



On the Permanent Nature of Affirmative Action Policies

Philippe Jehiel, Matthew V Leduc

► To cite this version:

Philippe Jehiel, Matthew V Leduc. On the Permanent Nature of Affirmative Action Policies. 2021. halshs-03359602v1

HAL Id: halshs-03359602

<https://shs.hal.science/halshs-03359602v1>

Preprint submitted on 3 Jan 2022 (v1), last revised 27 Feb 2023 (v2)

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



WORKING PAPER N° 2021 – 54

On the Permanent Nature of Affirmative Action Policies

**Philippe Jehiel
Matthew V. Leduc**

JEL Codes: D40; I28; I30; J15

Keywords: Affirmative Action, General Equilibrium, Loss Aversion, Prospect Theory, Moral Hazard, Game Theory



On the Permanent Nature of Affirmative Action Policies*

Philippe Jehiel[†]
Paris School of Economics
& University College London

Matthew V. Leduc[‡]
Paris School of Economics
& Université Paris 1

December 19, 2021

Abstract

Successive governments must decide whether to implement an affirmative action policy aimed at improving the performance distribution of future generations of a targeted group. Workers receive wages corresponding to their expected performance, suffer a feeling of injustice when getting less than their actual performance, and employers do not (perfectly) observe whether workers benefited from affirmative action. We find that welfare-maximizing governments choose to implement affirmative action *perpetually*, despite the resulting feeling of injustice that eventually dominates the purported beneficial effect on the targeted group's performance. This is in contrast with the first-best that requires affirmative action to be temporary.

Keywords: Affirmative Action, General Equilibrium, Loss Aversion, Prospect Theory, Moral Hazard, Game Theory

JEL codes: D40; I28; I30; J15

Online Appendix available at: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3982544

*We thank seminar participants at Université Paris 1, the Paris School of Economics, the Institut Henri Poincaré, Cambridge University, the Stockholm School of Economics and meetings of the Econometric Society, namely Agnieszka Rusinowska, Emily Tanimura, Gabrielle Demange, Jean-Marc Tallon, Leonardo Pejsachowicz, Francis Bloch, Frédéric Koessler, Olivier Tercieux, Tristan Tomala, Nicholas Vieille, Marie Laclau, Mikhail Safronov, Xin Gao. We also thank Jörgen Weibull, Andrea Galeotti, Matthew L. Elliott and Larry Samuelson for useful comments. Jehiel thanks the ERC (grant no 742816) for funding.

[†]Email: philippe.jehiel@gmail.com

[‡](Corresponding author) Email: mattvleduc@gmail.com. Mailing address: Paris Jourdan Sciences Economiques, 48 boulevard Jourdan, 75014 Paris, France. Tel.: (+33) 01 80 52 16 60

"I yield to no one in my earnest hope that the time will come when an affirmative action program is unnecessary and is, in truth, only a relic of the past. [...] within a decade at most, American society must and will reach a stage of maturity where acting along this line is no longer necessary."

Supreme Court justice Harry Blackmun, 1978

"The Court expects that 25 years from now, the use of racial preferences will no longer be necessary to further the interest approved today."

Supreme Court justice Sandra Day O'Connor, 2003

1 Introduction

The original rationale for affirmative action was to help underrepresented groups close achievement gaps and it was meant to be temporary. Decades after their inception, affirmative action policies however often remain in place. This article will attempt to provide an explanation by studying the incentives of successive governments to implement affirmative action policies.

In our approach, we take as given the initial distribution of talents where we equate talent with school performance and productivity in the labour market. We allow the distribution of talents to differ in the main group A and the targeted group B that may benefit from affirmative action. We analyze a setting in which governments believe that an affirmative action policy improves the talent distribution of group B in future periods. This belief, held by governments, is in line with popular "role model" theories (see, for example, Chung (2000)), according to which witnessing certain members of an underrepresented group achieving success would lead other group members to achieve higher success in the future. We abstract from some direct negative effects of affirmative action such as those voiced in Sowell (2005) which include, namely, mismatches between workers and jobs.

To highlight a new form of inefficiency, we consider successive governments assumed to be benevolent and we study the incentives of these to implement affirmative action policies during their tenure. We do so using a repeated game setting, with each successive government seeking to maximize the same discounted welfare with possibly different weights given to groups A and B .

In the main part of the paper, we suppose that employers cannot condition wages on group membership, which is in line with many anti-discrimination policies. In a perfectly competitive labor market, each employer pays a worker a wage equal to his expected performance. The employer does not observe whether the worker benefited from affirmative action or not and can only estimate this performance based on a curriculum vitae, which may be artificially improved by affirmative action. Paying workers a wage equal to their expected performance thus means that non-beneficiaries of affirmative action will get a wage below their true performance level. We postulate that in such a case, the worker suffers from a *feeling of injustice* that is proportional to the difference between his true performance (which the worker knows) and his wage. We believe such a feeling of injustice is very common, as suggested by recent so-called "populist" reactions. Note that this depressed wage can be understood, more broadly, as being associated with the devaluation of a worker's diplomas or career promotions, which results from the possibility that he may have benefited from affirmative action.

In a first-best scenario, this depressed wage given to non-beneficiaries of affirmative action (and the associated feeling of injustice) means that affirmative action should not last permanently. The optimal duration would be determined, namely, by the weights governments place on the welfare

of the main and the targeted groups, respectively. Indeed, it is essentially a tradeoff between the purported benefits to the targeted group, which stem from its improved performance distribution, and the feeling of injustice suffered by the non-beneficiaries, which stems from the depressed wage. If governments place relatively more weight on the welfare of the main group, then affirmative action would end more quickly. If they place relatively more weight on the welfare of the targeted group, then it would last longer. However, as long as the governments care no less about the welfare of the main group than that of the targeted group, affirmative action would necessarily be ended at some point, as soon as the main group suffers some (even very small) feeling of injustice.

To the contrary, we show that successive governments *always* choosing to implement an affirmative action policy is the unique equilibrium. The intuition is that whether a government actually implements an affirmative action policy is not observable (or is only imperfectly observable) by employers. Therefore a deviation towards implementing an affirmative action policy is perceived by a government as having no effect on decreasing wages, while it is also believed to improve the performance distribution of the targeted group (through a role model argument). This creates a moral hazard, by which each government necessarily chooses to implement an affirmative action policy and fails to internalize the effect that it has on devaluing diplomas and promotions (and thus on depressing wages). The feature that actual affirmative action policy decisions are not observed (or are only imperfectly observed) is justified by the fact that it is often very difficult in practice to determine whether a government actually implemented an affirmative action policy or not. For example, in the United States, these policies are complex, they vary from state to state and even when they are not officially implemented, they may actually take place through private channels (e.g. non-governmental diversity enhancement programs, etc.). Thus, an actual deviation from the official (equilibrium) policy may not be (perfectly) observed by employers in the labor market¹.

The paper is organized as follows. In Section 2, we introduce the basic setting and define the workers' utilities and welfare. In Section 3, we study how employers set the wages they pay to workers and show that it leads to a feeling of injustice felt by non-beneficiaries of affirmative action. In Section 4, we analyze each government's welfare maximization problem and present the two central results: (i) perpetual affirmative action as an equilibrium policy and (ii) the first-best policy where affirmative action is ultimately ended. In Section 5, we discuss how our assumptions can be relaxed, as well as model extensions. We also compare our model with the existing literature. Proofs are relegated to Section 6.

2 Setting

At each time $t \in \mathbb{N}$, a different government must decide whether to implement an affirmative action policy for the duration of its tenure (one period). That is, it chooses an action $\sigma_t \in \{0, 1\}$, where $\sigma_t = 0$ corresponds to no affirmative action and $\sigma_t = 1$ corresponds to affirmative action.

A population of workers consists of two groups: group A (the main group) and group B (the targeted group). A worker has a performance level $c \in [0, 1]$. This can be understood, for instance, as his result in a standardized university admission test.

¹Formally, this amounts in viewing affirmative action decisions as being implemented by many different agents, thereby making the observation by employers of an individual decision very hard.

At any time t , group A 's performance density is $f_A(c)$ while group B 's performance² density is $f_{B,n_t}(c)$, where $n_t = \sum_{s < t} \sigma_s$ is the number of times previous governments have implemented affirmative action policies. $f_A(c)$ and $f_{B,n_t}(c)$ have support $[0, 1]$. We will describe later how $f_{B,n_t}(c)$ varies with n_t but intuitively as n_t increases, $f_{B,n_t}(c)$ shifts lower values of c to higher values, resulting in first-order stochastic dominance. Each agent lives for only one period³. At each time t , a mass $|A|$ and a mass $|B|$ of new agents from groups A and B respectively are born, with performance levels drawn according to $f_A(c)$ and $f_{B,n_t}(c)$.

2.1 Effect of affirmative action policy

An affirmative action policy has two effects. First, it gives an immediate artificial boost to the curriculum vitae of a worker benefiting from it. This models the fact that a beneficiary of affirmative action has expanded opportunities in terms of university admissions or career promotions compared to a non-beneficiary, thereby artificially enhancing the quality of his curriculum vitae. Second, it is also believed by governments to have long-term, positive effects on the performance distribution of group B . This purported long-term effect is in line with popular "role model" theories (e.g. Chung (2000)). This second effect will be captured by the dependence of $f_{B,n_t}(c)$ on n_t .

It is important to note that an affirmative action policy can be interpreted⁴ as anything that artificially increases the quality of a curriculum vitae (immediate effect) and improves the performance distribution of future generations (purported role model effect).

2.1.1 Effect of affirmative action policy on curriculum vitae quality

When $\sigma_t = 1$, a member of group B with performance level $c \in [0, 1]$, but who obtained his university degree under that government, will have a curriculum vitae quality $\bar{c} = g(c)$, where g is an increasing function such that $g(c) > c$, $\forall c \in (0, 1)$, and $g(0) = 0$, $g(1) = 1$. The support of $\bar{c}$ is thus also $[0, 1]$. An affirmative action policy therefore increases the curriculum vitae quality of a beneficiary above his actual performance level. By contrast, an affirmative action policy has no effect on the curriculum vitae quality of members of group A , nor on those of members of group B who did not benefit from the affirmative action policy (i.e. in the case $\sigma_t = 0$). That is, their curriculum vitae quality corresponds to their actual performance level: $\bar{c} = c$.

2.1.2 Effect of affirmative action policy on actual performance

We suppose that if $\sigma_t = 1$, then the next period's performance distribution of group B is shifted so that $f_{B,n_{t+1}}(c) \succ f_{B,n_t}(c)$, where $\succ$ indicates strong first-order stochastic dominance. Note that the effect of the shift is permanent, i.e. the improvement remains in all future periods.

If $\sigma_t = 1$ for all t , then $f_{B,n_t}(c) \uparrow \bar{f}_B(c)$. Since $f_{B,n_t}(c)$ converges from below to a limiting distribution $\bar{f}_B(c)$, this implies that the distributional improvements become smaller and smaller

²In the applications we will have in mind, it is reasonable to think that group A 's performance distribution initially differs from that of group B , although this plays no role in our analysis.

³The model can easily be extended to allow agents to live for more than one period and to have overlapping generations. Since such elaborations would play no role in our analysis, we have chosen the simpler setting in which agents just live for one period.

⁴See Section 1 of the Online Appendix for a discussion of how our model can accommodate even more general interpretations of affirmative action.

as governments keep implementing affirmative action policies. Group A 's performance distribution $f_A(c)$ does not vary with t .

2.2 Utilities and welfare

A worker is of type $\theta = (c, \bar{c}, G)$, where c is his true performance level, $\bar{c}$ is his curriculum vitae quality and $G \in \{A, B\}$ is the group this worker belongs to. A time t worker knows his type and the wage function $\omega_t(\bar{c})$ set by employers, which is the wage the worker earns based on his curriculum vitae quality.⁵ This is formalized in the following definition.

Definition 1 *A wage function $\omega_t : [0, 1] \rightarrow [0, 1]$ determines, at time t , the wage a worker earns when presenting a curriculum vitae of quality $\bar{c}$ to an employer.*

2.2.1 Utility

The utility of a type $(c, \bar{c}, G)$ worker at time t is

$$u_{G,t}(\bar{c}, c) = \omega_t(\bar{c}) - \gamma_G \max\{c - \omega_t(\bar{c}), 0\} \quad (1)$$

where $\gamma_G \max\{c - \omega_t(\bar{c}), 0\}$, for some $\gamma_G > 0$, captures the fact that a feeling of “injustice” is suffered when a worker gets a salary that is below his true performance level. Note that we allow $\gamma_A \neq \gamma_B$ so as to capture that the feeling of injustice may differently affect groups A and B .

In particular, the utility of a type $(c, \bar{c}, G)$ worker who benefits from affirmative action has the form

$$u_{G,t}(\bar{c}, c) = \omega_t(g(c)) - \gamma_G \max\{c - \omega_t(g(c)), 0\}$$

since $\bar{c} = g(c)$, while the utility of a type $(c, \bar{c}, G)$ worker who does not benefit from affirmative action has the form

$$u_{G,t}(\bar{c}, c) = \omega_t(c) - \gamma_G \max\{c - \omega_t(c), 0\}$$

since $\bar{c} = c$. We will often denote by $u_{B,t}(g(c), c)$ (respectively, by $u_{B,t}(c, c)$) the utility of a group B worker benefiting (respectively, not benefiting) from affirmative action, while we will denote by $u_{A,t}(c, c)$ the utility of a group A worker.

In the above, we assume that workers have the correct perception of their performance level c . We also note that there are no extra positive effects on utility of receiving a wage greater than the performance level. Such an asymmetry in the utility assessment of wages above or below the performance level is in line with well documented psychological studies (see in particular the prospect theory of Kahneman and Tversky (1979)), which suggest a different assessment for payoff realizations above or below the reference point (here naturally identified with the performance level).

⁵If workers were to live several periods, we could envision a more elaborate model in which the wage earned in later periods would also depend on the true performance assumed to be partly observed then. Our qualitative insights would be unaffected.

2.2.2 Welfare

The welfare of each group at time t is defined by taking the aggregate utility of that group. We thus have,

$$W_{A,t} = |A| \int_0^1 u_{A,t}(c, c) f_A(c) dc$$

$$W_{B,t} = |B| \int_0^1 \left(\sigma_t u_{B,t}(g(c), c) + (1 - \sigma_t) u_{B,t}(c, c) \right) f_{B,n_t}(c) dc$$

where σ_t is the actual policy decision made by the time t government.

3 Effect of affirmative action policy on wage levels

3.1 Informational environment

We assume that employers at time t do not observe the sequence $\{\sigma_s\}_{s=1}^t$ of actual policy decisions made by specific governments. Moreover, they do not observe the actual performance distribution $f_{B,n_t}(c)$, otherwise they would be able to infer the sum of the actual σ_s , i.e. $n_t = \sum_{s < t} \sigma_s$. As a result, they cannot tell whether a worker benefited from affirmative action or not. In Section 5.1.2, we note that our main insights still hold as long as employers do not *perfectly* observe σ .

As usual, *in equilibrium* employers know the strategy $\sigma^* = \{\sigma_s^*\}_{s=1}^\infty$ that is chosen by governments and thus they can compute the probability that a worker benefited from affirmative action, conditional upon observing his curriculum vitae quality $\bar{c}$. Thus, employers are aware of the general policy on matters of affirmative action (i.e. σ^*), without being able to actually observe (perfectly) whether it is implemented or not in a given period (i.e. the actual decisions σ).

A justification for the unobservability (or imperfect observability) of the actual σ (which plays a key role in our results) is that it is often very difficult to observe whether an affirmative action policy is implemented or not. For example, in the United States, these policies vary from state to state and even when they are not officially implemented, they may actually take place through various channels, as reported in the introduction. More formally, we could envision our government as being composed of many different decision makers, all concerned with the same aggregate welfare, but whose specific affirmative action decisions could not be observed by employers. Thus, an actual deviation σ_t from the official policy σ_t^* may not be observed.⁶

3.2 Setting wages

We consider a perfectly competitive labor market, where an employer pays a worker a wage equal to his expected performance level. In Section 2.3 of the Online Appendix, this reduced-form approach is rationalized based on a Bertrand-type model of competition between employers.

We also assume, in the main part of the paper, that employers are not allowed to take group information (A or B) into account when giving a wage to a particular worker (see Section 5.1.1

⁶As said earlier, we discuss later on in Section 5.1.2 the robustness of our insights if the labor market only imperfectly observes the choices of σ_t . The unobservability of the performance distribution $f_{B,n_t}(c)$ is justified since it is difficult to estimate in practice and studies published on that topic are often imprecise or time-lagged. The only thing that needs to be understood (or believed) by employers is the purported dependence of $f_{B,n_t}(c)$ on $n_t = \sum_{s < t} \sigma_s^*$ along the equilibrium path.

for elaborations). This is consistent with anti-discrimination laws enacted in many countries and occupational areas⁷. That is, they must set a wage conditioned only on the curriculum vitae quality $\bar{c}$. The wage $\omega_t^*(\bar{c})$ paid to a worker of type $(c, \bar{c}, A)$ or to a worker of type $(c, \bar{c}, B)$ is thus the conditional expectation $\mathbb{E}_t[c|\bar{c}, \sigma^*]$ of the worker's true performance level c , expressed in the following lemma.

Lemma 1 *Given an equilibrium government policy strategy σ^* , the wage paid at time t to a worker with curriculum vitae quality $\bar{c}$ has the form*

$$\omega_t^*(\bar{c}) = \mathbb{P}_t^*(\{aa\}|\bar{c}) \cdot g^{-1}(\bar{c}) + (1 - \mathbb{P}_t^*(\{aa\}|\bar{c})) \cdot \bar{c}$$

where

$$\mathbb{P}_t^*(\{aa\}|\bar{c}) = \frac{|B|\sigma_t^* f_{B,n_t}(g^{-1}(\bar{c}))/g'^{-1}(\bar{c})}{|A|f_A(\bar{c}) + |B|(1 - \sigma_t^*)f_{B,n_t}(\bar{c}) + |B|\sigma_t^* f_{B,n_t}(g^{-1}(\bar{c}))/g'^{-1}(\bar{c})}$$

and $\{aa\}$ is the event that a worker benefited from affirmative action.

In words, $\omega_t^*(\bar{c})$ is a convex combination between $g^{-1}(\bar{c})$ and $\bar{c}$, where the weight assigned to $g^{-1}(\bar{c})$ is the overall probability that a worker with curriculum vitae $\bar{c}$ benefited from affirmative action at time t (taking into account the equilibrium policy σ_t^* , hence the expression for $\mathbb{P}_t^*(\{aa\}|\bar{c})$).

We make the following assumption on $\mathbb{E}_t[c|\bar{c}, \sigma^*]$ for simplicity of exposition. Our results do not depend on it, but it will allow us to present them in a simpler manner, since we can rule out strategic behavior by which an agent could present a curriculum vitae of lower⁸ quality than $\bar{c}$. An extension where a worker is allowed to present a curriculum vitae of a different quality than $\bar{c}$ is presented in the Online Appendix, where the robustness of our results to such strategic behavior is established in a more general context.

Assumption 1 *The conditional expectation $\mathbb{E}_t[c|\bar{c}, \sigma^*]$, and thus the wage function $\omega_t^*(\bar{c})$, is non-decreasing in $\bar{c}$.*

This assumption is easily satisfied under some conditions, e.g. when the likelihood ratio $\frac{f_A(c)}{f_B(g(c))}$ is increasing or when the mass $|A|$ is sufficiently larger than the mass $|B|$.

Relying on the expression of equilibrium wage derived in Lemma 1, we note that whether the earned wage lies above or below the performance level solely depends on whether or not the worker benefited from affirmative action:

Lemma 2

- (i) *Suppose $\sigma_t^* = 1$. Then any worker gets a wage lower than his curriculum vitae quality (i.e. $\bar{c} > \omega_t^*(\bar{c})$). Moreover, a worker benefiting from affirmative action gets a wage higher than his performance level (i.e. $c = g^{-1}(\bar{c}) < \omega_t^*(\bar{c})$), while a worker not benefiting from affirmative action gets a wage lower than his performance level (i.e. $c = \bar{c} > \omega_t^*(\bar{c})$).*
- (ii) *Suppose $\sigma_t^* = 0$. Then any worker gets a wage equal to his curriculum vitae quality and his performance level (i.e. $c = \bar{c} = \omega_t^*(\bar{c})$).*

⁷We show in Section 5.1.1, that our results actually hold even if group-based discrimination is allowed, as long as some members of the targeted group B do not benefit from affirmative action (a fairly weak assumption). In such a case, the feeling of injustice is suffered entirely by them (and not by group A members) and this is enough for our results to hold.

⁸Indeed, if the wage function $\omega_t^*(\bar{c})$ is decreasing on some parts of the support $[0, 1]$, a worker could earn a higher wage by presenting a curriculum vitae of lower quality than $\bar{c}$.

3.3 Feeling of injustice and broader interpretation of the depressed wage

In our model, the wage is depressed due to the possibility that a worker benefited from affirmative action. This represents the fact that a certain curriculum vitae quality is, in expectation, no longer associated with the same performance level as if there were no affirmative action policy. Indeed, an affirmative action policy has the effect of devaluing the diplomas or promotions that figure on a worker's curriculum vitae, if there is only some chance that the worker may have benefited from such a policy.

Using Lemma 2, we now make the following observation.

Observation 1 (Feeling of injustice)

(i) Suppose $\sigma_t^* = 1$.

The utility of a type $(c, \bar{c}, G)$ worker not benefiting from affirmative action can be written as

$$\begin{aligned} u_{G,t}(c, c) &= \omega_t^*(c) - \gamma_G \max\{c - \omega_t^*(c), 0\} \\ &= \omega_t^*(c) - \gamma_G(c - \omega_t^*(c)) \end{aligned}$$

since $\bar{c} = c$ and $\omega_t^(c) < c$. Such a worker gets a wage lower than his performance level and suffers a feeling of injustice.*

By contrast, the utility of a type $(c, \bar{c}, G)$ worker benefiting from affirmative action can be written as

$$\begin{aligned} u_{G,t}(g(c), c) &= \omega_t^*(g(c)) - \gamma_G \max\{c - \omega_t^*(g(c)), 0\} \\ &= \omega_t^*(g(c)) \end{aligned}$$

since $\bar{c} = g(c) > c$ and $\omega_t^(g(c)) > c$. Such a worker gets a wage higher than his performance level and does not suffer a feeling of injustice.*

(ii) Suppose $\sigma_t^* = 0$.

The utility of a type $(c, \bar{c}, G)$ worker can be written as

$$\begin{aligned} u_{G,t}(c, c) &= \omega_t^*(c) - \gamma_G \max\{c - \omega_t^*(c), 0\} \\ &= c \end{aligned}$$

since $\bar{c} = c$ and $\omega_t^(c) = c$. Such a worker gets a wage equal to his performance level and does not suffer a feeling of injustice.*

4 The government's decision problem

4.1 Informational environment

It is irrelevant whether a time t government observes the past policies $\sigma_{t'}, t' < t$, that were actually chosen by other governments or the actual performance distribution $f_{B,n_t}(c)$ prevailing at time t . This will have no impact on its decision to implement an affirmative action policy or not, as it

will become clear later. The government, however knows (or believes) the purported effect of n_t on $f_{B,n_t}(c)$ and thus believes that choosing an affirmative action policy $\sigma_t = 1$ will improve the performance distribution of group B : $f_{B,n_{t+1}}(c) \succ f_{B,n_t}(c)$.

4.2 Policy decisions

For all $t \geq 1$, a time t government wants to maximize the following objective function:

$$\max_{\sigma_t \in \{0,1\}} \sum_{s=t}^{\infty} \delta^{s-t} (W_{A,s} + \lambda_B W_{B,s}) \quad (2)$$

Here $\lambda_B \geq 0$ is a weight placed on the welfare of group B (the weight placed on the welfare of group A is normalized to 1). We suppose this weight is constant through time. $\delta \in (0,1)$ is the discount factor.

We will mostly consider the case when $\lambda_B \leq 1$. In particular, $\lambda_B = 1$ corresponds to the standard total welfare criterion and $\lambda_B < 1$ reflects a preference for the main group A in the governmental objective. We will also comment on the case when $\lambda_B > 1$, which reflects a preference for the targeted group B .

Since a time t government knows the equilibrium $\sigma^* = \{\sigma_s^*\}_{s=1}^{\infty}$ that is played by other governments, it is able to compute $W_{A,s}$ and $W_{B,s}$, for $s > t$. In other words, $\omega_s^*(\bar{c})$ and $f_{B,n_s}(c)$ are taken to be consistent with this equilibrium play⁹.

Our first main result, Proposition 1, states that different governments choosing to *perpetually* implement an affirmative action policy is the unique equilibrium.

Proposition 1 (Equilibrium policy) *Given any $\lambda_B > 0$, $\sigma_t^* = 1$ for all t is the unique equilibrium.*

The intuition behind Proposition 1 is that any government believes that implementing an affirmative action policy improves the performance distribution of future cohorts of B workers. Moreover, since a policy deviation is not observed by employers, it cannot have any improving impact on the wage function ω_t^* chosen by employers. Therefore, there is no reason why a particular government would deviate from an equilibrium $\sigma_t^* = 1$ in which it implements an affirmative action policy. Conversely, a deviation from a putative equilibrium in which $\sigma_t^* = 0$ to $\sigma_t = 1$ would increase the average performance of future cohorts of B workers (and thus the average wage they receive), without having a worsening impact on the wage function ω_t^* . This establishes $\sigma_t^* = 1$ as the unique equilibrium.

Our second main result, Proposition 2, states that in the first-best scenario, affirmative action policies *always* end after a finite number of periods.

Proposition 2 (First-best policy) *Suppose that at time $t = 0$, a single government announces (and commits to) the policy plan $\hat{\sigma} = \{\hat{\sigma}_t\}_{t=1}^{\infty}$ that maximizes the welfare function $\sum_{t=1}^{\infty} \delta^t (W_{A,t} + \lambda_B W_{B,t})$, and assume $\gamma_A \neq 0$. Then for any $\lambda_B \in [0,1]$, there exists $\bar{\delta} \in (0,1)$ such that for all $\delta \in (\bar{\delta},1)$, $\{\hat{\sigma}_t\}_{t=1}^{\infty}$ has a threshold form $\hat{\sigma}_t = 1$ for $t < \bar{T}$ and $\hat{\sigma}_t = 0$ for $t \geq \bar{T}$, for some (finite) $\bar{T} \in \mathbb{N}$.*

⁹In the welfare, we have not included the firms' profits. Note however that these profits are null on the equilibrium path, due to our assumption of perfect competition. Including firms' profits in the governmental objective would affect the assessment of deviations but the qualitative insights presented below would be unaffected.

Proposition 2 essentially means that if different governments were able to coordinate their actions over time, they would never choose to make affirmative action permanent. The intuition is quite simple: After a certain number of periods the improvement in the performance distribution becomes marginal, while the depressing effect on wages is not. As a matter of fact, $f_{B,n_t}(c)$ converges from below to a limiting distribution $\bar{f}_B(c)$, implying that the distributional improvements become smaller and smaller as affirmative action policies are implemented over time.

The optimal threshold $\bar{T}$, while always finite, depends on the relative weight placed by governments on the welfare of the targeted group B relative to the main group A , i.e. on λ_B .

Note that when $\lambda_B < 1$ (i.e. when the government cares relatively more about group A than group B), the parameter γ_A governing the feeling of injustice of group A can be 0 and the first-best policy will still prescribe stopping affirmative action after a finite number of periods, because the depressed wage penalizes group A sufficiently while the performance distribution of group B is only marginally improved.

When $\lambda_B = 1$ (i.e. when the government cares equally about group A and group B), then since the average wage is equal to the average performance level across the market (i.e. $\mathbb{E}_t[\omega_t^*(\bar{c})] = \mathbb{E}_t[c]$), an affirmative action policy effectively represents just a transfer of welfare from group A to group B . Indeed, this transfer of welfare takes place through group A workers receiving wages lower than their performance levels while group B workers receive wages higher than their performance levels. In this case, as long as the parameter γ_A is *strictly greater* than 0 (no matter how small it is), a first-best policy will prescribe stopping affirmative action after a finite number of periods because otherwise the feeling of injustice felt by group A would become worse than the improvement in the performance distribution of group B after sufficiently many implementations of the affirmative action policy.

Finally, if λ_B were to be strictly greater than 1 (i.e. when the government cares relatively more about group B than group A), then we might need γ_A to be sufficiently positive in order to justify stopping affirmative action. In such a case, $\gamma_A \geq \underline{\gamma}_A$ for some $\underline{\gamma}_A > 0$ would be a sufficient (but not always necessary) condition for our first-best result to hold.

5 Some extensions

5.1 Discussion of assumptions

In the next subsections, we show that our main results are quite robust and most often hold, even if we relax the assumptions made in the main part of the paper. We explain that all we really need for our main results to hold is that an employer cannot be certain that a group- B worker has not benefited from affirmative action.

5.1.1 Allowing for group-based discrimination in the labor market

In our main model, we have not allowed employers to condition wages on the group A or B to which a worker belongs. This was motivated on grounds that such discrimination is in general forbidden. If discrimination were allowed, then within our current model, governments would still choose $\sigma_t^* = 1$ in every period, but this time workers from both groups A and B would receive a wage equal to their

performance levels. As a result, the equilibrium would coincide with first-best¹⁰. Such a conclusion that unregulated discrimination leads to first-best would however be fragile to the introduction of natural frictions. For example, if we assume affirmative action only reaches a fraction $\xi \in (0, 1)$ of group B , then the share $1 - \xi$ of group B not benefiting from affirmative action would play a role similar to that of group A in the current model, thereby leading to similar insights.

5.1.2 Making the policy σ_t partially observable to employers

From another perspective, in our main model, we have assumed that employers make no observation at all about the chosen policies σ_t . We note that if employers at t were to observe a noisy signal about σ_t (with a support of signal realization that would be the same whether $\sigma_t = 0$ or 1), then $\sigma_t^* = 1$ for all t would remain an equilibrium. This holds because in such an equilibrium, the signal observed by employers would be perceived to be uninformative and thus would have no effect on the chosen wage, thereby providing the incentive to governments to choose $\sigma_t = 1$ in all periods, as in the main model¹¹. Since the first-best policy would be unaffected by this modification, this establishes the robustness of our main insights with respect to a broad class of observational environments.

5.1.3 Other elaborations

In an Online Appendix¹², we present further elaborations. Namely, we suggest how our model can also accommodate the case when affirmative action takes the form of a biased promotion process. We also present a generalized model where we allow for strategic behavior by workers, by which they can present a curriculum vitae of any chosen quality. This generalization formally removes the need for Assumption 1. We also introduce a labor market congestion externality caused by affirmative action, which allows us to capture at least to some extent the fact that jobs obtained by beneficiaries of affirmative action are no longer available to non-beneficiaries. We show that our main insights go through even in the presence of such elaborations. Finally, we microfound the wage-setting behavior of firms with Bertrand competition.

5.2 Comparisons with existing literature

We mainly depart from the existing literature on affirmative action by studying the incentives of governments to implement affirmative action policies. Indeed, most of the literature focuses on other incentives: those linked to hiring decisions made by employers or to investments in human capital made by workers, which may be reduced by an affirmative action policy (e.g. Lundberg and Startz (1983), Coate and Loury (1993a) or Coate and Loury (1993b); see also Fang and Moro (2011) for a survey on discrimination and affirmative action).

¹⁰It is worth noting that within our basic model, inefficiency arises because both affirmative action policies and anti-discrimination rules operate at the same time. This is in contrast to conventional views in which affirmative action and anti-discrimination rules are often regarded as necessary and complementary tools to achieve better equality and higher welfare.

¹¹This observation is related to one made by Bagwell (1995) in the context of Stackelberg interactions in which the action of the first-mover would be observed with noise. He notes that the Nash equilibrium of the normal form game is then a Perfect Bayesian Equilibrium of the two-stage interaction. Van Damme and Hurkens (1997) later note that there are other equilibria involving mixed strategies, but such equilibria can be regarded as being less robust to the extent that they rely on indifferences of the players.

¹²Online Appendix available at: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3982544

The existing literature on affirmative action is vast and often tries to describe or explain inequalities between groups. Early developments include taste-based theories of discrimination (e.g. Becker (1957)), which suppose that exogenous preferences generate wage differences between groups, although the latter are unlikely to persist in competitive markets. Statistical discrimination theories, on the other hand, mainly attempt to explain outcome differences using imperfect information about the workers' performance levels, which leads to different wages being rationally paid to workers of different groups (e.g. Phelps (1972), Arrow (1973), Lundberg and Startz (1983), Coate and Loury (1993a) or Coate and Loury (1993b)). Such models often also link these different wages to the workers' incentives to invest in human capital, thus sustaining a performance gap between groups.

Our argument is based on a novel moral hazard consideration on the government's part and it complements other (more direct) critiques such as those voiced by Sowell (2005).

6 Proofs

Proof of Lemma 1.

Note that $\omega_t^*(\bar{c})$ is the conditional expectation of the worker's actual performance level at time t when declaring a curriculum vitae of quality $\bar{c}$ and when the equilibrium policy is σ^* . Thus,

$$\begin{aligned}\omega_t^*(\bar{c}) &= \mathbb{E}_t[c|\bar{c}, \sigma^*] \\ &= \mathbb{P}_t^*(\{aa\}|\bar{c}) \cdot g^{-1}(\bar{c}) + (1 - \mathbb{P}_t^*(\{aa\}|\bar{c})) \cdot \bar{c}\end{aligned}$$

Now to express $\mathbb{P}_t^*(\{aa\}|\bar{c})$, we first express $\mathbb{P}_t^*(\{aa\}|\tilde{c} \in N(\bar{c}, \epsilon))$, where $N(\bar{c}, \epsilon)$ is an ϵ -neighborhood of $\bar{c}$:

$$\begin{aligned}\mathbb{P}_t^*(\{aa\}|\tilde{c} \in N(\bar{c}, \epsilon)) &= \frac{\mathbb{P}_t^*(\{\tilde{c} \in N(\bar{c}, \epsilon)\} \cap \{aa\})}{\mathbb{P}_t^*(\tilde{c} \in N(\bar{c}, \epsilon))} \\ &= \frac{\mathbb{P}_t^*(\tilde{c} \in N(\bar{c}, \epsilon) \cap B) \cdot \sigma_t^*}{\mathbb{P}_t^*(\tilde{c} \in N(\bar{c}, \epsilon))} \\ &= \frac{\mathbb{P}_t^*(\tilde{c} \in N(\bar{c}, \epsilon)|B) \cdot \mathbb{P}(B) \cdot \sigma_t^*}{\mathbb{P}_t^*(\tilde{c} \in N(\bar{c}, \epsilon))} \\ &= \frac{\sigma_t^* \int_{\tilde{c} \in N(g^{-1}(\bar{c}), \epsilon/g'^{-1}(\bar{c}))} f_{B, n_t}(g^{-1}(\tilde{c})) d\tilde{c} \frac{|B|}{|A|+|B|}}{\int_{\tilde{c} \in N(\bar{c}, \epsilon)} f_A(\tilde{c}) d\tilde{c} \frac{|A|}{|A|+|B|} + (1 - \sigma_t^*) \int_{\tilde{c} \in N(\bar{c}, \epsilon)} f_{B, n_t}(\tilde{c}) d\tilde{c} \frac{|B|}{|A|+|B|} + \sigma_t^* \int_{\tilde{c} \in N(g^{-1}(\bar{c}), \epsilon/g'^{-1}(\bar{c}))} f_{B, n_t}(g^{-1}(\tilde{c})) d\tilde{c} \frac{|B|}{|A|+|B|}} \\ &= \frac{|B| \sigma_t^* \int_{\tilde{c} \in N(g^{-1}(\bar{c}), \epsilon/g'^{-1}(\bar{c}))} f_{B, n_t}(g^{-1}(\tilde{c})) d\tilde{c}}{|A| \int_{\tilde{c} \in N(\bar{c}, \epsilon)} f_A(\tilde{c}) d\tilde{c} + |B| (1 - \sigma_t^*) \int_{\tilde{c} \in N(\bar{c}, \epsilon)} f_{B, n_t}(\tilde{c}) d\tilde{c} + |B| \sigma_t^* \int_{\tilde{c} \in N(g^{-1}(\bar{c}), \epsilon/g'^{-1}(\bar{c}))} f_{B, n_t}(g^{-1}(\tilde{c})) d\tilde{c}}\end{aligned}$$

Then, we take the limit as $\epsilon \rightarrow 0$:

$$\begin{aligned}\mathbb{P}_t^*(\{aa\}|\bar{c}) &= \lim_{\epsilon \rightarrow 0} \mathbb{P}_t^*(\{aa\}|\tilde{c} \in N(\bar{c}, \epsilon)) \\ &= \lim_{\epsilon \rightarrow 0} \frac{|B| \sigma_t^* f_{B, n_t}(g^{-1}(\bar{c})) 2\epsilon / g'^{-1}(\bar{c})}{|A| f_A(\bar{c}) 2\epsilon + |B| (1 - \sigma_t^*) f_{B, n_t}(\bar{c}) 2\epsilon + |B| \sigma_t^* f_{B, n_t}(g^{-1}(\bar{c})) 2\epsilon / g'^{-1}(\bar{c})} \\ &= \frac{|B| \sigma_t^* f_{B, n_t}(g^{-1}(\bar{c})) / g'^{-1}(\bar{c})}{|A| f_A(\bar{c}) + |B| (1 - \sigma_t^*) f_{B, n_t}(\bar{c}) + |B| \sigma_t^* f_{B, n_t}(g^{-1}(\bar{c})) / g'^{-1}(\bar{c})}\end{aligned}$$

■

Proof of Lemma 2.

From Lemma 1 we know that

$$\omega_t^*(\bar{c}) = \mathbb{P}_t^*(\{aa\}|\bar{c}) \cdot g^{-1}(\bar{c}) + (1 - \mathbb{P}_t^*(\{aa\}|\bar{c})) \cdot \bar{c}$$

Part (i): When $\sigma^* = 1$, then $\mathbb{P}_t^*(\{aa\}|\bar{c}) > 0$. Since $g^{-1}(\bar{c}) < \bar{c}$, it follows immediately that $g^{-1}(\bar{c}) < \omega_t^*(\bar{c}) < \bar{c}$. Thus, if the worker does not benefit from affirmative action (i.e. $c = \bar{c}$), then $\omega_t^*(\bar{c}) < c$ and he gets a wage lower than his performance level. On the other hand, if the worker benefits from affirmative action (i.e. $c = g^{-1}(\bar{c})$), then $c < \omega_t^*(\bar{c})$ and he gets a wage higher than his performance level.

Part (ii): When $\sigma^* = 0$, then $\mathbb{P}_t^*(\{aa\}|\bar{c}) = 0$. Thus, $\omega_t^*(\bar{c}) = \bar{c}$ and $\bar{c} = c$ since no one benefits from affirmative action. ■

Proof of Proposition 1 (Equilibrium policy).

We first show that $\{\sigma_s^*\}_{s=1}^\infty = \{1\}_{s=1}^\infty$ is an equilibrium.

Given some equilibrium decision profile $\sigma^* = \{\sigma_s^*\}_{s=1}^\infty$, any deviation σ'_t at some time t has no impact on the wage function since this deviation is unobserved by employers. Indeed, employers form a wage $\omega_t^*(\bar{c}) = \mathbb{E}_t[c|\bar{c}, \sigma^*]$ that depends on the *equilibrium* policy decisions σ^* .

Therefore, $\sum_{s=t}^\infty W_{A,s} \delta^{s-t}$, the discounted future welfare of group A , is completely unaffected by an unobserved deviation to σ'_t . Indeed, $W_{A,s} = |A| \int_0^1 u_{A,t}(c, c) f_A(c) dc$, where the density function $f_A(c)$ is constant through time and thus not impacted by σ_t , while $u_{A,t}(c, c) = \omega_t^*(c) - \gamma_A(c - \omega_t^*(c))$ and the wage $\omega_t^*(c)$ is unaffected by an unobserved deviation to σ'_t .

On the other hand, $\sum_{s=t}^\infty \lambda_B W_{B,s} \delta^{s-t}$, the discounted future welfare of group B , is strictly lower following an unobserved deviation from $\sigma_t^* = 1$ to $\sigma'_t = 0$. Indeed, at time t ,

$$\begin{aligned} W_{B,t|\sigma'_t=0} &= |B| \int_0^1 u_{B,t}(c, c) f_{B,n_t}(c) dc \\ &= |B| \int_0^1 (\omega_t^*(c) - \gamma_B(c - \omega_t^*(c))) f_{B,n_t}(c) dc \\ &< |B| \int_0^1 \omega_t^*(g(c)) f_{B,n_t}(c) dc \\ &= |B| \int_0^1 u_{B,t}(g(c), c) f_{B,n_t}(c) dc \\ &= W_{B,t|\sigma_t^*=1} \end{aligned}$$

where we have used Observation 1(i), the fact that the wage is not affected by an unobserved deviation and the facts that $\omega_t^*(c) < \omega_t^*(g(c))$ and that $c > \omega_t^*(c)$ by Lemma 2.

Moreover, at times $s > t$, $f_{B,n_s|\sigma'_t=0}(c) < f_{B,n_s|\sigma_t^*=1}(c)$ since a deviation to $\sigma'_t = 0$ has the effect of not changing the distribution of performance at time $t + 1$ compared to the previous period t . Thus, using Observation 1(i), and the fact that the wage is not affected by an unobserved deviation,

then for all $s > t$,

$$\begin{aligned}
W_{B,s|\sigma'_t=0} &= |B| \int_0^1 u_{B,s}(g(c), c) f_{B,n_s|\sigma'_t=0}(c) dc \\
&= |B| \int_0^1 \omega_s^*(g(c)) f_{B,n_s|\sigma'_t=0}(c) dc \\
&< |B| \int_0^1 \omega_s^*(g(c)) f_{B,n_s|\sigma'_t=1}(c) dc \\
&= |B| \int_0^1 u_{B,s}(g(c), c) f_{B,n_s|\sigma'_t=1}(c) dc \\
&= W_{B,s|\sigma'_t=1}
\end{aligned}$$

where the inequality follows from $f_{B,n_s|\sigma'_t=0}(c) \prec f_{B,n_s|\sigma'_t=1}(c)$ and Assumption 1.

It follows that as long as $\lambda_B > 0$, then $\sigma_t^* = 1$ for all t will be an equilibrium.

To show that this is the unique equilibrium, we now have to show that a deviation from $\sigma_t^* = 0$ to $\sigma'_t = 1$ is always desirable for a time- t government. For that purpose, suppose that $\sigma_t^* = 0$ for some t . Then, we must show that $\sum_{s=t}^{\infty} \lambda_B W_{B,s} \delta^{s-t}$ is strictly higher following an unobserved deviation from $\sigma_t^* = 0$ to $\sigma'_t = 1$.

Consider first the effect of this deviation on the welfare at time t of members of group B . The same argument as before can be used to show that $W_{B,t|\sigma'_t=1} > W_{B,t|\sigma'_t=0}$.

Consider now the effect of this deviation on the welfare, at any future time $s > t$, of members of group B . We know that $f_{B,n_s|\sigma'_t=0}(c) \prec f_{B,n_s|\sigma'_t=1}(c)$ for all $s > t$ since a deviation to $\sigma'_t = 1$ has the effect of shifting (in a strict first-order stochastic dominance sense) the future performance distributions of group B .

Then for all $s > t$,

$$\begin{aligned}
W_{B,s|\sigma'_t=0} &= |B| \int_0^1 (\sigma_s^* \omega_s^*(g(c)) + (1 - \sigma_s^*)c) f_{B,n_s|\sigma'_t=0}(c) dc \\
&< |B| \int_0^1 (\sigma_s^* \omega_s^*(g(c)) + (1 - \sigma_s^*)c) f_{B,n_s|\sigma'_t=1}(c) dc \\
&= W_{B,s|\sigma'_t=1}
\end{aligned}$$

where we made use of $\sigma_s^* \omega_s^*(g(c)) + (1 - \sigma_s^*)c$ being increasing in c (follows from Assumption 1) and $f_{B,n_s|\sigma'_t=0}(c) \prec f_{B,n_s|\sigma'_t=1}(c)$ for all $s > t$. ■

Proof of Proposition 2 (First-best policy).

We start with the following lemma.

Lemma 3 *Let $\sigma' = \{\sigma'_t\}_{t=1}^{\infty}$ be a policy plan with $\sigma'_\tau = 0$ and $\sigma'_{\tau+1} = 1$ for some τ . Let $\sigma = \{\sigma_t\}_{t=1}^{\infty}$ be another policy plan with $\sigma_\tau = 1$, $\sigma_{\tau+1} = 0$ and $\sigma'_t = \sigma_t$ for all other t . Then there exists $\bar{\delta} \geq 0$ such that for all $\delta \in (\bar{\delta}, 1)$, σ yields a strictly higher welfare than σ' .*

Proof of Lemma 3.

First note that for any group $G \in \{A, B\}$,

$$\sum_{t=1}^{\infty} \delta^t W_{G,t} = \sum_{t=1}^{\tau-1} \delta^t W_{G,t} + \delta^\tau W_{G,\tau} + \delta^{\tau+1} W_{G,\tau+1} + \sum_{t=\tau+2}^{\infty} \delta^t W_{G,t},$$

where only the terms $\delta^\tau W_{G,\tau}$ and $\delta^{\tau+1} W_{G,\tau+1}$ are different under policies σ versus σ' . We thus only need to compare these two terms under the two policies.

Suppose for now that $\delta = 1$.

For group A , the sum $\delta^\tau W_{A,\tau} + \delta^{\tau+1} W_{A,\tau+1}$ is the same under policies σ and σ' .

For group B , on the other hand, $\delta^\tau W_{B,\tau} + \delta^{\tau+1} W_{B,\tau+1}$ is strictly greater under plan σ than under σ' . To see this, note that under σ

$$\delta^\tau W_{B,\tau} + \delta^{\tau+1} W_{B,\tau+1} = \delta^\tau |B| \int_0^1 \omega_\tau^*(g(c)) f_{B,n_\tau}(c) dc + \delta^{\tau+1} |B| \int_0^1 \omega_{\tau+1}^*(c) f_{B,n_{\tau+1}}(c) dc$$

while under σ'

$$\delta^\tau W'_{B,\tau} + \delta^{\tau+1} W'_{B,\tau+1} = \delta^\tau |B| \int_0^1 \omega_\tau'^*(c) f_{B,n'_\tau}(c) dc + \delta^{\tau+1} |B| \int_0^1 \omega_{\tau+1}'^*(g(c)) f_{B,n'_{\tau+1}}(c) dc.$$

The fact that $\delta^\tau W_{B,\tau} + \delta^{\tau+1} W_{B,\tau+1} > \delta^\tau W'_{B,\tau} + \delta^{\tau+1} W'_{B,\tau+1}$, when $\delta = 1$, follows from the facts that $\omega_{\tau+1}^*(c) = \omega_\tau'^*(c) = c$, that $\omega_\tau^*(g(c)) = \omega_{\tau+1}'^*(g(c))$, that $f_{B,n_\tau}(c) = f_{B,n'_{\tau+1}}(c)$ and that $f_{B,n_{\tau+1}}(c) \succ f_{B,n'_\tau}(c)$.

By continuity, it then follows that there exists $\bar{\delta} \in (0, 1)$ such that for all $\delta \in (\bar{\delta}, 1)$, the total welfare is also higher under plan σ than under σ' . ■

Therefore, when δ is high enough, it follows by iterative application of Lemma 3 that the optimal policy has a threshold form $\hat{\sigma}_t = 1$ for $t < \bar{T}$ and $\hat{\sigma}_t = 0$ for $t \geq \bar{T}$ for some $\bar{T} \in \mathbb{N} \cup \infty$.

We will now rule out the case where $\bar{T}$ could be infinite and thus show that $\bar{T} \in \mathbb{N}$.

Let us thus compare the welfare of some (large) $\bar{T} < \infty$ to that of the case $\bar{T}' = \infty$. In what follows, the quantities with a prime (') will be the ones associated to $\bar{T}' = \infty$.

We need to show that

$$\sum_{t=1}^{\infty} \delta^t (W_{A,t} + \lambda_B W_{B,t}) > \sum_{t=1}^{\infty} \delta^t (W'_{A,t} + \lambda_B W'_{B,t}). \quad (3)$$

Equivalently, it will be convenient to multiply the welfare by the constant $\frac{1}{|A|+|B|}$ and verify that

$$\frac{1}{|A|+|B|} \left(\sum_{t=1}^{\infty} \delta^t (W_{A,t} + \lambda_B W_{B,t}) - \sum_{t=1}^{\infty} \delta^t (W'_{A,t} + \lambda_B W'_{B,t}) \right) > 0$$

$$\begin{aligned} \sum_{t=1}^{\infty} \frac{\delta^t}{|A|+|B|} ((W_{A,t} + \lambda_B W_{B,t}) - (W'_{A,t} + \lambda_B W'_{B,t})) &= \sum_{t=1}^{\infty} \delta^t \frac{1}{|A|+|B|} ((W_{A,t} + \lambda_B W_{B,t}) - (W'_{A,t} + \lambda_B W'_{B,t})) \\ &= \sum_{t=1}^{\infty} \delta^t \left(\frac{|A|}{|A|+|B|} \int \omega_t^*(c) f_A(c) dc \right. \\ &\quad + \frac{\lambda_B |B|}{|A|+|B|} \int [\sigma_t \omega_t^*(g(c)) + (1 - \sigma_t) \omega_t^*(c)] f_{B,n_t}(c) dc \\ &\quad - \frac{|A|}{|A|+|B|} \int \omega_t'^*(c) f_A(c) dc \\ &\quad - \frac{\lambda_B |B|}{|A|+|B|} \int \omega_t'^*(g(c)) f_{B,n_t}(c) dc \Big) \\ &\quad + \sum_{t=1}^{\infty} \delta^t \frac{|A|}{|A|+|B|} \gamma_A \int (\omega_t^*(c) - \omega_t'^*(c)) f_A(c) dc \end{aligned} \quad (4)$$

The case $\lambda_B = 1$ is interesting and worth examining first. In that case, note that the first two terms of the right-hand side of Eq. (4) rewrite as

$$\left(\frac{|A|}{|A| + |B|} \int \omega_t^*(c) f_A(c) dc + \frac{|B|}{|A| + |B|} \int [\sigma_t \omega_t^*(g(c)) + (1 - \sigma_t) \omega_t^*(c)] f_{B,n_t}(c) dc \right) = \mathbb{E}_t[c],$$

since the time t average wage under policy $\bar{T} < \infty$ is equal to the time t average performance level under policy $\bar{T} < \infty$ (here denoted by $\mathbb{E}_t[c]$).

Likewise, the third and fourth terms rewrite as

$$-\left(\frac{|A|}{|A| + |B|} \int \omega_t'^*(c) f_A(c) dc + \frac{|B|}{|A| + |B|} \int \omega_t'^*(g(c)) f_{B,n_t}(c) dc \right) = -\mathbb{E}_t'[c],$$

since the time t average wage under policy $\bar{T}' = \infty$ is equal to the time t average performance level under policy $\bar{T}' = \infty$ (here denoted by $\mathbb{E}_t'[c]$).

We then have that the right-hand side of Eq. (4) can be written as

$$\sum_{t=1}^{\infty} \delta^t (\mathbb{E}_t[c] - \mathbb{E}_t'[c]) + \sum_{t=1}^{\infty} \delta^t \frac{|A|}{|A| + |B|} \gamma_A \int (\omega_t^*(c) - \omega_t'^*(c)) f_A(c) dc$$

We must now verify if this is greater than 0. We first make the following observations:

- The first term is negative and converges to 0 as $\bar{T} \rightarrow \infty$. Indeed, $\mathbb{E}_t[c] < \mathbb{E}_t'[c]$ for $t \geq \bar{T}$, since the time t average performance level keeps increasing as affirmative action gets implemented for more periods. This term converges to 0 as $\bar{T} \rightarrow \infty$ since $\mathbb{E}_t[c] = \mathbb{E}_t'[c]$ for $t < \bar{T}$ and $\sup_{t \geq \bar{T}} |\mathbb{E}_t[c] - \mathbb{E}_t'[c]| \xrightarrow{\bar{T} \rightarrow \infty} 0$, reflecting the fact that the improvements in the performance distribution of group B become marginal after a while.
- The second term is positive and bounded away from 0 as $\bar{T} \rightarrow \infty$. Indeed, under a policy of permanent affirmative action $\bar{T}' = \infty$,

$$\begin{aligned} \omega_t'^*(c) &= \mathbb{P}_t'^*(\{aa\}|c) g^{-1}(c) + (1 - \mathbb{P}_t'^*(\{aa\}|c)) c \\ &< c \end{aligned}$$

since $\mathbb{P}_t'^*(\{aa\}|c) > 0$ for all t . Thus, for $t \geq \bar{T}$, $\omega_t^*(c) - \omega_t'^*(c) = c - \omega_t'^*(c) > \Delta$ for some $\Delta > 0$. This captures the gain to group A of stopping affirmative action after a finite number of periods.

From the above observations, we can formally state that $\forall \epsilon > 0$, there exists $\bar{T} < \infty$ large enough and $\bar{\delta}(\bar{T}) \in (0, 1)$ such that $\forall \delta \in (\bar{\delta}(\bar{T}), 1)$

$$\sum_{t=1}^{\infty} \delta^t |\mathbb{E}_t[c] - \mathbb{E}_t'[c]| < \epsilon,$$

and

$$\sum_{t=1}^{\infty} \delta^t \frac{|A|}{|A| + |B|} \gamma_A \int (\omega_t^*(c) - \omega_t'^*(c)) f_A(c) dc > 2\epsilon$$

from which it follows that the right-hand side of Eq. (4) is positive and thus that Eq. (3) is verified.

To complete the proof, we now turn to the case when $\lambda_B < 1$.

First note that when δ is high enough, unsurprisingly, group A gains from stopping affirmative action whereas group B loses. Thus, rearranging the left-hand side of Eq. (4) as follows

$$\sum_{t=1}^{\infty} \frac{\delta^t}{(|A| + |B|)} ((W_{A,t} - W'_{A,t}) + \lambda_B (W_{B,t} - W'_{B,t})),$$

we notice that decreasing the weight λ_B placed on the welfare of group B to values strictly smaller than 1 keeps this quantity positive. We can thus conclude that it will still be worth stopping affirmative action after $\bar{T} < \infty$ periods as opposed to continuing it forever. The first-best optimal policy $\bar{T}_{\lambda_B}$ for some $\lambda_B < 1$ will thus be such that $\bar{T}_{\lambda_B} \leq \bar{T}_{\lambda_B=1} < \infty$. ■

References

- ARROW, K. (1973): “The theory of discrimination. In: Discrimination in labor markets,” *Orley Achenfelter and Albert Rees (eds.), Princeton-New Jersey*.
- BAGWELL, K. (1995): “Commitment and observability in games,” *Games and Economic Behavior*, 8, 271–280.
- BECKER, G. S. (1957): *The economics of discrimination*, University of Chicago press.
- CHUNG, K.-S. (2000): “Role models and arguments for affirmative action,” *American Economic Review*, 90, 640–648.
- COATE, S. AND G. LOURY (1993a): “Antidiscrimination enforcement and the problem of patronization,” *American Economic Review*, 83, 92–98.
- COATE, S. AND G. C. LOURY (1993b): “Will affirmative-action policies eliminate negative stereotypes?” *American Economic Review*, 1220–1240.
- FANG, H. AND A. MORO (2011): “Theories of statistical discrimination and affirmative action: A survey,” *Handbook of social economics*, 1, 133–200.
- KAHNEMAN, D. AND A. TVERSKY (1979): “Prospect Theory: An Analysis of Decision under Risk,” *Econometrica*, 47, 263–292.
- LUNDBERG, S. J. AND R. STARTZ (1983): “Private discrimination and social intervention in competitive labor market,” *American Economic Review*, 73, 340–347.
- PHELPS, E. S. (1972): “The statistical theory of racism and sexism,” *American Economic Review*, 62, 659–661.
- SOWELL, T. (2005): *Affirmative Action Around the World: An Empirical Study*, Yale University Press.
- VAN DAMME, E. AND S. HURKENS (1997): “Games with imperfectly observable commitment,” *Games and Economic Behavior*, 21, 282–308.