

Reliability and Replicability of Implicit and Explicit Reinforcement Learning Paradigms in People With Psychotic Disorders

Danielle N. Pratt¹, Deanna M. Barch^{2,*}, Cameron S. Carter³, James M. Gold⁴, John D. Ragland³, Steven M. Silverstein⁵, and Angus W. MacDonald III.^{*,1}

¹Department of Psychology, University of Minnesota, Minneapolis, MN; ²Department of Psychology, Washington University, St. Louis, MO; ³Department of Psychiatry, University of California at Davis, Davis, CA; ⁴Maryland Psychiatric Research Center, University of Maryland School of Medicine, Baltimore, MD; ⁵Department of Psychiatry, University of Rochester, Rochester, NY

*To whom correspondence should be addressed; University of Minnesota, N218 Elliott Hall, 75 E. River Rd. Minneapolis, MN 55455, USA; tel: 612-624-3813, fax: 612-626-2079, e-mail: angus@umn.edu

Background: Motivational deficits in people with psychosis may be a result of impairments in reinforcement learning (RL). Therefore, behavioral paradigms that can accurately measure these impairments and their change over time are essential. **Methods:** We examined the reliability and replicability of 2 RL paradigms (1 implicit and 1 explicit, each with positive and negative reinforcement components) given at 2 time points to healthy controls ($n = 75$), and people with bipolar disorder ($n = 62$), schizoaffective disorder ($n = 60$), and schizophrenia ($n = 68$). **Results:** Internal consistency was acceptable (mean $\alpha = 0.78 \pm 0.15$), but test-retest reliability was fair to low (mean intraclass correlation = 0.33 ± 0.25) for both implicit and explicit RL. There were no clear effects of practice for these tasks. Largely, performance on these tasks shows intact implicit and impaired explicit RL in psychosis. Symptom presentation did not relate to performance in any robust way. **Conclusions:** Our findings replicate previous literature showing spared implicit RL and impaired explicit reinforcement in psychosis. This suggests typical basal ganglia dopamine release, but atypical recruitment of the orbitofrontal and dorsolateral prefrontal cortices. However, we found that these tasks have only fair to low test-retest reliability and thus may not be useful for assessing change over time in clinical trials.

Key words: practice effects/positive and negative reinforcement/schizophrenia

Introduction

Approximately, 75% of individuals with schizophrenia have motivational deficits.¹ These deficits are linked to social, occupational, and other functional impairments and, therefore, represent an important treatment target.^{1–3} Targeting symptoms requires a means to measure their

improvement, a role that is increasingly filled through performance-based tasks. Such clinically relevant tasks should be well-tolerated, sensitive to psychosis-related impairments, and reliable enough to track change with minimal practice effects.⁴ The current study examined these properties in 2 reward-based learning tasks, a critical construct for measuring motivation in people with psychosis. These tasks were identified as potential outcome measures in pharmaceutical clinical trials based on surveys given to experts that identified important cognitive domains⁵ and promising paradigms.^{6,7}

While motivational deficits were formerly attributed to an inability to experience pleasure, recent evidence does not support this hypothesis.⁸ People with schizophrenia report similar in-the-moment pleasure to healthy individuals and only showed reduced levels of anticipated pleasure.^{9,10} Therefore, a more compelling explanation for motivational deficits is the reward processing system impairments in people with psychosis.^{8,10} One aspect of the reward processing system is reinforcement learning (RL) or determining how to maximize a reward by exploring the environment and adapting the performance to exploit rewards and avoid losses.¹¹

A large literature has parcellated key components of RL. RL may be implicit (outside of conscious awareness) or explicit (overt representations about the potential reward associations).¹² Reinforcement may also be positive (learning which actions lead to reward) or negative (learning which actions avoid an undesirable outcome). Differing brain mechanisms underlie each of these systems. For implicit RL, the changes in dopamine (DA) release based on expected and unexpected rewards modify activity in the ventral and dorsal regions of the basal ganglia to support maximally adaptive responses to current stimuli.^{13,14} Unexpected rewards induce DA

firing and reinforce a behavior while nonoccurrences of rewards/losses reduce firing and inhibit that behavior. Implicit learning occurs slowly, without conscious awareness, until the DA neurons fire to the cues themselves.^{13–17} Explicit RL mechanisms appear to recruit the orbito-frontal cortex (OFC) and dorsolateral prefrontal cortex (DLPFC) for faster, more flexible top-down control of choices about reward and punishment.^{13,18–20}

The literature is mixed regarding RL deficits observed in people with psychotic disorders. Implicit RL appears relatively intact in schizophrenia^{8,10,21–23} (see exception),²⁴ though it may still be accompanied by abnormal basal ganglia function.^{25,26} Impairments in positive explicit RL are common, with potentially preserved negative explicit RL^{19,22,24,27–30} (see^{10,24,31–33} for deficits in both and³⁴ for deficits in neither). The deficit in positive explicit RL may be exaggerated in individuals with higher negative symptoms, in particular, amotivation/anhedonia.^{8,19}

For these reasons, reward processing and RL specifically are important treatment targets and, therefore, must be measured with precision over time.⁴ The goal of the Cognitive Neuroscience Test Reliability and Clinical applications for Serious mental illness (CNTRaCS) consortium is to psychometrically optimize tasks of disorder-relevant cognitive constructs, such as RL, for use in clinical trials.³⁵ Two RL tasks have been particularly useful to the field in recent years. One, developed by Pizzagalli et al,³⁶ is here called the Implicit Probabilistic Incentive Learning Tasks (IPILT). The second, developed by the Pessiglione et al,³⁷ is here called the Explicit Probabilistic Incentive Learning Tasks (EPILT). CNTRaCS adapted these tasks with new stimuli and conditions to better address relevant questions.¹² Our group previously found that bias toward reward or away from punishment did not differ by diagnosis for the IPILT, whereas people with schizophrenia were impaired on the EPILT.¹² The applicability for clinical trials has yet to be determined. Therefore, the goals of this study were to (1) establish task reliability, (2) examine practice effects, (3) replicate the findings in Barch et al,¹² and (4) assess relationships to symptom severity.

Methods

Participants

Seventy-five healthy controls (HC), 62 people with bipolar disorder (BP) with psychotic features, 85 people with schizoaffective disorder (SczA; 25 unmedicated), and 91 people with schizophrenia (Scz; 23 unmedicated) were recruited. Participants gave written informed consent. For the current analyses, only medicated patients were included. Other reasons for exclusion ($n = 5$) are described in the [supplemental methods](#). Data were collected nearly equally across all 5 sites of the CNTRaCS consortium (Washington University in St. Louis, University of California-Davis, Maryland Psychiatric Research Center

at the University of Maryland School of Medicine, Rutgers University, and University of Minnesota-Twin Cities), and each site's institutional review board approved this study. See [supplemental material](#) for inclusion and exclusion criteria. Groups were similar on age and parental socioeconomic status but differed on sex ratios ([table 1](#)).

Tasks and Procedures

Participants completed 3 visits. The first consisted of an IQ screen,³⁹ diagnostic interview,⁴⁰ and clinical and functional⁴¹ scales. Clinical symptoms were rated using the Brief Psychiatric Rating Scale (BPRS),⁴² Young Mania Rating Scale (YMRS),⁴³ Bipolar Depression Rating Scale (BDRS),⁴⁴ and Clinical Assessment Interview for Negative Symptoms (CAINS).⁴⁵ All raters achieved agreement (intraclass correlation [ICC] ≥ 0.80) on “gold standard” ratings with ongoing drift prevention interviews every 2–4 weeks.

Two sets of cognitive testing sessions were completed approximately 1 month apart. The cognitive tasks in this article are summarized in [supplemental table S1](#) and explained later. Other testing procedures are in the [supplemental material](#).

Implicit Probabilistic Incentive Learning Tasks. The IPILT ([figure 1a](#)) was modified^{36,46} to include both a positive (IPILT-P) and negative (IPILT-N) reinforcement version.¹² Participants made perceptual discriminations between 2 variants of a briefly shown (100 ms) line-drawn stimulus. On the IPILT-P, ~40% of correct responses received the feedback “Correct! You Win!” and participants gained \$0.05. On the IPILT-N, participants started with \$3.60 and lost \$0.05 on ~40% of incorrect trials being told “Sorry. You Lose.” One of the 2 stimulus variants was always associated with 3 times more feedback (RICH) than the other (LEAN). The IPILT-P and IPILT-N each had 3 blocks of 60 trials. The dependent measures included response bias, $\log b$, and discriminability, $\log d$ or d' (equations in [supplemental material](#)).

Explicit Probabilistic Incentive Learning Tasks. The EPILT ([figure 3a](#)) was adapted³⁷ to assess explicit learning from gain and avoiding loss incentives.¹² There were 2 phases: training and transfer. During training, participants were asked to learn value discriminations for 4 pairs of images over 160 trials. Two of the pairs were Gain conditions, where the optimal choice was associated with a gain of \$0.05 and the word “WIN!” Nonoptimal choice resulted in no gain of money and the feedback “Not a winner. Try again!” For one of the Gain pairs, the optimal response was reinforced 90% of the time and the other pair's optimal response was reinforced 80% of the time. The other 2 pairs of stimuli were Avoid Loss conditions, with a not lose/lose response pattern. These used the same reward probabilities, where the optimal choice resulted in no loss of money and the feedback “Keep

Table 1. Demographics and Clinical Characteristics

	HC (<i>n</i> = 75)		BP (<i>n</i> = 60)		SczA (<i>n</i> = 57)		Scz (<i>n</i> = 68)		Group Differences
	<i>M</i>	<i>SD</i>	<i>M</i>	<i>SD</i>	<i>M</i>	<i>SD</i>	<i>M</i>	<i>SD</i>	
Age	37.4	11.3	38.3	10.7	38.4	11.2	37.0	10.1	NS
Gender (% female)	45.3%		64.5%		45.0%		39.7%		BP > HC = Scz = SczA
Race (% Nonwhite)	48.0%		25.8%		65.0%		54.4%		SczA > Scz = HC > BP
Personal education	14.7	2.3	14.7	2.6	13.2	2.4	12.8	2.1	HC = BP > Scz = SczA
Parental SES ^a	13.9	2.7	14.2	2.5	13.8	2.8	13.1	3.4	NS
WTAR standard score	109.0	13.2	105.3	14.8	96.9	17.1	95.7	17.0	HC = BP > Scz = SczA
Time between test sessions	28.8	4.9	31.9	7.1	29.9	4.2	30.2	6.2	BP > HC
BPRS Positive	—	—	4.6	2.0	9.0	4.6	6.8	3.6	SczA > Scz > BP
BPRS Negative	—	—	5.9	1.9	7.6	2.8	7.7	2.5	Scz = SczA > BP
BPRS Disorganization	—	—	4.8	1.1	5.8	2.3	5.6	1.9	Scz = SczA > BP
BPRS Depression	—	—	9.6	3.9	10.4	4.5	7.7	3.0	SczA = BP > Scz
BPRS Mania	—	—	7.5	3.0	8.0	3.7	6.3	2.7	SczA = BP > Scz
YMRS	—	—	9.8	7.6	12.2	6.9	8.7	5.5	SczA > Scz; SczA = BP; Scz = BP
BDRS	—	—	5.5	5.3	7.1	7.0	4.0	4.6	SczA > Scz; SczA = BP; Scz = BP
CAINS Motivation & Pleasure	—	—	8.2	5.5	11.8	6.9	12.0	6.3	Scz = SczA > BP
CAINS Expression	—	—	0.8	1.7	3.3	3.1	4.1	2.8	Scz = SczA > BP
SLOF Self-Report	—	—	4.3	0.5	4.2	0.5	4.2	0.5	NS
SLOF Informant	—	—	4.4	0.4	4.1	0.4	4.2	0.5	NS
Typical antipsychotic	—	—	6.5%		16.7%		22.1%		Scz = SczA > BP
Atypical antipsychotic	—	—	46.8%		95.0%		88.2%		Scz = SczA > BP
Both typical and atypical	—	—	11.7%		10.3%		3.2%		NS

Note: HC, healthy controls; BP, bipolar disorder; SczA, schizoaffective disorder; Scz, schizophrenia; SES, Socioeconomic status; WTAR, Wechsler's Test of Adult Reading; BPRS, Brief Psychotic Rating Scale; YMRS, Young Mania Rating Scale; BDRS, Bipolar Depression Rating Scale; CAINS, Clinical Assessment Interview for Negative Symptoms; SLOF, Specific Levels of Functioning scale. Racial composition: 51.7% White, 38.9% Black, 2.3% Native American/Alaskan, 0.4% Pacific Islander, 5.3% Asian, and 5.7% other/unknown.

^aEstimated from the Hollingshead Index.³⁸

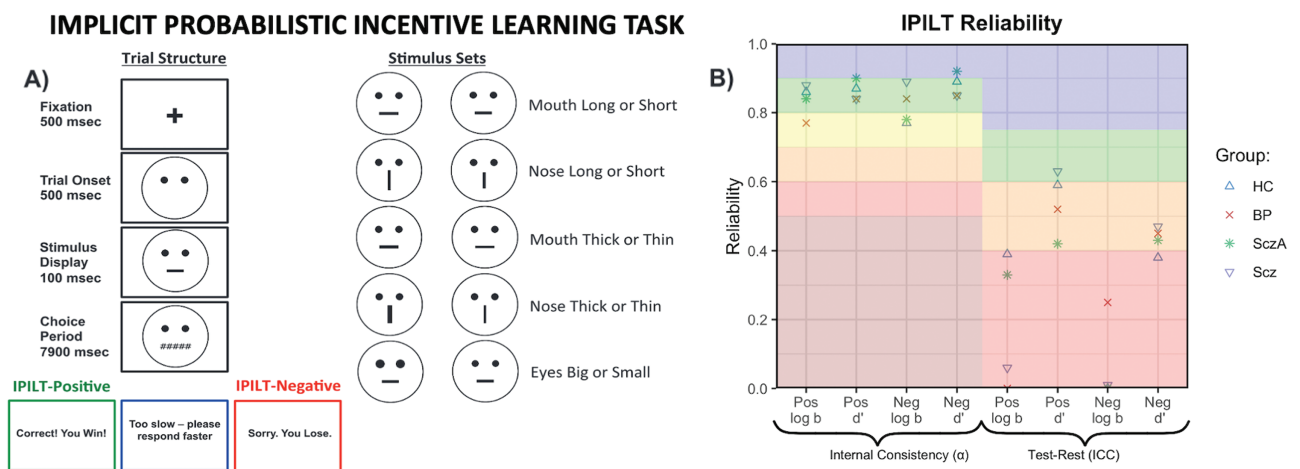


Fig. 1. The Implicit Probabilistic Incentive Learning Tasks (IPILT). (A) The trial structure schematic and stimulus sets for IPILT, adapted from Barch et al.¹² Only one stimulus set is presented per session, counterbalanced by participant and session. (B) Internal consistency (session 1 only) and test-retest reliability for the IPILT split by group for positive and negative bias (log *b*) and discriminability (*d'*).

your money!" and the nonoptimal choice resulted in a loss of \$0.05 and the feedback "LOSE!" The dependent measure for this phase was percent accuracy.

For the transfer phase, the original 4 training pairs were each presented 4 times, alongside 56 novel pairings, totaling 72 trials. Novel pairings included only trained images. Participants chose the image they

thought was "best," without feedback. The accuracy of selecting the more rewarded image in a pairing was the primary dependent measure. As described further in the [supplemental material](#), the pairings of interest were: (1) Frequent Winner vs Frequent Loser (FW vs FL), (2) Frequent Winner vs Infrequent Winner (FW vs IW), (3) Frequent Winner vs Frequent Loss Avoider

(FW vs FLA), and (4) Frequent Loss Avoider vs Infrequent Winner (FLA vs IW).

Data Analysis

Reliability. The internal consistency at each testing session for all dependent measures for both the IPILT and the EPILT was calculated using Cronbach's α ; data for session 2 are presented in the [supplemental results](#). Test-retest reliability for each task was assessed using ICCs. A 2-way random effects model assessing the degree of agreement between testing sessions was utilized (McGraw and Wong formula ICC(A,1)).^{47,48} Reliability for the IPILT was reassessed excluding a stimulus set that appeared more difficult and, therefore, noncomparable to others.

Practice Effects. Following previous CNTRaCS analyses,⁴⁹ multilevel modeling was utilized for each task's primary dependent measure with Session as a within-subjects predictor and Diagnosis, mean-centered Age, and Sex as between-subjects predictors as well as the higher-order interactions. The models had fixed slopes and random intercepts and used an unstructured covariance structure. Nonsignificant interactions were dropped.

Replication. Full analysis and data processing details can be found in Barch et al.¹² Results from session 1 are reported in the primary text with session 2 included in the [supplemental material](#). For the IPILT, greater response bias (log b) indicated a higher propensity to choose the more optimal stimulus and larger discriminability (d') indicated a greater ability to distinguish between stimuli. The IPILT-P and IPILT-N were separately analyzed with repeated-measures ANOVAs. Block was a within-subject factor and Stimulus Set and Diagnostic Group were between-subject factors.

For the training phase of the EPILT, a repeated-measures ANOVA with Block, Valence (Gain vs Avoid Loss), and Probability (90/10 vs 80/20) as within-subject factors and Stimulus Set and Diagnostic Group as between-subject factors was completed. The transfer phase was analyzed with Pairing as a within-subject factor and Stimulus Set and Diagnostic Group as between-subject factors in a repeated-measures ANOVA.

Clinical Relationships. To assess the relationship of clinical symptoms with RL, we implemented canonical correlations with 10-fold cross-validation to assess robustness. For the IPILT, we used the average bias across blocks and the change in bias from block 1 to block 3. For the EPILT, we used the average accuracy for the Gain and Avoid Loss training conditions and the accuracy of 3 transfer pairings (FWvsIW, FWvsFLA, and FLAvsIW). The clinical symptoms assessed were the BPRS scores for positive, negative, depression, mania,

and disorganized symptoms, YMRS total score, BDRS total score, CAINS Motivation and Pleasure, and CAINS Expression. Results for session 1 are reported here and those for session 2 are in the [supplemental material](#).

Results

IPILT

Reliability. The IPILT-P and IPILT-N each had very good internal consistency for both response bias and discriminability across groups at session 1 ([figure 1b](#)). Internal consistency for session 2 was similar ([supplemental table S2](#)). Test-retest reliability for the IPILT-P bias, IPILT-N bias, and IPILT-N discriminability was low, with IPILT-P discriminability being moderate.

Practice Effects. For the IPILT-P response bias, there was a Group $\times$ Session interaction; SczA improved across testing sessions compared with HC ($b = 0.384$, $t(248) = 2.14$, $P = .033$). For the IPILT-N, heightened bias away from punishment was seen in SczA in session 1, but SczA bias was similar to other groups in session 2. As such, compared with HC, there was a Group $\times$ Session interaction for SczA ($b = -0.447$, $t(242) = -2.10$, $P = .037$). Effects with Age and Sex are in [supplemental table S5](#).

Replication. For the IPILT-P ([figure 2a](#)), all Groups showed a positive response bias toward reward (model intercept: $F(1,240) = 78.6$, $P < .0001$, $\eta^2_p = 0.247$), as expected, with no significant Group differences ($F(3,240) = 1.3$, $P = .267$, $\eta^2_p = 0.016$), largely replicating our previous work.¹² We also observed a significant main effect of Block ($F(2,239) = 8.1$, $P = .0004$, $\eta^2_p = 0.053$), suggesting an increasing bias toward the rewarded response over blocks across groups, which was not seen previously.

For the IPILT-N ([figure 2b](#)), participants had a response bias away from punishment (model intercept: $F(1,232) = 13.5$, $P = .0003$, $\eta^2_p = 0.055$), as predicted.¹² A significant effect of Block was again observed ($F(2,231) = 32.5$, $P < .0001$, $\eta^2_p = 0.220$), where bias increased over blocks. A main effect of Group ($F(3,232) = 5.0$, $P = .002$, $\eta^2_p = 0.060$) was also observed, with SczA having significantly higher bias than all other groups, which was not reported previously.

While examining replication, one Stimulus Set (Eyes Big or Small) appeared to be more difficult than others while examining d' (IPILT-P: $F(4,224) = 2.6$, $P = .037$; IPILT-N: $F(4,232) = 10.7$, $P < .001$). Test-retest reliability was reexamined without this set and still found to be generally poor (see [supplemental results](#) and [figure S1](#)).

Relationship to Symptoms. Clinical symptoms did not relate to IPILT-P or IPILT-N average bias or change in bias (first canonical variate = 0.30, $P = .66$; cross-validated first variate = 0.07, $P = .99$).

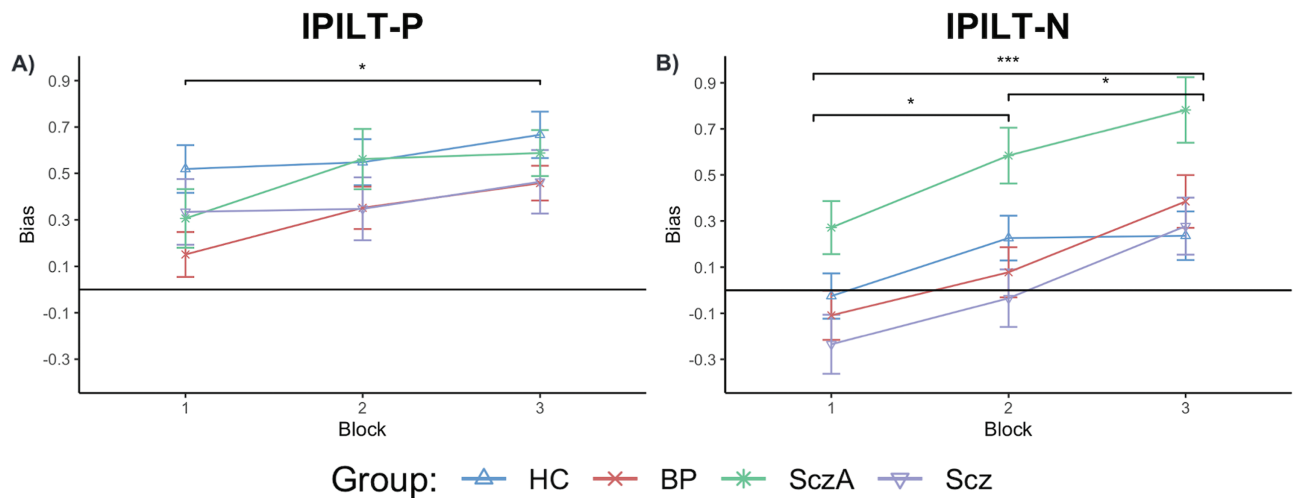


Fig. 2. Bias across blocks for (A) Implicit Probabilistic Incentive Learning Tasks-Positive (IPILT-P) and (B) IPILT-Negative (IPILT-N) split by group. * $P < .001$, *** $P < .00001$.

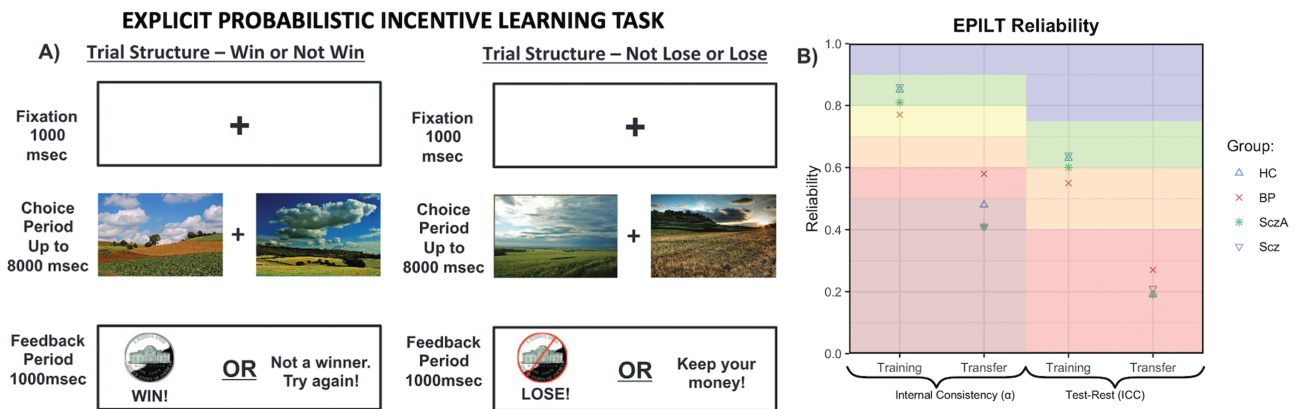


Fig. 3. The Explicit Probabilistic Incentive Learning Tasks (EPILT). (A) Trial structure for the EPILT. Participants were trained on 4 image pairs. For 2 pairs, they received the feedback on the left with either 90% or 80% of the feedback reinforcing the optimal image. For the other 2 pairs, they received on the right with either 90% or 80% of the feedback reinforcing the optimal image. Image adapted from Barch et al.¹² Only one stimulus set is presented per session, counterbalanced by participant and session. (B) Internal consistency (session 1 only) and test-retest reliability for the EPILT split by group for Training and Transfer accuracy.

EPILT

Reliability. EPILT-Training had very good internal consistency across groups at session 1 (figure 3b). However, internal consistency for the EPILT-Transfer phase was poor to fair. These results were repeated in session 2 (supplemental table S2). Test-retest reliability for the EPILT-Training phase was fair. For the EPILT-Transfer phase, test-retest reliability was poor.

Practice Effects. During training, there was a Group $\times$ Session interaction in the Gain conditions. For the 80% Gain condition, accuracy for HC and Scz went down from session 1 to session 2 (supplemental table S6); Scz did not differ from HC ($b = 0.003$, $t(225) = 0.084$, $P = .933$), whereas accuracy increased for BP ($b = .068$, $t(221) = 2.12$, $P = .035$) and SczA ($b = 0.088$, $t(226) = 2.63$, $P = .009$). For the 90% Gain condition, compared with HC (who

decreased in accuracy over time), there was a Group $\times$ Session interaction for BP who increased in accuracy at Session 2 ($b = 0.072$, $t(222) = 2.23$, $P = .027$). There was a main effect of Session for the 80% Avoid Loss ($b = 0.025$, $t(225) = 2.75$, $P = .007$) and 90% Avoid Loss ($b = 0.029$, $t(231) = 2.96$, $P = .003$) conditions. Session did not interact with Group, Age, or Sex for these conditions.

In the EPILT-Transfer (supplemental table S7), there were Group $\times$ Session interactions for FLAvsIW and FWvsIW where BP improved in accuracy compared with HC who performed worse at session 2 ($b = 0.114$, $t(232) = 2.30$, $P = .022$ and $b = 0.121$, $t(224) = 2.83$, $P = .005$, respectively). Interactions with Sex are in supplemental table S7.

Replication. In the EPILT-Training (figure 4a), we observed main effects of Block ($F(3,222) = 112.1$, $P <$

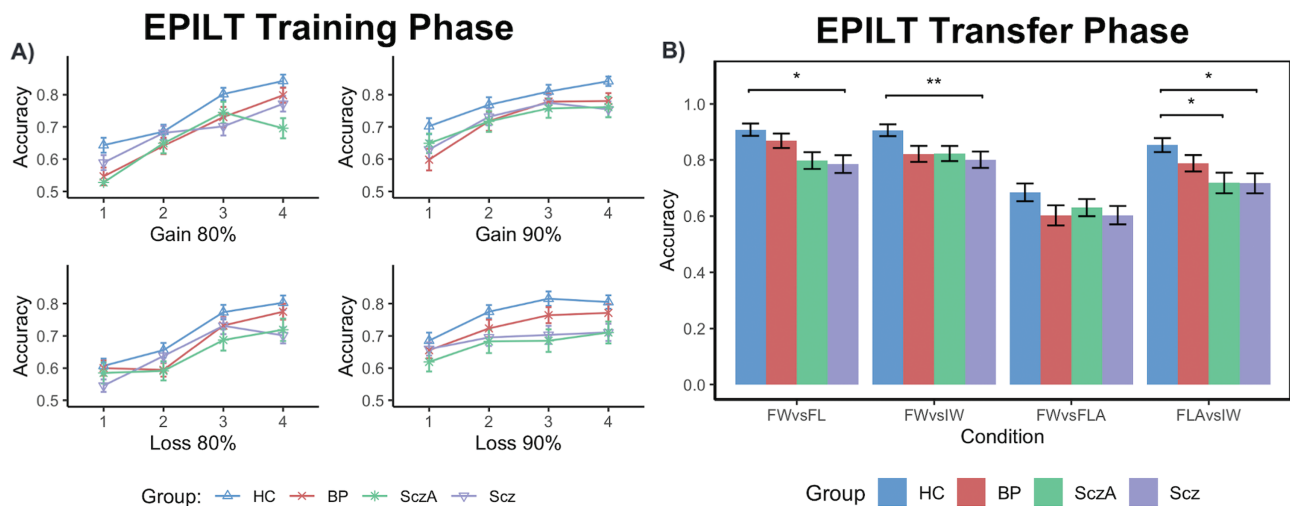


Fig. 4. (A) The Explicit Probabilistic Incentive Learning Tasks (EPILT)-Training accuracy across groups split by Gain or Loss conditions of the training phase across blocks and (B) EPILT-Transfer accurately identifying the most valuable item in a pair for each pairing split by group. * $P < .05$, ** $P < .01$.

.0001, $\eta^2_p = 0.602$), Probability ($F(1,224) = 39.3$, $P < .0001$, $\eta^2_p = 0.149$), Valence ($F(1,224) = 7.2$, $P = .008$, $\eta^2_p = 0.031$), and Group ($F(3,224) = 6.6$, $P = .0003$, $\eta^2_p = 0.082$). Post hoc analyses revealed that accuracy increased across Block in the 90% condition compared with the 80%, which replicated previous findings. The main effect of Valence also replicated; however, accuracy was higher for Gain conditions than Avoiding Loss, which was opposite previous findings. HC were more accurate than all patient groups, who did not differ in their accuracy from each other, replicating previous findings.

These main effects are qualified by significant interactions between Block $\times$ Probability ($F(3,672) = 15.46$, $P < .0001$, $\eta^2_p = 0.065$) and Block $\times$ Valence ($F(3,672) = 3.0$, $P = .028$, $\eta^2_p = 0.013$). For Block $\times$ Probability, the 80% condition improved in accuracy over blocks more than the 90% condition did. For Block $\times$ Valence, participants were more accurate in block 1 for the Loss Avoidance condition, but in block 4, participants were more accurate in the Gain condition. These interactions differ from the previous report.¹²

In the transfer phase (figure 4b), there was a main effect of Pairing ($F(3,221) = 84.5$, $P < .0001$, $\eta^2_p = 0.534$), but the interaction between Pairing $\times$ Group was not found in this dataset ($F(9,669) = 1.0$, $P = .419$, $\eta^2_p = 0.014$), thereby only partially replicating previous findings.¹² Post hoc contrasts showed that participants were most accurate for FWvsFL and FWvsIW, followed by FLAvsIW, and least accurate for FWvsFLA. We saw a main effect of Group ($F(3,223) = 6.6$, $P = .0003$, $\eta^2_p = 0.082$), not previously seen, where HC performed better than other groups (who performed equally).

Relationship to Symptoms. A significant first canonical variate revealed a relationship driven by training Avoiding Loss accuracy (standardized coefficient = 0.990) and a negative relationship with BPRS

positive symptoms and positive relationship with YMRS (standardized coefficients = -0.823 and 0.990 , respectively). However, this was not robust to cross-validation (supplemental table S8).

Discussion

To assess the usefulness of novel RL paradigms for measuring treatment effects in serious mental illness, we examined the reliability, practice effects, and replicability of tasks performed twice in the absence of a systematic intervention. The internal consistency of both the IPILT and EPILT was appropriate, but test-retest reliability was at best fair and poor for most measures. Further, there were no homogenous effects of practice for these paradigms; instead, groups differed on how they changed over time. Previous findings with these tasks largely replicated; a notable exception included an increase in bias toward reward across blocks in positively reinforced implicit learning. Lastly, performance generally did not relate to symptomatology in any robust way.

High internal consistency speaks to the unidimensional nature of the measures across blocks. If unidimensionality were sufficient for measuring change, the IPILT and EPILT would be sensitive to change. However, test-retest reliability is a better index of measurement noise and more directly translates into power calculations for (and the expense of) clinical trials.^{50,51} Perhaps our results were, therefore, impacted by meta-learning, the phenomenon of learning-to-learn. In the first testing session, participants are required to learn the right decisions in order to produce maximally efficacious outcomes on these tasks. However, there were no clear improvements in learning across testing sessions, with the exception of Avoiding Loss in explicit RL (supplemental tables S5–S7), suggesting that practice effects associated with meta-learning

were minimal. Some patient groups did change with practice: Scz developed an implicit bias toward reward more readily than HCs across sessions, individuals with a mood disorder had improved positive explicit RL, while other groups decreased in accuracy, etc. However, these instances did not result in differential ceiling effects across sessions. Another possibility that low reliability occurred is because these learning tasks used probabilistic feedback, which would facilitate learning sometimes and other times impede it; many more trials may be needed to overcome this. It may, therefore, be important to consider how RL might be approached using more deterministic feedback to reduce these sources of noise.

We did not observe much differentiation between patients and controls on our measure of implicit RL. This finding is consistent with the literature on RL in people with psychosis, with a general consensus that implicit RL is relatively intact, implying potentially normal basal ganglia DA firing. (see exception²⁴).^{8,10,21–23} While we found that SczA had a higher bias away from punishment than the other groups at session 1, this effect dissipated by session 2 and may not be robust.

Our findings are largely consistent with the literature for explicit RL. Previous studies have revealed a deficit in positive explicit RL in Scz but have been split about whether negative explicit RL is impaired in Scz (deficits in only positive,^{19,22,24,27–30} deficits in both,^{10,24,31–33} and deficits in neither).³⁴ In our sample, patient groups performed similarly and were significantly less accurate than HC for both the Gain and Avoid Loss conditions, though BP were intermediate for some conditions. We actually observed a greater deficit and more separation between groups in the Avoid Loss conditions lending support to a deficit in both positive and negative explicit RL and may point to negative RL being a better indication of diagnostic severity. However, this finding was the opposite of what our group previously found with this same paradigm in a different sample. All participants benefitted more from Gain feedback across blocks, despite initial accuracy being higher while Avoiding Loss. Patients were also impaired in their ability to distinguish novel pairings to previously learned reward stimuli. It appears that people with psychosis may fail to recruit the OFC and DLPFC for faster, more flexible decisions with top-down control for rewards and punishments.^{13,18–20}

Previous studies have observed that negative explicit RL performance decreases as negative symptoms increase, in particular amotivation/anhedonia.^{8,19} We did not find a relationship with negative symptoms and explicit RL but did see a worsening of negative explicit RL with increasing positive symptoms and decreasing mania symptoms that were not robust to cross-validation. Based on these findings, the EPILT might not be an appropriate task to measure negative symptoms and their mechanisms but might hold promise for measuring positive

symptoms. No other symptom dimension was related to explicit or implicit RL. The severity of symptoms was similar to those seen previously in our group when examining these tasks¹² and higher than other studies observing a relationship between negative symptoms and explicit RL performance.^{19,27}

There were a number of limitations to this study. As all patient groups in this analysis were on medications, with varying effects on DA receptors, we did not study the disorders in their natural state. Also, clinical symptoms were only assessed in the patient populations and examined at a separate time from either testing session, separated by ~2 weeks from the first cognitive testing session, perhaps attenuating stronger relationships. Further, the field is moving toward the inclusion of computational metrics for examining behavioral data; it is possible that these parameters may have better reliability and may assess the mechanisms involved in symptomatology more directly than traditional metrics. Finally, we do not recommend the inclusion of the “Eyes Big or Small” stimulus set for the IPILT, which complicated our results.

In conclusion, both the IPILT and the EPILT do not appear to be strongly affected by practice effects or symptomatology and are internally reliable. We were largely able to replicate previous findings in terms of patterns of spared implicit and impaired explicit RL in patients on these tasks. Some of the differences that we noted may reflect more power in this study, as we have larger sample sizes, and we typically found more parsimonious relationships between the effects of positive and negative reinforcers, with more main effects and 2-way interactions and less 3-way interactions. It will be important to address concerns with test-retest reliability before recommending these tasks be used in clinical trials for serious mental illness.

Supplementary Material

Supplementary material is available at *Schizophrenia Bulletin*.

Funding

This work was supported by the National Institute of Mental Health (R01s MH084840 to D.M.B., MH084826 to C.S.C., MH084821 to J.M.G., MH084828 to S.M.S., and MH084861 to A.W.M.).

Acknowledgments

The authors would like to thank the participants in this study, who generously gave their time. Parts of this manuscript were reported at the 2019 Society for Research in

Psychopathology conference. Author D.N.P. performed the data analysis, though would like to thank Michael Frank for consultation. D.M.B., C.S.C., J.M.G., S.M.S., and A.W.M. developed the study concept and design, aided in interpretation, and provided critical revisions. All authors approved the final version of the article for submission. The authors have declared that there are no conflicts of interest in relation to the subject of this study.

References

1. Fervaha G, Foussias G, Agid O, Remington G. Motivational deficits in early schizophrenia: prevalent, persistent, and key determinants of functional outcome. *Schizophr Res.* 2015;166(1-3):9-16.
2. Nakagami E, Xie B, Hoe M, Brekke JS. Intrinsic motivation, neurocognition and psychosocial functioning in schizophrenia: testing mediator and moderator effects. *Schizophr Res.* 2008;105(1-3):95-104.
3. Najas-Garcia A, Gómez-Benito J, Huedo-Medina TB. The relationship of motivation and neurocognition with functionality in schizophrenia: a meta-analytic review. *Community Ment Health J.* 2018;54(7):1019-1049.
4. Green MF, Nuechterlein KH, Gold JM, et al. Approaching a consensus cognitive battery for clinical trials in schizophrenia: the NIMH-MATRICES conference to select cognitive domains and test criteria. *Biol Psychiatry.* 2004;56(5):301-307.
5. Carter CS, Barch DM, Buchanan RW, et al. Identifying cognitive mechanisms targeted for treatment development in schizophrenia: an overview of the first meeting of the cognitive neuroscience treatment research to improve cognition in schizophrenia initiative. *Biol Psychiatry.* 2008;64(1):4-10.
6. Barch DM, Carter CS, Arnsten A, et al. Selecting paradigms from cognitive neuroscience for translation into use in clinical trials: proceedings of the third CNTRICS meeting. *Schizophr Bull.* 2009;35(1):109-114.
7. Ragland JD, Cools R, Frank M, et al. CNTRICS final task selection: long-term memory. *Schizophr Bull.* 2009;35(1):197-212.
8. Strauss GP, Waltz JA, Gold JM. A review of reward processing and motivational impairment in schizophrenia. *Schizophr Bull.* 2014;40(SUPPL. 2):107-116.
9. Gard DE, Kring AM, Gard MG, Horan WP, Green MF. Anhedonia in schizophrenia: distinctions between anticipatory and consummatory pleasure. *Schizophr Res.* 2007;93(1-3):253-260.
10. Gold JM, Waltz JA, Prentice KJ, Morris SE, Heerey EA. Reward processing in schizophrenia: a deficit in the representation of value. *Schizophr Bull.* 2008;34(5):835-847.
11. Sutton RS, Barto AG. *Reinforcement Learning: An Introduction*. Cambridge, MA, MIT Press; 2018.
12. Barch DM, Carter CS, Gold JM, et al. Explicit and implicit reinforcement learning across the psychosis spectrum. *J Abnorm Psychol.* 2017;126(5):694-711.
13. Frank MJ, Claus ED. Anatomy of a decision: striato-orbitofrontal interactions in reinforcement learning, decision making, and reversal. *Psychol Rev.* 2006;113(2):300-326.
14. Schultz W. Multiple dopamine functions at different time courses. *Annu Rev Neurosci.* 2007;30:259-288.
15. Schultz W. Activity of dopamine neurons in the behaving primate. *Semin Neurosci.* 1992;4(2):129-138.
16. Frank MJ. Dynamic dopamine modulation in the basal ganglia: a neurocomputational account of cognitive deficits in medicated and nonmedicated Parkinsonism. *J Cogn Neurosci.* 2005;17(1):51-72.
17. Centonze D, Picconi B, Gubellini P, Bernardi G, Calabresi P. Dopaminergic control of synaptic plasticity in the dorsal striatum. *Eur J Neurosci.* 2001;13(6):1071-1077.
18. O'Doherty J, Kringelbach ML, Rolls ET, Hornak J, Andrews C. Abstract reward and punishment representations in the human orbitofrontal cortex. *Nat Neurosci.* 2001;4(1):95-102.
19. Gold JM, Waltz JA, Matveeva TM, et al. Negative symptoms and the failure to represent the expected reward value of actions: behavioral and computational modeling evidence. *Arch Gen Psychiatry.* 2012;69(2):129-138.
20. Schoenbaum G, Roesch M. Orbitofrontal cortex, associative learning, and expectancies. *Neuron* 2005;47(5):633-636.
21. Heerey EA, Bell-Warren KR, Gold JM. Decision-making impairments in the context of intact reward sensitivity in schizophrenia. *Biol Psychiatry.* 2008;64(1):62-69.
22. Chang WC, Waltz JA, Gold JM, Chan TCW, Chen EYH. Mild reinforcement learning deficits in patients with first-episode psychosis. *Schizophr Bull.* 2016;42(6):1476-1485.
23. Somlai Z, Moustafa AA, Kéri S, Myers CE, Gluck MA. General functioning predicts reward and punishment learning in schizophrenia. *Schizophr Res.* 2011;127(1-3):131-136.
24. Waltz JA, Frank MJ, Wiecki TV, Gold JM. Altered probabilistic learning and response biases in schizophrenia: behavioral evidence and neurocomputational modeling. *Neuropsychology* 2011;25(1):86-97.
25. Reiss JP, Campbell DW, Leslie WD, et al. Deficit in schizophrenia to recruit the striatum in implicit learning: a functional magnetic resonance imaging investigation. *Schizophr Res.* 2006;87(1-3):127-137.
26. Weickert TW, Goldberg TE, Callicott JH, et al. Neural correlates of probabilistic category learning in patients with schizophrenia. *J Neurosci.* 2009;29(4):1244-1254.
27. Strauss GP, Frank MJ, Waltz JA, Kasanova Z, Herbener ES, Gold JM. Deficits in positive reinforcement learning and uncertainty-driven exploration are associated with distinct aspects of negative symptoms in schizophrenia. *Biol Psychiatry.* 2011;69(5):424-431.
28. Cheng GLF, Tang JCY, Li FWS, Lau EYY, Lee TMC. Schizophrenia and risk-taking: impaired reward but preserved punishment processing. *Schizophr Res.* 2012;136(1-3):122-127.
29. Waltz JA, Frank MJ, Robinson BM, Gold JM. Selective reinforcement learning deficits in schizophrenia support predictions from computational models of striatal-cortical dysfunction. *Biol Psychiatry.* 2007;62(7):756-764.
30. Reinen JM, Van Snellenberg JX, Horga G, Abi-Dargham A, Daw ND, Shohamy D. Motivational context modulates prediction error response in schizophrenia. *Schizophr Bull.* 2016;42(6):1467-1475.
31. Cicero DC, Martin EA, Becker TM, Kerns JG. Reinforcement learning deficits in people with schizophrenia persist after extended trials. *Psychiatry Res.* 2014;220(3):760-764.
32. Fervaha G, Agid O, Foussias G, Remington G. Effect of intrinsic motivation on cognitive performance in schizophrenia: a pilot study. *Schizophr Res.* 2014;152(1):317-318.

33. Farreny A, Del Rey-Mejías Á, Escartin G, et al. Study of positive and negative feedback sensitivity in psychosis using the Wisconsin Card Sorting Test. *Compr Psychiatry*. 2016;68:119–128.
34. Albrecht MA, Waltz JA, Cavanagh JF, Frank MJ, Gold JM. Increased conflict-induced slowing, but no differences in conflict-induced positive or negative prediction error learning in patients with schizophrenia. *Neuropsychologia* 2019;123:131–140.
35. Gold JM, Barch DM, Carter CS, et al. Clinical, functional, and intertask correlations of measures developed by the cognitive neuroscience test reliability and clinical applications for schizophrenia consortium. *Schizophr Bull*. 2012;38(1):144–152.
36. Pizzagalli DA, Jahn AL, O'Shea JP. Toward an objective characterization of an anhedonic phenotype: a signal-detection approach. *Biol Psychiatry*. 2005;57(4):319–327.
37. Pessiglione M, Seymour B, Flandin G, Dolan RJ, Frith CD. Dopamine-dependent prediction errors underpin reward-seeking behaviour in humans. *Nature* 2006;442(7106):1042–1045.
38. Hollingshead A, Redlich F. Social class and mental illness: community study. 1958. <https://psycnet.apa.org/record/2005-00650-000>. Accessed February 24, 2020.
39. Wechsler D. *Wechsler Test of Adult Reading: WTAR*. San Antonio, TX: Psychological Corporation; 2001.
40. First M, Spitzer R, Gibbon M, Williams J. *Structured Clinical Interview for DSM-IV-TR Axis I Disorders, Research Version, Patient Edition*. New York, NY: Biometrics Research; 2002.
41. Schneider LC, Struening EL, Elmer Struening TL. SLOF: A Behavioral Rating Scale for assessing the mentally ill. 1983. <https://academic.oup.com/swra/article-abstract/19/3/9/1627459>. Accessed February 10, 2020.
42. Ventura J, Nuechterlein KH, Subotnik K, Gilbert E. Symptom dimensions in recent onset schizophrenia: the 24-item expanded BPRS. *Schizophr Res*. 1995;15(1–2):22.
43. Young RC, Biggs JT, Ziegler VE, Meyer DA. A rating scale for mania: reliability, validity and sensitivity. *Br J Psychiatry*. 1978;133:429–435.
44. Berk M, Malhi GS, Cahill C, et al. The Bipolar Depression Rating Scale (BDRS): its development, validation and utility. *Bipolar Disord*. 2007;9(6):571–579.
45. Kring AM, Gur RE, Blanchard JJ, Horan WP, Reise SP. The Clinical Assessment Interview for Negative Symptoms (CAINS): final development and validation. *Am J Psychiatry*. 2013;170(2):165–172.
46. Pizzagalli DA, Iosifescu D, Hallett LA, Ratner KG, Fava M. Reduced hedonic capacity in major depressive disorder: evidence from a probabilistic reward task. *J Psychiatr Res*. 2008;43(1):76–87.
47. Shrout PE, Fleiss JL. Intraclass correlations: uses in assessing rater reliability. *Psychol Bull*. 1979;86(2):420–428.
48. McGraw KO, Wong SP. Forming inferences about some intraclass correlation coefficients. *Psychol Methods*. 1996;1(1):30–46.
49. Strauss ME, McLouth CJ, Barch DM, et al. Temporal stability and moderating effects of age and sex on CNTRACS task performance. *Schizophr Bull*. 2014;40(4):835–844.
50. Weir JP. Quantifying test-retest reliability using the intraclass correlation coefficient and the SEM. *J Strength Cond Res*. 2005;19(1):231–240.
51. Matheson GJ. We need to talk about reliability: making better use of test-retest studies for study design and interpretation. *PeerJ*. 2019;7:e6918–e6918.