

**Social Influence and General Practice Healthcare Quality Improvement in
England's Quality and Outcomes Framework (QOF): Evidence from General
Practitioner Movers**

Theodore L. Caputi MPH

MSc by Research

University of York

Health Sciences

January 2021

Abstract

Quality improvement is a central concern of England's National Health Service (NHS) and healthcare systems around the world. Governments and researchers have theorised and tested several methods for improving quality in healthcare, and some interventions have produced modest success. Still, improving healthcare quality remains an elusive task.

In this thesis, I focus on quality improvement in general practice. I provide an overview of the history of general practice quality initiatives within England's NHS, describe the theories and methods for testing healthcare quality improvement, and analyse the feasibility of an unexplored potential vehicle for quality improvement: social influence.

In 2004, the NHS introduced the Quality and Outcomes Framework (QOF) as a pay-for-performance scheme intended to improve quality in general practice. Under this scheme, general practices earn QOF points by meeting certain preset benchmarks for quality, and at the end of the year, practices receive a portion of their compensation in proportion to their QOF point totals. I test whether social influence could feasibly have an influence on a practice's QOF point achievement. Specifically, I use a difference-in-difference approach to assess whether general practitioner (GP) moves (i.e., GPs moving from a practice with either a lower or higher QOF score relative to their destination practice) are correlated with changes in QOF score at the destination practice. I find that, in all QOF domains except patient experience, having a GP move from a better (worse) practice is correlated with an increase (decrease) in the destination practice's QOF scores. These results suggest that social influence among GPs may play a role in quality-relevant decisions among GPs. I discuss the impact these findings, if confirmed with more granular data, could have on quality improvement theory and policy.

Acknowledgements

I would first like to thank my dissertation supervisors, Professors Karen Bloor and Tim Doran, and my Thesis Advisory Panel Committee member, Professor Hugh Gravelle, for their support throughout the dissertation writing process. I appreciate financial support from the Marshall Aid Commemoration Commission, without which I would not have had the tremendous experience of living and learning in York. I would also like to thank my fellow Marshall Scholarship cohort members, who have become like family to me and made the United Kingdom feel like home. Finally, I would like to thank my family for all they have done and continue to do in support of me and my education. I dedicate this dissertation to the healthcare workers in the United States, the United Kingdom, and the rest of the world, who are risking their own health to lead us safely through the COVID-19 pandemic.

Author's Declaration

I declare that this thesis is a presentation of original work and I am the sole author. This work has not previously been presented for an award at this, or any other, University. All sources are acknowledged as References.

Contents

1	Introduction	8
2	General Practice in England's NHS	9
2.1	History of General Practice Quality	9
2.1.1	The Collings Report Reveals Shortcomings in Patient Experience . .	9
2.1.2	The Hadfield Report Shifts Attention to Burden of GPs	11
2.1.3	Redressing Grievances with the NHS	13
2.1.4	Internal Regulation of Quality	14
2.1.5	The Managerial Revolution and Market-Based Mechanisms	16
2.1.6	The Current Quality and Outcomes Framework Scheme	18
2.2	Frameworks and Theories	20
2.2.1	Definitions and Frameworks of Healthcare Quality	20
2.2.2	Theoretical Approaches to Quality Improvement in Healthcare	22
2.2.3	Challenges to Measuring Quality in Healthcare	26
2.2.4	Best Practices for Selecting Quality Indicators	27
2.3	The Quality and Outcomes Framework	32
2.3.1	Selecting Performance Indicators for UK General Practitioners	32
2.3.2	The Purpose of the QOF: Pay Raise or Incentive to Improve?	36
2.3.3	The QOF Scheme in Practice	37
3	Social Influence in General Practice Quality	46
3.0.1	Introduction to Social Influence	46
3.0.2	Peer Influence in Quality Improvement	46
3.0.3	Within-Group and Between-Group Medical Practice Variation	49
3.0.4	Empirical Evidence of Social Influence Effect	50
3.1	Implications of A Mover Effect	52
3.1.1	Leveraging the Social Influence Effect	52

3.1.2	Credibility of Quality Measures	53
3.2	Factors Influencing GP Movement	54
3.2.1	Physician Preferences for Practices	54
3.2.2	Other Forces for Movement	56
4	Methods for Evaluating the Effectiveness of Health Policies	57
4.1	Theory of Causal Inference	57
4.2	Causal Inference in Randomised Studies	59
4.3	Causal Inference in Non-Randomised Studies	60
4.3.1	Regression Modelling	61
4.3.2	Multi-Level Models	63
4.3.3	Difference-in-Difference	70
4.3.4	Interrupted Time Series Analysis	75
5	The Association Between General Practitioner (GP) Entry with Health-care Quality in England’s National Health Service (NHS): Evidence from General Practitioner Movers	80
5.1	Introduction	80
5.1.1	Peer Effects and GP Movers	83
5.1.2	Possible Interventions	85
5.2	Aims	85
5.3	Methods	86
5.3.1	Data	86
5.3.1.1	NHS Quality and Outcomes Framework Data	86
5.3.1.2	GP-by-Practice Tenure Data	87
5.3.1.3	Merging the Datasets	89
5.3.1.4	Constructed Variables	89
5.3.2	Statistical Analysis	91

5.3.2.1	Descriptive Analysis	91
5.3.2.2	Base Case Analysis: Naive Difference-in-Difference	92
5.3.2.3	Difference-in-Difference with Fixed Effects	92
5.3.2.4	Difference-in-Difference with Practice-Level Random Slope Effects	94
5.3.2.5	Main Analysis: Difference-in-Difference with Fixed Effects (Movers Only)	94
5.3.2.6	Percentage Change Difference-in-Difference	95
5.3.3	Sensitivity Analysis: Random Assignment to Treatment Groups . . .	96
5.4	Results	96
5.4.1	Descriptive Results	96
5.4.2	Naive Difference-in-Difference	98
5.4.3	Difference-in-Difference with Fixed Effects	98
5.4.4	Difference-in-Difference with Practice-Level Random Slope Effects . .	98
5.4.5	Difference-in-Difference with Fixed Effects (Movers Only)	99
5.4.6	Sensitivity Analysis: Treatment Groups Randomly Assigned	100
5.4.7	Percentage Change Difference-in-Difference (Movers Only)	100
5.5	Discussion	100
5.5.1	The Mover Effect and Quality Improvement	100
5.5.2	Validity of Quality Measures	101
5.5.3	Limitations	103
5.5.4	Future Research	105

List of Figures

1	Illustration of Fixed Effects	67
2	Illustration of Difference-in-Difference Design	73
3	Illustration of One Sample Interrupted Time Series	76
4	Illustration of Two Sample Interrupted Time Series	78
5	Illustration of Interrupted Time Series Analysis with Fixed Effects	81
6	GP Moves by Distance	97
7	GP Moves by Year	107
8	GP Moves by Year	108

List of Tables

1	Domains and Areas of 2004-2005 QOF Scheme	38
2	Sample Performance Indicators and Point Value, QOF 2004-2005	40
3	Timeline of Changes to the QOF Scheme (England)	41
4	Average Payment Per QOF Point (England)	43
5	Distribution of Points Available across Domains, 2006-2012	45
6	Average Achievement across Domains, 2006-2012	45
7	Factors Driving GP Movement	54
8	Neyman's Notation for a Completely Randomised Experiment	57
9	Two Groups Identical at Baseline	71
10	Simple Difference in Difference	72
11	Overview of GP Starts, Retirements, and Moves	106
12	Distance of GP Moves in Miles	106
13	Summary of GP Moves	109
14	Summary of Moves in Terms of Domain-Specific QOF Scores	110
15	Naiive Difference-in-Difference (No Fixed Effects)	111
16	Difference-in-Difference with Practice and Time Fixed Intercept Effects . . .	112
17	Difference-in-Difference with Practice Random Slope Effects	113
18	Difference-in-Difference with Fixed Intercept Effects, Only Movers	114
19	Sensitivity Analysis: Treatment Groups Randomly Assigned	115
20	Percentage Change Difference-in-Difference	116

1 Introduction

Quality in health services has been the central concern of England’s National Health Service (NHS) since the 1998 publication of *A First Class Service* ([Secretary of State for Health, 1998](#)). *A First Class Service* formally placed the responsibility of improving quality on all NHS organisations and employees and laid out a new system of clinical governance, defined as “a framework through which NHS organisations are accountable for continuously improving the quality of their services and safeguarding high standards of care by creating an environment in which excellence in clinical care will flourish” ([Secretary of State for Health, 1998](#)), through which NHS service providers would be expected to continuously improve quality. Because of the centrality of general practice in England’s healthcare system and because previous contracts between England’s NHS and GPs made monitoring quality difficult, improving quality in general practice immediately became a top priority for the NHS ([Baker et al., 1999](#)).

The Qualities and Outcomes Framework (QOF) is the most comprehensive in a series of schemes adopted by the NHS to improve quality among general practitioners (GPs). Overall, the QOF is a voluntary but widely adopted pay-for-performance scheme in which GP practices are rewarded for scoring points through delivering a set of predetermined interventions and reaching a set of predetermined clinical, organizational, and other quality-related targets. The QOF scheme was first introduced to GPs in 2004 as part of the General Medical Services contract between the NHS and individual GP practices. At the time of its adoption, the QOF included 146 indicators within four domains: clinical, organisational, patient experience, and additional services ([Downing et al., 2007](#)). Practices could earn up to 1050 total QOF points depending upon their level of achievement on each of the indicators, with certain indicators weighted more heavily than others. QOF points were assessed by practice each year using the Quality Management and Analysis System, and GP practices were awarded a financial reward based upon their QOF points. The QOF was initially hailed as both offering “the promise of a quantum change in performance rather than an incremental

one” (Shekelle, 2003) as well as “medicine by numbers... threaten[ing] the professional basis of general practice” (Lipman et al., 2005).

Since its adoption, the QOF scheme has undergone several changes. For example, an expert panel involving participants from the University of Birmingham, the Society for Academic Primary Care and the Royal College of General Practitioners were appointed in 2005 to review the QOF indicators, which led to the addition of seven new clinical areas beginning in 2006 (Lester et al., 2006).

This dissertation leverages publicly-available, practice-level QOF scores as a benchmark to examine the possibility for another mechanism of healthcare quality improvement: social influence. General practitioners may offer higher (or lower) quality care depending upon their peers in the practice. QOF scores – a common set of indicators all GP practices are financially incentivized to meet – provide a comparable (albeit noisy) metric of quality across England’s general practitioners. To assess the role of social influence, I adopt a movers-based difference-in-difference design, comparing the QOF performance among practices that received a new GP moving from a practice with either higher QOF scores (a “better practice”) and those that received a new GP moving from a worse QOF scores (a “worse practice”). I find suggestive evidence that having a GP move from a better (or worse) practice improves (worsens) the destination’s QOF points in several domains.

2 General Practice in England’s NHS

2.1 History of General Practice Quality

2.1.1 The Collings Report Reveals Shortcomings in Patient Experience

Quality in General Practice has been a concern in the UK National Health Service (NHS) since its formation. When the NHS was founded in England in 1948, just three years after the end of the devastating World War II, general practice became formally enshrined as the

core of healthcare in the UK, responsible for all personal healthcare and acting as a gateway to specialised and hospital services (Goodwin et al., 2011). It did not take long thereafter for GPs to garner criticism. In 1950, Dr. Joseph S. Collings (Collings, 1950) published a long-form (30 page), ethnographic study in *The Lancet* entitled “General Practice in England Today - A Reconnaissance”. Collings had experience as a physician in Canada and New Zealand and was a research fellow at Harvard University’s School of Public Health (Petchey, 1995), but before conducting his study, had never visited the UK. In his study, Collings described his observations of day-to-day GP activity from (convenience sampled) 55 English GP practices that he had visited, representing industrial, urban-residential, and rural area practices. He concluded that “the overall state of general practice is bad and still deteriorating” and that the conditions of English GP varied markedly across industrial, urban-residential, and rural area practices. He described industrial practices as inhospitable, overcrowded, and untidy and remarked that consultations at industrial practices were cursory and incomplete. While he described rural practices as “an anachronism”, he found they were relatively better than industrial practices, applauding their safety and comfort and the organisational prowess of rural doctors. Urban-residential practices spanned the difference between industrial and rural area GPs. Collings vehemently dismissed the notion that quality would inherently be lower in industrial rather than rural practices, insisting instead that “[i]f the doctors’ surgeries were better from a functional point of view; if they were equipped at least to some minimum standard; if a little order were brought into the chaos which characterises the organization of the average practice (e.g., by employing clerical assistants); and if the work of the industrial practitioner were coordinated with that of other medical agencies-then the same volume of work could be handled at a much higher level and with due regard to the safety and comfort of the patient” (Collings, 1950). Collings concluded that the NHS had not planned for or incentivised good practice and that professional and administrative checks were insufficient. Ultimately, he recommended that a system be developed in which GP practices operate out of state-run facilities (Fry, 1988).

While Collings’ report was methodologically limited (many criticized Collings’ sampling strategy and his observational methods ([Petchey, 1995](#))), its rich descriptions resonated strongly with contemporary GPs ([Honigsbaum, 1979](#)). It spurred responses from academics and physicians within ([Hunt, 1955](#)) and outside of *The Lancet* ([Anthony, 1950](#); [Dawson, 1950](#); [Lask, 1950](#); [Pirrie, 1950](#); [Sanguinetti, 1950](#)) and, by extension, sparked widespread policy interest in the evaluation and improvement of quality among NHS GPs ([Petchey, 1995](#)). For example, the Collings’ report is specifically cited as the impetus for the 1952 creation of the College of General Practitioners (now the Royal College of General Practitioners) ([Royal College of General Practitioners, 2019](#)). As stated on the College’s website, the College of General Practitioners was founded under the “shared belief that what was needed [to address the issues raised in the Collings report] was an academic body to support good standards of practice, education and research, such as already existed in other medical colleges” ([Royal College of General Practitioners, 2019](#)).

2.1.2 The Hadfield Report Shifts Attention to Burden of GPs

Many scholars suggest that the 1950 Collings Report directly prompted another major investigation of GPs in the UK: the 1953 Hadfield report ([Wilkie, 2014](#)). The Hadfield Report ([Hadfield, 1953](#)), sponsored by the British Medical Association’s General Practice Review Committee and authored by Committee Secretary Dr. Stephen J. Hadfield¹, surveyed the experience of 200 GPs located in England, Scotland, and Wales between February 1951 and March 1952. Seemingly to improve upon Collings’ sampling strategy, GPs were selected near-randomly from 20 separate regions (the 19 regional hospital board areas and the Metropolitan Branch of the British Medical Association). Hadfield visited each practice, sat in on at least one consultation, and then accompanied the GP on his normal rounds. He documented his impressions on the lifestyle of the “average” GP (e.g., where they live, their daily schedule, their leisure time), the services that GPs provide (e.g., continuity of

¹It should be noted that the British Medical Association strongly supported the newly implemented National Health Service ([Anonymous, 1953](#)).

treatment, preventive medicine, diagnosis, treatment, providing information to specialists, non-medical advice, clinical assistantships in hospitals, and administrative duties), the professional relationships GPs hold (e.g., with other GPs, with consultants, with public health officers, with nurses), the organisation of GP practices (e.g., premises, equipment, appointment management, employment of non-medical staff, and record-keeping), and views that GPs hold on auxiliary services (e.g., effect of the NHS on general practice, services provided to health visitors, the adequacy of hospital services, and the utility of health centers). Hadfield described several shortcomings of the GP system throughout his report and painted general practice in the UK as underfunded and demoralising work. However, he ultimately found that 44% of practices had overall “good” quality, 44% had “adequate” quality, and only 7% had “inadequate” quality (5% were not assessed).

Hadfield concluded his report with three “adverse impressions” from his investigation. First, he criticized the lack of cohesion and continuity within the NHS, including weak relationships among GPs, specialists, and public health medical officers and discontinuities in administration. Second, he criticized inadequate and lagging hospital services, which place additional strain on GPs. Finally, he criticized GPs response to an increased demand for services. He described that the public’s expectations of entitlements from the NHS are already “deep-rooted” and argued that, given the increased responsibility of GPs and the increased expectations from patients, a strategy to handle increased demand was necessary. Ultimately, the Hadfield Report, while considered more optimistic than the Collings Report, affirmed that there was substantial room for improvement in UK general practice. An anonymous editorial appearing alongside the Hadfield Report in *British Medical Journal* stated: “[i]t is a picture that justifies a restrained optimism for the future of general practice – restrained because, unless remedies are applied soon and continuously, the present good service the public receives from general practitioners will deteriorate in time”.

To some degree, it appears that the legacy of the Hadfield Report was shifting public attention away from patient experiences, the focus of the Collings’ Report, and towards the

frustrations faced by GPs under the new NHS regime. For example, soon after the publication of the Hadfield Report, Dr. Stephen Taylor, a doctor and former Labour politician, authored a report entitled “Good General Practice” recommended that GPs could improve quality by reducing their personal burden (Taylor, 1954). Taylor concluded, “good general practice begins with the good GP... [s]o most of the conclusions are suggestions for self help” (Taylor, 1954).

2.1.3 Redressing Grievances with the NHS

By 1965, GPs had accumulated several grievances with their position within the NHS. GPs argued that the overwhelming demands of the job, combined with no formal provisions for quality improvement, left them incapable of providing good quality service to patients. The British Medical Association described these grievances in its “Charter for the Family Doctor Service”, which set forth what GPs needed in order to offer high-quality service to their patients. Published in *British Medical Journal*, the Charter demanded a lightened burden so that GPs could properly practice medicine (GMS Committee, 1965):

If the trained doctor who enters general practice is to be able to do those things he has been trained to do, then he must have adequate time, space, and assistance with which to do them. In order to have these basic requirements he must find money to provide the space, equip it properly, and engage the services of non-medical helpers who in their own special field of experience can take a good deal of work off his hands, and so leave him free to do that which he has been trained to do. Above all, the general practitioner must have time to take a full history and make a careful examination of those patients who come to him with other than trivial ailments, in the management of which many were content to take the advice of the local pharmacist... Medicine advances at such a rate that unless a doctor once qualified deliberately sets aside time for reading, and periodically time for retraining, he is bound to get out of date and to feel himself isolated from his colleagues who, by continuing to work in hospital, are confronted throughout the whole of their work with new ideas and with frequent interchange of these and of experience with other hospital colleagues. This continuing education needs time, and time is money.

The following year, the demands of the Charter were accepted and codified in the 1966 NHS General Practice Contract, which massively improved GP pay and conditions (Gillam, 2017). For example, the Contract reduced capitation-based income and increased practice

allowances and fee-for-service opportunities, increasing overall pay to GPs. In addition, the Contract offered reimbursements for staff-related expenses, allowing GPs to hire nurses and non-clinical staff, and subsidies for facility development. The Contract further placed a limit on a GP's patient list size at 2,000, reducing the maximum number of patients a GP would be responsible for. The advantageous provisions of the Contract attracted a new wave of people to become GPs ([Horder and Swift, 1979](#)), and the era immediately following 1966 has been considered a “golden age for general practice” ([Addicott and Ham, 2014](#)), at least for the general practitioners.

2.1.4 Internal Regulation of Quality

GP Associations had demonstrated their political dominance over the Government and the NHS in the 1966 General Practice Contract negotiations, and in subsequent decades, health-care quality was monitored, standardized, and addressed internally through professional medical organisations ([Howie, Heaney and Maxwell, 2004](#)). The College of General Practitioners had worked to establish standards of practice, offering “Foundation Membership” to GPs who satisfied a defined set of criteria ([Royal College of General Practitioners, 2019](#)). In 1972, the College of General Practitioners was awarded a formal Royal Charter, which established the organisation (renamed the Royal College of General Practitioners) into the first official professional organisation for GPs. Four years later, the Royal College of General Practitioners successfully lobbied Parliament to approve the 1976 NHS Vocational Training Bill, which required three years of vocational training to become a GP and established national standards organisations to monitor that training. In turn, the number of GPs enrolling in formal vocational training to become a GP increased 130% between 1976 and 1985 ([Petchey, Williams and Baker, 1997](#)).

By the 1980s, external researchers had accumulated substantial evidence that the quality of general practice in the UK was inconsistent and that certain groups of patients were receiving inadequate care ([Practice, 1985](#); [Baker, 1989](#)). The public grew dissatisfied ([Baker,](#)

1989), particularly in an era where other UK government agencies experienced increased oversight. To mitigate some of this criticism and to reduce the probability of a government-led reform, the Royal College of General Practitioners launched its “Quality Initiative” in 1983, which sought to improve professional development of doctors, practice management and team-work, quality assessment, contracts and incentives, and resource needs (Practice, 1985). Generally, the Initiative was intended to raise awareness about quality measures among GPs and to encourage GPs to consider how they could improve the quality of their own practice (Baker, 1989). Specific recommendations from the “Quality Initiative”, many of which focused on the practitioner rather than the patient, included the following (British Medical Journal, 1985):

- Higher training through appropriately supervised experience in the early years of general practice should now be introduced to promote the further professional development of young principals.
- The concept of “protected time” for learning should be extended to all principals and practice staff and the resources necessary for this should be made available.
- A credible entry standard for NHS general practice should be determined. In future, doctors entering the NHS list as principals should be encouraged to reach the standard available through the membership of the Royal College of General Practitioners examination.
- The completion of higher training should be denoted by accreditation given by the college, and should be based on the continuous assessment of clinical performance as in other specialties of medicine.
- Good quality practice needs good quality management. Every practice should provide the basic range and quality of services which every patient can expect anywhere.
- Standard setting and performance review are activities that must be incorporated into everyday clinical practice. They make essential contributions to effective practice management and to the continuing education of members of the practice team.
- Incentives should be developed to encourage doctors to participate in performance review. A solution should be found within the framework of the National Health Service to bridge the gap between how doctors are paid and the standard of services that they provide. Unacceptable levels of performance should be reflected in a doctor’s remuneration.
- Resources must be found for the implementation of the RCGP’s council’s strategy for quality.
- Membership of the college should be open to older doctors who submit successfully to a system of assessment that is no less rigorous than the MRCGP examination.

- Performance review should be incorporated into every member’s professional life, and such participation should contribute to the attainment of the fellowship of the college.

2.1.5 The Managerial Revolution and Market-Based Mechanisms

The light-touch quality improvement strategy implemented by the Royal College of General Practice was well-received by GPs ([Baker, 1989](#)) but failed to ward off criticism. Inspired by the Griffiths Report of 1983, policy-makers began calling for the government to take on tighter managerial responsibility over GPs and to apply the general management tools from private industry into healthcare ([Strong and Robinson, 1990](#)). For the first time, the Government – at the time run by political conservatives predisposed to market-based solution and the privatisation of previously state-owned industries – began discussing market-based interventions designed to incentivise high-quality treatment. In 1986, the Secretaries of State for Social Services released a “Green Paper” that recommended a pay-for-performance system called the “Good Practice Allowance”, which would provide financial rewards to GP practices that passed an inspection and achieved a certain rate of screening and preventive services ([Baker, 1989](#)). There was sufficient backlash from the medical community for the government to remove all mention of the “Good Practice Allowance” in its follow-up 1987 “White Paper”, but other market-based incentive mechanisms persisted. For example, the Government recommended that GP practices be encouraged to advertise their services, increasing competition among practices, and that GP practices would be given some financial compensation for increasing their list size. A small set of preventive services would also garner an additional fee.

These new managerial systems were formally adopted in the 1990 NHS General Practitioner Contract ([Williams et al., 1993](#)). By increasing the GP’s variable proportion of income (e.g., capitation-based and fee-for-service payments), the new contract was designed such that funding would “follow the patient”, meaning that GP practices (as well as acute care, community care, priority care, and health promotion units) were financially incentivised

to compete for patients. Importantly, the 1990 Contract included additional incentives for meeting or exceeding certain “performance related” or quality benchmarks, such as immunisation rates. In this way, GPs were, for the first time, financially incentivised to achieve quality standards.

Beginning in 1991, the Conservative Government instituted another market-based intervention intended to improve health care quality and slow the growth of healthcare costs: GP fundholding (Kay, 2002). Under the GP fundholding scheme, practices that voluntarily opted into the scheme were given a certain amount of money each year to purchase hospital and drug services (and beginning in 1993, community services) for their patient populations (Shortt, 2003). If the practice had any of this funding remaining at the end of the year, it was allowed to retain the excess funds for practice augmentation. In this way, practices were incentivised to find low-cost services for their patients, and providers would compete for the patronage of fundholding GPs. The voluntary scheme quickly became popular among GPs, and by 1995, 2,500 practices (encompassing 10,000 GPs and 41% of the population) had voluntarily enrolled (Ham, 1996). There was not much evidence that the GP fundholding scheme improved healthcare quality (Coulter, 1995), though some suggested it slowed the growth of healthcare costs (Harris and Scrivener, 1996) and elective surgery (Dusheiko et al., 2006). The scheme was abolished by the new Labour Government in 1998 (Kay, 2002). Notably, researchers decried that the decisions both to implement and to abolish the GP fundholding scheme were highly politicized and that both decisions were made without substantial, valid evidence of the program’s effects (Kay, 2002).

The Labour Government promised to create a National Institute for Clinical Excellence (NICE) to improve healthcare quality and quality monitoring. The need for these improvements was punctuated, for both politicians and their constituents, by several high-profile scandals in the UK healthcare system. For example, a formal inquiry revealed between 30 and 35 excess paediatric heart surgeon deaths between 1991 and 1995 due to inadequate care at the Bristol Royal Infirmary (Dyer, 2001). Even more shockingly, a separate inquiry

found that general practitioner Harold Shipman had murdered at least 215 patients ([Smith, 2002](#)). These scandals made it clear that quality monitoring in England’s healthcare system was incapable of identifying serial killers, let alone suboptimal care.

Subsequently, the Labour Government fulfilled its promise and launched the National Institute for Clinical Excellence (NICE) on 1 April 1999. NICE was formed with the understanding that physicians were unable to stay up-to-date with best practices for care and was tasked with providing physicians and other healthcare professionals in the NHS (including GPs) with the correct tools to bridge clinical care and best practice ([Rawlins, 1999](#)). Specifically, NICE was given three main responsibilities: (A) appraising new health technologies, (B) developing clinical guidelines from the literature, and (C) promoting clinical audits and inquiries. NICE’s actions on clinical guidelines and audits laid the groundwork for an expanded presence of Government in quality monitoring. The clinical guidelines set standards which could be adapted into performance indicators ([Campbell, 2002](#)), and NICE’s new oversight of clinical inquiries moved the balance of power away from the Royal Colleges, who remained responsible for conducting National Confidential Enquiries.

2.1.6 The Current Quality and Outcomes Framework Scheme

By the early 2000s, the field of general practice was in turmoil, facing decreases in both GP morale and in new GP recruitment ([Roland and Campbell, 2014](#); [Roland and Guthrie, 2016](#)). The NHS worried that these factors could lead to decreases in quality. To ameliorate the situation, the Labour Government, which had already committed to increase funding for the NHS, offered GPs a significant pay raise with fewer hours in return for greater accountability for the quality of their services. The GPs agreed. With the 2004 General Medical Services contract between the NHS and individual GP practices, the NHS’s initial forays into pay-for-performance grew into the dramatically expanded Qualities and Outcomes Framework (QOF) scheme, and average GP pay was increased by nearly 20%. The QOF was controversial among practitioners and policy-makers. Some described it as “the promise of a quantum

change in performance rather than an incremental one” ([Shekelle, 2003](#)), while others criticized it as “medicine by numbers... threaten[ing] the professional basis of general practice” ([Lipman et al., 2005](#)).

The QOF was and remains a voluntary but widely adopted pay-for-performance scheme in which GPs are awarded payment-linked points for delivering certain interventions and achieving certain quality targets. The QOF initially included 146 indicators within four domains: clinical, organisational, patient experience, and additional services ([Downing et al., 2007](#)). Indicators were weighted so that achievement on one indicator could have a larger impact on QOF scores than achievement on another indicator. In general, this weighting was assigned to reflect which indicators required the most work from the practice and its GPs (rather than reflecting the indicators most important for public or patient health). Practices could earn up to 1050 total QOF points depending upon their level of achievement on each of the indicators, with certain indicators weighted more heavily than others. QOF points were assessed by practice each year using the Quality Management and Analysis System, and GP practices were awarded a financial reward based upon their performance, with adjustments made for workload and the health of patients in the practice’s area.

The QOF scheme has changed several times since its adoption. For example, in 2005, an expert panel involving participants from the University of Birmingham, the Society for Academic Primary Care and the Royal College of General Practitioners were appointed in 2005 to review the QOF indicators, which led to the addition of seven new clinical areas beginning in 2006 ([Lester et al., 2006](#)). Since 2009, the National Institute for Health and Care Excellence, the UK’s governmental body responsible for facilitating the translation of research into practice, has been tasked with creating a menu of “suitable” indicators that could be included in the QOF scheme ([Sutcliffe et al., 2012](#)). As of 2018/2019, the QOF includes three domains (Clinical, Public Health, and Public Health – Additional Services) involving 65, seven, and five indicators, respectively ([National Health Service, 2018a](#)). Now nearly 15 years since its initial implementation and with adoption rates above 94% (parti-

pation of 94.8% in 2017-2018; ([National Health Service, 2018b](#))), the QOF scheme is a well ingrained component of general practice in the United Kingdom.

2.2 Frameworks and Theories

2.2.1 Definitions and Frameworks of Healthcare Quality

Quality is a nebulous term. Even within the context of healthcare, various conceptions of quality exist. Consequently, when evaluating quality, it is useful to be aware of the various definitions and frameworks that may be applied.

Definitions of quality vary substantially and emphasize different points. [Mosadeghrad \(2012\)](#) describes several definitions that each have different points of emphasis. For example, ([Donabedian, 1980](#)) focused on the benefit-risk ratio, defining quality as “the application of medical science and technology in a manner that maximises its benefit to health without correspondingly increasing the risk”. [Ovretveit and Townsend \(1992\)](#) extends the endgoal of quality from simply maximizing benefit and minimizing risk to also exceeding patient expectations: “Provision of care that exceeds patient expectations and achieves the highest possible clinical outcomes with the resources available”. [Schuster, McGlynn and Brook \(1998\)](#) define good quality care as “providing patients with appropriate services in a technically competent manner, with good communication, shared decision making, and cultural sensitivity”, matching their framework. [Lohr \(1990\)](#) emphasizes the role of professional standards in achieving quality, defining quality as “the degree to which healthcare services for individuals and population increases the likelihood of desired health outcomes and is consistent with the current professional knowledge”, a definition formally endorsed by the Institute of Medicine ([Schuster, McGlynn and Brook, 1998](#)) alongside the six domains of quality (e.g., safe, effective, patient-centered, timely, efficient, and equitable) ([, 2001](#)).

Most frameworks for understanding health quality attempt to break down the vague concept of “total quality” into domains or perspectives that evaluators can explore separately ([Mosadeghrad, 2012](#)). However, these domains or perspectives vary significantly. For ex-

ample, the framework put forth by [Ovretveit and Townsend \(1992\)](#) considers three different dimensions: professional quality, referring to the use of the correct procedures to meet patient need; client quality, referring to satisfaction from the patient; and management quality, referring to efficiency of healthcare delivery. [Joss and Kogan \(1995\)](#) consider quality in three different domains: technical quality, referring to the procedures performed; systemic quality, referring to the overarching systems and processes that the procedures fit into; and generic quality, referring to interpersonal relationships. [Schuster, McGlynn and Brook \(1998\)](#) focus on providing the appropriate level of care and assess shortcomings in terms of too little care, too much care, or the wrong care.

Perhaps the best known conceptual model for evaluating quality in healthcare was developed in 1966 by Dr. Avedis Donabedian of the University of Michigan on commission from the Health Services Research Study Section of the United States Public Health Service ([Donabedian, 1966](#)). Donabedian conceived of healthcare as a dynamic system that could be evaluated in terms of its processes, structures, and outcomes. He felt that quality was ultimately determined by important patient outcomes but recognized that health outcomes were exceedingly difficult to accurately observe. His framework, consequently, combines his focus on outcomes with detailed observation of structures and processes that are closely linked with clinical outcomes. When an evaluator uses the Donabedian framework to evaluate healthcare quality, he or she observes the structure of healthcare, i.e., the context in which medical care is delivered; the process of healthcare, i.e., the actual delivery of care; and outcomes, i.e., the effect of the healthcare delivered on the morbidity or mortality of an individual or population. Donabedian's conceptual model persists as perhaps the most common framework for evaluating healthcare quality ([Berwick and Fox, 2016](#)).

Aggregating definitions of and frameworks for quality by past scholars, [Mohammad \(2013\)](#) develops a particularly useful meta-framework. In this meta-framework, he distinguishes between quality evaluated from the perspective of healthcare services suppliers and quality evaluated from the perspective of healthcare services demanders. On the supply-side, quality

can be assessed in relation to a set of benchmarks predetermined by managers. Benchmarks are set, *before* any service is rendered, in relation to what past experience (e.g., research literature or expert opinion) would dictate is best practice. The definition for quality set forth by [Lohr \(1990\)](#) would fit into this category. On the demand-side, however, quality is measured in relation to patient-specific needs and expectations. In this way, quality is assessed as a matter of satisfaction or dissatisfaction by the patient. The definition of quality set forth by [Ovretveit and Townsend \(1992\)](#) reaches into this domain.

2.2.2 Theoretical Approaches to Quality Improvement in Healthcare

Over the past 70 years, scholars and policy-makers have debated the merits of various approaches to maintaining and improving quality in healthcare. The history of general practice in the UK has been dominated by four broad models, commonly referred to as the “Altruism”, “Hierarchy and Targets”, “Reputation”, and “Choice and Competition” models ([Bevan and Fasolo, 2013](#)).

The Altruism model, also known as the Trust model ([Le Grand, 2007](#)), is based on the notion that people will try their best and achieve the best possible outcomes given their resources. In that vein, the Altruism model suggests that the quality of healthcare will improve if GPs are given better resources and information. According to the Altruism model, it is mostly useless to either punish or incentivise bad or good performance, as any observed variation in performance is externally determined and cannot be changed by the GP. Instead, under this model, poor performance should be met with increased investment to improve those GP’s resources and information ([Powell, 1976](#)). For the first several decades of the NHS’s existence, this was a particularly compelling model for physicians, including GPs ([Bevan, 2010](#)). Medicine was one of the most highly-trusted professions, and people valued not just the care they received from their GP but the personal relationship they had with their GP. The public felt uncomfortable with questioning the motivation of physicians. The Altruism model is low-cost and amenable to GPs ([Le Grand, 2007](#)), but it requires the

strong assumption that all GPs are exerting maximum effort.

The Hierarchy and Targets model (also referred to as the Mistrust model ([Le Grand, 2007](#))), on the other hand, accounts for variation in GPs' motivation through external incentives. This model relates closely to the theoretical economic literature on principals and agents ([Jensen and Meckling, 1976](#)). The principal – in this case, the Government or the taxpayers – pays the agent – in this case, GPs or GP practices – to perform some work. With perfect information, the principal can offer to pay the agent only if she performs at her maximum capacity. However, the Hierarchy and Targets model assumes that there is asymmetric information; that is, the principal can only observe the work the agent completes but not the full-effort capacity of the agent. Conceptually, this asymmetrical information benefits the agent, who can shirk (reducing the agent's disutility of effort) without the principal knowing whether the limiting factor on output is the agent's capacity or his or her level of effort. If the principal wants the agent to maximize effort and output, then, he has to provide the agent with incentives to the output ([Conn, 1982](#)). Consequently, the Hierarchy and Targets model calls for financial incentives (or disincentives) on the quality of GP's services. The principal sets certain standards for the agent's output, and the agent is rewarded based upon the extent to which her work meets those standards. This model underlies pay-for-performance schemes, through which GPs are rewarded if they meet certain preset targets for quality. The Hierarchy and Targets model is intuitively appealing but also comes with high monitoring costs. Further, financial incentives for high-quality tend to be unpopular among physicians, who lose their ability to costlessly shirk ([Le Grand, 2007](#)). While the Hierarchy and Targets model relates closely to the principal-agent problem, it is important to note that, in this regard, the GP is incentivised to act in the interest of the Government and taxpayers rather than the patient. The effectiveness of this model on patient wellbeing, then, is dependent upon the alignment between government and patient interest. When these interests are misaligned (e.g., the best possible treatment for a patient would be expensive), the GP is incentivised to act in the interest of the government and

against the interest of his or her patient.

In some ways, the Reputation model combines aspects of the Altruism model and the Hierarchy and Targets model. According to the Reputation model, agents can be incentivised to maximize effort not only by financial incentives but also by threats to their reputation. This model assumes that physicians enjoy their status as being highly trusted, which, in turn, means that they'd be willing to exert some additional effort to protect their reputations ([Wachter, 2013](#)). Reputation models are often implemented by publicizing the quality of care for individual GPs or practices ([Jarvis, 2012](#)). For example, the NHS follows the Reputation model when it shares individual practice's QOF scores online. Ideally, information shared through this model will be widely disseminated and easily understood, such that exerting less than maximum effort has a meaningful reputational impact for the GP. Similarly to the Hierarchy and Targets model, the Reputation model is intuitive but also costly and unpopular among GPs ([Le Grand, 2007](#)). Further, when GPs are involved in the process of determining how to evaluate quality, it is possible for GPs to reduce quality standards so that all or nearly all practices meet standards or for GPs to establish upper thresholds for achievement that disincentivise any additional effort.

The Choice and Competition or Quasi-Markets ([Le Grand, 2007](#)) model is similar to the Hierarchy and Targets model in that it offers financial incentives (or disincentives) to GPs for exerting high effort (or low effort). However, instead of giving principals the responsibility of determining standards of quality, the Choice and Competition model allows consumers to decide what constitutes high quality service. This model assumes that patients – healthcare consumers – are sophisticated arbiters of quality and allows them to choose from a market of healthcare providers. The providers are, in turn, rewarded based upon the number of patients they serve and the extent of services they provide. In sum, GPs exert effort to maximize quality, so that they will attract patients, which, in turn, increases their payments. In a Choice and Competition model, information about GP quality is made publicly available (and ideally, easily accessible) and patients are given the autonomy to choose their providers.

The Choice and Competition model is theoretically appealing as it increases patient autonomy, avoids the issue that Governments may choose sub-optimal quality standards, and has low monitoring costs. However, it also increases the search costs on patients and relies on the strong assumption that individuals will be (or become) good arbiters of quality. While there is some evidence that patients' healthcare decisions respond to quality ([Howard, 2006](#)), these transaction costs may be substantial ([Magee, Davis and Coulter, 2003](#)). If these transaction costs are too burdensome for many patients, who then choose sub-optimally (e.g., prefer to remain at their original practice), the incentive for GPs to improve quality may be lost ([Fotaki, 2013](#)).

Each of these models has played a role at some point in the history of general practice within the NHS. The Altruism model matches early policies intended to improve quality. The Hadfield Report ([Hadfield, 1953](#)) and later the Charter for the Family Doctor Service ([GMS Committee, 1965](#)) argued that quality was low because doctors did not have the resources to adequately practice medicine. In turn, the 1966 NHS General Practice Contract was signed to increase GP pay and conditions ([Gillam, 2017](#)), with the understanding that increased resources would allow GPs to improve quality. Internal regulation of quality by bodies such as the Royal College of General Practitioners could be considered a form of the Reputation model. While there were no financial incentives or penalties, these professional organisations created a reputational incentive for adhering to certain standards by setting quality requirements for membership ([Royal College of General Practitioners, 2019](#)). When the public became wary that internal regulation may be too lenient, the Government called for market-based incentives. In line with the Choice and Competition model, the Government suggested that funding should “follow the patient”, incentivising practices to compete among themselves ([Baker, 1989](#)). These suggestions became part of the 1990 NHS General Practitioners Contract, along with some additional incentives for achieving certain targets, in line with the Hierarchy and Targets model ([Williams et al., 1993](#)). Since the early 2000s, the Government has invested heavily in the Hierarchy and Targets model with the adoption

of the Quality and Outcomes Framework in 2004. QOF, the largest pay-for-performance scheme in the world, sets certain quality targets and provides financial rewards to practices that meet or exceed those targets, giving GPs a financial incentive to achieve quality standards (Shekelle, 2003). The Hierarchy and Targets model remains the dominant model for quality improvement in the UK.

2.2.3 Challenges to Measuring Quality in Healthcare

There is a substantial literature documenting the challenges involved in evaluating quality in service industries (Proctor and Wright, 1998), and healthcare quality is no different (Mosadeghrad, 2014a,b; Cheng, Tang and Jackson, 1999). Frequently cited challenges to assessing quality in the healthcare industry include patient participation, simultaneity, perishability, intangibility, and heterogeneity (Ozcan, 2005).

- **Patient Participation:** Healthcare delivery involves highly complex and dynamic interactions between patients and an array of healthcare professionals (e.g., physicians, nurses, pharmaceutical developers, pharmacists, public health officials, health product producers), and so it can be difficult to trace a causal pathway from any one actor's actions and patient outcomes. If a patient resists the best practice care offered to them or chooses to ignore a physician's advice and suffers worse clinical outcomes, should that reflect poorly on the healthcare system?
- **Simultaneity:** Healthcare is simultaneously produced and consumed, making it challenging for quality control protocols to proactively promote good quality care. In the case of manufactured goods, a quality control officer could pick goods that fail to meet certain standards from the shelf; this is not possible with healthcare services.
- **Perishability:** Healthcare services are largely perishable. For example, physicians typically provide services during a prescribed set of hours – if some fraction of those hours are unoccupied, the capacity of that physician is essentially wasted. Given the

unpredictable nature of health, it may be impossible for physicians to accurately match the timing of patient demands for healthcare with their availability.

- **Intangibility:** As an intangible “good”, health is subjective, multifaceted, and fleeting, which makes outcomes in health more difficult to quantify than physical goods.
- **Heterogeneity:** Heterogeneity in the needs of the patient, even within certain well-defined conditions, makes it challenging to compare services rendered with uniform benchmarks. The practice of medicine requires countless “judgement calls”, in which the physician has to make decisions based upon nuances that could not be prescribed in professional guidelines.

It is potentially impossible to successfully disentangle several of these complications from any system of measuring healthcare quality, particularly given resource constraints on healthcare quality monitoring. Instead, healthcare quality monitors have largely accepted these limitations and focused on selecting and applying sets of performance indicators in ways that mitigate these concerns as best as possible ([Bennett et al., 2014](#)).

2.2.4 Best Practices for Selecting Quality Indicators

By the turn of the century, as institutions were becoming increasingly concerned with healthcare quality and performance indicators were becoming more fundamental to the healthcare system in the UK and internationally, there already existed decades worth of small-scale and anecdotal research on how managers chose and used healthcare quality indicators ([Starkweather, Gelwicks and Newcomer, 1975](#); [Grimes and Moseley, 1976](#)). Realizing that quality indicators were to become more systematically important to the healthcare system, scholars sought to synthesize the findings from these past experiences and inform policy-makers of the best practices for selecting performance indicators. [Freeman \(2002\)](#) conducted a systematic review of the literature on performance indicators. He found that summative performance indicators designed to reward or punish individuals tended to corrupt the in-

dicators themselves, presumably because they failed to account for heterogeneity and the number of “judgement calls” in clinical medicine. On the other hand, he found that formative indicators gave clinicians clues of how they could improve while allowing clinicians and their managers to make adjustments for local and clinical context. ([Mainz, 2003b](#)) used the experience of the Danish National Indicator Project to outline a specific process for developing the right indicators. He recommended that clinical indicators first be defined and prioritised based upon the scientific literature, then specific targets should be adopted with special consideration given to each indicator’s measurement and inclusion/exclusion criteria. He also characterized the features of an ideal healthcare indicator ([Mainz, 2003a](#)):

- Indicator is based on agreed definitions, and described exhaustively and exclusively
- Indicator is highly or optimally specific and sensitive, i.e. it detects few false positives and false negatives
- Indicator is valid and reliable
- Indicator discriminates well
- Indicator relates to clearly identifiable events for the user (e.g. if meant for clinical providers, it is relevant to clinical practice)
- Indicator permits useful comparisons
- Indicator is evidence-based

The concepts behind these criteria were and remained highly influential. For example, these same themes can be observed in the criteria for ideal indicators provided by [Campbell and Lester \(2010\)](#) almost a decade later.

- Acceptibility: Both GPs and patients will accept assessment of the indicator.
- Attributable: A GP’s performance on an indicator relates well to his or her own performance and not to external factors.
- Feasibility: Reliable data is available.
- Reliability: There is minimal measurement error in the indicator.

- Sensitivity to Change: Changes to clinical practice can be observed through the indicator.
- Predictive Value: The indicator is closely linked with important clinical Outcomes.
- Relevance: There is evidence of incongruity between actual practice and best practice.

Building upon these research studies, governmental bodies and medical organisations weighed in on how to best choose performance indicators. The Royal Statistical Society’s Working Party on Performance Monitoring the Public Services published a report in 2004 that included a list of fourteen guidelines for appropriate indicators in evaluating the quality of healthcare (as well as other government services) ([Bird et al., 2005](#)). Among other recommendations, the Working Group suggested that indicators should be directly linked with the NHS’s primary objectives, assessed through a common methodology, be collected by a group representative of the national population, and conform to international standards. In 2006, the United States Institute of Medicine issued guidance on the selection of valid indicators ([Institute of Medicine, 2006](#)). It identified good performance indicators as being scientifically sound, feasible, important, aligned across datasets, and comprehensive. In 2008, the Association of Public Health Observatories and the UK’s NHS Institute for Innovation and Improvement published “The Good Indicators Guide” ([Pencheon, 2007](#)). In this guide, the authors set forth a set of criteria for good performance indicators, including: importance and relevance, validity, possibility, meaning, and implications. Several of these criteria map well to many of the challenges of healthcare evaluation described by [Ozcan \(2005\)](#). For example, the Guide’s validity criteria ensure that the indicators account for patient participation and medical heterogeneity. Armed with well-defined guidelines, the process of selecting indicators became more systematic and, consequently, robust ([Boulkedid et al., 2011](#)).

As this area of research matured and performance indicators were tested and implemented in real-world settings, more specific, detailed, and practical guidance was issued for choosing the best performance indicators and using them appropriately. [Wollersheim et al.](#)

(2007) published a review of best practices for the selection and application of performance indicators. The review included practical advice, such as a seven-step plan for implementing a performance indicator system, specific databases to search for finding previously implemented performance indicators, a recommended structure for performance indicator selection committees, and vivid examples of the development and use of performance indicators in three clinical contexts. The authors emphasized that, when developing and selecting performance indicators, selection committee should first survey measures of quality in a broad set of clinical domains, match the indicators to the evidence base, and then prioritise the indicators through a consensus model with experts. In 2010, The Kings Fund issued a report on practical issues of the selection and use of performance indicators (Raleigh and Foot, 2010). In this report, the authors advised that only the best performance indicators be adopted into practice. They also advised using the Donabedian Model (Donabedian, 1980) to choose indicators appropriately balanced among the structure, process, and outcome domains. Using this framework, they raised awareness of potential pitfalls for indicators in each domain; for example, they describe that process indicators, while more responsive and less prone to several forms of bias than outcome indicators, were susceptible to gaming and may require new forms of data collection. The Kings Fund report further explained how indicators should be analysed. In particular, the authors strongly recommend that analysis of healthcare quality measures be adjusted for case-mix and determinants of health outside of the healthcare services industry.

Opinions splintered around the concept of thresholds. Traditionally, performance indicators were designed such that a given practice's performance on a specific metric was measured against an *a priori* fixed threshold(s). Performance indicators with fixed threshold – also known as absolute targets or criterion referenced targets – have several advantages (Doran et al., 2014). For example, physicians can develop specific strategies to reach a set target. Further, physicians can be assured that their efforts, if successful, will not go unrewarded. However, absolute targets also have some severe disadvantages. Healthcare payers cannot

know their costs upfront; if more or less physicians achieve the target than previously expected, it could lead to a budget surplus or deficit. Further, quality monitors may not have complete information on the full capabilities of physicians and, consequently, may set the thresholds too low or too high. To ameliorate these issues, some began advocating for relative quality measures, in which a healthcare provider received a reward depending upon their place in the distribution of all providers for a particular metric ([Chien and Rosenthal, 2013](#)). Notably, relative thresholds were adopted in the United States Medicare program through the Physician Value Based Payment Modifier provision of the Affordable Care Act. Some have advocated for a middle ground; in 2009, NICE recommended fixing performance targets to different points of the previous year's distribution of performance ([Doran et al., 2014](#)). In this way, practices would know *a priori* the targets they need to achieve, but thresholds would have a reasonable, empirical basis.

Despite the depth and rich history of the literature on performance indicators, there is no universally accepted method or international standard for developing performance indicators ([Shekelle, 2013](#); [Stelfox and Straus, 2013](#)). Some approaches vary only semantically, while others have more substantial differences. For example, [Stelfox and Straus \(2013\)](#) take a population-health approach by recommending that quality indicators be restricted to clinical problems with “large burden of illness to justify quality measurement and improvement efforts”. [Shekelle \(2013\)](#), on the other hand, takes a more individual-centric approach, noting that while quality of care for rare diseases may not meaningfully affect population health, it is only equitable that those unfortunate enough to have rare diseases receive the same protection from performance indicators. Instead, individual institutions need to make their own “judgement calls” to determine which best practices from this literature work best in their own clinical, organisational, and political context.

2.3 The Quality and Outcomes Framework

2.3.1 Selecting Performance Indicators for UK General Practitioners

In this vein, the UK did not adopt a formal performance indicator selection protocol prescribed by the literature in its Quality and Outcomes Framework pay-for-performance scheme, opting instead to form their own. However, common themes of the discourse related to selection and analysis of healthcare performance indicators laid the foundation for the initial implementation and subsequent restructures of the Quality and Outcomes Framework scheme, and references to this literature (both explicit and implicit) can be observed in QOF policy documents and in the QOF structure.

The first set of indicators were negotiated by representatives from the British Medical Association and the NHS over an 18-month period ([Doran et al., 2014](#)). While recollections of these initial negotiations are sparse in the literature, there were some set ground rules for developing the indicators, and these ground rules appear to be linked with the preceding literature. For example, GPs were given the ability to exclude certain patients from their indicators based upon refusal by the patient or clinical inappropriateness ([Roland and Guthrie, 2016](#)), mitigating concerns that patient participation and heterogeneity would interfere with quality measures (though increasing concern that GPs may use exclusions to game their evaluations). Results from patient experience surveys were included as measures of quality, emphasizing that the QOF would take a comprehensive approach to quality (albeit, initial indicators related to patient satisfaction were relatively weak until the adoption of the GP Patient Survey in 2008). The selected indicators were designed with a fixed threshold(s), as was most common in the literature.

By 2005, the performance indicators were selected more systematically. That year (and again in 2007), in line with the literature's emphasis on comprehensive performance indicators, topics for development were garnered from a wide swath of stakeholders, including physicians, patients, non-profit organisations, governmental agencies, and private industry

([Lester and Campbell, 2010](#)). These topics were then matched to the evidence base and prioritised by an expert panel including representatives from the British Medical Association’s General Practitioners’ Committee and the Department of Health. In 2007, the indicator selection protocol also included meetings with groups that had made suggestions that were ultimately prioritised and the use of a modified Delphi procedure, a formal consensus system based on both scientific literature and expert opinion.

Since 2009, the UK Government’s National Institute for Health and Clinical Excellence (NICE) has been formally tasked with managing the process by which performance indicators for the QOF scheme are developed ([Sutcliffe et al., 2012](#)). While the final indicators are still selected through a negotiation between the NHS the British Medical Association’s General Practitioners’ Committee, the NICE process has effectively set the agenda for which performance indicators should be added, retained, or removed each year by developing a menu of “appropriate” or “suitable” indicators ([NICE, 2019](#)).

Initially, NICE quality indicators were developed through a structured process led by a NICE-supported but formally independent Quality Standards Advisory Committee, which was made up of 21 standing members (including physicians, administrators, public health officials, and laypeople) and five specialists ([Bennett et al., 2014](#)). Today, the process is led by a similar group called the NICE Indicator Advisory Committee ([National Institute for Health and Care Excellence, 2017](#)). In the first stage of the process, referred to as “Topic Overview”, NICE staff engaged with topic-specific stakeholders for two weeks and prepared briefings on the areas discussed, using the Donabedian conceptual model to guide their exploration ([Donabedian, 1980](#)). Then, in the “Prioritising Quality Areas” stage, the Advisory Committee received the NICE briefings and met to agree on topic-specific priorities. Priority was determined according to (A) evidence that quality of care varies across GPs, (B) feasibility of improving quality in the area, and (C) possibility of constructing a valid measure around that topic. Specifically, the committee is advised to rate proposed indicators against the following criteria:

- The proposed indicator should focus either on a health outcome or on a process that is closely linked to improved outcomes.
- The proposed indicator should focus on a national priority, particularly one that is specifically identified by either NHS England, Public Health England, or the devolved administrations.
- The proposed indicator should be feasibly measurable, either through existing or possible data collection systems.
- The proposed indicator should relate to an area where there is currently variation in practice and evidence that adoption of best practices can improve outcomes.
- The proposed indicator should reflect upon the actors it intends to measure (e.g., general practitioners).

Although not formally recognized in NICE documentation, scholars and regulators have commented that additional criteria apply to the selection of QOF indicators ([Campbell and Lester, 2010](#); [Gillam and Steel, 2013](#)). They suggest that a clinical area’s “QOFability”, or suitability for inclusion in the QOF program, relate to (A) its prevalence in the population, (B) its effect on morbidity and mortality, and (C) the ease with which it is measured. Further, specific indicators are deemed more “QOFable” if they are easily accessed from the Quality Management and Analysis System (QMAS), based on Government guidelines, well defined, universally achievable, unlikely to be “gamed”, and directly related to the GP’s actions. For example, indicators involving medical procedures and drug interventions may be prioritised in the QOF over community health interventions simply because the former are more straightforward than the latter ([Gillam and Steel, 2013](#)). Further, some have noted that the NICE Indicator Advisory Committee, which includes a good number of GPs, may reject quality indicators that make it more challenging to achieve benchmarks from year to year.

With priorities set by the advisory committee, NICE and the NHS began the process of drafting the indicator. The two agencies (while accepting public consultation) to define the indicator, define a methodology that could be used to collect data on the indicator, perform feasibility and/or pilot testing on the indicator, assess cost-effectiveness of adopting the indicator, and finally conduct a review of possible thresholds for the indicator ([Campbell et al., 2011](#)). In cases where a relevant indicator already exists, it is reviewed in this stage, as well. Through all of these steps, NICE arrives at an indicator statement, which is then subject to a 4-week public review by stakeholders and members of the public. The draft statement and stakeholder comments are then presented to the Advisory Committee, which makes necessary adjustments. Next, NICE staffers check the validity of the statement and its consistency with other indicators. The final product is then passed along to the NICE Guidance Executive, a group of NICE directors, who gave final approval for the indicator statement. After a performance indicator is approved by NICE, it often becomes the subject of future academic and policy research, which can then be used to change or eliminate the indicator.

Once all NICE indicators are approved through this process, they are published on NICE's website (along with any related information collected during the selection process) as suitable for implementation. Importantly, NICE performance indicators are not automatically implemented in the QOF scheme but, instead, only enter the QOF scheme through contractual negotiation between the NHS and the British Medical Association's General Practitioners Committee. That is, NICE indicators are implemented (and corresponding point values set) when a new contract is agreed upon by the NHS and GPs ([NHS England, 2019](#)). Consequently, through the contract negotiation process, QOF performance indicators are subject to an additional political feasibility constraint beyond NICE's recommended indicators. Indeed, some believe that the General Practitioners Committee intentionally block reforms that would make QOF benchmarks more difficult to achieve.

2.3.2 The Purpose of the QOF: Pay Raise or Incentive to Improve?

An important quality of any performance metric is that it be aligned with overarching goals. Unfortunately, while the QOF's focus on quality and outcomes is clear, its overall purpose has not always been well defined. Consequently, it is important to note that the evaluation of any indicator may depend partially on the evaluator's own view of the purpose of the QOF scheme.

At the time the QOF was introduced, it was widely regarded as a bargaining chip for increased GP pay. The NHS had already agreed that an increase in pay was necessary to sustain the GP workforce, and participation in a pay-for-performance scheme was an at least politically necessary *quid pro quo* (Roland and Guthrie, 2016). It remained ambiguous whether the QOF scheme was simply a means to justify a pay raise – intended to reward GPs for continuing the work they were already doing – or whether the QOF was intended to become a mechanism to incentivise GPs to adopt higher-quality practices in their clinics. From the initial indicators selected, it appeared that the QOF was simply a mechanism to systematically observe GP quality while offering them higher pay. According to Roland (2007), most of the original indicators were based on existing guidance, such as the National Service Frameworks, and within the first year of implementation, GPs earned 91.3% of total QOF points (Dixon et al., 2011).

Subsequently, however, some indicators that were not already common practice were adopted in order to incentivise changes in practice (Roland and Guthrie, 2016). By 2006, the NHS stated that “[t]he QOF is not about performance management of general practice but about resourcing and then rewarding good practice.” (National Health Service, 2006a). Since the 2009 QOF scheme reforms, NICE has become more clear that the QOF indicators are intended to incentivise quality improvement rather than reward the status quo. For example, NICE asserts that its quality indicators are developed with the “aim to shape measureable quality improvements by identifying the key areas for improvement within a particular area of health, public health or social care” (Bennett et al., 2014). Further,

one of the criteria NICE uses to select performance indicators relates to the feasibility of improvement ([National Institute for Health and Care Excellence, 2017](#)).

Whether this is true, in practice, however, is debatable. As of 2013, most indicator's upper thresholds were held below the average practice's achievement rate ([Doran et al., 2014](#)), and the average practice participating in the QOF scheme achieves over 90% of available points. On the other hand, even in practices with high achievement rates relative to these thresholds, the QOF scheme provides an incentive to maintain a reasonable level of quality.

2.3.3 The QOF Scheme in Practice

When the QOF scheme was introduced as the largest pay-for-performance scheme in the world in 2004, it linked 146 indicators across four domains (clinical, organisational, patient experience, and additional services) with approximately 20-25% of overall practice income ([Guthrie, McLean and Sutton, 2006](#)). By participating in the (voluntary) program and meeting thresholds within these indicators, practices had the opportunity to earn up to 1050 total points. At the end of the year, point totals would be assessed and practices would be commensurately remunerated, with marginal adjustments made for workload and patient health. For reference, the average practice could expect to receive approximately 77 GBP for each point it earned (see Table 4).

The Clinical Domain commanded the greatest weight in the QOF with 76 indicators across 11 areas holding an allocation of 550 points (52.4% of total points). The organisational domain was the next most significant, with 56 indicators across five areas holding an allocation of 184 points. The patient experience domain had just four indicators across two areas but with an allocation of up to 100 points. The Additional Services domain, which included ten indicators over four areas, held the least weight in the scheme, with an allocation of only 36 points. In addition, there were some additional opportunities for practices to earn points, often called "depth of quality measures". A holistic care measure, which assessed achievement over the entire clinical domain, was worth up to 100 points. A quality practice

Table 1: Domains and Areas of 2004-2005 QOF Scheme

Clinical Domain Total Available Points: 550 (52.4%)	Organisational Domain Total Available Points: 184 (17.5%)
Coronary heart diseases Left ventricular dysfunction Stroke and transient ischaemic attack Hypertension Diabetes mellitus Chronic obstructive pulmonary disease Epilepsy Hypothyroidism Cancer Mental health Asthma	Records and information Patient communication Education and training Medicines management Clinical and practice management
Patient Experience Domain Total Available Points: 100 (9.5%)	Additional Services Domain Total Available Points: 36 (3.4%)
Patient survey Consultation length	Cervical screening Child health surveillance Maternity services Contraceptive services

NOTE: The data for this table was reported by [Lester and Campbell \(2010\)](#). Total Available Points do not add to 1050 (100%) because 180 points are designated to “depth of quality measures” outside the four domains.

measure, which assessed overall achievement in the non-clinical domains, was worth up to 30 points. Finally, an access measure, which assessed patient access to clinical care, was worth 50 points. The areas assigned to each of the four domains are shown in Table 1.

Most QOF indicators had fixed lower and upper thresholds related to the percentage of eligible patients for whom an indicator was achieved ([Doran et al., 2014](#)). If a practice’s performance was below the lower threshold, it would receive no QOF points and no remuneration. If it met or exceeded the upper threshold, it would receive the maximum QOF points and payment. If the practice’s performance fell between the lower and upper thresholds, it would receive points and payment between the minimum and maximum amount. All QOF lower thresholds were set at 25% of eligible patients, while upper thresholds ranged from 50%

to 90% depending upon the perceived difficulty of achieving the indicator. Other indicators took on binary values, for which practices received full points when achieved and 0 points if not achieved. Most indicators had a time frame of 15 months. Table 2 provides several examples of indicators from the 2004-2005 QOF scheme, including their point value, lower threshold, and upper threshold.

While points achieved were directly linked with payment, the QOF scheme had a few nuanced features that modified the relationship between points and payment through practice-level clinical condition prevalence ([Dixon et al., 2011](#)). For each clinical outcome, payment was calculated using the square root of prevalence among the practice's registered patients. For this reason, practices with high disease prevalence received lower payment per patient treated to the indicator's standard than practices with lower disease prevalence. Further, those practices with very low disease prevalence (under 5% of the maximum) were rewarded as if they had a higher prevalence even though they had fewer patients to care for. These adjustments disproportionately disadvantaged smaller practices from areas with lower average health relative to larger practices with higher average health.

The overall structure of the QOF scheme (i.e., practice-level pay-for-performance based upon predetermined threshold values) has remained mostly constant since its inception, but changes have been made to the domains, areas, indicators, and thresholds. A review of changes to the QOF for England is summarised in Table 3.

Table 2: Sample Performance Indicators and Point Value, QOF 2004-2005

Indicator	Point Value	Lower Threshold	Upper Threshold
CHD 7. The percentage of patients with coronary heart disease whose notes have a record of total cholesterol in the previous 15 months	7	25%	90%
LVD 1. The practice can produce a register of patients with CHD and left ventricular dysfunction	4	NA	NA
BP 2. The percentage of patients with hypertension whose notes record smoking status at least once	10	25%	90%
MH 2. The percentage of patients with severe long-term mental health problems with a review recorded in the preceding 15 months. This review includes a check on the accuracy of prescribed medication, a review of physical health and a review of co-ordination arrangements with secondary care	23	25%	90%
EPILEPSY 4. The percentage of patients age 16 and over on drug treatment for epilepsy who have been convulsion-free for last 12 months recorded in last 15 months	6	25%	90%
INFORMATION 1. The practice has a system to allow patients to contact the out-of-hours service by making no more than two telephone calls	0.5	NA	NA
EDUCATION 4. The practice has a system to allow patients to contact the out-of-hours service by making no more than two telephone calls	3	NA	NA

NOTE: These examples are reported from [National Health Service \(2005\)](#). Those indicators with an NA upper and lower threshold are binary indicators, where the practice receives full points if the indicator is satisfied and 0 if the indicator is not satisfied.

Table 3: Timeline of Changes to the QOF Scheme (England)

Contract Year	Summary of Changes
2004/2005	First year of QOF
2005/2006	Average payment per QOF point increases from £77.50 to £124.60
2006/2007	Lower thresholds raised from 25% to 40%. Some upper thresholds changed. 166 QOF points redistributed. 50-point access measure removed, reducing total QOF points to 1000.
2007/2008	
2008/2009	58 QOF points reallocated to an indicator that patients have access to a healthcare professional within 48 hours (PE7).
2009/2010	72 QOF points reallocated. Clinical area for cardiovascular disease primary prevention added, and new indicators were introduced for heart failure, chronic kidney disease, depression, diabetes areas. Two high-value patient experience indicators were removed. Prevalence adjustment (also known as “square rooting”) removed.
2010/2011	No changes due to Swine Flu.
2011/2012	Indicator that patients could access a healthcare professional within 48 hours (PE7) was removed. Upper thresholds for two indicators were increased by one percentage point.
2012/2013	New clinical areas for peripheral arterial disease and osteoporosis. Seven indicators were retired, and 12 new indicators were introduced. Lower thresholds were increased by 5-10 percentage points for all indicators, and upper thresholds were increased by 4-10 percentage points for 13 indicators.
2013/2014	Devolved governments begin negotiating their own QOF schemes. Upper thresholds were set at the 75 th percentile of achievement and lower thresholds were set at 40 percentage points below the 75 th percentile for 20 indicators. Base payment per point was increased from £133.76 to £156.92. The organisational domain was removed, and new domains were added for “Public Health” and “Quality and Productivity” (the Additional Services domain was renamed “Public Health – Additional Services”). The time frame for most indicators was reduced from 15 months to 12 months. The total number of available QOF points was reduced from 1000 to 900.
2014/2015	Several indicators that were unpopular with GPs were removed. The “Quality and Productivity” and “Patient Experience” domains were removed, as well as the hypothyroidism, child health surveillance, and maternity areas. Age restrictions for learning disability indicators were removed, while those for blood pressure were raised from 40+ to 45+. The total number of available QOF points was reduced from 900 to 559.

Contract Year	Summary of Changes
2015/2016	Upper thresholds were set at the 75 th percentile of achievement and lower thresholds were set at 40 percentage points below the 75 th percentile for remaining indicators. The average payment per QOF point was increased to £160.15. Minor changes were made to the atrial fibrillation, coronary heart disease, dementia, chronic kidney disease, and obesity areas.
2017/2018	No changes
2018/2019	No changes

NOTE: These changes are adapted from [Institute of Healthcare Management \(2018\)](#) with additions from [Doran et al. \(2014\)](#), [Stokes \(2014\)](#), and [National Health Service \(2014\)](#). For years 2004/2005 to 2012/2013, the QOF point universally applied to England, Wales, Scotland, and Northern Ireland. Those changes from 2013/2014 on apply only to England.

Table 4: Average Payment Per QOF Point (England)

Contract Year	Average Payment Per Point	Source
2004/2005	£77.50	Institute of Healthcare Management (2018)
2005/2006	£124.60	Gillam and Siriwardena (2011)
2006/2007		
2007/2008		
2008/2009		
2009/2010		
2010/2011	£127.29	Institute of Healthcare Management (2018)
2011/2012	£130.51	Institute of Healthcare Management (2018)
2012/2013	£133.76	Doran et al. (2014)
2013/2014	£156.92	Doran et al. (2014)
2014/2015		
2015/2016		
2017/2018		
2018/2019	£179.26	Institute of Healthcare Management (2018)

NOTE: For years 2004/2005 to 2012/2013, the QOF scheme universally applied to England, Wales, Scotland, and Northern Ireland. Those values from 2013/2014 on apply only to England. Payment per point is subject to adjustment for practice demographics; before 2009/2010, QOF payments per point was subject to “square rooting”. These estimates are based upon 1,000 total QOF points each year.

Perhaps the most dramatic revision of the QOF program occurred with the 2006 General Medical Services contract at the behest by the newly formed Expert Panel. Most lower thresholds for clinical indicators were raised from 25% to 40%, and a few upper thresholds were also adjusted (12 of 50 were raised by 5-20 percentage points and 1 was lowered by 10 percentage points) ([Doran et al., 2014](#)). Further, seven new clinical areas were added: dementia, depression, chronic kidney disease, atrial fibrillation, palliative care, obesity and learning disabilities ([Applebee, 2006](#)). The left ventricular dysfunction area was replaced by heart failure. The total available QOF points was reduced from 1050 to 1000, as the 50 point access measure was moved to the Directed Enhanced Services scheme, an optional program for practices to meet the needs of their local communities ([National Health Service, 2006b](#)). Indicators within existing clinical areas were amended, and point values were moved around, both within the existing indicators and from existing areas to indicators in new areas. The

2006 revision of the QOF scheme was seen as an increased burden on GPs. Even in cases where indicators had been removed in the revision, the contract negotiators had stated that those indicators were now part of standard, “good medical practice” ([Applebee, 2006](#)).

A minor revision of the QOF scheme occurred in 2009, when NICE took leadership over the indicator selection process. A new clinical area for cardiovascular disease primary prevention was introduced, and new indicators were added to the heart failure, chronic kidney disease, depression, and diabetes clinical areas ([National Health Service, 2010](#)). Additional changes were made to the patient experience and additional services domains. Importantly, the prevalence adjustment mechanisms were removed, which ensured the QOF would fairly reward practices in lower average health areas ([Dixon et al., 2011](#)).

In 2010, NICE recommended bench-marking thresholds against the distribution of the past year’s performance, but this suggestion was rejected during contract negotiation ([Doran et al., 2014](#)). Instead, both parties agreed to minor increases in the upper thresholds in 2011/2012 and slightly more substantial increases to lower and upper thresholds in 2012/2013. Minor changes to the indicators also occurred, including the addition of two new clinical areas for PAD and osteoporosis.

The next notable change for the QOF scheme occurred in 2013, when it was decided that devolved governments would negotiate their own QOF contracts. From the 2013/2014 contract forward, the QOF scheme for England, Wales, Scotland, and Northern Ireland were permitted to vary ([Institute of Healthcare Management, 2018](#)). England made a number of significant changes, including the addition of “Public Health” and “Quality and Productivity” domains and the removal of the “Organisational” domain. A reduction of the total number of available QOF points occurred in line with an increase in the average payment per QOF point. England also adopted a system in which upper thresholds would be set at 75th percentile achievement from the previous year ([Doran et al., 2014](#)). To give GPs time to adapt to changes in thresholds, upper thresholds were changed for only 20 indicators in 2013/2014, with the remainder indicators’ upper thresholds changed in 2015/2016

(postponed from 2014/2015).

Another set of changes occurred in 2014/2015. The “Quality and Productivity” and “Patient Experience” domains were removed, leaving just the “Clinical”, “Public Health”, and “Public Health – Additional Services” domains. The total number of available QOF points were reduced from 900 to 559.

Importantly, these changes to the QOF scheme corresponded with variation in the balance of weight among domains (Table 5) and the average QOF points of practices (Table 6). The weights among domains are relatively constant from 2006/2007 to 2011/2012, and general practices achieved, on average, a very high percentage of points among each of the domains (with the exception of the “Patient Experience” domain from 2008/2009 to 2010/2011) over the same time period. Notably, percent achievement accounts for patient exclusions (i.e., where GPs select certain patients to exclude from QOF calculations).

Table 5: Distribution of Points Available across Domains, 2006-2012

Contract Year	Additional Services	Clinical	Organisational	Patient Experience	Total
2006/2007	3.6%	65.5%	18.1%	10.8%	100.0%
2007/2008	3.6%	65.5%	18.1%	10.8%	100.0%
2008/2009	3.6%	65.0%	16.8%	14.6%	100.0%
2009/2010	4.4%	69.7%	16.8%	9.2%	100.0%
2010/2011	4.4%	69.7%	16.8%	9.2%	100.0%
2011/2012	4.4%	66.1%	26.2%	3.3%	100.0%

Table 6: Average Achievement across Domains, 2006-2012

Contract Year	Additional Services	Clinical	Organisational	Patient Experience	Total
2006/2007	96.5%	96.3%	92.5%	95.9%	95.5%
2007/2008	97.0%	97.5%	94.5%	97.2%	96.8%
2008/2009	97.3%	97.8%	95.8%	84.2%	95.4%
2009/2010	95.3%	95.9%	96.3%	71.5%	93.7%
2010/2011	97.1%	96.8%	97.4%	72.6%	94.7%
2011/2012	97.0%	97.0%	96.4%	99.0%	96.9%

As of 2018/2019, the QOF remains integral part of general practice in the UK. While the

program is still voluntary, over 95% of GP practices (N=6873) participated with an average achievement score of 539.2 points (of 559 possible) and 13.0% achieving the maximum score ([National Health Service, 2019](#)).

3 Social Influence in General Practice Quality

3.0.1 Introduction to Social Influence

Social influence may be an important and understudied determinant of quality in healthcare. The notion that physicians may influence the quality of the healthcare others deliver through social influence² is borne from two existing bodies of literature. The first field demonstrated that, while successful quality improvement interventions are rare ([Conry et al., 2012](#)), physicians respond well to information delivered by peers and peer reviews of their practice. The second field documented a specific pattern of medical practice variation in which variance is lower within groups than between groups. Peer influence was proposed as one possible mechanism for this pattern. To date, only a few studies have attempted to test the role of social influence, primarily using data on physicians who practice in multiple locations. If social influence has a significant impact on clinical decision making in regards to healthcare quality, it could become a useful policy lever for quality improvement.

3.0.2 Peer Influence in Quality Improvement

Peer influence, encompassing clinical peer review, peer feedback, and peer teaching, is an established vehicle for quality improvement ([Sargeant et al., 2015](#)). In fact, the Joint Commission on Accreditation of Healthcare Organizations required physician peer review as part

²The research literature has frequently used the terms “peer influence” and “social influence” interchangeably. For purposes of clarity in this dissertation, I distinguish between these terms. “Peer influence” will refer to the influence that all similar physicians hold over an individual physician, whereas “social influence” refers to how direct interactions with other physicians may affect a physician’s behaviour. For example, peer review by a government body of physicians or a formal, instructional visit led by another physician would fall under the category of “peer influence”. Influence among physicians who are friends or work in the same practice – those who interact often and on an on-going basis – would be considered “social influence”.

of its accreditation process as far back as 1952 ([Goldberg, 1984](#)). Since at least the 1970s, robust studies have documented the efficacy of peer-based interventions in improving quality in several domains of healthcare delivered by general practitioners and other physicians ([Mittman, Tonesk and Jacobson, 1992](#); [Jamtvedt et al., 2003](#)). Today, “peer review and clinical audit” are two of the key strategies for quality improvement endorsed by the World Health Organization ([World Health Organization, OECD and International Bank for Reconstruction and Development/The World Bank, 2018](#)).

Several studies and reviews have demonstrated that peer-based interventions are successful method of improving quality in general practice ([Wensing, van der Weijden and Grol, 1998](#)). For example, peer-based interventions have successfully improved preventive medicine in general practice ([Hulscher et al., 1999](#)). Using a randomised controlled trial design, [Inui \(1976\)](#) found that a tutorial session improved dietary counseling and hypertension monitoring practice by 44% relative to a control group with no intervention. Similarly, [Cockburn et al. \(1992\)](#) reported that, in a randomised controlled trial, an instructional visit by a facilitator or a courier improved GPs adherence to smoking counselling guidelines by 15% and 5% relative to sending the information by mail. Relatedly, several studies have shown that peer-to-peer interventions can be successful in reducing over-prescribing rates ([Arnold and Straus, 2006](#); [Avorn and Soumerai, 1983](#)). For example, a multi-institutional, quasi-experimental study of Mexican physicians found that an intervention involving a managerial peer review committee improved prescribing practices among up to 40% of physicians ([Pérez-Cuevas et al., 1996](#)). In a parallel randomised controlled trial using peer feedback groups to address over-prescribing by Dutch GPs in the treatment of asthma and urinary tract infection found that peer groups significantly improved clinical practice ([Veninga et al., 2000](#)). Peer-based interventions have also shown efficacy in reducing unnecessary referrals to secondary care. Six weekly “cluster group” meetings were able to successfully reduce NHS GP specialist referral from 5.5 patients per 1000 to 4.3 patients per 1000 ([Evans, Aiking and Edwards, 2011](#)).

Peer influence is also a well-tested and evidence-based vehicle for quality improvement for several medical fields outside of general practice. A study of the Mayo Clinic Health System found that the implementation of clinical peer review reduced the all-cause mortality rate in a critical access hospital by 50% ([Deyo-Svendsen et al., October/December 2016](#)). Two studies found that programs involving peer-to-peer site visits reduced mortality associated with coronary artery bypass grafting ([Holman et al., 2001](#); [O'Connor et al., 1996](#)). A common and well-studied form of peer-based interventions to improve healthcare quality, particularly within specialty care, are Communities of Practice (also known as clinical communities) ([Aveling et al., 2012](#); [Endsley, Kirkegaard and Linares, 2005](#); [Iedema, Meyerkort and White, 2005](#)). Generally self-organised groups of physicians in the same specialty, Communities of Practice share relevant clinical information and form peer review practices. Communities of Practice may involve interactive training sessions, story-sharing, and promotion of internal audits. These Communities have demonstrated efficacy for improving clinical practice. For example, a Michigan-based Community of Practice intervention reduced the number of infections per catheter-days in Intensive Care Unit settings by more than 65% ([Pronovost et al., 2016](#)).

Collectively, the empirical evidence in favor of peer-based quality improvement interventions has made these among the most popular quality improvement interventions in medicine. Scholars have suggested several theoretical reasons for the success of these interventions. Some suggest that the vividness and practical applicability of story-telling may play a role; physicians may communicate with one another in a narrative form that other physicians can directly apply to their practices ([Jordan et al., 2009](#)). Others hypothesize that peer-based interventions encourage a level of interactivity that facilitate learning ([Johnson, Manyika and Yee, 2005](#)) or that the trust in peer-relationships facilitate adoption ([Wenger, 1999](#)). Recently, some have suggested that peer-based interventions could lead to consistent relationships that impose social influence ([Meltzer et al., 2010](#)).

3.0.3 Within-Group and Between-Group Medical Practice Variation

Largely separate from the literature on healthcare quality improvement interventions, studies stretching over several decades have found consistent patterns of Medical Practice Variation in the UK. In the 1930s, Dr. J. Alison Glover reported in *Proceedings of the Royal Society of Medicine* that rates of tonsillectomies among children varied significantly among different geographic regions (Glover, 1938). Subsequent studies of small area variation in medical practice extended this finding to several areas of medical practice and surgery (McPherson et al., 1982; Wennberg and Gittelsohn, 1973, 1982; Paul-Shaheen, Clark and Williams, 1987; Jaffer et al., 2010; O'Connor et al., 1999).

More recently, researchers have used microlevel diagnosis and treatment data and multi-level modelling techniques to document a more essential feature of Medical Practice Variation: variance in medical practice tends to be greater between rather than within medical groups. Westert, G.P. (1992), for example, found that duration of hospital stays was more similar within hospital wards than across wards. This pattern in variation in duration of hospital stay has now been observed in several settings (van de Vijssel, Heijink and Schipper, 2015). Using data from the Dutch National Survey of General Practice and techniques from social network analysis, de Jong, Groenewegen and Westert (2003) found that GPs tend to be more similar within partnerships than between partnerships. For example, they found that self-reported practice behaviour (medical techniques used, handling work style elements), time spent on clinical activities, and practice choice in response to a given complaint were more similar within practices than between practices. It should be noted that not all published research has confirmed this trend (e.g., Gutacker et al., 2018).

Several theories were posed to explain this pattern of within versus between variation. Some suggested that selection of doctors to practices may play a role. Similar people tend to attract (Fehr, 1996), and physicians may choose to work with and recruit other physicians who are similar to them. Others suggested that (observed or unobserved) differences in circumstance may vary practice between but not within practice. Physicians within a given

practice presumably have the same incentive structure and the same access to information and resources (de Jong, Groenewegen and Westert, 2003), but this need not be the case for physicians in different practices. If, for example, physicians in practice A is judged on a different metric from those in practice B, physicians may respond accordingly (Krasnik et al., 1990) and exhibit medical practices that vary markedly between practices. However, all physicians within practice A or within practice B would be aligned.

Social influence was also hypothesized as an explanation of this pattern. Sociologists have observed that people desire social acceptance, avoid deviation from and adapt towards social norms (Eddy, 1984; Dambrun, Guimond and Duarte, 2002), and evaluate their own performance in relation to role equivalents (Fehr, 1996; Erickson, 1988). In this vein, “people both shape and are shaped by social networks”, even when adapting to others conflicts with rational choice (Pescosolido, 1992). Further, social cues may be a vehicle for physicians to avoid risk in the context of uncertainty (Bandura, 1986). Indeed, physicians report that “informal consultations” with other physicians are common (Keating, Zaslavsky and Ayanian, 1998) and even that they may affect practice and quality of care (Eisenberg, 1985; Keating et al., 2007; Geneau et al., 2008). Consequently, physicians may “follow the pack” with the understanding that deviations from the practices of their proximate peers could reduce social approval or increase risk (Westert and Groenewegen, 1999).

3.0.4 Empirical Evidence of Social Influence Effect

Only a small number of studies have adequately tested the hypothesis that social influence among physicians may affect their practice. These studies typically evaluate how physicians who practice in multiple locations vary their practice based upon location. In this way, the researchers exclude the possibility that how physicians select into practices impacts the within-group variance; because the physician is the same person in both settings, selection can be ruled out. Further, this approach mitigates (but does not eliminate) the possibility that differences in circumstance explain the pattern of within versus between group varia-

tion. It is true that a physician could work in different practice settings which, themselves, have many different circumstances, such as equipment availability, support staff, or incentive schemes. However, given that the physician is the same in both practice settings, the knowledge and intrinsic resources of the physician should be the same. Consequently, while acknowledging the limitations, a comparison of how physicians who practice in multiple locations vary their practice based upon location is a reasonable estimate of the impact of social influence on medical practice.

The earliest example comes from [Griffiths, Waters and Acheson \(1979\)](#), who studied physicians who performed elective repair of inguinal hernia in multiple hospital settings within the NHS. They found that physicians who performed the procedure at more than one hospital significantly varied the patient's length of postoperative stay among the hospitals, while physicians who performed the procedure at just one hospital were relatively consistent. [Westert, Nieboer and Groenewegen \(1993\)](#) applied modern statistical approaches in a similar setting in the Netherlands and found similar results; physicians who performed similar procedures across similar patients at more than one hospital varied their patients' length-of-stay to match each hospital. Subsequently, [de Jong et al. \(2006\)](#) found a similar pattern for length-of-stay among multi-hospital physicians practising in the United States. To date, no study involving doctors in multiple locations has evaluated this pattern for general practitioners.

Apart from these, only a few studies have even hinted at an effect of social influence among physicians. [Chung et al. \(2003\)](#) found that new physicians who spent more time in the same residency (i.e., training) rotation shared more similar practice styles. While this points to social influence, it is possible that rotations were subject to differential lessons that were reinforced over time. [Goodwin et al. \(2013\)](#) studied the effect of hospitals on length of stay and discharge destination among patients in Texas. In a two-level model involving admissions and hospitalists, there was significant variation in practice among hospitalists. However, this variation was greatly attenuated with the inclusion of a third level: hospitals.

In this way, the authors find that the hospitalist-level variation is largely explained by the hospitals within which the hospitalists' practice. Again, it is possible this is indicative of a social influence effect; however, the selection or circumstance hypotheses offer alternative explanations.

3.1 Implications of A Mover Effect

3.1.1 Leveraging the Social Influence Effect

If social influence is determined to be a meaningful and significant driver of medical practice variation, it could become a powerful policy lever for improving healthcare quality in general practice. Selection of certain GPs into practices is challenging to overcome without an administrative take-over of general practitioners offices, and while some variation in circumstance (e.g., incentive schemes) can be managed, others (e.g., availability of equipment or information) can be substantially more difficult to mitigate. Several actionable steps can be taken by both the government and individual GP practices in order to leverage social influence.

Quality improvement programs should focus on not just giving GPs information on best practices. Instead, these programs should focus on shifting social norms towards those best practices.

The Government may consider instituting rotational programs or sabbaticals that expose people in worse performing practices to those in better performing practices. For example, the Government may incentivise GPs who recently left top-performing practices to spend a few months at a time in low-performing practices, with the hope that the newcomer will use his or her social influence to shift quality of practice upwards.

Similarly, the Government may consider incentivising permanent moves to lower performing practices. Previous research from Australia suggests that physicians would be willing to move locations for an incentive ([Scott et al., 2013](#)). If a GP mover is likely to affect the quality of her incoming practice, the benefits to that practice may outweigh the costs of

an incentive program (depending upon the incentive necessary for the GP to move and the expected benefit to the target practice).

GPs may wish to consider the social influence of candidates when hiring a new GP. These practices may wish to hire GPs from higher-performing practices in the hopes that the newcomers' practices will shift the social norm towards higher quality of care.

Conversely, if a GP practice is hiring an individual from a worse performing practice, the existing partners may want to emphasize the need for the newcomer to quickly adapt to the practice's standards.

3.1.2 Credibility of Quality Measures

While quality measures like those in the QOF scheme are commonplace in healthcare, they are not universally accepted. Some feel strongly that the quality measures are invalid or unreliable. For example, the "Altruism" model poses that practices perform in accordance with their available resources and information. In this case, rewarding high-performing practices is counterproductive; QOF payments are undeserved by GPs and fail to help low-resource, low-performing practices improve.

If there is an observable "mover" effect, however, this strongly suggests that the makeup of physicians within a given practice determine its performance. If the addition of one GP from a better or worse practice into the target practice can affect the target practice's QOF points, then the quality scores are a function of the GPs.

It is possible that a "mover" effect exists within the context of the "Altruism" model – for example, if a doctor moves in with better information, that would elevate the information of the practice. However, in this case, a mover from a worse practice should not meaningfully affect quality; the new person would benefit from the shared information in the practice, which would stay constant. Therefore, if there is a *negative* effect from GPs entering from worse practices, this is unaccounted for in the "Altruism" model. Consequently, the presence of a "mover" effect enhances the credibility of quality measures.

Table 7: Factors Driving GP Movement

	Any Move	Destination
GP	What Prompts a GP to Want to Move	What Determines Where the GP Moves
Practice	What Prompts a Practice to Hire a New GP	What Does the New Practice Look for in a New GP

NOTE: This table was developed by the author.

3.2 Factors Influencing GP Movement

There are several factors that could potentially influence GP movement from one practice to another, and it is important to consider these factors as mechanisms that could be driving the “mover effect”. The factors influencing GP movement can be divided into four main categories: factors that prompt a GP to move, factors the GP considers when choosing where to move, factors that prompt a practice to hire a new GP, and factors that the practice considers when choosing a new GP (see Table 7). For example, the top left box may include poor working conditions at an original practice, while the top right box may include high quality working conditions at a target practice. There are some studies that consider which factors are most important in setting GP’s preferences for practices (i.e., the top-right box of Table 7), but the empirical and theoretical literature on the other categories is limited.

3.2.1 Physician Preferences for Practices

One of the first studies to consider the factors that drive physicians’ movement within a developed country was conducted by [Beardow, Cheung and Styles \(1993\)](#). The authors surveyed GP trainees in the UK on which factors influenced their decision of practice within the North West Thames Regional Health Authority, finding that the most important factors were (A) a good relationship with existing staff and GPs at the practice, (B) whether the practice had a practice nurse or practice manager, (C) whether the practice had attached

health authority staff, (D) whether the practice offered educational opportunities, and (E) whether the practice had a good relationship with hospitals. The authors recommended enhancing these factors in areas with labour shortages. Other practice factors, such as high deprivation allowance, opportunities to work privately, whether the practice was near where the trainee grew up, or whether the practice was where the trainee studied, were considered less important factors. A similar postal survey of 101 GP trainees in the south west region of England studied the factors that trainees viewed as important in making career decisions (Rowsell, Morgan and Sarangi, 1995). The authors of that study found that time for leisure activities, on-call rota, maternity/paternity leave, and weekend work were the four most important factors in career decision-making. Less important factors included flexible work patterns, partner's paid work, salary, and regular work schedules. Yet another survey of 52 trainees located in North Trent found that trainees hoped to join "medium-sized practices with full ancillary and attached staff" and that relatively few were willing to work in an inner city or deprived area (Webb and Hannay, 1996).

One research team used responses to actual GP vacancy advertisements to assess which factors were most important to GPs (Carlisle and Johnstone, 1996). Specifically, the authors investigated the association between GP responses to 489 vacancies posted in the *British Medical Journal* between January and April 1995 and practice-level characteristics. They found that practices that did their own on-call work, practices not in the inner-city, and practices without deprived patients received a higher number of applications. The authors noted that, given their findings, it is possible that physicians may not have a strong preference for or against treating deprived patients but may instead simply wish to live in a non-deprived area.

Later studies used discrete choice models and conjoint analysis to determine GP's preferences for certain practice characteristics (Lagarde and Blaauw, 2009; Mandeville, Lagarde and Hanson, 2014). For example, Scott (2001) conducted a discrete choice experiment among a sample of 848 British GPs, and found that the strongest determinant of preference was

out-of-work hours. [Wordsworth et al. \(2004\)](#) similarly conducted a discrete choice experiment to assess which practice attributes were most important among GPs, giving special consideration to differences in the preferences among principal and sessional GPs. The authors found that, while some attributes had different levels of importance among principal and sessional GPs, both types of GPs preferred practices with less out-of-hours work and an increase in annual earnings. [Gosden, Bowler and Sutton \(2000\)](#) designed a questionnaire using conjoint analysis in which a respondent's preferences over eight characteristics could be derived from 27 practice scenarios. The respondents – a sample of 172 GPs who had recently taken on a principal post – were then asked to compare one hypothetical practice to each of the other 26 hypothetical practices. The authors found that the most important factor for GPs was aversion to moving to a high deprivation area, and that GPs preferred practices that had an extended primary healthcare team, afforded opportunities for the GPs to develop outside interests, paid higher salaries, had shorter working hours, and had shorter list sizes. Similar experiments have been used to elicit preferences from physicians in other developed nations. For example, [Günther et al. \(2010\)](#) performed a discrete choice experiment to discern what pecuniary and non-pecuniary job benefits would offset the disutility German physicians experience when moving from an urban to rural practice. The authors found that, all else equal, physicians required over 9,000 Euro per month to compensate for the disutility of moving from an urban to a rural practice, but that requirement diminished with the addition of non-pecuniary job benefits such as on-site childcare and fewer on-call duties. Collectively, these studies indicate that pay can be a major factor influencing whether a physician prefers one practice over another but that other non-financial incentives exist.

3.2.2 Other Forces for Movement

Only a few studies have considered what prompts a GP to want to move. ([Andreassen, Di Tommaso and Strøm, 2013](#)) find that favorable wages and tax conditions induce Norwegian medical doctors to consider moving from other healthcare professions to full-time practice,

Table 8: Neyman’s Notation for a Completely Randomised Experiment

Unit	Potential Outcome		Individual-Level Causal Effect
	Treatment $Y(1)$	Control $Y(0)$	
1	$Y_1(1)$	$Y_1(0)$	$Y_1(1) - Y_1(0)$
...	...	...	...
i	$Y_i(1)$	$Y_i(0)$	$Y_i(1) - Y_i(0)$
...	...	...	...
N	$Y_N(1)$	$Y_N(0)$	$Y_N(1) - Y_N(0)$

NOTE: This table was adapted from [Rubin \(2005\)](#).

which may indicate that financial incentives can motivate physicians to move.

4 Methods for Evaluating the Effectiveness of Health Policies

4.1 Theory of Causal Inference

The core theory of causal inference, termed the “Potential Outcomes framework”, comes from the largely separate work of two researchers in the early 20th century: Ronald Fisher (1890-1962) and Jerzy Neyman (1894-1981) ([Rubin, 2005](#)).

Neyman, then a Ph.D. student at the University of Warsaw, introduced the theory and, notably, a notational structure for causal inference in his dissertation. As shown in Table 8, individuals indexed at $i \in 1, \dots, N$ would be observed in both the treatment and the control conditions. For each individual, the outcome in the control condition $Y_i(0)$ would be subtracted from the outcome of the same individual in the treatment condition $Y_i(1)$ to compute individual-level causal effects.

The sample Average Treatment Effect (ATE) would, by definition, be the mean of these individual-level causal effects (1)

$$\text{ATE} = \sum_{i=1}^N \frac{Y_i(1) - Y_i(0)}{N} \quad (1)$$

Unfortunately, it is impossible to observe the same individual having been placed in the treatment condition and in the control condition — an individual can only live one existence at any one time.³ Neymand studied the ATE in the context of a simple, completely randomised experiment, i.e., a study in which participants are randomly assigned to either a treatment condition or a control condition. He found that, in this context, difference between the sample mean of those in the treatment condition and the sample mean of those in the control condition was an unbiased estimator of the Average Treatment Effect (2).

$$\mathbb{E}(\bar{Y}_1 - \bar{Y}_0) = \mathbb{E}\left(\sum_{i=1}^N \frac{Y_i(1) - Y_i(0)}{N}\right) \quad (2)$$

That is, in the context of a completely randomised experiment, the difference in control group and treatment group sample means, which can be observed, is an unbiased estimator of the true ATE, which cannot be observed. Neymand also described that the estimate of variance for a difference in sample means, $s_1^2/n_1 + s_0^2/n_0$, is either equal to or greater than the variance of the true ATE. This means that a significant difference between the control and experimental group sample means indicates a significant ATE.

Near the same time, Roland Fisher similarly conceived of a theory of cause and effect that relied upon knowing the outcome of a given individual had he received a treatment and had he not received a treatment. For example, Fisher stated in 1918 (Fisher, 1919): “If we say, ‘This boy has grown tall because he has been well fed,’ we are not merely tracing out the cause and effect in an individual instance; we are suggesting that he might quite probably have been worse fed, and that in this case he would have been shorter”.

³Certain study designs, such as the cross-over study design, expose the same individual to multiple exposures over the study period. These are still not perfect counterfactuals. For example, an individual exposed to treatment A then B may respond differently than an individual exposed to treatment B then treatment A.

Fisher is generally credited as the founder of the randomised experiment ([Rubin, 2005](#)).⁴ However, instead of treating randomisation in a theoretical conceptualisation, as Neyman had done, Fisher viewed randomisation as a way to test hypotheses in empirical work. In his 1926 textbook *Statistical Methods for Research Workers*, Fisher formally proposed the “Fisher sharp null hypothesis” that $Y_i(0) = Y_i(0) \forall i \in 1...N$. Under this formulation, the outcomes for any given individual are known for both the experimental and control condition, which is impossible. Instead a hypothesis test can be applied. If the Fisher sharp null hypothesis were true, $Y(1) - Y(0) = 0$, and a formal hypothesis test can be used to calculate the probability (i.e., p-value) that, given the observed data, y the observed difference between sample means $y(1) - y(0) = 0$.

Importantly, both Neyman’s and Fisher’s models depend upon a condition called the Stable Unit Treatment Value Assumption (SUTVA). Under SUTVA, an each individual’s outcomes are independent; that is, assigning an individual to either the control condition or the treatment condition does not affect anyone else in the experiment. Further, the control and treatment conditions are assumed to be stable, regardless of the number of participants who are assigned to the condition. That is, an individual could expect the same experience regardless of whether many or few participants were assigned to her condition.

4.2 Causal Inference in Randomised Studies

It was not until the early 1970s that [Rubin \(1974\)](#) formally recognised and described the conceptual value of randomly assigning individuals to participate in the study and the conceptual value of randomly assigning participants to the control or treatment conditions. In a completely randomised experiment, an assignment vector $W = (W_1...W_N)$ is randomly generated such that each individual i in the population is assigned to either not participate in the study ($W_i = *$), participate in the study in the control condition ($W_i = 0$),

⁴While Neyman’s 1923 thesis used the concept of completely randomised experiments in his theory before Ronald Fisher proposed them in 1925 ([Fisher, 1925](#)). Neyman, himself, insisted that his theory relied on randomisation simply as a means to treat the outcomes probabilistically while Fisher demonstrated the necessity of randomisation in experiments beyond statistical implications ([Reid, 1998](#)).

or participate in the study in the treatment condition ($W_i = 1$). It is not necessary that the probability of assignment be the same (i.e., $P(W_i = *) = P(W_i = 0) = P(W_i = 1)$), so long as the probabilities are uniform for all individuals, i.e., $P(W_i = A) = P(W_j = A) \forall i, j \in 1 \dots N$ and $\forall A \in *, 0, 1$. In a randomised experiment, W exhibits two key features. The first, called “ignorability”, is that the two potential outcomes for each participant (one observed and one unobserved) is independent of the randomisation (Rubin, 1974). The second is that each individual could be assigned to either the control or the treatment condition, i.e., $0 < P(W_i | X, Y(), Y(1)) < 1$.

In this case, the vector W is an entirely exogenous random variable; that is, no other factors influence the values that it takes on. Consequently, even accounting for a vector (or matrix) of external factors X , the vector of (observed or unobserved) outcome values given the treatment $Y(1)$ and vector of (observed or unobserved) potential outcome values given the control $Y(0)$, an unbiased estimator for the ATE can be computed in a completely randomised trial using Neyman’s equality (2). This can be written formulaically as:

$$\mathbb{E}(\bar{y}_1 - \bar{y}_0 | X, Y(1), Y(0)) = \mathbb{E}(\bar{y}_1 - \bar{y}_0) = \mathbb{E}\left(\sum_{i=1}^N \frac{Y_i(1) - Y_i(0)}{N}\right) \quad (3)$$

$$V(\bar{y}_1 - \bar{y}_0 | X, Y(1), Y(0)) = V(\bar{y}_1 - \bar{y}_0) \geq \mathbb{E}(s_1^2/n_1 + s_0^2/n_0) \quad (4)$$

In sum, because the assignment vector W is randomly generated depending upon no other factors, it is trivial to compute an unbiased estimator of the causal effect and a positively biased estimator of the variance. Using these estimators, it is further possible to test the hypothesis that the effect is greater than a hypothesized value H (often $H = 0$).

4.3 Causal Inference in Non-Randomised Studies

By recognising the value of randomisation in experiments, Rubin (1974) was able to extend the Potential Outcomes Framework to account for non-randomised studies.

In non-randomised studies it is possible that the assignment vector W depends upon other factors. For example, a non-randomised study may show that people who drive sports cars are healthier than those who drive sedans. While there may be an effect of the car on health, it is also possible that an underlying third factor connects the treatment (i.e., car) with the outcome (i.e., health). Those who own sports cars may be wealthier than those who drive sedans, and wealth could be the causal driver of a difference in health. In terms of the equations above, W is not independent with respect to X and so $\mathbb{E}(\bar{y}_1 - \bar{y}_0 | X, Y(0), Y(1)) \neq \mathbb{E}(\bar{y}_1 - \bar{y}_0)$. This source of bias is commonly referred to as “confounding”. As another example, individuals might self-select into the treatment condition based upon their demographics, their circumstances, or their need for treatment. For example, it is known from randomised controlled trials that paracetamol reduces headaches. However, in an observational study of people using paracetamol and people not using paracetamol, it is likely that those using paracetamol are more likely to have a headache than those who are not. In this case, W is not independent with respect to $Y(0)$ and $Y(1)$ and so, again, $\mathbb{E}(\bar{y}_1 - \bar{y}_0 | X, Y(0), Y(1)) \neq \mathbb{E}(\bar{y}_1 - \bar{y}_0)$. This source of bias is commonly referred to as “self-selection”. Further, in studies without random assignment, it is possible that the decision of whether to participate depends upon other factors. Researchers will only observe data from those who elected to participate, which could be unrepresentative of the population, and it is possible that the treatment may have a different effect if imposed upon those who were not included in the study. This may be referred to as “sample-selection” bias.

It is impossible to compute a truly assumption-free causal estimate without true random assignment. However, several methods have been designed to formally adjust for and minimise the restrictions levied by these assumptions.

4.3.1 Regression Modelling

Perhaps the most commonly used modern statistical technique is regression. While not inherently a causal inference method, regression can be used to mitigate confounding bias

in non-randomised studies and forms the foundation for several significant causal inference methods ([Zohoori and Savitz, 1997](#)).

The simplest form of regression, used for continuous outcome data, is Ordinary Least Squares Regression (OLS). Other regression models exist to account for different distributions of the data (e.g., logistic or probit regression can be used for binary outcome data, Poisson or negative binomial regression can be used for count outcome data, beta regression can be used for outcome variables lie between 0 and 1), but the use of regression to account for confounding is similar. The model underlying OLS is shown as:

$$Y = X\beta + \epsilon \quad (5)$$

In this model, Y is a $n \times 1$ observable matrix of all outcome values, X is a $n \times k$ observable matrix constructed by binding columns of k independent variables together, and ϵ is a $n \times 1$ matrix of all errors. β represents the $k \times 1$ matrix of non-random but unobserved parameters that minimise the function $\sum_{i=1}^n \epsilon^2$. When expectation of the errors are 0, that errors have common variance (homoskedastic), and that errors are uncorrelated, it can be shown that the OLS estimators $\hat{\beta} = (X'X)^{-1}X'y$ are also lowest-variance unbiased estimators of the true coefficients β ([Puntanen and Styan, 1989](#)).

In a completely randomised experiment with Y as the outcome and X as an indicator variable equal to 1 if the individual was in the treatment condition and 0 if the individual was in the control condition, these assumptions on the error terms would likely be met. The distribution of covariates would be identical between those in the control and treatment conditions, and so any other factor that has an influence on the outcome would be evenly distributed between the control and treatment conditions.

However, in non-randomised studies, it is likely that the distribution of covariates would be unequal between groups, and the parameter estimate may be biased by confounders disproportionate to certain groups. Fortunately, if, *conditional on all other independent variables*, the expectation of the error term is 0, the errors have common variance, and the

errors are uncorrelated, the OLS estimators will still be the lowest-variance and unbiased estimators. Consequently, it is possible to reduce confounding bias by including variables for potential confounders in the regression model.

A crucial assumption for this method is that all potential confounders are included in the model ([Crémieux and Ouellette, 2001](#)). It is relatively simple to include observed confounders when computing the model; however, unobserved confounders can also bias the OLS estimators and, as the researcher cannot observe them, cannot be accounted for in regression. This form of bias is commonly referred to as “omitted variable bias”; without random assignment, it is impossible to entirely eliminate the possibility of omitted variable bias.

4.3.2 Multi-Level Models

While it is impossible to entirely eliminate the possibility of omitted variable bias, bias presented by certain types of unobserved confounders can be excluded when data is nested. In many empirical settings, observations are naturally grouped together ([Duncan, Jones and Moon, 1996](#)). For example, in longitudinal data, individuals may complete a survey at several different time points. Each individual-by-time observation then fits within a group of same-individual observations (i.e., all the observations from that same individual over all time points) and within a group of same-time observations (i.e., all observations from that same time point over all individuals). For another example, individual-level panel data may be grouped by jurisdiction (e.g., state) and by time, even if the individuals contributing data are different from year to year. In these settings with nested data, fixed and random effects can be used to account for confounding that is invariant within groups, even when the confounders are unknown ([Leyland and Groenewegen, 2003](#)).

The most basic tool for multi-level models is fixed effects. Consider, for example, a study aiming to assess whether a procedure changes length of hospital stay using de-identified, patient-level data on procedures from several hospitals. Observations are nested within

hospitals, such that each hospital contains some observations where the procedure was performed and some observations where the procedure was not performed. A naive regression would simply compare the length of hospital stay among those who had the procedure versus those who did not (equation (6)).

$$y_i = \beta_0 + \beta_1 * (\text{Procedure} = 1) + \epsilon_i \quad (6)$$

However, there are likely several variables that confound the relationship between the exposure (receiving the procedure) and the outcome (length of hospital stay). Some of these confounders may vary at the individual level; for example, if increased severity of the condition both causes the physician to use the procedure and leads to increased length of hospital stay after the procedure, then severity of the condition would be a confounder. If patient-level data on condition severity is available, it should, consequently, be added to the regression model. If data is not available for condition severity, the regression cannot be adequately adjusted, which could lead to omitted variable bias. While all confounders have at least some variance at the individual-level, some confounders may vary predominantly at the hospital level. For example, it is possible that some hospitals serve wealthier populations than others, which could influence both the exposure and the outcome. Even if individual-level data is unavailable for the wealth of the patient, accounting for the hospital through a hospital fixed effect could account for the hospital-level variation in wealth of patients. Further, a hospital fixed effect accounts for hospital-level variation in all potential confounders, including both observed and unobserved confounding. That is, if the researchers fail to consider an important confounder but that confounder varies primarily at the hospital level, then a hospital fixed effect may, at least partially, reduce the bias from that variable. Conceptually, the fixed effect groups together observations within the same hospital and estimates the effect of the exposure on the outcome within each hospital. The final effect estimate, then, is the weighted average of the effect of the procedure within each hospital.

A simple fixed effect, like that described in the previous example, is implemented by

setting a separate baseline length of stay (i.e., mean length of stay for participants who did not receive the first procedure) for each group j :

$$y_{ij} = \beta_1 * (\text{Procedure} = 1) + \alpha_j + \epsilon_{ij} \quad (7)$$

A common technique using fixed effects in regression is intercept fixed effects ([Bell, Fairbrother and Jones, 2019](#)). Intercept fixed effects essentially set a different intercept for the outcome variable within each group. (Equivalently and often more efficiently, fixed effects can be implemented by subtracting the group-level mean from each individual in each group). In this way, the fixed effects account for baseline differences among the groups of data. These fixed effects also account for any confounder that does not vary with the independent variable within the group. That is, if a given confounder is invariant within groups (e.g., sex of a participant or location of a hospital), then that effect on the outcome will be present at the intercept. Consequently, using a separate intercept for each group will eliminate the effect of that confounder. Intercept fixed effect correspond graphically to parallel, group-level lines on either side of the weighted average intercept model (Figure 1, Panel B). The slope of these lines is a pooled estimate, and so it is the same among all groups. The intercept, however, is individually determined, and so the lines may be parallel.

A common implementation of intercept fixed effects uses panel data, where data is collected for each individual i within several groups j over several years t , and a separate intercept is set for each group and year:

$$y_{ijt} = \beta_1 * \text{Exposure} + \alpha_j + \lambda_t + \epsilon_{ijt} \quad (8)$$

In this model, a separate intercept is set for each group, and each observation's outcome is measured in relation to the mean outcome within that time. The group fixed effect accounts for time-invariant confounding at the group level, while the time fixed effect accounts for group-invariant confounding at the year level. Notably, group-level fixed effects

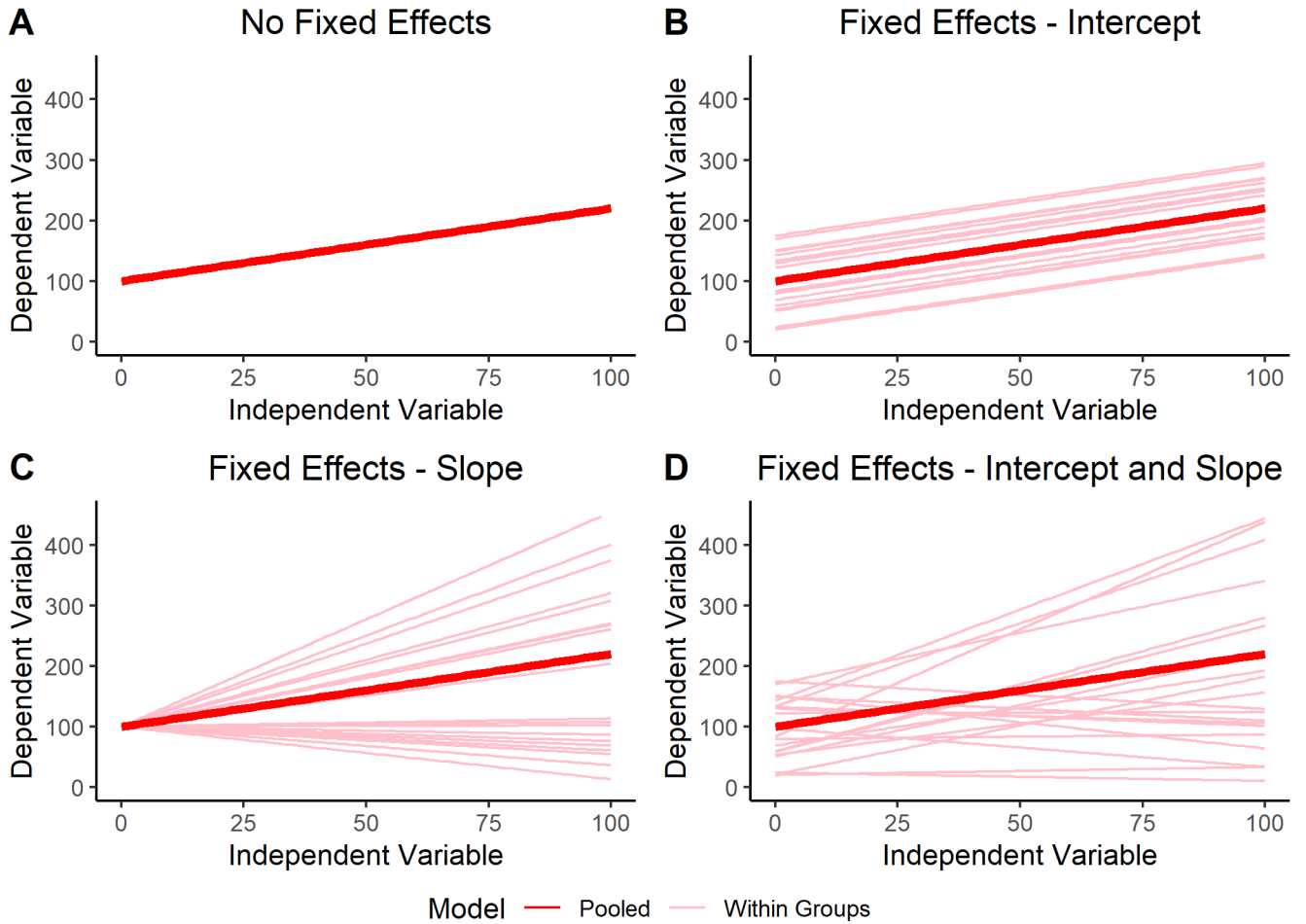
cannot account for confounders that vary within the group; this model does not account for confounder variation for individuals within groups or for confounders that change over time within groups.

While intercept fixed effects are perhaps the most common, the concept of fixed effects can be applied to any parameter in a regression model, as long as the model has sufficient degrees of freedom (i.e., the model matrix is not over-determined). For example, slope fixed effects assume that the relationship between an independent variable and the outcome is constant within groups. Slope fixed effects can be implemented by interacting the relevant independent variable with the group-level fixed effects. When using fixed slopes, the intercept is a pooled estimate that is constant across all groups, while the slopes are estimated independently. Consequently, a fixed slopes model corresponds graphically to a single origin point with several group-level lines that are steeper or less steep than the weighted average slope model (Figure 1, Panel C).

Further, in cases where there is sufficient data (and degrees of freedom), fixed effects can be used for both the intercept and the slope within the same model. Fixed intercept effects allow for the intercept to vary between groups, and fixed slope effects allow for the slope to vary between groups. Using both fixed intercept and fixed slope effects, both the intercept and the slope can vary between groups. Graphically, this corresponds to several non-parallel lines not emerging from a single point (Figure 1, Panel D). Depending upon the between-group variance in the intercept and the slope, these lines can be similar to or different from the model for the weighted average slope and intercept model.

An important feature of fixed effects is that they carry the assumption that the groups are truly independent with respect to the parameter being estimated. On one hand, this can be a valuable feature; a researcher does not need to assume a certain group-level distribution of the underlying parameter when using fixed effects. Because fixed effects are estimated entirely within groups, they can appropriately account for group-invariant confounding even if the impact of group on the parameter is highly heterogeneous. On the other hand, this

Figure 1: Illustration of Fixed Effects



NOTE: These figures were generated from simulations by the author. Group intercepts take on values randomly generated between 20 and 180, and group slopes take on values between -1.2 and 3.6.

feature limits the interpretability of any parameter estimated through fixed effects. Because parameters with fixed effects are estimated entirely within groups, the parameter estimate is only meaningful within that group and cannot be extrapolated to other groups. For this reason, fixed effects are typically used to rule out confounding by lurking variables. If fixed effects are applied to the independent variable of interest, the parameter estimates for that variable will be uninterpretable.

One major disadvantage of fixed effects is their inefficiency. Parameter estimates become more imprecise with decreasing degrees of freedom, and including a fixed effect reduces a model's degrees of freedom by the number of levels it takes on. In this respect, the related concept of random effects provide a distinct advantage. Random effects account for between-group variation in confounders with a much lesser impact on the model's efficiency. In exchange, however, random effects carry some assumptions on the group-level distribution of the parameter.

Like with fixed effects, the purpose of random effects is to adjust for observed and unobserved group-level variation in naturally nested data. Unlike fixed effects, however, the random effects method assumes that the group-level parameter estimates come from a known, typically normal distribution ([Puntanen and Styan, 1989](#)). In this way, the group-level parameter estimates are not strictly independent and can instead be modelled as a random variable. Consequently, these group-level parameters are estimated through “partial pooling”, accounting to some extent for the distribution of the other group-level parameters, and group-level parameter estimates that vary too far from the distribution are pulled towards the mean.

A simple random effects model is shown in equation (9). In this model, y represents a vector of $N \times 1$ matrix containing the values of the dependent variable, Z represents a $N \times q$ “design matrix” where q is the number of groups in the data, u is a $q \times 1$ matrix of random effects, and ϵ is a $N \times 1$ matrix of the errors.

$$y = \beta_0 + Zu + \epsilon \quad (9)$$

The design matrix Z is constructed such that each row represents an individual observation and each column represents a group. When observation i is in group j , $Z_{ij} = 1$. Otherwise, $Z_{ij} = 0$. This is a familiar construction; Z is simply a collection of dummy variables for whether the observation fits within a given group. Indeed, if the parameters in u were estimated through OLS, then u would be a vector of group-level fixed effects. However, as previously noted, random effects impose some assumptions on the estimators, distinguishing them from fixed effects. Instead of directly estimating u , it is assumed that $u \sim N(0, G)$. u is distributed around 0 because the intercept β_0 absorbs the pooled mean for y , and G is the variance-covariance matrix of u . In this case, the intercept is the only parameter estimated through random effects and u is a single-variable, $q \times 1$ matrix, so the variance-covariance matrix of u is simply the variance of the random intercept. In cases where more than one parameter is estimated by the random effects model, it may be necessary to place additional constraints on the covariance structure among random effects (e.g., independence) to efficiently estimate G .

In more concrete terms, a random effects model can be illustrated in a regression model, as in equation (10).

$$y_{ij} = (\gamma_{00} + u_{0j}) + \epsilon_{ij} \quad (10)$$

In this regression, y_{ij} represents the outcome for observation i in group j , γ_{00} represents a common intercept for all observations, u_{0j} is a random variable centered at 0 with variance G that takes on a different value for each group j , and ϵ_{ij} is the random error term. Collectively, $(\gamma_{00} + u_{0j})$ is the group-level intercept, centred at γ_{00} .

Like fixed effects, the random effects method can be applied to any parameter in a regression – not just the intercept. Consequently, the graphic representation in Figure 1

is roughly applicable to random effects. When random effects are applied to the intercept, group-level estimates will parallel the weighted average model. When random effects are applied to the slope, group-level estimates will emerge from the same intercept but have span outwards from the weighted average model. When random effects are applied to both the intercept and the slope, group-level estimates need not have anything in common with the weighted average model.

In sum, in cases where it is important to maximise the efficiency of a model and it can reasonably be assumed that group-level parameters are drawn from a random distribution, random effects are a suitable solution to addressing between-group confounding.

In many cases, it is appropriate to use both fixed and random effects. These models – referred to as mixed models – apply fixed effects and random effects methods to different parameters in the regression. For example, a regression could use intercept fixed effects and slope random effects. Mixed methods are particularly useful when reasonable assumptions can be made on the distribution of some group-level parameters but not others. A simple mixed model can be written using equation (11).

$$y = X\beta + Zu + \epsilon \quad (11)$$

In this model, y is the $N \times 1$ matrix of outcomes, X is a $N \times p$ matrix, where p is the number of fixed effects to be included in the model, β is a $p \times 1$ matrix of fixed effects parameter estimates, Z is a $N \times q$ matrix where q is the number of groups used for random effects, u is a $q \times 1$ matrix for the random effects estimates, and ϵ is a $N \times 1$ matrix of errors.

4.3.3 Difference-in-Difference

One of the most common methods for causal inference is the Difference-in-Difference design (Lechner, 2010). Recall that in a completely controlled experiment, it is appropriate to merely compare the sample mean outcome value between those who received the treatment $Y(1)$ and those who did not receive the treatment $Y(0)$ to estimate a causal effect. The

Table 9: Two Groups Identical at Baseline

	Outcome_{t=1}	Outcome_{t=2}
Control $Y(0)$	10	20
Treatment $Y(1)$	10	40

NOTE: This table was created by the author. Notice that each group have the same outcome value at $t = 1$. This is important, as it was assumed that the two groups were highly alike.

Difference-in-Difference design creates a parallel to this reasoning for non-randomised studies.

In its simplest possible form, a researcher wishing to use Difference-in-Difference attempts to find two samples of individuals that are highly alike on all observable characteristics (including the outcome of interest) at baseline (i.e., before the intervention), with one sample experiencing an intervention and the other not experiencing the intervention. Under the assumption that the two groups were highly similar at baseline and extrapolating that the two groups would have experienced the same trajectory in the outcome variable had neither group experienced the intervention, the difference between the two groups after the intervention could be considered an estimate of the causal effect of the intervention. As an example (Table 9), the difference between the dependent variable in the treatment group and the control group ($40 - 20 = 20$) could be considered a causal effect of the treatment.

The assumption that the two groups were virtually identical at baseline is very strong. In virtually all non-randomised settings, there will be some baseline differences between the treatment and control samples. The Difference-in-Difference approach accounts for these issues using two differencing methods to construct a counterfactual. The first difference occurs between the outcome values for the treatment group at baseline and the control group at baseline, accounting for differences at baseline. In this way, the control group becomes a more realistic, comparable counterfactual for the treatment group. The second difference occurs between the outcome values for the treatment group after the intervention and before the intervention. This difference allows the treatment group's pre-intervention outcome values to serve as a counterfactual for the treatment group's post-intervention outcome values.

Table 10: Simple Difference in Difference

	Outcome_{t=1}	Outcome_{t=2}	Changes
Control $Y(0)$	$y_{21} = 30$	$y_{22} = 25$	$y_{21} - y_{22} = 5$
Treatment $Y(1)$	$y_{11} = 60$	$y_{12} = 40$	$y_{11} - y_{12} = 20$
Difference	$y_{11} - y_{21} = 30$	$y_{12} - y_{22} = 15$	$(y_{12} - y_{22}) - (y_{11} - y_{21}) = -15$

NOTE: This table was created by the author.

In sum, the Difference-in-Difference approach estimates the causal effect of the treatment as shown in equation (12), where y_{11} is the sample mean outcome for the treatment group in period 1, y_{21} is the sample mean outcome for the control group in period 1, y_{12} is the sample mean outcome for the treatment group in period 2, and y_{22} is the sample mean outcome for the control group in period 2.

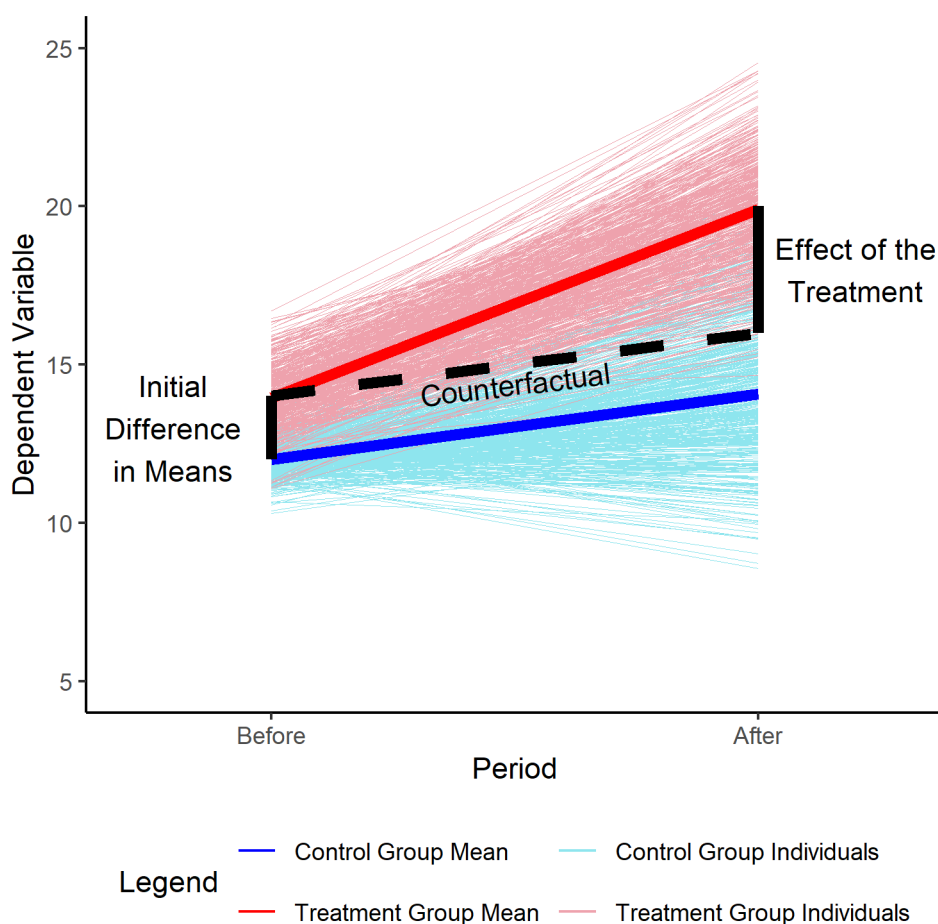
$$DD = (y_{12} - y_{22}) - (y_{11} - y_{21}) \quad (12)$$

Table 10 shows an example of how the Difference-in-Difference estimator is calculated. The Difference-in-Difference estimator from this example is $(y_{12} - y_{22}) - (y_{11} - y_{21}) = -15$.

The causal effect from the Difference-in-Difference method can be shown graphically (Figure 2). From this representation, it is shown that the trajectory of the control group combined with the baseline difference between the treatment and control group create a counterfactual. The causal effect is illustrated by the black, vertical bar on the right of the figure, which highlights the sample mean for the treatment group in excess of what would be predicted by the counterfactual. This surplus of the dependent variable can be considered the causal effect of the treatment.

While this simple Difference-in-Difference model accounts for baseline sample mean differences in the outcome variable between the treatment and control group, it is still likely that the treatment and control group vary on other observable confounders. The Difference-in-Difference model can be extended to further account for confounding using regression, as

Figure 2: Illustration of Difference-in-Difference Design



NOTE: These figures were generated from simulations by the author. The control group had a mean of 12 before the intervention and an increase of 2 ($SD=2$) by period 2. The treatment group had a mean of 14 before the intervention and an increase of 6 ($SD=1.5$) by period 2. The causal effect can be calculated as $DD = (20 - 14) - (14 - 12) = 4$, which is approximately the height of the bar on the right side of the figure.

shown in equation (13).⁵

$$y_{it} = \beta_0 + \beta_1 * \text{Period} + \beta_2 * \text{Condition} + \beta_3 * \text{Period} * \text{Condition} + X_{it}\beta + \epsilon_{it} \quad (13)$$

In this model, y_{it} refers to the outcome value for individual i during period t , period is an indicator variable equal to 0 before the intervention and 1 after the intervention, condition is an indicator variable equal to 0 for the control group and 1 for the treatment group, and X_{it} is a $N \times p$ matrix constructed by binding p confounder variables by column. β_0 is the parameter for the baseline outcome value among the control group before the intervention. β_2 is the parameter for the difference between the baseline outcome value for the treatment and the control groups. β_3 is the parameter of interest, corresponding with the Difference-in-Difference estimator. β is a column vector corresponding to parameter estimates for the p confounders, and ϵ_{it} is the random error term. When the same individuals are sampled in each group before and after the intervention, the regression model can also account for individual-level fixed intercept effects (see Section 4.3.2), which essentially transforms the outcome to be the individual-level difference in the outcome before and after the intervention.

The underlying assumption for all Difference-in-Difference models is that, absent the intervention, the treatment group would have had the same trajectory as the control group in terms of the outcome variable. That is, there is no selection between the control and treatment group such that they would experience a different change in the outcome even if both were not exposed to the treatment. This is a strong assumption. Choosing samples that are similar at baseline and using regression to adjust for observable confounders somewhat mitigates this concern, though it is still possible that unobserved confounders could bias the Difference-in-Difference estimator. Another assumption of the Difference-in-Difference model is that there are no spillovers. When estimating a Difference-in-Difference model, the

⁵For an example of how this method can be used in practice alongside propensity score matching, see (Erlangga, Ali and Bloor, 2019)

econometrician assumes that only those in the treatment group are affected by the treatment and that the treatment group's exposure does not affect the environment faced by the control group.

4.3.4 Interrupted Time Series Analysis

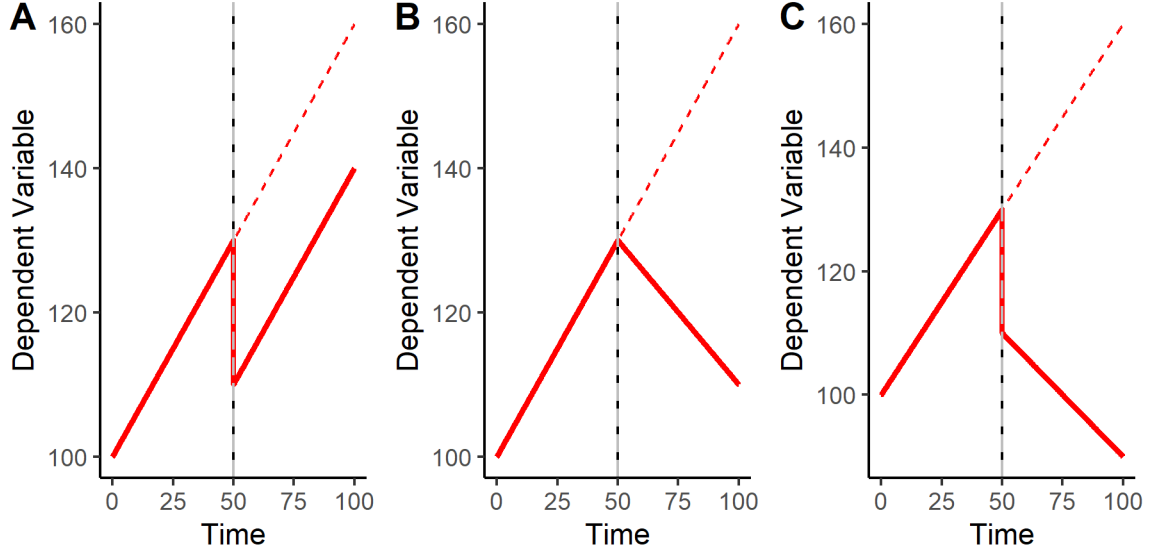
Like Difference-in-Difference, Interrupted Time Series Analysis (ITSA) uses differencing to construct a counterfactual. There are two main forms of ITSA. One-sample ITSA uses the trajectory of the outcome within a group *before* an intervention to construct a counterfactual, against which post-intervention outcome values can be compared. This is similar to the first difference in the Difference-in-Difference approach ($y_{11} - y_{21}$). Two-sample ITSA uses both the trajectory of the treatment group before the intervention and the trajectory of an untreated control group to construct a counterfactual, corresponding with the full Difference-in-Difference model ($(y_{12} - y_{22}) - (y_{11} - y_{21})$). Both differ from Difference-in-Difference, however, in their treatment of time. In ITSA, the outcome variable is observed at several time periods before the intervention and at least one period after the intervention ([Craig et al., 2012](#)). Consequently, it is possible to observe not just the level of the outcome variable in each group before the intervention but also the trends over time before the interruption.

Because ITSA accounts for several time periods, it is almost always computed using regression. The regression model for a one-sample ITSA is shown in equation (14) ([Lopez Bernal, Cummins and Gasparrini, 2016](#)).

$$y_{it} = \beta_0 + \beta_1 * \text{Time} + \beta_2 * \text{Post} + \beta_3 * \text{Time} * \text{Post} + \epsilon_{it} \quad (14)$$

In this model, y_{it} is the outcome for individual i in time t , Time is a linear time term (e.g., 0 in period 1, 1 in period 2, etc.), and Post is an indicator variable if the observation is after the intervention. β_0 is the parameter for the intercept (the average outcome at period 1) and β_1 is the parameter for the slope on time before the interruption. Notably, as opposed to the Difference-in-Difference model which only reports mean differences, the interrupted

Figure 3: Illustration of One Sample Interrupted Time Series



NOTE: These figures were generated from simulations by the author. The dotted red line in each panel shows the counterfactual, while the solid red line presents observed values for the sample studied. The interruption is set at time 50, noted with a dashed gray line. The total, observation-level effect of the intervention over the entire study period is represented by the space between the counterfactual (dashed red line) and the observed time series (solid red line) after the intervention. The initial intercept is set at 100, and the initial slope is set at 0.6. The level change is set to 20, and the post-intervention slope is set to -0.4 .

time series analysis model decomposes the effect into an immediate effect and a gradual effect: β_3 corresponds to the change in level after the intervention or the immediate effect of the intervention. β_4 corresponds to the change in slope after the interruption or the gradual effect of the intervention. ϵ_{it} refers to the random errors in the model.

The effects represented by β_2 and β_3 can be visualised graphically. Panel A of Figure 3 shows an immediate effect (β_2) on the dependent variable following the intervention. Panel B shows a change in slope at the point of the interruption, representing a gradual effect on the dependent variable (β_3). Panel C shows both these effects in the same time series.

The key identifying assumption of the one-sample ITSA is that, in the absence of the intervention, the trajectory would have continued with the initial slope and intercept (Penfold and Zhang, 2013). In cases where the trend before the intervention – sometimes called the pre-trend – is consistent, precise, and linear and there are no other confounding events occurring near the time of the intervention, this may be a reasonable assumption. However,

if the pre-trend is not consistent or linear or if there are other confounding events occurring at the time of the intervention, it can be difficult to convincingly argue that this assumption would hold, and causal estimates from the one-sample ITSA may be biased.

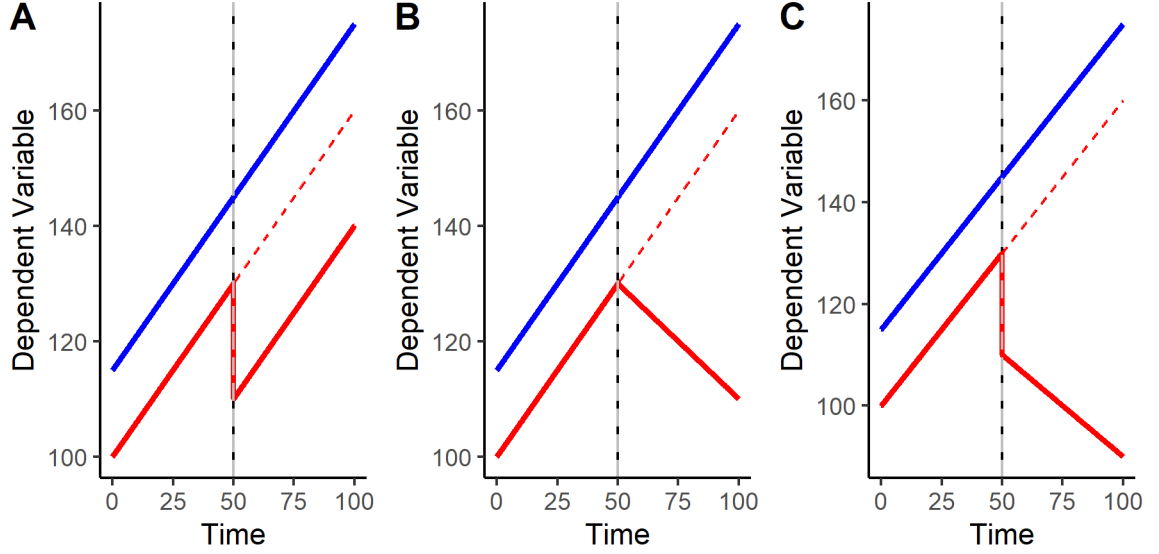
The two-sample ITSA provides the researcher with an opportunity to empirically argue this assumption (Lopez Bernal, Cummins and Gasparrini, 2018). Like the Difference-in-Difference model, the two-sample ITSA uses a control group that was not affected by the intervention, as well as the treatment group pre-trend, to construct a counterfactual. Consequently, the control group provides the researcher with an opportunity to demonstrate that, absent the intervention, the trajectory would have continued from the pre-trend. That is, if the researcher can show that the control group and the treatment group have parallel pre-trends and that, after the intervention, the control group's trajectory follows the pre-trend, it is often reasonable to assume that, absent the intervention, the treatment group also would have followed the same pre-trend. The assumption, then, is that the control group and the treatment group would have had parallel trends throughout the study period in the absence of the intervention. This is commonly called the “parallel trends assumption”. The design of the two-sample ITSA is shown in Figure 4.

The regression for two-sample ITSA follows the same basic form as that for one-sample ITSA (14) but with each parameter interacted with condition, as shown in equation (15).

$$\begin{aligned}
y_{it} = & \beta_0 + \beta_1 * \text{Time} + \beta_2 * \text{Post} + \beta_3 * \text{Condition} \\
& + \beta_4 * \text{Time} * \text{Post} + \beta_5 * \text{Time} * \text{Condition} \\
& + \beta_6 * \text{Post} * \text{Condition} + \beta_7 * \text{Time} * \text{Post} * \text{Condition} \\
& + \epsilon_{it}
\end{aligned} \tag{15}$$

In this model, y_{it} is the outcome for individual i at time t , Time and Post are defined as in equation (14), and Condition is an indicator variable equal to 0 if the individual is in

Figure 4: Illustration of Two Sample Interrupted Time Series



NOTE: These figures were generated from simulations by the author. The dotted red line in each panel shows the counterfactual, while the solid red line presents observed values for the treatment group and the solid blue line represents observed values for the control group. The interruption is set at time 50, noted with a dashed gray line. The total, observation-level effect of the intervention over the entire study period is represented by the space between the counterfactual (dashed red line) and the observed time series (solid red line) after the intervention. For the treatment group, The initial intercept is set at 100, and the initial slope is set at 0.6. The level change is set to 20, and the post-intervention slope is set to -0.4 .

the control group and 1 if the individual is in the treatment group. β_0 , β_1 , β_2 , and β_4 are parameter estimates corresponding to the intercept, initial slope with time, post-intervention level change, and post-intervention slope change for the control group. β_3 , β_5 , β_6 , and β_7 are parameter estimates corresponding to the difference in intercept, difference in initial slope with time, difference in post-intervention level change, and difference in post-intervention slope change for the treatment group, relative to the control group. The parameters of interest are β_6 , which corresponds to the immediate effect of the intervention, and β_7 , which corresponds to the gradual effect of the intervention. Parallel pre-trends can be empirically tested through the estimate for β_5 ; if the estimate for β_5 is significantly different from 0 (usually with an $\alpha = 0.10$), then the pre-trends are not parallel between the control and treatment group. The estimate for β_3 , while less important than parallel pre-trends, is often also reported to demonstrate that the control and treatment group had significantly different values of the outcome at baseline.

While the two-sample ITSA is a stronger design than the one-sample ITSA owing to the use of a contemporaneous control group, it still has several limitations. While strongly parallel pre-trends signals (and is a prerequisite of) the parallel trends assumption, it is not conclusive, and the researcher is forced to make some assumptions about what the trend of the treatment group would have been if those individuals had not experienced the intervention. Some of these concerns may be addressed by adding confounders into the regression model, though the possibility of unobserved confounding and omitted variable bias remains. For example, it remains possible that individuals selected into the treatment or control groups based upon some unobserved covariate(s), and without controlling for that selection, causal effect estimates will be biased. Further, both one-sample and two-sample ITSA assume that there were no other relevant interventions occurring in either of the two groups at around the same time as the intervention. If two relevant interventions occur at nearly the same time, it is nearly impossible to disentangle their effects through an ITSA. Relatedly, both assume that spillovers are not present between the groups; that is, only those individuals in the treatment group are affected by the intervention, and individuals in the control group are not affected by spillovers from the treatment group.

Another limitation of this approach is its high dimensionality. Even with a large dataset, dimensionality can be a concern if the panel is short because we include fixed effects for each individual and time period, which use up many of the data's degrees of freedom and increase standard errors. This concern limits the number of confounders we can adjust for in $X\beta$.

An alternative method to weakening the assumptions imposed by the one-sample ITSA is through the use of fixed and random effects. ITSA often uses data from individuals collected over the course of the study period (i.e., before and after the study period). In this case, the individual-by-time observations are naturally grouped by individual, and individual-level fixed or random effects can be applied when estimating the ITSA parameters. As discussed in Section 4.3.2, fixed effects or random effects can be applied to any of the parameters in the regression (though it does not make sense to apply fixed effects or random effects to

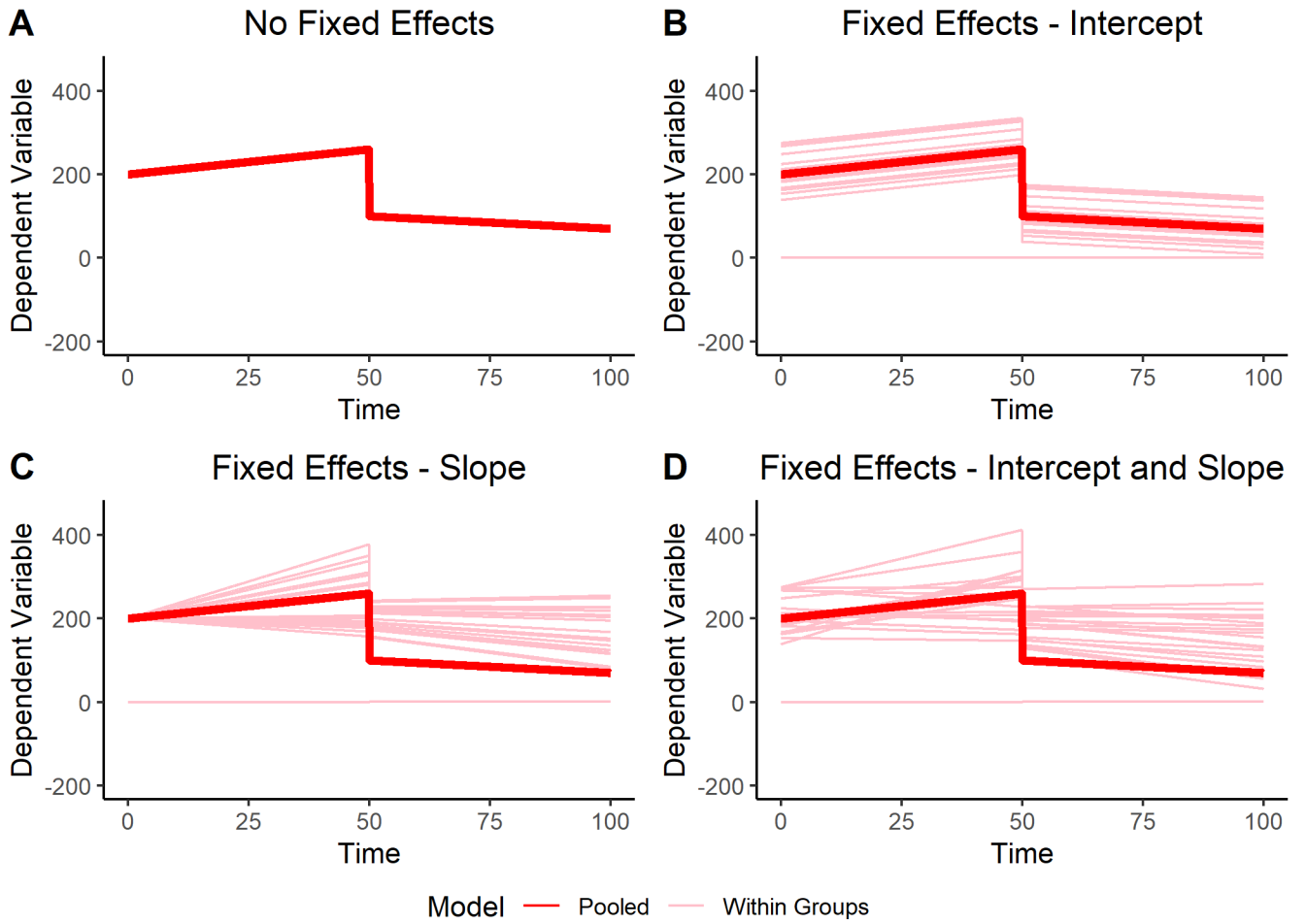
the parameters of interest). Applying fixed effects or random effects to different parameters weakens some underlying assumptions of ITSA and allows the model to more precisely estimate the effect of the intervention on a given individual. For example, adding intercept fixed effects to the ITSA regression removes the assumption that all individuals experience the same immediate effect of the treatment and allows the pooled ITSA immediate effect estimate to account for within-individual immediate effects of the intervention. Fixed effects in ITSA is shown graphically in Figure 5, Panel B. Fixed slope effects, shown graphically in Panel C, remove the assumption that all individuals experience the same change in slope before and after the treatment, allowing the ITSA gradual effect estimate to account for within-individual changes in trajectory before and after the treatment. Fixed intercept and slope effects, as shown in Panel D, remove both these assumptions. Using random intercept and slope effects have largely the same goals as using fixed intercept and slope effects, though with better efficiency and additional assumptions on the underlying distribution of the parameter estimates (see Section 4.3.2). For an example using fixed and random effects in a two-sample interrupted time series analysis, see ([Castro-Avila, Bloor and Thompson, 2019](#)).

5 The Association Between General Practitioner (GP) Entry with Healthcare Quality in England’s National Health Service (NHS): Evidence from General Practitioner Movers

5.1 Introduction

As discussed in Section 2, improving quality in general practice – interventions to improve patient outcomes, system performance, and professional development – has been a major

Figure 5: Illustration of Interrupted Time Series Analysis with Fixed Effects



NOTE: These figures were generated from simulations by the author. Group intercepts take on values randomly generated between 120 and 280. Initial slopes take on values between -1.2 and 3.6 . Post-interruption intercepts take on values between 180 and 280. Post-interruption slopes take on values between -1.5 and 0.3 .

priority for the NHS for the past several decades ([Batalden and Davidoff, 2007](#)). Advocates, academics, and policy makers have designed and tested the effectiveness of several quality improvement interventions, including financial incentives, standard-setting, reminders, inspections, public reporting, educational programs, conferences, and feedback mechanisms ([Oxman et al., 1995](#)), and while many of these interventions have garnered evidence of effectiveness, experts agree that meaningfully improving quality in healthcare is still an elusive task ([Braithwaite, 2018](#)). It is reasonable, then, to consider new, innovative methods for healthcare quality improvement. One potentially understudied policy lever for quality improvement, based upon the behaviour change literature ([Rubenstein et al., 2000](#); [Mittman, Tonesk and Jacobson, 1992](#)), is labor supply moves, i.e., using the movement of physicians around the healthcare system to improve the dissemination and uptake of best practices. A physician entering a new practice could influence the behavior of other physicians at that practice, either through sharing new information on best practices or by exerting social influence (e.g., physicians may wish to match the practice styles of their proximate peers ([de Jong et al., 2006](#))). Past researchers have considered these “GP movers” in the context of improving healthcare capacity in under-served (particularly rural) areas ([Dolea, Stormont and Braichet, 2010](#); [Li et al., 2014](#); [Yong et al., 2018](#)), but virtually no studies have considered physician movement in terms of affecting quality, as measured by QOF points. This chapter uses publicly available data on how GPs have moved from practice to practice to provide preliminary evidence that GP movement may be a meaningful determinant of general practices’ ability to maximize QOF points. While this observational study cannot provide causal evidence of an effect, this feasibility test and its results may encourage investment in future interventional studies to test the effect of moving GPs on QOF points and, by extension, healthcare quality.

5.1.1 Peer Effects and GP Movers

Over several decades, psychologists, sociologists, economists, and marketers have identified peer effects – the social influence that individuals place on their proximate peers – as one of the strongest forces for behavior change. For example, studies have demonstrated that peer effects influence adolescent’s health behaviors ([Gaviria and Raphael, 2001](#); [Lundborg, 2006](#); [Ali and Dwyer, 2011](#)), college students’ academic achievement ([Sacerdote, 2001](#)), the diffusion of technology ([Bollinger and Gillingham, 2012](#)), paternity leave participation ([Dahl, Løken and Mogstad, 2014](#)), consumer decisions ([Moretti, 2011](#)), financial decisions ([Bursztyn et al., 2014](#)), and labor productivity ([Herbst and Mas, 2015](#)). It is unsurprising, then, that many health services researchers have endorsed peer effects as a possible method for improving healthcare quality ([Goodpastor and Montoya, 1996](#)). These scholars theorised that social influence, operationalised in such forms as role models ([Kenny, Mann and MacLeod, 2003](#)), opinion leaders ([Locock et al., 2001](#)), socially central individuals ([Meltzer et al., 2010](#)), and peer feedback ([Pronovost and Hudson, 2012](#)), could be a powerful mechanism to encourage physicians to implement evidence-based practices ([Phelps and Mooney, 1993](#)).

Indeed, there is some empirical evidence that healthcare quality improvement interventions are more successful when they incorporate a component of social influence. For example, [Ivers et al. \(2012\)](#) systematically reviewed the literature on audit and feedback interventions (i.e., quality interventions where a physician is informed of their performance relative to a set standard or the average) and found that feedback and audit interventions are more than three time as effective when the source of the feedback is a supervisor or colleague (i.e., a peer) (adjusted risk difference [ARD]: 16.50) as opposed to a standards review board or employer (ARD: 2.37) or study investigators (ARD: 5.04). Evidence from several qualitative studies have affirmed that the use of peers in quality improvement interventions increases physicians’ receptiveness to change ([Ferguson, Wakeling and Bowie, 2014](#)). Aside from audit and feedback interventions, a Cochrane Review conducted by [Flodgren et al. \(2019\)](#) found that, compared with no intervention, local opinion leader based interventions

improved adjusted median compliance with evidence-based practice by 10.8%. Further, quality interventions involving opinion leaders improved adjusted median compliance with evidence-based practice by 7.1% to 13.7% relative to other types of quality improvement interventions. Medical education also appears to be more effective when interacted with an element of social influence. An early study found that a quality improvement intervention which paired educational outreach visits with feedback had a greater effect on physician's future performance when the outreach was conducted by a peer physician rather than a trained practice assistant ([Van den Hombergh et al., 1999](#)).

Physician moves present another potential quality improvement intervention involving social influence, but to my knowledge, no studies have considered the impact of physician movement on quality improvement. Only a few studies have considered the clinical relevance of entering or exiting physician movers, at all. One study of primary physicians in Norway found that a physician's "practice style" was stable even after he or she moved to a different practice in a different location, suggesting that physicians' preferences play a meaningful role in medical practice variation apart from differences in clientele ([Grytten and Sørensen, 2003](#)). This same method was used by [Epstein and Nicholson \(2009\)](#), who investigated whether individual obstetricians changed their rate of Cesarean-subsection (C-section) procedures in response to moving to a new peer group. While the authors found some evidence of peer influence, the effect size was small; a 2.4 percentage point increase in a physician's peer group's C-section rate driven by moving physicians was associated with a 0.16 percentage point change in the physician's C-section rate. A similar study of cardiologists who moved between two geographic regions in the United States Medicare system found that these physicians adapted to physician behavior in their target region; indeed, the cardiologists adapted their practice, on average, by 0.6 to 0.8 percentage points for every percent point difference between their origin and target region ([Molitor, 2018](#)).

5.1.2 Possible Interventions

There are several potential mover-based interventions that could be used to first test and eventually leverage physician peer effects on healthcare quality. Past studies have implemented interventions incentivising physicians to move from urban to rural settings. These same interventions may be extended to attract physicians from areas with high quality social norms to areas with low quality social norms. Similarly, rotational programs that have traditionally been used to expose training physicians to different areas of medicine (Bell-Dzide, Gokula and Gaspar, 2014) could be modified to expose low-performing practices to physicians from high-performing practices with the intention of improving quality. From the individual practice’s perspective, hiring managers may consider the potential peer-based impact that a new physician may have on a practice. Such a hiring practice may wish to find ways to intensify the social influence of a physician moving from practices with higher quality social norms (e.g., ensuring that the new physician interacts with the practice’s existing physicians) or mitigate the social influence of a physician moving from a practice with lower quality social norms (e.g., increasing training for the new physician).

5.2 Aims

The purpose of this study is to assess whether preliminary evidence exists in support of a GP “mover effect”, i.e., that GPs who move from practice to practice impact the practice’s ability to maximize QOF points. The “mover effect” may provide some validation to of the QOF scheme; if movers can affect QOF scores, under the assumption that context variables remain constant,⁶ then QOF scores are affected by physician effort. In addition, this study aims to describe how GPs have moved within England and Wales’ NHS.

⁶This assumption should not be overlooked. The identifying assumption of the difference-in-difference design is that, conditional on the covariates controlled for in the regression, the treatment and control group would experience a similar trajectory in the outcome if exposed to the same treatment. This analysis controls for fixed time and practice effects, as well as the capacity of the practice, but it is not possible with observational data to entirely rule out the possibility that other covariates that differentially influence the treatment and control groups’ trajectories.

5.3 Methods

5.3.1 Data

5.3.1.1 NHS Quality and Outcomes Framework Data

The NHS provides a set of publicly available datasets⁷ that describe each GP's QOF points in terms of the Quality and Outcomes Framework for each year between 2004/2005 and 2012/2013. The data aggregates QOF points from April 1 of year t to March 31 of year $t + 1$, and so while there is overlap in the calendar years of the individual datasets, they refer to strictly different periods. For the datasets between 2006/2007 and 2011/2012, overall QOF point totals are reported for five domains: Total, Patient Experience, Clinical, Organisational, and Additional Services.

I combine these datasets so that I have panel data of Total, Patient Experience, Clinical, Organisational, and Additional Services quality for each GP practice for each period. I standardize each quality variable to the maximum for each year and domain and then scale the quality variables to 100. Consequently, each quality variable is comparable across years. These data also provide the list size for each practice in each year.

In this analysis, I focus on QOF scores rather than quality achievement. It should be noted that, as explained in the Literature Review and elsewhere (Santos, Gravelle and Propper, 2017; de Wet, McKay and Bowie, 2012), QOF scores are an imperfect measure of quality improvement achievement. For example, some practices may list certain patients as exceptions to maximize their QOF scores without improving quality. Further, QOF points for several clinical indicators are awarded on a nonlinear basis (e.g., lower thresholds under which no points are awarded for that indicator, upper thresholds over which practices receive full credit for that indicator). However, I use QOF points because practice incentives are based upon QOF points rather than QOF achievement, because practices employ strategies to maximize QOF points (Sharma, 2017), and because percentage of points is more comparable

⁷<https://digital.nhs.uk/data-and-information/publications/statistical/quality-and-outcomes-framework-achievement-data>

than achievement metrics across years.

5.3.1.2 GP-by-Practice Tenure Data

The NHS also provides a publicly available dataset⁸ that describes each GP’s tenure with each practice they’ve worked at, updated quarterly. The dataset includes only full GPs (i.e., excludes trainees). Each row of this dataset contains an identifier for the GP, an identifier for the Practice, the date that the GP started working at that practice, and the date the GP stopped working at that practice (if applicable). Importantly, the GP tenure is described as a headcount; that is, individuals are associated with practices, but the amount of work they conduct at that practice (e.g., full-time equivalence) is not described in the publicly available dataset.

	DocID	PRACTICE	parent_org_type	join_date	left_date	amended
1	G8608145	L83647	P	19740401	20160331	1
2	G8435051	C84033	P	20110801	20150211	0
3	G7124307	B81009	P	20140801		0
4	G9043488	A85005	P	20151101		0
5	G9510238	M84037	P	19950904		0
6	G8506841	M85062	P	19740401	19910630	1
7	G9041118	D82012	P	20120601		0
8	G3181182	A86012	P	19740401	19990110	0
9	G8935843	E81065	P	20090401	20120331	0
10	G8509648	A83029	P	19850811	20160222	1

From this dataset, it is possible to construct a dataset of GP moves. First, I filter the dataset by only taking those GPs with more than two rows – these are GPs that have moved practices at least once in their careers. For each GP, I order the observations such that the GP’s earliest appointment is first and their latest appointment is last. When a GP has moved more than once, i.e., they have three or more rows in the dataset, I duplicate all rows that are neither their first nor last appointment. This creates pairs of observations in the dataset, where the first row is where the GP started and the second is where the GP

⁸<https://digital.nhs.uk/services/organisation-data-service/data-downloads/gp-and-gp-practice-related-data>

went. I assign each pair a unique identifier and then transform the data to describe each individual move, with data related to quality at the first practice ($p = 1$) and the second practice ($p = 2$). This includes variables describing when the GP entered her first practice, when she exited her first practice, when she entered her second practice, and when she exited her second practice (if applicable).

I focus on movers who move within a single annual (“April to March”) time period. I then test whether the GP’s date of leaving her first practice and the GP’s date of entering her second practice both fall within each time period between April 1, 2006 and March 31, 2012. If so, I treat that period as the period of the move. I exclude moves that occur where the GP left her first practice in one period and entered her second practice in another period. This is to ensure that the sample focuses only on movers between practices, i.e., switching jobs, and not individuals who are taking extended breaks in their practice. This is important for several reasons. First, it acts as a reality check that we are observing job movers who are likely to be influenced by their most recent job. A physician who leaves one practice and joins another 10 years later is unlikely to be strongly influenced by their original practice, and so it is more challenging to assert that the physician comes from higher or lower quality relative to the destination practice. Further, from a policy standpoint, moving physicians between jobs is a somewhat feasible policy lever while forcing extended breaks from medicine is not. Because several moves occurred right at the March 31 to April 1 cutoff, I extended each time period by 5 days on either side of April 1. Therefore, a mover who left a practice on March 30 in year y and entered a practice in year $y + 1$ would be included in the dataset as an individual who moved in the $y, y + 1$ period. Ultimately, I exclude 1759 moves due to extended breaks.

Because practices with several moves may differ fundamentally from those with zero or one moves (e.g., high turnover), introducing selection bias, and because discerning an effect could be challenging when the treatment occurs more than once, I focus on practices that had only one GP move within my study period and exclude those practices with more than

one entrant from further analysis.

Additionally, I use this dataset to calculate how many GPs work at any given practice on the first date of every month. From this, I calculate the average number of GPs at each practice for each one-year period.

5.3.1.3 Merging the Datasets

I first merge the quality panel data onto the movers data by the time period of the move and the mover's first practice. From this, I extract the QOF points for the practice that each mover left during the period that they left. I will use this data to determine whether the mover came from a better (worse) practice in each quality domain.

I then restrict the quality panel dataset to just those practices that had zero or one mover within the time period. I merge the movers dataset back onto the quality panel dataset by the mover's second practice. This panel dataset now contains, for each practice, the quality metrics for the mover's old practice and the date that the GP moved.

I finally merge the panel dataset with the number of GPs for each period. I divide the number of GPs in the practice by the practice's list size to assess the capacity of the providers in terms of physicians per 1000 patients.

5.3.1.4 Constructed Variables

Using this merged dataset, I construct several variables pertaining to the second GP practice ($p = 2$) relative to the time period of the move (t_0) that will be useful for the analysis and data visualization.

- “time_to_movement $_{p=2, t}$ ” refers to the difference in years between the current period t and the period where the mover entered t_0 (possible values between 1 and 5). For those practices with no movers, t_0 takes on a random value from 1 to 5. This random “movement” period will be used to create a placebo test. If the effect is due to movers

and not to some statistical artifact, then one would expect to observe an effect for movers but not for randomly generated movement periods.

$$\text{time_to_movement}_t = \begin{cases} t - t_0 & \text{Practices with Movers} \\ t - \text{randint}(1,5) & \text{Practice had no movers} \end{cases}$$

- “post_{*t*}” refers to the time that the mover influenced the destination practice, measured as the fraction of time spent at the destination practice in that year. Following past work using an interrupted time series design ([Shover et al., 2019](#)), the “post” variable is defined such that it can take on values between 0 and 1 in time periods $t \geq t_0$ for practices that had movers. If the mover moves in midway through the period, post will equal the fraction of days the mover spent in that practice in that period after she arrived. Similarly, if the mover leaves the destination practice, post will equal the fraction of days the mover spent in the destination practice in that period before she left. For practices with no movers, I assign random movement dates, i.e., random interruptions, to create a placebo counterfactual.

$$\text{post}_t = \begin{cases} 0 & t < t_0 \text{ and Practice had Mover} \\ \frac{\text{Days at Incoming Practice}}{\text{Days in Period } t} & t \geq t_0 \text{ and Practice had Mover} \\ 0 & t < \text{randint}(1,5) \text{ and Practice had No Movers} \\ 1 & t \geq \text{randint}(1,5) \text{ and Practice had No Movers} \end{cases}$$

- For each domain d , “move_{*d*}” takes on values depending upon whether the mover came from practice with better, worse, or the same QOF scores during the time period of the move. The “move” variable is constant within practices (e.g., it doesn’t change within practices across years). Therefore, where $\text{quality}_{d,t_0,p}$ is the normalized QOF score for

domain d at the movement period t_0 for either the destination ($p = 2$) or origin ($p = 1$) practice, the move is defined according to the following:

$$\text{move}_d = \begin{cases} \text{better} & \text{quality}_{d, t_0, p=1} > \text{quality}_{d, t_0, p=2} \\ \text{worse} & \text{quality}_{d, t_0, p=1} < \text{quality}_{d, t_0, p=2} \\ \text{same} & \text{quality}_{d, t_0, p=1} = \text{quality}_{d, t_0, p=2} \\ \text{no move} & \text{Practice had no movers} \end{cases}$$

- For each domain d , “percent_change $_d$ ” is defined using the following formula:

$$\text{percent_change}_d = \frac{\text{quality}_{d, t_0, p=1} - \text{quality}_{d, t_0, p=2}}{\text{quality}_{d, t_0, p=2}} * 100\%$$

Note that when $\text{percent_change}_d > 0$, the mover came from a practice with higher scores in domain d , and when $\text{percent_change}_d < 0$, the mover came from a practice with lower scores in domain d .

- For each move, distance travelled for a particular move is calculated as the (direct) distance between the origin and destination practice. The distance travelled is computed by merging the post code to a database of latitude and longitudes⁹ and then calculating the direct distance in miles (mi).

5.3.2 Statistical Analysis

5.3.2.1 Descriptive Analysis

First, I describe the movement of general practitioners across practices in England’s NHS using the epracmem dataset. For each year, I describe the number of practicing doctors, as well as how many doctors enter, exit, or move by year. For all moves, I describe moves

⁹<https://www.doogal.co.uk/postcodedownloads.php>

by period, including the number of moves; the difference travelled; and the average length of physicians' stay at the origin practice, the destination practice, and between the two practices. Then, I focus on the analytical sample, which included those GPs who moved between two practices with reported QOF scores, and excluded any whose destination practice had more than one entrant over the study period. For the analytical sample, I report the QOF scores of the movers' origin and destination practices by year of move. I provide visualisations of the moves across England and Wales (Figures 7 and 8).

5.3.2.2 Base Case Analysis: Naive Difference-in-Difference

Next, I attempt to provide some evidence that GP movers may affect GP practice QOF scores. The basis for my approach is a Difference-in-Difference design. In a simple two-group, two-period Difference-in-Difference design, the treatment group's outcome value before the treatment and the trend of the control group are used to estimate a counterfactual. Formally, a simple Difference-in-Difference design is estimated with the regression in equation (16).

$$y_{i,t} = \beta_0 + \beta_1 * \text{Period}_t + \beta_2 * 1(\text{Treated})_i + \beta_3 * \text{Period}_t * 1(\text{Treated})_i + X\beta + \epsilon_{i,t} \quad (16)$$

In this equation, $\text{Period}_{i,t}$ takes the value 0 before the intervention and 1 after the intervention, $1(\text{Treated})_i$ takes the value 0 for all participants in the control group and 1 for all participants in the treatment group, and β is the parameter estimate for confounders in the matrix X .

5.3.2.3 Difference-in-Difference with Fixed Effects

A major concern of the Difference-in-Difference design is uncontrolled confounding. While

known confounders can be adjusted for in the regression, there is a possibility for unknown confounders, i.e., unobserved differences between practices in the control and the experimental groups that cause the groups to have different trajectories after the intervention. To ameliorate this concern, I extend the naive Difference-in-Difference design to include practice-level and period-level intercept fixed effects. The practice-level intercept fixed effect accounts for time invariant confounding and period-level intercept fixed effects account for practice invariant confounding.

Fixed intercept effects can be incorporated in the Difference-in-Difference model with the formula in equation (17).

$$\text{quality}_{it} = \beta_0 + \beta_1 * \text{Period}_t + \beta_{2-4} * \text{move}_i + \beta_{5-7} * \text{Period}_t * \text{move}_i + X\beta + \alpha_i + \lambda_t + \epsilon_{it} \quad (17)$$

In this model, α_i takes on a different value for each practice i , and λ_t takes on a different value for each period t . Because these fixed effects account for all time-invariant confounding and all practice-invariant confounding (respectively), the matrix of confounders X should include only those confounders that vary within practices across time periods.

One potential confounder that may vary within practices and periods is capacity. Practices may (naturally) change the number of physicians working at the practice when hiring a new entrant, and this may vary among the treatment groups and with the outcome. For example, practices with a lower doctor-to-patient ratio may be willing to hire a broader range of physicians (including physicians from worse practices) than practices with a higher doctor-to-patient ratio, and these strained practices may also be less able to invest resources into improving quality. Therefore, I include capacity (in terms of GPs per 1,000 patients) as a confounder in the model. Notably, this accounts for practices that have a GP leave (e.g., practices that replace a leaving GP with an incoming GP would have no change in the number of GPs, while practices that add a GP would have an increase in the number of GPs

at the practice).

5.3.2.4 Difference-in-Difference with Practice-Level Random Slope Effects

The difference-in-difference with practice-level fixed intercept effects (equation (17)) assumes that, while the intercept for each practice can be different, the slope (i.e., the relationship between time and the outcome variable) must be parallel. It is possible, however, that individual practices have non-parallel trajectories in the outcomes.

It is not possible to completely eliminate this possibility, but it is possible to add some (albeit parametric) flexibility to the slope parameter. Specifically, I incorporate random slope effects into the model. While fixed effects set a different parameter estimate for each group, random effects use data across groups (i.e., partial pooling) to estimate a distribution for a parameter, thereby treating the parameter as a random variable in the model. Consequently, random effects can account for confounding more efficiently than fixed effects but add the assumption that the underlying parameter follows a normal distribution. In this case, random slope effects ensure that the model is adequately efficient to discern meaningful associations.

5.3.2.5 Main Analysis: Difference-in-Difference with Fixed Effects (Movers Only)

While the Difference-in-Difference design with practice-level and time-level fixed intercept effects accounts for some forms of confounding, it is still possible that practices that accept new GPs are unobservably different from practices that do not accept new GPs. For example, individual practices decide when they need to hire an additional physician to address capacity constraints, and some of the factors for this decision may be unobservable. These differences between practices that hire a new GP and practices that do not hire a new GP may confound the model. Consequently, in the main analysis, I restrict the sample to only those practices that accepted an incoming GP. That is, I compare those practices that had a GP enter

from a practice with better QOF scores only to those practices that had a GP enter from a practice with worse QOF scores (or practices that had a GP enter from a practice with the same QOF scores). In this way, only practices that decided they needed a new physician and successfully hired one are compared.

5.3.2.6 Percentage Change Difference-in-Difference

The main analysis assesses whether having a physician enter from a practice with better QOF scores, a practice with worse QOF scores, or a practice with the same QOF scores is associated with changes in the destination practice's QOF points. It is plausible that these relationships are determined not only by the direction of the difference between the physician's origin and destination practice but also by its magnitude. If having a physician enter from a practice with lower QOF scores causally decreases the QOF scores of the practice, then it is likely that having a physician move from a practice with much lower QOF scores will have a more detrimental effect than having a physician move from a practice with only somewhat lower QOF scores.

To assess this possibility, I estimate the relationship between the outcome (a change in the destination practice's QOF score) and the percentage difference between the moving physician's origin and destination practice. This percentage difference is positive when the physician moves from a practice with higher QOF scores and negative when the physician moves from a practice with lower QOF scores.

$$\begin{aligned}
\text{quality}_{ijt} = & \beta_0 + \beta_1 * \text{Period}_t + \beta_2 * \text{PercentChange}_{ij} + \beta_3 * \text{Post}_t \\
& + \beta_4 * \text{Post}_t * \text{PercentChange}_{ij} + X\beta + \\
& + \alpha_j + \lambda_t + \epsilon_{ijt}
\end{aligned} \tag{18}$$

5.3.3 Sensitivity Analysis: Random Assignment to Treatment Groups

It is possible that a significant finding from the Main Analysis could be due to an overpowered analysis rather than a meaningful association. Consequently, I run a sensitivity placebo-type analysis in which practices are randomly assigned to treatment groups (similar, in principle, to the placebo test applied for synthetic controls ([Abadie, Diamond and Hainmueller, 2010](#))). If the results of this analysis are statistically significant, it is possible that any result from the main analysis is a statistical artefact.

5.4 Results

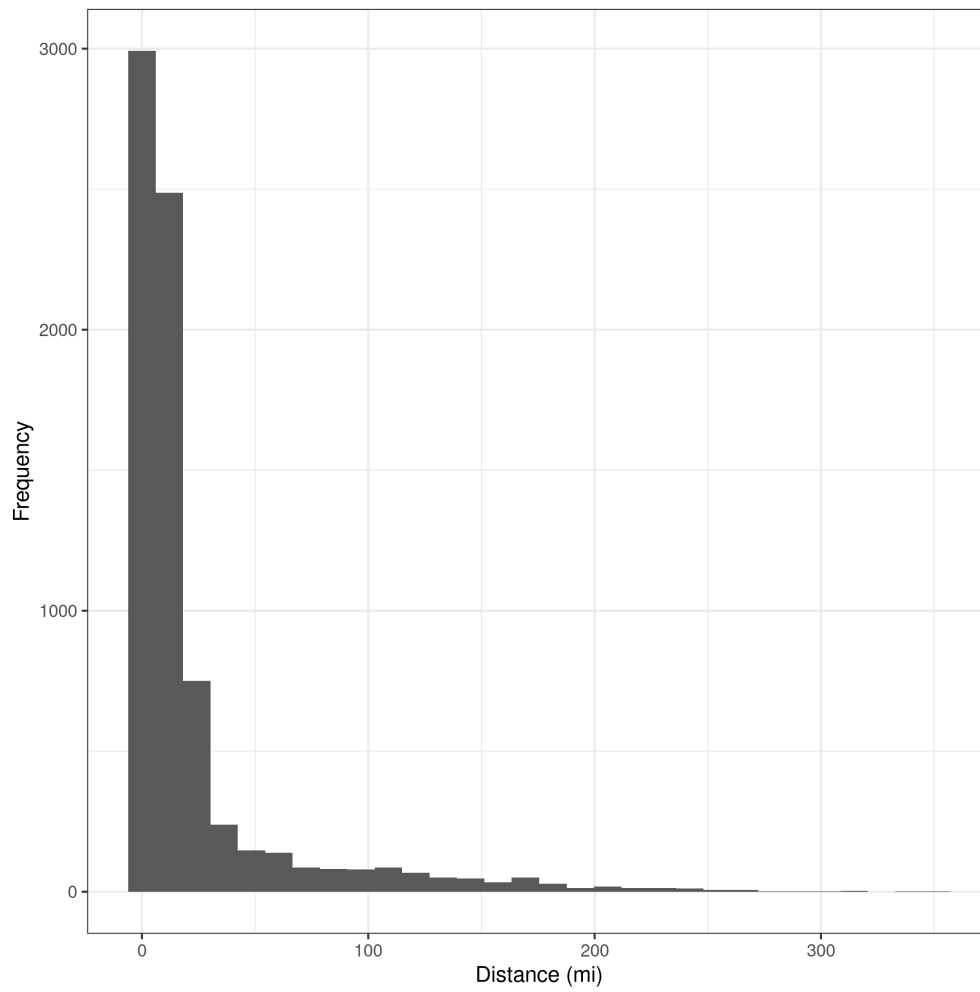
5.4.1 Descriptive Results

The number of practicing GPs has increased steadily from 2004 to 2013, increasing from under 35,000 to over 50,000 (Table 11). The number of GPs starting their career has also increased and consistently outpaces the number of GPs who retire. The number of GPs moving from one practice to another has also increased markedly.

There were a total of 8892 physician moves between 1 April 2004 and 30 March 2013 (Table 12). Physician moves became increasingly common over the course of the study period, nearly doubling from 772 in 2004-2005 to 1483 in 2012-2013. The distance of moves has a positive skew in each period; while the mean distance between a physicians' origin and destination practice is between 18.8 and 26.7 miles, the median distance is between 7.3 and 9.5 miles (Figure 6). Still, in all periods but one, over 20% of moves were over 20 miles and over 10% of moves were over 50 miles. In terms of career trajectory, the median mover had been a physician for between 1000-1500 days at the time of his move and spent less time at their origin practice than they ended up spending at their destination practice (Table 13) within this dataset.

The analytical sample included 2147 GP practices that had exactly one physician move. Within this sample of movers (Table 14), Physicians were slightly more likely to move to

Figure 6: GP Moves by Distance



NOTE: This figure shows GP moves the distance of GP moves occurring between 1 April 2004 and 30 March 2013 in miles. Distances are computed using the latitude and longitude of the practice post codes.

practices with higher total, clinical, and patient experience QOF scores, while physicians did not show a clear leaning towards practices with higher or lower organisational or additional services QOF scores. The median percentage difference between the QOF scores at the physicians origin and destination practice is negligible for nearly all periods and QOF domains. Overall, there is little evidence that physicians are systematically moving to practices with higher or lower QOF scores.

5.4.2 Naive Difference-in-Difference

The naive Difference-in-Difference analysis does not show a clear association between whether a practice had a GP move from a better or worse practice and changes to its QOF scores (Table 15). However, the assumption that the control and treatment groups are relatively similar at baseline is violated, and the naive Difference-in-Difference fails to account for several potential forms of confounding. Therefore, the results of the naive Difference-in-Difference are inconclusive.

5.4.3 Difference-in-Difference with Fixed Effects

Adding practice-level and time-level fixed intercept effects to the analysis accounts for several potential sources of confounding. According to this model, having a GP move from a worse practice significantly decreases patient experience and organisational QOF scores but significantly increases clinical QOF scores relative to practices with no moves (Table 16). Further, having a GP move from a better practice significantly increases total, clinical, organisational, and additional services QOF scores but significantly decreases patient experience QOF scores relative to practices with no moves.

5.4.4 Difference-in-Difference with Practice-Level Random Slope Effects

An analysis using practice-level random slope effects demonstrates that having a GP enter from a better practice is correlated with a significant increase in total, clinical, organisational,

and additional services QOF scores (Table 17). Having a GP enter from a worse practice is correlated with a significant decrease in organisational QOF scores. Having a GP enter from a worse or better practice is associated with a significant decline in patient experience QOF scores.

5.4.5 Difference-in-Difference with Fixed Effects (Movers Only)

When restricting the sample to only those practices with movers, thereby using relatively comparable groups of practices, the associations between having a mover from a better practice and increases in practice QOF become more clear. Relative to those practices that had a GP move from a worse practice, practices that had a GP enter from a better practice saw increases in total, clinical, organisational, and additional services QOF scores (Table 18). For example, when a GP enters a destination practice from an original practice with higher scores, the destination practice experiences, on average, a 1.195 QOF point increase in total QOF points, a 0.676 point increase in clinical domain QOF points, a 2.131 point increase in organizational QOF points, and a 0.771 point increase in additional services QOF points, all relative to practices that had a mover from a worse QOF practice.

It should be noted that these estimates do not account for lags (or leads) of the outcome. Lags may be particularly relevant given that some QOF indicators involve performance over longer than a given fiscal year, and so a practice's performance in the past year could be relevant in the succeeding year. For example, a 2012-2013 clinical indicator is "The percentage of patients with diabetes whose notes record BMI in the preceding 15 months". Therefore, if a practice were to record all diabetic patients' BMI on March 1st of year $t - 1$, then this would earn them the appropriate points for year $t - 1$ and year t . Unfortunately, the short panel analyzed here makes the inclusion of lags challenging. Including one lag term does not qualitatively change the results for total QOF points, though the effects of some specific domains are diminished. For example, the effect of a move from a better practice on clinical (relative to a move from a worse practice) QOF points is diminished from a

statistically significant 0.676 points to a statistically insignificant 0.123 points.

5.4.6 Sensitivity Analysis: Treatment Groups Randomly Assigned

In contrast, when the treatment groups are randomly assigned, there are few observed significant effects (Table 19). The results of this sensitivity analysis suggests that a meaningful association exists and that observed correlations are not statistical artefacts.

5.4.7 Percentage Change Difference-in-Difference (Movers Only)

When accounting for the percentage change between the entrant’s old and new practice, there is a significant positive correlation between the percentage change and the destination practices’ total, clinical, organisational, and additional services QOF scores. These results suggest a sort of “dose-dependent” response – the greater the difference between the entering GP’s origin practice and the destination practice in terms of QOF points, the greater the effect.

5.5 Discussion

5.5.1 The Mover Effect and Quality Improvement

These results present suggestive albeit preliminary evidence that GP movement may have an effect on changes in general practice-level QOF scores. Specifically, the main analysis suggests that practices that had an incoming GP from a practice with a higher QOF scores saw significant increases in total, clinical, organisational, and additional services QOF scores relative to those practices that had an incoming GP from a practice with worse QOF scores. If these results signal a causal effect, GP movement may be an effective policy lever for a centralized health service such as England’s NHS to improve quality scores or to reduce disparities in quality among practices. Given that other policy levers have shown limited effectiveness in affecting quality, the mover effect may be worth pursuing even if it is more difficult to manipulate than other quality improvement interventions (e.g., pay-for-performance

schemes). In addition, there appears to be some evidence that an incoming GP decreases practice-level patient experience scores, regardless of whether the GP comes from a better or worse practice. This result may suggest that GPs adopt the practices of those around them, even if those practices are suboptimal (e.g., not just learning new best practices). In turn, this may suggest the need for GP practices to carefully monitor and take steps to address changes in patient experience after accepting an incoming GP.

5.5.2 Validity of Quality Measures

In addition to the potential for this analysis to inform future quality and QOF score improvement interventions, investigation of a “mover effect” could also serve to inform theoretical approaches to quality improvement and/or to validate the QOF performance indicators.

The dominant framework for improving quality in general practice in England’s NHS is the “Hierarchy and Targets” theoretical model ([Bevan and Fasolo, 2013](#)). This model, in which physicians are rewarded based upon the practice’s achievement relative to preset performance indicators (i.e., QOF points awarded), is predicated on the notion that physicians have some agency over the practice’s performance on those performance indicators. That is, if a practice’s performance on a set of indicators is entirely (or heavily) dependent upon factors outside of a physician’s control (e.g., practice size, deprivation), then offering incentives to the physician will not improve the practice’s performance. Indeed, in this case, incentives may even reduce healthcare equity, as the most fortunate practices – those practices with the best external conditions – will disproportionately reap the incentives.

Physicians point to this possibility when criticising the current “Hierarchy and Targets” theoretical model ([Oliver, 2009](#)), and although targets are currently well-ingrained in the NHS’s scheme for quality improvement, it is not the only model that has been proposed or implemented throughout the history of the NHS. Other models of quality improvement account for this possibility. For example, the “Altruism” theoretical model, which assumes that performance is entirely dictated by outside forces, implies that quality can be improved by

allocating resources to the worst performing practices in hopes that the additional resources will improve the circumstances and performance of those practices (Bevan and Fasolo, 2013).

Consequently, in order to determine whether the capacity for the current targets-based model to achieve its goal of improving healthcare quality, it is necessary to evaluate whether (and the extent that) GPs have meaningful agency over their practice's performance on the QOF indicators. Given the complexity of the health production function, however, it can be challenging to disentangle the effects of physicians from those of external factors, such as the health and health behaviors of the population, community health priorities, and practice-level resources. Focusing on movers into practices, then, presents an opportunity to plausibly observe the effect of physicians rather than other context variables. That is, if a physician moves into a new practice, it is unlikely that the context variables for that practice would meaningfully or systematically change – for example, the patient population, community priorities, and practice-level resources should be constant. If this assumption is met, a change in the practice's QOF point score after an entering GP can reasonably be ascribed to the physician, including her medical practice choices and social influence. It should be noted that further investigation is necessary to ensure that the practice's resources are unchanged (e.g., that the number of GPs to list size is unchanged) and that this effect relates to legitimate quality improvement (e.g., not simply increased attention to completing the appropriate paperwork to maximize QOF points).

The direction of a “mover effect” can then be used to help resolve which model is most appropriate for quality improvement. Under the “Hierarchy and Targets” model, physicians are generally capable of improving their performance given adequate motivation. However, under the “Altruism” theoretical model, physicians generally act to the best of their ability given the information and resources available to them. While a mover is unlikely to meaningfully change the practice's tangible resources, she may bring new information to the practice. In this case, given that physicians are well aware of the quality metrics, it is possible that if a new physician with better information enters the practice, that information may be shared,

and the other physicians in the practice may begin incorporating the new information to improve the quality of their care. On the other hand, under the “Altruism” model, when a physician who enters the practice from a worse performing practice, only the entering GP will adapt her practices to the new information available at the new practice; there should, at least theoretically, be no effect of the entering GP, who now has access to the target practice’s information and resources, on the quality of the practice. Indeed, the addition of any new physician, even one from a worse performing practice, should relieve the patient burden of the other physicians in the practice and, in turn, improve the practice’s quality. Consequently, if GPs entering from better practices positively impact the quality of their incoming practices and GPs entering from worse practices negatively impact the quality of their incoming practices,¹⁰ this would contradict the tenets of the “Altruism” theoretical model. This finding would, instead, constitute evidence then that the physician’s actions – beyond the resources or information available to the practice and the increased capacity from the new physician – affect quality, thereby providing support to the “Hierarchy and Targets” model.

5.5.3 Limitations

This study introduces the concept of a “mover effect” into the field of healthcare quality improvement. However, it is unable to conclusively determine whether and how GP movement causally affects the quality outcomes of practices. This is an early investigation into a “mover effect” and should be considered a feasibility test or preliminary analysis rather than a conclusive analysis.

Perhaps the most prominent limitation of this study is the aggregated nature of practice-level quality data. The QOF regime assigns quality scores to practices (i.e., aggregates of

¹⁰It is at least hypothetically possible that physicians moving from better or worse practices move the target practice’s quality in an unexpected direction. For example, it is possible that physicians may resent a high-performing newcomer and resist adopting her practices. This finding would still suggest that physicians adjust their quality of care in reaction to social influences, but the implications for interventions would be different. In that instance, it might be worthwhile to consider exposing high-quality practices to low-quality GPs through incentivised moves or rotational programs.

GPs), rather than individual GPs. Consequently, I can only observe whether a physician comes from a practice with better or worse quality scores, not whether the moving physician has higher quality than the other physicians in the destination practice. Further, it is not possible to disentangle the effects of the new physician, herself, from her effect on the other physicians in the destination practice. For example, if a practice's quality score is simply the average of an individual-level quality score for each physicians in that practice, it is possible that a relatively high quality physician moving to a relatively low quality practice could increase that practice's aggregate quality scores without changing any of the other physicians' quality scores. That is, the incoming GP's own work cannot be separated from his affect on others. In addition, the publicly available data does not clarify the hours worked by GPs at each practice, and so it is difficult to (A) discern the extent of influence a GP has on her destination practice and (B) to correctly adjust for physician capacity using full-time equivalent physicians.

Additionally, the statistical analyses employed in this study can only provide causal estimates conditional on some key assumptions. For example, the Difference-in-Difference design with fixed practice and period effects assumes no uncontrolled confounders that vary within both practice and period. If such a confounder exists, it could bias these results. I mitigate that concern somewhat by including a random practice-level slope effect in my main model. However, even this model imposes parametric assumptions on the distribution of the QOF trajectories of practices. Overall, whether a mover comes from a better or worse practice is not truly or quasi-randomly assigned, and so any causal estimates rest upon the assumption that all confounders not orthogonal to the treatment are controlled for.

Further, QOF scores are an imperfect measure of quality. As discussed in Section 2, upper thresholds, established in part through a political process, are set relatively low such that most practices obtain a high percentage of available points. Because the average percentage of QOF points achieved is high, QOF points may not capture meaningful variance in quality (e.g., two practices with very different quality may receive the same QOF points if both

practices are above the upper threshold).

5.5.4 Future Research

This study will hopefully serve as a springboard for more research into the newly introduced concept of a “mover effect” in the context of healthcare quality improvement. Researchers may partner with policy-makers to collect more granular information on physicians and their contribution to their practice’s quality. Physician-level data on quality could enable research into whether movers exert peer influence on the other members of their destination practice. Richer data on the physicians, themselves, could reveal which physician-level attributes (e.g., sex, status within the practice) enhance or mitigate a physician’s influence on her peers. Data collected at shorter intervals over a longer stretch of time could demonstrate how physician moves influence the full trajectory of a practice’s quality scores. More discerning and sensitive quality metrics compared with QOF points (e.g., unbounded measures with higher variance) could better show the effect of physician movers. This study shows the feasibility, with greater investment in data collection, of these types of studies.

In addition, random assignment of the treatment (in this case, whether a physician moves from a better or worse practice) could reveal the true, causal effect of a physician mover on destination practice quality. Quasi-random assignment is unlikely in this context, but researchers and policy-makers could design an interventional study to achieve random assignment and mitigate the possibility of unmeasured confounding. The suggestive results of this preliminary study using already collected, publicly available data imply that a more robust study may be worth the investment.

GP movement is a potentially important and understudied feature of GP labor supply, even though GP movement has become more common in recent decades. This study suggests that GP movement may play a meaningful role in healthcare quality improvement. If further research confirms the “mover effect”, it could become a powerful policy lever in England’s NHS and in other healthcare systems.

Table 11: Overview of GP Starts, Retirements, and Moves

Year	No. Practicing GPs	No. GP Starts	No. GP Retirements	No. GP Movers
2004	34,787	2,374	1,060	733
2005	35,994	2,364	999	721
2006	38,222	3,075	1,282	796
2007	40,013	2,997	1,352	771
2008	41,608	2,878	1,433	830
2009	43,611	3,392	1,657	1,109
2010	45,437	3,409	2,002	1,362
2011	46,681	3,242	2,112	1,424
2012	47,831	3,164	2,417	1,543
2013	50,128	4,528	3,087	1,808

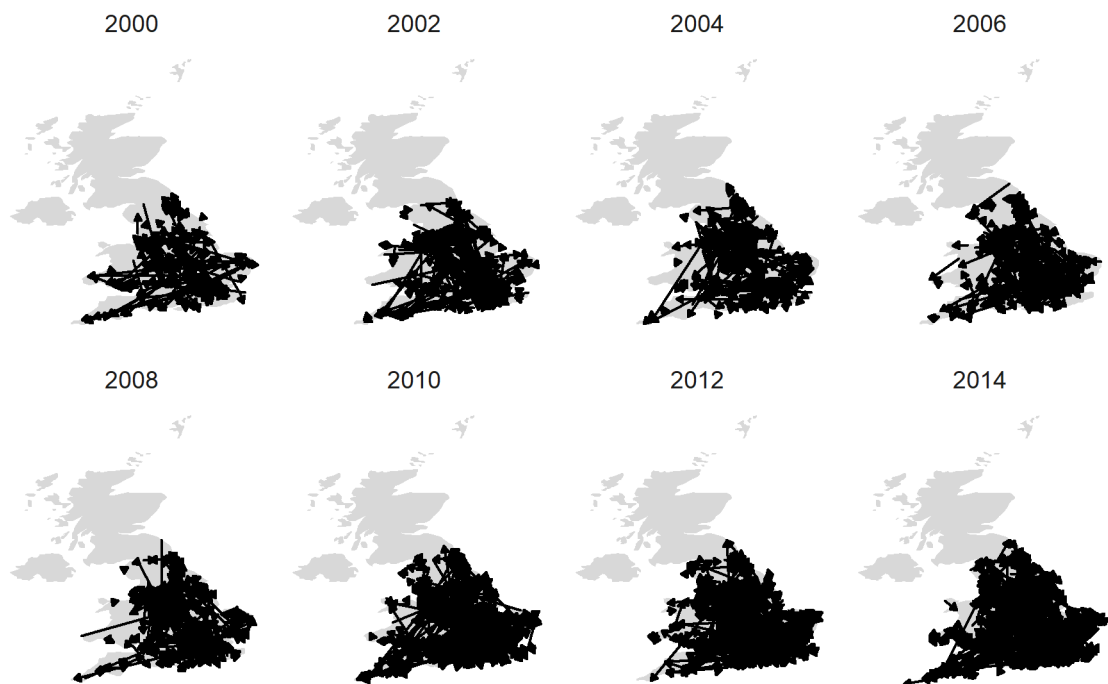
NOTE: This table presents an overview of the starts, retirements, and moves of GPs. “No. GP Starts” refers to the number of GPs who began their careers in a given calendar year. “No. GP Retirements” refers to the number of GPs who stopped practicing until at least 2016 in a given calendar year. (Note that a doctor who stopped practicing in 2012 but continued practicing in 2018 would be considered “retired” in 2012 by this definition). “No. GP Moves” refers to the number of physicians who, within a given calendar year, left at least one practice and joined at least one practice.

Table 12: Distance of GP Moves in Miles

Period	N	Mean Distance	Median Distance (IQR)	Over 20 Miles (%)	Over 50 Miles (%)
2004-2005	772	23.6	7.6 (2.8 to 17.9)	22.9	14.1
2005-2006	665	23.4	8 (2.9 to 21.3)	26.2	14.2
2006-2007	721	24.2	8.6 (3.3 to 20.4)	25.4	15.3
2007-2008	678	26.7	9.3 (3.2 to 21.5)	27.1	14.6
2008-2009	819	25	9.1 (3 to 21.8)	27.3	14.1
2009-2010	1125	23.4	9.5 (3.4 to 20.9)	27	12.4
2010-2011	1254	23.1	8.9 (3.4 to 20.2)	25.2	11.5
2011-2012	1375	21.2	8.5 (3.2 to 18.5)	23.2	10.5
2012-2013	1483	18.8	7.3 (2.7 to 16.1)	19	8.9

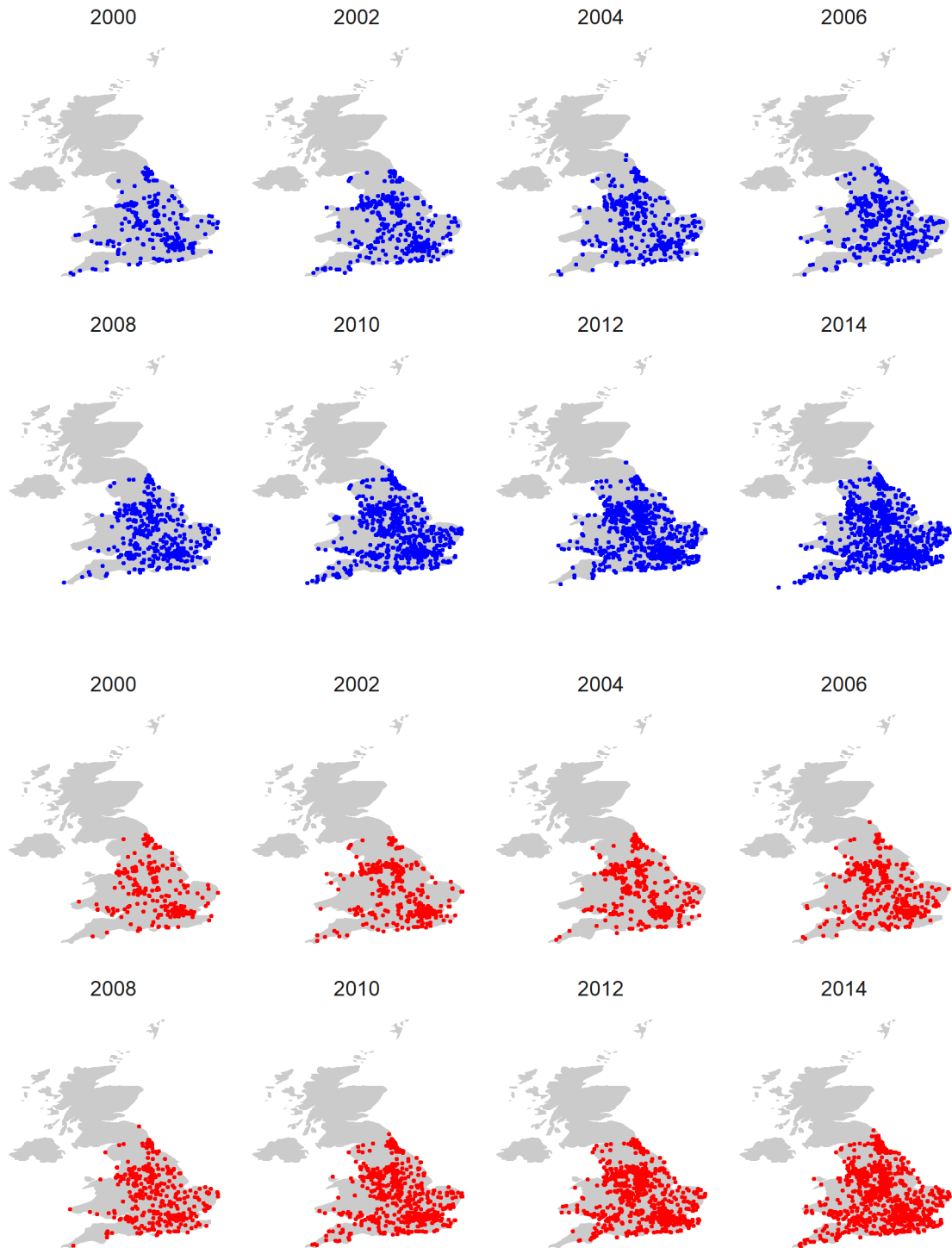
NOTE: This table presents data on the distance of moves for each period. Distance was calculated using the latitude and longitude of the GP’s origin and destination practice. “N” is the number of moves occurring in that period. “Mean Distance” is the average number of miles for each move within that period. “Median Distance” is the median number of miles for each move within that period, as well as the interquartile range. “Over 20 Miles” and “Over 50 Miles” are the proportion of moves occurring within each period that exceeded 20 and 50 miles, respectively.

Figure 7: GP Moves by Year



This figure shows GP moves by year from 2000 to 2014. The arrows point in the direction of the mover. GPs who move from or to a postcode that cannot be geocoded are omitted.

Figure 8: GP Moves by Year



This figure shows GP moves by year from 2000 to 2014. Panel A shows where GPs have moved from using red points. Panel B shows where GPs have moved to using blue points. GPs who move from or to a postcode that cannot be geocoded are omitted.

Table 13: Summary of GP Moves

Period	N	Entry to Move (Median Days)	Time at Origin (Median Days)	Time at Destination (Median Days)
2004-2005	772	1806	1036	3220.5
2005-2006	665	1186	717	2918
2006-2007	721	1099	641	2454
2007-2008	678	992.5	668	2493
2008-2009	819	1005	716	2067
2009-2010	1125	1086	720	1816
2010-2011	1254	1279	777.5	1841
2011-2012	1375	1443	851	1552
2012-2013	1483	1513	762	2831

NOTE: This table presents data on the career trajectories of GP movers. “Entry to Move” is the median number of days between the date the physician became a registered NHS physician and the date they joined a new practice from an old practice. “Time at Origin” is the median number of days the mover spent at their original practice, and “Time at Destination” is the median number of days the mover spent at their destination practice.

Table 14: Summary of Moves in Terms of Domain-Specific QOF Scores

Period	To Better	To Same	To Worse	Mean Origin Score	Mean Destination Score	Median Pct. Difference (IQR)
Total						
2006-2007	176 (66.4%)		89 (33.6%)	95.6	95.9	0% (-2.2 to 2.11)
2007-2008	182 (75.8%)		58 (24.2%)	96.8	97.4	-0.1% (-1.4 to 1.39)
2008-2009	109 (39.8%)		165 (60.2%)	95	95.6	0.8% (-1.3 to 3.28)
2009-2010	128 (32.5%)		266 (67.5%)	93.8	93.3	0.3% (-3.5 to 3.38)
2010-2011	148 (31.8%)		317 (68.2%)	94.3	95.3	1% (-1.8 to 3.98)
2011-2012	346 (68%)		163 (32%)	97.1	97.3	0% (-1.4 to 1.66)
Clinical						
2006-2007	127 (47.9%)	11 (4.2%)	127 (47.9%)	96.5	96.7	0% (-1.6 to 1.35)
2007-2008	124 (51.7%)	13 (5.4%)	103 (42.9%)	97.3	97.8	0% (-1.7 to 0.66)
2008-2009	116 (42.3%)	20 (7.3%)	138 (50.4%)	98	98.3	0% (-0.7 to 0.83)
2009-2010	191 (48.5%)	8 (2%)	195 (49.5%)	96.6	95.8	0% (-2.4 to 2.25)
2010-2011	202 (43.4%)	7 (1.5%)	256 (55.1%)	96.8	97.5	0.2% (-1 to 2.54)
2011-2012	263 (51.7%)	10 (2%)	236 (46.4%)	97.1	97.3	-0.1% (-1.6 to 1.75)
Organisational						
2006-2007	130 (49.1%)	12 (4.5%)	123 (46.4%)	92.4	92.6	0% (-4 to 2.95)
2007-2008	110 (45.8%)	7 (2.9%)	123 (51.2%)	94.5	95.4	0.2% (-2.2 to 3.22)
2008-2009	129 (47.1%)	11 (4%)	134 (48.9%)	96.4	95	0% (-2.7 to 1.6)
2009-2010	196 (49.7%)	27 (6.9%)	171 (43.4%)	97.3	96.3	0% (-2.1 to 1.41)
2010-2011	202 (43.4%)	34 (7.3%)	229 (49.2%)	97.5	98	0% (-1.4 to 1.54)
2011-2012	224 (44%)	15 (2.9%)	270 (53%)	96.6	97	0.1% (-1.3 to 1.61)
Patient Experience						
2006-2007	225 (84.9%)	22 (8.3%)	18 (6.8%)	95.8	96.7	0% (0 to 0)
2007-2008	221 (92.1%)	10 (4.2%)	9 (3.8%)	97.5	98.6	0% (0 to 0)
2008-2009	108 (39.4%)	5 (1.8%)	161 (58.8%)	79.5	84.3	6.6% (-4.2 to 22.53)
2009-2010	108 (27.4%)	4 (1%)	282 (71.6%)	65.4	67.4	5.5% (-23.5 to 40.61)
2010-2011	131 (28.2%)	1 (0.2%)	333 (71.6%)	67.2	71.9	3.3% (-19 to 44.2)
2011-2012	463 (91%)	43 (8.4%)	3 (0.6%)	99.4	99.4	0% (0 to 0)
Additional Services						
2006-2007	60 (22.6%)	145 (54.7%)	60 (22.6%)	97	97.2	0% (0 to 0)
2007-2008	49 (20.4%)	131 (54.6%)	60 (25%)	97.2	98.1	0% (0 to 0.01)
2008-2009	47 (17.2%)	173 (63.1%)	54 (19.7%)	98.4	97.4	0% (0 to 0)
2009-2010	179 (45.4%)	66 (16.8%)	149 (37.8%)	96.5	95.6	0% (-3.8 to 2.92)
2010-2011	150 (32.3%)	151 (32.5%)	164 (35.3%)	97.8	97.8	0% (-1 to 1.42)
2011-2012	175 (34.4%)	162 (31.8%)	172 (33.8%)	97.7	97.6	0% (-1.3 to 1.22)

NOTE: This table summarises the moves that are included in the analytical sample, i.e., moves into practices that had only one entrant between April 2006 and March 2012. “Period” is the April 1 to March 30 year. “To Better” are those GPs who moved from a practice with lower scores to a practice with higher scores. “To Worse” are those GPs who moved from a practice with a higher score to a practice with a lower score. “To Same” are those GPs who moved from a practice with the same score as the practice they moved from. “Mean Origin Score” is the average domain-specific QOF score for the GP mover’s original practice. “Mean Destination Score” is the average domain-specific QOF score for the GP mover’s destination practice. “Mean Difference” is the difference between the mover’s original practice QOF score and the mover’s destination practice QOF score.

Table 15: Naïve Difference-in-Difference (No Fixed Effects)

	<i>Dependent variable:</i>				
	Total (1)	Patient Experience (2)	Clinical (3)	Organizational (4)	Additional Services (5)
Common Intercept	95.216*** (95.127, 95.306)	88.663*** (88.378, 88.949)	96.561*** (96.474, 96.647)	94.361*** (94.228, 94.494)	96.140*** (96.007, 96.274)
Δ Intercept for Practices with Incoming GP from Same		11.363*** (9.093, 13.632)	2.618*** (1.815, 3.420)	3.066*** (2.150, 3.981)	2.906*** (2.532, 3.281)
Δ Intercept for Practices with Incoming GP from Worse	2.198*** (1.986, 2.410)	1.597*** (0.833, 2.360)	1.907*** (1.700, 2.115)	2.665*** (2.347, 2.983)	2.552*** (2.168, 2.936)
Δ Intercept for Practices with Incoming GP from Better	-0.808*** (-1.025, -0.592)	-3.102*** (-3.748, -2.456)	-0.572*** (-0.785, -0.359)	-0.704*** (-1.038, -0.371)	-1.340*** (-1.723, -0.958)
Any Move	-0.042 (-0.181, 0.097)	-4.026*** (-4.471, -3.581)	-0.059 (-0.194, 0.076)	1.753*** (1.546, 1.960)	-0.018 (-0.226, 0.190)
Effect of Move of GP from Same		4.032* (-0.221, 8.286)	-1.156* (-2.419, 0.107)	-0.176 (-1.898, 1.546)	-0.053 (-0.650, 0.543)
Effect of Move of GP from Worse	-0.193 (-0.567, 0.180)	-0.133 (-1.480, 1.214)	0.212 (-0.146, 0.571)	-0.493* (-1.047, 0.062)	-0.096 (-0.797, 0.604)
Effect of Move of GP from Better	1.098*** (0.731, 1.465)	1.635*** (0.538, 2.733)	0.741*** (0.371, 1.111)	0.416 (-0.151, 0.984)	-0.284 (-1.006, 0.438)
Practice FE	NO	NO	NO	NO	NO
Year FE	NO	NO	NO	NO	NO
Practice RE	NO	NO	NO	NO	NO
Observations	43,193	43,193	43,193	43,193	43,193

NOTE: All regressions are adjusted for confounding by the doctor-to-patient ratio.

Table 16: Difference-in-Difference with Practice and Time Fixed Intercept Effects

	<i>Dependent variable:</i>				
	Total (1)	Patient Experience (2)	Clinical (3)	Organizational (4)	Additional Services (5)
Any Move	0.136** (0.0002, 0.271)	0.558** (0.054, 1.063)	0.031 (−0.095, 0.157)	0.199 (−0.038, 0.436)	0.139 (−0.095, 0.374)
Effect of Move of GP from Same		1.976 (−2.741, 6.693)	0.055 (−0.902, 1.012)	−0.114 (−1.623, 1.394)	−0.121 (−0.674, 0.432)
Effect of Move of GP from Worse	−0.060 (−0.331, 0.212)	−1.103** (−2.170, −0.035)	0.407*** (0.144, 0.669)	−0.558** (−1.064, −0.052)	0.029 (−0.583, 0.642)
Effect of Move of GP from Better	1.096*** (0.791, 1.401)	−1.853*** (−2.943, −0.763)	1.078*** (0.800, 1.356)	1.612*** (1.093, 2.131)	0.741** (0.117, 1.364)
Practice FE	YES	YES	YES	YES	YES
Year FE	YES	YES	YES	YES	YES
Practice RE	NO	NO	NO	NO	NO
Observations	43,193	43,193	43,193	43,193	43,193

NOTE: All regressions are adjusted for confounding by the doctor-to-patient ratio.

Table 17: Difference-in-Difference with Practice Random Slope Effects

	<i>Dependent variable:</i>				
	Total (1)	Patient Experience (2)	Clinical (3)	Organizational (4)	Additional Services (5)
Δ Intercept of Practices with Incoming GP from Same		11.931*** (9.398, 14.465)	2.566*** (1.137, 3.995)	3.096*** (1.679, 4.513)	3.242*** (2.632, 3.851)
Δ Intercept of Practices with Incoming GP from Worse	2.376*** (1.992, 2.761)	2.616*** (1.808, 3.424)	2.310*** (1.919, 2.701)	2.501*** (2.006, 2.997)	2.777*** (2.127, 3.427)
Δ Intercept of Practices with Incoming GP from Better	-0.556*** (-0.937, -0.175)	-1.363*** (-2.067, -0.659)	-0.651*** (-1.049, -0.254)	-0.782*** (-1.298, -0.266)	-1.755*** (-2.400, -1.109)
Any Move	0.131** (0.002, 0.260)	0.375 (-0.082, 0.832)	0.057 (-0.063, 0.178)	0.209* (-0.002, 0.420)	0.075 (-0.143, 0.294)
Effect of Move of GP from Same		0.732 (-3.261, 4.724)	-0.208 (-1.249, 0.833)	-0.281 (-1.839, 1.278)	-0.296 (-0.859, 0.268)
Effect of Move of GP from Worse	-0.100 (-0.403, 0.203)	-0.970* (-2.034, 0.094)	0.197 (-0.095, 0.488)	-0.525** (-1.041, -0.009)	-0.167 (-0.805, 0.471)
Effect of Move of GP from Better	0.856*** (0.536, 1.176)	-1.116** (-2.100, -0.133)	1.116*** (0.811, 1.421)	1.171*** (0.639, 1.703)	0.666** (0.020, 1.313)
Practice FE	NO	NO	NO	NO	NO
Year FE	YES	YES	YES	YES	YES
Practice Intercept RE	YES	YES	YES	YES	YES
Practice Slope RE	YES	YES	YES	YES	YES
Observations	43,193	43,193	43,193	43,193	43,193
Akaike Inf. Crit.	253,368.600	359,030.200	247,656.500	295,798.800	298,463.500
Bayesian Inf. Crit.	253,507.400	359,186.400	247,812.600	295,954.900	298,619.600

NOTE: All regressions are adjusted for confounding by the doctor-to-patient ratio.

Table 18: Difference-in-Difference with Fixed Intercept Effects, Only Movers

	<i>Dependent variable:</i>				
	Total (1)	Patient Experience (2)	Clinical (3)	Organizational (4)	Additional Services (5)
Effect for Move of GP from Worse	0.032 (−0.235, 0.299)	0.221 (−0.776, 1.218)	0.365*** (0.119, 0.611)	−0.410* (−0.880, 0.060)	0.155 (−0.234, 0.545)
Δ Effect for Move of GP from Better	1.195*** (0.855, 1.535)	−1.009 (−2.240, 0.221)	0.676*** (0.366, 0.986)	2.131*** (1.543, 2.720)	0.771*** (0.189, 1.353)
Practice FE	YES	YES	YES	YES	YES
Year FE	YES	YES	YES	YES	YES
Practice RE	NO	NO	NO	NO	NO
Observations	12,521	12,272	12,521	12,521	12,521

NOTE: All regressions are adjusted for confounding by the doctor-to-patient ratio.

Table 19: Sensitivity Analysis: Treatment Groups Randomly Assigned

	<i>Dependent variable:</i>				
	Total (1)	Patient Experience (2)	Clinical (3)	Organizational (4)	Additional Services (5)
Any Move	0.249** (0.047, 0.451)	0.489 (-0.273, 1.252)	0.010 (-0.178, 0.198)	0.361** (0.001, 0.721)	0.235 (-0.122, 0.592)
Effect of Move of GP from Same	0.038 (-0.224, 0.300)	0.181 (-0.792, 1.153)	0.225* (-0.015, 0.466)	0.137 (-0.316, 0.591)	-0.053 (-0.505, 0.398)
Effect of Move of GP from Worse	-0.035 (-0.294, 0.223)	-0.416 (-1.381, 0.549)	0.252** (0.011, 0.493)	-0.276 (-0.736, 0.183)	-0.157 (-0.608, 0.293)
Effect of Move of GP from Better	-0.032 (-0.289, 0.226)	-0.768 (-1.742, 0.206)	0.247** (0.006, 0.488)	-0.096 (-0.556, 0.363)	0.002 (-0.453, 0.458)
Practice FE	YES	YES	YES	YES	YES
Year FE	YES	YES	YES	YES	YES
Practice RE	NO	NO	NO	NO	NO
Observations	43,193	43,193	43,193	43,193	43,193

NOTE: All regressions are adjusted for confounding by the doctor-to-patient ratio.

Table 20: Percentage Change Difference-in-Difference

	<i>Dependent variable:</i>				
	Total (1)	Patient Experience (2)	Clinical (3)	Organizational (4)	Additional Services (5)
Any Move * % Δ with GP's Last Practice	0.171*** (0.147, 0.195)	-0.017* (-0.038, 0.003)	0.055*** (0.034, 0.076)	0.178*** (0.159, 0.197)	0.001*** (0.001, 0.002)
Practice FE	YES	YES	YES	YES	YES
Year FE	YES	YES	YES	YES	YES
Practice RE	NO	NO	NO	NO	NO
Observations	12,521	12,272	12,521	12,521	12,521

NOTE: All regressions are adjusted for confounding by the doctor-to-patient ratio.

References

- Abadie, Alberto, Alexis Diamond, and Jens Hainmueller.** 2010. "Synthetic Control Methods for Comparative Case Studies: Estimating the Effect of California's Tobacco Control Program." *Journal of the American Statistical Association*, 105(490): 493–505.
- Addicott, Rachael, and Christopher Ham.** 2014. *Commissioning and Funding General Practice: Making the Case for Family Care Networks*.
- Ali, Mir M., and Debra S. Dwyer.** 2011. "Estimating Peer Effects in Sexual Behavior among Adolescents." *Journal of Adolescence*, 34(1): 183–190.
- Andreassen, Leif, Maria Laura Di Tommaso, and Steinar Strøm.** 2013. "Do Medical Doctors Respond to Economic Incentives?" *Journal of Health Economics*, 32(2): 392–409.
- Anonymous.** 1953. "General Practice Today and Tomorrow." *British Medical Journal*, 717–719.
- Anthony, E.** 1950. "The G.P. at the Crossroads." *British Medical Journal*, 1(4661): 1077–1079.
- Applebee, Kathie.** 2006. "Understanding the Revised QOF." *Independent Nurse*, 2006(5).
- Arnold, S. R., and S. E. Straus.** 2006. "Interventions to Improve Antibiotic Prescribing Practices in Ambulatory Care." *Evidence-Based Child Health: A Cochrane Review Journal*, 1(2): 623–690.
- Aveling, Emma-Louise, Graham Martin, Natalie Armstrong, Jay Banerjee, and Mary Dixon-Woods.** 2012. "Quality Improvement through Clinical Communities: Eight Lessons for Practice." *Journal of Health Organization and Management*, 26(2): 158–174.
- Avorn, Jerry, and Stephen B. Soumerai.** 1983. "Improving Drug-Therapy Decisions through Educational Outreach." *New England Journal of Medicine*, 308(24): 1457–1463.

- Baker, R. H.** 1989. “The Quality Initiative of the Royal College of General Practitioners.” *Quality Assurance in Health Care: The Official Journal of the International Society for Quality Assurance in Health Care*, 1(1): 23–29.
- Baker, R., M. Lakhani, R. Fraser, and F. Cheater.** 1999. “A Model for Clinical Governance in Primary Care Groups.” *BMJ*, 318(7186): 779–783.
- Bandura, Albert.** 1986. *Social Foundations of Thought and Action: A Social Cognitive Theory*. Social Foundations of Thought and Action: A Social Cognitive Theory, Englewood Cliffs, NJ, US:Prentice-Hall, Inc.
- Batalden, P. B, and F. Davidoff.** 2007. “What Is ”Quality Improvement” and How Can It Transform Healthcare?” *Quality and Safety in Health Care*, 16(1): 2–3.
- Beardow, R., K. Cheung, and W. M. Styles.** 1993. “Factors Influencing the Career Choices of General Practitioner Trainees in North West Thames Regional Health Authority.” *British Journal of General Practice*, 43(376): 449–452.
- Bell, Andrew, Malcolm Fairbrother, and Kelvyn Jones.** 2019. “Fixed and Random Effects Models: Making an Informed Choice.” *Quality & Quantity*, 53(2): 1051–1074.
- Bell-Dzide, Dodzi, Murthy Gokula, and Phyllis Gaspar.** 2014. “Effect of a Long-Term Care Geriatrics Rotation on Physician Assistant Students’ Knowledge and Attitudes towards the Elderly.” *The Journal of Physician Assistant Education: The Official Journal of the Physician Assistant Education Association*, 25(1): 38–40.
- Bennett, Brian, Emilene Coventry, Nicola Greenway, and Mark Minchin.** 2014. “The NICE Process for Developing Quality Standards and Indicators.” *Zeitschrift für Evidenz, Fortbildung und Qualität im Gesundheitswesen*, 108(8): 481–486.
- Berwick, Donald, and Daniel M. Fox.** 2016. ““Evaluating the Quality of Medical Care”: Donabedian’s Classic Article 50 Years Later.” *The Milbank Quarterly*, 94(2): 237–241.

- Bevan, Gwyn.** 2010. “Performance Measurement of “Knights” and “Knaves”: Differences in Approaches and Impacts in British Countries after Devolution.” *Journal of Comparative Policy Analysis: Research and Practice*, 12(1-2): 33–56.
- Bevan, Gwyn, and Barbara Fasolo.** 2013. “Models of Governance of Public Services: Empirical and Behavioural Analysis of ‘econs’ and ‘Humans’.” In *Behavioural Public Policy..* , ed. Oliver Angus, 38–62. Cambridge, UK:Cambridge University Press.
- Bird, Sheila M., Cox Sir David, Vern T. Farewell, Goldstein Harvey, Holt Tim, and Smith Peter C.** 2005. “Performance Indicators: Good, Bad, and Ugly.” *Journal of the Royal Statistical Society: Series A (Statistics in Society)*, 168(1): 1–27.
- Bollinger, Bryan, and Kenneth Gillingham.** 2012. “Peer Effects in the Diffusion of Solar Photovoltaic Panels.” *Marketing Science*, 31(6): 900–912.
- Boulkedid, Rym, Hendy Abdoul, Marine Loustau, Olivier Sibony, and Corinne Alberti.** 2011. “Using and Reporting the Delphi Method for Selecting Healthcare Quality Indicators: A Systematic Review.” *PLOS ONE*, 6(6): e20476.
- Braithwaite, Jeffrey.** 2018. “Changing How We Think about Healthcare Improvement.” *BMJ*, k2014.
- British Medical Journal.** 1985. ““Towards Quality in General Practice”: GMSC Discusses RCGP’s Paper.” *Br Med J (Clin Res Ed)*, 291(6500): 987–988.
- Bursztyn, Leonardo, Florian Ederer, Bruno Ferman, and Noam Yuchtman.** 2014. “Understanding Mechanisms Underlying Peer Effects: Evidence From a Field Experiment on Financial Decisions.” *Econometrica*, 82(4): 1273–1301.
- Campbell, S M.** 2002. “Research Methods Used in Developing and Applying Quality Indicators in Primary Care.” *Quality and Safety in Health Care*, 11(4): 358–364.

- Campbell, Stephen, and Helen Lester.** 2010. “Developing Indicators and the Concept of QOFability.” In *The Quality and Outcomes Framework: QOF - Transforming General Practice.* , ed. Stephen Gillam and A. Niroshan Siriwardena, 16–27. Radcliffe Publishing.
- Campbell, Stephen M., Evangelos Kontopantelis, Kerin Hannon, Martyn Burke, Annette Barber, and Helen E. Lester.** 2011. “Framework and Indicator Testing Protocol for Developing and Piloting Quality Indicators for the UK Quality and Outcomes Framework.” *BMC Family Practice*, 12(1): 85.
- Carlisle, R., and S. Johnstone.** 1996. “Factors Influencing the Response to Advertisements for General Practice Vacancies.” *BMJ*, 313(7055): 468–471.
- Castro-Avila, Ana, Karen Bloor, and Carl Thompson.** 2019. “The Effect of External Inspections on Safety in Acute Hospitals in the National Health Service in England: A Controlled Interrupted Time-Series Analysis.” *Journal of Health Services Research & Policy*, 24(3): 182–190.
- Cheng, Lim Puay, Nelson K.H. Tang, and Peter M. Jackson.** 1999. “An Innovative Framework for Health Care Performance Measurement.” *Managing Service Quality: An International Journal*, 9(6): 423–433.
- Chien, Alyna T., and Meredith B. Rosenthal.** 2013. “Medicare’s Physician Value-Based Payment Modifier — Will the Tectonic Shift Create Waves?” *New England Journal of Medicine*, 369(22): 2076–2078.
- Chung, Paul J, Jeanette Chung, Manish N. Shah, and David O. Meltzer.** 2003. “How Do Residents Learn? The Development of Practice Styles in a Residency Program.” *Ambulatory Pediatrics*, 3(4): 166–172.
- Cockburn, J., D. Ruth, C. Silagy, M. Dobbin, Y. Reid, M. Scollo, and L. Naccarella.** 1992. “Randomised Trial of Three Approaches for Marketing Smoking

- Cessation Programmes to Australian General Practitioners.” *British Medical Journal*, 304(6828): 691–694.
- Collings, Joseph S.** 1950. “General Practice in England Today - A Reconnaissance.” *The Lancet*, 255(6604): 555.
- Conn, David.** 1982. “Effort, Efficiency, and Incentives in Economic Organizations.” *Journal of Comparative Economics*, 6(3): 223–234.
- Conry, Mary C, Niamh Humphries, Karen Morgan, Yvonne McGowan, Anthony Montgomery, Kavita Vedhara, Efharis Panagopoulou, and Hannah Mc Gee.** 2012. “A 10 Year (2000–2010) Systematic Review of Interventions to Improve Quality of Care in Hospitals.” *BMC Health Services Research*, 12: 275.
- Coulter, Angela.** 1995. “Evaluating General Practice Fundholding in the United Kingdom.” *European Journal of Public Health*, 5(4): 233–239.
- Craig, Peter, Cyrus Cooper, David Gunnell, Sally Haw, Kenny Lawson, Sally Macintyre, David Ogilvie, Mark Petticrew, Barney Reeves, Matt Sutton, and Simon Thompson.** 2012. “Using Natural Experiments to Evaluate Population Health Interventions: New Medical Research Council Guidance.” *J Epidemiol Community Health*, 66(12): 1182–1186.
- Crémieux, Pierre-Yves, and Pierre Ouellette.** 2001. “Omitted Variable Bias and Hospital Costs.” *Journal of Health Economics*, 20(2): 271–282.
- Crossing the Quality Chasm: A New Health System for the 21st Century.** 2001. *Crossing the Quality Chasm: A New Health System for the 21st Century*. Washington, D.C.:National Academies Press.
- Dahl, Gordon B., Katrine V. Løken, and Magne Mogstad.** 2014. “Peer Effects in Program Participation.” *American Economic Review*, 104(7): 2049–2074.

- Dambrun, Michaël, Serge Guimond, and Sandra Duarte.** 2002. “The Impact of Hierarchy-Enhancing vs. Attenuating Academic Major on Stereotyping: The Mediating Role of Perceived Social Norm.” *Current Research in Social Psychology*, 7(8): 114–136.
- Dawson, G. de H.** 1950. “The G.P. at the Crossroads.” *British Medical Journal*, 1(4667): 1431.
- de Jong, Judith D., Gert P. Westert, Ronald Lagoe, and Peter P. Groenewegen.** 2006. “Variation in Hospital Length of Stay: Do Physicians Adapt Their Length of Stay Decisions to What Is Usual in the Hospital Where They Work?” *Health Services Research*, 41(2): 374–394.
- de Jong, Judith D, Peter P Groenewegen, and Gert P Westert.** 2003. “Mutual Influences of General Practitioners in Partnerships.” *Social Science & Medicine*, 57(8): 1515–1524.
- de Wet, Carl, John McKay, and Paul Bowie.** 2012. “Combining QOF Data with the Care Bundle Approach May Provide a More Meaningful Measure of Quality in General Practice.” *BMC Health Services Research*, 12(1): 351.
- Deyo-Svendsen, Mark E., Michael R. Phillips, Jill K. Albright, Keith A. Schilling, and Karl B. Palmer.** October/December 2016. “A Systematic Approach to Clinical Peer Review in a Critical Access Hospital.” *Quality Management in Healthcare*, 25(4): 213.
- Dixon, Anna, Artak Khachatryan, Andrew Wallace, Stephen Peckham, Tammy Boyce, and Stephen Gillam.** 2011. “The Quality and Outcomes Framework (QOF): Does It Reduce Health Inequalities?”
- Dolea, Carmen, Laura Stormont, and Jean-Marc Braichet.** 2010. “Evaluated Strategies to Increase Attraction and Retention of Health Workers in Remote and Rural Areas.” *Bulletin of the World Health Organization*, 88: 379–385.

- Donabedian, Avedis.** 1966. "Evaluating the Quality of Medical Care." *The Milbank Memorial Fund Quarterly*, 44(3): 166–206.
- Donabedian, Avedis.** 1980. *Explorations in Quality Assessment and Monitoring: The Definition of Quality and Approaches to Its Assessment: 1.* . Spi edition ed., Ann Arbor, Mich:Health Administration Press.
- Doran, Tim, Evangelos Kontopantelis, David Reeves, Matthew Sutton, and Andrew M. Ryan.** 2014. "Setting Performance Targets in Pay for Performance Programmes: What Can We Learn from QOF?" *BMJ*, 348.
- Downing, Amy, Gavin Rudge, Yaping Cheng, Yu-Kang Tu, Justin Keen, and Mark S. Gilthorpe.** 2007. "Do the UK Government's New Quality and Outcomes Framework (QOF) Scores Adequately Measure Primary Care Performance? A Cross-Sectional Survey of Routine Healthcare Data." *BMC Health Services Research*, 7(1): 166.
- Duncan, Craig, Kelvyn Jones, and Graham Moon.** 1996. "Health-Related Behaviour in Context: A Multilevel Modelling Approach." *Social Science & Medicine*, 42(6): 817–830.
- Dusheiko, Mark, Hugh Gravelle, Rowena Jacobs, and Peter Smith.** 2006. "The Effect of Financial Incentives on Gatekeeping Doctors: Evidence from a Natural Experiment." *Journal of Health Economics*, 25(3): 449–478.
- Dyer, C.** 2001. "Bristol Inquiry." *BMJ*, 323(7306): 181–181.
- Eddy, David M.** 1984. "Variations in Physician Practice: The Role of Uncertainty." *Health Affairs*, 3(2): 74–89.
- Eisenberg, J. M.** 1985. "Physician Utilization. The State of Research about Physicians' Practice Patterns." *Medical Care*, 23(5): 461–483.

- Endsley, Scott, Margaret Kirkegaard, and Anthony Linares.** 2005. "Working Together: Communities of Practice in Family Medicine." <https://www.aafp.org/fpm/2005/0100/p28.html?printable=fpm>.
- Epstein, Andrew J., and Sean Nicholson.** 2009. "The Formation and Evolution of Physician Treatment Styles: An Application to Cesarean Sections." *Journal of Health Economics*, 28(6): 1126–1140.
- Erickson, Bonnie H.** 1988. "The Relational Basis of Attitudes." In *Social Structures: A Network Approach*, ed. Barry Wellman, S. D. Berkowitz and Mark Granovetter. CUP Archive.
- Erlangga, Darius, Shehzad Ali, and Karen Bloor.** 2019. "The Impact of Public Health Insurance on Healthcare Utilisation in Indonesia: Evidence from Panel Data." *International Journal of Public Health*, 64(4): 603–613.
- Evans, Elizabeth, Harry Aiking, and Adrian Edwards.** 2011. "Reducing Variation in General Practitioner Referral Rates through Clinical Engagement and Peer Review of Referrals: A Service Improvement Project." *Quality in Primary Care*, 19(4): 263–272.
- Fehr, Beverley.** 1996. *Friendship Processes*. SAGE.
- Ferguson, Julie, Judy Wakeling, and Paul Bowie.** 2014. "Factors Influencing the Effectiveness of Multisource Feedback in Improving the Professional Practice of Medical Doctors: A Systematic Review." *BMC Medical Education*, 14(1): 76.
- Fisher, R. A.** 1919. "The Causes of Human Variability." *The Eugenics Review*, 10(4): 213–220.
- Fisher, R.A.** 1925. *Statistical Methods for Research Workers*. London: Oliver & Boyd.

- Flodgren, Gerd, Mary Ann O'Brien, Elena Parmelli, and Jeremy M. Grimshaw.** 2019. "Local Opinion Leaders: Effects on Professional Practice and Healthcare Outcomes." *Cochrane Database of Systematic Reviews*, , (6).
- Fotaki, Marianna.** 2013. "Is Patient Choice the Future of Health Care Systems?" *International Journal of Health Policy and Management*, 1(2): 121–123.
- Freeman, Tim.** 2002. "Using Performance Indicators to Improve Health Care Quality in the Public Sector: A Review of the Literature." *Health Services Management Research*, 15(2): 126–137.
- Fry, John.** 1988. "General Practice and Primary Health Care 1940s-1980s." <https://www.nuffieldtrust.org.uk/research/general-practice-and-primary-health-care-1940s-1980s>.
- Gaviria, Alejandro, and Steven Raphael.** 2001. "School-Based Peer Effects and Juvenile Behavior." *Review of Economics and Statistics*, 83(2): 257–268.
- Geneau, Robert, Pascale Lehoux, Raynald Pineault, and Paul Lamarche.** 2008. "Understanding the Work of General Practitioners: A Social Science Perspective on the Context of Medical Decision Making in Primary Care." *BMC Family Practice*, 9(1): 12.
- Gillam, Stephen.** 2017. "The Family Doctor Charter: 50 Years On." *British Journal of General Practice*, 67(658): 227–228.
- Gillam, Stephen, and Aloysius Niroshan Siriwardena,** ed. 2011. *Quality and Outcomes Framework: QOF - Transforming General Practice*. Abingdon:Radcliffe Pub.
- Gillam, Stephen, and Nicholas Steel.** 2013. "The Quality and Outcomes Framework—Where Next?" *BMJ*, 346.
- Glover, J. Alison.** 1938. "The Incidence of Tonsillectomy in School Children." *Proceedings of the Royal Society of Medicine*, 31(10): 1219–1236.

- GMS Committee.** 1965. "Charter for General Practice." *Br Med J*, 1(5436): 669–670.
- Goldberg, B. A.** 1984. "The Peer Review Privilege: A Law in Search of a Valid Policy." *American Journal of Law & Medicine*, 10(2): 151–167.
- Goodpastor, Walter A., and Isaac D. Montoya.** 1996. "Motivating Physician Behaviour Change: Social Influence versus Financial Contingencies." *International Journal of Health Care Quality Assurance*, 9(6): 4–9.
- Goodwin, James S., Yu-Li Lin, Siddhartha Singh, and Yong-Fang Kuo.** 2013. "Variation in Length of Stay and Outcomes among Hospitalized Patients Attributable to Hospitals and Hospitalists." *Journal of General Internal Medicine*, 28(3): 370–376.
- Goodwin, Nick, Anna Dixon, Teresa Poole, and Veena Raleigh.** 2011. *Improving the Quality of Care in General Practice: Report of an Independent Inquiry Commissioned by the King's Fund*. London:King's Fund.
- Gosden, Toby, Isobel Bowler, and Matthew Sutton.** 2000. "How Do General Practitioners Choose Their Practice? Preferences for Practice and Job Characteristics." *Journal of Health Services Research & Policy*, 5(4): 208–213.
- Griffiths, M, W E Waters, and E D Acheson.** 1979. "Variation in Hospital Stay after Inguinal Herniorrhaphy." *British Medical Journal*, 1(6166): 787–789.
- Grimes, R M, and S K Moseley.** 1976. "An Approach to an Index of Hospital Performance." *Health Services Research*, 11(3): 288–301.
- Grytten, Jostein, and Rune Sørensen.** 2003. "Practice Variation and Physician-Specific Effects." *Journal of Health Economics*, 22(3): 403–418.
- Günther, Oliver H., Beate Kürstein, Steffi G. Riedel-Heller, and Hans-Helmut König.** 2010. "The Role of Monetary and Nonmonetary Incentives on the Choice of Prac-

- tice Establishment: A Stated Preference Study of Young Physicians in Germany.” *Health Services Research*, 45(1): 212–229.
- Gutacker, Nils, Karen Bloor, Chris Bojke, and Kieran Walshe.** 2018. “Should Interventions to Reduce Variation in Care Quality Target Doctors or Hospitals?” *Health Policy (Amsterdam, Netherlands)*, 122(6): 660–666.
- Guthrie, Bruce, Gary McLean, and Matt Sutton.** 2006. “Workload and Reward in the Quality and Outcomes Framework of the 2004 General Practice Contract.” *British Journal of General Practice*, 56(532): 836–841.
- Hadfield, Stephen J.** 1953. “A Field Survey of General Practice, 1951-2.” *British Medical Journal*, 2(4838): 683–706.
- Ham, Chris.** 1996. “Managed Markets in Health Care: The UK Experiment.” *Health Policy*, 35(3): 279–292.
- Harris, Conrad M., and Glen Scrivener.** 1996. “Fundholders’ Prescribing Costs: The First Five Years.” *BMJ*, 313(7071): 1531–1534.
- Herbst, D., and A. Mas.** 2015. “Peer Effects on Worker Output in the Laboratory Generalize to the Field.” *Science*, 350(6260): 545–549.
- Holman, William L., Richard M. Allman, Monique Sansom, Catarina I. Kiefe, Eric D. Peterson, Kevin J. Anstrom, Steadman S. Sankey, Steve G. Hubbard, Robert G. Sherrill, and for the Alabama CABG Study Group.** 2001. “Alabama Coronary Artery Bypass Grafting Project: Results of a Statewide Quality Improvement Initiative.” *JAMA*, 285(23): 3003.
- Honigsbaum, F.** 1979. “The Division in British Medicine : A History of the Separation of General Practice from Hospital Care, 1911-1968.” PhD diss. London School of Economics and Political Science (University of London).

- Horder, J. P., and G. Swift.** 1979. "The History of Vocational Training for General Practice." *The Journal of the Royal College of General Practitioners*, 29(198): 24–32.
- Howard, David H.** 2006. "Quality and Consumer Choice in Healthcare: Evidence from Kidney Transplantation." *The B.E. Journal of Economic Analysis & Policy*, 5(1).
- Howie, J. G. R., D. Heaney, and M. Maxwell.** 2004. "Quality, Core Values and the General Practice Consultation: Issues of Definition, Measurement and Delivery." *Family Practice*, 21(4): 458–468.
- Hulscher, M E, M Wensing, R P Grol, T van der Weijden, and C van Weel.** 1999. "Interventions to Improve the Delivery of Preventive Services in Primary Care." *American Journal of Public Health*, 89(5): 737–746.
- Hunt, J. H.** 1955. "The Scope and Development of General Practice in Relation to Other Branches of Medicine: A Constructive Review." *The Lancet*, 266(6892): 681–687.
- Iedema, Rick, Shannon Meyerkort, and Les White.** 2005. "Emergent Modes of Work and Communities of Practice." *Health Services Management Research*, 18(1): 13–24.
- Institute of Healthcare Management.** 2018. "Focus on GP Quality Indicators: General Practitioners Committee."
- Institute of Medicine.** 2006. *Performance Measurement: Accelerating Improvement (Pathways to Quality Health Care Series)*. Washington, D.C.:National Academies Press.
- Inui, Thomas S.** 1976. "Improved Outcomes in Hypertension After Physician Tutorials: A Controlled Trial." *Annals of Internal Medicine*, 84(6): 646.
- Ivers, Noah, Gro Jamtvedt, Signe Flottorp, Jane M Young, Jan Odgaard-Jensen, Simon D French, Mary Ann O'Brien, Marit Johansen, Jeremy Grimshaw, and Andrew D Oxman.** 2012. "Audit and Feedback: Effects on Professional Practice and Healthcare Outcomes." *Cochrane Database of Systematic Reviews*.

- Jaffer, Amir K., Daniel J. Brotman, Lori D. Bash, Syed K. Mahmood, Brooke Lott, and Richard H. White.** 2010. "Variations in Perioperative Warfarin Management: Outcomes and Practice Patterns at Nine Hospitals." *The American Journal of Medicine*, 123(2): 141–150.
- Jamtvedt, G., J. M. Young, D. T. Kristoffersen, M. A. Thomson O'Brien, and A. D. Oxman.** 2003. "Audit and Feedback: Effects on Professional Practice and Health Care Outcomes." *The Cochrane Database of Systematic Reviews*, , (3): CD000259.
- Jarvis, William R.** 2012. "What Can Canada Learn from the USA's Experience in Reducing Healthcare-associated Infections?" *Clinical Governance: An International Journal*, 17(2): 149–154.
- Jensen, Michael C., and William H. Meckling.** 1976. "Theory of the Firm: Managerial Behavior, Agency Costs and Ownership Structure." *Journal of Financial Economics*, 3(4): 305–360.
- Johnson, Bradford C., James M. Manyika, and Lareina A. Yee.** 2005. "The next Revolution in Interactions." McKinsey Quarterly.
- Jordan, Michelle E., Holly J. Lanham, Benjamin F. Crabtree, Paul A. Nutting, William L. Miller, Kurt C. Stange, and Reuben R. McDaniel.** 2009. "The Role of Conversation in Health Care Interventions: Enabling Sensemaking and Learning." *Implementation Science*, 4(1): 15.
- Joss, Richard, and Maurice Kogan.** 1995. *Advancing Quality: Total Quality Management in the NHS*. Buckingham ; Bristol, PA, USA:Open University Press.
- Kay, Adrian.** 2002. "The Abolition of the GP Fundholding Scheme: A Lesson in Evidence-Based Policy Making." *The British Journal of General Practice*, 52(475): 141–144.

- Keating, Nancy L., John Z. Ayanian, Paul D. Cleary, and Peter V. Marsden.** 2007. "Factors Affecting Influential Discussions among Physicians: A Social Network Analysis of a Primary Care Practice." *Journal of General Internal Medicine*, 22(6): 794–798.
- Keating, N. L., A. M. Zaslavsky, and J. Z. Ayanian.** 1998. "Physicians' Experiences and Beliefs Regarding Informal Consultation." *JAMA*, 280(10): 900–904.
- Kenny, Nuala P., Karen V. Mann, and Heather MacLeod.** 2003. "Role Modeling in Physicians??? Professional Formation: Reconsidering an Essential but Untapped Educational Strategy:." *Academic Medicine*, 78(12): 1203–1210.
- Krasnik, A., P. P. Groenewegen, P. A. Pedersen, P. von Scholten, G. Mooney, A. Gottschau, H. A. Flierman, and M. T. Damsgaard.** 1990. "Changing Remuneration Systems: Effects on Activity in General Practice." *British Medical Journal*, 300(6741): 1698–1701.
- Lagarde, Mylene, and Duane Blaauw.** 2009. "A Review of the Application and Contribution of Discrete Choice Experiments to Inform Human Resources Policy Interventions." *Human Resources for Health*, 7(1): 62.
- Lask, M.** 1950. "The G.P. at the Crossroads." *British Medical Journal*, 1(4658): 902.
- Lechner, Michael.** 2010. "The Estimation of Causal Effects by Difference-in-Difference Methods." *Foundations and Trends® in Econometrics*, 4(3): 165–224.
- Le Grand, Julian.** 2007. *The Other Invisible Hand: Delivering Public Services through Choice and Competition*. Princeton:Princeton University Press.
- Lester, Helen, and Stephen Campbell.** 2010. "Developing Quality and Outcomes Framework (QOF) Indicators and the Concept of 'QOFability'." *Quality in Primary Care*, 18(2): 103–109.

- Lester, Helen, Deborah J. Sharp, F. D. R. Hobbs, and Mayur Lakhani.** 2006. "The Quality and Outcomes Framework of the GMS Contract: A Quiet Evolution for 2006." *British Journal of General Practice*, 56(525): 244–246.
- Leyland, Alastair H., and Peter P. Groenewegen.** 2003. "Multilevel Modelling and Public Health Policy." *Scandinavian Journal of Public Health*, 31(4): 267–274.
- Li, Jinhu, Anthony Scott, Matthew McGrail, John Humphreys, and Julia Witt.** 2014. "Retaining Rural Doctors: Doctors' Preferences for Rural Medical Workforce Incentives." *Social Science & Medicine*, 121: 56–64.
- Lipman, Toby, Chris Johnstone, Martin Roland, Patricia Wilkie, and David Jewell.** 2005. "So How Was It for You? A Year of the GMS Contract Never Offer GPs Money, They Will Just Take it An Important Step forwards Is the GMS Contract Just for Doctors? Or Do Patients Benefit as Well? Careful with the Unintended Consequences: INTO THE SUNLIT UPLANDS?" *British Journal of General Practice*, 55(514): 396–396.
- Locock, L., S. Dopson, D. Chambers, and J. Gabbay.** 2001. "Understanding the Role of Opinion Leaders in Improving Clinical Effectiveness." *Social Science & Medicine* (1982), 53(6): 745–757.
- Lohr, Kathleen N.,** ed. 1990. *Medicare: A Strategy for Quality Assurance: Volume 1*. Washington (DC): National Academies Press (US).
- Lopez Bernal, James, Steven Cummins, and Antonio Gasparrini.** 2016. "Interrupted Time Series Regression for the Evaluation of Public Health Interventions: A Tutorial." *International Journal of Epidemiology*, dyw098.
- Lopez Bernal, James, Steven Cummins, and Antonio Gasparrini.** 2018. "The Use of Controls in Interrupted Time Series Studies of Public Health Interventions." *International Journal of Epidemiology*, 47(6): 2082–2093.

- Lundborg, Petter.** 2006. "Having the Wrong Friends? Peer Effects in Adolescent Substance Use." *Journal of Health Economics*, 25(2): 214–233.
- Magee, Helen, Lucy-Jane Davis, and Angela Coulter.** 2003. "Public Views on Healthcare Performance Indicators and Patient Choice." *Journal of the Royal Society of Medicine*, 96(7): 338–342.
- Mainz, J.** 2003*a*. "Defining and Classifying Clinical Indicators for Quality Improvement." *International Journal for Quality in Health Care*, 15(6): 523–530.
- Mainz, Jan.** 2003*b*. "Developing Evidence-Based Clinical Indicators: A State of the Art Methods Primer." *International Journal for Quality in Health Care: Journal of the International Society for Quality in Health Care*, 15 Suppl 1: i5–11.
- Mandeville, Kate L., Mylene Lagarde, and Kara Hanson.** 2014. "The Use of Discrete Choice Experiments to Inform Health Workforce Policy: A Systematic Review." *BMC Health Services Research*, 14(1): 367.
- McPherson, K., J. E. Wennberg, O. B. Hovind, and P. Clifford.** 1982. "Small-Area Variations in the Use of Common Surgical Procedures: An International Comparison of New England, England, and Norway." *The New England Journal of Medicine*, 307(21): 1310–1314.
- Meltzer, David, Jeanette Chung, Parham Khalili, Elizabeth Marlow, Vineet Arora, Glen Schumock, and Ron Burt.** 2010. "Exploring the Use of Social Network Methods in Designing Healthcare Quality Improvement Teams." *Social Science & Medicine*, 71(6): 1119–1130.
- Mittman, Brian S., Xenia Tonesk, and Peter D. Jacobson.** 1992. "Implementing Clinical Practice Guidelines: Social Influence Strategies and Practitioner Behavior Change." *QRB - Quality Review Bulletin*, 18(12): 413–422.

- Mohammad, Mosadeghrad Ali.** 2013. "Healthcare Service Quality: Towards a Broad Definition." *International Journal of Health Care Quality Assurance*, 26(3): 203–219.
- Molitor, David.** 2018. "The Evolution of Physician Practice Styles: Evidence from Cardiologist Migration." *American Economic Journal: Economic Policy*, 10(1): 326–356.
- Moretti, Enrico.** 2011. "Social Learning and Peer Effects in Consumption: Evidence from Movie Sales." *The Review of Economic Studies*, 78(1): 356–393.
- Mosadeghrad, Ali Mohammad.** 2012. "A Conceptual Framework for Quality of Care." *Materia Socio-Medica*, 24(4): 251–261.
- Mosadeghrad, Ali Mohammad.** 2014a. "Factors Affecting Medical Service Quality." *Iranian Journal of Public Health*, 43(2): 210–220.
- Mosadeghrad, Ali Mohammad.** 2014b. "Factors Influencing Healthcare Service Quality." *International Journal of Health Policy and Management*, 3(2): 77–89.
- National Health Service.** 2005. "National Quality and Outcomes Framework Statistics for England - 2004-05." <https://digital.nhs.uk/data-and-information/publications/statistical/quality-and-outcomes-framework-achievement-data/national-quality-and-outcomes-framework-statistics-for-england-2004-05>.
- National Health Service.** 2006a. "National Quality and Outcomes Framework Achievement Data - England, 2005-06." <https://digital.nhs.uk/data-and-information/publications/statistical/quality-and-outcomes-framework-achievement-data/national-quality-and-outcomes-framework-achievement-data-england-2005-06>.
- National Health Service.** 2006b. "NHS Employers: Directed Enhanced Services." <https://web.archive.org/web/20060511071457/http://www.nhsemployers.org/primary/primary-893.cfm>.

- National Health Service.** 2010. “Quality and Outcomes Framework - 2009-10.” <https://digital.nhs.uk/data-and-information/publications/statistical/quality-and-outcomes-framework-achievement-data/quality-and-outcomes-framework-2009-10>.
- National Health Service.** 2014. “Quality and Outcomes Framework (QOF) - 2013-14 - NHS Digital.” <https://digital.nhs.uk/data-and-information/publications/statistical/quality-and-outcomes-framework-achievement-prevalence-and-exceptions-data/quality-and-outcomes-framework-qof-2013-14>.
- National Health Service.** 2018a. “QOF 2018/19.” <https://qof.digital.nhs.uk/>.
- National Health Service.** 2018b. “Quality and Outcomes Framework, Achievement, Prevalence and Exceptions Data - 2017-18 [PAS].” <https://digital.nhs.uk/data-and-information/publications/statistical/quality-and-outcomes-framework-achievement-prevalence-and-exceptions-data/2017-18>.
- National Health Service.** 2019. “Quality and Outcomes Framework, Achievement, Prevalence and Exceptions Data 2018-19 [PAS] - NHS Digital.” <https://digital.nhs.uk/data-and-information/publications/statistical/quality-and-outcomes-framework-achievement-prevalence-and-exceptions-data/2018-19-pas>.
- National Institute for Health and Care Excellence.** 2017. “Health and Social Care Directorate Indicator Process Guide.”
- NHS England.** 2019. “2019/20 General Medical Services (GMS) Contract Quality and Outcomes Framework (QOF).”
- NICE.** 2019. “Quality and Outcomes Framework Indicators.” <https://www.nice.org.uk/standards-and-indicators/qofindicators>.
- O’Connor, Gerald T., Stephen K. Plume, Elaine M. Olmstead, Jeremy R. Morton, Christopher T. Maloney, William C. Nugent, Felix Hernandez, Robert**

Clough, Bruce J. Leavitt, Laurence H. Coffin, Charles A. S. Marrin, David Wennberg, John D. Birkmeyer, David C. Charlesworth, David J. Malenka, Hebe B. Quinton, Joseph F. Kasper, Felix Hernandez, Hebe Quinton, Deborah Carey-Johnson, Cynthia M. Downs, Joseph J. Hessel, Robert M. Hoffman, Edward R. Johnson, Helen McKinnon, Cathy Mingo, Craig Pedersen, Wendy Perkins, Matthew L. Rowe, Katrina Sargent, Ted Silver, Peter Ver Lee, Craig Warren, Shelley Barber, Pamela Brown, Richard G. Brandenburg, Steve Colmanaro, Walter D. Gundel, Richard S. Jackson, David Johnson, Ann Laramée, William C. Paganelli, Diane Pappalardo, Daniel S. Raabe, Christopher Terrien, Paul Vaitkus, Matthew Watkins, Virginia Beggs, William Burke, Edward Catherwood, Lawrence J. Daeey, Candis Darcey, Gordon De-foe, Thomas Dodds, Mary Fillinger, Bruce Friedman, Bruce Hettleman, Terry Kneeland, Elizabeth Maislen, Nathaniel Niles, Nancy J. O'Connor, John Robb, William Schults, William F. Sullivan, Jon Wahrenberger, Beth Wolf, Yvon Baribeau, Ann Becker, Craig C. Berry, Kevin Berry, William A. Bradley, S. Cuddy, Robert C. Dewey, Frank Fedele, Louis I. Fink, Erik J. Funk, Alan E. Garstka, Dan Halstead, Michael J. Hearne, J. Beatty Hunter, Alan D. Kaplan, Peggy Lambert, Patrick M. Lawrence, Jeffery Lockhart, Kathy McNeil, Edward Palank, M. Judith Porelle, Donna Pulsifer, Joanne Robichaud, James Schnitz, Shirley Shea, Benjamin M. Westbrook, Thomas P. Wharton, Kirke W. Wheeler, Dee White, Robert B. Keller, David C. Soule, Mary Abbott, Lawrence Adrian, Warren D. Alpern, Eric Anderson, Richard A. Anderson, Linda Banister, Claire Berg, Seth Blank, Carl E. Bredenberg, Michael Brennan, Linda Brewster, David Burkey, Alice Cirillo, Cantwell Clark, Jane Cleaves, Deborah Courtney, Joshua Cutler, Desmond Donegal, Pat Fallo, Daniel Hanley, Jane Kane, Saul Katz, Mirle A. Kellett, Robert Kramer, Costas T. Lambrew, F. Stephen Larned, Chris A. Lutes, Edward R. Now-

- icki, John R. O'Meara, Harold Osher, Patricia Peasley, Cathy Prouty, Reed D. Quinn, Dennis Redfield, Karen Reynolds, Thomas Ryan, Jean Saunders, Alyce Schultz, Susan Seekins, Russell Stogsdill, Paul W. Sweeny, Karen Tolan, Nancy Tooker, Joan F. Tryzelaar, Marie Turcotte, Kathy Viger, Cynthia Westlund, Richard L. White, Wanda Whittet, and Carol Zografos. 1996. "A Regional Intervention to Improve the Hospital Mortality Associated With Coronary Artery Bypass Graft Surgery." *JAMA*, 275(11): 841–846.
- O'Connor, G. T., H. B. Quinton, N. D. Traven, L. D. Ramunno, T. A. Dodds, T. A. Marciniak, and J. E. Wennberg. 1999. "Geographic Variation in the Treatment of Acute Myocardial Infarction: The Cooperative Cardiovascular Project." *JAMA*, 281(7): 627–633.
- Oliver, A. 2009. "Can Financial Incentives Improve Health Equity?" *BMJ*, 339(sep24 2): b3847–b3847.
- Ovretveit, John, and Christina Townsend. 1992. *Health Service Quality: An Introduction to Quality Methods for Health Services*. Oxford:Blackwell Science.
- Oxman, A. D., M. A. Thomson, D. A. Davis, and R. B. Haynes. 1995. "No Magic Bullets: A Systematic Review of 102 Trials of Interventions to Improve Professional Practice." *CMAJ: Canadian Medical Association journal = journal de l'Association medicale canadienne*, 153(10): 1423–1431.
- Ozcan, Yasar A. 2005. *Quantitative Methods in Health Care Management: Techniques and Applications*. John Wiley & Sons.
- Paul-Shaheen, P., J. D. Clark, and D. Williams. 1987. "Small Area Analysis: A Review and Analysis of the North American Literature." *Journal of Health Politics, Policy and Law*, 12(4): 741–809.

- Pencheon, David.** 2007. “The Good Indicators Guide: Understanding How to Use and Choose Indicators.” <https://www.england.nhs.uk/improvement-hub/publication/the-good-indicators-guide-understanding-how-to-use-and-choose-indicators/>.
- Penfold, Robert B., and Fang Zhang.** 2013. “Use of Interrupted Time Series Analysis in Evaluating Health Care Quality Improvements.” *Academic Pediatrics*, 13(6, Supplement): S38–S44.
- Pérez-Cuevas, Ricardo, Hector Guiscafré, Onofre Muñoz, Hortensia Reyes, Patricia Tomé, Vita Libreros, and Gonzalo Gutiérrez.** 1996. “Improving Physician Prescribing Patterns to Treat Rhinopharyngitis. Intervention Strategies in Two Health Systems of Mexico.” *Social Science & Medicine*, 42(8): 1185–1194.
- Pescosolido, Bernice A.** 1992. “Beyond Rational Choice: The Social Dynamics of How People Seek Help.” *American Journal of Sociology*, 97(4): 1096–1138.
- Petchey, R., J. Williams, and M. Baker.** 1997. “Ending up a GP’: A Qualitative Study of Junior Doctors’ Perceptions of General Practice as a Career.” *Family Practice*, 14(3): 194–198.
- Petchey, Roland.** 1995. “Collings Report on General Practice in England in 1950: Unrecognised, Pioneering Piece of British Social Research?” *BMJ*, 311(6996): 40–42.
- Phelps, Charles, and Cathleen Mooney.** 1993. “Variations in Medical Practice Use: Causes and Consequences.”
- Pirrie, Ivan.** 1950. “The G.P. at the Crossroads.” *British Medical Journal*, 2(4671): 169–170.
- Powell, J. Enoch.** 1976. *Medicine and Politics: 1975 and After*. . New ed ed., Tunbridge Wells: Pitman Medical.

- Practice, British Journal of General.** 1985. “Quality Assured.” *The Journal of the Royal College of General Practitioners*, 35(278): 411–411.
- Proctor, and Wright.** 1998. “Can Services Marketing Concepts Be Applied to Health Care?” *Journal of Nursing Management*, 6(3): 147–153.
- Pronovost, Peter J, and Daniel W Hudson.** 2012. “Improving Healthcare Quality through Organisational Peer-to-Peer Assessment: Lessons from the Nuclear Power Industry.” *BMJ Quality & Safety*, 21(10): 872–875.
- Pronovost, Peter J., Sam R. Watson, Christine A. Goeschel, Robert C. Hyzy, and Sean M. Berenholtz.** 2016. “Sustaining Reductions in Central Line–Associated Bloodstream Infections in Michigan Intensive Care Units: A 10-Year Analysis.” *American Journal of Medical Quality*, 31(3): 197–202.
- Puntanen, Simo, and George P. H. Styan.** 1989. “The Equality of the Ordinary Least Squares Estimator and the Best Linear Unbiased Estimator.” *The American Statistician*, 43(3): 153–161.
- Raleigh, Veena, and Catherine Foot.** 2010. *Getting the Measure of Quality: Opportunities and Challenges*. London:King’s Fund.
- Rawlins, Michael.** 1999. “In Pursuit of Quality: The National Institute for Clinical Excellence.” *The Lancet*, 353(9158): 1079–1082.
- Reid, Constance.** 1998. “Neyman—from Life.” In *Neyman.* , ed. Constance Reid, 1–293. New York, NY:Springer.
- Roland, Martin.** 2007. “The Quality and Outcomes Framework: Too Early for a Final Verdict.” *The British Journal of General Practice*, 57(540): 525–527.
- Roland, Martin, and Bruce Guthrie.** 2016. “Quality and Outcomes Framework: What Have We Learnt?..” *BMJ*, i4060.

- Roland, Martin, and Stephen Campbell.** 2014. “Successes and Failures of Pay for Performance in the United Kingdom.” *New England Journal of Medicine*, 370(20): 1944–1949.
- Rowsell, R., M. Morgan, and J. Sarangi.** 1995. “General Practitioner Registrars’ Views about a Career in General Practice.” *British Journal of General Practice*, 45(400): 601–604.
- Royal College of General Practitioners.** 2019. “History of the College.” <https://www.rcgp.org.uk/about-us/the-college/who-we-are/history-heritage-and-archive/history-of-the-college.aspx>.
- Rubenstein, Lisa V., Brian S. Mittman, Elizabeth M. Yano, and Cynthia D. Mulrow.** 2000. “From Understanding Health Care Provider Behavior to Improving Health Care: The QUERI Framework for Quality Improvement.” *Medical Care*, 38(6): I129–I141.
- Rubin, Donald B.** 1974. “Estimating Causal Effects of Treatments in Randomized and Nonrandomized Studies.” *Journal of Educational Psychology*, 66(5): 688–701.
- Rubin, Donald B.** 2005. “Causal Inference Using Potential Outcomes: Design, Modeling, Decisions.” *Journal of the American Statistical Association*, 100(469): 322–331.
- Sacerdote, B.** 2001. “Peer Effects with Random Assignment: Results for Dartmouth Roommates.” *The Quarterly Journal of Economics*, 116(2): 681–704.
- Sanguinetti, Harold H.** 1950. “The G.P. at the Crossroads.” *British Medical Journal*, 1(4664): 1270–1271.
- Santos, Rita, Hugh Gravelle, and Carol Propper.** 2017. “Does Quality Affect Patients’ Choice of Doctor? Evidence from England.” *The Economic Journal*, 127(600): 445–494.
- Sargeant, Joan, Jocelyn Lockyer, Karen Mann, Eric Holmboe, Ivan Silver, Heather Armson, Erik Driessen, Tanya MacLeod, Wendy Yen, Kathryn**

- Ross, and Mary Power.** 2015. “Facilitated Reflective Performance Feedback: Developing an Evidence- and Theory-Based Model That Builds Relationship, Explores Reactions and Content, and Coaches for Performance Change (R2C2).” *Academic Medicine*, 90(12): 1698–1706.
- Schuster, Mark A., Elizabeth A. McGlynn, and Robert H. Brook.** 1998. “How Good Is the Quality of Health Care in the United States?” *The Milbank Quarterly*, 76(4): 517–563.
- Scott, Anthony.** 2001. “Eliciting GPs’ Preferences for Pecuniary and Non-Pecuniary Job Characteristics.” *Journal of Health Economics*, 20(3): 329–347.
- Scott, Anthony, Julia Witt, John Humphreys, Catherine Joyce, Guyonne Kalb, Sung-Hee Jeon, and Matthew McGrail.** 2013. “Getting Doctors into the Bush: General Practitioners’ Preferences for Rural Location.” *Social Science & Medicine*, 96: 33–44.
- Secretary of State for Health.** 1998. “A First Class Service.”
- Sharma, Anita.** 2017. *Maximising Quality and Outcomes Framework Quality Points: The QOF Clinical Domain*.
- Shekelle, Paul.** 2003. “New Contract for General Practitioners: A Bold Initiative to Improve Quality of Care, but Implementation Will Be Difficult.” *BMJ*, 326(7387): 457–458.
- Shekelle, Paul G.** 2013. “Quality Indicators and Performance Measures: Methods for Development Need More Standardization.” *Journal of Clinical Epidemiology*, 66(12): 1338–1339.
- Shortt, SED.** 2003. “General Practice Fundholding in the United Kingdom.” *Canadian Family Physician*, 49(March 2003).

- Shover, Chelsea L., Corey S. Davis, Sanford C. Gordon, and Keith Humphreys.** 2019. "Association between Medical Cannabis Laws and Opioid Overdose Mortality Has Reversed over Time." *Proceedings of the National Academy of Sciences*, 201903434.
- Smith, Denis.** 2002. "Not by Error, But by Design - Harold Shipman and the Regulatory Crisis for Health Care." *Public Policy and Administration*, 17(4): 55–74.
- Starkweather, D. B., L. Gelwicks, and R. Newcomer.** 1975. "Delphi Forecasting of Health Care Organization." *Inquiry: A Journal of Medical Care Organization, Provision and Financing*, 12(1): 37–46.
- Stelfox, Henry T., and Sharon E. Straus.** 2013. "Measuring Quality of Care: Considering Measurement Frameworks and Needs Assessment to Guide Quality Indicator Development." *Journal of Clinical Epidemiology*, 66(12): 1320–1327.
- Stokes, Tim.** 2014. "QOF Changes Are Ground-Breaking and Will Need Monitoring." <https://www.guidelinesinpractice.co.uk/non-clinical-best-practice/qof-changes-are-ground-breaking-and-will-need-monitoring/347592.article>.
- Strong, P, and J Robinson.** 1990. *The NHS - Under New Management*. . 1st Paperback Printing edition ed., Milton Keynes England ; Philadelphia:Open University Press.
- Sutcliffe, Daniel, Helen Lester, John Hutton, and Tim Stokes.** 2012. "NICE and the Quality and Outcomes Framework (QOF) 2009-2011." *Quality in Primary Care*, 20(1): 47–55.
- Taylor, Stephen James Lake.** 1954. *Good General Practice. A Report of a Survey by Stephen Taylor*. Oxford University Press.
- Van den Hombergh, P., R. Grol, H. van den Hoogen, and W. van den Bosch.** 1999. "Practice Visits as a Tool in Quality Improvement: Mutual Visits and Feedback

- by Peers Compared with Visits and Feedback by Non-Physician Observers.” *Quality and Safety in Health Care*, 8(3): 161–166.
- van de Vijzel, Aart R., Richard Heijink, and Maarten Schipper.** 2015. “Has Variation in Length of Stay in Acute Hospitals Decreased? Analysing Trends in the Variation in LOS between and within Dutch Hospitals.” *BMC Health Services Research*, 15(1): 438.
- Veninga, C. C. M, P Denig, R Zwaagstra, and F. M Haaijer-Ruskamp.** 2000. “Improving Drug Treatment in General Practice.” *Journal of Clinical Epidemiology*, 53(7): 762–772.
- Wachter, Robert M.** 2013. “Personal Accountability in Healthcare: Searching for the Right Balance.” *BMJ Quality & Safety*, 22(2): 176–180.
- Webb, R., and D. Hannay.** 1996. “Career Choices of Trainees in General Practice.” *BMJ*, 312(7026): 314–314.
- Wenger, Etienne.** 1999. *Communities of Practice: Learning, Meaning, and Identity*. Cambridge University Press.
- Wennberg, J., and null Gittelsohn.** 1973. “Small Area Variations in Health Care Delivery.” *Science (New York, N.Y.)*, 182(4117): 1102–1108.
- Wennberg, John, and Alan Gittelsohn.** 1982. “Variations in Medical Care among Small Areas.” *Scientific American*, 246(4): 120–135.
- Wensing, M., T. van der Weijden, and R. Grol.** 1998. “Implementing Guidelines and Innovations in General Practice: Which Interventions Are Effective?” *British Journal of General Practice*, 48(427): 991–997.
- Westert, Gert P., and Peter P. Groenewegen.** 1999. “Medical Practice Variations: Changing the Theoretical Approach.” *Scandinavian Journal of Public Health*, 27(3): 173–180.

- Westert, G.P.** 1992. *Variation in use of hospital care*. Van Gorcum.
- Westert, G. P., A. P. Nieboer, and P. P. Groenewegen.** 1993. "Variation in Duration of Hospital Stay between Hospitals and between Doctors within Hospitals." *Social Science & Medicine* (1982), 37(6): 833–839.
- Wilkie, Veronica.** 2014. "British General Practice: Another Collings Moment?" *BMJ*, 349.
- Williams, Simon J., Michael Calnan, Sarah L. Cant, and Joanne Coyle.** 1993. "All Change in the NHS? Implications of the NHS Reforms for Primary Care Prevention." *Sociology of Health and Illness*, 15(1): 43–67.
- Wollersheim, H., R. Hermens, M. Hulscher, J. Braspenning, M. Ouwens, J. Schouten, H. Marres, R. Dijkstra, and R. Grol.** 2007. "Clinical Indicators: Development and Applications." *The Netherlands Journal of Medicine*, 65(1): 15–22.
- Wordsworth, Sarah, Diane Skåtun, Anthony Scott, and Fiona French.** 2004. "Preferences for General Practice Jobs: A Survey of Principals and Sessional GPs." *The British Journal of General Practice*, 54(507): 740–746.
- World Health Organization, OECD, and International Bank for Reconstruction and Development/The World Bank.** 2018. *Delivering Quality Health Services: A Global Imperative for Universal Health Coverage*. World Health Organization.
- Yong, Jongsay, Anthony Scott, Hugh Gravelle, Peter Sivey, and Matthew McGrail.** 2018. "Do Rural Incentives Payments Affect Entries and Exits of General Practitioners?" *Social Science & Medicine*, 214: 197–205.
- Zohoori, Namvar, and David A. Savitz.** 1997. "Econometric Approaches to Epidemiologic Data: Relating Endogeneity and Unobserved Heterogeneity to Confounding." *Annals of Epidemiology*, 7(4): 251–257.