

# Multilateral Security

*Privacy is a transient notion. It started when people stopped believing that God could see everything and stopped when governments realised there was a vacancy to be filled.*

— Roger Needham

*You have zero privacy anyway. Get over it.*

— Scott Mcnealy

## 9.1 Introduction

---

Often our goal is not to prevent information flowing ‘down’ a hierarchy but to prevent it flowing ‘across’ between departments. Relevant applications range from healthcare to national intelligence, and include most applications where the privacy of individual customers’, citizens’ or patients’ data is at stake. They account for a significant proportion of information processing systems but their protection is often poorly designed and implemented. This has led to a number of expensive fiascos.

The basic problem is that if you centralise systems containing sensitive information, you risk creating a more valuable asset and simultaneously giving more people access to it. This is now a pressing problem in the world of ‘Web 2.0’ as online applications amass petabytes of people’s private information. And it’s not just Google Documents; a number of organisations plan to warehouse your medical records online. Microsoft has announced HealthVault, which will let your doctors store your medical records online in a data centre and give you some control over access; other IT firms have broadly similar plans. Yet privacy activists point out that however convenient this

may be in an emergency, it gives access to insurance companies, government agencies and anyone else who comes along with a court order [1332]. So what are the real issues with such systems, should they be built, if so how should we protect them, and are there any precedents from which we can learn?

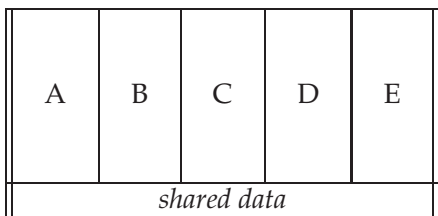
One lesson comes from banking. In the old days, a private investigator who wanted copies of your bank statements had to subvert someone at the branch where your account was kept. But after banks hooked all their branches up online in the 1980s, they typically let any teller enquire about any customer's account. This brought the convenience of being able to cash a check when you are out of town; but it's also meant that private eyes buy and sell your bank statements for a few hundred dollars. They only have to corrupt one employee at each bank, rather than one at each branch. Another example comes from the UK Inland Revenue, the tax collection office; staff were caught making improper access to the records of celebrities, selling data to outsiders, and leaking income details in alimony cases [129].

In such systems, a typical requirement will be to stop users looking at records belonging to a different branch, or a different geographical region, or a different partner in the firm — except under strict controls. Thus instead of the information flow control boundaries being horizontal as we saw in the Bell-LaPadula model as in Figure 9.1, we instead need the boundaries to be mostly vertical, as shown in Figure 9.2.

These lateral information flow controls may be organizational, as in an intelligence organization which wants to keep the names of agents working in one foreign country secret from the department responsible for spying on another. They may be privilege-based, as in a law firm where different clients' affairs, and the clients of different partners, must be kept separate. They may even be a mixture of the two, as in medicine where patient confidentiality



**Figure 9.1:** Multilevel security



**Figure 9.2:** Multilateral security

is based in law on the rights of the patient but usually enforced by limiting medical record access to a particular hospital department.

The control of lateral information flows is a very general problem, of which we'll use medicine as a clear and well-studied example. The problems of medical systems are readily understandable by the nonspecialist and have considerable economic and social importance. Much of what we have to say about them goes across with little or no change to the practice of other professions, and to government applications where access to particular kinds of classified data are restricted to particular teams or departments.

One minor problem we face is one of terminology. Information flow controls of the type we're interested in are known by a number of different names; in the U.S. intelligence community, for example, they are known as *compartmented security* or *compartmentation*. We will use the European term *multilateral security* as the healthcare application is bigger than intelligence, and as the term also covers the use of techniques such as anonymity — the classic case being de-identified research databases of medical records. This is an important part of multilateral security. As well as preventing overt information flows, we also have to prevent information leakage through, for example, statistical and billing data which get released.

The use of de-identified data has wider applicability. Another example is the processing of census data. In general, the relevant protection techniques are known as *inference control*. Despite occasional differences in terminology, the problems facing the operators of census databases and medical research databases are very much the same.

## **9.2 Compartmentation, the Chinese Wall and the BMA Model**

---

There are (at least) three different models of how to implement access controls and information flow controls in a multilateral security model. These are compartmentation, used by the intelligence community; the *Chinese Wall* model, which describes the mechanisms used to prevent conflicts of interest in professional practice; and the *BMA model*, developed by the British Medical Association to describe the information flows permitted by medical ethics. Each of these has potential applications outside its initial field.

### **9.2.1 Compartmentation and the Lattice Model**

For many years, it has been standard practice in the United States and allied governments to restrict access to information by the use of codewords as well as classifications. The best documented example is the codeword *Ultra* in World

War 2, which referred to British and American decrypts of German messages enciphered using the Enigma cipher machine. The fact that the Enigma had been broken was so important that it was worth protecting at almost any cost. So Ultra clearances were given to only a small number of people — in addition to the cryptanalysts and their support staff, the list included the Allied leaders, their senior generals, and hand-picked analysts. No-one who had ever held an Ultra clearance could be placed at risk of capture; and the intelligence could never be used in such a way as to let Hitler suspect that his principal cipher had been broken. Thus when Ultra told of a target, such as an Italian convoy to North Africa, the Allies would send over a plane to ‘spot’ it and report its position by radio an hour or so before the attack. This policy was enforced by special handling rules; for example, Churchill got his Ultra summaries in a special dispatch box to which he had a key but his staff did not. Because such special rules may apply, access to a codeword is sometimes referred to as an *indoctrination* rather than simply a clearance. (Ultra security is described in Kahn [677] and in Welchman [1336].)

Much the same precautions are in place today to protect information whose compromise could expose intelligence sources or methods, such as agent names, cryptanalytic successes, the capabilities of equipment used for electronic eavesdropping, and the performance of surveillance satellites. The proliferation of codewords results in a large number of compartments, especially at classification levels above Top Secret.

One reason for this is that classifications are inherited by derived work; so a report written using sources from ‘Secret Desert Storm’ and ‘Top Secret Umbra’ can in theory only be read by someone with a clearance of ‘Top Secret’ and membership of the groups ‘Umbra’ and ‘Desert Storm’. Each combination of codewords gives a compartment, and some intelligence agencies have over a million active compartments. Managing them is a significant problem. Other agencies let people with high level clearances have relatively wide access. But when the control mechanisms fail, the result can be disastrous. Aldrich Ames, a CIA officer who had accumulated access to a large number of compartments by virtue of long service and seniority, and because he worked in counterintelligence, was able to betray almost the entire U.S. agent network in Russia.

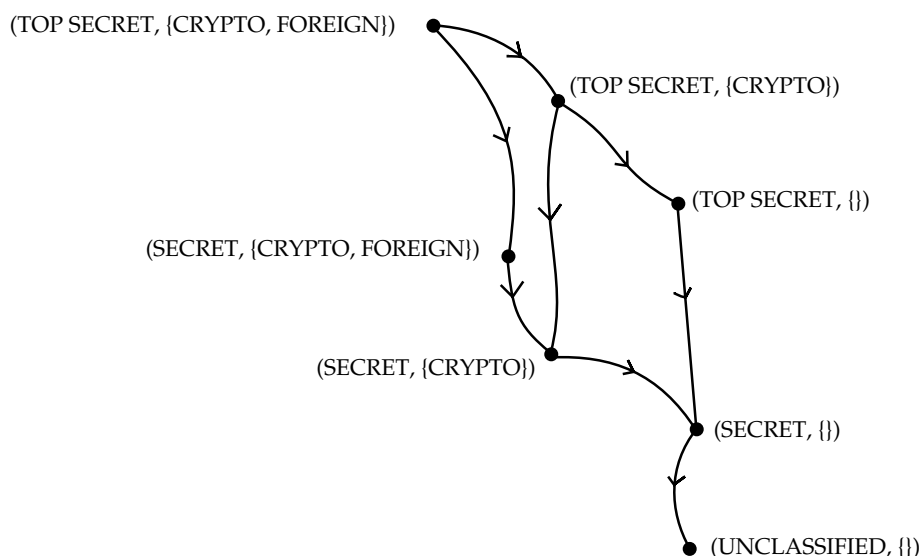
Codewords are in effect a pre-computer way of expressing access control groups, and can be dealt with using a variant of Bell-LaPadula, called the *lattice model*. Classifications together with codewords form a lattice — a mathematical structure in which any two objects  $A$  and  $B$  can be in a dominance relation  $A > B$  or  $B > A$ . They don’t have to be:  $A$  and  $B$  could simply be incomparable (but in this case, for the structure to be a lattice, they will have a least upper bound and a greatest lower bound). As an illustration, suppose we have a codeword, say ‘Crypto’. Then someone cleared to ‘Top Secret’ would be entitled to read files classified ‘Top Secret’ and ‘Secret’, but would have no

access to files classified ‘Secret Crypto’ unless he also had a crypto clearance. This can be expressed as shown in Figure 9.3.

In order for information systems to support this, we need to distill the essence of classifications, clearances and labels into a security policy that we can then use to drive security targets, implementation, and evaluation. As it happens, the Bell-LaPadula model goes across more or less unchanged. We still have information flows between High and Low as before, where High is a compartment that dominates Low. If two nodes in a lattice are incompatible — as with ‘Top Secret’ and ‘Secret Crypto’ in the above diagram — then there should be no information flow between them at all.

In fact, the lattice and Bell-LaPadula models are essentially equivalent, and were developed at the same time.

- Roger Schell, Peter Downey, and Gerald Popek of the U.S. Air Force produced an early lattice model in 1972 [1119].
- A Cambridge PhD thesis by Jeffrey Fenton included a representation in which labels were managed using a matrix [464].
- About this time, the Pentagon’s World Wide Military Command and Control System (WWMCCS) used a primitive lattice model, but without the \*-property. The demonstration that a fielded, critical, system handling Top Secret data was vulnerable to attack by Trojans caused some consternation [1118]. It meant that all users had to be cleared to the highest level of data in the machine.



**Figure 9.3:** A lattice of security labels

- Kenneth Walter, Walter Ogden, William Rounds, Frank Bradshaw, Stan Ames, and David Shumway of Case Western University produced a more advanced lattice model as well as working out a lot of the problems with file and directory attributes, which they fed to Bell and LaPadula [1312, 1313]<sup>1</sup>.
- Finally, the lattice model was systematized and popularized by Denning [368].

Most products built for the multilevel secure market can be reused in compartmented mode. But, in practice, these products are not as effective as one might like. It is easy to use a multilevel operating system to keep data in different compartments separate — just give them incompatible labels ('Secret Tulip', 'Secret Daffodil', 'Secret Crocus', ...). But the operating system has now become an isolation mechanism, rather than a sharing mechanism; the real problem is how to control information sharing.

One solution is to impose least upper bounds in the lattice using some algorithm. An example comes from the system used by the government of Saudi Arabia to manage the Haj, the annual pilgrimage to Mecca [606]. While most compartments are by default Confidential, the combination of data from different compartments is Secret. Thus 'Haj-visas' and 'Gov-guest' are confidential, but their combination is Secret.

In many intelligence systems, where the users are already operating at the highest level of clearance, data owners don't want a further classification level at which everything is visible. So data derived from two compartments effectively creates a third compartment using the lattice model. The proliferation of millions of compartments is complex to manage and can be intertwined with applications. So a more common solution is to use a standard multilevel product, such as a mail guard, to ensure that 'untrustworthy' email goes to filters. But now the core of the trusted computing base consists of the filters rather than the guard.

Worse, the guard may lose some of the more important functionality of the underlying operating system. For example, the Standard Mail Guard [1193] was built on top of an operating system called LOCK whose basic mechanism is type enforcement, as described in the previous chapter. Later versions of LOCK support role-based access control, which would be a more appropriate mechanism to manage the relationships between compartments directly [612]. Using it merely as a platform to support BLP may have been wasteful.

In general, the real problems facing users of intelligence systems have to do with combining data in different compartments, and downgrading it after

<sup>1</sup>Walter and his colleagues deserve more credit than history has given them. They had the main results first [1312] but Bell and LaPadula had their work heavily promoted by the U.S. Air Force. Fenton has also been largely ignored, not being an American.

sanitization. Multilevel and lattice security models offer little help here. Indeed one of the biggest problem facing the U.S. intelligence community since 9/11 is how to handle search over systems with many compartments. A search done over many agencies' databases can throw up results with many codewords attached; if this were to be aggregated in one place, then that place would in effect possess all clearances. What new systems do is to send out search queries bound with the clearance of the user: 'Show me everything that matches Uzbek and Peshawar and weapons and motorcycle, and can be seen by someone with a clearance of Top Secret Umbra'. Here, local labels just get in the way; but without them, how do you forestall a future Aldritch Ames?

There's a also sobering precedent in the Walker spy case. There, an attempt to keep naval vessels in compartments just didn't work, as a ship could be sent anywhere on no notice, and for a ship to be isolated with no local key material was operationally unacceptable. So the U.S. Navy's 800 ships all ended up with the same set of cipher keys, which got sold to the Russians [587].

### 9.2.2 The Chinese Wall

The second model of multilateral security is the Chinese Wall model, developed by Brewer and Nash [224]. Its name comes from the fact that financial services firms from investment banks to accountants have internal rules designed to prevent conflicts of interest, which they call Chinese Walls.

The model's scope is wider than just finance. There are many professional and services firms whose clients may be in competition with each other: software vendors and advertising agencies are other examples. A typical rule is that 'a partner who has worked recently for one company in a business sector may not see the papers of any other company in that sector'. So once an advertising copywriter has worked on (say) the Shell account, he will not be allowed to work on any other oil company's account for some fixed period of time.

The Chinese Wall model thus features a mix of free choice and mandatory access control: a partner can choose which oil company to work for, but once that decision is taken his actions in that sector are completely constrained. It also introduces the concept of *separation of duty* into access control; a given user may perform transaction A or transaction B, but not both.

Part of the attraction of the Chinese Wall model to the security research community comes from the fact that it can be expressed in a way that is fairly similar to Bell-LaPadula. If we write, for each object  $c$ ,  $y(c)$  for  $c$ 's company and  $x(c)$  for  $c$ 's conflict-of-interest class, then like BLP it can be expressed in two properties:

- The *simple security property*: a subject  $s$  has access to  $c$  if and only if, for all  $c'$  which  $s$  can read, either  $y(c) \notin x(c')$  or  $y(c) = y(c')$

- The *\*-property*: a subject  $s$  can write to  $c$  only if  $s$  cannot read any  $c'$  with  $x(c') \neq \emptyset$  and  $y(c) \neq y(c')$ .

The Chinese Wall model made a seminal contribution to the theory of access control. It also sparked a debate about the extent to which it is consistent with the BLP tranquility properties, and some work on the formal semantics of such systems (see, for example, Foley [480] on the relationship with non-interference). There are also some interesting new questions about covert channels. For example, could an oil company find out whether a competitor which used the same investment bank was planning a bid for a third oil company, by asking which specialists were available for consultation and noticing that their number had dropped suddenly?

In practice, however, Chinese Walls still get implemented using manual methods. One large software consultancy has each of its staff maintain an ‘unclassified’ curriculum vitae containing entries that have been sanitized and agreed with the customer. A typical entry might be:

**Sep 97 — Apr 98:** consulted on security requirements for a new branch accounting system for a major U.S. retail bank

This is not the only control. A consultant’s manager should be aware of possible conflicts and not forward the CV to the client if in doubt; if this fails the client can spot potential conflicts himself from the CV; and if this also fails then the consultant is duty bound to report any potential conflicts as soon as they appear.

### 9.2.3 The BMA Model

Perhaps the most important, interesting and instructive example of multilateral security is found in medical information systems. The healthcare sector spends a much larger share of national income than the military in developed countries, and although hospitals are still less automated, they are catching up fast. A 2006 study for the U.S. Department of Health and Human Services (DHHS) showed that investments in health IT were recouped in from three to thirteen years, and could make health care safer as well as more efficient [1160].

Healthcare safety and (especially) privacy have become hot-button issues in many countries. In the USA, the Health Insurance Portability and Accountability Act (HIPAA) was passed by Congress in 1996 following a number of privacy failures. In one notorious case, Mark Farley, a convicted child rapist working as an orthopedic technician at Newton-Wellesley Hospital in Newton, Massachusetts, was caught using a former employee’s password to go through the records of 954 patients (mostly young females) to get the phone numbers of girls to whom he then made obscene phone calls [223]. He ended up doing jail time, and the Massachusetts senator Edward Kennedy was one of HIPAA’s

sponsors. There are many more incidents of a less dramatic nature. Also in 1995–96, the UK government attempted to centralise all medical records, which led to a confrontation with the British Medical Association (BMA). The BMA hired me to devise a policy for safety and privacy of clinical information, which I'll discuss below.

The controversy continued. In the late 1990s, a project in Iceland to build a national medical database incorporating not just medical records but also genetic and genealogical data, so that inherited diseases can be tracked across generations, caused an uproar. Eleven percent of the population opted out; eventually the Icelandic Supreme Court decided that the database had to be opt-in rather than opt-out, and now about half the population participate.

In 2002, President Bush rewrote and relaxed the HIPAA regulations, known as the 'Privacy Rule'; this was followed by further 'administrative simplification' in 2006. The U.S. situation is now that, although medical data must still be protected in hospitals, clinics and insurers, its use outside the immediate care setting (for example, by researchers, employers and welfare agencies) is outside the regulations and so much less controlled. No-one's completely happy: health privacy advocates consider the regime to be quite inadequate; hospitals complain that it adds unnecessarily to their costs; and patient advocates note that HIPAA is often used by hospital staff as an excuse to be unhelpful [560]. At the time of writing (2007), Atlanta's Piedmont Hospital has just become the first institution in the USA to be audited for compliance with the security and privacy regulations, which came into force in 2005. This audit covered topics from physical and logical access to systems and data through Internet usage to violations of security rules by employees, and helped many other healthcare providers decide to invest in encryption and other protection technologies [1295]. In addition, the Government Accountability Office (GAO) has just reported that the DHHS needs to do a lot more to ensure patient privacy, particularly by defining an overall strategy for privacy and by adopting milestones for dealing with nationwide health data exchange (which is not just a matter of inadequate technical protection but also of varying state laws) [735].

In various European countries, there have been debates about the safety and privacy tradeoffs involved with emergency medical information. The Germans put data such as current prescriptions and allergies on the medical insurance card that residents carry; other countries have held back from this, reasoning that if data currently held on a human-readable MedAlert bracelet, such as allergies, are moved to a machine-readable device such as a smartcard, then there is a risk to patients who fall ill in locations where there is no reader available, such as on an airplane or a foreign holiday. In the UK, the government is creating a 'summary care record' of prescriptions and allergies that will be kept on a central database and will be available to many health-care workers, from emergency room clinicians to paramedics and the operators of out-of-hours medical helpline services. One problem is that a patient's current

medications often reveal highly sensitive information — such as treatment for HIV, depression or alcoholism — and making such information available to hundreds of thousands of people carries substantial risks of abuse. Patients have been offered the right to opt out of this system.

There have also been debates about privacy and ethical issues relating to secondary uses of medical information, such as in research. First, there are worries about privacy failures, for example, when a research professor loses a laptop containing the records of millions of patients. Although records used in research often have names and addresses removed, it is a seriously hard job to de-identify records properly; I'll discuss this in detail below. Second, there are ethics issues related to consent. For example, a devout Catholic woman might object to her gynaecological data being used to develop a better morning-after pill. Third, there are economic issues; if my data get used to develop a drug from which a company makes billions of dollars, shouldn't I get a share?

The protection of medical information is thus an interesting case history for the security engineer. It has a lot of rich and complex tradeoffs; it's important to all of us; and it's frequently in the news.

Medical privacy is also a model for protecting personal information of other kinds, such as the information held on individual customers by companies and government agencies. In all European countries (and in many others, such as Canada and Australia) there are *data protection* laws that restrict the dissemination of such data. I'll discuss data protection law in Part III; for present purposes, it's enough to note that some classes of data (affecting health, sexual behavior, political activity and religious belief) the *data subject* must either consent to information sharing, or have a right of veto, or there must be a specific law that permits sharing for the public interest in circumstances that are well enough defined for the data subject to predict them. This raises the issue of how one can construct a security policy in which the access control decisions are taken not by a central authority (as in Bell-LaPadula) or by the system's users (as in discretionary access control) but by the data subjects.

Let's look first at the access control aspects.

### 9.2.3.1 The Threat Model

The main threat to medical privacy is abuse of authorised access by insiders, and the most common threat vector is social engineering. The typical attack comes from a private detective who phones a doctor's office or health insurer with a plausible tale:

Hello, this is Dr Burnett of the cardiology department at the Conquest Hospital in Hastings. Your patient Sam Simmonds has just been admitted here in a coma, and he has a funny looking ventricular arrhythmia. Can you tell me if there's anything relevant in his record?

This kind of attack is usually so successful that in both the USA and the UK there are people who earn their living doing it [411]. (It's not restricted to health records — in June 2000, Tony Blair's fundraiser Lord Levy was acutely embarrassed after someone called the tax office pretending to be him and found out that he'd only paid £5000 in tax the previous year [1064]. But the medical context is a good one in which to discuss it.)

As I mentioned briefly in Chapter 2, an experiment was done in the UK in 1996 whereby the staff at a health authority (a government-owned insurer that purchases health care for a district of several hundred thousand people) were trained to screen out false-pretext telephone calls. The advice they were given is described in [36] but the most important element of it was that they were to always call back — and not to a number given by the caller, but to the number in the phone book for the hospital or other institution where the caller claimed to work. It turned out that some thirty telephone enquiries a week were bogus.

Such *operational security* measures are much more important than most technical protection measures, but they are difficult. If everyone was as unhelpful as intelligence-agency staff are trained to be, the world would grind to a halt. And the best staff training in the world won't protect a system where too many people see too much data. There will always be staff who are careless or even crooked; and the more records they can get, the more harm they can do. Also, organisations have established cultures; we have been simply unable to embed even lightweight operational-security measures on any scale in healthcare, simply because that's not how people work. Staff are focussed on delivering care rather than questioning each other. The few real operational improvements in the last few years have all followed scares; for example, maternity units in Britain now have reasonable entry controls, following incidents in which babies were stolen from nurseries. Also, geriatric wards are often locked to stop demented patients from wandering off. However, most hospital wards are completely open; anyone can wander in off the street to visit their relatives, and the clinical benefits of frequent visits outweigh the occasional violent incidents. PCs are left unattended and logged on to the hospital network. Recently, a health IT investment programme in the UK has tried to standardise access control and issued clinical staff with smartcards to log on to hospital systems; but since logging off as Nurse Jones and on again as Nurse Smith takes several seconds, staff don't bother.

A more general problem is that even where staff behave ethically, a lack of technical understanding — or, as we might more properly describe it, poor security usability — causes leaks of personal information. Old PCs sold on the second hand market or given to schools often have recoverable data on the hard disk; most people are unaware that the usual 'delete' command does not remove the file, but merely marks the space it occupies as re-usable. A PC sold on the second hand market by investment bank Morgan Grenfell

Asset Management had recoverable files containing the financial dealings of ex-Beatle Paul McCartney [254]: there have been similar problems with old health records. Equipment also gets stolen: some 11% of UK family doctors have experienced the theft of a practice PC, and in one case two prominent society ladies were blackmailed over terminations of pregnancy following such a theft [37]. The UK government response to this threat is to try to persuade family doctors to move to 'hosted' systems, where the practice data are kept on regional server farms; but it's quite unclear that there's a net privacy gain. Data theft may be harder, but once data are centralised you can expect access creep; more and more public agencies will come up with arguments why they need access to the data. Even if all the access cases are individually sound, the net effect over time can be quite destructive of privacy.

The fundamental problem is this. The likelihood that a resource will be abused depends on its value and on the number of people who have access to it. Aggregating personal information into large databases increases both these risk factors at the same time. Put simply, we can live with a situation in which a doctor's receptionist has access to 2,000 patients' records: there will be abuse from time to time, but at a tolerably low level. However, if the receptionists of the 5,000 family doctors who might work with a large American HMO, or in one of the five regions of England's National Health Service, all have access to the records of maybe ten million patients, then abuse becomes likely. It only takes one insider who learns to walk up to a PC that's logged on using someone else's smartcard, read a file, and pass the information on to a private eye in exchange for cash. It's not just doctors; in England, each region has tens of thousands of people with access, from nurses and programmers and receptionists to drivers and caterers and cleaners. Many of the staff are temporary, many are foreign, and many are earning close to the minimum wage. And privacy issues aren't limited to organizations that treat patients directly: some of the largest collections of personal health information are in the hands of health insurers and research organizations. I'll discuss their special problems below in section 9.3.

In such an environment, lateral information flow controls are required. A good example of what can go wrong without them comes from an early UK hospital system whose designers believed that for reasons of safety, all staff should have access to all records. This decision was influenced by lobbying from geriatricians and pediatricians, whose patients are often treated by a number of specialist departments in the hospital. They were frustrated by the incompatibilities between different departmental systems. The system was fielded in 1995 in Hampshire, where the then health minister Gerry Malone had his parliamentary seat. The system made all lab tests performed for local doctors at the hospital's pathology lab visible to most of the hospital's staff. A nurse who had had a test done by her family doctor complained to him after she found the result on the hospital system at Basingstoke where she

worked; this caused outrage among local medics, and Malone lost his seat in Parliament at the 1997 election (by two votes) [46].

So how can we avoid letting everyone see every record? There are many ad-hoc things you can do: one fairly effective measure is to keep the records of former patients in a separate archive, and give only a small number of admissions staff the power to move records from there to the main system. Another is to introduce a *honey trap*: one Boston hospital has on its system some bogus 'medical records' with the names of Kennedy family members, so it can identify and discipline staff who browse them. A particularly ingenious proposal, due to Gus Simmons, is to investigate all staff who consult a patient record but do not submit a payment claim to the insurer within thirty days; this aligns the patient's interest in privacy with the hospital's interest in maximizing its income.

However, a patchwork of ad-hoc measures isn't a good way to secure a system. We need a proper access control policy, thought through from first principles and driven by a realistic model of the threats. What policy is appropriate for healthcare?

### **9.2.3.2 The Security Policy**

This question faced the BMA in 1995. The UK government had introduced an IT strategy for the National Health Service which involved centralizing a lot of data on central servers and whose security policy was multilevel: the idea was that AIDS databases would be at a level corresponding to Secret, normal patient records at Confidential and administrative data such as drug prescriptions and bills for treatment at Restricted. It was soon realised that this wasn't going to work. For example, how should a prescription for AZT be classified? As it's a drug prescription, it should be Restricted; but as it identifies a person as HIV positive, it must be Secret. So all the 'Secret' AZT prescriptions must be removed from the 'Restricted' file of drug prescriptions. But then so must almost all the other prescriptions as they identify treatments for named individuals and so should be 'Confidential'. But then what use will the file of prescriptions be to anybody?

A second problem is that the strategy was based on the idea of a single *electronic patient record* (EPR) that would follow the patient around from conception to autopsy, rather than the traditional system of having different records on the same patient at different hospitals and doctors' offices, with information flowing between them in the form of referral and discharge letters. An attempt to devise a security policy for the EPR, which would observe existing ethical norms, quickly became unmanageably complex [558].

In a project for which I was responsible, the BMA developed a security policy to fill the gap. The critical innovation was to define the medical record not as the total of all clinical facts relating to a patient, but as the maximum

set of facts relating to a patient and to which the same staff had access. So an individual patient will have more than one record, and this offended the ‘purist’ advocates of the EPR. But multiple records are dictated anyway by law and practice. Depending on the country (and even the state) that you’re in, you may have to keep separate medical records for human fertilization, sexually transmitted diseases, prison medical services, and even birth records (as they pertain to the health of the mother as well as the child, and can’t simply be released to the child later without violating the mother’s confidentiality). This situation is likely to get more complex still as genetic data start being used more widely.

In many countries, including all signatories to the European Convention on Human Rights, a special status is given to patient consent in law as well as in medical ethics. Records can only be shared with third parties if the patient approves, or in a limited range of statutory exceptions, such as tracing contacts of people with infectious diseases like TB. Definitions are slightly fluid; in some countries, HIV infection is notifiable, in others it isn’t, and in others the data are collected stealthily.

The goals of the BMA security policy were therefore to enforce the principle of patient consent, and to prevent too many people getting access to too many identifiable records. It did not try to do anything new, but merely to codify existing best practice. It also sought to express other security features of medical record management such as safety and accountability. For example, it must be possible to reconstruct the contents of the record at any time in the past, so that for example if a malpractice suit is brought the court can determine what information was available to the doctor at the time. The details of the requirements analysis are in [37].

The policy consists of nine principles.

1. Access control: each identifiable clinical record shall be marked with an access control list naming the people or groups of people who may read it and append data to it. The system shall prevent anyone not on the access control list from accessing the record in any way.
2. Record opening: a clinician may open a record with herself and the patient on the access control list. Where a patient has been referred, she may open a record with herself, the patient and the referring clinician(s) on the access control list.
3. Control: One of the clinicians on the access control list must be marked as being responsible. Only she may alter the access control list, and she may only add other health care professionals to it.
4. Consent and notification: the responsible clinician must notify the patient of the names on his record’s access control list when it is opened, of all subsequent additions, and whenever responsibility is transferred. His

consent must also be obtained, except in emergency or in the case of statutory exemptions.

5. Persistence: no-one shall have the ability to delete clinical information until the appropriate time period has expired.
6. Attribution: all accesses to clinical records shall be marked on the record with the subject's name, as well as the date and time. An audit trail must also be kept of all deletions.
7. Information flow: Information derived from record A may be appended to record B if and only if B's access control list is contained in A's.
8. Aggregation control: there shall be effective measures to prevent the aggregation of personal health information. In particular, patients must receive special notification if any person whom it is proposed to add to their access control list already has access to personal health information on a large number of people.
9. Trusted computing base: computer systems that handle personal health information shall have a subsystem that enforces the above principles in an effective way. Its effectiveness shall be subject to evaluation by independent experts.

This policy may seem to be just common sense, but is surprisingly comprehensive and radical in technical terms. For example, it is strictly more expressive than the Bell-LaPadula model of the last chapter; it contains a BLP-type information flow control mechanism in principle 7, but also contains state. (A fuller discussion from the point of view of access control, and for a technical audience, can be found at [38].)

Similar policies were developed by other medical bodies including the Swedish and German medical associations; the Health Informatics Association of Canada, and an EU project (these are surveyed in [732]). However the BMA model is the most detailed and has been subjected to the most rigorous review; it was adopted by the Union of European Medical Organisations (UEMO) in 1996. Feedback from public consultation on the policy can be found in [39].

### **9.2.3.3 Pilot Implementations**

In a top-down approach to security engineering, one should first determine the threat model, then write the policy, and then finally test the policy by observing whether it works in real life.

BMA-compliant systems have now been implemented both in general practice [585], and in a hospital system developed in Hastings, England, that enforces similar access rules using a mixture of roles and capabilities. It has rules such as 'a ward nurse can see the records of all patients who have within

the previous 90 days been on her ward', 'a doctor can see the records of all patients who have been treated in her department', and 'a senior doctor can see the records of all patients, but if she accesses the record of a patient who has never been treated in her department, then the senior doctor responsible for that patient's care will be notified'. (The hospital system was initially designed independently of the BMA project. When we learned of each other we were surprised at how much our approaches coincided, and reassured that we had captured the profession's expectations in a reasonably accurate way.)

The lessons learned are discussed in [366, 367, 585]. One was the difficulty of constructing a small trusted computing base. The hospital records system has to rely on the patient administrative system to tell it which patients, and which nurses, are on which ward. A different prototype system at a hospital in Cambridge, England, furnished staff with certificates in smartcards which they used to log on.

### 9.2.4 Current Privacy Issues

In 2002, Prime Minister Tony Blair was persuaded to allocate £6bn to modernise health service computing in England. This led to a scramble for contracts with security being something of an afterthought. The original vision was for much improved communications in each local health community; so that if a diabetic patient was being seen by a family doctor, a hospital diabetologist, a community nurse and an optician, they would all be able to see each others' notes and test results. The patient herself would also be able to upload data such as blood glucose levels, see her medical notes, and participate in her care. This vision had been pioneered in the Wirral near Liverpool.

When the dust of the contracting process had settled, the local empowerment vision had been replaced with a much more central approach. Contracts were let for five regions, each with about 10 million people, calling for all hospital systems to be replaced during 2004–2010 with standard ones. The number of system suppliers has been whittled down to two — Cerner and iSoft — and the security policy has been the subject of much debate. The current policy is for three main mechanisms.

1. The workhorse of access control will be role-based access controls, similar to those pioneered at Hastings, but much more complex; rather than a dozen or so roles the plan is now for there to be over three hundred.
2. In order to access patient data, a staff member will also need a *legitimate relationship*. This is an abstraction of the Hastings idea of 'her department'.
3. By default each patient has a single electronic patient record. However, patients will also be able to declare that certain parts of their records are either 'sealed' or 'sealed and locked'. In the latter case, the records will only be visible to a particular care team. In the former, their existence will

be visible to other staff who look at the patient record, and who will be able to break the seal in an emergency.

Initial implementations have thrown up a whole host of detailed problems. For example, patients receiving outpatient psychiatric care at a hospital used to have their notes kept in paper in the psychiatrist's filing cabinet; all the receptionist got to know was that Mrs Smith was seen once a month by Dr Jones. Now, however, the receptionist can see the notes too. Her role had to be given access to patient records so that she could see and amend administrative data such as appointment times; and if she's working reception in the hospital wing where Dr Jones has his office, then she has a legitimate relationship. Record sealing and locking aren't implemented yet. Thus she gets access to everything. This is a good example of why the 'EPR' doctrine of one record per patient was a bad idea, and the BMA vision of multiple linked records was better; it now looks like all records in psychiatry, sexual health etc may have to be sealed (or even sealed-and-locked) by default. Then the care of such patients across different departments will start to cause problems. As with multilevel secure systems, the hard thing isn't so much separating systems, but managing information flows across levels, or across compartments.

Perhaps the toughest problems with the new English systems, however, concern patient consent. The health service is allowing people to opt out of the summary care record — the central database of emergency medical information, containing things like medications, allergies and major medical history. This is not such a big deal; most people have nothing stigmatising in there. (Indeed, most people under the retirement age have no significant chronic conditions and could do perfectly well without a summary record.) The bigger deal is that the new hospital systems will make detailed records available to third parties as never before, for research, health service management and even law enforcement.

Previously, your medical privacy was protected by the fact that a hospital might have had over seventy different departmental record systems, while your records at your family doctor were protected by being partly on paper and partly on a PC that was switched off at six every evening and to which outsiders had no access. Once everything sits in standard systems on a regional health server farm, the game changes. Previously, a policeman who wanted to see your medical records needed to persuade a judge that he had reasonable grounds to believe he would find actual evidence of a crime; he then had to take the warrant along to your family doctor, or your hospital's medical director. The costs of this procedure ensured that it was invoked only rarely, and in cases like terrorism, murder or rape. A server farm, though, is a much easier target — and if it contains data of everyone who's confessed illegal drug use to their doctor, it's a tempting target. Indeed, from June 2007 all UK doctors are supposed to complete a 'treatment outcomes profile' for drug users, asking

them whether they've committed any crimes in the past four weeks, including theft, assault and selling drugs. It's hard to believe that this information won't eventually find its way into police hands. But what are the consequences for public health when people can no longer trust their doctors — especially the most vulnerable and marginalised members of society? We already have cases of immigrants with TB absconding, since health service demographic data started being used to find illegal immigrants.

Thus even if the security policy in centralised systems amounts to a faithful implementation of the BMA policy — with the exception of the eighth principle of non-aggregation — we may expect problems. There are some aspects of security policy that just don't scale. Creating large databases of sensitive personal information is intrinsically hazardous. It increases the motive for abuse, and the opportunity for abuse, at the same time. And even if the controls work perfectly to prevent unlawful abuse (whether by outsiders or insiders) the existence of such databases can lead to lawful abuse — powerful interests in society lobby for, and achieve, access to data on a scale and of a kind that sensible people would not permit.

There are some advantages to standard central systems. In the USA, the Veterans' Administration runs such systems for its hospital network; after Hurricane Katrina, veterans from Louisiana who'd ended up as refugees in Texas or Florida, or even Minnesota, could go straight to local VA hospitals and find their notes there at the doctor's fingertips. Patients of many other hospitals and clinics in New Orleans lost their notes altogether. But centralization can definitely harm privacy. In May 2006, the personal information on all 26.5 million U.S. veterans — including names, social security numbers and in some cases disabilities — was stolen from the residence of a Department of Veterans Affairs employee who had taken the data home without authorization. And it's not enough just to compartmentalise the medical records themselves: in the Netherlands, which has carefully avoided record centralization, there is still a 'Vecozo' database that contains medical insurance details on citizens, and almost 80,000 people had access to it, from doctors and pharmacists to alternative healers and even taxi firms. There was a scandal when journalists found it was easy to get the private addresses and ex-directory phone numbers of a number of famous politicians, criminals and personalities [126]. (After the scandal broke, the insurers and their database operator each tried to blame the other — neither would accept responsibility for the fact that it made too much information available to too many people.)

So if a political decision is taken to have a large centralised database, the aggregation issue will haunt the detailed design and continued operation: even if some people (or applications) are allowed to look at everything, it's an extremely bad idea not to control the principals that actually do so. If you find that most physicians at your hospital look at a few thousand out of the several million records in the database, and one looks at all of them, what does

that tell you? You'd better find out<sup>2</sup>. But many fielded systems don't have rate controls, or effective alarms, and even where alarms exist they are often not acted on. Again in the UK, over 50 hospital staff looked at the records of a footballing personality in hospital, despite not being involved in his care, and none of them was disciplined.

And even apart from controversial uses of medical records, such as police access, there are serious problems in protecting relatively uncontroversial uses, such as research. I'll turn to that next.

## 9.3 Inference Control

---

Access control in medical record systems is hard enough in hospitals and clinics that care for patients directly. It is much harder to assure patient privacy in secondary applications such as databases for research, cost control and clinical audit. This is one respect in which doctors have a harder time protecting their data than lawyers; lawyers can lock up their confidential client files and never let any outsider see them at all, while doctors are under all sorts of pressures to share data with third parties.

### 9.3.1 Basic Problems of Inference Control in Medicine

The standard way of protecting medical records used in research is to remove patients' names and addresses and thus make them anonymous. Indeed, privacy advocates often talk about 'Privacy Enhancing Technologies' (PETs) and de-identification is a frequently cited example. But this is rarely bullet-proof. If a database allows detailed queries, then individuals can still usually be identified, and this is especially so if information about different clinical episodes can be linked. For example, if I am trying to find out whether a politician born on the 2nd June 1946 and treated for a broken collar bone after a college football game on the 8th May 1967, had since been treated for drug or alcohol problems, and I could make an enquiry on those two dates, then I could very probably pull out his record from a national database. Even if the date of birth is replaced by a year of birth, I am still likely to be able to compromise patient privacy if the records are detailed or if records of different individuals can be linked. For example, a query such as 'show me the records of all women aged 36 with daughters aged 14 and 16 such that the mother and exactly one daughter have psoriasis' is also likely to find one individual out of

<sup>2</sup>In November 2007, a former DuPont scientist was sentenced for theft of trade secrets after they noticed he was downloading more internal documents than almost anyone else in the firm, and investigated [294]. It's not just hospitals and spooks that need to keep an eye on data aggregation!

millions. And complex queries with lots of conditions are precisely the kind that researchers want to make.

For this reason, the U.S. Healthcare Finance Administration (HCFA), which pays doctors and hospitals for treatments provided under the Medicare program, maintains three sets of records. There are complete records, used for billing. There are *beneficiary-encrypted* records, with only patients' names and social security numbers obscured. These are still considered personal data (as they still have dates of birth, postal codes and so on) and so are only usable by trusted researchers. Finally there are *public-access* records which have been stripped of identifiers down to the level where patients are only identified in general terms such as 'a white female aged 70–74 living in Vermont'. Nonetheless, researchers have found that many patients can still be identified by cross-correlating the public access records with commercial databases, and following complaints by privacy advocates, a report from the General Accounting Office criticised HCFA for lax security [520].

U.S. law, which comes under the HIPAA privacy rule, now recognizes *de-identified information* as medical data that has been 'properly' de-identified. This means either that 18 specific identifiers have been removed and the database operator has no actual knowledge that the remaining information can be used alone or in combination with other data to identify the subject; or that a qualified statistician concludes that the risk is substantially limited. Where such data are inadequate for research, it also recognises *limited data sets* that contain more information, but where the users are bound by contractual and technical measures to protect the information and not to try to re-identify subjects.

Many other countries have healthcare monitoring systems that use similar approaches. Germany has very strict privacy laws and takes the 'de-identified information' route; the fall of the Berlin Wall forced the former East German cancer registries to install protection mechanisms rapidly [192]. New Zealand takes the 'limited data sets' approach with a national database of encrypted-beneficiary medical records; access is restricted to a small number of specially cleared medical statisticians, and no query is answered with respect to less than six records [955]. In Switzerland, some research systems were replaced at the insistence of privacy regulators [1137].

In other countries, protection has been less adequate. Britain's National Health Service built a number of centralized databases in the 1990s that make personal health information widely available within government and that led to confrontation with doctors. The government set up a committee to investigate under Dame Fiona Caldicott; her report identified over sixty illegal information flows within the health service [46, 252]. Some research datasets were de-identified; others (including data on people with HIV/AIDS) were re-identified afterwards, so that people and HIV charities whose data had been collected under a promise of anonymity were deceived. Parliament then passed a law giving ministers the power to regulate secondary uses of

medical data. Data kept for secondary uses are kept with postcode plus date of birth, and as UK postcodes are shared by at most a few dozen houses, this means that most records are easily identifiable. This remains a cause of controversy. In 2007, Parliament's Health Select Committee conducted an inquiry into the Electronic Patient Record, and heard evidence from a wide range of viewpoints — from researchers who believed that the law should compel information sharing for research, through to physicians, human-rights lawyers and privacy advocates who argued that there should only be the narrowest exceptions to medical privacy<sup>3</sup>. The Committee made many recommendations, including that patients should be permitted to prevent the use of their data in research [624]. The Government rejected this.

The most controversial of all was a genetic database in Iceland, which I'll discuss in more detail below.

Stripping personal information is important in many other fields. Under the rubric of *Privacy Enhancing Technology* (PET) it has been promoted recently by regulators in Europe and Canada as a general privacy mechanism [447]. But, as the medical examples show, there can be serious tension between the desire of researchers for detailed data, and the right of patients (or other data subjects) to privacy. Anonymisation is much more fragile than it seems; and when it fails, companies and individuals that relied on it can suffer serious consequences.

AOL faced a storm of protest in 2006 when it released the supposedly anonymous records of 20 million search queries made over three months by 657,000 people. Searchers' names and IP addresses were replaced with numbers, but that didn't help. Investigative journalists looked through the searches and rapidly identified some of the searchers, who were shocked at the privacy breach [116]. This data was released 'for research purposes': the leak led to complaints being filed with the FTC, following which the company's CTO resigned, and the firm fired both the employee who released the data and the employee's supervisor.

Another example is in movie privacy. The DVD rental firm Netflix ships over a million DVDs a day to over 6 million U.S. customers, has a rating system to match films to customers, and published the viewer ratings of 500,000 subscribers with their names removed. (They offered a \$1m prize for a better recommender algorithm.) In November 2007, Arvind Narayanan and Vitaly Shmatikov showed that many subscribers could be reidentified by comparing the anonymous records with preferences publicly expressed in the Internet Movie Database [928]. This is partly due to the 'long tail' effect: once you disregard the 100 or so movies everyone watches, people's viewing preferences are pretty unique. Anyway, U.S. law protects movie rental privacy, and the attack was a serious embarrassment for Netflix.

<sup>3</sup>Declaration of interest: I was a Special Adviser to the Committee.

So it is important to understand what can, and what cannot, be achieved with this technology.

### 9.3.2 Other Applications of Inference Control

The inference control problem was first seriously studied in the context of census data. A census collects a vast amount of sensitive information about individuals, then makes statistical summaries of it available by geographical (and governmental) units such as regions, districts and wards. This information is used to determine electoral districts, to set levels of government funding for public services, and as inputs to all sorts of other policy decisions. The census problem is somewhat simpler than the medical record problem as the data are rather restricted and in a standard format (age, sex, race, income, number of children, highest educational attainment, and so on).

There are two broad approaches, depending on whether the data are de-identified before or during processing — or equivalently whether the software that will process the data is untrusted or trusted.

An example of the first kind of processing comes from the treatment of U.S. census data until the 1960's. The procedure then was that one record in a thousand was made available on tape — minus names, exact addresses and other sensitive data. There was also noise added to the data in order to prevent people with some extra knowledge (such as of the salaries paid by the employer in a company town) from tracing individuals. In addition to the sample records, local averages were also given for people selected by various attributes. But records with extreme values — such as very high incomes — were suppressed. The reason for this is that a wealthy family living in a small village might make a significant difference to the per-capita village income. So their income might be deduced by comparing the village's average income with that of other villages nearby.

In the second type of processing, identifiable data are retained in a database, and privacy protection comes from controlling the kind of queries that may be made. Early attempts at this were not very successful, and various attacks were proposed on the processing used at that time by the U.S. census. The question was whether it was possible to construct a number of enquiries about samples containing a target individual, and work back to obtain supposedly confidential information about that individual.

If our census system allows a wide range of statistical queries, such as 'tell me the number of households headed by a man earning between \$50,000 and \$55,000', 'tell me the proportion of households headed by a man aged 40–45 years earning between \$50,000 and \$55,000', 'tell me the proportion of households headed by a man earning between \$50,000 and \$55,000 whose children have grown up and left home', and so on, then an attacker can quickly home in on an individual. Such queries, in which we add additional

circumstantial information in order to defeat averaging and other controls, are known as *trackers*. They are usually easy to construct.

A problem related to inference is that an opponent who gets hold of a number of unclassified files might deduce sensitive information from them. For example, a New Zealand journalist deduced the identities of many officers in GCSB (that country's equivalent of the NSA) by examining lists of service personnel and looking for patterns of postings over time [576]. Intelligence officers' cover postings might also be blown if an opponent gets hold of the internal phone book for the unit where the officer is supposed to be posted, and doesn't find his name there. The army list might be public, and the phone book 'Restricted'; but the fact that a given officer is involved in intelligence work might be 'Secret'. Combining low level sources to draw a high level conclusion is known as an *aggregation attack*. It is related to the increased risk to personal information that arises when databases are aggregated together, thus making more context available to the attacker and making tracker and other attacks easier. The techniques that can be used to counter aggregation threats are similar to those used for general inference attacks on databases, although there are some particularly difficult problems where we have a multilevel security policy and the inference or aggregation threats have the potential to subvert it.

### 9.3.3 The Theory of Inference Control

A theory of inference control was developed by Denning and others in late 1970s and early 1980s, largely in response to problems of census bureaux [369]. The developers of many modern privacy systems are often unaware of this work, and repeat many of the mistakes of the 1960s. (Inference control is not the only problem in computer security where this happens.) The following is an overview of the most important ideas.

A *characteristic formula* is the expression (in some database query language) that selects a set, known as the *query set*, of records. An example might be 'all female employees of the Computer Laboratory at the grade of professor'. The smallest query sets, obtained by the logical AND of all the attributes (or their negations) are known as *elementary sets* or *cells*. The statistics corresponding to query sets may be *sensitive statistics* if they meet criteria which I'll discuss below (such as the set size being too small). The objective of inference control is to prevent the disclosure of sensitive statistics.

If we let  $D$  be the set of statistics that are disclosed and  $P$  the set which are sensitive and must be protected, then we need  $D \subseteq P'$  for privacy, where  $P'$  is the complement of  $P$ . If  $D = P'$  then the protection is said to be *precise*. Protection which is not precise will usually carry some cost in terms of the range of queries which the database can answer and may thus degrade its usefulness to its owner.

### 9.3.3.1 Query Set Size Control

The simplest protection mechanism is to specify a minimum query size. As I mentioned, New Zealand's National Health Information System databases will reject statistical queries whose answers would be based on fewer than six patients' records. But this is not enough in itself. An obvious tracker attack is to make an enquiry on six patients' records, and then on those records plus the target's. Rather than reduce the effectiveness of the database by building in more restrictive query controls, the designers of this system opted to restrict access to a small number of specially cleared medical statisticians.

Even so, one extra control is needed, and is often forgotten. You must prevent the attacker from querying all but one of the records in the database. In general, if there are  $N$  records, query set size control with a threshold of  $t$  means that between  $t$  and  $N - t$  of them must be the subject of a query for it to be allowed.

### 9.3.3.2 Trackers

Probably the most important attacks on statistical databases come from trackers. There are many simple examples. In our laboratory, only one of the full professors is female. So we can find out her salary with just two queries: 'Average salary professors?' and 'Average salary male professors?'.

This is an example of an *individual tracker*, a custom formula that allows us to calculate the answer to a forbidden query indirectly. There are also *general trackers* — sets of formulae which will enable any sensitive statistic to be revealed. A somewhat depressing discovery made in the late 1970s was that general trackers are usually easy to find. Provided the minimum query set size  $n$  is less than a quarter of the total number of statistics  $N$ , and there are no further restrictions on the type of queries that are allowed, then we can find formulae that provide general trackers [372]. So tracker attacks are easy, unless we place severe restrictions on the query set size or control the allowed queries in some other way. (In fact results like this caused the research community to largely lose interest in inference security as being 'too hard', and this is one of the reasons that many system designers are not aware of the problems and build databases vulnerable to trackers and other attacks.)

### 9.3.3.3 More Sophisticated Query Controls

There are a number of alternatives to simple query set size control. The U.S. census, for example, uses the ' $n$ -respondent,  $k\%$ -dominance rule': it will not release a statistic of which  $k\%$  or more is contributed by  $n$  values or less. Other techniques include, as I mentioned, suppressing data with extreme values. A census bureau may deal with high-net-worth individuals in national statistics

but not in the local figures, while some medical databases do the same for less common diseases. For example, a UK prescribing statistics system suppresses sales of the AIDS drug AZT from local statistics [847]. When it was designed in the late 1990s, there were counties with only one single patient receiving this drug.

#### 9.3.3.4 Cell Suppression

The next question is how to deal with the side-effects of suppressing certain statistics. UK rules, for example, require that it be ‘unlikely that any statistical unit, having identified themselves, could use that knowledge, by deduction, to identify other statistical units in National Statistics outputs’ [953]. To make this concrete, suppose that a university wants to release average marks for various combinations of courses, so that people can check that the marking is fair across courses. Suppose now that the table in Figure 9.4 contains the number of students studying two science subjects, one as their major subject and one as their minor subject.

The UK rules imply that our minimum query set size is 3 (if we set it at 2, then either of the two students who studied ‘geology-with-chemistry’ could trivially work out the other’s mark). Then we cannot release the average mark for ‘geology-with-chemistry’. But if the average mark for chemistry is known, then this mark can easily be reconstructed from the averages for ‘biology-with-chemistry’ and ‘physics-with-chemistry’. So we have to suppress at least one other mark in the chemistry row, and for similar reasons we need to suppress one in the geology column. But if we suppress ‘geology-with-biology’ and ‘physics-with-chemistry’, then we’d also better suppress ‘physics-with-biology’ to prevent these values being worked out in turn. Our table will now look like Figure 9.5.

This process is called *complementary cell suppression*. If there are further attributes in the database schema — for example, if figures are also broken down by race and sex, to show compliance with anti-discrimination laws — then even more information may be lost. Where a database scheme contains  $m$ -tuples, blanking a single cell generally means suppressing  $2^m - 1$

| Major:    | Biology | Physics | Chemistry | Geology |
|-----------|---------|---------|-----------|---------|
| Minor:    |         |         |           |         |
| Biology   | –       | 16      | 17        | 11      |
| Physics   | 7       | –       | 32        | 18      |
| Chemistry | 33      | 41      | –         | 2       |
| Geology   | 9       | 13      | 6         | –       |

**Figure 9.4:** Table containing data before cell suppression

| Major:    | Biology | Physics | Chemistry | Geology |
|-----------|---------|---------|-----------|---------|
| Minor:    |         |         |           |         |
| Biology   | –       | blanked | 17        | blanked |
| Physics   | 7       | –       | 32        | 18      |
| Chemistry | 33      | blanked | –         | blanked |
| Geology   | 9       | 13      | 6         | –       |

**Figure 9.5:** Table after cell suppression

other cells, arranged in a hypercube with the sensitive statistic at one vertex. So even precise protection can rapidly make the database unusable.

Sometimes complementary cell suppression can be avoided, as when large incomes (or rare diseases) are tabulated nationally and excluded from local figures. But it is often necessary when we are publishing microstatistics, as in the above tables of exam marks. Where the database is open for online queries, we can get much of the same effect by *implied queries control*: we allow a query on  $m$  attribute values only if all of the  $2^m$  implied query sets given by setting the  $m$  attributes to true or false, have at least  $k$  records.

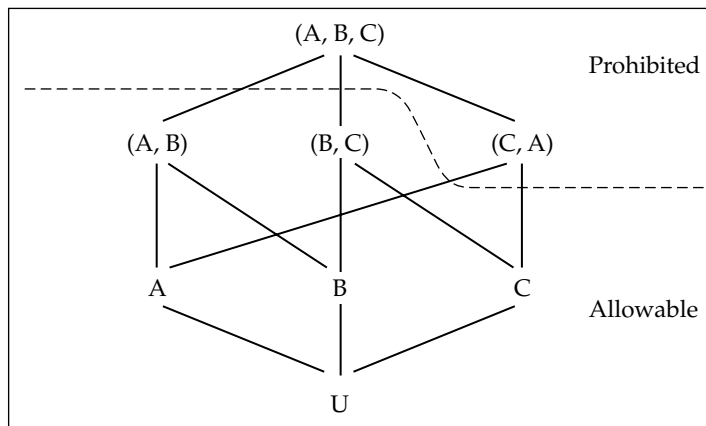
### 9.3.3.5 Maximum Order Control and the Lattice Model

The next thing we might try in order to make it harder to construct trackers is to limit the type of inquiries that can be made. *Maximum order control* limits the number of attributes that any query can have. However, to be effective, the limit may have to be severe. One study found that of 1000 medical records, three attributes were safe while with four attributes, one individual record could be found and with 10 attributes most records could be isolated. A more thorough approach (where it is feasible) is to reject queries that would partition the sample population into too many sets.

We saw how lattices can be used in compartmented security to define a partial order to control permitted information flows between compartments with combinations of codewords. They can also be used in a slightly different way to systematize query controls in some databases. If we have, for example, three attributes  $A$ ,  $B$  and  $C$  (say area of residence, birth year and medical condition), we may find that while enquiries on any one of these attributes are non-sensitive, as are enquiries on  $A$  and  $B$  and on  $B$  and  $C$ , the combination of  $A$  and  $C$  might be sensitive. It follows that an enquiry on all three would not be permissible either. So the lattice divides naturally into a ‘top half’ of prohibited queries and a ‘bottom half’ of allowable queries, as shown in Figure 9.6.

### 9.3.3.6 Audit Based Control

As mentioned, some systems try to get round the limits imposed by static query control by keeping track of who accessed what. Known as *query overlap control*,



**Figure 9.6:** Table lattice for a database with three attributes

this involves rejecting any query from a user which, combined with what the user knows already, would disclose a sensitive statistic. This may sound perfect in theory but in practice it suffers from two usually unsurmountable drawbacks. First, the complexity of the processing involved increases over time, and often exponentially. Second, it's extremely hard to be sure that your users aren't in collusion, or that one user has registered under two different names. Even if your users are all honest and distinct persons today, it's always possible that one of them will take over another, or get taken over by a predator, tomorrow.

### 9.3.3.7 Randomization

Our cell suppression example shows that if various kinds of query control are the only protection mechanisms used in a statistical database, they can often have an unacceptable performance penalty. So query control is often used in conjunction with various kinds of randomization, designed to degrade the signal-to-noise ratio from the attacker's point of view while impairing that of the legitimate user as little as possible.

The simplest such technique is *perturbation*, or adding noise with zero mean and a known variance to the data. One way of doing this is to round or truncate the data by some deterministic rule; another is to swap some records. Perturbation is often not as effective as one would like, as it will tend to damage the legitimate user's results precisely when the sample set sizes are small, and leave them intact when the sample sets are large (where we might have been able to use simple query controls anyway). There is also the worry that suitable averaging techniques might be used to eliminate some of the added noise. A modern, sophisticated variant on the same theme is *controlled tabular adjustment* where you identify the sensitive cells and replace their

values with ‘safe’ (sufficiently different) ones, then adjust other values in the table to restore additive relationships [330].

Often a good randomization technique is to use *random sample queries*. This is another of the methods used by census bureaux. The idea is that we make all the query sets the same size, selecting them at random from the available relevant statistics. Thus all the released data are computed from small samples rather than from the whole database. If this random selection is done using a pseudorandom number generator keyed to the input query, then the results will have the virtue of repeatability. Random sample queries are a natural protection mechanism for large medical databases, where the correlations being investigated are often such that a sample of a few hundred is sufficient. For example, when investigating the correlation between a given disease and some aspect of lifestyle, the correlation must be strong before doctors will advise patients to make radical changes to their way of life, or take other actions that might have undesirable side effects. If a teaching hospital has records on five million patients, and five thousand have the disease being investigated, then a randomly selected sample of two hundred sufferers might be all the researcher could use.

This doesn’t work so well where the disease is rare, or where for other reasons there is only a small number of relevant statistics. A possible strategy here is *randomized response*, where we randomly restrict the data we collect (the subjects’ responses). For example, if the three variables under investigation are obesity, smoking and AIDS, we might ask each subject with HIV infection to record whether they smoke or whether they are overweight, but not both. Of course, this can limit the value of the data.

### 9.3.4 Limitations of Generic Approaches

As with any protection technology, statistical security can only be evaluated in a particular environment and against a particular threat model. Whether it is adequate or not depends to an even greater extent than usual on the details of the application.

An instructive example is a system used for analyzing trends in drug prescribing. Here, prescriptions are collected (minus patient names) from pharmacies. A further stage of de-identification removes the doctors’ identities, and the information is then sold to drug company marketing departments. The system has to protect the privacy of doctors as well as of patients: the last thing a busy family doctor wants is to be pestered by a drug rep for prescribing a competitor’s brands.

One early prototype of this system merely replaced the names of doctors in a cell of four or five practices with ‘doctor A’, ‘doctor B’ and so on, as in Figure 9.7. We realised that an alert drug rep could identify doctors from prescribing patterns, by noticing, for example, “Well, doctor B must be Susan

| Week:    | 1  | 2  | 3  | 4  |
|----------|----|----|----|----|
| Doctor A | 17 | 26 | 19 | 22 |
| Doctor B | 25 | 31 | 9  | 29 |
| Doctor C | 32 | 30 | 39 | 27 |
| Doctor D | 16 | 19 | 18 | 13 |

**Figure 9.7:** Sample of de-identified drug prescribing data

Jones because she went skiing in the third week in January and look at the fall-off in prescriptions here. And doctor C is probably her partner Mervyn Smith who'll have been covering for her'' The fix was to replace absolute numbers of prescriptions with the percentage of each doctor's prescribing which went on each particular drug, to drop some doctors at random, and to randomly perturb the timing by shifting the figures backwards or forwards a few weeks [847].

This is a good example of the sort of system where the inference control problem can have a robust solution. The application is well-defined, the database is not too rich, and the allowable queries are fairly simple. Indeed, this system was the subject of litigation; the UK government's Department of Health sued the database operator, alleging that the database might compromise privacy. Their motive was to maintain a monopoly on the supply of such information to industry. They lost, and this established the precedent that (in Britain at least) inference security controls may, if they are robust, exempt statistical data from being considered as 'personal information' for the purpose of privacy laws [1204].

In general, though, it's not so easy. For a start, de-identification doesn't compose: it's easy to have two separate applications, each of which provides the same results via anonymized versions of the same data, but where an attacker with access to both of them can easily identify individuals. In the general case, contextual knowledge is extremely hard to quantify, and is likely to grow over time. Latanya Sweeney has shown that even the HCFA's 'public-use' files can often be reidentified by cross-correlating them with commercial databases [1235]: for example, most U.S. citizens can be identified by their ZIP code plus their gender and date of birth. Such *data detective* work is an important part of assessing the level of protection which an actual statistical database gives, just as we only have confidence in cryptographic algorithms which have withstood extensive analysis by capable motivated opponents. The emergence of social networks since 2004 has made inference control much harder wherever they can be brought to bear; I will discuss this when we get to social networks in section 23.3.3. And even without cross-correlation, there may be contextual information available internally. Users of medical research databases are often doctors who have normal access to parts of the patient record databases from which the statistical data are drawn.

### 9.3.4.1 Active Attacks

*Active attacks* are particularly powerful. These are where users have the ability to insert or delete records into the database. A user might add records to create a group that contains the target's record plus those of a number of nonexistent subjects created by himself. One (imperfect) countermeasure is add or delete new records in batches. Taking this to an extreme gives *partitioning* — in which records are added in groups and any query must be answered with respect to all of them or none. However, this is once more equivalent to publishing tables of microstatistics.

Active attacks are not limited to data, but can also target metadata. A nice example, due to Whit Diffie, is the *chosen drug attack*. Suppose a drug company has access through a statistical system to the amounts of money spent on behalf of various groups of patients and wishes to find out which patients are receiving which drug, in order to direct its marketing better (there was a scandal in Quebec about just such an inference attack). A possible trick is to set the drug prices in such a way as to make the resulting equations easy to solve.

A prominent case at the turn of the century was a medical research database in Iceland. The plan was for three linked databases: one with the nation's medical records, a second with the genealogy of the whole population, and a third with genetic data acquired from sequencing. The rationale was that since Iceland's population is largely descended from a few founding families who settled there about a thousand years ago, there is much less genic variance than in the general human population and so genes for hereditary illnesses should be much easier to find. A Swiss drug company bankrolled the construction of the database, and the Reykjavik government embraced it as a means of modernising the country's health IT infrastructure and simultaneously creating a few hundred high-tech jobs in medical research. Iceland's doctors, however, mostly reacted negatively, seeing the system as a threat both to patient privacy and professional autonomy.

The privacy problem in the Icelandic database was more acute than in the general case. For example, by linking medical records to genealogies, which are in any case public (genealogy is a common Icelandic hobby), patients can be identified by such factors as the number of their uncles, aunts, great-uncles, great-aunts and so on — in effect by the shape of their family trees. There was much debate about whether the design could even theoretically meet legal privacy requirements [47], and European privacy officials expressed grave concern about the possible consequences for Europe's system of privacy laws [349]. The Icelandic government pressed ahead with it anyway, with a patient opt-out. Many doctors advised patients to opt out, and 11/population did so. Eventually, the Icelandic Supreme Court found that European privacy law required the database to be opt-in rather than opt-out. In addition, many Icelanders had invested in the database company, and lost money when its

share value sank at the end of the dotcom boom. Nowadays about half the population have opted in to the system and the controversy is defused.

My own view, for what it's worth, is that patient consent is the key to effective medical research. This not only allows full access to data, without the problems we've been discussing in this section, but provides motivated subjects and much higher-quality clinical information than can be harvested simply as a byproduct of normal clinical activities. For example, a network of researchers into ALS (the motor-neurone disease from which Cambridge astronomer Stephen Hawking suffers) shares fully-identifiable information between doctors and other researchers in over a dozen countries with the full consent of the patients and their families. This network allows data sharing between Germany, with very strong privacy laws, and Japan, with almost none; and data continued to be shared between researchers in the USA and Serbia even when the USAF was bombing Serbia. The consent model is spreading. Britain's biggest medical charity is funding a 'Biobank' database in which several hundred thousand volunteers will be asked to give researchers not just answers to an extensive questionnaire and full access to their records for the rest of their lives, but also to lodge blood samples so that those who develop interesting diseases in later life can have their genetic and proteomic makeup analysed.

### 9.3.5 The Value of Imperfect Protection

So doing de-identification right is hard, and the issues can be politically fraught. The best way to solve the inference control problem is to avoid it, for example by recruiting volunteers for your medical research rather than recycling data collected for other purposes. But there are applications where it's used, and applications where it's all that's available. An example was the epidemic of HIV/AIDS; in the 1980s and 1990s researchers struggling to understand what was going on had little choice but to use medical data that had been originally collected for other purposes. Another example, of course, is the census. In such applications the protection you can provide will be imperfect. How do you cope with that?

Some kinds of security mechanism may be worse than useless if they can be compromised. Weak encryption is a good example. The main problem facing the world's signals intelligence agencies is *traffic selection* — how to filter out interesting nuggets from the mass of international phone, fax, email and other traffic. A terrorist who helpfully encrypts his important traffic does this part of the police's job for them. If the encryption algorithm used is breakable, or if the end systems can be hacked, then the net result is worse than if the traffic had been sent in clear.

Statistical security is not generally like this. The main threat to databases of personal information is often *mission creep*. Once an organization has access to

potentially valuable data, then all sorts of ways of exploiting that value will be developed. Some of these are likely to be highly objectionable; one topical U.S. example is the resale of medical records to banks for use in filtering loan applications. However, even an imperfect de-identification system may destroy the value of medical data to a bank's loan department. If only five percent of the patients can be identified, and then only with effort, then the bank may decide that it's simpler to tell loan applicants to take out their own insurance and let the insurance companies send out medical questionnaires if they wish. So de-identification can help prevent mission creep, even if the main effect is prophylaxis against future harm rather than treatment of existing defects.

As well as harming privacy, mission creep can have safety implications. In the UK, diabetic registers were set up in the 1990s to monitor the quality of diabetes care; they were databases to which GPs, hospital consultants, nurses and ophthalmologists could upload test results, so that important indicators would not be missed. As hospitals had no working email system, they were promptly abused to provide a rudimentary messaging system between hospitals and general practice. But as the diabetes registers were never designed as communications systems, they lacked the safety and other mechanisms that such systems should have had if they were to be used for clinical data. Even rudimentary de-identification would have prevented this abuse and motivated diabetologists to get email working instead.

So in statistical security, the question of whether one should let the best be the enemy of the good can require a finer judgment call than elsewhere.

## 9.4 The Residual Problem

---

The above two sections may have convinced you that the problem of managing medical record privacy in the context of immediate care (such as in a hospital) is reasonably straightforward, while in the context of secondary databases (such as for research, audit and cost control) there are statistical security techniques which, with care, can solve much of the problem. Somewhat similar techniques can be used to manage highly sensitive commercial data such as details of forthcoming mergers and acquisitions in an investment bank, and even intelligence information. (There was a lot of interest in the BMA model from people designing police intelligence systems.) In all cases, the underlying concept is that the really secret material is restricted to a compartment of a small number of identified individuals, and less secret versions of the data may be manufactured for wider use. This involves not just suppressing the names of the patients, or spies, or target companies, but also careful management of contextual and other information by which they might be re-identified.

But making such systems work well in real life is much harder than it looks. First, determining the sensitivity level of information is fiendishly difficult,

and many initial expectations turn out to be wrong. You might expect, for example, that HIV status would be the most sensitive medical data there is; yet many HIV sufferers are quite open about their status. You might also expect that people would rather entrust sensitive personal health information to a healthcare professional such as a doctor or pharmacist rather than to a marketing database. Yet many women are so sensitive about the purchase of feminine hygiene products that, rather than going into a pharmacy and buying them for cash, they prefer to use an automatic checkout facility in a supermarket — even if this means they have to use their store card and credit card, so that the purchase is linked to their name and stays on the marketing database forever. The actual embarrassment of being seen with a packet of tampons is immediate, and outweighs the future embarrassment of being sent discount coupons for baby wear six months after the menopause.

Second, it is extraordinarily difficult to exclude single points of failure, no matter how hard you try to build watertight compartments. The CIA's Soviet assets were compromised by Aldrich Ames — who as a senior counterintelligence man had access to too many compartments. The KGB's overseas operations were similarly compromised by Vassily Mitrokhin — an officer who'd become disillusioned with communism after 1968 and who was sent to work in the archives while waiting for his pension [77]. And in March 2007, historians Margo Anderson and William Seltzer found, that contrary to decades of denials, census data was used in 1943 to round up Japanese-Americans for internment [1142]. The single point of failure there appears to have been Census Bureau director JC Capt, who unlawfully released the data to the Secret Service following a request from Treasury Secretary HC Morgenthau. The Bureau has since publicly apologised [893].

In medicine, many of the hard problems lie in the systems that process medical claims for payment. When a patient is treated and a request for payment sent to the insurer, it has not just full details of the illness, the treatment and the cost, but also the patient's name, insurance number and other details such as date of birth. There have been proposals for payment to be effected using anonymous credit cards [191], but as far as I am aware none of them has been fielded. Insurers want to know which patients, and which doctors, are the most expensive. In fact, during a debate on medical privacy at an IEEE conference in 1996 — just as HIPAA was being pushed through the U.S. Congress — a representative of a large systems house declared that the medical records of 8 million Americans were one of his company's strategic assets, which they would never give up. This holds whether the insurer is a private insurance company (or employer) or a government-owned health authority, such as HCFA, the VA, or Britain's National Health Service. Once an insurer possesses large quantities of personal health information, it becomes very reluctant to delete it. Its potential future value, in all sorts of applications

from cost control through research to marketing, is immediate and obvious, while patients' privacy concerns are not.

In the USA, the retention of copies of medical records by insurers, employers and others is widely seen as a serious problem. Writers from such widely different political viewpoints as the communitarian Amitai Etzioni [441] and the libertarian Simson Garfinkel [515] agree on this point, if on little else. As mentioned, HIPAA only empowered the DHHS to regulate health plans, healthcare clearinghouses, and healthcare providers, leaving many organizations that process medical data (such as lawyers, employers and universities) outside its scope. In fact, Microsoft's recent announcement that it would set up a 'HealthVault' to guard your medical records was met with a sharp retort from privacy activists that since Microsoft isn't a 'covered entity' as specified by HIPAA, putting your medical data there would place it outside HIPAA's protection [81].

What lessons can be drawn from other countries?

Medical privacy is strongly conditioned by how people pay for healthcare. In Britain, the government pays for most healthcare, and the attempts of successive British governments to centralise medical records for cost control and management purposes have led to over a decade of conflict with doctors and with patients' associations. In Germany, the richer people use private insurers (who are bound by tight data protection laws), while the poor use state health insurers that are run by doctors, so non-doctors don't have access to records. Singapore residents pay into compulsory savings accounts from their wages and use them to pay for healthcare; the government steps in to insure expensive procedures, but most doctor visits are paid by the patient directly. Patients who stay healthy and accumulate a surplus can add some of it to their pension and pass the rest to their heirs. The most radical solution is in Japan, where costs are controlled by regulating fees: doctors are discouraged from performing expensive procedures such as heart transplants by pricing them below cost. In the mid-1990s, healthcare took up some 3% of GNP in Japan, versus 7–9% for the typical developed country and 15% for America; since then the figures have risen by a percent or so, but the general rankings remain the same. Japanese (and Singaporeans) pay less for healthcare than Europeans, and Americans pay more. The curious thing is that Japanese (and Singaporeans) live longer than Europeans, who live longer than Americans. Life expectancy and medical costs seem to be negatively correlated.

To sum up, the problem of health record privacy is not just a socio-technical one but socio-technico-political. Whether large quantities of medical records accumulate in one database depends on how the health care system is organized, and whether these are destroyed — or de-identified — after payment has been processed is more to do with institutional structures, incentives and regulation than technology. In such debates, one role of the security engineer is to get policymakers to understand the likely consequences of their actions.

Privacy is poorest in countries that fail to align incentives properly, and as a result have detailed cost oversight of individual treatments — whether by insurers / employers, as in the USA, or by bureaucrats as in Britain.

In the UK, a scandal broke in November 2007 when the tax authorities lost the records of 25 million people. The records of all the nation's children and their families — including names, addresses, phone numbers and the parents' bank account details — were burned on two CDs for dispatch to the National Audit Office, and lost in the post. The Prime Minister had to apologise to Parliament and promised to make good any resulting 'identify theft' losses. In the aftermath, there has been wide public questioning of his government's programme to build ever-large central databases of citizens' personal information — not just for taxation but for medical research, health-service administration, and child welfare. As I write in December 2007, the feeling in London is that plans for a national ID card are effectively dead, as is a proposal to build a database of all vehicle movements to facilitate road pricing. The National Health Service is continuing to build central health databases against growing medical resistance, but the opposition Conservative Party (which now has a clear lead in the polls) have promised to abolish not just the ID card system but proposed children's databases if they win the next election.

Other privacy problems also tend to have a serious political entanglement. Bank customer privacy can be tied up with the bank's internal politics; the strongest driver for privacy protection may come from branch managers' reluctance to let other branches learn about their customers. Access to criminal records and intelligence depends on how law enforcement agencies decide to share data with each other, and the choices they make internally about whether access to highly sensitive information about sources and methods should be decentralized (risking occasional losses), or centralized (bringing lower-probability but higher-cost exposure to a traitor at head office). The world since 9/11 has moved sharply towards centralisation; expect a high-profile traitor like Aldrich Ames to come along sometime soon.

---

## 9.5 Summary

---

In this chapter, we looked at the problem of assuring the privacy of medical records. This is typical of a number of information security problems, ranging from the protection of national intelligence data through professional practice in general to the protection of census data.

It turns out that with medical records there is an easy problem, a harder problem, and a really hard problem.

The easy problem is setting up systems of access controls so that access to a particular record is limited to a sensible number of staff. Such systems can be designed largely by automating existing working practices, and role-based

access controls are currently the technology of choice. The harder problem is statistical security — how one designs databases of medical records (or census returns) so as to allow researchers to make statistical enquiries without compromising individuals' privacy. The hardest problem is how to manage the interface between the two, and in the specific case of medicine, how to prevent the spread of payment information. The only realistic solution for this lies in regulation.

Medical systems also teach us about the limits of some privacy enhancing technologies, such as de-identification. While making medical records anonymous in research databases can help mitigate the consequences of unauthorised access and prevent mission creep, it's by no means bulletproof. Rich data about real people can usually be re-identified. The mechanisms used in healthcare to deal with this problem are worth studying.

## Research Problems

---

In the near future, a lot of medical treatment may involve genetic information. So your medical records may involve personal health information about your parents, siblings, cousins and so on. How can privacy models be extended to deal with multiple individuals? For example, in many countries you have the right not to know the outcome of a DNA test that a relative has for an inheritable disease such as Huntington's Chorea, as it may affect the odds that you have the disease too. Your relative does have a right to know, and may tell others. This is a problem not just for technology, but also for privacy law [1231]

Are there any ways of linking together access control policies for privacy with statistical security? Can there be such a thing as seamless privacy where everything fits neatly together? Or would you end up giving patients an extremely complex set of access control options — like Facebook's but worse — in which each patient had to wade through dozens of pages of options and approve or deny permission for her data to be used in each of dozens of secondary applications and research projects? In short, are there any useful and useable abstractions?

What other ways of writing privacy policies are there? For example, are there useful ways to combine BMA and Chinese Wall? Are there any ways, whether technical or economic, of aligning the data subject's interest with those of the system operator and other stakeholders?

## Further Reading

---

The literature on compartmented-mode security is somewhat scattered: most of the public domain papers are in the proceedings of the NCSC/NISSC and

ACSAC conferences cited in detail at the end of Chapter 8. Standard textbooks such as Amoroso [27] and Gollmann [537] cover the basics of the lattice and Chinese Wall models.

For the BMA model see the policy document itself — the Blue Book [37], the shorter version at [38], and the proceedings of the conference on the policy [43]. See also the papers on the pilot system at Hastings [366, 367]. For more on Japanese healthcare, see [263]. For a National Research Council study of medical privacy issues in the USA, see [951]; there is also an HHS report on the use of de-identified data in research at [816].

As for inference control, this has become an active research field again in the last few years, with regular conferences on ‘Privacy in Statistical Databases’; see the proceedings of these events to catch up with current frontiers. Denning’s book [369] is the classic reference, and still worth a look; there’s an update at [374]. A more modern textbook on database security is the one by Castano et al [276]. The most comprehensive resource, though, from the practical viewpoint — with links to a vast range of practical literature across a number of application areas — may be the website of the American Statistical Association [26]. The standard reference for people involved in government work is the Federal Committee on Statistical Methodology’s ‘*Report on Statistical Disclosure Limitation Methodology*’ which provides a good introduction to the standard tools and describes the methods used in various U.S. departments and agencies [455]. As an example of a quite different application, Mark Allman and Vern Paxson discuss the problems of anonymizing IP packet traces for network systems research in [23].

Finally, Margo Anderson and William Seltzer’s papers on the abuses of census data in the USA, particularly during World War 2, can be found at [31].

