



**DEPARTMENT OF ECONOMICS
DISCUSSION PAPER SERIES**

DYNAMIC RAWLSIAN POLICY

Charles Brendon and Martin Ellison

Number 595
March 2012

Manor Road Building, Oxford OX1 3UQ

Dynamic Rawlsian Policy

Charles Brendon and Martin Ellison*

University of Oxford

March 14, 2012

Abstract

A well-known time-inconsistency problem hinders optimal decision-making when policymakers are constrained in their present choices by expectations of future outcomes. The time-inconsistency problem is caused by differences in the preferences of policymakers who exist at different points in time. Adapting the arguments of Rawls (1971), we propose that these differences can be eliminated if policy is set from behind a ‘veil of ignorance’, without knowledge of when the policy will be implemented. We set up a well-defined choice problem that captures this normative perspective. The policies that it generates have a number of appealing properties.

Keywords: macroeconomic policy, Rawls, time inconsistency, veil of ignorance

JEL classification: E52, E61.

*Department of Economics, University of Oxford, Manor Road, Oxford, OX1 2UQ, United Kingdom. E-mails *charles.brendon@economics.ox.ac.uk* and *martin.ellison@economics.ox.ac.uk*.

1 Introduction

This paper introduces a new approach to solving time-inconsistency problems, and shows that the resulting policy has a number of desirable properties. Time-inconsistency problems occur whenever an economic agent's view of the best action to take in a particular time period changes solely because of the passage of time. The problem is present in a large class of models for which expectations of future policy affect current constraints, as first identified by Kydland and Prescott (1977). The problem arises in these models because of conflicts between the preferences of policymakers existing at different points in time. For example, the current policymaker may have an incentive to make promises that future policymakers do not want to honour. Almost all contemporary monetary or fiscal policy models feature conflicts of this type in one form or another.

Our idea is to remove conflict between policymakers by denying them knowledge of when they will be called upon to set policy. To do this we invoke the famous 'veil of ignorance' concept of Rawls (1971), where the veil is used to conceal all information relevant to the time period in which policy choices will be implemented. The aim is to capture the normative idea that institutional design ought to be conducted in a manner that abstracts from the special circumstances of conflicting parties, and should ideally lead to a policy that can be agreed upon by all. In models with time-inconsistency problems, concealing time behind a veil of ignorance removes the source of conflict between policymakers existing at different points in time. As Rawls advocates, they find it "impossible to tailor principles to the circumstances of their own case".

Two approaches to the time inconsistency problem dominate the existing literature. Neither is especially satisfactory. The 'commitment' solution sets policy to maximise the preferences of the particular policymaker that exists at the start of time, and assumes that these preferences can be imposed on all future generations of policymakers. It is not explained why the policymaker at the start of time is uniquely able to impose their preferences on policymakers in the future. As Svensson (1999) puts it, "Why is period 0 special?". The 'discretionary' solution instead assumes that choices can only be made contemporaneously, and defines policy as the outcome of a non-cooperative game between policymakers existing at different points in time. This is likely to yield outcomes that are undesirable for every policymaker. The most prominent alternatives to commitment and discretion are the timeless perspective of Woodford (2003) and the unconditional optimality approach first suggested by Taylor (1979), and more recently explored by Damjanovic et al. (2008). We show that our approach has a number of

appealing properties relative to these alternatives.

The paper is organised as follows. In Section 2 we specify the economic environment and discuss existing solutions to the time-inconsistency problem it contains. We introduce the idea of dynamic Rawlsian policy in Section 3, and define it for purely forward-looking models in Section 4. The final part of the section presents an example to illustrate how dynamic Rawlsian policy differs from other approaches in a simple New Keynesian model. The general case for models with both backward and forward looking constraints is in Section 5, which finishes with an example based on the problem of optimal redistribution with participation constraints in the spirit of Marcat and Marimon (1992) and Kocherlakota (1996). A final Section 6 concludes.

2 Specification of the environment

We consider a general problem in which a policymaker operating at time 0 has preferences over the expected present discounted value of a time-invariant objective function:

$$W_0 = \sum_{t=0}^{\infty} \beta^t \int_{H^\infty} \pi(x(h, t), \varepsilon) dF_t(h|h^0). \quad (1)$$

The function $x : H^\infty \times \mathbb{Z} \rightarrow X$ is a mapping from the history of endogenous and exogenous variables $h \in H^\infty$ and time period t to the t -dated realisation of a vector of endogenous variables $x \in X$. History h includes the vector of contemporary exogenous variables $\varepsilon \in E$, and $H^\infty = X^\infty \times E^\infty$. The infinite history of exogenous variables is denoted $h_\varepsilon \in E^\infty$, and we define F as an exogenous probability measure on E^∞ that characterises the evolution of these variables. The mapping $x(h, t)$, probability measure F , and initial history h^0 together induce a time-specific conditional probability measure on H^∞ . We denote this $F_t(A|h^0)$. It gives the probability that time t is characterised by a history in the set $A \subseteq H^\infty$, conditional on initial history h^0 . For simplicity, the time-invariant objective function $\pi(\cdot)$ is assumed to be continuously differentiable in all its arguments, and strictly quasi-concave in x for all ε . β is the discount factor of the policymaker. The preferences W_s of a policymaker in period s are defined analogously.

The policymaker at time 0 chooses the function $x(h, t)$ subject to two types of constraints. The first are the n ‘forward-looking’ restrictions:

$$\sum_{r=0}^T \int_{H^\infty} g_r(x(h, t+r), \varepsilon) dF_r(h|h^0) \geq 0, \quad (2)$$

which hold for all $h' \in H^\infty$ and t . T is the maximum horizon at which expectations of future variables influence current constraints, and may be infinite in some models of interest. The function $g_r : X \times E \rightarrow \mathbb{R}^n$ is vector-valued if $n > 1$. Each component of the function is assumed to be continuously differentiable and weakly concave in x , so the constraint set is convex for given ε . Examples of forward-looking constraints in macroeconomics include the standard New Keynesian Phillips curve (where $T = 1$), and fiscal solvency requirements (where T is typically infinite).

The second type of constraints are ‘backward-looking’:

$$m(x(h, t), z(h \setminus 1, t - 1), \varepsilon) \geq 0. \quad (3)$$

These similarly hold for all h and t . The function $z : H^\infty \times \mathbb{Z} \rightarrow Z$ is equivalent to the function x but with its output restricted to the model’s predetermined variables, defined as those variables that enter (3) with a lag. We denote by $y \in Y$ the remaining non-predetermined variables, so $X = Y \times Z$, and assume that the predetermined variables do not feature as arguments in any of the g_r functions for $r \geq 1$. This is without loss of generality: we can always define new contemporaneous auxiliary variables to substitute out for any predetermined variables in these functions. $h \setminus r$ is the history obtained by deleting the last r observations in h . We assume that $m(\cdot)$ is continuously differentiable and weakly concave in (x, z) , preserving the convexity of the constraint set. An example of a backward-looking constraint in macroeconomics is the law of motion for capital accumulation.

2.1 Time-0 optimality and time-inconsistency

The optimal choice of $x(h, t)$ from the time 0 policymaker’s perspective can be determined by standard techniques. The row vectors $\lambda_g(h, t)$ and $\lambda_m(h, t)$ are multipliers associated with the forward and backward looking constraints at time t after history h . The first order condition for maximising (1) subject to (2)-(3) with respect to the generic variable x_i is:

$$\begin{aligned} & \frac{\partial \pi(x(h, t), \varepsilon)}{\partial x_i(h, t)} + \sum_{r=0}^{\min\{t, T\}} \beta^{-r} \lambda_g(h \setminus r, t - r) \frac{\partial g_r(x(h, t), \varepsilon)}{\partial x_i(h, t)} \\ & + \lambda_m(h, t) \frac{\partial m(x(h, t), z(h \setminus 1, t - 1), \varepsilon)}{\partial x_i(h, t)} \\ & + \beta \int_{H^\infty} \lambda_m(h', t + 1) \frac{\partial m(x(h', t + 1), z(h, t), \varepsilon')}{\partial x_i(h, t)} dF_1(h' | h) \\ & = 0, \end{aligned} \quad (4)$$

for almost all feasible histories h at each time period $t \geq 0$.

Under the regularity conditions that we imposed, there is a unique function $x_o^0(h, t)$ that satisfies this condition and the constraint set. The subscript denotes that this is an ‘optimal’ solution and the superscript indicates the time period from which it is considered optimal, in this case period 0. $x_o^0(h, t)$ is time-inconsistent because, in general, $x_o^0(h, t) \neq x_o^s(h, t)$ for $0 < s \leq t$. This follows from the first order condition. If $T > 0$ then the upper limit of the summation in (4) depends upon t for at least one x_i , at least when moving from $t = 0$ to $t = 1$. This means that if optimisation were taking place in a different period than time 0 then the set of first-order conditions would be different. This is the familiar ‘time-inconsistency of optimal plans’ first highlighted by Kydland and Prescott (1977). How a unique $x(h, t)$ should be chosen for all (h, t) pairs is the problem addressed by our paper.

2.2 Discretion

The discretionary approach deals with the problem of time-inconsistency by assuming that each policymaker only chooses contemporaneous variables. Responses of future policymakers to current choices are treated as known, and discretionary policy is defined as an outcome of a non-cooperative game between policymakers existing at different points in time. For all t , a ‘discretionary’ policy function $x_d^t(h, t)$ maximises W_t subject to constraints (2)-(3) holding at time t , and the outcomes for $s > t$ being determined by known functions $x_d^s(h, s)$.¹

The resulting policy is not easy to characterise analytically, since there may be complex equilibrium dependence on past choices. Even in simple linear-quadratic models, if lagged endogenous variables feature in the constraint set then Blake and Kirsanova (2012) show that multiple Markov-stationary discretionary equilibria may exist. There is also the general possibility of non-Markov ‘reputational’ equilibria of the type in Ireland (1997). For this reason we cannot compactly state a general set of necessary optimality conditions. We do know that current policymakers are never directly concerned about the impact of their choices on past expectations, so the $\frac{\partial g_r(x(h, t), \varepsilon)}{\partial x_i(h, t)}$ terms for $r \geq 1$ in the first order conditions characterising commitment (4) will certainly not feature in any optimality conditions for discretion. This generally makes equilibrium discretionary outcomes undesirable to policymakers in every time

¹The superscript is as before the time period within which options are being assessed, whilst the time argument of the function denotes the period for which the assessment is being made. Since the defining characteristic of discretionary choice is that decisions are taken by contemporaneous policymakers subject to future responses, there is no natural value to attach to $x_d^t(h, s)$ when $s > t$.

period.

2.3 Commitment

The commitment approach assumes that the objectives of future policymakers can be amended in such a way that their preferences and policy actions cohere with those desired by the policymaker at time 0. Marcet and Marimon (1998) assume that this can be done contractually, with the period s policymaker legally obliged to maximise a welfare objective V_s^{c0} that incorporates a concern for honouring past promises. The objective of the policymaker at time $s > 0$ is thus changed to:

$$V_s^{c0} = W_s + \sum_{r=1}^{\min\{s,T\}} \sum_{t=0}^T \beta^{-r} \lambda_g(h \setminus r, s-r) \int_{H^\infty} g_{r+t}(x(h', s+t), \varepsilon) dF_t(h'|h), \quad (5)$$

for the given history h realised up to time s . The lagged multipliers $\lambda_g(h \setminus r, s-r)$ are held equal to the values they were assigned when previously optimising in period $s-r$. A comparison of (5) and (4) shows that repeated implementation of this objective will maximise W_0 subject to all relevant constraints.

We denote by $x_{c0}^s(h, s)$ and $x_{c0}^s(h, t)$ the optimal choices for periods s and $t > s$ made by the period- s policymaker under the objective V_s^{c0} . The contractual approach ensures that $x_{c0}^s(h, t) = x_{c0}^r(h, t)$ for all $t \geq r, s \geq 0$, so there is coherence between the views of policymakers existing in different periods regarding what is the optimal policy. Notice though that $x_{c0}^t(h, t) = x_o^0(h, t)$ by construction, so the fact that the policymaker at time t chooses t -dated variables is really a matter of semantics: to all intents and purposes the time-0 policymaker is a dictator. This aspect of the commitment approach is questionable. Why is the policymaker at time 0 uniquely placed to break free from all prior commitments and impose its will on all future policymakers? Why select $x_{c0}^t(h, t)$ instead of $x_{c1}^t(h, t)$ or $x_{c2}^t(h, t)$?

2.4 Timeless optimality

The timeless perspective method of Woodford (2003) proposes changing the objectives of all policymakers, including those of the policymaker that exists at time 0. The idea is to incorporate the impact of current decisions on past constraints as in (5), but so that policy will be optimal under the preferences of a hypothetical policymaker that existed infinitely long

ago. Accordingly, the timeless perspective imposes the objective:

$$V_s^{tp} = W_s + \sum_{r=1}^T \sum_{t=0}^T \beta^{-r} \lambda_g(h \setminus r, s-r) \int_{H^\infty} g_{r+t}(x(h', s+t), \varepsilon) dF_t(h'|h), \quad (6)$$

for the given history h realised up to time s .² It remains to choose the values of $\lambda_g(h \setminus r, s-r)$ for $r > s$, which cannot be taken from a previous optimisation problem if time 0 is the first period in which the timeless perspective is applied. Woodford (2003) makes the choice by first solving for policy as a function of these multipliers, and then setting them according to functions of h that are also satisfied by subsequent multipliers.³ This ensures that $\lambda_g(h, t) = \lambda_g(h, s) = \lambda_g(h)$ for all t, s , and h , thereby making the timeless perspective solution time-invariant. We denote by $x_{tp}^s(h, s)$ any policy that solves this problem, as assessed from the perspective of a policymaker in period s following history h . With $x_{tp}^s(h, t)$ defined for $t > s$ as before, we achieve coherence of preferences across different policymakers in the sense that $x_{tp}^s(h, t) = x_{tp}^r(h, t)$ for $t \geq r, s \geq 0$.

The timeless approach gives no ‘dictatorial’ status to the policymaker at time 0, but it has other problematic features. The basic time-inconsistency problem is that $x_o^r(h, t) \neq x_o^s(h, t)$ for $0 \leq r < s \leq t$. However, rather than actively seeking to resolve the conflict between policymakers existing at different points in time, the timeless perspective proceeds by imposing the preferences of a policymaker who never actually existed. There is nothing to ensure that the resulting policy need even be in the ‘Pareto set’ of policies for which it is impossible to improve the outcome of one policymaker without worsening the outcome of another. Indeed, Section 4.6 presents a monetary policy problem in which an alternative policy is strictly preferred to the timeless policy by every policymaker in every period.

Another concern is that the objective function (6) requires inverse discounting of the Lagrange multipliers at the rate β^{-1} . This means that the sum in the objective may not converge in models where the horizon T of the forward-looking constraints is infinite. Convergence requires that $\lim_{r \rightarrow \infty} \beta^{-r} g_r(x(h, t), \varepsilon) = 0$, but in many models the infinite sums relate to net present values of revenue or utility streams, ensuring that the g_r functions are themselves approximately proportional to β^r . Hence the sum may not converge, and the timeless objective (6) may not be well defined. The participation constraints that feature in the example of Section 5 present this problem.

²In the commitment problem, if date 0 were in the infinite past then $\min\{s, T\} \rightarrow T$, and (5) would indeed coincide with this.

³A slight problem with this prescription is that there are generally many admissible functions of this form. Giannoni and Woodford (2002) present a particular criterion that delivers uniqueness.

2.5 Unconditional optimality

The unconditional optimality method restricts policymakers existing at all points in time to be concerned only about steady-state outcomes. Such a restriction eliminates the short-term incentive to capitalise on fixed prior expectations that is at the heart of the time-inconsistency problem. The method dates back at least to Taylor (1979), and was explored recently by Damjanovic *et al.* (2008). The proposal is that policy should maximise the unconditional expectation of the within-period objective criterion in stochastic steady state, which implies replacing W_0 with:

$$E(W_0) \equiv \frac{1}{1-\beta} \int_{H^\infty} \pi(x(h, t), \varepsilon) dF_u(h). \quad (7)$$

The expectation in (7) is taken with respect to a steady-state probability distribution F_u that is defined over Borel subsets of H^∞ in a manner consistent with the chosen policy.⁴ In other words, if policy renders certain values of the endogenous variables more likely in steady state then F_u endogenously adjusts to reflect this. In general it is also assumed that the time argument in $x(h, t)$ is redundant, so that the same choice is made regardless of when a particular history is observed. We denote by $x_{uo}^s(h, t)$ the policy desired for period t by the policymaker at s following history h . The unconditional approach is consistent with $x_{uo}^s(h, t) = x_{uo}^r(h, t)$ for $t \geq r, s \geq 0$, though it does not guarantee it as preferences are not well defined out of steady state.

The exclusive focus of the unconditional approach on steady-state outcomes means it is equivalent to assuming that the policymaker does not discount due to pure time preference. There is no incentive for optimal policy to distort the steady state in return for more benign transition dynamics. This implies, for example, that an unconditionally optimal policy in the Ramsey growth model maximises steady-state consumption rather than following the modified golden rule. The pure focus of unconditional optimality on steady-state welfare ultimately proves quite restrictive. In Section 5 we present a simple problem of optimal redistribution subject to participation constraints for which the ‘unconditionally optimal’ policy cannot be defined, since steady-state welfare is not bounded.

⁴Damjanovic *et al.* (2008) implicitly rule out dependence of x on lags of endogenous variables beyond the first. This would mean that F_u need only be defined on subsets of $E^\infty \times X$, rather than H^∞ .

3 Rawlsian policymaking

This paper shares the idea of the commitment, timeless and unconditional optimality approaches that preferences should be amended to ensure coherence between the objectives of policymakers existing at different points in time. We furthermore agree with the timeless and unconditional optimality approaches that there may be better ways of doing this than imposing the preferences of the time 0 policymaker on all other policymakers, as the commitment approach does. Where we disagree is about the best way of amending preferences to ensure that policymakers have coherent objectives. There is plenty of scope for us to disagree, given that it is not obvious what objectives to use if we do not impose the time 0 policymaker's preferences. The timeless approach suggests adopting the preferences of a hypothetical policymaker that existed in the infinite past, whereas the unconditional optimality approach constructs new objectives based solely on steady-state outcomes. In neither case is the normative justification particularly clear. We seek to redress this by proposing a new approach that has both normative appeal and a number of desirable features that improve upon existing approaches.

Our proposal attacks the time-inconsistency problem at source. We do this by recognising that the extent to which $x_o^r(h, t)$ differs from $x_o^s(h, t)$ for $t \geq r, s \geq 0$ is determined solely by the differing amounts of time that will pass before period t arrives, from the perspective of policymakers making assessments about period t in different periods r and s . The policymakers in periods r and s disagree about what is best to do in period t , because the differing amounts of time before period t arrives mean they face different expectational constraints. The idea of our approach is to remove the information from which the disagreement flows, i.e., we ask what happens if policymakers do not know how much time will pass before period t arrives.

The motivation for our approach is taken from Rawls (1971). His theory of justice argues that institutional design should take place behind a 'veil of ignorance' that prevents policy from being tailored to the particular circumstances of any of its designers. It should be as if this veil conceals information about the position a person will take in the society whose basic structures are being determined. For example, it will be as if each designer does not know where they will be in the distribution of income across individuals. The idea is that removing information about particular circumstances means that choices must be made on the basis of more general principles than basic self-interest. In our problem, the particular circumstance we wish to conceal is the time between the period a given policymaker exists and the period for which choice is being made. This can be achieved by applying a veil of ignorance to conceal

information about the passage of time. Given that it is the only source of disagreement, we expect that denying policymakers knowledge of time should lead to coherence in their objectives, whilst ensuring that policy is based on the ‘true’ preferences to the greatest extent possible.

The aim of this paper is to capture what it means to set policy behind a veil of ignorance that denies policymakers knowledge of time. It is not immediately obvious how this can be done. A natural restriction is that choices should not be a function of time itself, but policies that condition on the history of endogenous variables in ways that allow the policymaker to keep track of time also need to be ruled out. We additionally need to prevent the policymaker from using initial knowledge about the relative likelihoods of different current and future histories, because if the policymaker knows that a particular history is more likely to be observed at one date than another then conditioning policy on that history could effectively restore the policymaker’s knowledge of time. In general, we assert that the information set of policymakers must be orthogonal to time, at least insofar as policy choices relate to variables that feature in expectational constraints.

The outcome of our dynamic Rawlsian policy procedure is a desired outcome for period t for a policymaker assessing policy at $s \leq t$. We denote the policy $x_R^s(h, t)$, with corresponding desired outcomes of $y_R^s(h, t)$ in the non-predetermined variables Y and $z_R^s(h, t)$ in the predetermined variables Z . Our task is to characterise this procedure and obtain a desired choice $x_R^s(h, t)$ for all h and all $t \geq s \geq 0$.

4 Purely forward-looking models

We begin by considering models with only forward-looking constraints. Whilst this is fairly restrictive, it sets the scene for the more complex description of dynamic Rawlsian policy in a general setting. The advantage of purely forward-looking models is that the history of endogenous variables is irrelevant to the constraints of the policy problem, and this simplifies matters considerably.

4.1 Restricting dependence

The information set of policymakers needs to be orthogonal to time under the veil of ignorance. An obvious first step is to prevent policy directly conditioning on the time period:

Condition 1 $y_R^r(h, s) = y_R^r(h, t)$ for any history h and all $s, t \geq r \geq 0$.

When Condition 1 is applied, the vector of endogenous variables y no longer depends directly on the time period and we can write $y_R^r(h)$ as the optimal choice for the policymaker in period r following history h . We also need to restrict policy so that it does not depend on variables from which the passage of time could be indirectly inferred, due to them being chosen by the policymaker. There is a natural way to do this in models that are purely forward-looking, because the history of endogenous variables in such models does not affect the primitives of the policy problem in any way. This means there can be no possible benefit from allowing policy to react to endogenous variables, except as a means to exploit the time-specific information that they convey. This is exactly the sort of indirect time dependency we wish to rule out. The easiest way to prevent it is to require that choice does not condition on the history of non-predetermined endogenous variables.⁵

Condition 2 $y_R^s(h, t) = y_R^s(h', t)$ for any histories $h = (h_\varepsilon, h_y)$ and $h' = (h_\varepsilon, h'_y)$ and all $t \geq s \geq 0$.

We believe this is uncontroversial. The history of endogenous variables in any model is completely determined by the history of exogenous variables and the policy in place at the time the endogenous variables were selected. Allowing policy to depend on the reasoning applied by past institutions would introduce unnecessary arbitrariness. Conditions 1 and 2 together imply that endogenous variables in purely forward-looking models must be chosen as a function only of the history of exogenous variables, and the optimal choice for the policymaker in period s following history h_ε can thus be denoted $y_R^s(h_\varepsilon)$.

4.2 History weighting

In almost all cases, imposing Conditions 1 and 2 is not in itself sufficient to make policy independent of the period in which optimisation takes place. The problem is that the period s policymaker has knowledge of the initial history h^s , and can use it in conjunction with future histories h^r to keep track of time. For example, the policymaker can successively eliminate the most recent observations from a future history h^r until they obtain the initial history h^s . In our earlier notation where $h^r \setminus r - s$ is the history obtained by deleting the last $r - s$ observations from h , this must satisfy $h^r \setminus r - s = h^s$. The number of observations eliminated indicates how many periods have passed, because the probability that $h^r \setminus r - s^* = h^s$ for some

⁵We define h_y and h_z as the histories of control and state variables respectively. Unlike h_ε , these do not include contemporaneous entries.

other date $s^* \neq s$ is zero in a stochastic model featuring continuous random variables. Once the policymaker can keep track of time, they have an incentive to make policy depend on the date in which optimisation takes place, even if policy functions themselves are time-invariant. For example, policy choices will give greater weight to future histories that are more likely given the initial history h^s . The resulting policy would not be Rawlsian in the sense that we have described.

We rule out unwanted history dependence of this type by denying the policymaker knowledge of the particular initial history that prevails at the time optimisation takes place. The period- s policymaker then has to set policy that optimises over a distribution of initial histories $G(h^s)$ rather than optimising to a known initial history h^s .

$$W_s = \int_{H^\infty} \left[\sum_{t=s}^{\infty} \beta^{t-s} \int_{H^\infty} \pi(x(h, t), \varepsilon) dF_{t-s}(h|h^s) \right] dG(h^s). \quad (8)$$

The choice of weighting function G should ensure that observing a particular history conveys no information about time. Under Condition 2 the policymaker only observes h_ε , and the unconditional distribution of exogenous variables $F(h_\varepsilon)$ is the unique weighting function consistent with history being uninformative about time. This restriction fits our Rawlsian perspective that policy should be set behind a veil of ignorance. It is summarised by:

Condition 3 *The policymaker in period $s \geq 0$ uses the unconditional distributions of exogenous variables $F(h_\varepsilon)$ as a ‘weighting function’ over initial histories when assessing current and future policy choices.*

In words, the policymaker considers a history h_ε no more likely to characterise any one period than another.

4.3 Dynamic Rawlsian policy

Taken together, Conditions 1 to 3 imply that the period s policymaker acting behind a veil of ignorance in a purely forward-looking model assesses outcomes according to the function:

$$W_s^R = \int_{E^\infty} \left[\sum_{t=s}^{\infty} \beta^{t-s} \int_{E^\infty} \pi(y(h_\varepsilon), \varepsilon) dF_{t-s}(h_\varepsilon|h'_\varepsilon) \right] dF(h'_\varepsilon). \quad (9)$$

This can be maximised subject to (2), which must hold for all possible histories. Appendix A outlines the mechanics for solving this problem. Defining Lagrange multipliers in the usual

way, the characteristic first order conditions are:

$$\frac{\partial \pi(y(h_\varepsilon), \varepsilon)}{\partial y_i(h_\varepsilon)} = - \sum_{r=0}^T \lambda_g(h_\varepsilon \setminus r) \frac{\partial g_r(y(h_\varepsilon), \varepsilon)}{\partial y_i(h_\varepsilon)}, \quad (10)$$

which together with the constraint set and complementary slackness conditions are sufficient to characterise $y_R^s(h, t)$.

The object on the left-hand side of (10) can be thought of as the ‘direct’ marginal benefit to the policymaker from changing the value of the criterion function π following history h_ε . The sum on the right-hand side collects the shadow costs of this, through the impact on forward-looking constraints binding under predecessor histories up to the maximum horizon T . The most notable feature of (10) is the absence of the discount factor β from the terms on the right-hand side. This implies that dynamic Rawlsian policy equates the direct marginal benefit of changing endogenous variables in any given time period to the shadow costs of changing expectations that are also formed in that time period, even though in general the expectations pertain to outcomes in future periods. This may seem surprising at first, for it appears that the policymaker is simultaneously varying outcomes in both the present and the future. But on reflection, from a Rawlsian policymaker’s point of view it is equally likely for history h_ε to occur at time s as at time $s + r$, so marginal changes to policy in response to that history have to be treated as occurring in both periods at once with equal probability. The time preference structure of (10) follows from this.⁶ This contrasts with the commitment approach, where outcomes are period-specific by design and any expectational effects in first-order conditions reflect changed constraints in time periods prior to the associated direct effects, thereby yielding ‘inverse discounting’ in the summation.

4.4 Comparison to timeless optimality

It is useful to compare our dynamic Rawlsian policy to the timeless method suggested by Woodford (2003) and discussed in Section 2.4. In a purely forward-looking model, the first order conditions for maximising the timeless objective (6) subject to the forward-looking

⁶That is, if the policymaker is optimising in period s then outcomes are changed with equal probability in all periods from s onwards. This implies a change to expectations formed in s regarding outcomes in all periods up to $s + T$, as well as the ‘direct’ effects accruing in s . Symmetrically, there are direct effects accruing in $s + 1$, along with changes to that period’s expectations of outcomes at every horizon to $s + 1 + T$. The same is true for $s + 2$, and so on. Appendix A shows how the associated first-order condition reduces to (10).

constraints (2) can be written as:

$$\frac{\partial \pi(y(h, t), \varepsilon)}{\partial y_i(h, t)} = - \sum_{r=0}^T \beta^{-r} \lambda_g(h \setminus r, t - r) \frac{\partial g_r(y(h, t), \varepsilon)}{\partial y_i(h, t)}. \quad (11)$$

The main difference between (11) and the first order conditions for dynamic Rawlsian policy (10) is the β^{-r} term in the summation on the right hand side. Consistent with the discussion in Section 4.3, this is because timeless policy allows choice over variables that are known to be particular to a given time period. What matters when choosing these is the relative sizes of direct marginal benefits that obtain in period t and the shadow benefits from relaxing expectational constraints that were binding prior to period t . The direct marginal benefits accrue at the same time as the variables are changed, whereas the shadow values of relaxing prior constraints are given a time preference weighting consistent with the period in which those constraints bind, i.e., r periods ago for some $r \in \{0, \dots, T\}$. The shadow values of relaxing constraints are therefore ‘inverse discounted’ by a factor β^{-r} relative to the direct effects.

More subtly, the timeless approach assumes that the shadow values for all periods back to $t - T$ are always relevant when choosing for period t . This is why there is no *min* in the upper limit of the summation in (11), in contrast to the equivalent commitment condition implied by (4). It reflects the fact that choice is being tailored to the preferences of a policymaker in the infinite past, to whom all such horizons would indeed be relevant. By contrast, the veil of ignorance approach only considers marginal effects accruing onwards from the period in which optimisation takes place. The fact that the summation in (10) also runs to T results from the fact that all future horizons up to $s + T$ are relevant to optimality assessments in period s , rather than an assumption in the timeless case that all past horizons back to $s - T$ should be considered.

4.5 Comparison to unconditional optimality

A second point of comparison is the unconditionally optimal approach, described in Section 2.5. In a purely forward-looking model, maximising the unconditional objective (7) subject to the forward-looking constraints (2) gives first order conditions of the form:

$$\frac{\partial \pi(y(h, t), \varepsilon)}{\partial y_i(h, t)} = - \sum_{s=0}^T \lambda_g(h \setminus s, t - s) \frac{\partial g_s(y(h, t), \varepsilon)}{\partial y_i(h, t)}. \quad (12)$$

This must hold only for F_u -almost all histories h , where F_u is the unconditional steady-state distribution on H^∞ induced by policy. This is quite a limited characterisation of policy,

since only a relatively small subset of H^∞ is likely to be consistent with the steady state. The policymaker is implicitly indifferent about changes to policy prior to steady-state being achieved, provided these do not stop F_u from characterising that steady-state. That is, the policymaker is indifferent about policies after histories with no density under F_u .

In its important details, condition (12) coincides with the first-order condition of our dynamic Rawlsian policy. In particular, its treatment of time preference is identical and there is no ‘inverse discounting’ in the summation. The rationale behind it is different though. Unconditionally optimal policy restricts the policymaker to care only about outcomes that obtain once stochastic steady-state has been achieved, however long it may take to get there. This means considering marginal changes to policy as if expectations about future outcomes are changed at the same time as contemporary outcomes. This is what it means to vary the steady state of the model, and thus is an implication of the focus on steady-state welfare. The Rawlsian approach instead imposes that the policymaker must treat outcomes as if direct and expectational effects are induced simultaneously, as a consequence of the policymaker being denied knowledge of when outcomes are to occur.

The Rawlsian approach induces an ‘unconditionally optimal’ stochastic steady state immediately in purely forward-looking models, but note that any policy ultimately resulting in the same steady state is considered equally good under the unconditional objective. This includes policies that depend on lagged endogenous variables, and policies that induce convergence to the steady state only in the limit as time tends to infinity. This seems a considerable disadvantage of the unconditionally optimal approach.

4.6 Example

A monetary policy problem in the linearised New Keynesian model provides a suitable setting to illustrate how dynamic Rawlsian policy differs from other approaches in purely forward-looking models. The policy problem is to maximise an objective function defined over inflation and output, subject to the constraint of a purely forward-looking New Keynesian Phillips curve. We assume that the market power of firms is not corrected by any output subsidy, so output tends to be below the socially efficient level. This gives the policymaker an incentive to stimulate output, even at the cost of some inflation. We work with a purely deterministic model to make the analysis as simple as possible.

The preferences of the policymaker at time 0 are defined over inflation $\pi(h, t)$ and output

$y(h, t)$:

$$-\frac{1}{2} \sum_{t=0}^{\infty} \beta^t [\pi(h, t)^2 + \omega(y(h, t) - y^*)^2], \quad (13)$$

where the target level of output y^* is strictly positive. The policymaker has discount factor β and places weight ω on deviations of output from its target level. The model is deterministic, so there is no expectations operator and the history vector h relates only to endogenous variables. The conventional New Keynesian Phillips curve provides a purely forward-looking constraint:

$$\pi(h, t) = \beta \pi(h', t+1) + \gamma y(h, t), \quad (14)$$

where γ is the slope of the Phillips curve conditional on forward-looking expectations and h' is the history $\{h, \pi(h, t), y(h, t)\}$.

Values of inflation and output under discretion, time-0 commitment, timeless perspective and dynamic Rawlsian policy are presented in Table 1. We assume a Markov-stationary discretionary equilibrium, meaning that endogenous variables can only be conditioned on the model's state vector and exogenous shocks. Since in this model there are no state variables and shocks, this implies that the discretionary choices $y_d^t(h, t)$ and $\pi_d^t(h, t)$ must be scalars independent of h and t . A standard 'inflation bias' result follows and $\pi_d^t(h, t)$ is suboptimally high. The time-0 commitment choices $y_{c0}^t(h, t)$ and $\pi_{c0}^t(h, t)$ involve a substantial initial inflation that stimulates output closer to its target, followed by a gradual decline in both inflation and output to zero at rate $\phi < 1$.⁷ The timeless perspective implements the long-run commitment outcome immediately, with zero inflation and output in all periods.

⁷It can be shown that $\phi \equiv \frac{1+\beta+\frac{\gamma^2}{\omega} - \sqrt{\left(1+\beta+\frac{\gamma^2}{\omega}\right)^2 - 4\beta}}{2\beta}$.

<i>Policy</i>	<i>Inflation</i>	<i>Output</i>
Discretion	$\frac{\gamma}{1-\beta+\frac{\gamma^2}{\omega}}y^*$	$\frac{1-\beta}{1-\beta+\frac{\gamma^2}{\omega}}y^*$
Commitment	$\frac{\omega}{\gamma}(1-\phi)y^*\phi^t$	$y^*\phi^{t+1}$
Timeless	0	0
Dynamic Rawlsian	$\frac{(1-\beta)\gamma}{(1-\beta)^2+\frac{\gamma^2}{\omega}}y^*$	$\frac{(1-\beta)^2}{(1-\beta)^2+\frac{\gamma^2}{\omega}}y^*$

Table 1: Inflation and output under discretion, commitment, timeless and dynamic Rawlsian policy

Since this is a purely deterministic forward-looking model, the dynamic Rawlsian policy in period $s \geq 0$ defines the time-invariant levels of inflation and output y_R^s and π_R^s that would be best to implement from period s onwards in perpetuity. Consistent with our aim of achieving coherence between policymakers existing in different periods, Table 1 shows that these choices are independent of s . The Rawlsian policymaker selects a positive level of inflation that is nonetheless strictly lower than the discretionary outcome. The policy internalises the reaction of forward-looking expectations to policy, but unlike in the commitment case it is not possible to exploit the fixity of past expectations at time 0 and tailor a distinct policy to each time period. This is because the Rawlsian policymaker has no knowledge of time to exploit.

The dynamic Rawlsian policy has desirable properties when compared with the timeless approach.⁸ Specifically, it dominates the timeless policy in a Pareto sense, because policymakers assessing outcomes under W_s would strictly favour a switch from the timeless perspective to the dynamic Rawlsian policy for all $s \geq 0$. To see this, note that the per-period loss under both approaches will be constant since endogenous variables themselves take constant values. Under the timeless policy this per-period loss is ωy^{*2} because inflation and output are both zero. The corresponding per-period loss under the dynamic Rawlsian policy is:

$$\frac{(1-\beta)^2}{(1-\beta)^2+\frac{\gamma^2}{\omega}}\omega y^{*2} < \omega y^{*2},$$

⁸We do not compare with unconditionally optimal policy, since dynamic Rawlsian policy is unconditionally optimal in purely forward-looking models, as discussed in Section 4.5.

and the timeless policy is dominated. The fundamental time inconsistency problem is that policymakers in different periods disagree about what it is best to do in any given set of circumstances, but this does not preclude agreement about what might be better to do. It is questionable to pursue the timeless approach when a policy exists that is Pareto superior in this way. More generally, it is straightforward to show that our dynamic Rawlsian policy will always be best among the set of time-invariant choices in purely deterministic forward-looking models such as this.

5 General models

The general model features both forward and backward looking constraints. In this case, it is usually desirable to condition policy on the lagged values of predetermined variables, since these affect the constraint set faced by the policymaker. This needs to be done in a manner consistent with our general principle that dynamic Rawlsian policy should be based on information orthogonal to time, at least insofar as knowledge of time would otherwise cause disagreement among policymakers who exist in different time periods.

The main problem in the general model is that the lagged values of predetermined variables are themselves outcomes of purposive policy choice. In principle, a policymaker could always plan to set them in a way that allowed some inference about time for any given initial weighting structure. Unlike the distribution over E^∞ , the distribution over the predetermined variables is not one of the model's fundamentals, so if lagged predetermined variables are to be included in the information set then imposing a particular weighting structure *ex ante* is no longer guaranteed to ensure orthogonality of the information set with respect to time. Orthogonality could always be confounded by subsequent policy choices.

To overcome the problem, we exploit the fact that policymakers in different periods only disagree about the appropriate choices for non-predetermined variables. If a model does not feature non-predetermined variables then there are no expectational constraints, and preferences are coherent in the sense that $x_o^s(h, t) = x_o^r(h, t)$ for all $t \geq r, s \geq 0$. So, if we choose the predetermined variables distinctly from the non-predetermined variables, then the choice of predetermined variables ought not to imply disagreement, provided that the realisations of the non-predetermined variables are treated as exogenous to this choice. Equally, if the non-predetermined variables are chosen under the assumption that the predetermined variables evolve according to an exogenously fixed probability distribution over time, then we

should in principle be able to extend the ‘history weighting’ procedure to choices of the non-predetermined variables that condition on both h_ε and the lagged predetermined variables z . This will ensure that the policymaker cannot infer time when choosing those variables that are subject to a time-inconsistency problem.

We proceed by dividing up the choice process and defining separate problems for choosing the predetermined and non-predetermined variables. The non-predetermined variables will be chosen from the perspective of period s according to a function $y_R^s(h_\varepsilon, z)$, with $y_R^s : E^\infty \times Z \rightarrow Y$. The second argument of the function relates to the realised values of lagged predetermined variables. This choice problem will treat the realisation of the states of the world (h_ε, z) as a stochastic process, and will hold the predetermined variables constant for each (h_ε, z) pair. The predetermined variables for period t will be chosen from the perspective of period s according to a function $z_R^s(h_\varepsilon, z_{s-1}, t)$, with $z_R^s : E^\infty \times Z \times \mathbb{Z} \setminus \{0, \dots, s-1\} \rightarrow Y$. h_ε is the history of exogenous variables up to period t ; z_{s-1} is the lagged vector of predetermined variables at the start of time. This choice will be made subject to a given $y_R^s(h_\varepsilon, z_{s-1}, t)$ function. The realisation of this function at each horizon will be entirely independent of the choice of $z_R^s(h_\varepsilon, z_{s-1}, t)$, ensuring that the basic time-consistency of ‘backward-looking’ choice will not be undermined by an incentive to influence expectations.⁹

5.1 Restricting dependence

The first task is to formally define the set of variables upon which policy can depend. We assume that Conditions 1 and 2 continue to apply, so we can write $y_R^s(h_\varepsilon, h_z)$ as the desired outcome for the non-predetermined variables in period t on the part of a policymaker assessing outcomes in period $s \leq t$. Moreover, only the most recent lag of the endogenous variables affects the constraints of the model, so we can further rule out dependence on irrelevant past choices:

Condition 4 $y_R^s(h, t) = y_R^s(h', t)$ for any histories $h = (h_\varepsilon, h_y, h_z)$ and $h' = (h_\varepsilon, h_y, h'_z)$ such that h_z and h'_z do not differ in their last entry, and all $t \geq s \geq 0$.

⁹This would not be the case if the ‘backward-looking’ problem were defined instead on the same evaluative space as the ‘forward-looking’ problem, i.e. $E^\infty \times Z$. This is because changes to any $z_R^s(h_\varepsilon, z)$ function so defined would change the relative likelihood of future realisations of the (h_ε, z) pair, and thus render more or less likely the responses implied by the associated $y_R^s(h_\varepsilon, z)$. Choosing the predetermined variables would be an indirect way to choose the non-predetermined variables. We need the likelihood of different points in the evaluative space itself to be independent of choice at all horizons.

Together with Conditions 1 and 2, this implies that we can write the function for non-predetermined variables as $y_R^s(h_\varepsilon, z)$.

There is no intrinsic time-inconsistency problem when choosing the predetermined variables, but we need to ensure that this choice treats the non-predetermined variables as exogenous for every stochastic contingency. This means choosing outcomes that are measurable with respect to the exogenous state of the world, i.e., the initial state vector and the evolution of history from the initial period onwards. Realisations of the non-predetermined variables can then be treated as exogenous and measurable with respect to the same space when choosing the predetermined variables. Hence we impose:

Condition 5 $z_R^s(h, t) = z_R^s(h', t)$ for all $t \geq s \geq 0$ and any histories $h = (h_\varepsilon, h_y, h_z)$ and $h' = (h_\varepsilon, h'_y, h'_z)$ such that h_z and h'_z do not differ in their entries corresponding to period $s - 1$.

Condition 5 implies that we can express as $z_R^s(h_\varepsilon, z_{s-1}, t)$ the predetermined variables desired for period t by a Rawlsian policymaker assessing outcomes in period $s \leq t$. Note that we achieve coherence between policymakers in choice over Z if $z_R^s(h_\varepsilon, z_R^r(h_\varepsilon, z_{r-1}, s - 1), t) = z_R^r(h_\varepsilon, z_{r-1}, t)$ for all $r < s \leq t$.

5.2 Linking the choice problems

The proposal is that the non-predetermined and predetermined variables should be determined by separate choice problems. The non-predetermined variables are chosen according to some function $y_R^s(h_\varepsilon, z)$, under the assumption that the predetermined variables evolve according to some exogenous statistical model. The predetermined variables are set for period t according to a function $z_R^s(h_\varepsilon, z_{-1}, t)$, under the assumption that the non-predetermined variables are fixed in advance for each $(h_\varepsilon, z_{-1}, t)$ triple. The two choice problems are subject to a common set of constraints, which restricts the joint choices that can be made. We therefore need an appropriate way of incorporating the common constraints into the choice problems. Our approach adapts the Lagrangian method, by attaching the same state-contingent shadow value to marginal relaxations of a given constraint in both problems, irrespective of whether the choice is for non-predetermined or predetermined variables. The method to do this is specified below, but an important implication is that multipliers must feature in any representation of the problems that determine choice of non-predetermined and predetermined variables. A separate problem can then be defined to characterise these multipliers.

5.3 History weighting

With multipliers duly included, the ‘forward-looking’ policy problem in period s takes the form:

$$\begin{aligned}
& \max_{y(h_\varepsilon, z)} \sum_{t=s}^{\infty} \beta^{t-s} \int_{E^\infty \times Z} \int_{E^\infty \times Z} \{ \pi(x(h_\varepsilon, z), \varepsilon) \\
& + \lambda_g(h_\varepsilon, z) \sum_{r=0}^T \int_{E^\infty \times Z} g_r(x(h'_\varepsilon, z'), \varepsilon') dG_r(h'_\varepsilon, z' | h_\varepsilon, z) \\
& + \lambda_m(h_\varepsilon, z) m(x(h_\varepsilon, z), z, \varepsilon) \} dG_{t-s}(h_\varepsilon, z | h''_\varepsilon, z'') dG(h''_\varepsilon, z'')
\end{aligned} \tag{15}$$

subject to a given function for the predetermined variables $z(h_\varepsilon, z)$, given functions for the Lagrange multipliers $\lambda_g(h_\varepsilon, z)$ and $\lambda_m(h_\varepsilon, z)$, some set of time-specific conditional distributions $G_s(h_\varepsilon, z | h'_\varepsilon, z')$, and a weighting function over the initial information set $G(h_\varepsilon, z)$. We need to assign values to these distribution functions in a manner consistent with our Rawlsian perspective.

In the purely forward-looking case, it was appropriate to use the unconditional distribution of exogenous variables $F(h_\varepsilon)$ as the weighting function over initial histories $G(h_\varepsilon)$, and use the exogenous conditional distributions $F_s(h_\varepsilon | h'_\varepsilon)$ as the time-specific conditional distributions $G_s(h_\varepsilon | h'_\varepsilon)$. Following identical reasoning, we retain these aspects to weight initial histories of the exogenous variables in the general case when policy can also depend on predetermined variables. It remains to characterise the weighting function over the initial values of the predetermined variables $G(z | h_\varepsilon)$, and the associated conditional distributions $G_s(z | h_\varepsilon, h'_\varepsilon, z')$. This weighting function should ideally satisfy:

1. **Stationarity:** Any given (h_ε, z) pair should be considered equally likely to occur in all periods, so that the ‘forward-looking’ policymaker is unable to infer anything about time from knowledge of (h_ε, z) .
2. **Agreement:** The function $z(h_\varepsilon, z)$ should agree with the choice of conditional distributions, in the sense that $G_s(z' | h'_\varepsilon, h_\varepsilon, z)$ places positive probability mass only on those combinations of predetermined variables that can occur after applying $z(h_\varepsilon, z)$ for s successive time periods, from an initial history (h_ε, z) with shocks evolving according to h'_ε .
3. **Full support:** The weighting function $G(h_\varepsilon, z)$ should have full support on the space $E^\infty \times Z$, otherwise choice in some regions of that space could be made completely

arbitrarily, at no cost under the assumed objective.

These requirements are generally mutually inconsistent. For example, if a policy induces convergence to a stochastic steady state then we can have stationarity and agreement, but not full support. To see this, note that the only weighting function consistent with stationarity and agreement is one that places positive probability weight only on initial histories (h_ε, z) that are consistent with the economy being in stochastic steady state. This means there is a unique z for every h_ε . The absence of probability weight on other initial histories violates the requirement of full support. If full support is required, there is a problem in that it is more likely a policymaker will observe predetermined variables away from steady state the earlier is the time period. The only situation in which this does not arise is when all variables are assumed to start at the steady state with certainty, an assumption that is not very useful when formulating appropriate policy for the transition to steady state.

We can only satisfy two of the three requirements. Stationarity is essential, given our aim is to deny policymakers knowledge of time and ensure that time is orthogonal to the policymaker's information set. As we want our method to be useful for policy purposes, we also want to satisfy full support and derive optimal strategies for the complete space $E^\infty \times Z$. We therefore drop the requirement for agreement. This implies a forced separation when choosing the non-predetermined variables, in that the stochastic model assumed to govern the evolution of lagged predetermined variables will not typically be in agreement with the contemporaneous response of predetermined variables implied by the policy function $z(h_\varepsilon, z)$. Condition 6 summarises the restrictions we place on the weighting function and associated conditional distributions. It implies Condition 3 as a special case when the model is purely forward-looking.

Condition 6 *When choosing $y_R^s(h_\varepsilon, z)$, the policymaker in period $s \geq 0$ uses a ‘weighting function’ G and associated conditional distributions G_{t-s} that satisfy:*

$$\begin{aligned} & \int_{E^\infty \times Z} \int_{E^\infty \times Z} v(h_\varepsilon, z) dG_{t-s}(h_\varepsilon, z | h'_\varepsilon, z') dG(h'_\varepsilon, z') \\ &= \int_{E^\infty \times Z} \int_{E^\infty \times Z} v(h_\varepsilon, z) dG_{r-s}(h_\varepsilon, z | h'_\varepsilon, z') dG(h'_\varepsilon, z'), \end{aligned} \quad (16)$$

for all $t, r \geq s$ and all functions $v : E^\infty \times Z \rightarrow \mathbb{R}$, when assessing current and future policy choices for the non-predetermined variables. When considering the marginal components of these distributions over E^∞ alone, they apply the weighting function $G(A) = F(A)$ for all

$A \subseteq E^\infty$, and the conditional distributions $G_{t-s}(A|h_\varepsilon, z) = F_{t-s}(A|h_\varepsilon)$ for all $t \geq s$, all $(h_\varepsilon, z) \in E^\infty \times Z$ and all $A \subseteq E^\infty$.

5.4 Choice problems

The policy problem (15) under Condition 6 can be rewritten as:

$$\begin{aligned} & \max_{y(h_\varepsilon, x^b)} (1 - \beta)^{-1} \int_{E^\infty \times Z} \{ \pi(x(h_\varepsilon, z), \varepsilon) \\ & + \sum_{r=0}^T \lambda_g^r(h_\varepsilon, z) g_r(x(h_\varepsilon, z), \varepsilon) \\ & + \lambda_m(h_\varepsilon, z) m(x(h_\varepsilon, z), z, \varepsilon) \} dG(h_\varepsilon, z), \end{aligned} \quad (17)$$

where we define:

$$\lambda_g^r(h_\varepsilon, z) \equiv \int_Z \lambda_g(h_\varepsilon \setminus r, z') dG_{-r}(z' | h_\varepsilon \setminus r, h_\varepsilon, z).$$

The conditional distribution $G_{-r}(z' | h_\varepsilon \setminus r, h_\varepsilon, z)$ is the probability that the predetermined variables r periods ago had values z' , given contemporary values for the lagged predetermined variables z and a current shock history h_ε , under distribution functions that are consistent with Condition 6.

The Lagrange multiplier function $\lambda_g^r(h_\varepsilon, z)$ is the aggregate shadow value of relaxing expectational constraints across all possible prior states of the world. We ensure consistency in policymaking by using the same function in both the policymaker's forward and backward looking problems. By seeking consistency only with respect to this aggregate multiplier function, we have that optimal choice under (17) is invariant to the choice of G .¹⁰ The object under the integral can be maximised piecewise for each (h_ε, z) , irrespective of the precise weighting function applied. In this way, we prevent policy from being affected by our need to specify G_s functions that are not necessarily in agreement with $z(h_\varepsilon, z)$. This would not be the case if consistency were to be required with respect to the individual Lagrange multiplier functions $\lambda_g(h_\varepsilon \setminus r, z')$, holding $G_{-r}(z' | h_\varepsilon \setminus s, h_\varepsilon, z)$ fixed.

The choice of predetermined variables is not affected by a time-inconsistency problem, so there is no need for further informational constraints on the 'backward-looking' problem. The

¹⁰This is true provided the G function has full support on $E^\infty \times Z$.

problem assessed in period s is:

$$\begin{aligned}
& \max_{z(h_\varepsilon, z_{s-1}, t)} \sum_{t=s}^{\infty} \beta^{t-s} \int_{E^\infty} \{ \pi(x(h_\varepsilon, z_{s-1}, t), \varepsilon) \\
& + \lambda_g(h_\varepsilon, z_{s-1}, t) \sum_{r=0}^T \int_{E^\infty} g_r(x(h'_\varepsilon, z_{s-1}, t+r), \varepsilon') dF_r(h'_\varepsilon | h_\varepsilon) \\
& + \lambda_m(h_\varepsilon, z_{s-1}, t) m(x(h_\varepsilon, z_{s-1}, t), z(h_\varepsilon \setminus 1, z_{s-1}, t-1), \varepsilon) \} dF_{t-s}(h_\varepsilon | h''_\varepsilon), \quad (18)
\end{aligned}$$

subject to initial conditions $(h''_\varepsilon, z_{s-1})$, a given function for the non-predetermined variables $y(h_\varepsilon, z_{s-1}, t)$, given Lagrange multiplier functions $\lambda_g(h_\varepsilon, z_{s-1}, t)$ and $\lambda_m(h_\varepsilon, z_{s-1}, t)$, and the exogenous time-specific conditional distribution over E^∞ , $F_t(h_\varepsilon | h'_\varepsilon)$.

The Lagrange multiplier functions need defining in a way that guarantees the constraints of the model will always be satisfied. We state the problem these multipliers solve using their representation in the ‘backward-looking’ policymaker’s problem, i.e., as functions of the triple $(h_\varepsilon, z_{s-1}, t)$. Their representation as functions of (h_ε, z) in the ‘forward-looking’ problem then follows by an appropriate mapping. The Lagrange multiplier functions $\lambda_g(h_\varepsilon, z_{s-1}, t)$ and $\lambda_m(h_\varepsilon, z_{s-1}, t)$ solve the following problems:

$$\min_{\lambda_g(h_\varepsilon, z_{s-1}, t) \geq 0} \lambda_g(h_\varepsilon, z_{s-1}, t) \sum_{r=0}^T \int_{E^\infty} g_r(x(h'_\varepsilon, z_{s-1}, t+r), \varepsilon') dF_r(h'_\varepsilon | h_\varepsilon), \quad (19)$$

$$\min_{\lambda_m(h_\varepsilon, z_{s-1}, t) \geq 0} \lambda_m(h_\varepsilon, z_{s-1}, t) m(x(h_\varepsilon, z_{s-1}, t), z(h_\varepsilon \setminus 1, z_{s-1}, t-1), \varepsilon). \quad (20)$$

5.5 Defining a solution

Definition 7 *Dynamic Rawlsian policy in period s is characterised by the functions: $y(h_\varepsilon, z_{s-1}, t)$, $y(h_\varepsilon, z)$, $z(h_\varepsilon, z_{s-1}, t)$, $z(h_\varepsilon, z)$, $\lambda_m(h_\varepsilon, z_{s-1}, t)$, $\lambda_m(h_\varepsilon, z)$, $\lambda_g(h_\varepsilon, z_{s-1}, t)$ and $\{\lambda_g^r(h_\varepsilon, z)\}_{r=0}^T$, if and only if:*

1. (a) $y(h_\varepsilon, z)$ solves (17), given $z(h_\varepsilon, z)$, $\lambda_m(h_\varepsilon, z)$ and $\{\lambda_g^r(h_\varepsilon, z)\}_{r=0}^T$.
(b) $z(h_\varepsilon, z_{s-1}, t)$ solves (18), given $y(h_\varepsilon, z_{s-1}, t)$, $\lambda_m(h_\varepsilon, z_{s-1}, t)$ and $\lambda_g(h_\varepsilon, z_{s-1}, t)$.
(c) $\lambda_g(h_\varepsilon, z_{s-1}, t)$ solves (19), given $y(h_\varepsilon, z_{s-1}, t)$ and $z(h_\varepsilon, z_{s-1}, t)$.
(d) $\lambda_m(h_\varepsilon, z_{s-1}, t)$ solves (20), given $y(h_\varepsilon, z_{s-1}, t)$ and $z(h_\varepsilon, z_{s-1}, t)$.
2. The following equivalences hold for all $(h_\varepsilon, z_{s-1}, t) \in E^\infty \times Z \times \mathbb{Z} \setminus \{0, \dots, s-1\}$ and all $r \in \{0, \dots, \min\{t-s, T\}\}$.¹¹

¹¹We normalise $z(h_\varepsilon \setminus 1, z_{s-1}, s-1) \equiv z_{s-1}$ where appropriate.

- (a) $y(h_\varepsilon, z_{s-1}, t) = y(h_\varepsilon, z(h_\varepsilon \setminus 1, z_{s-1}, t-1))$.
- (b) $z(h_\varepsilon, z_{s-1}, t) = z(h_\varepsilon, z(h_\varepsilon \setminus 1, z_{s-1}, t-1))$.
- (c) $\lambda_g^r(h_\varepsilon, z) = \int_Z \lambda_g(h_\varepsilon \setminus r, z_{s-1}, t) dG_{-r}(z(h_\varepsilon \setminus (r+1), z_{s-1}, t-1) | h_\varepsilon \setminus r, h_\varepsilon, z)$ for some distribution function $G_{-r}(z' | h_\varepsilon \setminus r, h_\varepsilon, z)$ that agrees with $z(h_\varepsilon, z)$.
- (d) $\lambda_m(h_\varepsilon, z_{s-1}, t) = \lambda_m(h_\varepsilon, z(h_\varepsilon \setminus 1, z_{s-1}, t-1))$.

Part 1 states the problems that the policy and multiplier functions solve, whilst Part 2 ensures consistency across the different representations of these functions in the different policy problems. 2(c) is the most complex element. As discussed in Section 5.4, it states that the aggregate shadow values associated with relaxing past expectational constraints must be equal when choosing both the non-predetermined and predetermined variables. The aggregate in the ‘backward-looking’ problem can be established under a model-consistent distribution function without jeopardising our Rawlsian perspective, because there is no need to make time orthogonal to the information set when solving for the predetermined variables. Stationarity in particular is not an issue as it is imposed when solving for non-predetermined variables solely to prevent inference about time. The conditional distribution function $G_{-r}(z' | h_\varepsilon \setminus r, h_\varepsilon, z)$ provides an appropriate weighting over multipliers featuring in the ‘backward-looking’ problem, which can be equated to the corresponding aggregate in the ‘forward-looking’ problem.

5.6 Dynamic Rawlsian policy

The detailed derivation of dynamic Rawlsian policy in general models is presented in Appendix B. A consolidated first order condition characterises dynamic Rawlsian policy as the desired outcome for a generic non-predetermined or predetermined variable x_i , as assessed by the Rawlsian policymaker in any time period for any time period:

$$\begin{aligned}
& \frac{\partial \pi(x(h_\varepsilon, z), \varepsilon)}{\partial x_i(h_\varepsilon, z)} + \sum_{r=0}^T \lambda_g^r(h_\varepsilon, z) \frac{\partial g_r(x(h_\varepsilon, z), \varepsilon)}{\partial x_i(h_\varepsilon, z)} \\
& + \lambda_m(h_\varepsilon, z) \frac{\partial m(x(h_\varepsilon, z), z, \varepsilon)}{\partial x_i(h_\varepsilon, z)} \\
& + \beta \int_{E^\infty} \lambda_m(h'_\varepsilon, z(h_\varepsilon, z)) \frac{\partial m(x(h'_\varepsilon, z(h_\varepsilon, z)), z(h_\varepsilon, z), \varepsilon')}{\partial x_i(h_\varepsilon, z)} dF_1(h'_\varepsilon | h_\varepsilon) \\
& = 0.
\end{aligned} \tag{21}$$

The first order condition is time-invariant, reflecting the orthogonality of time imposed when choosing the non-predetermined variables and the time-consistency of choice for the predeter-

mined variables. This implies $x_R^s(h, t) = x_R^r(h, t)$ for all $s, r \leq t$, and all policymakers agree on the appropriate policy for each period.

There is asymmetry in the discounting structure of the first order condition for dynamic Rawlsian policy. There is no inverse discounting associated with forward-looking constraints, for the same intuitive reasons as were discussed in the purely forward-looking case. With no knowledge about when a particular history will occur, the policymaker compares the direct marginal benefit of changing variables at date t to the shadow values of changing variables at some future date $t + s$. There is, however, discounting associated with the backward-looking constraints. Policymakers in all time periods know that changes to the predetermined variables cannot have a contemporaneous effect on past values of the predetermined variables; they can only influence the second argument of the m function with a lag. This contrasts with changes to expected future policy, which influence expectational constraints immediately. Discounting is therefore retained when determining optimal choices for the predetermined variables, but not when choosing the non-predetermined variables. This treatment of backward-looking constraints is desirable, because it ensures that dynamic Rawlsian policy coincides with standard optimal policy in purely backward-looking models, for which no time-inconsistency problem arises.

5.7 Comparison to other approaches

The dynamic Rawlsian policy is unique in its asymmetric treatment of forward and backward looking constraints. The commitment and timeless perspective approaches discount both backward and forward looking constraints, consistent with the first order optimality condition (4). The unconditional optimality approach discounts neither, because of its exclusive focus on the long run stochastic steady state. More generally, the dynamic Rawlsian policy is the only approach to satisfy the following four properties:

1. Time-invariance: $x_R^s(h, t)$ is independent of t for all $t \geq s$.
2. Position-invariance: $x_R^s(h, t)$ is independent of s for all $s \leq t$.
3. Optimal time-invariant outcome in deterministic purely forward-looking models: $x_R^s(h, t)$ always takes the best constant value in this setting.
4. Appropriate limit in purely backward-looking models: $x_R^s(h, t)$ always takes the time-consistent optimal value in this setting.

Of the other policies considered in this paper, the time 0 optimal policy $x_o^s(h, t)$ only satisfies property 4, the ‘contractual’ time 0 commitment policy satisfies properties 2 and 4, the timeless perspective satisfies properties 1, 2 and 4, and unconditionally optimality satisfies properties 1, 2 and 3. Discretionary policy $x_d^s(h, t)$ is not well defined when $s \neq t$, and may not be uniquely defined even when $s = t$. Nonetheless, whilst discretionary policy will satisfy property 4, it almost always violates property 3. There is nothing uniquely desirable about the ability of dynamic Rawlsian policy to satisfy these four properties, but we consider it a virtue that it does. In addition, the next example shows that there are important limits to the implementability of both the timeless perspective and unconditionally optimal approaches. Broader applicability is a further benefit of our approach.

5.8 Example

We illustrate the properties of dynamic Rawlsian policy in the general case using a simple model of redistributive social insurance. The problem is adapted from Marcet and Marimon (1992) and Kocherlakota (1996), and involves a utilitarian policymaker redistributing consumption goods across agents. The policymaker is subject to ‘forward-looking’ constraints, in that agents will opt out of the insurance scheme if at any time they consider themselves better off leaving and living in autarky for the rest of their lives. The ‘backward-looking’ constraints arise because the policymaker can save and accumulate resources, subject to an intertemporal budget constraint.

There are N infinitely-lived agents in the economy, each endowed with a stochastic income stream through time. The income of agent n in period t is $y^n(t) \in \Upsilon$. A social planner has weighted-utilitarian preferences across the welfare of these agents, with an objective at time s given by:

$$W_s = \sum_{n=1}^N \alpha^n U_s^n, \quad (22)$$

where U_s^n , the welfare of agent n in period s , is a standard function of agent n ’s consumption:

$$U_s^n = \sum_{t=s}^{\infty} \beta^{t-s} \int_{H^\infty} u(c^n(h, t)) dF_{t-s}(h|h'). \quad (23)$$

$c^n : H^\infty \times \mathbb{Z} \rightarrow \mathbb{R}_+$ is the consumption of agent n at time t following history h ; $u : \mathbb{R}_+ \rightarrow \mathbb{R}$ is a within-period utility function that is continuously differentiable and concave.

The planner must allocate income to the different agents, subject to the ‘forward-looking’ constraints that agents have the option of permanently leaving the insurance scheme and

consuming under autarky in perpetuity. To prevent agents leaving, the planner must ensure that every agent in every period anticipates a higher utility from staying than leaving. The participation constraint for each n and all h, t is thus:

$$U_t^n \geq \sum_{r=t}^{\infty} \beta^{r-t} \int_{\mathcal{Y}} u(y^n(t)) dF_{r-t}(y^n(t)|h). \quad (24)$$

This implies time inconsistency, because the best level of utility to promise agent n at time $r > t$ to prevent them leaving is likely to be different when comparing assessments made in periods r and t . By the time period r arrives, the t -dated participation constraint is no longer a concern and the policymaker faces a different set of constraints.

The planner allocating income to different agents is allowed to save and accumulate resources by purchasing real bonds, subject to a ‘backward-looking’ aggregate intertemporal budget constraint for all t :

$$\sum_{n=1}^N y^n(t) + B(h \setminus 1, t-1) \geq \sum_{n=1}^N c^n(h, t) + R^{-1} B(h, t). \quad (25)$$

$R \in \mathbb{R}_{++}$ is an exogenous, constant real interest rate, and $B : H^\infty \times \mathbb{Z} \rightarrow \mathbb{R}$ gives the quantity of bonds purchased for the next period as a function of history and time. An additional transversality condition rules out Ponzi schemes and ensures dynamic solvency.

5.8.1 Commitment policy

The character of optimal commitment policy is well known in this model.¹² The ratio of marginal utilities between agents m and n evolves over time according to:

$$\frac{u'(c^n(h, t))}{u'(c^m(h, t))} = \frac{\alpha^m + \mu^m(h, t)}{\alpha^n + \mu^n(h, t)}, \quad (26)$$

where the agent-specific cumulative multiplier $\mu^n(h, t)$ is defined as:

$$\mu^n(h, t) \equiv \sum_{r=0}^{t-s} \lambda_P^n(h \setminus r, t-r), \quad (27)$$

and $\lambda_P^n(h, t) \geq 0$ is the multiplier on constraint (24) in period t . The cumulative multiplier has a recursive representation for $t > s$:

$$\mu^n(h, t) \equiv \mu^n(h \setminus 1, t-1) + \lambda_P^n(h, t) \quad (28)$$

¹²See Chapter 19 of Ljungqvist and Sargent (2004) for a textbook treatment.

The term $\alpha^n + \mu^n(h, t)$ corresponds to agent n 's 'Pareto weight' in the within-period allocation problem. An important feature of the commitment solution is that these weights are non-stationary and non-decreasing. To see this, suppose the solution involves participation constraints binding with non-zero probability for all agents in all periods.¹³ As time progresses, the policymaker's underlying preference weights $\{\alpha^n\}_{n=1}^N$ then have less and less impact on within-period distributions, which instead are dominated by a need to make good on past utility promises. This 'tyranny of the past' is a direct result of the time preference weighting that the commitment solution applies.

A further implication of the commitment solution is that the policymaker over-saves relative to the first best. The commitment policy implies:

$$u'(c^n(h, t)) = \beta R \int_{H^\infty} \frac{\alpha^n + \mu^n(h', t+1)}{\alpha^n + \mu^n(h, t)} u'(c^n(h', t+1)) dF_1(h'|h). \quad (29)$$

Since $\alpha^n + \mu^n(h', t+1) > \alpha^n + \mu^n(h, t)$ in all $t+1$ histories for which the participation constraint binds, we have that:

$$u'(c^n(h, t)) > \beta R \int_{H^\infty} u'(c^n(h', t+1)) dF_1(h'|h), \quad (30)$$

and the marginal utility of consumption at time t is sub-optimally high whenever there is some risk of the participation constraint for agent n binding in period $t+1$. This implies that the policymaker is over-saving.

5.8.2 Dynamic Rawlsian policy

We solve for the dynamic Rawlsian policy in Appendix C. As in the commitment case, the ratio of marginal utilities between agents must equal the ratio of augmented Pareto weights:

$$\frac{u'(c^n(h_\varepsilon, B))}{u'(c^m(h_\varepsilon, B))} = \frac{\alpha^m + \mu^m(h_\varepsilon, B)}{\alpha^n + \mu^n(h_\varepsilon, B)}, \quad (31)$$

where the notation reflects the restriction of dependence to the history of the income vector across agents, the history of real interest rates, and the inherited stock of real bonds B .

¹³This can only be the case if the exogenous real interest rate satisfies $R < \beta^{-1}$. In models where $R = \beta^{-1}$ and Υ has a maximal element that is drawn with non-zero probability, there is no need for further adjustments to the Pareto weight of an agent once they draw the maximal income. Ljungqvist and Sargent (2004) give more details.

However, the equation defining the agent-specific cumulative multiplier μ^n is now:

$$\mu^n(h_\varepsilon, B) = \sum_{r=0}^{\infty} \beta^r \lambda_P^{n,r}(h_\varepsilon, B), \quad (32)$$

which has the recursive representation:

$$\mu^n(h_\varepsilon, B) \equiv \beta \mu^n(h_\varepsilon \setminus 1, B^{-1}(h_\varepsilon, B)) + \lambda_P^{n,0}(h_\varepsilon, B), \quad (33)$$

with $B^{-1}(h_\varepsilon, \cdot)$ the inverse of the policy function for bonds $B(h_\varepsilon \setminus 1, \cdot)$.¹⁴ The agent-specific cumulative multiplier $\mu^n(h_\varepsilon \setminus 1, B^{-1}(h_\varepsilon, B))$ is the counterpart to $\mu^n(h \setminus 1, t-1)$ in (28), a lagged value attached to the shadow cost of varying expectational constraints under predecessor histories.

The important difference between dynamic Rawlsian and commitment policies is the presence of β in the recursive representation (33). This means that the Pareto weights are stationary in the corresponding within-period problem. Any adjustments to the Pareto weights made to incentivise an agent to remain in the insurance scheme will therefore decay over time. The Pareto weights retain a non-negligible link to the ‘fundamental’ weights $\{\alpha^n\}_{n=1}^N$, regardless of the number of times an individual has seen their participation constraint bind.

The first-order condition for bond holdings under the dynamic Rawlsian policy is:

$$u'(c^n(h_\varepsilon, B)) = \beta R \int_{E^\infty} \frac{\alpha^n + \mu^n(h'_\varepsilon, B(h_\varepsilon, B))}{\alpha^n + \mu^n(h_\varepsilon, B)} u'(c^n(h'_\varepsilon, B(h_\varepsilon, B))) dF_1(h'_\varepsilon | h_\varepsilon). \quad (34)$$

Although the first order condition is identical in form to that under commitment (29), the μ^n weights now decline through time and we can no longer assert that $\alpha^n + \mu^n(h'_\varepsilon, B(h_\varepsilon, B)) \geq \alpha^n + \mu^n(h_\varepsilon, B)$. This means it is quite possible that the policymaker will induce a consumption path that is relatively ‘front-loaded’. We also have that bond holdings will be constant whenever $R = \beta^{-1}$ and the population of agents is sufficiently large to make aggregate risk negligible. This contrast with the commitment policy, where the policymaker always has an incentive to accumulate additional bonds.

5.8.3 Alternative approaches

The timeless policy is not defined in this example. The problem is that it requires policy to maximise the objective of a policymaker who existed in the infinite past, but a policymaker

¹⁴The inverse policy function satisfies $B(h_\varepsilon \setminus 1, B^{-1}(h_\varepsilon, B')) = B'$ for all $(h_\varepsilon, B') \in E^\infty \times \mathbb{R}$. That is, $B^{-1}(h_\varepsilon, B)$ gives the prior stock of bonds required for B to be chosen when the shock history is $h_\varepsilon \setminus 1$. To ease notation, we assume this inverse is uniquely defined.

from the infinite past has had more time to accumulate bonds than the current policymaker. The bond holdings of the current policymaker will generally not be sufficient to implement optimal policy from the timeless perspective; current policy is constrained by not yet having done what was best from the perspective of the infinite past. For instance, in the simple case where $R = \beta^{-1}$ and Υ has a maximal element drawn with non-zero probability by a population large enough to make aggregate risk negligible, the commitment policy assigns constant high consumption in perpetuity to any agent who draws the maximal income.¹⁵ If we consider the infinite past then all agents will have drawn the maximal income at some time with probability one, so optimality from the timeless perspective requires the policymaker to give constant high consumption in perpetuity to all agents. This will generally not be feasible with the bond holdings of the current policymaker.

The unconditionally optimal policy is also not defined, because it is always possible to increase steady-state welfare by deferring more current consumption when $R > 1$. There is no ‘optimal’ steady-state as additional savings always improve welfare. This argument results from the linear savings technology, and applies irrespective of whether the model features participation constraints.

6 Conclusion

Any attempt to improve on discretionary equilibria under time inconsistency must impose changes on the choice procedures of at least some policymakers. In this paper we set out a novel method for doing so. We adopt a Rawlsian perspective, arguing that disagreements between policymakers are best addressed by considering what all parties would prefer were they forced to choose from behind a veil of ignorance. This veil is intended to deny each policymaker knowledge of any particular circumstances to which they might otherwise be tempted to tailor policy, in our case the time period associated with any given choice. We outline a set of conditions that capture the idea of a veil of ignorance, ultimately ensuring that choices over non-predetermined variables are made with reference to an information set that is orthogonal to time. We derive our dynamic Rawlsian policy in a general setting and for two specific examples, showing that it has a number of desirable properties relative to alternative approaches in the literature.

¹⁵Ljungqvist and Sargent (2004) provide a detailed treatment.

References

- [1] Blake, A. P. and T. Kirsanova (2012), ‘Discretionary Policy and Multiple Equilibria in LQ RE Models’, *Review of Economic Studies*, forthcoming
- [2] Damjanovic, T., V. Damjanovic, and C. Nolan (2008), ‘Unconditionally optimal monetary policy’, *Journal of Monetary Economics*, 55, 491–500.
- [3] Giannoni, M. P., and M. Woodford (2002), ‘Optimal Interest-Rate Rules 1: General Theory’, NBER Working Paper No. 9419.
- [4] Ireland, P. N. (1997), ‘Sustainable Monetary Policies’, *Journal of Economic Dynamics and Control*, 22, 87–108.
- [5] Kocherlakota, N. R. (1996), ‘Implications of Efficient Risk Sharing without Commitment’, *Review of Economic Studies*, 63, 595–609.
- [6] Kydland, F. E. and E. C. Prescott (1977), ‘Rules rather than discretion: the inconsistency of optimal plans’, *Journal of Political Economy*, 85, 473–491.
- [7] Ljungqvist, L. and T. J. Sargent (2004), *Recursive Macroeconomic Theory (Second Edition)*, MIT Press
- [8] Marcet, A. and R. Marimon (1992), ‘Communication, Commitment and Growth’, *Journal of Economic Theory*, 58, 219–249
- [9] Marcet, A. and R. Marimon (1998), ‘Recursive Contracts’, CEP Discussion Paper No 1055, London School of Economics
- [10] Rawls, J. (1971), *A Theory of Justice*, Oxford: Clarendon Press.
- [11] Svensson, L. E. O. (1999), ‘How Should Monetary Policy Be Conducted In An Era Of Price Stability?’, CEPR Discussion Paper 2342.
- [12] Taylor, J. B. (1979), ‘Estimation and Control of a Macroeconomic Model with Rational Expectations’, *Econometrica*, 47, 1267–1286.
- [13] Woodford, M. (2003), *Interest and Prices: Foundations of a Theory of Monetary Policy*, Princeton: Princeton University Press.

A Optimality conditions in purely forward-looking models

The policymaker in purely forward-looking models maximises the objective in (9), subject to forward-looking constraints of the form (2). Assuming the integrals and sum in (9) are well-behaved, the objective can be written as:

$$W_s^R = \int_{E^\infty} \sum_{t=s}^{\infty} \beta^{t-s} \pi(y(h_\varepsilon), \varepsilon) dF(h_\varepsilon).$$

The forward-looking constraints must bind following every history $h_\varepsilon \in E^\infty$. We denote the Lagrange multipliers on these constraints $\lambda_g(h_\varepsilon)$. Maximising piecewise under the integral, the first order condition for optimality with respect to the non-predetermined variable y_i is:

$$\sum_{t=s}^{\infty} \left(\beta^{t-s} \frac{\partial \pi(y(h_\varepsilon), \varepsilon)}{\partial y_i(h_\varepsilon)} + \sum_{r=0}^{\min\{t-s, T\}} \beta^{t-s-r} \lambda_g(h_\varepsilon \setminus r) \frac{\partial g_r(y(h_\varepsilon), \varepsilon)}{\partial y_i(h_\varepsilon)} \right) = 0.$$

Together with the constraint set, this is sufficient to characterise the solution to the policy problem. It follows by collecting terms in β^t that dynamic Rawlsian policy is characterised by the condition:

$$\frac{1}{1-\beta} \frac{\partial \pi(y(h_\varepsilon), \varepsilon)}{\partial y_i(h_\varepsilon)} = - \frac{1}{1-\beta} \sum_{r=0}^T \lambda_g(h_\varepsilon \setminus r) \frac{\partial g_r(y(h_\varepsilon), \varepsilon)}{\partial y_i(h_\varepsilon)},$$

and the terms in β cancel.

B Optimality conditions in general models

We solve for dynamic Rawlsian policy in the general model by taking the first order conditions to each problem separately, before combining them into a set of consolidated first order conditions. The optimal set of non-predetermined variables for a Rawlsian policymaker in period s solves (17) subject to given $z(h_\varepsilon, z)$, $\lambda_m(h_\varepsilon, z)$ and $\{\lambda_g^r(h_\varepsilon, z)\}_{r=0}^T$ functions. The associated first-order conditions are:

$$\begin{aligned} & \frac{\partial \pi(x(h_\varepsilon, z), \varepsilon)}{\partial y_i(h_\varepsilon, z)} + \sum_{r=0}^T \lambda_g^r(h_\varepsilon, z) \frac{\partial g_r(x(h_\varepsilon, z), \varepsilon)}{\partial y_i(h_\varepsilon, z)} \\ & + \lambda_m(h_\varepsilon, z) \frac{\partial m(x(h_\varepsilon, z), z, \varepsilon)}{\partial y_i(h_\varepsilon, z)} \\ & = 0. \end{aligned} \tag{B.1}$$

This holds for G -almost all $(h_\varepsilon, z) \in E^\infty \times Z$, where $G(h_\varepsilon, z)$ is the stationary distribution in (17).

The optimal set of predetermined variables for a Rawlsian policymaker in period s solves (18) subject to given $y(h_\varepsilon, z_{s-1}, t)$, $\lambda_m(h_\varepsilon, z_{s-1}, t)$ and $\lambda_g(h_\varepsilon, z_{s-1}, t)$ functions, with first order conditions:

$$\begin{aligned}
& \frac{\partial \pi(x(h_\varepsilon, z_{s-1}, t), \varepsilon)}{\partial z_i(h_\varepsilon, z_{s-1}, t)} + \lambda_g(h_\varepsilon, z_{s-1}, t) \frac{\partial g_0(x(h_\varepsilon, z_{s-1}, t), \varepsilon)}{\partial z_i(h_\varepsilon, z_{s-1}, t)} \\
& + \lambda_m(h_\varepsilon, z_{s-1}, t) \frac{\partial m(x(h_\varepsilon, z_{s-1}, t), z(h_\varepsilon \setminus 1, z_{s-1}, t-1), \varepsilon)}{\partial z_i(h_\varepsilon, z_{s-1}, t)} \\
& + \beta \int_{E^\infty} \lambda_m(h'_\varepsilon, z_{s-1}, t+1) \frac{\partial m(x(h'_\varepsilon, z_{s-1}, t+1), z(h_\varepsilon, z_{s-1}, t), \varepsilon')}{\partial z_i(h_\varepsilon, z_{s-1}, t)} dF_1(h'_\varepsilon | h_\varepsilon) \\
& = 0.
\end{aligned} \tag{B.2}$$

This holds for all $t \geq s$, and F_{t-s} -almost all $h_\varepsilon \in E^\infty$, given an initial (h_ε, z_{s-1}) pair. The mappings in Part 2 of the veil of ignorance solution definition allow this to be rewritten as:

$$\begin{aligned}
& \frac{\partial \pi(x(h_\varepsilon, z), \varepsilon)}{\partial z_i(h_\varepsilon, z)} + \lambda_g^0(h_\varepsilon, z) \frac{\partial g_0(x(h_\varepsilon, z), \varepsilon)}{\partial z_i(h_\varepsilon, z)} \\
& + \lambda_m(h_\varepsilon, z) \frac{\partial m(x(h_\varepsilon, z), z, \varepsilon)}{\partial z_i(h_\varepsilon, z)} \\
& + \beta \int_{E^\infty} \lambda_m(h'_\varepsilon, z(h_\varepsilon, z)) \frac{\partial m(x(h'_\varepsilon, z(h_\varepsilon, z)), z(h_\varepsilon, z), \varepsilon')}{\partial z_i(h_\varepsilon, z)} dF_1(h'_\varepsilon | h_\varepsilon) \\
& = 0.
\end{aligned} \tag{B.2a}$$

where $z = z(h_\varepsilon \setminus 1, z_{s-1}, t-1)$. This provides the desired time-invariant representation.

The characterisation of dynamic Rawlsian policy is completed by complementary slackness conditions for the Lagrange multipliers:

$$\lambda_g(h_\varepsilon, z_{s-1}, t) \sum_{r=0}^T \int_{E^\infty} g_r(x(h'_\varepsilon, z_{s-1}, t+r), \varepsilon') dF_r(h'_\varepsilon | h_\varepsilon) = 0, \tag{B.3}$$

$$\lambda_m(h_\varepsilon, z_{s-1}, t) m(x(h_\varepsilon, z_{s-1}, t), z(h_\varepsilon \setminus 1, z_{s-1}, t-1), \varepsilon) = 0. \tag{B.4}$$

These map straightforwardly to the corresponding functions on the space $E^\infty \times Z$, as outlined in the solution definition. Note that (B.3) ensures the forward-looking constraints will be satisfied under model-consistent expectations, so the potential deviation from these when

deriving appropriate forward-looking policy has not affected the validity of the Lagrangean approach.

Conditions (B.1) and (B.2a) are consolidated in the first order conditions (21) of the dynamic Rawlsian policy in the main text.

C Rawlsian redistribution policy with imperfect commitment

The ‘forward-looking’ problem can be written as:

$$\begin{aligned} & \max_{\{c^n(h_\varepsilon, B)\}_{n=1}^N} (1 - \beta)^{-1} \int_{E^\infty \times \mathbb{R}} \left\{ \sum_{n=1}^N \alpha^n u(c^n(h_\varepsilon, B)) \right. \\ & + \sum_{n=1}^N \sum_{r=0}^{\infty} \lambda_P^{n,r}(h_\varepsilon, B) \beta^r [u(c^n(h_\varepsilon, B)) - u(y^n)] \\ & \left. + \lambda_R(h_\varepsilon, B) \left[\sum_{n=1}^N y^n(t) + B - \sum_{n=1}^N c^n(h_\varepsilon, B) - R^{-1} B(h_\varepsilon, B) \right] \right\} dG(h_\varepsilon, B). \end{aligned} \quad (35)$$

This is solved treating the functions $B(h_\varepsilon, B)$, $\lambda_P^{n,s}(h_\varepsilon, B)$ and $\lambda_R(h_\varepsilon, B)$ as fixed, and for a distribution G with full support on $E^\infty \times \mathbb{R}$. Note that the coefficient β^s in the second line here comes directly from the participation constraint.¹⁶ The ‘backward-looking’ problem is:

$$\begin{aligned} & \max_{B(h_\varepsilon, B_{s-1}, t)} \sum_{t=s}^{\infty} \beta^{t-s} \int_{E^\infty} \left\{ \sum_{n=1}^N \alpha^n u(c^n(h_\varepsilon, B_{s-1}, t)) \right. \\ & + \sum_{n=1}^N \lambda_P^n(h_\varepsilon, B_{s-1}, t) \sum_{r=0}^{\infty} \beta^r \int_{E^\infty} [u(c^n(h'_\varepsilon, B_{s-1}, t+r)) - u(y^n(t+r))] dF_r(h'_\varepsilon | h_\varepsilon) \\ & + \lambda_R(h_\varepsilon, B_{s-1}, t) \left[\sum_{n=1}^N y^n(t) + B(h_\varepsilon \setminus 1, B_{s-1}, t-1) \right. \\ & \left. \left. - \sum_{n=1}^N c^n(h_\varepsilon, B_{s-1}, t) - R^{-1} B(h_\varepsilon, B_{s-1}, t) \right] \right\} dF_{t-s}(h_\varepsilon | h_\varepsilon^0). \end{aligned} \quad (36)$$

for given $c^n(h_\varepsilon, B_{s-1}, t)$, $\lambda_P^n(h_\varepsilon, B_{s-1}, t)$ and $\lambda_R(h_\varepsilon, B_{s-1}, t)$ functions. The multiplier functions are then chosen in a manner that ensures complementary slackness.

¹⁶In terms of our general notation, the function g_s here is given by $\beta^s [u(c^n(h_\varepsilon, B)) - u(y^n(t))]$.

A first-order condition from (35) with respect to $c^n(h_\varepsilon, B)$ gives:

$$\left[\alpha^n + \sum_{r=0}^{\infty} \beta^r \lambda_P^{n,r}(h_\varepsilon, B) \right] u'(c^n(h_\varepsilon, B)) - \lambda_R(h_\varepsilon, B) = 0. \quad (37)$$

Optimal choice of $B(h_\varepsilon, B_{-1}, t)$, meanwhile, requires generically:

$$-\lambda_R(h_\varepsilon, B_{s-1}, t) + \beta R \int_{E^\infty} \lambda_R(h'_\varepsilon, B_{s-1}, t+1) dF_1(h'_\varepsilon | h_\varepsilon) = 0. \quad (38)$$

The second part of our definition of dynamic Rawlsian policy implies the multipliers $\lambda_R(h_\varepsilon, B)$ defined on the reduced space $E^\infty \times \mathbb{R}$ must satisfy:

$$-\lambda_R(h_\varepsilon, B) + \beta R \int_{E^\infty} \lambda_R(h'_\varepsilon, B(h_\varepsilon, B)) dF_1(h'_\varepsilon | h_\varepsilon) = 0. \quad (39)$$

The expressions discussed in the text follow immediately from (37) and (39).