

Supplemental Appendix

**Do Workfare Programs Live Up to Their
Promises?
Experimental Evidence from Côte d'Ivoire**

Marianne Bertrand, Bruno Crépon, Alicia Marguerie, Patrick Premand

Appendix A Additional Tables

Table A1: Overview of Public Works Programs

(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)	(13)
Country	Program	Public works program funding (\$m)	# of beneficiaries	% of country labor force	Max # of times the same beneficiary can participate	# of work days per year	Uses self-targeting (individuals apply)	Uses other targeting methods (1=no, universal, 2=community targeting, 3=proxy means test, 4=Lottery)	Daily wage (\$, PPP)	Total transfer per beneficiary per year (\$, PPP)	Has an objective to improve post-program outcomes	Delivers complementary training
Panel A. Active Public Works Programs in sub-Saharan Africa (with World Bank support, 2020)												
Benin	Community and Local Government Basic Social Services Project	1	16,200	0.35	2	48	0	3	4.7	225	1	1
Cameroon	Social Safety Nets for Crisis Response	6	30,000	0.32	1	60	1	3	4.1	249	1	1
Congo, Dem. Rep.	Eastern Recovery Project	19	12,000	0.05	1	80	1	2, 4	5.2	414	1	1
Comoros	Social Safety Net Project	3	5,890	2.80	3	60	1	2	4.9	295	1	1
Côte d'Ivoire	Emergency Youth Employment & Skills Development Project	18	12,500	0.18	1	132	1	4	10.1	1,365	1	1
Ghana	Productive Safety Net Project	28	30,000	0.24	2	90	1	2, 3	7.2	646	1	1
Ethiopia	Urban Productive Safety Net	171	1,564,416	3.24	3	40	1	2, 3	7.4	298	1	1
Liberia	Youth, Employment, Skills Project	9	45,000	0.43	1	40	1	2	6.0	239	1	1
Liberia	Youth Opportunities Project	4	10,000	0.08	1	100	1	2	5.8	579	1	1
Madagascar	Productive Safety Net Project	15	32,500	0.27	3	80	1	2, 3	3.4	271	1	0
Malawi	Social Action Fund	20	600,000	10.54	1	15	0	2	3.4	51	1	0
Niger	Safety Net Project	11	60,000	0.88	1	60	1	3	4.4	265	0	0
Niger	Adaptive Safety Net Project 2	19	66,000	0.78	2	60	1	2	5.1	307	1	1
Nigeria	Youth Employment and Social Support Operation	200	250,000	0.47	2	240	0	2, 3	3.5	831	1	1
Tanzania	Productive Social Safety Net	25	230,000	1.04	1	60	0	2, 3	4.6	277	1	1
Tanzania	Productive Social Safety Net II	75	837,573	3.08	1	60	0	2, 3	3.3	201	1	0
Uganda	Third Northern Uganda Social Action Fund	49	499,000	3.54	2	60	0	2	3.6	213	1	1
Average for projects in Panel A:		39	253,005	1.66	1.65	76	0.65		5.23	396	0.94	0.76
Total for projects in Panel A:		671	4,301,079									
Panel B. Most Studied Programs in Quasi-Experimental Literature												
India	NREGA National Rural Employment Guarantee	10,533	111,900,000	23.72	unlimited	35	1	1	9.1	317	0	0
Ethiopia	PSNP Public Works sub program	282	6,803,770	12.89	unlimited	60	0	2, 3	4.9	292	0	1

Panel A provides an overview of public works programs supported by the World Bank in Sub-Saharan Africa (in 2020). This provides a lower bound for investments in public works programs as it excludes top-up financing for these projects, domestic financing, and investments in cash-for-works or labor-intensive public works from other agencies and donors.

Table A2: Impacts *during* and *post* program, robustness to alternative definition of earnings

	Total earnings			Self-employment earnings			Wage employment earnings		
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)
	Ln, winsorized at 97%	Ln, winsorized at 99%	Levels, winsorized at 99%	Ln, winsorized at 97%	Ln, winsorized at 99%	Levels, winsorized at 99%	Ln, winsorized at 97%	Ln, winsorized at 99%	Levels, winsorized at 99%
Panel A: Impacts <i>during</i> the program (around 4,5 months after program starts)									
Public Works Treatment (ITT)	2.95*** (0.19)	2.95*** (0.20)	24380.2*** (5998.7)	-1.04*** (0.20)	-1.06*** (0.20)	-11303.4*** (4091.6)	5.93*** (0.25)	5.93*** (0.25)	37181.6*** (4033.5)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	7.87	7.89	54626.16	3.23	3.26	23508.66	4.62	4.64	28163.73
Observations	2912	2912	2912	2912	2912	2912	2912	2912	2912
Perm. p-value: no effects	0.000	0.000	0.000	0.000	0.000	0.010	0.000	0.000	0.000
Panel B: <i>Post</i> program impacts (pooled treatment) (12 to 15 months after program ends)									
Public Works Treatment (ITT)	-0.037 (0.18)	-0.028 (0.18)	7597.5*** (2380.5)	0.22 (0.23)	0.23 (0.23)	7802.7*** (2250.6)	-0.20 (0.19)	-0.20 (0.19)	-542.8 (1145.3)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	8.42	8.43	45690.44	3.56	3.57	20487.91	5.04	5.05	22036.94
Observations	3934	3934	3934	3934	3934	3934	3934	3934	3934
Perm. p-value: No effects	0.849	0.877	0.002	0.347	0.322	0.002	0.300	0.323	0.607
Panel C: <i>Post</i> program impacts (by treatment arm) (12 to 15 months after program ends)									
Public Works Treatment (PW)	-0.18 (0.22)	-0.18 (0.22)	5558.1** (2665.5)	0.13 (0.27)	0.14 (0.28)	5098.3** (2288.8)	-0.22 (0.24)	-0.23 (0.24)	-21.1 (1486.7)
Self-Empl. training (SET)	0.22 (0.26)	0.23 (0.26)	5171.9 (3709.3)	0.28 (0.34)	0.29 (0.35)	7999.6** (3763.3)	-0.065 (0.26)	-0.064 (0.26)	-1327.7 (1531.2)
Wage-Empl. training (WET)	0.24 (0.21)	0.24 (0.21)	1188.7 (3179.3)	-0.0027 (0.32)	-0.0025 (0.32)	439.3 (3699.5)	0.15 (0.27)	0.15 (0.28)	-299.5 (1544.4)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	8.42	8.43	45690.44	3.56	3.57	20487.91	5.04	5.05	22036.94
p-value PW+SET=0	0.878	0.840	0.003	0.184	0.168	0.000	0.246	0.246	0.318
p-value PW+WET=0	0.818	0.787	0.049	0.706	0.689	0.165	0.764	0.770	0.826
p-value SET=WET	0.963	0.972	0.392	0.462	0.451	0.173	0.456	0.452	0.476
Observations	3934	3934	3934	3934	3934	3934	3934	3934	3934
Perm. p-value PW+SET=0	0.898	0.832	0.005	0.185	0.169	0.000	0.249	0.250	0.349
Perm. p-value PW+WET=0	0.822	0.778	0.054	0.727	0.698	0.176	0.771	0.787	0.818
Perm. p-value SET=WET	0.955	0.961	0.405	0.486	0.474	0.169	0.459	0.438	0.502

ITT estimates in panels A and B based on specification in equation 1. Estimates in panel C based on specification in equation 2. Robust standard errors clustered at (broad) brigade level in parentheses. For variables (y) in logarithms, we take $\ln(y+1)$. Permutation tests use 1000 permutations for each hypothesis. * $p < .1$, ** $p < .05$, *** $p < .01$.

Table A3: Estimated impacts *during* and *post* program on economic outcomes, with baseline controls

	(1) Employed	(2) Wage employed (in at least 1 activity)	(3) Self- employed (in at least 1 activity)	(4) Total hours worked (weekly)	(5) Hours worked in wage empl. (weekly)	(6) Hours worked in self-Empl. (weekly)	(7) Total earnings in CFA (monthly)	(8) Ln total earnings (monthly)	(9) Earnings in wage empl. in CFA (monthly)	(10) Earnings in self-empl. in CFA (monthly)	(11) Total ex- penditures in CFA (monthly)	(12) Savings in CFA (stock)	(13) Well-being index (z-score)
Panel A: Impacts <i>during</i> the program (around 4.5 months after program starts)													
Public Works Treatment (ITT)	0.14*** (0.01)	0.48*** (0.02)	-0.09*** (0.02)	5.34*** (1.23)	15.89*** (1.24)	-5.60*** (0.93)	27485.77*** (2608.25)	2.92*** (0.19)	36799.02*** (1463.69)	-5567.13*** (1196.97)	14417.38*** (1314.87)	40162.30*** (2305.56)	0.19*** (0.05)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Baseline controls	Yes	No	Yes	Yes	Yes	Yes	Yes	Yes	No	Yes	Yes	Yes	Yes
Mean in Control	0.85	0.49	0.35	39.18	20.79	11.28	42841.22	7.87	20188.33	12753.65	47233.52	19250.05	-0.03
Observations	2958	2958	2958	2958	2958	2958	2912	2912	2912	2912	2945	2958	2934
Panel B: <i>Post</i> program impacts (pooled treatment) (12 to 15 months after program ends)													
Public Works Treatment (ITT)	0.01 (0.02)	0.00 (0.02)	0.02 (0.02)	0.52 (1.33)	-1.42 (1.22)	2.02* (1.14)	5155.88*** (1902.00)	-0.01 (0.19)	-665.49 (1079.42)	4783.59*** (1852.73)	2387.97 (1466.46)	10143.10*** (3316.16)	0.12*** (0.04)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Baseline controls	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.86	0.51	0.36	40.49	22.06	13.26	43481.10	8.42	20706.18	18872.95	50700.71	46348.14	-0.05
Observations	3934	3934	3934	3934	3934	3934	3934	3934	3934	3934	3814	3934	3932
Panel C: <i>Post</i> program impacts (by treatment arms) (12 to 15 months after program ends)													
Public Works Treatment (PW)	0.01 (0.02)	0.00 (0.03)	0.02 (0.03)	-1.61 (1.72)	-1.55 (1.60)	0.20 (1.31)	3406.18 (2116.14)	-0.15 (0.22)	69.13 (1319.73)	3071.78 (1901.47)	2131.18 (1596.09)	8598.92** (3644.89)	0.15*** (0.05)
Self-Empl. training (SET)	0.01 (0.02)	-0.02 (0.03)	0.02 (0.03)	3.36* (1.95)	0.49 (1.77)	2.62 (1.71)	4226.69 (3044.51)	0.20 (0.24)	-1668.02 (1287.18)	5203.42* (2826.96)	-207.36 (1422.57)	8143.34* (4479.69)	-0.01 (0.05)
Wage-Empl. training (WET)	0.00 (0.02)	0.01 (0.03)	0.00 (0.03)	3.05 (2.17)	-0.16 (1.61)	2.93 (1.82)	947.80 (1963.81)	0.24 (0.19)	-796.06 (1313.08)	210.73 (2173.59)	881.90 (1599.02)	-3206.44 (4242.97)	-0.08* (0.05)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Baseline controls	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.86	0.51	0.36	40.49	22.06	13.26	43481.10	8.42	20706.18	18872.95	50700.71	46348.14	-0.05
Observations	3934	3934	3934	3934	3934	3934	3934	3934	3934	3934	3814	3934	3932

ITT estimates in panels A and B based on specification in equation 1. Estimates in panel C based on specification in equation 2. The set of baseline controls differs for each outcome. Variables are selected from a pool of 1257 covariates using post-double selection lasso. Control variables include information about individual characteristics, education, household composition, experience of violence, household expenditure, asset ownership, and access to infrastructure. Robust standard errors clustered at (broad) brigade level in parentheses. Hours, earnings, expenditures, and savings winsorized at the 97th percentile. * $p < .1$, ** $p < .05$, *** $p < .01$.

Table A4: Estimated impacts *during* and *post* program on economic outcomes, without clustering standard errors

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)
	Employed	Wage employed (in at least 1 activity)	Self employed (in at least 1 activity)	Total hours worked (weekly)	Hours worked in wage empl. (weekly)	Hours worked in self-empl. (weekly)	Total earnings in CFA (monthly)	Ln total earnings (monthly)	Earnings in wage empl. in CFA (monthly)	Earnings in self empl. in CFA (monthly)	Total ex- penditures in CFA (monthly)	Savings in CFA (stock)
Panel A: Impacts <i>during</i> the program (around 4,5 months after program starts)												
Public Works Treatment (ITT)	0.14*** (0.013)	0.48*** (0.017)	-0.096*** (0.019)	4.89*** (1.16)	15.5*** (1.00)	-5.69*** (0.80)	27082.9*** (2575.7)	2.95*** (0.16)	36799.0*** (1302.7)	-5715.4*** (1100.6)	14529.3*** (1490.0)	39785.7*** (2017.9)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.85	0.49	0.35	39.18	20.79	11.28	42841.22	7.87	20188.33	12753.65	47233.52	19250.05
Observations	2958	2958	2958	2958	2958	2958	2912	2912	2912	2912	2945	2958
Perm. p-value: no effects	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
Panel B: <i>Post</i> program impacts (pooled treatment) (12 to 15 months after program ends)												
Public Works Treatment (ITT)	0.015 (0.015)	0.0068 (0.019)	0.010 (0.019)	1.34 (1.20)	-0.61 (1.07)	1.70* (1.01)	4360.6** (2038.5)	-0.037 (0.18)	-452.7 (1007.9)	4005.2** (1763.0)	1361.7 (1394.0)	11505.2*** (3022.1)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.86	0.51	0.36	40.49	22.06	13.26	43481.10	8.42	20706.18	18872.95	50700.71	46348.14
Observations	3934	3934	3934	3934	3934	3934	3934	3934	3934	3934	3814	3934
Perm. p-value: No effects	0.310	0.755	0.646	0.293	0.610	0.139	0.027	0.844	0.652	0.028	0.349	0.001
Panel C: <i>Post</i> program impacts (by treatment arm) (12 to 15 months after program ends)												
Public Works Treatment (PW)	0.011 (0.018)	0.0081 (0.023)	0.0035 (0.023)	-0.76 (1.41)	-0.71 (1.30)	-0.12 (1.17)	2800.5 (2308.5)	-0.18 (0.22)	312.2 (1289.6)	2168.7 (1964.9)	925.7 (1701.6)	10429.5*** (3639.5)
Self-Empl. training (SET)	0.011 (0.017)	-0.018 (0.023)	0.021 (0.025)	3.42** (1.55)	0.46 (1.31)	2.77* (1.44)	4229.3 (3664.3)	0.22 (0.24)	-1591.8 (1266.1)	5595.5* (3051.3)	278.1 (1774.0)	7169.5* (4079.4)
Wage-Empl. training (WET)	0.0018 (0.018)	0.014 (0.024)	0.000048 (0.024)	3.12** (1.49)	-0.14 (1.33)	2.89** (1.30)	637.5 (2641.8)	0.24 (0.22)	-792.6 (1302.2)	135.8 (2223.9)	1077.6 (1775.7)	-3798.3 (3736.3)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.86	0.51	0.36	40.49	22.06	13.26	43481.10	8.42	20706.18	18872.95	50700.71	46348.14
p-value PW+SET=0	0.231	0.666	0.322	0.092	0.849	0.066	0.048	0.868	0.285	0.010	0.492	0.000
p-value PW+WET=0	0.485	0.344	0.883	0.121	0.522	0.034	0.168	0.813	0.698	0.286	0.254	0.074
p-value SET=WET	0.617	0.175	0.408	0.854	0.651	0.940	0.342	0.956	0.510	0.086	0.661	0.008
Observations	3934	3934	3934	3934	3934	3934	3934	3934	3934	3934	3814	3934
Perm. p-value PW+SET=0	0.273	0.665	0.391	0.159	0.859	0.168	0.015	0.865	0.303	0.004	0.457	0.001
Perm. p-value PW+WET=0	0.474	0.375	0.914	0.211	0.568	0.126	0.246	0.815	0.710	0.427	0.297	0.170
Perm. p-value SET=WET	0.675	0.224	0.511	0.914	0.732	0.965	0.366	0.959	0.562	0.180	0.634	0.050

ITT estimates in panels A and B based on specification in equation 1. Estimates in panel C based on specification in equation 2. Robust standard errors are in parentheses. Standard errors are not clustered by brigade. *Hours*, *earnings*, *expenditures*, and *savings* winsorized at the 97th percentile. For variables (y) in logarithms, we take $\ln(y + 1)$. Permutation tests use 1000 permutations for each hypothesis. * $p < .1$, ** $p < .05$, *** $p < .01$.

Table A5: Impacts *during* and *post* program, with full baseline sample at follow-up (ITT)

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)
	Employed	Wage employed (in at least 1 activity)	Self employed (in at least 1 activity)	Total hours worked (weekly)	Hours worked in wage empl. (weekly)	Hours worked in self-empl. (weekly)	Total earnings in CFA (monthly)	Ln total earnings (monthly)	Earnings in wage empl. in CFA (monthly)	Earnings in self-empl. in CFA (monthly)	Total ex- penditures in CFA (monthly)	Savings in CFA (stock)
Panel A: Impacts <i>during</i> the program (around 4,5 months after program starts)												
Public Works Treatment (ITT)	0.14*** (0.015)	0.48*** (0.024)	-0.096*** (0.020)	4.89*** (1.25)	15.5*** (1.29)	-5.69*** (0.94)	27082.9*** (2824.9)	2.95*** (0.19)	36799.0*** (1472.5)	-5715.4*** (1214.6)	14529.3*** (1441.4)	39785.7*** (2389.2)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.85	0.49	0.35	39.18	20.79	11.28	42841.22	7.87	20188.33	12753.65	47233.52	19250.05
Observations	2958	2958	2958	2958	2958	2958	2912	2912	2912	2912	2945	2958
Perm. p-value: no effects	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
Panel B: <i>Post</i> program impacts (pooled treatment) (12 to 15 months after program ends)												
Public Works Treatment (ITT)	0.0079 (0.014)	0.0069 (0.019)	0.020 (0.024)	0.73 (1.28)	-0.59 (1.15)	1.59 (1.06)	5030.5** (1996.1)	-0.072 (0.19)	0.93 (978.4)	4504.1** (1907.9)	1950.4 (1452.2)	8953.4*** (3390.6)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.86	0.51	0.35	41.02	21.97	13.34	42690.67	8.44	20132.39	18357.50	50010.79	48816.10
Observations	3910	3910	3910	3910	3910	3910	3910	3910	3910	3910	3788	3910
Perm. p-value: No effects	0.590	0.726	0.392	0.578	0.610	0.131	0.013	0.726	0.999	0.017	0.184	0.016
Panel C: <i>Post</i> program impacts (by treatment arm) (12 to 15 months after program ends)												
Public Works Treatment (PW)	0.0037 (0.019)	0.0085 (0.025)	0.013 (0.028)	-1.39 (1.68)	-0.69 (1.53)	-0.24 (1.19)	3461.2 (2236.8)	-0.22 (0.23)	770.0 (1229.9)	2655.2 (1955.1)	1498.0 (1627.1)	7875.4** (3354.3)
Self-Empl. training (SET)	0.011 (0.017)	-0.018 (0.026)	0.021 (0.031)	3.43* (1.83)	0.46 (1.77)	2.78* (1.59)	4220.4 (3135.6)	0.22 (0.25)	-1604.8 (1239.2)	5597.6** (2782.6)	272.2 (1483.8)	7103.4 (4515.6)
Wage-Empl. training (WET)	0.0022 (0.020)	0.013 (0.026)	0.00081 (0.031)	3.17 (2.18)	-0.14 (1.61)	2.92 (1.85)	674.6 (2169.2)	0.24 (0.21)	-792.5 (1247.4)	172.0 (2372.0)	1134.1 (1779.3)	-3725.1 (4493.3)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.86	0.51	0.35	41.02	21.97	13.34	42690.67	8.44	20132.39	18357.50	50010.79	48816.10
p-value PW+SET=0	0.388	0.677	0.218	0.245	0.884	0.121	0.009	0.998	0.474	0.002	0.280	0.001
p-value PW+WET=0	0.760	0.348	0.675	0.357	0.543	0.148	0.150	0.930	0.985	0.369	0.170	0.403
p-value SET=WET	0.682	0.208	0.530	0.914	0.699	0.955	0.359	0.935	0.518	0.174	0.586	0.034
Observations	3910	3910	3910	3910	3910	3910	3910	3910	3910	3910	3788	3910
Perm. p-value PW+SET=0	0.409	0.687	0.218	0.261	0.885	0.141	0.008	0.998	0.456	0.004	0.301	0.001
Perm. p-value PW+WET=0	0.761	0.351	0.686	0.345	0.571	0.148	0.151	0.935	0.990	0.357	0.185	0.386
Perm. p-value SET=WET	0.696	0.240	0.555	0.915	0.712	0.952	0.379	0.945	0.490	0.196	0.613	0.046

ITT estimates in panels A and B based on specification in equation 1. Individuals from the control group that were able to apply and participate in the third or fourth wave of the program are kept in the sample, and control individuals who applied to those waves but were not selected are not given larger weights. Robust standard errors clustered at brigade level in parentheses. *Hours*, *earnings*, *expenditures*, and *savings* winsorized at the 97th percentile. For variables (y) in logarithms, we take $\ln(y + 1)$. Permutation tests use 1000 permutations for each hypothesis. * $p < .1$, ** $p < .05$, *** $p < .01$.

Table A6: Impacts *post* program, LATE accounting for non-compliance in control between midline and endline

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)
	Employed	Wage employed (in at least 1 activity)	Self- employed (in at least 1 activity)	Total hours worked (weekly)	Hours worked in wage empl. (weekly)	Hours worked in self-empl. (weekly)	Total earnings in CFA (monthly)	Ln total earnings (monthly)	Earnings in wage empl. in CFA (monthly)	Earnings in self-empl. in CFA (monthly)	Total ex- penditures in CFA (monthly)	Savings in CFA (stock)
Panel A: <i>Post</i> program impacts (pooled treatment) (12 to 15 months after program ends)												
Public Works Treatment (LATE)	0.0089 (0.016)	0.0078 (0.022)	0.023 (0.027)	0.83 (1.45)	-0.66 (1.29)	1.80 (1.19)	5684.9** (2252.2)	-0.082 (0.21)	1.05 (1100.7)	5090.0** (2153.1)	2202.8 (1634.7)	10118.0*** (3820.0)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.86	0.51	0.35	41.02	21.97	13.34	42690.67	8.44	20132.39	18357.50	50010.79	48816.10
Observations	3910	3910	3910	3910	3910	3910	3910	3910	3910	3910	3788	3910
Perm. p-value: No effects	0.128	0.780	0.543	0.383	0.273	0.064	0.186	0.629	0.333	0.090	0.496	0.000

Estimates of Local Average Treatment Effects (LATE) accounting for one-sided non-compliance, whereby some control individuals were able to participate in the program between midline and endline. All treatment individuals are considered as participating in the program. Participation in the program is instrumented by the randomized treatment assignment. Robust standard errors clustered at brigade level in parentheses. *Hours*, *earnings*, *expenditures*, and *savings* winsorized at the 97th percentile. For variables (y) in logarithms, we take $\ln(y + 1)$. Permutation tests use 1000 permutations for each hypothesis. * $p < .1$, ** $p < .05$, *** $p < .01$.

Table A7: Estimated impacts *during* program on economic outcomes for household members other than the beneficiary

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)
	Employed	Wage- employed (in at least 1 activity)	Self-employed (in at least 1 activity)	Total hours worked (weekly)	Hours worked in wage empl. (weekly)	Hours worked in self-empl. (weekly)	Total earning in CFA (monthly)	Earnings in wage empl. in CFA (monthly)	Earnings in self-empl. in CFA (monthly)
Panel A: Impacts <i>during</i> the program (around 4,5 months after program starts)									
Public Works Treatment (ITT)	-0.0079 (0.015)	-0.010 (0.013)	0.0091 (0.012)	0.0059 (0.88)	-0.20 (0.65)	0.39 (0.58)	-28.1 (1483.6)	-340.2 (628.4)	663.7 (757.9)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.59	0.27	0.25	25.03	11.67	9.88	23673.33	8493.04	9578.04
Observations	9929	9895	9895	9929	9895	9895	9772	9739	9739
Perm. p-value: no effects	0.634	0.404	0.423	0.993	0.759	0.488	0.984	0.586	0.377

ITT estimates in panel A based on specification in equation 1. Based on information of households that were included at baseline, even if the beneficiary moved at midline. Collective households and attritors for the household survey at midline were excluded. Robust standard errors clustered at (broad) brigade level in parentheses. *Hours* and *earnings* winsorized at the 97th percentile. Permutation tests use 1000 permutations for each hypothesis. * $p < .1$, ** $p < .05$, *** $p < .01$.

Table A8: Estimated impacts *during* and *post* program on well-being index components

	(1)	(2)	(3)	(4)	(5)	(6)	(7)
	Well-being index (z-score)	Self-esteem (Rosenberg scale) (z-score)	Positive affect (CES-D sub scale) (z-score)	Positive attitude towards the future (ZTPI sub scale) (z-score)	Present fatalism (ZTPI sub scale) (z-score)	Happiness in daily activities (z-score)	Pride in daily activities (z-score)
Panel A: Impacts <i>during</i> the program (around 4.5 months after program starts)							
Public Works Treatment (ITT)	0.20*** (0.05)	0.14*** (0.04)	0.18*** (0.04)	0.086** (0.04)	0.021 (0.04)	0.15*** (0.05)	0.15*** (0.05)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	-0.03	-0.02	-0.05	-0.01	0.03	0.01	-0.00
Observations	2934	2951	2958	2951	2955	2950	2949
Perm. p-value: no effects	0.000	0.000	0.000	0.032	0.605	0.002	0.000
Panel B: <i>Post</i> program impacts (pooled treatment) (12 to 15 months after program ends)							
Public Works Treatment (ITT)	0.11*** (0.04)	0.10** (0.04)	0.041 (0.04)	0.061 (0.05)	-0.093** (0.04)	0.076* (0.04)	0.053 (0.04)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	-0.05	-0.05	-0.03	-0.01	0.05	-0.02	-0.02
Observations	3932	3933	3932	3933	3933	3933	3933
Perm. p-value: No effects	0.005	0.024	0.332	0.196	0.030	0.078	0.182
Panel C: <i>Post</i> program impacts (by treatment arm) (12 to 15 months after program ends)							
Public Works Treatment (PW)	0.14*** (0.05)	0.12** (0.05)	0.063 (0.05)	0.094* (0.05)	-0.052 (0.05)	0.11** (0.05)	0.089* (0.05)
Self-Empl. training (SET)	-0.0068 (0.05)	-0.044 (0.06)	-0.0078 (0.05)	-0.11* (0.06)	-0.11** (0.05)	-0.018 (0.05)	-0.025 (0.05)
Wage-Empl. training (WET)	-0.075 (0.05)	-0.019 (0.05)	-0.062 (0.05)	0.0034 (0.05)	-0.020 (0.05)	-0.082* (0.05)	-0.087* (0.05)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	-0.05	-0.05	-0.03	-0.01	0.05	-0.02	-0.02
p-value: PW+SET=0	0.011	0.165	0.275	0.816	0.001	0.059	0.177
p-value: PW+WET=0	0.187	0.065	0.977	0.046	0.201	0.598	0.972
p-value: SET=WET	0.169	0.703	0.232	0.038	0.073	0.138	0.216
Observations	3932	3933	3932	3933	3933	3933	3933
Perm. p-value PW+SET=0	0.012	0.167	0.285	0.819	0.000	0.066	0.184
Perm. p-value PW+WET=0	0.204	0.076	0.971	0.042	0.222	0.593	0.963
Perm. p-value SET=WET	0.191	0.695	0.227	0.044	0.092	0.145	0.235

ITT estimates in panels A and B based on specification in equation 1. Estimates in panel C based on specification in equation 2. The definition of the *well-being index* and variables entering the index is detailed in Appendix D. *Present fatalism* enters as an inverted measure in the index (a negative impact in column (5) is associated with a positive impact on the index in column (1)). Robust standard errors clustered at (broad) brigade level in parentheses. Permutation tests use 10000 permutations for each hypothesis. * $p < .1$, ** $p < .05$, *** $p < .01$.

Table A9: Estimated impacts *during* and *post* program on behavior index components

	(1) Behavior index (z-score)	(2) Conduct problems (SDQ sub scale) (z-score)	(3) Pro-social behavior (SDQ sub scale) (z-score)	(4) Impulsiveness (DERS sub scale) (z-score)	(5) Anger in daily activities (z-score)
Panel A: Impacts <i>during</i> the program (around 4.5 months after program starts)					
Public Works Treatment (ITT)	0.12*** (0.04)	-0.031 (0.04)	0.023 (0.04)	-0.095** (0.04)	-0.13*** (0.04)
Strata f.e.	Yes	Yes	Yes	Yes	Yes
Mean in Control	-0.02	0.01	-0.00	0.02	-0.01
Observations	2946	2957	2956	2954	2950
Perm. p-value: no effects	0.008	0.444	0.573	0.040	0.004
Panel B: <i>Post</i> program impacts (pooled treatment) (12 to 15 months after program ends)					
Public Works Treatment (ITT)	-0.012 (0.04)	0.013 (0.04)	-0.0032 (0.04)	0.0050 (0.04)	0.014 (0.04)
Strata f.e.	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.01	-0.01	-0.01	-0.01	-0.03
Observations	3933	3933	3933	3933	3933
Perm. p-value: No effects	0.775	0.769	0.925	0.904	0.728
Panel C: <i>Post</i> program impacts (by treatment arm) (12 to 15 months after program ends)					
Public Works Treatment (PW)	0.025 (0.05)	0.034 (0.05)	0.0066 (0.06)	-0.051 (0.04)	-0.013 (0.05)
Self-Empl. training (SET)	-0.039 (0.04)	-0.062* (0.04)	-0.054 (0.06)	0.054 (0.04)	0.012 (0.05)
Wage-Empl. training (WET)	-0.077 (0.05)	-0.0041 (0.05)	0.024 (0.07)	0.12*** (0.04)	0.072 (0.05)
Strata f.e.	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.01	-0.01	-0.01	-0.01	-0.03
p-value: PW+SET=0	0.762	0.530	0.317	0.939	0.978
p-value: PW+WET=0	0.293	0.574	0.546	0.147	0.222
p-value: SET=WET	0.410	0.197	0.066	0.165	0.229
Observations	3933	3933	3933	3933	3933
Perm. p-value PW+SET=0	0.763	0.516	0.326	0.943	0.971
Perm. p-value PW+WET=0	0.288	0.564	0.556	0.157	0.228
Perm. p-value SET=WET	0.401	0.215	0.065	0.159	0.242

ITT estimates in panels A and B based on specification in equation 1. Estimates in panel C based on specification in equation 2. The definition of the *behavior index* and variables entering the index is detailed in Appendix D. *Conduct problems*, *impulsiveness* and *anger in daily activities* enter as inverted measures in the index (a negative impact in columns (2), (4) or (5) is associated with a positive impact on the index in column (1)). Robust standard errors clustered at (broad) brigade level in parentheses. Permutation tests use 10000 permutations for each hypothesis. * $p < .1$, ** $p < .05$, *** $p < .01$.

Table A10: Estimated impacts *during* and *post* program on risky behaviors

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)
	Stealing	Assaulting someone	Believing smuggling is necessary (to earn a living)	Prostituting	Threatening someone	Taking illicit drugs	Smuggling stolen objects	Ties with a smuggling network	Keeping fire arms at home
Panel A: Impacts <i>during</i> the program (around 4.5 months after program starts)									
Public Works Treatment (ITT)	0.09*	-0.11**	0.00	0.01	-0.02	-0.11***	-0.05	0.01	0.01
	(0.05)	(0.04)	(0.04)	(0.05)	(0.04)	(0.04)	(0.04)	(0.04)	(0.04)
Mean in Control	0.05	0.18***	0.10***	0.07*	0.11***	0.03	0.05*	0.02	0.01
	(0.04)	(0.03)	(0.03)	(0.04)	(0.03)	(0.04)	(0.03)	(0.04)	(0.04)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
p-value: no design effect in list A	1.000	1.000	0.446	0.008	0.156	0.015	0.314	0.485	0.152
p-value: no design effect in list B	1.000	1.000	0.655	0.373	1.000	0.425	0.645	0.107	0.856
Impact in list A	0.09	-0.19	-0.07	0.03	-0.11	-0.08	-0.14	0.04	-0.01
Impact in list B	0.08	-0.03	0.08	-0.02	0.07	-0.15	0.04	-0.01	0.03
p-value: Impact A= Impact B	0.930	0.097	0.126	0.694	0.070	0.501	0.046	0.587	0.677
Observations	2956	2956	2956	2956	2953	2955	2954	2954	2955
Panel B: <i>Post</i> program impacts (pooled treatment) (12 to 15 months after program ends)									
Public Works Treatment (ITT)	0.03	-0.01	0.01	-0.01	-0.01	0.01	-0.09***	0.00	0.00
	(0.03)	(0.03)	(0.03)	(0.04)	(0.03)	(0.03)	(0.03)	(0.03)	(0.03)
Mean in Control	0.04	0.09***	0.08***	0.10***	0.08***	0.03	0.09***	0.06**	0.05**
	(0.03)	(0.03)	(0.02)	(0.03)	(0.03)	(0.03)	(0.03)	(0.03)	(0.03)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
p-value: no design effect in list A	0.083	—	0.317	—	1.000	0.157	0.063	0.510	0.981
p-value: no design effect in list B	—	1.000	1.000	0.083	0.214	0.171	0.612	0.425	0.611
Impact in list A	0.01	0.04	0.02	-0.02	-0.01	-0.04	-0.09	0.07	0.05
Impact in list B	0.05	-0.07	-0.00	-0.00	-0.01	0.06	-0.08	-0.07	-0.04
p-value: Impact A=Impact B	0.592	0.108	0.752	0.743	0.945	0.144	0.915	0.058	0.190
Observations	3933	3933	3933	3933	3933	3932	3933	3933	3933
Panel C: <i>Post</i> program impacts (by treatment arm) (12 to 15 months after program ends)									
Public Works Treatment (PW)	0.01	0.03	0.03	-0.05	0.00	0.02	-0.06	0.01	0.01
	(0.04)	(0.04)	(0.03)	(0.04)	(0.04)	(0.04)	(0.04)	(0.04)	(0.04)
Self-Empl. training (SET)	0.03	-0.08	-0.04	0.10***	0.02	-0.05	-0.07*	0.01	-0.02
	(0.04)	(0.05)	(0.04)	(0.04)	(0.05)	(0.03)	(0.04)	(0.04)	(0.04)
Wage-Empl. training (WET)	0.01	-0.06*	-0.02	0.01	-0.06	0.01	-0.01	-0.03	0.01
	(0.04)	(0.04)	(0.04)	(0.03)	(0.04)	(0.04)	(0.04)	(0.04)	(0.04)
Mean in Control	0.04	0.09***	0.08***	0.10***	0.08***	0.03	0.09***	0.06**	0.05**
	(0.03)	(0.03)	(0.02)	(0.03)	(0.03)	(0.03)	(0.03)	(0.03)	(0.03)
Strata f.e.	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
p-value PW+SET=0	0.275	0.344	0.662	0.197	0.534	0.498	0.001	0.708	0.681
p-value PW+WET=0	0.539	0.419	0.820	0.316	0.194	0.450	0.072	0.571	0.668
p-value SET=WET	0.667	0.689	0.613	0.008	0.077	0.188	0.185	0.416	0.388
Observations	3933	3933	3933	3933	3933	3932	3933	3933	3933

Variables measured using a double list experiment (Droitcour et al., 1991), whereby each respondent was assigned to a list A with sensitive items and a list B without sensitive items, or vice versa (see Appendix D). Difference-in-means estimation (Miller, 1984) was used to estimate the mean in control and treatment effects. In Panels A and B, the specification $Y_i = \alpha + \gamma_1 L_i + \gamma_2 W_i + \gamma_3 (L_i \times W_i) + \delta_i X_{i,l} + \varepsilon_i$ was used; where γ_1 is the mean in the control group, γ_3 is the treatment effect, and $X_{i,l}$ is a vector of stratification variables. Similarly, in Panel C we used $Y_i = \alpha + \gamma_1 L_i + \gamma_2 W_i + \gamma_3 (W_i \times T1_i) + \gamma_4 (W_i \times T2_i) + \gamma_5 (L_i \times W_i) + \gamma_6 (L_i \times W_i \times T1_i) + \gamma_7 (L_i \times W_i \times T2_i) + \delta_i X_{i,l} + \varepsilon_i$; where γ_1 is the mean in the control group, γ_5 the effect of "pure" public works, and γ_6 (γ_7) the additional effect of self-employment training (wage job search training). The test for the presence of design effects is based on the likelihood ratio test (Blair and Imai, 2012). The null hypothesis is no design effect. We report Bonferroni adjusted p-values. Weights are used for estimation but not for the design effect test (because it is not supported). The dash symbol indicates that the test statistics could not be processed due to a lack of variance in estimated probabilities: $P(C = 4, S = 1) = P(C = 4, S = 0) = 0$. However, none of the point estimates of joint probabilities were negative in such cases. * $p < .1$, ** $p < .05$, *** $p < .01$.

Table A11: Impacts *during* and *post* program on time use

	(1) Rest at 6 am of prev. day	(2) Leisure at 6 am of prev. day	(3) Work at 6 am of prev. day	(4) Rest at 10 pm of prev. day	(5) Leisure at 10 pm of prev. day	(6) Work at 10 pm of prev. day
Panel A: Impacts <i>during</i> the program (around 4,5 months after program starts)						
Public Works Treatment (ITT)	-0.14*** (0.018)	-0.044*** (0.011)	0.33*** (0.032)	0.071*** (0.019)	-0.054*** (0.015)	-0.028** (0.012)
LocXGender control	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.24	0.09	0.20	0.68	0.18	0.08
Observations	2955	2955	2955	2953	2953	2953
Perm. p-value: no effects	0.000	0.001	0.000	0.001	0.001	0.014
Panel B: <i>Post</i> program impacts (pooled treatment) (12 to 15 months after program ends)						
Public Works Treatment (ITT)	0.025 (0.019)	-0.0029 (0.012)	0.0080 (0.016)	-0.000063 (0.018)	0.0011 (0.015)	0.0061 (0.011)
LocXGender control	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.26	0.09	0.24	0.68	0.19	0.09
Observations	3933	3933	3933	3933	3933	3933
Perm. p-value: No effects	0.206	0.804	0.624	0.997	0.932	0.561
Panel C: <i>Post</i> program impacts (by treatment arms) (12 to 15 months after program ends)						
Public Works Treatment (ITT)	0.021 (0.022)	0.0036 (0.014)	0.021 (0.021)	-0.00080 (0.023)	0.0037 (0.019)	0.0041 (0.014)
Self-empl.training (SET)	-0.0075 (0.032)	-0.0036 (0.013)	-0.016 (0.023)	0.0027 (0.023)	-0.0085 (0.019)	0.0030 (0.012)
Wage-empl. training (WET)	0.022 (0.026)	-0.017 (0.014)	-0.025 (0.022)	-0.00038 (0.022)	0.00041 (0.020)	0.0033 (0.015)
LocXGender control	Yes	Yes	Yes	Yes	Yes	Yes
Mean in Control	0.26	0.09	0.24	0.68	0.19	0.09
p-value PW+SET=0	0.648	0.995	0.798	0.929	0.790	0.579
p-value PW+WET=0	0.076	0.383	0.834	0.957	0.839	0.583
p-value SET=WET	0.230	0.385	0.652	0.895	0.680	0.981
Observations	3933	3933	3933	3933	3933	3933
Perm. p-value PW+SET=0	0.649	0.999	0.795	0.938	0.793	0.616
Perm. p-value PW+WET=0	0.079	0.375	0.838	0.951	0.840	0.585
Perm. p-value SET=WET	0.236	0.401	0.644	0.894	0.692	0.985

ITT estimates in panels A and B based on specification in equation 1. Estimates in panel C based on specification in equation 2. Robust standard errors clustered at (broad) brigade level in parentheses. Permutation tests use 1000 permutations for each hypothesis. * $p < .1$, ** $p < .05$, *** $p < .01$.

Table A12: Estimated impacts *post* program on aspirations and reservation wage

	(1)	(2)	(3)	(4)
	Aspires to wage job in 20 years	Aspires to self- employment in 20 years	Expected earnings in wage empl. in CFA (monthly)	Expected earnings in self-empl. in CFA (monthly)
Panel A: <i>Post</i> program impacts (pooled treatment) (12 to 15 months after program ends)				
Public Works Treatment (ITT)	0.0013 (0.01)	-0.0017 (0.01)	-3090.0 (2159.72)	-5702.2 (3529.50)
Strata f.e.	Yes	Yes	Yes	Yes
Mean in Control	0.17	0.83	107539.41	118806.05
Observations	3934	3934	3934	3934
Perm. p-value: no effects	0.923	0.913	0.151	0.115
Panel B: <i>Post</i> program impacts (by treatment arm) (12 to 15 months after program ends)				
Public Works Treatment (PW)	0.011 (0.02)	-0.012 (0.02)	-1398.7 (2540.13)	-1975.7 (4293.28)
Self-Empl. training (SET)	-0.036** (0.02)	0.035** (0.02)	-2422.2 (2903.83)	-5451.8 (4423.61)
Wage-Empl. training (WET)	0.0044 (0.01)	-0.0039 (0.01)	-2846.1 (2271.34)	-6156.1 (4346.66)
Strata f.e.	Yes	Yes	Yes	Yes
Mean in Control	0.17	0.83	107539.41	118806.05
p-value PW+SET=0	0.190	0.209	0.177	0.107
p-value PW+WET=0	0.339	0.346	0.104	0.047
p-value SET=WET	0.022	0.028	0.882	0.871
Observations	3934	3934	3934	3934
Perm. p-value PW+SET=0	0.203	0.224	0.193	0.112
Perm. p-value PW+WET=0	0.343	0.375	0.130	0.061
Perm. p-value SET=WET	0.035	0.033	0.886	0.880

ITT estimates in panels A and B based on specification in equation 1. Estimates in panel C based on specification in equation 2. Robust standard errors clustered at (broad) brigade level in parentheses. Permutation tests use 10000 permutations for each hypothesis. * $p < .1$, ** $p < .05$, *** $p < .01$.

Table A13: Comparison of Machine Learning algorithms to predict impacts on earnings *during* and *post* program

	Estimates in logs					Estimates in levels				
	(1) Elastic net	(2) Generalized Random forest	(3) Gradient boosting	(4) R-Learner	(5) Random forest	(6) Elastic net	(7) Generalized Random forest	(8) Gradient boosting	(9) R-Learner	(10) Random forest
Panel A: Predicted average treatment effect and heterogeneity loading parameter <i>during</i> program										
ATE (β_1)	2.631 (2.9,2.3) [0.000]	2.642 (2.9,2.3) [0.000]	2.643 (2.9,2.3) [0.000]	2.635 (2.9,2.3) [0.000]	2.634 (2.9,2.3) [0.000]	24577.6 (31323.2,17831.5) [0.000]	24351.6 (31214.6,17669.0) [0.000]	24238.5 (30976.0,17471.1) [0.000]	24549.6 (31423.5,17789.9) [0.000]	24363.0 (31150.3,17618.5) [0.000]
HET (β_2)	0.980 (1.3,0.6) [0.000]	1.231 (1.6,0.9) [0.000]	0.420 (0.6,0.2) [0.000]	0.920 (1.2,0.6) [0.000]	0.849 (1.1,0.6) [0.000]	0.405 (0.9,-0.07) [0.157]	1.070 (2.1,-0.05) [0.121]	0.288 (0.7,-0.03) [0.160]	0.429 (1.3,-0.2) [0.353]	0.390 (0.8,0.009) [0.089]
Λ	0.9	1.0	0.7	0.9	0.9	6007.2	6401.7	5988.1	4621.1	6785.3
Panel B: Predicted average treatment effect and heterogeneity loading parameter <i>post</i> program										
ATE (β_1)	-0.0551 (0.4,-0.5) [1.000]	-0.0508 (0.4,-0.5) [1.000]	-0.0484 (0.4,-0.5) [1.000]	-0.0570 (0.4,-0.5) [1.000]	-0.0297 (0.4,-0.5) [1.000]	3217.5 (8428.4,-1918.3) [0.439]	3527.3 (8737.3,-1706.6) [0.374]	3546.3 (8685.8,-1630.9) [0.358]	3314.0 (8523.0,-1842.7) [0.414]	3474.5 (8642.7,-1641.7) [0.370]
HET (β_2)	-0.0581 (0.6,-0.7) [1.000]	0.156 (1.6,-1.3) [1.000]	0.0376 (0.3,-0.2) [1.000]	-0.00219 (3.1,-2.8) [1.000]	0.111 (0.7,-0.5) [0.970]	0.0837 (0.6,-0.5) [1.000]	0.391 (1.8,-1.0) [1.000]	0.0220 (0.4,-0.3) [1.000]	0.416 (5.4,-1.1) [0.756]	0.0800 (0.5,-0.4) [1.000]
Λ	0.10	0.1	0.1	0.1	0.2	1625.5	1742.1	1679.7	2182.6	1746.7

Heterogeneity analysis based on the approach in [Chernozhukov et al. \(forthcoming\)](#) (see discussion in Section 5.1.2 and Appendix E). Estimates are based on equation 3. P-value (in brackets) for β_1 tests the hypothesis of no effect. P-value (in brackets) for β_2 tests the hypothesis of no heterogeneity. Panel A (respectively Panel B) shows estimates of β_1 and β_2 at midline (respectively endline). The Λ (lambda) statistic is displayed at the bottom of each panel: the larger Λ gets, the stronger the correlation between $\widehat{s}_0(Z)$, noted $S(Z)$, and $s_0(Z)$. Predictions are estimated for each sample split, reported values are the medians across 100 sample splits. Adjusted confidence intervals at 90% are provided in parentheses and adjusted p-values for partition uncertainty are provided in brackets. Earnings variable is in CFA, winsorized at the 97th percentile. For variables (y) in logarithms, we take $\ln(y+1)$.

Table A14: Heterogeneity in impacts on earnings *during* and *post* program, machine learning results for an extended set of covariates

	(1) Ln total earnings (Monthly)	(2) Ln total earnings (Monthly)	(3) Total earnings (Monthly)	(4) Total earnings (Monthly)
	<i>Midline</i>	<i>Endline</i>	<i>Midline</i>	<i>Endline</i>
Panel A: Predicted average treatment effect and heterogeneity loading parameter				
ATE (β_1)	2.311 (1.918,2.714) [0.000]	-0.136 (-0.723,0.449) [0.928]	27252.0 (17338.2,37013.8) [0.000]	4011.4 (-4014.3,11850.4) [0.631]
HET (β_2)	2.336 (1.156,3.577) [0.000]	0.272 (-0.306,0.814) [0.637]	2.297 (0.387,4.180) [0.032]	0.136 (-1.855,2.883) [0.888]
Best ML method	Generalized Random forest	Random forest	Generalized Random forest	R-Learner
Panel B: By quartile of predicted impacts on earnings <i>during</i> program (using <i>midline</i>) (GATES)				
Mean in quartile 1 (0 to 25%)	1.430 (0.596,2.233)	-0.695 (-2.047,0.682)	9616.7 (-9485.5,28647.1)	1309.1 (-17605.5,20036.9)
Mean in quartile 2 (25 to 50%)	2.063 (1.260,2.854)	-0.589 (-1.953,0.810)	31539.2 (12393.6,51035.5)	3483.3 (-15126.9,22326.7)
Mean in quartile 3 (50 to 75%)	2.602 (1.763,3.397)	-0.126 (-1.515,1.306)	32238.6 (12793.6,51515.6)	6594.5 (-11856.0,24930.0)
Mean in quartile 4 (75 to 100%)	3.312 (2.521,4.094)	-0.0495 (-1.417,1.326)	33266.5 (13826.0,52380.5)	8086.5 (-11723.2,27267.8)
P-value all coefficients are equal	0.007	1.000	0.292	1.000
Best ML method	Generalized Random forest	Generalized Random forest	Generalized Random forest	Generalized Random forest
Panel C: By quartile of predicted impacts on earnings <i>post</i> program (using <i>endline</i>) (GATES)				
Mean in quartile 1 (0 to 25%)	1.935 (1.102,2.761)	-0.338 (-1.512,0.819)	21497.8 (3457.5,39527.0)	2309.3 (-12952.7,18286.3)
Mean in quartile 2 (25 to 50%)	1.989 (1.164,2.868)	-0.359 (-1.520,0.828)	27897.3 (10231.8,45734.2)	5194.6 (-10461.8,20933.5)
Mean in quartile 3 (50 to 75%)	2.246 (1.456,3.038)	-0.432 (-1.562,0.736)	28716.5 (10596.6,46321.6)	3064.9 (-12565.0,18074.0)
Mean in quartile 4 (75 to 100%)	2.555 (1.768,3.370)	0.368 (-0.800,1.488)	29256.8 (11482.7,46980.5)	5016.5 (-10414.7,21248.5)
P-value all coefficients are equal	0.093	0.682	0.685	0.998
Best ML method	Random forest	Random forest	R-Learner	R-Learner

Heterogeneity analysis based on the approach in [Chernozhukov et al. \(forthcoming\)](#) (see discussion in Section 5.1.2 and Appendix E). All baseline variables in the balance table (Table 1) are used as covariates. Columns (1) and (3) (respectively columns (2) and (4)) focus on outcomes at midline (respectively endline). Estimates in Panel A are based on equation 3. P-value (in brackets) for β_1 tests the hypothesis of no effect. P-value (in brackets) for β_2 tests the hypothesis of no heterogeneity. GATES estimates are based on the specification in equation 4. Panel B (respectively Panel C) shows impacts per quartile of the predicted treatment effects at midline (respectively endline). The best predictions are reported. The chosen algorithm is the one with the highest Λ (lambda): the larger Λ gets, the stronger the correlation between $\hat{s}_0(Z)$, noted $S(Z)$, and $s_0(Z)$. Predictions are estimated for each sample split, reported values are the medians across 100 sample splits. Adjusted confidence intervals at 90% are provided in parentheses and adjusted p-values for partition uncertainty are provided in brackets. Earnings variable is in CFA, winsorized at the 97th percentile. For variables (y) in logarithms, we take $\ln(y+1)$.

Table A15: Estimated impacts on (ln) earnings *post* program, by treatment arm

	(1) Public Works only (PW)	(2) PW and Self-Employment Training (SET)	(3) PW and Wage-Employment Training (WET)
	<i>Endline</i>	<i>Endline</i>	<i>Endline</i>
Panel A: Predicted average treatment effect and heterogeneity loading parameter			
ATE (β_1)	-0.0753 (-0.601,0.447) [1.000]	-0.0116 (-0.545,0.529) [1.000]	0.120 (-0.406,0.648) [0.930]
HET (β_2)	0.0866 (-0.211,0.382) [0.818]	-0.572 (-1.431,0.299) [0.425]	0.450 (-0.595,1.935) [0.722]
Best ML method	Gradient boosting	Elastic net	R-Learner
Panel B: By quartile of predicted impacts on earnings <i>post</i> program (<i>using endline</i>) (GATES)			
Mean in quartile 1 (0 to 25%)	-0.220 (-1.270,0.826)	0.408 (-0.663,1.487)	-0.0540 (-1.112,1.004)
Mean in quartile 2 (25 to 50%)	-0.141 (-1.188,0.923)	0.132 (-0.958,1.206)	0.138 (-0.913,1.187)
Mean in quartile 3 (50 to 75%)	-0.172 (-1.217,0.873)	-0.116 (-1.193,0.962)	0.238 (-0.805,1.279)
Mean in quartile 4 (75 to 100%)	0.148 (-0.909,1.214)	-0.522 (-1.587,0.543)	0.168 (-0.912,1.223)
P-value all coefficients are equal	0.902	0.478	0.950
Best ML method	Gradient boosting	Elastic net	R-Learner

Heterogeneity analysis based on the approach in [Chernozhukov et al. \(forthcoming\)](#) (see discussion in Section 5.1.2 and Appendix E). Columns (1-3) show estimated impacts for each treatment arm compared to the control group. Estimates in Panel A are based on equation 3. P-value (in brackets) for β_1 tests the hypothesis of no effect. P-value (in brackets) for β_2 tests the hypothesis of no heterogeneity. Estimates in panel B are based on the specification in equation 4. They show impacts per quartile of the predicted treatment effects at endline. The best predictions are reported. The chosen algorithm is the one with the highest Λ (lambda): the larger Λ gets, the stronger the correlation between $\hat{s}_0(Z)$, noted $S(Z)$, and $s_0(Z)$. Predictions are estimated for each sample split, reported values are the medians across 100 sample splits. Adjusted confidence intervals at 90% are provided in parentheses and adjusted p-values for partition uncertainty are provided in brackets. Earnings variable is in CFA, winsorized at the 97th percentile. For variables (y) in logarithms, we take $\ln(y + 1)$.

Table A16: Estimated impacts on earnings (in levels) *post* program, by treatment arm

	(1)	(2)	(3)
	Public Works only (PW)	PW and Self-Employment Training (SET)	PW and Wage-Employment Training (WET)
	<i>Endline</i>	<i>Endline</i>	<i>Endline</i>
Panel A: Predicted average treatment effect and heterogeneity loading parameter			
ATE (β_1)	2155.1 (-3648.0,8097.0) [0.799]	4166.5 (-2085.9,10463.6) [0.397]	3970.3 (-2233.2,10007.5) [0.412]
HET (β_2)	0.711 (-0.793,2.217) [0.645]	0.167 (-2.294,3.175) [0.822]	0.859 (-0.383,2.666) [0.359]
Best ML method	Generalized Random forest	R-Learner	R-Learner
Panel B: By quartile of predicted impacts on earnings <i>post</i> program (<i>using endline</i>) (GATES)			
Mean in quartile 1 (0 to 25%)	-804.5 (-12141.7,10585.9)	4252.3 (-8217.4,16697.1)	954.7 (-11358.2,13202.6)
Mean in quartile 2 (25 to 50%)	314.7 (-11327.4,12027.1)	5075.0 (-7721.7,17162.5)	3213.2 (-8929.3,15509.9)
Mean in quartile 3 (50 to 75%)	2580.6 (-9034.2,14329.5)	4686.6 (-8013.1,17365.6)	3843.2 (-8184.3,15983.8)
Mean in quartile 4 (75 to 100%)	6306.2 (-5642.9,17819.6)	3651.1 (-8919.7,15948.0)	5448.1 (-6760.2,17345.3)
P-value all coefficients are equal	0.715	0.772	0.828
Best ML method	Generalized Random forest	R-Learner	R-Learner

Heterogeneity analysis based on the approach in [Chernozhukov et al. \(forthcoming\)](#) (see discussion in Section 5.1.2 and Appendix E). Columns (1-3) show estimated impacts for each treatment arm compared to the control group. Estimates in Panel A are based on equation 3. P-value (in brackets) for β_1 tests the hypothesis of no effect. P-value (in brackets) for β_2 tests the hypothesis of no heterogeneity. Estimates in panel B are based on the specification in equation 4. They show impacts per quartile of the predicted treatment effects at endline. The best predictions are reported. The chosen algorithm is the one with the highest Λ (lambda): the larger Λ gets, the stronger the correlation between $\hat{s}_0(Z)$, noted $S(Z)$, and $s_0(Z)$. Predictions are estimated for each sample split, reported values are the medians across 100 sample splits. Adjusted confidence intervals at 90% are provided in parentheses and adjusted p-values for partition uncertainty are provided in brackets. Earnings variable is in CFA, winsorized at the 97th percentile.

Table A17: Baseline characteristics of the bottom and top quartiles of predicted impacts on earnings (in levels) *during* program

	(1) Mean in 1st quartile	(2) Mean in 4th quartile	(3) Test (1)-(2) (p-value)
Individual characteristics			
Female	0.18	0.43	0
Lives in urban area	0.91	0.97	0
Age	24.91	24.41	0.012
Number of children	0.84	0.80	0.065
Education			
Years of education	10.16	10.16	0.065
Primary education not completed	0.45	0.46	0.095
Has participated in vocational training	0.53	0.27	0
Household characteristics and assets			
Household size	6.75	5.52	0
Is head of household	0.28	0.20	0.005
Total number of assets	0.59	0.51	0
Number of transportation assets	1.22	0.42	0
Number of agricultural assets	7.74	2.42	0
Number of household durables	2.31	1.26	0
Number of communication assets	8.42	5.70	0
Employment			
Has an activity	0.91	0.66	0
Is wage employed	0.42	0.28	0
Is self-employed	0.52	0.22	0
Number of activities	1.23	0.72	0
Total earnings (monthly, CFA)	34385.3	8298.2	0
Searched for a job (last month)	0.62	0.52	0.519
Savings, constraints and expenditures			
Has saved (last 3 months)	0.52	0.45	0.033
Of which: share of formal savings	0.28	0.24	0.467
Has a savings account	0.15	0.10	0.006
Savings Stock (CFA)	51724.9	14685.6	0
Self-reported constraints to repay loans	0.29	0.11	0
Self-reported constraints to access credit	0.41	0.59	0
Self-reported constraints for basic needs expenditures	0.69	0.72	0.510
Transportation expenditures (last 7 days, CFA)	3066.3	1065.9	0
Communication expenditures (last 7 days, CFA)	2810.9	895.8	0
Cognitive skills and risk preference			
Cognitive (Raven Test)	0.23	0.23	0.798
Dexterity (Nuts Test)	0.37	0.38	0.104
Dexterity (Bolts Test)	0.34	0.33	0.900
Positive affect (CES-D scale, No. of positive days)	6.31	6.22	0.856
Positive attitude towards the future (ZTPI scale)	29.32	29.20	0.772
Is Risk averse (based on hypothetical lotteries)	0.71	0.72	1
Patience (scale 0 to 10, 10=very patient)	3.30	3.34	0.945
Preference for present (actualization rate for 1 month)	0.57	0.58	0.917

Heterogeneity analysis based on the approach in [Chernozhukov et al. \(forthcoming\)](#) (see discussion in Section 5.1.2 and Appendix E). Column (1) (respectively (2)) displays average characteristics of the bottom (respectively top) quartile of the distribution of predicted impacts on earnings during the program. Column (3) reports p-values for a test of equality between the top and bottom quartile. Reported results are based on the algorithm with best predictions for midline : Generalized Random forest. The chosen algorithm is the one with the highest Λ (lambda): the larger Λ gets, the stronger the correlation between $\widehat{s}_0(Z)$, noted $S(Z)$, and $s_0(Z)$ (see appendix Table A13 for comparisons across algorithms). Means by quartile are estimated for each sample split, reported values are the medians across 100 sample splits. Adjusted p-values for partition uncertainty are provided. *Household assets* and *savings stock* variables winsorized at the 99th percentile.

Table A18: Estimated impacts *during* and *post* program on savings

	(1)	(2)	(3)	(4)
	Ln Savings (Monthly)	Ln Savings (Monthly)	Savings in CFA (Monthly)	Savings in CFA (Monthly)
	<i>Midline</i>	<i>Endline</i>	<i>Midline</i>	<i>Endline</i>
Panel A: Predicted average treatment effect and heterogeneity loading parameter				
ATE (β_1)	3.543 (2.981,4.100) [0.000]	0.472 (-0.0260,0.977) [0.127]	36768.4 (30901.6,42641.1) [0.000]	9492.9 (1192.6,17761.4) [0.050]
HET (β_2)	1.303 (0.454,2.169) [0.005]	0.247 (-1.230,2.565) [0.809]	0.440 (-0.0100,0.889) [0.110]	0.193 (-0.101,0.505) [0.388]
Best ML method	Generalized Random forest	R-Learner	Random forest	Gradient boosting
Panel B: By quartile of predicted impacts on savings <i>during</i> program (GATES)				
Mean in quartile 1 (0 to 25%)	2.481 (1.355,3.611)	0.748 (-0.440,1.920)	29852.0 (18249.3,41593.9)	10275.5 (-9226.9,29636.8)
Mean in quartile 2 (25 to 50%)	3.174 (2.052,4.286)	0.302 (-0.870,1.471)	33873.9 (21920.6,45466.9)	7074.7 (-12459.8,27015.8)
Mean in quartile 3 (50 to 75%)	3.785 (2.671,4.889)	0.563 (-0.619,1.746)	38335.9 (26568.8,49862.4)	10085.7 (-9870.0,29882.8)
Mean in quartile 4 (75 to 100%)	4.778 (3.659,5.896)	0.316 (-0.885,1.497)	45983.1 (34465.3,57655.4)	9897.8 (-9570.1,28961.8)
P-value all coefficients are equal	0.031	1.000	0.269	1.000
Best ML method	Generalized Random forest	Generalized Random forest	Random forest	Random forest
Panel C: By quartile of predicted impacts on savings <i>post</i> program (GATES)				
Mean in quartile 1 (0 to 25%)	3.615 (2.557,4.647)	0.316 (-0.703,1.346)	36489.5 (26070.4,47031.0)	3831.7 (-12411.8,19889.7)
Mean in quartile 2 (25 to 50%)	3.551 (2.488,4.582)	0.336 (-0.705,1.367)	37213.5 (26671.1,47659.7)	7044.9 (-9401.7,23694.9)
Mean in quartile 3 (50 to 75%)	3.574 (2.503,4.652)	0.483 (-0.527,1.506)	37366.7 (26941.2,47912.0)	9436.6 (-7461.4,25802.8)
Mean in quartile 4 (75 to 100%)	3.585 (2.532,4.646)	0.731 (-0.286,1.740)	39460.8 (28965.9,50086.8)	16964.8 (816.3,33311.6)
P-value all coefficients are equal	0.507	0.909	1.000	0.377
Best ML method	R-Learner	R-Learner	Gradient boosting	Gradient boosting

Heterogeneity analysis based on the approach in [Chernozhukov et al. \(forthcoming\)](#) (see discussion in Section 5.1.2 and Appendix E). Columns (1) and (3) (respectively columns (2) and (4)) focus on outcomes at midline (respectively endline). Estimates in Panel A are based on equation 3. P-value (in brackets) for β_1 tests the hypothesis of no effect. P-value (in brackets) for β_2 tests the hypothesis of no heterogeneity. GATES estimates are based on the specification in equation 4. Panel B (respectively Panel C) shows impacts per quartile of the predicted treatment effects at midline (respectively endline). The best predictions are reported. The chosen algorithm is the one with the highest Λ (lambda): the larger Λ gets, the stronger the correlation between $\hat{s}_0(Z)$, noted $S(Z)$, and $s_0(Z)$. Predictions are estimated for each sample split, reported values are the medians across 100 sample splits. Adjusted confidence intervals at 90% are provided in parentheses and adjusted p-values for partition uncertainty are provided in brackets for Panel A. *Savings* variable is in CFA, winsorized at the 97th percentile. For variables (y) in logarithms, we take $\ln(y+1)$.

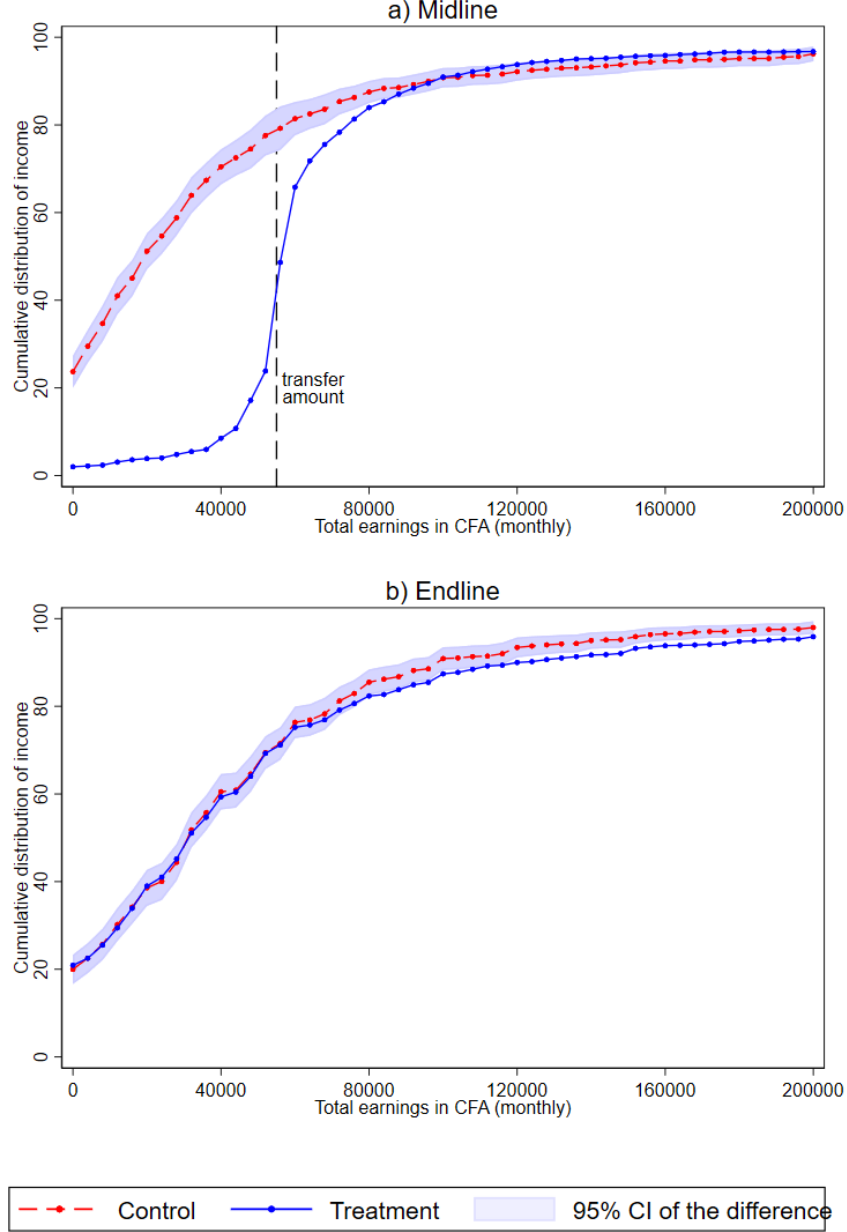
Table A19: Baseline characteristics of the bottom and top quartiles of predicted impacts on (ln) savings *during* program

	(1) Mean in 1st quartile	(2) Mean in 4th quartile	(3) Test (1)-(2) (p-value)
Individual characteristics			
Female	0.14	0.46	0
Lives in urban area	0.95	0.95	0.598
Age	24.65	24.29	0.080
Number of children	0.66	0.81	0.070
Education			
Years of education	11.76	9.71	0
Primary education not completed	0.36	0.49	0
Has participated in vocational training	0.47	0.34	0.001
Household characteristics and assets			
Household size	5.35	6.46	0
Is head of household	0.37	0.11	0
Total number of assets	0.65	0.44	0
Number of transportation assets	0.68	0.70	0.331
Number of agricultural assets	3.88	4.44	0.402
Number of household durables	1.75	1.70	0.125
Number of communication assets	6.95	6.72	0.014
Employment			
Has an activity	1	0.45	0
Is wage employed	0.55	0.12	0
Is self-employed	0.48	0.14	0
Number of activities	1.30	0.49	0
Total earnings (monthly, CFA)	38482.1	658.7	0
Searched for a job (last month)	0.63	0.53	0.756
Savings, constraints and expenditures			
Has saved (last 3 months)	0.74	0.24	0
Of which: share of formal savings	0.34	0.18	0.002
Has a savings account	0.19	0.06	0
Savings Stock (CFA)	63881.6	5518.0	0
Self-reported constraints to repay loans	0.22	0.20	0.709
Self-reported constraints to access credit	0.39	0.57	0
Self-reported constraints for basic needs expenditures	0.66	0.72	0.252
Transportation expenditures (last 7 days, CFA)	2754.9	967.5	0
Communication expenditures (last 7 days, CFA)	2760.5	820.6	0
Cognitive skills and risk preference			
Cognitive (Raven Test)	0.24	0.23	0.781
Dexterity (Nuts Test)	0.37	0.38	0.300
Dexterity (Bolts Test)	0.34	0.33	1
Positive affect (CES-D scale, No. of positive days)	6.42	6.16	0.212
Positive attitude towards the future (ZTPI scale)	29.34	29.10	0.586
Is Risk averse (based on hypothetical lotteries)	0.73	0.71	1
Patience (scale 0 to 10, 10=very patient)	3.39	3.31	1
Preference for present (actualization rate for 1 month)	0.57	0.59	0.507

Heterogeneity analysis based on the approach in [Chernozhukov et al. \(forthcoming\)](#) (see discussion in Section 5.1.2 and Appendix E.3). Column (1) (respectively (2)) displays average characteristics of the bottom (respectively top) quartile of the distribution of predicted impacts on (ln) savings during the program. Column (3) reports p-values for a test of equality between the top and bottom quartile. Reported results are based on the algorithm with best predictions for midline : Generalized Random forest. The chosen algorithm is the one with the highest Λ (lambda): the larger Λ gets, the stronger the correlation between $\widehat{s}_0(Z)$, noted $S(Z)$, and $s_0(Z)$. Means by quartile are estimated for each sample split, reported values are the medians across 100 sample splits. Adjusted p-values for partition uncertainty are provided. *Household assets* and *savings stock* variables winsorized at the 99th percentile.

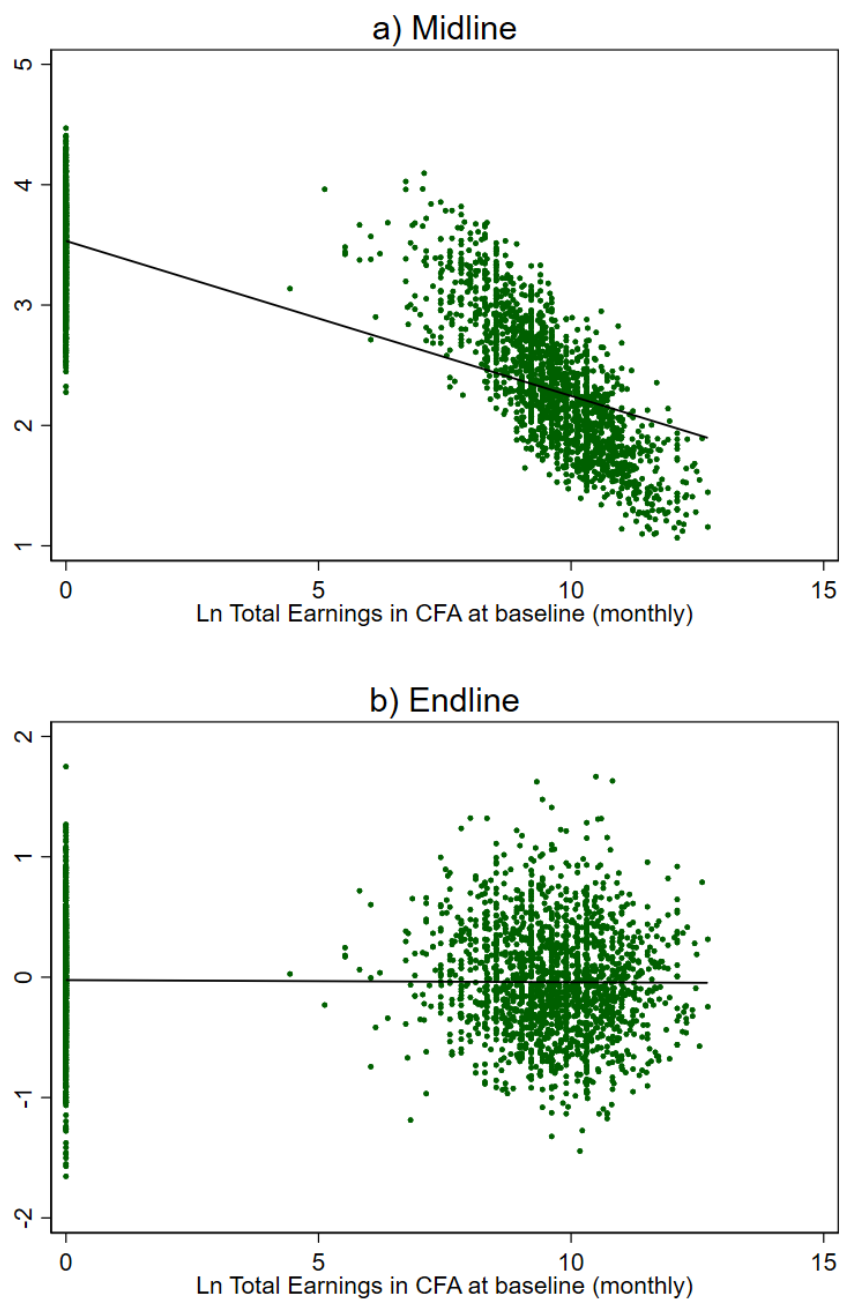
Appendix B Additional Figures

Figure B1: Cumulative distribution of earnings *during* and *post* program



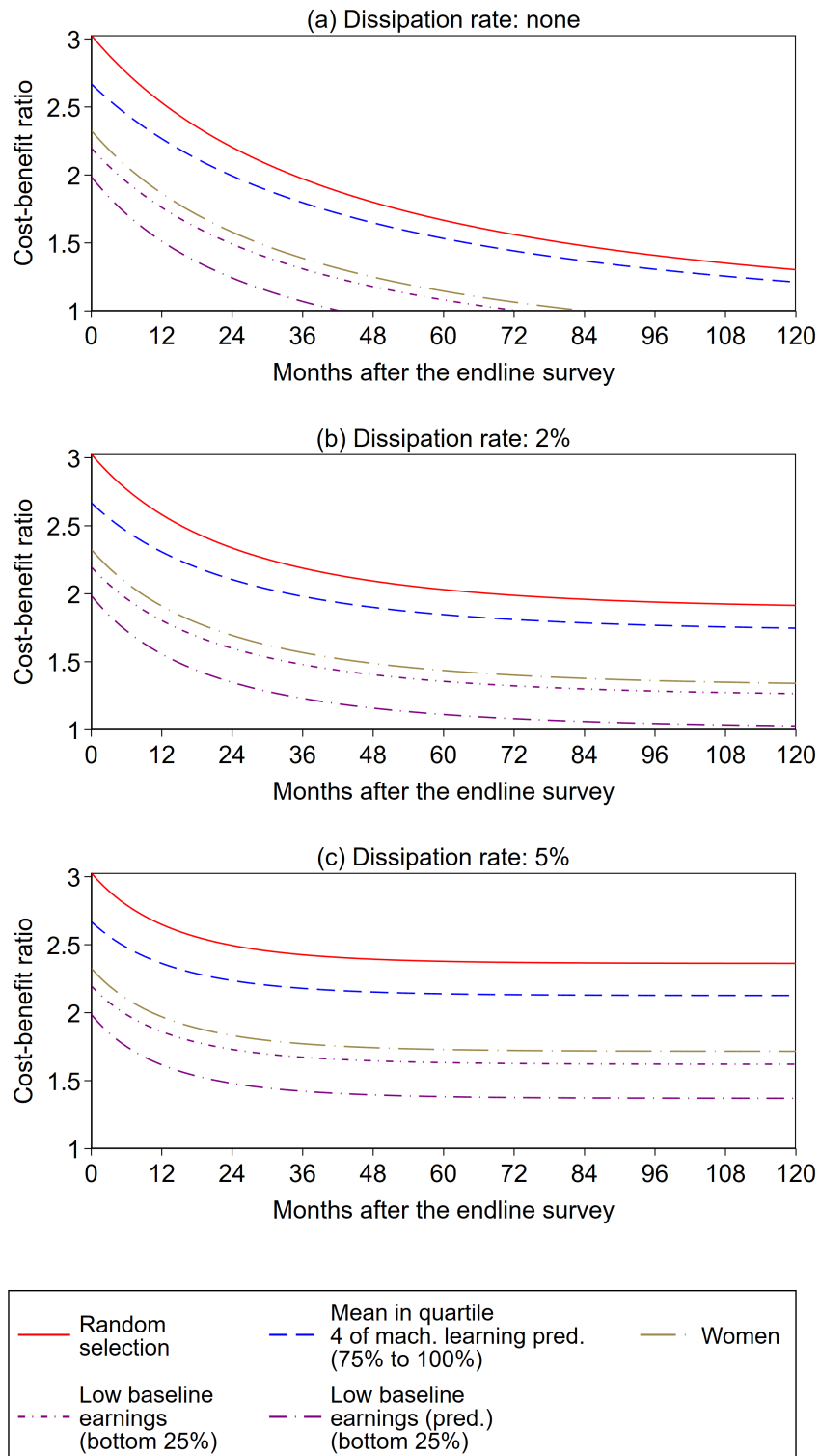
Note: Total monthly earnings variable winsorized at the 97th percentile. The figures show the results of the estimation of the cumulative distribution of potential outcomes in the treatment and control groups. They are based on the estimation of equation 1, with variables defined as $1(y < t)$ for t varying over the support of y . The red curve shows the average of those variables in the control group. The blue curve adds to this average the estimated treatment coefficient. The grey area around the red curve represents a band of ± 1.96 times the standard error of the estimated treatment effect.

Figure B2: Reported baseline earnings vs predicted impacts *during* and *post* program



Note: Panel (a) compares baseline labor income to best predictions at midline, using Generalized Random Forest. Panel (b) compares baseline labor income to best predictions at endline using Random Forest. The correlation coefficient is -0.805 at midline and -0.018 at endline. In both panels, linear regression lines are added (corresponding to a R^2 of 0.648 in Panel (a)). For variables (y) in logarithms, we take $\ln(y + 1)$

Figure B3: Cost-benefit ratios over time under alternative targeting rules, depending on the sustainability of post-program impacts



Note: The discounted sum of post-program impacts is assumed to continue beyond what we measure at the endline survey, 15 months after the end of the program. The top panel assumes no dissipation of impacts. The middle and bottom panels assume a dissipation rate on top of the discount rate: a 2% monthly dissipation rate (respectively 5%) is equivalent to a 22% decrease in impact in one year (respectively 49%).

Appendix C Description of complementary training

Randomized subsets of beneficiaries received complementary training on basic entrepreneurship or job search skills. Each training lasted approximately 80-100 hours over two two-week periods. Field exercises were undertaken between the training periods, in parallel to the public works jobs (typically in the afternoons). The training was delivered by work brigades, i.e. in groups of 25 youths. Participants did not have to work during the training, but still received their corresponding daily wage.⁶⁶ The curricula for the complementary skills training were tailored for low-skill populations that may not be able to read and write, in particular by relying on drawings and visuals.

The basic entrepreneurship training aimed to build skills to help youth set-up and manage a small micro-enterprise. The training lasted 100 hours and focused on providing cross-cutting business skills and practical guidance to develop simple business plans for small-scale activities that can be set-up using savings from the public works program. A first phase (40 hours over two weeks) started with basic considerations on how to choose and set up an activity. It then covered business skills on topics such as simple accounting practices, managing stocks or setting up prices. A second phase included field research for youths to collect information, undertake basic market research and outline a business plan. A third phase (40 hours over two weeks) included feedback on a basic business plan, and reviewed related topics from the curriculum with an operational focus. The final phase (20 hours) included post-training follow-up. The curriculum and training material were specifically developed for youths with limited education, relying on practical exercises and visuals.

The training on wage job search skills provided information on wage jobs opportunities, skills on jobs search techniques. It also offered a more professional environment during the public works programs and skills certification to facilitate signaling upon exit from the program. The training itself lasted 80 hours. The first phase (40 hours over two weeks) discussed how to identify wage jobs opportunities (either locally or through migration), how to search for wage jobs, prepare a CV, apply for a job and prepare a job interview. It included a self-evaluation of strength and weaknesses, and guidance on identifying the types of wage jobs suitable to participants' profiles. The second phase included field exercises to collect information on potential opportunities, identify

⁶⁶Some youth were offered the second half of the training shortly after their exit from the public works program. While these youths were not paid during that time, they received a small stipend to cover transportation costs.

and visit potential employers or professional networks, etc. During this phase, beneficiaries started implementing a “job search action plan” they had previously developed with their trainer. The third phase (40 hours over two weeks) provided individual and group feedback on field exercises, reviewed part of the curriculum (search, job applications and interviews) and provided additional practical guidance on how to search for a job.

In addition, and prior to the training on job search, supervisors of the brigades who were offered the wage job search training were also trained (over one week) on how to manage teams and provide feedback to workers, with the objective to mimic the professional experience one would have in a more formal wage job. Supervisors then periodically rated youths on a range of skills to identify their strength and weaknesses. (This occurred 3 times per participant on average, over the remaining 2.5 months of public works), and these evaluations were later used to issue a work certificate that signaled between one and five competencies identified as strengths for each participant. It was recommended to participants to use those certificates in their job search, to better signal their skills. However, certificates were handed to them with a delay, a few months after the end of the public works.⁶⁷

⁶⁷The impact evaluation policy report contains additional details on these two types of training, including on the curricula (Bertrand et al., 2016).

Appendix D Definition of Key Outcome Variables

Total monthly earnings are expressed in CFA francs. They are aggregated over up to three activities undertaken by an individual in the 30 days preceding the survey. They include payments received in cash and the monetary equivalent for in-kind payments. The variable is winsorized at 97% (unless stated otherwise). Total monthly earnings are decomposed in total (monthly) earnings from wage employment and self-employment (as well as earnings from other occupations, which are generally small hence not shown separately). When shown in log, the log transformation is applied to earnings plus one.

Has an Activity is a dummy taking a value of 1 if the individual has worked at least one hour over the 7 days preceding the survey, consistent with the official employment indicators used in Côte d'Ivoire. We assign a value of 0 for inactive and unemployed individuals. To provide information on the composition of employment, we also analyze having at least one wage job (*Wage employed*), or at least one self-employment activity (*Self-employed*).

Weekly hours worked capture the total number of hours worked over the 7 days preceding the survey. It aggregates information from up to three activities undertaken by an individual across all occupations (wage employment, self-employment or other types of activity). The variable is winsorized at 97%. Weekly hours worked are decomposed in (*hours worked in wage employment*) and (*hours worked in self-employment*) (as well as hours worked in other occupations, which are generally small and not displayed separately).

Savings stock is the total amount of savings in CFA francs at the time of the survey. It aggregates savings from formal or informal sources. The variable is winsorized at 97%. When shown in log, the log transformation is applied to savings plus one.

Total expenditures is expressed in CFA francs and aggregates several types of expenditures, both for the individual and for other household members. It includes basic expenditures (health, clothing, sanitation, and accommodation), communication expenditures (mobile, internet, and media), investments (education, training, maintenance of assets), transportation expenditures, temptation goods (alcohol, tobacco, gambling, and luxury goods) and social expenditures (celebrations and charity). The variable is winsorized at 97%. When shown in log, the log transformation is applied to expenditures plus one.

The *well-being index* aggregates 6 measures: two measures of *happiness* and *pride* in daily ac-

tivities from a time-use module,⁶⁸ the Rosenberg *self-esteem* scale,⁶⁹ the *positive affect sub-scale* from the CESD scale,⁷⁰ the sub-scale of (*positive*) *attitude towards the future* and the inverted sub-scale of *present fatalism* from the ZTPI scale.⁷¹ The well-being index is a z-score, with a mean of zero and a standard deviation of one in the control group, so that estimated coefficients can be interpreted in standard deviations. A positive impact on the well-being index is interpreted as an overall improvement in well-being.⁷²

The *behavior index* aggregates 4 measures: an inverted measure of *anger* in daily activities taken from the time-use module,⁷³ an inverted measure of *impulsiveness* from the DERS scale,⁷⁴ the *conduct problems sub-scale* (inverted) and the *pro-social behavior sub-scale* from the Strengths and Difficulties Questionnaire (SDQ).⁷⁵ As for the well-being index, the behavior index is a z-score with a mean of 0 and standard deviation of 1 in the control group, so that estimated coefficients can be interpreted in standard deviations. An increase in the index corresponds to an overall improvement in behavior and attitude.⁷⁶

The *risky behavior index* is the mean of 9 risky behaviors measured through list experiments. They include stealing, assaulting someone, believing smuggling is necessary (to earn a living), prostitution, threatening someone, taking illicit drugs, smuggling, having ties with a smuggling

⁶⁸The time use module measured which activities the respondent performed at different times of the last “business day” (at 6am, 10am, 3pm, 7pm and 10pm). Respondents were also asked whether they felt happy, proud or angry while performing those activities. The measure of happiness (respectively pride) is the number of times (out of the 5 times in the last day) respondents reported feeling happy (respectively proud). A z-score of the measure is included in the well-being index.

⁶⁹The Rosenberg self-esteem scale includes 10 items that measure self-esteem or self-worth. We use a validated version of the instrument in French (Vallieres and Vallerand, 1990).

⁷⁰The Center for Epidemiologic Studies Depression (CESD) scale includes an inverted subscale that measures positive feelings (“Positive Affects”). We use a validated version in French (Morin et al., 2011).

⁷¹The Zimbardo Time Perspective Inventory (ZTPI) captures different dimensions of time perspectives. We use the two subscales of “future” (to have a positive attitude towards future) and “present fatalism” which is very close to the concept of external locus of control, in the sense that one feels *no* control over life events. The inverted “present fatalism” measure is therefore similar to a measure of internal locus of control. We use a validated version of the instrument in French (Apostolidis and Fieulaine, 2004).

⁷²The index adds up the 6 measures described above, out of which one is inverted (*present fatalism*). Therefore a negative impact on the *present fatalism* measure induces an improvement in the overall well-being index, corresponding to greater well-being.

⁷³This was built as in footnote 68.

⁷⁴The Difficulties in Emotion Regulation Scale (DERS) is used to measure socio-emotional regulation, in particular the difficulties of regulation of emotions in adults. Three of the six questions of the “difficulties to control impulsive behavior” scale were retained, based on a validated French version of the instrument (Côté et al., 2013).

⁷⁵The Strengths and Difficulties Questionnaire (SDQ) measures behavioral difficulties in young people, initially among children and adolescents from 3 to 16 years old (Goodman et al., 1998). The instrument was slightly adapted for an older age group 18 to 30 years old. We use two of five sub-scales from a validated questionnaire in French available at www.sdqinfo.com.

⁷⁶The index adds up the 4 measures described above, which are all inverted in the index except *pro-social behavior* measure. A negative impact on inverted measures, for example *conduct problems*, corresponds to a positive behavior and leads to an improvement in the overall behavior index.

network, and having firearms at home. Because respondents may not respond truthfully to direct questions about these sensitive behaviors, we used list experiments instead. Rather than asking directly a sensitive question about a risky behavior (e.g. stealing), 5 affirmations are read to respondents, and respondents are asked how many of these affirmations (between 0 and 5) are true for them. To estimate the proportion of individuals for which the sensitive question is true in a sample, the sample is (randomly) assigned to two lists. The first list includes 5 affirmations with the risky behavior, and the second list only includes the other 4 affirmations (without the risky behavior). We implemented a “double” list experiment to avoid losing statistical power: each half of the sample answered both a list with sensitive questions, and a (different) list with control questions corresponding to the other sample. List experiments were piloted extensively to ensure a good understanding by respondents. In the analysis, we use the likelihood ratio test introduced by [Blair and Imai \(2012\)](#) to test for the existence of design effects.

Appendix E Applying Machine Learning to Study Heterogeneity in Treatment Effects

This appendix complements section 5.2 by providing details on the application of machine learning methods to analyze treatment effect heterogeneity. The application is based on Chernozhukov et al. (forthcoming). Section E.1 provides an overview of the approach. Section E.2 describes the sample used to train the models and make predictions.

Section E.3 presents the machine learning algorithms and their parameters. Finally, Section E.4 describes how we adapted the procedure in Chernozhukov et al. (forthcoming) to our experimental setting.

We use similar notations as in Chernozhukov et al. (forthcoming). The Baseline Conditional Average (BCA) writes:

$$b_0(Z) := \mathbb{E}[Y(0)|Z]$$

The Conditional Average Treatment Effect (CATE) is defined as :

$$s_0(Z) := \mathbb{E}[Y(1)|Z] - \mathbb{E}[Y(0)|Z]$$

E.1 Overview

We are looking to estimate treatment effects for specific subgroups in the population, defined by their (observable) characteristics z . We would like to estimate the conditional average treatment effect (CATE) for some subgroups, corresponding to $s_0(Z) = E[Y(1) - Y(0)|Z^K = z]$ where Z^K is a vector of K baseline covariates (features) and $Y(k)$ the potential outcome of interest for treatment ($k = 1$) and control ($k = 0$). We use ML methods to obtain an estimate $S(Z)$.

A key challenge when using machine learning techniques to study heterogeneity is to derive confidence intervals and perform inference. In this paper, we use the inference framework developed by Chernozhukov et al. (forthcoming), who present a general approach using machine learning estimators as a proxy predictor to make inference on *key features* of the CATE function (rather than the whole function).⁷⁷

⁷⁷This contrasts with the approach in Wager and Athey (2018) who derive point-wise confidence intervals in the specific case of causal forests.

This allows us to (i) formally test for the existence of heterogeneity, (ii) compute confidence intervals around the conditional treatment effect for groups of interest (such as those at the top and bottom of the distribution of $S(Z)$), and (iii) compare the characteristics of the population who benefit the most or the least from the program. The approach in Chernozhukov et al. (forthcoming) is “generic” in the sense that it applies to any machine learning algorithm used to estimate heterogeneous treatment effect, including the causal forest and generalized random forest estimators proposed by Wager and Athey (2018) and Athey et al. (2019).

In our own implementation, we consider several alternative machine learning algorithms (detailed in Appendix E.3). We present the results for the best performing algorithm in the main tables of the paper, and provide robustness checks in Table A13.

When applying machine learning methods, we split our data so that separate sub-samples are used to either build the model (the *auxiliary sample*, on which machine learning predictors are trained and constructed) or make inference (the *main sample*, to which the model is applied, and on which we estimate the different key features of the CATE function). In our application, this procedure is repeated 100 times on random sub-samples.⁷⁸ Chernozhukov et al. (forthcoming) offer a procedure to aggregate results across simulations and construct valid confidence intervals and p-values.⁷⁹

We test for the presence of heterogeneity by estimating the β_2 coefficient in the following equation:

$$Y = \alpha_1 + \alpha_2 B(Z) + \beta_1(T - P(Z)) + \beta_2(T - P(Z))(S(Z) - \hat{E}(S(Z))) + \varepsilon \quad (\text{F1})$$

Machine learning is used to get $S(Z)$, a relevant *proxy* predictor of $s_0(Z)$, as well as $B(Z)$, a machine learning predictor for $Y(0)$ (both constructed on the auxiliary sample). T is the treatment variable, and $P(Z) = \hat{E}(T|Z)$. We use weights $w(Z) = \{P(Z)(1 - P(Z))\}^{-1}$.

β_1 captures the average treatment effect (ATE) while β_2 is the heterogeneity loading parameter

⁷⁸Iterations of the data-splitting process are necessary to identify how much variation is induced by specific data splits. It also ensures that each observation will be used on average both for construction and prediction (depending on the data-split), so all the information contained in survey data is used. This is especially important given our rather small sample size.

⁷⁹Their procedure takes into account the uncertainty coming from both the estimation of the parameters and the data splitting process when aggregating the results (p-values, confidence interval bounds) across simulations. It takes the median of the estimated parameters over all splits, as well as the median of p-values which is then adjusted by a factor of 2. Confidence intervals computed at 95% significance ($\alpha = 0.05$) have to be re-adjusted for split uncertainty. After adjustment, the procedure provides confidence interval bounds at 90%.

(HET).⁸⁰ We are particularly interested in β_2 , which offers a test for heterogeneity in the treatment effect. Rejecting the null hypothesis that $\beta_2 = 0$ means that (i) there is heterogeneity, and (ii) that our machine learning predictor is a good approximation of $s_0(Z)$. On the contrary, if β_2 is not statistically different from zero, it means either that our machine learning predictor is uncorrelated with $s_0(Z)$ (our predictor is not able to capture heterogeneity correctly), or that there is no heterogeneity. In our application, we test for heterogeneity in impacts on earnings both *during* and *post* program.

Besides detecting heterogeneity, we are also interested in the magnitude of the treatment effects along the distribution. In our application, we consider the top and bottom quartiles of the distribution, corresponding to the 25% of individuals who benefit the most and the least in terms of impacts on earnings. This is because around 25% of total applicants were selected to participate in the program we study.⁸¹

We recover the parameters of interest $E(s_0(Z)|G_k)$, also referred as Group Average Treatment Effects (GATES) — where groups are quartiles of the distribution of predicted treatment effects — through the following weighted linear projection:

$$Y = \alpha_1 + \alpha_2 B(Z) + \sum_{k=1}^4 \gamma_k (T - p(Z)) 1(G_k) + \nu \quad (\text{F2})$$

The projection coefficients γ_k are the GATES parameters. The groups are defined as $G_k = \{S(Z) \in I_k\}$ with $I_k = [q_{k-1}, q_k)$, where q_k are the quartiles of $S(Z)$, and $q_0/q_4 = -/+ \infty$. We again use weights $w(Z) = \{P(Z)(1 - P(Z))\}^{-1}$. The estimated parameter γ_4 (corresponding to the top quartile of the predicted distribution of impact, group G_4) can be interpreted as the average treatment effect among the 25% of individuals identified by the ML procedure as those who benefit the *most* from the program. Similarly, γ_1 can be interpreted as the average treatment effect among the 25% of individuals identified by the ML procedure as those who benefit the *least* from the program (group G_1).

Note that the question of targeting individuals for whom the effect of a program is the highest is only imperfectly resolved by using GATES for at least two reasons. The first reason is that

⁸⁰In this framework, the quantity $BLP(Z) = \beta_1 + \beta_2(S(Z) - \hat{E}(S(Z)))$ can be interpreted as the best linear predictor of $s_0(Z)$ based on $S(Z)$. Also $\beta_1 = \mathbb{E}[s_0(Z)]$ is the average treatment effect (ATE) and $\beta_2 = \frac{Cov(S(Z), s_0(Z))}{Var(S(Z))}$ is the heterogeneity loading parameter (HET).

⁸¹Chernozhukov et al. (forthcoming) consider quintiles. We adapted the procedure to quartiles in the context of our application, as the rate of success of the lottery to assign applicants to the program is roughly 25%.

there is heterogeneity in the treatment effect beyond what is explained by the covariates $s_0(z)$. [Buhl-Wiggers et al. \(2022\)](#) shows that the proportion of the variance captured by estimators of treatment effect heterogeneity does not exceed a few percent of the lower bound of the treatment effect variance, even when resorting to ML estimators. The second reason is that the estimated parameters γ_k identify the expectation of the treatment effect conditionally on belonging to the group G_k . However, the group G_4 , for example, identifies individuals in the upper quartile of the estimated treatment effect $S(Z)$. Nothing describes how these individuals compare to those in the upper quartile of the true heterogeneous treatment effect $s_0(Z)$. For these two populations to be considered as approaching each other when the sample size becomes large, a property of convergence of the point estimators $S(Z)$ is needed, which is precisely what the analysis of [Chernozhukov et al. \(forthcoming\)](#) accomplishes.

A natural next step is to study the characteristics of the groups of interest (i.e. $\mathbb{E}[g(Z)|G_k]$, where $g(Z)$ is the vector of characteristics of an observation).

In particular, we can compare baseline characteristics between the top and bottom quartile of the distribution of predicted impacts, namely groups G_4 and G_1 . Although machine learning methods do not allow us to exactly identify which characteristics matter the most for heterogeneous treatment effects, learning about the characteristics of those who benefit the most and the least provides insights about the variables that could be used for targeting.

In the empirical section, we also assess how belonging to a particular group (“heterogeneity group”) for a given outcome Y affects treatment effect on another outcome $\tilde{Y}$. In other words, we seek to identify $\mathbb{E}[S_{\tilde{Y}}(Z)|G_{Y,k}]$, where $S_{\tilde{Y}}(Z)$ is the treatment effect on variable $\tilde{Y}$ conditional on Z and $G_{Y,k}$ is the k^{th} heterogeneity group for the treatment effect on the variable Y conditional on Z . This is useful to determine whether there are trade-offs between optimizing selection into the program to maximize during-program impacts and post-program impacts. This is also useful for our understanding of mechanisms for longer-term impacts: for example, we can assess whether individuals that benefit most from the program in terms of during-program earnings are also those with the greatest post-program savings or post-program well-being. In practice, we can use equation (4) to perform this analysis, replacing Y as a dependent variable with the alternative outcome variable $\tilde{Y}$.

E.2 Sample for Machine Learning Implementation

Supervised machine learning algorithms require samples for which both covariates (features) and outcomes are observed. In our case, this requires baseline covariates (a set of K covariates, Z^K) and midline or endline outcome of interest (Y). As discussed in the text, our study data has two specificities. First, our midline sample is a subsample of the baseline, while the full baseline sample is included at endline. Second, some control individuals entered subsequent waves of the public works program between midline and endline surveys, and are excluded from the endline sample used for analysis. As a result, the algorithms can use three potential samples: a ‘midline’ (Z^K, Y^{During}, T) (respectively ‘endline’ (Z^K, Y^{Post}, T)) sample can be used to build and apply the model to predict ‘during’ (respectively ‘post’) conditional treatment effects, where T corresponds to the treatment variable. A third (marginally smaller) sample can be used to study how effects vary ‘during’ and ‘post’ program by taking the intersection of non-attriters and non-missing outcomes for both surveys.⁸² When applying the algorithm on the endline data, we drop control individuals who applied to a later wave of the public works program (as in the main analysis).⁸³

The final sample size depends on the number of missing variables for the outcome considered. The total sample we use ranges between 2,884 and 2,958 units for midline and between 3,745 and 3,910 units for endline.

We use a set (Z^K , with $K = 21$) of features (covariates) measured at baseline (Table F1). They include both individual and household characteristics, as well as main indicators on employment, financial situation and self-reported constraints on basic needs expenditures. We also show the robustness of the main results to the inclusion of all baseline variables in the balance check table (Table 1).

E.3 Machine Learning Algorithms

We consider five alternative machine learning algorithms to estimate the proxy predictors and apply the procedure in Chernozhukov et al. (forthcoming): elastic net, boosted trees, random forest, Rlearner (based on elastic net and proposed by Nie and Wager (2020)) and generalized

⁸²There is also some attrition between survey rounds, and some missing values in baseline covariates. We exclude from each sample the attriters from follow-up surveys (since outcomes are not observed).

Missing values among baseline covariates are replaced by the mean in the same strata. Individuals with missing values for the outcome of interest (among nonattriters) are dropped from the sample.

⁸³Recall that 200 individuals were sampled to be added to the control group at endline to compensate for these observations: because these individuals were not part of the baseline survey, the machine learning model cannot be applied to them since predictions rely on observed Z^K .

random forest (proposed by [Wager and Athey \(2018\)](#)). All algorithms are implemented in **R**, and we adapted the codes provided by the authors.⁸⁴ These machine learning methods can be divided in two groups based on the way they approach the CATE function ([Künzel et al. \(2017\)](#)).

1. Two learners

The first group of machine learning methods includes Elastic Net, Random Forest and Boosted Trees. They predict separately $\mathbb{E}[Y(1)]$ in the treatment group and $\mathbb{E}[Y(0)]$ in the control group. In practice, a first model is fitted on the treatment group and a second on the control group, using an auxiliary sample.⁸⁵ The two fitted models are then used to predict potential outcomes $\hat{Y}(1)$ and $\hat{Y}(0)$ for each individual in the remaining sample (main sample). In order to obtain $S(Z)$, the difference between the two predictions for each individual are computed. All models are implemented using the **caret** package ([Kuhn \(2008\)](#)) (respectively named `glmnet`, `ranger` and `gbm`).

Tuning parameters

For each split, the tuning parameters are chosen separately for the model on the control and the treatment group. There are no set rule to choose these parameters. In our case, we let **caret** define a default search grid and we set a relatively high tuning length for all models based on our computational capacities. Tuning parameters were all selected based on the mean squared error estimates and 5-folds cross validation. For all methods, we pre-process outcomes and covariates and center-scale them before feeding the model.

For each method we have the following tuning parameters :

- Elastic net : *alpha* (Mixing Percentage), *lambda* (Regularization Parameter)
- Boosted trees : *n.trees* (# Boosting Iterations), *interaction.depth* (Max Tree Depth), *shrinkage* (Shrinkage), *n.minobsinnode* (Min. Terminal Node Size),
- Random Forest : *mtry* (# Randomly Selected Predictors), *splitrule* (Splitting Rule), *n.minobsinnode*

2. Single learners

⁸⁴<https://github.com/demirermert/MLInference>

⁸⁵The sample is split between an auxiliary sample where machine learning predictors are trained and constructed and a main sample where they are used for prediction and on which we estimate the different key features of the CATE function.

The two alternative models we consider are Rlearner and Generalized Random Forest (with their variations). They are “single learners” and use a different approach to approximate $s_0(Z)$. Instead of fitting a model on the treatment and on the control group separately to estimate $\mathbb{E}[Y(1)]$ and $\mathbb{E}[Y(0)]$, they directly fit a model to estimate $\mathbb{E}[Y(1)] - \mathbb{E}[Y(0)]$. [Athey and Imbens \(2016\)](#) discuss the benefits of this approach compared to the two-learners approach. One remaining quantity, the Baseline Conditional Average $b_0(Z)$, is needed. For Rlearner with boosting, we use boosted trees fitted on the control group to estimate $b_0(Z)$ and symmetrically elastic net for Rlearner based on Lasso. For Generalized Random Forest we predict $b_0(Z)$ using the random forest already fitted on the control group. We rely on the **grf** package to implement Generalized Random Forest and on the **rlearner**⁸⁶ package for Rlearner.

Tuning parameters

For each split of the data, we choose the tuning parameters separately for $S(Z)$ and $B(Z)$. Again, there is no theoretical basis to determine the choice of search grid parameters. We keep the package default grid and put a convenient length of parameters combination according to our computational capabilities.

Parameters were selected based on the mean squared error estimates and 5-folds cross validation.

E.4 Adaptation to the Experimental Setting

In our application, we repeat $S = 100$ times the procedure developed by [Chernozhukov et al. \(forthcoming\)](#).⁸⁷

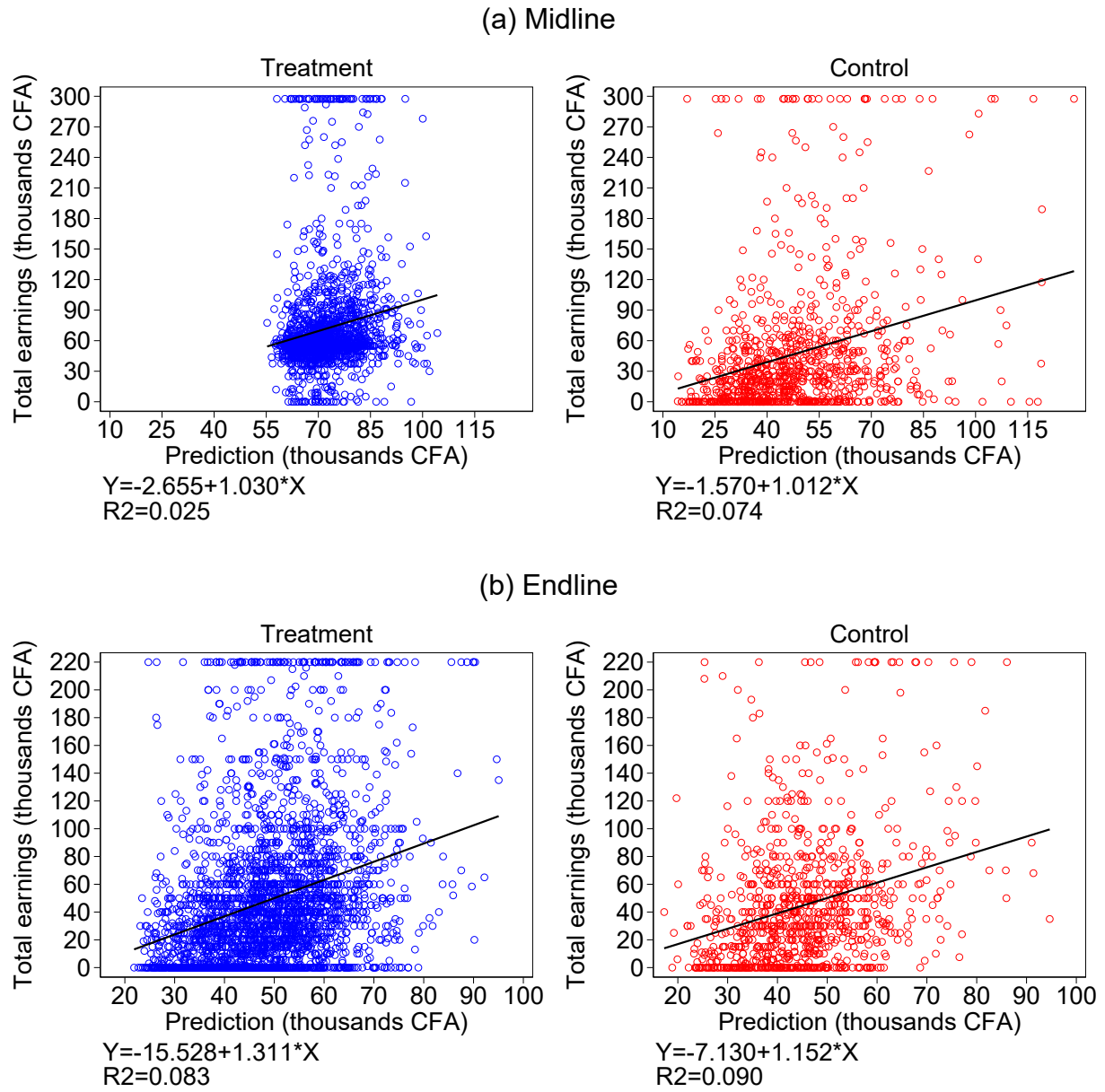
The first step of the method requires partitioning our dataset into an auxiliary and a main sample. We adapted the algorithm so that the sample splits are stratified by our randomization blocks (*locality * gender*), which represents 32 strata. This is important to preserve the identification strategy when estimating directly the CATE for the single learners, since they fit a model on different splits of the data.

⁸⁶<https://github.com/xnie/rlearner>

⁸⁷Figure F1 shows the scatter plot of earnings y and $\hat{y}$ for midline and endline as well as the regression line of y on $\hat{y}$. The figure shows that the slope coefficients are close to 1 and that the R^2 remain low. Note that a low R^2 on $y(1)$ and $y(0)$ does not mean that our algorithm cannot identify heterogeneity in the treatment effects, which is related to difference between $y(1)$ and $y(0)$. This is illustrated in the paper where we find heterogeneity in impacts on midline earnings.

Lastly, we introduce two adjustments in the linear projections of Best Linear Predictor (BLP) and Group Average Treatment Effects (GATE), along with predicted baseline effect $B(Z)$ and predicted treatment effect $S(Z)$. We add locality-gender fixed effects, corresponding to the randomization stratification variables. We also adjust the weights used. In the main specification of the paper, we use weights to take into account randomized assignment by lotteries, survey attrition, and sub-sampling at midline. Since our survey weights will be multiplied by inverse propensity score weights as recommended by [Chernozhukov et al. \(forthcoming\)](#), we make sure not to incorporate two times the inverse propensity score.

Figure F1: Relation between predictions and actual earnings (in level) (Random forest)



Note: Predictions estimated using equation $\hat{Y} = \hat{B}(Z) + T * \hat{\tau}(Z)$. Each value plotted is the median across 100 sample splits.

Table F1: Baseline variables used in Machine Learning algorithms

Variable description	Type
<i>Individual characteristics</i>	
Female	Binary
Age	Continuous
Number of children	Continuous
Live in urban area	Binary
<i>Education</i>	
Education (total number of years)	Continuous
Has participated in vocational training	Binary
<i>Household characteristics</i>	
Household size	Continuous
Is head of household	Binary
<i>Household assets</i>	
Total number of assets ¹	Continuous
<i>Employment</i>	
Total number of activities	Continuous
Total number of wage-employment activities	Continuous
Total number of self-employment activities	Continuous
Is engaged in (at least one) casual activity	Binary
Total Earnings (monthly)	Continuous
<i>Savings, Expenditures and Constraints</i>	
Has Saved (last 3 months)	Binary
Savings Stock (FCFA)	Continuous
Has a Savings Account	Binary
Self-reported constraints to repay loans	Binary
Self-reported constraints to access credit	Binary
Transportation expenditures (last 7 days, CFA)	Continuous
Communication expenditures (last 7 days, CFA)	Continuous

[1] Assets include livestock, chicken, other animals, plows, field sprayer, carts, wheelbarrows, bicycles, motorcycles, pirogues, refrigerators, freezers, air conditioning units, fans, stoves, computers, radios, television, TV antenna, video players, landline, mobile phones, cars.

Appendix F Weights

This appendix describes the weights used in the analysis. Table G1 summarizes the weights used with midline data, and Table G2 with endline data. In general results are robust to weights not being included.

Randomization weights

We consider two sets of randomization weights. First, for both midline and endline, we consider weights that account for variations in selection probability by lottery location and gender. There are K different public lotteries ($K = 32$) with N_k individuals participating to each lottery. Denote N_{k1} the individuals from lottery k selected in the program ('treated') and N_{k0} those who are not selected, with $N_k = N_{k1} + N_{k0}$. Among the N_{k0} , N_{k0s} are randomly drawn to be surveyed and constitute the 'control group'. The size of the population of lottery participants is N_P , with $N_P = \sum_k N_k = N_1 + N_0$. The size of the survey sample is $N_E = \sum_k N_{k1} + N_{k0s} = N_1 + N_{0s}$. We use weight w_{ki} ($i = 0, 1$ according to treatment status) for individuals in the survey sample, with $w_{k1} = N_k / N_{k1} \times N_1 / N_P$ and $w_{k0s} = N_k / N_{k0s} \times N_0 / N_P$. This means that we put a higher weight on lotteries where the demand for the program (total population participating in the lottery) was higher, compared with other lotteries.

Second, when estimating treatment effects by arm using endline survey data, we also consider that the number of brigades assigned to each treatment arm varies by locality. Brigades of treated individuals (N_1) are assigned to 3 treatment options T_a , T_b and T_c . We use the following notation: $N_k = N_{a,k} + N_{b,k} + N_{c,k} + N_{0,k}$ with $N_{1,k} = N_{a,k} + N_{b,k} + N_{c,k}$, and $N_P = \sum_k N_k = N_0 + N_a + N_b + N_c$ with $N_1 = N_a + N_b + N_c$. We put a weight $w_{j,k}$ to treated individuals from lottery k who were assigned to treatment T_j , and a weight w_{k0s} for control individuals in the survey sample, with:

- $w_{j,k} = N_{k1} / N_{j,k} \times N_j / N_1$ with $j = a, b, c$ ⁸⁸
- $w_{k0s} = N_{k1} / N_{k0s} \times N_0 / N_1$

Sub-sampling weights (midline survey only)

The sample for the midline survey includes the control group (N_{0s}) and a sub-sample of the treatment group. Consider that we draw a random sub sample of group l in proportion $P_l = N_l^S / N_l$.

⁸⁸Note : $\sum_j w_{j,k} = w_{k1} = 1$, which is the weight used for midline data when there is only one treatment group.

To take sub-sampling into account, original weights are multiplied by S/P_l . Therefore, in group $l = k, 1$ we draw N_{k1}^S individuals out of N_{k1} , and the original weight w_{k1} becomes $\omega_{k1}^S = w_{k1} \times N_{k1}/N_{k1}^{S_{k1}}$. All control units are included in the midline sample so that their weights w_{k0s} are unchanged.

Control group and subsequent enrollment in the program (endline survey only)

When using endline data, we adjust weights for control individuals because some of them were able to apply (and sometimes get selected) in waves 3 and 4 of the program, as discussed in section 4.⁸⁹ Weights for control individuals depend on their status in wave 3 and wave 4, which is one of the following 7 situations:

1. Group $C_3T_3\bar{C}_4$: Applied to wave 3 (C_3), was selected as ‘beneficiary’ of wave 3 after public lotteries (T_3) and was therefore not allowed to apply to wave 4 ($\bar{C}_4$).
2. Group $C_3\bar{T}_3C_4T_4$: Applied to wave 3 (C_3), was not selected after public lotteries ($\bar{T}_3$), applied to wave 4 (C_4) and was selected as ‘beneficiary’ of wave 4 after lotteries (T_4).
3. Group $C_3\bar{T}_3C_4\bar{T}_4$: Applied to wave 3 (C_3), was not selected after public lotteries ($\bar{T}_3$), applied to wave 4 (C_4) and was not selected after lotteries ($\bar{T}_4$).
4. Group $C_3\bar{T}_3\bar{C}_4$: Applied to wave 3 (C_3), was not selected after public lotteries ($\bar{T}_3$) and did not apply to wave 4 ($\bar{C}_4$).
5. Group $\bar{C}_3C_4T_4$: Did not apply to wave 3 ($\bar{C}_3$), applied to wave 4 (C_4) and was selected as ‘beneficiary’ of wave 4 after public lotteries (T_4).
6. Group $\bar{C}_3C_4\bar{T}_4$: Did not apply to wave 3 ($\bar{C}_3$), applied to wave 4 (C_4) and was not selected after public lotteries ($\bar{T}_4$).
7. Group $\bar{C}_3\bar{C}_4$: Did not apply to wave 3 ($\bar{C}_3$), and did not apply to wave 4 ($\bar{C}_4$).

We introduce a new multiplicative weight for control units ($\tilde{w}_{k0s,j}$). We do not include control units that have benefited from subsequent waves of the program (waves 3 and 4) in the estimation. This means we assign a weight of 0 to groups $C_3T_3\bar{C}_4$, $C_3\bar{T}_3C_4T_4$ and $\bar{C}_3C_4T_4$.⁹⁰ To compensate, we put a higher weight on individuals who also applied in subsequent phases (waves 3 and 4) but were not selected during the lotteries. The weights for the remaining four groups are:

⁸⁹Recall that the study focuses on wave 2 (out of 4 waves) of the public works program

⁹⁰Hence $\tilde{w}_{k0s,C_3T_3\bar{C}_4} = 0$; $\tilde{w}_{k0s,\bar{C}_3C_4T_4} = 1 \times 0 = 0$; $\tilde{w}_{k0s,C_3\bar{T}_3C_4T_4} = \frac{N_{k0s,C_3}}{N_{k0s,C_3\bar{T}_3}} \times 0 = 0$.

- $\tilde{w}_{k0s,C_3\bar{T}_3C_4\bar{T}_4} = \frac{N_{k0s,C_3}}{N_{k0s,C_3\bar{T}_3}} \times \frac{N_{k0s,C_3\bar{T}_3C_4}}{N_{k0s,C_3\bar{T}_3C_4\bar{T}_4}}$
- $\tilde{w}_{k0s,C_3\bar{T}_3\bar{C}_4} = \frac{N_{k0s,C_3}}{N_{k0s,C_3\bar{T}_3}} \times 1 = \frac{N_{k0s,C_3}}{N_{k0s,C_3\bar{T}_3}}$
- $\tilde{w}_{k0s,\bar{C}_3C_4\bar{T}_4} = 1 \times \frac{N_{k0s,\bar{C}_3C_4}}{N_{k0s,\bar{C}_3C_4\bar{T}_4}}$
- $\tilde{w}_{k0s,\bar{C}_3\bar{C}_4} = 1$

Tracking weights

Lastly, we add a weight taking into account the differential response rate of individuals during each survey (midline and endline). More precisely, each survey consisted in two phases a and b :

- A main data collection phase (a), during which the response rate is $R_{a,j}$ for group $j = 1, 0$.
- An additional tracking phase (b), targeting attriters from the main phase. We note $R_{b,j}$, the response rate of the tracking phase for group $j = 1, 0$.

To determine the tracking sample, we first define a sub-sample of ‘eligible’ attriters.⁹¹ $E_{b,j}$ from which a random sub-sample is drawn in proportion $\pi_j = NE_{b,j}^S / NE_{b,j}$ (j is an index for treatment status * locality). Individuals interviewed during the tracking phase take a different weight than those interviewed during the main survey phase. Tracking respondents are weighted by $\omega_j^T = R_{a,j}^S + \lambda_j s_j R_{b,j}^S (1 - R_{a,j}^S) E_{b,j}^S$, with λ_j , so that the final weight is $\omega_j^{S,f} = \omega_j^S \times \omega_j^T$.

The sum of the weights on population j is therefore : $\omega_j \times (N_{a,j}^S + \lambda_j NER_{s,b,j}^S)$, with $NER_{s,b,j}^S$ the number of individuals from the tracking sample who responded during tracking phase. We make the hypothesis that residual non-response $R_{b,j}^S$ is random. We seek to be representative of the respondent population of phases a and b . Therefore, we take $\lambda_j = NE_{b,j}^S / NER_{s,b,j}^S$

In group j , weights will be set such as:⁹²

- $\omega_j^S \times 1$ for phase a respondents
- $\omega_j^S \times NE_{b,j}^S / NER_{s,b,j}^S$ for phase b respondents

⁹¹Among the attriters of phase (a) some individuals were considered ‘ineligible’ for tracking as they were (quasi) impossible to reach: dead individuals, individuals who migrated to another country, (for endline) individuals who were already impossible to find at baseline.

⁹²In theory, ω_j should be adjusted so that it does not use correction N_j / N_j^S but rather the correction corresponding to the total of eligibles $N_{a,j} + NE_{b,j}$. However, this number is only known for selected units $S_j = 1$. Therefore we will ignore this aspect, which is fair considering that units were randomly drawn. Finally, it means that we estimate the unknown amount $N_{a,j} + NE_{b,j}$ by $N_{a,j}^S + NE_{b,j}^S \times N_j / N_j^S$

Table G1: Summary of weights used with midline data

Randomization weights w_k		Sub Sampling weights ω_k^S		Tracking weights ω_j^T	
Treated	$w_{k1} = N_k / N_{k1} \times N_1 / N_P$	Treated	$w_{k,1}^S = N_{k1} / N_{k1}^{S_{k1}}, k=\text{locality}$	Respondents main phase ($R_a = 1$)	$\omega^T = 1$
Control	$w_{k0s} = N_k / N_{k0s} \times N_0 / N_P, k=\text{locality} * \text{gender}$	Control	$w_{k,0}^S = 1$	Non Respondents main phase ($R_a = 0$)	$\omega_j^T = NE_{b,j}^S / NER_{s,b,j}^S$ if respondent in tracking phase ($E_b = 1$ and $R_b = 1$), $j=\text{locality} * \text{treatment status}$
					$\omega^T = 0$ if non respondent (but sampled) in tracking phase ($E_b = 1$ et $R_b = 0$)
					$\omega^T = 0$ if not sampled for tracking phase ($E_b = 0$)
Final weight: $w_{k,i}^F = w_{k,i} \times \omega_{k,i}^S \times \omega_{k,i}^T, i = 0, 1$ (treatment status), $k \in \llbracket 1, 32 \rrbracket$ (locality * gender)					

Table G2: Summary of weights used with endline data

Randomization weights $w_{j,k}$		Post-enrollment weights $\tilde{\omega}_{k,j}$		Tracking weights ω_j^T	
Treatment arm T_a, T_b or T_c	$w_{j,k} = N_k / N_{j,k} \times N_j / N_P, j = a, b, c$	Selected to participate to wave 3 or 4 (groups $C_3 T_3 \bar{C}_4, C_3 \bar{T}_3 C_4 T_4$ et $\bar{C}_3 C_4 T_4$)	0	Respondents main phase ($R_a = 1$)	$\omega^T = 1$
Control	$w_{k0s} = N_k / N_{k0s} \times N_0 / N_P$, k =locality * gender	Group $C_3 \bar{T}_3 C_4 \bar{T}_4$	$\frac{N_{k0s,C3}}{N_{k0s,C3\bar{T}_3}} \times \frac{N_{k0s,C3\bar{T}_3C4}}{N_{k0s,C3\bar{T}_3C4\bar{T}_4}}$	Non Respondents main phase ($R_a = 0$)	$\omega_j^T = NE_{b,j}^S / NER_{s,b,j}^S$ if respondent in tracking phase ($E_b = 1$ and $R_b = 1$), j =locality * treatment status
		Group $C_3 \bar{T}_3 \bar{C}_4$	$\frac{N_{k0s,C3}}{N_{k0s,C3\bar{T}_3}}$		$\omega^T = 0$ if non respondent (but sampled) in tracking phase ($E_b = 1$ et $R_b = 0$)
		Group $\bar{C}_3 C_4 \bar{T}_4$	$\frac{N_{k0s,\bar{C}_3C4}}{N_{k0s,\bar{C}_3C4\bar{T}_4}}$		$\omega^T = 0$ if not sampled for tracking phase ($E_b = 0$)
		Group $\bar{C}_3 \bar{C}_4$	1		
Final weight: $w_{k,i}^F = w_{j,k} \times \tilde{\omega}_{i,l} \times \omega_{k,i}^T, j = 0, a, b, c$ (treatment status), $i = 1, 0s, l$ post-enrollment group, $k \in \llbracket 1, 32 \rrbracket$ (locality * gender)					

References

- ALATAS, V., A. BANERJEE, R. HANNA, B. A. OLKEN, AND J. TOBIAS (2012): “Targeting the poor: Evidence from a field experiment in Indonesia,” *American Economic Review*, 102, 1206–40.
- ALATAS, V., R. PURNAMASARI, M. WAI-POI, A. BANERJEE, B. A. OLKEN, AND R. HANNA (2016): “Self-targeting: Evidence from a field experiment in Indonesia,” *Journal of Political Economy*, 124, 371–427.
- ALIK-LAGRANGE, A., N. BUEHREN, M. GOLDSTEIN, AND J. HOOGVEEN (2023): “Welfare impacts of public works in fragile and conflict affected economies: The Londö public works in the Central African Republic,” *Labour Economics*, 81, 102293.
- ALIK-LAGRANGE, A. AND M. RAVALLION (2018): “Workfare versus transfers in rural India,” *World Development*, 112, 244–258.
- ALMEIDA, R. K. AND E. GALASSO (2010): “Jump-starting self-employment? Evidence for welfare participants in Argentina,” *World Development*, 38, 742–755.
- AMARAL, S., S. BANDYOPADHYAY, AND R. SENSARMA (2015): “Employment programmes for the poor and female empowerment: The effect of NREGS on gender-based violence in India,” *Journal of Interdisciplinary Economics*, 27, 199–218.
- APOSTOLIDIS, T. AND N. FIEULAIN (2004): “Validation française de l’échelle de temporalité: The Zimbardo Time Perspective Inventory (ZTPI),” *European Review of Applied Psychology*, 54, 207–217.
- ATHEY, S. AND G. IMBENS (2016): “Recursive partitioning for heterogeneous causal effects,” *Proceedings of the National Academy of Sciences*, 113, 7353–7360.
- ATHEY, S., J. TIBSHIRANI, AND S. WAGER (2019): “Generalized random forests,” *Annals of Statistics*, 47, 1148–1178.
- ATHEY, S. AND S. WAGER (2021): “Policy Learning With Observational Data,” *Econometrica*, 89, 133–161.
- BAGGA, A., M. HOLMLUND, N. KHAN, S. MANI, E. MVUKIYEHE, AND P. PREMAND (2024): “Do Public Works Programs Have Sustained Impacts?: A Review of Experimental Studies from LMICs,” *World Bank Research Observer*.
- BANERJEE, A., R. HANNA, B. A. OLKEN, AND D. SVERDLIN LISKER (2024): “Social Protection in the Developing World,” *Journal of Economic Literature*, 62, 1349–1421.
- BANERJEE, A., P. NIEHAUS, AND T. SURI (2019): “Universal Basic Income in the Developing World,” *Annual Review of Economics*, 11, 959–983.
- BASURTO, M. P., P. DUPAS, AND J. ROBINSON (2020): “Decentralization and efficiency of subsidy targeting: Evidence from chiefs in rural Malawi,” *Journal of Public Economics*, 185.
- BEEGLE, K., E. GALASSO, AND J. GOLDBERG (2017): “Direct and indirect effects of Malawi’s public works program on food security,” *Journal of Development Economics*, 128, 1–23.
- BERG, E., S. BHATTACHARYYA, D. RAJASEKHAR, AND R. MANJULA (2018): “Can public works increase equilibrium wages? Evidence from India’s National Rural Employment Guarantee,” *World Development*, 103, 239–254.

- BERHANE, G., D. O. GILLIGAN, J. HODDINOTT, N. KUMAR, AND A. S. TAFFESSE (2014): “Can Social Protection Work in Africa? The Impact of Ethiopia’s Productive Safety Net Programme,” *Economic Development and Cultural Change*, 63, 1–26.
- BERTRAND, M., B. CRÉPON, A. MARGUERIE, AND P. PREMAND (2016): “Impacts à Court et Moyen Terme sur les Jeunes des Travaux à Haute Intensité de Main d’œuvre (THIMO): Résultats de l’évaluation d’impact de la composante THIMO du Projet Emploi Jeunes et Développement des compétences (PEJEDEC) en Côte d’Ivoire,” Washington DC: Banque Mondiale et Abidjan: BCP-Emploi.
- (2025a): “Côte d’Ivoire - Projet Emploi Jeune et Développement des Compétences (PEJEDEC) Public Works Impact Evaluation - Baseline Survey, 2013 (PEJEDEC PWP BL 2013).” Downloaded from - <https://doi.org/10.48529/stxm-dw72>.
- (2025b): “Côte d’Ivoire - Projet Emploi Jeune et Développement des Compétences (PEJEDEC) Public Works Impact Evaluation - Endline Survey, 2015 (PEJEDEC PWP EL 2015).” Downloaded from - <https://doi.org/10.48529/6zxm-x259>.
- (2025c): “Côte d’Ivoire - Projet Emploi Jeune et Développement des Compétences (PEJEDEC) Public Works Impact Evaluation - Midline Survey, 2013 (PEJEDEC PWP ML 2013).” Downloaded from - <https://doi.org/10.48529/p1ec-xe78>.
- BESLEY, T. AND S. COATE (1992): “Workfare versus welfare: Incentive arguments for work requirements in poverty-alleviation programs,” *The American Economic Review*, 82, 249–261.
- BHATTACHARYA, D. AND P. DUPAS (2012): “Inferring welfare maximizing treatment assignment under budget constraints,” *Journal of Econometrics*, 167, 168–196.
- BITLER, M. P., J. B. GELBACH, AND H. W. HOYNES (2008): “Distributional impacts of the Self-Sufficiency Project,” *Journal of Public Economics*, 92, 748–765.
- BLAIR, G. AND K. IMAI (2012): “Statistical analysis of list experiments,” *Political Analysis*, 20, 47–77.
- BLATTMAN, C. AND S. DERCON (2018): “The Impacts of Industrial and Entrepreneurial Work on Income and Health: Experimental Evidence from Ethiopia,” *American Economic Journal: Applied Economics*, 10, 1–38.
- BREDA, T., J. GRENET, M. MONNET, AND C. VAN EFFENTERRE (2023): “How Effective are Female Role Models in Steering Girls Towards STEM? Evidence from French High Schools,” *The Economic Journal*, 133, 1773–1809.
- BRYAN, G., D. KARLAN, AND A. OSMAN (2024): “Big Loans to Small Businesses: Predicting Winners and Losers in an Entrepreneurial Lending Experiment,” *American Economic Review*, 114, 2825–60.
- BUHL-WIGGERS, J., J. T. KERWIN, J. MUÑOZ-MORALES, J. SMITH, AND R. THORNTON (2022): “Some children left behind: Variation in the effects of an educational intervention,” *Journal of Econometrics*.
- CHERNOZHUKOV, V., M. DEMIRER, E. DUFLO, AND I. FERNANDEZ-VAL (forthcoming): “Fisher-Schultz Lecture: Generic Machine Learning Inference on Heterogenous Treatment Effects in Randomized Experiments, with An Application to Immunization in India,” *Econometrica*.

- CHRISTIAENSEN, L. AND P. PREMAND (2017): “Cote d’Ivoire Jobs Diagnostic : Employment, Productivity, and Inclusion for Poverty Reduction.” World Bank, Washington DC.
- CRÉPON, B. AND P. PREMAND (2024): “Direct and Indirect Effects of Subsidized Dual Apprenticeships,” *The Review of Economic Studies*, rdae094.
- CÔTÉ, G., P. GOSSELIN, AND I. DAGENAIS (2013): “Évaluation multidimensionnelle de la régulation des émotions : propriétés psychométriques d’une version francophone du Difficulties in Emotion Regulation Scale,” *Journal de Thérapie Comportementale et Cognitive*, 23.
- DATT, G. AND M. RAVALLION (1994): “Transfer benefits from public-works employment: Evidence for rural India,” *The Economic Journal*, 1346–1369.
- DEININGER, K., H. K. NAGARAJAN, AND S. K. SINGH (2016): “Short-term effects of India’s employment guarantee program on labor markets and agricultural productivity,” Policy Research Working Paper No. 7665, The World Bank, Washington DC.
- DJEBBARI, H. AND J. SMITH (2008): “Heterogeneous impacts in PROGRESA,” *Journal of Econometrics*, 145, 64–80.
- DROITCOUR, J., R. A. CASPAR, M. L. HUBBARD, T. L. PARSLEY, W. VISSCHER, AND T. M. EZZATI (1991): “The Item Count Technique as a method of indirect questioning: A review of its development and a case study application,” in *Measurement errors in surveys*, ed. by P. P. Biemer, R. M. Groves, L. E. Lyberg, N. A. Mathiowetz, and S. Sudman, John Wiley & Sons Inc., 185–210.
- FETZER, T. (2020): “Can Workfare Programs Moderate Conflict? Evidence from India,” *Journal of the European Economic Association*, 18, 3337–3375.
- FILMER, D., L. FOX, K. BROOKS, A. GOYA, T. MENGISTAE, P. PREMAND, D. RINGOLD, S. SHARMA, AND S. ZORYA (2014): “*Youth Employment in sub-Saharan Africa*”, World Bank, Africa Development Series.
- FRANKLIN, S., C. IMBERT, G. ABEBE, AND C. MEJIA-MANTILLA (2024): “Urban Public Works in Spatial Equilibrium: Experimental Evidence from Ethiopia,” *American Economic Review*, 114, 1382–1414.
- GALASSO, E. AND M. RAVALLION (2004): “Social protection in a crisis: Argentina’s Plan Jefes y Jefas,” *The World Bank Economic Review*, 18, 367–399.
- GALASSO, E., M. RAVALLION, AND A. SALVIA (2004): “Assisting the transition from workfare to work: A randomized experiment,” *ILR Review*, 58, 128–142.
- GAZEAUD, J., E. MVUKIYEHE, AND O. STERCK (2023): “Cash Transfers and Migration: Theory and Evidence from a Randomized Controlled Trial,” *The Review of Economics and Statistics*, 105, 143–157.
- GEHRKE, E. AND R. HARTWIG (2018): “Productive effects of public works programs: What do we know? What should we know?” *World development*, 107, 111–124.
- GENTILINI, U., M. GROSH, J. RIGOLINI, AND R. YEMTSOV (2020): “Exploring Universal Basic Income : A Guide to Navigating Concepts, Evidence, and Practices,” Washington, DC: World Bank.
- GILLIGAN, D. O., J. HODDINOTT, AND A. S. TAFFESSE (2009): “The impact of Ethiopia’s

- Productive Safety Net Programme and its linkages,” *The Journal of Development Studies*, 45, 1684–1706.
- GOODMAN, R., H. MELTZER, AND V. BAILEY (1998): “The Strengths and Difficulties Questionnaire: A pilot study on the validity of the self-report version,” *European child & adolescent psychiatry*, 7, 125–130.
- HANNA, R. AND B. A. OLKEN (2018): “Universal Basic Incomes versus Targeted Transfers: Anti-Poverty Programs in Developing Countries,” *Journal of Economic Perspectives*, 32, 201–26.
- HAUSHOFER, J., P. NIEHAUS, C. PARAMO, E. MIGUEL, AND M. WALKER (2025): “Targeting Impact versus Deprivation,” *American Economic Review*, 115, 1936–74.
- HECKMAN, J. J., J. SMITH, AND N. CLEMENTS (1997): “Making the most out of programme evaluations and social experiments: Accounting for heterogeneity in programme impacts,” *The Review of Economic Studies*, 64, 487–535.
- HODDINOTT, J., G. BERHANE, D. O. GILLIGAN, N. KUMAR, AND A. SEYOUM TAFFESSE (2012): “The Impact of Ethiopia’s Productive Safety Net Programme and Related Transfers on Agricultural Productivity,” *Journal of African Economies*, 21, 761–786.
- HUSSAM, R., E. M. KELLEY, G. LANE, AND F. ZAHRA (2022): “The Psychosocial Value of Employment: Evidence from a Refugee Camp,” *American Economic Review*, 112, 3694–3724.
- IMBERT, C. AND J. PAPP (2015): “Labor market effects of social programs: Evidence from india’s employment guarantee,” *American Economic Journal: Applied Economics*, 7, 233–63.
- (2019): “Short-term Migration, Rural Public Works, and Urban Labor Markets: Evidence from India,” *Journal of the European Economic Association*, 18, 927–963.
- IVASCHENKO, O., D. NAIDOO, D. NEWHOUSE, AND S. SULTAN (2017): “Can public works programs reduce youth crime? Evidence from Papua New Guinea’s Urban Youth Employment Project.” *IZA Journal of Development and Migration*, 7.
- JALAN, J. AND M. RAVALLION (2003): “Estimating the benefit incidence of an antipoverty program by propensity-score matching,” *Journal of Business & Economic Statistics*, 21, 19–30.
- KITAGAWA, T. AND A. TETENOV (2018): “Who should be treated? Empirical welfare maximization methods for treatment choice,” *Econometrica*, 86, 591–616.
- KNAUS, M. C., M. LECHNER, AND A. STRITTMATTER (2020): “Machine learning estimation of heterogeneous causal effects: Empirical Monte Carlo evidence,” *The Econometrics Journal*, utaa014.
- KUHN, M. (2008): “Building Predictive Models in R Using the caret Package,” *Journal of Statistical Software*, 28, 1–26.
- KÜNZEL, S., J. SEKHON, P. BICKEL, AND B. YU (2017): “Meta-learners for Estimating Heterogeneous Treatment Effects using Machine Learning,” *Proceedings of the National Academy of Sciences*, 116.
- MANSKI, C. F. (2004): “Statistical treatment rules for heterogeneous populations,” *Econometrica*, 72, 1221–1246.

- McKENZIE, D. (2021): “Small business training to improve management practices in developing countries: re-assessing the evidence for ‘training doesn’t work’,” *Oxford Review of Economic Policy*, 37, 276–301.
- MILLER, J. D. (1984): “A new survey technique for studying deviant behavior,” George Washington University, Washington, DC.
- MORIN, A., G. MOULLEC, C. MAÏANO, L. LAYET, J.-L. JUST, AND N. G. (2011): “Psychometric properties of the Center for Epidemiologic Studies Depression Scale (CESD) in French Clinical and Non-Clinical Adults,” *Epidemiology and Public Health/Revue d’Épidémiologie et de Santé Publique*, 59, 327–340.
- MURALIDHARAN, K., P. NIEHAUS, AND S. SUKHTANKAR (2016): “Building state capacity: Evidence from biometric smartcards in India,” *American Economic Review*, 106, 2895–2929.
- (2023): “General equilibrium effects of (improving) public employment programs: Experimental evidence from India,” *American Economic Review*, 91, 1261–1295.
- MURGAI, R., M. RAVALLION, AND D. VAN DE WALLE (2015): “Is Workfare Cost-effective against Poverty in a Poor Labor-Surplus Economy?” *The World Bank Economic Review*, 30, 413–445.
- NIE, X. AND S. WAGER (2020): “Quasi-oracle estimation of heterogeneous treatment effects,” *Biometrika*, asaa076.
- PEROVA, E., E. JOHNSON, A. MANNAVA, S. REYNOLDS, AND A. TEMAN (2021): “Public Work Programs and Gender-Based Violence: Evidence from Lao PDR,” Policy Research Working Paper No. 9691, World Bank, Washington, DC.
- PREMAND, P. AND P. SCHNITZER (2020): “Efficiency, Legitimacy, and Impacts of Targeting Methods: Evidence from an Experiment in Niger,” *The World Bank Economic Review*, lhaa019.
- QUINN, S. AND C. WOODRUFF (2019): “Experiments and Entrepreneurship in Developing Countries,” *Annual Review of Economics*, 11, 225–248.
- RAVALLION, M. (1991): “Reaching the Rural Poor through Public Employment: Arguments, Evidence, and Lessons from South Asia,” *The World Bank Research Observer*, 6, 153–175.
- RAVALLION, M., G. DATT, AND S. CHAUDHURI (1993): “Does Maharashtra’s Employment Guarantee Scheme guarantee employment? Effects of the 1988 wage increase,” *Economic Development and Cultural Change*, 41, 251–275.
- RAVALLION, M., E. GALASSO, T. LAZO, AND E. PHILIPP (2005): “What can ex-participants reveal about a program’s impact?” *Journal of Human Resources*, 40, 208–230.
- RAVI, S. AND M. ENGLER (2015): “Workfare as an effective way to fight poverty: The case of India’s NREGS,” *World Development*, 67, 57–71.
- ROSAS, N. AND S. SABARWAL (2016): “Public Works as a Productive Safety Net in a Post-Conflict Setting : Evidence from a Randomized Evaluation in Sierra Leone,” Policy research working paper no. 7580, World Bank, Washington, DC.
- SUBBARAO, K., C. DEL NINNO, C. ANDREWS, AND C. RODRÍGUEZ-ALAS (2012): *Public works as a safety net: Design, evidence, and implementation*, World Bank, Washington DC.

- VALLIERES, E. F. AND R. J. VALLERAND (1990): “Traduction et Validation Canadienne-Française de L’échelle de L’estime de Soi de Rosenberg,” *International Journal of Psychology*, 25, 305–316.
- WAGER, S. AND S. ATHEY (2018): “Estimation and Inference of Heterogeneous Treatment Effects using Random Forests,” *Journal of the American Statistical Association*, 113, 1228–1242.