



**COMMONWEALTH OF KENTUCKY
FINAL REPORT ON
RISK LIMITING AUDIT
PILOT PROGRAM
MARCH 21, 2023**



KENTUCKY STATE BOARD OF ELECTIONS

Special Thanks to:

Jennifer Morrell

Michelle Forney

Trey Grayson

Anderson County Clerk, Jason Denny and Staff

Fayette County Clerk, Don Blevins Jr. and Staff

Henderson County Clerk, Renesa Abner and Staff

Johnson County Clerk, Sallee Holbrook and Staff

Kenton County Clerk, Gabe Summe and Staff

Madison County Clerk, Kenny Barger Jr. and Staff

Kentucky Secretary of State's Office

Election Systems & Software (ES&S)

Harp Enterprises Inc.

Table of Contents

Executive Summary	4
Introduction	
What is a Risk Limiting Audit (RLA)?	5
What is the Purpose of an RLA?.....	5
Different Types of RLAs	5-6
Development of Kentucky's RLA Pilot Program	
Initial Contact with Subject Matter Expert and Planning	7
Working Group Selection and Participants	7
Working Group Education	7
Decisions During Development	8
Preparation and Implementation of Kentucky's RLA Pilot Program	
Preparation Steps taken prior to Implementation	9
Implementation Delay	9
RLA Implementation Launch	9-10
RLA Process	10-11
Cost and Analysis of Kentucky's RLA Pilot Program	
Time and Cost of Kentucky's RLA Pilot Program	12
Results and Analysis of Kentucky's RLA Pilot Program	12
Lessons Learned and Conclusions	
Lessons Learned	13
Conclusions	13

Executive Summary

In 2021, the Kentucky General Assembly enacted House Bill 574, which included language establishing a Risk-Limiting Audit (RLA) Pilot Program. Due to a non-election year in 2021, the State Board of Elections (SBE) was unable to implement a pilot RLA in the first year. In the election year of 2022, a pilot program was developed and implemented by SBE, which further incorporated language passed by that year's House Bill 564 to require five percent of the Commonwealth's counties to participate in the pilot and for the results to be made available to the public. This report details how the pilot RLA was established, developed, and implemented.

In 2022, SBE and the six counties selected for participation continued to develop and refine the process of the pilot RLA. Following the 2022 general election, final implementation of the pilot RLA launched on Thursday, January 19, 2023. While this date was some two months after the general election, a lesson was learned as to the timing of any future RLA, as one pilot county (Fayette) experienced both an automatic recount in a state House race, as well as, a candidate-initiated recount in Circuit Court for a judicial race – both of which tied up the ballots needed for the pilot RLA.

During the development of the pilot RLA, SBE formed a working group consisting of a volunteer expert in the development and administration of RLAs, six (6) county clerks and their staff, members of the Office of the Secretary of State, a former Secretary of State and current Kentucky County Clerks Association (KCCA) representative, along with representatives for the two certified voting equipment vendors in the Commonwealth. Using institutional knowledge, in combination with procedural expertise, SBE sought to strengthen public trust in voting equipment accuracy and the aggregating of ballots in the Kentucky's election process.

The results of Kentucky's pilot RLA met the original expectations of risk and accuracy, however, much was learned about the intricacies, operations, and timelines required for the successful full-scale implementation of a truly meaningful Risk Limiting Audit. As a result of the pilot RLA, SBE recommends the expansion of the pilot program, including incorporating more counties into the process, and exploring ways to expand capabilities before further discussions about the establishment of a permanent Risk Limiting Audit Program in the Commonwealth.

Introduction

What is a Risk-Limiting Audit? What is its purpose?

A risk-limiting audit (RLA) is a post-election tabulation audit in which a random sample of voted ballots is manually examined for evidence that the originally reported outcome of the election is correct. As its name suggests, an RLA limits the risk of certifying a contest with the wrong winner (Note: Kentucky's pilot RLA was performed after certification).

An RLA gives statistical evidence that the machine-tabulated results are consistent with what a full hand count of ballots would reveal. Unlike fixed percentage audits, an RLA limits the risk that the wrong election result will be certified because of a tabulation error. They also allow jurisdictions to strategically allocate resources to check more ballots when needed in close contests, and fewer ballots in contests with wider margins.

What are the types of RLAs that can be done?

There are three main methods for conducting an RLA. Where and how ballots are scanned will be factored into the decision of which method(s) will work best.

	Ballot Comparison	Ballot Polling	Batch Comparison
Where do ballots get scanned?	Central location	Individual polling locations or a central location	Individual polling locations or a central location
Ballot information needed from the voting system	Cast vote records generated by the election management system (EMS)	None	Batch sub-total reports generated by the election management system (EMS)
Ballot batch size	The smaller the better. 100 ballots are a good average size	Not so large that a person cannot hold them on their own	The smaller the better
Ballots need unique, printed identifier	Yes	No	No
How are ballots validated?	RLA software	RLA software and/or manual tally sheets	RLA software and manual tally sheets

In a **ballot comparison audit**, specific ballots are identified and retrieved. The audit team examines the ballot and enters the voter markings for the audited contest(s) exactly the way they appear on the ballot. In some cases, hand-marked paper ballots may require the audit team to make decisions about voter intent. The RLA software compares the voter markings entered by the audit team to the cast vote record created by the voting system. The audit is looking for discrepancies between the two.

In a **ballot polling audit**, individual ballots are retrieved. The audit team examines the ballot and records the voter markings for the audited contest(s) on a tally sheet. Once all the ballots have been examined and voter markings recorded, the votes are totaled and the margin of victory for the winner(s) is compared to the margin of victory originally reported. The audit is looking for a similar or greater margin.

In a **batch comparison audit**, specific batches of ballots are identified and retrieved. The audit team examines the ballot and records the voter markings for the audited contest(s) on a tally sheet. Once all the ballots in the batch have been examined and voter markings recorded, the votes are totaled. The audit compares the manually recorded subtotals to the originally reported subtotals from the voting system. The audit is looking for discrepancies between the two.

Development of Kentucky’s RLA Pilot Program

Initial Contact with Subject Matter Expert and Planning:

In 2021, SBE began conversations with Jennifer Morrell, a renowned expert in election audits and risk limiting audits, who volunteered to assist Kentucky with its pilot RLA. Formal planning conversations started in January 2022 with the establishment of milestones and a calendar generated for the project that was finalized in February (See Appendix A). For the remainder of the year, meetings occurred between SBE, Ms. Morrell, and the various stakeholders chosen to participate in the working group.

Working Group Selection and Participants:

House Bill 564 mandated that at least five-percent of Kentucky’s counties participate in the pilot RLA. To encompass the Commonwealth’s election system, SBE selected two large, two medium, and two small counties based on voter registration. Both vendors of certified voting equipment in the state, Election Systems & Software (ES&S) and Harp Enterprises, Inc. (resellers of Hart Intercivic Voting Machines), were also invited to participate.

After seeking volunteer counties, SBE selected the following participants:

Fayette County	Large Jurisdiction with Hart Intercivic
Kenton County	Large Jurisdiction with ES&S
Henderson County	Medium Jurisdiction with Hart Intercivic
Madison County	Medium Jurisdiction with ES&S
Anderson County	Small Jurisdiction with Hart Intercivic
Johnson County	Small Jurisdiction with ES&S

The county clerks and their staff, in combination with the representatives of the voting equipment vendors, brought a full understanding of the tallying and aggregating of ballots. Vendors explained in detail their system’s capabilities in the working parameters of an RLA. Members of the current Office of the Secretary of State, along with former Secretary of State Trey Grayson, were selected to join the working group to provide institutional and functional knowledge throughout the process.

Working Group Education:

After the establishment of the working group, a better understanding of an RLA was needed. With the assistance of Ms. Morrell, the group was provided with resources to determine the most beneficial type of RLA for Kentucky. Through multiple sessions with all members, the working group developed a Standard Operating Procedure Guide that was invaluable in conducting the pilot RLA.

Decisions During Development:

1. After getting a well-rounded education in Risk-Limiting Audits, the group decided on the type of RLA the state would conduct. **Ballot Polling** was determined to be the best method available to suit Kentucky's election processes; because the Commonwealth does not perform centralized scanning of ballots or use Cast Vote Records in its local procedures, Ballot Comparison RLAs were ruled out. Batch Comparison was eliminated due to the need for additional software and the lack of specifically organized storage of voted ballots after the election in most counties. Ballot Polling was the optimal choice for the pilot program.
2. Next, SBE needed to decide which race in the general election to audit. The decision was made to audit the race for Kentucky's junior seat in the United States Senate, as this race was included on all ballots throughout the Commonwealth.
3. After much discussion, and out of abundance of caution, the working group decided not to perform the pilot RLA until after the thirty-day impoundment period for General Elections mandated under KRS 117.295(1). While the statute does allow for opening of ballot boxes in conjunction with an *actual* RLA, the working group did not want to open the door for any issues to be raised in prospective recount litigation around the fact that this was a *pilot* RLA, a situation not explicitly accounted for in the statute. Further, the pilot was intended to show proof of concept, that a ballot polling RLA would work for the Commonwealth, and to determine whether there was a need for additional piloting.
4. Risk limiting audits have gained support in the elections community based on the white paper entitled BRAVO: Ballot-polling Risk-limiting Audits to Verify Outcomes (Appendix E). Using the statistical formula contained within that white paper, the working group developed a spreadsheet tool that Kentucky could use in its pilot for determining an accurate sample size needed to achieve our statistical risk limit with accurate results after one round of ballot counting and produce a random sampling from all machines involved in the pilot. Another reason for this choice was the fact that the tool was developed at no cost, was easy to use, and adaptable for differing risk levels.
5. A final Standard Operating Procedure Guide for the entire process of a pilot RLA was then agreed upon and adopted by the working group.

Preparation and Implementation of Kentucky's RLA Pilot Program

Preparation Steps taken prior to Implementation:

Initially, the County Clerks were to complete the **Ballot Manifests** (see below for detailed description); however, SBE explored handling the process for a few of the participants for a time study. These counties reviewed and signed off on the work prior to the day of implementation. The remaining counties completed their manifests and submitted those to SBE to be used for the random sampling of ballots.

***Ballot Manifests** are lists produced with each county's name, a batch name (which for our pilot was the voting location name and if more than one scanner was used at that location, a corresponding number was assigned), and the number of ballots scanned into that device. These manifests are created for all scanners used for mail-in absentee, excused in-person absentee, no-excuse in-person absentee, and on Election Day. These manifests are input to the spreadsheet tool, which then generates the random sampling of ballots to be pulled for the RLA (see example on page 21).*

Leading up to the final RLA implementation, supplies were purchased by SBE including, colored paper, a set of 10-sided dice, and index cards used to make scanner identification cards. The final process was tested by extracting Fayette County and running a specific RLA for that county alone. The remaining five counties were run as another test by combining their ballots versus the results of their combined county totals (see examples in Appendix C and D).

Implementation Delay:

Initially, the working group decided to perform the pilot RLA immediately after the thirty-day impoundment period for which all voting equipment is to remain locked and sealed. This decision was made in part, to avoid adding any issues to any potential recounts that may occur following the general election, as the language of KRS 117.295(1) allows for the opening of ballot boxes pursuant to a risk limiting audit, but not a pilot RLA.

As fate would have it, Fayette County found itself entwined in both an automatic, Clerk controlled, state House recount, as well as, a candidate-initiated recount of a judicial election overseen by the Circuit Court. With the decision having been made to hold the pilot RLA following the conclusion of the recounts, the pilot was bumped from December of 2022 to January of 2023.

RLA Implementation Launch:

Following the completion of Fayette County's recounts, the working group met and established a final implementation date for the pilot RLA. Then, during a publicly broadcast SBE board meeting on Tuesday, January 17, 2023, the fifteen-digit seed for the random sample generator was set. This crucial step illustrated the randomization of the sampling of ballots to be pulled for the audit. Using

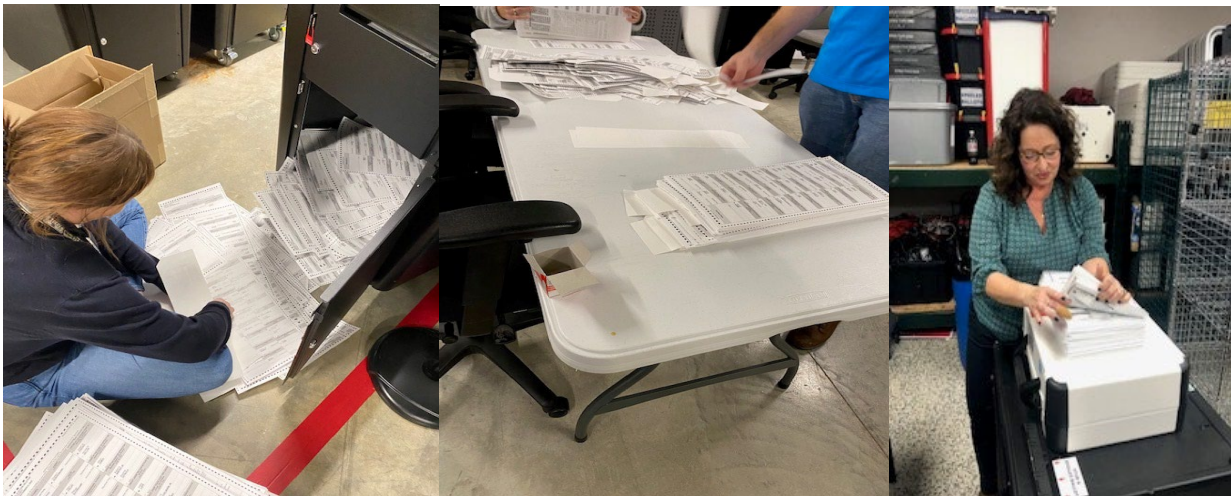
this process eliminated any prior knowledge of which ballots would be pulled from the locked machines (all counties' machines remained locked following the general election and maintained their original seals with the exception of Fayette, on account of their two recount proceedings).

The fifteen-digit seed was set when each participant in the board meeting rolled a 10-sided dice and those numbers were sequenced. The combination resulted in the number, 089268122289305, which was then entered into the random generator to produce the sampling of ballots for each county to pull during the pilot RLA. On January 18, 2023, SBE finalized the lists and tally sheets for the hand-counted results for each county to be listed on the final tabulation.

On Thursday, January 19, 2023, final instructions were given via video call at 10:00 AM ET, and thereafter, each participating county began retrieving ballots as indicated on their lists.



RLA Process:



The retrieving of selected ballots from each machine was to be conducted by one Republican party member and one Democratic party member. The procedure for processing is outlined below:

1. Ballot scanner containers unlocked and opened;
2. Ballots organized into stacks (except for Fayette, which already had its ballots organized due to recounts);

3. Random ballots pulled from stacks based on the list provided by the random ballot generator;
4. Immediately following the pulling of a ballot, placeholder sheet inserted into the stack;
5. Selected ballot placed into a folder or container to be tallied;
6. Remaining ballots placed back inside machine, machine relocked (after all ballots were retrieved from a scanner - multiple ballots were often required to be pulled from an individual machine);
7. All pulled ballots hand counted and tallied using only red ink pens to ensure transparency and establish protocol;
8. Results scanned and emailed to SBE for analysis.

During this process, it was understood that changes needed to take place in Kenton County and Madison County to complete the pilot RLA in one day. The method of ballot pulling was altered in these counties. Ballots were taken from the indicated scanners; however, rather than pulling all of the ballots out of the machine and then counting down to the specific ballots indicated by the RLA tool, the ballot retrieval teams estimated where in the pile to select ballots. It was discussed by members of the working group and determined that the sample size was the most important factor in the equation. Therefore, while ballots were still randomly drawn, they were not drawn with an exact numerical accuracy as occurred in other counties. This potential issue was discussed in a previous meeting of the working group, and the selection of a Ballot Polling RLA allowed for this because one goal of a ballot pulling RLA is to guarantee that auditors are not pulling ballots from machines that may have all been cast at the same time of day.

Because no rule was established in the development of the calculation tool to prevent the same ballot from being selected more than once, the established random sampling also chose three ballots twice. This resulted in inconsistencies as, Fayette County pulled one less ballot than the sampling size indicated, Madison County pulled their twice-selected ballot but only tallied it once, and Henderson County pulled their twice-selected ballot and tallied it twice.

Additionally, after the pilot RLA was completed, it was noted that six additional ballots were inadvertently not pulled in Madison County. Even with these missteps, the results derived still found a margin acceptably within the established risk-limit.



Cost and Analysis of Kentucky's RLA Pilot Program

Time and Cost of Kentucky's RLA Pilot Program:

Over the thirteen-month planning period, the Commonwealth incurred no direct costs. Operational meetings were conducted during regular business hours. On implementation day, SBE, Clerks, their staff, and paid citizen-auditors completed the pulling and tallying of a total of one-thousand sixty-five (1,065) ballots. The process across all counties took approximately ninety hours (number of people involved with total hours worked) to complete, with most of that time spent on ballot organization. Office supplies were nominal in cost.

Kenton County was the only jurisdiction to directly pay the individuals tasked with performing the pilot RLA. Kenton County used four citizen-auditors for 5.5 hours each, for a total of 22 hours. This group pulled and tallied 282 ballots at total combined cost of \$264. They would not have finished this process in the given timeframe had they not modified the way ballots were randomly pulled and had SBE and County Clerk staff not assisted them. This, unfortunately, does not paint a completely accurate picture of time and cost for the process in a large county (as has been stated, our other large county, Fayette, benefited in having their ballots pre-sorted on account of the previous recount proceedings, so a true cost analysis for large counties remains altogether inconclusive).

Supply costs were negligible for the pilot RLA. SBE spent approximately \$200, and counties were required to only provide red ink pens, paper clips, folders, and a container to place extracted ballots into.

Results and Analysis of Kentucky's RLA Pilot Program:

Even with last-minute changes and the failure to account for pulling and tallying the eight ballots, as described on page 11, the results indicated in Appendix C and D show that the random sampling of ballots statistically proved that the voting machines accurately accounted for the margin of victory for the winning candidate in Fayette, as a single county, and in the other five counties, for a combined margin within 1%.

Specifically, in Fayette County, Charles Booker (D) won with approximately 59.8% of the vote, with Rand Paul (R) garnering 39.0%. The random sampling pulled from their machines during the RLA indicated a margin of 59.14% Booker to 39.14% Paul - which places the RLA under the allowable 1% margin of error. Statistically, the results provide a confidence level that machines accurately reported the correct winner in this one county.

In the results of the other five counties, the certified results showed a Paul victory with approximately 61.59% to Booker's 37.25%. The random sampling (minus the seven ballots for the five-county audit described on page 11) still showed a 60% Paul victory with Booker taking 38% of the count. Statistically, the results once again provided a confidence level that the machines accurately reported the correct winner in this combination of counties.

Lessons Learned and Conclusions

Lessons Learned:

1. Waiting for the impoundment period to end alleviated litigation and recount concerns.
2. Extra checks must be put in place to ensure that all of the required ballots are pulled and tallied correctly, to prevent a second round of counting.
3. Larger retrieval teams are needed.
4. More working space is needed, dependent on a county's layout and number of machines to be pulled from.
5. ES&S counties have an even harder time with the organization of ballots due to the two different sizes of ballots that are scanned (preprinted ballots and the ballots generated from the Express Vote).
6. We need to have more discussion about implementation if this process were to be held pre-certification with the possibility of dealing with recounts and other litigation.
7. Legislation would need to be passed to push back certification to allow for an RLA to be conducted before the winning candidates were certified.
8. More discussion and planning would need to occur before RLAs were used to audit more local races.

Conclusions:

For the most part, Kentucky's first RLA Pilot Program went well and all participants learned a great deal about the process. Costs were held to a minimum; however, with the additional need for workers to handle the unforeseen challenges in the organization of ballots and the possibility of closer races increasing the number of ballots required to be pulled, the financial impact of the program will be guaranteed to increase.

The spreadsheet tool worked well in getting an accurate sampling of ballots that would yield similar results compared to the certified results; however, in the case of the five counties, not having the entire suggested sample size pulled, the outcome was slightly skewed. In reference to the Fayette County pilot, the outcome that the group saw was nearly identical to election night results. SBE and the working group in both environments learned some helpful lessons and discovered issues to be resolved in the future.

SBE's recommendation at this time would be to at least explore continuation of the pilot program to include more counties. The purpose of every step taken by all the stakeholders involved is to strengthen faith and confidence in Kentucky's Election System. RLAs can be a tool to build on the transparency of the Commonwealth's voting machines and help prove the accuracy of these products.

In closing, SBE would like to thank all the participants in this pilot for their commitment and assistance in carrying out this project.

Appendix A

Kentucky RLA Pilot Original Working Group Agenda & Milestones

March

- Establish which counties will participate in the working group and conduct a pilot in November
- Decide if you want to have vendors participate on all calls with the working group
 - I realized after our conversation that the vendor support staff in KY may have no knowledge of RLAs so it might be worthwhile inviting them to participate.
- Establish a date and time for standing working group calls
 - I recommend a 60-minute meeting held twice a month with meetings canceled in the lead up to each election.
- Consider assigning someone from BOE office to take notes which can be used to create a report at the end of the year for providing feedback to lawmakers.
- Jennifer will send out reading materials

April 1st Meeting - Goals for the Pilot and Basics for a Tabulation Audit

- Establish group norms
- Purpose of tabulation audits
- Fixed % vs RLA
- Goals of pilot
 - Safe learning environment
 - Creating a method that works in KY

April 2nd Meeting - Risk-Limiting Audits 101

May 17th - PRIMARY ELECTION

June 3rd Meeting - Methods for Conducting RLAs

- Ballot comparison
- Batch comparison
- Ballot polling
- Hybrid

June 4th Meeting - Ballot Accounting and Storage

- Mail/Absentee
- Early In-Person
- Election Day In-Person

July 5th Meeting - Ballot Manifest

- Format
- Populating
- Reconciling

July 6th Meeting - Voting System Requirements

- Standards for contest names
- CVRs and imprinting
- Batch subtotal reports
- Summary reports
- Good to have vendors on this call if they aren't already participating

August 7th Meeting- Ballot Reconciliation

- Mail ballots
- In-person Early Voting
- In-person Election Day

August 8th Meeting - RLA Procedure Guide Part 1

Sep 9th Meeting - RLA Procedure Guide Part 2

Sep 10th Meeting- Voting Works & Overview of Arlo

- Meet Voting Works team
- Overview of Arlo

Oct 11th Meeting-Final Ballot Storage and Organization

- Ensure everyone has what they need to create a ballot manifest in preparation for pilot
- Review reconciliation process
- Jennifer possibly on-site with some counties

Oct 12th Meeting- Final Checklist

- Ballot manifest
- CVR/Batch reports (zero report from L&A if available)
- Storage and staging
- Assigned roles and responsibilities
- Tally sheet

Nov 13th Meeting - BOE Staff and Voting Works (counties not needed for this call)

- Test login credentials
- Test file uploads
- Mock RLA

November 8th - GENERAL ELECTION

Date TBD - RLA Pilot

December - Report of Findings

(Meeting Session were combined as the process progressed.)

Appendix B

PILOT DRAFT PROCEDURES: Ballot Polling Risk-Limiting Audit

A risk-limiting audit (RLA) is a post-election tabulation audit in which a random sample of voted ballots is manually examined for evidence that the originally reported outcome of the election is correct. As its name suggests, an RLA limits the risk of certifying a contest with the wrong winner. An RLA gives statistical evidence that the machine-tabulated results are consistent with what a full hand count of ballots would reveal.

There are three main methods for conducting an RLA: ballot comparison, ballot polling, and batch comparison. For purposes of this pilot, the ballot polling method will be used.

In a **ballot polling audit**, individual ballots are retrieved. The audit team examines the ballot and manually records the voter markings for the audited contest(s) on a tally sheet. Once all the ballots have been examined and voter markings recorded, the votes are totaled and the margin of victory for the winner(s) is compared to the margin of victory originally reported. **The audit is looking for a similar or greater margin.**

Keep in mind that there is some tolerance for discrepancies depending on the margin of the target contest(s) and the risk limit that has been set.

As an overview, there are a few steps to conducting a risk-limiting audit and each of these stages will be outlined in greater detail:

1. **Pre-Election Preparation** - At the state level, this means determining which counties will participate in the pilot, setting the time frame for conducting the pilot, determining which contest(s) will be audited, and determining the risk limit for the audit. For the counties, pre-election preparation requires some planning as to how ballots will be stored and accounted for. This part of the process is key to successfully conducting a risk-limiting audit. Spending quality time in figuring out ballot storage and accounting will make the process of creating a ballot manifest and retrieving ballots go smoothly.
2. **Ballot Retrieval and Hand Tally** - After all ballots have been counted and a final results report generated, the counties will each submit a ballot manifest to the state as well as a summary results file. The state will combine the ballot manifests from all of the counties and determine how

many and which ballots must be retrieved for auditing and communicate that information to the counties. The counties will retrieve the designated ballots and will tally the votes cast for the audited contest(s) and then communicate those hand-tallied results back to the state.

3. **Audit Results** - The state will collect the results from all of the participating counties and determine whether the hand-tallied margin of victory for the audited contest(s) was equal to or greater than the margin for the contest(s) in the reported results.

Throughout these instructions the words “container” and “ballot storage container” are used to represent the ballot collection bin that is attached to the voting equipment used to scan ballots. All ballots will remain in the bin they were originally scanned into for this RLA pilot.

PRE-ELECTION PREPARATION

Completed prior to early voting and absentee ballot processing

Supplies

- Count verification form (for ballots scanned in-person)
- Ballot batch control sheet (for absentee or other ballots scanned centrally)
- Ballot storage container labels
- Ballot manifest

Staffing

- Staff assigned to verify ballot container labels are completed correctly
- Staff assigned to enter information from ballot container labels into the ballot manifest
- Staff assigned to validate data in the ballot manifest by performing a reconciliation

Ballot Accounting

- Review the ballot accounting practices conducted in your jurisdiction.

- This includes absentee ballot batch tracking forms, count verification forms, and chain of custody forms and procedures.
 - Should also include a review of how and where ballots will be stored.
- These practices are the foundation of your RLA paper trail and ensure ballots have not been lost or added as a result of human error. **They provide evidence the paper trail being audited is trustworthy.**
- Voting location count verification forms, absentee ballot batch tracking forms, and chain of custody logs should be reviewed, or audited, prior to an RLA or in conjunction with it.
 - This includes verifying the information from these forms gets transferred to container labels.

Ballot polling RLAs rely on jurisdictions locating a specific container of ballots and having an accurate piece count of the number of individual ballots in each container.

You must have a reliable system to:

- Verify the total number of ballots in each container independent of the voting system
- Assign a unique number to each container that ballots are stored in.

Accounting Process for Ballots Scanned at Polling Locations

1. Poll workers must complete a count verification form.
 - a. Reconciliation forms should validate the number of ballots issued and/or voters checked in equals the number of ballots scanned.
 - b. Provisional or emergency ballots that are segregated for scanning at a later time should be included somewhere on the reconciliation form.
 - c. Each reconciliation form should have a place for poll workers to record information explaining any discrepancies in the reconciliation to aid with additional research.
2. When ballots are transferred from the voting location to the central election facility, they should be locked and sealed in a storage container which is labeled with the following:
 - a. polling location name or number
 - b. scanner/container ID number (unique number assigned to the ballot storage container)
 - c. total number of ballots sealed in the container (taken from the count verification form or closing reports)
 - d. container security seal number
 - e. name or initials of the individual(s) who verified the quantity of ballots and sealed the container

The example below illustrates what a ballot storage container label might look like.

CONTAINER LABEL
Polling Location #: North Library Scanner #: 60
TOTAL BALLOTS: 984
Seal # A95162
Poll Worker Initials: TP
Poll Worker Initials: JM

Accounting Process for Ballots Scanned Centrally

1. Ballots (mail/absentee or walk-in absentee) should have a ballot envelope batch control form to account for ballots – from the time they are initially received in the facility to the point they are received in the scanning room.
 - a. Record the total number of ballots and who took custody on the batch control form each time they are moved or change hands.
 - b. Batch tracking should include any ballots removed from the original batch and sent for duplication.
2. Once received for scanning, verify the total number of ballot cards scanned equals the number of ballot cards in the batch (as indicated by the batch control sheet).
 - a. This may be done by verifying the start of day count and end of day count on each scanner to ensure that the number of ballots scanned in each batch is accurate.
 - b. Election workers tasked with scanning should be trained to take their time and ensure that only one ballot is scanned at a time.
 - c. Keep each envelope batch control form for the ballots scanned in each scanner (if multiple scanners are used) in one place and in successive order, i.e. date order.
3. Scanned ballots should be stored in a container that is labeled with:
 - a. a unique ID number
 - b. total number of ballot cards scanned - aggregate all batch control forms and verify that it matches the count from the scanner
 - c. container security seal number
 - d. name or initials of the individual(s) who verified the quantity of ballots and sealed the container

Absentee Ballot Container Label Example

Container #: A	
Batch Name	Total Ballots
01	20
02	37
03	63
04	42
TOTAL BALLOTS:	162
Seal # A95162	
Staff Initials: TP/JM	

Ballot Manifest

A ballot manifest is a useful and important document that makes reconciliation in preparation for your canvass much easier. The manifest form should be created prior to the election. Entries in the manifest are made during and after the election. A reconciliation of the manifest is done prior to the start of the audit. The ballots to be audited are randomly selected from the manifest once the audit begins.

- A ballot manifest is used to randomly select the ballots to be audited and indicates where the ballots are physically stored for easy retrieval.
- The ballot manifest should never be generated by the voting system.
- The ballot manifest will be a simple spreadsheet provided by the state.

Building a ballot manifest will follow a process similar to the following:

1. Verify the total number of ballots in each container matches what is recorded on the container label.
2. Seal the container (if not done already) and record the security seal number and individual(s) who verified the quantity and sealed the container.

3. Enter the container name, scanner ID, and total number of ballots into the ballot manifest spreadsheet. (Data can be written on a printed, paper version of the ballot manifest and then transferred to the spreadsheet.)
 - a. If the container name and scanner ID are the same, just enter the container name.
4. Place a checkmark on the container label to indicate the information has been transferred to the ballot manifest.

Example Ballot Manifest Spreadsheet:

County Name	Location/Container Name	Scanner ID	Total Number of Ballots
Madison	Central Count - Mail	20378	162
Madison	Central Count - Absentee Walk-ins	20379	15
Madison	Vote Center 1	101	150
Madison	Vote Center 1	102	206
Madison	Vote Center 1	103	172
Madison	Vote Center 2	121	100
Madison	Vote Center 2	122	58
Madison	Vote Center 3	123	194
Madison	Vote Center 3	124	167
Madison	Vote Center 4	131	25
Madison	Vote Center 4	132	61
Madison	Precinct 19	161	219
Madison	Precinct 24	182	132

Note: When ballots are scanned prior to election day, perform a daily reconciliation by comparing the totals from the ballot manifest to the cast vote record (CVR), the precinct counters, or some other sub-totals report generated by the voting system.

AUDIT PREPARATION

Completed after the election, prior to the audit.

State Responsibilities

- Determine the risk limit.
- Enter the risk limit into the audit spreadsheet.
- Set up target contest in the audit spreadsheet.

Room Preparation

- Designate a secure area for staging ballot storage containers for all scanned ballots that have been verified, sealed, and added to the ballot manifest.
- Consider whether access to an off-site facility is necessary to have a large enough space for all ballot storage containers along with tables to accommodate personnel performing the audit. If so, make the necessary security arrangements and create chain of custody documentation to handle the transport of ballots.
- For an official audit, consider whether there is space for observers or consider using AV equipment to complete all tasks on camera via Zoom or other conferencing tool.

Supplies

- Summary results report from the tabulation system
- Ballot manifest

Staffing

- Resources and staff should be assigned well before the RLA date.
 - Extra staffing may be required to help retrieve and review ballots.
- The number of ballots being audited will help determine the number of teams required.
- Teams of “auditors” are required; they should be bipartisan teams of two. For the pilot, the auditors may be county staff.
- All staff participating should be trained prior to the day of the audit.

Reconciliation

1. Finish tabulating all valid ballots that will be included in the audit.

2. Generate a summary results report from the voting system for all ballots that will be audited.
 - a. Be sure to include overvotes, undervotes, blank-voted contests, and valid write-in votes.
3. Verify the total number of ballots shown in the summary results report equals the aggregate number of ballots in the ballot manifest.
 - a. Research any discrepancies.
 - b. Correct when possible.
 - c. Do not force the numbers to match.
4. (Optional Additional Reconciliation)
 - a. Verify the total number of ballots in the ballot manifest equals the number of vote histories in the voter database.

Prepare for Ballot Retrieval

(Note: for purposes of the January 2023 RLA pilot, the steps below have already been completed).

1. Email the ballot manifest to the state.
 - a. Ballots are randomly selected for audit from the ballot manifest.
2. Email the summary results report to the state.
3. Publish the ballot manifest on the state or jurisdiction website. (NOTE: This will not be done for the pilot but is good practice when doing an official audit.)

CONDUCTING THE AUDIT

State Responsibilities

Hold a public meeting to generate a random number.
 Launch audit.
 Distribute ballot retrieval lists and tally spreadsheets to each county.

Room Preparation

- Generally, the audit should be conducted in the location where ballots are stored.
- Ensure there is enough room in the facility to accommodate both staff and observers while retrieving and examining ballots.

- If space is limited, consider retrieving ballots where they are stored and transferring the ballots selected for audit to an alternate location for the examination and recording portion of the audit. Be sure to follow all chain of custody protocols.
- Each audit team will need a table with room for the stack of ballots that needs to be reviewed, those that have been reviewed, and for the person marking the tally sheet.

Supplies

In addition to the ballots and ballot accounting documentation mentioned in the previous sections, you will need the following for the day of the audit:

- Chain of custody logs and extra seals for verifying sealed ballot containers (if required), resealing ballot containers, and recording new seal numbers
 - In some jurisdictions, the label on the container has been designed to double as the chain of custody log.
- Scissors (if needed to cut security seals on ballot containers)
- Voter intent guides/uniform definition of a vote (used by audit teams to make decisions about voter intent)
- Printer (for printing ballot retrieval lists, placeholders, and tally sheets)
- Pens for checking off ballots retrieved for audit
 - Pens and ballots in the same work area can be viewed as a security risk. Consider limiting any pens used during the audit to something unique, like green or red, that may not be recognized as a mark by the ballot scanner.
- Envelopes, tubs or folders to house the ballots selected for audit
- Colored paper to be used as cover sheets by the audit teams to identify ballots removed from storage containers for audit. (The quantity of paper needed can be determined by the projected sample size.)
- Paper clips to attach each cover sheet to the retrieved ballot
- Rubber fingers (optional)

Staffing

- Ballots should be retrieved, examined and tallied by an audit team consisting of at least two people.
- Depending on the number of ballots to audit, you may want multiple audit teams.

Print Retrieval Documents

The state will identify the ballots to be examined and will provide a list to each county along with the corresponding cover sheets.

Election staff should print the following documents:

- The list of individual ballots to be audited
 - List should include the corresponding location/container name
 - List should include a unique name or identifier for each audit team when using more than one team, e.g. Audit Team 1, Audit Team 2, etc.
 - Be formatted in a way that allows the retrieval teams to check off the ballots that have been retrieved for the audit
- Placeholder sheets (useful if ballots will need to be returned to their original storage location)
 - Printed on colored paper so that they stand out.
 - Identifying information from the retrieval list should be printed on each placeholder sheet.
 - Print 2 copies of each sheet. One will go in the container in place of the audited ballot and one will get paper-clipped to the audited ballot as a cover sheet.
- Hand tally sheet for the contest(s) being audited
 - Should include a place for the audit boards to initial or sign

Retrieve Ballots for Audit

Election staff overseeing the audit should determine before retrieval begins whether the ballots will be replaced in their original container after the audit or will be stored in a separate container.

- Provide each audit team with their corresponding ballot retrieval list, placeholder sheets and envelope/folder designated to hold the retrieved ballots.
 - Audit team retrieves ballots together using the steps outlined below.
 - Ballots should be kept together with the retrieval list in the designated audit envelope/folder.
 - The steps for retrieving ballots are repeated until all ballots have been retrieved and checked off the list.
1. Locate the container for the ballot(s) you are looking for.
 2. Verify the seals on the ballot storage container match the seals recorded on the chain-of-custody log (if a separate log is maintained).
 3. Organize the ballots into a neat pile first. (There is no need to ensure that ballots are all in the same direction or face up.)
 4. Locate the ballot(s) you are looking for within the batch by counting down to the appropriate ballot. (E.g. if the retrieval list says ballot 29, count down to the 29th ballot in the stack and retrieve that one.) **Note: more than 1 ballot may be selected for a given container.**
 5. Paperclip one of the pre-printed, placeholder sheet onto the ballot selected to provide identifying information. Place the other placeholder sheet in the batch of ballots where the ballot was removed.
 6. Place the retrieved ballot in the designated audit tub/envelope/folder.
 7. Check or initial the ballot retrieval list to indicate the ballot has been pulled for the audit.
 8. Repeat steps 4 through 7 for any additional ballots that must be retrieved from this container. Double check that you have retrieved all ballots for this container before moving on to step 9.

9. Re-seal the ballot container and record the seal numbers on the chain-of-custody log. (Some jurisdictions opt to wait until the conclusion of the audit to reseal ballot containers in the event additional rounds of auditing are required).

Ballot Review and Verification

- Hand-marked paper ballots may require the audit team to make decisions about voter intent.
 - Each auditing team should have a copy of the approved voter intent guidelines (uniform definition of a vote) to use when making that determination.
 - This ensures ballots are adjudicated during the audit the same way they were adjudicated for the election.
- Ensure all audit teams are aware which contest(s) need to be examined and recorded.
- Both audit team members should sit close enough to watch each other.
- If the voter's choices are not clear, and the audit team cannot agree on what constitutes a valid mark, they can indicate "disagreement" on the tally sheet.

Step	Ballot Polling
1	Auditor#1: Read out loud the voter selection(s) for each audited contest.
2	Auditor#2: Record the voter selections on the tally sheet and verbally repeat the choice.
3	Auditor#1: Verify has been recorded on the tally sheet matches what is marked on the ballot.
4	Audit Board: Initial the tally sheet indicating that the ballot has been audited.
5	Audit Board: Tally all the votes for each candidate once all ballots have been audited and recorded.
12	Audit Board: Return all tally sheets to the Clerk.

Conclude the Audit

1. Each county will add up the votes for each candidate from all tally sheets and submit them to the state using the designated spreadsheet.
2. Each county will scan their audit board tally sheets and email them to the state.

3. The state will aggregate all of the counties' tallies and determine whether the risk limit has been met. The state will communicate the result to the counties.
 - a. If the risk limit has been met, the state will produce an audit report and distribute that to the participating counties.
 - b. If the risk limit has not been met, the state will generate another list of ballots to be audited and the procedures for conducting the audit will be repeated.
4. State and/or counties should make the final audit reports available on their website. (NOTE: This will not be done for the pilot but is good practice when doing an official audit.)

Appendix E

BRAVO: Ballot-polling Risk-limiting Audits to Verify Outcomes

Mark Lindeman

Philip B. Stark¹

Vincent S. Yates¹

¹Department of Statistics
University of California, Berkeley

Abstract

Risk-limiting post-election audits guarantee a high probability of correcting incorrect electoral outcomes, regardless of why the outcomes are incorrect. Two types of risk-limiting post-election vote tabulation audits are *comparison audits* and *ballot-polling audits*. Comparison audits check some of the subtotals reported by the vote tabulation system, by hand-counting votes on the corresponding ballots. Ballot-polling audits select ballots at random and interpret those ballots by hand until there is strong evidence that the outcome is right, or until all the votes have been counted by hand: They directly assess whether the outcome is correct, rather than assessing whether the tabulation was accurate. Comparison audits have advantages, but make large demands on the vote tabulation system. Ballot-polling audits make no such demands. For small margins, they can require large samples, but the total burden may still be modest for large contests, such as county-wide or state-wide races. This paper describes BRAVO, a flexible protocol for risk-limiting ballot-polling audits. Among 255 state presidential contests between 1992 and 2008, the median expected sample size to confirm the plurality winner in each state using BRAVO was 307 ballots (per state). Ballot-polling audits can improve election integrity immediately at nominal incremental cost to election administration.

1 Introduction and background

Voting systems in use in the U.S. are known to be vulnerable to misinterpretation of voter intent, mechanical failure, misconfiguration, and deliberate subversion [Secretary of State, 2007; McDaniel et al., 2007]. The vote totals they produce should not be assumed to be accurate, nor even sufficiently accurate to produce correct outcomes. The possibility of gross error is not merely theoretical: In March 2012, a routine post-election audit

in Palm Beach County, Florida led to a recount that found the wrong winners had been declared in two city council contests [Sorentrude et al., 2012]. Wrong outcomes in U.S. elections are thought to be rare, but the true incidence is unknown.

Risk-limiting post-election audits have been widely endorsed as a means to check whether voting systems find the correct outcomes [ElectionAudits.org, 2008]. A risk-limiting audit checks some voted ballots (or voter-verifiable records) in search of strong evidence that the reported election outcome was correct. If the reported outcome is *incorrect*, a risk-limiting audit has a large, pre-specified minimum chance of leading to a full hand count, which reveals the correct outcome. The *risk limit* is the maximum chance that an audit of an election with an incorrect outcome will *not* lead to a full hand count. A burgeoning literature (e.g., Checkoway et al. [2010]; Hall et al. [2009]; Stark [2008a,b, 2009a,c,b, 2010]; Benaloh et al. [2011]; Lindeman and Stark [2012]) explores methods and concomitants for risk-limiting audits.

Two types of risk-limiting audits have been proposed: *comparison audits* and *ballot-polling audits* (Lindeman and Stark [2012]; Johnson [2004] makes an analogous distinction). Comparison audits check outcomes by comparing hand counts¹ to voting system counts for clusters of ballots². In a *ballot-level* comparison audit, each cluster is a single ballot. In a *batch-level* comparison audit, each cluster comprises multiple ballots: for instance, all the ballots cast in a particular precinct. In either case, a risk-limiting comparison audit examines randomly selected clusters to assess whether the voting system subtotals are sufficiently accurate to confirm the reported out-

¹We construe “counting” the votes in a cluster of ballots to have two parts: (i) Interpreting individual ballots or records in the cluster to identify valid votes. (ii) Tallying the valid votes on those ballots. A hand count of a cluster that consists of a single ballot cast in a plurality contest yields a one or a zero for each candidate, a one if the ballot shows a valid vote for the candidate or a zero if not.

²*Transitive audits* use counts from a secondary system, as briefly discussed below.

come, at the specified risk limit. If not, a full hand count is conducted.

In contrast, a ballot-polling audit does not use voting system subtotals. Instead, it examines the votes on a random sample of individual ballots.³ When and if the vote shares in the sample provide sufficiently strong evidence to confirm the reported outcome, the audit stops.

Comparison risk-limiting audits have important benefits, but they can be hard to implement. The most efficient comparison audits—those at the ballot level—require information that current certified vote tabulation systems do not provide: interpretations of individual ballots (cast-vote records or CVRs) that can be associated with the paper ballots they purport to represent. Preparing for a comparison audit requires exporting a complete list of auditable subtotals from the voting machine and the ability to retrieve (and hand count) the ballots corresponding to each subtotal. It sometimes requires translating voting-system reports that were intended for printing into formats that software can read, a labor-intensive, error-prone process. And it requires summing the auditable subtotals to verify that they match the reported contest totals.

To conduct a ballot-level comparison audit of a current certified system generally requires *transitive auditing* (see Calandrino et al. [2007], although they do not call it by that name). A transitive audit uses a secondary system to make a CVR for each ballot in a way that allows the CVR to be matched to the ballot it purports to represent. So far, transitive audits have relied on digital images of all the ballots cast in the contest (produced by the voting system or by rescanning the ballots); creating CVRs from those images using a combination of software and hand labor; and maintaining a mapping between the physical ballots and the CVRs by keeping the ballots in the order in which they were scanned or by marking individual ballots with an identifier. If the outcome according to the secondary system matches the voting system’s reported outcome, then a risk-limiting audit of the secondary system outcome is a risk-limiting audit of the official system, at the same risk limit: The reported outcome is incorrect if and only if the secondary system’s outcome is incorrect. Hence, if the reported outcome is incorrect, the chance that the audit will lead to a full hand count of the physical ballots is at least 100% minus the risk limit used in the transitive audit. While the transitive audit strategy reaps the statistical efficiency advantages of ballot-level comparison audits, the logistical obstacles may be substantial, especially in large jurisdictions.

Ballot-polling audits, in contrast, require virtually nothing from the voting system and no extra preparation

beyond the sorts of ballot accounting that local election officials generally do. They are an immediate option in any jurisdiction with an auditable paper trail. Moreover, for most margins of victory, they typically require examining substantially fewer ballots than batch-level comparison audits to attain a particular risk limit. Thus, we propose ballot-polling audits as a practical way to conduct risk-limiting audits of selected contests—such as presidential elections—immediately.

When the reported count is accurate or nearly accurate, the amount of auditing to be done in a risk-limiting ballot-polling audit is less predictable than the amount to be done in a risk-limiting comparison audit of the same contest. Therefore it is advantageous to be able to expand the audit flexibly depending on early results, using *sequential sampling*. The best-known sampling methods use a single sample and cannot easily be extended to risk-limiting sequential audits. The method we describe here is based on a robust sequential sampling method that can draw one or more ballots at a time, as the auditors may desire.

2 Previous work

Johnson [2004] presents the first ballot-polling election audit of which we are aware. Johnson calls his ballot-polling method a “statistical recount.”⁴ To conduct a statistical recount, an initial random sample of n ballots is drawn without replacement. If the reported winner’s margin in the sample is larger than a “critical margin” (essentially half the width of a confidence interval at some predetermined confidence level), the audit stops and the election is “validated.” If the sample margin is negative and larger in absolute value than the critical margin—giving evidence that the loser actually won—the audit stops and the election is “invalidated.” If the absolute value of the sample margin is smaller than the critical margin, the audit expands the random sample in stages. At each stage, a random sample of n ballots is drawn without replacement from the as-yet-unexamined ballots. The cumulative sample margin is compared with a critical margin (updated to reflect the cumulative sample size), until the election outcome has been validated or invalidated, or all ballots have been examined. Johnson asserts that the method can be used to validate election outcomes with a specified level of “statistical confidence.” However, the distribution of the maximum of a collection of (normalized) hypergeometric random variables with nested samples is not (normalized) hypergeometric, which the method assumes implicitly. This could make the chance that the method mistakenly validates an

³Batch-level ballot-polling audits are possible, but less efficient than all of the approaches discussed here. They do not seem to have any advantage over ballot-level ballot-polling audits.

⁴Johnson [2004] also discusses a ballot-level comparison audit, which he calls a “statistical error count.”

incorrect outcome much larger than 1 minus the confidence level. Hence, a “statistical recount” is not risk-limiting.

Simon and O’Dell [2006] propose to verify federal election outcomes through “Universal Precinct-based Sampling” (UPS). UPS entails drawing and hand-counting a 10% sample of the ballots cast in each precinct, and each other “pile” of ballots—preferably on election night. Simon and O’Dell assert that these samples (treated as simple random samples) combined provide vote share estimates “within 1% of an accurate vote count” with 99% statistical confidence in competitive House contests—and even more accurate estimates in larger contests. Discrepancies of more than 1% in a reported winner’s vote share would trigger further scrutiny, i.e., further sampling or perhaps a full hand count. UPS is not risk-limiting, for a variety of reasons.⁵ UPS emphasizes immediacy, whereas BRAVO emphasizes rigor, limiting risk, and efficiency: BRAVO requires dramatically smaller samples in many large contests, as we demonstrate below.

Lindeman and Stark [2012] sketch a risk-limiting ballot-polling audit method similar to BRAVO. BRAVO improves on that method in a number of ways: (1) Auditors are never required to escalate to a full hand count based on a test statistic. (2) That method applies only to vote-for-one contests with a majority winner, while BRAVO works with vote-for-one or vote-for- k plurality, majority, or supermajority contests. (3) Because BRAVO uses pairwise comparisons instead of “pooling” votes for reported losing candidates, it can be substantially more efficient.

There is by now a large literature on sequential testing, but to our knowledge, developments since Wald [1945] are not helpful for risk-limiting ballot-polling audits. For instance, there are sequential procedures for identifying the “best” multinomial category by sampling until one category has occurred substantially more frequently than any other [Gibbons et al., 1977; Ng and Panchapakesan, 2007; Ramey and Alam, 1979]. These are not suited to risk-limiting audits: They give a final P -value associated with the “best” category, but provide no mechanism to

⁵(1) As described, the chance UPS leads to a full hand count when the outcome is incorrect could be zero, so the method is not risk-limiting. (2) The fixed 10% sample size does not ensure that the standard error of the sample margin will be small enough to draw a conclusion about the sign of the true margin. UPS could fall far short of its nominal 99% confidence level. (3) UPS uses a stratified sample, so Simon and O’Dell’s margin of error calculation is at best approximate. An exact calculation would require knowing how many ballots are in each “pile,” and we suspect that the calculation would be intractable, although conservative bounds might be obtained. (4) One sampling method Simon and O’Dell propose—examine every tenth ballot, starting with the n th ballot, where n is a random integer between 1 and 10 inclusive—further complicates the risk calculation compared with taking independent simple random samples from each “pile.”

assure that this P -value is below the desired risk limit α .

Rivest and Shen [2012] propose a Bayesian approach to election auditing that can be used with ballot polling. It is not clear whether this approach is risk-limiting, nor how the risk limit might depend on the prior probability distribution.

3 Overview of BRAVO

We assume that the election generated an indelible, voter-verifiable audit trail, which might consist of a combination of voter-marked paper ballots and voter-verifiable paper records (VVPRs). We refer to such records as “ballots,” even though they might not be paper ballots. We also assume that a *compliance audit* [Benaloh et al., 2011; Lindeman and Stark, 2012; Stark and Wagner, 2012] has provided convincing evidence that this audit trail is adequately accurate to reflect who actually won; otherwise, a full hand count—the recourse of risk-limiting audits when the outcome is in doubt—would not necessarily reveal the actual winner.

The ballot-polling risk-limiting audit we describe here, in its simplest form, can be described as a loop:

1. Randomly select a ballot and examine it for a valid vote. (Ballots can be selected more than once; they are counted as many times as they are selected.)
2. If the ballots selected so far provide sufficiently strong evidence that the reported outcome is right, stop the audit and accept the reported outcome.
3. If, after examining a large number of ballots, the audit has not given strong evidence that the reported outcome is right—or if it provides strong evidence that the reported outcome is wrong—conduct a full hand count of all the ballots to determine the correct outcome.
4. If neither condition 2 nor 3 holds, return to step 1.

For the audit to be risk-limiting, condition 2—the criterion of “sufficiently strong evidence”—must be defined correctly before the audit. Condition 3 can be less sharply defined; what matters is that the audit ends either by meeting condition 2 or by conducting a full hand count.

The arithmetic involved in BRAVO is simple enough to do on a four-function calculator. The ballotPollAuditTools web page (<http://statistics.berkeley.edu/~stark/Java/Html/ballotPollTools.htm>), or custom tools, can be used to help with the sampling logistics and calculations involving multiple candidates.

A sequential audit of one ballot at a time can reduce the expected number of ballots that must be inspected

when the apparent outcome is correct, because the audit can stop as soon as there is strong evidence that the apparent winner(s) actually won. However, it may be more efficient to select and inspect more than one ballot at a time, especially to audit contests that cross jurisdictional boundaries. Accordingly, we suggest methods to conduct the audit in groups or “stages” of ballots so that jurisdictions can audit concurrently.

4 Applicability of BRAVO

We address simple measures, measures that require a super-majority, and plurality contests, including contests such as city council contests and school board contests in which each voter may select more than one candidate. The method we present is not suitable for auditing ranked-choice voting (RCV) or instant-runoff voting (IRV), because whether an RCV or IRV ballot ultimately is counted for a particular candidate typically depends on the other ballots and cannot be determined in isolation.

There are C candidates in the contest; there are $k \geq 1$ reported winners. In a simple measure, $C = 2$ and $k = 1$. In a city council contest, we might have $C = 10$ and $k = 5$. For plurality contests, determining whether the k reported winners among the C candidates in the contest really won amounts to determining whether each of the k reported winners received more votes than all of the $C - k$ reported losers. As elaborated in [Stark 2008a], it is natural to take the null hypothesis to be that the outcome is wrong; in this case, that means that at least one of the reported losers received at least as many votes as at least one of the reported winners. To reject that hypothesis is to conclude that all k candidates reported to have won really did win. BRAVO tests this hypothesis by testing the union of the hypotheses that apparent loser ℓ received at least as many votes as apparent winner w , for all pairs (w, ℓ) of apparent winners and apparent losers. Testing this composite hypothesis at significance level α limits the risk to α , if a full hand count is required whenever the test does not (eventually) reject all of the pairwise hypotheses.⁶

For contests that require an outright majority, testing whether the apparent winner truly won amounts to testing the null hypothesis that no more than half of the valid votes are for the reported winning position. To reject that hypothesis is to conclude that the apparent majority position really received a majority ($> 50\%$) of the votes. As before, testing at significance level α limits the risk to α , if a full hand count is required whenever the test does not ultimately reject the null hypothesis. The test can be

⁶As elaborated below, the fact that counting continues until *all* the pairwise hypotheses are rejected results in a very simple test that does not need any statistical adjustment for multiplicity.

generalized to a supermajority requirement by replacing 50% with the appropriate proportion.

There are many possible approaches to ballot-polling risk-limiting audits. The approach we present here is based on Wald’s Sequential Probability Ratio Test (SPRT) [Wald 1945]. In Wald’s method, sampling continues until either the null hypothesis or the alternative hypothesis is accepted. In BRAVO, the null hypothesis that the reported election outcome is wrong is never be accepted based on anything short of a full hand count. For plurality contests with more than two candidates, BRAVO performs several SPRTs in parallel using the same sample, to test whether any loser is in fact a winner. Methods sharper than BRAVO could be devised, although we doubt that the computations would be as simple.

5 Sampling framework and notation conventions

We assume that the contest under audit is a plurality contest with C candidates and $k \geq 1$ winners. The set $\mathcal{W}$ denotes the the apparent winners and the set $\mathcal{L}$ denotes the apparent losers. The set $\mathcal{W}$ contains k candidates. (Generally, a voter is permitted to vote for up to k candidates if the contest has k winners, but that restriction matters only to determine whether a given ballot has valid votes in the contest: A ballot marked for too many candidates would be treated as an overvote rather than as a valid vote for each of those candidates.) Measures that require a majority or supermajority are treated incidentally in section 6. We assume that every candidate in $\mathcal{W}$ was reported to have more votes than every candidate in $\mathcal{L}$; that is, the reported winner with the fewest votes is not apparently tied with the reported loser with the most votes.⁷

The true, unknown proportion of ballots that record votes for candidate c is π_c . We consider ballots that do not show a valid vote in the contest (for instance, ballots with overvotes) to be votes for a fictitious candidate 0; candidate 0 is neither in $\mathcal{W}$ nor in $\mathcal{L}$. Thus, the vector $\boldsymbol{\pi}_c \equiv (\pi_c)_{c=0}^C$. The reported proportion of ballots that record votes for candidate c is p_c . Typically, instead of reporting p_c directly, election officials report s_c , the fraction of valid votes cast for candidate c ; the denominator is not the total number of ballots but rather the total number of reported valid votes. (In a vote-for-one contest, s_c typically exceeds p_c , because some ballots do not show a vote for any candidate in the contest.) The relationship between them is

$$s_c \equiv \frac{p_c}{\sum_{j=1}^C p_j}. \quad (1)$$

⁷Statistical sampling is not well suited to confirm that two candidates received exactly the same number of votes.

To check whether the set $\mathcal{W}$ is really the set of winners, we will draw ballots uniformly at random and interpret the votes on those ballots manually. The probability that a ballot selected at random records a vote for candidate c is π_c . The apparent winners really won if

$$\min_{c \in \mathcal{W}} \pi_c > \max_{c \in \mathcal{L}} \pi_c. \quad (2)$$

Equivalently,

$$\pi_w > \pi_\ell, \forall w \in \mathcal{W}, \ell \in \mathcal{L}. \quad (3)$$

We assume that the value of π_0 , the proportion of ballots that do not show a valid vote, is irrelevant to the outcome.⁸ Thus, the parameter of interest is $\boldsymbol{\pi} \equiv (\pi_c)_{c=1}^C$, the vector of conditional probabilities that a ballot bears a valid vote for candidate c given that it bears a valid vote for some candidate. Henceforth, subscripted values of $\boldsymbol{\pi}$ (e.g., π_i) refer to values in $\boldsymbol{\pi}$, not in $\boldsymbol{\pi}_c$. (However, $\boldsymbol{\pi}_c$ matters for expected sample sizes, especially in the vote-for-one case.)

6 Special case: Contests with two candidates (and majority contests)

To begin, consider the case of a vote-for-one contest with two candidates: the reported winner, w , and the reported loser, ℓ . (A vote-for-one majority contest in which candidate w is reported to have received more than a majority of valid votes can also be audited as described in this section, by treating votes for all other candidates as a vote for ℓ ; however, it is more efficient to keep the losing candidates separate, as proved below.) We want to test the null hypothesis that $\pi_w \leq 1/2$. To incorrectly reject the null hypothesis is to keep the reported outcome when it is wrong, so the significance level α is the risk limit.

We want to avoid unnecessary full hand counts, especially when the original vote totals are correct or very nearly so. By the same token, if the audit produces inconclusive results after extensive sampling—or if it produces strong evidence that the reported outcome in fact is incorrect—then we should proceed to a full hand count. Several mechanisms for triggering a full hand count are possible.⁹ Here we assume that the auditors optionally set M , a maximum number of ballots to audit, beyond which they intend to proceed to a full hand count if the audit has not produced sufficiently strong evidence that the reported outcome is correct. The value M is a soft

limit: It is acceptable to continue the audit after M has been reached¹⁰ and it is certainly acceptable to proceed to a full hand count before M is reached.

Given s_w (the reported proportion of valid votes cast for w) and the risk limit α , the following steps define a risk-limiting audit based on Wald's SPRT (Wald [1945]):

1. Set the test statistic $T = 1$, and the cumulative sample size $m = 0$. (Optionally, pick M , the maximum number of ballots to audit before requiring a full hand count.)
2. Draw a ballot at random from the set of ballots that include the contest under audit. (A ballot can be drawn more than once.) If the ballot shows a vote for w , multiply T by $s_w/0.5$. If the ballot shows a vote for ℓ , multiply T by $(1 - s_w)/0.5$. (If the ballot does not show a valid vote, T is unchanged, or equivalently, multiplied by 1.) Increment m .
3. If $T \geq 1/\alpha$, stop the audit: The reported outcome is confirmed at risk limit α .
4. If $m = M$, or at the auditors' discretion, stop and perform a full hand count; the outcome according to the hand count replaces the reported outcome.
5. If neither (3) nor (4) holds, go back to step (2).

Because $s_w/0.5 > 1$ and $(1 - s_w)/0.5 < 1$, T increases each time the audit finds a vote for w and decreases each time it finds a vote for ℓ . (The critical value $1/\alpha$ follows from equation (3.24) in Wald [1945], with $\beta = 0$.) If the reported vote shares are correct, values of T close to 0 are unlikely: The chance of observing $T \leq p$ is no greater than p .¹¹ Thus, very small values of T provide a rationale for performing a full hand count before M is reached.

If the m ballots drawn so far show m_w votes for w and m_ℓ votes for ℓ , then

$$T = (2s_w)^{m_w} (2 - 2s_w)^{m_\ell}. \quad (4)$$

Suppose that the reported vote shares are correct. Table 1 gives estimated means and selected percentiles for the expected number of ballots with valid votes to inspect in order to confirm reported outcomes at risk limit $\alpha = 10\%$, assuming that the reported vote shares are correct and that every audit proceeds until the reported outcome is confirmed.¹² The simulations summarized in Table 1 sequentially draw a ballot at random and multiply

⁸If the voting rules dictate otherwise, the method here can readily be adjusted.

⁹Lindeman and Stark [2012] describes the use of β , which strictly controls the chance of a full hand count on the condition that the winner's vote share is accurate within a certain specified tolerance, such that the reported winner in fact won. Wald's SPRT (Wald [1945]) uses both α and β ; in the exposition here, implicitly we have set $\beta = 0$.

¹⁰Auditors may strongly prefer to continue sampling if the audit is close to reaching the risk limit.

¹¹This follows from inequality (3.17) in Wald [1945].

¹²If the reported vote shares are even approximately correct, then the expected value of the multiplier in step (2) is positive, so $p(T \geq 1/\alpha) \rightarrow 1$ as $m \rightarrow \infty$.

the test statistic T by $s_w/0.5$ if the ballot shows a vote for the winner or by $(1 - s_w)/0.5$ if it shows a vote for the loser. This is repeated until the test statistic $T \geq 1/\alpha$; the number of draws required is recorded. The entire process was repeated 10^7 times for each margin considered. Table 1 reports the sample mean and empirical percentiles of the 10^7 simulations for each margin considered.

Most of the mean sample sizes reported in the table are modest in the context of a large jurisdiction with hundreds of thousands or millions of ballots cast. The mean sample size varies roughly as the inverse square of the margin, $\pi_w - \pi_\ell$; it increases sharply as π_w approaches 50% from above. Notice that the number of ballots to be audited is unpredictable: A small fraction of audits inspect several times more ballots than average.

Table 1 also reports an estimate of the Average Sample Number (ASN) [Wald, 2004], the expected number of draws required either to accept or to reject the null hypothesis. The entries in the table are computed on the assumption that the reported vote totals are correct. Our estimate is

$$ASN \approx \frac{\ln(1/\alpha) + z_w/2}{p_w z_w + p_\ell z_\ell}, \quad (5)$$

where $z_w \equiv \ln(2s_w)$ and $z_\ell \equiv \ln(2 - 2s_w)$.¹³

Let $s_\ell \equiv 1 - s_w$ and $x = s_w - s_\ell$, the reported margin as a proportion of valid votes. Suppose $p_0 = 0$. Then the denominator of expression (5) becomes

$$\begin{aligned} & \frac{(1+x)\ln(1+x) + (1-x)\ln(1-x)}{2} \\ &= x^2/2 + x^4/12 + x^6/30 + \dots \end{aligned} \quad (6)$$

The first term dominates when x is small, so

$$ASN \approx 2 \ln(1/\alpha)/x^2. \quad (7)$$

That is, the ASN is roughly inversely proportional to the margin squared.

When there is a positive probability p_0 that a draw will give an invalid ballot or a ballot with an overvote, that “thins” the rate of drawing informative ballots, increasing ASN by the factor $1/(1 - p_0) > 1$. For instance, if 10% of ballots contain invalid votes, then the ASN increases by $1/(1 - 0.1) - 1 \approx 11.1\%$.

When $\pi_c \neq p_c$, expression (5) holds if p_w and p_ℓ are replaced with π_w and π_ℓ , provided the winner’s true share of valid votes, $\pi_w/(\pi_w + \pi_\ell)$, is a bit greater than $(s_w + 0.5)/2$.¹⁴ Otherwise—if the actual margin is not

¹³This estimate is derived from expression (3.57) in [Wald, 2004], setting B and $L(\theta)$ to 0. (Expression (4.8) in [Wald, 1945] is equivalent.) The additional term $z_w/2$ allows for the fact that the final value of T ordinarily exceeds $1/\alpha$.

¹⁴The denominator of equation (5) vanishes if the winner’s true share of valid votes equals $\ln(1-x)/(\ln(1-x) - \ln(1+x))$. This value exceeds $(s_w + 0.5)/2$ by less than 0.001 for $s_w \leq 0.642$.

more than half the reported margin—the ASN is undefined: The audit may have a positive probability of continuing indefinitely, even if the reported outcome is correct. Of course, the auditors may elect to perform a full hand count at any time, and a full hand count might be prudent if the true margin is half or less than half of the reported margin, even if the outcome is correct.

7 Plurality contests with $C > 2$ candidates

Many contests have more than two candidates. Vote-for-one plurality contests in which one candidate receives a majority of the vote can be audited using the method described in the previous section by “pooling” votes cast for all the reported losers as if they were votes for a single hypothetical loser ℓ . Here, we present a more flexible and efficient approach that works for k -winner plurality contests with $k \geq 1$.

Consider each pair of an apparent winner and an apparent loser, (w, ℓ) , $w \in \mathcal{W}$, $\ell \in \mathcal{L}$. The approach tests the $k(C - k)$ null hypotheses $\{\pi_w \leq \pi_\ell\}$, using the same sample but different test statistics $\{T_{w\ell}\}$. Each pairwise comparison relies only on valid votes cast for the candidates w and ℓ . The multipliers that update $T_{w\ell}$ after each ballot is drawn depend only on the reported votes for those two candidates.

In the most general case we consider, each ballot may have valid votes for zero or more reported winners, and for zero or more reported losers. Candidate w really beat candidate ℓ if and only if more than half the ballots that show a vote for either w or ℓ —but not both—show votes for candidate w . (Ballots that show votes for both w and ℓ , or for neither w nor ℓ , favor neither candidate.) Testing whether w really beat ℓ won therefore amounts to testing whether candidate ℓ got half or more of the votes on such ballots.

Let $s_{w\ell} \equiv s_w/(s_w + s_\ell)$ be the fraction of votes w was reported to have received among ballots reported to show a vote for w or ℓ or both. The value of $s_{w\ell}$ can be calculated from standard reported election results, whereas the fraction of ballots reporting votes for w but not ℓ cannot be. Hence, our test for the pair (w, ℓ) should not require knowing that fraction.

Suppose that w reportedly beat ℓ , so that $s_{w\ell} > 0.5$, and suppose that a fraction s of ballots reportedly showed votes for both w and ℓ . Then the fraction of ballots reported to have votes for w but not ℓ among those that show votes for w or ℓ but not both is

$$\frac{s_w - s}{s_w + s_\ell - 2s} > s_{w\ell}. \quad (8)$$

That is, the reported margin for w among ballots with votes for exactly one of w and ℓ is larger than the reported margin among ballots reported to show votes for one or

Winner's True Share	Quantiles					Mean	ASN
	25 th	50 th	75 th	90 th	99 th		
70%	12	22	38	60	131	30	30
65%	23	38	66	108	236	53	53
60%	49	84	149	244	538	119	119
58%	77	131	231	381	840	184	185
55%	193	332	587	974	2,157	469	469
54%	301	518	916	1,520	3,366	730	731
53%	531	914	1,619	2,700	5,980	1,294	1,295
52%	1,188	2,051	3,637	6,053	13,455	2,900	2,902
51%	4,725	8,157	14,486	24,149	53,640	11,556	11,562
50.5%	18,839	32,547	57,838	96,411	214,491	46,126	46,150

Table 1: Estimated means and percentiles of the number of ballots with valid votes to inspect for 10% risk limit using BRAVO, as a function of the winner’s share of vote (estimated using 10^7 replications), as well as Wald’s ASN.

both of w and ℓ . Hence, it is conservative to use $s_{w\ell}$ as the basis for the multiplier in the test. This leads to the complete BRAVO procedure:

1. Set $m = 0$ and set $T_{w\ell} = 1$ for all $w \in \mathcal{W}$ and $\ell \in \mathcal{L}$.
2. Draw a ballot uniformly at random with replacement from those cast in the contest and increment m .
3. If the ballot shows a valid vote for a reported winner w , then for each ℓ in $\mathcal{L}$ that did not receive a valid vote on that ballot multiply $T_{w\ell}$ by $s_{w\ell}/0.5$. Repeat for all such w .
4. If the ballot shows a valid vote for a reported loser ℓ , then for each w in $\mathcal{W}$ that did not receive a valid vote on that ballot multiply $T_{w\ell}$ by $(1 - s_{w\ell})/0.5$. Repeat for all such ℓ .
5. If any $T_{w\ell} \geq 1/\alpha$, reject the corresponding null hypothesis for each such $T_{w\ell}$. Once a null hypothesis is rejected, do not update its $T_{w\ell}$ after subsequent draws.
6. If all null hypotheses have been rejected, stop the audit: The reported results stand. Otherwise, if $m < M$, return to step 2.
7. Perform a full hand count; the results of the hand count replace the reported results.

This method limits the overall risk to α : Stopping short of a full hand count is an error only if at least one of the null hypotheses is in fact true. The audit stops only if all of the null hypotheses are rejected. Consider the set of null hypotheses that are true. The chance BRAVO erroneously rejects *all* of those and stops without a full hand count is at most the smallest chance of erroneously rejecting any of them individually. Hence, by testing every

(winner, loser) pair individually at level α , the chance of stopping short of a full hand count if any of the $C - k$ apparent losers actually won is at most α .

For any given risk limit, the expected number of draws to confirm a correct outcome using BRAVO generally depends primarily upon the smallest margin of decision—the difference in vote shares between the winner with the smallest vote share and the loser with the largest vote share.¹⁵ Call these candidates w^* and ℓ^* . In the vote-for-one case, if no other margin of decision is very close to the smallest one (between the reported winner and runner-up), then the expected number of ballots with valid votes to inspect is very close to the expression for the ASN above, setting $p_w = p_{w^*}$, $p_\ell = p_{\ell^*}$, $s_w = p_{w^*}/(p_{w^*} + p_{\ell^*})$, and $s_\ell = 1 - s_w$ —as if all ballots not cast for one of the two leading candidates were invalid.¹⁶

However, if one or more other margins of decision are close to the smallest one, then the expected number of ballots may be substantially larger, as it becomes harder to reject all the pairwise null hypotheses at once. For instance, in a three-candidate contest where the candidate vote shares are 40%, 30%, and 30%, the average sample size (determined by simulation with 10^6 trials) is approximately 433 ballots. This number is modest, but about 31% larger than the ASN formula indicates;

¹⁵This generalization holds (with the caveat described next) when discrepancies between reported and actual vote shares are small relative to the differences among candidates’ actual vote shares; in contests that allow a voter to vote for more than one candidate, it also may depend on the fraction of ballots that show votes for both members of each (winner, loser) pair.

¹⁶For instance, consider a three-candidate contest where the vote shares are 49.5%, 40.5%, and 10%. The apparent winner’s share of the top-two vote is $p_{w^*}/(p_{w^*} + p_{\ell^*}) = 0.495/(0.495 + 0.405) = 0.55$. The average number of ballots inspected, as determined by simulation with 10^6 trials, is about 521. This is essentially equal to the value of ASN (expression (5)) with $p_w = 0.495$, $p_\ell = 0.405$, $s_w = 0.55$, and $s_\ell = 0.45$, or 10/9 the value of ASN with $s_w = p_w = 0.55$ and $w_\ell = p_\ell = 0.45$.

however, simulating the workload is still quite tractable, even for rather small margins. In contests that allow each voter to select more than one candidate, the situation is even more complicated, and may depend on the number of ballots with valid votes for each (winner, loser) pair. Then, even simulating the workload becomes knotty, because it requires assumptions about those “overlaps,” and there is little data to support the assumptions. However, these details affect only the workload, not the risk limit: BRAVO does not require any assumptions about the margins or the overlap to limit the risk to α rigorously.

8 Historical examples

To explore the potential applicability of BRAVO, we examined state-level¹⁷ vote totals from the five U.S. presidential elections from 1992 through 2008. In practice, neither BRAVO nor any other risk-limiting audit method could have been applied in all states in all these elections: Some states have used (and some states now use) voting systems that do not provide a voter-verifiable audit trail. Nevertheless, it seems worthwhile to consider the empirical distribution of reported vote shares in all these states in recent elections.

We extend the analysis to 1992 because Ross Perot’s candidacy that year offers the most interesting recent example of a viable third-party candidacy; as we expected, it had little impact on the results.¹⁸ We use vote counts from Dave Leip’s Atlas of U.S. Presidential Elections (uselectionatlas.org), including counts of invalid ballots where available.¹⁹ We estimated the expected sample sizes needed to confirm the reported outcomes at a risk limit of 10% by simulation. The simulations assume that the reported vote shares are correct. Each simulation estimates the expected number of ballots required to confirm the plurality winner in one state in one election, running BRAVO 10^6 times for each such contest.

Of the 255 state presidential contests in this period, we set aside 23 that had margins smaller than 2 percent. For a 2-percent margin, the expected number of ballots to inspect is over 11,000 (see Table 1 for winner’s share 51%).

¹⁷We treat the District of Columbia as a state. For purposes of this analysis, we do not consider the Maine and Nebraska electors chosen at the congressional district level.

¹⁸By way of example: In Maine, Perot narrowly edged out George H.W. Bush for second place: The vote shares were Bill Clinton 38.77%; Perot 30.44%; Bush 30.39%; other candidates 0.41%. The resulting expected sample size was about 610 ballots—more than the 470 ballots one would expect if only the Clinton-Perot comparison were considered, but not burdensome. In most states the impact was far smaller, essentially undetectable.

¹⁹Invalid ballot counts were not available for 1992 and 1996. Also, in 2008, Connecticut reported fewer ballots cast than presidential votes counted; we assumed that no invalid votes were cast. Imputing invalid vote counts would not materially affect the results.

We think that many states would find it onerous to individually sample tens of thousands of ballots, especially with no clear expectation of when they can stop. Eight of these 23 contests had margins under 0.25 percent, below the threshold for an automatic recount in some states; hand counts (where feasible) arguably would be the most efficient and trustworthy means of confirming those outcomes. In the intermediate cases, alternatives to BRAVO, such as ballot-level comparison audits, might be preferable.

In the remaining 232 state-level presidential contests (91% of the contests under examination), the total expected number of ballots to inspect in order to confirm all outcomes was approximately 225,000, out of almost 512 million ballots cast in those contests. In 49 cases, the expected sample size was under 100 ballots; in 179 cases (70% of all contests under examination), the expected sample size was under 1,000 ballots. The median expected sample size (treating the 23 closest contests as if they required infinitely large samples) was about 307 ballots.

Broadly speaking, we think BRAVO at a 10% risk limit would not be onerous for most states in most elections. However, logistical challenges of statewide audits and alternative methods for the hardest cases require more attention than we offer here.

9 Logistics

9.1 Drawing a random sample

Drawing a simple random sample of ballots is not as straightforward as drawing marbles from an urn. Even if it were possible physically to mix the ballots, it is imprudent. A better approach starts by implicitly enumerating the ballots using a *ballot manifest*. A simple ballot manifest lists the physical containers of ballots in which the ballots are stored and how many ballots are in each container. The auditors can sequentially generate random numbers, convert each one to a ballot number (from 1 to the total number of ballots in the election), and convert each number to a particular ballot, for instance, “the 142nd ballot in box 41.” If the contests under audit are on almost every ballot, we can simply treat ballots that do not contain the contest as having no valid vote, but if the contests under audit are on only a small fraction of the ballots, many draws will yield useless ballots. Retrieving those ballots can be an expensive waste of time.

If the contests under audit appear only on some of the ballots in the election, we can reduce the number of useless draws if we can construct a ballot manifest specific to those elections from the comprehensive manifest; see, e.g., [Lindeman and Stark \[2012\]](#).

Unless the ballot manifest uses specific ballot identifiers, the audit needs a trustworthy way to identify a particular physical ballot in the sample. Innocent unpredictable errors in identifying, say, the 142nd ballot in a box should not, in principle, bias the results. However, reasonable suspicion that systematic error exists or that the auditors can exercise discretion over which ballots are selected undermines the value of the audit. Therefore, care should be taken to specify and follow credible procedures.

If the ballot manifest has errors, the chance of drawing each ballot will not be equal. Some ballots might have no chance at all of being drawn, and the audit might attempt to draw ballots that do not in fact exist—or that exist but are not where the manifest claims they are. As a result, the sampling distribution of the test statistics $\{T_{w\ell}\}$ are not what they were assumed to be, which could make the actual risk limit larger than claimed. Fortunately, there is an easy remedy.

Bañuelos and Stark [2012] show that if there is an upper bound on the number of ballots, a simple modification to the procedure makes the audit *more* conservative if the manifest has errors. That is, the actual risk limit will be even smaller than the claimed risk limit. The modification is simple: If the upper bound on the number of ballots is larger than the total listed in the manifest, create a fictitious group of ballots that contains the extras, and append it to the manifest. In each draw in the audit, sample uniformly from 1 to the upper bound on the number of ballots. If the ballot drawn cannot be found—either because it is in the fictitious group or because the manifest lists more ballots in a batch than are actually there—pretend that a ballot was found, and that the ballot showed a valid vote for every loser ℓ .²⁰ Perhaps surprisingly, it is not necessary to adjust the margin to account for missing ballots or ballots not listed in the manifest, only to treat the ballots actually selected in the course of the audit in the most pessimistic way.

9.2 Group sequential sampling versus item-by-item audits

For practical reasons, a jurisdiction may prefer to retrieve some number of ballots and inspect them together, rather than retrieving and inspecting one ballot at a time. For instance, it may be convenient to select, say, 100 ballots, sort the selections based on the containers in which the ballots are stored, and retrieve those ballots. If any such sorting occurs—if ballots are audited in a different order than the order in which they were randomly selected—

²⁰The rules of the election might allow a ballot to have valid votes for fewer candidates than there are losers ℓ ; nonetheless, for this method to result in a conservative procedure, the auditors should pretend that the ballot gives a valid vote to every loser.

then all the ballots in a selected group *must* be audited. Using groups in this manner increases the expected work, because the audit cannot end in the middle of inspecting a group.²¹ However, as long as sorted groups are audited completely, using group sampling does not compromise the risk limit α .

Moreover, the group size can be varied at will throughout the audit. For instance, if a jurisdiction can audit additional groups rather easily and would like to limit the amount of ballots it has to inspect, it might calculate the ASN and begin the audit by auditing half that many ballots, which provides a non-trivial (on the order of 25%) chance of completing the audit in the first group. Conversely, if limiting the number of stages is more important than limiting the number of ballots counted, the jurisdiction might use the ASN, or even some multiple of the ASN, as the sample size in the first group. The jurisdiction can then adjust subsequent group sizes based on how far the test statistics are from $1/\alpha$ and how highly it values limiting the number of ballots inspected versus limiting the number of groups. For instance, in the two-candidate case, the conditional ASN to confirm the reported outcome given that the results so far yield a test statistic $T = T^*$, assuming that the reported vote shares are correct, is derived from expression (5) by replacing $\ln(1/\alpha)$ with $\ln((1/\alpha)/T^*)$ in the numerator. The jurisdiction can use this conditional ASN to help it decide how many ballots to sample in the next group.²²

9.3 Coordinating audits across multiple jurisdictions

Ballot-polling audits show particular promise in large, multi-jurisdictional elections because the expected number of ballots to inspect often will be a tiny fraction of the ballots in the election, and because the work can be divided among many election officials. However, they also pose distinctive logistical challenges. It is one matter for a particular jurisdiction to conduct a sequential audit ballot-by-ballot, or to choose a convenient alternative. It is another matter for dozens or hundreds of local election officials to participate in a ballot-by-ballot sequential audit. Given modern communications, a coordinated ballot-by-ballot audit could be feasible: At each step, a ballot would be randomly selected from all the ballots in the contest, regardless of jurisdiction, and the appropriate officials would locate and inspect that ballot.

²¹The test statistic may exceed $1/\alpha$ in the course of a group but drop below $1/\alpha$ by the end of the group. In that case, one or more additional groups must be counted before the audit ends.

²²For finer control, it is possible to estimate quantiles of the expected number of ballots to sample: The jurisdiction can select a sample size that is expected to provide, say, a 90% chance of completing the audit if the reported vote shares are correct. We do not explore these details here.

Even if this approach is feasible, it is likely to become tedious for all but the smallest audits. Therefore, multi-jurisdictional audits are likely to be conducted by group sampling rather than item-by-item.

If each jurisdiction can provide its ballot manifests to a central election authority in a machine-readable common format, the central authority can produce a contest-wide master manifest, and then use this manifest to conduct the sampling. However, if it is not immediately practical for all jurisdictions to produce ballot manifests in a common format, it suffices for each jurisdiction to provide the number of ballots enumerated in the manifest. Then, if desired, the central authority can conduct the sample, in each case telling the appropriate jurisdiction to inspect its x th ballot. A more complex alternative, which allows jurisdictions to draw their own samples, is a two-stage sample: The central authority draws a sample to determine how many ballots each jurisdiction should inspect (which will, in general, vary across jurisdictions). Each county then separately samples and inspects as many ballots as required. Regardless of the method used, the jurisdictions report their results to the central authority, which combines them to determine whether the audit can stop. Each of these sampling methods can be repeated as often as necessary.

Auditors in multi-jurisdictional audits are likely to place a higher priority on limiting the number of groups to be audited than limiting the number of ballots to inspect. This is true because of the challenges of coordinating an unpredictable number of groups across multiple jurisdictions, and also because widely distributing the auditing work reduces the efficiency bind: Many hands make work light. For instance, if one jurisdiction conducts an audit with an ASN of 500 ballots, it might prefer to audit 500 (or fewer) ballots in the initial group, rather than to audit 1000 ballots or more simply to increase its chance of finishing in one group. In contrast, if the State of California conducts a statewide audit with an ASN of 500 ballots, it might prefer an initial sample of 1000 ballots or more, knowing that the work will be distributed across 58 counties. Again, the number of ballots to be audited in each group can be informed by the audit results for previous groups, to reduce the chance of needing to audit additional groups, or to reduce the expected number of ballots to audit.

9.4 Pairwise comparisons versus “pooling”

Imagine auditing a vote-for-one plurality contest. If the reported winner apparently received more than half the votes, there are at least two approaches we could take to auditing: (i) “pool” some or all the losing candidates and test the null hypothesis that each “pooled” group received more votes than the reported winner, and (ii) test

the $C - 1$ hypotheses that each of the $C - 1$ apparent losers got more votes than the reported winner. We show in this section that the latter approach is more efficient; the proof also applies to the general case BRAVO considers: plurality contests with k winners in which each voter may select zero or more candidates.

Suppose that the reported results are correct. We will show that, for any (winner, loser) pair (w, ℓ) and any number n of draws, $\Pr\{T_{w\ell} \geq 1/\alpha\}$ is larger than $\Pr\{T_{wl} \geq 1/\alpha\}$, where l is any composite “pooled” candidate consisting of ℓ and any subset of the other losers. Since we stop the audit only when all the test statistics are greater than $1/\alpha$, it follows that the expected sample size is smaller if candidates are not pooled.

Recall that s_w is w ’s share of the valid ballots and s_ℓ is the loser’s share. Let s_m be the combined share of any other losers with which we are considering pooling ℓ . To pool ℓ with m , we require $s_w > s_\ell + s_m$; otherwise, w was not reported to get more votes than the pooled loser. Consider the j th draw. Let I_{wj} denote the event that the ballot drawn shows a vote for w , $I_{\ell j}$ the event that the ballot shows a vote for ℓ , and I_{mj} the event that the ballot shows a vote for some other loser we are grouping with. Then $I_{wj} + I_{\ell j} + I_{mj} \leq 1$, $\Pr(I_{wj} = 1) = s_w$, and so on.

We focus on a single draw (since the draws are independent and identically distributed) and condition on the event that the ballot drawn shows a valid vote; otherwise, it cannot help, whether we pool or not. The event $T \geq 1/\alpha$ is the same as the event $\ln T \geq \ln 1/\alpha$. The random variable $\ln T$ is the sum of independent, identically distributed increments—a random walk with drift. Without pooling, the increment to $\ln T$ from the j th draw is

$$Z_{j\ell} = I_w \ln 2 \frac{s_w}{s_w + s_\ell} + I_\ell \ln 2 \frac{s_\ell}{s_w + s_\ell}. \quad (9)$$

With grouping, the increment is

$$Z_{jl} = I_w \ln 2 \frac{s_w}{s_w + s_\ell + s_m} + (I_\ell + I_m) \ln 2 \frac{s_m + s_\ell}{s_w + s_\ell + s_m}. \quad (10)$$

A result in [Wald \[1945\]](#) implies that the expected number of observations necessary to reach a decision is $1/\alpha$ divided by the expected increment to the test statistic: The larger the expected increment, the smaller the expected number of draws. Hence, it suffices to show that the expected value of $D \equiv Z_{j\ell} - Z_{jl}$, the difference in the expected increments without and with grouping, is non-

negative. Observe:

$$\begin{aligned}
D &\equiv s_w \ln \frac{s_w + s_\ell + s_m}{s_w + s_\ell} \\
&\quad + s_\ell \ln \frac{s_\ell (s_w + s_\ell + s_m)}{(s_w + s_\ell)(s_m + s_\ell)} \\
&\quad - s_m \ln 2 \frac{s_m + s_\ell}{s_w + s_\ell + s_m} \\
&= s_w \ln \frac{s_w + s_\ell + s_m}{s_w + s_\ell} \\
&\quad + \ln \frac{s_\ell^s (s_w + s_\ell + s_m)^{s_\ell + s_m}}{2^{s_m} (s_w + s_\ell)^{s_\ell} (s_m + s_\ell)^{s_m + s_\ell}}.
\end{aligned}$$

Since $\frac{s_\ell}{s_w + s_\ell} < \frac{1}{2}$,

$$\begin{aligned}
D &\geq s_w \ln \frac{s_w + s_\ell + s_m}{s_w + s_\ell} \\
&\quad + \ln \left[\frac{1}{2^{s_m + s_\ell}} \frac{(s_w + s_\ell + s_m)^{s_\ell + s_m}}{(s_m + s_\ell)^{s_m + s_\ell}} \right] \\
&= (s_w + s_\ell + s_m) \ln [s_w + s_\ell + s_m] \\
&\quad - s_w \ln [s_w + s_\ell] \\
&\quad - (s_\ell + s_m) \ln [2(s_m + s_\ell)]
\end{aligned}$$

Finally since $s_w + s_\ell + s_m > s_w + s_\ell$, and $s_w > s_\ell + s_m$,

$$\begin{aligned}
D &\geq (s_w + s_\ell + s_m) \ln [s_w + s_\ell + s_m] \\
&\quad - s_w \ln [s_w + s_\ell + s_m] \\
&\quad - (s_\ell + s_m) \ln [s_w + s_\ell + s_m] \\
&= 0. \square
\end{aligned}$$

10 Discussion

BRAVO provides a way to perform a risk-limiting audit of majority, super-majority, and plurality contests, including contests with more than one winner and contests in which voters may cast votes for more than one candidate. BRAVO places minimal demands on the voting system: It requires the reported contest results, an audit trail, and a ballot manifest that explains how the ballots are stored, so that ballots can be selected at random. In contrast, comparison audits require detailed reports from the voting system that are not produced in machine-readable form by current vote tabulation systems.

The number of ballots that must be audited using BRAVO when the reported results are correct can be quite small, and that workload is distributed over all the jurisdictions involved in the contest. That makes BRAVO immediately practical for jurisdictions that have an audit trail, for contests with margins down to a few percent. In particular, the median expected sample size for states' presidential contests from 1992 through 2008 is about 307 ballots.

BRAVO does not check the accuracy of the voting system, only the correctness of the electoral outcome. The voting system could get the right outcome through fortuitous cancellation of errors, which a comparison risk-limiting audit might detect. The workload for BRAVO becomes prohibitive when the margin is small; in such cases, ballot-level or batch-level comparison audits might be preferable, despite their higher set-up costs.

11 Acknowledgments

We are grateful to Jennie Bretschneider, Joseph Lorenzo Hall, Ronald L. Rivest, and Pamela Smith for helpful comments.

References

- Bañuelos, J. and Stark, P. (2012). Limiting risk by turning manifest phantoms into evil zombies. Technical report, arXiv.org. <http://arxiv.org/abs/1207.3413>. Retrieved 17 July 2012.
- Benaloh, J., Jones, D., Lazarus, E., Lindeman, M., and Stark, P. (2011). SOBA: Secrecy-preserving observable ballot-level audits. In *Proceedings of the 2011 Electronic Voting Technology Workshop / Workshop on Trustworthy Elections (EVT/WOTE '11)*. USENIX.
- Calandrino, J., Halderman, J., and Felten, E. (2007). Machine-assisted election auditing. In *Proceedings of the 2007 USENIX/ACCURATE Electronic Voting Technology Workshop (EVT 07)*. USENIX.
- Checkoway, S., Sarwate, A., and Shacham, H. (2010). Single-ballot risk-limiting audits using convex optimization. In *Proceedings of the 2010 Electronic Voting Technology Workshop / Workshop on Trustworthy Elections (EVT/WOTE '10)*. USENIX. http://www.usenix.org/events/evtwote10/tech/full_papers/Checkoway.pdf. Retrieved April 20, 2011.
- ElectionAudits.org (2008). Principles and best practices for post-election audits. http://www.electionaudits.org/files/best%20practices%20final_0.pdf. Retrieved May 10, 2012.
- Gibbons, J. D., Olkin, I., and Sobel, M. (1977). Selecting and Ordering Populations : A New Statistical Methodology. In *Selecting and Ordering Populations: A New Statistical Methodology*, chapter 6, pages 158–186. SIAM.

- Hall, J., Miratrix, L., Stark, P., Briones, M., Ginnold, E., Oakley, F., Peadar, M., Pellerin, G., Stanionis, T., and Webber, T. (2009). Implementing risk-limiting post-election audits in California. In *Proc. 2009 Electronic Voting Technology Workshop/Workshop on Trustworthy Elections (EVT/WOTE '09)*, Montreal, Canada. USENIX.
- Johnson, K. (2004). Election certification by statistical audit of voter-verified paper ballots. http://papers.ssrn.com/sol3/papers.cfm?abstract_id=640943. Retrieved 6 March 2011.
- Lindeman, M. and Stark, P. B. (2012). A gentle introduction to risk-limiting audits. *IEEE Security and Privacy*, page to appear.
- McDaniel, P., Blaze, M., Vigna, G., and et al. (2007). EVEREST: Evaluation and Validation of Election-Related Equipment, Standards and Testing. Retrieved 17 March 2012.
- Ng, H. K. T. and Panchapakesan, S. (2007). Is the Selected Multinomial Cell the Best? *Sequential Analysis*, 26(4):415–423.
- Ramey, J. and Alam, K. (1979). A sequential procedure for selecting the most probable multinomial event. *Biometrika*, pages 171–173.
- Rivest, R. and Shen, E. (2012). A Bayesian method for auditing elections. In *Proceedings of the 2012 Electronic Voting Technology Workshop / Workshop on Trustworthy Elections (EVT/WOTE '12)*. USENIX.
- Secretary of State, C. (2007). Top-to-bottom review. Retrieved 9 July 2012.
- Simon, J. D. and O'Dell, B. (2006). An end to 'faith-based' voting: universal precinct-based handcount sampling to check computerized vote counts in federal and statewide elections. <http://electiondefensealliance.org/files/UPSEndFaithBasedVoting.pdf>. Retrieved July 6, 2012.
- Sorentrue, J., Kam, D., and Bennett, G. (2012). Recount shows wrong winners declared in two Wellington election races. Retrieved 3 May 2012.
- Stark, P. (2008a). Conservative statistical post-election audits. *Ann. Appl. Stat.*, 2:550–581.
- Stark, P. (2008b). A sharper discrepancy measure for post-election audits. *Ann. Appl. Stat.*, 2:982–985.
- Stark, P. (2009a). CAST: Canvass audits by sampling and testing. *IEEE Transactions on Information Forensics and Security, Special Issue on Electronic Voting*, 4:708–717.
- Stark, P. (2009b). Efficient post-election audits of multiple contests: 2009 California tests. <http://ssrn.com/abstract=1443314>. 2009 Conference on Empirical Legal Studies.
- Stark, P. (2009c). Risk-limiting post-election audits: P -values from common probability inequalities. *IEEE Transactions on Information Forensics and Security*, 4:1005–1014.
- Stark, P. (2010). Super-simple simultaneous single-ballot risk-limiting audits. In *Proceedings of the 2010 Electronic Voting Technology Workshop / Workshop on Trustworthy Elections (EVT/WOTE '10)*. USENIX.
- Stark, P. B. and Wagner, D. A. (2012). Evidence-based elections. *IEEE Security and Privacy*, page to appear.
- Wald, A. (1945). Sequential tests of statistical hypotheses. *Ann. Math. Stat.*, 16:117–186.
- Wald, A. (2004). *Sequential Analysis*. Dover Publications, Mineola, New York.