2011–2014—Con.

Race and Hispanic origin-income-sex-age sampling domain ¹	Projected population	Target number of examinations for 1 year	Projected number of households screened to have one examined person	Projected number of households screened to attain target number of examinations over 4 years in self-weighting area sample
Non-Hispanic white and other, low income				
Male and female:				
Under age 1	391,169	31	334	41,375
1–2.	887,794	31	148	18,397
3–5.	1,168,203	31	119	14,809
Male:				
6–11	1,093,558	31	131	16,208
12–19	1,278,739	31	107	13,210
20–29	1,717,072	32	84	10,805
30–39	920,739	32	157	20,077
40–49	1,169,242	32	123	15,698
50–59	1,251,108	32	120	15,382
60–69	1,041,702	32	138	17,613
70–79	630,528	32	250	32,007
80 and over	377,839	21	544	45,681
Female:				
6–11	983,467	31	138	17,092
12–19	1,354,427	31	100	12,389
20–29	2,514,596	32	59	7,501
30–39	1,485,968	32	92	11,824
40–49	1,370,541	32	98	12,504
50–59	1,426,091	32	104	13,312
60–69	1,621,617	32	93	11,843
70–79	1,270,799	32	126	16,147
80 and over	1,259,338	32	176	22,487
Non-Hispanic white and other, non-low income				
Male and female:				
Under age 1	1,969,723	57	77	17,544
1–2.	3,824,644	57	40	9,135
3–5.	5,879,919	57	26	5,944
Male:				
6–11	6,142,447	57	26	5,990
12–19	8,456,481	57	18	4,006
20–29	11,308,888	57	15	3,399
30–39	11,180,561	57	17	3,902
40–49	12,432,334	57	15	3,387
50–59	14,052,788	57	13	2,884
60–69	10,659,189	57	17	3,827
70–79	5,616,501	57	31	7,136
80 and over	3,096,878	57	67	15,360
Female:				
6–11	5,906,145	57	27	6,109
12–19	7,985,401	57	19	4,281
20–29	10,452,963	57	15	3,498
30–39	10,778,324	57	17	3,788
40–49	12,442,842	57	13	3,013
50–59	14,378,358	57	12	2,632
60–69	10,945,786	57	17	3,925
70–79	6,110,217	57	31	7,164
80 and over	4,051,299	57	58	13,264
Total	312,366,116	5,000	...	...

... Category not applicable.

¹Age in years.

Table II. Final sampling rates and base weights: National Health and Nutrition Examination Survey, 2011–2014

Race and Hispanic origin-income-sex-age sampling domain ¹	2011–2012 ²		2013–2014 ³	
	Numerator of sampling rate ⁴	Base weight	Numerator of sampling rate ⁴	Base weight
Non-Hispanic black				
Male and female:				
Under age 1	1.00	1,430.08	1.00	1,890.34
1–2.	0.82	1,748.32	1.00	1,890.34
3–5.	0.58	2,471.85	0.76	2,471.85
Male:				
6–11	0.58	2,458.66	0.77	2,458.66
12–19	0.45	3,159.91	0.60	3,159.91
20–39	0.26	5,491.79	0.34	5,491.79
40–49	0.31	4,588.14	0.41	4,588.14
50–59	0.34	4,201.29	0.45	4,201.29
60 and over	0.70	2,037.30	0.93	2,037.30
Female:				
6–11	0.62	2,319.15	0.82	2,319.15
12–19	0.45	3,212.60	0.59	3,212.60
20–39	0.21	6,887.17	0.27	6,887.17
40–49	0.25	5,647.05	0.33	5,647.05
50–59	0.30	4,740.20	0.40	4,740.20
60 and over	0.50	2,861.99	0.66	2,861.99
Hispanic				
Male and female:				
Under age 1	0.98	1,458.48	1.00	1,890.34
1–2.	0.51	2,829.21	0.68	2,769.02
3–5.	0.38	3,798.67	0.51	3,678.71
Male:				
6–11	0.37	3,858.13	0.51	3,736.30
12–19	0.29	4,870.33	0.41	4,667.40
20–39	0.16	8,821.38	0.23	8,197.64
40–49	0.19	7,696.97	0.28	6,680.39
50–59	0.26	5,460.40	0.38	4,925.07
60 and over	0.69	2,064.02	0.94	2,020.57
Female:				
6–11	0.38	3,775.37	0.52	3,656.15
12–19	0.31	4,619.96	0.43	4,427.46
20–39	0.15	9,676.17	0.21	8,902.08
40–49	0.18	7,785.45	0.28	6,757.18
50–59	0.26	5,539.49	0.38	4,996.40
60 and over	0.56	2,544.90	0.76	2,491.33
Non-Hispanic non-black Asian				
Male and female:				
Under age 1	1.00	1,430.08	1.00	1,890.34
1–2.	1.00	1,430.08	1.00	1,890.34
3–5.	1.00	1,430.08	1.00	1,890.34
Male:				
6–11	1.00	1,430.08	1.00	1,890.34
12–19	1.00	1,430.08	1.00	1,890.34
20–39	0.44	3,270.54	0.58	3,270.54
40–49	0.62	2,295.99	0.82	2,295.99
50–59	0.68	2,093.59	0.90	2,093.59
60 and over	1.00	1,430.08	1.00	1,890.34
Female:				
6–11	1.00	1,430.08	1.00	1,890.34
12–19	0.93	1,540.84	1.00	1,890.34
20–39	0.47	3,068.76	0.62	3,068.76
40–49	0.58	2,449.97	0.77	2,449.97
50–59	0.63	2,283.92	0.83	2,283.92
60 and over	1.00	1,430.08	1.00	1,890.34

See footnotes at end of table.

Table II. Final sampling rates and base weights: National Health and Nutrition Examination Survey, 2011–2014—Con.

Race and Hispanic origin-income-sex-age sampling domain ¹	2011–2012 ²		2013–2014 ³	
	Numerator of sampling rate ⁴	Base weight	Numerator of sampling rate ⁴	Base weight
Non-Hispanic white and other, low income				
Male and female:				
Under age 1	0.90	1,589.94	1.00	1,890.34
1–2.	0.40	3,575.80	0.55	3,464.05
3–5.	0.32	4,442.21	0.44	4,303.39
Male:				
6–11	0.35	4,058.72	0.48	3,931.88
12–19	0.29	4,979.74	0.39	4,824.12
20–29	0.23	6,088.01	0.33	5,729.90
30–39	0.44	3,276.56	0.59	3,177.27
40–49	0.34	4,190.58	0.47	4,063.59
50–59	0.33	4,276.60	0.46	4,147.01
60–69	0.38	3,734.97	0.52	3,621.79
70–79	0.70	2,055.29	0.95	1,993.01
80 and over	1.00	1,430.08	1.00	1,890.34
Female:				
6–11	0.37	3,848.80	0.51	3,728.53
12–19	0.27	5,309.69	0.38	4,987.89
20–29	0.16	8,770.22	0.24	8,018.48
30–39	0.26	5,563.59	0.36	5,236.32
40–49	0.27	5,261.14	0.38	4,951.67
50–59	0.29	4,941.72	0.39	4,791.97
60–69	0.26	5,554.53	0.36	5,227.79
70–79	0.35	4,074.04	0.48	3,950.58
80 and over	0.49	2,925.40	0.67	2,836.75
Non-Hispanic white and other, non-low income				
Male and female:				
Under age 1	0.38	3,749.62	0.53	3,562.14
1–2.	0.20	7,201.28	0.29	6,620.53
3–5.	0.13	11,066.36	0.19	9,704.34
Male:				
6–11	0.13	10,981.69	0.20	9,630.10
12–19	0.09	16,419.59	0.14	13,564.01
20–29	0.07	19,355.25	0.12	15,538.73
30–39	0.08	16,858.95	0.14	13,926.96
40–49	0.07	19,422.44	0.12	15,592.66
50–59	0.06	22,809.90	0.11	17,810.47
60–69	0.08	17,190.55	0.13	14,200.89
70–79	0.16	9,218.96	0.23	8,210.64
80 and over	0.33	4,282.86	0.46	4,068.72
Female:				
6–11	0.13	10,768.92	0.20	9,443.51
12–19	0.09	15,367.79	0.15	12,881.82
20–29	0.08	18,803.49	0.12	15,311.41
30–39	0.08	17,365.14	0.13	14,345.11
40–49	0.07	21,833.52	0.11	17,048.09
50–59	0.06	24,990.82	0.10	18,993.03
60–69	0.09	16,761.14	0.14	13,846.16
70–79	0.16	9,182.40	0.23	8,178.08
80 and over	0.29	4,959.41	0.41	4,634.20

¹Age in years.²Sampling rates may be calculated by dividing the numerator by 1,430.³Sampling rates may be calculated by dividing the numerator by 1,890.⁴Rates correspond to 180% sample.

Table III. Description of subsamples: National Health and Nutrition Examination Survey, 2011–2014

Characteristic of interest	Sample collected	Age group (years)	Sample fraction (of age group)	Random groups included ¹
Persistent pesticide residues and metabolites, dioxins, furans, PCBs ² , organic fluorochemicals, nonpersistent pesticides, and thyroid function (2011–2012 only).	Blood	12 and over	1/3	1, 2, 5, 11 (2011–2012); 0, 3, 7, 10 (2013–2014)
Brominated flame retardants	Blood	12 and over	1/3	1, 2, 5, 11 (2011–2012); 4, 6, 8, 9 (2013–2014)
Volatile organic compounds (lead, cadmium, mercury, selenium, and magnesium/mercury; ethyl and methyl (2013–2014 only)	Blood	12 and over	1/2	1, 2, 4, 5, 10, 11
Selenium, copper, zinc	Blood	6 and over	1/3	1, 2, 5, 11
Aldehyde (2013–2014 only)	Blood	12 and over	1/3	1, 2, 5, 11
Pesticide residue and metabolites (2011 only).	Blood plasma	12 and over	1/3	4, 6, 8, 9
Persistent organochlorine pesticides, nonpersistent pesticides including organophosphate pesticide residues, caffeine	Urine	6 and over	1/3	4, 6, 8, 9
Phthalates	Urine	6 and over	1/3	1, 2, 5, 11 (2011–2012); 0, 3, 7, 10 (2013–2014)
Polyaromatic hydrocarbons, heavy metals, speciated arsenic, mercury, iodine, perchlorate, and, for 2013–2014 only, aromatic amines, heterocyclic nitrosamines, nicotine, and volatile nitrosamines	Urine	6 and over	1/3	1, 2, 5, 11
Acrylamide (2013–2014 only)	Washed cells	6 and over	1/3	1, 2, 5, 11

¹Each group is a random 1/12 sample.²Polychlorinated biphenyls.

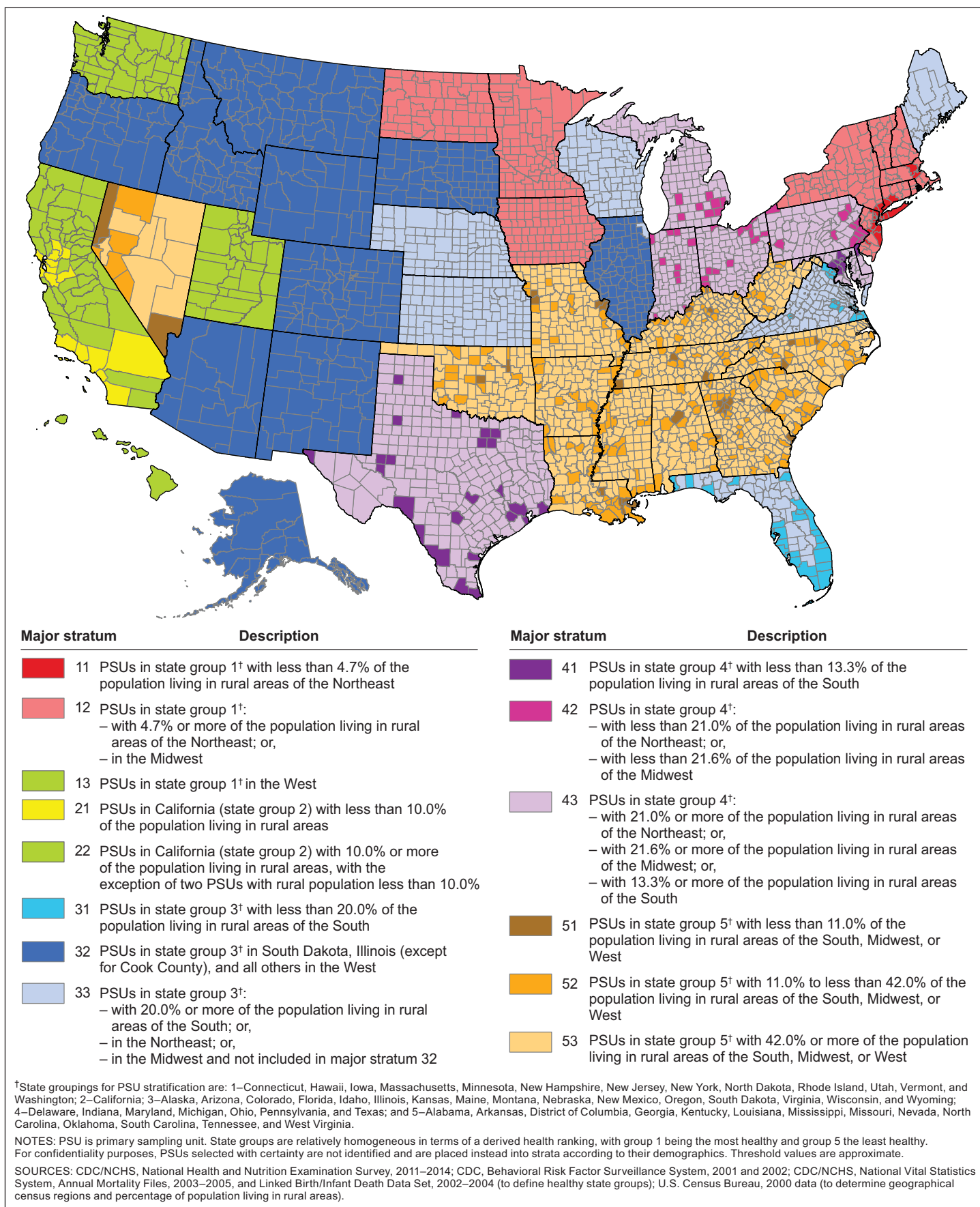


Figure. Major strata formed for selection of 2011–2014 primary sampling units: National Health and Nutrition Examination Survey, 2011–2014

Vital and Health Statistics Series Descriptions

ACTIVE SERIES

- Series 1. **Programs and Collection Procedures**—This type of report describes the data collection programs of the National Center for Health Statistics. Series 1 includes descriptions of the methods used to collect and process the data, definitions, and other material necessary for understanding the data.
- Series 2. **Data Evaluation and Methods Research**—This type of report concerns statistical methods and includes analytical techniques, objective evaluations of reliability of collected data, and contributions to statistical theory. Also included are experimental tests of new survey methods, comparisons of U.S. methodologies with those of other countries, and as of 2009, studies of cognition and survey measurement, and final reports of major committees concerning vital and health statistics measurement and methods.
- Series 3. **Analytical and Epidemiological Studies**—This type of report presents analytical or interpretive studies based on vital and health statistics. As of 2009, Series 3 also includes studies based on surveys that are not part of continuing data systems of the National Center for Health Statistics and international vital and health statistics reports.
- Series 10. **Data From the National Health Interview Survey**—This type of report contains statistics on illness; unintentional injuries; disability; use of hospital, medical, and other health services; and a wide range of special current health topics covering many aspects of health behaviors, health status, and health care utilization. Series 10 is based on data collected in this continuing national household interview survey.
- Series 11. **Data From the National Health Examination Survey, the National Health and Nutrition Examination Surveys, and the Hispanic Health and Nutrition Examination Survey**—In this type of report, data from direct examination, testing, and measurement on representative samples of the civilian noninstitutionalized population provide the basis for (1) medically defined total prevalence of specific diseases or conditions in the United States and the distributions of the population with respect to physical, physiological, and psychological characteristics, and (2) analyses of trends and relationships among various measurements and between survey periods.
- Series 13. **Data From the National Health Care Survey**—This type of report contains statistics on health resources and the public's use of health care resources including ambulatory, hospital, and long-term care services based on data collected directly from health care providers and provider records.
- Series 20. **Data on Mortality**—This type of report contains statistics on mortality that are not included in regular, annual, or monthly reports. Special analyses by cause of death, age, other demographic variables, and geographic and trend analyses are included.
- Series 21. **Data on Natality, Marriage, and Divorce**—This type of report contains statistics on natality, marriage, and divorce that are not included in regular, annual, or monthly reports. Special analyses by health and demographic variables and geographic and trend analyses are included.
- Series 23. **Data From the National Survey of Family Growth**—These reports contain statistics on factors that affect birth rates, including contraception and infertility; factors affecting the formation and dissolution of families, including cohabitation, marriage, divorce, and remarriage; and behavior related to the risk of HIV and other sexually transmitted diseases. These statistics are based on national surveys of women and men of childbearing age.

DISCONTINUED SERIES

- Series 4. **Documents and Committee Reports**—These are final reports of major committees concerned with vital and health statistics and documents. The last Series 4 report was published in 2002. As of 2009, this type of report is included in Series 2 or another appropriate series, depending on the report topic.
- Series 5. **International Vital and Health Statistics Reports**—This type of report compares U.S. vital and health statistics with those of other countries or presents other international data of relevance to the health statistics system of the United States. The last Series 5 report was published in 2003. As of 2009, this type of report is included in Series 3 or another series, depending on the report topic.
- Series 6. **Cognition and Survey Measurement**—This type of report uses methods of cognitive science to design, evaluate, and test survey instruments. The last Series 6 report was published in 1999. As of 2009, this type of report is included in Series 2.
- Series 12. **Data From the Institutionalized Population Surveys**—The last Series 12 report was published in 1974. Reports from these surveys are included in Series 13.
- Series 14. **Data on Health Resources: Manpower and Facilities**—The last Series 14 report was published in 1989. Reports on health resources are included in Series 13.
- Series 15. **Data From Special Surveys**—This type of report contains statistics on health and health-related topics collected in special surveys that are not part of the continuing data systems of the National Center for Health Statistics. The last Series 15 report was published in 2002. As of 2009, reports based on these surveys are included in Series 3.
- Series 16. **Compilations of Advance Data From Vital and Health Statistics**—The last Series 16 report was published in 1996. All reports are available online, and so compilations of Advance Data reports are no longer needed.
- Series 22. **Data From the National Mortality and Natality Surveys**—The last Series 22 report was published in 1973. Reports from these sample surveys, based on vital records, are published in Series 20 or 21.
- Series 24. **Compilations of Data on Natality, Mortality, Marriage, and Divorce**—The last Series 24 report was published in 1996. All reports are available online, and so compilations of reports are no longer needed.

For answers to questions about this report or for a list of reports published in these series, contact:

Information Dissemination Staff
National Center for Health Statistics
Centers for Disease Control and Prevention
3311 Toledo Road, Room 5419
Hyattsville, MD 20782

Tel: 1-800-CDC-INFO (1-800-232-4636)

TTY: 1-888-232-6348

Internet: <http://www.cdc.gov/nchs>

Online request form: <http://www.cdc.gov/cdc-info/requestform.html>

For e-mail updates on NCHS publication releases, subscribe online at: <http://www.cdc.gov/nchs/govdelivery.htm>.

**U.S. DEPARTMENT OF
HEALTH & HUMAN SERVICES**

Centers for Disease Control and Prevention
National Center for Health Statistics
3311 Toledo Road, Room 5419
Hyattsville, MD 20782

OFFICIAL BUSINESS
PENALTY FOR PRIVATE USE, \$300

FIRST CLASS MAIL
POSTAGE & FEES PAID
CDC/NCHS
PERMIT NO. G-284