

American Community Survey Accuracy of the Data (2010)

INTRODUCTION

The data contained in these data products are based on the American Community Survey (ACS) sample interviewed from January 1, 2010 through December 31, 2010. The ACS sample is selected from all counties and county-equivalents in the United States. In 2006, the ACS began collection of data from sampled persons in group quarters (GQs) – for example, military barracks, college dormitories, nursing homes, and correctional facilities. Persons in group quarters are included with persons in housing units (HUs) in all 2010 ACS estimates that are based on the total population. All ACS population estimates from years prior to 2006 include only persons in housing units. The ACS, like any other statistical activity, is subject to error. The purpose of this documentation is to provide data users with a basic understanding of the ACS sample design, estimation methodology, and accuracy of the ACS data. The ACS is sponsored by the U.S. Census Bureau, and is part of the 2010 Decennial Census Program.

Additional information on the design and methodology of the ACS, including data collection and processing, can be found at: http://www.census.gov/acs/www/methodology/methodology_main/.

The **2010 Accuracy of the Data from the Puerto Rico Community Survey** can be found at http://www.census.gov/acs/www/Downloads/data_documentation/Accuracy/PRCS_Accuracy_of_Data_2010.pdf.

TABLE OF CONTENTS

INTRODUCTION	1
DATA COLLECTION	3
SAMPLING FRAME	4
SAMPLE DESIGN	4
WEIGHTING METHODOLOGY	10
CONFIDENTIALITY OF THE DATA.....	16
ERRORS IN THE DATA.....	17
CALCULATION OF STANDARD ERRORS.....	20
Sums and Differences of Direct Standard Errors.....	21
Ratios	22
Proportions/Percents	22
Percent Change	22
Products.....	23
Comparing 1-Year Period Estimates with Overlapping 3-Year Period Estimates.....	23
Comparing 1-Year Period Estimates with Overlapping 5-Year Period Estimates.....	23
TESTING FOR SIGNIFICANT DIFFERENCES.....	23
EXAMPLES OF STANDARD ERROR CALCULATIONS.....	24
CONTROL OF NONSAMPLING ERROR.....	27
ISSUES WITH APPROXIMATING THE STANDARD ERROR OF LINEAR COMBINATIONS OF MULTIPLE ESTIMATES	31

DATA COLLECTION

Housing Units

The ACS employs three modes of data collection:

- Mailout/Mailback
- Computer Assisted Telephone Interview (CATI)
- Computer Assisted Personal Interview (CAPI)

With the exception of addresses in Remote Alaska, the general timing of data collection is:

- Month 1: Addresses in sample that are determined to be mailable are sent a questionnaire via the U.S. Postal Service.
- Month 2: All mail non-responding addresses with an available phone number are sent to CATI.
- Month 3: A sample of mail non-responses without a phone number, CATI non-responses, and unmailable addresses are selected and sent to CAPI.

Note that mail responses are accepted during all three months of data collection.

All Remote Alaska addresses are assigned to one of two data collection periods, January-April, or September-December and are sampled for CAPI at a rate of 2-in-3. Data for these addresses are collected using CAPI only and up to four months are given to complete the interviews in Remote Alaska for each data collection period.

Group Quarters

Group Quarters data collection spans six weeks, except in Remote Alaska and for Federal prisons, where the data collection time period is four months. As is done for HUs, Group Quarters in Remote Alaska are assigned to one of two data collection periods, January-April, or September-December and up to four months is allowed to complete the interviews. Similarly, all Federal prisons are assigned to September with a four month data collection window.

Field representatives have several options available to them for data collection. These include completing the questionnaire while speaking to the resident in person or over the telephone, conducting a personal interview with a proxy, such as a relative or guardian, or leaving paper questionnaires for residents to complete for themselves and then pick them up later. This last option is used for data collection in Federal prisons.

SAMPLING FRAME

Housing Units

The universe for the ACS consists of all valid, residential housing unit addresses in all county and county equivalents in the 50 states, including the District of Columbia. The Master Address File (MAF) is a database maintained by the Census Bureau containing a listing of residential and commercial addresses in the U.S. and Puerto Rico. The MAF is updated twice each year with the Delivery Sequence Files (DSF) provided by the U.S. Postal Service. The DSF covers only the U.S. These files identify mail drop points and provide the best available source of changes and updates to the housing unit inventory. The MAF is also updated with the results from various Census Bureau field operations, including the ACS.

Group Quarters

The main source of group quarters (GQ) in the ACS GQ sampling frame is the 2000 Census. Updates to the 2000 Census GQ inventory come from various operations such as the GQ Incomplete Information Operation (IIO) at the Census Bureau's National Processing Center in Jeffersonville, Indiana, the Census Questionnaire Resolution (CQR) Program and GQs closed on Census day. GQs that were closed on Census day were added from a preliminary inventory file obtained from Decennial Systems and Contract Management Office (DSCMO) since it was possible that while these GQs were closed on Census day, they could be open when the ACS contacts them. The GQ frame also underwent an unduplication process. Updates identified during various field operations, including ACS GQ data collection operations and the Demographic Area Address Listing (DAAL), are also reflected in the GQ frame. Headquarters Staff uses the state Department of Corrections websites to update the inventory and attributes of state prisons. The inventory of federal prisons in the ACS GQ frame is kept current through a partnership with the Bureau of Prisons where they provide ACS with an annual updated list of federal prisons. After the frame was put together from these different sources, it was then sorted geographically.

SAMPLE DESIGN

Housing Units

The ACS employs a two-phase, two-stage sample design. The ACS first-phase sample consists of two separate samples, Main and Supplemental, each chosen at different points in time. Together, these constitute the first-phase sample. Both the Main and the Supplemental samples are chosen in two stages referred to as first- and second-stage sampling. Subsequent to second-stage sampling, sample addresses are randomly assigned to one of the twelve months of the sample year. The second-phase of sampling occurs when the CAPI sample is selected (see Section 2 below).

The Main sample is selected during the summer preceding the sample year. Approximately 99 percent of the sample is selected at this time. Each address in sample is randomly assigned to one of the 12 months of the sample year. Supplemental sampling occurs in January/February of the sample year and accounts for approximately 1 percent of the overall first-phase sample. The Supplemental sample is allocated to the last nine months of the sample year. A sub-sample of non-responding addresses and of any addresses deemed unmailable is selected for the CAPI data collection mode.

Several of the steps used to select the first-phase sample are common to both Main and Supplemental sampling. The descriptions of the steps included in the first-phase sample selection below indicate which are common to both and which are unique to either Main or Supplemental sampling.

1. First-phase Sample Selection

- First-stage sampling (*performed during both Main and Supplemental sampling*) – First stage sampling defines the universe for the second stage of sampling through two steps. First, all addresses that were in a first-phase sample within the past four years are excluded from eligibility. This ensures that no address is in sample more than once in any five-year period. The second step is to select a 20 percent systematic sample of “new” units, i.e. those units that have never appeared on a previous MAF extract. Each new address is systematically assigned to either the current year or to one of four back-samples. This procedure maintains five equal partitions of the universe.
- Assignment of blocks to a second-stage sampling stratum (*performed during Main sampling only*) – Second-stage sampling uses seven sampling strata in the U.S. The stratum level rates used in second-stage sampling account for the first-stage selection probabilities. These rates are applied at a block level to addresses in the U.S. by calculating a measure of size for each of the following entities:
 - Counties
 - Places
 - School Districts (elementary, secondary, and unified)
 - American Indian Areas
 - Tribal Subdivisions
 - Alaska Native Village Statistical Areas
 - Hawaiian Homelands
 - Minor Civil Divisions – in Connecticut, Maine, Massachusetts, Michigan, Minnesota, New Hampshire, New Jersey, New York, Pennsylvania, Rhode Island, Vermont, and Wisconsin (these are the states where MCDs are active, functioning governmental units)
 - Census Designated Places – in Hawaii only

The measure of size for all areas except American Indian Areas, Tribal Subdivisions, and Alaska Native Village Statistical Areas is an estimate of the number of occupied HUs in the area. This is calculated by multiplying the number of ACS addresses by an estimated occupancy rate at the block level. A measure of size for each Census Tract is also calculated in the same manner.

For American Indian, Tribal Subdivisions, and Alaska Native Village Statistical Areas, the measure of size is the estimated number of occupied HUs multiplied by the proportion of people reporting American Indian or Alaska Native (alone or in combination) in Census 2000.

Each block is then assigned the smallest measure of size from the set of all entities of which it is a part. The second-stage sampling strata and the overall first-phase sampling rates are shown in Table 1 below.

- Calculation of the second-stage sampling rates (*performed during Main sampling only*) – The overall first-phase sampling rates given in Table 1 are calculated using the distribution of ACS valid addresses by second-stage sampling stratum in such a way as to yield an overall target sample size for the year of approximately 3,000,000 in the U.S. These rates also account for expected growth of the HU inventory between Main and Supplemental of roughly 1 percent. The first-phase rates are adjusted for the first-stage sample to yield the second-stage selection probabilities.
- Second-stage sample selection (*performed in Main and Supplemental*) – After each block is assigned to a second-stage sampling stratum, a systematic sample of addresses is selected from the second-stage universe (first-stage sample) within each county and county equivalent.
- Sample Month Assignment (*performed in Main and Supplemental*) – After the second stage of sampling, all sample addresses are randomly assigned to a sample month. Addresses selected during Main sampling are allocated to each of the 12 months. Addresses selected during Supplemental sampling are assigned to the months of April-December.

Table 1. First-phase Sampling Rate Categories for the United States

Sampling Rate Category	Sampling Rates
Blocks in smallest governmental units (MOS ¹ < 200)	10.00%
Blocks in smaller governmental units (200 ≤ MOS < 800)	6.67%
Blocks in small governmental units (800 ≤ MOS ≤ 1200)	3.33%
Blocks in large tracts (MOS >1200, TRACTMOS ² ≥ 2000) where Mailable addresses ³ ≥ 75% and predicted levels of completed mail and CATI interviews prior to second-stage sampling > 60%	1.50%
Other Blocks in large tracts (MOS >1200, TRACTMOS ≥ 2000)	1.63%
All other blocks (MOS >1200, TRACTMOS < 2000) where Mailable addresses ≥ 75% and predicted levels of completed mail and CATI interviews prior to second-stage sampling > 60%	2.04%
All other blocks (MOS >1200, TRACTMOS < 2000)	2.22%

¹MOS = Measure of size.

²TRACTMOS = Census Tract measure of size.

³Mailable addresses: Addresses that have sufficient information to be delivered by the U.S. Postal Service (as determined by ACS).

2. Second-phase Sample Selection – Subsampling the Unmailable and Non-Responding Addresses

All addresses determined to be unmailable are subsampled for the CAPI phase of data collection at a rate of 2-in-3. Unmailable addresses, which include Remote Alaska addresses, do not go to the CATI phase of data collection. Subsequent to CATI, all addresses for which no response has been obtained prior to CAPI are subsampled based on the expected rate of completed interviews at the tract level using the rates shown in Table 2.

Table 2. Second-Phase (CAPI) Subsampling Rates for the United States

Address and Tract Characteristics	CAPI Subsampling Rate
United States	
Unmailable addresses and addresses in Remote Alaska	66.7%
Mailable addresses in tracts with predicted levels of completed mail and CATI interviews prior to CAPI subsampling between 0% and less than 36%	50.0%
Mailable addresses in tracts with predicted levels of completed mail and CATI interviews prior to CAPI subsampling greater than 35% and less than 51%	40.0%
Mailable addresses in other tracts	33.3%

Group Quarters

The GQ sampling frame is divided into three strata: one for small GQs (having 15 or fewer people according to Census 2000 or updated information), one for GQs that were closed on Census Day 2000, and one for large GQs (having more than 15 people according to Census 2000 or updated information). GQs in the first two strata are sampled using the same procedure, while GQs in the large stratum are sampled using different a method. The small GQ stratum and the stratum for GQs closed on Census Day are combined into one sampling stratum and sorted geographically¹.

1. First-phase Sample Selection for Small GQ Stratum

- First-stage sampling - Small GQs are only eligible to be selected for the ACS once every five years. To accomplish this, the first stage sampling procedure systematically assigned all small GQs to one of five partitions of the universe. Each partition was assigned to a particular year (2010-2014) and the one assigned to 2010 became the first stage sample. In future years, each new GQ will be systematically assigned to one of the five samples. These samples are rotated over five year periods and become the universe for selecting the second stage sample.
- Second-stage sampling – During the second stage, GQs are selected from the first stage sample in a systematic sample of 1-in-x where x is dependent upon the state's target sampling rate. Since the first stage sample is one fifth of the universe, x can be calculated as $x = (1/5) \times (1 / \text{rate})$ where rate is the state's target sampling rate. For example, suppose a state had a target sampling rate of 2.5%. The systematic sample

¹Note that all references to the small GQ stratum include both small GQs and GQs closed on Census day.

would then be 1-in-8 since $(1 / 5) \times (1 / 0.025) = 8$. Regardless of their actual size, all GQs in the small stratum have the same probability of selection.

2. Sample Selection for the Large GQ Stratum

Unlike housing unit address sampling and the small GQ sample selection, the large GQ sampling procedure has no first-stage in which sampling units are randomly assigned to one of five years. All large GQs are eligible for sampling each year. The large GQ samples are selected using a two-phase design.

- **First-phase Sampling** - In the large GQ stratum, GQ hits are selected using a systematic PPS (probability proportional to size) sample, with a target sampling rate that varies according to state. A hit refers to a grouping of 10 expected interviews. GQs are selected with probability proportional to its most current count of persons or capacity. For stratification, and for sampling the large GQs, a GQ measure of size (GQMOS) is computed, where GQMOS is the expected population of the GQ divided by 10. This reflects that the GQ data is collected in groups of 10 GQ persons. People are selected in hits of 10 in a systematic sample of 1-in-x where $x = 1 / \text{rate}$ (one divided by the state's target sampling rate). For example, suppose a state had a target sampling rate of 2.5%. The hits would then be selected in a systematic sample of 1-in-40, since $1 / 0.025 = 40$.

All GQs in this stratum are eligible for sampling every year, regardless of their sample status in previous years. For large GQs, hits can be selected multiple times in the sample year. For most GQ types, the hits are randomly assigned throughout the year. Some GQs may have multiple hits with the same sample date if more than 12 hits are selected from the GQ. In these cases, the person sample within that month is unduplicated.

The following table summarizes the 2010 state target sampling rates for the U.S.

Table 3. 2010 State Targeted Sampling Rates for the U.S.

Alaska	5.53%	New Hampshire	3.17%
Delaware	4.86%	New Mexico	3.06%
District of Columbia	3.08%	North Dakota	4.59%
Hawaii	3.24%	Rhode Island	2.75%
Idaho	3.48%	South Dakota	3.91%
Maine	3.14%	Utah	2.79%
Montana	4.38%	Vermont	4.95%
Nevada	3.36%	Wyoming	7.11%
All Other States			2.5%

3. Sample Month Assignment

In order to assign each hit to a panel month, all of the GQ samples from a state are combined and sorted by small/large stratum and second-phase order of selection. Consecutive samples are assigned to the twelve panel months in a predetermined order, starting with a randomly determined month, except for Federal prisons and remote Alaska. Remote Alaska GQs are assigned to January and September based on where the GQ is located. Correctional facilities have their sample clustered. All Federal prisons hits are assigned to the September panel. In non-Federal correctional facilities, all hits for a given GQ are assigned to the same panel month. However, unlike Federal prisons, the hits in state and local correctional facilities are assigned to randomly selected panels spread throughout the year.

4. Second Phase Sample: Selection of Persons in Small and Large GQs

Small GQs in the second phase sampling are “take all,” i.e., every person in the selected GQ is eligible to receive a questionnaire. If the actual number of persons in the GQ exceeds 15, a field subsampling operation is performed to reduce the total number of sample persons interviewed at the GQ to 10. If the actual number of people is 15 or less, all people in the GQ will receive the questionnaire.

For each hit in the large GQs, the automated instrument uses the population count at the time of the visit and selects a subsample of 10 people from the roster. The people in this subsample receive the questionnaire.

WEIGHTING METHODOLOGY

The estimates that appear in this product are obtained from a raking ratio estimation procedure that results in the assignment of two sets of weights: a weight to each sample person record and a weight to each sample housing unit record. Estimates of person characteristics are based on the person weight. Estimates of family, household, and housing unit characteristics are based on the housing unit weight. For any given tabulation area, a characteristic total is estimated by summing the weights assigned to the persons, households, families or housing units possessing the characteristic in the tabulation area. Each sample person or housing unit record is assigned exactly one weight to be used to produce estimates of all characteristics. For example, if the weight given to a sample person or housing unit has a value 40, all characteristics of that person or housing unit are tabulated with the weight of 40.

The weighting is conducted in two main operations: a group quarters person weighting operation which assigns weights to persons in group quarters, and a household person weighting operation which assigns weights both to housing units and to persons within housing units. The group quarters person weighting is conducted first and the household person weighting second. The household person weighting is dependent on the group quarters person weighting because

estimates for total population which include both group quarters and household population are controlled to the Census Bureau's official 2010 total resident population estimates.

Group Quarters Person Weighting

Each GQ person is first assigned to a Population Estimates Program Major GQ Type Group (the type groups used by the Population Estimates Program). The major type groups used are:

Table 4: Population Estimates Program Major GQ Type Groups

Major GQ Type Group	Definition	Institutional / Non-Institutional
1	Correctional Institutions	Institutional
2	Juvenile Detention Facilities	Institutional
3	Nursing Homes	Institutional
4	Other Long-Term Care Facilities	Institutional
5	College Dormitories	Non-Institutional
6	Military Facilities	Non-Institutional
7	Other Non-Institutional Facilities	Non-Institutional

The procedure used to assign the weights to the GQ persons is performed independently within state. The steps are as follows:

- **Base Weight**—The initial base weight after the first phase of sampling is the inverse of its first-phase sampling rate. The initial base weight is equal to 40 for sample cases in most states in 2010, though in 15 states and the District of Columbia the initial base weights are smaller. This initial base weight is then adjusted for the second-phase sampling that occurs at the time of interview.
- **Non-Interview Factor**—This factor adjusts the weight of all responding GQ persons to account for the non-responding GQ persons including those persons contained in whole non-responding GQs. The non-interview factor is computed and assigned using the following groups:

State × Major GQ Type Group × County

- **GQ Person Post-stratification Factor**—This factor adjusts the GQ person weights so that the weighted sample counts equal independent population estimates from the Population Estimates Program by Major Type Group at the state level in the U.S. Because of collapsing of groups in applying this factor, only total GQ population is assured of agreeing with the Census Bureau's official 2010 population estimates at the state level. The GQ person post-stratification factor is computed and assigned using the following groups:

State × Major GQ Type Group

- Rounding—The final GQ person weight is rounded to an integer. Rounding is performed so that the sum of the rounded weights is within one person of the sum of the unrounded weights for any of the groups listed below:

Major GQ Type Group
Major GQ Type Group × County

Housing Unit and Household Person Weighting

The housing unit and household person weighting uses two types of geographic areas for adjustments: weighting areas and subcounty areas. Weighting areas are county-based and have been used since the first year of the ACS. Subcounty areas are based on incorporated place and minor civil divisions (MCD). Their use was introduced into the ACS in 2010.

Weighting areas are built from collections of whole counties. Census 2000 data are used to group counties of similar demographic and social characteristics. The characteristics considered in the formation include:

- Percent in poverty
- Percent renting
- Percent in rural areas
- Race, ethnicity, age, and sex distribution
- Distance between the centroids of the counties
- Core-based Statistical Area status

Each weighting area is also required to meet a threshold of 400 expected person interviews in the 2005 ACS. The process also tries to preserve as many counties that meet the threshold to form their own weighting areas. In total, there are 1,951 weighting areas formed from the 3,141 counties and county equivalents.

Subcounty areas are built from incorporated places and MCDs, with MCDs only being used in the 20 states where MCDs serve as functioning governmental units. Each subcounty area formed has a total population of at least 24,000, as determined by the July 1, 2010 Population Estimates data, which are based on the 2010 Census estimates of the population on April 1, 2010, extrapolated forward. The subcounty areas can be incorporated places, MCDs, place/MCD intersections (in counties where places and MCDs are not coexistent), ‘balance of MCD,’ and ‘balance of county.’ The latter two types group together unincorporated areas and places/MCDs that do not meet the population threshold. If two or more subcounty areas cannot be formed within a county, then the entire county is treated as a single area. Thus all counties whose total population is less than 48,000 will be treated as a single area since it is not possible to form two areas that satisfy the minimum size threshold.

The estimation procedure used to assign the weights is then performed independently within each of the ACS weighting areas.

1. Initial Housing Unit Weighting Factors - This process produced the following factors:

- Base Weight (BW) - This initial weight is assigned to every housing unit as the inverse of its block's sampling rate.
- CAPI Subsampling Factor (SSF) - The weights of the CAPI cases are adjusted to reflect the results of CAPI subsampling. This factor is assigned to each record as follows:

Selected in CAPI subsampling: $SSF = 2.0, 2.5, \text{ or } 3.0$ according to Table 2

Not selected in CAPI subsampling: $SSF = 0.0$

Not a CAPI case: $SSF = 1.0$

Some sample addresses are unmailable. A two-thirds sample of these is sent directly to CAPI and for these cases $SSF = 1.5$.

- Variation in Monthly Response by Mode (VMS)—This factor makes the total weight of the Mail, CATI, and CAPI records to be tabulated in a month equal to the total base weight of all cases originally mailed for that month. For all cases, VMS is computed and assigned based on the following groups:

Weighting Area $\times$ Month

- Noninterview Factor (NIF)—This factor adjusts the weight of all responding occupied housing units to account for nonresponding housing units. The factor is computed in two stages. The first factor, NIF1, is a ratio adjustment that is computed and assigned to occupied housings units based on the following groups:

Weighting Area $\times$ Building Type $\times$ Tract

A second factor, NIF2, is a ratio adjustment that is computed and assigned to occupied housing units based on the following groups:

Weighting Area $\times$ Building Type $\times$ Month

NIF is then computed by applying NIF1 and NIF2 for each occupied housing unit. Vacant housing units are assigned a value of $NIF = 1.0$. Nonresponding housing units are assigned a weight of 0.0.

- Noninterview Factor - Mode (NIFM) - This factor adjusts the weight of the responding CAPI occupied housing units to account for CAPI nonrespondents. It is computed as if NIF had not already been assigned to every occupied housing unit

record. This factor is not used directly but rather as part of computing the next factor, the Mode Bias Factor.

NIFM is computed and assigned to occupied CAPI housing units based on the following groups:

Weighting Area $\times$ Building Type (single or multi unit) $\times$ Month

Vacant housing units or non-CAPI (mail and CATI) housing units receive a value of NIFM = 1.0.

- Mode Bias Factor (MBF)—This factor makes the total weight of the housing units in the groups below the same as if NIFM had been used instead of NIF. MBF is computed and assigned to occupied housing units based on the following groups:

Weighting Area $\times$ Tenure (owner or renter) $\times$ Month $\times$ Marital Status of the Householder (married/widowed or single)

Vacant housing units receive a value of MBF = 1.0. MBF is applied to the weights computed through NIF.

- Housing unit Post-stratification Factor (HPF)—This factor makes the total weight of all housing units agree with the 2010 independent housing unit estimates at the subcounty level.
2. Person Weighting Factors—Initially the person weight of each person in an occupied housing unit is the product of the weighting factors of their associated housing unit ($BW \times \dots \times HPF$). At this point everyone in the household has the same weight. The person weighting is done in a series of three steps which are repeated until a stopping criterion is met. These three steps form a raking ratio or raking process. These person weights are individually adjusted for each person as described below.

The three steps are as follows:

- Subcounty Area Controls Raking Factor (SUBEQRF) – This factor is applied to individuals based on their geography. It adjusts the person weights so that the weighted sample counts equal independent population estimates of total population for the subcounty area. Because of later adjustment to the person weights, total population is not assured of agreeing exactly with the official 2010 population estimates at the subcounty level.
- Spouse Equalization/Householder Equalization Raking Factor (SPHHEQRF)—This factor is applied to individuals based on the combination of their status of being in a married-couple or unmarried-partner household and whether they are the

householder. All persons are assigned to one of four groups:

1. Householder in a married-couple or unmarried-partner household
2. Spouse or unmarried partner in a married-couple or unmarried-partner household (non-householder)
3. Other householder
4. Other non-householder

The weights of persons in the first two groups are adjusted so that their sums are each equal to the total estimate of married-couple or unmarried-partner households using the housing unit weight ($BW \times \dots \times HPF$). At the same time the weights of persons in the first and third groups are adjusted so that their sum is equal to the total estimate of occupied housing units using the housing unit weight ($BW \times \dots \times HPF$). The goal of this step is to produce more consistent estimates of spouses or unmarried partners and married-couple and unmarried-partner households while simultaneously producing more consistent estimates of householders, occupied housing units, and households.

- Demographic Raking Factor (DEMORF)—This factor is applied to individuals based on their age, race, sex and Hispanic origin. It adjusts the person weights so that the weighted sample counts equal independent population estimates by age, race, sex, and Hispanic origin at the weighting area. Because of collapsing of groups in applying this factor, only total population is assured of agreeing with the official 2010 population estimates at the weighting area level.

This uses the following groups (note that there are 13 Age groupings):

Weighting Area $\times$ Race / Ethnicity (non-Hispanic White, non-Hispanic Black, non-Hispanic American Indian or Alaskan Native, non-Hispanic Asian, non-Hispanic Native Hawaiian or Pacific Islander, and Hispanic (any race)) $\times$ Sex $\times$ Age Groups.

These three steps are repeated several times until the estimates at the national level achieve their optimal consistency with regard to the spouse and householder equalization. The effect Person Post-Stratification Factor (PPSF) is then equal to the product ($SUBEQRF \times SPHHEQRF \times DEMORF$) from all of iterations of these three adjustments. The unrounded person weight is then the equal to the product of PPSF times the housing unit weight ($BW \times \dots \times HPF \times PPSF$).

3. Rounding—The final product of all person weights ($BW \times \dots \times HPF \times PPSF$) is rounded to an integer. Rounding is performed so that the sum of the rounded weights is within one person of the sum of the unrounded weights for any of the groups listed below:

County
County $\times$ Race

County $\times$ Race $\times$ Hispanic Origin
 County $\times$ Race $\times$ Hispanic Origin $\times$ Sex
 County $\times$ Race $\times$ Hispanic Origin $\times$ Sex $\times$ Age
 County $\times$ Race $\times$ Hispanic Origin $\times$ Sex $\times$ Age $\times$ Tract
 County $\times$ Race $\times$ Hispanic Origin $\times$ Sex $\times$ Age $\times$ Tract $\times$ Block

For example, the number of White, Hispanic, Males, Age 30 estimated for a county using the rounded weights is within one of the number produced using the unrounded weights.

4. Final Housing Unit Weighting Factors—This process produces the following factors:
 - Householder Factor (HHF)—This factor adjusts for differential response depending on the race, Hispanic origin, sex, and age of the householder. The value of HHF for an occupied housing unit is the PPSF of the householder. Since there is no householder for vacant units, the value of HHF = 1.0 for all vacant units.
 - Rounding—The final product of all housing unit weights ($BW \times \dots \times HHF$) is rounded to an integer. For occupied units, the rounded housing unit weight is the same as the rounded person weight of the householder. This ensures that both the rounded and unrounded householder weights are equal to the occupied housing unit weight. The rounding for vacant housing units is then performed so that total rounded weight is within one housing unit of the total unrounded weight for any of the groups listed below:

County
 County $\times$ Tract
 County $\times$ Tract $\times$ Block

CONFIDENTIALITY OF THE DATA

The Census Bureau has modified or suppressed some data on this site to protect confidentiality. Title 13 United States Code, Section 9, prohibits the Census Bureau from publishing results in which an individual's data can be identified.

The Census Bureau's internal Disclosure Review Board sets the confidentiality rules for all data releases. A checklist approach is used to ensure that all potential risks to the confidentiality of the data are considered and addressed.

- Title 13, United States Code: Title 13 of the United States Code authorizes the Census Bureau to conduct censuses and surveys. Section 9 of the same Title requires that any information collected from the public under the authority of Title 13 be maintained as confidential. Section 214 of Title 13 and Sections 3559 and 3571 of Title 18 of the

United States Code provide for the imposition of penalties of up to five years in prison and up to \$250,000 in fines for wrongful disclosure of confidential census information.

- **Disclosure Avoidance:** Disclosure avoidance is the process for protecting the confidentiality of data. A disclosure of data occurs when someone can use published statistical information to identify an individual that has provided information under a pledge of confidentiality. For data tabulations, the Census Bureau uses disclosure avoidance procedures to modify or remove the characteristics that put confidential information at risk for disclosure. Although it may appear that a table shows information about a specific individual, the Census Bureau has taken steps to disguise or suppress the original data while making sure the results are still useful. The techniques used by the Census Bureau to protect confidentiality in tabulations vary, depending on the type of data.
- **Data Swapping:** Data swapping is a method of disclosure avoidance designed to protect confidentiality in tables of frequency data (the number or percent of the population with certain characteristics). Data swapping is done by editing the source data or exchanging records for a sample of cases when creating a table. A sample of households is selected and matched on a set of selected key variables with households in neighboring geographic areas that have similar characteristics (such as the same number of adults and same number of children). Because the swap often occurs within a neighboring area, there is no effect on the marginal totals for the area or for totals that include data from multiple areas. Because of data swapping, users should not assume that tables with cells having a value of one or two reveal information about specific individuals. Data swapping procedures were first used in the 1990 Census, and were used again in Census 2000 and the 2010 Census.
- **Synthetic Data:** The goals of using synthetic data are the same as the goals of data swapping, namely to protect the confidentiality in tables of frequency data. Persons are identified as being at risk for disclosure based on certain characteristics. The synthetic data technique then models the values for another collection of characteristics to protect the confidentiality of that individual.

ERRORS IN THE DATA

- **Sampling Error** — The data in the ACS products are estimates of the actual figures that would have been obtained by interviewing the entire population using the same methodology. The estimates from the chosen sample also differ from other samples of housing units and persons within those housing units. Sampling error in data arises due to the use of probability sampling, which is necessary to ensure the integrity and representativeness of sample survey results. The implementation of statistical sampling

procedures provides the basis for the statistical analysis of sample data. Measures used to estimate the sampling error are provided in the next section.

- **Nonsampling Error** — In addition to sampling error, data users should realize that other types of errors may be introduced during any of the various complex operations used to collect and process survey data. For example, operations such as data entry from questionnaires and editing may introduce error into the estimates. Another source is through the use of controls in the weighting. The controls are designed to mitigate the effects of systematic undercoverage of certain groups who are difficult to enumerate as well as to reduce the variance. The controls are based on the population estimates extrapolated from the previous census. Errors can be brought into the data if the extrapolation methods do not properly reflect the population. However, the potential risk from using the controls in the weighting process is offset by far greater benefits to the ACS estimates. These benefits include reducing the effects of a larger coverage problem found in most surveys, including the ACS, and the reduction of standard errors of ACS estimates. These and other sources of error contribute to the nonsampling error component of the total error of survey estimates. Nonsampling errors may affect the data in two ways. Errors that are introduced randomly increase the variability of the data. Systematic errors which are consistent in one direction introduce bias into the results of a sample survey. The Census Bureau protects against the effect of systematic errors on survey estimates by conducting extensive research and evaluation programs on sampling techniques, questionnaire design, and data collection and processing procedures. In addition, an important goal of the ACS is to minimize the amount of nonsampling error introduced through nonresponse for sample housing units. One way of accomplishing this is by following up on mail nonrespondents during the CATI and CAPI phases. For more information, please see the section entitled “Control of Nonsampling Error”.

MEASURES OF SAMPLING ERROR

Sampling error is the difference between an estimate based on a sample and the corresponding value that would be obtained if the estimate were based on the entire population (as from a census). Note that sample-based estimates will vary depending on the particular sample selected from the population. Measures of the magnitude of sampling error reflect the variation in the estimates over all possible samples that could have been selected from the population using the same sampling methodology.

Estimates of the magnitude of sampling errors – in the form of margins of error – are provided with all published ACS data. The Census Bureau recommends that data users incorporate this information into their analyses, as sampling error in survey estimates could impact the conclusions drawn from the results.

Confidence Intervals and Margins of Error

Confidence Intervals – A sample estimate and its estimated standard error may be used to construct confidence intervals about the estimate. These intervals are ranges that will contain the average value of the estimated characteristic that results over all possible samples, with a known probability.

For example, if all possible samples that could result under the ACS sample design were independently selected and surveyed under the same conditions, and if the estimate and its estimated standard error were calculated for each of these samples, then:

1. Approximately 68 percent of the intervals from one estimated standard error below the estimate to one estimated standard error above the estimate would contain the average result from all possible samples;
2. Approximately 90 percent of the intervals from 1.645 times the estimated standard error below the estimate to 1.645 times the estimated standard error above the estimate would contain the average result from all possible samples.
3. Approximately 95 percent of the intervals from two estimated standard errors below the estimate to two estimated standard errors above the estimate would contain the average result from all possible samples.

The intervals are referred to as 68 percent, 90 percent, and 95 percent confidence intervals, respectively.

Margin of Error – Instead of providing the upper and lower confidence bounds in published ACS tables, the margin of error is provided instead. The margin of error is the difference between an estimate and its upper or lower confidence bound. Both the confidence bounds and the standard error can easily be computed from the margin of error. All ACS published margins of error are based on a 90 percent confidence level.

$$\text{Standard Error} = \text{Margin of Error} / 1.645$$

$$\text{Lower Confidence Bound} = \text{Estimate} - \text{Margin of Error}$$

$$\text{Upper Confidence Bound} = \text{Estimate} + \text{Margin of Error}$$

Note that for 2005 and earlier estimates, ACS margins of error and confidence bounds were calculated using a 90 percent confidence level multiplier of 1.65. Beginning with the 2006 data release, we are now employing a more accurate multiplier of 1.645. Margins of error and confidence bounds from previously published products will not be updated with the new

multiplier. When calculating standard errors from margins of error or confidence bounds using published data for 2005 and earlier, use the 1.65 multiplier.

When constructing confidence bounds from the margin of error, the user should be aware of any “natural” limits on the bounds. For example, if a characteristic estimate for the population is near zero, the calculated value of the lower confidence bound may be negative. However, a negative number of people does not make sense, so the lower confidence bound should be reported as zero instead. However, for other estimates such as income, negative values do make sense. The context and meaning of the estimate must be kept in mind when creating these bounds. Another of these natural limits would be 100 percent for the upper bound of a percent estimate.

If the margin of error is displayed as ‘*****’ (five asterisks), the estimate has been controlled to be equal to a fixed value and so it has no sampling error. When using any of the formulas in the following section, use a standard error of zero for these controlled estimates.

Limitations –The user should be careful when computing and interpreting confidence intervals.

- The estimated standard errors (and thus margins of error) included in these data products do not include portions of the variability due to nonsampling error that may be present in the data. In particular, the standard errors do not reflect the effect of correlated errors introduced by interviewers, coders, or other field or processing personnel. Nor do they reflect the error from imputed values due to missing responses. Thus, the standard errors calculated represent a lower bound of the total error. As a result, confidence intervals formed using these estimated standard errors may not meet the stated levels of confidence (i.e., 68, 90, or 95 percent). Thus, some care must be exercised in the interpretation of the data in this data product based on the estimated standard errors.
- Zero or small estimates; very large estimates — The value of almost all ACS characteristics is greater than or equal to zero by definition. For zero or small estimates, use of the method given previously for calculating confidence intervals relies on large sample theory, and may result in negative values which for most characteristics are not admissible. In this case the lower limit of the confidence interval is set to zero by default. A similar caution holds for estimates of totals close to a control total or estimated proportion near one, where the upper limit of the confidence interval is set to its largest admissible value. In these situations the level of confidence of the adjusted range of values is less than the prescribed confidence level.

CALCULATION OF STANDARD ERRORS

Direct estimates of the standard errors were calculated for all estimates reported in this product. The standard errors, in most cases, are calculated using a replicate-based methodology that takes into account the sample design and estimation procedures. Excluding the base weights, replicate

weights were allowed to be negative in order to avoid underestimating the standard error. Exceptions include:

1. The estimate of the number or proportion of people, households, families, or housing units in a geographic area with a specific characteristic is zero. A special procedure is used to estimate the standard error.
2. There are either no sample observations available to compute an estimate or standard error of a median, an aggregate, a proportion, or some other ratio, or there are too few sample observations to compute a stable estimate of the standard error. The estimate is represented in the tables by “-” and the margin of error by “**” (two asterisks).
3. The estimate of a median falls in the lower open-ended interval or upper open-ended interval of a distribution. If the median occurs in the lowest interval, then a “-” follows the estimate, and if the median occurs in the upper interval, then a “+” follows the estimate. In both cases the margin of error is represented in the tables by “***” (three asterisks).

Sums and Differences of Direct Standard Errors

The standard errors estimated from these tables are for individual estimates. Additional calculations are required to estimate the standard errors for sums of or the differences between two or more sample estimates.

The standard error of the sum of two sample estimates is the square root of the sum of the two individual standard errors squared plus a covariance term. That is, for standard errors $SE(\hat{X}_1)$ and $SE(\hat{X}_2)$ of estimates $\hat{X}_1$ and $\hat{X}_2$:

$$SE(\hat{X}_1 \pm \hat{X}_2) = \sqrt{(SE(\hat{X}_1))^2 + (SE(\hat{X}_2))^2 \pm covariance} \quad (1)$$

The covariance measures the interactions between two estimates. Currently the covariance terms are not available. Data users should use the approximation:

$$SE(\hat{X}_1 \pm \hat{X}_2) \approx \sqrt{(SE(\hat{X}_1))^2 + (SE(\hat{X}_2))^2} \quad (2)$$

However, this method will underestimate or overestimate the standard error if the two estimates interact in either a positive or negative way.

The approximation formula (2) can be expanded to more than two estimates by adding in the individual standard errors squared inside the radical. As the number of estimates involved in the

sum or difference increases, the results of formula (2) become increasingly different from the standard error derived directly from the ACS microdata. Care should be taken to work with the fewest number of estimates as possible. If there are estimates involved in the sum that are controlled in the weighting then the approximate standard error can be increasingly different. Several examples are provided starting on page 31 to demonstrate issues associated with approximating the standard errors when summing large numbers of estimates together.

Ratios

The statistic of interest may be the ratio of two estimates. First is the case where the numerator *is not* a subset of the denominator. The standard error of this ratio between two sample estimates is approximated as:

$$SE\left(\frac{\hat{X}}{\hat{Y}}\right) = \frac{1}{\hat{Y}} \sqrt{[SE(\hat{X})]^2 + \frac{\hat{X}^2}{\hat{Y}^2} [SE(\hat{Y})]^2} \quad (3)$$

Proportions/Percents

For a proportion (or percent), a ratio where the numerator *is* a subset of the denominator, a slightly different estimator is used. If $\hat{P} = \hat{X}/\hat{Y}$, then the standard error of this proportion is approximated as:

$$SE(\hat{P}) = \frac{1}{\hat{Y}} \sqrt{[SE(\hat{X})]^2 - \frac{\hat{X}^2}{\hat{Y}^2} [SE(\hat{Y})]^2} \quad (4)$$

If $\hat{Q} = 100\% \times \hat{P}$ (P is the proportion and Q is its corresponding percent), then $SE(\hat{Q}) = 100\% \times SE(\hat{P})$. Note the difference between the formulas to approximate the standard error for proportions (4) and ratios (3) - the plus sign in the previous formula has been replaced with a minus sign. If the value under the radical is negative, use the ratio standard error formula above, instead.

Percent Change

Calculating the percent change from one time period to another. For example, computing the percent change of a 2008-2010 estimate to a 2007-2009 estimate. Normally, the current estimate is compared to the older estimate.

Let the current estimate = X and the earlier estimate = Y, then the formula for percent change is:

$$SE\left(\frac{\hat{X} - \hat{Y}}{\hat{Y}}\right) = SE\left(\frac{\hat{X}}{\hat{Y}} - 1\right) = SE\left(\frac{\hat{X}}{\hat{Y}}\right) \quad (5)$$

This reduces to a ratio. The ratio formula above may be used to calculate the standard error. As a caveat, this formula does not take into account the correlation when calculating overlapping time periods.

Products

For a product of two estimates - for example if you want to estimate a proportion's numerator by multiplying the proportion by its denominator - the standard error can be approximated as:

$$SE(\hat{X} \times \hat{Y}) = \sqrt{\hat{X}^2 \times [SE(\hat{Y})]^2 + \hat{Y}^2 \times [SE(\hat{X})]^2} \quad (6)$$

Comparing 1-Year Period Estimates with Overlapping 3-Year Period Estimates

It should be noted that the 1-year and 3-year estimates represent period estimates. Due to the difficulty in interpreting the “difference” in period estimates of different lengths, the Census Bureau currently discourages users from making such comparisons. The standard error of this difference is approximated as:

$$SE(\hat{X}_{1\text{-Year}} - \hat{Y}_{3\text{-Year}}) = \sqrt{\frac{1}{3} [SE(\hat{X}_{1\text{-Year}})]^2 + [SE(\hat{Y}_{3\text{-Year}})]^2} \quad (7)$$

Comparing 1-Year Period Estimates with Overlapping 5-Year Period Estimates

It should be noted that the 1-year and 5-year estimates represent period estimates. Due to the difficulty in interpreting the “difference” in period estimates of different lengths, the Census Bureau currently discourages users from making such comparisons. The standard error of this difference is approximated as:

$$SE(\hat{X}_{1\text{-Year}} - \hat{Y}_{5\text{-Year}}) = \sqrt{\frac{3}{5} [SE(\hat{X}_{1\text{-Year}})]^2 + [SE(\hat{Y}_{5\text{-Year}})]^2} \quad (8)$$

TESTING FOR SIGNIFICANT DIFFERENCES

Significant differences – Users may conduct a statistical test to see if the difference between an ACS estimate and any other chosen estimates is statistically significant at a given confidence level. “Statistically significant” means that the difference is not likely due to random chance alone. With the two estimates (Est_1 and Est_2) and their respective standard errors (SE_1 and SE_2), calculate a Z statistic:

$$Z = \frac{Est_1 - Est_2}{\sqrt{(SE_1)^2 + (SE_2)^2}} \quad (9)$$

If $Z > 1.645$ or $Z < -1.645$, then the difference can be said to be statistically significant at the 90 percent confidence level.¹ Any estimate can be compared to an ACS estimate using this method, including other ACS estimates from the current year, the ACS estimate for the same characteristic and geographic area but from a previous year, 2010 Census counts, estimates from other Census Bureau surveys, and estimates from other sources. Not all estimates have sampling error (2010 Census counts do not, for example), but they should be used if they exist to give the most accurate result of the test.

Users are also cautioned to *not* rely on looking at whether confidence intervals for two estimates overlap or not to determine statistical significance, because there are circumstances where that method will not give the correct test result. If two confidence intervals do not overlap, then the estimates will be significantly different (i.e. the significance test will always agree). However, if two confidence intervals do overlap, then the estimates may or may not be significantly different. The Z calculation above is recommended in all cases.

Here is a simple example of why it is not recommended to use the overlapping confidence bounds rule of thumb as a substitute for a statistical test.

Let: $X_1 = 5.0$ with $SE_1 = 0.2$ and $X_2 = 6.0$ with $SE_2 = 0.5$.

The Upper Bound for $X_1 = 5.0 + 0.2 * 1.645 = 5.3$ while the Lower Bound for $X_2 = 6.0 - 0.5 * 1.645 = 5.2$. The confidence bounds overlap, so, the rule of thumb would indicate that the estimates are not significantly different at the 90% level.

However, if we apply the statistical significance test we obtain:

$$Z = \frac{5 - 6}{\sqrt{0.2^2 + 0.5^2}} = 1.857$$

$Z = 1.857 > 1.645$ which means that the difference is significant (at the 90% level).

All statistical testing in ACS data products is based on the 90 percent confidence level. Users should understand that all testing was done using *unrounded* estimates and standard errors, and it may not be possible to replicate test results using the rounded estimates and margins of error as published.

EXAMPLES OF STANDARD ERROR CALCULATIONS

We will present some examples based on the real data to demonstrate the use of the formulas.

¹ The ACS Accuracy of the Data document in 2005 used a Z statistic of +/-1.65. Data users should use +/-1.65 for estimates published in 2005 or earlier.

Example 1 - Calculating the Standard Error from the Confidence Interval

The estimated number of males, never married is 42,741,179 from summary table B12001 for the United States for 2010. The margin of error is 100,944.

$$\text{Standard Error} = \text{Margin of Error} / 1.645$$

Calculating the standard error using the margin of error, we have:

$$SE(42,741,179) = \frac{100,944}{1.645} = 61,364.$$

Example 2 - Calculating the Standard Error of a Sum or a Difference

We are interested in the number of people who have never been married. From Example 1, we know the number of males, never married is 42,741,179. From summary table B12001 we have the number of females, never married is 36,898,838 with a margin of error of 85,983. So, the estimated number of people who have never been married is $42,741,179 + 36,898,838 = 79,640,017$. To calculate the approximate standard error of this sum, we need the standard errors of the two estimates in the sum. We have the standard error for the number of males never married from example 1 as 61,364. The standard error for the number of females never married is calculated using the margin of error:

$$SE(36,898,838) = 85,983 / 1.645 = 52,269.$$

So using formula (2) for the approximate standard error of a sum or difference we have:

$$SE(79,640,017) = \sqrt{61,364^2 + 52,269^2} = 80,608$$

Caution: This method will underestimate or overestimate the standard error if the two estimates interact in either a positive or negative way.

To calculate the lower and upper bounds of the 90 percent confidence interval around 79,640,017 using the standard error, simply multiply 80,608 by 1.645, then add and subtract the product from 79,640,017. Thus the 90 percent confidence interval for this estimate is $[79,640,017 - 1.645(80,608)]$ to $[79,640,017 + 1.645(80,608)]$ or 79,507,417 to 79,772,617.

Example 3 - Calculating the Standard Error of a Proportion/Percent

We are interested in the percentage of females who have never been married to the number of people who have never been married. The number of females, never married is 36,898,838 and the number of people who have never been married is 79,640,017. To calculate the approximate standard error of this percent, we need the standard errors of the two estimates in the percent. We have the approximate standard error for the number of females never married from example 2 as 52,269 and the approximate standard error for the number of people never married calculated from example 2 as 80,608.

The estimate is $(36,898,838 / 79,640,017) * 100\% = 46.33\%$

So, using formula (4) for the approximate standard error of a proportion or percent, we have:

$$SE(46.33\%) = 100\% * \left(\frac{1}{79,640,017} \sqrt{52,269^2 - 0.4633^2 * 80,608^2} \right) = 0.05$$

To calculate the lower and upper bounds of the 90 percent confidence interval around 46.33 using the standard error, simply multiply 0.05 by 1.645, then add and subtract the product from 46.33. Thus the 90 percent confidence interval for this estimate is $[46.33 - 1.645(0.05)]$ to $[46.33 + 1.645(0.05)]$, or 46.25% to 46.41%.

Example 4 - Calculating the Standard Error of a Ratio

Now, let us calculate the estimate of the ratio of the number of unmarried males to the number of unmarried females and its approximate standard error. From the above examples, the estimate for the number of unmarried men is 42,741,179 with a standard error of 61,364, and the estimate for the number of unmarried women is 36,898,838 with a standard error of 52,269.

The estimate of the ratio is $42,741,179 / 36,898,838 = 1.158$.

Using formula (3) for the approximate standard error of this ratio we have:

$$SE(1.158) = \frac{1}{36,898,838} \sqrt{61,364^2 + 1.158^2 * 52,269^2} = 0.00234$$

The 90 percent margin of error for this estimate would be 0.00234 multiplied by 1.645, or about 0.004. The 90 percent lower and upper 90 percent confidence bounds would then be $[1.158 - 0.004]$ to $[1.158 + 0.004]$, or 1.154 and 1.162.

Example 5 - Calculating the Standard Error of a Product

We are interested in the number of 1-unit detached owner-occupied housing units. The number of owner-occupied housing units is 74,873,372 with a margin of error of 216,091 from subject table S2504 for 2010, and the percent of 1-unit detached owner-occupied housing units is 82.0% (0.820) with a margin of error of 0.1 (0.001). So the number of 1-unit detached owner-occupied housing units is $74,873,372 * 0.820 = 61,396,165$. Calculating the standard error for the estimates using the margin of error we have:

$$SE(74,873,372) = 216,091/1.645 = 131,362$$

and

$$SE(0.820) = 0.001/1.645 = 0.0006079$$

The approximate standard error for number of 1-unit detached owner-occupied housing units is calculated using formula (5) for products as:

$$SE(61,396,165) = \sqrt{74,873,372^2 * 0.0006079^2 + 0.820^2 * 131,362^2} = 116,938$$

To calculate the lower and upper bounds of the 90 percent confidence interval around 61,396,165 using the standard error, simply multiply 116,938 by 1.645, then add and subtract the product from 61,396,165. Thus the 90 percent confidence interval for this estimate is $[61,396,165 - 1.645(116,938)]$ to $[61,396,165 + 1.645(116,938)]$ or 61,203,802 to 61,588,528.

CONTROL OF NONSAMPLING ERROR

As mentioned earlier, sample data are subject to nonsampling error. This component of error could introduce serious bias into the data, and the total error could increase dramatically over that which would result purely from sampling. While it is impossible to completely eliminate nonsampling error from a survey operation, the Census Bureau attempts to control the sources of such error during the collection and processing operations. Described below are the primary sources of nonsampling error and the programs instituted for control of this error. The success of these programs, however, is contingent upon how well the instructions were carried out during the survey.

- **Coverage Error** — It is possible for some sample housing units or persons to be missed entirely by the survey (undercoverage), but it is also possible for some sample housing units and persons to be counted more than once (overcoverage). Both the undercoverage and overcoverage of persons and housing units can introduce biases into the data, increase respondent burden and survey costs.

A major way to avoid coverage error in a survey is to ensure that its sampling frame, for ACS an address list in each state, is as complete and accurate as possible. The source of addresses for the ACS is the MAF. An attempt is made to assign all appropriate geographic codes to each MAF address via an automated procedure using the Census Bureau TIGER (Topologically Integrated Geographic Encoding and Referencing) files. A manual coding operation based in the appropriate regional offices is attempted for addresses which could not be automatically coded. The MAF was used as the source of addresses for selecting sample housing units and mailing questionnaires. TIGER produced the location maps for CAPI assignments. Sometimes the MAF has an address that is the duplicate of another address already on the MAF. This could occur when there is a slight difference in the address such as 123 Main Street versus 123 Maine Street.

In the CATI and CAPI nonresponse follow-up phases, efforts were made to minimize the chances that housing units that were not part of the sample were interviewed in place of units in sample by mistake. If a CATI interviewer called a mail nonresponse case and was not able to reach the exact address, no interview was conducted and the case was eligible for CAPI. During CAPI follow-up, the interviewer had to locate the exact address for each sample housing unit. If the interviewer could not locate the exact sample unit in a multi-unit structure, or found a different number of units than expected, the interviewers were instructed to list the units in the building and follow a specific procedure to select a replacement sample unit. Person overcoverage can occur when an individual is included as a member of a housing unit but does not meet ACS residency rules.

Coverage rates give a measure of undercoverage or overcoverage of persons or housing units in a given geographic area. Rates below 100 percent indicate undercoverage, while rates above 100 percent indicate overcoverage. Coverage rates are released concurrent with the release of estimates on American FactFinder in the B98 series of detailed tables. Further information about ACS coverage rates may be found at http://www.census.gov/acs/www/methodology/methodology_main/.

- Nonresponse Error — Survey nonresponse is a well-known source of nonsampling error. There are two types of nonresponse error – unit nonresponse and item nonresponse. Nonresponse errors affect survey estimates to varying levels depending on amount of nonresponse and the extent to which nonrespondents differ from respondents on the characteristics measured by the survey. The exact amount of nonresponse error or bias on an estimate is almost never known. Therefore, survey researchers generally rely on proxy measures, such as the nonresponse rate, to indicate the potential for nonresponse error.
 - o Unit Nonresponse — Unit nonresponse is the failure to obtain data from housing units in the sample. Unit nonresponse may occur because households are unwilling

or unable to participate, or because an interviewer is unable to make contact with a housing unit. Unit nonresponse is problematic when there are systematic or variable differences between interviewed and noninterviewed housing units on the characteristics measured by the survey. Nonresponse bias is introduced into an estimate when differences are systematic, while nonresponse error for an estimate evolves from variable differences between interviewed and noninterviewed households.

The ACS makes every effort to minimize unit nonresponse, and thus, the potential for nonresponse error. First, the ACS used a combination of mail, CATI, and CAPI data collection modes to maximize response. The mail phase included a series of three to four mailings to encourage housing units to return the questionnaire. Subsequently, mail nonrespondents (for which phone numbers are available) were contacted by CATI for an interview. Finally, a subsample of the mail and telephone nonrespondents was contacted by personal visit to attempt an interview. Combined, these three efforts resulted in a very high overall response rate for the ACS.

ACS response rates measure the percent of units with a completed interview. The higher the response rate, and consequently the lower the nonresponse rate, the less chance estimates may be affected by nonresponse bias. Response and nonresponse rates, as well as rates for specific types of nonresponse, are released concurrent with the release of estimates on American FactFinder in the B98 series of detailed tables. Further information about response and nonresponse rates may be found at http://www.census.gov/acs/www/methodology/methodology_main/.

- o Item Nonresponse — Nonresponse to particular questions on the survey questionnaire and instrument allows for the introduction of error or bias into the data, since the characteristics of the nonrespondents have not been observed and may differ from those reported by respondents. As a result, any imputation procedure using respondent data may not completely reflect this difference either at the elemental level (individual person or housing unit) or on average.

Some protection against the introduction of large errors or biases is afforded by minimizing nonresponse. In the ACS, item nonresponse for the CATI and CAPI operations was minimized by the requirement that the automated instrument receive a response to each question before the next one could be asked. Questionnaires returned by mail were edited for completeness and acceptability. They were reviewed by computer for content omissions and population coverage. If necessary, a telephone follow-up was made to obtain missing information. Potential coverage errors were included in this follow-up.

Allocation tables provide the weighted estimate of persons or housing units for which a value was imputed, as well as the total estimate of persons or housing units that

were eligible to answer the question. The smaller the number of imputed responses, the lower the chance that the item nonresponse is contributing a bias to the estimates. Allocation tables are released concurrent with the release of estimates on American Factfinder in the B99 series of detailed tables with the overall allocation rates across all person and housing unit characteristics in the B98 series of detailed tables. Additional information on item nonresponse and allocations can be found at http://www.census.gov/acs/www/methodology/methodology_main/.

- Measurement and Processing Error — The person completing the questionnaire or responding to the questions posed by an interviewer could serve as a source of error, although the questions were cognitively tested for phrasing and detailed instructions for completing the questionnaire were provided to each household.
 - Interviewer monitoring — The interviewer may misinterpret or otherwise incorrectly enter information given by a respondent; may fail to collect some of the information for a person or household; or may collect data for households that were not designated as part of the sample. To control these problems, the work of interviewers was monitored carefully. Field staff were prepared for their tasks by using specially developed training packages that included hands-on experience in using survey materials. A sample of the households interviewed by CAPI interviewers was reinterviewed to control for the possibility that interviewers may have fabricated data.
 - Processing Error — The many phases involved in processing the survey data represent potential sources for the introduction of nonsampling error. The processing of the survey questionnaires includes the keying of data from completed questionnaires, automated clerical review, follow-up by telephone, manual coding of write-in responses, and automated data processing. The various field, coding and computer operations undergo a number of quality control checks to insure their accurate application.
 - Content Editing — After data collection was completed, any remaining incomplete or inconsistent information was imputed during the final content edit of the collected data. Imputations, or computer assignments of acceptable codes in place of unacceptable entries or blanks, were needed most often when an entry for a given item was missing or when the information reported for a person or housing unit on that item was inconsistent with other information for that same person or housing unit. As in other surveys and previous censuses, the general procedure for changing unacceptable entries was to allocate an entry for a person or housing unit that was consistent with entries for persons or housing units with similar characteristics. Imputing acceptable values in place of blanks or unacceptable entries enhances the usefulness of the data.

ISSUES WITH APPROXIMATING THE STANDARD ERROR OF LINEAR COMBINATIONS OF MULTIPLE ESTIMATES

Several examples are provided here to demonstrate how different the approximated standard errors of sums can be compared to those derived and published with ACS microdata.

- A. Suppose we wish to estimate the total number of males with income below the poverty level in the past 12 months using both state and PUMA level estimates for the state of Wyoming. Part of the collapsed table C17001 is displayed below with estimates and their margins of error in parentheses.

Table A: 2010 Estimates of Males with Income Below Poverty from table C17001: Poverty Status in the Past 12 Months by Sex by Age

Characteristic	Wyoming	PUMA 00100	PUMA 00200	PUMA 00300	PUMA 00400
Male	23,001 (3,309)	5,264 (1,624)	6,508 (1,395)	4,364 (1,026)	6,865 (1,909)
Under 18 Years Old	8,479 (1,874)	2,041 (920)	2,222 (778)	1,999 (750)	2,217 (1,192)
18 to 64 Years Old	12,976 (2,076)	3,004 (1,049)	3,725 (935)	2,050 (635)	4,197 (1,134)
65 Years and Older	1,546 (500)	219 (237)	561 (286)	315 (173)	451 (302)

2010 American FactFinder

The first way is to sum the three age groups for Wyoming:

$$\text{Estimate(Male)} = 8,479 + 12,976 + 1,546 = 23,001.$$

The first approximation for the standard error in this case gives us:

$$SE(\text{Male}) = \sqrt{\left(\frac{1,874}{1.645}\right)^2 + \left(\frac{2,076}{1.645}\right)^2 + \left(\frac{500}{1.645}\right)^2} = 1,727.1$$

A second way is to sum the four PUMA estimates for Male to obtain:

$$\text{Estimate(Male)} = 5,264 + 6,508 + 4,364 + 6,865 = 23,001 \text{ as before.}$$

The second approximation for the standard error yields:

$$SE(Male) = \sqrt{\left(\frac{1,624}{1.645}\right)^2 + \left(\frac{1,395}{1.645}\right)^2 + \left(\frac{1,026}{1.645}\right)^2 + \left(\frac{1,909}{1.645}\right)^2} = 1,851.9$$

Finally, we can sum up all three age groups for all four PUMAs to obtain an estimate based on a total of twelve estimates:

$$Estimate(Male) = 2,041 + 2,222 + \dots + 451 = 23,001$$

And the third approximated standard error is

$$SE(Male) = \sqrt{\left(\frac{920}{1.645}\right)^2 + \left(\frac{778}{1.645}\right)^2 + \dots + \left(\frac{302}{1.645}\right)^2} = 1,649.0$$

However, we do know that the standard error using the published MOE is $3,309 / 1.645 = 2,011.6$. In this instance, all of the approximations under-estimate the published standard error and should be used with caution.

- B. Suppose we wish to estimate the total number of males at the national level using age and citizenship status. The relevant data from table B05003 is displayed in table B below.

Table B: 2010 Estimates of males from B05003: Sex by Age by Citizenship Status

Characteristic	Estimate	MOE
Male	151,375,321	27,279
Under 18 Years	38,146,514	24,365
Native	36,747,407	31,397
Foreign Born	1,399,107	20,177
Naturalized U.S. Citizen	268,445	10,289
Not a U.S. Citizen	1,130,662	20,228
18 Years and Older	113,228,807	23,525
Native	95,384,433	70,210
Foreign Born	17,844,374	59,750
Naturalized U.S. Citizen	7,507,308	39,658
Not a U.S. Citizen	10,337,066	65,533

2010 American FactFinder

The estimate and its MOE are actually published. However, if they were not available in the tables, one way of obtaining them would be to add together the number of males under 18 and over 18 to get:

$$Estimate(Male) = 38,146,514 + 113,228,807 = 151,375,321$$

And the first approximated standard error is

$$SE(Male) = SE(151,375,321) = \sqrt{\left(\frac{24,365}{1.645}\right)^2 + \left(\frac{23,525}{1.645}\right)^2} = 20,588.8$$

Another way would be to add up the estimates for the three subcategories (Native, and the two subcategories for Foreign Born: Naturalized U.S. Citizen, and Not a U.S. Citizen), for males under and over 18 years of age. From these six estimates we obtain:

$$\begin{aligned} Estimate(Male) \\ &= 36,747,407 + 268,445 + 1,130,662 + 95,384,433 + 7,507,308 \\ &+ 10,337,066 = 151,375,321. \end{aligned}$$

With a second approximated standard error of:

$$\begin{aligned} SE(Male) &= SE(151,375,321) \\ &= \sqrt{\left(\frac{31,397}{1.645}\right)^2 + \left(\frac{10,289}{1.645}\right)^2 + \left(\frac{20,228}{1.645}\right)^2 + \left(\frac{70,210}{1.645}\right)^2 + \left(\frac{39,658}{1.645}\right)^2 + \left(\frac{65,533}{1.645}\right)^2} \\ &= 67,413.0 \end{aligned}$$

We do know that the standard error using the published margin of error is $27,279 / 1.645 = 16,583.0$. With a quick glance, we can see that the ratio of the standard error of the first method to the published-based standard error yields 1.24; an over-estimate of roughly 24%, whereas the second method yields a ratio of 4.07 or an over-estimate of 307%. This is an example of what could happen to the approximate SE when the sum involves a controlled estimate. In this case, it is sex by age.

- C. Suppose we are interested in the total number of people aged 65 or older and its standard error. Table C shows some of the estimates for the national level from table B01001 (the estimates in gray were derived for the purpose of this example only).

Table C: Some Estimates from AFF Table B01001: Sex by Age for 2010

Age Category	Estimate, Male	MOE, Male	Estimate, Female	MOE, Female	Total	Estimated MOE, Total
65 and 66 years old	2,492,871	20,194	2,803,516	23,327	5,296,387	30,854
67 to 69 years old	3,029,709	18,280	3,483,447	24,287	6,513,225	30,398
70 to 74 years old	4,088,428	21,588	4,927,666	26,867	9,016,094	34,466
75 to 79 years old	3,168,175	19,097	4,204,401	23,024	7,372,576	29,913
80 to 84 years old	2,258,021	17,716	3,538,869	25,423	5,796,890	30,987
85 years and older	1,743,971	17,991	3,767,574	19,294	5,511,545	26,381
Total	16,781,175	NA	22,725,473	NA	39,506,648	74,932

To begin we find the total number of people aged 65 and over by simply adding the totals for males and females to get $16,781,175 + 22,725,542 = 39,506,717$. One way we could use is summing males and female for each age category and then using their MOEs to approximate the standard error for the total number of people over 65.

$$MOE(\text{Age 65 and 66}) = MOE(5,296,387) = \sqrt{20,194^2 + 23,327^2} = 30,854$$

$$MOE(\text{Age 67 to 69}) = MOE(6,513,225) = \sqrt{18,280^2 + 24,287^2} = 30,398$$

... etc. ...

Now, we calculate for the number of people aged 65 or older to be 39,506,648 using the six derived estimates and approximate the standard error:

$$SE(39,506,717) = \sqrt{\left(\frac{30,854}{1.645}\right)^2 + \left(\frac{30,398}{1.645}\right)^2 + \left(\frac{34,466}{1.645}\right)^2 + \left(\frac{29,913}{1.645}\right)^2 + \left(\frac{30,987}{1.645}\right)^2 + \left(\frac{26,381}{1.645}\right)^2} = 45,552$$

For this example the estimate and its MOE are published in table B09017. The total number of people aged 65 or older is 39,506,648 with a margin of error of 20,689. Therefore the published-based standard error is:

$$SE(39,506,648) = 20,689/1.645 = 12,577.$$

The approximated standard error, using six derived age group estimates, yields an approximated standard error roughly 3.6 times larger than the published-based standard error.

As a note, there are two additional ways to approximate the standard error of people aged 65 and over in addition to the way used above. The first is to find the published MOEs for the males age 65 and older and of females aged 65 and older separately and then combine to find the approximate standard error for the total. The second is to use all twelve of the published estimates together, that is, all estimates from the male age categories and female age categories, to create the SE for people aged 65 and older. However, in this particular example, the results from all three ways are the same. So no matter which way you use, you will obtain the same approximation for the SE. This is different from the results seen in example A.

- D. For an alternative to approximating the standard error for people 65 years and older seen in part C, we could find the estimate and its SE by summing all of the estimate for the ages less than 65 years old and subtracting them from the estimate for the total population. Due to the large number of estimates, Table D does not show all of the age groups. In addition, the estimates in part of the table shaded gray were derived for the purposes of this example only and cannot be found in base table B01001.

Table D: Some Estimates from AFF Table B01001: Sex by Age for 2010:

Age Category	Estimate, Male	MOE, Male	Estimate, Female	MOE, Female	Total	Estimated MOE, Total
Total Population	151,375,321	27,279	155,631,235	27,280	307,006,556	38,579
Under 5 years	10,853,263	15,661	10,355,944	14,707	21,209,207	21,484
5 to 9 years old	10,273,948	43,555	9,850,065	42,194	20,124,013	60,641
10 to 14 years old	10,532,166	40,051	9,985,327	39,921	20,517,493	56,549
...	...	...	...	...		
62 to 64 years old	4,282,178	25,636	4,669,376	28,769	8,951,554	38,534
Total for Age 0 to 64 years old	134,594,146	117,166	132,905,762	117,637	267,499,908	166,031
Total for Age 65 years and older	16,781,175	120,300	22,725,473	120,758	39,506,648	170,454

An estimate for the number of people age 65 and older is equal to the total population minus the population between the ages of zero and 64 years old:

Number of people aged 65 and older: $307,006,556 - 267,499,908 = 39,506,648$.

The way to approximate the SE is the same as in part C. First we will sum male and female estimates across each age category and then approximate the MOEs. We will use that information to approximate the standard error for our estimate of interest:

$$MOE(\text{Total Population}) = MOE(307,006,556) = \sqrt{27,279^2 + 27,280^2} = 38,579$$

$$MOE(\text{Under 5 years}) = MOE(21,209,207) = \sqrt{15,661^2 + 14,707^2} = 21,484$$

... etc. ...

And the SE for the total number of people aged 65 and older is:

$$\begin{aligned}
SE(\text{Age 65 and older}) &= SE(39,506,648) \\
&= \sqrt{\left(\frac{38,579}{1.645}\right)^2 + \left(\frac{21,484}{1.645}\right)^2 + \left(\frac{60,641}{1.645}\right)^2 + \left(\frac{56,549}{1.645}\right)^2 + \dots + \left(\frac{38,534}{1.645}\right)^2} \\
&= 103,620
\end{aligned}$$

Again, as in Example C, the estimate and its MOE are we published in B09017. The total number of people aged 65 or older is 39,506,648 with a margin of error of 20,689. Therefore the standard error is:

$$SE(39,506,648) = 20,689 / 1.645 = 12,577.$$

The approximated standard error using the thirteen derived age group estimates yields a standard error roughly 8.2 times larger than the actual SE.

Data users can mitigate the problems shown in examples A through D to some extent by utilizing a collapsed version of a detailed table (if it is available) which will reduce the number of estimates used in the approximation. These issues may also be avoided by creating estimates and SEs using the Public Use Microdata Sample (PUMS) or by requesting a custom tabulation, a fee-based service offered under certain conditions by the Census Bureau. More information regarding custom tabulations may be found at

http://www.census.gov/acs/www/data_documentation/custom_tabulations/.