

Theory and Long Term Play in Experimental Economics

David K. Levine

State of the Art

- Seventy years of experimental economics, what have we learned?
- What is our best theory predicting in advance how participants will play in an experiment?

I'm going to talk about long-term play

If people play many times we observe that their play stabilizes

But what does it stabilize to?

- Even today, the only widely used theory is some variant of Nash equilibrium
- This does poorly in a wide variety of experiments

in a typical paper I see an experiment, followed by Nash equilibrium did/did not work, and if it didn't work some special pleading about why it didn't work

Overview

things we agree on

- people enter the laboratory with preconceived notions of fairness and efficiency
- some reward pro-social behavior and punish anti-social behavior
- some strive for socially good outcomes

the question posed here

- how preconceived are these notions?
- do they adjust to circumstances in the laboratory?

Two Types of Theories

- psychological: hard-wired preferences, Fehr-Schmidt
- behavioral mechanism design: try to achieve a socially desirable outcome in the face of incentive constraints induced by selfish behavior
- problem with Fehr-Schmidt and theories of selfish players
 - people don't always manage to coordinate on good equilibria
 - behavioral mechanism design has explicit trembling, and this does a good job of predicting when coordination will be successful

I will mostly discuss behavioral mechanism design, but also Fehr-Schmidt

The Setting

finite game with n player roles denoted by i

strategy spaces $s^i \in S^i$, mixed σ^i

monetary payoffs $m^i(s^i, s^{-i})$

utility $u^i(m)$, $u^i(s)$

Behavioral Mechanism Design

three equally likely types of player, selfish, noise and ethical

selfish are “Nash”

noise player play (or tremble) (sort of) uniformly

ethical try to create the welfare maximizing mechanism

(“pick the best equilibrium”)

ethical willing to sacrifice up to a point (\$1.00 total over all paid rounds)

risk aversion (taken from lottery experiments)

$$u(m^i) = 1 - (1 + m^i/40)^{-8}$$

formulated as the mechanism design of maximizing welfare subject to incentive compatibility

Alternatives

refined perfection: welfare maximizing subgame perfect equilibrium in weakly undominated strategies and selfish risk neutral preferences

Fehr-Schmidt: refined perfection with Fehr-Schmidt preferences

$$u^i(m) = m^i - \left(\frac{1}{n-1} \right) \left(\sum_{j \neq i} (\alpha^i \max\{m^j - m^i, 0\} + \beta^i \max\{m^i - m^j, 0\}) \right)$$

α^i	β^i	ϕ^i	two-type
0.00	0.00	0.30	0.60
0.50	0.25	0.30	
1.0	0.60	0.30	
2.0	0.60		0.40
4.0	0.60	0.10	

“Standard” Experiments

- college students in laboratory for “normal stakes”
- experienced players, played the game at least nine times
- parameters calibrated to data on risk aversion and dictator giving

Overview

three classes of games, incentive games, coordination games and selfish games

29 treatments of 9 different games

- basic measure of fit: welfare, economically meaningful measure
- deviations between theory and data: relative error
 - proxy for dollars lost per hour
 - difference between theoretical and empirical welfare divided by empirical welfare
 - assumption: treatments try to pay similar empirical welfare per hour

Incentive Games

game	treatment	period(s)	paid	theory	data	errr	RP errr	FS errr	N
ult	no obs	10 – 40	1	\$3.45	\$3.44	0.00	0.02	0.57	64
	obs	10 – 40	1	\$3.45	\$3.43	0.01	0.03	0.41	80
pubp	pun 1	10	10	\$1.80	\$1.64	0.10	−0.09	−0.09	24
	pun 2	10	10	\$1.88	\$1.78	0.06	−0.16	0.35*	24
	pun 3	10	10	\$1.91	\$1.99	−0.04	−0.25	0.21*	24
	pun 4	10	10	\$1.92	\$1.91	0.01	−0.21	0.26*	24
gift	human	10	10	\$1.03	\$0.86	0.20	−0.80	0.53	44
	robot	10	10	\$1.05	\$1.15	−0.09	−0.30	−0.03	78

TABLE 2.4.1. Welfare in Incentive Games

paid: paid periods

errr: welfare difference divided by empirical welfare

RP: refined perfection

FS: Fehr-Schmidt

N: number of participants

* indicates that there were multiple equilibria and the welfare best was chosen

bold face: an anomaly defined as an error of more than 10% in absolute value

FS does a credible job of predicting effort levels and offers

Coordination Games

game	treatment	period(s)	paid	theory	data	errr	RP/FS errr	N
min	$n = 2p$	7	30	\$1.18	\$1.18	0.00	0.19	28
	$n = 14$	10	30	\$0.64	\$0.60	0.07	1.17	28
	$n = 15$	10	30	\$0.60	\$0.66	-0.09	0.97	30
	$n = 16$	10	30	\$0.60	\$0.61	-0.02	1.13	48
stag	$0.6R$	10 – 75	75	\$0.37	\$0.27	0.37	0.49	64
	R	10 – 75	75	\$0.37	\$0.27	0.37	0.49	64
	$2R$	10 – 75	75	\$0.37	\$0.36	0.06	0.29	64
IPD	$dal_fre:B;\delta = 0.50$	16+	60+	0.14	0.15	-0.01	0.47	32
	bru_kam	16+	60+	0.21	0.35	-0.09	0.43	40
	$dal_fre:A;\delta = 0.75$	16+	60+	0.27	0.25	0.01	0.20	24
	$dal_fre:C;\delta = 0.50$	16+	60+	0.67	0.39	0.19	0.41	24
	kag_sch	16+	60+	0.67	0.58	0.08	0.39	24
	she_tar_sai	16+	60+	0.67	0.69	-0.02	0.33	24
	$dal_fre:B;\delta = 0.75$	16+	60+	0.69	0.64	0.02	0.16	24
	$dal_fre:C;\delta = 0.75$	16+	60+	0.74	-0.69	0.03	0.17	44

TABLE 2.4.2. Welfare in Coordination Games

paid: paid periods

errr: welfare difference divided by empirical welfare

RP/FS: refined perfection and Fehr-Schmidt

for IPD payoffs are normalized to be zero for mutual defection and one for mutual cooperation

bold face: an anomaly defined as an error of more than 10% in absolute value
all these games have multiple subgame perfect equilibria and the welfare best was used

Selfish Games

(warning, not out of sample)

game	treatment	period(s)	paid	theory	data	errr	RP/FS errr	N
pub	no pun	10	10	\$1.51	\$1.51	0.00	−0.01	24
PD	PD1	9	30	\$0.19	\$0.19	0.00	−0.10	164
	PD2	9 – 10	30	\$0.23	\$0.23	0.00	−0.04	222
mar		10	1	\$0.44	\$0.44	0.00	0.00	160
IPD	<i>are_cou_ran</i> ; $\delta = 0.18$	16+	60+	0.17	0.19	−0.04	−0.36	66
	<i>dal_fre:A</i> ; $\delta = 0.50$	16+	60+	0.27	0.14	0.04	−0.04	197

TABLE 2.4.3. Welfare in Selfish Games

paid: paid periods

errr: welfare difference divided by empirical welfare

RP: refined perfection

FS: Fehr-Schmidt

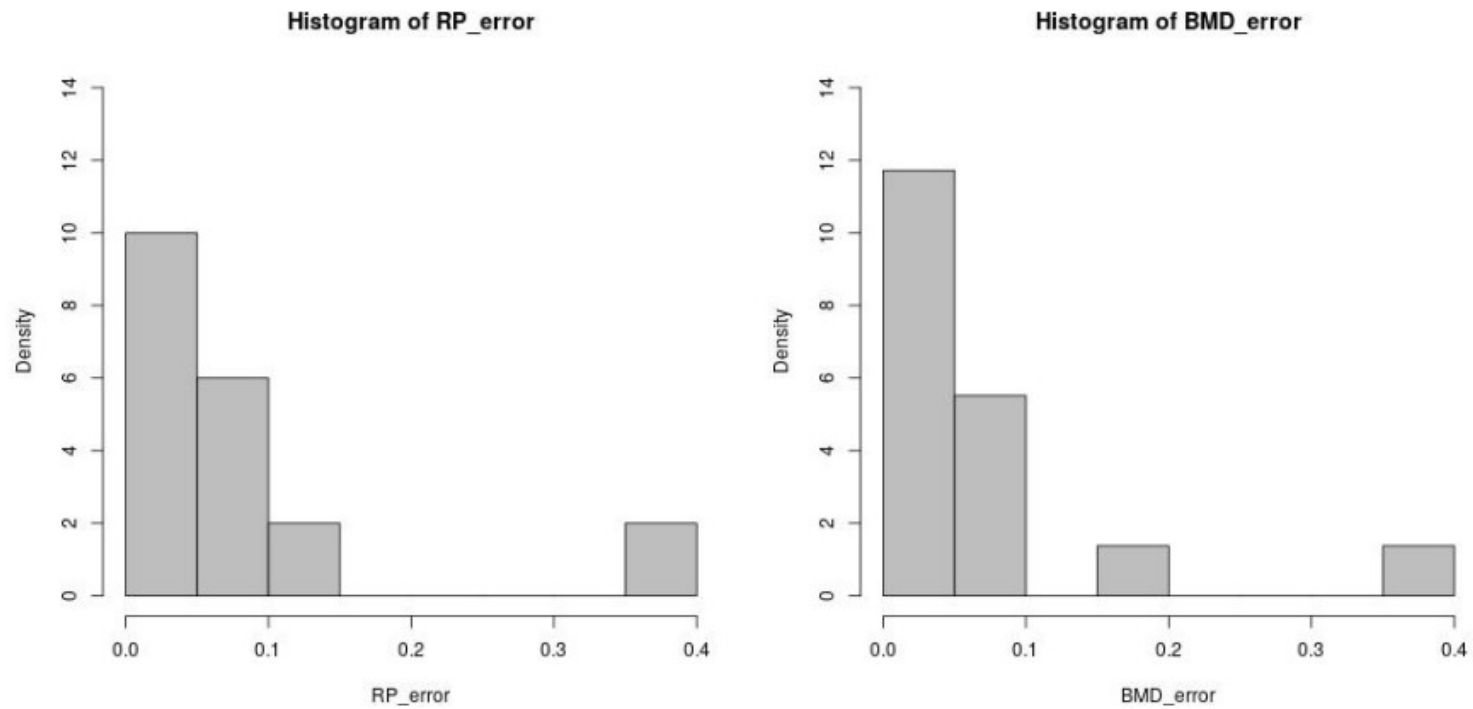
N: number of participants

for IPD payoffs are normalized to be zero for mutual defection and one for mutual cooperation

Games Where RP Is Thought To Do Well

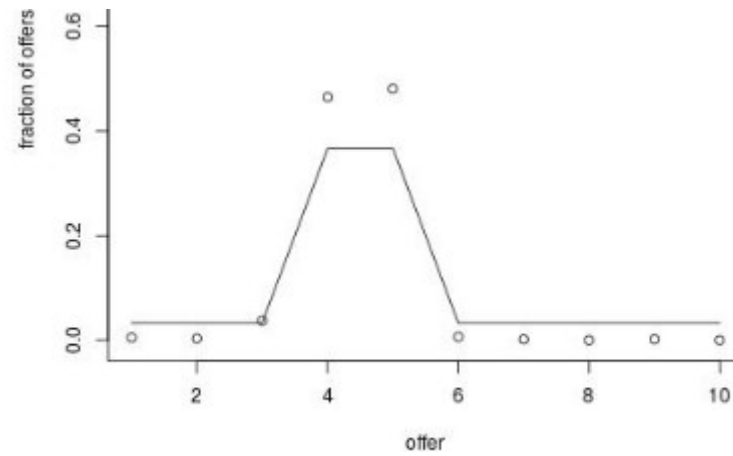
game	abs RP errr
dal_bo PD1: $H = 1$	0.10
dal_bo PD2: $H = 1$	0.04
dal_bo PD1: $H = 2$	0.09
dal_bo PD2: $H = 2$	0.08
dal_bo PD1: $H = 4$	0.02
dal_bo PD2: $H = 4$	0.14
dal_fre:A; $\delta = 0.5$	0.04
are_cou_ran	0.36
pub no pun	0.01
market auction	0.00

How Big is Big?

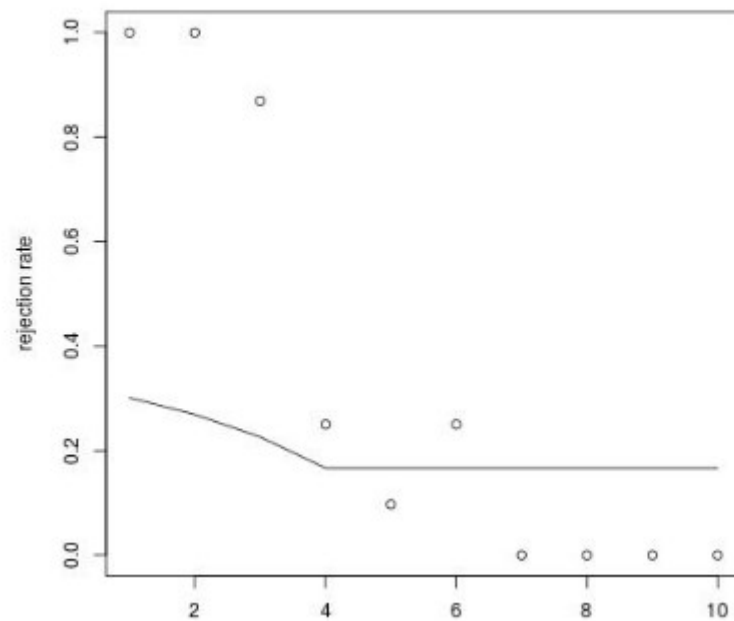


right stochastically dominates left

Ultimatum Bargaining



Rejection Rates



Gift Exchange

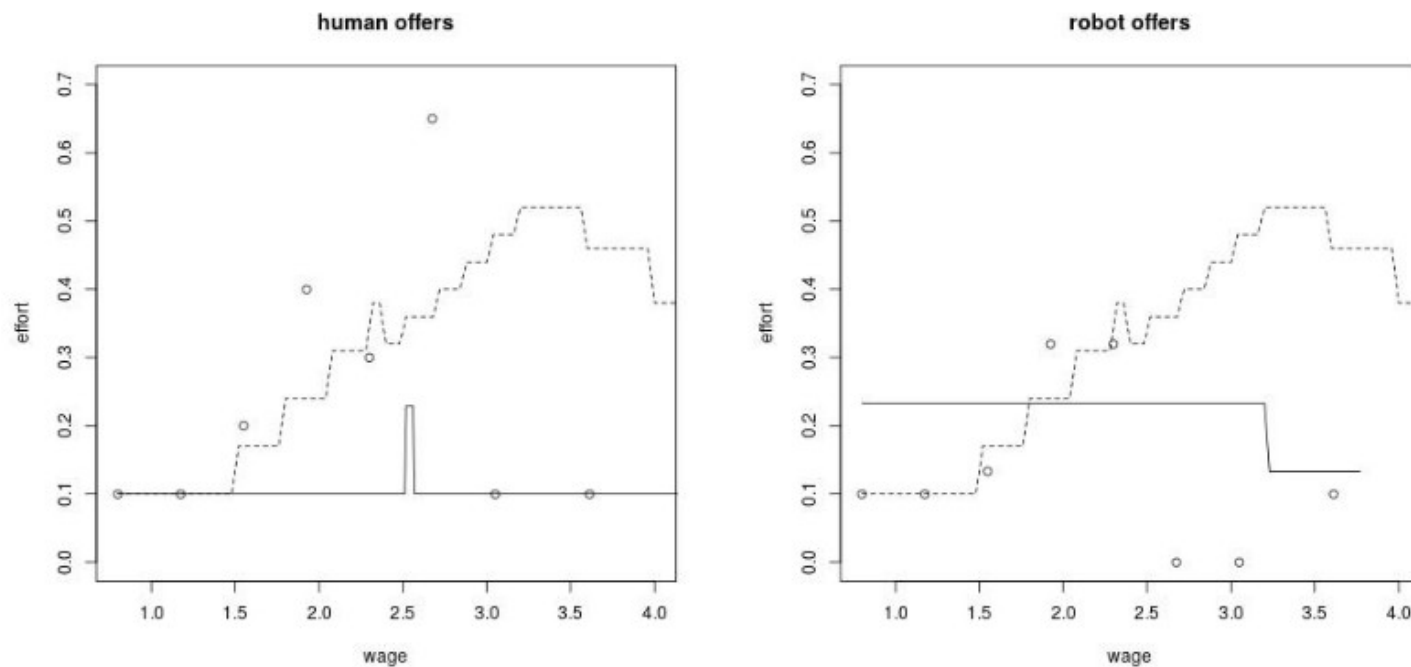


FIGURE 3.4.2. Effort vs. Wage; left panel human, right panel robot; dots data, solid line theory; dashed line Fehr-Schmidt

Fairness and Reciprocity in Ultimatum Bargaining

	obs	nobs
frequency of offers	4%	30%
rejection rate	87%	14%

Heroic Model: Fourteen Primitive Societies

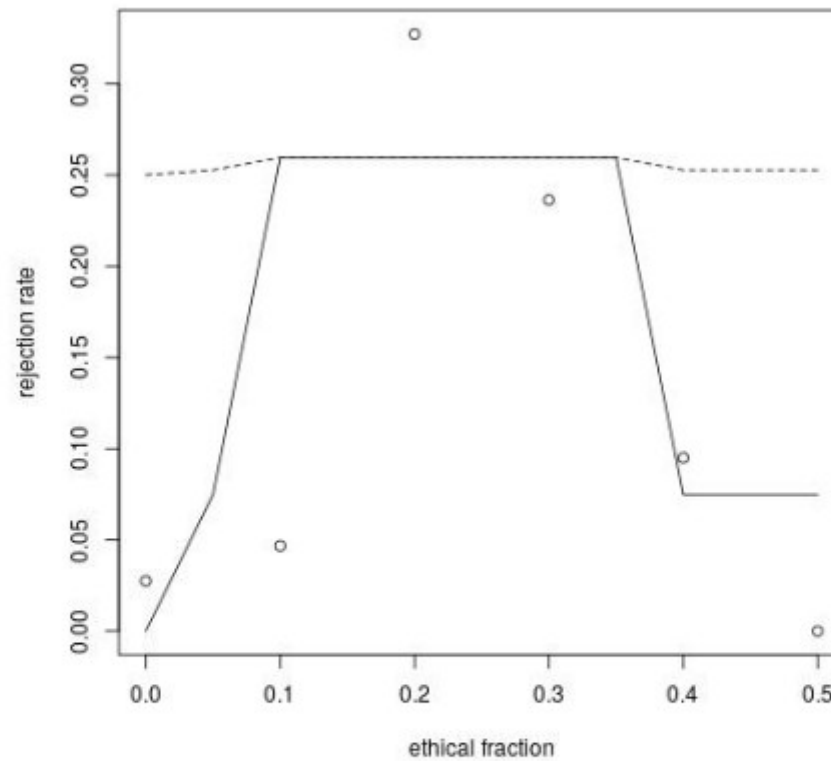
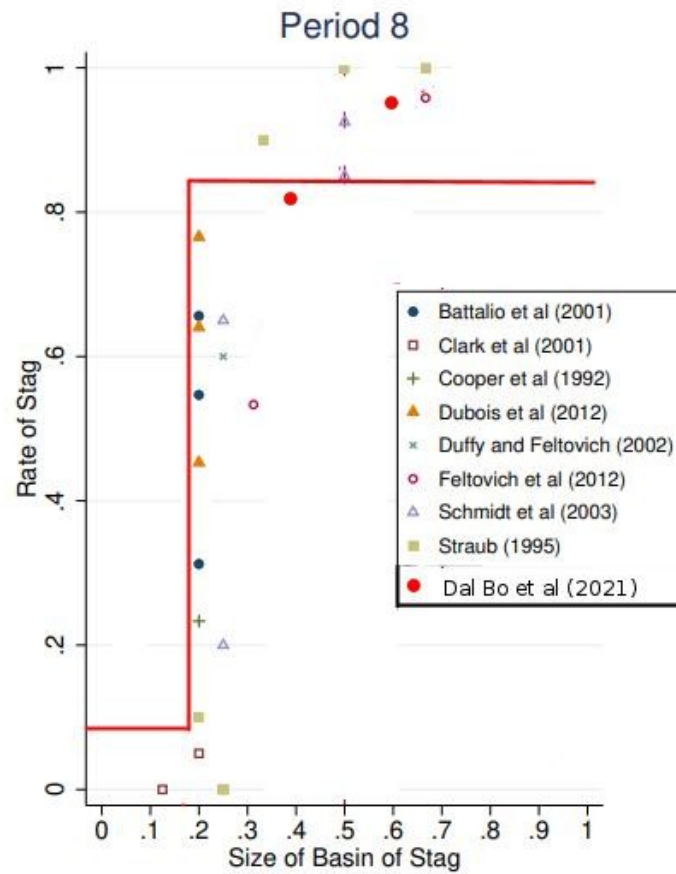
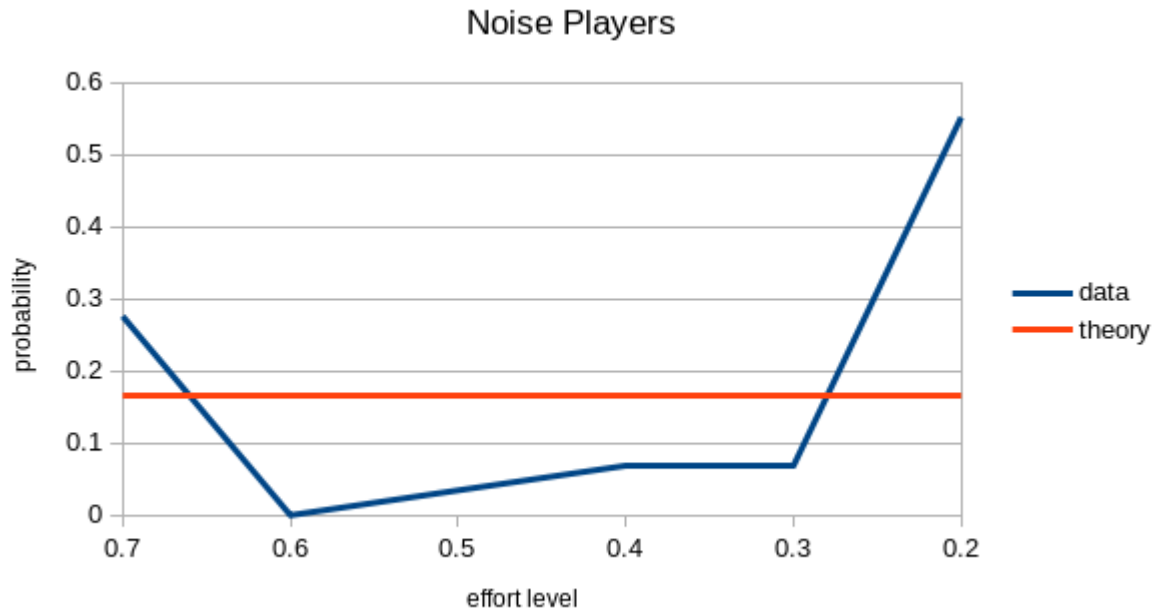


FIGURE 4.11.2. Rejection Rates
dots: averages within each category

Stag Hunt



Dissidents in the Minimum Game



The Indefinitely Repeated PD

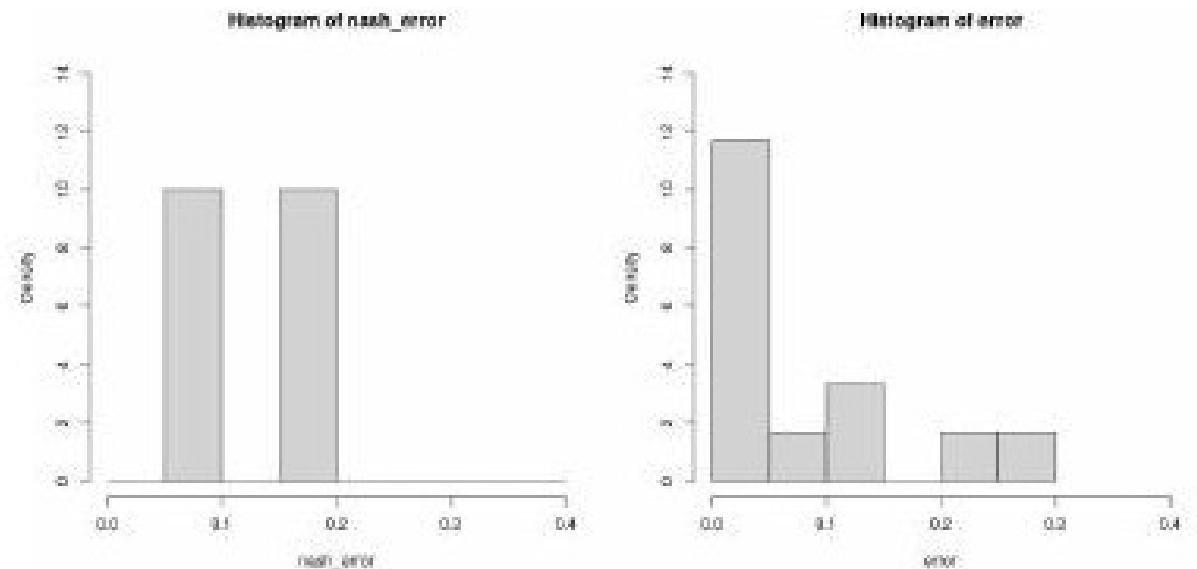


Figure 4.1: Error Histograms

left are the Nash cooperation rate errors from Table 4.2

right are the welfare errors in this paper from Table 4.1

left is only for games where defection is the only equilibrium

Risk Dominance and the Basin

limit choices to mutual-defection and grim-trigger

(out of 32 memory one strategies)

the basin: maximum fraction defecting for which grim-trigger is optimal

$$\beta = \max \left\{ 0, \frac{(1 - \delta)(u^1(DC) - 1) - \delta}{(1 - \delta)(u^1(DC) + u^1(CD) - 1) - \delta} \right\}.$$

or zero if negative

prediction: cooperate if greater than 0.5; defect if less than 0.5

Critical Discount Factor

Blonski, Ockenfels and Spagnolo (2011)

for any given payoff matrix there is a critical discount factor above which there will be cooperation

additional less persuasive axioms concerning additive separability and so forth; they compute

$$\Delta = \delta - \frac{u^1(DC) - u^1(CD) - 1}{u^1(DC) - u^1(CD)}$$

also the discount factor that makes players indifferent between grim-trigger and always-defect when the population is 50 – 50 between the two, so this has a similar flavor to the basin.

prediction: cooperate if positive; defect if negative

Reinforcement Learning

Fudenberg and Rehbindler (2024)

six parameter reinforcement learning model using Δ as an explanatory variable

try to predict play in every round of every match

I generated artificial data using the parameters from their Table 4 using Monte Carlo simulation with 1000 trials to get predictions of welfare in long-term play $\hat{\omega}$

Comparison

payoff matrix	δ	$\hat{\omega}$	β	Δ	$\hat{\hat{\omega}}$	data $\bar{\omega}$
<i>dal_fre:B</i>	0.50	0.14	0.28	-0.11	0.22	0.15
<i>dre_et_al_08:A</i>	0.75	0.17	0.33	-0.05	0.30	0.14
<i>are_cou_ran</i>	0.13	0.17	0.00	-0.27	0.24	0.19
<i>bru_kam</i>	0.80	0.21	0.77	0.13	0.58	0.35
<i>dal_fre:A</i>	0.50	0.27	0.00	-0.32	0.12	0.14
<i>dal_fre:A</i>	0.75	0.27	0.19	-0.07	0.31	0.25
<i>dal_fre:C</i>	0.50	0.67	0.61	0.10	0.47	0.39
<i>dre_et_al_08:B</i>	0.75	0.67	0.67	0.08	0.50	0.46
<i>kag_sch</i>	0.75	0.67	0.80	0.15	0.60	0.58
<i>she_tar_sai</i>	0.75	0.67	0.89	0.19	0.67	0.69
<i>dal_fre:B</i>	0.75	0.69	0.73	0.14	0.57	0.64
<i>dal_fre:C</i>	0.75	0.74	0.84	0.35	0.79	0.69

Table 1: Models and Values

italics: $\hat{\omega}$ equal to minimum with noise

normalized not relative welfare, one period memory

β, Δ **50%**, $\hat{\omega}, \hat{\hat{\omega}}$ **15%**

Model Fit

	parameters	mean abs	R^2	intercept (se)	slope (se)
$\hat{\omega}$	0	0.09	0.75	0.00	1.00
β	2	0.09	0.77	0.09 (0.06)	0.58 (0.10)
Δ	2	0.09	0.75	0.36 (0.03)	0.98 (0.18)
$\hat{\hat{\omega}}$	6	0.10	0.75	0.00	1.00

TABLE 6.5.2. Predictors of Welfare $\bar{\omega}$

parameters: number of parameters estimated from repeated game data

mean abs: mean absolute error

R^2 not adjusted

Bifurcation

payoff matrix	δ	bifurcation range for ϕ
<i>dal_fre:B</i>	0.50	(0.1, 0.2)
<i>dre_et_al_08:A</i>	0.75	(0.1, 0.2)
<i>are_cou_ran</i>	0.13	none
<i>bru_kam</i>	0.80	(0.332, 0.333)
<i>dal_fre:A</i>	0.50	none
<i>dal_fre:A</i>	0.75	(0.0, 0.1)
<i>dal_fre:C</i>	0.50	(0.5, 0.6)
<i>dre_et_al_08:B</i>	0.75	(0.339, 0.340)
<i>kag_sch</i>	0.75	(0.338, 0.339)
<i>she_tar_sai</i>	0.75	(0.338, 0.339)
<i>dal_fre:B</i>	0.75	(0.35, 0.4)
<i>dal_fre:C</i>	0.75	(0.6, 0.7)

TABLE 6.4.1. Bifurcation Ranges
 for ϕ at the top mutual-defection is the only equilibrium
 for ϕ at the bottom the best equilibrium is cooperative
 none indicates for all ϕ mutual-defection is the only equilibrium

Selfish Games

Definition: A game is a *selfish game* if it satisfies two properties:

1. Uniqueness: there is a unique strong subgame perfect equilibrium outcome, σ^*
2. Incentives do not matter: There is an optimal behavioral mechanism in which selfish players follow the same strategy as in the unique strong subgame perfect equilibrium, that is $\sigma_S = \sigma^*$

Subgame Perfection in Selfish Games

game	treatment	errr	unmod	RP errr	participants
dictator		0.07	0.00	1.00	many
pub	no pun	0.00	0.00	−0.01	24
PD	PD1	0.00	0.16	−0.10	164
	PD2	0.00	0.09	−0.04	222
mar		0.00	− 0.20	0.00	160
IPD	<i>are_cou_ran</i> ; $\delta = 0.18$	−0.04	−0.04	−0.36	66
	<i>dal_fre:A</i> ; $\delta = 0.50$	0.04	0.04	−0.04	197

TABLE 8.1.1. Unmodified Welfare in Selfish Games

errr: welfare difference divided by empirical welfare

unmod: errr for unmodified theory

RP: err for refined perfection

for IPD payoffs are normalized

bold face: an anomaly defined as an error of more than 10% in absolute value

Real Anomalies

centipede: failure of backward induction?

- the entire problem is people dropping out too soon, not staying to long
- lack of proper baseline obscures this

best-shot: triumph of backwards induction?

- all theories do an abysmal job of explaining what is going on
- 61% play subgame perfect? If 39% error is success, I'd hate to see failure

these get little attention because people wrongly think they understand them

Thank You