

Earning While Learning

How to Run Batched Bandit Experiments

Davud Rostam-Afschar

University of Mannheim

Milan, 17 September 2026

Why adaptive experiments?

Main research question

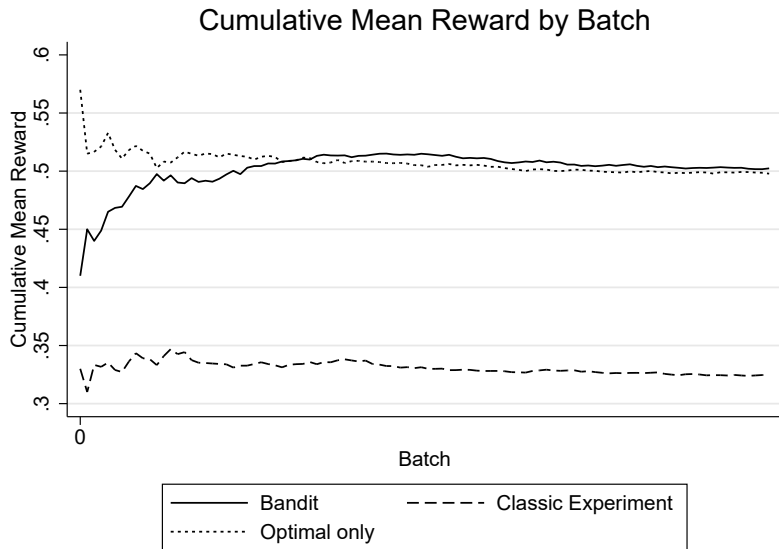
- ▶ In many settings, data arrive sequentially and treatments can be updated over time.



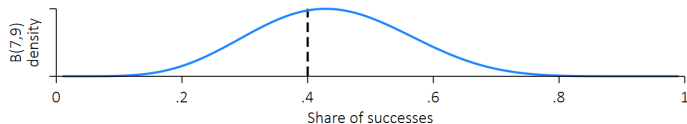
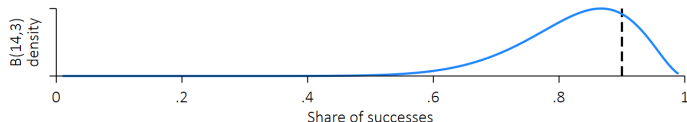
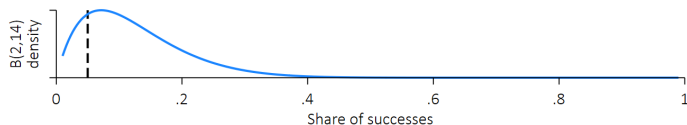
How can we improve decisions *during* an experiment, not only after it?

- ▶ Adaptive experiments allocate more observations to promising treatments while still learning about all options.
- ▶ This is useful in economics, survey research, online experiments, education, political science, sociology, health, ...

Adaptive vs fixed design



Thompson (1933, 1935) sampling



- ▶ Beta-Bernoulli Thompson sampling
- ▶ Models uncertainty about the shape of the distribution and the expected outcome R explicitly

[Click to watch!](#)

Thompson sampling

- ▶ Beta distribution $B(R_{k,t}|\alpha_k, \beta_k)$ denotes the density of the beta distribution for random variable R_t with parameters α_k and β_k
- ▶ Posterior distribution $P(\theta|R_t)$ is also beta with parameters that can be updated according to a simple rule:

$$(\alpha_k, \beta_k) = \begin{cases} (\alpha_k, \beta_k) & \text{if chosen arm} \neq k, \\ (\alpha_k, \beta_k) + (R_t, 1 - R_t) & \text{if chosen arm} = k. \end{cases}$$

- ▶ α_k or β_k increases by one with each observed success or failure
- ▶ Distribution more concentrated as $\alpha_k + \beta_k$ grows
- ▶ Mean $\alpha_k/(\alpha_k + \beta_k)$ and variance

$$\frac{\alpha_k}{\alpha_k + \beta_k} \frac{\beta_k}{\alpha_k + \beta_k} \frac{1}{\alpha_k + \beta_k + 1} = \frac{\alpha_k \beta_k}{(\alpha_k + \beta_k)^2 (\alpha_k + \beta_k + 1)}$$

Valid inference for adaptive experiments

Main challenge

Standard OLS inference can fail when treatment assignment depends on past outcomes.

- ▶ Adaptive assignment changes treatment probabilities over time.
- ▶ This breaks the usual asymptotics behind standard confidence intervals and tests.
- ▶ The problem is especially important when treatment differences are small or assignment becomes imbalanced.

Valid inference for adaptive experiments

Main challenge

Standard OLS inference can fail when treatment assignment depends on past outcomes.

- ▶ Adaptive assignment changes treatment probabilities over time.
- ▶ This breaks the usual asymptotics behind standard confidence intervals and tests.
- ▶ The problem is especially important when treatment differences are small or assignment becomes imbalanced.

What we do instead

Use **batched OLS (BOLS)** with heteroskedasticity-robust weighting:

$$\hat{\Delta}^{\text{BOLS}} = \sum_{t=1}^T \frac{w_t}{S} \hat{\Delta}_t, \quad w_t = \frac{1}{\sqrt{V_t}}, \quad S = \sum_{t=1}^T w_t.$$

Valid inference for adaptive experiments

Main challenge

Standard OLS inference can fail when treatment assignment depends on past outcomes.

- ▶ Adaptive assignment changes treatment probabilities over time.
- ▶ This breaks the usual asymptotics behind standard confidence intervals and tests.
- ▶ The problem is especially important when treatment differences are small or assignment becomes imbalanced.

What we do instead

Use **batched OLS (BOLS)** with heteroskedasticity-robust weighting:

$$\hat{\Delta}^{\text{BOLS}} = \sum_{t=1}^T \frac{w_t}{S} \hat{\Delta}_t, \quad w_t = \frac{1}{\sqrt{V_t}}, \quad S = \sum_{t=1}^T w_t.$$

i) Normalize batchwise variances to 1,

Valid inference for adaptive experiments

Main challenge

Standard OLS inference can fail when treatment assignment depends on past outcomes.

- ▶ Adaptive assignment changes treatment probabilities over time.
- ▶ This breaks the usual asymptotics behind standard confidence intervals and tests.
- ▶ The problem is especially important when treatment differences are small or assignment becomes imbalanced.

What we do instead

Use **batched OLS (BOLS)** with heteroskedasticity-robust weighting:

$$\hat{\Delta}^{\text{BOLS}} = \sum_{t=1}^T \frac{w_t}{S} \hat{\Delta}_t, \quad w_t = \frac{1}{\sqrt{V_t}}, \quad S = \sum_{t=1}^T w_t.$$

- Normalize batchwise variances to 1,
- Std. normal batchwise z-statistic leads to agg. z-statistic $Z \sim N(0, 1)$,

Valid inference for adaptive experiments

Main challenge

Standard OLS inference can fail when treatment assignment depends on past outcomes.

- ▶ Adaptive assignment changes treatment probabilities over time.
- ▶ This breaks the usual asymptotics behind standard confidence intervals and tests.
- ▶ The problem is especially important when treatment differences are small or assignment becomes imbalanced.

What we do instead

Use **batched OLS (BOLS)** with heteroskedasticity-robust weighting:

$$\hat{\Delta}^{\text{BOLS}} = \sum_{t=1}^T \frac{w_t}{S} \hat{\Delta}_t, \quad w_t = \frac{1}{\sqrt{V_t}}, \quad S = \sum_{t=1}^T w_t.$$

- Normalize batchwise variances to 1,
- Std. normal batchwise z-statistic leads to agg. z-statistic $Z \sim N(0, 1)$,
- Invert Z statistic to get confidence intervals.

Valid inference for adaptive experiments

Main challenge

Standard OLS inference can fail when treatment assignment depends on past outcomes.

- ▶ Adaptive assignment changes treatment probabilities over time.
- ▶ This breaks the usual asymptotics behind standard confidence intervals and tests.
- ▶ The problem is especially important when treatment differences are small or assignment becomes imbalanced.

What we do instead

Use **batched OLS (BOLS)** with heteroskedasticity-robust weighting:

$$\hat{\Delta}^{\text{BOLS}} = \sum_{t=1}^T \frac{w_t}{S} \hat{\Delta}_t, \quad w_t = \frac{1}{\sqrt{V_t}}, \quad S = \sum_{t=1}^T w_t.$$

- Normalize batchwise variances to 1,
- Std. normal batchwise z-statistic leads to agg. z-statistic $Z \sim N(0, 1)$,
- Invert Z statistic to get confidence intervals.

This gives asymptotically valid inference in batched adaptive experiments.

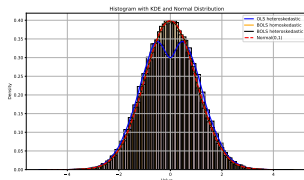
Weight structure

Obs	Selected arm	Batch	Outcome	True expected outcome	OLS	Batchwise OLS	ω_t
1	A	0	0	0.5	0.600	0.500	$\sqrt{\frac{2 \times 2}{2+2}}$
2	B	0	0	0.2	0.167	0.000	$\sqrt{\frac{2 \times 2}{2+2}}$
3	A	0	1	0.5	0.600	0.500	$\sqrt{\frac{2 \times 2}{2+2}}$
4	B	0	0	0.2	0.167	0.000	$\sqrt{\frac{2 \times 2}{2+2}}$
5	A	1	0	0.5	0.600	0.500	$\sqrt{\frac{2 \times 2}{2+2}}$
6	B	1	1	0.2	0.167	0.500	$\sqrt{\frac{2 \times 2}{2+2}}$
7	A	1	1	0.5	0.600	0.500	$\sqrt{\frac{2 \times 2}{2+2}}$
8	B	1	0	0.2	0.167	0.500	$\sqrt{\frac{2 \times 2}{2+2}}$
9	A	2	0	0.5	0.600	0.667	$\sqrt{\frac{1 \times 3}{1+3}}$
10	A	2	1	0.5	0.600	0.667	$\sqrt{\frac{1 \times 3}{1+3}}$
11	A	2	1	0.5	0.600	0.667	$\sqrt{\frac{1 \times 3}{1+3}}$
12	B	2	0	0.2	0.167	0.000	$\sqrt{\frac{1 \times 3}{1+3}}$
13	A	3	1	0.5	0.600	0.667	$\sqrt{\frac{1 \times 3}{1+3}}$
14	A	3	0	0.5	0.600	0.667	$\sqrt{\frac{1 \times 3}{1+3}}$
15	A	3	1	0.5	0.600	0.667	$\sqrt{\frac{1 \times 3}{1+3}}$
16	B	3	0	0.2	0.167	0.000	$\sqrt{\frac{1 \times 3}{1+3}}$

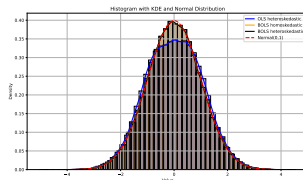
Weight structure

	Non-adaptive	Adaptive
	<i>Homoskedastic case: $\sigma_{1,t}^2 = \sigma_{0,t}^2 = \sigma^2$</i>	
Realized share $\hat{\pi}_t$	$\hat{\pi}_t$	$\hat{\pi}_t$ depends on adaptive assignments
Variance v_t	$v_t = \frac{\sigma_t^2}{n_t \hat{\pi}_t (1 - \hat{\pi}_t)}$	$v_t = \frac{\sigma_t^2}{n_t \hat{\pi}_t (1 - \hat{\pi}_t)}$
Weight $\frac{w_t}{S}$	$\frac{\sqrt{n_t \hat{\pi}_t (1 - \hat{\pi}_t)}}{\sum_{s=1}^T \sqrt{n_s \hat{\pi}_s (1 - \hat{\pi}_s)}}$	$\frac{\sqrt{n_t \hat{\pi}_t (1 - \hat{\pi}_t)}}{\sum_{s=1}^T \sqrt{n_s \hat{\pi}_s (1 - \hat{\pi}_s)}}$
	<i>Heteroskedastic case: $\sigma_{1,t}^2 \neq \sigma_{0,t}^2$</i>	
Realized share $\hat{\pi}_t$	$\hat{\pi}_t$	$\hat{\pi}_t$ depends on adaptive assignments
Variance v_t	$v_t = \frac{1}{n_t} \left(\frac{\sigma_{1,t}^2}{\hat{\pi}_t} + \frac{\sigma_{0,t}^2}{1 - \hat{\pi}_t} \right)$	$v_t = \frac{1}{n_t} \left(\frac{\sigma_{1,t}^2}{\hat{\pi}_t} + \frac{\sigma_{0,t}^2}{1 - \hat{\pi}_t} \right)$
Weight $\frac{w_t}{S}$	$\frac{\sqrt{n_t} \left(\frac{\sigma_{1,t}^2}{\hat{\pi}_t} + \frac{\sigma_{0,t}^2}{1 - \hat{\pi}_t} \right)^{-1/2}}{\sum_{s=1}^T \sqrt{n_s} \left(\frac{\sigma_{1,s}^2}{\hat{\pi}_s} + \frac{\sigma_{0,s}^2}{1 - \hat{\pi}_s} \right)^{-1/2}}$	$\frac{\sqrt{n_t} \left(\frac{\sigma_{1,t}^2}{\hat{\pi}_t} + \frac{\sigma_{0,t}^2}{1 - \hat{\pi}_t} \right)^{-1/2}}{\sum_{s=1}^T \sqrt{n_s} \left(\frac{\sigma_{1,s}^2}{\hat{\pi}_s} + \frac{\sigma_{0,s}^2}{1 - \hat{\pi}_s} \right)^{-1/2}}$

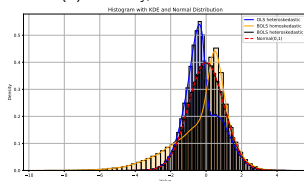
Inference patterns and applications



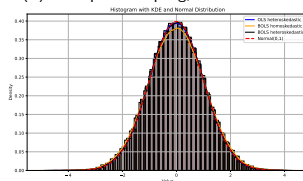
(a) ϵ -Greedy, homoskedastic



(b) Thompson sampling, homoskedastic



(c) ϵ -Greedy, heteroskedastic



(d) Thompson sampling, heteroskedastic

Why this is practical: With the provided bbandits code, adaptive experiments are easy to simulate, run, visualize, and analyze in Stata/Python.

Kemper and Rostam-Afschar (2025)

Kemper and Rostam-Afschar (2026)

Monte Carlo Simulations

- ▶ *Click to watch*: OLS fails normality when margin is small
- ▶ *Click to watch*: BOLS normal even when margin is small
- ▶ Run own simulations with

```
bbandits_sim 0.5 0.4 0.3, size(200) batch(10) clipping(0.1)  
thompson plot_thompson
```

The bbandits command

Syntax & Options

[Click to download!](#)

- ▶ `bbandits reward assignedarm batch, options`

Returned results

- ▶ OLS margins
- ▶ BOLS margins
- ▶ z statistics
- ▶ p-values
- ▶ BOLS 95% confidence intervals
- ▶ observations of the reference arm
- ▶ observations of the treatment arm

Empirical application (Kasy and Sautmann, 2021)

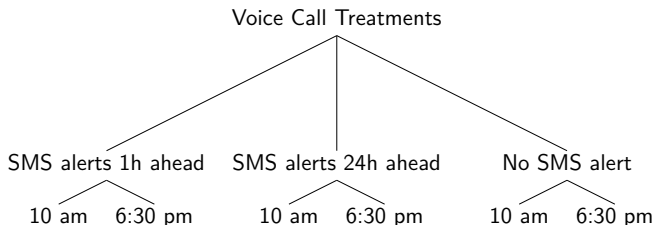
Six call methods to enroll rice farmers

- ▶ Kasy and Sautmann (2021) designed an experiment using *exploration sampling* for Precision Agriculture for Development
- ▶ NGO that works with government partners to provide a phone-based personalized agricultural extension service to farmers in India
- ▶ Aim is to choose best call methods to enroll rice farmers in one state

Empirical application (Kasy and Sautmann, 2021)

Six call methods to enroll rice farmers

- ▶ The outcome (reward) is a binary variable for call completion:
 - ▶ = 1 if call recipient answered five questions asked during call
 - ▶ = 0 otherwise



Empirical application (Kasy and Sautmann, 2021)

- ▶ *Exploration sampling* replaces the Thompson assignment shares
- ▶ modification shifts weight away from the best performing option to competing treatments
- ▶ 10,000 valid phone numbers randomly assigned to one of 16 batches
- ▶ batch size was 600 numbers each (and one with 400)
- ▶ From June 3, 2019 batches run every other day, completed next day

Empirical application (Kasy and Sautmann, 2021)

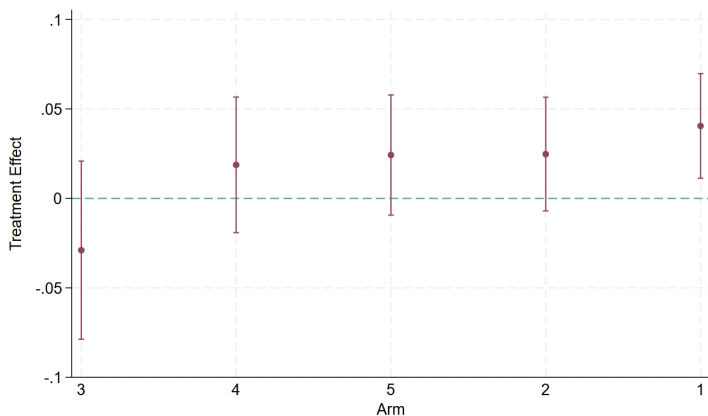
```
. use "example data\kasy_sautmann_2021.dta", clear
. bbandits outcome treatment date
```

```
Number of obs          =      10000
Est. Rewards only best arm      =      1926   Mean reward best arm      =      0.1926
Actual total reward          =      1804   Actual mean reward      =      0.1804
Est. reward uniformly chosen arms =      1709   Mean reward uniform      =      0.1709
```

Arm b	Mean Reward	Share arm b					
	0.1606	0.0903					
k v. b	Margin OLS	Margin BOLS	z	P> z	[95% Conf. Interval]		Share arm k
1-0	0.0320	0.0406	2.61	0.009	0.0101	0.0711	0.3931
2-0	0.0185	0.0249	1.51	0.132	-0.0075	0.0572	0.2234
3-0	-0.0158	-0.0289	-1.12	0.262	-0.0795	0.0216	0.0366
4-0	0.0078	0.0188	0.97	0.330	-0.0191	0.0568	0.1081
5-0	0.0192	0.0243	1.40	0.161	-0.0097	0.0582	0.1485

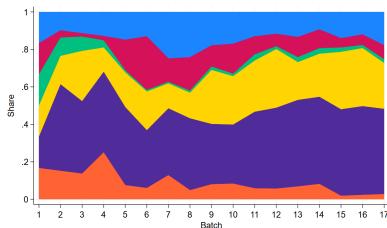
Treatment	0	1	2	3	4	5
SMS	—	1h ahead	24h ahead	—	1h ahead	24h ahead
Call time	10 am	10 am	10 am	6:30 pm	6:30 pm	6:30 pm

Empirical application (Kasy and Sautmann, 2021)

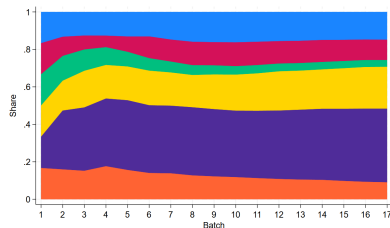


The figure was generated using `kasy_sautmann_2021.dta` and running
`bbandits outcome treatment date`

Empirical application (Kasy and Sautmann, 2021)



(a) Batchwise shares



(b) Cumulative shares

The figure was generated using `kasy_sautmann_2021.dta` and running
`bbandits outcome treatment date`

Empirical application (Kasy and Sautmann, 2021)

Takeaways

Clear best and worst arms

- ▶ Best: Calling farmers at 10 am after a message an hour ahead of time
- ▶ Worst: Calling at 6:30 pm without a text message alert

Improvement of success rate

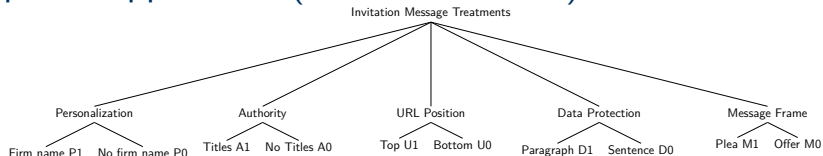
- ▶ 18.04% success rates within the experiment
- ▶ 17.09% success rate with equal assignment

Empirical application (Gaul et al., 2025)

32 invitation messages for business survey

- ▶ Gaul et al. (2025) designed an experiment using *Thompson sampling* to support the German Business Panel (GBP)
- ▶ Aim is to select among a variety of different invitation messages to survey firm decision makers in Germany
- ▶ The GBP is a web-based survey study of firm decision makers in Germany that invites participants each work day (see Bischof et al. (2024); Hack and Rostam-Afschar (2024))
- ▶ The outcome (reward) is a binary variable for the start of the survey:
 - ▶ = 1 if email invitation recipient started the survey
 - ▶ = 0 otherwise

Empirical application (Gaul et al., 2025)



- ▶ Five components of invitation letters and their full interactions
→ $2^5 = 32$ treatments
 - ▶ **personalization** by mentioning or not mentioning the firm name
 - ▶ **authority** of the sender by listing the official full academic titles along with the senders' names or their names only
 - ▶ **URL position** to start the survey at the top or bottom of the invitation
 - ▶ **data protection** in a separate paragraph with two strongly phrased sentences or in a single sentence
 - ▶ **message frame** by including phrases that plea for support in the survey's cause or to simply offer to participate

Empirical application (Gaul et al., 2025)

- ▶ 11,000 randomly selected contacts from firms in Germany
- ▶ Assigned to each of 15 batches from a list of 176,000 contacts
- ▶ Each batch corresponds to a week between August 16, 2022 and November 25, 2022
- ▶ First four batches used fixed and balanced burn-in phase with treatment probability $1/32$
- ▶ From batch 5, Thompson assignment rule for each consecutive batch

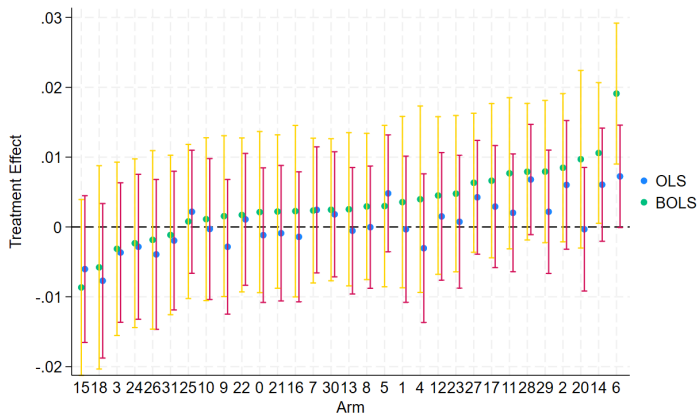
Empirical application (Gaul et al., 2025)

```
. use "example data\gaul_et_al_2024.dta", clear
. bbandits reward selected trial
```

Number of obs	=	176000		
Est. Rewards only best arm	=	8623	Mean reward best arm	= 0.0490
Actual total reward	=	7833	Actual mean reward	= 0.0445
Est. reward uniformly chosen arms	=	7430	Mean reward uniform	= 0.0422

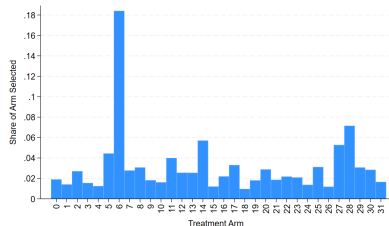
Arm b	Mean Reward	Share arm b					
	0.0417	0.0181					
k v. b	Margin OLS	Margin BOLS	z	P> z	[95% Conf. Interval]		Share arm k
1-0	-0.0014	0.0023	0.37	0.712	-0.0100	0.0146	0.0220
2-0	-0.0003	0.0097	1.50	0.135	-0.0030	0.0224	0.0288
3-0	-0.0028	-0.0023	-0.37	0.710	-0.0144	0.0098	0.0137
4-0	-0.0030	0.0039	0.57	0.567	-0.0095	0.0174	0.0125
5-0	0.0022	0.0008	0.14	0.888	-0.0103	0.0119	0.0312
6-0	-0.0060	-0.0086	-1.32	0.187	-0.0214	0.0042	0.0121
7-0	0.0018	0.0025	0.47	0.638	-0.0078	0.0127	0.0284
8-0	-0.0003	0.0036	0.57	0.570	-0.0087	0.0158	0.0141
9-0	0.0048	0.0030	0.50	0.619	-0.0088	0.0147	0.0444
10-0	-0.0003	0.0011	0.19	0.851	-0.0105	0.0128	0.0162
11-0	-0.0000	0.0029	0.55	0.583	-0.0075	0.0134	0.0308
12-0	-0.0077	-0.0058	-0.73	0.466	-0.0213	0.0097	0.0097
13-0	-0.0009	0.0022	0.40	0.692	-0.0088	0.0132	0.0186
14-0	0.0068	0.0078	1.51	0.130	-0.0023	0.0180	0.0715
15-0	0.0029	0.0066	1.16	0.245	-0.0045	0.0177	0.0331
16-0	0.0011	0.0017	0.30	0.762	-0.0093	0.0127	0.0219
17-0	0.0042	0.0064	1.23	0.218	-0.0038	0.0165	0.0527
18-0	-0.0005	0.0025	0.45	0.650	-0.0084	0.0135	0.0255
19-0	0.0015	0.0045	0.78	0.438	-0.0068	0.0158	0.0256
20-0	0.0061	0.0107	2.02	0.044	0.0003	0.0210	0.0571
21-0	0.0060	0.0084	1.53	0.126	-0.0024	0.0191	0.0271
22-0	0.0072	0.0194	3.61	0.000	0.0089	0.0300	0.1840
23-0	-0.0020	-0.0011	-0.19	0.847	-0.0126	0.0103	0.0166
24-0	-0.0012	0.0021	0.36	0.718	-0.0094	0.0137	0.0190
25-0	0.0022	0.0079	1.52	0.129	-0.0023	0.0182	0.0308
26-0	-0.0039	-0.0018	-0.27	0.787	-0.0147	0.0111	0.0119
27-0	-0.0037	-0.0031	-0.48	0.628	-0.0155	0.0094	0.0155
28-0	-0.0028	0.0015	0.26	0.797	-0.0100	0.0130	0.0183
29-0	0.0020	0.0077	1.38	0.168	-0.0032	0.0186	0.0400
30-0	0.0007	0.0048	0.83	0.405	-0.0064	0.0160	0.0210
31-0	0.0024	0.0024	0.44	0.658	-0.0081	0.0128	0.0278

Empirical application (Gaul et al., 2025)

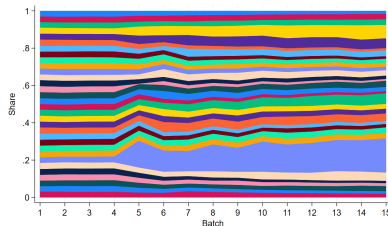


The figure was generated using `gaul_et_al_2024.dta` and running
`bbandits reward selected trial`

Empirical application (Gaul et al., 2025)



(a) Total frequency of treatment assignment



(b) Cumulative shares of arms played

The figure was generated using `gaul_et_al_2024.dta` and running `bbandits reward selected trial`

Empirical application (Gaul et al., 2025)

Takeaways

Clear best and worst arms

- ▶ Using personalization, authority, and pleading for support has greatest success

Interaction effects important, too

More results in the paper...

Running own bandit experiments interactively

- ▶ `. bbandits_initialize, batches(10) arms(3)`
`exploration_phase(2)`

	ID	reward	chosen_a...	batch
1	school_1	.	1	1
2	school_2	.	1	1
3	school_3	.	3	1
4	school_4	.	1	1
5	school_5	.	3	1
6	school_6	.	2	1
7	school_7	.	2	1
8	school_8	.	1	1
9	school_9	.	2	1
10	school_10	.	1	1

Running own bandit experiments interactively

- ▶ `. bbandits_update reward chosen_arm batch, thompson clipping(0.2) excel("mypath")`

	A	B	C	D	E	F
1	ID	rand	reward	chosen_arm	batch	chosen_arm_numeric
192	school_19	1	0	1	2	2
193	school_19	1	1	3	2	0
194	school_19	0	1	3	2	0
195	school_19	1	0	3	2	0
196	school_19	0	1	2	2	1
197	school_19	0	1	1	2	2
198	school_19	1	0	2	2	1
199	school_19	1	0	3	2	0
200	school_19	0	1	2	2	1
201	school_20	0	1	1	2	2
202	school_20	1		3	3	0
203	school_20	1		3	3	0
204	school_20	1		3	3	0
205	school_20	1		2	3	1
206	school_20	0		2	3	1

Running own bandit experiments interactively

- ▶ `. bbandits_update reward chosen_arm batch, thompson
clipping(0.2) excel("mypath")`

Running own bandit experiments interactively

- ▶ `. bbandits_update reward chosen_arm batch, thompson
clipping(0.2) excel("mypath")`
- ▶ ...

Running own bandit experiments interactively

- ▶ `. bbandits_update reward chosen_arm batch, thompson
clipping(0.2) excel("mypath")`
- ▶ ...
- ▶ analyse with
`. bbandits reward chosen_arm batch`

Best Practices

- ▶ Report OLS and BOLS
- ▶ BOLS inference in the small margin case **correct** but...
- ▶ OLS inference in the large margin case **more precise**
- ▶ Check batch-wise OLS estimates

Best Practices

- ▶ Report OLS and BOLS
- ▶ BOLS inference in the small margin case **correct** but...
- ▶ OLS inference in the large margin case **more precise**
- ▶ Check batch-wise OLS estimates
- ▶ **At least 50 observations** per batch and arm
- ▶ From *statistical testing* perspective:
more observations per batch and arm better
- ▶ from *regret optimization* perspective
fewer observations and thus fails are better

Best Practices

- ▶ Report OLS and BOLS
- ▶ BOLS inference in the small margin case **correct** but...
- ▶ OLS inference in the large margin case **more precise**
- ▶ Check batch-wise OLS estimates
- ▶ **At least 50 observations** per batch and arm
- ▶ From *statistical testing* perspective:
more observations per batch and arm better
- ▶ from *regret optimization* perspective
fewer observations and thus fails are better
- ▶ use **bbandits** to simulate, visualize, and analyse bandit experiments

Thank you!

rostam-afschar@uni-mannheim.de

Kemper, J., & Rostam-Afschar, D. (2026). **Inference for batched adaptive experiments.**

Kemper, J., & Rostam-Afschar, D. (2026). **Earning while learning: How to run batched bandit experiments.** *The Stata Journal*.

Gaul, J. J., Keusch, F., Rostam-Afschar, D., & Simon, T. (2025). **Invitation messages for business surveys: A multi-armed bandit experiment.** *Survey Research Methods*, 19(4), 409–429.

References I

- BISCHOF, J., P. DOERRENBURG, D. ROSTAM-AFSCHAR, D. SIMONS, AND J. VOGET (2024): "The German Business Panel: Firm-Level Data for Accounting and Taxation Research," *European Accounting Review*.
- GAUL, J. J., F. KEUSCH, D. ROSTAM-AFSCHAR, AND T. SIMON (2025): "Invitation Messages for Business Surveys: A Multi-Armed Bandit Experiment," *Survey Research Methods*, 19(4), 409–429.
- HACK, L., AND D. ROSTAM-AFSCHAR (2024): "Understanding Firm Dynamics with Daily Data," CRC TR 224 Discussion Paper Series 593, University of Bonn and University of Mannheim, Germany.
- KASY, M., AND A. SAUTMANN (2021): "Adaptive Treatment Assignment in Experiments for Policy Choice," *Econometrica*, 89(1), 113–132.
- KEMPER, J., AND D. ROSTAM-AFSCHAR (2025): "Inference for Batched Adaptive Experiments," Discussion paper, University of Mannheim.
- (2026): "Earning While Learning: How to Run Batched Bandit Experiments," Discussion Paper 3.
- THOMPSON, W. R. (1933): "On the likelihood that one unknown probability exceeds another in view of the evidence of two samples," *Biometrika*, 25(3-4), 285–294.
- THOMPSON, W. R. (1935): "On the Theory of Apportionment," *American Journal of Mathematics*, 57(2), 450–456.