

Fundamental concepts

This chapter provides a summary of the main types of trials and their key design features. A checklist for critical appraisal of trial reports is on page 199, and a glossary of common abbreviations is on page 201.

1.1 What is a clinical trial?

The two distinct study designs used in health research are **observational** and **experimental**. Observational studies (usually cross-sectional, retrospective case-control or prospective cohort) do not intentionally involve intervening in the way individuals live their lives or how they are treated.¹ However, clinical trials (experimental) are specifically designed to intervene, and then evaluate health-related outcomes, with the following objectives:

- To diagnose or detect a disease
- To prevent a disease or early death (prolong life)
- To treat or cure an existing disorder, including reducing or managing symptoms
- To change behaviour or lifestyle habits, including reducing risk factors.

Countries (low, middle and high income areas) can successfully conduct clinical trials that reflect local health issues. The fundamental features of design and analysis are similar but the conduct and delivery will vary (especially how interventions are administered and how follow-up and outcome data are collected).

An intervention could be a single therapy involving a substance that is injected, infused, swallowed, inhaled or absorbed through the skin; medical device; surgical procedure; radiotherapy; behavioural or psychological therapy; something to improve health service delivery or promote health education; or an alternative or complementary therapy.

A new generation of **biological and targeted therapies** (small molecules, monoclonal antibodies, immunotherapies and genetic and cell therapies) has revolutionised the treatment of several disorders, in which the choice of a therapy is influenced by the presence (or absence) of a biomarker, genetic abnormality or imaging marker. There are also **vaccines** that can be used for disease prevention or to reduce the risk of disease progression.

A combination of interventions can be referred to as a **regimen**, such as chemotherapy and surgery in treating cancer.

Any drug or micronutrient that is examined in a clinical trial with the specific purpose of treating, preventing or diagnosing disease is usually referred to as an **Investigational Medicinal Product (IMP)** or **Investigational New Drug (IND)**.[#] Most clinical trial regulations cover studies using an IMP and several medical devices.

[#] IMP in the UK and European Union, and IND in the United States, Canada and Japan.

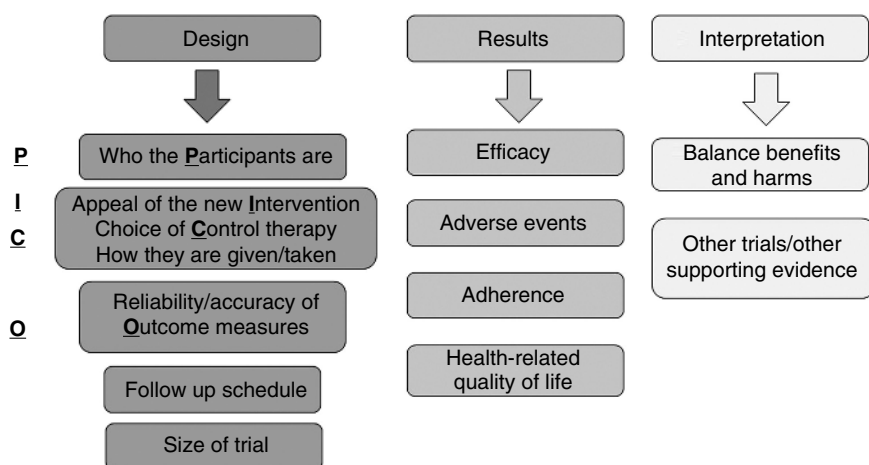


Figure 1.1 Overall view of trial design, types of results and interpretation. The acronym PICO (Participants/Population, Intervention, Control and Outcome) focuses on the four key design elements that must always be clearly defined (examples of trials in later chapters use the PICO list). Translational research (bio- and imaging markers) can also be examined.

New drugs and some medical devices usually require a **licence** or **marketing authorisation** for human use from a national regulatory authority. They can then be made available to the target population after review by a **health technology assessment (HTA)** or **payer/reimbursement** organisation through a process referred to as **market access**.

Throughout this book, the terms **intervention**, **treatment** and **therapy** are used interchangeably. People who take part in a trial are referred to as **participants** if they are healthy individuals or **patients** if they are already ill with the disorder of interest.

Figure 1.1 is an overall view of trial design and types of results (covered in more detail in other chapters).

Patient and Public Involvement and Engagement (PPIE) is a key activity in which lay members (e.g. past patients, carers and members of the public) can help with trial design (e.g. agree that the new therapy is appealing), conduct (create the participant-facing information materials) and interpret trials in a way that can be easily understood.

Artificial intelligence is also expected to be used, for example in identifying eligible participants from electronic medical records and analysing clinical data and multiple biomarkers.

Decentralised trials may be increasingly used where many or all of the processes from participant selection and eligibility, allocation of treatments (usually licensed products already in use), through to data and outcome collection are done electronically including remote assessments of participants.

1.2 Early trials

James Lind, a Scottish naval physician, is regarded as conducting the first clinical trial.² During a sea voyage in 1747, he chose 12 sailors with *similarly* severe cases of scurvy and examined 6 treatments, each given to 2 sailors: cider, diluted sulphuric acid, vinegar, seawater, a mixture of foods including nutmeg and garlic, and oranges and lemons. These sailors were made to live in the *same* part of the ship and given the *same* basic diet. Lind understood the importance of standardising their living conditions to ensure that any

change in their disease would unlikely be influenced by other factors. After about a week, both sailors given fruit had almost completely recovered unlike the other sailors. This dramatic effect led to the conclusion that eating fruit cured scurvy, without knowing that it was due to vitamin C.

Two important features of his trial were a **comparison** between two or more interventions and an attempt to ensure that the participants had **similar characteristics** (see confounding below). The requirement for these two features has not changed for more than 270 years, indicating how essential they are to evaluating interventions.

One key element missing from Lind's trial was the process of **randomisation**. The Medical Research Council trial of streptomycin and tuberculosis in 1948 is regarded as the first to use random allocation.³

1.3 Why clinical trials are needed

Statin therapy is effective in treating coronary heart disease. However, why do some patients who have had a heart attack and been given statin therapy have a second attack, while others do not? The answer is that people *vary*. People have different body characteristics (e.g., weight, blood pressure and blood measurements), genetic make-up and lifestyles (e.g., diet, exercise, and smoking and alcohol consumption habits). These lead to **variability** or **natural variation**. People respond to the same exposure or treatment in different ways. When a new intervention is evaluated, it is essential to consider whether the observed responses are consistent with natural variation (i.e. chance) or whether there really is a treatment effect. This is a principal concern of medical statistics.

1.4 Alternatives to clinical trials

Examining interventions can be done using a clinical trial and in particular a randomised controlled trial (RCT), observational study or trial with historical controls. They have fundamentally different designs. Some observational studies are used as supporting evidence for the effectiveness and safety of an intervention under the topic **real-world evidence** (RWE) or **real-world data** (RWD); see Chapter 9.

Although studies other than RCTs can provide useful information about an intervention, care is needed over their interpretation. Observational studies, for example, can give the same or conflicting conclusion to RCTs:

- A review of 20 observational studies indicated that giving the influenza vaccine to the elderly could halve the risk of developing respiratory and flu-like symptoms.⁴ Practically the same effect was found in a large double-blind RCT.⁵
- A review of 6 observational studies indicated that people with a high β -carotene intake, by eating lots of fruit and vegetables, had a 31% reduction in the risk of cardiovascular death than those with a low intake.⁶ However, 4 RCTs together showed that a high intake increases the risk by 12%.⁶

Observational (non-randomised) studies

Observational studies may be useful in evaluating treatments with large effects, although there may still be uncertainty over the actual size of the effect.^{7,8} They can have a larger number of participants than RCTs and therefore provide more evidence on side effects,

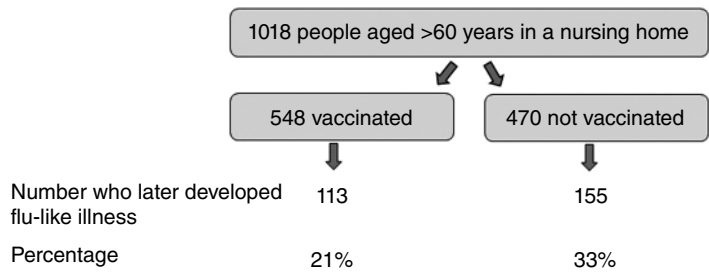


Figure 1.2 Example of an observational study of the flu vaccine. Source: Adapted from⁹.

particularly uncommon ones. However, when the treatment effect is small or moderate, the potential design limitations of observational studies can make it difficult to establish whether a new intervention is truly effective. These limitations are called **confounding** and **bias**.

Several observational studies have examined the effect of the influenza vaccine in preventing flu and respiratory disease in elderly individuals. Such a study would involve taking a group of people aged over 60 years, then ascertaining whether each participant had had the influenza vaccine or not, and who subsequently developed flu or flu-like illnesses. An example is given in Figure 1.2.⁹ The chance of developing flu-like illness was lower in the vaccine group than in the unvaccinated group: 21 versus 33%. However, did the flu vaccine really work?

Assume that vitamin C protects against acquiring influenza. People who choose to have the vaccine might also happen to eat more fruit than those who are unvaccinated (Table 1.1). The difference in flu rates of 5 versus 10% could be due to the vaccine only, the difference in fruit intake (80 vs 15%) only or both together. But we are not interested in fruit intake. If fruit intake had not been measured, it could be incorrectly concluded that the difference in flu rates is due to the vaccination.

When the association between an intervention (e.g. flu vaccine) and a disorder (e.g. flu) is examined, a spurious relationship could be created through a third factor, called a **confounding factor** (e.g. eating fruit). A confounder needs to be correlated with both the intervention *and* the disorder of interest. Even though there are methods of design and analysis that can allow for their effects, there could exist unknown confounders for which no adjustment can be made because they were not measured. There is also **residual confounding** which occurs when the statistical adjustment for a confounder has been insufficient.

There may also be **bias**, where the actions of participants or researchers produce a value of the trial endpoint that is *systematically* under- or over-reported in one trial group. In the flu example, the clinician or carer could deliberately choose fitter people to be vaccinated,

Table 1.1 Hypothetical observational study of the flu vaccine.

	1000 people aged ≥60 years	
	Vaccinated N = 200	Not vaccinated N = 800
Eat fruit regularly	160 (80%)	120 (15%)
Developed flu	10 (5%)	80 (10%)

Box 1.1 Confounding and bias

Confounding represents the natural relationships between the physical and biochemical characteristics, genetic make-up and lifestyle/behavioural habits that may affect how an individual responds to a treatment.

- It cannot be removed from a research study, but known confounders can be measured and therefore allowed for in a statistical analysis.

Bias is a study design feature that affects how participants are selected, treated, managed or assessed, which systematically distorts the results in one trial group more than another.

- It can be prevented, but human nature sometimes makes this difficult.
- It is difficult to allow for bias in statistical analyses because it often cannot be measured.

Randomisation within a clinical trial minimises the effect of confounding and some biases.

believing they would benefit the most. Also, the vaccinated group may be people who chose to go to their family doctor and request the vaccine. It is therefore possible that people who were vaccinated had different lifestyles and characteristics than unvaccinated people. The effect of the vaccine could be over-estimated if the vaccinated people are less likely to acquire the flu than the unvaccinated ones.

When designing trials, it is a useful exercise to imagine being a participant, someone who allocates/administers the trial interventions, or someone who assesses outcome measures, and consider whether there is anything that can be done that can distort the results. Then consider how this could be minimised or avoided by the trial design and conduct.

Confounding is sometimes called a form of bias because both affect the results. However, it is useful to distinguish them (Box 1.1).

1.5 Types of clinical trials

Clinical trials are broadly categorised into four types (Phase I, II, III and IV), largely depending on the main aim (Box 1.2).

<i>Phase I</i>	<i>Phase II</i>	<i>Phase III</i>	<i>Phase IV (real-world data)</i>
<i>Safety/toxicity</i>	<i>Efficacy</i>	<i>Efficacy</i>	Effectiveness
Pharmacology	Safety	Safety	Long term safety
	Adherence	Adherence	Uncommon side effects
		Quality of life	New indications
		Health economics	

Words in italics indicate the common primary focus.

Efficacy and effectiveness are used interchangeably. However, efficacy is sometimes used in a clinical trial setting where the participants could be a select (often motivated) group with high adherence to the trial intervention. Effectiveness applies to routine practice, which better reflects the use of the treatment in the target group and the extent

of adherence among them. The magnitude of benefit associated with a new therapy is sometimes higher when using efficacy than effectiveness.

Traditionally, there has been a sequence from phase I to phase III involving separate trials, particularly for pharmaceutical drugs. To reduce the entire evaluation period, phases can be combined within the same protocol (for example, phase I/II or phase II/III).

Box 1.2 Types of trials

Phase I

- First time a new drug or regimen is examined in humans ('first-in-human' studies), and also used to examine a licenced drug for a new indication (disorder) or a new combination of therapies.
- Few participants (often <50 but can sometimes be >100 if covering multiple disease subtypes).
- Primary aims: check that the toxicity profile is acceptable; find a dose (e.g. drug or radiotherapy) that is tolerable; examine the biological and pharmacological effects.
- Patients or healthy volunteers are monitored closely.

Phase II

Often 30–100 people (but may be larger in common disorders).

- Aims to obtain a *preliminary* estimate of efficacy and further evidence of harms.
- May be single group or randomised (multiple groups), including a control therapy.
- Results are used to help design a confirmatory phase III trial; or the efficacy evidence is good enough to change practice for rare disorders or when there is clear unmet need

Phase III

- Should (must) be randomised and with a comparison (control) group.
- Relatively large (usually several hundred or thousand people).
- Aims to provide a *definitive* answer on whether a new treatment is better than the control group (**superiority**), or similarly effective (**equivalence**) or not materially worse (**non-inferiority**) but with other advantages
- Used to obtain a marketing authorisation from a regulatory agency or for market access for a new drug or medical device (pivotal trial).

Phase IV (post-marketing, surveillance or real-world studies).

- Relatively large (usually several hundred or thousand people).
- Used to continue to monitor efficacy and safety in the population once the new treatment has been adopted into routine practice.
- Not usually randomised, but there are pragmatic randomised phase IV studies that compare currently used interventions.
- Based on participants in the general target population, rather than the selected group who agreed to participate in a phase III trial.

1.6 Key design features

Clinical trials have common fundamental design characteristics (Figure 1.3).

Inclusion and exclusion criteria

Specifying which participants are recruited is done using an **eligibility list**: a set of **inclusion and exclusion** criteria which each participant has to fulfil before they can take part. Common criteria include age range, having no serious co-morbid conditions, the ability to give consent, having no known contraindications to the therapy, and no previous or recent exposure to the trial treatment or others that are similar to it. They should have unambiguous definitions to make recruiting participants easier.

Having many criteria produces a homogenous group that is more likely to respond to the treatment in a similar manner, making it easier to detect an effect if it exists, particularly small or moderate effects. However, the trial results might not be easily generalisable. A trial with few criteria will have greater generalisability but more variability, perhaps making it more difficult to show that the treatment is effective and sometimes only large effects can be detected easily.

Studies increasingly use biomarkers or molecular/genetic testing (profiling) of blood, urine or tissue samples to select only patients who are more likely to benefit from a new treatment that has been developed specifically for that type of patient (**precision/personalised medicine**).

Experimental/investigational treatment group

The new intervention could have been developed using prior studies or laboratory research. Importantly, its description and delivery should be clear enough for other people to replicate it. Investigators expect to show that a new intervention is more effective than the control group (**superiority**), or it has an effect that is similar (**equivalence**) or 'not much worse' (**non-inferiority**).

Control (comparator) group

The outcome of participants given the new intervention can be compared with that in a control group. A **control** group normally receives the current standard of care, no intervention or placebo (see Blinding below). Treatment effects from RCTs are therefore always comparative. The choice of the control intervention depends on the availability of alternative treatments, what is recommended in local or national clinical guidelines, or the requirements of regulatory agencies or organisations responsible for market access. When an established treatment exists, it can be unethical to give a placebo instead because this deprives some participants of a known benefit.

Randomisation and allocating participants

The randomisation process ensures that each participant has the same chance of being allocated to a group as everyone else (Box 1.3). Neither the participant or research team can influence which intervention is given. This minimises the effect of both known *and unknown* confounders, and thus has a distinct advantage over observational studies in which statistical adjustments can only be made for known confounders that have been measured. Randomisation does not produce *identical* groups; there will always be small differences because of chance variation.

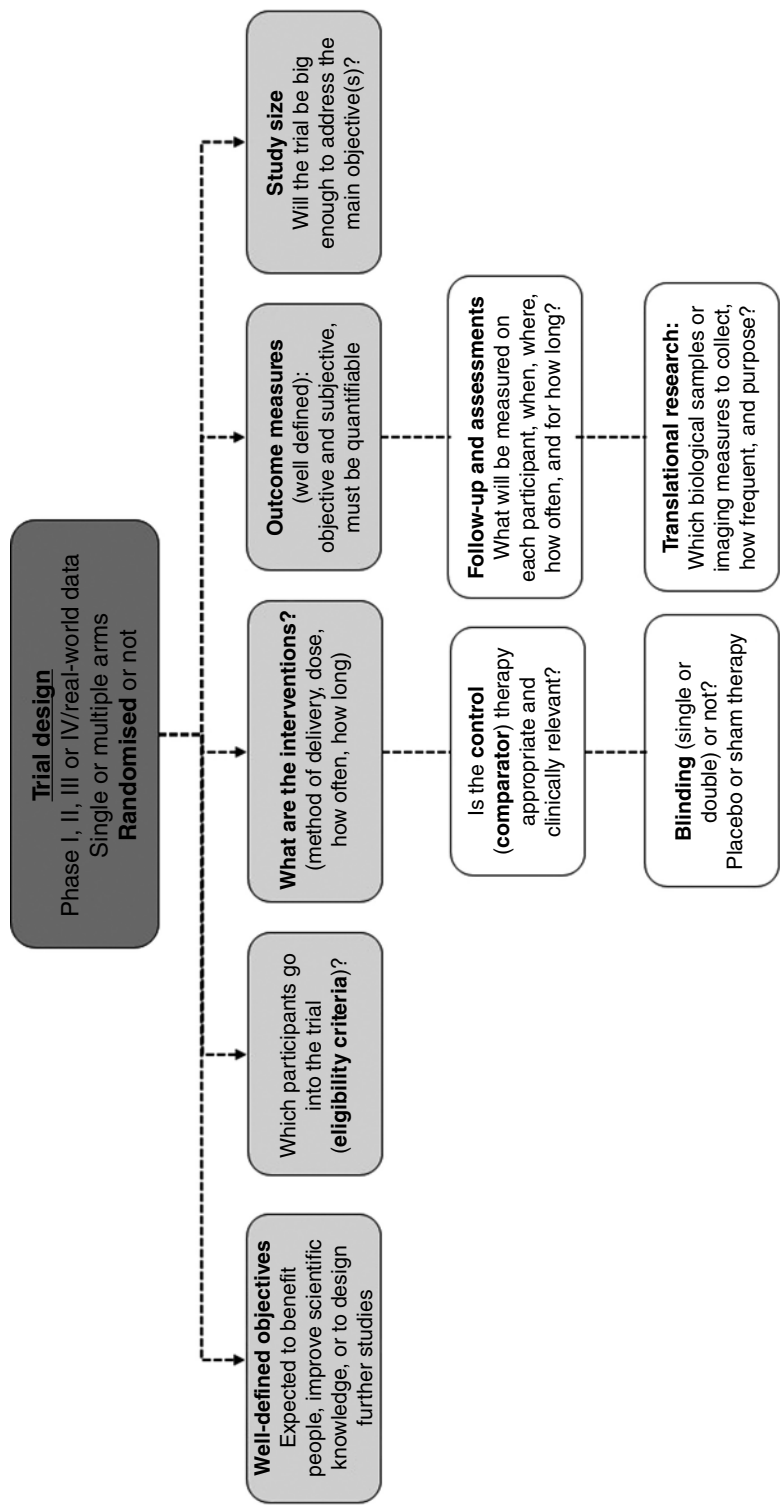


Figure 1.3 Overview of trial design features. Examples of objectives are: "To determine whether Intervention A reduces the risk of dying among people who have had a stroke"; "To evaluate whether Therapy B is non-inferior to standard of care in people with chronic lung disease"; "To investigate whether Drug A has potential efficacy and acceptable toxicities in patients with multiple sclerosis".

Box 1.3 ‘Make everything the same’ except the interventions: randomisation

- If the trial groups being compared have fundamental different characteristics, behaviours and disease features (confounding and bias), it is not easy to determine whether an experimental intervention has been effective or not.
- The most reliable approach to evaluate an intervention is to make the trial groups *the same*.
- Randomly allocating participants produces groups that are similar with regard to all characteristics *except the trial interventions*.
- The only systematic difference between the two groups should be the treatment given.
- Therefore, the trial results should only be due to the effect of the new treatment i.e. confounding or bias have not spuriously produced or hidden an effect.
- In addition to ensuring baseline characteristics are similar, how participants are managed, followed up and assessed should also be the same between the trial arms to avoid/minimise bias.

Randomisation minimises confounding and some bias. Examples of bias are:

- **Selection bias** can occur if choosing a particular participant for the trial is influenced by knowing the next treatment allocation.
- **Allocation bias** involves giving the trial treatment that the investigator or participant feels might be most beneficial. This can be avoided if randomisation is done through a central office or computer system because the researcher has no control over the process (called **allocation concealment**).

Random number lists can be generated using, for example Microsoft Excel; there are commercial or freely available software that perform randomisations; or someone with computing skills can create a bespoke system. Box 1.4 shows randomisation methods. Many investigators use 1:1 allocation (half the individuals receive the experimental intervention, and half receive the control intervention). In some cases, unequal allocation (e.g. 2:1) is used, in order to get more information, such as side effects, about the new intervention.

The choice of method depends on the size of the trial and number of stratification factors. Simple randomisation or permuted block randomisation should produce well-balanced characteristics in very large trials (several thousand). Stratified randomisation using permuted blocks is acceptable if the investigators wish to use random numbers and there are few (prognostic) stratification factors. But in small trials (e.g. <100 participants) or when many factors are used, this method could produce groups where one or more factors are noticeably imbalanced by chance. Minimisation handles several stratification factors easily.

Outcome measures (endpoints) and collecting data

An **outcome measure (endpoint)** is something that enables a disorder, symptom or risk factor to be quantified. Outcome measures can reflect efficacy (benefits), harms, health-related quality of life and adherence.

Phase I trials will have several outcome measures usually with a focus on adverse events. Most phase II and III trials typically have one to three **primary endpoints**[#]

[#] Generally used to justify a change in practice or further studies.

Box 1.4 Methods for allocating participants within RCTs

Random allocation

- **Simple randomisation:** analogous to tossing a coin (heads: Treatment A; tails: Treatment B). In reality a randomly generated list of numbers between 0 and 9 is used: give Treatment A if 0 to 4, and Treatment B if 5 to 9.
- **Permuted block randomisation:** simple randomisation performed using varying even-numbered block sizes to achieve complete balance in each block (e.g. with block size 6, there will be three participants given Treatment A and three given Treatment B).
- **Stratified randomisation:** either of the above methods are used within each level of at least one (stratification) factor known to be strongly associated with the trial outcomes. Investigators want to be assured that these factors are well balanced (e.g. with sex, there are separate randomisation lists for men and women, ensuring a similar number of men and women appear in each trial group).

Dynamic allocation or minimisation

- **Minimisation** also ensures balance for several factors. The treatment allocation for the first few participants (e.g. 20) can be made using simple randomisation. Then the allocation for each subsequent participant depends on the distribution of the stratification factors among those who have already been randomised: e.g. if there are more patients already allocated to Treatment A, the method allocates the next person to Treatment B.
- Minimisation is sometimes considered to not reflect true randomisation because it does not use random numbers. This is addressed by using a high probability (80 or 90%) of being allocated to the next treatment, instead of with certainty.
- Minimisation should not be used for single centre trials.

What matters is that it is not possible for anyone to predict or influence the next treatment allocation: randomness = unpredictability.

(considered the most important and clinically relevant), each corresponding to a **primary objective**. There are also several **secondary endpoints** and **objectives** (considered as supporting evidence), and perhaps some exploratory endpoints. There are five key areas of outcome measures (see Box 2.1, page 16).

Outcome measures come from various sources (Box 1.5). In many situations, there will already be established key endpoints (e.g., the occurrence of coronary heart disease and stroke in type 2 diabetes). In other areas such as in Alzheimer's disease and dementia, there may be no single established primary endpoint because there is an array of imaging tests and clinical assessments. For chronic disorders such as depression, respiratory lung disease and psoriasis, the outcomes need to reflect patient-relevant changes in symptoms due to the new intervention.

Collecting data electronically is increasingly used. Participants enter their own data in real time using personal electronic devices at home, instead of staying longer in clinics to complete paper forms. Many hospitals use electronic records so patient and outcome data can be extracted from those systems; and several countries have regional or national databases and registries that contain outcomes such as cancer occurrence and deaths.

Box 1.5 Sources of data

- Case report forms completed by research staff
- Face-to-face or telephone interviews with participants
- Self-completed questionnaires (handed or posted back to the researchers)
- Face-to-face interviews or questionnaires completed by a proxy for the study participant (e.g. carer or relative)
- Personal electronic devices (mobile phones, apps, internet weblinks and wearables)
- Biological samples (e.g. blood, urine, saliva or tissue)
- Imaging tests (e.g. X-ray, ultrasound, CT, PET or MRI scans)
- Health records from family or primary care clinics, or hospitals
- Local, regional or national registries/databases* that routinely record population data on, for example deaths and cause of death, occurrence of a disorder or how a specific disorder is treated and managed. Trial participants would need to be 'linked' to these databases.

*These can come from governmental, commercial or academic institutions.

Blinding (placebo)

Randomisation minimises some bias, but not all. Participants and researchers often have expectations associated with a particular treatment that affect how people respond to it, or how the researcher manages or assesses them. In participants, this bias is called the **placebo effect**: the psychological ability to affect their own health status. The effect of bias could result in participants who receive the new intervention appearing to do better than the control group, but the difference is not really due to the action of the new treatment.

Clinical trials are described as **double-blind** if neither the participant nor anyone involved in giving the treatment, or managing or assessing the participant, is aware of which treatment was given. In **single-blind** trials, usually only the participant is blind to the treatment he/she has received, but occasionally the participant is aware but the assessor is blind to treatment.

A placebo has no known active component. It could be a tablet, capsule, saline injection, sham surgical procedure, sham medical device or any other intervention that is meant to resemble the experimental intervention, but it has no known effect on the disorder of interest.

An example illustrating the importance of placebos comes from a trial of patients with osteoarthritis of the knee. At the time, such patients often underwent knee surgery (arthroscopic lavage or débridement). Patients typically reported less pain afterwards but there was no clear biological rationale for this. A single-blinded randomised trial¹⁰ comparing these two surgical procedures with sham surgery (skin incision to the knee) provided no evidence that they reduced knee pain.

Using placebos needs to be justified in any clinical trial, approved by an independent ethics board, and participants are fully aware that they may be assigned to the sham group.

When blinding is not possible, it is best to use outcome measures that do not rely on human judgement. For example, in a trial evaluating hypnotherapy for smoking cessation, a **subjective endpoint** would be to ask participants whether they stopped smoking at 1 year. However, some continuing smokers deliberately misreport their smoking status,

leading to bias. An **objective endpoint** would be to measure urinary cotinine, as a marker of current smoking status, which is specific to tobacco smoke inhalation and less prone to bias than self-reported habits.

Participant follow-up

Once participants enrol in a trial, outcome measures must be obtained directly from them and/or via health records. The type and timing of assessments which produce the endpoints should be:

- similar between the trial groups (to avoid bias) when possible;
- able to detect changes in the disorder at clinically important time points that reflect its natural history (e.g., assessments could be frequent early on for trials of advanced/acute disorders, or infrequent for preventive trials when the disorder can take years to occur); and
- pragmatic, where possible, to reflect routine practice, and not be arduous for participants or local research staff by requiring too many extra clinic visits or tests.

Following are examples of bias that may occur:

Withdrawal bias: participants in one trial group are more likely to withdraw from the study than other arms (e.g. due to side effects).

Follow-up bias: there is less follow-up on one trial arm than another, or the reasons for lack of follow-up are very different between the arms and associated with the trial interventions.

Study size

Justifying how large a trial needs to be is an important component of a grant application or review by an ethics board or regulatory agency. Phase I and II trials tend to be relatively small, while phase III trials must be large enough to provide a clear answer. It is better to have an overly large study than one that is not quite big enough which produces equivocal results from which no clear recommendation can be made (resources could then have been wasted).

There is nothing wrong with small trials as long as they are well-designed and conducted. They can be quick and relatively cheap, only require a few centres and produce an answer in a short time frame. This avoids spending too many resources on looking for a treatment effect when there really is none. However, if a positive result is found a larger confirmatory study is often needed.

Contrary to common belief, sample size estimation is not an exact science. If the target is 100 patients, researchers often believe that recruiting 95 or 105 is a problem: it is not. Sample sizes provide an *approximate* trial size to aim for, and reaching a few less or more than the target is not an issue at all.

Translational research (bio- and imaging markers)

There have been major technological advances in laboratory and imaging techniques, including cheaper and quicker genetic sequencing of normal and abnormal cells (e.g. cancer). This has led to the discovery of many biomarkers measured in tissue, blood, serum, skin, urine and saliva, and markers from imaging methods using modern and sophisticated CT, PET and MRI scans. The marker could be a biochemical measure in biospecimens, familial genetic trait or abnormality (germline), or genetic abnormality/mutation

in for example tumour tissue or circulating tumour cells (somatic). New treatments have thus been developed and evaluated in clinical trials that are targeted to a specific marker measured at baseline (i.e. before treatment), in which patients with the marker (e.g. their cancer cells have a particular genetic mutation) benefit much more than with standard (non-targeted therapies). This is the fundamental premise of **precision (personalised or risk-stratified) medicine**, i.e. there is no longer a 'one-size-fits-all' approach to treating several disorders but instead subsets of patients derive greater benefit from treatments that target a specific feature (marker) of their disease. Biomarker-driven trials will become more common as we obtain a greater biological understanding of various disorders, including how they develop and their natural history. They can be used for:

- correlative science
- biomarker discovery
- prognostic and predictive markers
- genomics
- identifying drug resistance
- disease monitoring

Biomarkers can be used to determine eligibility for a trial and also which trial treatment a person receives, or they can be evaluated at the end of the study to see if outcomes are influenced by the marker (**predictive marker**). Biomarkers can also be used to gain a better understanding of the mechanisms of action of a new therapy or evaluated to see whether they can be used to forecast patient prognosis or outcomes (**prognostic marker**).

Several funding organisations expect planned translational research to be included within a clinical trial proposal.

1.7 Summary points

- Clinical trials are essential for evaluating new methods of disease detection, prevention and treatment.
- Observational studies can provide useful supporting evidence on the effectiveness and safety of an intervention, but they have limitations.
- Clinical trials, especially when randomised, are considered to provide the strongest evidence.
- Randomisation minimises the effect of confounding and bias, and blinding further reduces the potential for bias.
- Understand the key design features: Participants/Population, Intervention, Control and Outcome measures (PICO).
- PPIE should be used to help design, conduct and interpret trials.
- Translational research (measurement of bio- and imaging markers) can play an important role in precision medicine as well as understanding mechanisms of action and predicting patient outcomes.

References

1. Hackshaw AK. *A concise guide to observational studies*. 1st edn, Wiley-Blackwell, 2015.
2. <https://www.jameslindlibrary.org/lind-j-1753/>
3. Medical Research Council. Streptomycin treatment of pulmonary tuberculosis. *BMJ* 1948; **2**:769–782.

4. Gross PA, Hermogenes H, Sacks HS, Lau J, Levandowski RA. The efficacy of influenza vaccine in elderly persons. *Ann Intern Med* 1995; **123**:518–527.
5. Govaert TME, Thijs CTMCN, Masurel N *et al*. The efficacy of influenza vaccination in elderly individuals. *JAMA* 1994; **272**(21):1661–1665.
6. Egger M, Schneider M, Davey Smith G. Meta-analysis: spurious precision? Meta-analysis of observational studies. *BMJ* 1998; **316**:140–144.
7. Collins R, MacMahon S. Reliable assessment of the effects of treatment on mortality and major morbidity, I: clinical trials. *Lancet* 2001; **357**:373–380.
8. MacMahon S, Collins R. Reliable assessment of the effects of treatment on mortality and major morbidity, II: observational studies. *Lancet* 2001; **357**:455–462.
9. Patriarca PA, Weber JA, Parker RA *et al*. Efficacy of influenza vaccine in nursing homes. Reduction in illness and complications during an influenza A (H3N2) epidemic. *JAMA* 1985; **253**:1136–1139.
10. Moseley JB, O'Malley K, Petersen NJ *et al*. A controlled trial of arthroscopic surgery for osteoarthritis of the knee. *N Engl J Med* 2002; **347**(2):81–88.