

Editors' Summary

THE BROOKINGS PANEL ON Economic Activity held its seventy-eighth conference in Washington, D.C., on September 9 and 10, 2004. This issue of *Brookings Papers on Economic Activity* includes the papers and discussions presented at the conference. The first paper evaluates unconventional measures available to monetary policymakers for stimulating the economy when interest rates are already near zero, a situation that may arise with price stability or negative inflation. The second paper presents empirical evidence on the effects of taxes, federal spending, and deficits on national saving, interest rates, and growth. The third paper explores the impacts on U.S. employment in recent years from conventional foreign trade in goods and from the rise in offshoring of service jobs. The fourth paper examines the effect of tax changes, such as those passed since 2000, on business capital formation.

CENTRAL BANKS USUALLY implement monetary policy by setting the short-term nominal interest rate that the bank controls, such as the federal funds rate in the United States. However, the success of many industrial countries over the years in reducing inflation and, consequently, average nominal interest rates has increased the likelihood that, during a recession, the policy rate will approach its lower bound of zero. When rates are at or near zero, a central bank can no longer stimulate aggregate demand by further rate reductions and must rely instead on “nonstandard” policy alternatives. An extensive literature examines these alternatives, but for the most part from a theoretical or historical perspective. Few studies have presented empirical evidence on their potential effectiveness in modern economies. Such evidence not only would help central banks plan for the contingency of the policy rate approaching zero, but also would bear directly on the choice of the appropriate inflation objective in normal times: the greater the confidence of central bankers that tools exist to help the economy escape

the liquidity trap that occurs at the zero bound, the less need there is to maintain an inflation “buffer.” Hence evidence of effective alternative policies would bolster the argument for a lower inflation objective. In the first article of this issue, Ben Bernanke, Vincent Reinhart, and Brian Sack apply the tools of modern empirical finance to the recent experiences of the United States and Japan to look for such evidence.¹

Following earlier work by Bernanke and Reinhart, the authors group nonstandard policy alternatives into three classes: official communications designed to shape public expectations about the future course of interest rates; quantitative easing, which increases both assets (holdings of government securities) and liabilities (unborrowed reserves) on the central bank’s balance sheet; and changes in the composition of that balance sheet through, for example, targeted purchases of long-term bonds aimed at reducing long-term interest rates.

The authors’ investigation employs two approaches. First, they perform event-study analysis, measuring and analyzing the behavior of selected asset prices and yields over short periods surrounding central bank statements or other financial or economic news. Second, they estimate “no-arbitrage” models of the term structure of interest rates for both the United States and Japan. For any given set of macroeconomic conditions and stance of monetary policy, these models allow the authors to predict interest rates at all maturities. Using the predicted term structure as a benchmark, they are then able to assess whether factors not included in the model—such as the Bank of Japan’s quantitative easing policy that began in 2001—have economically significant effects on interest rates.

Bernanke, Reinhart, and Sack begin with a discussion of nonstandard policies that might be effective in stimulating the economy when short-term rates are at the zero bound; the discussion draws on the historical experience with such policies in the United States and Japan as well as on existing theories of potential policy channels and previous empirical analysis. The first type of policy they consider is the use of central bank communications to influence the market’s expectations about future policy and hence future short-term rates. According to some theories, shaping expectations about future short-term rates is essentially the only tool central bankers have. But the authors take a broader view, arguing that private sector bor-

1. The editors thank Bernanke, Reinhart, and Sack for providing an excellent nontechnical summary of their paper, on which this summary draws extensively.

rowing and investment decisions are more sensitive to longer-term yields than to short-term rates, and considering the possibility that long-term rates can be moved independently of expectations of short-term rates. This leads to a discussion of the potential importance of policy statements, credibility, and policy rules. Although the authors see “rule-like” central bank behavior, particularly state-contingent behavior, as an important means of shaping the public’s policy expectations, they believe that a central bank would find it particularly difficult to establish in advance how it would react to highly unusual circumstances, such as when the short-term rate is near the zero bound. Hence in such cases statements about policy intentions and commitments are likely to be particularly important.

The second type of nonstandard policy, quantitative easing, involves purchasing government securities beyond what is required to drive the short-term rate to zero. The authors discuss three channels through which such a policy might operate to escape the liquidity trap. First, such purchases may lead to private sector rebalancing of portfolios, which in turn would raise the prices of other assets. However, the authors observe that there will be little incentive to rebalance if money is a good substitute for the short-term bills it replaces when the latter are paying close to zero interest. Second, a larger outstanding stock of money raises the prospect of higher seigniorage in the event of future inflation, substituting for direct taxes. The effectiveness of this “fiscal” channel requires that the public, in the midst of a deflation, expect future inflation and expect that the central bank will not withdraw the injected money when that inflation arrives. The authors believe that the fiscal channel could work if pursued aggressively enough and with a clear commitment not to reverse course. Third, the visible signal that quantitative easing provides makes it more believable that the central bank will hesitate to reverse such large purchases soon, perhaps because of the possible shock to money markets.

The third type of nonstandard policy involves altering the composition of the central bank’s balance sheet by participating in all segments of the market in government debt, including inflation-indexed debt, so as to influence term, risk, and liquidity premiums. Using emergency provisions dormant since the 1930s, the Federal Reserve could even accept private financial and real assets as collateral for discount window loans. Although many economists are skeptical about the potential effectiveness of this channel, the authors note a number of historical examples of central banks effectively pegging long-term rates.

Bernanke, Reinhart, and Sack's own empirical investigation begins with an event study measuring the influence of Federal Reserve policy announcements by the response of three market-based indicators following selected decisions of the Federal Open Market Committee (the Federal Reserve's principal policymaking body) since 1991. The indicators are the current-month federal funds futures contract, the Eurodollar futures contract expiring in about a year, and the yield on Treasury securities of five years' maturity. The authors measure the responses of each indicator observed during the forty-five minutes following FOMC announcements; the very short time interval is chosen to minimize the extent to which other factors could be affecting rates. The change in the current-month federal funds futures contract simply measures the markets' reaction to news about the Federal Reserve's near-term funds target. The change in the year-ahead Eurodollar futures contract presumably incorporates both the effect of this funds rate surprise and the effect of any accompanying announcement on market expectations about policy actions and the economy over the coming year. The change in the five-year Treasury yield presumably includes both those effects plus those arising from revisions in expectations beyond a year. The authors decompose the second indicator into the part explained by the first factor, the funds "surprise," and the orthogonal residual, which they label a second factor. The change in the five-year rate not explained by the first two is designated a third factor. The authors find that only about 20 percent of the variance in the one-year-ahead rate during the forty-five-minute policy "window" is explained by the current policy surprise; unexplained movements in the year-ahead rate—the second factor—make up the remaining 80 percent. Again, this factor presumably captures any revisions in the private sector's expectations of future short-term rates due to information contained in the policy statement that accompanies the change in the funds rate.

Perhaps the most striking revelation of the authors' analysis is the high correlation between unexplained movements in the future one-year rate and the five-year Treasury yield: the second factor accounts for 68 percent of the variability of the five-year yield during the event window, and the policy surprise itself explains another 12 percent, leaving only 20 percent unexplained. Informal inspection of the historical behavior of the second factor reveals that it becomes increasingly important in the latter part of the sample, when policy statements came into regular use. In contrast, larger realizations of the first and third factors do not seem to line up with

dates of policy statements. This leads the authors to a more formal investigation of the link between FOMC statements and the three factors. First, they regress the squared values of each of the factors on dummy variables indicating dates when statements were issued and the characteristics of those statements: The dummy variable STATEMENT takes a value of 1 on any date on which a statement was released. STATEMENT SURPRISE takes a value of 1 when, in the authors' judgment, the statement included information about the economy or the path of policy that most market participants would not have expected. Finally, PATH SURPRISE is set equal to 1 when, in the authors' judgment, the statement revealed new information about the likely future path of monetary policy. The authors attempt to assign "surprises" as objectively as possible, using commentaries written before and after each statement was released (including those of a leading financial firm that specializes in monitoring FOMC actions), internal Federal Reserve staff analyses of market reactions, pre-FOMC meeting surveys about expectations for the balance-of-risks part of the statement, and the results of a survey of expectations about the statement conducted by the New York Federal Reserve Bank. Of the 116 policy decisions in their sample, statements accompanied 56, and the authors identify 31 of them as having a significant element of surprise, and 9 of these, in turn, as path surprises.

In the regression explaining the first factor, the coefficient on STATEMENT is positive and significant, and the authors attribute this result to the fact that, for much of the sample, statements were made only on days when the federal funds target was changed. These are days on which the policy rate surprise tends to be relatively large. The coefficient on STATEMENT SURPRISE in this regression is negative and significant, suggesting that the FOMC viewed policy rate surprises and statements as substitutes, possibly because the FOMC was reluctant to issue surprising statements at the same time that it was also surprising the markets with the policy action. The PATH SURPRISE variable is insignificant.

The regressions explaining the second and third factors are of more interest, since they provide information about how FOMC statements affect market expectations. The mere issuance of a statement has essentially no effect on the variance of the second factor; in contrast, in periods when there is a statement surprise, the variance is nearly 200 basis points, and when there is also a path surprise, the variance is roughly 230 basis points. When the first two factors are controlled for, the five-year yield is not

noticeably affected by policy statements, surprising or otherwise; thus statements do not appear to provide information about long-term yields independent of their influence on expectations about rates one year in the future.

FOMC statements often contain language that suggests the direction of future policy actions. Bernanke, Reinhart, and Sack investigate whether the direction of the response of the factors is consistent with whether a “surprise” statement appears hawkish, dovish, or neutral with regard to interest rates. They define two dummy variables, for statement and path surprises, each of which takes a value of +1 for surprises that are judged to be hawkish, -1 for those judged to be dovish, and zero otherwise. They find no significant response of either the first or the third factor to either of these dummies. However, they find that hawkish statement surprises increase, and dovish surprises decrease, the year-ahead rate by 12 basis points. The response to path surprises is an even greater 16 basis points. Both responses are highly significant. The authors note that, in addition to the official FOMC statements they study, the speeches and congressional testimony of FOMC members may be important in shaping expectations.

Beyond their direct effects, statements that are conditional—that is, that make the central bank’s commitments contingent on specific economic developments—are likely to affect the market response to news about those events. For example, the statement that “[monetary] policy accommodation can be maintained for a considerable period,” first introduced in Federal Reserve Chairman Alan Greenspan’s semiannual report to Congress in July 2003, was, in subsequent FOMC statements, typically tied to labor market conditions and “slack” in the economy. A regression relating changes in the ten-year Treasury yield to surprises in the monthly payroll report shows greater sensitivity after August 2003 than in the preceding twelve years; this is consistent with the claim that this FOMC language heightened the market’s attention to employment growth.

Although the event studies succeed in isolating the effects of FOMC policy announcements, they measure only very short term effects and do not take advantage of the restrictions linking yields of various maturities normally incorporated in term structure equations. Particularly when one is evaluating unusual policies, such as buybacks of longer-term Treasuries, it is useful to measure changes in yields against a “benchmark” term structure that incorporates these linkages and takes into account the observable

macroeconomic factors that can affect the level and shape of the yield curve. For this purpose the authors use a “no-arbitrage” term structure model in which long-term yields are determined by two components: the expected future path of one-period interest rates, and a term premium in yields at various maturities that investors require as compensation for the risk in holding longer-term investments. The authors estimate a quarterly vector autoregression (VAR) for five observable factors that, taken together, provide a reasonable summary of economic conditions relevant to the term structure: the employment gap, the previous year’s core inflation, expected inflation for the coming year (taken from the Blue Chip survey), the federal funds rate, and the year-ahead Eurodollar futures rate. The first component—the expected future path of one-period interest rates—can be computed by simply iterating the estimated VAR forward. The risks that are priced in the second component—the term premium—are the innovations in the VAR process, with prices that are assumed to depend only on the contemporaneous values of the variables in the VAR. The authors estimate these prices by fitting the model to zero-coupon Treasury yields at maturities of six months and one, two, three, four, five, seven, and ten years. The estimated effects of risks on the term structure are of interest in themselves. The authors find that the term premium for longer maturities has declined over time, presumably reflecting greater stability in the economy and in policy. The estimated model predicts the Treasury yield curve reasonably well at all maturities: the standard deviation of prediction error is 33 basis points at six months and increases to around 80 basis points for longer maturities. The authors’ results show that the year-ahead futures rate makes an important contribution to the accuracy of the model’s predictions, reducing the standard deviation of prediction error for the one-year yield by about 30 basis points, or 35 percent; smaller contributions are found for longer maturities, decreasing from about 20 basis points for the two- and three-year yields to a mere 7 basis points for the ten-year yield.

How large are the innovations in the year-ahead futures contracts according to the VAR, how does their magnitude compare with the earlier event-study estimates of the impact of policy surprises, and by how much do these innovations shift Treasury yields at various maturities? The authors show that the event-study estimates of the effect of policy on the futures rate are only a small fraction of the VAR innovations before July 2003, but that since then the standard deviation of VAR innovations has fallen and that of

event-study shocks has risen, amounting to roughly 45 percent of the VAR innovations. Some of this difference may reflect communications effects not captured in the statements, and some must reflect responses of expectations to developments in the economy unrelated to FOMC communications. The authors examine in some detail the comparison of VAR innovations and event-study shocks during the period from August to December 2003, during which the FOMC used the “considerable period” language. During that period the results from the event study suggested that FOMC communications pushed down the Eurodollar futures rate by a cumulative 19 basis points, whereas the VAR innovation lowered that rate by 63 basis points. The model predicts that a 63-basis-point decline in the futures rate would shift the yield curve down 25 basis points at a two-year horizon and 7 basis points at a ten-year horizon. The authors guess that the true communications effect is somewhere between these two extremes.

The effectiveness of the third type of nonconventional policy, changes in the composition of the central bank’s balance sheet, depends on whether substitution among assets is sufficiently imperfect that large purchases of a specific class of asset might affect its yield relative to the short-term interest rate. If composition matters, the FOMC has the ability to stimulate the economy even when the federal funds rate is constrained by the zero floor. Because the Federal Reserve has not undertaken large purchases of longer-term assets in recent years, the authors instead examine three recent episodes in which it seems plausible that many market participants came to anticipate significant changes in the relative supplies of Treasury securities. The first of these was the Treasury’s announcement in 1999 of a plan to buy back government debt in the face of prospective budget surpluses; the second was the investment in Treasury securities by Asian official institutions intervening in exchange markets after the currency crises of 1998; the third was in the summer of 2003, when some financial market participants came to believe the Federal Reserve might purchase long-term Treasury securities in order to combat incipient deflation. For each episode the authors examine the movement of a number of yields in narrow windows surrounding important announcements and the movements of the residuals from the benchmark term structure model.

The dating of each episode is necessarily imprecise. The gradual elimination of government interest-bearing debt began with the surpluses that emerged in the early 1990s, and by the end of the decade many observers were forecasting that Treasury debt would disappear by 2010. Bernanke,

Reinhart, and Sack focus on two events that they believe were particularly salient in investors' minds: the mid-quarter refunding announcements of February 2000, when the under secretary of the Treasury suggested that the ten-year note would replace the thirty-year bond as the benchmark long-term security; and the November 2001 refunding announcement, when the Treasury confirmed it would stop selling the long bond. The authors report that in both cases the Treasury yield curve rotated down dramatically. Examination of their estimated term structure model that controls for variations in the economy and policy shows that the yield on twenty-year bonds dropped roughly 80 basis points below what was expected for this period. The authors view these results as only suggestive of the efficacy of decreasing the relative supply of long-term government debt, recognizing that the term structure model is unlikely to capture all the other factors that changed expectations about the supplies of bonds, and that it is hard to know the probability that investors assigned to a sizable paydown. They also note that the reaction of yields on long-term Treasuries did not translate into a reduction in interest rates on privately issued debt—long-term swap spreads widened noticeably during the period.

The beginning of the accumulation of Treasury securities by foreign official institutions is likewise difficult to date with any precision. Since 1998 the value of securities held in custody for foreign governments at the New York Federal Reserve Bank has almost doubled, reaching \$1¼ trillion. Japanese authorities in particular intervened heavily in foreign exchange markets in 2003, when their purchases totaled \$177 billion, and the first quarter of 2004, when they purchased \$138 billion. Regressions of changes in two-, five-, and ten-year Treasury yields on the volume of dollar interventions, which the authors assume were immediately recognized by market participants, show small (less than 1 basis point per \$1 billion of intervention) but statistically significant changes in the three-day period beginning with the intervention. Interestingly, swap spreads did not increase during this period, suggesting that the effects on the Treasury yields were transmitted to yields on private debt. Yields on five- and ten-year Treasuries remained 50 to 100 basis points below the predictions of the benchmark term structure model during this period. Although these findings are suggestive, the authors note that yields had begun to move down before the sizable Japanese interventions, perhaps because of fears of deflation.

The authors refer to the period between the fall of 2002 and the summer of 2003 as the “2003 deflation scare,” a period when various Federal

Reserve officials spoke explicitly about the risks of deflation and the possible response of monetary policy should the funds rate approach its lower bound of zero. Although officials consistently described this possibility as remote, the authors observe that some market participants believed the Federal Reserve took the danger seriously enough to consider purchasing large amounts of longer-term Treasuries. The perceived likelihood seemed to peak with the May 2003 FOMC meeting statement. Such action was seen to be “taken off the table” when the June FOMC statement did not mention deflation and when Chairman Greenspan, in his July testimony, said, “situations requiring special policy actions are most unlikely to arise.” The ten-year Treasury yield moved sharply with these events, falling 20 basis points with the May FOMC statement and rising abruptly with the June statement (10 basis points) and the chairman’s July testimony (20 basis points). Residuals from the benchmark term structure were consistent with these responses. The authors caution against making too much of these results, particularly since the period of the “scare” overlapped that of large-scale purchases of Treasuries by foreign official institutions. But they express confidence that, if the Federal Reserve were willing to purchase an unlimited amount of a particular asset, they could establish the asset’s price.

The authors’ analysis of the recent experience in Japan focuses on two nonstandard policies recently employed by the country’s central bank, the Bank of Japan. The first is the zero-interest-rate policy, under which the central bank, beginning in 1999, committed to keep its policy rate, the call rate, at zero until deflation had been eliminated; the second is the quantitative easing policy, announced in 2001, which consists of providing bank reserves at levels much greater than needed to maintain a policy rate of zero. The evidence for the effectiveness of these policies is more mixed than in the case of the United States. The authors’ event-study analyses, which may be less informative in Japan because of small sample sizes, do not show any reliable relationship over the past few years between policy statements by the Bank of Japan and markets’ one-year-ahead policy expectations. (The latter are measured in Japan’s case as the portion of movements in the nearest Euroyen futures contract not explained by unexpected changes in the current policy setting.) But the residuals from an estimated term structure model for Japan similar to that used for the United States make a stronger case for nonconventional policies. The Bank of Japan’s use of quantitative easing, with which the United States has no recent expe-

rience, appears to have lowered longer-term yields, as did the zero-interest-rate policy. Simulations of the model also indicate that interest rates at all maturities were noticeably lower under both nonstandard policies than they would have been otherwise.

Bernanke, Reinhart, and Sack believe that their results provide some grounds for optimism about the likely efficacy of two of their nonstandard policies. In particular, they confirm a potentially important role for central bank communications in shaping public expectations of future policy actions. In the United States, market expectations about the trajectory of the federal funds rate over the next year appear strongly linked to Federal Reserve policy statements. If such central bank “talk” does affect policy expectations, then policymakers retain some leverage over long-term yields, even if the current policy rate is at or near zero. Based on the recent episodes they examine, the authors further suggest that large changes in the relative supplies of securities may have economically significant effects on their yields.

Yet despite their relatively encouraging findings concerning the potential efficacy of nonstandard policies at the zero bound, the authors are cautious in making policy prescriptions because of the considerable uncertainty about the size and reliability of these effects. They therefore counsel a conservative approach—maintaining a sufficient inflation buffer and applying preemptive easing as necessary to minimize the risk of hitting the zero bound. But, recognizing that such policies cannot ensure that the zero bound will never be hit, they argue that refining our understanding of the potential usefulness of nonstandard policies for escaping the zero bound should remain a high priority.

IF, IN PRACTICE, FISCAL policy were conducted so as to keep federal budget balances within a moderate range, the implications of sustained large deficits would be of only academic interest. And in fact the budget balance averaged a modest deficit during the first three postwar decades, with mostly countercyclical fluctuations around that average. Fiscal discipline has eroded sharply twice since then, first in the early 1980s and then, after recovering, again in the early 2000s. The standardized budget balance, which adjusts for the effects of the business cycle, reached 5 percent of GDP in the mid-1980s and recovered to about half that size by the early 1990s, after the second Reagan tax reform. The deficit was steadily reduced between 1992 and 2000, turning into a surplus by the end of the period. At

that point analysts found themselves wrestling with the problems that a vanishing national debt might bring. Today, through a combination of large tax cuts, weaker growth, higher defense spending, and the end of windfall revenue from the booming stock market, the budget is over \$400 billion in deficit, and most private analysts project that deficits will remain a large fraction of GDP for the indefinite future under likely policy scenarios—a situation without precedent in U.S. economic history. In the second article of this issue, William Gale and Peter Orszag analyze the likely effects of tax cuts and deficits and apply their findings to the fiscal situation under present and projected policies.

The authors examine the consequences of deficits for the economy by conducting a careful empirical analysis of two channels through which those consequences should operate. The first is through saving, where the issue is the extent to which government saving or dissaving (the government budget balance) affects national saving, which is the sum of government and private saving. In what Gale and Orszag call the conventional case, private saving offsets little or none of a reduction in government saving, and so future national income and GNP are reduced by the returns forgone due to the reduction in national saving and investment. In the polar alternative, which, following the literature, the authors refer to as the Ricardian equivalence case, an increase in the deficit resulting from a tax cut induces an equal increase in private saving, so that national saving and investment are unchanged, and future GNP is unaffected. The second channel works through interest rates. If, as in the conventional case, the deficit reduces national saving, and that in turn raises rates, domestic investment will be reduced, lowering future GNP. In the absence of increased foreign investment into the United States, this will result in a corresponding reduction in GDP. To the extent foreign investment makes up for the reduced national saving, domestic investment, the domestic capital stock, and GDP will be unaffected, but national income and GNP will fall relative to GDP by the amount of the higher earnings paid to foreigners on their investment. In the Ricardian equivalence case, since there is no effect of tax cut-induced deficits on national income, deficits are not predicted to affect interest rates.

Gale and Orszag survey the extensive empirical literature that uses time-series data to estimate the effect of deficits on private saving; they also reexamine and extend some of the earlier work and perform their own, new tests. They first reestimate the consumption function used in a 1990 paper by Roger Kormendi and Philip Meguire, which relates aggregate consump-

tion to the various flows that make up disposable income—NNP (net national product, equal to GNP minus capital depreciation) plus transfers and government interest payments, minus taxes and retained earnings of corporations—and to private net worth, the federal budget surplus, the stock of government debt, and government spending on goods and services. These authors and others interpret zero coefficients on taxes and transfers in this equation as support for the Ricardian case. Gale and Orszag note that this interpretation rests on the questionable assumption that neither changes in taxes nor changes in transfers are related to changes in future government purchases, so that higher taxes or lower transfers today imply lower taxes or higher transfers in the future. Nonetheless, they agree that the coefficients on taxes and transfers are relevant to the policy issue of how deficits affect future economic performance. By extending the data period to include the 1990s and the start of the 2000s, the authors can examine how the specification fits outside the original data period. These additional years may be particularly informative because deficits have fluctuated widely.

Using annual data, the authors estimate the consumption function for two measures of the dependent variable: the first (that used by Kormendi and Meguire) includes services of durables in consumption in addition to spending on nondurables and services; the second is the simple sum of nondurables and services, scaled up to preserve the mean level of total consumption. They try three transformations of the dependent and independent variables: first differences, first differences scaled by last year's NNP, and first differences of ratios to NNP. And they try some specifications of the independent variables that differ from those in previous work: separating federal from state and local spending, and adding marginal tax rates constructed so as to largely reflect legislative changes. The separation of federal from state and local taxes is motivated by the fact that state and local revenue come largely from sales taxes, which vary positively with consumption, and this change proves to be the most important. When the authors use Kormendi and Meguire's specification of the variables in ordinary least squares (OLS) regressions and end the sample period in 1992, they replicate the earlier authors' finding of no significant effect of taxes on consumption. However, that finding is reversed when the data period is extended, with the estimated effect of federal taxes largest when the estimation extends through 2002. For this period, and with the refined variables just described, estimates of the consumption response to federal tax changes in the year they occur range from

–0.34 to –0.46, with additional lagged consumption responses in subsequent years. The effects of transfers on consumption are generally even larger. These results with the Kormendi-Meguire equation thus overturn those its originators found.

Gale and Orszag turn next to an Euler equation formulation for estimating consumer behavior. In this formulation, as first developed by Robert Hall, farsighted and rational consumers who are not subject to borrowing constraints smooth their marginal utility over time; the path of their consumption is a random walk with drift. Building on modifications suggested by other researchers and their own work, the authors introduce a number of other possible influences on consumption. They allow for a fraction of consumers to be constrained by liquidity, so that their consumption is limited to current-period disposable income. They allow for wealth effects and for the possibility that consumers distinguish between private net worth and government bonds and that they adjust their consumption for expected changes in government purchases. They also allow for the possible incentive effects of marginal tax rates on personal and corporate income. The authors first examine this main case and then explore the effects of adding other explanatory variables that might be expected to influence consumption.

In the main Euler equation specification, when variables are expressed as the first difference in ratios to NNP, OLS estimates of tax effects are larger than those found using the Kormendi-Meguire specification. The coefficients range between –0.5 and –0.7, depending on whether marginal tax rates are included in the estimation equation. They are similar to the Kormendi-Meguire results for the other two expressions of the variables.

To address the possible bias from the endogeneity of the right-hand-side variables in OLS estimation, Gale and Orszag also use values of most of the right-hand-side variables with two or three lags as instruments for their current values, and similarly lagged values of consumption as an instrument for current income. They also present estimates in which all the explanatory variables are instrumented, and estimates in which private wealth and government bonds, measured at the end of the previous period, are not instrumented. The authors find the effect of federal taxes on consumption to be large and significant in all the regressions. With all variables included and instrumented, the coefficients are –1.0 and –0.9 for the two measures of consumption in the first-differences regressions, –0.6 and –0.5 in the first-differences-of-ratios regressions, and –0.7 in regressions with either consumption variable in ratios. The authors see these results as con-

clusive evidence of a substantial effect of changes in current taxes on current consumption. Before-tax income is the only other variable that significantly affects consumption across all the instrumental variables regressions. With two of the three transformations of the variables, its coefficient has a similar size (and opposite sign) as federal taxes; with the third, the transformation using first differences of ratios to NNP, before-tax income has a coefficient of about 0.9.

Turning to the second channel of fiscal policy—the impact of government deficits on interest rates—Gale and Orszag take pains to avoid contaminating their estimates with the well-established cyclical interrelationships among output, interest rates, and deficits, and to focus on fiscal developments to which forward-looking financial markets would be expected to respond. Their review of recent research confirms the downward bias that results from estimations that ignore these problems. For their main regressions the authors focus on the real interest rate on ten-year Treasury bonds for the period five years ahead, and they use five-year-ahead projections of several federal fiscal variables, all as percentages of GDP, as the principal explanatory variable. The fiscal variables are publicly held debt, the unified budget deficit, the primary deficit (which excludes interest payments), and revenue and primary outlays. Real interest rates are calculated using the zero-coupon yield curve for the period five to fourteen years ahead, adjusted for expected inflation. The fiscal variables are taken from projections of the Congressional Budget Office. All regressions also include expected GDP growth over the relevant period and a constant term.

The OLS estimates in the simplest model, using annual data for 1976–2004, show substantial and statistically significant effects on future interest rates using either the debt variable or any of the deficit variables: each of the deficit variables explains roughly twice as much of the variation in rates as does debt. A sustained increase in the unified deficit equal to 1 percent of GDP raises the forward long-term interest rate by 29 basis points. The same increase in the primary deficit raises the interest rate by 40 basis points. Consistent with this result, the same changes due solely to reduced revenue or solely to higher primary outlays raise rates by 42 and 37 basis points, respectively. The authors also add to the equation terms interacting their fiscal variables with dummy variables indicating recessions. When these are included, the estimated effects of the fiscal variables on interest rates rise modestly, suggesting that some business cycle effects remain in

the initial estimates. Adding an assortment of other variables, some suggested by earlier research, reduces the coefficients on the fiscal variables somewhat, but their interpretation is not clear. The variable with the most significant effects, which the authors label the equity premium, is calculated from real long-term rates and GDP growth, which are already on the left- and right-hand sides of the estimating equation.

Gale and Orszag go on to estimate a number of additional specifications as a way of checking the robustness of their main results. They use the debt and unified deficit variables together in equations explaining forward real long-term rates, taking the projected debt at the end of four years and the projected deficit in the fifth year to avoid double counting. In both OLS and autoregressive moving average regressions for this specification, deficits always dominate debt: the coefficient on the latter has the wrong sign and is insignificant, and the deficit coefficient is larger than in the regressions using deficits alone. In regressions explaining current, rather than forward, real long-term interest rates, the estimated effects of the fiscal variables are, not surprisingly, somewhat smaller than in the main equation. In regressions explaining nominal forward long-term rates and including expected inflation as an additional explanatory variable, the latter consistently has a coefficient near 1.0, and the deficit effects are close to those in the main regressions. The least successful specification attempts to explain current nominal long-term rates. In these regressions the coefficient on expected inflation, the most significant explanatory variable, ranges nearer 2.0 than 1.0; the coefficients on the fiscal variables are smaller and often insignificant. The authors also vary the data period and demonstrate that their main results are not sensitive to the exclusion of either 1976-81 or 2000-2004, two periods that some analysts view as atypical.

Gale and Orszag conclude by addressing the current U.S. fiscal situation. They reason that realistic budget projections should assume that the 2001 and 2003 tax cuts are made permanent and that the alternative minimum tax is amended so as to leave only 5 million households affected by it. (They do not include the president's plan to partially privatize the Social Security system, which would significantly further enlarge deficits for decades.) Amending the official Congressional Budget Office projections accordingly, they calculate that the projected unified budget deficit will average 3.5 percent of GDP over the coming decade. Their empirical results imply that, compared with a decade of balanced budgets, sustained deficits of this size will reduce national saving by 2 to 3 percent of GDP and will

raise long-term interest rates by 80 to 120 basis points. At the end of ten years, the lower saving rate will have reduced the assets owned by Americans by 20 to 30 percent of GDP from what they would have been. And, at a 6 percent rate of return on capital, this smaller capital stock would reduce national income by 1 to 2 percent by 2015 and in each year thereafter.

NONFARM PAYROLL EMPLOYMENT continued to decline for nearly two years after the recession trough in November 2001, and by mid-2004 it was barely back to its trough level. Employment in manufacturing, the sector most affected by international trade, has been much weaker. As of November 2004 it was 3.2 million below its March 1998 peak and still 1.4 million below its level at the recession trough. The fact that this weak employment recovery was accompanied by a large expansion in the trade deficit has led many to identify trade as the cause. What is more, the substantial public attention given to the offshoring of jobs to India in recent years has raised the specter that jobs in the services sector, traditionally less vulnerable than manufacturing jobs to foreign competition, might now be increasingly exposed. Although economists typically focus on the long-run gains from trade for the nation as a whole, the loss of important markets to foreign competition is costly to the affected workers and firms and adds to political pressure for trade protection. In the third article of this issue, Martin Baily and Robert Lawrence examine employment in the current recovery and analyze the role of foreign trade and of services offshoring in its performance.

The authors first review developments in manufacturing during the three years ending in 2003, the last for which detailed annual industry data are available. Over this period, manufacturing employment declined 16 percent, far more than the 1.4 percent decline in employment of the nonfarm business sector. The authors focus on two measures of the importance of trade to manufacturing's performance in this period. Measured as a share of manufacturing value added (the portion of GDP originating in the sector), the manufacturing trade deficit rose from 21.3 percent to 28.3 percent; as a share of gross output—sales by manufacturing to other sectors—the deficit rose from 11.9 percent to 15.6 percent. The explanation for this development lies more in export weakness than in import strength: manufactured exports fell 8.8 percent, while manufactured imports rose only 2.3 percent. Only one of nineteen broad categories of manufacturing industry, primary metals, saw its trade balance improve, and all nineteen suffered employment declines,

the largest occurring in computer and electronic products, machinery, and fabricated metal products.

To examine these developments more closely at the industry level, Baily and Lawrence make use of the U.S. input-output tables for 1997, the most recent ones available from the Bureau of Economic Analysis at the requisite level of detail. Typically a sector's output goes both for final use and as an input to production in other sectors. For a dollar of final use of any product—consumption, exports, investment, or government expenditure—the coefficients in the input-output tables permit calculation of the required imports and output from each domestic sector, including the imports and domestic output used as inputs to other sectors. Both industry and final-use categories are specified at a highly detailed level and are taken from (three-digit) North American Industry Classification System trade data, plus domestic use. The authors aggregate the detailed industry categories in the input-output tables into the nineteen two-digit manufacturing industries, transform the gross output coefficients to value-added coefficients using the 1997 ratios, and calculate employment growth for each industry as the difference between growth in its value added and growth in its productivity. This allows them to trace the effect of any change in final use—exports less imports plus domestic spending—on employment in manufacturing and in each industry.

The use of input-output tables for calculating value added by industry reveals a picture of trade that may surprise observers accustomed to looking only at the direct exports of each industry. Because the input-output analysis traces through all intermediate as well as final uses of each industry's output, it accounts for the indirect exports of an industry (goods sold to other domestic firms for use as inputs in export production) as well as its direct exports. It also accounts for the direct and indirect import content of the industry's production, including its production for export. For 2003, this accounting reveals that primary metals is the most export-intensive U.S. manufacturing industry, with 54 percent of its jobs coming from exports. Shares for other heavily export-dependent industries are 37 percent in computers, 30 percent in textiles, 28 percent in chemicals and in machinery, 26 percent in "other transportation" (mainly aircraft), and 25 percent in petroleum and coal products.

Turning to manufacturing as a whole, where employment declined by 2.74 million between 2000 and 2003, Baily and Lawrence calculate the employment changes attributable to exports and imports. They take into

account both the changes in exports and imports and the employment implications of productivity increases in the exporting and import-competing sectors. For a given level of exports, rising productivity reduces over time the number of “export-created” jobs. Similarly, raising productivity reduces the number of domestic jobs displaced by a given level of imports. Using this accounting, Baily and Lawrence attribute a drop of 742,000 jobs, or 28 percent of the total decline, to weak export growth (exports grew more slowly than productivity in those sectors producing exports) and a rise of 429,000 jobs to imports (imports grew more slowly than productivity in import-displaced sectors). Combining the export and import figures results in only 12 percent of the total decline of manufacturing employment being attributed to trade. The remaining 88 percent is associated with weaker domestic demand, the residual category in the allocation.

The authors recognize the conceptual issues that complicate any such allocation and that alternative assumptions would lead to somewhat different interpretations of the employment decline. For example, a different decomposition could attribute all of the employment decline to productivity, because productivity rose about as much as output in this period. Or, since imports respond predictably to changes in domestic demand, one could attribute to domestic demand the change in imports that the change in domestic demand would predict, and identify only the unpredicted part as the “shock” from imports. A decomposition closer to public perceptions would relate to the number of jobs lost because of lower exports or higher imports without taking into account the implications of domestic productivity growth. The authors briefly discuss such alternatives but regard their own calculations as the most useful for examining the role of trade in the employment decline.

The predominance of weak domestic demand is again evident when the authors examine the nineteen individual manufacturing industries. In this analysis only chemical products experience an employment increase due to increased domestic use. The importance of weak exports rather than strong imports is also evident in the individual industry results. Several industries experienced substantial employment declines due to weak exports. None experienced large employment declines due to imports, and eleven of the nineteen experienced increased employment because imports rose by less than productivity.

Baily and Lawrence next examine several potential explanations for the weakness in exports. Although world output growth slowed in the period

they examine, accounting for some of the slowdown in U.S. export growth, they show that the decline in the United States' share of world trade also contributed importantly. Measured in dollars, the U.S. share of world exports rose in the 1990s, particularly after 1995, but has fallen sharply since 2000. When U.S. exports are valued by an index of other major currencies, this pattern is still evident, although not as strongly. The authors go on to break down the changes in U.S. exports into four components. The first reflects the change in exports that would have been required to maintain a constant U.S. share of world trade, and the second and third show the effects of changes in the composition of U.S. exports by commodity and destination, respectively. The fourth, residual component includes the effects of changes in competitiveness, along with any measurement errors and other factors not captured by the first three.

By the authors' calculations, keeping the U.S. share of world trade constant at its 2000 level would have required a \$152 billion increase in exports rather than the \$46 billion decline that occurred. Changes in the commodity composition of exports had little effect, but changes in destination country composition predicted (coincidentally) about a \$46 billion decline. The residual is a decline of \$156 billion, which is presumed to mainly reflect deteriorating competitiveness. Using established rules of thumb to trace the effects of exchange rate variations through time, the authors conclude that the strength of the dollar through much of this period can account for most, if not all, of this residual weakness in exports and the negative impact of trade on employment.

Foreign competition in manufactured goods has been a fact of economic life since countries first began to industrialize in the eighteenth century. The shifting of U.S. service jobs "offshore" is a recent development, however, made possible by the technological revolution in communications. To gain perspective on the importance of offshoring to the U.S. job market, Baily and Lawrence focus on jobs offshored to India, the country benefiting most from this new development. They start by confronting a serious data problem. Whereas economic analysis of goods trade can draw on a wealth of detailed data going back several decades, data on offshoring of service jobs are relatively scarce, and the two main sources tell seemingly very different stories. One source is NASSCOM, a trade association for emerging business services industries in India; the other is the Bureau of Economic Analysis (BEA), whose conventional trade data include data on services imports and exports between the United States and India. The BEA

data, which do not detail the types of services traded, show that services imports from India more than doubled between 1995 and 2000, to \$1.9 billion, and then fell to \$1.7 billion in 2002. But the NASSCOM data indicate that exports to the United States in just two categories—software services (such as computer programming) and business process services (such as call centers and back office processing)—totaled \$6 billion that year, or more than three times the BEA total. Possible sources of error in the BEA data include the difficulty of distinguishing hardware from the software bundled with it; the recording of packaged software as a good rather than a service; and the BEA's reliance on infrequently revised company surveys, which may cause it to miss services imports by sectors not traditionally associated with foreign trade.

The authors also observe that the decline in services imports recorded by the BEA in recent years is not consistent with the growing visibility of such imports. Some might consider it an error in the NASSCOM export data that some fraction of the services reported as Indian exports are actually performed in the United States—for example, by workers on assignment to U.S. client firms. However, they find no support in U.S. immigration data for an influx of Indian temporary workers large enough to explain the rise in employment reported by NASSCOM. And they note that such workers would not be counted in the U.S. payroll employment data they are trying to explain, and hence should be counted as displacing U.S. jobs regardless of where their work is located. The authors therefore choose to base their analysis on the NASSCOM data, which, if anything, set an upper bound on jobs lost to offshoring.

According to NASSCOM, employment in Indian services for export to the United States rose by 91,500 a year between (roughly) 2000 and 2003, divided about evenly between software and business processing services. Although this is only a small fraction of average annual growth in total U.S. services employment (2.1 million) during the 1990s, it is more than a quarter of the slow growth in U.S. services employment in 2000–03. Comparison with U.S. employment trends in services occupations that are closely related to those being offshored is complicated by the surge of jobs associated with Y2K. But, for the period from 1999 to 2003, employment rose by an average of 58,000 a year in high-wage IT occupations and fell by an average of 107,000 a year in low-wage IT-enabled occupations—categories comparable to software and business processing services, respectively, in the Indian data. Although each Indian job gained

corresponds to one U.S. job lost only if productivity in both jobs is the same, offshoring has clearly been an important cause of the weak job markets in these sectors, particularly in the low-wage occupations.

To examine the kinds of changes in trade and in the composition of domestic output that can be expected from offshoring in the longer run, the authors use a macroeconomic model developed by the consulting firm Macroeconomics Advisers, together with projections by Forrester Research that indicate an additional 3.1 million service jobs will be lost to offshoring by 2015. The baseline macroeconomic model assumes that lower fiscal deficits and dollar devaluation will reduce the U.S. current account deficit to 0.5 percent of GDP in that year, with monetary policy adjusting so as to maintain full employment. The authors adjust the baseline for the assumed loss of 3.1 million additional service jobs to offshoring and examine what difference it makes. The authors' preferred way of modeling this rise in offshoring is through a shift in foreign supply resulting from improved productivity abroad that lowers the price the United States pays for services imports. In this case productivity growth quickens, and, by 2015, GDP is 2.6 percent higher and manufacturing employment 62,000 higher, all relative to baseline. Real compensation rises and real profits rise even more. Although the authors see these projected changes as no more than illustrative, they note that they are qualitatively consistent with the changes one would expect from a more conventional opening of trade such as might arise from reducing trade barriers.

In concluding, Baily and Lawrence broaden their field of view to consider the possible role of factors other than trade in the disappointing recent growth of U.S. employment. Productivity growth has been rapid in this same period, leading some to regard it as part of the explanation. But the authors note that, although higher labor productivity at a given level of output translates into lower employment, it may instead raise output by improving competitiveness, leaving the effect on employment ambiguous. They observe that the slowdown in productivity in the 1970s was accompanied by slower employment growth, whereas the acceleration of productivity in the second half of the 1990s was accompanied by faster employment growth and a reduction of unemployment to the lowest level in decades. As they see it, rapid productivity growth after 2000 meant that aggregate demand also would have had to grow rapidly to maintain employment growth. Although the 2001 recession was mild, the expansion of aggregate demand that followed was not rapid enough. As principal fac-

tors contributing to this inadequate growth in demand, the authors cite the uncertainties following 9/11, the war in Iraq, higher oil prices, and the rise in the dollar's exchange value, all of which also help explain the weakness in exports and the worsening trade balance.

THE STOCK MARKET BOOM of the late 1990s and the record rates of business investment that accompanied it were followed by an unusually large decline in investment during the recession of 2001. Although, historically, investment has declined during recessions, this time the decline was extraordinary—larger than that of GDP itself. Just as extraordinary has been the slow recovery of investment: two years after the recession trough, investment was still only slightly below its peak value of the second quarter of 2000. Many observers, positing a link between the stock market and investment booms, have blamed the depth of the investment decline and its slow recovery on an “overhang” of capital from the earlier period. The sluggish performance of investment has been used to justify sharp cuts in corporate taxes, in the form of accelerated depreciation for most types of investment and lower tax rates on dividends, to stimulate investment by lowering the cost of capital. In the fourth article of this issue, Mihir Desai and Austan Goolsbee look for evidence in support of the overhang story at the firm, asset, and industry levels and evaluate the effectiveness of the tax stimulants that were introduced in recent years.

Although the overhang story fits the aggregate behavior of the stock market and investment in the 1990s and early 2000s, the authors do not consider this evidence conclusive. It could be that the decline in investment during the recession took place in a different set of firms and industries than those that experienced the earlier run-up, in which case any overhang from a previous binge cannot explain the current situation. To examine this possibility, they investigate investment behavior at the industry and firm levels. They begin by performing a simple regression relating the change in the rate of investment from 2000 to 2002 to that between 1994 and 1999 for a sample of eighty-one nonoverlapping industries. Although the resulting point estimate of the effect of investment growth in the earlier period is negative—industries in which investment grew rapidly in the first five-year period did tend to have slower investment growth in the second—the effect is small and statistically insignificant. For manufacturing, the estimated negative relation is stronger: a 1-percentage-point faster rate of growth in investment during the earlier period is associated with approximately a

half-percentage-point slower growth in the later period. (These results are for investment in both equipment and structures; the results are similar for investment in equipment alone and are highly significant.) The authors note, however, that manufacturing represents only about 20 percent of total investment, and so these results are consistent with a lack of significant mean reversion in the growth of total investment.

Desai and Goolsbee report similar results on mean reversion of the growth of investment by asset type across all industries. Examining twenty-five different categories of equipment and nine categories of structures, they again find that above-average increases in the rate of investment between 1994 and 1999 are followed by only slightly above average declines between 2000 and 2002, and the relationship is insignificant.

The authors examine in greater detail a set of data for all firms reported by Compustat; aggregate capital expenditure of the firms in this data set constitutes 85 to 90 percent of private nonresidential investment in the United States. Average growth rates of capital aggregated from the firm level into three broad sectors—manufacturing, computer and information businesses, and nonmanufacturing—display the same pattern as do the aggregate data for the overall economy, with large increases in investment rates during the 1990s followed by declines in 2000–02. The behavior of the computer and information sector was most dramatic, with investment more than doubling in the 1990s before falling back to less than 75 percent of its initial level by the end of 2002. However, regressions for the entire sample of firms show that changes in the growth rate of capital between 1994 and 1999 have only slightly negative effects on growth rates in the same firms in 2000–02. Without information from other periods, it is hard to know whether there was more or less mean reversion in the growth of investment, by firm or industry or asset type, in the early 2000s than is typical during a slowdown. But the authors find the lack of evidence of strong reversion suggestive, indicating that overhang may not be the dominant factor influencing investment in recent years.

Desai and Goolsbee do not explore the extent to which firms currently find themselves with excess capacity, or, for those that do, the extent to which it reflects excessive growth in capital in the 1990s rather than a reduction in demand for their product. In either case, excess capacity would be expected to result in lower investment and a lower market valuation relative to replacement cost—that is, a lower Tobin's q . Instead the authors are interested in knowing whether the extraordinary increases in firms' valua-

tions or investment in the 1990s have reduced the responsiveness of investment to changes in q in the current recovery. Such a reduction would help explain the apparent ineffectiveness of the recent cuts in taxes on corporations and dividends in stimulating investment. Using Q (Tobin's q adjusted for the effect of the corporate profits tax and the tax treatment of depreciation) they first estimate an investment equation for the panel of Compustat firms covering 1962–2003. On the right-hand side, in addition to the firm's Q and its ratio of cash flow to capital, they include fixed year effects and, typically, fixed firm effects and, for the period 2000–03, a term that interacts Q with the change in q during the four-year period ending three years earlier. This interaction is expected to capture whether firms that experienced large increases in q during the stock market boom have a reduced sensitivity in subsequent years to Q and the tax effects it embodies. The estimated coefficients on Q are small but highly significant, whereas the interaction terms are unimportant. The same qualitative results are found for two subsets of firms, information and manufacturing. Although these results leave open the question of whether a capital overhang from the 1990s depressed Q , and therefore investment, after 2000, they do not suggest that investment became less responsive to Q .

The results are different, however, when the authors replace the past change in q in the interaction term with the past percentage increase in the firm's net capital stock. This time the interaction terms are significant, indicating that firms that had larger accumulations of capital in the 1990s were indeed less sensitive to Q in the 2000s. For the firm with the largest past increase in capital, the coefficient on Q falls by 0.004, which corresponds to roughly a 30 percent reduction in the sensitivity of investment to Q . For the median firm the reduction in sensitivity resulting from the investment boom is about 9 percent.

Before turning to an examination of the impact of the recent tax cuts on investment, the authors briefly review tax-adjusted q theory, pioneered by Lawrence Summers in the 1980s. According to the original q theory, investment in the absence of taxes is proportional to the difference between the marginal value of capital and its replacement cost, and the proportionality depends, inversely, on costs of adjustment. Summers recognizes four important features of the tax law: the corporate profits tax, investment tax credits and noneconomic depreciation, dividend taxes, and taxes on capital gains, each of which modifies one or the other of the terms in the non-tax investment equation. The corporate profit tax scales both q and the cost

of capital goods by $1/(1 - t)$, where t is the corporate profit tax rate, leaving the form of the equation unchanged. Introduction of an investment tax credit or accelerated depreciation simply modifies the cost-of-capital term and is equivalent to any other change in the price of capital goods.

Dividend and capital gains taxes are more complicated. Assuming that investors require a given after-tax rate of return, both dividend and capital gains taxes raise the required before-tax rate of return and hence lower market value, but their effect on investment is ambiguous because it depends on how firms finance their investment. If firms finance investment by issuing new equity, the dividend tax is relevant. In the short run a reduction in the dividend tax increases q and encourages investment; in long-run equilibrium the capital stock is larger than it would have been without the tax cut, earning a lower before-tax return but the same after-dividend-tax return, and q returns to its initial value of 1.0 (in the absence of an investment tax credit or noneconomic tax depreciation). In this case the dividend tax does not appear in the q investment equation, its effects being completely reflected in the firm's market value, which in turn affects investment. The authors label this case the "traditional" view and contrast it with what they call the "new" view of dividend taxes, according to which dividends affect firm market value but do not affect investment. This situation arises if a firm's earnings are more than sufficient to finance the desired level of investment, so that retained earnings are the marginal source of funds. It would make no sense for such a firm to sell new shares to finance investment while at the same time paying out earnings as dividends to shareholders, who would have to pay tax on those dividends.

Retaining and investing earnings increases the value of the firm's stock and delivers part of the return to investors through capital gains. Because investors typically hold stocks a considerable time before selling, the effective tax rate on capital gains is below the statutory gains rate, and well below the rate on dividends. A firm maximizing its value to shareholders invests retained earnings up to the point where the after-tax value of a marginal dollar of earnings invested equals that of a dollar paid as a dividend. A dollar of retained earnings that is used to purchase capital adds q to the value of the firm. Assuming they are acting only on behalf of the shareholders, a firm will invest until the value of q is significantly below 1.0. If it stopped adding to capital when q was 1.0, the after-tax value of the next dollar invested would be greater than the after-tax value of that dollar paid out as a dividend.

On this new view, a reduction in the dividend tax rate does not affect the level of either investment or dividends. It increases q by the same proportion as the increase in the after-tax value of dividends, so that a firm that previously was optimally balancing retained earnings and dividends will not want to change either. In contrast, a reduction in the capital gains rate, which also increases q in the short run, will lead the optimizing firm to increase investment and reduce dividends. In the long run the market value of the firm will be higher still, but q will be driven below its initial value.

Which of these two views, the “new” or the “traditional,” is a more accurate description of reality? Given the importance to tax policy of knowing the answer, Desai and Goolsbee find it surprising that, with the exception of studies of aggregate investment in the United Kingdom by James Poterba and Summers, no one has attempted to directly estimate the effect of dividend taxes on aggregate investment or to use firm data to control for aggregate factors. In the case of the United States, part of the reason may be that, until the 2003 tax cut, dividend taxes were never cut in isolation from other tax changes.

Before turning to a direct performance comparison of new and traditional investment equations, the authors examine the performance of several other variants of the q equation in explaining investment of firms in their Compustat sample for 1962–2003. All of the equations include year and firm fixed effects. The authors try two different measures of q : one with and the other without adjustment for the corporate profits tax. Like most investment studies that, following theory, relate investment to q minus the cost of capital, they find that the coefficient is generally highly significant but quite small, implying unrealistically high adjustment costs. The authors believe that this small coefficient is likely to reflect measurement error in q that biases the coefficient toward zero. This leads them to estimate an investment equation that allows separate estimates of the response to q and the cost-of-capital terms for equipment and structures. They find that the coefficient on the cost of equipment (which represents roughly 80 percent of total investment) implies much more reasonable costs of adjustment. The coefficient on the cost of structures is insignificant, which the authors do not find surprising given the traditional difficulty in understanding the incentives for investment in structures. Using the earnings estimates of equity analysts as an instrument for q is another way to deal with the possibility of measurement error, and it likewise results in more reasonable estimates of the cost of adjustment.

Tax-adjusted q and unadjusted q perform comparably when entered singly, but when they are entered together, the coefficient on tax-adjusted q is much larger and of the correct sign, whereas unadjusted q takes on the wrong sign. Both estimates are highly significant. This is consistent with the presence of measurement error in q . The resulting errors in the tax-adjusted term can be offset by a negative coefficient on the q variable entered separately, allowing the variations in the tax rates to dominate the estimation of the coefficient on tax-adjusted q .

The authors regard these estimated investment equations as quite successful. The coefficient on the cost of equipment, which incorporates an adjustment for the tax treatment of depreciation, can be used to predict the effect of the substantial increases in depreciation allowances enacted in 2002 and expanded in 2003 on virtually every type of equipment. But since these equations do not account for the possible effects of capital gains and dividend taxes, they cannot be used to evaluate the effects on investment of the reduction in the capital gains tax rate from 20 percent to 15 percent in 2003, or of the even more dramatic reduction in dividend taxes that same year, from a maximum of 38.6 percent (the maximum tax rate on ordinary income) to the new capital gains rate. As discussed above, the effect of these changes should depend, crucially, on the financing margin facing firms. Desai and Goolsbee rerun the regressions on the panel of firms, this time including two q terms. The first one is appropriate for firms using equity financing of investment and reflects the traditional view, according to which the dividend tax rate should affect investment. The second is appropriate for a firm using retained earnings, reflecting the new view according to which the dividend tax rate should not affect investment. The new view q performs much better at predicting investment. For the entire sample its coefficient is positive and highly significant, whereas the coefficient on traditional-view q is insignificant and of the wrong sign. The same result holds for the subperiods 1962–96 and 1997–2003. The authors view the results for the latter period as especially relevant for assessing current tax policy, both because the current financial structure of firms is what currently matters and because this is the period when the tax rates on capital gains and dividends changed the most. Indeed, for this period the superiority of the new-view q is even greater: its coefficient is roughly twice that estimated using the entire sample. One awkwardness that arises with estimates for this period, however, is that the cost-of-capital terms, which were significant using the entire sample, are now insignificant.

The enactment of immediate expensing of 30 percent of investment in 2002 and 50 percent of investment in 2003 looks like a dramatic reduction that could have a major effect on investment incentives. But, given that investment in most forms of equipment could already be substantially written off in the first few years, the effects of immediate expensing are actually quite modest. The authors calculate the effect of the partial expensing on the cost of equipment, and on the cost of equipment plus structures, for industries at the three-digit classification level. Not surprisingly, the effects differ substantially across industries. Firms that invest mostly in long-lived assets, such as airlines, receive the largest benefit: expensing reduces the cost of their capital by over 3 percent. In industries such as real estate and hotels, in contrast, the reductions amount to only a fraction of a percent. Overall these reductions average about 3 percent for equipment and 2 percent for total investment. In comparison with earlier changes, such as the investment tax credit of 1962 or the Reagan depreciation allowance increases of 1981, all of which reduced capital costs by about 10 percent, the effects of these changes are quite small. The authors offer two reasons why: the corporate tax rate is lower today, so that the benefit of further reductions is smaller; and investment has shifted toward shorter-lived equipment such as computers. Even before the recent tax cuts, the present value of depreciation allowances for equipment was already only 10 percent less than that of full expensing.

Given their belief that dividend tax cuts do not significantly affect the required rate of return on capital, and given the modest effect that the 2002 and 2003 tax changes had on the tax-adjusted cost of equipment and structures, Desai and Goolsbee are not surprised that the Bush tax cuts have had little apparent effect on investment. Assuming a Cobb-Douglas production function, a 3 percent reduction in the cost of capital would increase the desired capital stock by 3.0 percent in the long run if, as the authors assume, output is held constant. If instead employment of labor is assumed to be unchanged or to respond to the higher labor productivity from capital deepening, the increase of the capital stock would be substantially greater. The authors estimate that costs of adjustment typically lead firms to move only about a third of the way toward the long-run goal each year. Using their assumption of constant output, the authors calculate the increase in capital, for a representative firm in each of the three-digit industries, using the capital share and reduction in the cost of capital that they estimate for each industry. They find that, on average, the increase in the

capital stock between 2001 and 2003 attributable to the tax cuts is only 1 to 2 percent, and the average total increase is still less than 2 percent by the end of 2004.

Desai and Goolsbee draw three broad conclusions from their study. First, they see little evidence that a capital overhang from the 1990s plays a dominant role in explaining the differences in investment across industries, asset types, or firms in the 2000s. Second, the tax cuts of the early 2000s, despite their high revenue cost, had minimal, if any, impact on marginal investment incentives. The new view of investment, according to which dividend tax cuts are capitalized in share prices but do not affect investment, is a better description of firm behavior than the traditional view. Third, the partial expensing provisions passed in 2002 and 2003 were not large enough to provide much counterweight to the large declines in aggregate investment.