

ending shall be sufficient to completely above the device



ILLINOIS TRAFFIC AND PEDESTRIAN STOP STUDY 2024 ANNUAL REPORT TRAFFIC STOP ANALYSIS

SUBMITTED BY THE MOUNTAIN-WHISPER-LIGHT: STATISTICS AND DATA SCIENCE

The
Mountain-Whisper-Light



Illinois Traffic and Pedestrian Stop Study

2024 ANNUAL REPORT: TRAFFIC STOPS

Part I Executive Summary and Appendices

Prepared for the Illinois Department of Transportation

By

The Mountain-Whisper-Light: Statistics & Data Science



In Cooperation with SC-B Consulting Inc.



Acknowledgements

The Mountain-Whisper-Light, Inc. and SC-B Consulting, Inc. gratefully acknowledge the contributions of Jessica Keldermans and Sean Berberet, Bureau of Data Collection, Illinois Department of Transportation, and their associates at IDOT. Many thanks to the office of the Illinois Secretary of State for providing summary statistics on driver's licenses (without any personally identifying information). We are also grateful to the law enforcement agencies and their staff for their diligence and patience in collecting and processing the data used in this study.

2024 study team, Illinois Traffic and Pedestrian Stop Study

The Mountain-Whisper-Light: Statistics & Data Science

- Nayak L. Polissar, PhD, President & Principal Investigator
- Nirnaya Miljacic, MS, PhD
- Medur Wilson, MS
- Cindy Elder, MA
- Daniel S. Hippe, MS

SC-B Consulting

- Sharie Carter-Bane, President

Table of Contents

Executive Summary	1
I. Background	1
Key Findings.....	2
II. Introduction.....	2
What is racial profiling?.....	2
How is this report structured?.....	2
What is the source of the data?.....	3
How were the data analyzed?.....	3
References (for Section II).....	4
III. Guide to Using Traffic Tables	4
IV. Interpretation of Traffic Tables	10
95% Confidence Interval.....	10
Ratios.....	10
Limitations.....	10
Statistics based on stops only.....	11
Can stop rates be compared across years?.....	11
V. Benchmarks	12
Population Benchmarks.....	13
Mileage Benchmarks.....	14
Initial comparison of the two benchmark models.....	15
VI. Selected Findings.....	17
Agency reporting status.....	18
Number of stops.....	18
Statewide stop rates and rate ratios.....	21
Distributions of stop rate ratios.....	21
Reason for Stop.....	27
Outcome of Stop: Citation.....	29
Searches.....	29
Dog Sniffs.....	31
Multiple stopping of individual drivers.....	31
Officer-assigned drivers race analysis.....	33
VII. Considerations for Interpreting the Data	36
VIII. Looking Ahead	37
Appendix A. Traffic Stop Data Collection Form in use during 2024.....	39
Appendix B. Technical Notes on Rates, Percentages and Ratios.....	40
Appendix C. Technical Notes on Population Benchmarks.....	44
Appendix D. Technical Notes on Mileage Benchmarks.....	47
Appendix E. Additional Notes on Illinois Law Concerning the Stops Study.....	50

THIS PAGE INTENTIONALLY BLANK

Executive Summary

I. Background

In October 2019, The Mountain-Whisper-Light Inc. (aka The Mountain-Whisper-Light: Statistics & Data Science, and hereafter, “TMWL”) was awarded a contract to conduct a statistical study of the traffic and pedestrian stop data provided by Illinois law enforcement agencies. Agencies provided their stop data to the Illinois Department of Transportation, pursuant to the Illinois Vehicle Code, 625 ILCS 5/11-212 Traffic and Pedestrian Stop Statistical Study. TMWL is carrying out the project in cooperation with SC-B Consulting Inc., an Illinois firm (hereafter, SC-B). Reports have already been issued on 2019-2023 traffic and pedestrian stops in Illinois and are available online at <https://www.idot.illinois.gov/transportation-system/local-transportation-partners/law-enforcement/illinois-traffic-stop-study>. (Click on “Studies.”)

According to the IDOT website, “On July 18, 2003, Senate Bill 30 was signed into law to establish a four-year statewide study of data from traffic stops to identify racial bias. The study began on January 1, 2004, and was originally scheduled to end December 31, 2007. However, the legislature extended the data collection several times, and also expanded the study to include data on pedestrian stops. Public Act 101-0024, which took effect on June 21, 2019, eliminated the study's scheduled end date of July 1, 2019, and extended the data collection.”

Under that provision of the Illinois Vehicle Code, IDOT is responsible for providing a standardized law enforcement data compilation form (see Appendix A below) and analyzing the data and submitting a report of the previous year's findings to the Governor, General Assembly, the Racial Profiling Prevention and Data Oversight Board, and each law enforcement agency (receiving results on their own data) no later than July 1 of each year. In May 2025, TMWL and SC-B, in cooperation with IDOT's Bureau of Data Collection, provided copies of statistical tables for 789 law enforcement agencies in the state of Illinois, based on data collection provided by the respective agencies on traffic and pedestrian stops. These 789 agencies reported at least one traffic or pedestrian stop. Among these agencies, 788 reported on traffic stops or on both traffic stops and pedestrian stops. One agency reported only on pedestrian stops. The agencies were invited to review and comment on the tables. Some agencies provided comments and the comments—from each agency that commented—are included with their tables in Part II of this report. We responded to some comments with additional information, and the readers of this report may wish to peruse the agency comments and our responses. Comments on the Traffic stop tables (or general comments) and comments on the Pedestrian stop tables are included in the Part II Traffic or Pedestrian tables, respectively. Readers can be assured that the statistical results presented in this report are valid.

We are pleased to submit this 2024 Annual Report for the Illinois Traffic and Pedestrian Stop Study. The Executive Summary in this document covers the traffic stop study and a companion volume with a similar format contains an Executive Summary for the pedestrian stop study.

Key Findings

1. The total number of reported traffic stops in 2024 was 2,050,405, a 9.4% decrease from 2023, and being close to 2015 and 2022 values (Figure 2a).

2. 79% of agencies provided complete stop data for 2024, which is similar to 2023 (78%). Agencies not collecting a full year of stop data (incomplete) or not submitting existing stop data (non-compliant) comprised 20%, an increase from 18% in 2023 (Table 2).
3. Traffic stop rates in 2024 for each of the six racial groups decreased slightly from 2023. Black drivers statewide rate ratio decreased to 1.39, from 1.73 in 2023. (Figure 3).
4. A number of drivers were stopped multiple times in 2024, with Black drivers being involved the most in the repeated stops, although somewhat less prominently than in 2023. (Figure 7). The count of Black drivers that were stopped over 10 times decreased from 865 in 2023 to 322 in 2024.
5. In 78% of all cases where an agency reported at least 50 stops with drivers of any particular Minority group and 50 stops with White drivers, the stop rate was higher for the Minority drivers.
6. In 93% of all cases where an agency reported at least 50 stops with Black drivers and 50 stops with White drivers, the stop rate was higher for Black drivers.
7. In 79% of all cases where an agency reported at least 50 stops with Hispanic drivers and 50 stops with White drivers, the stop rate was higher for Hispanic drivers.

II. Introduction

What is racial profiling?

The Illinois Criminal Justice Information Authority describes racial profiling as “police-led action that is initiated based on a person’s race or ethnicity.”¹ (*References can be found at the end of this section.*) In 2003, legislation called the Illinois Traffic and Pedestrian Stop Statistical Study Act was passed requiring officers to document who/why they stopped individuals for traffic violations. These data are reported annually to the Illinois Department of Transportation for review. In 2019, this Act became permanent and supports a Task Force to compile and analyze the resulting data.² This analysis provides statistical results for use in those ongoing efforts. The statistical results can be used to detect potentially “statistically significant aberrations” in traffic stops, pedestrian stops and searches of drivers and pedestrians (see Section I and Appendix E for more details). Findings are made available to the public and shared with law enforcement agencies to increase their awareness of potential racial profiling in their stops, providing a basis for an agency to reduce or eliminate bias, if it is occurring. The IDOT Racial Profiling Prevention and Data Oversight Board meets regularly to oversee these efforts in an advisory capacity and provide recommendations to the Governor's Office.

How is this report structured?

The report is presented in two parts. **Part I** is this Executive Summary, which includes appendices with detailed technical information on the statistical methodology and analysis. **Part II** includes extensive tables (one set of tables for each law enforcement agency that collected data for all stops reported in 2024). The tables show stop rates for each racial group, along with other statistics that cover activity during the stops, such as citations or warnings, searches and contraband found.

To obtain the greatest benefit from this report, readers are encouraged to read the full Executive Summary. In addition to the information on data collection, the statistical analysis team has provided a sample Traffic Table and a Guide to Using Traffic Tables that includes definitions of statistical terms used in this report. This Executive Summary also provides an explanation of the data presented in each panel

of the Part II tables. The Executive Summary also includes an Interpretation section with additional details on the numeric results presented in the tables and a plain-language description of how the analysis was implemented. Finally, the section on Selected Findings highlights some statewide results. The Appendices include technical material and describe the statistical methods and calculations in detail. The information in the appendices is provided for readers who wish to have a deeper understanding of the methodology.

What is the source of the data?

As noted above, per Illinois law, officers from law enforcement agencies are required to fill in a report when they stop a driver or a pedestrian. Separate templates are provided for traffic stops and pedestrian stops.

To follow the convention of previous reporting on the Illinois Traffic and Pedestrian Stop Study, we are submitting two separate reports, the Illinois **Traffic** Stop Study and the Illinois **Pedestrian** Stop Study. The above-mentioned data collection templates (known as Traffic Stop or Pedestrian Stop Data Forms) are shown in Appendix A of the ITSS and IPSS. There is an instruction manual that accompanies the traffic stops data collection form—available online at <https://www.idot.illinois.gov/transportation-system/local-transportation-partners/law-enforcement/illinois-traffic-stop-study>. (Click on “Forms.”)

How were the data analyzed?

The results of the data collection are that 788 agencies generated data on 2,050,405 traffic stops and 244 agencies generated data on 88,019 pedestrian stops in 2024. A total of 789 agencies provided data on either traffic stops or pedestrian stops, with 545 agencies providing data on traffic stops only. One agency provided pedestrian stop data only, and 243 agencies provided both traffic and pedestrian stop data. Among 788 agencies that provided traffic stops, 786 were considered compliant with the study, and two agencies were deemed non-compliant. None of the reported traffic stops were missing the race designation. Only three pedestrian stops (0.003% of pedestrian stops) were missing the race designation and did not enter into the analysis. All 244 agencies that reported pedestrian stops were deemed compliant with the study. Further analysis was carried out to provide statistics that may be helpful in determining if there is potential bias against minorities in initiating a stop or in the activities that occur during a stop.

As specified by the Illinois statute for this study, the tables report on the stops and subsequent experience of individuals stopped. The stopped individuals are classified into one of six racial groups. The law enforcement officer filling in the data collection form must use their judgment to classify an individual into one of the following six groups.

- Black or African American
- Hispanic or Latino
- Asian
- American Indian or Alaska Native
- Native Hawaiian or Other Pacific Islander
- White

The data collection forms are extensive. There are more than 60 data items listed for traffic stops and more than 20 data items listed for pedestrian stops. Some items are left blank unless there are further actions beyond a stop, such as a search.

Data collected by local agencies for traffic stops include:

- Information about the driver (including race) and the officer
- The location of the stop (using location designations developed by each agency)
- Reason for the stop
- Outcome of the stop
- Search activity and search findings of contraband.

References (for Section II)

1. Green, E., & Lavery, T. (2022). *2020-2021 Illinois Traffic and Pedestrian Stop Data Use and Collection Task Force Findings*. Illinois Criminal Justice Information Authority.
2. Illinois General Assembly. (2022, May 13). *Traffic and Pedestrian Stop Statistical Study*. Website. <https://www.ilga.gov/legislation/ilcs/fulltext.asp?DocName=062500050K11-212>

III. Guide to Using Traffic Tables

While many readers of this report previously reviewed traffic and pedestrian stop tables for their respective jurisdictions, here are some brief explanations of the statistics presented in the tables of this report.

Table 1 (below) is included as an example to show stop rates, along with certain percentages and ratios. A ratio compares either a rate or a percentage for a Minority to the corresponding rate or percentage for Whites. The ratios are intended to make it easier to compare a Minority to Whites on stop rates and other statistics that may suggest the possibility of racial profiling. The word “possibility” is very important, because racial profiling cannot be proved by the numeric results in this report alone. Some of the inherent uncertainties and limitations of the statistics are explained throughout this report and should be considered during the review of the statistical results presented.

The following section includes an example of traffic tables and offers a guide to the numbers in the tables, explained for each panel in the table. The table reproduced here (Table 1, below)) refers to all traffic stops reported in 2024 from law enforcement agencies in the state of Illinois. The counts, rates, percentages, and ratios are for purposes of illustration only and are not tied to any individual agency.

Before using the tables: Following the tables there is an important section on interpretation of the rates, ratios, percentages and 95% confidence intervals (margin of error). Reading that section is important for readers of this report to make a proper assessment of what the numbers represent.

Rates, percentages, and ratios: The terms “rate,” “percentage” and “ratio” are used throughout this report. A brief explanation of the terms is provided here.

A rate in one context is the number of individuals (such as the number of individuals stopped) divided by the population the individuals came from, also known in this report as the population “benchmark.” “Benchmark” is a term that will be used repeatedly. A rate in another context is the number of

individuals (such as the number of individuals stopped) divided by the number of miles that the population the individuals came from have driven within a jurisdiction during an average day, also known in this report as the mileage “benchmark.” For example, in Illinois in 2024 there were 439,060 traffic stops of individuals whom the officer assigned to the category “Hispanic or Latino.” The estimated population benchmark population of Hispanic or Latino drivers in Illinois in 2024 was 2,138,124. Dividing the 439,060 by 2,138,124 yields the stop rate of 0.2053. That is, there was an average of 0.2053 stops per driving member of the Hispanic or Latino population. The decimal value 0.2053 does not mean that 20.53% of Hispanic or Latino drivers had a stop. Some drivers may have been stopped more than once.

A **percentage** in this context has the usual meaning. For example, in Illinois in 2024 there were 1,000,699 stops of drivers whom the officer assigned to the category “White.” There were 605,638 of those stops with a citation for a moving violation. The number of stops with citations (605,638) divided by the number of stops (1,000,600) yields the decimal fraction 0.605. That fraction represented as a percentage is 60.5%. In Illinois in 2024, 60.5% of stops of drivers assessed as being White resulted in a citation of the driver.

The **ratio** used in this report is either the ratio of a Minority rate to a White rate or the ratio of a Minority percentage to a White percentage. If the ratio is 2.0, for example, it means that the Minority rate (or percentage) is twice the White rate (or percentage).

Table 1 shows the Illinois statewide results for illustration of traffic stop reporting. Following is a guide to each panel of the table. Note that only the statistics given in Panel 1 involve benchmarks.

Panel 1 (shaded region) presents the traffic stops and the results of the two benchmark models. After the stops row, the next three rows show the population benchmark and the statistics based on that benchmark: the stop rate by racial group and stop rate ratio for each minority group compared to White drivers. The following three rows (darker shaded rows) show the mileage benchmark and the statistics based on that benchmark: the stop rate by racial group and stop rate ratio for each minority group compared to White drivers. Each benchmark value is followed by the percentage (in parentheses) of the total benchmark for all race groups. Ninety-five percent confidence intervals are shown (in parentheses) for rates and rate ratios. The 95% confidence interval is a “margin of error,” and it is explained in a short section with that heading, below.

Panel 2 shows the number, percentage (in parentheses), and 95% confidence interval [in square brackets, like this] for selected reasons for traffic stops (moving violation, equipment, licensing/registration and commercial vehicle) for each racial group. The label for the panel includes the note “Percentage of All Stops for the Racial Group with the Noted Reason for Stop.” This tells us that the number of stops for a given reason, such as “Moving Violation,” is divided by the total number of stops for the racial group — to convert it to a percentage (after multiplication by 100%). For example, drivers assessed as being Asian had 53,377 stops noted by the officer as “Moving Violation,” and the Asian category had 84,589 total stops in 2024. Hence the percentage of stops noted as “Moving Violation” for drivers classified as Asian was $100\% \times (53,377/84,589) = 63.1\%$ (rounded).

Panel 3 shows the outcomes of traffic stops including written warning, verbal warning and citation for each racial group. The number, percentage (in parentheses) and 95% confidence interval [in brackets] are shown for each outcome. The ratio and 95% confidence interval (in parentheses)

comparing each Minority group to White drivers are shown for citations, the most serious outcome recorded for the stop on the traffic data collection form.

Panel 4 shows vehicle searches and outcomes of vehicle searches during traffic stops, including consent searches, all searches, and, for each racial group, whether contraband was found during any search. The number, percentage (in parentheses) and 95% confidence interval [in brackets] are shown for each outcome. The label for each row shows the basis for calculation of the percentages. The contraband-found percentage is calculated based on all vehicle searches. The ratio and 95% confidence interval (in parentheses) comparing each Minority group to White drivers are shown for contraband-found for all vehicle searches. (Note: searches following a dog sniff are not included in Panel 4. See Panel 6 for the statistics on stops with a dog sniff.)

Panel 5 shows driver and passenger searches and outcomes of these searches during traffic stops including consent searches, all searches and whether contraband was found during any search for each racial group. The number, percentage (in parentheses) and 95% confidence interval [in brackets] are shown for each outcome. The label for each row shows the basis for calculation of the percentages. The contraband found percentage is calculated based on all driver or passenger searches. The ratio and 95% confidence interval (in parentheses) comparing each Minority group to White drivers are shown for contraband found for all driver or passenger searches. (Note: searches following a dog sniff are not included in Panel 5. See Panel 6 for the statistics on stops with a dog sniff.)

Panel 6 shows dog sniffs, searches and outcomes of these searches during traffic stops, including dog alerts during a dog sniff, vehicle searches after a dog sniff and whether contraband was found after any vehicle search for each racial group. The number, percentage (in parentheses) and 95% confidence interval [in brackets] are shown for each outcome. The label for each row shows the basis for calculation of the percentages. The percentage of dog sniffs with a dog alert and the percentage of vehicle searches after a dog sniff are calculated based on all dog sniffs. The percentage for contraband found after a vehicle search is calculated based on all vehicle searches after a dog sniff, and the ratio and 95% confidence interval (in parentheses) are shown for contraband found for all vehicle searches after a dog sniff.

The top-right corner of Table 1 indicates the type of population benchmark used. Crash-based population benchmarks utilize Illinois crash report data and distance-based population benchmarks combine population statistics from surrounding ZIP codes while accounting for distance of the ZIP code area to the agency. The note at the bottom-left of the table indicates the type of population benchmark (crash-based or distance-based), and if the benchmark is crash-based, the note states the number of crashes that were utilized in calculations. The note also lists the primary area of the population benchmark, which captures the jurisdiction of the agency. These areas can be one or more cities (or towns or villages), counties or the state. All traffic population benchmarks also include areas outside of the primary area. The percentage of the population benchmark that comes from ZIP codes within the primary area is provided, and an indication of the overall area of the benchmark is provided by a radius around the primary area (in miles). Section V on benchmarks provides more information on how the benchmarks were constructed.

Note that since drivers from outside of an agency jurisdiction may travel through the jurisdiction, and thus may be stopped in the jurisdiction, the benchmark for the jurisdiction includes population residing outside of a jurisdiction. A common error of interpretation is that the benchmark for traffic stops involves only the jurisdiction population.

A ratio of 1.0 for Whites: For all rows showing comparisons of Minority groups to Whites, a value of 1.0 is shown in the White racial group column, the reference group. In this column for Whites, the Whites are being compared to themselves, so the ratio of rates must be 1.0. The column is included to make it clear that the Whites are the reference group to which each Minority is compared.

Zero stops or zero benchmark: For some agencies, the number of stops or the benchmark value or the number of outcomes may be zero for a racial group. When it is not possible to calculate a rate, percentage or ratio and an associated 95% confidence interval because of zero stops or zero benchmarks or zero outcomes, an “NA” is reported in the table. When reporting information such as searches following stops or contraband found, there are cases when all racial groups have entries of zero in the row. That is, there were no searches of any racial group, or no contraband found for any racial group. In that case, the row is omitted. Similarly, when making comparisons to Whites, if all minorities have counts of zero or the Whites have a count of zero, the ratios comparing each Minority to Whites cannot be computed and the row of ratios is omitted.

Table 1. Example of a table of traffic stops: Counts, Rates, Percentages and Ratios

Summary of Traffic Stops for 2024 - ILLINOIS STATEWIDE RESULTS					Population Benchmark: Crash-based*	
	White	Black or African American	Hispanic or Latino	Asian	American Indian or Alaska Native	Native Hawaiian or Other Pacific Islander
Panel: 1 Summary of Stops, Rates, and Rate Ratios with 95% Confidence Intervals. Total stops: 2,050,405. Total population benchmark: 9,911,249. Total mileage benchmark: 242,686,003.						
Stops (% of Total)	1,000,699 (49%)	511,178 (25%)	439,060 (21%)	84,589 (4.1%)	9,736 (0.5%)	5,143 (0.3%)
Population Benchmark (% of Total)	5,202,375 (52%)	1,912,683 (19%)	2,138,124 (22%)	619,735 (6.3%)	32,625 (0.3%)	5,707 (0.06%)
Population Stop Rate (95% Confidence Interval)	0.1924 (0.192 - 0.1927)	0.2673 (0.2665 - 0.268)	0.2053 (0.2047 - 0.206)	0.1365 (0.1356 - 0.1374)	0.298 (0.293 - 0.304)	0.9 (0.88 - 0.93)
Population R. Ratio vs Wh. (95% Confidence Interval)	1.0	1.39 (1.38 - 1.4)	1.07 (1.06 - 1.08)	0.71 (0.7 - 0.72)	1.55 (1.52 - 1.58)	4.7 (4.6 - 4.8)
Mileage Benchmark (% of Total)	179,412,400 (74%)	26,734,750 (11%)	22,680,050 (9.3%)	13,027,780 (5.4%)	708,445 (0.3%)	122,637 (0.05%)
Mileage Stop Rate (95% Confidence Interval)	0.00558 (0.00557 - 0.00559)	0.01912 (0.01907 - 0.01917)	0.01936 (0.0193 - 0.01942)	0.00649 (0.00645 - 0.00654)	0.0137 (0.0135 - 0.014)	0.042 (0.041 - 0.043)
Mileage R. Ratio vs White (95% Confidence Interval)	1.0	3.43 (3.42 - 3.44)	3.47 (3.46 - 3.48)	1.164 (1.156 - 1.172)	2.46 (2.41 - 2.51)	7.5 (7.3 - 7.7)
Panel: 2 Summary of Reason for Stop - Number (Percentage of All Stops for the Racial Group with the Noted Reason for Stop) [95% Confidence Interval]						
Moving Violation	605,638 (60.5%) [60.4% - 60.7%]	237,569 (46.5%) [46.3% - 46.7%]	219,630 (50%) [49.8% - 50.2%]	53,377 (63.1%) [62.6% - 63.6%]	5,966 (61%) [60% - 63%]	3,199 (62%) [60% - 64%]
Equipment	155,070 (15.5%) [15.4% - 15.6%]	98,868 (19.3%) [19.2% - 19.5%]	95,443 (21.7%) [21.6% - 21.9%]	15,555 (18.4%) [18.1% - 18.7%]	1,957 (20%) [19% - 21%]	920 (18%) [17% - 19%]
Licensing/Registration	231,663 (23.15%) [23.06% - 23.24%]	171,453 (33.5%) [33.4% - 33.7%]	119,072 (27.1%) [27% - 27.3%]	15,201 (18%) [17.7% - 18.3%]	1,730 (18%) [17% - 19%]	956 (19%) [17% - 20%]
Commercial Vehicle	8,322 (0.83%) [0.81% - 0.85%]	3,240 (0.63%) [0.61% - 0.66%]	4,890 (1.11%) [1.08% - 1.15%]	455 (0.54%) [0.49% - 0.59%]	83 (0.85%) [0.68% - 1.1%]	68 (1.3%) [1% - 1.7%]
Panel: 3 Summary of Outcome of Stop - Number (Percentage of All Stops for the Racial Group with the Noted Outcome of Stop) [95% Confidence Interval]						
Verbal Warning	262,708 (26.3%) [26.2% - 26.4%]	220,035 (43%) [42.9% - 43.2%]	171,364 (39%) [38.8% - 39.2%]	31,023 (36.7%) [36.3% - 37.1%]	3,935 (40%) [39% - 42%]	2,071 (40%) [39% - 42%]
Written Warning	395,300 (39.5%) [39.4% - 39.6%]	130,443 (25.5%) [25.4% - 25.7%]	109,851 (25%) [24.9% - 25.2%]	27,044 (32%) [31.6% - 32.4%]	3,029 (31%) [30% - 32%]	1,395 (27%) [26% - 29%]
Citation	342,691 (34.2%) [34.1% - 34.4%]	160,700 (31.4%) [31.3% - 31.6%]	157,845 (36%) [35.8% - 36.1%]	26,522 (31.4%) [31% - 31.7%]	2,772 (28%) [27% - 30%]	1,677 (33%) [31% - 34%]
Citation Ratio vs White (95% Confidence Interval)	1.0	0.918 (0.913 - 0.923)	1.05 (1.04 - 1.06)	0.92 (0.9 - 0.93)	0.83 (0.8 - 0.86)	0.95 (0.91 - 1)
Panel: 4 Summary of Vehicle Search Events - Number (Percentage for the Racial Group) [95% Confidence Interval]						
Consent Search (% of Stops)	8,193 (0.82%) [0.8% - 0.84%]	7,448 (1.46%) [1.42% - 1.49%]	4,682 (1.07%) [1.04% - 1.1%]	435 (0.51%) [0.47% - 0.56%]	73 (0.75%) [0.59% - 0.94%]	62 (1.2%) [0.92% - 1.5%]

Summary of Traffic Stops for 2024 - ILLINOIS STATEWIDE RESULTS					Population Benchmark: Crash-based*	
	White	Black or African American	Hispanic or Latino	Asian	American Indian or Alaska Native	Native Hawaiian or Other Pacific Islander
All Searches (% of Stops)	33,080 (3.31%) [3.27% - 3.34%]	29,155 (5.7%) [5.6% - 5.8%]	14,369 (3.27%) [3.22% - 3.33%]	941 (1.1%) [1% - 1.2%]	153 (1.6%) [1.3% - 1.8%]	132 (2.6%) [2.1% - 3%]
Contraband Found (% of All Searches)	10,504 (31.8%) [31.1% - 32.4%]	14,875 (51%) [50% - 52%]	5,445 (38%) [37% - 39%]	209 (22%) [19% - 25%]	50 (33%) [24% - 43%]	34 (26%) [18% - 36%]
Contraband Found Ratio vs White (95% Confidence Interval)	1.0	1.61 (1.57 - 1.65)	1.19 (1.15 - 1.23)	0.7 (0.61 - 0.8)	1 (0.76 - 1.4)	0.81 (0.56 - 1.1)
Panel: 5 Summary of Driver or Passenger Search Events - Number (Percentage for the Racial Group) [95% Confidence Interval]						
Consent Search (% of Stops)	5,846 (0.58%) [0.57% - 0.6%]	5,366 (1.05%) [1.02% - 1.08%]	2,840 (0.65%) [0.62% - 0.67%]	146 (0.17%) [0.15% - 0.2%]	40 (0.41%) [0.29% - 0.56%]	35 (0.68%) [0.47% - 0.95%]
All Searches (% of Stops)	18,459 (1.84%) [1.82% - 1.87%]	20,332 (3.98%) [3.92% - 4.03%]	9,984 (2.27%) [2.23% - 2.32%]	447 (0.53%) [0.48% - 0.58%]	85 (0.87%) [0.7% - 1.1%]	70 (1.4%) [1.1% - 1.7%]
Contraband Found (% of All Searches)	2,747 (14.9%) [14.3% - 15.4%]	2,937 (14.4%) [13.9% - 15%]	894 (9%) [8.4% - 9.6%]	43 (9.6%) [7% - 13%]	10 (12%) [5.6% - 22%]	7 (10%) [4% - 21%]
Contraband Found Ratio vs White (95% Confidence Interval)	1.0	0.97 (0.92 - 1)	0.6 (0.56 - 0.65)	0.65 (0.47 - 0.87)	0.79 (0.38 - 1.5)	0.67 (0.27 - 1.4)
Panel: 6 Summary of Dog Sniff Events - Number (Percentage for the Racial Group) [95% Confidence Interval]						
Dog Sniff (% of Stops)	2,839 (0.28%) [0.27% - 0.29%]	1,314 (0.26%) [0.24% - 0.27%]	701 (0.16%) [0.15% - 0.17%]	73 (0.086%) [0.068% - 0.11%]	11 (0.11%) [0.056% - 0.2%]	12 (0.23%) [0.12% - 0.41%]
Dog Alert after Dog Sniff (% of Dog Sniffs)	2,098 (74%) [71% - 77%]	924 (70%) [66% - 75%]	449 (64%) [58% - 70%]	48 (66%) [48% - 87%]	9 (82%) [37% - 100%]	5 (42%) [14% - 97%]
Vehicle Search after Dog Sniff (% of Dog Sniffs)	2,055 (72%) [69% - 76%]	902 (69%) [64% - 73%]	429 (61%) [56% - 67%]	47 (64%) [47% - 86%]	9 (82%) [37% - 100%]	5 (42%) [14% - 97%]
Contraband Found (% of Vehicle Searches, preceding row)	1,278 (62%) [59% - 66%]	591 (66%) [60% - 71%]	173 (40%) [35% - 47%]	16 (34%) [19% - 55%]	7 (78%) [31% - 100%]	2 (40%) [4.8% - 100%]
Contraband Found Ratio vs White (95% Confidence Interval)	1.0	1.1 (0.95 - 1.2)	0.65 (0.55 - 0.76)	0.55 (0.31 - 0.89)	1.3 (0.5 - 2.6)	0.64 (0.078 - 2.3)
*Population Benchmark Definition Population Benchmark Type: Crash-based (169,954 crash reports used). Primary Benchmark Area (State): Illinois. 93.7% of the benchmark comes from ZIP codes within the primary area. 95.1% of the benchmark comes from ZIP codes within 10 miles of the primary area, including the primary area.						

IV. Interpretation of Traffic Tables

95% Confidence Interval

Table 1 includes a “95% confidence interval” for each rate, percentage, or ratio. Here, the 95% confidence interval represents that part of uncertainty (not the whole of uncertainty) in estimating the rate, percentage or ratio, which is due to sampling variability. Most generally, the 95% confidence interval provides a range of plausible values. The “95%” figure can roughly be interpreted in the following way: if the year could be repeated many times with nothing essentially changing in the real world (in traffic patterns and behavior of drivers and officers) and with methods of analysis kept the same, 95% of the time the repeated result would be expected to be found inside that given interval. Because there is an element of chance involved in being stopped, being searched, etc., the value of a rate or percentage or ratio would be somewhat different with each repetition. The 95% confidence interval reflects that particular aspect of the overall problem, the ever-present play of chance. It uses widely accepted statistical methods to express that uncertainty in the estimated rate, percentage or ratio. There is another important source of uncertainty that reflects the methodology used and, most importantly, the limitations of the datasets that inform the analysis. A confidence interval does not cover that aspect of the problem. Yet, it is useful to be aware of that part of uncertainty that is due to simple play of chance, which gets more and more prominent as agencies become smaller, have fewer stops or smaller benchmark values that are used for calculating rates, percentages or ratios.

Ratios

A ratio of rates or percentages with a value of 1.0 indicates that the rates or percentages are equal between the Minority group and Whites. Ratios above or below 1.0 show greater or lesser stop activity with minorities, respectively. Comparisons of Minority groups to White drivers or White pedestrians where the 95% confidence interval lies above 1.0 are **bolded** in the stops tables. One can say that the value of 1.0 does not fall within the 95% confidence interval of the estimated ratio. These **bolded** ratios are statistical deviations and may be the basis for further consideration of potential racial disparities related to stops. A **bolded** ratio does not prove that there is racial profiling but may be taken as the basis for further inquiry. In addition to whether or not a ratio is **bolded**, the absolute magnitude of the ratio should be considered. For example, a **bolded** ratio of 5.0 is a higher priority to investigate than a small, bolded ratio of 1.2. A larger ratio implies that the potential impact on individuals is larger, and it is less likely that the elevated ratio is only due to limitations of the chosen benchmark than when the ratio is closer to 1.0.

Limitations

There is a limitation in the use of ratios to determine potential racial disparities. As explained, the 95% confidence intervals for stop rates and stop rate ratios do not involve the systematic error in estimating the driver and pedestrian benchmarks. Note that each law enforcement agency has a “jurisdiction,” which is the geographic area that the agency is responsible for policing. In this report “agency” and “jurisdiction” are sometimes used interchangeably.

The benchmarks attempt to estimate the driving population or the total mileage driven within the jurisdiction of each agency using a combination of data sources, including surveys by the U.S. Census Bureau, Illinois crash reports (collected by IDOT), Illinois driver’s license counts (provided by the

Secretary of State's office), tracking of cellular phones of drivers who opted-in, employment data, traffic counts and other sources. However elaborate, these datasets can only approximate local traffic dynamics and necessarily rely on particular assumptions, which may not always be accurate. Thus, all benchmarks will always come with some uncertainty, and the extent of this uncertainty is unknown inasmuch the underlying "truth" is unknown. The problem of accurately modeling the dynamic population of drivers simultaneously across hundreds of police agencies, which are vastly different in their sizes and in the levels of locally relevant detail, is a problem of enormous complexity in principle. At the same time, the statistics that can be estimated using the current level of modeling are relevant numerical relations. Their careful interpretation may prove useful to many agencies in various ways, as long as the expectations of what these statistics can offer to particular agencies are qualified by the awareness that some limitations and the need to rely on assumptions are unavoidable facts of life.

Another limitation that may affect the rates, percentages and ratios is the designation of race by the law enforcement officer conducting the stop. That designation of race might not correspond to the driver's or pedestrian's own racial identity.

Also, the stop rate for a racial group will depend on a) the assignment of beats (geographic surveillance area) to officers in a jurisdiction and b) the degree of association of those beats with the driving patterns of each racial group. If there is higher (or lower) surveillance of an area with a high driving concentration of a racial group, then that can lead to a higher (or lower) stop rate for the racial group, compared to areas where surveillance is constant across all racial groups.

Studies of this kind do not collect information on what happened BEFORE a stop occurred. They only see that a stop did occur and that it did report some categorized outcomes. Therefore, they cannot directly "calculate" if a larger than expected number of stops of drivers of a particular group was justified or involves profiling. Under such constraint, these statistical findings call for a careful and thoughtful interpretation and perhaps a communal dialogue, as a step following review of this report.

Statistics based on stops only

Some percentages and ratios of percentages in the tables are based on stop counts and stop activity only, and these statistics do not use benchmarks in the calculations. These percentages and ratios are given in Panels 2-6 of the statistical tables. They do not have the potential benchmarking errors noted above.

It is important to note that these percentages are calculated with reference to a specific activity. For example, in the traffic tables, the percentage of searches for a racial group is a percentage of *stops* leading to a search. The percentage of contraband found in a vehicle is the percentage of *vehicle searches* leading to contraband found. For percentages, each row label (or the heading for the panel) indicates the basis for the percentage.

Can stop rates be compared across years?

The methodology used for calculating traffic stop rates in this study, using a population benchmark, differs from studies of stops in 2019-2020 and, for a different range of years, for stops in 2018 and earlier. However, this is the fourth report in which the methodology is essentially the same since the 2021 stops report. See Section V below for specific details on the benchmarks. While the newest methodology provides more accurate estimates of the racial composition of the driving population, the

changes impact comparisons of results from the 2021-2024 stops analysis to the analyses in 2019-2020 and in years prior to 2019. Comparisons of 2024 to 2019-2020 are not perfect but are more acceptable than comparisons of 2024 to 2004-2018. The table formats are very similar when comparing 2024 to 2019-2020, even though there are some underlying methodological differences.

These and other changes have improved the estimate of the benchmark populations and the accuracy of stop rate ratios. Thus, any difference in rate ratios between 2021-2024 stops reports and earlier stops reports (2019-2020 and 2004-2018) may be at least partly due to the improvement in statistical methods used in this report rather than a real change in stop rates. The newest method is intended to estimate the population benchmark more accurately. Another factor making it difficult to compare 2024 stop rates to 2018 rates (and earlier) is that the 2024 report presents rates, percentages and rate ratios separately for each of the six individual races — rather than with all minorities combined into one category, as used in the 2018 and earlier reports. Perusal of tables in Part II of this report will show the reader that the five Minority races do have different stop rates, which can largely distort the interpretation of their rates if combined. The statewide rates in Table 1, Panel 1, above, show a diversity of stop rates among the six races as well as among the five Minority races. The 2019-2020 reports also reported results separately for each individual race, making comparisons of 2019-2020 to 2021-2023 more straightforward.

Certain percentages will be comparable across years, because the percentages are based on stops data only without involving a benchmark, and percentages are calculated in the same manner as in previous years. However, to compare a percentage based on 2024 stops data to a percentage reported in a year prior to 2019, some additional calculations and reference to the original datasets will be needed. To calculate a percentage for 2024 stops of all minorities combined (not recommended), the user will need to add together, across the five Minority racial groups, all of the numerators and, separately, all of the denominators and then divide the numerator sum by the denominator sum, then multiply by 100% to get the all-Minority percentages.

In addition to combining the minorities together, prior to 2018 the stop rate ratios were not calculated relative to White drivers, but relative to all drivers. Such a measure (then called the “Ratio”, also sometimes known as “Disparity Index”) is inherently less sensitive to changes, numerically always closer to (biased toward) the value 1.0 than the current rate ratios are, and this measure depends on the relative sizes of racial groups compared to each other. This is all due to fact that this measure involves the same drivers in both its numerator (being drivers of a particular race) and denominator (being among all drivers). This older measure is methodologically (and epidemiologically) not used in modern statistical practice and should not be compared to the current rate ratios in any way.

V. Benchmarks

The number of stops for each racial group and each agency is compared to a “benchmark” in order to calculate the agency’s stop rate for the racial group.

In the abstract mathematical sense, benchmarks can be seen as any set of numerical values, one number for each racial group, that are proportional to each other exactly as counts of police stops would be proportional to each other in a mathematically ideal setting. We see this idealized setting as a

world where stops are perfectly unbiased to the race of drivers that are stopped and where drivers behave perfectly independently of their race. Benchmarks then reflect this mathematical ideal. That is, if stop counts were proportional to their benchmark numbers, all rates would be equal and all rate ratios would equal 1. Then, we can look at the actual stop counts (presented in the tables) and observe and measure how different these counts are from the idealized case.

Apart from being proportional to each other, which is of primary importance, the actual benchmark values can carry additional meanings. For example, they may represent counts of drivers, as far as the ideal counts of stops and those population counts can be expected to have identical racial distributions.

In all previous reports since 2004, benchmarks were synonymous with driver population counts. This year we expand this view by introducing benchmarks that have a different added meaning: the total distance that drivers of a particular race group have driven collectively within a jurisdiction and during an average day. In an idealized world, these distances and the traffic stops would be expected to be equally distributed by race. Here we expect that, in the ideal world, if Asian drivers accounted for 20% of all miles driven, they would also have accounted for 20% of idealized traffic stops.

It should be emphasized that the new benchmark model is not meant to replace the old, but to serve as a meaningful alternative, since it is informed by different and potentially very rich data sources. In this report, the new model is mentioned in this section V and in the appendix D. Here, the new model is introduced and we make some initial comparison with the old methodology. The new benchmark values and their rates and rate ratios are also shown in the individual tables in part II. However, in this report the new benchmarks are not used for any other analysis or summary, nor are they used for selected findings nor key findings. A reader who wishes to ignore the new distance-based benchmarks and the rates and rate ratios based on them can readily do so.

Population Benchmarks

The first benchmark model is the same one used since the report on 2021 stops. That benchmark methodology is now distinguished by naming it “Population Benchmarks”. Units of this model are individuals who drive at least once within a given jurisdiction during the year and can potentially be stopped in that jurisdiction by a law enforcement officer of that agency. By this method, benchmarks are seen as the counts of such drivers, not only the actual resident drivers of the jurisdiction but also drivers living in neighboring areas, which may be close or distant, and who may pass through the jurisdiction. For agencies that want to look into time trends in their data, this population benchmark provides continuity with results in our previous reports for 2021 to 2023 stops.

The population benchmark provides an estimated population count for each of the six racial groups. These population counts are then compared to the traffic stop counts of each racial group to assess and compare the stop rates (stops per unit of population) of each racial group. See Appendix C of our previous year’s report on 2023 stops, Technical Notes on Benchmarks, for a detailed discussion of population benchmarks and associated calculations, including important limitations.

The methods for calculating the population benchmark for each agency for this report on 2024 stops are essentially the methods used in reports on 2021-2023 stops. Briefly, traffic stop population benchmarks are based on the U.S. Census Bureau’s most recent American Community Survey population statistics tabulated at the ZIP code level. For agencies with a sufficient number of crash reports available in their

jurisdiction, the Illinois traffic crash report data (based on 2019-2023 SR 1050 crash reports¹) were used to build the traffic stop crash-based population benchmarks. For the other agencies (without sufficient crash reports), the traffic stop population benchmarks were constructed by combining ZIP code data from the surrounding area, weighted by the distance from the agency's jurisdiction (distance-based population benchmarks). Both types of population benchmarks (crash-based and distance-based) combine populations from ZIP codes directly associated with an agency (e.g., the ZIP codes of a city for a city police agency) as well as populations from ZIP codes from the surrounding area. Also, note that the traffic stop and pedestrian stop population benchmark methodologies differ because of the different data sources available to generate them. Thus, it is not unusual for there to be notable differences between the traffic and pedestrian population benchmarks for the same agency.

Mileage Benchmarks

The second model is based on the estimated number of miles driven or walked within a jurisdiction. The meaning of this benchmark is, precisely: the total number of miles driven (or walked) by all individuals of a particular race group, within the boundaries of a jurisdiction during a single average day. In this view, a stop can happen within each driven (or walked) mile, and only some of these miles experienced a stop. Thus, the stop rate here is count of stops divided by the count of miles driven (or walked) within a jurisdiction for a particular race group. We refer to this benchmark as “Mileage Benchmark”, as its numerical values estimate counts of miles of driving (walking) within a jurisdiction. An individual driving or walking in a jurisdiction has the possibility of being stopped by a law enforcement officer of the agency and, thus, the individual does contribute, in proportion to the distance traversed, to the agency's traffic or pedestrian benchmark, respectively.

These benchmarks are constructed from the various outputs of the “Replica” data platform — a massive computer simulation of traffic dynamics at a nationwide level with high granularity of geolocation detail, which is being maintained and updated using near real-time data originating from extensive sources. Replica is used by various state and local government agencies, including IDOT. In this aspect, the project of investigating potential bias in traffic and pedestrian police stops becomes integrated with the software infrastructure already used by IDOT.

For traffic stops, the mileage benchmark provides an estimated cumulative “mileage” (total distance driven, expressed in miles) for each of the six racial groups. These mileage counts are then compared to the traffic stop counts of each racial group to assess and compare the stop rates (stops per unit of distance driven) of each racial group. See Appendix D in this document for more information on the methodologies, the advantages and the limitations of the mileage benchmark model.

The methods for calculating the mileage benchmark for each agency for this report are as follows. Briefly, traffic stop mileage benchmarks are based on five broad data categories: mobile location data (tracking cell phones of opted-in individuals, cellular network data, vehicle GPS data, etc.), consumer and residential data (five-year ACS population statistics, Public Use Microdata, employment data, etc.), built environment data (street maps, building footprints, points of interest data like restaurants and theaters, etc.), economic activity data (financial transactions data), and “ground truth” or various observed counts (car/traffic counts, bike and pedestrian counts, transit ridership counts, etc.) All these

¹ <https://idot.illinois.gov/content/dam/soi/en/web/idot/documents/transportation-system/manuals-guides-and-handbooks/safety/illinois-traffic-crash-report-sr-1050-instruction-manual-2019.pdf> (last accessed June 13th, 2024).

sources are used to create a virtual population of individuals in their daily commute that must conform to the facts seen in all the datasets. This richness and variety of data sources combined with extensive computational capabilities holds a great potential, and that is the primary reason that prompted this team to use the Replica platform for creating an alternative set of benchmarks. The model simulates individual trips taken during a typical weekday or a typical weekend day. Each trip is fully geolocated and partitioned into short street segments. The mileage benchmark is calculated as the sum of the lengths of all the segments that fall within a given jurisdiction and belonging to all recorded daily trips made by drivers of a particular race. Thus, a driver who covers more miles within a jurisdiction will influence the benchmark more than a driver who just passes through.

Initial comparison of the two benchmark models

As already mentioned, the introduction of the mileage benchmarks is mainly to provide an alternative, well informed and largely independent description of the same process of local traffic dynamics. Both models start with the same demographic information on where potential drivers reside (ACS, PUMS data), and then proceed to make the same statistical estimates. Namely, population benchmarks are based on crashes that supposedly reveal drivers at random: by being a not-at-fault victim of a crash, a driver's presence is revealed at a random location and at a random hour. With each new mile driven, a driver can possibly experience a crash, a traffic stop or both. A driver that covers more miles within a jurisdiction will be that much more likely to be involved in a crash or a traffic stop. Therefore, by their essential setup and assumptions, crash-based population benchmarks and mileage-based benchmarks should capture the same probabilities of being stopped, and thus should share the same distribution by race; that is, provide the same stop rate ratios. However, the Replica model is not informed by crash reports, and the population benchmarks do not have access to the multitude of the data sources that Replica uses. Both methodologies attempt to reach the same endpoint by following different paths, and both methods have considerable strengths.

Observing the rate ratios across hundreds of individual tables in Part II of this report, it can be seen that mileage-based rate ratios are typically further away from 1.0 than population-based rate ratios and typically larger. Crash-based benchmarks do possess one serious, known limitation: crash reports do not record the driver's race, so the actual race is always unknown. As a necessary approximation, we use the driver's ZIP code of residence and assume that the driver is a typical racial representative of his/her ZIP code demographics: if, for example, half of the drivers in that ZIP code are White and a third are Black, the driver increases the benchmark of the White race group by half of a person and the benchmark of the Black group by a third of a person. The question is then how much is this approximation affecting the rate ratios, and is this at least partially the reason why the population-based rate ratios differ from mileage-based rate ratios?

To estimate this, we used the fact that the Replica model also provides the ZIP code of residence of each virtual driver. We then constructed another set of Replica-based benchmarks by repeating the approximation necessary for the crashes: the actual race of a virtual driver is deemed unknown, and the driver is seen as a typical racial representative of his/her zip code demographics.

We compared the three benchmark models (crash-based, Replica-based with driver race known, Replica-based with driver race unknown) for 20 specially selected agencies. These are a few of the largest cities in Illinois (avoiding Chicago as a whole), several suburbs of Chicago and four Chicago police districts. They are selected for being spread out across Illinois and being relatively large — because: a)

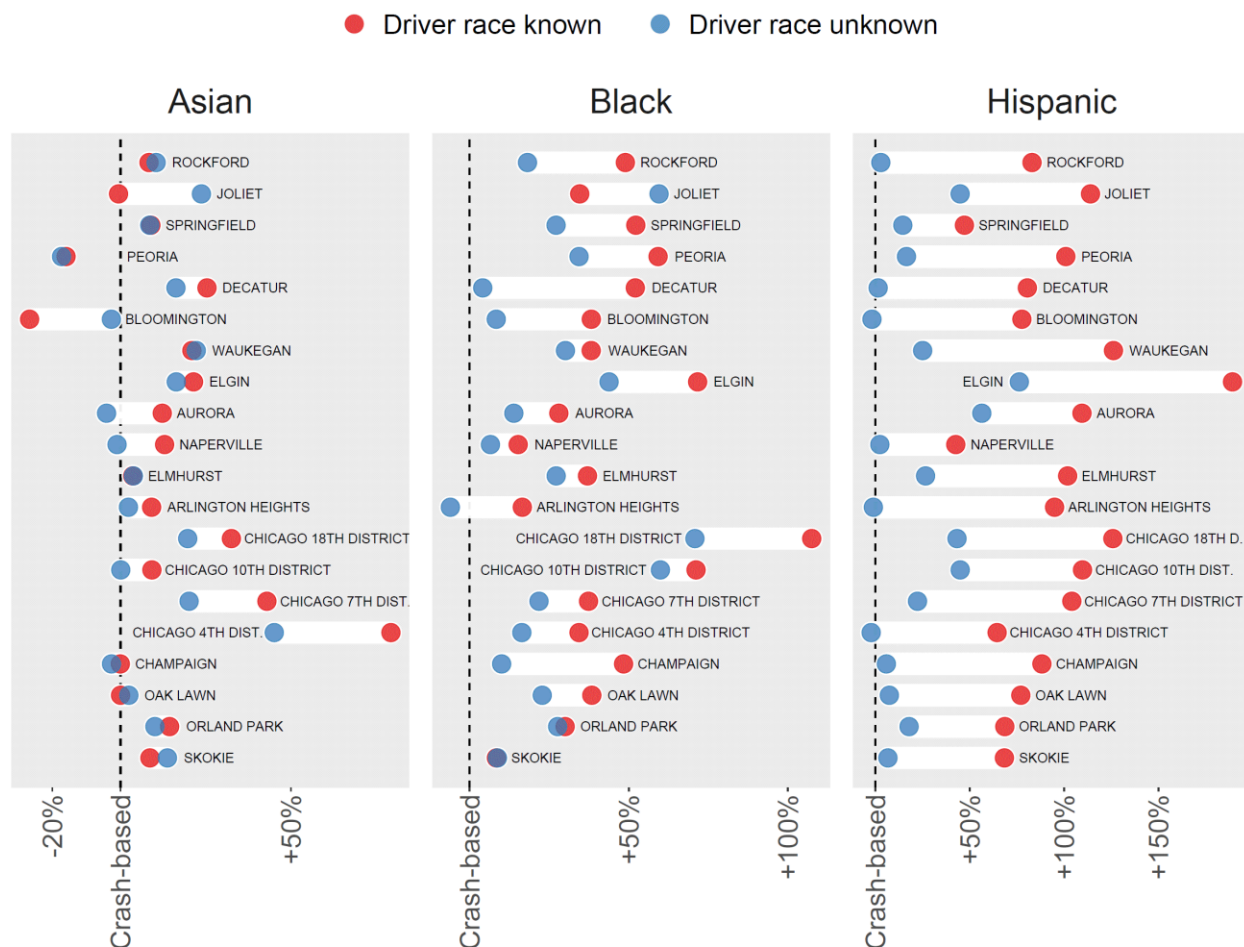
the boundaries of their jurisdictions can be well approximated as collections of ZIP code areas and b) they have a large number of recorded crashes. Also, we only looked at rate ratio of the largest minority groups (Black/Hispanic/Asian). In this way, in selecting large agencies, the analysis tried to isolate the effect on rate ratios of driver race being known or not known. The analysis thus avoided the effect of details that can also influence rate ratios: jurisdictions not aligning well with their ZIP codes, agencies not having enough crashes, small counts of rare minorities, etc.

The results are shown in Figure 1. Here, for each race group, agencies are vertically stacked (annotated by their names), and their Replica-based rate ratios (red and blue dots) are shown on the x-axis relative to their crash-based rate ratios (vertical dashed line). Here, “+50%” on the horizontal axis means that a Replica-based rate ratio is 50% higher than the crash-based rate ratio for the same agency and minority group. To help guide the eye, two dots of the same agency are connected by a white-colored strip. Across the three panels, each agency is always placed at the same vertical position.

In nearly all cases, not taking account of the drivers’ actual race makes the Replica-based rate ratio closer to the crash-based rate ratio. In the Asian group, the average (across 20 agencies) absolute distance from Replica-based to crashed-based rate ratio is $\pm 18\%$ when driven race is known, and $\pm 12\%$ when it is unknown. In the Black group, these values are $\pm 43\%$ and $\pm 26\%$, respectively. In the Asian group, they are $\pm 94\%$ and $\pm 12\%$, respectively. The fact that the two benchmark models (crash-based and Replica-based) agree considerably better when the same approximation is introduced in both (driver race unknown), testifies to the large consistency between them. Indeed, following very different data paths the two methods do produce a similar endpoint. If rate ratios indeed differed by only 12% to 26% on average across agencies and minority groups, we could conclude that the two models, for most of the practical purposes, give fairly similar results. For example, if in one model a rate ratio is 2.4, in the other model it could be 3.0 (26% larger). Also, we can see how far and in what direction rate ratios change when the approximation is removed, at least in the Replica model. Rate ratios being larger when a driver’s race is known is consistent with the speculative (and unproven) interpretation that drivers tend to drive more within localities where residents of their own race are more frequent. In any case, the approximation made when a driver’s race is not known assumes that any two drivers of different races who reside within the same ZIP code A will spend the same time driving inside any other ZIP code B.

This was an initial analysis, because the new Replica-based benchmark model is introduced only this year (for analysis of 2024 stops). Going forward, the team will continue investigating the contrasts between crash-based and Replica-based benchmarks, for use in future reports, and make improvements.

Figure 1. Comparison of rate ratios for the three benchmark models. Twenty selected agencies are vertically stacked (height is arbitrary) and annotated by their names. Red and blue dots are two different Replica-based rate ratios given relative to (divided by) their crash-based versions, calculated when driver race is known (red dots) or unknown (blue dots). White stripes serve as eye guides only, connecting two rate ratios of a same agency.



VI. Selected Findings

This section of the report shows some tables and figures that present results on the agencies and their stops from the entire state for 2024. Some results are contrasted with their corresponding 2022 and/or 2023 values. All benchmark-related results in this section are solely based on population benchmarks.

Agency reporting status

Among the 978 agencies that were active at the end of 2024 and could submit stop data to IDOT, 78.9% of the agencies had stops and provided complete stop data to IDOT (Table 2, top numeric row), which is similar to 2023. Further descriptive statistics on the agencies' reporting on 2024 stops are:

- 1.5% of agencies had no traffic stops, a decrease from 3.7% in 2023.
- 1.4% of agencies collected stop data for less than a year ("incomplete").
- 18.1% of agencies had stops but did not submit any stop data ("Non-compliant"), which is an increase from 15.8% in 2023 and similar to 18.7% in 2022.

Table 2. Agency status on reporting. Illinois, all agencies, Traffic stops, 2023 and 2024.

Status of Agency	2023		2024	
	Number of agencies	Percent of agencies	Number of agencies	Percent of agencies
Complete reporting ^a	779	78.1%	772	78.9%
Zero stops ^b	37	3.7%	15	1.5%
Incomplete ^c	23	2.3%	14	1.4%
Non-compliant ^d	158	15.8%	177	18.1%
All agencies combined	997	100%	978	100%
^a Agency with one or more stops that were completely reported. ^b Agency performed no stops over the year. ^c Agency submitted some but not all of their stops for the year. ^d Agency made stops, but no stops data was submitted.				

Number of stops

The total number of reported traffic stops in 2024 was 2,050,405. The number of stops per agency was generally substantial. Hundreds of agencies (about 80%) had more than 100 stops during 2024 (Table 3). In 2024, Chicago Police reported a notable 45% decrease in the number of stops from 2023.

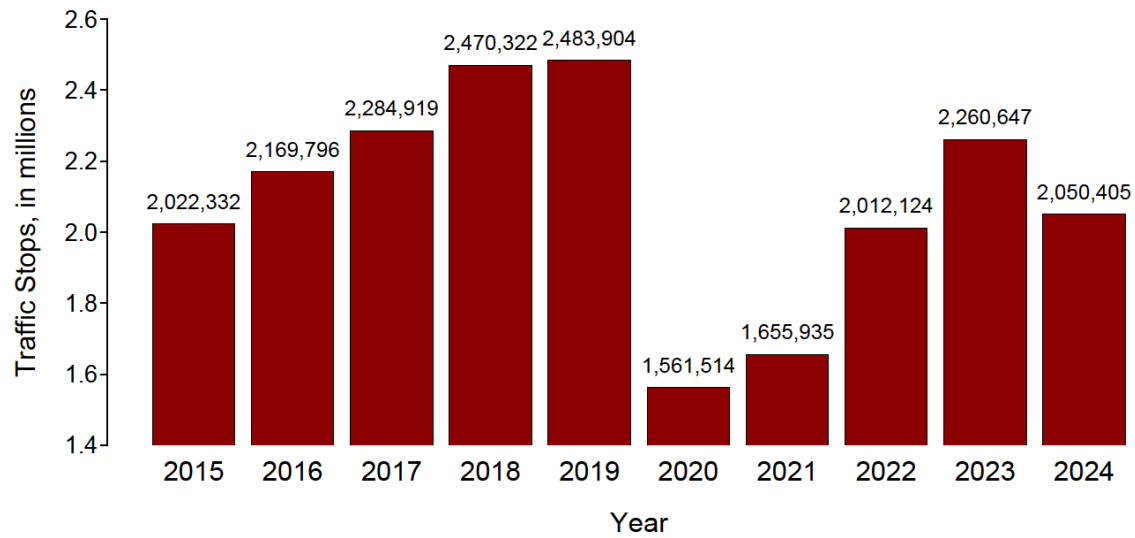
Table 3. Number of Traffic stops for agencies with at least one stop. Illinois, all agencies, Traffic stops, 2023 and 2024.

Number of stops	2023		2024	
	Number of agencies	Percent of agencies	Number of agencies	Percent of agencies
1-10	59	7.4	50	6.3
11-100	134	16.7	104	13.2
101-1,000	288	35.9	294	37.3
1,001-10,000	301	37.5	316	40.1
10,001-100,000	18	2.2	22	2.8
More than 100,000	2	0.2	2	0.3
All compliant agencies with ≥ 1 stop	802	100%	788	100%
Notes: (1) Includes only agencies with at least one stop and includes all agencies with either complete or incomplete reporting of 2024 stops. (2) Chicago Police: 535,088 in 2023; 293,150 in 2024. (Chicago is also represented in the Table above).				

In 2024 there were no reported stops with missing information about the race of the driver. In 2023 there were 78 such stops.

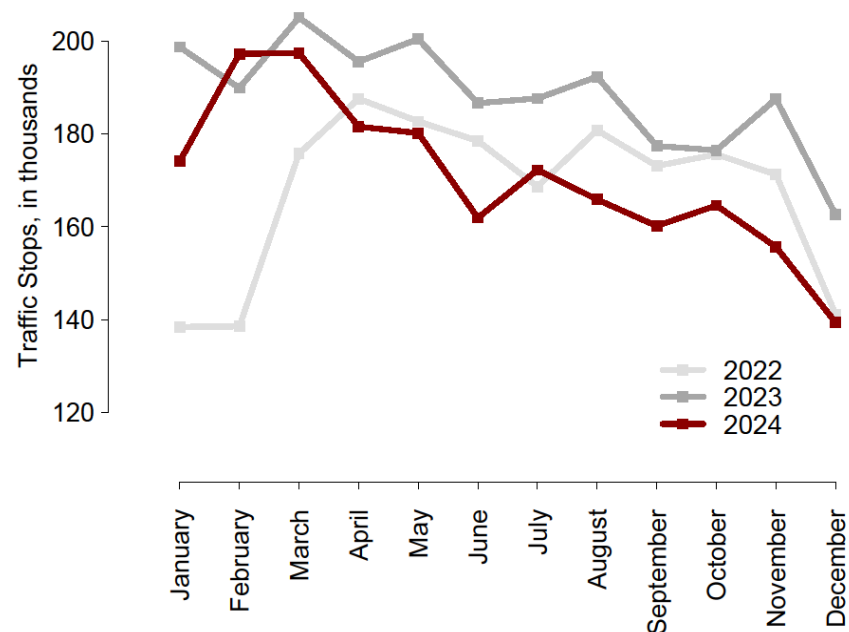
The number of reported stops per year has steadily grown 2015-2019 (Figure 2a). There was a 23% increase in the number of stops during this period. Due to COVID-19, in 2020 the number of reported stops decreased by 37% from 2019. During 2020-2023 this number steadily increased again, in 2023 by 45% from 2020. In 2024, the number decreased 9.3% from 2023 and approximately returned to its 2015 value.

Figure 2a. Illinois, number of traffic stops, 2015-2024.



The monthly pattern of stops shows that in 2024 the number of stops was steadily declining after March (Figure 2b). Until March, the number of stops in 2024 was comparable to the same period in 2023. After March, stops in 2024 were more comparable to the same period in 2022. In the beginning of 2022, the number of stops was still notably reduced due to COVID-19.

Figure 2b. Illinois, number of traffic stops per month, 2022 (light gray line), 2023 (gray line), and 2024 (dark red line).

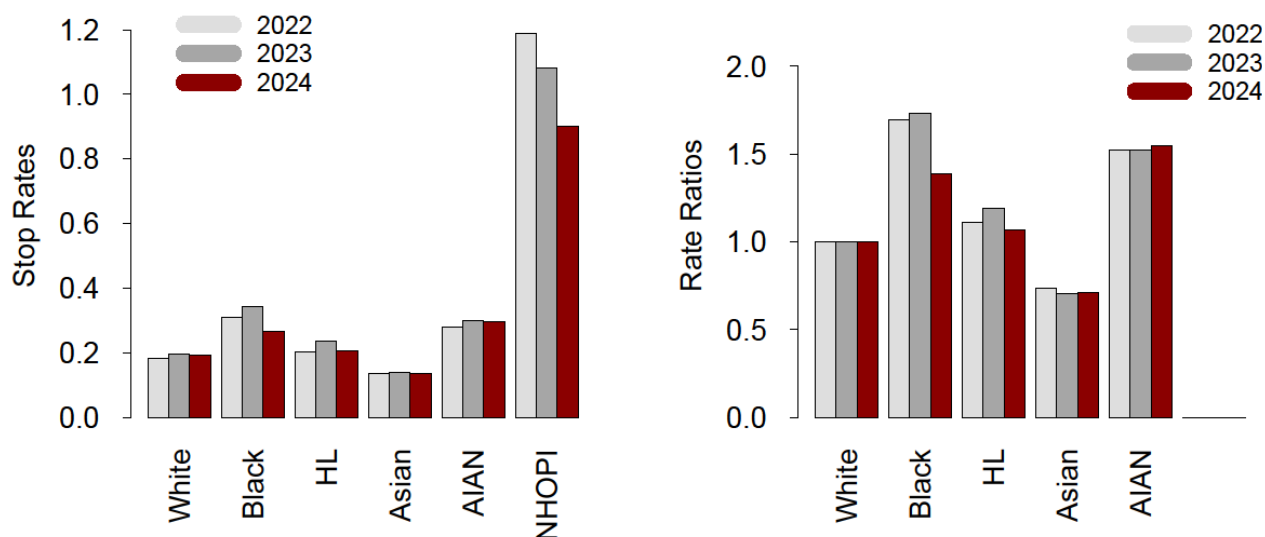


Statewide stop rates and rate ratios

The statewide stop rates are diverse among the six racial groups (Figure 3, left panel). Of interest, the smallest Minority group (Native Hawaiian or Other Pacific Islander) had the highest stop rate. Being by far the smallest minority, it is the hardest group to have any statistics accurately estimated, even at the statewide level. In addition, its stop rate may potentially be an anomaly due to a mismatch between the officer-identified race of stopped individuals and the self-identified race reported in the U.S. census survey data, which is used as part of the population benchmark calculations in this study. Stops rates of all racial groups have slightly decreased from 2023, likely reflecting the decreased number of stops in 2024.

The statewide stop rate ratios seem fairly constant within the last three years, 2022-2024. Asian drivers have their rate ratio less than 1, Black and AIAN drivers larger than 1, and Hispanic or Latino drivers are close to 1. Rate ratio for NHOPI group is not shown in the figure, being too high to show on this scale. In 2024, Black drivers have a notably smaller rate ratio (1.39), compared to 2023 (1.73).

Figure 3. Stop rates (left panel) and rate ratios (right panel) for each racial group, 2022 (light gray bars), 2023 (gray bars), and 2024 (dark red bars). Illinois, Traffic stops, 2022-2024.



Abbreviations for racial groups: Black = “Black or African American”, HL = “Hispanic or Latino”, AIAN = “American Indian or Alaska Native”, NHOPI = “Native Hawaiian or Other Pacific Islander”.

Distribution of stop rate ratios

Table 4.a shows the numbers of comparisons of stop rates of a Minority racial group and Whites carried out in the traffic stops study. Any comparison yields a rate ratio — the Minority stop rate divided by the White stop rate. Each agency might contribute up to five such comparisons (five Minority groups, each compared to Whites on their stop rates). For this analysis there were fewer than five comparisons when

White drivers had zero stops or when a population benchmark value was zero for either a Minority racial group or Whites, thus making some comparison rate ratios numerically undefined.

The first column under “A” in Table 4.a illustrates all comparisons: each Minority/White rate ratio from each agency has been compiled across all agencies. Table 4.a then categorizes the rate ratios by their magnitude and shows the percentage distribution across categories. The columns under “B” restricts the comparisons to agencies with stop rates based on at least 50 White stops and 50 stops of the Minority group being compared to Whites. The 50 stops would provide a more precise rate ratio than a smaller number of stops. The large percentage of stops in the category “<0.25” in panel A for both 2022 and 2023 is due to the presence of many small agencies that have a small number of stops and zero stops for one or more minorities.

We note a drastic reduction — nearly 4-fold from Panel A to Panel B — in the total number of rate ratios, from 3,920 (all comparisons) down to 1,001 (more precise comparisons). From the more precise comparisons (Panel B, based on 50 or more stops of Whites and 50 or more stops of the Minority group compared) we estimate that in 78.0% of these rate ratios, Minority drivers were stopped at a higher rate than White drivers (rate ratio > 1). This suggests, as a possibility but does not prove, that racial profiling was a factor in a number of traffic stops throughout Illinois.

The overall distribution of rate ratios seems rather similar in 2023 and 2024. The 95% confidence intervals provided in the tables of Part II should be used as a guide to the precision of rates, percentages and rate ratios when interpreting the numeric results for a specific agency.

Table 4.a Distribution of stop rate ratios. (Each Non-White racial group compared to Whites for an agency.) Illinois, Traffic stops, 2023 and 2024.

A. All agencies and racial groups*			B. Agencies and racial groups with at least 50 stops**	
Stop rate ratios	2023	2024	2023	2024
<0.25	32.8%	30.3%	0.9%	0.9%
0.25 to <0.5	8.1%	7.4%	4.1%	4.6%
0.5 to <1.0	14.4%	13.9%	16.7%	16.5%
1.0 to <2.0	18.1%	19.9%	35.2%	34.6%
2.0 to <4.0	15.7%	16.0%	33.0%	31.9%
≥4.0	10.9%	12.5%	10.1%	11.6%
All ratios***	100%	100%	100%	100%
* All comparisons of Whites and a racial group for all agencies. Excludes ratios from agencies with zero stops of White drivers or a population benchmark value of zero for either a Minority racial group or Whites. ** All comparisons of Whites and a racial group for all agencies; all comparisons must have at least 50 stops of Whites and 50 stops of the compared racial group. Excludes undefined rate ratios, or where either Whites or the compared racial group have less than 50 stops. ***The number of ratios that were included in the analysis in column A and B respectively, were 3,999 and 958 in 2023; 3,920 and 1,001 in 2024. Each ratio involves a comparison of one non-White racial group vs. Whites for one agency.				

Table 4.b shows the distribution of stop rate ratios in 2024 among the three most populous Minority groups. Since each agency provides only a single stop rate ratio for a single Minority group, here, a proportion of stop ratios equates to a proportion of agencies. From the more precise comparisons (Panel B, based on agencies with more stops) we estimate that in 93.3% of agencies with at least 50 stops for both Whites and Blacks, Black drivers are stopped at a higher rate than White drivers (rate ratio > 1). For Hispanic drivers, this value is 79.2%. Similar to the note on Table 4.a, this suggests as a possibility, but does not prove, that racial profiling was a factor in a number of traffic stops. This finding does not occur among stopped Asian drivers, who are stopped at a higher rate than White drivers in only 35% of agencies with at least 50 stops for both Whites and Asians.

Table 4.b Distribution of stop rate ratios for Black, Hispanic and Asian drivers. (Each noted non-White racial group compared to Whites for an agency). Illinois, Traffic stops, 2024.

A. All agencies and racial groups				B. Agencies and racial groups with at least 50 stops*		
Stop rate ratios	Black	Hispanic	Asian	Black	Hispanic	Asian
<0.25	8.0%	15.4%	35.1%	0	1.7%	1.7%
0.25 to <0.5	4.5%	5.9%	19.4%	0.7%	3.3%	17.2%
0.5 to <1.0	12.0%	18.1%	24.4%	5.9%	15.8%	46.1%
1.0 to <2.0	26.1%	36.6%	13.8%	24.7%	50.8%	30.0%
2.0 to <4.0	35.8%	19.3%	4.3%	50.8%	25.6%	2.8%
≥4.0	13.5%	4.7%	3.1%	17.8%	2.8%	2.2%
All ratios**	100%	100%	100%	100%	100%	100%
*All comparisons of Whites and a racial group for all agencies; all comparisons must have at least 50 stops of Whites and 50 stops of the compared racial group. Excludes undefined rate ratios, or where either Whites or the compared racial group have less than 50 stops. **The number of ratios that were included in the analysis in column A was 784 for Black, 784 for Hispanic, and 784 for Asian group; in column B this was 421 for Black, 360 for Hispanic, and 180 for Asian group. Each ratio involves a comparison of one non-White racial group vs. Whites for one agency.						

Table 4.c shows the distribution of citation ratios among the three Minority groups, and all the racial groups collectively in 2024. A citation is the most severe outcome among the three outcomes noted on the data collection form: verbal warning, written warning and citation. Here we estimate that in 71.8% of all agencies with at least 50 stops for both Whites and Blacks, Black drivers are getting citations at a higher rate than White drivers (citation ratio > 1). For Hispanic drivers, this value is 86.6%. Similar to the note on Table 4.a, this suggests as a possibility, but does not prove, that racial profiling was a factor in a number of citations. This finding does not occur among Asian drivers, whose citation rate is higher than among White drivers in 54.2% of all agencies with at least 50 stops for both Whites and Asians. Overall, in 72.4% of all citation ratios Minority drivers are receiving citations at a higher rate than White drivers.

Table 4.c Distribution of citation ratios. (Each ratio that enters into the computation involves each noted non-White racial group compared to Whites for an agency.) Illinois, Traffic stops, 2024.

Citation rate ratios*	Black	Hispanic	Asian	All racial groups
<0.25	0	0.3%	2.2%	0.5%
0.25 to <0.5	0	0	1.7%	0.7%
0.5 to <1.0	28.2%	13.1%	41.9%	26.4%
1.0 to <2.0	69.9%	82.1%	54.2%	69.9%
2.0 to <4.0	1.9%	3.6%	0	2.2%
≥4.0	0	0.8%	0	0.3%
All ratios**	100%	100%	100%	100%
*All comparisons of Whites and a racial group for all agencies; all comparisons must have at least 50 stops of Whites and 50 stops of the compared racial group. Excludes undefined ratios, or ratios where either Whites or the compared racial group have less than 50 stops. **The number of ratios that were included in the analysis for 2024 stops is 996. Each ratio that enters into the computation involves a comparison of one non-White racial group to Whites for one agency.				

Table 4.d shows the distribution of contraband-found ratios in vehicle searches among the three more populous Minority groups, and all the racial groups collectively in 2024. Here we estimate that in 63.2% of all agencies with at least 50 stops for both Whites and Blacks, contraband is found in Black drivers' vehicle searches at a higher rate than in White drivers (ratio > 1). For Hispanic drivers, this value is 36.5%, for Asian drivers it is 25.4%, and the overall percentage for all racial groups is 46.9%. This result does not suggest a presence of racial profiling related to the contraband aspect of traffic stops.

Table 4.d *Distribution of contraband-found ratios in vehicle searches.* (Each ratio that enters into the computation involves each noted non-White racial group compared to Whites for an agency.) Illinois, Traffic stops, 2024.

Contraband rate ratios*	Black	Hispanic	Asian	All racial groups
<0.25	7.6%	8.8%	47.4%	14.9%
0.25 to <0.5	3.3%	7.7%	2.6%	5.0%
0.5 to <1.0	25.9%	47.0%	24.6%	33.2%
1.0 to <2.0	56.2%	31.2%	21.1%	40.9%
2.0 to <4.0	6.8%	4.9%	2.6%	5.4%
≥4.0	0.3%	0.4%	1.8%	0.7%
All ratios**	100%	100%	100%	100%
*All comparisons of Whites and a racial group for all agencies; all comparisons must have at least 50 stops of Whites and 50 stops of the compared racial group. Excludes undefined ratios, or ratios where either Whites or the compared racial group have less than 50 stops. **The number of ratios that were included in the analysis for 2024 stops is 765. Each ratio that enters into the computation involves a comparison of one non-White racial group to Whites for one agency.				

Table 4.e shows the distribution of contraband-found ratios in searches of individual drivers or passengers among three Minority groups individually, and all the racial groups collectively in 2024. Here we estimate that in 42.5% of all agencies with at least 50 stops for both Whites and Blacks, contraband is found while searching Black drivers or their passengers at a higher rate than in White drivers or their passengers (ratio > 1). For Hispanic drivers or their passengers, this number is 25.8%, for Asian drivers it is 14.6%, and the overall percentage for all racial groups is 32.0%. This result does not suggest a presence of racial profiling related to this aspect of traffic stops.

Table 4.e Distribution of contraband-found ratios from searches of individuals: driver or passengers.
(Each ratio that enters into the computation involves each noted non-White racial group compared to Whites for an agency.) Illinois, Traffic stops, 2024.

Contraband rate ratios*	Black	Hispanic	Asian	All racial groups
<0.25	26.3%	35.7%	81.7%	38.6%
0.25 to <0.5	7.3%	13.1%	2.4%	8.6%
0.5 to <1.0	23.9%	25.4%	1.2%	20.8%
1.0 to <2.0	28.2%	14.6%	8.5%	19.6%
2.0 to <4.0	12.4%	8.0%	2.4%	9.3%
≥4.0	1.9%	3.3%	3.7%	3.1%
All ratios**	100%	100%	100%	100%
*All comparisons of Whites and a racial group for all agencies; all comparisons must have at least 50 stops of Whites and 50 stops of the compared racial group. Excludes undefined ratios, or ratios where either Whites or the compared racial group have less than 50 stops. **The number of ratios that were included in the analysis for 2024 stops is 572. Each ratio that enters into the computation involves a comparison of one non-White racial group to Whites for one agency.				

NOTE: in the next few sections (from “Reason for Stop” to “Dog Sniffs”), the results presented do not involve a benchmark.

Reason for Stop

The reason for each stop is summarized in Figure 4a. The percentage of stops for each reason varied substantially by racial group (Figure 4b). As a side note, “Commercial Vehicle” is not a reason to be stopped. Commercial vehicles have a different set of regulations/violations that may not apply to passenger vehicles. Therefore, commercial vehicles have unique reasons for being stopped, such as weight overages and unsecured loads.

Figure 4a. Percentage of stops by reason for stop. Illinois, Traffic stops, 2024.

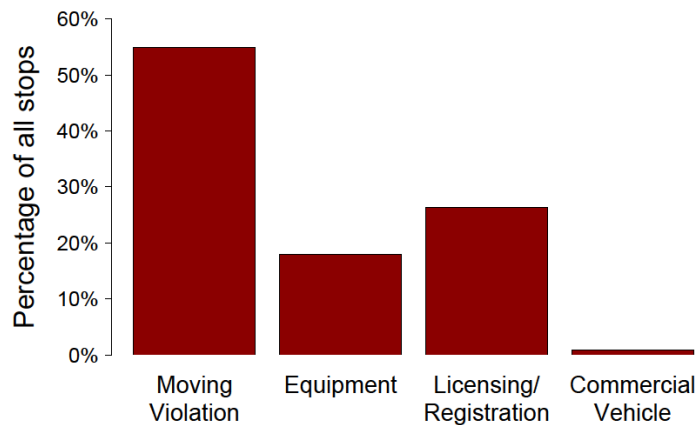
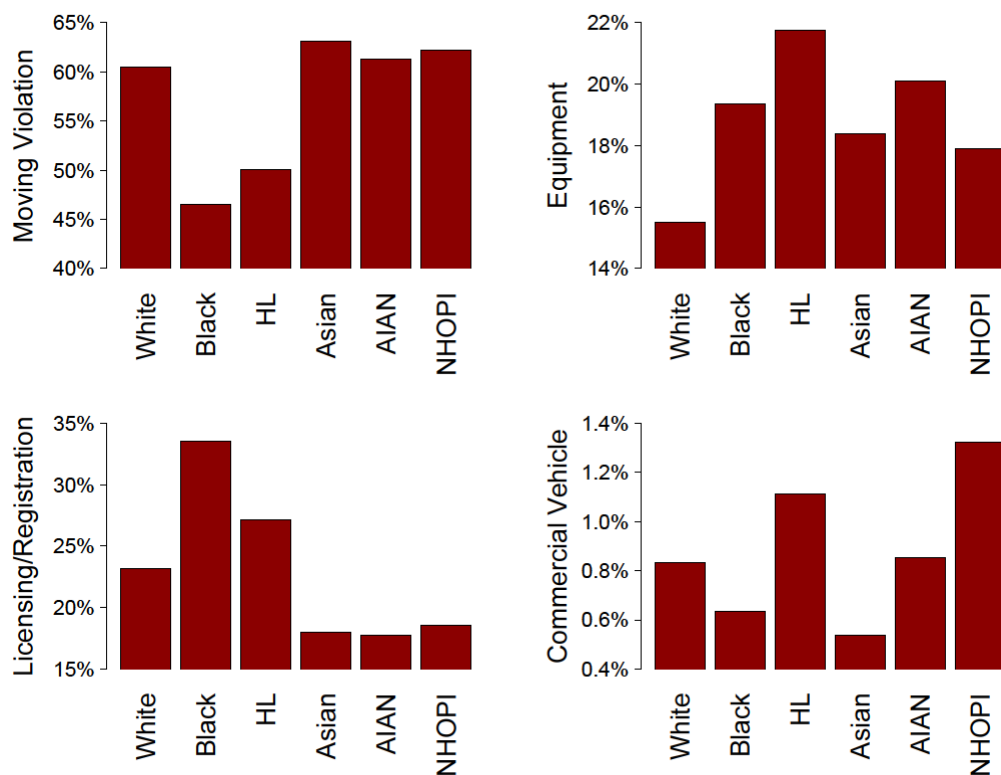


Figure 4b. Percentage of stops for the noted reason, by race. For each race, the percentages sum to 100% across the four noted reasons. Note that the upper and lower limits of the y-axis vary across the four panels. Illinois, Traffic stops, 2024.

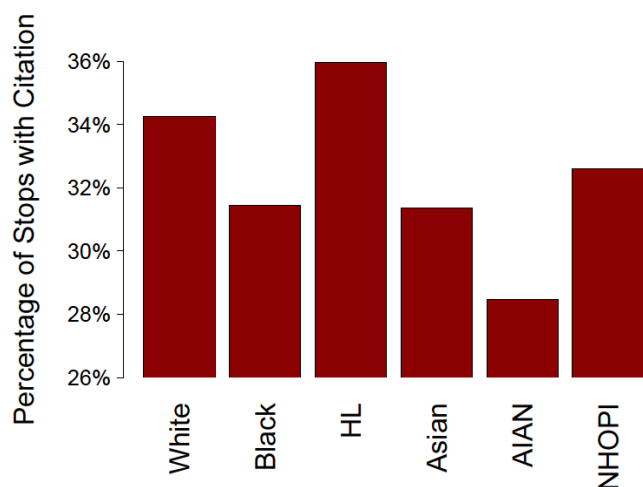


Abbreviations for racial groups: Black = “Black or African American,” HL = “Hispanic or Latino,” AIAN = “American Indian or Alaska Native,” NHOPI = “Native Hawaiian or Other Pacific Islander.”

Outcome of Stop: Citation

For the six racial groups, the percentage receiving a citation as the outcome of the stop ranges between 28% and 36% (Figure 5). “Citation” is the most serious result of the three outcomes recorded on the traffic stop data collection form: citation, written warning or verbal warning/stop card.

Figure 5. Percentage of stops with a citation, by race. Illinois, Traffic stops, 2024.

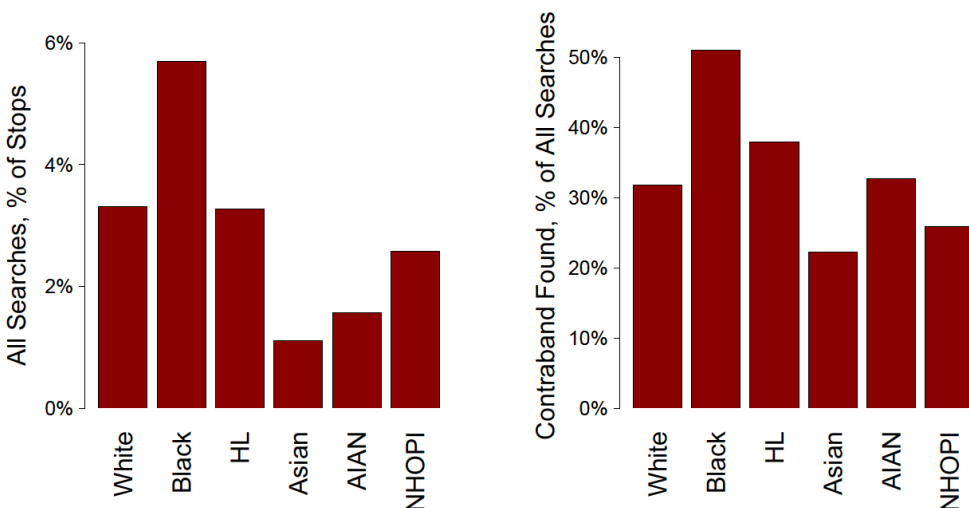


Abbreviations for racial groups: Black = “Black or African American,” HL = “Hispanic or Latino,” AIAN = “American Indian or Alaska Native,” NHOPI = “Native Hawaiian or Other Pacific Islander.”

Searches

Figure 6a shows that the vehicle search rate was moderately low for all of the racial groups (approximately 1-6% of stops, left panel), but given a vehicle search, the contraband yield was not low (22-51% of searches, right panel). There is variation among the races’ percentages in both panels.

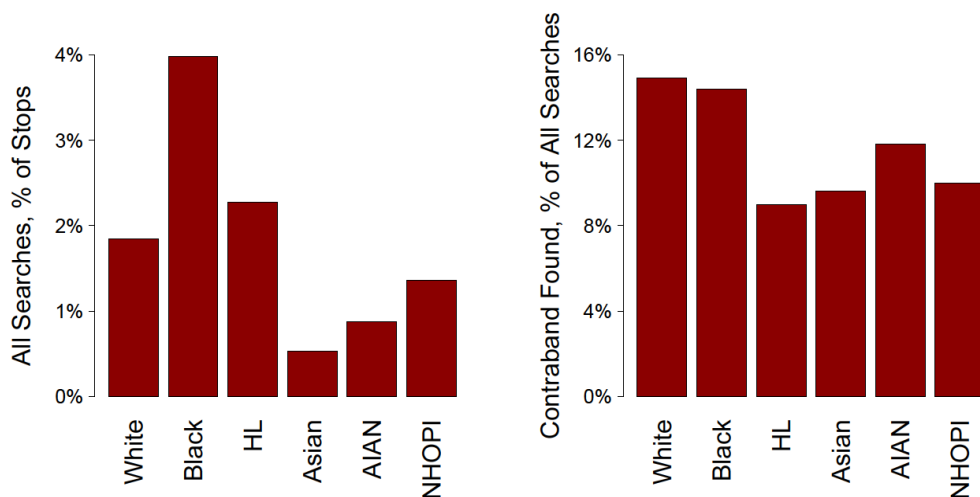
Figure 6a. Percentage of stops with vehicle searches; percentages of vehicle searches with Contraband Found, by race. Note that the upper and lower limits of the vertical axis vary across the two panels. Illinois, Traffic stops, 2024.



Abbreviations for racial groups: Black = “Black or African American,” HL = “Hispanic or Latino,” AIAN = “American Indian or Alaska Native,” NHOPI = “Native Hawaiian or Other Pacific Islander.”

Figure 6b shows that the driver or passenger search rate (searching an individual) was low for all of the racial groups (approximately 0.5-4% of stops, left panel), and given a driver or passenger search, the contraband yield was somewhat higher (9-15% of searches, right panel). As noted for other figures, there is variation among the races’ percentages in both panels.

Figure 6b. Percentage of stops with driver or passenger searches; percentages of vehicle searches with Contraband Found, by race. Note that the upper and lower limits of the vertical axis vary across the two panels. Illinois, Traffic stops, 2024.



Abbreviations for racial groups: Black = “Black or African American,” HL = “Hispanic or Latino,” AIAN = “American Indian or Alaska Native,” NHOPI = “Native Hawaiian or Other Pacific Islander.”

Dog Sniffs

While there were 4,950 dog sniffs performed statewide in 2024 (about the same as 4,969 dog sniffs in 2023), it was still relatively rare compared to the total number of stops by Illinois state law enforcement agencies. Only one in 414 stops in 2024 had a dog sniff. Not all agencies conduct dog sniffs, because the trained dogs are not available in each agency. While the frequency of dog sniffs is low statewide (0.09%-0.28% of stops across the six racial groups), the finding of contraband following a vehicle search after a dog sniff is substantial, at 34-66% of vehicle searches across the four racial groups, excluding the American Indian and Native Hawaiian groups, which had too few stops for this comparison. These two groups have very small numbers of stops with dog sniffs (11 and 12, respectively), making them unreliable for more detailed contrasts.

Table 5. Number of stops with a dog sniff and their percentage among all stops. Given that a dog sniff occurred, number and percentage of stops with contraband found. Illinois, Traffic stops, 2024.

Racial Group	Stops with Dog Sniff		Contraband Found	
	Number	Percentage of stops	Number	Percentage of vehicle searches*
White	2,839	0.28%	1,278	62%
Black or African American	1,314	0.26%	591	66%
Hispanic or Latino	701	0.16%	173	40%
Asian	73	0.09%	16	34%
American Indian or Alaska Native	11	0.11%	7	78%
Native Hawaiian or Other Pacific Islander	12	0.23%	2	40%
All groups combined	4,950	0.24%	2,067	41.8%
*The vehicle search occurred after a dog sniff.				

Multiple stops of individual drivers

Here we analyze the number of times each stopped individual driver was stopped. All stopped drivers are grouped according to their race, and for each group we calculate the proportion of drivers stopped exactly once during 2024, stopped 2-3 times, 4-10 times, and over 10 times. In each racial group these proportions sum up to 100%.

As in the previous year’s report, individual drivers are recognized in the data by their unique combinations of name, year of birth, zip code of residence and gender. The same amount of data cleaning as in the previous report was performed on officer-recorded names so that the most frequent patterns in the way drivers’ names are entered into the dataset are captured and standardized. By these adjustments, “John Doe,” “Doe, John,” “John L. Doe Junior,” etc., are recognized as the same name. There may be instances of an individual name being written in nonstandard ways that the algorithm failed to recognize as the same name. Thus, multiple stops may be somewhat more prevalent in the data than what was detected in this analysis, and more sophisticated name-handling techniques would capture more name matches. This analysis found 1,594,718 individual drivers whose race was recorded in at least one stop. If an individual driver was assigned different races in different stops, we set their race as their most frequent race assignment (in the case of a tie, we randomly selected between several most frequent assignments).

A summary finding is that 82.0% of the individual drivers were stopped exactly once, 15.8% were stopped 2-3 times, 2.1% were stopped 3-10 times and 0.04% were stopped over 10 times. These numbers are similar to 2023 results, except in the last category (0.07% in 2023). More detailed results are shown in Figure 7. The peak associated with Black drivers is not as prominent as it was in 2023 (compare to Figure 3 in 2023 traffic stops report, page 17) in categories 4-10 and over 10 times stopped.

These stopped drivers are not necessarily representative of their driver source populations, because this analysis is only about the drivers who have been stopped at least once. Stopped drivers may not accurately represent all drivers, a large fraction of whom were not stopped at all during 2024: roughly statewide, about 2 million traffic stops are distributed over 10 million drivers.

With this in mind, Black drivers still stand out as the racial group whose stopped individual members had the highest occurrence of being stopped multiple times.

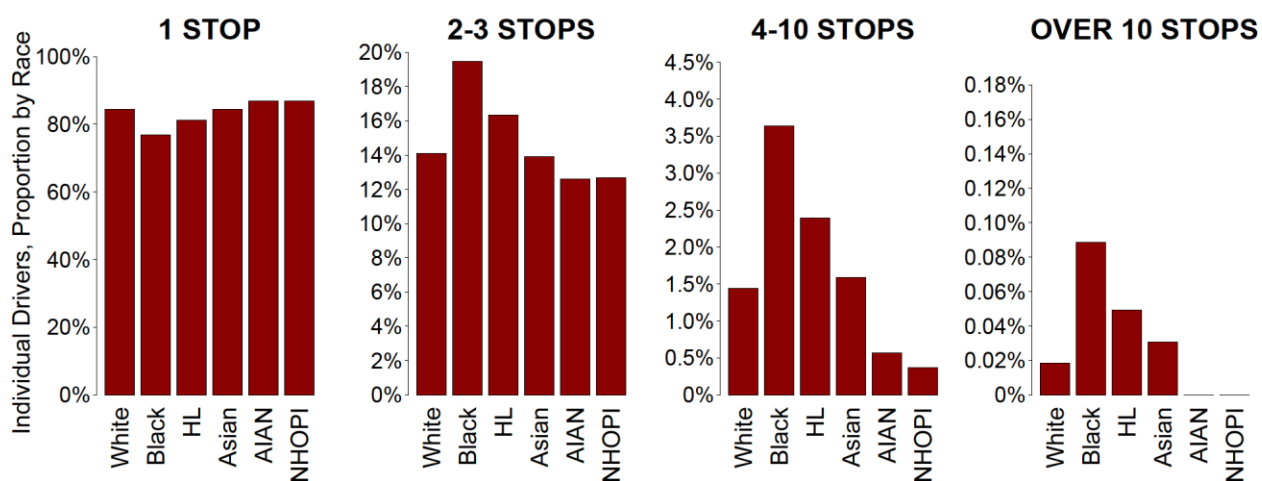
Compared to an average (stopped) White driver, an average (stopped) Black driver had a 38% higher chance of having been stopped 2-3 times, 2.5 times higher chance of having been stopped 4-10 times, and 4.7 times higher chance of having been stopped over 10 times. To illustrate the stops comparison from another perspective, in 2024, although estimated Black drivers’ statewide population (1,912,683) was 2.7 times smaller than the estimated White drivers’ population (5,202,375), the number of individual Black drivers stopped more than 10 times (322) was 2.1 times larger than the number of individual White drivers stopped more than 10 times (151). It may be of interest to note that in 2023, this ratio was more than twice larger: 5.1. In 2024 there was a notable reduction of drivers stopped 10 times or more of all races and particularly among Black drivers, see Table 6.

Table 6. Counts of individual drivers that were stopped 10 or more times, by race. Illinois, Traffic stops, 2023 and 2024.

Year	White	Black	Hispanic or Latino	Asian	AIAN & NHOPI	All Races
2023	171	865	250	14	0	1,300
2024	151	322	168	21	0	662

This analysis was partially motivated by interviews with police officers, see section VIII of the 2023 traffic stops report. Some officers stated that they may seek specific vehicles or behaviors, which entails recognition and/or tailing of individual drivers and, potentially, their multiple stops throughout the year. Our analysis suggests that this practice may indirectly involve a racial aspect.

Figure 7. Proportion of individual drivers stopped a particular number of times, among all stopped drivers of a particular racial group. First panel is drivers stopped exactly once, second panel is drivers stopped 2-3 times, third panel is drivers stopped 4-10 times, fourth panel is drivers stopped over 10 times. Illinois, Traffic stops, 2024.



Abbreviations for racial groups: Black = “Black or African American,” HL = “Hispanic or Latino,” AIAN = “American Indian or Alaska Native,” NHOPI = “Native Hawaiian or Other Pacific Islander.”

We repeated the analysis from the previous report of drivers with multiple stops that were assigned to different race groups on different occasions. The results were essentially the same as in 2023, and thus are not reported here. The interested reader can find them in the report on 2023 traffic stops.

Officer-assigned driver’s race analysis

In this section we continue the exploration of the issue of the driver’s race assignment by the officer, from previous reports. This work is both exploratory and self-contained: we never actually change the officer-assigned race when calculating and presenting anything outside this section of the report.

In our past reports on 2022-2023 stops, we introduced a way to statistically analyze the officer-assigned race, looking for potential signs of misclassification. We employ Bayesian Improved Surname Geocoding (BISG), a statistical methodology commonly used in social sciences that estimates race based on a person’s first name, last name and zip code. See our 2022 traffic stop report, section VI, for more details.

In essence, for a given full name and a zip code of residence, BISG returns a set of probabilities that a person belongs to each race group. Table 7 provides some illustration. A few fictional names were

created by recombining most common first and most common last names found among the 2024 stops, with none of the combinations being an actual name of a driver among the reported stops. Table 7 shows examples of probability (first column, probability ranges) that BISG will assign to a name and a race. For example, a person named “Nathaniel Woodruff” would be given a medium probability (60-70%) of being Hispanic/Latino.

Table 7. Probabilities estimated by BISG that the given fictional individual residing in zip area 60490, is of a particular race. API = ‘Asian or Native Hawaiian or Other Pacific Islander.’

BISG estimates	White	Black	Hispanic/ Latino	API
High Probability (>95%)	Alistair Smith	Eddie Jefferson	Julian Salazar	Ahmad Patel
Medium Probability (60-70%)	Jeremiah Gray	Nathaniel Woodruff	Willie Rodriguez	Hamza Clarence
Low Probability (<40%)	Destiny Cadle	Fatima Nelson	Diamond Cotts	Mohammad Williams

For a traffic stop, a “mismatch” occurs when the officer-assigned race differs from the one BISG gave the highest probability to. If this highest probability was over 95%, a “high-probability (HP) mismatch” occurs. Simply put, these HP-mismatches represent unlikely pairings of a name and a race, which makes them suitable as a statistical analysis tool. For example, an officer reporting “Alistair Smith” as Asian would constitute an HP-mismatch.

In 2024, about 5% of all traffic stops were HP-mismatched. Among 708 agencies that reported at least 20 traffic stops, 14 agencies had over 15% of their reported stops HP-mismatched. These were largely the same individual agencies that had similarly high HP-mismatch rates for 2023 stops. These agencies also tend to have a large Hispanic/Latino driving population — the median proportion of Hispanic/Latino drivers among all drivers in those 14 agencies is 3.5 times as high as the median of all 708 agencies.

It does not seem unusual if agencies with a high presence of minority and mixed-race drivers get more opportunity to encounter the less common drivers’ names, triggering HP-mismatches more often. However, some stop statistics could be expected to be largely independent of the frequency of HP-mismatches, and that could be investigated using BISG estimation.

For a hypothetical example, consider the situation when contraband is found or not found after a vehicle or driver/passenger search. A recent paper (Luh 2022) suggested that officers may be incentivized to fake race designation if a search with a minority driver did not find contraband. If true, agencies with high minority populations would be expected to have a higher rate of HP-mismatches when contraband was not found, compared to searches when contraband was found.

To potentially observe this in the stop data, we partitioned all 786 compliant agencies into groups according to the number of HP-mismatches their stops have. See Table 8. The groups were formed so that they also collectively reported similar numbers of searches. Next, using those collective searches,

we observed two HP-mismatch rates: first in searches where contraband was not found, and second in searches where contraband was found. We then formed the Rate Ratio of the rate with contraband not found vs. the rate when contraband was found. To potentially correlate this with the presence of minorities, for each group of agencies we calculated the average proportion of White drivers per agency. See Table 8.

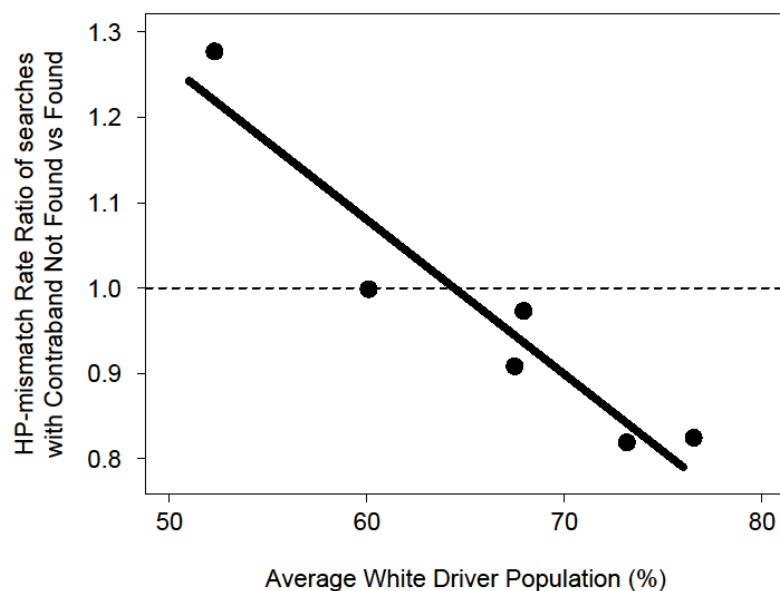
Table 8. Partitioning of all 786 compliant agencies into six groups according to their individual numbers of HP-mismatches. Shown for each group: cumulative number of searches and the ratio of HP-mismatch rates when contraband was not found vs. when it was found. Average White driver population % per agency group was calculated using population benchmarks. Illinois, Traffic stops, 2024.

HP-mismatches per Agency	0-4	5-9	10-19	20-39	40-99	100 or more
Total Searches per Group	14970	14131	13021	13272	12759	13786
HP-mismatch Rate Ratio: Not Found/Found	0.82	0.97	0.82	0.91	1.00	1.28
Average White Driver Population % per Agency	76.5%	67.9%	73.2%	67.5%	60.1%	52.3%

Figure 8 shows a strong linear trend between the average proportion of White drivers and the HP-mismatch rate ratio. Pearson correlation coefficient is -0.95 (p-value 0.005), which suggests a nearly perfect correlation. This trend can be summarized in the following way: among searches, the smaller the proportion of White drivers, the more likely HP-mismatch will occur after contraband was not found than after it was found. Depending on the average proportion of White drivers being low or high, searches with contraband not found could have a 20% higher or lower chance of having an HP-mismatch, as compared to searches where contraband was found.

This result may indeed suggest that race designation of searches is not independent of contraband being found or not, and that with more minority drivers present in the local driver population the misclassifications occur more after failed searches. However, these findings are of observational nature and could serve only as a starting point for a further, more rigorous analysis. No particular agency can be named as introducing errors into their stop data.

Figure 8. Observed linear trend (regression line) between the average White driver proportion and the HP-mismatch rate ratio (dots), among the six groups of agencies. Illinois, Traffic stops, 2024.



References (for race analysis section)

1. Tzioumis, K. (2018) Demographic Aspects of First Names. *Sci Data* 5, 180025. <https://doi.org/10.1038/sdata.2018.25>.
2. Haas A, Adams JL, Haviland AM, et al. (2022). The Contribution of First-name Information to the Accuracy of Racial-and-Ethnic Imputations Varies by Sex and Race-and-Ethnicity Among Medicare Beneficiaries. *Medical Care*, 60(8):556-562.
3. Elliott MN, Morrison PA, Fremont A, et al. (2009) Using the Census Bureau's surname list to improve estimates of race/ethnicity and associated disparities. *Health Serv Outcomes Res Method*, 9(2):69–83.
4. Voicu, I. (2018) Using first name information to improve race and ethnicity classification. *Statistics and Public Policy*, 5(1):1-13.
5. Luh, E. (2022) Not So Black and White: Uncovering Racial Bias from Systematically Misreported Trooper Reports. *SSRN*. <http://dx.doi.org/10.2139/ssrn.3357063>.
6. Argyle, Lisa and Michael Barber "Misclassification and Bias in the Predictions of Individual Ethnicity from Administrative Records." *American Political Science Review*, Volume 118, Issue 2, May 2024, pp. 1058 – 1066.

VII. Considerations for Interpreting the Data

Here we reiterate some general considerations mentioned throughout the report, for interpreting the results shown in the statistical tables of individual agencies.

A considerable number of agencies have a relatively small number of stops for one or more of the racial groups. The limited stop counts yield a wide 95% confidence interval, which means high uncertainty in the corresponding rate, percentage, or ratio, due to play of chance. The uncertainty from potential benchmark issues (discussed earlier) or race classification issues (also discussed earlier) add somewhat to the uncertainty implied by the confidence intervals. Any investigation of racial profiling that is initiated based on this report should consider both confidence intervals and those other potential sources of uncertainty.

In Part II of this report (agency tables) each agency has ratios of rates or ratios of percentages. Some of them are bolded as a “statistical deviation.” The bolded ratios and their meaning and interpretation are topics covered elsewhere in this report. In addition to whether or not a ratio is bolded, the absolute magnitude of the ratio should be considered when interpreting the results, as discussed earlier.

If a ratio is not bolded, it does not prove that there is no racial profiling in the agency. It is worth looking at the upper and lower bound of the 95% confidence interval to see what the uncertainty is. That interval quantifies the uncertainty due to the play of chance and shows the largest ratio and the smallest ratio that are reasonably plausible, given the data.

For example, consider a ratio of **1.0** for a specific Minority percentage of stops with a search, compared to the corresponding White percentage of stops with a search — in a particular agency. The ratio of 1.0 indicates that the percentage of stops with a search was the same for both the Whites and for the specific Minority group. However, the counts of searches are very small in this example, and the 95% confidence interval for the ratio is **0.025** up to **5.8**. (This is very similar to an actual agency result.) That is, it is plausible that the true search percentage of the Minority group is anywhere from one-fortieth of the White percentage up to almost six times the White percentage.

Clearly, in a case like the one described above, we do not know enough about the ratio to draw any conclusion except that we are uncertain. Thus, a confidence interval for a ratio that includes 1.0 and is very wide (encompassing values well above the calculated ratio and also well below the ratio) usually means that presence or absence of potential racial profiling cannot be determined from the data in hand.

Lastly, while there is a considerable focus on the stop rate ratios reported in Panel 1 of the tables in Part II of this report (detailed tables), the other panels provide valuable complementary information on the outcomes of stops and how the outcome statistics compare between racial groups. As noted earlier, the stop outcome results do not rely on any external benchmark and thus avoid the issue of benchmark accuracy.

Ultimately, stop results for an agency should be interpreted holistically, considering all panels together; different panels may suggest different interpretations when viewed individually.

VIII. Looking Ahead

TMWL is continuing to review the current statistical methodology and consider refinements and improvements. In our analysis of 2021 stops we made a major update to our benchmarking approach

that was carried forward to the 2024 stop study. In this report we introduce as an alternative a benchmark calculated from a very different model and informed by very different sources of data. Our striving for ever-more-accurate benchmarks will likely continue as new relevant datasets and methodologies emerge.

Appendix A. Traffic Stop Data Collection Form in use during 2024



Illinois Department
of Transportation

Traffic Stop Data Sheet



Agency Code

Section A - Traffic Stop Information

Date of Stop (MM/DD/YYYY) Time of Stop (Military Time) Duration of Stop (Minutes)

Officer Name Officer Badge Number

Name of Driver

Address City State Zip Code

Vehicle Make Vehicle Year Driver's Year of Birth (ex: 1957)

Driver Sex
1 ☐ Male 2 ☐ Female

Driver Race
1 ☐ White 2 ☐ Black or African American 3 ☐ American Indian or Alaska Native 4 ☐ Hispanic or Latino 5 ☐ Asian
6 ☐ Native Hawaiian or Other Pacific Islander

Reason for Stop
1 ☐ Moving Violation 2 ☐ Equipment 3 ☐ License Plate / Registration 4 ☐ Commercial Vehicle

If Moving, Type of Violation
1 ☐ Speed 2 ☐ Lane Violation 3 ☐ Seat Belt 4 ☐ Traffic Sign or Signal 5 ☐ Follow too Close 6 ☐ Other

Result of Stop
1 ☐ Citation 2 ☐ Written Warning 3 ☐ Verbal Warning / Stop Card

Beat of Location Stop

Section B - Searches

Vehicle	Consent Search Requested?	Consent Given?	Search Conducted?	Search Conducted By?
	1 ☐ Yes 2 ☐ No	1 ☐ Yes 2 ☐ No	1 ☐ Yes 2 ☐ No	1 ☐ Consent 2 ☐ Other

If yes, what was found: 1 ☐ Drugs 2 ☐ Drug Paraphernalia 3 ☐ Alcohol 4 ☐ Weapon 5 ☐ Stolen Property 6 ☐ Other

If a search of the Vehicle was conducted, was contraband found? 1 ☐ Yes 2 ☐ No

If the contraband found was drugs, what was the amount? 1 ☐ < 2 grams 2 ☐ 2-10 grams 3 ☐ 11-50 grams 4 ☐ 51-100 grams 5 ☐ > 100 grams

Driver	Consent Search Requested?	Consent Given?	Search Conducted?	Search Conducted By?
	1 ☐ Yes 2 ☐ No	1 ☐ Yes 2 ☐ No	1 ☐ Yes 2 ☐ No	1 ☐ Consent 2 ☐ Other

Passenger(s)	Consent Search Requested?	Consent Given?	Search Conducted?	Search Conducted By?
	1 ☐ Yes 2 ☐ No	1 ☐ Yes 2 ☐ No	1 ☐ Yes 2 ☐ No	1 ☐ Consent 2 ☐ Other

If a search of the Driver or Passenger(s) was conducted, was contraband found? 1 ☐ Yes 2 ☐ No

If yes, what was found: 1 ☐ Drugs 2 ☐ Drug Paraphernalia 3 ☐ Alcohol 4 ☐ Weapon 5 ☐ Stolen Property 6 ☐ Other

If the contraband found was drugs, what was the amount? 1 ☐ < 2 grams 2 ☐ 2-10 grams 3 ☐ 11-50 grams 4 ☐ 51-100 grams 5 ☐ > 100 grams

Section C - Police Dog Sniff Searches

Did a police dog perform a sniff of the vehicle? 1 ☐ Yes 2 ☐ No

If a police dog performed a sniff of the vehicle, did the dog alert to the presence of contraband? 1 ☐ Yes 2 ☐ No

If an alert occurred, was the vehicle searched? 1 ☐ Yes 2 ☐ No

If the vehicle was searched, was contraband found? 1 ☐ Yes 2 ☐ No

If yes, what was found: 1 ☐ Drugs 2 ☐ Drug Paraphernalia 3 ☐ Alcohol 4 ☐ Weapon 5 ☐ Stolen Property 6 ☐ Other

If the contraband found was drugs, what was the amount? 1 ☐ < 2 grams 2 ☐ 2-10 grams 3 ☐ 11-50 grams 4 ☐ 51-100 grams 5 ☐ > 100 grams

Printed 05/19/20

BDC 2581 (Rev. 02/21/17)
Formerly TS 2581

Appendix B. Technical Notes on Rates, Percentages and Ratios

B.1. Overview

This technical appendix includes a detailed explanation of the rate, post-stop outcomes, and ratio calculations used in constructing the statewide and agency tables that appear in Part II of this report. We explain how comparisons of each Minority group to White drivers or pedestrians are carried out. We also explain how the confidence interval is calculated based on known sources of uncertainty in the data.² Further, this section describes how an agency may be designated (by a bold font in the tables) as potentially standing out beyond an assumption of no racial profiling. An agency that is designated as standing out might use this report as a basis for further inquiry. As stated elsewhere and repeated here, there is nothing in this report that proves an agency is practicing racial profiling. We provide some limitations for interpreting the findings based on the available data and methods.

B.2. Stop rates, post-stop outcomes and ratio calculations

We performed all calculations for the entire state of Illinois and for each agency.

B.2.1 Stop rates and rate ratios

We calculated stop rates separately for each racial group by dividing the number of stops in the racial group by the population benchmark estimate of the driving population in the racial group. A description of the methods used to estimate the population benchmark was explained at length in Appendix C of the previous year's report (pages 34-89).³

We assumed the number of stops followed a Poisson distribution, used in previous examination of racial disparities in traffic stops (Gelman et al. 2007, Ridgeway 2007) and calculated 95% confidence intervals for the rates using exact methods (Garwood 1936). (References are at the end of this section.) When the benchmark estimate of the population was zero, no rate or confidence interval could be calculated. A benchmark population of zero for a specific Minority group happens when the census population estimate for the Minority is zero.

We compared each Minority group to White drivers or pedestrians using the ratio of the Minority group stop rate to the White group stop rate. We calculated a 95% confidence interval for each rate ratio by conditioning on the sum of the numbers of stops in the two racial groups being compared. Assuming the number of stops in each group followed a Poisson distribution, conditioning on the sum of the number of stops creates a binomial variable. For distance-based population benchmarks, an exact confidence was calculated using binomial methods (Lehmann and Romano 2005). If it was impossible to calculate a rate because of a zero benchmark, or if the number of stops in the White group was zero, no rate ratio or confidence interval was reported.

We calculated the 95% confidence intervals for rate ratios from crash-based population benchmarks in a different way than for distance-based population benchmarks in order to incorporate the number of crashes used in the population benchmark (see Appendix C of the previous year's report for how crash-

² The estimated benchmark population is an example of a component of the methodology that has uncertainty that could not be quantified for this study.

³ <https://idot.illinois.gov/content/dam/soi/en/web/idot/documents/transportation-system/reports/safety/traffic-stop-studies/final--part-i-executive-summary-traffic--6-30-23.pdf>.

based and distance-based benchmarks were defined and calculated). For each Minority group, the proportion of Minority stops out of the sum of the Minority and White stops (p_{stops}) and the proportion of the Minority group in the benchmark population out of the Minority and White groups ($p_{benchmark}$) were calculated. The rate ratio (for a given Minority compared to Whites) can be calculated from these proportions using the following formula:

$$\frac{\left(\frac{p_{stops}}{1-p_{stops}} \right)}{\left(\frac{p_{benchmark}}{1-p_{benchmark}} \right)}.$$

However, the corresponding 95% confidence interval for the rate ratios requires the effective sample sizes (the numerator and denominator) corresponding to $p_{benchmark}$, which is related to the number of crashes used in the population benchmark.

The stops proportion was treated as a binomial variable, as above. The population benchmark proportion was initially treated as an over- or under-dispersed binomial with the number of crashes used as the denominator. The variance of the population benchmark proportion was estimated using the parametric bootstrap, where the number of crashes per ZIP code was drawn from a multinomial distribution for each bootstrap iteration. The dispersion parameter of the population benchmark proportion was estimated as the ratio of the bootstrap variance divided by the variance that is estimated assuming a standard binomial proportion (i.e., using the classic formula: $p[1 - p]/N$, where p is the population benchmark proportion and N is the number of crashes). The dispersion parameter indicates how much more variable (dispersion > 1) or less variable (dispersion < 1) the proportion is than expected for a standard binomial variable if the denominator was the number of crashes. The effective denominator for the population benchmark proportion, which is the denominator that would produce the same variance as expected using the standard binomial formula, was then calculated as the number of crashes divided by the dispersion parameter. Similarly, the effective numerator of the population benchmark proportion was calculated as the benchmark proportion times the effective denominator. Using the number of Minority stops, White stops, effective benchmark numerator and effective benchmark denominator, the 95% confidence of the rate ratio was calculated using exact binomial methods as carried out above for distance-based population benchmarks. This method of calculating 95% confidence intervals tends to produce wider intervals than if they were calculated the same way as for distance-based population benchmarks, because the effective benchmark numerator and denominator based on the number of crashes are each less than the corresponding benchmark population counts. This methodology is used to account for additional variability in the benchmark population estimates related to the number of crashes, which is generally smaller than the number of stops.

A rate ratio of 1.0 indicates the Minority group and White drivers or pedestrians had equal rates of stops. If the 95% confidence interval lies entirely above 1.0, the rate ratio is statistically significantly greater than 1.0 and may require agency inquiry. These statistically significant rate ratios are bolded in the summary tables. These bolded ratios are statistical deviations and the basis for further consideration of potential racial disparities. Comparisons of Minority groups to White drivers or pedestrians where the 95% confidence lies below 1.0 are not bolded because the intent of this study is to identify potential racial profiling that discriminates against Minority drivers or pedestrians.

B.2.2 Post-stop outcomes

For all calculations, we assumed the population benchmark accurately captured the population of drivers or pedestrians. The population benchmark used to calculate each rate is itself an estimate of the population of drivers or pedestrians for a racial group. Confidence intervals of rates and rate ratios assumed only sampling error and thus do not account for this additional source of error in population benchmark estimates. Accounting for population benchmark error would increase the width of the confidence intervals reported for rates and rate ratios and would likely reduce the number of agencies that appear to stand out as needing further inquiry. We calculated post-stop outcome percentages separately for each racial group. Table B1 shows the type of numerator and denominator used to calculate each percentage shown in the traffic tables.

Table B1. Numerators and denominators for traffic stop outcomes.

Outcome	Numerator	Denominator
CATEGORY: Reasons for Stop		
Moving Violation	Number of moving violation stops	Number of stops
Equipment	Number of equipment stops	Number of stops
Licensing/Registration	Number of licensing/registration stops	Number of stops
Commercial Vehicle	Number of commercial vehicle stops	Number of stops
CATEGORY: Outcomes of Stop		
Verbal Warning	Number of verbal warnings	Number of stops
Written Warning	Number of written warnings	Number of stops
Citation	Number of citations	Number of stops
CATEGORY: Vehicle Searches		
Consent Search	Number of consent searches	Number of stops
All Searches	Number of searches	Number of stops
Contraband Found	Number of searches where contraband was found	Number of searches
CATEGORY: Driver or Passenger Searches		
Consent Search	Number of stops with a consent search*	Number of stops
All Searches	Number of stops with a driver or passenger search*	Number of stops
Contraband Found	Number of stops with a driver or passenger search where contraband was found*	Number of stops with a driver or passenger search*
CATEGORY: Dog Sniff Searches		
Dog Sniff	Number of dog sniffs	Number of stops
Dog Alert after Dog Sniff	Number of dog alerts	Number of dog sniffs
Vehicle Search after Dog Sniff	Number of vehicle searches after a dog sniff	Number of dog sniffs
Contraband Found after Vehicle Search	Number of vehicle searches after a dog sniff, where contraband was found	Number of vehicle searches following a dog sniff
*Although a stop may result in the search of more than one individual (e.g., both the driver and a passenger are searched), multiple individuals searched (from one vehicle) are counted here as one stop with a driver or passenger search or both.		

We assumed that percentages follow a binomial distribution and can be approximated by a Poisson distribution (Serfling 1978), and we calculated confidence intervals for the rates using exact methods (Garwood 1936). When the denominator of the percentage was zero (for example, an agency had a benchmark of zero for a specific racial group), no percentage or confidence interval could be calculated.

For selected outcomes we compared each Minority group to White drivers, using the ratio of the Minority group percentage to the White group percentage. We calculated a 95% confidence interval for each ratio using exact methods (Lehmann and Romano 2005). If it was impossible to calculate a percentage because of a zero denominator, or if the numerator of the White group percentage was zero, no ratio or confidence interval was reported.

B.3 Durations

We calculated the median durations of stops separately for each racial group. The median represents the value such that about half of stops have a shorter duration than the median and half of stops have a longer duration than the median.

B.4 Limitations

For all calculations, we assumed that the driver or pedestrian was assigned to the correct racial group. However, an officer's assessment of the race of a driver may be in error compared to the driver's self-assessed race. Because police officers made the racial group assignment, there is a potential misclassification bias on the race of drivers or pedestrians. If misclassification resulted in a Minority driver or pedestrian frequently being categorized in a different Minority group, the stop rates of some Minority groups may be underestimated, while others are overestimated. Consequently, the rate ratios of some Minority groups may be underestimated, while others are overestimated. This is a limitation that would be difficult to correct based on the available information. Section IV of this report consider — in more detail — the issue of determining race of drivers.

Some of the alerts to rate ratios (**bolded font** in the tables) may be “false positives.” This can happen as follows. Within the statewide or individual agency tables for traffic and pedestrian stops, we calculated five Minority group comparisons with the White group. There were five of these comparisons for each ratio analysis. For example, there are five ratios comparing the stop rate for each of the five minorities to the stop rate for Whites.⁴ Thus, we constructed five 95% confidence intervals — one each for the five stop-rate ratios. That is, each agency was checked in each of the five Minority groups for potential profiling. For each Minority comparison with White drivers or White pedestrians there was the potential to make what statisticians refer to as a “type I error.” That is, by chance, a rate ratio and its confidence interval may have incorrectly indicated the need for inquiry for profiling. While we set a 5% type I error rate for each Minority comparison, the multiple (five) comparisons inflate the possibility of making such an error overall to more than 5%. We chose not to correct for these multiple comparisons, viewing each Minority comparison to Whites as an independent examination of profiling.

⁴ There may be fewer than five ratios depending on the occurrence of zero stops for Whites or zero benchmark for a Minority. These are cases where a ratio cannot be calculated.

References (for Appendix B)

- Garwood, F (1936). Fiducial limits for the Poisson distribution. *Biometrika*, Vol. 28, Issue 3-4: 437-442.
- Gelman, A, Fagan, J, and Kiss, A (2007). An analysis of the New York City Police Department's 'stop-and-frisk' policy in the context of claims of racial bias. *Journal of the American Statistical Association*, Vol. 102, No. 479, 813–823.
- Lehmann, EL, and Romano, JP (2005). *Testing Statistical Hypotheses*, Third edition. Springer: New York.
- Ridgeway, G (2007). *Analysis of Racial Disparities in the New York Police Department's Stop, Question, and Frisk Practices*. Santa Monica, CA: RAND Corporation.
https://www.rand.org/pubs/technical_reports/TR534.html
- Serfling, RJ (1978). Some elementary results on Poisson approximation in a sequence of Bernoulli trials. *SIAM Review*, Vol. 20, No. 3, 567-579.

Appendix C. Technical Notes on Population Benchmarks

C.1. Overview

In the analysis of potential racial profiling, the number of stops by each agency of each racial group is compared to a “population benchmark” of the racial group. The rate of stops per benchmark population for the racial group can be compared to the same rate for Whites. The population benchmark provides an expected racial distribution of the local population of drivers.

This distribution would be approximately equal to the expected racial distribution of the stops if the stops were conducted in a completely randomized way, blind to the race and the behavior of the driver. That is, the stop rates calculated using a perfectly accurate population benchmark would be approximately constant across all racial groups if there were no profiling and if there were no difference in the general behavior of drivers across all racial groups.

This report shares the same methodology of calculating the population benchmarks as our previous year's report. The only difference is that our data sources were updated to their most recent available versions, and that there were some changes in the selection of data sources to be used this year. Details on this are covered below. Details on how racial categories were defined, how population benchmark regions were determined and other population benchmark calculations, the differences in population benchmark methodology employed now compared with prior years, and limitations and strengths of the methodology are described at length in the Appendix C of our previous year's report.

C.2. Data Sources

Multiple data sources were combined to calculate population benchmarks, including multiple datasets provided by the U.S. Census Bureau, Illinois Department of Transportation and Illinois Secretary of State. The U.S. Census Bureau datasets used include those from the decennial census, the American Community Survey and Gazetteer files, depending on the year and type of population benchmark (traffic stops or pedestrian stops).

C.2.1. Data from the U.S. Census Bureau

The American Community Survey is an ongoing survey conducted by the U.S. Census Bureau that collects information on the U.S. population in all 50 states, the District of Columbia and Puerto Rico.⁵ The information collected is similar to that collected by the U.S. decennial census, but the ACS results are released on an annual basis rather than every 10 years. Another difference between the ACS and census is that the ACS is based on a random sample of about 3.5 million individuals while the census attempts to reach every person living in the U.S. and its territories.

Besides the 1-year (1Y) ACS releases, there are also 5-year (5Y) releases. These 5Y releases combine 5 consecutive years, primarily to increase the sample size of relatively small areas or groups of individuals. It would be challenging to estimate the population of small communities reliably with only one survey year of data. In addition to standard tabulations, the ACS also provides individual-level data, referred to as the public use microdata sample (PUMS). The PUMS data allows more detailed and complex analyses involving multiple variables. Due to privacy concerns, there are restrictions on the level of geographic identification provided with each type of release of ACS data.

The Gazetteer files provide geographic information, such as geographic area, latitude and longitude, for different relevant regions in the U.S., including ZIP codes, places (a city, town, or village, referred to simply as city hereafter), counties, and states. These files are updated annually.

The U.S. Census Bureau approximates ZIP codes (defined by the U.S. Postal Service) with ZIP code tabulation areas (ZCTAs).⁶ Throughout this report, the term “ZIP code” will be used to refer both to ZCTAs and U.S. Postal Service ZIP code for simplicity.

Table C.1 lists the U.S. Census Bureau datasets used for different purposes, for both the traffic and pedestrian stop population benchmarks. More detail on pedestrian stop population benchmarks can be found in the corresponding Illinois pedestrian stops study report, 2024 stops, Part I. Of note, as can be seen from the table, different datasets were used for traffic and pedestrian population benchmarks, which is different than in past years. The primary reason is that pedestrian population benchmarks are based on city-, county- or state-level population statistics, while the traffic stop population benchmarks are based on ZIP-code-level population statistics.

The reader who compares this appendix to the corresponding appendix in the 2024 pedestrian stops report will note that the decennial 2020 census data is not used either in this traffic analysis, nor in the 2024 pedestrian stops analysis. The reason is that the newest 2023 5-year ACS release covers years 2019-2023, with the year 2020 being earlier in that time interval, so the census data are now less “current” than ACS data. We thus plan to keep using the newest 5Y releases until the next decennial census becomes available.

⁵ <https://www.census.gov/programs-surveys/acs>. Last accessed 5/25.

⁶ <https://www.census.gov/programs-surveys/geography/guidance/geo-areas/zctas.html>. Last accessed 5/25.

Table C.1. U.S. Census Bureau datasets used for benchmarks.

Information Needed	Traffic Stop Benchmarks	Pedestrian Stop Benchmarks
Age distribution in Illinois	1Y ACS PUMS 2023	N/A
Age distribution by race/ethnicity*	5Y ACS PUMS 2019-2023	5Y ACS PUMS 2019-2023
Individual race groups to reallocate residents with more than one race*	5Y ACS PUMS 2019-2023	5Y ACS PUMS 2019-2023
Population counts for each race/ethnicity		
By ZIP code†	5Y ACS 2019-2023	5Y ACS 2019-2023‡
By city	N/A	5Y ACS 2019-2023
By county	N/A	5Y ACS 2019-2023
For Illinois	N/A	5Y ACS 2019-2023
Geographic area of each city in Illinois	Gazetteer Files 2024	N/A
Geographic area of each county in Illinois	Gazetteer Files 2024	N/A
Latitude and longitude of each ZIP code	Gazetteer Files 2024	N/A
1Y = 1-year; 5Y = 5-year; ACS = American Community Survey; DEC = decennial census; PUMS = public-use microdata sample; *Includes Illinois and 24 states within 400 miles of Illinois; †ZIP codes approximated using ZIP code tabulation areas (ZCTAs) defined by the U.S. Census Bureau; ‡ZIP-code-level data was used for Chicago Police District benchmarks.		

For this report, multiple ACS releases were used, all corresponding to 2023 as the most recent year of data available. The first was the 2023 1Y PUMS, which was used to estimate the age distribution of the entire population of Illinois in 2023. The second release used was the 2019-2023 5Y PUMS, which was used to 1) estimate the state-level age distribution for each racial group and 2) estimate reallocation factors for individuals reporting multiple races. The 5Y release was used instead of the 1Y release to achieve a larger sample size for those racial groups which had fewer individuals in Illinois. The third release used was the 2019-2023 5Y detailed table of race and ethnicity for each ZIP code in Illinois or any of 24 surrounding states within 400 miles of Illinois (Alabama, Arkansas, Georgia, Indiana, Iowa, Kansas, Kentucky, Louisiana, Michigan, Minnesota, Mississippi, Missouri, Nebraska, North Carolina, Ohio, Oklahoma, Pennsylvania, South Carolina, South Dakota, Tennessee, Texas, Virginia, West Virginia and Wisconsin). In general, the 2023 ACS datasets were used for both traffic stop and pedestrian stop benchmarks instead of the 2020 decennial census because, starting this year, ACS datasets (2019-2023) are more current than the decennial dataset (2020), and will be more current in the coming years until a new decennial update becomes available.

C.2.2. Data from Illinois Traffic Crash Reports

On behalf of this study, the Illinois Department of Transportation's Bureau of Data Collection in the Office of Planning and Programming provided a report of data extracted from Illinois SR 1050 traffic crash reports from 2019-2023. Information in the crash reports included the date and time of the crash, the location of the crash (latitude, longitude, city and county), the number of vehicles involved, the ZIP code of each driver's address, the type of roadway on which the crash occurred

and the type of law enforcement agency filing the report. This information was used to estimate driver benchmark populations for agencies with a sufficient number of usable reports available. In particular, the crash data were used to estimate the proportion of drivers originating from each ZIP code directly associated with an agency's jurisdiction as well as ZIP codes from the surrounding area.

C.2.3. Data from the Illinois Secretary of State

On behalf of this study, IDOT's Bureau of Data Collection in the Office of Planning and Programming requested and received a report from the Office of the Illinois Secretary of State with counts of licensed drivers in Illinois for each single year of age. The report was run in early March 2025. This was combined with ACS estimates of the population count of each age in Illinois (2023 1Y PUMS) to determine the proportion of individuals who are potential drivers based on having a driver's license as a function of age.

Appendix D. Technical Notes on Mileage Benchmarks

D.1. Overview

Mileage benchmarks for drivers and pedestrians were generated for each agency using the Replica data platform (<https://www.replicahq.com>), of which IDOT is a pre-existing client among other state and local government agencies using their traffic models.

This platform provides an elaborate simulation of traffic dynamics at a nationwide level with high granularity of detail. The model creates a virtual population in their daily commute, calibrated to match the national population at an individual level. The population data is generated based on U.S. census, while the detailed commute patterns are modeled using data from a diverse set of third-party data from public and private-sector sources.

D.2 Data Sources

Replica's model is derived from sources which fall into five broad categories:

- 1) Mobile Location Data
 - a. Location-based Services (cellphones apps of opted-in users)
 - b. Cellular Networks Data (cellphone interactions with tracking hardware)
 - c. Vehicle in-dash GPS Data (cellular hardware in vehicles for routing)
 - d. Point-of Interest Aggregates (detecting mobile devices in public venues)
- 2) Population Data
 - a. 5-Year American Community Survey (ACS)
 - b. Public Use Microdata Sample (PUMS)
 - c. Longitudinal Employer-Household Dynamics (LEHD)
 - d. Census Transportation Planning Products (CTPP)
 - e. The U.S. Census LEHD data
- 3) Built Environment Data (transportation network/land use/real estate maps)
 - a. OpenStreetMap (OSM)
 - b. General Transit Feed Specification (GTFS)

- c. Various proprietary sources (building footprints and floor areas/restaurant/stadiums/theaters/parking/hotel data)
- 4) Economic Activity Data (financial transactions data)
- 5) Ground Truth Data (actual observed counts)
 - a. Auto/Traffic Counts
 - b. Transit Ridership Counts
 - c. Bike and Pedestrian Counts

D.3 Construction of the benchmarks

The Replica model simulates individual trips taken during one typical weekday (Thursday) and one typical weekend day (Saturday), created in separate quarterly seasons. Each trip involves detailed information about its route, timing, demographics of the driver, mode of travel (car drive, bike ride, walking) etc. Geographically, each trip is broken into discrete street segments roughly 20-300 meters in length. Each segment is fully geolocated and can be readily determined if it falls within a given polygon of any complicated shape, such as the collection of ZIP codes in a jurisdiction. Thus, when setting the boundary of each jurisdiction as a specific polygon on the map, it is straightforward to sum the lengths of all the street segments on a trip's route that are located within the jurisdiction. This gives the distance traversed by the trip inside the jurisdiction. Summing these distances over all daily trips where the driver was of a particular race provides the mileage benchmark for a particular police agency and a particular race group on that day. Mileage driven on a typical day was the average of the miles driven over five typical weekdays and two typical weekend days.

Individuals of Hispanic ethnicity were automatically assigned to the Hispanic racial group. The drivers of the "two or more races" groups were partitioned into other racial groups according to a scheme similar to one the used in creating population benchmarks. First equal fractions reallocation factors for Illinois residents who self-identify as not Hispanic or Latino and more than one race, based on the 2023 5Y ACS PUMS, were found (representation of 54% White, 22% Black, 15% Asian, 8% AIAN, 1% NHOPI). These fractions were solely used while creating population benchmarks, when allocating "two or more races" drivers. Next, we found the prevalences of single-race groups among the agency's trips (their total mileage percentages across races). As the last step, we constructed the reallocation weights by multiplying the statewide fractions with the local prevalences then normalizing the weights. In this scheme, miles of multi-race drivers were allocated not only according to their state-level averages, but also according to the local presence of single-race drivers. In other words, American Indian fraction coming from multi-race drivers is allocated a little more into localities where single-race American Indians drive. (NOTE: In earlier versions of individual tables sent to agencies for comments, this allocation scheme was based only on an agency's prevalences of single-race drivers. These finer details influence mainly the benchmarks of rare minorities, AIAN and NHOPI, and mostly in smaller agencies.) We eliminated the trips where the driver belonged to the "Other race alone" group.

In calculating mileage benchmarks, we selected trips of three specific travel modes: walking (for the pedestrian benchmark), and biking and private auto (for traffic benchmarks).

We made a distinction between two types of road segments:

- 1) Special: belonging to interstate, federal and state highways

2) Regular: all other road segments.

In calculating the mileage benchmarks, we associated the special segments solely with the Illinois State Police. All other agencies were associated with regular road segments, meaning that their benchmarks were calculated as the total length of those parts of all the trips travelled within their jurisdictions only over regular road segments.

Finally, we collected the data for two most recent available seasons, Fall 2023 and Spring 2024. The mileage benchmark value for each race group was first calculated independently for the two seasons and then averaged over the two seasonal values.

D.4 Advantages

The Replica model presents two major advantages for calculating benchmarks over the methodology used in the past.

First, it is informed by large, relevant and diverse datasets that include tracking and observing various aspects of the commute and of the activities of the actors that create and evolve the traffic patterns.

Second is its extreme spatial resolution. While in calculating population benchmarks, we had to resort to treating jurisdictions as collections of ZIP areas or census block group areas that are often too large for smaller agencies, with mileage benchmarks and using Replica, the geography can be made as specific as desired.

For city- and county-level agencies, their jurisdictions were assigned as their natural city or county boundaries. State-level agencies other than the Illinois State Police were assigned a jurisdiction over the whole state (except on the special road segments). University-level agencies were assigned with their specific campus boundaries extended by 200-500 meters in all directions. Other agencies were assigned boundaries on a case-by-case basis. One particular agency, Round Lake Park Police, was assigned only the northern part of its village area. The spatial resolution of the Replica model also opens the possibility of placing different emphasis on particular locations within a jurisdiction. For example, a particular section of a street or a specific intersection could be given more weight when calculating the agency's benchmarks.

D.5 Limitations

The main limitation in using Replica software is the inability to independently check the accuracy of the model. We are simply the users of the platform. In particular, it is not clear how the model performs in localities where traffic density is very low, with very few cellphones to track, especially of drivers of the smallest minority groups.

Because of the enormous size of the simulation in terms of datasets involved and the many details of their software and hardware implementation, not all open to the public, the only way to establish the accuracy of the model is via communal effort and a consensus among various investigators to be reached over time, which is a process that has gained some momentum over the years.

Replica has provided some validation results from their clients on their data validation page (<https://www.replicahq.com/data-validations>). These reports involve comparing results from the Replica service with equivalent “ground truth” data obtained by the client via after-the-fact surveys, transit service tables, etc. The following examples are from the linked webpage:

- MassDOT used Replica for travel demand model improvements for the Boston Metro area. Actual counts from the region showed strong correlation with the modeled values.
- Florida DOT explored using Replica for travel demand model improvements. The Replica-generated results were compared with actual counts in the subject area with confidence.
- Puget Sound Regional Council looked into telecommuting rates and compared the Replica output with a travel survey. The results “looked similar.”
- Valley Metro, Arizona, found origin-destination flows on transit trips in Replica matched well with surveys on transit trips by origin districts.
- Metropolitan Council, Minnesota, showed Replica data matched well with surveys on transit trips.
- Atlanta Regional Commission, Georgia, has a proprietary synthetic model for their planning department. The Replica site has an extensive report on the comparisons between Replica, the ARC and the census showing the similarities between these sources across population, demand, traffic and transit models. Replica’s data matched well with surveys on transit trips.

Appendix E. Additional Notes on Illinois Law Concerning the Stops Study

The Illinois General Assembly has promulgated laws that require the collection and analysis of data on traffic stops by law enforcement agencies in the state. The statutes relating to the statistical analysis of traffic and pedestrian stops are found in the Compiled Statutes of the Illinois General Assembly, 625 ILCS 5/11-212, effective 6/21/2019. See also Public Act 101-0024.

Section 11-212 of the Illinois statute authorizes the “Traffic and pedestrian stop statistical study.” This section also requires that when a police officer stops an individual, a specific set of information is to be recorded. This information includes name, address, gender, race (six specific categories: White, Black or African American, Hispanic or Latino, Asian, American Indian or Alaska Native, and Native Hawaiian or Other Pacific Islander), the violation, vehicle information, date, time, location, search information, whether contraband was found, disposition of the stop (warning, citation or arrest — arrest recorded only for pedestrian stops⁷) and the name and badge number of the officer. This information is to be obtained whether the police officer makes a traffic stop or a pedestrian stop and either issues a citation or a warning (or arrest for a pedestrian stop). In addition, the length of the contact in minutes is to be recorded for traffic stops. These data items are recorded using the data collection form included in Appendix A. The law further specifies that the collected data are to be sent to the Illinois Department of Transportation by a specific date each year for the stop data collected in the preceding year.

⁷ The pedestrian stop data collection form in use during 2024 has a provision for recording an arrest. The traffic stop data collection form in use during 2024 does not provide a means of recording an arrest.

The Illinois Department of Transportation is further directed by statute to analyze the data and submit summary reports to the governor, the General Assembly and the racial profiling agency. IDOT is authorized to contract with an outside entity for the analysis of the data. That analysis is the purpose of this report. Moreover, the reporting entity is directed to scrutinize the data for evidence of “statistically significant aberrations.” An illustrative list of possible aberrations recorded in the statute include: (1) a higher-than-expected number of minorities stopped, (2) a higher-than-expected number of citations issued to minorities, (3) a higher-than-expected number of minorities stopped by a specific police agency and (4) a higher-than-expected number of searches conducted on Minority drivers or pedestrians.

The relevant statute, 625 ILCS 5/11-212 and subsection (a) provides that the law enforcement officer “...shall record at least the following....” The statute seems to suggest the current data collection form includes a minimum level of information, and leaves open the possibility of gathering additional information in the future.

There are a few additional data items that could be collected during traffic stops to enhance the analysis effort. Some additional data items might include: (1) arrest for DUI, (2) officer’s race (which has been shown to affect stop rates; see Ba et al. *Science*. 2021 Feb 12:696-702), (3) occurrence of a physical arrest in a traffic stop (the arrest outcome is currently included only in the pedestrian stop data collection form) and (4) latitude and longitude of the stop (which can be used to more precisely determine the benchmark for drivers or pedestrians but might need some technological changes). Additionally, there is a section in this report on estimating the accuracy of race designation by the stopping officer. The findings of that research suggest to us that obtaining the self-reported race from the driver may improve accuracy of reported race.