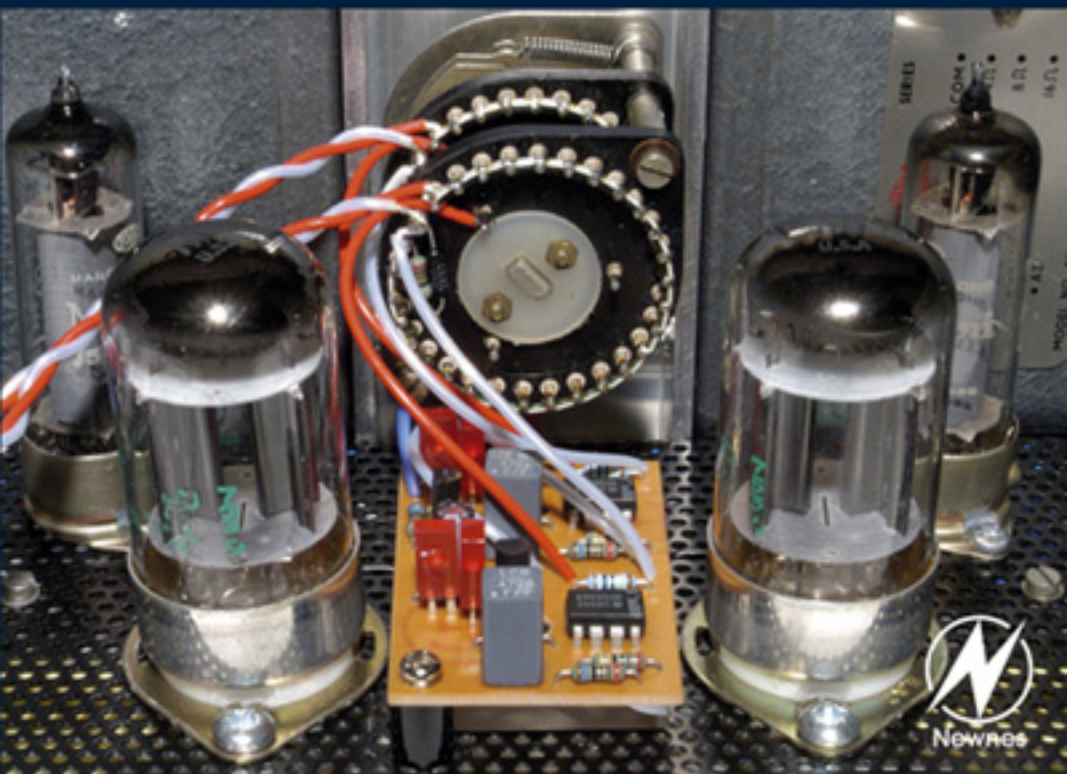




Building Valve Amplifiers

MORGAN JONES



Building Valve Amplifiers

This Page is Intentionally Left Blank

Building Valve Amplifiers

Morgan Jones



ELSEVIER

AMSTERDAM • BOSTON • HEIDELBERG • LONDON
NEW YORK • OXFORD • PARIS • SAN DIEGO
SAN FRANCISCO • SINGAPORE • SYDNEY • TOKYO

Newnes is an imprint of Elsevier



Newnes

Newnes

An imprint of Elsevier

Linacre House, Jordan Hill, Oxford OX2 8DP

200 Wheeler Road, Burlington, MA 01803

First published 2004

Copyright © 2004, Morgan Jones. All rights reserved

The right of Morgan Jones to be identified as the author of this work has been asserted in accordance with the Copyright, Designs and Patents Act 1988

All photographs (including front cover) by Morgan Jones

No part of this publication may be reproduced in any material form (including photocopying or storing in any medium by electronic means and whether or not transiently or incidentally to some other use of this publication) without the written permission of the copyright holder except in accordance with the provisions of the Copyright, Designs and Patents Act 1988 or under the terms of a licence issued by the Copyright Licensing Agency Ltd, 90 Tottenham Court Road, London, England W1T 4LP. Applications for the copyright holder's written permission to reproduce any part of this publication should be addressed to the publisher.

Permissions may be sought directly from Elsevier's Science and Technology Rights Department in Oxford, UK: phone: (+44) (0) 1865 843830; fax: (+44) (0) 1865 853333; e-mail: permissions@elsevier.co.uk. You may also complete your request on-line via the Elsevier homepage (<http://www.elsevier.com>), by selecting 'Customer Support' and then 'Obtaining Permissions'.

British Library Cataloguing in Publication Data

Jones, Morgan

Building valve amplifiers

1. Audio amplifiers – Design and construction
2. Audio amplifiers – Maintenance and repair

I. Title

621.3'81535

ISBN 0750656956

Library of Congress Cataloguing in Publication Data

A catalogue record for this book is available from the Library of Congress

For information on all Newnes publications
visit our website at <http://books.elsevier.com>

Typeset by Integra Software Services Pvt. Ltd, Pondicherry, India.

www.integra-india.com

Printed and bound in Great Britain

Working together to grow
libraries in developing countries

www.elsevier.com | www.bookaid.org | www.sabre.org

ELSEVIER

BOOK AID
International

Sabre Foundation

CONTENTS

<i>Preface</i>	vi
<i>Acknowledgements</i>	viii
Part 1 Construction	1
1 Planning	3
2 Metalwork for poets	57
3 Wiring	94
Part 2 Testing	169
4 Test equipment principles	171
5 Faultfinding to fettling	258
6 Performance testing	306
<i>Index</i>	349

PREFACE

This practical book is intended to bridge the gap between a theoretical circuit diagram and relaxing to the sound of a reliable, well-built amplifier. It is the complement to “Valve Amplifiers” (which has **equations** in it).

One possible way of building an amplifier is to choose the most expensive components on offer, then assemble them on a custom-built chassis using that year’s choice of designer wire. However, the author assumes that you are holding this book because you want to know how to build a valve amplifier that is significantly better than you could afford for the same price ready-made. For that reason, some of the physics that you slept through at school will reappear, but with the valuable bonus that it will allow you to make reasoned choices that improve quality or save money – possibly both.

It is undoubtedly easier to do metalwork in a fully equipped machine shop, complete with folding machine, drill press, lathe, mill, and copious hand tools. Nevertheless, it is perfectly possible to do good work with a power drill in a stand and a few carefully selected hand tools. In addition to the standard techniques, a number of “cheats” will be shown that allow you to produce work of a standard that appears to have been done by a precision machine shop. This will enable your creation to be a thing of beauty that can be proudly displayed.

The rules for good audio construction are not complex. It’s just that there are rather a lot of them. Once the logic is understood, a good layout comes naturally.

Even the most carefully considered designs need a little fettling once built. Test equipment ranges from cheap DVMs to oscilloscopes and spectrum analysers with Gigahertz of bandwidth, to PC-based virtual instruments. They all cost money, but if you understand the operating principles, you can choose which features are worth paying for, which can be safely ignored, and how to use what you can afford to its best advantage.

Startlingly, years of experience don't make the author any less frightened at the instant of first switch-on. Accidents do happen, but there are ways of minimising the quantity of smoke. Sometimes, an amplifier is stubborn and just doesn't **quite** work properly, requiring genuine faultfinding.

This book is distilled from years of bludgeoning recalcitrant electronics, thumping metal, and sucking teeth at the price of test equipment . . .

ACKNOWLEDGEMENTS

Many individuals and organizations have assisted in the preparation of this book. Some even knew that they were doing so. Inspiration came from a variety of sources. Some, simply because they asked significant questions, others because they displayed good technique, others because they generously made equipment available for photography, and yet others because they demonstrated techniques to be avoided . . .

In particular, the author would like to thank Paul Leclercq for proofreading (once again), and for the varied and generous assistance given by Brian Terrell.

PART I

CONSTRUCTION

This Page is Intentionally Left Blank

CHAPTER I

PLANNING

In this first chapter, we will investigate how to plan the mechanical layout of a valve amplifier. At this stage, freedom of choice is unlimited, whereas later it will be restricted. So it is important that the choices and compromises made now are the best ones. Whilst good planning will not save a poor design, poor planning can certainly ruin a good one.

Chassis layout

Valve amplifiers use a number of large components that must be positioned relative to one another such that the connecting wires between each component are as short as possible, but that the components and their associated wires do not interfere with each other. Chassis layout breaks down into the following considerations:

- Electromagnetic induction: Minimising hum induction from chokes and transformers into each other and into valves.
- Heat: Output valves, etc. are hot and must be cooled. Conversely, capacitors run cool, and must be kept that way.

- Unwanted voltage drops: All wires have resistance, so the wiring must be arranged to minimise any adverse effects of these voltage drops.
- Electrostatic induction: Minimising hum from AC power wiring is not often a problem, because even thin conductive foil provides perfect electrostatic screening.
- Mechanical/safety: Achieving an efficient chassis arrangement that is easily made, maintained, and used.
- Acoustical: Almost all components are microphonic, but valves are the worst. We should consider which components are most sensitive to vibration, and minimise their exposure to it.
- Aesthetic: The highest expression of engineering is indistinguishable from art. If you have a superb chassis layout, it will probably look good. Conversely, if it looks poor, it is probably a poor layout...

We have a seven-dimensional problem. A poor transistor amplifier might be able to hide behind the fence of negative feedback, but valve amplifiers using an output transformer cannot usually tolerate more than 25 dB of feedback before their stability becomes distinctly questionable. Consequently, layout is critical to performance.

The large components are generally the mains transformers, output transformers, power supply chokes, power supply capacitors, and valves. The traditional way of deciding how to position them is to cut out pieces of paper of the same size as the components and shuffle them around on a piece of graph paper. Alternatively, the lumps themselves can be arranged and glanced at for a few days until the best layout presents itself.

Even better, components can be shuffled around and a chassis designed using an engineering drawing package on a computer, with the enormous bonus that a template of the layout can be printed with all the fixing holes precisely positioned, saving errors in marking out. Although it takes time to draw a valve holder or a transformer precisely, it has only to be done once

and you will quickly build up a library of mechanical parts. In consequence, the author has almost forgotten how to perform traditional marking out using a scribe, ruler, and square . . .

It is vital to make the chassis large enough!

This point cannot be emphasised too strongly. Achieving neat construction on a cramped chassis requires a great deal more skill and patience than on a spacious chassis. There are many considerations that must be taken into account, so it is vital that this stage is not rushed. Each of the following design considerations might not make a great difference in itself, but the sum of their effects is the difference between a winner and an “also ran”.

Electromagnetic induction

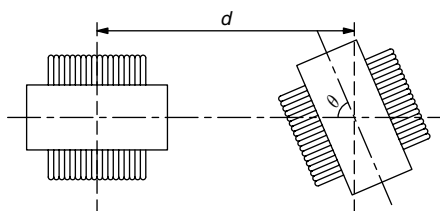
Almost all of the larger components either radiate a magnetic field or are sensitive to one. Not all of a transformer’s primary flux reach the secondary. Leakage flux can induce currents into wires such as valve grids. Whether, or not, these currents are significant depends on the signal level and source impedance at that point, so output valves are less of a problem than the input stage.

Coupling between wound components

Wound components such as transformers and chokes can easily couple into one another, so hum can be produced by a mains transformer inducing current directly into an output transformer. Fortunately, the cure is reasonably simple, and may be summarised by a simple ratio whose value must be minimised.

$$\text{induction} \propto \frac{\cos \theta}{d^3}$$

The angle θ and distance d are shown in the diagram (see [Figure 1.1](#)).

**Figure 1.1**

Orienting transformers for minimum coupling

Rotating transformer cores by 90° ($\cos 90^\circ = 0$), so that the coil of one transformer (or choke) not aligned with the other is very effective, and typically results in an immediate 25 dB of practical improvement. Even better, if one coil is driven from an oscillator whilst the interference developed in the other is monitored (oscilloscope or amplifier/loudspeaker), careful adjustment of relative angles can often gain a further 25 dB.

Because coupling decays with the cube of distance [1], as the distance between offending items is increased, the interference falls away rapidly. However, simply increasing the **gap** between two adjacent transformers from 6 to 25 mm does not materially reduce the interference, because the transformers are typically 75 mm cubes, and the spacing that applies is the distance between centres, which has only changed from 81 to 100 mm, resulting in only 5.5 dB of theoretical improvement. In practice, when the transformers are this close, coupling does not obey the cube law very well because the transformers do not see each other as point sources, so a 3 dB reduction, or less, is more likely.

Although smoothing chokes are gapped, and therefore inevitably leaky, they do not generally have an appreciable AC voltage across them. So their AC leakage is low, and they can often be used to screen output transformers from the mains transformer. The exception to this rule is the choke input power supply which has a substantial AC voltage across its

choke, so its leakage field is capable of inducing currents into surrounding circuitry.

The core of a poorly designed mains transformer can easily be saturated by the large current pulses drawn by a large reservoir capacitor in combination with a semiconductor rectifier, producing a particularly noisy leakage flux, and this can be quickly identified by a search coil (see [Figure 1.2](#)).

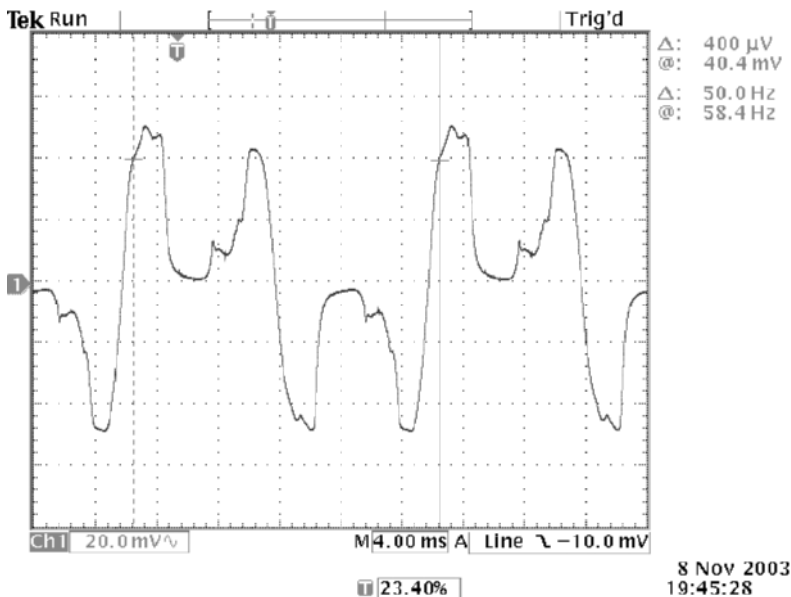


Figure 1.2

Leakage flux caused by an unscreened saturating mains transformer

A screening can that totally encloses a transformer significantly attenuates the high frequency content of any leakage flux. The second guilty party is a fifty-year-old mains transformer whose core material has deteriorated, but this transformer is totally enclosed by a thin steel screening can. Note that the waveform has smoothly rounded edges, indicating far less high frequency content than the previous example (see [Figure 1.3](#)).

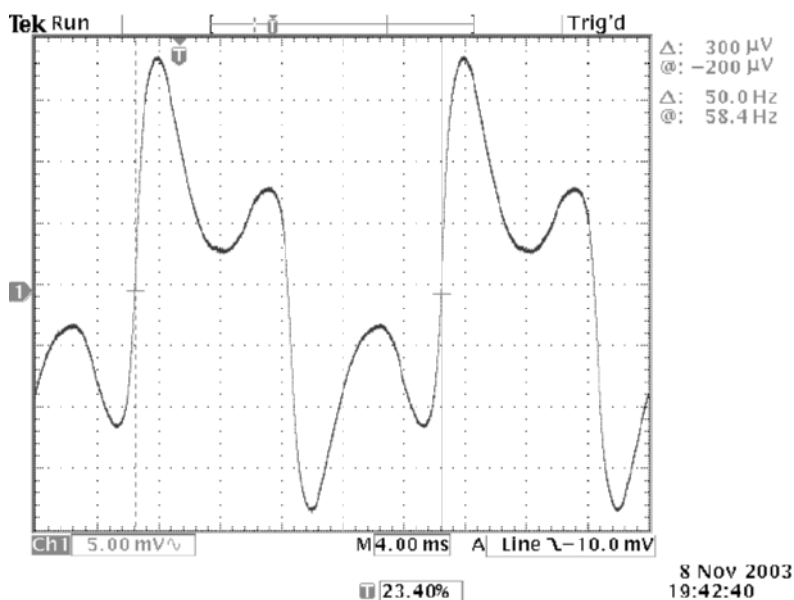


Figure 1.3

Leakage flux caused by a screened saturating mains transformer

Transformers and the chassis

All transformers leak flux. The question is whether that is a problem. If an output transformer leaks flux into the aluminium chassis of a power amplifier, it probably isn't a problem because aluminium doesn't conduct magnetic flux. But a mains transformer leaking flux into the steel chassis of a pre-amplifier **is** a problem because the steel chassis passes the flux into sensitive signal circuitry. Fortunately, because $\mu \approx 0$ for non-magnetic materials, but $\mu > 5000$ for iron, even a small gap is able to prevent flux leaking into the chassis. A 1–2 mm paxolin sheet is ideal, but this may be hard to find. So a (more expensive) alternative is the polystyrene sheet sold by model shops. (Although un-etched FR4 PCB material would be ideal, if the copper side were to be placed in contact with an aluminium chassis, it would electrolytically corrode it, whereas the face-up copper would oxidise quickly and become unsightly.)

Beware that practical toroidal transformers leak flux and that mounting a toroid directly onto a steel chassis is just asking for hum problems. However, the danger of accidentally creating a shorted turn is even greater. Toroids are secured by a conductive screw pulling a large conductive washer onto the top of core to clamp the transformer tightly to the chassis. Accidentally connecting the washer or screw to the chassis by any means other than the bottom of the central mounting screw would form a shorted turn that could **destroy** a power transformer.

Beam valves and mains transformers

Beam valves focus their current into thin sheets that pass largely unintercepted between the horizontal wires of g_2 . This means that a vertical beam deflection would affect g_2 current, and because g_2 is typically supplied from a finite source resistance, Ohm's law ensures that this would change V_{g_2} , thus changing I_a . One way of deflecting electrons is with a magnetic field, such as the hum field from a transformer. Hum due to beam deflection can be minimised by applying Fleming's left hand rule, and ensuring that the electron beam is never at right angles to the leakage flux from the transformer. When considering induction between two transformers, it does not matter which transformer is rotated, so long as the coils are at 90° to one another. With beam valves, only one orientation is ideal with respect to a nearby mains transformer (see [Figure 1.4](#)).

The valve is shown in two positions, both the same distance from the centre of the mains transformer, and both with correct beam orientation relative to the leakage flux from the transformer. However, leakage flux tends to be concentrated on the axis of the coil, and would also induce hum into the control grid's circuit, whereas the alternate position has much lower flux density. (Diagrams of this form portray higher flux density by having more lines in a given area.)

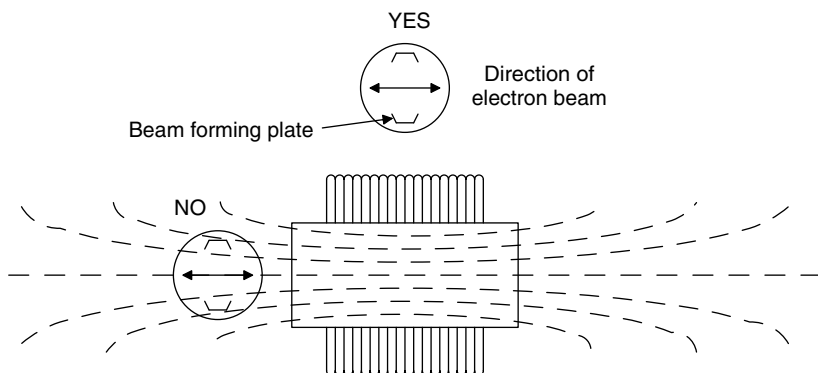


Figure 1.4
Beam valves and mains transformers

Input valves are very sensitive to hum fields, and should be placed at the opposite end of the chassis to the mains transformer.

In theory, output transformers should leak less because they operate at a lower flux density (to reduce distortion). In practice, probing output transformers and mains transformers with a small search coil failed to show the expected difference. The quality of the transformer seems to be the over-riding consideration, rather than its use. Thus, a Leak TL12+ push-pull output transformer leaked more flux than a good-quality modern output transformer in a single-ended amplifier, despite the latter being gapped.

Although, as expected, leakage flux at 90° to the coil's axis cancels to zero, leakage at the edges of the coil can be comparable with that on axis because the coil's outermost turn is so far away from the (flux-concentrating) core (see [Figure 1.5](#)).

Heat

Heat is the enemy of electronics. Output transformers and chokes are usually quite cool, so they can move towards

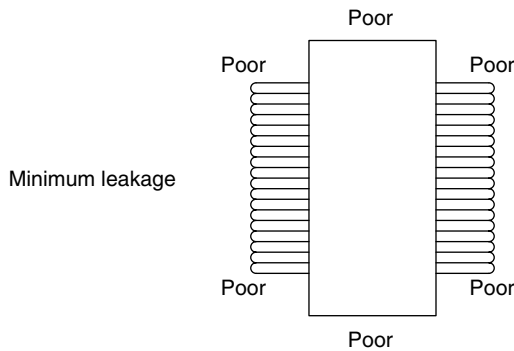


Figure 1.5

Transformers leak most flux along the axis of their coils and where the edges of their coils are furthest from the core

the centre of the chassis if necessary (creating a mechanical problem, but we will consider this later). Mains transformers are generally warm, and it is usually best to mount them towards the edge of the chassis.

At best heat shortens component life and causes components to drift in value. At worst, it causes fires. And we intend to use valves, which are deliberately heated . . .

Modes of cooling

There are three methods of transferring heat from one place to another.

Conduction is the most efficient method of heat transfer and requires a conducting material to bond the heat source physically to its destination. An ideal conductor would transfer the heat without any temperature drop between source and destination, and materials having free electrons (electrical conductors) such as copper and silver are particularly good. As an example, the body of a power transistor must conduct the heat generated by the (very much smaller) silicon device to an external heatsink efficiently (see [Figure 1.6](#)).

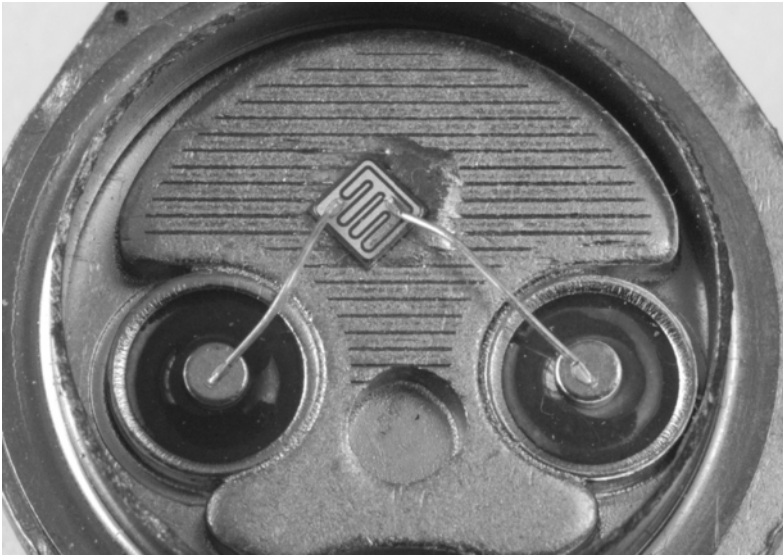


Figure 1.6

The inside of a 2N3055 15 A 115 W power transistor. Note the relative size of the silicon device and its wires compared to the case

Convection relies on the unrestricted movement of a fluid (gas or liquid) between source and destination. Fluid is heated at the source, expands, and is displaced by denser cooler fluid, forming a convection current which continuously pushes hot fluid away and draws cool fluid towards the source. Convection efficiency can be increased in two ways:

- If a greater volume of fluid moves per second, more heat can be lost, so a fan or pump greatly improves heat transfer. The pumped water in a water-cooled internal combustion engine or central heating system is an example of forced convection.
- We choose a fluid that requires more energy for a given temperature rise. Liquids are better than gases, and water is particularly good. However, although some transmitter valves have water-cooled anodes, practical pipe diameters and pump power mean that there is a limit to the maximum flow rate, and thus the heat that can be transferred. Even better, changing

the **state** of a material requires a great deal of energy, and converting water at 100°C to steam at the **same** temperature requires **six** times as much energy as heating the same mass of water from 20 to 100°C, so the largest valves are steam cooled.

Radiation (strictly, electromagnetic radiation) does not require a physical medium between heat source and destination, but it is the least efficient means of transferring heat. Radiation losses are governed by Stefan's law:

$$E \simeq \sigma T^4$$

Where:

E = power per unit area

σ = Stefan's constant $\approx 5.67 \times 10^{-8} \text{ W/K}^4/\text{m}^2$

T = absolute temperature = °C + 273.16

Because heat loss is proportional to the fourth power of temperature, particularly hot bodies, such as the Sun, can transfer heat quite effectively by radiation.

Valve cooling and positioning

Output valves are hot, and must be allowed to cool properly. Typical audio valves place the anode within an evacuated glass envelope, so the anode cannot lose heat by convection. The supporting wires from the anode to outside connectors are quite thin, so the anode cannot lose heat by conduction. The only remaining method of heat transfer is radiation.

Radiation obeys reciprocity, in that a good reflector is a good insulator, so domestic kettles are shiny metal or white plastic. Conversely, matt black absorbs well, so anodes are often blackened to allow them to radiate more efficiently.

Although we think of glass as being optically transparent, some light is inevitably lost in transmission. Similarly, glass is

imperfectly transparent to infrared radiation. Radiation that is not transmitted is absorbed and heats the glass. So some of the received heat from the anode can be lost by convection if air is allowed to flow freely past the valve envelope. Very roughly, the envelope splits heat losses from the valve equally between convection and radiation. Thus, efficient cooling requires that we consider how a valve can radiate, and how easily convection currents can flow past the envelope.

Power valves should be separated from one another by a spacing of $1\frac{1}{2}$ envelope widths or more, otherwise they heat each other by radiation. But a useful trick can be applied to the radiant heat received by a valve. Many valves, particularly small-signal valves, have an anode with quite a narrow cross-section. Received radiant heat is proportional to the area seen at the destination, so rotating a valve to present a narrow cross-section to the source can reduce heating from adjacent power valves. By reciprocity, if the source also has a narrow cross-section, rotating the source to present a small area to the destination also reduces transmission. This technique is particularly useful for circuits such as differential pairs or phase splitters where electrical considerations dictate that the two valves be close together, but it can also be used to reduce heat received from a nearby power valve (see [Figure 1.7](#)).

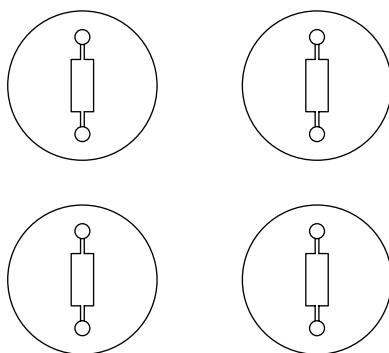


Figure 1.7

Careful anode orientation allows closer spacing along the anode's narrow axis

Placing power valves in the middle of a chassis is not likely to be a good idea because the chassis severely restricts convection currents. Mounting a valve horizontally can improve convection efficiency because it exposes the envelope to a larger cross-section of cooling air. However, this could conceivably cause a hot control grid to sag onto the nearby cathode, with disastrous results. So check the manufacturer's full data sheet to see if there are any strictures about mounting position. If it doesn't cause other problems, align the socket so that the plane of the grid wires is vertical [2], preventing them from being able to sag onto the cathode.

Early valves have cylindrical electrodes, so viewed down their axis, electrons strike the anode equally from all points of the compass. Since the electron density is equal at all angles, the anode temperature is also equal at all angles. However, beam tetrodes do not have axial symmetry, and direct their beam of electrons along a single diameter, causing the anode to heat unequally. Bearing in mind that the glass envelope converts half of the radiant heat loss from the anode to convection loss from the envelope, a horizontally mounted beam tetrode should be rotated on its axis so that the hottest parts are at the sides, allowing them to be efficiently cooled by convection, rather than at top and bottom. As an example, see GEC's recommended alignment for the KT66 [3] (see Figure 1.8).

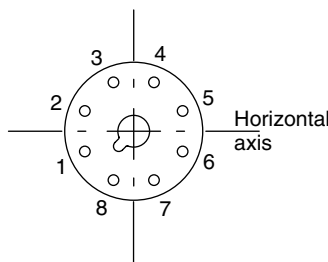


Figure 1.8

GEC recommended orientation of valve base for horizontally mounted KT66

Because the sections of a KT66's envelope between pins 1 and 2, and 5 and 6 are the hottest, when mounting a push-pull pair of KT66 vertically, it makes sense to ensure that pins 7 and 8 of one valve face pins 7 and 8, or 3 and 4, rather than 1 and 2, or 5 and 6.

Mounting valve sockets on perforated sheet greatly assists cooling by allowing a convection current to flow past the valve, but as the hole area of such sheet is typically only $\approx 40\%$, it is still not perfect. If better cooling is needed, the valve socket can be centrally mounted on a wire fan guard, and if even that isn't sufficient to keep the envelope temperature below the manufacturer's specified maximum, a low-noise fan can be mounted on pillars underneath the chassis using the same screws that secure the fan guard (see [Figure 1.9](#)).

Using the chassis as a heatsink

Some small components such as power resistors and regulator ICs, unavoidably generate significant heat. Resistors are

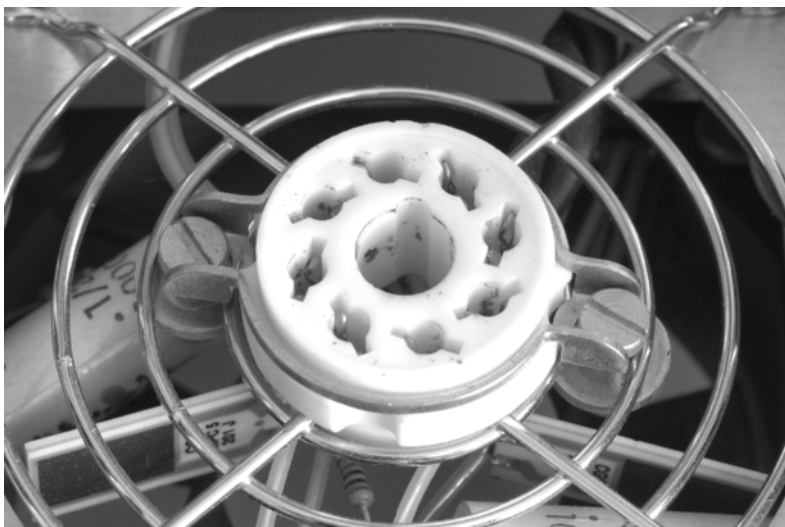


Figure 1.9

With care (and large washers), a valve socket can be mounted on a fan guard, allowing almost unrestricted airflow

commonly mounted on stand-offs to allow an unimpeded air flow, and regulators are often fitted with small finned aluminium heatsinks. Neither of these strategies is ideal because they attempt to lose heat by convection to still air enclosed by the chassis.

Convection cooling only works if there is a free flow of cooler air past the hot component. Once the cooling flow stops, the hot component is surrounded by still air, which is a good insulator, and its temperature quickly rises. Eventually, the still air begins to lose its heat by conduction to the surrounding chassis, and an equilibrium results with a high internal air temperature and a hot component.

A high air temperature within the chassis is undesirable because:

- The components causing the high air temperature are unnecessarily hot, and even though they may have been designed to withstand heat, their working life is inevitably reduced.
- Electrolytic capacitors are extremely sensitive to heat, and a very rough rule of thumb is that their working life halves for each 10 °C rise in temperature. Capacitor manufacturers' data sheets include extremely useful charts that allow lifetime predictions to be made from ambient temperature and capacitor ripple current, so it is well worth looking up the full data sheet for your particular capacitor at the manufacturer's website.
- Components having a critical value, such as in equalisation or biasing networks will drift away from their optimum value as a consequence of heating from the air.

Ultimately, we can only lose heat to the surrounding air, and the larger the surface presented to the air, the more efficiently it will cool. This means that the best way to cool components is to ensure that they are thermally bonded to the chassis, using a thin smear of heatsink compound.

Heatsink compound is not a particularly good conductor of heat, but it is far better than air. The purpose of heatsink

compound is to fill the tiny insulating **air** gaps that result from placing two imperfectly smooth surfaces together. Too much compound worsens cooling. Most people (and this includes manufacturers) use too much. A thin, even smear applied to both mating surfaces is all that is required. Beware that the screws bonding the bracket to the chassis can loosen at operational temperature. So check them for tightness when fully warmed-up, otherwise the compromised thermal bond can cause the semiconductors to be hotter than expected, and worse, the erratic bond can cause drift.

A useful secondary advantage of mounting aluminium-clad or TO-220 resistors directly onto the chassis is that they provide convenient mounting tags for other components. It may seem unnerving to touch a chassis with hotspots due to local heat-sinking, but this technique minimises the internal air temperature, and thus minimises the heating of sensitive components.

Efficient convection requires a free flow of cool air to replace hot air. Although most designers recognise the importance of allowing adequate ventilation by providing holes in the top of a chassis near hot components, air must also be free to enter the chassis from the bottom if an efficient convection current is to flow. Thus, the ideal solution is to make the **entire** underside of the chassis from perforated steel or aluminium, and support it on feet ≥ 20 mm high (to allow air to flow into the underside of the chassis) (see [Figure 1.10](#)).

If necessary, individual components can be cooled even more efficiently by bonding them directly to a finned heatsink fitted to the outside of the chassis, as is common with transistor amplifiers. Although this technique most obviously springs to mind when considering large power amplifiers or power supplies, precision pre-amplifiers should have their internal temperature rise minimised in order to prevent equalisation networks drifting in value, so it could be worth bonding anode load resistors to an external heatsink.



Figure 1.10

Perforated sheet and tall feet allow excellent cooling

When finned heatsinks are used, it is most important to orient them correctly. Heatsink manufacturers specify thermal resistance ($^{\circ}\text{C}/\text{W}$) with the fins vertical in free air because this maximises the surface area available to the natural cooling convection current flowing up the fins. Despite this, the author has lost count of the number of commercial amplifiers having horizontal heatsinks. Heatsinks cost money, so why degrade their performance? The prettiest example that the author can find of the importance of the correct fin direction is a motorcycle (see [Figure 1.11](#)).

The engine is a V-twin. The pistons are identical and run in barrels that are detachable from the crankcase. Despite the

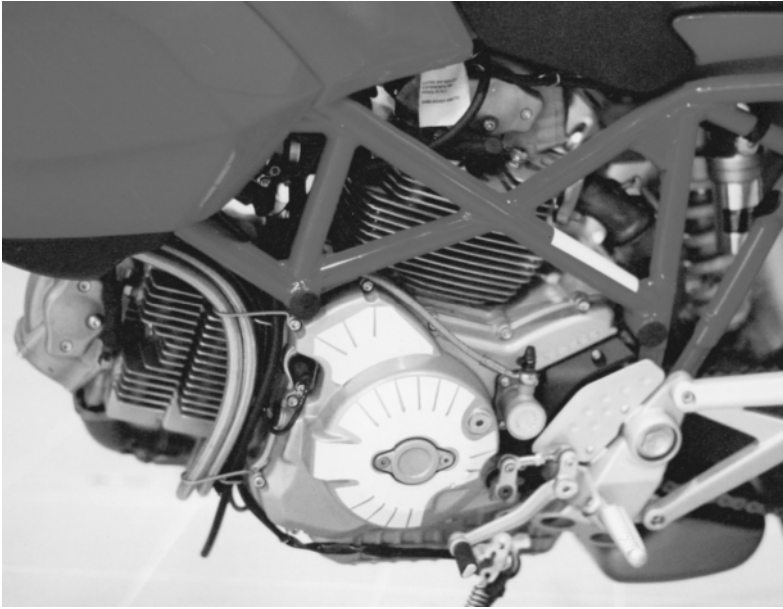


Figure 1.11

Ducati Monster has barrels with fins aligned in the direction of air flow

increased production cost, the two barrels are different, one having longitudinal fins, the other latitudinal, and this is done solely to optimise cooling. Your amplifier's fins do not have a forced 100 mph horizontal convection current, so they need to be vertical.

Cold valve heater surge current

The resistance of a conductor such as a valve heater filament changes significantly with temperature in accordance with the following equation:

$$R_t = R_0(1 + \alpha t)$$

Where:

R_0 = cold resistance

R_t = resistance at temperature t

α = thermal coefficient of electrical conductivity (0.0045 per °C for pure tungsten)

t = temperature change

Normally the temperature rise of conductors in electronics is too small for the effect to be noticeable, but a thoriated tungsten filament operates at ≈ 1975 K, so that at an ambient temperature of 20°C (293 K), its resistance is much lower, and theory predicts that it very briefly draws 8.6 times the operating current. In practice, the surge current is limited by the current capability of the supply, perhaps to only half the predicted value.

Indirectly heated valves operate their filaments at ≈ 1650 K, resulting in a theoretical surge current of ≈ 7 times the operating current, although measurements suggest that a ratio of $\approx 5:1$ is more appropriate. More significantly, the thermal inertia of the cathode sleeve slows heating, so the surge current lasts for a few seconds, and could be sufficient to blow a poorly chosen mains fuse. Further, it would not be prudent to use heater wiring of a rating only just sufficient to cope with the steady-state current if long-term reliability were required.

Wire ratings

Wires have resistance, so passing a current causes self-heating ($P = I^2 R$). If the wire becomes too hot, the insulation may catch fire, so it is important to ensure that the wiring is rated appropriately for the current to be passed. This means that wire current ratings are determined by ambient temperature, ability to cool, and the temperature rating of the insulator. As a very rough guide to wires having PVC insulation:

<i>Conductor diameter (mm)</i>	<i>Maximum current (A)</i>
0.6	1.5
1.0	3.0
1.7	4.5
2.0	6.0

Component catalogues are good sources of information on the suitability of a particular wire.

High-temperature insulators (such as PTFE) allow a higher current rating than might be expected from a given conductor diameter, but the penalty is higher resistance, so voltage drops along that wire will be proportionately higher, and this can become significant in the capacitor/rectifier/transformer loop of a power supply.

Arcing and insulation breakdown should not be a problem at the voltages found within most valve amplifiers, but it is still advisable to maintain 2–3 mm separation between conductors with a high voltage between them, unless each conductor is insulated. As an example, hardwiring with bare wire touching the surface of a capacitor with wound polypropylene tape insulation is not advised, but a sleeved wire touching the same capacitor would be unlikely to constitute a safety hazard.

Unwanted voltage drops

All wires have resistance which is proportional to their length, and inversely proportional to their cross-sectional area. Although the currents in valve circuitry are typically quite low, making the voltage drops proportionately low, once we consider that we want a signal to noise ratio of >90 dB, the voltage drops caused by small resistances become significant.

The highest currents, and therefore highest voltage drops occur in the loop from transformer via rectifier to reservoir capacitor and back again. A capacitor input filter draws pulses of current at twice mains frequency from the transformer that are typically four to six times greater than the DC load current. It is essential that the wires carrying these pulses are as low resistance, and therefore, as short as possible, which means that the rectifier and associated reservoir capacitor should be close

to its mains transformer. It's the same logic that puts the battery in a car's engine compartment. (The original Mini was intended to be rear wheel drive from a 500 cc longitudinally mounted engine, but this proved to be underpowered, so the only way to fit a larger engine was to mount it transversely, but that left insufficient room for the battery, forcing it into the boot, requiring a long, very thick wire to the starter motor.)

When an output stage enters Class B, it draws current pulses at twice the audio signal's frequency from the power supply. To prevent these pulses breaking into driver circuitry and increasing distortion, the loop from audio load to reservoir capacitor should be made as small as possible. This means that output transformers should be close to their HT capacitor.

The potential effects of unwanted voltage drops usually determine the 0 V signal earth scheme, often known colloquially as earthing or grounding. There are fundamentally two methods of dealing with earthing:

- “Earth follows signal”: The 0 V signal earth wire follows the path of the signal. In order to minimise unwanted voltage drops along this (necessarily long) wire, it has a large cross-section so this brute force strategy leads to 1.6 mm (16 swg) tinned copper bus-bars.
- Star earth: All connections to the 0 V signal earth are made at a single point. Because the distance between individual connections is so small, the impedance is small, so unwanted voltage drops are also small.

The significance of these two wiring schemes is that the choice between them needs to be made at a very early stage. “Earth follows signal” tends to produce long slim mechanical layouts, whereas an ideally implemented star earth tends to produce square or even circular layouts centred about the star earth. The significance of unwanted voltage drops is so great that it

dominates the gross mechanical layout of an example later in this chapter, whilst detailed electrical implications will be covered in Chapter 3.

More powerful amplifiers are heavier, and in an effort to split the weight into two manageable chassis, you might consider having a remote power supply, necessitating an umbilical cable to connect between the chassis. Beware that the heater wiring linking the two chassis needs to be of a much higher current rating than normal in order to reduce the unwanted voltage drop over this increased distance, possibly necessitating an umbilical connector with a larger diameter cable entry...

Electrostatic induction

Electrostatic coupling is capacitive coupling. Minimising the capacitance between two circuits minimises the interference. Remembering the equation for the parallel plate capacitor:

$$C = \frac{A\epsilon_0\epsilon_r}{d}$$

Where:

A = area of plates

d = distance between plates

ϵ_0 = permittivity of free space $\approx 8.854 \times 10^{-12}$

ϵ_r = relative permittivity of dielectric between the plates

We should aim to attack all parts of this equation to minimise capacitance. Reducing plate area means keeping wires short and crossing them at right angles, whilst increasing plate distance means keeping parallel wires apart. In terms of chassis layout, this means that the output valves should be reasonably close to the output transformer and the driver circuitry should be reasonably close to the output valves.

In a high-impedance circuit such as a constant current sink, a chassis-mounted transistor causes a conflict between thermal and electrostatic considerations because the collector is invariably connected to the transistor case. Screwing a TO-126 transistor such as an MJE340 to an earthed chassis using an insulating kit adds $\approx 6\text{ pF}$ shunt capacitance, completely negating any efforts we may have made in designing a wide-bandwidth constant current sink. In this instance, we are forced to compromise our thermal considerations and lose heat directly to the air within the chassis using a small finned heatsink.

Fortunately, capacitance to chassis from aluminium-clad or transistor-style resistors used as anode or cathode loads is not a problem at audio frequencies. Because one end of the resistor is at AC earth, but the capacitance is distributed along the resistor, the total capacitance to earth is effectively halved, and because the impedances are lower than in a constant current sink, the stray capacitance becomes insignificant.

Valve sockets

There is a variety of different types of valve socket available for a given valve type (see [Figure 1.12](#)).

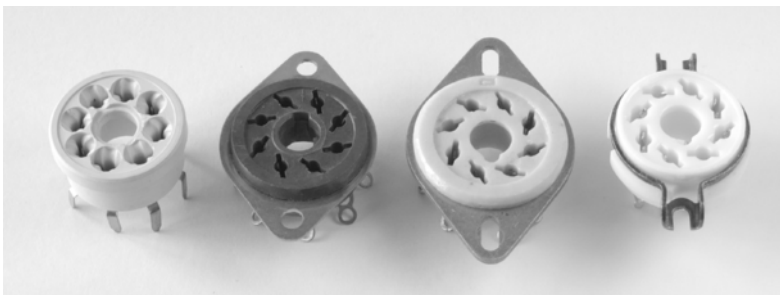


Figure 1.12

Octal sockets. Note that the chassis-mount types all have different mounting centers

All three of the chassis-mounting International Octal sockets in the photograph have different spacings for their securing screws, so it is vital to make a firm decision about which socket type is to be used, and whether the circuit is to be hard-wired or PCB. One point that may be worth considering when choosing Octal sockets is that NOS McMurdo phenolic sockets have the same hole spacings as Loctal sockets, allowing an easy change from 6SN7 to 7N7 at a later date.

Another consideration is positioning of heater wiring. Taking the B9A base as an example, most valves using this base have their heaters connected between pins 4 and 5, although heating the popular ECC83/12AX7 from 6.3 V requires a **link** between pins 4 and 5, and the other heater wire must be taken across the centre of the socket to pin 9. McMurdo B9A valve sockets were designed to that if the axis of the socket was aligned at 45° to the edge of the chassis, pins 4 and 5 were closest to the chassis edge, minimising hum from heater wiring (see [Figure 1.13](#)).

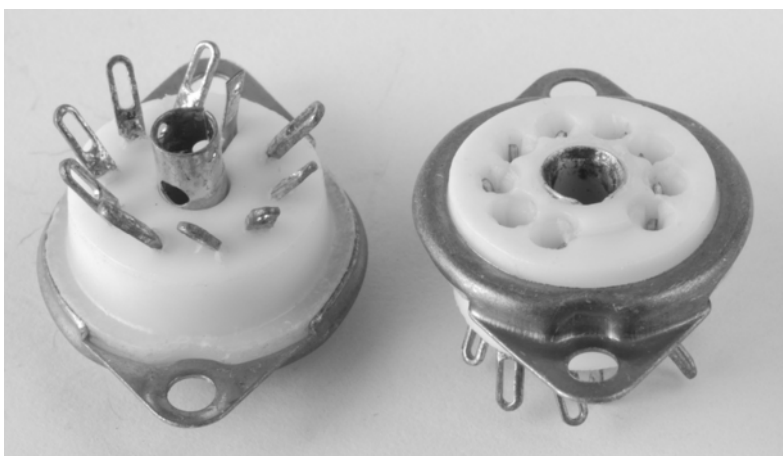


Figure 1.13

McMurdo B9A sockets position the heater pins (4, 5) close to the edge of the chassis provided that the mounting holes are at a 45° angle to the edge

Traditional valve amplifiers always had valve sockets aligned so that their heater pins were closest to the chassis edge, thus minimising the length of exposed heater wiring...

How electrostatic screening works

Electrostatic screening works by placing an earthed conductive barrier between the source of interference and the sensitive circuit (see Figure 1.14).

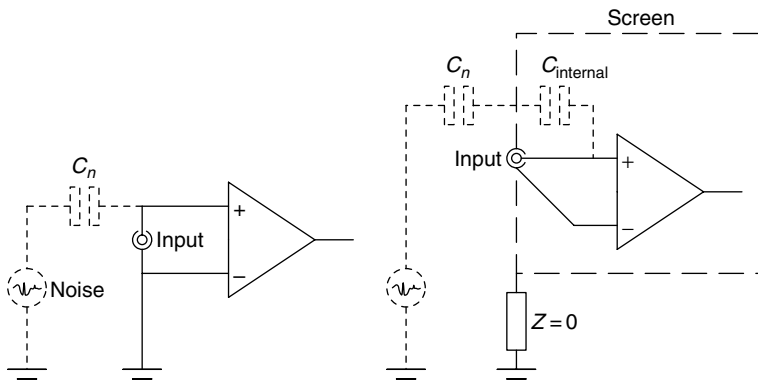


Figure 1.14

Screening breaks one capacitance into two, centre-tapped by an earth

As can be seen from the diagram, the screen diverts noise currents from the source to earth. If there is any impedance between the screen and RF earth, the noise current can develop a noise voltage across it, causing the screen to induce noise into its enclosed (sensitive) circuitry.

Alternatively, the impedance from the screen to RF earth can be considered to be the lower leg of a potential divider, and the capacitance from screen to noise source the upper leg. The lower leg could be a length of conductor, which has inductance, thus forming a 12 dB/octave filter in conjunction with the capacitor. To maximise the effect of screening, the screen must have a low-resistance, low-inductance path to chassis.

When resistances and inductances are connected in parallel, the total value falls, so the screen should ideally contact the chassis at multiple points to minimise impedance and maximise screening. This is why RF designers cut the flanges of the lids containing their circuitry – it ensures that each finger firmly contacts the case, and minimises impedance at RF (see [Figure 1.15](#)).

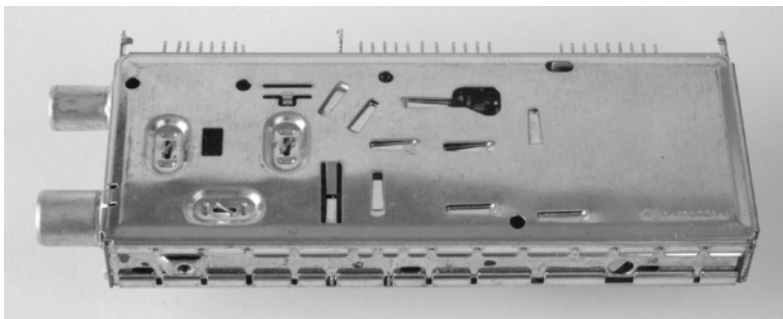


Figure 1.15

Lids with fingers allow multiple low resistance, low inductance contacts

From an audio point of view, screening cans should be firmly screwed to the chassis at multiple points using shakeproof washers to ensure a gas-tight connection that does not deteriorate over the years. Pre-amplifiers using choke interstage smoothing should ideally use oil-filled chokes, not because the oil confers any advantage, but because the metal can needed to contain the oil provides electrostatic screening.

Valves that are sensitive to electrostatic hum can be enclosed by an earthed metal screening can. Valve screening cans are not created equal. Perfect screening would enclose the valve completely, but that would prevent heat from escaping, making the valve significantly hotter, and reducing its life [2]. It is easy to raise a valve's temperature, but lowering it is much harder. Screening cans are therefore a compromise between screening and cooling (see [Figure 1.16](#)).



Figure 1.16
Selection of valve screening cans

From left to right

The first “screening can” was actually manufactured as a heat-sink, and slightly reduces envelope temperature. All that is needed to turn it into a screen is a short wire to bond it electrically to the chassis. The dungaree screening can is a bayonet fitting onto a skirted base and has holes to allow air to flow, which reduces its screening efficiency and temperature rise. Note that the inside of the can is painted matt black to absorb radiant heat from the valve, which is then reradiated by the black outside surface. The third can is for small-signal valves only as airflow is very restricted, but the black paint inside and outside assists cooling. The fourth can is an abomination and deserves to be crushed because not only does it restrict airflow, but it insulates the radiant heat from the valve.

Very occasionally, screening cans are deliberately made to aid heat losses, and have fingers to contact the glass envelope (see [Figure 1.17](#)).

Input sockets

Input sockets should be kept apart from mains wiring (hum) and loudspeaker terminals (instability), but as both of these

problems are caused by capacitance from the input socket to the offending connectors, screening the input socket cures the problem easily if space is limited.

An extremely useful facility on a power amplifier is to add a muting switch that applies a short circuit to the input of the power amplifier (see [Figure 1.18](#)).

Not only does the muting switch enable the amplifier to be plugged to different sources without having to switch it off and then on again, but it is very handy when the amplifier is being tested. The $1\text{ k}\Omega$ resistor protects the output of pre-amplifiers, etc. from the short circuit.



Figure 1.17

Screening can with internal fingers to aid heat dissipation



Figure 1.17 (Contd)

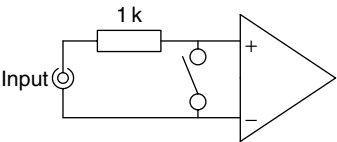


Figure 1.18
This muting switch at the input of the amplifier enables audio leads to be swapped without the necessity of switching the amplifier off

Mechanical/safety

The purpose of the chassis is to support the components mechanically, some of which are heavy, and to **enclose all the dangerous voltages**, thus eliminating the risk of electric shock. The safest form of chassis for the home constructor is the totally enclosed earthed metal chassis. Placing a mains transformer in the centre of a chassis made of folded aluminium is asking for the chassis to sag. Heavy items should be placed towards the edges, where the nearby vertical section adds substantial bracing.

If a heavy item **must** be moved towards the centre of a chassis, a bulkhead or two can be added across the chassis to brace it, giving the further advantage of breaking the inside of the chassis into electrostatically screened compartments. This allows noisy circuitry (rectifiers and smoothing) to be placed in a “noisy” compartment, and sensitive audio in a “quiet” compartment (see [Figure 1.19](#)).

When planning the positions of the major parts, it is important that one part should not obscure electrical connections or securing screws of another unless it can be easily moved aside to gain access. Think carefully about accessibility of screws and nuts. It's time to reconsider if you can't easily fit a spanner or nut-driver onto a nut, or a screwdriver into a screw head. If there's really no choice about nuts being inaccessible, consider securing the part with a single 3 mm plate having tapped holes in lieu of individual nuts. Failure to observe these considerations can make subsequent maintenance extremely difficult. Typical consumer electronic design places maintainability very low on its list of priorities, but not only do you want to be able to maintain your creations, you also want to be able to modify them as your knowledge expands or when better parts become available.

It is often useful to take a modular approach. If a block such as a complete heater supply including mains transformer, rectifi-

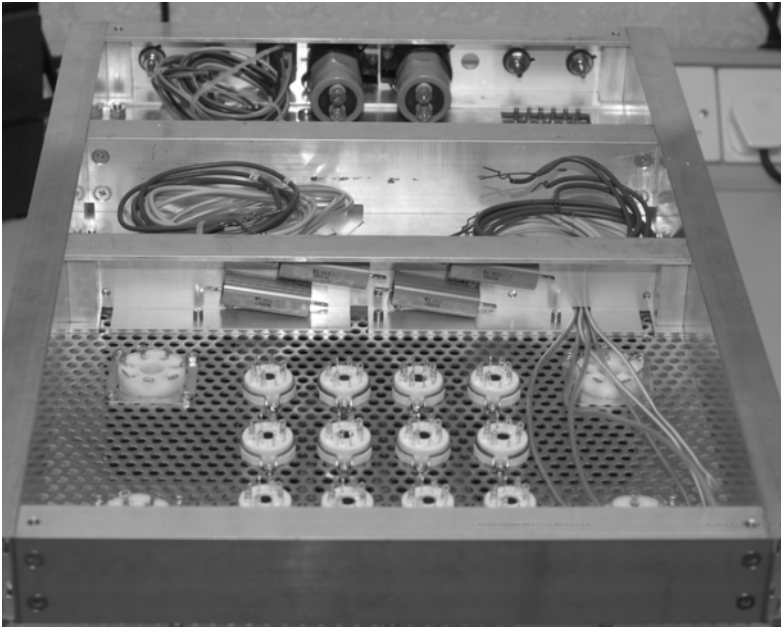


Figure 1.19

Chassis bulkheads allow segmentation into “dirty” or “clean” compartments

cation/smoothing and regulation is made as a module on a sub-chassis, it can be conveniently built and tested outside the main chassis, then installed when ready. A printed circuit board (PCB) is another example of this technique.

Components directly soldered to valve sockets and tags are known as **hard-wired**, and a less obvious technique for power amplifiers is to hard-wire the entire driver circuitry onto a small rectangular plate that fits over a correspondingly sized hole in the main chassis. The great advantage of this technique over mounting valveholders directly into the main chassis is that a complete driver redesign requiring a different number of valves bases of different types simply means replacing the small plate. Even better, the plate can be made of perforated aluminium to assist cooling.

Bear in mind that some components can be quite large. Capacitors vary greatly in size depending on the choice of dielectric, so 100 nF could range from the size of a fingernail to the entire finger! (see [Figure 1.20](#)).

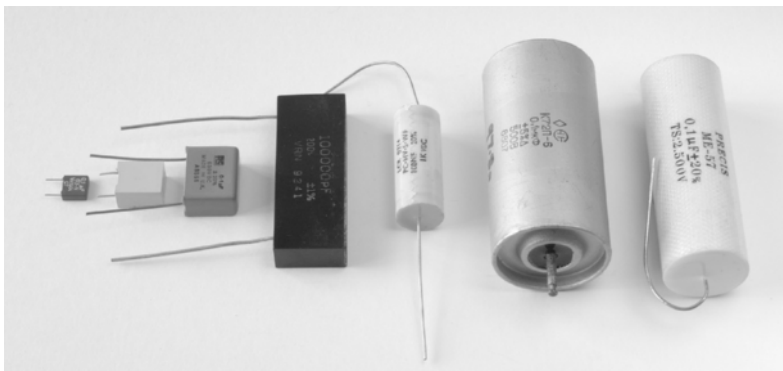


Figure 1.20

A selection of 100 nF non-polarised capacitors. Left to right: 63 V polyester, 400 V polyester, 630 V polycarbonate, 200 V silvered mica, 1500 V polypropylene, 500 V polytetrafluoroethylene, 2000 V mixed dielectric

Although we will look at electrical safety in detail later, electrolytic smoothing capacitors should be considered during planning. The voltage on the can of the capacitor is indeterminate, but is generally near to the potential on the negative terminal. Although the can could be bonded to chassis if the negative terminal is at 0 V, this invariably causes a hum loop, so the can is usually insulated from the chassis. The capacitor therefore has only a single layer of insulation between the HT supply and the outside world, but safety calls for either a double layer of insulation, or one layer plus an earthed metal shroud. (See later for explanation of Class I and II equipment.)

Preventative maintenance means being able to see the signs of impending failure, and catching it **before** it goes bang. This is one reason for having all the power valves clearly visible – if you spot a red-hot anode, you probably can get to the “on/off”

switch before the fault destroys even more components. Internally, if all the components are clearly visible, you might spot a charring resistor or bulging capacitor **before** it destroys other components in addition to itself at the first application of power.

A rather more depressing aspect of preventative maintenance is to consider the consequences of failure of specific components. When traditional electrolytic capacitors fail, they spray soggy paper and foil from their bases. Although modern capacitors are rather tamer and vent in a more controlled fashion, it is a good idea to consider where any vented material might land. It would be unfortunate if a failed LT capacitor vented a conductive spray over a perfectly innocent (powered) HT regulator PCB. Similarly, silicon rectifiers can sometimes fail explosively, and it would be particularly poor planning if the short-lived volcano was directly under a complex wiring loom. Although the author has not tested carbon resistors to destruction, industry wisdom is that they also fail pyrotechnically...

When you work on the amplifier, you will inevitably turn it upside down, so it is helpful if the taller and heavier components can be arranged so that they allow the chassis to stand firmly without rocking, and without tipping onto the delicate valves.

Acoustical

It is unusual to consider acoustical problems in a power amplifier, but the input valves are microphonic, and a flimsy chassis does not help. The author favours a rigid chassis with plenty of bracing, held together with plenty of large screws, even if the result is somewhat stronger than strictly required by the supported weight.

Valves are microphonic because any movement of their control grid structure alters the local electric field and therefore the

number of electrons reaching the anode. When the grid structure is placed closer to the cathode to increase mutual conductance, the same movement becomes relatively greater, so these valves are intrinsically more microphonic, although their more rigid frame-grid structure often tames the problem.

Pre-amplifiers may need deliberate acoustic isolation from structure-borne vibration, and anti-microphonic valve sockets with integral rubber suspension mounts used to be readily available. However, we still need to take wires to the valve base, and unless these are flexible, they form an acoustic short-circuit. A more effective way of achieving isolation uses knicker elastic to float a sub-chassis supporting an entire circuit in the same way that a trampoline is supported. Each anchoring point requires a small rubber grommet to prevent chaffing of the elastic and to prevent the elastic slapping metal when the sub-chassis moves (see [Figure 1.21](#)).

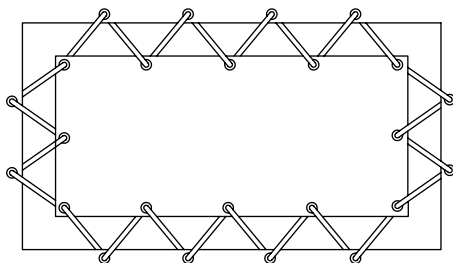


Figure 1.21

Trampoline suspension of sub-chassis using knicker elastic

The internal mechanical resonances of a valve are typically above 500 Hz [4], so these are the frequencies that need to be isolated. The suspended sub-chassis is effectively a 12 dB/oct low-pass filter with $f_{-3\text{dB}}$ at the resonant frequency of the suspension, so if it resonated at 63 Hz (three octaves down), it would theoretically attenuate by 36 dB at 500 Hz. In practice, damping in the springs reduces this figure, but best acoustic

isolation is still achieved by minimising the suspension resonant frequency. The resonant frequency is given by the standard resonance equation:

$$f = \frac{1}{2\pi\sqrt{Cm}}$$

Where:

C = suspension compliance

m = suspended mass

Thus, deliberately adding mass (perhaps lead sheet) to the sub-chassis can improve isolation. Although the resulting suspension is likely to be stiffer than an anti-microphonic valve socket, it is still important to minimise the number of wires to the sub-chassis, use flexible wires, and dress them in loops that are large enough to avoid short-circuiting the suspension acoustically.

Air-borne vibration is far harder to eliminate. The ideal solution would be to enclose the circuitry within a solidly built closed box, internally lined with an acoustic absorbent such as bonded acetate fibre (BAF) or fibreglass building insulation. Unfortunately, there would then be no way for the heat to escape...

Aesthetic

Although placed last on the list, this consideration is probably close to the forefront of your mind as you will want to be proud of the results. The finished project does **not** need to look like a collision between a rat's nest and a supermarket trolley. Glowing glass is pretty, so almost all power amplifiers have their valves at the front where they can be seen. And rightly so.

Form follows function, so right-handed people place the most important control (volume) at the right. Symmetry is usually

pleasing, so a rotary input selector switch could be placed at the left. “Retro” looks are currently very much in fashion, so moving coil meters and hexagonal bakelite knobs are popular.

Blue LEDs are currently popular and will doubtless be considered to be deeply embarrassing in twenty years time. The author is rather fond of dual-colour LEDs (red and green) with relative currents adjusted to give the same orange glow as a valve heater. If the green LED is powered from the heater supply, and the red from the HT, the LED performs a useful monitoring function as well as being pretty.

Practical examples

Application of the previous rules is best demonstrated by examples. Bear in mind that there are many ways of skinning cats, so if you can find a better solution, use it...

Power amplifiers

Let us suppose that we want to make a mono push–pull EL84 amplifier, perhaps a rebuild of a Leak TL12+.

The major items to be positioned are:

- Mains transformer
- Output transformer
- Rectifier valve
- HT reservoir capacitor
- Output valves
- Driver circuitry

Although negative feedback reduces hum, it would be best if there were no hum induced into the output transformer from the mains transformer. If we draw a line between the centres of

the two transformers, and rotate them so that their coils are at 90° to each other, this minimises hum (see [Figure 1.22](#)).

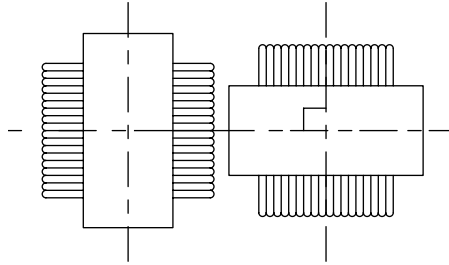


Figure 1.22

Planning a layout: A. Output transformer and mains transformer at right angles to minimise coupling

We can also reduce induced hum by spacing the two transformers apart. As soon as we separate major parts, we ought to think about whether we can use the space in between them, perhaps for the HT reservoir capacitor. The reservoir capacitor draws current through the rectifier from the transformer in a train of 100 Hz high current pulses (see [Figure 1.23](#)).

The resistance of the capacitor/rectifier/transformer wiring loop should be as low as possible, so the ideal position for the rectifier valve is adjacent to the mains transformer, with the reservoir capacitor nearby. The centre tap of the output transformer needs a low-impedance HT supply, so it should be close to the smoothing capacitor. If the HT capacitor is the traditional dual capacitor, it now makes sense for the two transformers to be moved apart just sufficiently that the rectifier and capacitor can be fitted in between with adequate room for cooling (see [Figure 1.24](#)).

The capacitor is heated by radiation from the rectifier, but provided that there is a little separation between the two, the capacitor can cool by convection and only receives radiation proportional to the angle of arc subtended by the capacitor (see [Figure 1.25](#)).

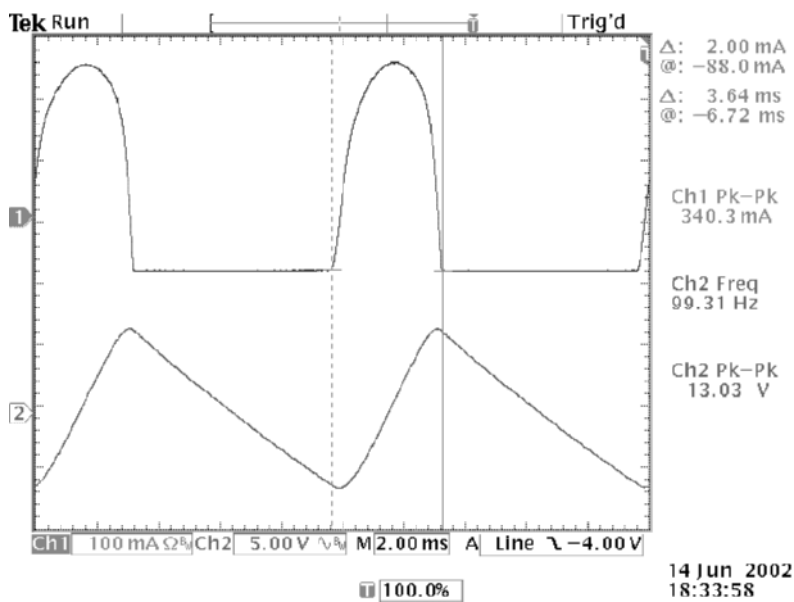


Figure 1.23
Capacitor ripple current is typically 4–6 times DC load current

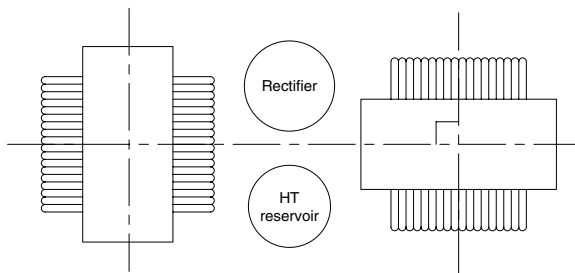


Figure 1.24
Planning a layout: B. Adding rectifier and reservoir capacitor

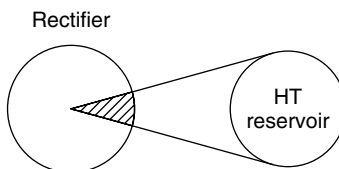


Figure 1.25
Planning a layout: C. Radiated heat from rectifier to capacitor

The output valves must be able to cool, so they should be separated by a distance of one and a half times the valve envelope diameter. But they also need to be reasonably close to the output transformer, and it would be convenient if they were also reasonably close to the driver circuitry. The logical position for the output valves is thus on the side of the output transformer opposite to the rectifier, with the driver circuitry a little further beyond. The driver circuitry is now as far as possible from the mains transformer, and is partly screened from it by the output transformer. The two transformers have their coils aligned at 90° to one another, but this can be achieved with either the coil of the mains transformer or the output transformer pointed at the driver circuitry.

Because the output valves are widely spaced, even if the output transformer's coil axis is pointed away from them, they are likely to be in the leakage field from the edges of the coil. A separation of one valve envelope diameter from the output transformer significantly reduces the localised leakage field from the edges of the coil. This generalisation works because larger valves produce higher power and require a larger transformer. If the output valves were beam tetrodes such as KT66, an even greater spacing might be desirable.

Given the previous caveats about output transformer orientation, it makes sense to align the mains transformer's coil at 90° to the audio circuitry.

The valves in the driver circuitry should be separated from the output valves in order to keep them cool, and room is needed for their associated components, so moving them to the edges of the chassis allows their heater wiring to be pushed into the corners of the chassis, and coupling capacitors (which can often be quite large) can be conveniently placed between driver and output valves. We have now arrived at the final layout (see [Figure 1.26](#)).

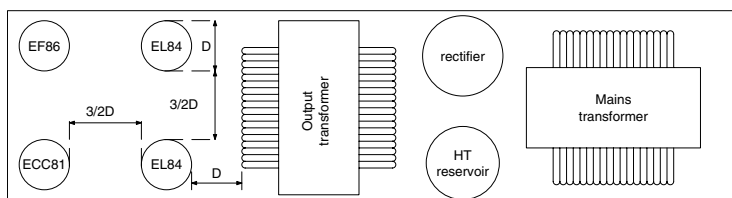


Figure 1.26

Planning a layout: D. Final layout

Readers who are familiar with the Leak TL12+ will recognise that this layout is much longer and slimmer than the original Leak. The Leak layout requires longer wires, but a square chassis is more accessible for wiring, so the reduced wiring time would have cut production costs.

Meters and monitoring points

Power valves are inevitably used by circuits capable of destroying them. To illustrate this point, consider the output valves in a power amplifier. They are connected to a substantial power supply via an output transformer having low DC resistance. If the valves are cathode biased, the cathode resistor will protect them in the event of a fault, but grid bias provides no protection. Grid bias inevitably requires some means of monitoring cathode current.

If you are in a “retro” mood, you will fit a handsome round black moving coil meter and a rotary switch with a black pointer knob allowing each individual cathode current to be monitored, and if you are feeling really keen, you will add positions on the switch allowing key HT voltages to be monitored. The monitoring is important, so you will want to fit the meter and associated switch towards the front of the amplifier. The meter has a stray magnetic field, so you won’t want it to be close to beam valves such as KT66, 6L6, KT88, or 6550 because it might affect screen grid current. The switch needs to be operated easily without obscuring the meter, and its selection

and the meter reading need to be seen simultaneously, so they have to be mounted on the same plane. If you decide to put a meter and switch on the top of the chassis, make sure there's room to operate the switch without your fingers being burned by a nearby (hot) output valve.

Another possibility is to fit 4 mm chassis sockets as low-voltage monitoring points, and it makes sense to put these sockets as close as possible to the voltage being monitored. Again, make sure you don't put them so close to hot valves that you burn your fingers when connecting the lead from your DVM. If the amplifier is push-pull, output currents must be balanced, and you will probably want to be able to connect two DVMs simultaneously, so remember to provide a 0 V test point for **each** DVM.

Alternatively, the traditional method was to use one meter to measure total current, and one to measure imbalance. Finding two suitable moving coil meters is tricky, but postage stamp sized DVMs that only need a $\frac{1}{4}$ " hole for mounting are now readily available, so they can be fitted far more easily than a moving coil meter, or you could simply provide dedicated test points for total current and imbalance.

Sympathetic power amplifier recycling

Rather than build an amplifier from scratch, you might prefer to recycle an old amplifier's chassis and transformers, but use driver circuitry of your own design, saving an awful lot of metalwork (but not much money). If you take this approach, please do it sympathetically. Randomly gouging holes and leaving others unused looks unsightly – if it's worth doing to please the ear, it's worth taking a little trouble to please the eye (see [Figure 1.27](#)).

The amplifier started life as a Leak Stereo 20, but the coupling capacitors were leaky and the electrolytics had dried out, so the author decided to recycle the chassis with a simpler driver



Figure 1.27

Sympathetic recycling allows a new amplifier to retain its classic looks

stage, using one less valve. This left a vacant $\frac{3}{4}$ " hole, which was conveniently filled by a threaded DIN socket mounted on a small plate which was secured by the holes for the original B9A valve socket. Similarly, the holes for the original dual electrolytic capacitors were carefully enlarged for their polypropylene replacements, and a hole was punched in front of the mains transformer for the third polypropylene capacitor. The transformer links for selecting mains voltage and loudspeaker impedance were hardwired and their connectors replaced with blank panels. The transformers were mounted on $\frac{1}{16}$ " paxolin board to prevent leakage flux from reaching the (steel) chassis.

Classic amplifiers invariably had appalling connectors, ranging from lethal mains plugs to awkward loudspeaker connections, and this Stereo 20 was no exception. Fortunately, the hole for the large 3-pin Bulgin connector could be enlarged to take a flange mounting IEC inlet with integral fuse, and the hole for the mains outlets was enlarged to take an IEC outlet. Once unscrewed from the chassis, the tinned brass tags on the loudspeaker outlet boards were removed by squeezing their inside retaining lugs together using a pair of heavy duty pliers, allowing three-way binding posts to be neatly fitted. A blanking grommet masked the hole for the original mains fuse, and the hole for the switch cable was opened out to take a mains switch (see [Figure 1.28](#)).

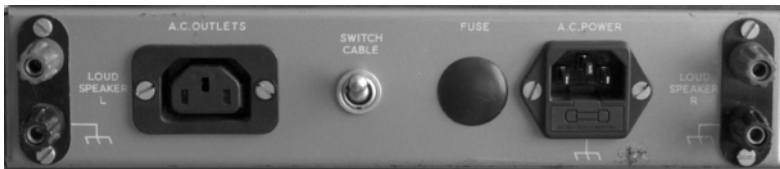


Figure 1.28

A little ingenuity allows the fitting of safe and convenient connectors

Output-transformerless (OTL) amplifiers

The significance of an OTL amplifier is that the absence of an output transformer forces the output valves to pass significant currents, and that these amplifiers often employ substantial global feedback, requiring unwanted voltage drops and stray capacitances to be minimised to ensure stability. The combination of these two factors requires true star earthing encompassing the entire amplifier and power supply, and this dominates mechanical design. As an example, a generic OTL headphone amplifier known at the Headwize forum as the Broskie-Cavalli-Jones (BCJ) will be considered (see [Figure 1.29](#)).

Despite initial appearances, the first two valves are not a differential pair because the first valve has its anode decoupled to

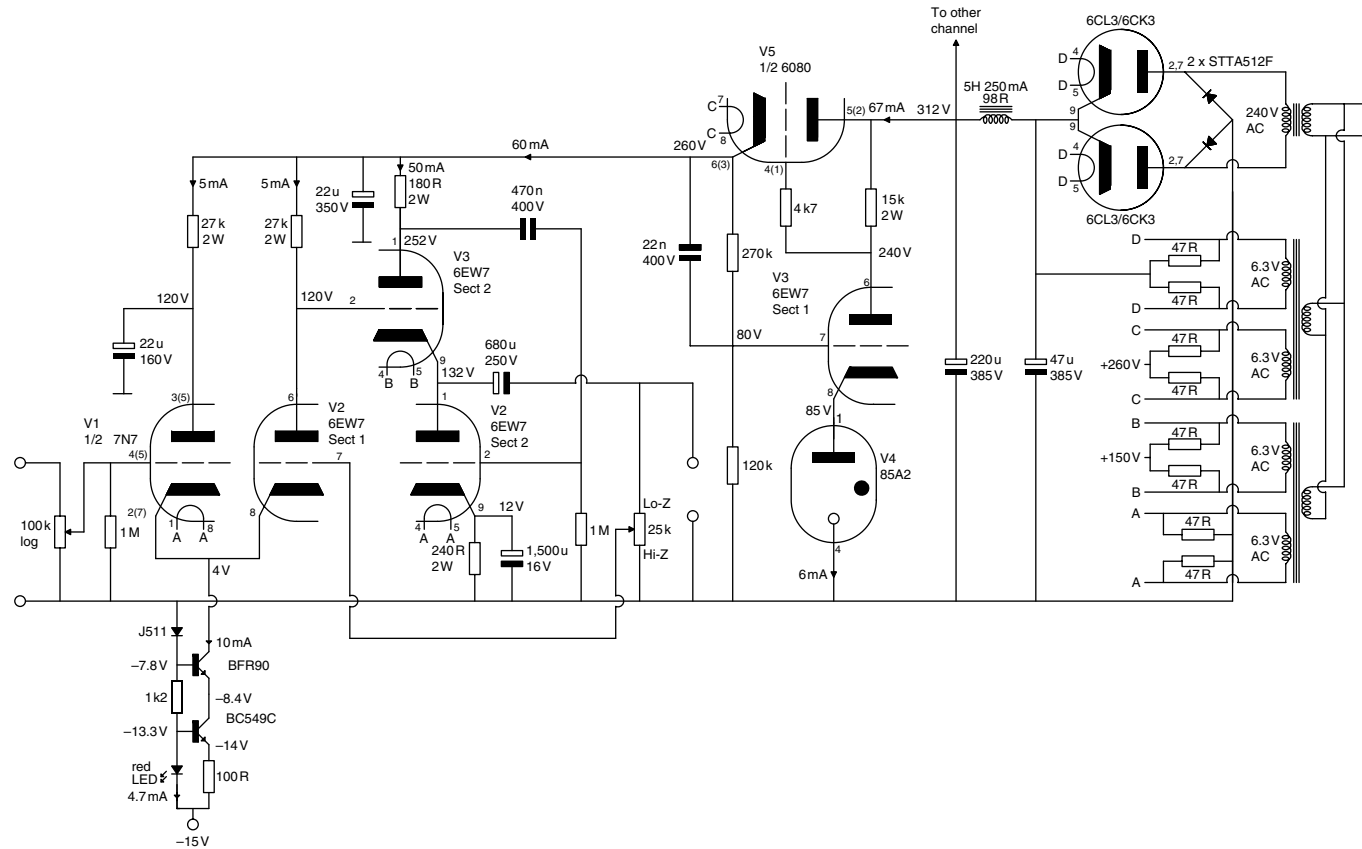


Figure 1.29
Output-transformerless headphone amplifier

ground by a capacitor. The first stage is a cathode follower with a semiconductor constant current sink as its load. The output of the cathode follower is direct coupled to the cathode of the common grid stage, which is direct coupled to an optimised White cathode follower output stage [5]. User-adjustable global feedback to suit the particular headphone in use is taken from the output terminals of the amplifier to the grid of the common grid stage [6].

Because optimised White cathode followers develop their correction voltage across the series combination of the intentional regulating resistor at the anode of the upper valve and the output impedance of the supply, a regulator is almost mandatory, and because the 6080 is a twin triode, a separate regulator can easily be used for each channel.

All 0 V connections from the amplifier and power supply must be brought back to the single star earth and each channel's regulated HT supply must also be a star point. These two requirements enforce a circular layout and all other design considerations are secondary (see [Figure 1.30](#)).

The valves must lose considerable heat, yet they are close together. The only way that this can be achieved is by mounting them on perforated sheet.

Pre-amplifiers

Pre-amplifiers present different challenges from power amplifiers. It is unusual for the power supply to be on-board, and heat is far less of an issue. Conversely, signal levels are far lower . . .

A typical pre-amplifier might have five inputs selected by a front panel rotary switch. The traditional solution was to mount the switch on the front panel, and take each input up to the switch with screened lead (unscreened wire would cause hum and

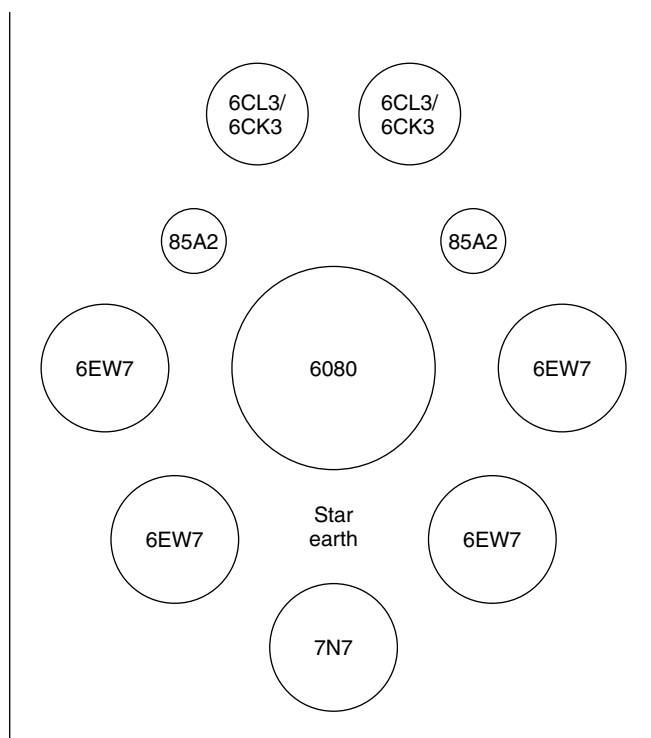


Figure 1.30

Optimum star earthing dominates mechanical planning

crosstalk problems). This was a very poor solution. Screened lead is expensive to buy and to fit, yet with five inputs, this solution only ever uses a fifth of the wire at any given time. The place for the input selector switch is on a bracket at the back of the pre-amplifier, sufficiently close to the input sockets that unscreened wire can go directly from each socket to switch contacts.

The output of the selector switch feeds the volume control, so the obvious position is nearby, but it is often best positioned diagonally opposite (see [Figure 1.31](#)).

The layout assumes a right-handed operator, so you might want to mirror it if you are belligerently left-handed. The selector

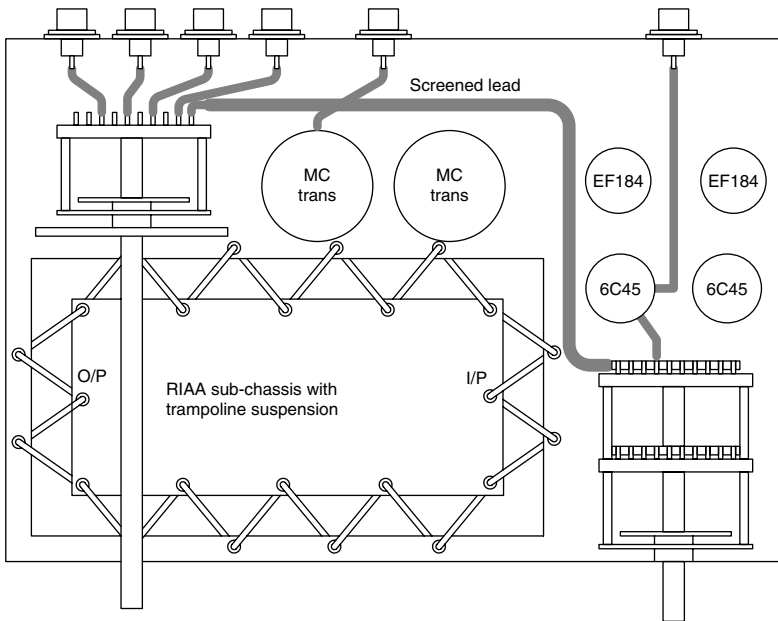


Figure 1.31
Pre-amplifier layout

switch is mounted at the back left, and extended to the front with a coupling and a length of 6 mm or $\frac{1}{4}$ " shaft as appropriate. (13" lengths of stainless or silver steel rod are cheap at engineering suppliers.) If necessary, a coupling could be boded from a short length of motorcycle fuel hose and associated Jubilee clamps, but a proper coupler is better (see [Figure 1.32](#)).

This pre-amplifier has buffering between the volume control and its output, so this circuitry is logically best positioned at the back right, forcing the volume control towards the front of the chassis. This layout also assumes a moving coil RIAA stage using input transformers, and since they generally have flying leads, it makes sense to position them between the input sockets and the first gain stage, because it allows their (flexible) leads to link the (floating) sub-chassis to the connectors on the (fixed) main chassis. The RIAA stage is suspended on knicker elastic, and the necessary gap all around its sub-chassis allows cooling air

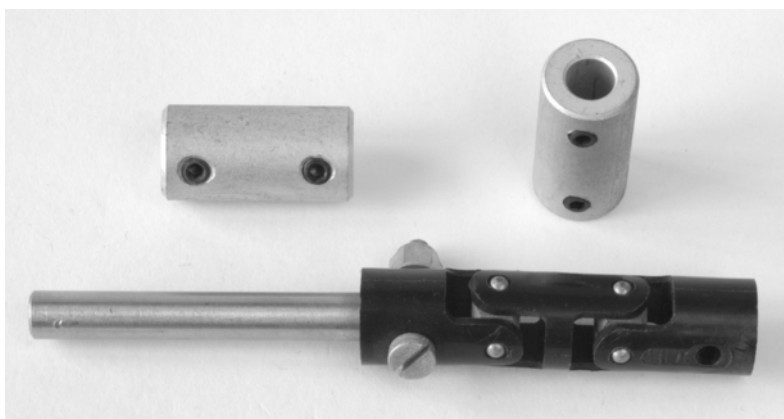


Figure 1.32

Shaft couplers allow potentiometer or switch shafts to be extended cheaply

past its valves. The stage runs from right to left, so its output is close to the selector switch. The output of the selector switch to the volume control is a comparatively long run, so screened lead is necessary. An electrostatic screen could be added between the RIAA stage and the volume control, but this is probably unnecessary.

Umbilical leads

Since a pre-amplifier almost invariably has a remote supply, it needs an umbilical lead to connect the two together. To allow either item to be moved, the umbilical needs to be unpluggable. The obvious solution is to provide a connector at each end of the umbilical lead. Once you have seen the cost of multipole connectors, you will quickly decide that a connector is only needed at one end of the lead. The problem is to decide which end of the lead should have the connector.

It makes better sense to make the lead part of the pre-amplifier rather than the power supply. The reason for this is that a pre-amplifier tends to be designed for a specific application, perhaps to match a particular cartridge in a particular turntable,

and to drive particular power amplifiers. That being the case, its proposed physical positioning is also known, so required lead lengths are also known. The previous argument also applies to its audio leads, so why not hard-wire the outputs, rather than adding expensive and unnecessary connectors?

Personality

A typical power supply provides a HT supply, one, or more LT supplies, and some means of remote power switching. All pre-amplifiers need some, or all, of these facilities, so it would be very useful if a given power supply could be used by any pre-amplifier. A pre-amplifier could then be quickly and easily replaced without having to modify the power supply. One obvious consequence of this approach is that the power supply's multipole connector must be chosen so as to leave pins for future expansion and that extra connectors for future pre-amplifiers must be readily available at a later date. Thus, if you spot a bargain connector, buy as many cable plugs to fit the chassis socket as you think you will need for future pre-amplifiers. (Remember that the power supply's chassis connector must be a socket to make it impossible for your fingers to contact potentially live pins.)

Pre-amplifiers often need elevated heater supplies to keep V_{hk} within acceptable limits in circuits such as cathode followers or cascode where a cathode could easily be at 200 V. The required elevating voltages are determined purely by the pre-amplifier, so any circuit for elevating the heater supplies should be within the pre-amplifier chassis, not the power supply. In this way, the power supply's personality is determined purely by the pre-amplifier plugged into it, and it remains universal. Thus, **all** individual supplies (whether HT or LT) within the power supply should float from the chassis, so that their connections to chassis and to one another are determined purely by the associated pre-amplifier.

Mains/chassis earth has to be carried from the power supply to the pre-amplifier, and it is vital that this connection is reliable

and of low resistance. If at all possible, use more than one pin on the connector to carry this connection – this might mean that you need a larger connector with more pins, but it’s safer. We want a low resistance bond between the two chassis, so that implies a good cross-section of copper conductor in the umbilical lead. The best way of achieving this is to use an umbilical cable having a braided screen, and use the screen as the earth bond connector. If you make up your own cables, use two, or perhaps even three, layers of screens because this will not only ensure that there are no gaps in the composite screen but it constitutes a robust cable having a low earth resistance. By the time your umbilical has all this screening, plus a protective nylon braid over the top, it will be quite thick, so make sure that your chosen connector can accommodate and clamp the cable securely.

Power supplies

Fortunately, power supply planning is much simpler than that required for amplifiers. Nevertheless, the function of a power supply is to take AC mains, rectify it and deliver clean DC to another point. It therefore makes sense to lay a power supply out so that one end is deemed to be “dirty” (incoming AC mains, AC outlets, and rectification) and the far end “clean” (regulators, DC outlets).

Regulator positioning

The ideal position for a regulator is adjacent to the load because this minimises unwanted voltage drops, yet there are powerful arguments against this positioning, particularly if the load is a pre-amplifier:

- Regulators are inevitably hot, and pre-amplifiers should stay cool.
- The raw input to a regulator is noisy.

- Valve HT regulators require AC heater wiring, which is noisy. Obviously, semiconductor HT regulators do not suffer from this problem, and can be more easily placed adjacent to the load.
- Valve heaters draw significant current, causing a voltage drop down the umbilical cable. However, a remote heater regulator can circumvent this problem because its output voltage can be adjusted to compensate whilst measuring the voltage at the valve pins.

Once these considerations have been taken together, the outcome is generally that the regulators are typically within the power supply, even though this does not at first sight seem to be ideal.

Professional power supplies often use **remote sensing**, also known as **four-wire connection**, so that each output terminal has a “sense” terminal associated with it which is connected to the (remote) load, thus automatically compensating for cable resistance. Although this technique requires extra pins on the connector and wires in the umbilical lead, the extra leads do not pass any significant current (see [Figure 1.33](#)).

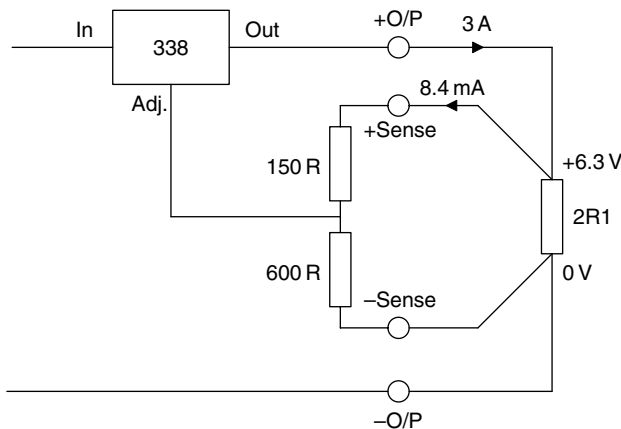


Figure 1.33

Remote sensing overcomes the voltage drop down an umbilical cable at the expense of two extra wires

Fitting mains outlets

Mains outlets are extremely useful, and allow the entire Hi-Fi to be switched on at the touch of a single switch. The most common (and safest) outlet is the IEC 10A 3 pin. The following points make life easier later on:

- Align the sockets vertically. This uses space more efficiently on a slim chassis and allows more sockets.
- Vertical sockets allow straight tinned copper wire buss bars to be threaded from pin to pin – making wiring much easier and tidier.
- Orient the sockets so that the live pin is closest to the chassis top plate rather than the open bottom. The dangerous pin is now buried deep inside the equipment where you are far less likely to touch it accidentally during testing.
- Don't bother attempting to mark out the fixing holes. Once the hole for the body has been cut, an outlet can be inserted and you can mark the precise position of the holes with a scribe.
- 3 mm aluminium is thick enough to be tapped 4BA or M3.5 – but positioning nuts is fiddly. (Remember to lubricate the tap with methylated spirits.)
- Provided that the holes **are** tapped, shuttered IEC sockets will just fit between the flanges of 2" channel if they are fitted vertically (see [Figure 1.34](#)).

Fitting TO-3 regulators and transistors

Power supplies often need TO-3 style components to be fitted to the chassis. You don't need to measure and mark out all the holes precisely. Simply use the mica washer from an insulating kit as a template to mark out the holes. If you are using power transistors, they will be hot, so calculate how many watts of heat must be dissipated, assume a heatsink temperature of 50 °C (30 °C above ambient), and calculate

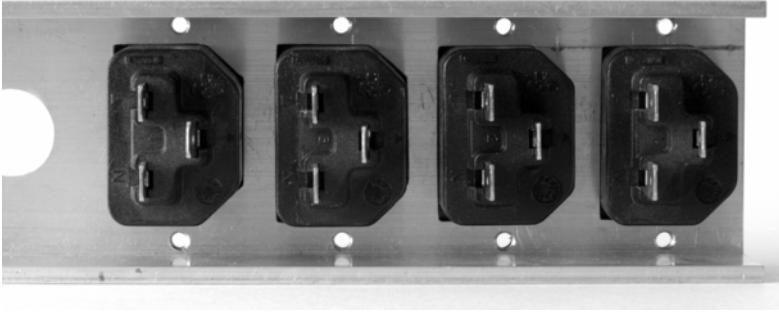


Figure 1.34

IEC mains outlets will just fit vertically between the flanges of 2" channel provided that the aluminium is tapped to take the securing screws

whether you need a dedicated heatsink, or whether the chassis will do:

$$R_{\theta} = \frac{\Delta T}{q}$$

Where:

R_{θ} = required thermal resistance of heatsink ($^{\circ}\text{C}/\text{W}$)

ΔT = temperature difference ($^{\circ}\text{C}$)

q = power to be dissipated (W)

As an example, we might have an LM338 heater regulator passing 3 A, with a drop of 6 V across it, corresponding to 18 W of dissipation. As a very rough approximation, if a reasonably large chassis has sides made of 2" $\times$ 1" 3 mm thick channel, at any given point, the channel has a thermal resistance of $\approx 1^{\circ}\text{C}/\text{W}$. Thus, mounting the regulator on the channel would keep the temperature rise to $\approx 18^{\circ}\text{C}$, resulting in a local chassis temperature of $\approx 38^{\circ}\text{C}$, which is acceptable. However, if two such regulators were nearby, the chassis temperature could easily exceed 50°C , so it would be worth either fitting an additional finned heatsink or separating the heat sources by a few inches (50–70 mm).

References

1. “Microphone Engineering Handbook” Michael Gayford (Editor). Focal (1994). ISBN 0-7506-1199-5.
2. “Components Group Mobile Exhibition” Brimar Valves (November 1959).
3. “Type KT66” Osram Valve Technical Publication TP3. GEC (February 1949).
4. “Mullard Technical Communications” (November 1962). Available at: <http://www.thevalvepage.com/valvetek/microph/microph.htm>.
5. “The White cathode follower” John Broskie. “*TubeCAD journal*” (October 1999). Available at: <http://www.tubecad.com/october99/page4.html>.
6. A post by Alex Cavalli on page 5 of the “Building the Morgan Jones” thread. Available at: <http://headwize2.powerpill.org/>.

Further reading

- “Noise Reduction Techniques in Electronic Systems” 2nd edn. Henry W Ott. Wiley (1988).
- “The technique of radio design” E E Zepler. Chapman & Hall (1943).
- “Foundations of Wireless (and electronics)” 7th edn. M G Scroggie. Liffé.
- “Radio communication handbook” 5th edn. Radio Society of Great Britain (1976).
- “Electronic Classics” Andrew Emmerson. Newnes (1998). ISBN 0 7506 3788 9.

CHAPTER 2

METALWORK FOR POETS

Many electronics enthusiasts hate metalwork. If they were to take a longer and more thoughtful look at their pet hate, they would realise that what they actually hate is attempting to do metalwork with **poor tools**. A skilled worker can produce good work even with poor tools, but would far rather use the best. Beginners do not have this level of expertise, and need all the help that they can get, so they **need** good tools. Buy the best that you can't quite afford. Good tools last a lifetime, and not only are they cheaper in the long run, but they are a pleasure to use.

Marking out

This is where the mistakes are made, so don't rush, ten minutes saved here could cost hours later. "Measure twice, cut once."

Measurements

Before you can mark out, you need detailed measurements of the parts you need to fit. Commercial manufacturers are supplied with detailed drawings including tolerances of parts by their

sub-contractors, but you might use salvaged parts, which **never** come with drawings. You need to confidently make accurate measurements. Digital callipers are far cheaper than they used to be, and make life so much easier. Even better, they have a button that changes them from inches to millimetres, or vice versa (see [Figure 2.1](#)).

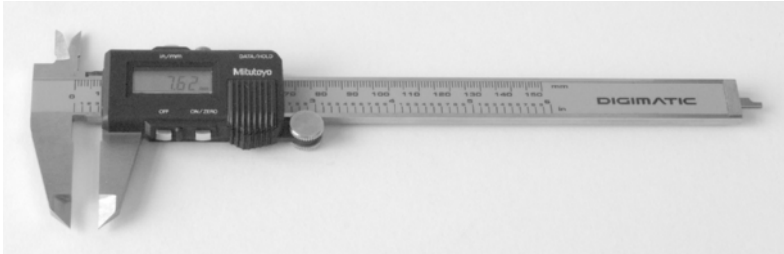


Figure 2.1
Digital callipers are now cheap...

Locations on metalwork are traditionally found using a clean grey steel rule in conjunction with an engineer's try square (see [Figure 2.2](#)).

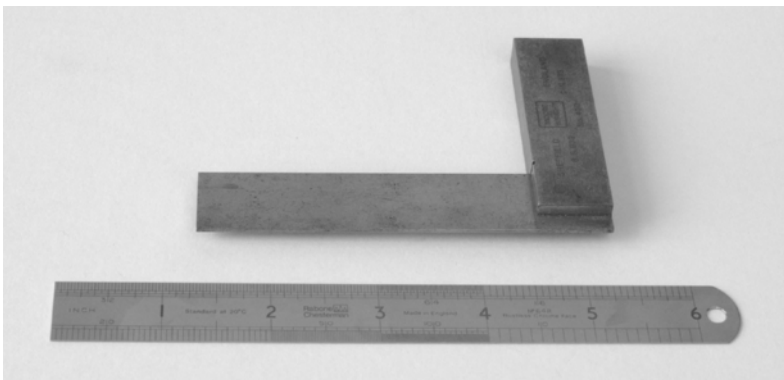


Figure 2.2
Engineer's rule and try square

Although you will need both a 150 mm and a 300 mm rule, it is better to use the shorter rule whenever possible because it is

thinner, and reduces the parallax errors that occur as a result of looking at graduations some distance away from the work. When buying rules, carefully check the engraving quality of the graduations. Some cheaper rules etch the graduations, making them less clearly defined, and difficult to read. Likewise, a rule with a grey finish is easier to read than a shiny (or even worse, rusty) rule.

A small (3") try square is easier to use on sub-assemblies, but a large (6") square is needed for full-size chassis.

Marking out manually

With the best will in the world, your marking out will never be perfect, so choose a reference edge from which to measure, and only use a try square from this edge to minimise errors.

The centre of drilled holes is marked by the intersection of two lines, and these lines are made using a sharp scribe (see [Figure 2.3](#)).



Figure 2.3

Traditional scribe (lower) and modern (upper). Both are equally good

With all the construction lines that you will need, there will be a lot of these intersections, so when you are marking the position of a hole, use a scribe to draw a circle around the intersection of roughly the same diameter as the hole. This will prevent you from drilling holes in the wrong place, and may stop you drilling the hole oversize. If you have a pair of dividers, use

them to mark the size of larger holes accurately – any mistakes in marking out will become apparent instantly.

You will often have to cut irregularly shaped holes for transformer connections. Cross-hatch the metal to be removed with a marker pen to avoid confusing construction lines with cutting lines. The reason for using a marker pen rather than a scribe is that if you make a mistake, it can be removed with methylated spirits, whereas scribed marks are permanent.

The CAD solution

Alternatively, you can rely on the precision of a printer to do your marking out for you, although it is a good idea to check the accuracy of your printer in both horizontal and vertical axes against a steel rule. Modern printers can print with a resolution of at least 300 dots per inch (dpi), and often far more, and are much better than you are at positioning lines, so an engineering drawing package can easily produce a precise template that just needs to be spotted through with a scribe to ensure perfect marking out.

Admittedly, learning to use a CAD drawing package takes time, but it will be justified by the improved quality of your metalwork.

Checking and punching

At this point it is still possible to correct mistakes, so offer up the various parts to be fitted, and check that the marking out looks sensible.

Given half a chance, a drill will skid randomly around the surface of metal before cutting – probably in the wrong place. Punching a dimple to start the drill helps it to cut in the correct place. Having checked your marking out, use a centre-punch to

indent the centres of all the drilled holes. The modern punch is the automatic spring-loaded punch, whereas the older hand punch has to be struck with a hammer. Although the autopunch is superficially attractive, you will find that it is less accurate than the hand punch (see [Figure 2.4](#)).

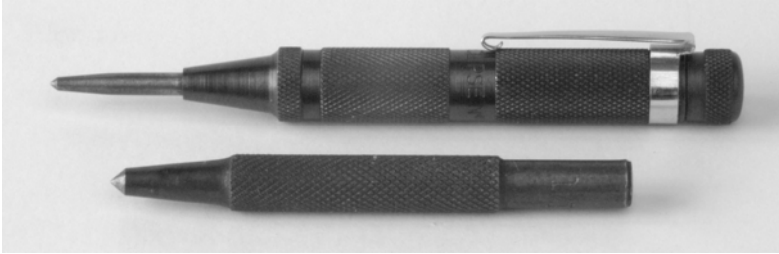


Figure 2.4

Traditional centre punch (lower) and autopunch (upper)

Whichever punch you use, it must be **sharp**, and should be ground to an included angle of 90° . A punch is easily sharpened on a bench grinder, or a grinding wheel in a drill. (Always wear goggles when using a grinder.)

Centre-punching sheet metal is noisy. Noise is minimised, and accuracy increased, by supporting the area to be punched directly above a leg of the bench or table. This work is being done for the pleasure of your ears, so buy some ear-defenders (cheaper than you would think), and wear them.

Safety and drilling

Drilling triangular holes in sheet metal is remarkably easy, but drilling round holes in the correct position takes a little more care.

A drill mounted in a stand, bolted to a bench, invariably means that the operator bends or sits to carry out the work. Drilling produces swarf which flies out from the drill; this swarf may be

hot (although if it is hot whilst you are drilling aluminium, then something is wrong). You are now at eye level to the swarf, so wear safety glasses. Whether you do your metalwork in a dedicated workshop or on the kitchen table, you will get hot. If perspiration runs into your eyes, it stings and may cause you to blink – perhaps with disastrous results. Steal the sweatband that your wife or daughter wore when she was feeling self-conscious about her figure. Your ears are near your eyes, and drills are noisy, so also wear ear protection.

When you drill sheet metal, the work will try to vibrate, and if it is allowed to do so, the drill will snatch at the work, and you will suddenly find the work spinning on the end of the drill. This is most alarming and can be very dangerous.

Drilling round holes in the correct position without snatching boils down to:

- Preventing relative movement between the drill and work
- Lubrication
- Correct drill speed.

Preventing relative movement

Support the work. Don't attempt to drill into the middle of an unsupported chassis. Always support the work on a scrap of wood, often known as a **drilling block**, otherwise, the force required to make the drill cut will make it suddenly bite deeply and snatch. It is important that the work is level when drilling, so use off-cuts of MDF as drilling blocks, rather than natural wood that may have warped or may not have been planed accurately. Ensure that the work remains firmly in contact with the drilling block by deburring the underside of each hole immediately after it is drilled.

Ideally, clamp the work, but if you must hold it down by hand, press firmly, and avoid having your fingers near a part of the

work that would cut you if the work were to spin. If the worst comes to the worst, do not try to fight the drill; it is much stronger than you are. Let go, and switch the drill off. Ideally, you would have a foot operated stop switch, but few amateur workshops can afford this level of sophistication.

Use a drill stand bolted to your bench. Again, this is much easier than it used to be, because drill manufacturers have standardised on a 43 mm collar to grip the drill in the stand, so a choice of stands is available. Alternatively, small drill presses are now available at remarkably low prices, and although they do not withstand comparison with a genuine workshop drill press, they are perfectly good for amateur work. If you have never used a drill in a stand before, you will not believe how much easier it makes your work!

The most important specification when buying a drill press, or stand, is the depth of the **throat**. The throat is the distance between the drill axis and the outside of the pillar that supports the drill. A shallow throat renders the drill unable to reach the centre of your chassis, forcing you to use a handheld drill. In this instance, a deep throat is very desirable.

Lubrication

Lubrication greatly aids cutting efficiency. If you use a lubricated 3 mm centre drill to start each hole, you will find that you can immediately use the final size of lubricated drill, without needing pilots. This not only speeds your work, but actually produces more precisely aligned holes (see [Figure 2.5](#)).

For aluminium, the author keeps a small glass jar containing sufficient depth of methylated spirits to keep the bristles of a $\frac{1}{2}$ " paintbrush wet. Before drilling, wet the tip and flutes of the drill, and spot a droplet at each hole to be drilled. In time, you will knock the jar over and discover why you only wanted a shallow



Figure 2.5

Centre drills enable holes to start in the correct place

depth of spirits. When you have finished work for the day, screw the lid on the jar, otherwise the spirits will evaporate – leaving you with an empty jar and a lingering fragrance. (Don't use a jar with a plastic lid as the vapour will eventually melt it.)

The traditional **cutting fluid** for steel was a 50/50 mix of lubricating oil and water with a healthy dash of washing-up liquid. The water cools, the oil aids cutting, and the washing-up liquid allows the two to mix to form an emulsion. Unfortunately, it is extremely messy and rusts tools. Fortunately, modern proprietary cutting fluids such as Templer's "Temaxol" are less destructive.

Only a little lubrication is required – too much will spray you and your surroundings.

Correct drill speed

Use the correct drill speed. A 1 mm drill needs to run fast to clear the swarf or it will break, so it should run at 2500 rpm, whereas a 12 mm drill should be run at the slowest possible speed, 200 rpm or lower, if possible. This is not nearly as much

of a problem as it used to be for the amateur, because the better quality power drills have a two-speed mechanical gearbox and electronic speed control.

Use sharp, correctly ground drills – **never** attempt to resharpen drills (if you genuinely know how to do this correctly, you don't need to read this section). The cheapest way to buy decent drills is to look for mail order suppliers in a model engineering magazine or on the Internet – most of them will gladly send you a free catalogue. All of them stock sets of drills by good manufacturers, but individual drills are expensive, so buy a set of 1–6 mm in 0.1 mm steps – another from 1–12 mm in 0.5 mm steps is useful. These are virtually all the drills you will ever need, and they will come in a protective steel box with each size marked (see [Figure 2.6](#)).

Whatever you do, don't buy drills at your local DIY supermarket. There is a world of difference between the inaccurately ground rubbish made of muckite that is adequate for knocking up the occasional shelf and a true engineering drill.

Deburring

When you drill a hole in metal, there will always be a small burr on the upper surface, and a larger burr on the lower surface. The best way of deburring is not to produce a burr in the first place. Provided that the work is firmly pressed down and supported directly below a well-lubricated drill, very little burr will be produced.

All metalworkers have their preferences for burr removal, but a rose countersink in a handle, or a specialised “wiggly” deburring tool both work well (see [Figure 2.7](#)).

You could simply use a large drill without a handle to deburr holes, but you will find that the flutes of the drill tend to cut the surface of your skin as you grip it, and the drill tries to bite

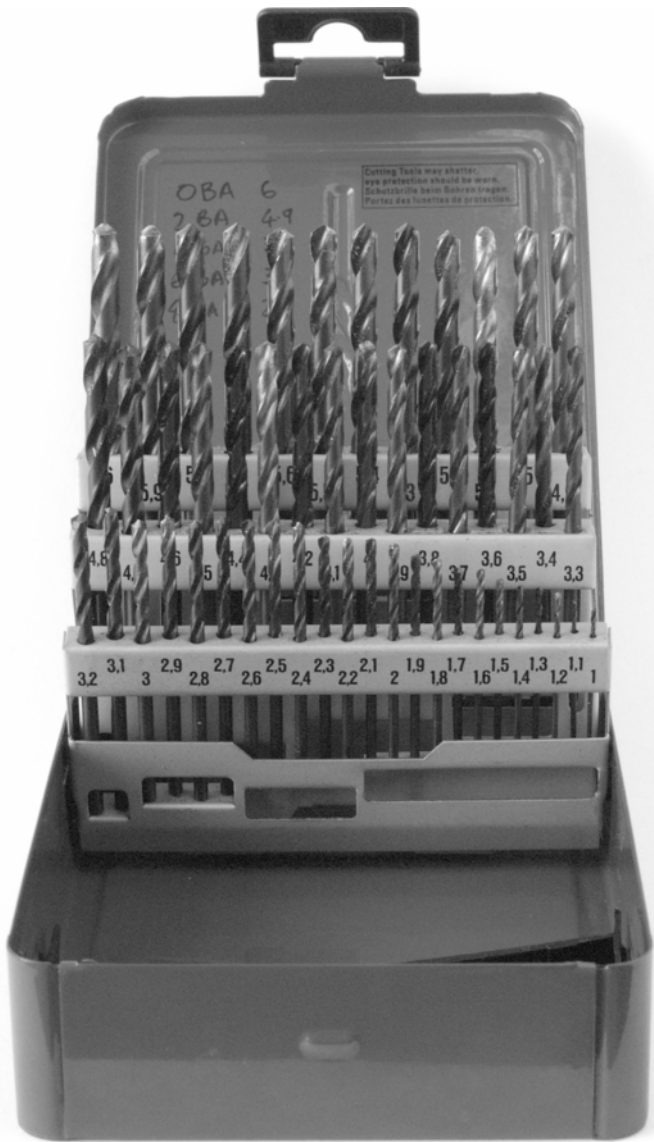
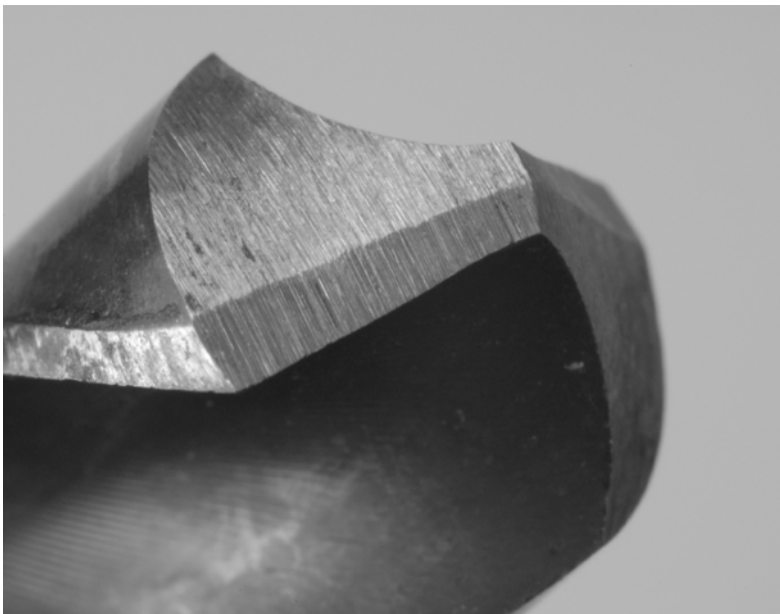


Figure 2.6
Set of engineer's twist drills in case

**Figure 2.7**

“Wiggly” deburring tool (lower) and rose deburring tool (upper)

deep into the work. If you are lucky enough to find a toolmaker who sharpens their own drills, politely ask them to regrind a $\frac{1}{2}$ " and a $\frac{1}{4}$ " drill for deburring. The leading edge is ground to a negative rake of $\approx 1^\circ$, so that when gently rotated by hand against a drilled hole, it doesn't quite bite. Correctly ground, these tools deburr superbly (see [Figure 2.8](#)).

**Figure 2.8**

$\frac{1}{2}$ " drill ground to form deburring tool

Countersinking, and why to avoid it

Stainless steel countersunk screws do look nice, and are sometimes essential to avoid the screw head fouling other parts, but countersinking produces problems.

The included angle of screw heads is 90° , so we can forget any ideas of using a twist drill (typically 118° or 135°) or centre drill (60°). There are various different types of countersinks (see [Figure 2.9](#)).



Figure 2.9
Selection of countersinks

It is essential to lubricate countersinks before cutting. Unlubricated rose cutters clog, and when forced into the work, suddenly bite and cut an octagonal countersink. Countersinks with three cutting edges are less likely to clog but are not quite as controllable. Conical countersinks are the most controllable and produce the best finish, but they are quite expensive.

Although countersinks are nominally self-centring, they easily move off-course, causing the screw head not to seat properly, so be careful to align the countersink precisely before cutting.

Nobody produces perfect metalwork, and cap head screws are remarkably forgiving of slight inaccuracies. Countersunk holes are not in the least forgiving.

Tapping holes

Sometimes, despite careful planning, we know that a fastener will be inaccessible from the far side, so a screw and a nut cannot be used. Although self-tapping screws can be used to secure parts, the sharp protruding screw can easily nick wires on the other side, and the weak thread it cuts in the chassis doesn't really withstand repeated removal and reinsertion. Be aware that a tapped hole does not allow for any subsequent manoeuvring of the screw, so you need to be confident about the accuracy of your drilling.

A far better alternative to a self-tapping screw is to cut a proper engineering thread into the chassis, using a **tap**. Tapping and clearance sizes are given (in mm) for British Association (BA) and ISO metric screws. The table shows that standard sizes are comparable, except for 4BA which does not really have a metric equivalent. The significance of the comparison is that many NOS parts were intended for BA fasteners.

<i>ISO metric</i>	<i>Clearance</i>	<i>Tapping</i>	<i>British Association (BA)</i>	<i>Clearance</i>	<i>Tapping</i>
M6	6.2	5.1	0BA	6.2	5.1
M5	5.2	4.3	2BA	4.9	4.0
M4	4.2	3.4			
M3.5	3.7	2.9	4BA	3.9	3.0
M3	3.2	2.5	6BA	2.9	2.3

A rule of thumb for the threads in the table is that the thickness of the material to be tapped should not be less than half the thread diameter. Thus, up to M6 can be tapped into 3 mm

channel, but the largest thread that should be tapped into 1.6 mm sheet is M3.

Taps are typically sold in sets of three for a particular thread size: starting tap, second tap, and bottoming tap. For aluminium, you only really need the starting tap and bottoming tap (see [Figure 2.10](#)).

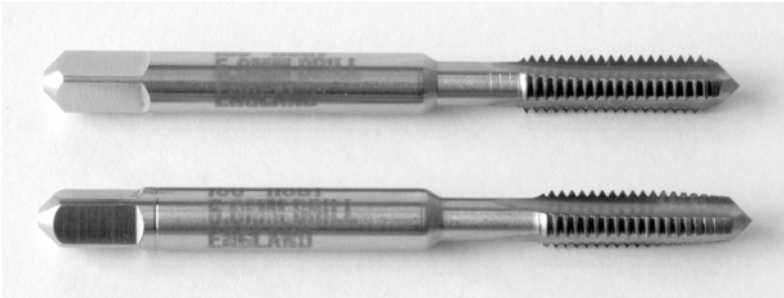


Figure 2.10
Starting (lower) and bottoming (upper) taps

Taps are held in a **tap wrench**. The chuck tap wrench is better for small taps and is easier to keep aligned, whereas the larger bar wrench allows more force to be applied, and is better for larger taps (see [Figure 2.11](#)).

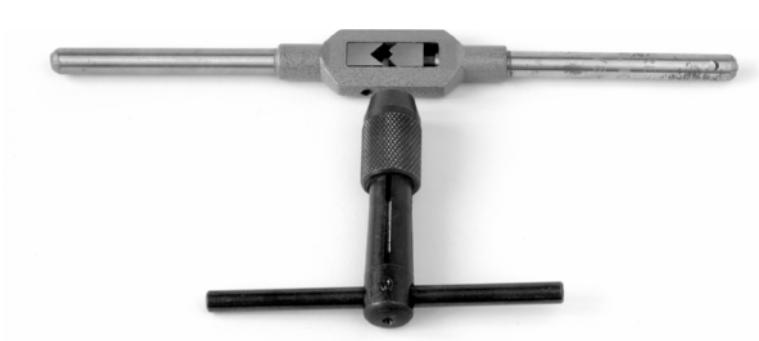


Figure 2.11
Chuck and bar tap wrenches

Although taps are slightly self-aligning, it is important that the tap is first aligned with the axis of the hole – check from side to side and from front to back. The way that you use a tap to cut the thread depends on whether the hole to be tapped is clear (you can see through it) or blind. A clear hole allows swarf to clear easily, but a blind hole quickly clogs the tap.

To tap a clear hole, align the tap so that it is vertical, lubricate it, and press down gently as you screw it into the hole. Never force a tap. If everything has been done correctly, the tap will glide through the work and emerge on the other side covered in swarf. Use a brush to remove all the swarf, **then** unscrew the tap from the work. If the metal is only a few millimetres thick, you will only need a starting tap to cut a perfect thread (see [Figure 2.12](#)).

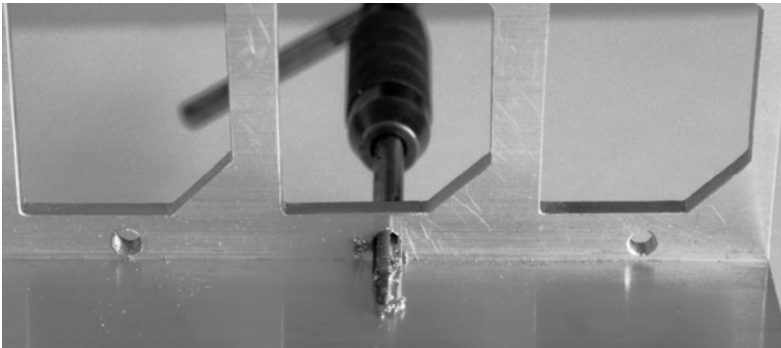


Figure 2.12

The tap emerges from tapping 3 mm aluminium covered in swarf. Remove the swarf with a brush, **then** unscrew the tap

If the hole is blind, the swarf can't clear easily when tapping, so use plenty of lubricant, and tap by turning the tap $\frac{1}{2}$ turn forwards followed by $\frac{1}{4}$ turn backwards to clear the swarf. It may be necessary to withdraw the tap and clear the swarf from tap and hole several times before the tap meets the bottom of the hole. It is most important that swarf is cleared regularly, otherwise you may not notice when the tap touches the bottom of the hole, and may break the tap. Removing a broken tap from a blind hole is almost

impossible. Blind holes invariably require both a starting tap and a bottoming tap to be used, so they take time. Nevertheless, don't be tempted to hurry – it is so easy to wreck the work.

Tapping invariably leaves a very small burr that is best removed by a rose deburring tool. Be very gentle when deburring, the deburring tool can easily dig in and remove a lot of thread.

Before inserting a screw, clear any remaining swarf by squirting the hole thoroughly with a cleaning solvent such as isopropyl alcohol.

Sheet metal punches

You should not attempt to drill holes larger than 9 mm in sheet metal, it is simply asking for trouble.

The solution for round holes is to use a sheet metal punch. Although these are not often available from high street shops, electronics factors and engineering suppliers stock a reasonable range. If the distributor doesn't have the size you need, contact the manufacturer directly (see [Figure 2.13](#)).



Figure 2.13
Selection of sheet metal punches

The punches are of a two-part construction that are drawn together by an Allen bolt. Provided that they are kept well greased on all surfaces, they cut a beautifully neat hole, and last for years. Useful sizes are:

- $\frac{3}{8}$ " : Imperial potentiometers and rotary switches, grommets for small cables.
- $\frac{1}{2}$ " : Imperial toggle switches (but measure first – modern switches are often even smaller than they quote!), 32A loudspeaker terminals, larger grommets.
- 16 mm: B7G sockets, DIN sockets.
- $\frac{3}{4}$ " : NOS B9A valve sockets, some cable clamps. Modern ceramic sockets are 21.8 mm, so use $\frac{7}{8}$ " (22.2 mm).
- $1\frac{1}{8}$ " : NOS International Octal and Loctal sockets. Modern ceramic sockets are 26.1 mm, but the nearest available punch is $1\frac{1}{16}$ " (27 mm).

Larger punches require quite a large hole for the bolt, so there is no reason why you should not use the $\frac{3}{8}$ " punch to cut the bolt hole for a 35 mm punch. Large punches have to sheer a considerable circumference of metal and need a correspondingly large force. Surprisingly, most of the force you expend actually goes into overcoming friction at the shoulder of the Allen bolt against the top of the punch. Distributors do not usually stock the punch manufacturer's entire range of punches, and they very rarely stock the thrust bearings that can be retrofitted to large punches – contact the manufacturer. Thrust bearings take all the strain out of punching holes >30 mm, and are thoroughly recommended, not only for laziness, but because by putting less force into the work, you bend it less.

If, as was suggested earlier, you have used a pair of dividers to draw the exact position of the finished circle, you can align the punch precisely before tightening up the Allen bolt and cutting the hole.

Note that although chassis punches produce very little burr, they do slightly deform the surface from which the cut began. It is

therefore usual to punch from the inside of the chassis to the face side to avoid this deformation being visible. Additionally, the pressure of the opposing face on the work can mark decorative surfaces (such as brushed anodised aluminium), but this can be avoided by placing a thin cardboard washer between the opposing punch face and the decorative surface before punching.

Sometimes, you need to enlarge an existing hole, perhaps to replace the original phenolic B9A valve sockets ($\frac{3}{4}$ ") on a Leak Stereo 20 with new ceramic sockets (22 mm). If you were very patient, you could carefully open out all seven holes with a file. It would be much nicer to use a punch, but you have the problem of centring the punch. The way that this can be overcome is to fit a thick $\frac{3}{4}$ " washer to the punch's Allen bolt so that it locates in the $\frac{3}{4}$ " hole of the chassis. If the chassis is horizontal, the washer will fall onto the cutting edge of the punch and locate itself in the existing hole. You then start cutting, so that the tips bend into the chassis, release the punch, and remove the washer. You now have correctly aligned indentations to locate the cutting edge and you can complete the hole. Easy.

With care, chassis punches can be used on 3 mm aluminium, although the bolt needs to be well greased, and a thrust bearing makes life easier. The metal needs to be weakened by drilling a series of small ≈ 2 mm holes just inside the circumference of the hole to be punched (see [Figure 2.14](#)).

Sadly, sheet metal punches eventually wear out, and when they do, they require more effort and produce more burr. The author has just replaced his $\frac{3}{8}$ " punch. It only lasted 25 years...

Making small holes in thin sheet

If a C-core transformer were to be dropped through a chassis, the laminations would contact the (almost certainly conductive)

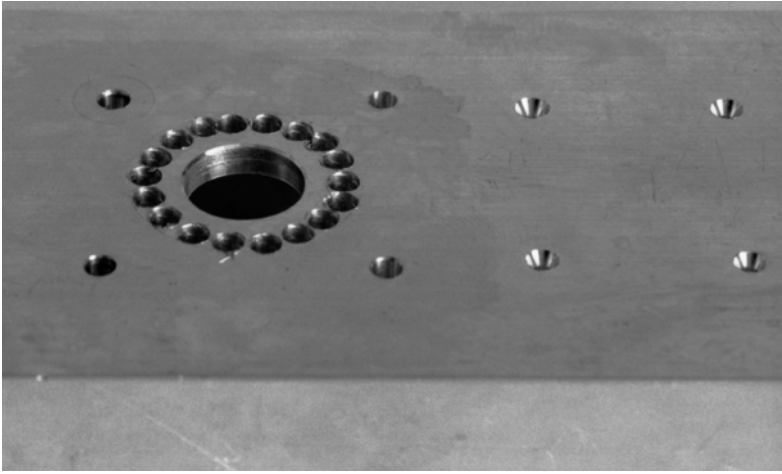


Figure 2.14

Sheet metal punches can be used on thicker metal if it is weakened beforehand by drilling a ring of 2 mm holes

chassis, greatly increasing transformer losses. The solution is to fit an insulating spacer. Although the insulating material can be cut neatly to size with a scalpel or sharp scissors, making clean small holes for the screws to pass through is a problem. This quandary can often be solved by a paper punch. Mark the position of the hole, take the bottom off the punch, tip out the chads, and use the punch upside down. Gripped gently, the cutter holds the material in place, and the marking out can be clearly seen through the exit hole of the punch. The material can be moved precisely into position, whereupon a beautifully clean hole in the correct position can be punched.

Sawing metal

It might be thought that to cut a piece of metal, it is only necessary to take a few wild swings at the work with a hacksaw whilst the room rings to the screech and shudder of the saw.

This is an excellent way to ruin a perfectly good hacksaw blade, deafen yourself, and produce work of an appallingly low standard. Before using a hacksaw, check:

- Is the blade inserted the correct way round? (It should cut on the forward stroke) (see [Figure 2.15](#)).
- Does it have a complete set of teeth? If **any** are missing, discard the blade, blades are cheap – your time is not.
- Does it have the right number of teeth per inch (TPI)? A saw should have three teeth in contact with the work at all times. Alternatively, use 18TPI for aluminium, 24TPI for mild steel, 32TPI for harder materials such as stainless steel, and adjust cutting angle if necessary.
- Is the blade properly tensioned? (The wingnut should be as tight as possible.)

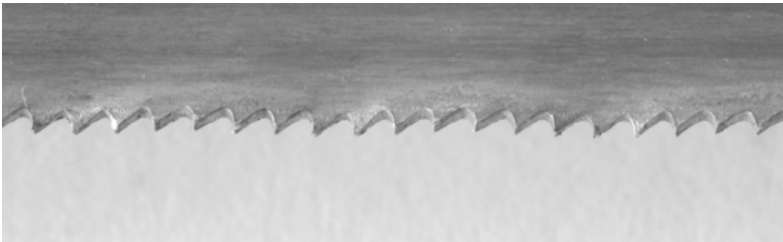


Figure 2.15

Detail of standard hacksaw blade showing cutting direction (Blade cuts when moving from right to left.)

Junior hacksaws are very useful for smaller, more precise work, and the wire frame type is very common. However, **proper** junior hacksaws have a rigid frame and a screw to tension the blade. Hacksaws are commonly available with pistol grip handles because this makes it easier to apply maximum force, which is important for a full-size hacksaw, but you will find that the more traditional handle allows greater precision (see [Figure 2.16](#)).

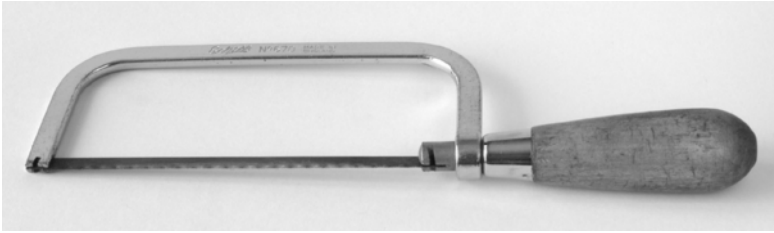


Figure 2.16

A junior hacksaw is ideal for more precise work

Sheet metal

Cutting sheet metal is a problem because a hacksaw blade is insufficiently fine to cut at right angles to the work, so the only way to cut metal sheet is to cut at an extreme angle, and if this means crouching on the floor whilst the work is held vertically in a vice, so be it. A better method is to clamp the work horizontally to the bench, and use a **panel** saw, which looks like a wood saw, but takes a hacksaw blade (see [Figure 2.17](#)).

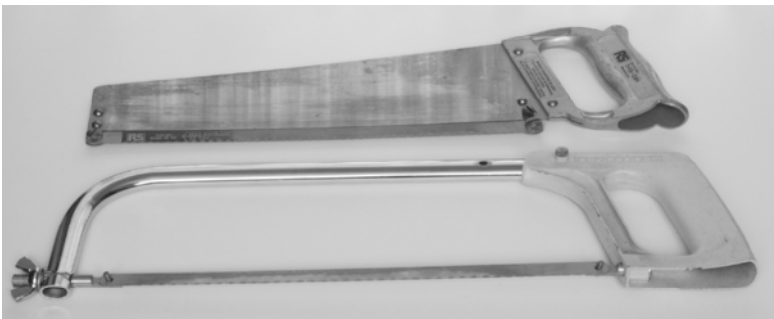


Figure 2.17

Standard hacksaw (lower) and panel saw (upper)

A drop or two of lubricant whilst sawing does wonders for your cutting efficiency but generates a terrible mess and makes it difficult to see your cutting line. The author only saws with lubricant under duress.

Cutting irregular holes

The best way to cut irregular holes is by hand, using a tension file in a standard hacksaw frame or in a coping saw frame (see [Figure 2.18](#)).



Figure 2.18

Tension file fitted to standard hacksaw frame (lower) and coping saw frame (upper)

The process starts by drilling a hole in the material that is to be removed.

The file cuts on the forward stroke, so run the file gently through your fingers to determine the cutting direction (see [Figure 2.19](#)).

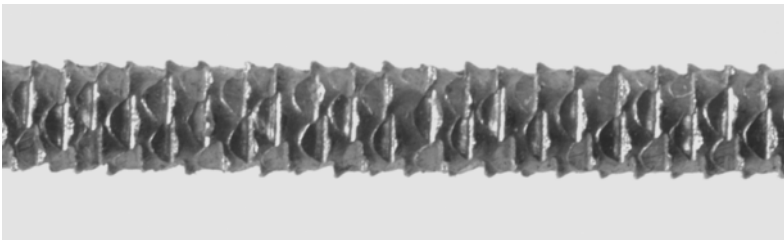


Figure 2.19

Detail of tension file blade showing cutting direction (right to left)

Now fit the file to the handle end of the frame, but pass the file blade through the hole, before fitting it to the other end

of the frame and tensioning it. With care, the hole can be cut so accurately that very little remedial filing is needed (see [Figure 2.20](#)).

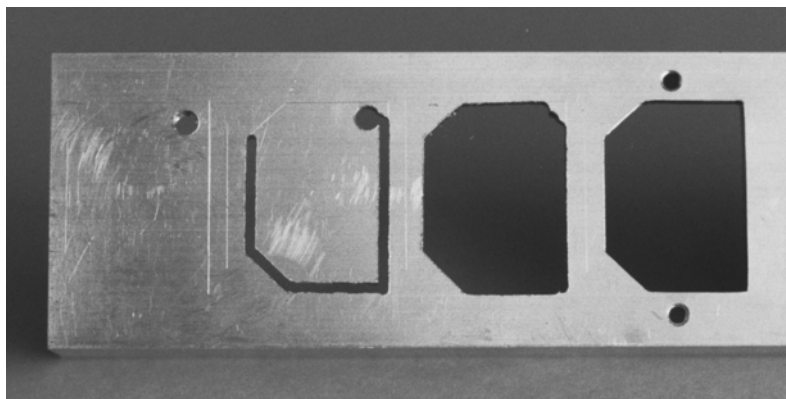


Figure 2.20

The sequence of events for cutting an irregular hole. Left to right: mark out and drill a pilot hole, cut irregular hole with tension file, tidy hole with file, drill and tap fixing holes

Irritatingly, some connectors require a “D” shaped hole, and there’s nothing for it but to cut the hole by hand, so it’s just as well that a tension file does such a good job (see [Figure 2.21](#)).

Sometimes the frame of the hacksaw will be unable to reach the proposed hole from any direction . . .

Hole-cutting files are also available fitted with a handle. To stop them bending, they have to be of a larger diameter than tension files, so you have to remove more material to cut a given distance, making the work harder. Nevertheless, because a frame is not needed, they can reach any part of the work and produce a neat result (see [Figure 2.22](#)).

A faster alternative is to use an electric jigsaw at low speed with a **metal** cutting blade having the finest possible teeth. This is not

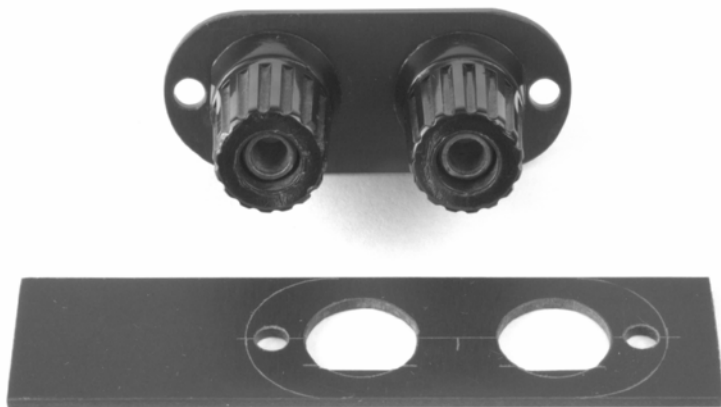


Figure 2.21

“D” shaped holes are a nuisance, but are quite easily cut using a tension file or jeweller’s saw

nearly so easy to control and is potentially dangerous. It is also extremely noisy and ear protection is essential. The drumming of the saw on chippings leaves marks on the work unless a piece of thin cardboard is carefully fitted to cover the sole of the saw (see [Figure 2.23](#)).

As before, a hole is drilled in the material to be removed and the saw blade passed through. Ensure that the teeth of the blade are **not** touching the edge of the hole, and whilst pressing down firmly, start the saw. When you reach the end of a cut, back the saw off a little before switching off, otherwise the teeth will snatch and the saw will try to jump up from the work, deforming it.

A technique that can be useful when a rectangular hole with radiused corners is needed is to use a chassis punch at each corner, then use a saw to cut straight lines between the punched holes, leaving a rectangle with neatly radiused corners.

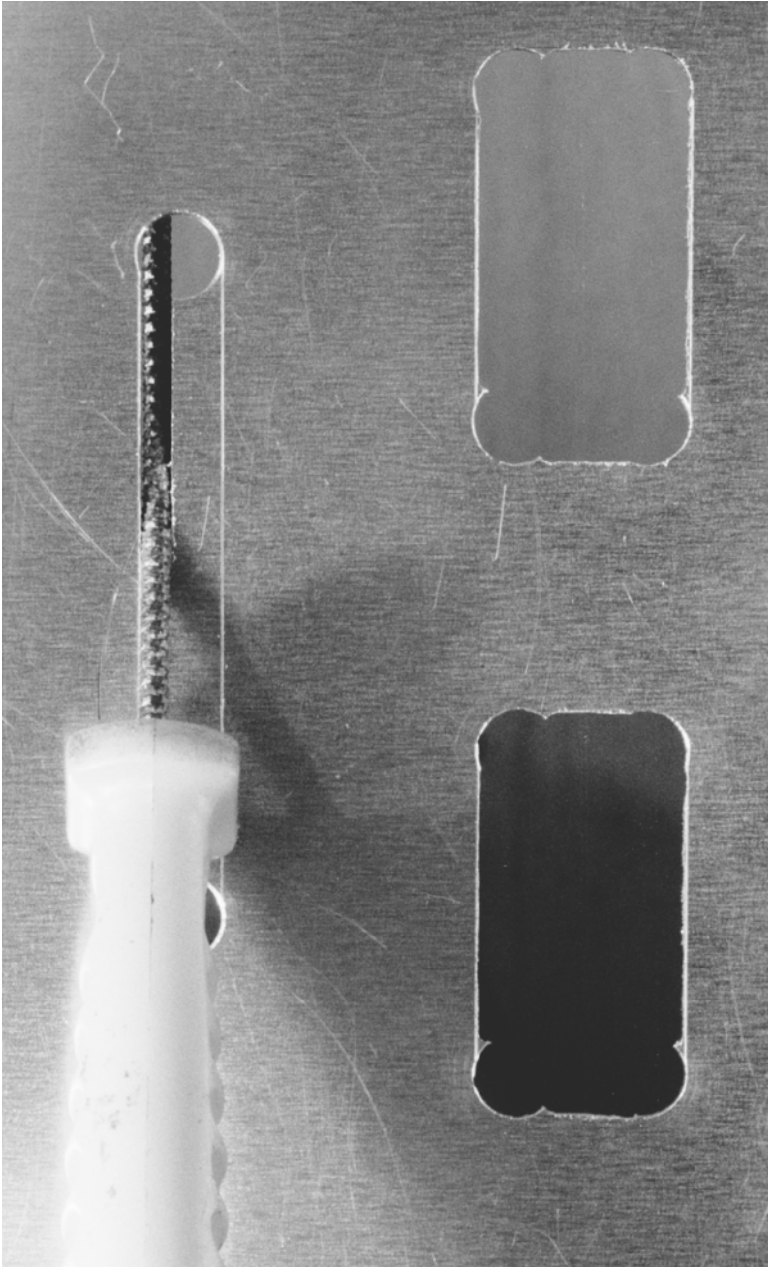


Figure 2.22

A single-ended Abraflex can cut holes any distance from the edge of the chassis

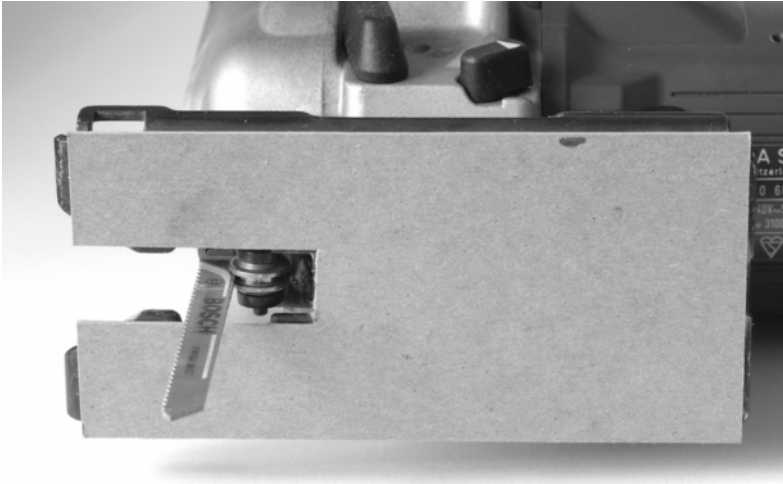


Figure 2.23

Fitting a cardboard sole to a jigsaw reduces surface damage when cutting metal

Making round holes without a punch

Mark out the hole using a scribe, cut it out as carefully as possible, and remove obvious errors with a half-round file. Take a piece of 160 grade emery cloth and wrap it tightly round a short length of broom handle, and use it to scour a few turns inside the hole as if you were turning a handle. Rotate the work by 90°, and repeat, then twice more, so that the work has been scoured in all four orientations. (The reason for rotating the work is that if all the cutting is done in one orientation, a slightly irregular hole results because you are able to apply slightly more force in one direction than another.) If you now inspect the hole, you will find a significant burr on both sides that needs to be removed carefully using a needle file. Change to 240 grade silicon carbide paper, and repeat the scouring process, and deburr for the final time. With only a little effort, it is perfectly possible to produce a hole by hand that appears to have been made by a precision-boring machine.

Although the author uses this technique mainly for very large holes, it can be used for smaller holes by substituting a round file. Round files will cut using the scouring action, but they clog quickly, so it is best to move them gently back and forth as you scour.

Making holes in perforated sheet

Perforated sheet is wonderful for cooling, and makes subsequent wiring easy, but it is very difficult to make small holes that are not aligned with the existing holes. For this reason, you might choose to allow the geometry of the sheet to determine mounting positions of valves, etc., rather than enforcing your own geometry. Forget about using a twist drill. The only way to make small holes (≈ 4 mm, or so) is to mark them out with a scribe and carefully file them using a needle file. If this sounds tedious, it's because it is, although a single-ended Abrafile in a handle speeds work (see [Figure 2.24](#)).

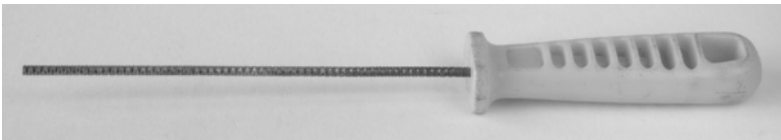


Figure 2.24
Single-ended Abrafile

A cheat that you can use (that does your tools no good at all, but saves time) is to put a round needle file in power drill, and plunging it gently back and forth, allow this to cut the hole. Used this way, needle files don't last very long, but you may feel that the cost of a new needle file is an acceptable trade against your time.

Fortunately, larger holes can be marked out and sawn using a tension file or a jeweller's saw. They're much easier.

Chassis punches still work well on perforated sheet, but it's important to align the cutting edge so that it is supported equally at its cutting tips before tightening the bolt and beginning cutting, otherwise it can pull itself out of alignment.

Deburring can **only** be done with a needle file, and takes time.

Cutting perforated sheet to size

Hand saws aren't really suitable because they try to follow the path of least resistance – which is usually straight down the middle of the holes, and not necessarily where you want the cut. A bandsaw is best, and a jigsaw will do, but the essential is a sharp fine blade running quickly.

It's also very difficult to file the cut edge because the file tends to be trapped by the holes. However, 160 grade emery cloth supported on a block of wood 3" by 5" does a superb job very quickly. Top and bottom burrs can be removed using 240 grade silicon carbide paper (again supported by the block), but the vertical burrs at each of the (many) holes have to be removed individually with a needle file.

Despite all these caveats about perforated sheet and its associated metalwork, the end results are well worthwhile.

The traditional chassis

The traditional method of construction was the folded aluminium chassis, and classic designs included beautiful engineering drawings complete with exact dimensions, folding lines, and all holes marked and dimensioned. The author can only assume that there were many more folding machines available in the early 60s, and that all constructors had access to a full sheet metal workshop.

Steel is not suitable for the chassis of valve amplifiers. Steel is magnetic, and allows leakage flux from transformers to flow through the chassis and induce currents into the pins of the valves. If a steel chassis is unavoidable, induction into the chassis can be greatly reduced by fitting a non-ferrous gasket between transformers and the chassis; 1.5 mm Paxolin is ideal.

Although it is sometimes possible to buy an undrilled folded aluminium chassis, even a small chassis needs to be 1.6 mm (16 swg) thick in order to be able to support the weight of the transformers. Although a pre-folded chassis might seem convenient to use, it is always awkward to drill holes in the sides of the chassis, because it is difficult to support the metal whilst it is being drilled.

Making a chassis using extruded aluminium channel and sheet

A far better alternative is to make the chassis out of separate pieces (see [Figure 2.25](#)).

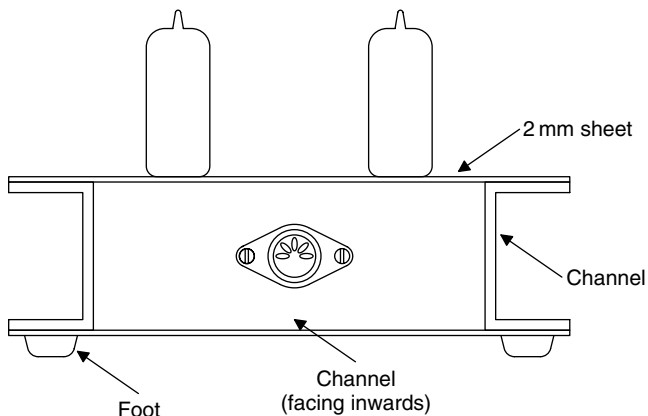


Figure 2.25

A very strong chassis can be made from aluminium channel

The top plate, to which most of the components will be fitted, is made of 2 mm aluminium, which is readily available either as off-cuts from an aluminium stockholder or from one of the electronic factors. Whilst it is tempting to use even thicker metal, many of the holes will be cut using chassis punches that can be damaged by thicker metal. Additionally, most valve holders were designed for 1.6 mm chassis, and whilst they can tolerate 2 mm, clearances become problematic if the plate is thicker.

The front, back, and sides are made from extruded aluminium channel cut to length. The sides have the channel facing out, thus providing convenient handles with which to lift the chassis, with the front and back fitting into the remaining space between the sides. The whole construction is then fastened with engineering screws and nuts. This form of construction has many advantages over the folded chassis:

- The chassis can easily be made to any convenient size using hand tools. It need not even be rectangular!
- Cutting holes in the chassis is now easy, because each surface can be properly supported whilst it is being worked upon.
- If modifications are required later (not uncommon), individual parts can be replaced if necessary.
- Aluminium channel tends to be quite thick (3 or 6 mm), making it a good heatsink, and threaded holes may be tapped into it, which is often convenient.
- If access is needed at one edge, that piece of channel can be temporarily removed.
- Looking at the bottom of the chassis, the channels are rigid load-bearing members to which feet and the safety cover plate can be easily fixed (which should ideally be perforated, to allow a cooling air flow). Fitting a cover plate to the bottom of a folded chassis is usually rather more difficult.

There are only two minor disadvantages. Firstly, the total top area is a little larger than the folded aluminium chassis, because some space is wasted at the sides by the outward facing

channel. Secondly, for safety, each separate piece of aluminium should be reliably earth-bonded to the top plate with star washers at one or more of the fixing points.

Corner pillars

Larger chassis are less rigid, yet they are usually called upon to support more weight. Although a chassis can be constructed by simply screwing channels to the top plate, a huge increase in rigidity can be gained by fitting 16mm square section corner pillars so that the channels form a self-supporting rigid picture frame onto which the top plate, or plates, are screwed.

The author faces each end of the pillars in his lathe – which allows the pillars to be made a snug fit inside the channel. Assuming you **can** gain access to a lathe, the sequence of events for machining a pillar to length with maximum precision and minimum tears is as follows:

- Carefully measure the internal width of the channel at the corners – not the open ends. The reason for this is that channel is rarely square (usually as a result of being clamped in a vice whilst you cut it to length) (see [Figure 2.26](#)).

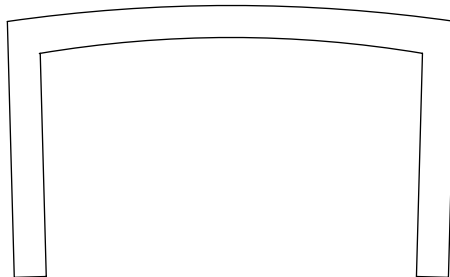


Figure 2.26

Clamping channel for sawing tends to deform it slightly

- Cut four lengths of 16 mm square bar to be 1–2 mm longer than your measured internal dimension.
- Take a pillar and face **both** ends in the lathe with minimum wastage. (Centring is not important.)
- Remove the pillar from the lathe, and measure its length with calipers.
- Subtract the required length from the measured length to find out by how much it is oversize. If you use digital calipers (easiest), you will probably obtain fractionally different answers depending on where you measure – indicating that the faces are not perfectly square. This doesn't matter, it's still far better than can be achieved by hand.
- Add 0.1 mm to the previous answer. In theory, this means that the finished pillar will fit into the channel with a 0.1 mm gap. In practice, it means that the pillar is guaranteed to fit in the channel, and because the channel is usually slightly distorted, it will probably grip the pillar snugly.
- Put the pillar back in the lathe, engage power, and gently ease the lathe tool up to the face until it **just** begins to cut.
- If you now set the dial on the lathe slide to zero, its graduations can now guide the removal of the required amount of material quickly and precisely.

Fitting the pillars

The neatest way of securing the pillars to the channels is with a pair of M5 screws from each channel into tapped holes in each pillar. This is done as follows:

- Two of your channels will have pillars inside them. Decide which these channels are, and drill M5 clearance holes at each end (5 mm, or perhaps 5.1 mm to allow adjustment on assembly) (see [Figure 2.27](#)).
- Gently ease a pillar into the channel. You may need to use a finger and thumb to spread the channel slightly to ease it in. Provided it is slightly in, the pillar can be gently tapped with

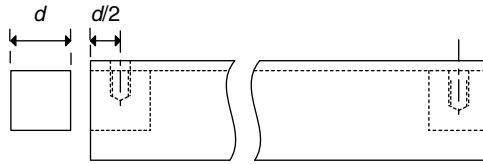


Figure 2.27

Positioning drill holes for pillars inside channel

a soft hammer until it is perfectly flush with the end of the channel, or press the two together on a flat surface. Running a finger over the join checks correct alignment far better than peering at it.

- Without disturbing alignment, use a small G-clamp to hold the pillar securely in place **without** obscuring the 5 mm holes.
- You can now mark through the 5 mm holes to determine where the pillar should be drilled and tapped. The ideal way of marking through is to make a dedicated centre-punch out of a piece of scrap metal having a 5 mm diameter stub faced by a poorly aligned tool so that a slight dimple is left at the centre. The punch can then be dropped through each hole in the channel and tapped lightly with a hammer to leave a precisely positioned dimple ready for drilling (see [Figure 2.28](#)).
- It is **vital** to **scribe** an identifying number onto each pillar and its corresponding position so that you know both where it belongs and its orientation. (Methylated spirits and/or finger grease removes pen or pencil.)
- Remove the pillar, use a 3 mm centre drill to start each hole, then drill a 4.3 mm (M5 tapping size) hole to a depth of 10–12 mm, and tap it M5.
- Slide the pillars into their correct positions in the channels, and secure them with M5 screws. Their alignment will be near-perfect, and a tiny amount of nudging will achieve perfection.

You now have a pair of channels with pillars secured at each end, and want to fit these to your other two channels to form the frame. It is very easy to make a mistake at this point, so be careful.

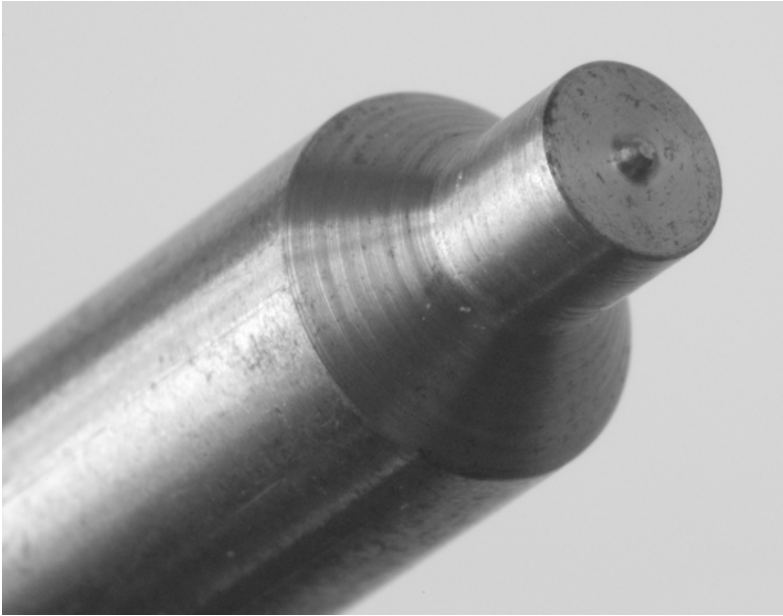


Figure 2.28

A 5 mm centre-punch allows perfect marking out and drilling

- Looking down onto the finished frame, the M5 clearance holes must be offset from the end of the channel by half the pillar cross-section, plus the channel thickness (see [Figure 2.29](#)).
- Having drilled the M5 clearance holes, lay out the entire frame on a clean flat surface, and use small G-clamps to secure it at the corners without obscuring the M5 holes. Check alignment carefully and adjust as necessary.

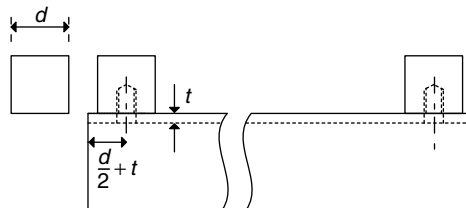


Figure 2.29

The pillars are offset by the thickness of the adjacent channel so the 5 mm holes must also be offset

- Mark through the M5 holes to the pillars, ideally using the dedicated M5 centre-punch.
- Remove the clamps and unscrew the pillars from the first pair of channels.
- Drill and tap the pillars.

You can now proudly assemble a perfectly aligned picture frame.

Fitting the top plate

The top plate inevitably needs lots of holes for all the valves and subsidiary components. It also needs to be secured to the picture frame. The easiest way to do this is to drill securing holes in the top plate, then clamp it to the picture frame with a couple of G-clamps and punch through to mark where holes need to be drilled in the channel. Rather than use nuts and screws, the author prefers to tap the channel, but you might have a different opinion.

Fitting carpet-piercing spikes

The corner pillars provide an ideal depth of material to be tapped M6 and fitted with loudspeaker carpet-piercing spikes. The spikes might not confer any sonic advantage, but they certainly allow a heavy amplifier to sit on a carpet without leaving a mark. The obvious way of drilling the required axial hole in the pillars is in the lathe during facing, but it is fiddly to centre square stock perfectly (even in a four-jaw chuck), so drilling the channel so that it aligns correctly with the (blind) hole becomes even more difficult. The author has found it easier to face the pillars without worrying about centring, then offer the assembled frame up to the drill and drill the M6 tapping holes (5.1 mm) straight through the channels and into the pillars. Because two channels at right angles are resting on the base of the drill stand, this ensures that the drilled hole is square.

You could tap M6 straight through the channel and into the pillars, but you might prefer to remove the pillars for tapping, and drill the channel M6 clearance. This reduces the amount of thread supporting the finished amplifier, but because the clearance hole guides the spike into the thread, it makes it slightly easier when fitting a spike blindly from underneath.

If you don't have access to a lathe...

Right-angled brackets cut by hand from 25 mm extruded aluminium angle make perfectly acceptable corner braces. The channel can be drilled and tapped to take carpet-piercing spikes, but be aware that because the channel only allows 3 mm depth of thread as opposed to 10 mm, this is not quite as strong as tapping into pillars.

Finishing

Aluminium can be spray painted, but paint tends not to stick very well to aluminium, and tends to chip off unless etching primer is used. Buying aerosol cans of car primer and top coat is quite expensive, and the fumes are most unhealthy. Nevertheless, this is one way of finishing the chassis.

A better method, but one that requires rather more planning, is to have the chassis anodised by a professional anodiser. Note that only aluminium can be handed to an anodiser; no foreign substances whatsoever are allowed. Surprisingly, this is actually quite cheap, because the pieces that you will hand over will be very small compared to the main batch that is being anodised. It may mean that you need to wait until a batch of your chosen colour is to be anodised, but the finished result will be far superior, provided that you have prepared the work properly.

An awkward problem with anodising is that it slightly increases the size of your work, and reduces the size of holes. If you have done a really nice piece of metalwork that snugs together beautifully, it is somewhat disconcerting to discover that it no longer fits. Commercial manufacturers use their prototypes to discover the amount by which their work should be undersize before anodising...

Both painting and anodising show up every imperfection of the underlying surface, so the surface cannot be too well prepared. A “brushed” finish can be obtained using reducing grades of silicon carbide (often known as “wet and dry”) rubbed along one direction only. If the final stage is lubricated with soap and water, a very smooth finish can be obtained. Alternatively, soap-filled wire wool scouring pads soaked in hot water can be very effective.

It is far better to begin with a good surface, so most stockholders keep aluminium that has one face protected with a plastic film. Keep the film on for as long as possible, and do not allow objects to fall on the sheet; aluminium is soft and easily marked. If using a scribe, keep your construction lines to a minimum, and score lightly with a **sharp** scribe – the marks that a blunt scribe makes are much more difficult to remove.

Whether you paint or anodise your chassis, make sure that you really have drilled **all** the holes you need. Drilling holes afterwards is invariably messy and easily spoils your finish.

CHAPTER 3

WIRING

The purpose of wiring is to translate a theoretical circuit diagram into a practical, working circuit. Unfortunately, there are many pitfalls, and careless implementation can ruin a good design. **Fortunately**, most of the pitfalls can be predicted easily, and therefore, avoided.

Tools

The cheapest tools are the most expensive ones. Cheap tools make it difficult to do a good job, so they waste time, and they need periodic replacement. Conversely, good tools mean that the quality of the job is limited purely by your own skill, they're nicer to use, and they last a lifetime.

Soldering irons

Obviously you need a soldering iron, but what is the difference between them and why are some so cheap?

The job of the iron is to heat the parts to be soldered to a temperature such that once the solder is applied, it melts

quickly and flows to form a perfect joint. Almost anything will do this, but the component may not work afterwards. Two thermal properties characterise the iron; thermal mass and temperature. Thermal mass is simply the mass of the hot part of the iron, and the greater this is, the more difficult it is for the proposed joint to cool it.

A cheap iron determines its temperature by only generating sufficient heat to match its losses to the environment, and to keep the tip hot enough to melt solder. As soon as it is touched to the joint, it begins to cool. If the work is not to cool the iron down so much that it is unable to melt solder, then the iron must contact the joint at a rather higher temperature. A higher thermal mass helps, but makes the tip clumsy to use. The upshot of all this is that the iron usually runs **too** hot, burns the flux in the solder, and may well damage the components. It will almost certainly cause tracks on printed circuit boards to lift if used for desoldering.

A better iron is thermostatically temperature controlled, and has an over-sized element (typically 50 W as opposed to 12–25 W). The iron is at the correct temperature all the time, but when the joint cools the tip, the thermostat trips and the over-sized element quickly restores the correct temperature. Additionally, temperature-controlled irons are generally low voltage (usually 24 V), which makes them longer lasting (the wire of the heating element is thicker, and less liable to break). Because the tip doesn't overheat, it is less likely to damage fragile components like polystyrene capacitors, and it lasts longer whilst being easier to keep clean.

All soldering irons suffer from leakage current between the heating element and the tip. The electrical insulation between the element and the tip must be thin in order to transfer heat quickly and efficiently, but because it is thin and hot, this insulator cannot be perfect (insulators become more leaky as temperature rises). Ohm's Law dictates that the leakage current

is determined by the electrical resistance of the insulation and the voltage applied to the element, so low-voltage irons have far lower leakage than mains irons. This point is significant because semiconductors such as CMOS digital ICs and low-noise discrete transistors or ICs can be damaged by the leakage current from a mains iron.

Because low-voltage temperature-controlled irons need a mains transformer and are more complex, they are invariably more expensive than cheap mains irons. However, they pay for themselves in time, because elements and bits last far longer, and they are less likely to damage a printed circuit board or an IC. And, in all that time, they are far nicer to use.

Remember, it is the **board** (whether PCB or tagstrip) that is expensive, not the components. Individual components can always be replaced, but if the board is wrecked, everything has to be replaced.

Earth bus-bars have a large cross-sectional area to reduce electrical resistance, but this inevitably makes them difficult to solder because they conduct heat away so efficiently that it takes considerable time for the iron to raise a part of the bar to soldering temperature, by which time the nearby polystyrene capacitor or IC has already been destroyed. The solution is to use a larger iron which can heat the bar faster – the author uses a 200 W temperature-controlled mains iron (see [Figure 3.1](#)).



Figure 3.1

A 200 W iron makes short work of heavy soldering

Tips

The part of the iron that contacts the work is the **tip**, or in more old-fashioned parlance, the **bit**. Old-fashioned irons had solid copper bits whose working surface would gradually be dissolved by the solder to become concave and would then need to be filed flat. Filing the bit was also the accepted way of cleaning these irons.

Modern irons use iron-coated tips to protect the copper, and should **never** be filed. The normal method of keeping the tip clean is to wipe it on a moistened sponge (specially made for the purpose by most soldering iron manufacturers). It is most important to keep the sponge moist, and most engineers keep an old washing-up liquid bottle of water near their iron for this purpose. If the tip becomes sufficiently contaminated by old, dusty solder that a quick wipe on the moistened sponge cannot clean it, then a wipe on one of the proprietary tip cleaners should do the job. If that fails, then careful scraping with a knife or wire wool will cure the problem. Do not be tempted to use silicon carbide or glass paper as the heat melts the glue and makes the tip even dirtier.

Tips come in many different shapes, sizes, and temperatures. The best tip is conical, with an oblique cut across the end to produce an elliptical soldering surface. These tips are usually specified by the width across the minor axis of the ellipse, and a good general purpose width is 2.4 mm. A wider tip allows you to get more heat into the work, and is better for heavier jobs but clumsy, whereas a fine 1.2 mm tip is excellent for pick-up arm wires or surface mount ICs but is unable to heat larger jobs. Ideally, you need a range of tips, and should be prepared to change tip with each soldered joint if necessary (see [Figure 3.2](#)).

Irons that use a magnetic thermostat, such as the traditional Weller “Magnastat” have tips that are available in different temperatures. The tips rely on the Curie effect whereby a magnet temporarily loses its magnetism at its Curie temperature, and this



Figure 3.2
Selection of useful tips

releases a spring-loaded ferrous shaft coupled to a micro-switch in the handle. For most work, a No. 6 (315°C) tip is ideal, but when working inside old amplifiers, a No. 7 (370°C) is better at burning away the dirt that insulates the old solder. Very occasionally, these irons overheat, and the cause is a dud tip. In the unlikely event of this happening to you, try changing the tip before diving in and replacing the switch (much harder).

The most recent irons have their temperature electronically controlled by a thermocouple, allowing adjustment during use – but this facility is a luxury for amateur work and although these irons are rapidly becoming cheaper, they are not really worth **extra** expense. A far more important question when you spot a “bargain” iron is the future availability of tips – if this is at all questionable, buy lots of tips at the time.

Gas irons

Portable gas irons burn butane gas without a flame in a tip containing a catalytic convertor (see [Figure 3.3](#)).

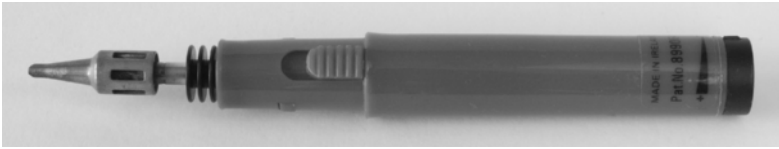


Figure 3.3

Although crude, portable gas irons can be useful

Temperature is crudely set by adjusting the gas flow. Because of their haphazard temperature control, gas irons have no place in quality work, yet there are times when they are invaluable. Very occasionally, a single joint needs to be made (or broken) but a mains socket isn't nearby, so rather than finding an extension lead and the mains iron, a gas iron can be quite handy. The other use is when wires need to be soldered to a replacement heating element for the mains iron...

Solder

The traditional electronic solder was 60/40 self fluxing solder. The 60/40 referred to the ratio of tin to lead, and the flux is a chemical, that when heated by the iron, cleans the surfaces to be soldered and allows a good joint. Industrially, the use of lead is being heavily discouraged, so pure tin solder is now available, but this requires a rather higher soldering temperature, making component damage more likely.

Other solders are available, some of which have sufficiently powerful fluxes to cut through surface aluminium oxide and enable soldering to aluminium. **Never** attempt to use aluminium solder for electronic work, and once a soldering iron tip has been contaminated by aluminium solder, it should never be used for electronic work.

It is most important to keep your solder **clean**; keep it in an airtight jar when you're not using it. Drawing it through a clean

cloth moistened with isopropyl alcohol before use removes surface contamination, and will significantly improve the quality of your soldered joints. Alternatively, a light rub using metal polish intended for brass easily strips surface contamination, but needs careful buffing to remove all traces of the polish.

In small quantities, solder is expensive, so buy a 500 g reel of solder. In this quantity, there are various different sorts of solder in various thicknesses. An excellent general purpose solder is 0.7mm low melting point solder with 2% silver content. For surface mount components you **must** use silver loaded solder; otherwise the silver in the plating of the components leaches out and they won't solder. Silver loaded solder will allow you to make far better soldered joints than ordinary 60/40, and you will be seduced by its superior properties.

Soldering and shrink back

Assuming that you have an iron at the correct temperature, with a correctly sized clean tip, and some clean solder of the appropriate type, how do you ensure that you make a perfect soldered joint?

The surfaces to be soldered must be clean, and the soldering iron tip must be clean and free of dross. Soldering works by the solder combining intimately with the surface metal of the components, and dirt hinders this process. The solder should be applied to the point of contact between the work and the iron such that it melts and flows immediately; the tip of the iron should then be wiped clean on the moistened sponge (see [Figure 3.4](#)).

Clean surfaces will **wet** perfectly, and surface tension causes the solder to flow instantly across the work to form a perfect joint. Dirt causes the solder to form globules on the surface that do

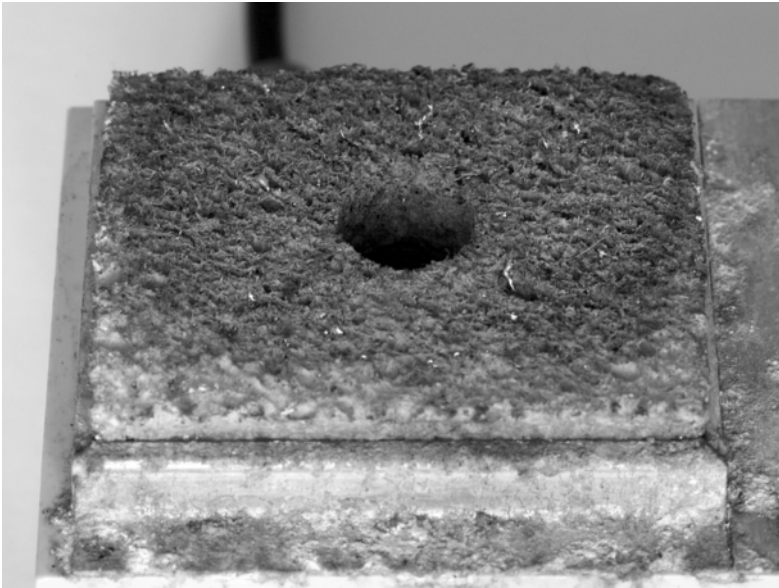


Figure 3.4

The tip should be wiped on a moistened sponge before and after every joint

not wet the joint, and defective joints are therefore known as **dry** joints. There are many variations between these extremes, but it can generally be said that good joints are made quickly, whereas dry joints are more likely to occur when the iron has been in contact with the joint for more than a second or two. Don't dab at the solder. Press the iron firmly into the work so that the heat flows quickly, apply the solder firmly, let it flow, and release the iron immediately.

The best joints are **mechanical** joints. If the parts to be soldered are already unable to move relative to one another, then they will not move as the solder solidifies, and a perfect, shiny, joint should result. If there is any movement whilst the solder is cooling, a dull, dry joint will result.

The best joint is the **first** joint. Any subsequent resoldering of a joint degrades the joint because it allows further oxidisation of

the heated materials. If you are forced to resolder an old joint, remove the old solder and replace with new. The fresh flux will ensure clean surfaces, and the fresh solder will be not be contaminated with dross.

When you solder insulated wires, you are likely to encounter a phenomenon known as **shrink-back**. When you strip the wire, you make a circumferential cut into the insulation, but to avoid nicking the conductor, the cut is not made sufficiently deep to cut all of the insulation. The remaining insulation is broken when the unwanted insulation is pulled away. Inevitably, this has the effect of stretching the remaining insulation along the wire. When the conductor is heated for tinning, the stresses in the plastic insulation relieve, and the insulation apparently shrinks back, leaving more exposed conductor. There are various ways of reducing shrink-back:

- Use wire that doesn't suffer from shrink-back. Silver threaded down PTFE tubing doesn't suffer from shrink-back because the coefficient of friction of PTFE on silver is so low that the moment the insulation is broken, it springs back to its original length. Again, use a sharp scalpel to make the cut, but instead of rolling the scalpel around, simply press it firmly into the PTFE, and twist the insulation to make the cut.
- Minimise the stress that causes shrink-back. Make the cut in the insulation using a sharp scalpel. To avoid nicking the conductor, the blade has to be rotated around the wire without any cutting motion. The wire can then be bent very slightly at the cut to break any remaining insulation, rotated slightly, and bent slightly until the insulation is completely broken. The act of bending to break the insulation stretches one side of the insulation but compresses the other, minimising the total stretching, and therefore the amount of shrink-back.
- The degree of shrink-back is proportional to the time the iron is in contact with the work and its temperature. This is why it is so important to make fast joints using the correct

size tip. A tightly wrapped mechanical joint allows surface tension to make the solder flow quickly.

- Accept that there will be shrink-back, and pre-tin the conductor, so that shrink-back occurs, **then** solder the wire in place. Unfortunately, this method precludes mechanical joints using stranded wire, but is satisfactory with solid core wire.

Solder tags

Solder tags are used to make electrical connections to items that have too much thermal mass to be soldered directly. A soldered earth bond directly to a chassis would not only cool any reasonable soldering iron, but the chassis is probably made of aluminium, necessitating special solder, so this is a legitimate use of a solder tag.

Be warned that as parts age, they corrode and become more difficult to solder, and printed circuit board manufacturers refer to this as **solderability**. Solder tags inevitably lie in drawers for decades before being used, so you may occasionally find a solder tag that flatly refuses to solder. A great deal of time can be wasted abrading the surface in a futile attempt to persuade the tag to solder. If you are unlucky enough to find a tag that won't immediately solder, discard it and pick another, ideally from a different batch.

A solder tag inevitably adds a mechanical contact in series with the soldered joint, so it is always the second-best choice, and should be avoided if at all possible. As an example, traditional amplifier output terminals often came with solder tags, but could, and should, be soldered directly. Part of the thread needs to be removed to form an exposed brass pillar suitable for soldering, this is most easily done with a file, and although a prettier job can be done in a lathe, the setting-up takes time. The wire should be tightly wrapped one complete turn round the exposed pillar and the iron applied to the pillar. Once the pillar is hot enough to easily melt solder, the joint can be finished by sliding the iron so

that it touches the wire directly, and a perfect joint will be formed. The thermal mass of these terminals can be greatly reduced by unscrewing the part that grips the loudspeaker cable so that it wobbles freely on its thread, and thus does not have to be heated by the iron. The small 4 mm combination terminals can be soldered by a conventional 50 W using a large tip, but the large 32 A terminals are best soldered with a 200 W iron.

Desoldering

Sometimes you will need to desolder a joint. If the joint is mechanical, it is best to cut the wires away, otherwise the prolonged heat whilst bending wires will damage the wire or component. The remaining joint has fragments of wire in solder and these must be removed.

The solder can be removed by one of two methods:

- Desolder **wick** uses surface tension to wick the solder into copper braid which is discarded once contaminated by solder. Solder wick must be kept clean if it is to work, so store it in an airtight container. This method causes the least damage to the work, but is wasteful and expensive.
- Alternatively, solder can be drawn off with a vacuum. In industry, vacuum-desoldering stations are common, but they are extremely expensive to buy, and must be maintained and used carefully. The far cheaper alternative is the handheld, spring loaded, solder sucker (see [Figure 3.5](#)).



Figure 3.5
Handheld solder sucker

The tool looks like an oversized pen, and has a PTFE tip that is placed directly in contact with the molten solder, whereupon the trigger is depressed and the sucker sucks up the solder, hopefully. Even handheld solder suckers need to be looked after if they are to work. The inside of the sucker needs to be cleaned periodically, or it will jam, and even a temperature-controlled iron eventually damages the PTFE tip, necessitating replacement because it can no longer seal against the solder to draw it into the sucker.

The main problem with solder suckers is the recoil. The sucker works by accelerating a plunger within the sucker away from the work. The recoil from this drives the PTFE tip into the work with sufficient force to break ceramic stand-offs, or to kick tracks off old printed circuit boards. For this reason, solder suckers should be used very carefully, and you will want to revert to braid on delicate work.

Once the solder has been removed from a joint, the remaining fragments of wire can be easily curled away with fine nose pliers even when the work is cold. But don't ever do this on a printed circuit board, if you apply force to a track on a printed circuit board it is liable to lift. Wire fragments should be delicately removed whilst the remaining solder is still molten. Fine tweezers can be handy here as they usefully limit the force that can be applied.

Hand tools

In addition to a soldering iron, solder, and some means of desoldering, you need hand tools to dress leads and fit components. It is easy to be seduced by all the wonderful pictures of tools in a catalogue, but you will find that for day-to-day use, you need only a few tools, provided that they are of excellent quality. Good hand tools cost more, but they last far longer, and are cheaper in the long run...

Cutters

You only **really** need two sizes, one for cutting cable from wires to heavy mains cable, and one for cutting component leads precisely. The author has many different cutters – most of which are used only rarely.

Most electronic factors stock superb cable cutters that at first sight look terrible, but are well made from decent materials and are a delight to use. Although these cutters will slice cleanly through any multicore copper cable that will fit in their jaws, whatever you do, don't use them for cutting steel-reinforced cables, or even bicycle brake cable. One attempt destroys them. Curiously, vets also sell these cutters (at twice the price) for cutting dog's toe-nails. If you have these, and use them for most work, then the only other cutters you need are box-jointed Lindström "Supreme" semi-flush micro cutters (see [Figure 3.6](#)).



Figure 3.6
Cable cutters and micro-cutters

Lindström semi-flush micro cutters seem to last the author about ten years before becoming too blunt to use, but cutters with tungsten carbide cutting edges are now available, and should last even longer. Flush cutters are also available, and they are particularly useful when replacing components. Unfortunately, the necessarily smaller included angle of their cutting edges makes them sharper but far less durable (see [Figure 3.7](#)).

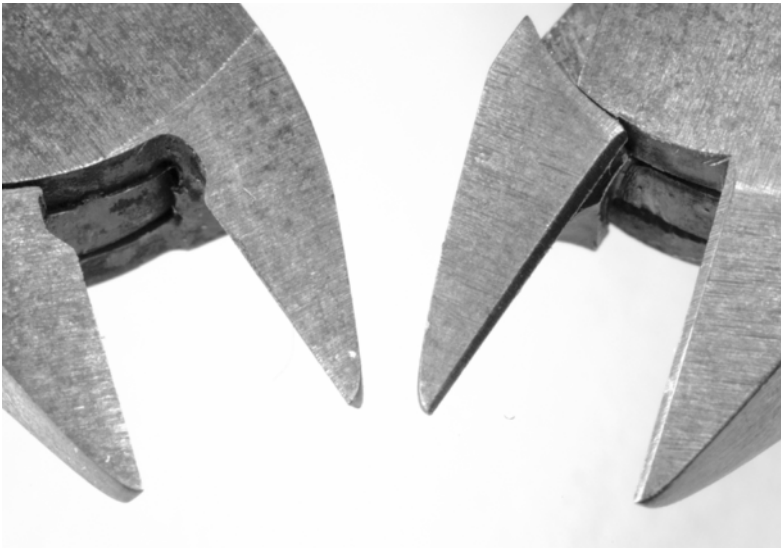


Figure 3.7

Note the difference between flush and semi-flush cutting edges

Once the cutting edge has been decided, choose an ergonomically designed handle and springing system if possible – you will use these cutters a lot. If you buy flush cutters, only use them when nothing else will do, or you will wear them out very quickly.

Pliers

Again Lindström “Supreme”, short jaw (21 mm). These are for dressing component leads, not for removing your car exhaust!

A pair of short-jawed heavy duty combined pliers/cutters is handy too, but you probably already have a pair for dealing with your car or bike (see [Figure 3.8](#)).



Figure 3.8
Heavy duty pliers and micro-pliers

Occasionally, neither of the previous pliers will do, and you need a pair of long-nosed pliers. They are not nearly as nice to use, and you usually use them because you are being forced into bodging (see [Figure 3.9](#)).



Figure 3.9
Long-nosed pliers

Wire strippers

There are many different sorts of wire strippers, and personal preference is important. The chosen wire stripper should be able to strip cleanly without nicking or cutting the strands of the wire beneath the insulation. The author strips most wire using a scalpel or traditional wire strippers (see [Figure 3.10](#)).

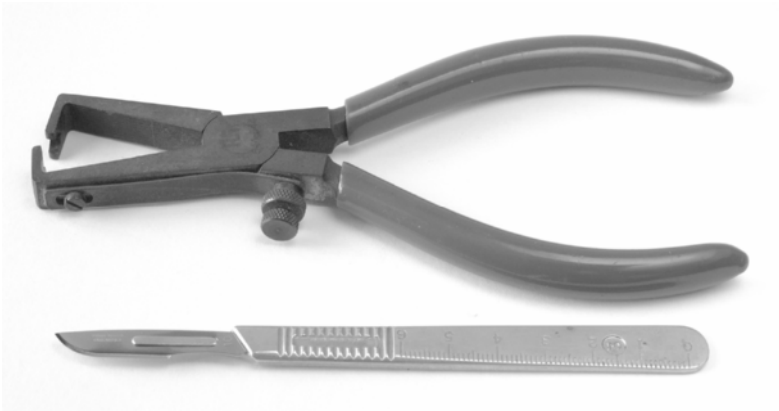


Figure 3.10
Scalpel and PO strippers

If you use wire-wrap wire you must use a dedicated wrapping and stripping tool. These are outrageously expensive for what they are, but nothing else strips the insulation cleanly without nicking (see [Figure 3.11](#)).

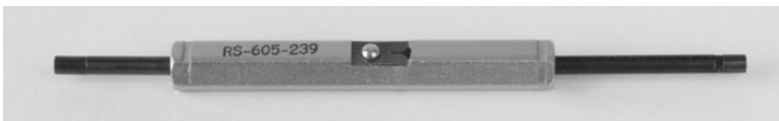


Figure 3.11
Wirewrap tool with integral stripper

Screwdrivers

You really can't have too many good quality accurately ground screwdrivers, and with care, they last forever, so buy the right ones. Stubby screwdrivers are not recommended because they cause inaccurate screwing and damage screw heads. Poor-quality screwdrivers are inaccurately ground and made of inferior metal. The result is that they slip and damage themselves and the screw head. That doesn't sound too bad until you realise that it happens when you have to apply most force – when the screw is seized. Once you have damaged its head, removing a screw becomes far, far harder – it's so much easier to avoid damaging the screw head in the first place.

Flat-bladed screwdrivers

You need a 3 mm flat blade screwdriver, often known as an “electrician's” screwdriver. An extra-long screwdriver (250 mm shaft length) with a 5 mm flat blade is extremely useful, and the bigger screwdrivers can be useful too, particularly 7 mm. You also need a very small screwdriver; it usually has a yellow handle and is about 60 mm long in total, with a blade width of about 1.5 mm. Buy two – they disappear.

Pozidriv and Phillips screwdrivers

There is a world of difference between these two “crosshead” screwdrivers, so you will want to avoid damaging screws and screwdrivers. A Phillips screwdriver has a sharper, more pointed tip, whereas Pozidriv is rather stubbier (see [Figure 3.12](#)).

Supadriv screws are almost identical to Pozidriv, but have a secondary set of splines between the primary splines. The extra splines enable greater torque to be transferred to Supadriv

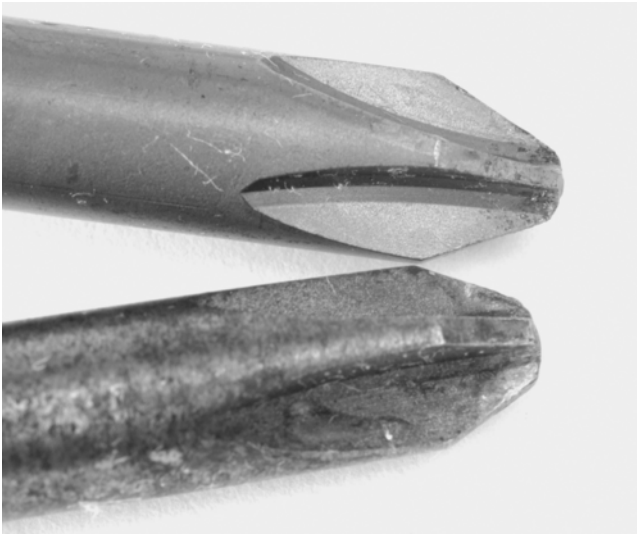


Figure 3.12

Note that the Phillips tip (upper) has a sharper angle and is longer than the equivalent size of Pozidriv tip (lower)

screws, but they are perfectly compatible with older Pozidriv screwdrivers. Choose the extra-long version with an ergonomic handle as these are useful for computer monitors and aid precise screwing. Sizes 0, 1, 2 are useful, 3 and 4 are a luxury. European Pozidriv and Supadriv screws are identified by the identifying radial lines between the splines of the screw slots (see [Figure 3.13](#)).

The author had to tip out his large box of non-BA and non-Metric screws and search diligently before finally unearthing half a dozen Phillips screws. Since Phillips screws are so rare, why are there so many Phillips screwdrivers ready and waiting to mangle Pozidriv screws? Sadly, Pozidriv screws without identifying radial lines are becoming common, particularly on Oriental equipment, so try a Pozidriv screwdriver for fit **before** a Phillips, even if the identifying radial lines are absent.

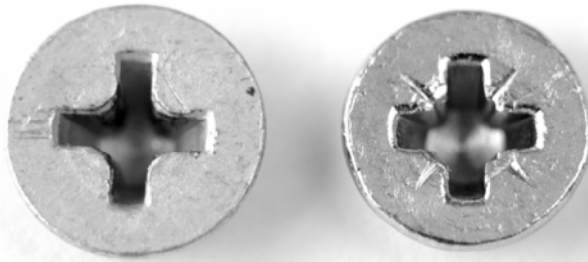


Figure 3.13

Phillips (left) versus Supadriv/Pozidriv (right) screw head. Note the identifying radial lines on the face of the Supadriv head between the splines

Allen (hex) keys and drivers

Buy a good quality set of long arm Allen keys in Imperial **and** Metric sizes as it is essential to have the correct size. Be careful when using long arm Allen keys. Although they make using chassis punches much easier, they are capable of splitting screw heads made of inferior metal if over-torqued (see [Figure 3.14](#)).

Be wary of ball-nosed keys – they allow easy access, but unless they are made of excellent metal, the reduced contact area at the head wears quickly and is easily damaged.



Figure 3.14

A 6 mm long arm versus normal 6 mm Allen key

If you use a lot of Allen screws, you might want to consider buying a set of drivers – they are much faster to use than keys.

Nuts and associated tools

Nuts are undone with spanners or nutrunners, **not** with pliers! A set of open-ended BA is needed for traditional British valve amplifiers, but modern electronics uses metric fasteners. Electronic equipment uses very small nuts, and if the spanner is inaccurately ground, it will slip and chew the nut, so it is essential to use top quality spanners (see [Figure 3.15](#)).



Figure 3.15
Selection of BA spanners

Nutrunners are available in all sizes, but one for potentiometer nuts is particularly useful, and prevents visible gouging of front panels by spanners (see [Figure 3.16](#)).



Figure 3.16

A nutrunner for potentiometers is particularly useful

A Bahco 6" adjustable spanner will do nicely for everything else provided it is adjusted to grip the nut as tightly as possible before applying torque, and is sufficiently good to be used for careful roadside repairs of bikes with parts labelled "Ducati" or "Campagnolo".

A nutlauncher is an incredibly useful tool. Looking like a syringe, when the plunger is pressed, three wire hooks spread out of the far end and grip the nut allowing you to spin it lightly onto the end of an otherwise inaccessible screw. Wonderful! (see [Figure 3.17](#)).

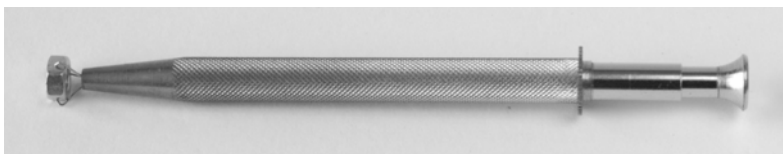


Figure 3.17

A nutlauncher makes short work of starting nuts buried deep in wiring

Once the nut is engaged, a nutrunner or spanner can finish the job.

Sometimes, with the best will in the world, a part will drop off and fall deep inside a heavy amplifier. Rather than lifting the amplifier upside down and shaking it until the part falls out, a magnet is extremely useful for retrieving small steel parts. Many tool shops sell a tool that is effectively a powerful magnet fitted to the end of a telescopic aerial and packaged to look like a pen (see [Figure 3.18](#)).



Figure 3.18

A magnet on telescopic mount makes retrieval of small steel parts easy

Scalpel

Scalpels are extremely useful, and the small handle version is best suited for electronics, with either No. 10, or No. 10A blades. Be careful with scalpels – they were designed for cutting flesh. You will find that the blades lose their edge very quickly, so buy blades in bulk, and be prepared to fit a fresh blade the moment you notice a lack of keenness. If you have the option, buy unsterilised blades – they are cheaper, and you don't intend surgery. A guard made from layers of heatshrink sleeving can be made to cover the blade when not in use (see [Figure 3.19](#)).



Figure 3.19

Scalpel with guard made from heatshrink sleeving

Hot air gun

Tool catalogues stock all sorts of expensive hot air guns for shrinking heatshrink sleeving, but a hot air gun intended for stripping paint works just as well, and is far cheaper. Never attempt to shrink sleeving by heating it from only one side, move the gun continuously around the work to heat it evenly

from all sides. The airflow from hot air guns is somewhat indiscriminate, and can easily damage surrounding components such as electrolytic capacitors, but a heat shield of aluminium cooking foil works wonders (see [Figure 3.20](#)).



Figure 3.20

A temporary aluminium foil heat shield protects other components when using hot air gun on heatshrink

Not only does the shield protect other components from the hot air blast, but the turbulence caused by blowing air at the shield means that the heatshrink sleeving is heated evenly from all sides, and shrinks in half the time, which further reduces the risk to surrounding components.

Marker pens and digital cameras

It may sound obvious, but if you label things **before** you take them apart, life becomes so much simpler. Alternatively, a digital camera is a very useful tool because if you take shots of the work as it is disassembled, you have the perfect guide to reassembly.

Lighting

You can't have too much light to work by. If you are able to have a dedicated bench, try to position it near a window, and fit dedicated lighting over the bench. A 100 W fluorescent tube provides plenty of light, but you may find that the increased flicker compared with tungsten is more tiring to your eyes.

If you aren't able to have a bench with dedicated lighting, at least have an adjustable desk lamp that can be positioned over your work. In addition, a powerful torch can be very useful for finding nuts and washers that fall into the depths of an amplifier.

Toolbox

Your precision hand tools should not be thrown in a toolbox together with old spark plugs and oil filters. Keep them in a clean partitioned box of their own and don't lend them to **anybody**.

Techniques

All AC power wiring generates an external field that can induce audible hum into signal wiring. Heater wiring is the obvious problem, because it will unavoidably be close to sensitive signal wiring, but AC mains and the high-voltage AC to rectifiers can also cause problems.

To save time on assembly and any subsequent maintenance, any fixing screw should be just long enough to use all of the thread in the nut and no longer, and this dictum is particularly important when securing valve bases. Overlong screws on valve bases make heater wiring particularly difficult. Unfortunately, if you decouple your heaters to the chassis at the valve base with 10 nF capacitors (ideal), the necessary solder tags and star washers inevitably require a longer screw. You just have to grin and bear it.

Electromagnetic fields and heater wiring

The electromagnetic field is due to the **current** flowing in the power wires, which induces currents in any nearby signal wiring. This means that not only are valves with 12.6 V heaters cheaper (contrast the price and availability of NOS 6SN7 with 12SN7), but they're better, too! Halved heater current means 6 dB less heater-induced hum.

Heater wiring is usually taken from a winding on the mains transformer to the nearest valve, and looped through, from one valve to the next, until each valve has heater power. The input valve is the most sensitive stage, so this should be the last in the heater chain, in order that the wiring leading to this valve carries the least current.

To minimise the external electromagnetic field, the heater wire should be tightly twisted. This means that although any given twist induces a current of one polarity, the twists either side of it induce opposite polarity, and so the fields tend to cancel. This twist should be maintained as close up to the pins of the valve as possible, and when one phase of the heater wire has to go to the opposite side of the base and return, as is the case when wiring past an ECC83/12AX7 to the next valve, the wire should go **across** the base and be twisted as it passes across. Admittedly, the return current is less than the outgoing current, but some cancellation is better than none.

The worst way to wire heaters would take the incoming pair connecting to the two heaters pins from one side, then loop around the opposite side to form a circle of heater wiring around the valve base (see [Figure 3.21](#)).

Heater wiring leading to valves using B9A sockets such as EL84/6BQ5, etc., is best twisted from 0.6 mm (conductor diameter) insulated solid core wire, which is rated at 1.5 A. Octal valves generally require more heater current, so the larger tags on their

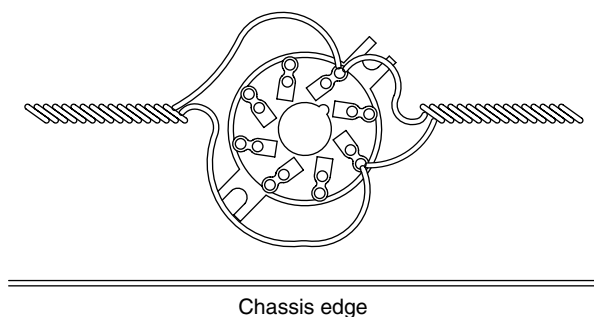


Figure 3.21

Poor heater wiring at valve socket: $\text{Hum} \propto \text{enclosed loop area}$. Note that the wiring should have been butted up against the chassis edge

sockets can accommodate thicker wire, which could not have been connected to the pins of a B9A socket. When wiring to valves other than rectifiers, it is useful to use a different colour for each phase, and the author has always used black and blue. When wiring to a push–pull output stage, if the same colour goes to the same pin on each valve, then the hum induced within each valve will be the same phase, and will be cancelled in the output transformer. (This argument assumes that both valves were made by the same manufacturer to the same pattern.)

Valve rectifiers such as GZ34 or GZ37 not only have a dedicated 5 V heater supply, but they also have the incoming high-voltage AC, so the two should be clearly distinguished. The author uses a red twisted pair for the HT, and a blue twisted pair for the heater.

Twisting wire is easy. Cut equal lengths of wire to be twisted, pair the wires together at one end and clamp them in a vice. Gently tension both wires equally at the far end, and grip them in the chuck of a drill. Hold the wires reasonably taut by pulling on the drill and start twisting. When the wire begins to accelerate you towards the vice it will have about ten twists per inch. Switch off, and **whilst maintaining tension** by holding the wire with your fingers, undo the chuck. The wire will now try to untwist, and if allowed to do so suddenly, it will tie itself

in knots. Gently release the tension in the wire and release from the vice. You now have perfectly twisted wire.

You will find that it is easier to achieve a perfect twist on longer lengths of wire than short ones, because it is easier to equalise the tension between the wires. Equal tension is important because if one wire is slack compared to the other, it tends to wrap itself around the tighter wire (which remains straight). For this reason, it is worth twisting 4m or even 6m at a time, but these longer lengths are more easily twisted with a power drill (ideally, a battery drill because they tend to run slower), whereas shorter lengths can be twisted with a hand drill.

Although perfectly satisfactory with commercial sleeved copper wire, the previous method breaks 99.99% pure (fine) silver wire posted down PTFE sleeving. Sadly, silver wire in PTFE sleeving has to be carefully wrapped by hand. Instead of a drill starting a light twist evenly along the entire length of the wire and tightening it, the required wrap must be locally applied with force carefully applied by fingers either side of each wrap, starting at one end of the wire and progressing towards the other.

It is extremely difficult to twist even slightly different gauges of wire together, and you might wonder why this should ever be attempted. The most likely scenario is that a flying lead from a mains transformer primary needs to be taken to a single pole mains switch at the far side of the chassis, return, and then go to the mains inlet. In this instance, current flows from the mains inlet to the switch and returns along an adjacent wire, so it makes a great deal of sense to twist these two wires together to reduce hum induction. Unfortunately, unless the mains transformer has **very** long leads, it is unlikely to be able to reach to the mains switch and back, so you are forced to twist it with a wire from your stock that is almost certainly slightly different. When you do this, one wire almost invariably remains almost straight, and the other wraps around it. The solution is to gently tension the wire that wanted to wrap, and wrap around

it the wire that wanted to be straight. This method is a little fiddly, but enables two different wire types to be twisted evenly, thus ensuring cancellation of hum.

Although solid core wire perfectly retains a tight twist, multi-core wire tends to separate, reducing cancellation. The way to avoid this problem is to pre-tension each wire equally by individually twisting it in the **opposite** direction of the final twisted pair. Without releasing each wire's pre-tension, grip the wires in the vice and chuck, and twist them together. The effect of this is that as the wires are twisted together, the pre-tension is relieved, so there is no latent force trying to separate the final twist. Although this method works very well, solid core wire retains the tightest twist and can be positioned more precisely, so it is superior for heater wiring.

Electromagnetic fields decay with the square of distance, so heater wiring runs should be as far away as possible from signal circuitry as possible, and only come up to the valve at the last possible moment and in the most direct manner possible.

Valve sockets should be oriented so that the pins receiving heater wiring are as close to the chassis wall as possible, and heater wire should **never** loop round a valve (except for rectifier valves, where hum is not an issue).

Electrostatic fields and heater wiring

The electrostatic field is due to the **voltage** on the wiring.

Heater wiring should be pushed firmly into the corners of the (conductive) chassis, since the electrostatic mirror at the corner tends to null some of the electrostatic field. Heater wiring must not run exposed from one valve to the next, but should return to the corner of the chassis to re-emerge at the next valve. These strictures mean that good heater wiring requires considerable

time/cost, so modern commercial amplifiers sometimes skimp on the quality of their heater wiring, yet good wiring allows an RIAA stage to use AC heaters...

AC heater wiring should be connected to the transformer in a balanced fashion. Unfortunately, heater wiring **must** have a DC path to HT 0 V in order to define the heater to cathode voltage, and this can be achieved in various ways (see Figure 3.22).

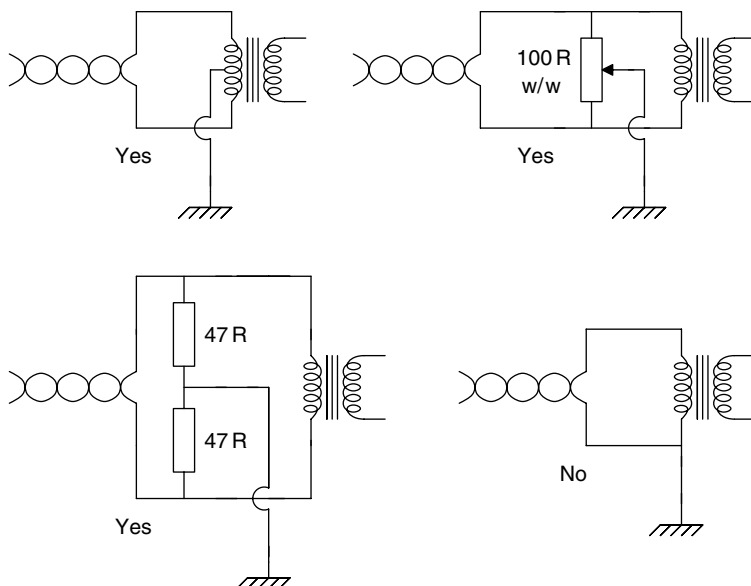


Figure 3.22
Defining DC heater potential

The worst way to define the DC path is simply to connect one side of a transformer winding to 0 V. This ensures that one phase of the wire induces no hum, whilst the other phase induces maximum hum.

The ideal way of defining the DC path is to use a transformer with a centre tap on the heater winding, but if this is not available, fixed or variable resistors can be used to derive a

midpoint. Accurately matched resistors used to be rare, so a variable resistor known as a **humdinger** control used to be fitted, and adjusted for minimum hum. Once an LT midpoint has been derived, and connected to HT 0 V, each wire has equal voltage (but opposite phase) hum, and the electrostatic fields tend to cancel.

The previous cancellation cannot be perfect, and for ultimate reduction of heater-induced hum, we should screen the heater wiring with braid or thin-walled aluminium tubing (available from modelling shops), and/or use DC heater supplies. Even when using DC heater supplies, it is worth treating the heater wiring as if it were carrying AC, as this will ensure that the finished project has **no** heater-induced hum. The output of a DC heater regulator has virtually no AC present, so the LT 0 V can be connected directly to the HT 0 V.

Heater wiring is the first piece of wiring to go into a project, thereafter, it is obscured by signal wiring. Once all the other wiring is in place, it is impossible to replace the heater wiring, so it **must** be installed correctly (see [Figure 3.23](#)).



Figure 3.23

Good heater wiring is pushed deep into the corners of the chassis and maintains its tight twist all the way to the valve socket

Because it is so difficult to make changes to heater wiring, it is a very good idea to immediately do the mains wiring and test it all by plugging all the valves in and making sure that they glow. Quite apart from saving tears, seeing a set of glowing heaters gives an important psychological boost before tackling the more complex wiring.

Mains wiring

Mains wiring should be as short and direct as possible because the wire's insulation is often quite thick, so it cannot be twisted well. Mains wiring inevitably generates considerable electrostatic interference fields, so put the mains switch near the back panel, and if necessary, add a mechanical linkage to bring its control to the front.

Modern semiconductor equipment **sleeves** all exposed mains wiring with rubber or PVC sleeving such that it is moderately safe to rummage inside a piece of powered equipment. Valve amplifiers operate on such high voltages that it is **never** safe to rummage in powered equipment, and even unpowered equipment should be approached with caution. Safety is therefore not greatly improved by sleeving mains wiring, but it is still good practice to sleeve mains wiring with heatshrink sleeving or purpose-made rubber boots to fit over IEC sockets or fuseholders (see [Figure 3.24](#)).

The **only** approved 3-pin domestic mains connector is the IEC plug/socket, which is familiar to many as the connector fitted to computers and electric kettles. They are available with integral fuseholders which allow much safer construction and are thoroughly recommended.

The mains connectors used on classic valve amplifiers are, without exception, outrageously dangerous.

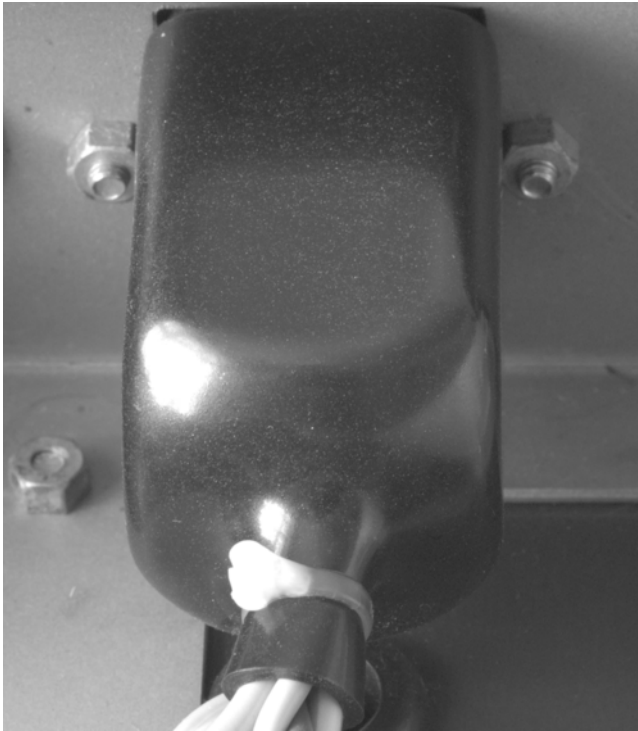


Figure 3.24
IEC input connector sleeved with insulating boot

Mains switching

To switch a piece of equipment off, all we need to do is to break the circuit from the source of power. A mains switch could therefore equally well be inserted in the live or neutral wire, and still perform the job, and this is known as single pole switching. However, a switch in the neutral leaves all internal mains wiring live and constitutes a shock hazard within the equipment. Single pole switching should therefore **always** switch the live circuit to minimise shock hazard.

Double pole switching switches both the live and the neutral, and ensures safety even if the live and neutral wires are reversed.

CD players often use non-polarised “shaver” plugs on their rear panels, which allow live and neutral to be reversed, so they commonly use double pole mains switching.

Where there is **no possibility** of live/neutral reversal, single-pole switching is safer, and more reliable, because failure of the switch ensures a break in the live connection to circuitry. A double-pole mains switch has twice as many contacts to fail, and if the neutral contact fails, the equipment could appear to be safe, even though the mains wiring is connected to live mains, and still constitutes a shock hazard.

Fuses

A fuse is a piece of fine wire having resistance, connected in series with the circuit to be protected. If excessive current passes through the wire, it heats in accordance with I^2R , and heats sufficiently that it melts, or ruptures. A fuse is a single-pole switch, and should therefore be connected in the live wire, **before** any other circuitry, such as a mains switch.

Clearly, the wire in a fuse must lose some heat to its surroundings, so a mild overload could allow much of the heat that should have melted the fuse to escape, whereas a short duration gross overload cannot lose so much heat and ruptures the fuse quickly. As an example, a 13 A fuse to BS1362 (as fitted to a UK 13 A domestic plug), ruptures in 0.4 s at 100 A, but needs 10 s at 50 A.

To calculate the mains fuse rating, the total power consumption of each individual load on the mains transformer should be summed to find the total load taken from the mains. The current drawn from the mains may now be found, and the fuse rating should be the next rating above this. This calculation contains many sweeping assumptions and approximations, but fuses are not accurate either, so the method will be found to be satisfactory.

Some equipment, particularly if it includes a toroidal mains transformer, draws a large inrush current, but its working current is much lower. To cope with these requirements, **anti-surge** or **timed** fuses are available, which can withstand short overloads. These fuses usually have their values preceded by a “T”, so T3.15 A refers to a timed 3.15 A fuse.

Protecting each output of a multiple winding mains transformer is difficult for the following reasons:

- Fuses are very rarely fitted to HT supplies because they offer only very limited protection to the output valves. In a Class A amplifier, the output valves are usually run at precisely their maximum anode rating, so a doubling of anode current quickly causes damage. However, a fuse may not blow with an overload as small as this, so little protection is offered. Fuses can be fitted to Class AB amplifiers, and are advisable for OTL designs, but their non-constant resistance can cause distortion.
- Fuses are never fitted to heater supplies because heater circuitry is normally so simple as to not warrant a fuse. Failure of heater supplies often causes damage elsewhere in a DC coupled amplifier, as valves switch off and anode voltages rise to the full HT voltage.
- Grid bias supplies to output valves should **never** be fused because failure of this supply would immediately destroy the output valves.

For these reasons, most valve amplifiers do not have any fuses other than a fuse on the primary of the mains transformer.

Glass-bodied fuses should **never** be used to protect high-voltage circuits such as AC mains. A short circuit causes the fuse to rupture instantly and vaporise, thus depositing a conductive metal film on the inside of the glass which continues to pass a small current, but heats the glass. If glass is heated sufficiently, it becomes a conductor in its own right, and so the fuse has

failed to protect the circuit. Fuses suitable for high-voltage use have ceramic bodies filled with sand to prevent the creation of a continuous conductive film.

To prevent tampering, mains and high-voltage fuseholders accessible from the outside of the chassis should require a tool to release the fuse.

Class I and Class II equipments

Class II appliances have all hazardous voltages ($>50\text{ V}$) **double insulated** from contact with the operator, and use a 2-core mains lead. Double insulation requires two insulating barriers, one of which may be air, each independently capable of withstanding the shrouded voltages and protecting the operator. It is possible to make a Class II appliance that has exposed metal, but rigorous testing is required to ensure that the appliance meets the full technical standard. A symbol consisting of two concentric squares signifies double insulation.

Class I appliances require only one layer of insulation from hazardous voltages, but this layer must be totally shrouded by a conductive layer **bonded** to mains earth via a low-resistance path. It is far easier to make equipment that conforms to Class I than Class II, so amateur equipment should **always** be built to the Class I standard to ensure safety.

These classes of insulation do not merely apply to the appliance, but also to the mains cable from the wall socket. Since domestic power leads cannot conveniently (and cheaply) include an earthed conductive sleeve, the cable is normally double insulated, so a single layer of insulation over each conductor is insufficient, which is one reason why three-core mains cable has a thick external sheath.

Earthing

Earthing is the cause of many problems in amateur-constructed equipment, but if thought about logically, there is no need for it to cause any problems whatsoever. Colloquially, the term “earthing” refers to the mains earth safety bond to the metal chassis, and also to the 0 V signal wiring, but the two are quite distinct.

Earth safety bonding

The three wires leaving a domestic supply are line, neutral, and earth. Neutral and earth are connected together at the substation or possibly at the electricity supply company’s cable head within the house. This means that if line contacts earth, a fault current flows, determined by the **earth loop resistance**, which is the entire resistance around the loop, including the resistance of the line wires (see [Figure 3.25](#)).

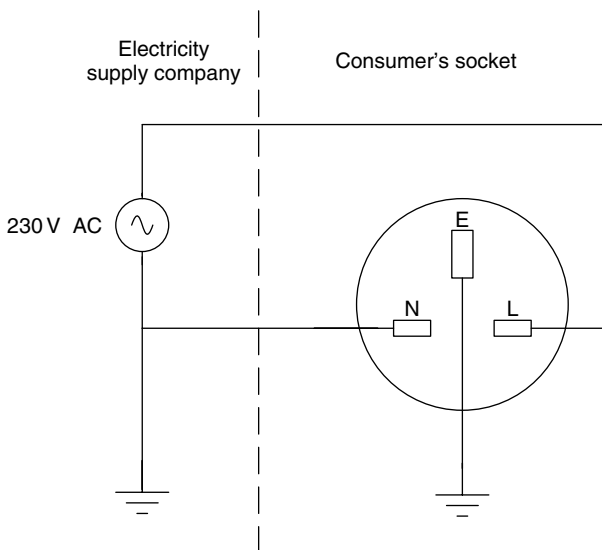


Figure 3.25

How an earthed chassis and a fuse protect against shock

The purpose of the earth safety bond is to provide a sufficiently low-resistance path to earth that if the line wire of the mains should come into contact with the exposed metalwork (which would then be a shock hazard), the resulting line to earth fault current is sufficiently great to rupture the fuse **quickly**. The time taken for the fuse to rupture is proportional to the earth loop resistance, so there is no such thing as an earth loop resistance that is too low.

Although exposed valves may appear to conform to Class II, because the electrodes are insulated by a vacuum and the glass envelope, if the envelope is broken, the secondary layer of insulation also disappears. To ensure conformity, valves on the top of the chassis should either be enclosed by a perforated metal cover to meet Class I, or insulating barriers to meet Class II.

If we build an amplifier on a chassis with exposed metal, then the construction must be to Class I and all hazardous voltages must be insulated from, and totally enclosed by, earth-bonded metalwork. The ideal place to achieve this bonding is close to the entry of the power cable. The bond should be made using a solder tag bolted to the chassis with a shakeproof (star) washer **between** the washer and the chassis because this bites into the metal of the chassis and the tag to provide a gas-tight joint. If the chassis is anodised aluminium, the surface anodising must be thoroughly scraped away underneath the tag to ensure a good bond.

The nut and bolt should be prevented from loosening using further shakeproof washers and/or locknuts. The earth bond bolt should **never** pass through plastic, such as the mains input socket, because plastic quickly creeps and causes the bond to loosen.

The ideal earth bond would weld the incoming earth wire from the mains cable directly to the chassis, but this is not very practical. A perfectly reasonable alternative takes the incoming earth wire directly to an M6 or 0BA solder tag, where it should

form a mechanical soldered joint. The tag is then screwed to the chassis with a shakeproof washer either side of the chassis, and another shakeproof washer above the earth tag, followed by a flat washer (to prevent the tag rotating when the bolt is tightened), then secured with a locknut. The nut and screw should be firmly tightened with a large screwdriver and a spanner **after** soldering otherwise the secure bond to the chassis prevents the iron from heating the tag (see [Figure 3.26](#)).

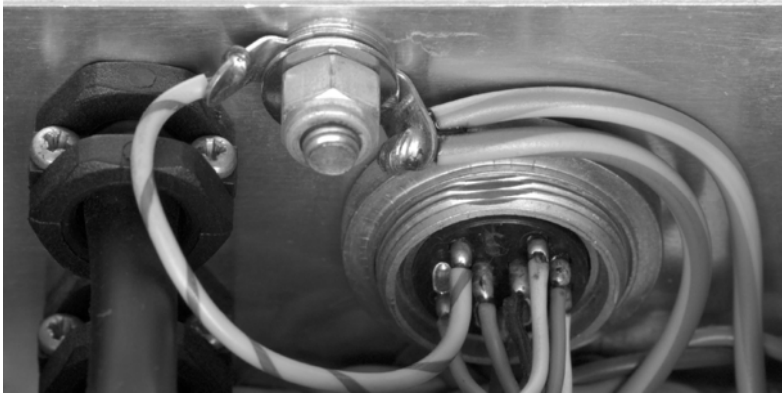


Figure 3.26

A good earth bond maintains low contact resistance for the life of the equipment

A thick cable should be used to bring mains earth to the bond point in order to reduce earth resistance. Although it is permissible for 3 A rated equipment to have $0.5\ \Omega$ of resistance from the pin of the mains plug to the chassis (**not** measured directly at the bond point), reducing this resistance to $0.1\ \Omega$, or less, by using $2.5\ \text{mm}^2$ mains cable, reduces the likelihood of hum and further improves safety.

The preceding arguments apply to equipment that is directly powered from the mains, but pre-amplifiers often have remote power supplies. Nevertheless, the same argument should still be applied, and a substantial cable should be used to transfer mains earth to the pre-amplifier chassis, and the bonding

technique should be the same. Likewise, turntables should be firmly earth bonded via their mains cable, and not via a flimsy pick-up arm lead to a possibly indifferent mains earth.

Sometimes there will be metal that could come into contact with mains voltages but does not have a guaranteed electrical path to the mains earth bond. Examples of this are:

- The acoustically suspended motor on a turntable.
- A mains transformer or HT choke that is acoustically isolated because of vibration.
- Any separate anodised aluminium panel supporting a mains connection, such as a front panel mains switch, or mains transformer on the baseplate.

Each of these should have a connection, such as a wire or a screw, via a star washer, to bond them to the main earth bond, but in the first two examples it is important that the wire should not be so stiff that it short circuits the acoustical isolation, so a loop or short helix of wire is ideal.

0 V system earthing

It is the 0 V signal earth connection to chassis that causes the hum due to hum loops between multiple earths, **not** the safety bond.

Hum loops are circuits within earth paths that can have hum currents induced into them by mains transformers. Because there is resistance in the circuit, an unwanted voltage is developed, and this causes the audible hum (see [Figure 3.27](#)).

To remove the hum, the loop must be broken, and this is often done by removing the earth wire from within the mains plug of one of the affected pieces of equipment, **but this is extremely dangerous**. The loop should be broken by removing the 0 V signal earth bond to chassis from one of the pieces of equipment.

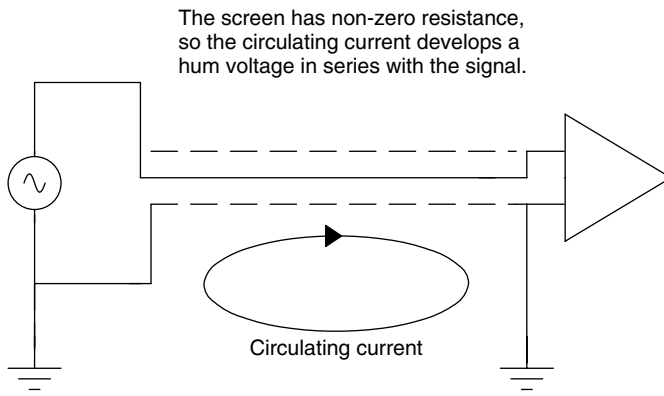


Figure 3.27

How a hum loop causes audible hum

Fortunately, most modern equipment is double insulated, so hum loops do not often occur, but a modern improvement on Class 1 equipment is to provide a **ground lift** switch or pluggable link that can make, or break, the 0 V signal earth to chassis connection at will on each piece of equipment. This method allows the optimum 0 V system earthing arrangement to be determined quickly and safely.

Breaking the 0 V signal earth to chassis bond

Switches certainly start with low-resistance contacts, but the resistance rises as the contacts tarnish, especially if they are not cleaned by periodic use, as would be the case for a 0 V signal earth bond switch, so a pluggable link is far better. The best pluggable link is the **U-link** used by telecommunications companies as the interface between their equipment and the customer's equipment in private telephone exchanges. The British U-link has a pair of linked 4 mm plugs on $\frac{1}{2}$ " centres and provides a low-resistance path that does not deteriorate (see [Figure 3.28](#)).

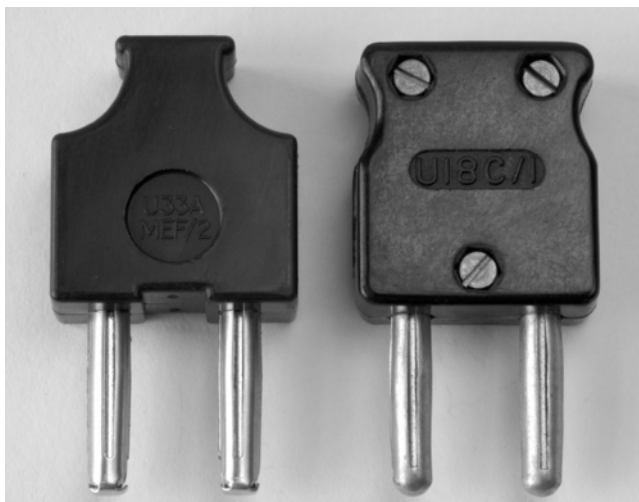


Figure 3.28

A U-link intended for telecommunications use make an ideal removable 0 V signal earth to chassis bond

The U-link plugs into two 4 mm chassis sockets, one insulated (connected to the 0 V signal earth), and the other uninsulated and screwed firmly to the chassis with shakeproof washers.

Ideally positioning the 0 V signal earth to chassis bond

We have seen that to avoid earth loops, there must be only one 0 V signal earth bond to chassis, and hence mains earth. We should now consider the optimum position for this single bond.

An amplifier amplifies the **difference** between its two input terminals. We tend not to think of the 0 V signal earth terminal as an input, but it is. In order to screen an amplifier, we surround it with a conductive casing/chassis connected to the amplifier's 0 V signal earth. Inevitably, there is capacitance from the chassis to mains. Similarly, there is capacitance from 0 V signal earth wiring to mains. Both capacitances return currents to the point where they are connected together. The

further we are from the bond, the larger the unwanted voltage drop developed across the non-zero resistance.

A moving magnet RIAA stage having 2 mV nominal sensitivity has a sensitivity of 284 μ V at 50 Hz, so it cannot tolerate any unwanted voltage drop. Thus, the optimum place to bond the 0 V signal earth to chassis is at the input of the RIAA stage.

The bond should be made with as short and thick a wire as possible, in order to reduce its inductance and ensure that it is a good bond even at RF.

Interconnects, screens, and their earths

The 0 V signal earth is a **signal** wire because it allows the signal current to return to the source. It is therefore most important that we treat it with the same care and consideration that we would apply to the more obvious signal wires.

It is worth treating signal line and earth as though they are true balanced signals, even in an unbalanced system. The author has used twisted pair/overall screen interconnect cables on unbalanced systems since 1976 because of their superior rejection of external fields, but it is important that the screen is connected to 0 V signal earth at the source end.

The screen should be earthed at the source end because the source has a low output impedance and can firmly define its output voltage as the difference between the two output wires. One of these wires is connected to the screen of the output cable. The screen picks up RF interference which it superimposes onto the commoned signal wire. The RF is also superimposed onto the other signal wire via the output impedance of the source, and if the source has a truly zero output impedance at RF, both output wires now have the full RF superimposed on them. This might seem undesirable, but . . .

At the amplifier end, the input stage responds to the **difference** in signal between the two input wires, and therefore rejects the RF that is identical on both input wires.

If instead we connect the cable so that the screen is commoned to the earthy side of the signal at the destination end, the induced RF picked up by the screen now has to travel down the entire length of the cable to the source before it can be coupled via the output impedance of the source to an inner wire. Once coupled, the RF then has to travel the entire length of the cable before it can arrive at the amplifier input.

The RF signal on one wire has now had to travel twice the length of the cable, but has suffered the effects of series inductance and shunt capacitance which form a low-pass filter. This means that one of the wires at the input to the amplifier has the full RF signal, and the other has an attenuated signal, resulting in a difference signal to which the amplifier is sensitive.

The previous argument assumes equal impedances to RF earth (whatever that might be) at the pre-amplifier input. Whilst this is not necessarily true for domestic equipment, the quasi-balanced method of interconnecting cables described is a considerable improvement on coaxial cable or “screened lead” because domestic impedances cause the two wires (the inner and screen) to have wildly different signals induced into them and there is no possibility of common mode rejection at the receiving amplifier.

Internal earth wiring of amplifiers

Once within an amplifier, the 0 V signal earth path can either travel in the same way as it did between equipment, in which case, it is known as **earth follows signal** [1] or it can be **star earthed**.

Earth follows signal is the traditional method of wiring valve amplifiers, and is the easiest to do properly. The traditional method uses an **earth bus bar** which is a thick (at least 1.6 mm diameter, or 16 swg) tinned copper wire 0 V signal earth connected directly from input sockets to the 0 V signal earth, or source, of the power supply (commonly the reservoir capacitor, or if fitted, the 0 V output of the regulator).

The repetition of “0 V signal earth” may seem pedantic, but it reminds us that the signal earth also carries power supply currents. This last factor is extremely important, since the power currents are often many times larger than their associated signal currents. Because individual stages amplify voltage differences referenced to their local earth, the earth bus bar needs low resistance to minimise the spurious voltages developed by power supply or signal currents passing through it.

Even the precaution of having a low-resistance bus bar is not sufficient, so connections must be made to the bus bar in the correct order such that voltages are not developed in sensitive input circuitry. The author remembers an RIAA disc pre-amplifier that had been constructed to the Mullard 2-valve design, but had considerable hum. The hum was cured by moving **one** wire 150 mm along the (1.6 mm) earth bus bar.

The correct order from the input socket is input circuitry (such as shunt capacitors, etc.), grid leak resistor, cathode bypass capacitor (if fitted), cathode resistor, any anode signal circuitry (such as equalisation), next valve’s grid leak resistor, etc. If in doubt, think **current** rather than voltage. Decide whether you would be happy for a particular current to develop a voltage at a particular point along the bus bar and whether that voltage would then be amplified.

Depending on earthing arrangements, the earth bus bar may need to be firmly bonded to the chassis via a solder tag. Traditionally, one of the screws retaining the input valve socket was

used as an earth bond, but all of the rules regarding mains safety earths still apply, and a few milliohms of contact resistance is sufficient to cause an irritating hum, so a large screw that can be tightened down firmly with proper shake-proof washers is best. Centre spigots of the valveholders (if fitted) should also be bonded to the bus bar as they help reduce capacitance between pins of the valve socket.

To make a neat earth bus bar, the thick wire needs to be straightened, and this is not a trivial task. The traditional way to achieve this is as follows:

Grip one end of the wire in a vice, and then grip the other end in a substantial pair of pliers, such as would be used for working on a car. The wire is then wrapped one turn round the jaws of the pliers, and the pliers are firmly gripped with both hands whilst one foot is braced against the vice. The wire is then firmly pulled until it can be felt to stretch, and **without moving the position of the pliers** is cut at the vice end. A beautifully straight piece of wire is now the result and this can now be cut away from the pliers.

It should be realised that considerable force is required to achieve this result, and this can apply dangerous forces to your back if you do not position yourself correctly. If you are in any doubt as to how to position yourself, or have back problems, **do not attempt to use this method**. It is far better to have tatty wiring that can be seen, than beautiful wiring that cannot be seen because you are lying down with your back in traction.

Unfortunately, tinned copper wire oxidises over time, and becomes difficult if not impossible to solder, but a beautifully shiny and easily soldered surface can be restored by a quick polish using metal polish intended for brass. A wipe with isopropyl alcohol or methylated spirits will remove the last traces of polish and enable perfect soldered joints.

The ultimate expression of the earth bus bar is the RF earth plane, which is a two-dimensional conducting earthed surface to which earth connections are made (a wire is only one-dimensional). This construction is now common on audioprinted circuit boards as designers have realised how important RF immunity is to audio circuitry. On a PCB, the entire upper surface of the board can be used as an earth plane, which has low inductance because it is so wide, and therefore guarantees a good RF earth at every point that a contact is made.

An alternative hard-wired approach is to use a strip of $10\text{ mm} \times 1\text{ mm}$ silver or copper as the bus bar. The very large cross-sectional area ensures low resistance, whilst the width lowers the inductance and makes soldering fractionally easier. Holes can be drilled into the bus bar to suit component wire diameter, or V notches filed at the edge to locate the component whilst soldering. The iron should be held to the strip first to bring it up to temperature, then moved so that it touches the component, and more solder applied to form the joint. A 200 W iron is needed, as a smaller iron takes so long to warm the bus bar to soldering temperature that it damages adjacent components.

Some classic valve amplifiers (such as the Rogers Cadet III) approximated to a earth plane by soldering to conductive posts pressed directly into the (steel) chassis, and using this as the 0 V signal earth. Neither this method nor a steel chassis is recommended for a new design.

Star earthing is achieved by having a **single** earth point, perhaps bonded directly to chassis, to which all 0 V connections are brought individually. Ideally, all the connections to this point are made with short leads to minimise inductance and noise, but building an entire amplifier using this method can become a little fiddly (see [Figure 3.29](#)).

Because of the difficulty of making so many 0 V connections to a single tag, many constructions use a combination of star and

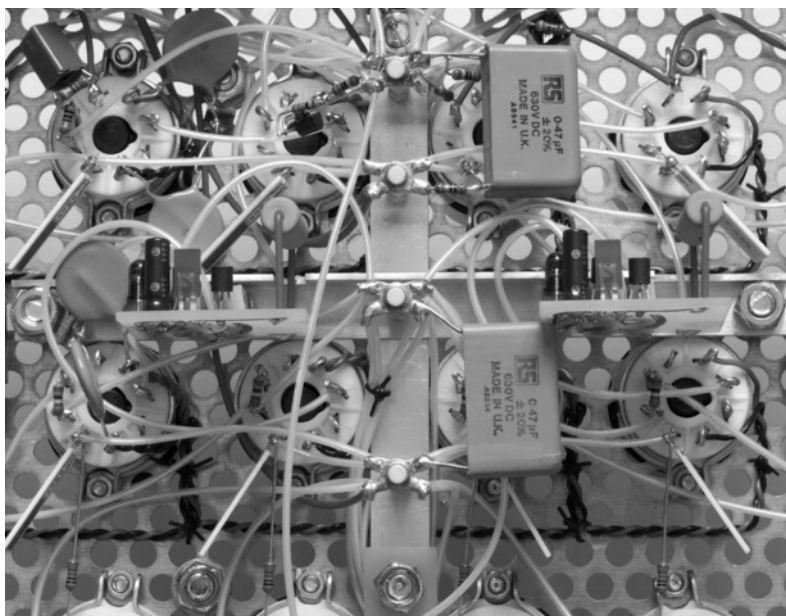


Figure 3.29

True star wiring uses many wires, but is manageable with care. Stars from top to bottom; 0 V signal earth, +270 V, +160 V, -335 V

bus bar earthing, with the input stage star earthed, and following stages earthed to the bus bar.

Rectifiers and high-current circuitry

Although low-voltage bridge rectifiers are commonly available, they typically use standard diodes rather than Schottky or soft-recovery, so we often need to make our own bridge rectifier and if we choose diodes that are in an insulated package (such as STTA512F, rather than STTA512D) fitting them to an aluminium angle bracket (to allow heatsinking) is easy (see [Figure 3.30](#)).

Some parts of an amplifier unavoidably carry high currents. In a capacitor input power supply, the loop from the mains



Figure 3.30

Insulated diodes make bridge rectifier construction easy

transformer via the rectifier to the reservoir capacitor carries the capacitor ripple current, so it is essential that no connections are made to the 0 V signal earth **within** this loop because the large current would develop a noise voltage across the unavoidable wiring resistance. The two circuits are split at the reservoir capacitor. Assuming that the reservoir capacitor has screw terminals, it is best to connect the solder tag directly contacting the capacitor to go to the rectifier/transformer and the one further from it to go to the load (This means that the high currents in the capacitor/rectifier/transformer loop cannot develop a voltage in the load circuitry.) (see [Figure 3.31](#)).

In order to minimise the electromagnetic field caused by the passage of transformer to reservoir capacitor currents, this



Figure 3.31

Star connections to a reservoir capacitor reduces interaction between incoming and outgoing currents

loop should ideally be made of twisted pair to minimise its area, and be as short and thick as possible to reduce its resistance and voltage drop.

A power amplifier's output stage is effectively supplied from a single capacitor, so output valve cathodes should be brought individually to the 0 V terminal to form a star. Similarly, individual wires from output transformers should be brought individually to the HT terminal to form a star (see [Figure 3.32](#)).

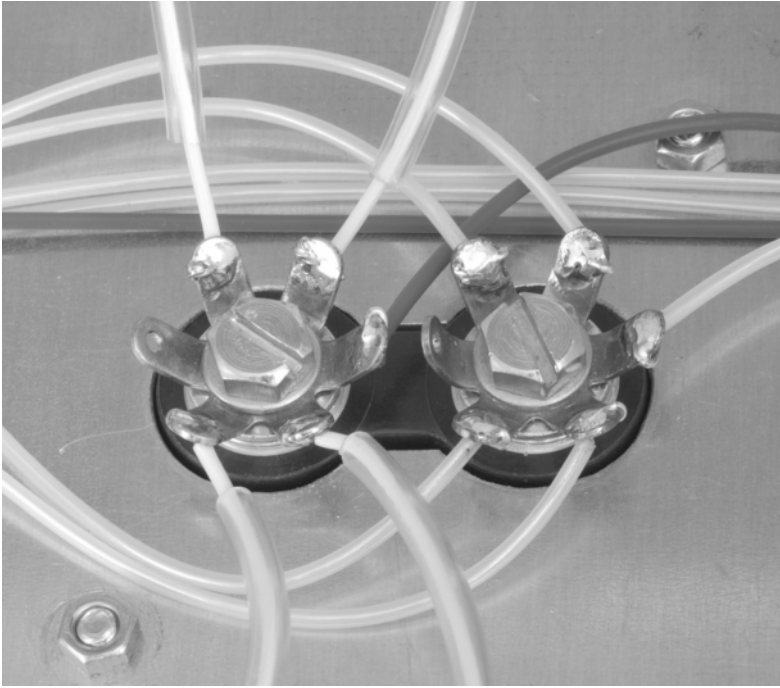


Figure 3.32

Star connections to the capacitor feeding the output stage minimise interactions between valves

The output to the loudspeaker from a power amplifier is also a high-current loop, and additional connections to sense this voltage, such as global negative feedback, should be made very carefully. The ideal method is to connect a screened twisted pair to the output terminals, with the screen connected to chassis at one end only (to prevent loop currents flowing in the 0V signal earth). This cable is then routed to the input stage, where one side is connected to the lower end of the cathode resistor, and the other is connected via a series resistor to the cathode (assuming cathode feedback). This method ensures that the feedback voltage is derived from the correct point and presented to the correct point (see [Figure 3.33](#)).

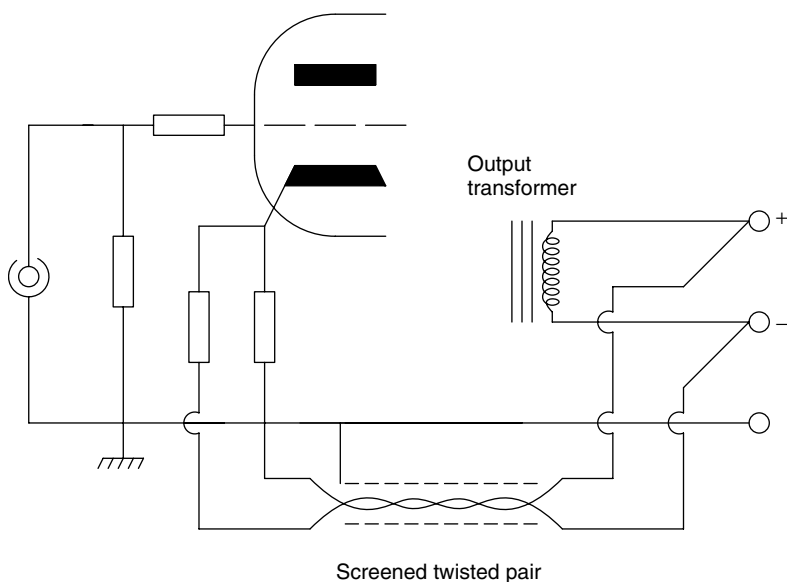


Figure 3.33

Screened twisted pair derives the feedback from the correct point (amplifier output terminals) and prevents induction from transformers entering the feedback loop, then applies the feedback at the correct point

Layout of components

There are various approaches for positioning the smaller components such as resistors and capacitors, ranging from one extreme to the other:

1. Position the components in a regimented manner in neat lines and use neatly harnessed wiring to make the complex interconnections. The parallel wires in the harness inevitably increase stray capacitances, but this method was used in the Quad II power amplifier (see [Figure 3.34](#)).
2. Position the components in a regimented manner in neat lines, and either use hidden wiring which can be untidy (Leak), or exposed neat wiring (Mullard recommendations) to make the interconnections between components. This

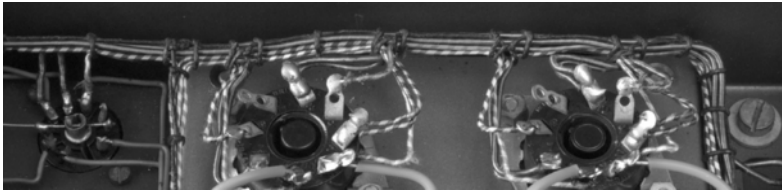


Figure 3.34

The Quad II harnessed its wiring into a loom, increasing capacitance

technique reduces capacitances compared to method (1), but can take up significant space (see [Figures 3.35](#) and [3.36](#)).

3. Methods (1) and (2) had regimented rows of components enforced upon them by the layout of the tags on the boards. If that restriction is waived, then a small resistor no longer takes up the same amount of space as a large coupling capacitor and space is saved. The hidden wiring becomes copper tracks and the board becomes a printed circuit board (PCB).

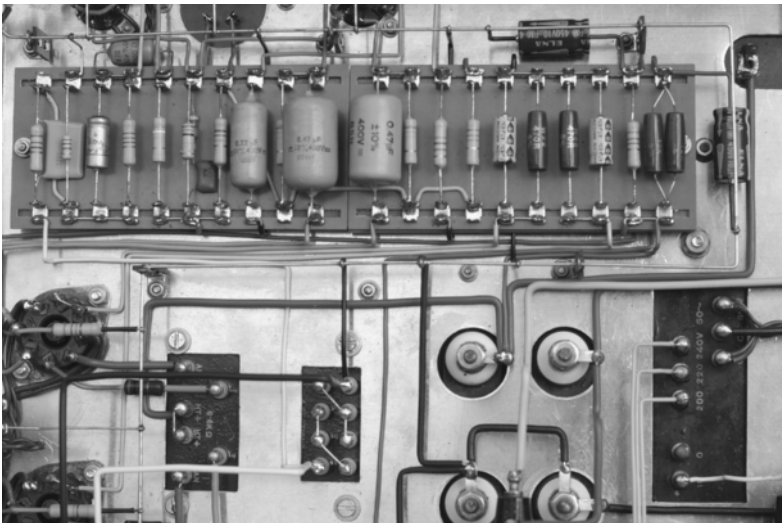


Figure 3.35

This fine example of a Mullard 5-20 was spotted at a radio fair. Ideally, the 0 V signal earth bus bar would contact the chassis at a single point, rather than at each support. Courtesy of Ian Lavender

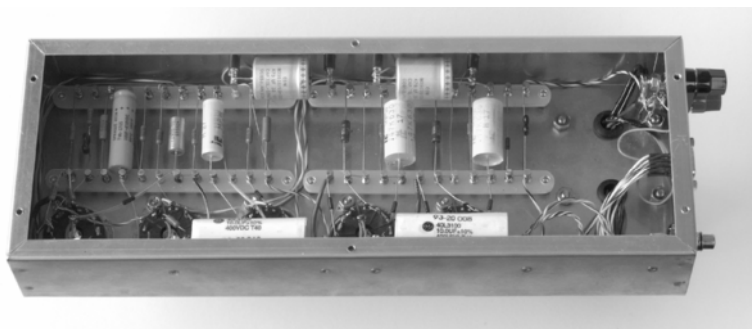


Figure 3.36

This 10 W amplifier is one of a pair used in an “old technology” TV monitor project and receives power from the central power unit rather than having its own supply. Courtesy of Brian Terrell

4. The preceding methods are two-dimensional. Interconnecting wire length can be further reduced if the assembly becomes three-dimensional. Components are now soldered directly to valve sockets and if necessary, to nearby stand-offs, or perhaps a tagstrip. This is potentially the best construction method, but it needs careful thought to keep the components even reasonably tidy and accessible for testing/maintenance. Tagboards, tagstrips, and stand-offs are all useful ways of supporting components and joints insulated from earth (see [Figure 3.37](#)).

At valve impedances, practical inductances are completely insignificant except in the 0 V bus bar, so good layout in a valve amplifier requires low stray capacitances, and the best way of reducing capacitance between wires or components is to cross them at right angles. Capacitance to earth is minimised by short leads. One approach that is not immediately obvious is that sleeved wires have higher stray capacitance than self-supporting bare wires because $\epsilon_r \geq 1$. It is the combination of these requirements that forces component layouts that appear to have been thrown together by an avant-garde sculptor. We thus arrive at the surprising conclusion that a good component layout probably looks untidy, but the converse is unlikely to be true (see [Figure 3.38](#)).

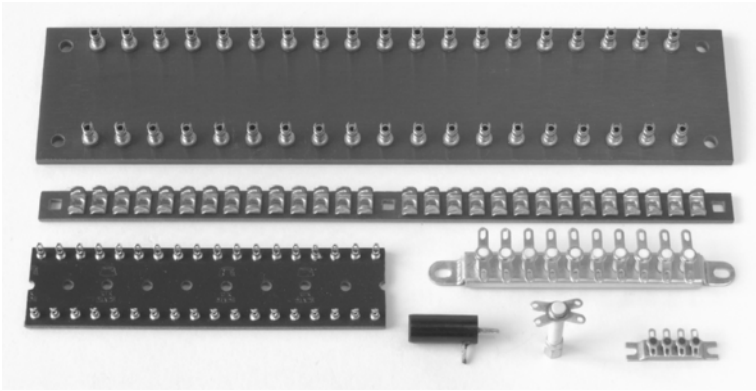


Figure 3.37

A selection of tagstrips, tagboards, and stand-offs

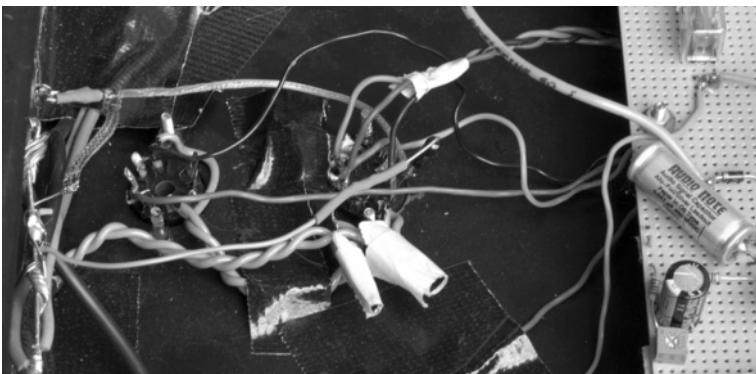


Figure 3.38

Bad practice: How many howlers can you spot?

This appalling example features a variety of bad practices:

- Gaffer tape: Gaffer tape is for bodging on stages and studio floors, not internal electronic use.
- Terrible soldering: As a single example, the thickest wire on the earth bus bar was clearly soldered by an iron having an insufficiently large tip, resulting in poor solder flow and damaged insulation.
- Earth bus bar: The earth bus bar should carry the 0 V signal earth from the valve to the input socket, not the screened

lead. Thus, the bus bar is heading in the wrong direction – it should pass over the valve sockets.

- Screws: The screws securing the valve sockets are too long. It's very difficult to install good heater wiring when your fingers keep hitting obstacles.
- Heater wiring: The wire is too thick and stranded, and because the wire wasn't pre-tensioned before twisting, these factors resulted in a loose twist just where a tight twist is most important (near small-signal valves).
- Diagonal components on matrix board: Diagonal components do not themselves cause problems, but they **do** indicate a poorly planned layout.

Whether the layout is a “tidy” PCB or an optimised hardwired layout, good layout requires considerable care, and thermal considerations must also be accommodated.

Making PCBs

Printed circuit boards have been mentioned several times, and the author uses them frequently, but they are not ideal for the novice. This is because they are really a production method of construction, and it requires considerable confidence to design a theoretical circuit and commit it directly to a PCB. Whilst it cannot be argued that a good PCB gives a thoroughly professional appearance to the finished project, hardwiring will often give superior performance. Despite these caveats, making a single-sided PCB is not nearly as difficult as you might think, and it doesn't need specialist equipment.

PCB layout

The hardest part is working out the layout. Unfortunately, achieving a good layout is very much like learning to ride a bicycle – easily done once you have experience, but tricky to describe to a novice.

Nevertheless, points to remember are:

- If the circuit diagram was drawn well (connections in the right order), a good PCB will be laid out very much like the diagram.
- Most components (except valve sockets) fit a 0.1" grid.
- You are working from underneath (the foil side), so if you use ICs, remember to work out their pin-out as viewed from underneath. (IC manufacturers' diagrams give the pin-out view from above, whereas valve manufacturers give the view from underneath.)
- A simple audio PCB should not require any links. If you need links, you probably haven't tried hard enough. (Or you used an autorouter in a PCB design program.)
- PCB foil is thin, so it is vital to think about currents and ensure that you don't pass heavy currents through sensitive areas. Remember that you can always widen the track, or star tracks to a central point to minimise interaction between currents (see [Figure 3.39](#)).



Figure 3.39

Star earth and power on a PCB

- PCBs inhibit cooling, so try to put hot resistors or transistors near the edge where they can be bolted to a heatsink or can cool naturally. If you can't do that, raise them well clear of the board and give them space from other components to allow cooling air to flow round them.
- Don't put power valves on a PCB – it's just asking for thermal problems.
- If you put valves on a PCB, don't even attempt to do heater wiring using PCB tracks, it makes layout of the audio almost impossible, and heater wiring is far better done afterwards with twisted pair spaced away from the board so that it rests snugly against the chassis (see [Figure 3.40](#)).

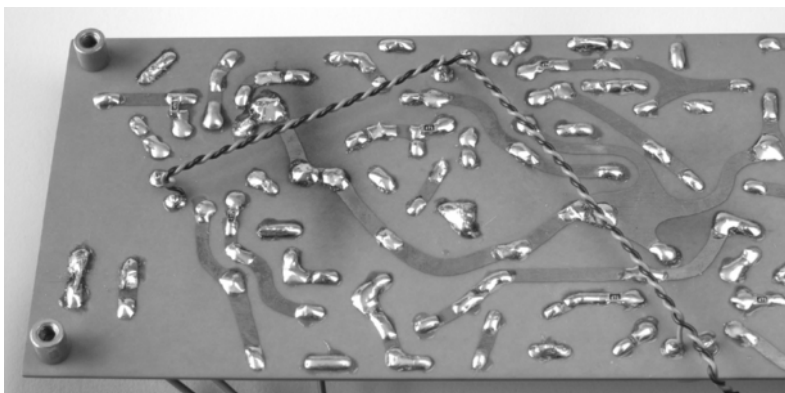


Figure 3.40

Heater wiring on a PCB is best achieved with twisted pair

- Surface mount resistors make excellent grid-stoppers, and really aren't that difficult to use, but in order that surface tension aids positioning, rather than hindering, it's important that your PCB tracks are the same width as the component. You need the smallest tip on your iron and ≤ 0.5 mm silver-loaded solder. Tin the tips of the two tracks with the absolute minimum of solder, position the resistor with tweezers, and with a slightly wetted iron, reflow the solder at one end to form a badly soldered joint. Solder the other end properly with a little fresh solder, then return to the first

joint and solder it properly with a little fresh solder. It's actually more difficult to describe the process than to do it!

- Pre-amplifiers tend to need a lot of connections to earth, so a double-sided PCB with the component side used as an earth plane not only helps layout, but improves screening. If you use an earth plane, radial-leaded capacitors such as Wima FKP1 obscure their leads on the component side, so you need to bend their leads horizontally (without stressing the component) and solder directly to the earth plane. They don't, therefore, need a hole to be drilled through the board for this connection.
- A power amplifier's driver stage usually has its valves showing on the top of the chassis, but PCB capacitors can sometimes be quite tall, forcing the PCB to be mounted so far below the chassis that the valve hardly peeps through the chassis, restricting cooling. In this instance, it is better to mount the valve sockets on the foil side (beware that from the point of view of the tracks, this reverses the pin order). Unfortunately, mounting the socket on the foil side is mechanically weak because removal of a valve tends to tear tracks off the PCB. The solution is to choose a flanged PCB socket and use PCB mounting spacers of just the right length to press the flange against the chassis plate (see [Figure 3.41](#)).
- If you are forced to have two tracks adjacent to one another that are hostile to each other (perhaps the input and output of an amplifier), you can always put an earthed track in between to guard them.

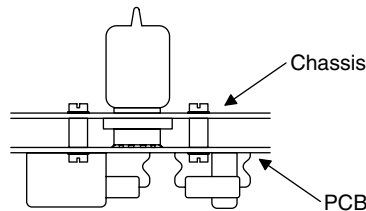


Figure 3.41

The chassis can help in supporting valve sockets soldered to the foil side of a PCB

There are two ways of working out the layout:

Method 1: You buy a pad of graph paper marked in 0.1" squares. This is usually described as "10ths, $\frac{1}{2}$ & 1 inch". You then sit down with a 2H pencil and a plastic eraser and work out your layout. After much rubbing out and redrawing, you obtain an efficient layout, which you then copy onto a virgin sheet of graph paper.

Method 2: You use a familiar computer drawing program with "snap" set to 0.1". After much shuffling and redrawing, you end up with an efficient layout and you struggle to make your printer print it the correct size. This process initially takes longer than Method 1, but it allows changes and copies later on.

Transferring the layout to the PCB

Either way, you now have a drawing the same size as your final PCB. Next, you cut the PCB material to size and clean the edges. Whether you use glass reinforced plastic (GRP) commonly known as FR4 (fire resistant type 4), or synthetic resin bonded paper (SRBP), the best way of tidying the board edges is to rub them down with 160 emery cloth held on a sanding block (for some reason, traditional cork sanding blocks are now expensive, but a scrap of MDF is fine). For a perfect finish, the edges can be smoothed with 300 grade silicon carbide paper. The inevitable burrs are best removed with the flat side of a half-round needle file (half-round needle files tend to be finer than flat ones).

Some constructors, having used a computer to produce their layout, print the layout onto a transparency. Ensure that the transparencies you use in a laser printer are suitable for photocopiers and/or laser printers, otherwise you will wreck your printer! You then buy some fresh PCB material coated with UV sensitive etch resist, put the PCB and transparency into a UV light box and take a contact photograph of the layout

which is later developed. Note that the drawing must be “mirrored” before printing so that the printed side is in contact with the board. If the printed side is not in contact with the board, the image on the board becomes defocussed, leading to a poor PCB. Because this is a photographic process, exposure time is important, as is the density of the black that your printer can produce on the transparency. Additionally, film deteriorates with time, which is why you need fresh PCB material . . .

Alternatively, you could buy UV etch resist and spray the board yourself. If you buy the etch resist, carefully check the “use by” date and reject it if there is less than a year remaining. Treat etch resist, developer, and unused (coated) board like 35 mm film, and store them in a fridge to maximise their shelf life.

Correctly made photographically produced PCBs have a perfect finished appearance, but severely restrict your choice of PCB material unless you are prepared to coat boards yourself. Additionally, UV light boxes are not cheap, and practice is required to achieve correct exposure. Because of these photographic problems, the author still uses a very low technology method . . .

Cut the board of your choice to size and wash/scrub the foil with a soap-filled scouring pad until the copper gleams. Stick the paper layout gently to the foil side. Use a scribe to mark all the component and fixing holes through the paper (see [Figure 3.42](#)).

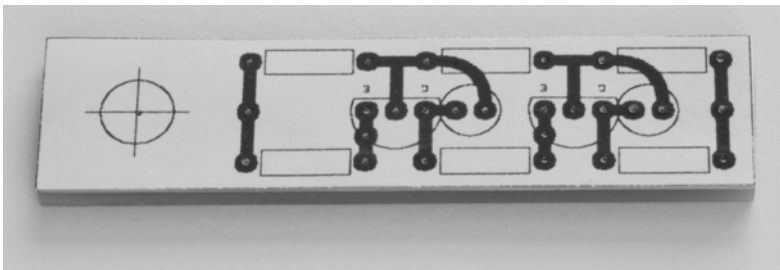


Figure 3.42

Marking “join the dots” holes through the paper template

Carefully remove the paper. Wash the board again under very hot water so that no trace of glue remains. Wash your hands thoroughly. Dry the board, and handle it by the edges only. You now have a board covered in dots. Sit down at a well-lit table, put some Bach on (greatly helps concentration), take a 00 sable paint brush, a pot of well-shaken enamel paint, and “join the dots”. It is not necessary to be a painter of the calibre of Michelangelo, just to have a steady enough hand to paint smooth curves that do not touch. Preparing the board for painting tenses the muscles, and you will find that you produce better work if the board is prepared the day before, so that you have a relaxed, steady hand. It is most important not to touch the foil, as grease/sweat could impede the subsequent etching process, so an artist’s Mahl stick to support your wrist can be very useful when painting large boards (see [Figure 3.43](#)).



Figure 3.43

This easily-made Mahl stick prevents accidental contact with the PCB

The author uses red paint because it is easy to see errors against the copper backing (see [Figure 3.44](#)).

If you are impatient, once painted, you will arrange for a gentle draught of warm air over the board to dry the paint quickly.



Figure 3.44

The PCB is painted and ready for etching

Whilst this is happening, wash your brush carefully in white spirit followed by warm soapy water.

Etching

Printed circuit boards are etched using ferric chloride solution. Ferric chloride granules are available from all electronics factors – just add water. Bear in mind that it is intended for etching, so it is highly corrosive, and should be stored in a PVC container. Fruit juice bottles are a possibility, but label the container conspicuously and keep it away from inquiring hands.

Assuming that you have a board covered with developed etch resist or dry paint, it is time for etching. Warm etching fluid

works far faster than cold (chemical reactions double in speed with each 10 °C rise in temperature). The author fills his PVC etching tray (an old lunch box) to a depth of $\frac{1}{4}$ " (≈ 6 mm) and discretely pops it in the microwave for a couple of minutes to warm it gently before etching. Have plenty of kitchen roll handy at all times in case you accidentally spill a drop.

Ease the board gently into the etching fluid foil side up, and rock the tray gently and continuously to clear etched debris from the surface of the board. After a few minutes you will see the board begin to etch around the edges. Continue agitating the fluid and watching. Depending on temperature, foil thickness, and age of fluid, the process takes between ten and twenty minutes. It is most important not to leave the board in the etchant for too long or it will undercut the tracks beneath the resist. When etching appears to be just complete, give the board another minute to be certain. Place your bottle of etching fluid in the sink, turn the cold tap on, and pour the fluid back from the tray into the bottle. If you should spill any etchant, it will immediately be diluted by the running water and the sink will not be stained. Put the bottle of etchant away, and run water gently into the etching tray to dilute the remaining etchant and allow it to overflow until clean water results. You can now retrieve the board (see [Figure 3.45](#)).

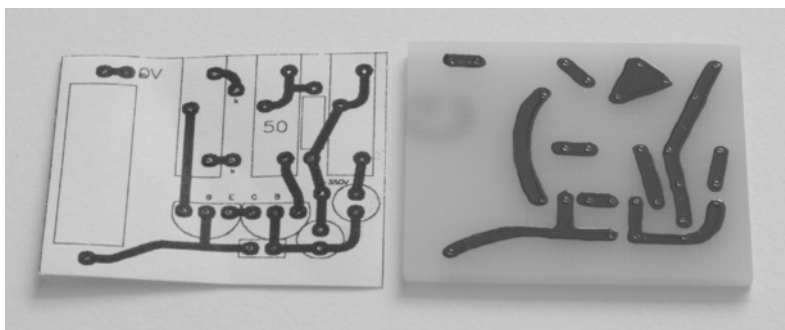


Figure 3.45
PCB etched

The etch resist can now be removed. If you painted the board, pour chemical paint stripper generously onto the board, wait a minute or two, and wash it (and the paint) off with an old toothbrush under running cold water. Don't press hard with the brush or you will scrub paint into the board never to be removed. By now, your hands are probably becoming greasy again, so wash them thoroughly. The board is now ready to be drilled (see [Figure 3.46](#)).

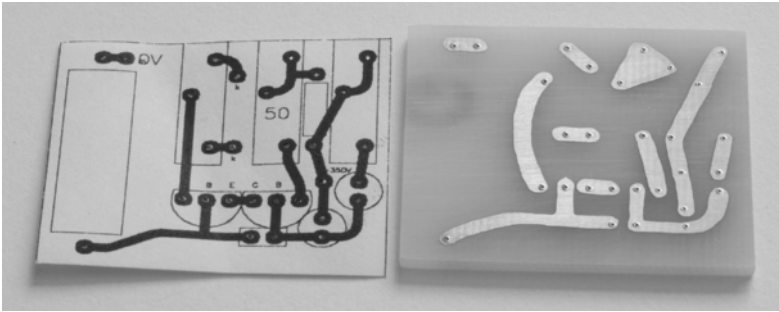


Figure 3.46
PCB stripped

Drilling

If you use FR4, you need silicon carbide drills. Ordinary high speed steel (HSS) drills **do** work, but they blunt quickly. Typical hole size is 0.8mm, but you may wish to measure components individually and drill holes precisely sized for each component to ensure the best possible solder joint. Small drills must be treated carefully and ideally need very high drilling speeds. A drill in a stand is essential. Provided that the stand is well greased and adjusted, even a power drill can be used to drill <1 mm holes, but it will probably need a pin chuck to hold the drills (see [Figure 3.47](#)).

Stuffing and soldering

Once drilled, the board can be stuffed. Component leads should be cropped short **before** soldering, otherwise the mechanical shock from cropping disturbs the soldered joint (see [Figure 3.48](#)).



Figure 3.47

A pin chuck allows very small drills (<1.5 mm) to be used

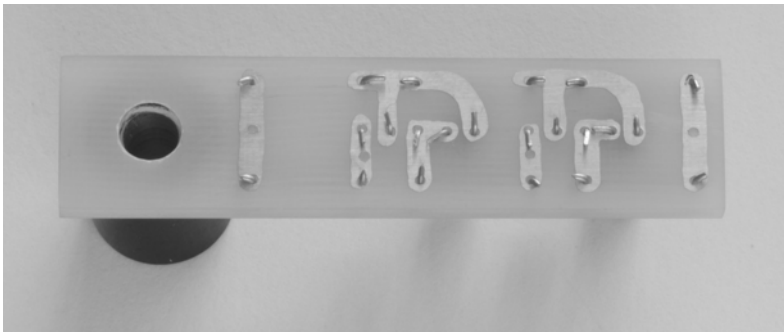


Figure 3.48

The board has been stuffed, leads cropped, and is ready for soldering

Once soldered, it is well worth removing the flux from a PCB. The act of soldering causes a small explosion as the flux heats, causing droplets of solder and flux to be sprayed over the surrounding area. On a PCB, this can cause electrical leakage, so although flux removing fluid is messy stuff, it ensures correct operation of the completed board (see [Figure 3.49](#)).

Although a board is a sub-assembly, there is no reason why it should not form part of a larger sub-assembly. In this parti-

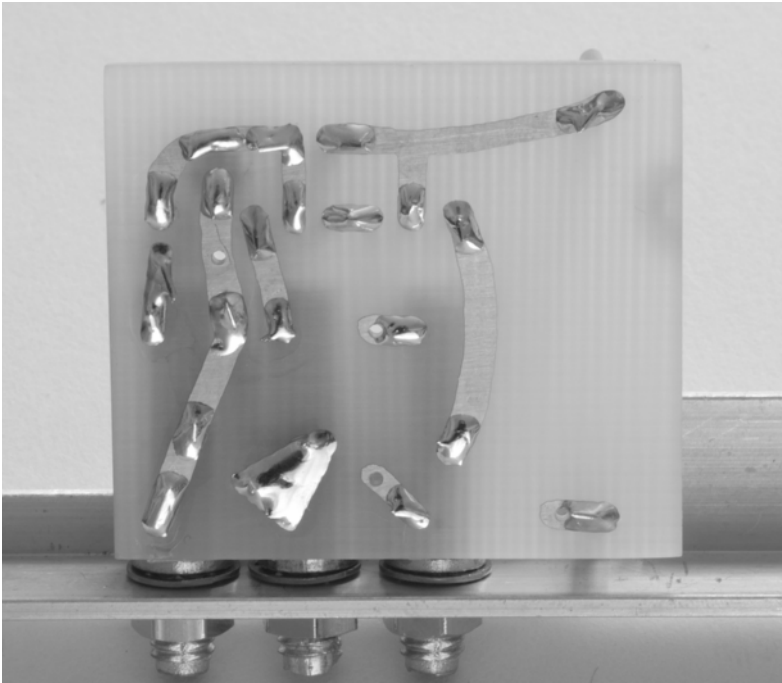


Figure 3.49

The joints have been soldered and defluxed

cular case, each board is a constant current sink that ultimately controls the total anode current and balance of a pair of output valves in a push–pull stereo power amplifier (see [Figure 3.50](#)).

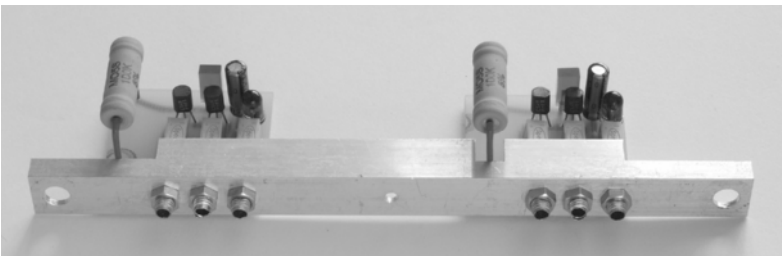


Figure 3.50

A pair of boards on a bracket form a larger bias sub-assembly

The aluminium bracket supports a pair of PCBs by their trimmer potentiometers, which have been positioned so that their shafts align with the holes in the perforated aluminium sheet to which this bias assembly will be fitted.

Modifying PCBs

Although it might seem that a circuit constructed on a PCB is set in stone, modification is perfectly possible – it just takes care.

Phenolic boards are far more fragile than glass-reinforced plastic boards (FR4, etc.), so decide carefully how you will remove a component before attacking the board.

Never attempt to salvage components. If you bent the wire over before soldering, the component will not come out simply by heating the solder and pulling from the component side. Cut the component away from the component side and use a sharp point such as a needle or a dedicated PCB rework tool to push each lead through whilst heating the solder on the track side. Once a wire fragment stands proud, it can easily be removed by a pair of fine nose pliers or tweezers. Never tug on wire fragments – you could tear a track away from the board. Sometimes it helps to remove some of the solder from the track side before removing a wire fragment. Paradoxically, you don't want to remove all the solder because this makes it difficult to get the heat into the remaining solder that bonds the unwanted component to the track. Desolder braid is a safer option than a desolder gun – the recoil can kick tracks away.

The replacement component may not be the same length as the original. Don't attempt to force the wires on the new component to fit the old holes, use one old hole and drill a new hole to suit (make sure that you don't accidentally drill into another track!) Pass the wires through, bend the wire from the blank hole so that it contacts the track you want, and solder it.

Sometimes you need to cut and remove a section of track. A sharp scalpel easily cuts the thin copper foil, so you don't need to murder the board. Make a cut either side of the section of track to be removed, tin the track, heat it with the iron, and use the tip of the scalpel to ease the unwanted track away. If you are careful, it is possible to remove sections of track without anyone knowing you've been there.

Sometimes you need to add a track from one side of the board to the other. The standard industry technique is to use green (any colour would do, but green is used) wirewrap wire to make the link. Cut the wire to length, and strip $\frac{1}{2}$ " from each end. Tin each end to allow the stress of stripping to be relieved and shrink back the insulation, and trim the exposed wire to 1–2 mm as appropriate for the job, then solder. If the wire is long, a few small spots of glue can be used to secure it.

If you need to add some circuitry and there's a blank piece of board, it's surprising how neat a job can be achieved by drilling holes for the components and bending and soldering their wires together on the track side in lieu of tracks.

Hardwiring modules

A high-voltage bridge rectifier with each arm composed of three diodes and snubber capacitors can be rather bulky if hardwired on tagstrip or terminal board, and making it on a PCB is inviting **tracking**. A method that can be useful is to glue the capacitors together using a hot melt glue gun, then fit the diodes in appropriate positions. It can be quite tricky to determine how to position the diodes to minimise the wiring, but starting from the circuit diagram of the rectifier and gradually metamorphosing it to the required shape helps (see [Figure 3.51](#)).

Once the layout has been decided, the compact rectifier can be neatly constructed, then potted (see [Figure 3.52](#)).

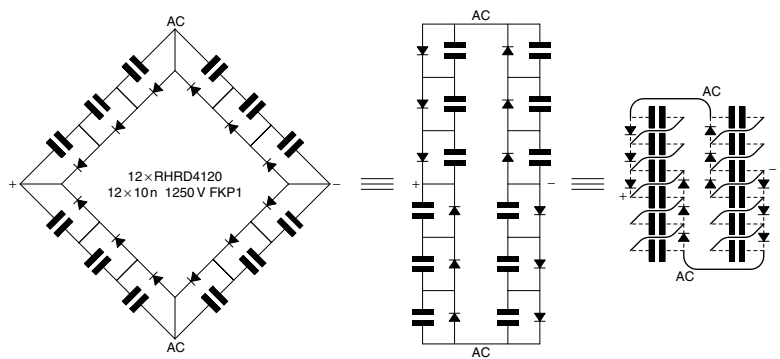


Figure 3.51
A compact mechanical layout for an HT bridge rectifier can be gradually developed from the circuit diagram



Figure 3.52
Once laid out, the rectifier can be constructed quickly and neatly, then potted

Problems and solutions

The following section covers techniques that you may find useful but that don't sit neatly in previous sections.

Identifying the outer foil of a capacitor

Electrostatic noise pick-up is proportional to the capacitance between the noise source and the affected part, so large objects like coupling capacitors are especially vulnerable. Because capacitors are invariably wound as a coil of foil, the outer foil screens the inner foil, and should ideally be connected to earth. Obviously, a capacitor coupling two stages together has neither end connected to earth, but the output of the preceding stage is a far closer approximation to earth than the grid of the next stage. Thus, coupling capacitors should have their outer foil connected to the output of the previous stage to minimise noise pick-up.

Unfortunately, the outer foil is not often marked, but it can usually be identified (see [Figure 3.53](#)).

The capacitor under test is connected as a reservoir capacitor to crudely rectified AC, and a strip of metal foil is wrapped

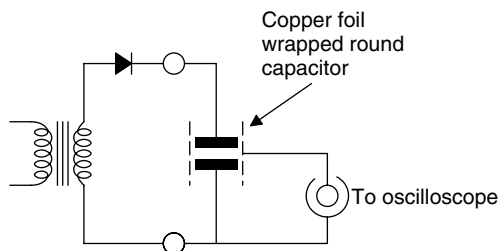


Figure 3.53
Identifying the outer foil of a capacitor

around the capacitor. The metal foil is connected to the input of an oscilloscope or an amplifier. If the outer foil of the reservoir capacitor is connected to earth, very little noise will be picked up, but if it is connected to the unearthed output of the rectifier, the capacitance between it and the added foil will easily couple the higher harmonics of the rectified AC into the amplifier or oscilloscope. Thus, the outer foil can be identified by finding which connection creates the most noise. Unfortunately, this method doesn't work for polypropylene capacitors because they are often made from two capacitors in series placed end to end, so they don't strictly have outer and inner foils.

High-current (>2A) heater regulators that shut-down at switch-on

Single-ended amplifiers cannot cancel heater-induced hum from directly heated valves in their output transformer, so one solution is to power the heater from regulated DC. The 5 A LM338 is the ideal three-terminal regulator for the job, but its comprehensive protection can sometimes be tripped by the (much lower) cold resistance of a valve heater. If this occurs, the solution is to add a resistor between input and output of the regulator to bleed current directly into the heater. At switch-on, the regulator shuts down, but the bleed resistor passes sufficient current to warm the heater, raising its resistance, eventually allowing the regulator to come out of shut-down.

The value of the bleed resistor is not critical, and setting it to pass 10% of the final current generally solves the problem. Thus, if we had a 10 V 3.25 A heater and dropped 3 V across the LM338 under load:

$$R = \frac{V}{0.1 \times I} = \frac{3 \text{ V}}{0.1 \times 3.25 \text{ A}} = 9.23 \Omega \approx 10 \Omega$$

Under working conditions, the $10\ \Omega$ resistor will only dissipate $0.9\ \text{W}$, but at switch-on, it must pass a much higher current, so an aluminium-clad resistor bolted directly to the chassis is ideal.

Tarnished silver-plated stand-offs that won't solder

Silver tarnishes when exposed to air, so 30-year old silver-plated stand-offs are unlikely to be solderable. However, dipping the tag into Goddard's Silver Dip converts the tarnish back to silver, enabling soldering. As with all chemicals, Silver Dip should be handled with care, and only the silver-plated tag of the stand-off should be allowed to touch the liquid, but mounting them on a scrap of cardboard neatly solves this problem (see [Figure 3.54](#)).

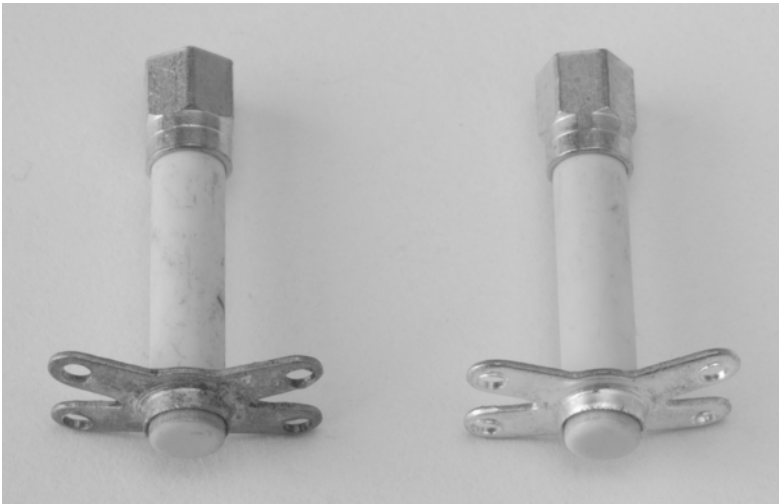


Figure 3.54

Silver-plated tags before and after dipping in "Silver Dip"

Rinse the stand-offs in clean water, and dry them with forced hot air from a hair dryer. Although this technique is very

effective, it is rather smelly and the tags tend to tarnish soon afterwards, so it is best to treat the stand-offs just prior to use.

NOS valve sockets that won't solder

NOS valve sockets are likely to be at least 25 years old, and the (usually tinned) tags will be reluctant to solder – this seems to be a particular problem with Loctal and Septar PTFE sockets. Taking a wire brush to the pins is not a good idea because even a fine wire brush scratches the insulation supporting the tags and embeds those scratches with particles of (conductive) solder, making the base much leakier. One solution is to soak them in a jar of isopropyl alcohol for a day or two, periodically giving the jar a good shake. Really recalcitrant sockets can be periodically removed and brushed with a toothbrush until the tags become clean and shiny. Finish with a good rinse in clean alcohol, followed by clean water, and dry them with forced hot air from a hair dryer.

The alcohol quickly becomes dirty, but can easily be recovered by filtering it through a coffee filter paper supported in a stainless steel kitchen funnel (the alcohol could damage a typical plastic funnel). Be aware at all times that isopropyl alcohol is highly inflammable, so don't smoke whilst handling the stuff!

Enlarging the hole for the wire in a solder tag

You might think this is easy, but solder tags are made of brass, and brass has a nasty habit of snatching when you drill it. Given that tags are small, you've only just discovered the problem, and you want to get on with wiring, this is the perfect recipe for slicing the tip of your fingers. The solution is a small taper reamer (see [Figure 3.55](#)).



Figure 3.55
Tapered reamer for enlarging small holes

This is a very handy tool to keep with your wiring tools. Because it cuts a tapered hole, it needs to be used from both sides of the hole. It throws up a burr, but this isn't a problem, because when you fold your wires around the joint to form a mechanical joint before soldering, the high pressure at the burr makes an excellent contact even before you solder.

Reference

1. "Designing a professional missing console: Part Six – When is a Ground not a Ground?" Steve Dove, Studio Sound. March 1981. pp 56-60.

This Page is Intentionally Left Blank

PART II

TESTING

This Page is Intentionally Left Blank

CHAPTER 4

TEST EQUIPMENT PRINCIPLES

A fully equipped electronic workshop has the entire gamut of test gear from spectrum analysers and oscilloscopes to insulation testers, variable power supplies, and component analysers. Sooner or later, you will have to bite the bullet and buy some test equipment. Your purchase might range from a used moving-coil meter found at an amateur radio fair to a shiny new digital phosphor oscilloscope. Either way, your money is a finite resource and you will want maximum bang per buck.

In this chapter, we will investigate how each instrument works, enabling us to understand its inevitable limitations thus putting us in a position to judge which features are worth paying for, and which can be neglected.

The moving coil meter and DC measurements

All test equipment needs some form of indicator to display the results of the measurement. In the valve era, the choice was between a moving coil meter or a cathode ray tube (CRT). Since CRTs and their necessary support circuitry were (and still are) expensive, the vast majority of test equipment

used a moving coil meter. This has two significant consequences. Firstly, it is fashionable for modern studio and consumer audio equipment to be styled to look “retro”, using hexagonal black bakelite knobs and round black meters having gothic pointers and cream scales. Secondly, there are lots of second-hand moving coil meters available for peanuts.

How a moving coil meter works

If we place one magnet near another, they will leap together so that their unlike **poles** touch (North to South). Conversely, if we try to force like poles together (North to North, or South to South), they repel. When we pass a current through a wire, it generates a weak magnetic field, and if we wind the wire into a tight coil, this concentrates the magnetic field. The coil is then known as an **electromagnet** because electricity is used to produce the magnetic field. Unsurprisingly, the strength of the electromagnet is proportional to the number of turns and the current flowing.

If we were to place the electromagnet near a **permanent** magnet, it would be attracted or repelled with a force proportional to the current and, if we could measure this force, we would effectively have measured the current. Moving coil meters use the force to extend a spring, and because the extension of a spring is proportional to the force (Hooke’s Law), a scale that actually measures extension can be **calibrated** in units of current.

It is quite difficult to make a low-friction **linear** bearing (one that moves in a straight line), but a low-friction rotating bearing is easily made. Unfortunately, although we did not mention it earlier, the attraction and repulsion between the electromagnet and permanent magnet are inversely proportional to the **square** of the distance between them, resulting in

a scale cramped in some sectors but unnecessarily expanded in others. Ideally, we would like a scale with equal-sized graduations throughout its range. Fortunately, shaping the permanent magnet so that it applies a constant radial magnetic field at all positions of the rotating electromagnet produces a linear scale (see [Figure 4.1](#)).

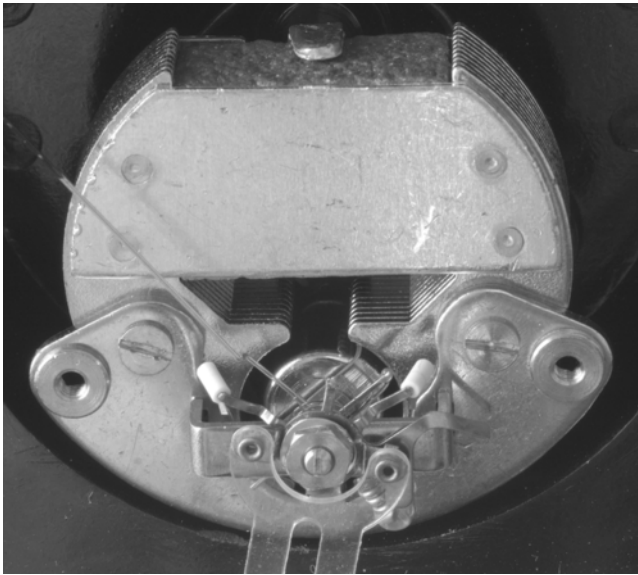


Figure 4.1

Internal view of moving coil meter (Note the shape of the pole pieces to produce a radial magnetic field.)

We now have an instrument whose pointer deflection is directly proportional to the current flowing through the coil.

We can make the meter more sensitive by winding more turns of finer wire on the electromagnet, using a more powerful permanent magnet, or weakening the spring. In practice, it is difficult to make a robust meter that requires less than $50\mu\text{A}$ for full scale deflection (**FSD**), and 1 mA FSD meters are particularly common.

Measuring larger currents

A manufacturer could make different meters to measure different currents, but it is much cheaper to make a standard meter movement and adapt it to measure larger currents. Suppose that we need a 100 mA meter, perhaps for monitoring current in an output valve. Our meter movement still only needs its design current, perhaps 1 mA, to drive its pointer to full scale, so we must bypass, or **shunt**, 99 mA away from the delicate meter movement (see Figure 4.2).

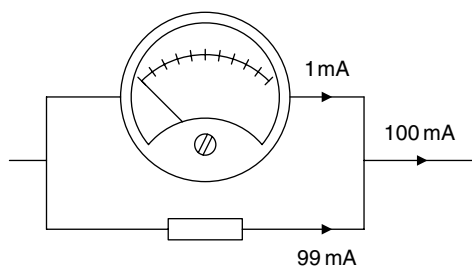


Figure 4.2

Converting a 0–1 mA meter to read 100 mA

Our problem is to determine the correct value of shunt resistor that will allow 100 mA of total current to be shared correctly between the meter (1 mA) and the shunt resistor (99 mA). Fortunately, now that we have stated the problem, the solution is quite easy. The meter's internal resistance and electrical sensitivity are usually stated at the bottom of the scale plate, even if the scale itself is calibrated in foot/Lamberts per fortnight.

If we know the current passing through a known resistance, we can use Ohm's law to calculate the voltage developed across that resistance. As an example, suppose that our 1 mA movement has an internal (coil) resistance of 65 Ω :

$$V = IR = 0.001 \text{ A} \times 65 \Omega = 65 \text{ mV}$$

The shunt resistor is in parallel, so it sees exactly the same voltage as the meter, and if we want to shunt 99 mA, we simply use Ohm's law to find the required resistance:

$$R = \frac{V}{I} = \frac{0.065}{0.099} = 0.657 \, \Omega$$

Sadly, this calculation demonstrates two important points. Firstly, there is no hope of the required resistor being a standard value, and secondly, the resistor is such a low value that the resistance of its wires becomes critical. Fortunately, this is a sufficiently common application that dedicated meter shunts may be available, or a reasonably close value can itself be shunted by another resistor until the correct value is obtained. One solution would be to use five $3.3 \, \Omega$ resistors in parallel ($0.66 \, \Omega$), and trim this combination with a $500 \, \Omega$ variable resistor in parallel.

Measuring small resistances ($<10 \, \Omega$) accurately is difficult, so it is better to set up a test circuit to pass the required 100 mA (perhaps measured by a reliable DVM), and finely adjust the shunt resistance until the meter reads full scale. Nevertheless, adapting a meter to measure different currents is always awkward...

What to do if your meter is unspecified

If you are very unlucky, your meter might **not** specify its sensitivity and internal resistance. Sensitivity can be found by placing it in series with a variable resistor in series with a reliable DVM set to read current, all connected across a battery. Gently reduce the value of the variable resistor from $100 \, \text{k}\Omega$ until the meter reads precisely full scale, then read off the current measured by the reliable DVM (see [Figure 4.3](#)).

The internal resistance could be found using the resistance range of the DVM, but they tend to sink rather a lot of current, and

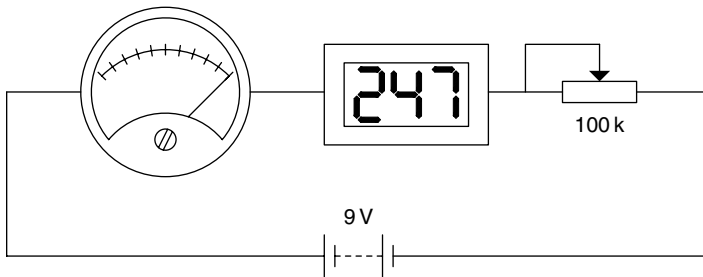


Figure 4.3
Determining the sensitivity of an unknown meter

that might damage the meter under test. A safer method is to leave the variable resistor set to cause FSD, but move the DVM to measure the voltage across the meter under test, and use Ohm's law to calculate its internal resistance (see [Figure 4.4](#)).

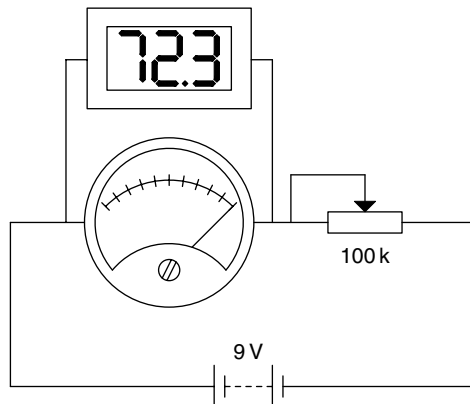


Figure 4.4
Determining the internal resistance of an unknown meter

As an example, perhaps the first test required $247\mu\text{A}$ to attain FSD and when the DVM was used to measure meter voltage, 72.3mV was seen across the meter at FSD.

$$R_{\text{internal}} = \frac{V}{I} = \frac{72.3 \times 10^{-3}}{247 \times 10^{-6}} = 293\ \Omega$$

Measuring voltages

Moving coil meters can be very easily adapted to measure voltages. We know the required current to drive the meter's pointer to FSD, and we know the voltage that we want to cause FSD, so we apply Ohm's law. Perhaps we want to convert our 1 mA meter movement to read 40 V FSD:

$$R = \frac{V}{I} = \frac{40 \text{ V}}{1 \text{ mA}} = 40 \text{ k}\Omega$$

This calculation effectively states that a 40 k Ω resistor would pass 1 mA if 40 V were applied across its terminals. If we placed a perfect 1 mA meter in series with the 40 k Ω resistor, it would achieve FSD when 40 V was applied. You will notice that the final voltmeter has an FSD of 40 V and a resistance of 40 k Ω , or 1 k Ω per volt. It is common for multimeters to specify their loading resistance in this manner because it allows the user to easily assess the loading, the meter imposes on any given voltage range (see [Figure 4.5](#)).

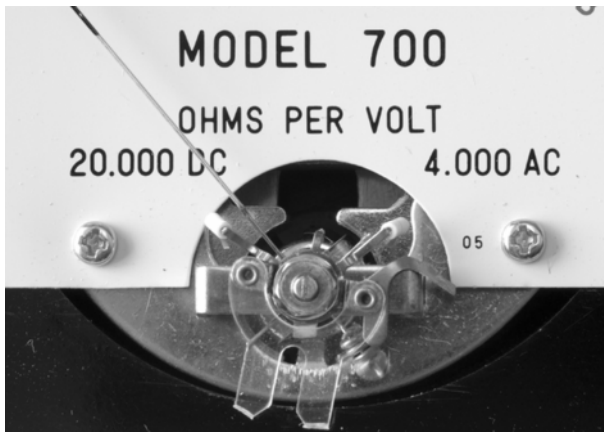


Figure 4.5

Multimeters often specify their sensitivity in Ω/V on their scale plate

But all meters have internal resistance, and our example 1 mA meter has an internal resistance of $65\ \Omega$, so we ought to subtract this from the series resistor to make the total resistance $40\ \text{k}\Omega$:

$$R_{\text{series}} = 40\ \text{k}\Omega - 65\ \Omega = 39.935\ \text{k}\Omega$$

A freshly calibrated laboratory standard moving coil meter with a 5" mirror scale could attain an accuracy at FSD of $\pm 0.5\%$, or to put it another way, there is little point in worrying about precisely trimming the series resistance by 0.16% when the meter itself contributes a greater error. The author simply picks suitable theoretical resistor values from his stock of 1% preferred values and doesn't even bother measuring them.

If this attitude offends your sensibilities, consider why we are fitting the meter in the first place. We very rarely measure to confirm a theoretical value – we generally make a measurement to display **aberrations** from the required value. In short, we want to know if the voltage/current is **not** what it should be. As an example, car manufacturers do not provide a precisely calibrated 0–15 V voltmeter, but simply show red (rather than green) to warn that the battery voltage is not what is ought to be. Once the required value is observably wrong, the precision of the measurement is irrelevant. Distinguishing between the need to identify a fault condition rather than making a precision laboratory measurement is important because it makes identifying fault conditions much cheaper.

Measuring output stage cathode current

It is often necessary to measure an output valve's cathode current. This can be done by inserting a $1\ \Omega$ resistor in the cathode circuit and measuring the voltage across it. Postage stamp-sized DVMs tend to be 200 mV FSD, and 1 mA through a $1\ \Omega$ resistor produces 1 mV, so these meters can be connected directly across the $1\ \Omega$ resistor to give a reading scaled in mA.

We don't need to use a powerful $1\ \Omega$ resistor for current sensing and it's actually a disadvantage to do so. The current sense resistor needs precision and fragility, so a $1\ \Omega \pm 1\%$ 0.6 W metal film resistor is ideal because it allows a reasonably accurate measurement, is almost non-inductive, and might serve as a fuse in the event of catastrophic failure. However, $\pm 0.1\%$ DVMs are cheap, so it seems a shame to throw away that accuracy with a 1% sense resistor. In a push-pull amplifier, the balance of output currents is more important than their absolute value, so we need matched $1\ \Omega$ resistors.

Matching low-value resistors

$1\ \Omega$ resistors can't be measured reliably using a DVM because the resistance of the DVM leads is typically $0.2\ \Omega$ and variable contact resistance makes it difficult to make a measurement. However, matched low-value resistors can be found by soldering a chain of them in series and applying a constant voltage from a regulated power supply across the far ends. Then, using a DVM, measure and record the voltage drop across each resistor and pick resistors with the same measured voltage drop. Accuracy can be improved by:

- Rather than applying a constant voltage across the ends of the chain, drive a constant current through the chain using the constant current regulator described later in this chapter, and power this combination from a regulated voltage. Because the constant current regulator has a constant voltage across it, it is far more able to maintain a constant current.
- Before making the first measurement, adjust the voltage or current to give a reading on your DVM that is just under that required to make it change range. For a $3\frac{1}{2}$ digit DVM, 180 mV is a good voltage, because it will read 180.0 mV, whereas 200 mV is a poor choice because it would read 200 mV, which is less precise (fewer significant figures).

Matching high-value resistors

We can match 1% high-value resistors ($>500\ \Omega$) fairly well simply by measuring their resistance with a DVM. Matching 0.1% resistors is trickier. One way would simply be to buy a more accurate DVM. But $4\frac{1}{2}$ digit DVMs with 0.05% accuracy, or better, are expensive. A much cheaper way is to make a bridge (see Figure 4.6).

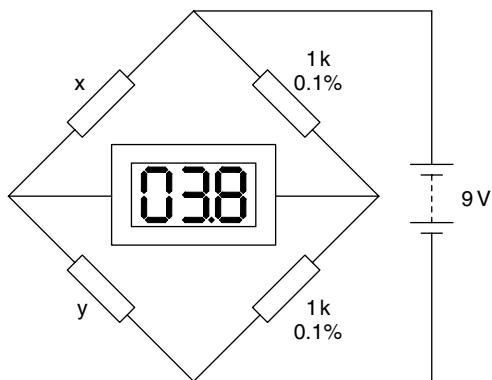


Figure 4.6

This simple bridge circuit enables precision comparison of resistors

Although the bridge was actually invented by Samuel Hunter Christie, this circuit is often known as a Wheatstone bridge because Sir Charles Wheatstone found so many uses for it. However, Wheatstone **did** invent the concertina, or accordion, which is perhaps less forgivable...

We have a pair of equal value resistors on the right hand side and an upper resistor on the left hand side to which we want to find a matching lower resistor. If the ratio of upper to lower resistor on each side is identical, the meter will measure 0 V. If not, the imbalance will cause a voltage, which can be recorded. Resistors with equal imbalance voltage are matched to one another, but not necessarily to the upper resistor. This is a very sensitive test, and allows a cheap $3\frac{1}{2}$ digit DVM set to its 200 mV range to easily match resistors to 0.01%.

If we want to be able to match to a specific resistor, we need the right hand resistors in the bridge to be perfectly matched. Using the previous method, we can find a pair of perfectly matched resistors, and substitute them into the right hand side of the bridge. These reference resistors should be manufactured for low temperature coefficient, so they will probably be 0.1% metal film, or perhaps even 0.01% metal foil. Because the meter measures ratios, rather than an absolute voltage, power supply voltage is not critical, and the bridge can even be powered by a 9 V battery.

Meters and AC measurements

A moving coil meter can only respond to DC, so to measure AC, a rectifier must be added to convert the AC to DC. But AC waveforms can be of any shape, and we are not always able to view them on an oscilloscope, so we need a means of describing important parameters. Once we know the different ways of describing AC, we can choose the rectifier that delivers the measurement we need.

Peak voltages

A convenient way to specify complex waveforms is by peak voltages (see [Figure 4.7](#)).

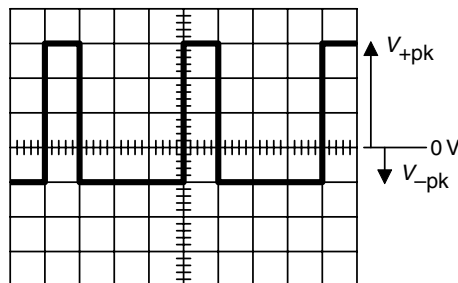


Figure 4.7
Measuring peak voltages relative to 0 V

Note that because the waveform is asymmetric about 0 V, it is necessary to specify V_{+pk} and V_{-pk} separately. In this example, the vertical scale is 1 V/div, so $V_{+pk} = 3$ V, and $V_{-pk} = -1$ V.

If required, V_{pk-pk} could be specified. This could be measured directly as $4 V_{pk-pk}$, or it could be calculated from the previous measurements:

$$V_{pk-pk} = V_{+pk} - V_{-pk} = +3 \text{ V} - (-)1 \text{ V} = 4 \text{ V}$$

V_{pk} and V_{pk-pk} are particularly useful measures when testing a stage for overload because the numbers directly correlate with the numbers predicted by loadlines on a valve's anode characteristics, making it easy to check theory against practice.

Mean level

It is often useful to know the average, or mean level (either voltage or current) of an AC waveform. Since the waveform is assumed to be repetitive, the mean level is determined by finding the mean of one cycle of the waveform.

Conventionally, to find a mean, we sum individual data and divide by the number of items of data. Because an oscilloscope display is a graph of voltage against time, we sum voltages multiplied by their time (producing area) and divide by the total time of the waveform's cycle (see [Figure 4.8](#)).

Note that although this waveform has exactly the same **shape** as the previous one, $V_{+pk} = 4$ V and $V_{-pk} = 0$ V. Over one cycle:

First unit of time:	$V = +4$ V
Second unit:	$V = 0$ V
Third unit:	$V = 0$ V
Fourth unit:	$V = 0$ V

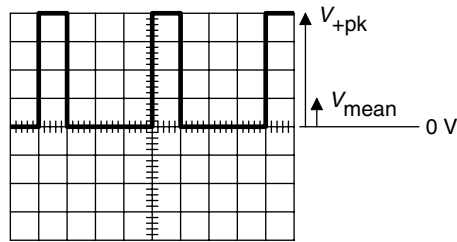


Figure 4.8

Mean voltage can be intuitively seen because it is proportional to the **area** under the waveform. Equal areas above and below 0 V would force $V_{\text{mean}} = 0 \text{ V}$

We want to find the mean of four data items, so:

$$V_{\text{mean}} = \frac{4 \text{ V} \times 1 + 0 \text{ V} \times 3}{4} = 1 \text{ V}$$

We have calculated the mean level over one cycle of a waveform and assigned it a fixed value, which implies that it is unchanging. If the waveform is repetitive, each cycle is the same as the last, so the mean level must also be constant from one cycle to the next, and it must therefore be a DC voltage (or current).

Any waveform may be considered to be an **AC component** riding on a constant level of DC known as the **DC component** (V_{mean}).

When a single triode stage distorts, it produces predominately 2nd harmonic distortion, and a side effect of this is that it changes the mean level of the signal. A cathode-biased stage typically uses a resistor bypassed by a capacitor, and the capacitor integrates the audio signal to produce a voltage that is partly determined by the mean level of the applied audio. The significance of this is that if we had an accurate DVM, but didn't have a means of directly measuring distortion, we could infer distortion by measuring the change in the DC voltage across the capacitor with and without the signal applied.

Later, we will find that mean level crops up in all sorts of odd places.

Power and RMS

It is often necessary to be able to determine the heating power that a waveform would dissipate in a resistor using $P = V^2/R$ or $P = I^2R$. For DC, this is simply a matter of measuring the appropriate quantities and performing the calculation.

For AC, we need a means of specifying amplitude such that the previous power equations give the correct power. This new value is known as “root of the mean of the squares”, which refers to the way it is calculated. It is invariably abbreviated to RMS or V_{RMS} .

For this example, we have reverted to our original waveform, which we could now describe as having a $4 V_{\text{pk-pk}}$ AC component, but zero DC component. Because the DC component is zero, the waveform has dropped down the oscilloscope screen so that part of it swings negatively (see Figure 4.9).

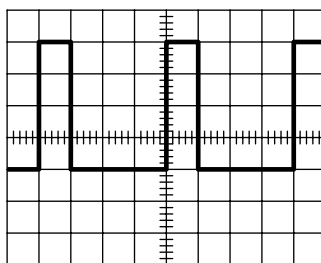


Figure 4.9

RMS voltage: No intuitive result is possible simply by viewing the waveform

We calculate V_{RMS} in a very similar way to V_{mean} , but we first square each voltage, find the mean (of the squares), and finally take the square root of this mean.

$$V_{\text{RMS}} = \sqrt{\frac{3^2 \times 1 + (-1)^2 \times 3}{4}} = \sqrt{3} \approx 1.732 \text{ V}$$

Fortunately, we rarely have to calculate V_{mean} or V_{RMS} manually because a digital oscilloscope is really an application-specific computer in disguise, so it is ideally suited to performing hideous number crunching in the twinkling of an eye (see Figure 4.10).

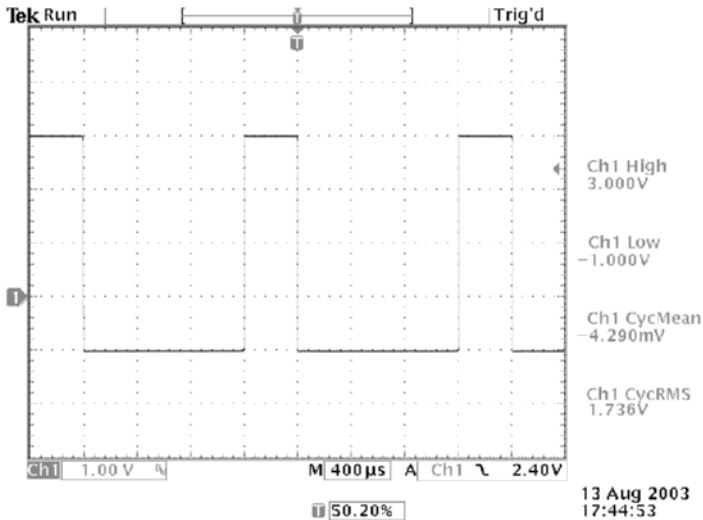


Figure 4.10

Oscilloscope automated measurements. Wonderful!

The slight discrepancy in V_{mean} and V_{RMS} measured by the oscilloscope compared to the calculated values simply reflects the author's inability to set the controls precisely on his rather cheap (but perfectly adequate) function generator.

Some waveforms such as the sine wave are used so often that conversion factors have been calculated for determining V_{RMS} when V_{pk} is known. **For the sine wave only:**

$$V_{\text{RMS}} = \frac{V_{\text{pk}}}{\sqrt{2}}$$

The most obvious use of V_{RMS} is in determining the output power of an amplifier from a measurement of voltage across a

known load resistance, but since we are only concerned with undistorted output power, a mean reading meter calibrated RMS sine wave is perfectly acceptable. A simple moving coil meter or cheap DVM responds to the mean voltage, **assumes that the waveform is a sine wave**, and applies a correction factor that corresponds to the power that the sinusoidal waveform would dissipate in a resistor (see Figure 4.11).

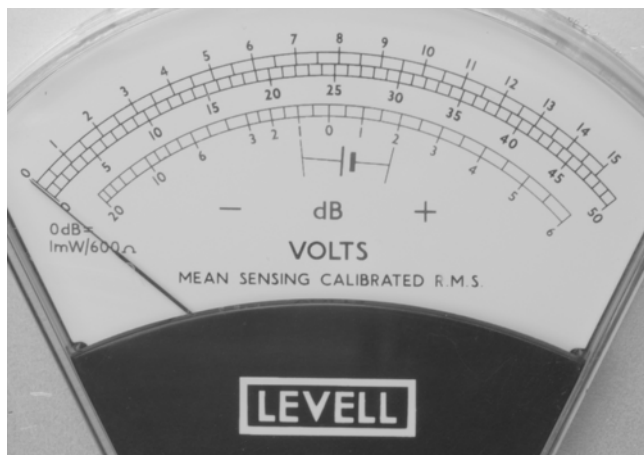


Figure 4.11

This AC microvoltmeter explicitly states that it responds to the mean level of (the rectified) AC waveform and is calibrated to give a correct RMS reading for sine waves only

Peak to mean ratio

It is often convenient to consider the ratio of V_{pk} (either positive or negative peak) to V_{mean} .

$$\text{Peak to mean ratio} = \frac{V_{pk}}{V_{mean}}$$

The full significance of peak to mean ratio will become apparent when we investigate dedicated audio test sets and measurement of noise.

Crest factor

A waveform composed mainly of high-amplitude narrow spikes is difficult to measure, so this property is defined.

$$\text{Crest factor} = \frac{V_{\text{pk}}}{V_{\text{RMS}}}$$

Unlike a cheap DVM, a **true RMS** DVM does **not** assume a sine wave and apply a conversion factor, instead, it actually calculates V_{RMS} for the waveform from first principles. Although you and I can calculate V_{RMS} for **any** waveform, true RMS DVMs are somewhat more limited, and they specify their limitation by stating the maximum crest factor that can be tolerated whilst accurately determining V_{RMS} . As an example, the Fluke 89 IV true RMS DVM can compute RMS accurately provided that the crest factor is $\leq 3 \times \text{FSD}$.

Although DVM manufacturers make a big fuss about true RMS, there is only one time when we really need the facility, and that is when measuring AC heater voltages when the waveform has become distorted and we are truly concerned about the precise heating effect. Otherwise, we can knock true RMS off our shopping list when choosing a good DVM. In practice, if we decide that we want a particularly accurate DVM, we often find that the facility is included whether we want it or not.

Speed of measurement

Despite the prevalence of DVMs, moving coil meters still have their uses. A moving coil multimeter flicks quickly to its reading, and this may warn you to switch off sufficiently quickly to avoid damage.

The industry standard meter in the valve era was the AVO Model 8 moving coil multimeter – which is still available.

You will find many circuit diagrams stating that voltages were measured using a $20\text{ k}\Omega/\text{V}$ meter, which refers to the loading that the AVO 8 imposes on the circuit. Unfortunately, the AVO 8 has a heavily damped movement, making it frustratingly slow to use, so the author prefers a **really** cheap moving coil meter (see [Figure 4.12](#)).



Figure 4.12

This cheap moving coil multimeter responds quickly, so it is still useful!

Digital voltmeters (DVMs)

We have mentioned DVMs in passing, but it is now time to consider them in detail. The main advantage of a DVM is that it can be designed to achieve arbitrary accuracy, whereas moving coil meters require precision mechanical engineering (springs, bearings, individual calibration), powerful magnets (with a substantial proportion of expensive cobalt) to achieve an accuracy of $\pm 0.5\%$ at best.

Despite the fact that we don't actually know what time really **is**, we can measure it stunningly accurately very cheaply. (You would be annoyed if the digital watch given away with a gallon of engine oil had an error of one minute in a week, yet this corresponds to an error of <10 parts per million.)

Thus, whenever possible, digital measuring systems convert the measurement of the required parameter into a measurement of time. Unfortunately, if time is used to **make** the measurement, we must inevitably **take** time to make that measurement – usually by counting the pulses of a master clock. Thus, the 0–2 V range of a $3\frac{1}{2}$ digit DVM might correspond to counting from 0 to 2000 pulses, and displaying that count scaled by an appropriately positioned decimal point. More significantly, if we wanted a $4\frac{1}{2}$ digit DVM, it would need to count 20 000 pulses, which would either take ten times as long, or require a master clock ten times faster.

We are now in a position to compare DVMs. The natural attitude is to say, “I want maximum accuracy with minimum cost, and I'm not really bothered about measurement time.” Expensive experience forces the author to state that you **are** bothered about measurement time. It is not sufficient for the DVM to make the measurement once and stop, it must make the measurement repeatedly. As an example, we might want to set the anode current in an output stage to precisely 140 mA by

adjusting a variable resistor. When we adjust the resistor, anode current changes instantaneously, but it takes time for the DVM to make the new measurement that reflects this change, by which time, we may have overshoot our adjustment. Remember that we usually measure to check that a parameter **isn't** wrong. If it **is** wrong, we must minimise the time that it is wrong, to minimise the smoke – so the addition of measurement time to human reaction time might risk the survival of a set of output valves worth £200. You can buy a very good DVM for £200, yet **one** measurement might recoup that cost. Are you still unconcerned about measurement time?

A DVM is fundamentally a precision analogue to digital convertor (ADC) preceded by switchable attenuators and amplifiers. **All** ADCs have the following relationship:

$$\text{cost} \propto \text{sample rate} \times \text{number of bits}$$

In terms of a DVM, the sample rate is the number of measurements per second, and the number of bits reflects the basic accuracy of the DVM – usually best on the 2 V or 5 V DC range.

Although DVM manufacturers are open about errors and price, they tend not to mention measurement speed. In an effort to beat the price/performance equation, DVMs often add an “analogue” bar graph display to their digital display. The bar graph display is not accurate, but it **is** fast. Because any bar graph display must be composed of discrete segments, it is usually supported by auto-ranging electronics to allow it to display **small** changes. The upshot is that the display faithfully displays instantaneous **changes**, but the precise value is unknown.

In short, although bar graphs are very useful for revealing unstable parameters, nothing beats a fast accurate measurement, and this is one of the fundamental differences between cheap and expensive DVMs.

Minor points when choosing a DVM

Some models may not measure current, but this is not a great loss, since to measure current you must break a wire and later reconnect it. Others may measure capacitance, but they are not usually particularly accurate and do not measure small capacitances very well, so you may feel that it is better to put money aside towards a second-hand component bridge.

Some DVMs are the size of an (overgrown) pen, allowing you to read the display without looking away from the position of the probe tip. One very useful feature on a DVM is a resistance range that is guaranteed not to switch diodes on, because this allows you to measure resistors in circuit without semiconductors upsetting the measurement. The better DVMs usually **don't** have this facility because restricting the applied voltage to 200 mV makes it more difficult to achieve a precision measurement. If you are forced to have only one digital multimeter, you need a fast meter that probably includes a bar graph to allow changing values to be clearly seen.

Summarising, there is no such thing a single perfect meter, and each type has its advantages and disadvantages. Once you start testing, you will quickly discover that you can't have too many meters, so it's worth having different types. At the last count, the author had five DVMs and three analogue meters.

Analogue oscilloscopes

There is a great deal of software available for converting a computer with a soundcard into a clumsy oscilloscope. Forget it. A 192 kHz soundcard has a maximum theoretical bandwidth of only 96 kHz, and you need far more than that.

The author has not previously recommended oscilloscopes because they were expensive and require skilled use and

interpretation, but prices are falling and people are becoming more knowledgeable. However, there is a catch. Analogue oscilloscopes use a cathode ray tube (CRT) to produce the display, which requires ≈ 10 kV of EHT to accelerate the electrons to the phosphors. High-voltage circuits are always fragile, and oscilloscope EHT supplies are no exception. Although many parts are generic, the step-up transformer and voltage multiplier chain are invariably specific to the instrument, and may no longer be available. If you buy a cheap second-hand 20 MHz oscilloscope, and it fails after a few years, it will probably have paid for itself, but second-hand >100 MHz oscilloscopes are more expensive, so an irreparable failure is galling.

Whether you are using the latest technological confection from Tektronix, complete with all possible bells and whistles, or a 5 MHz dinosaur from a junk shop, the principles are the same. An oscilloscope is simply an electronic graph drawing machine that draws graphs of voltage (vertical, or “Y” axis) against time (horizontal, or “X” axis).

An analogue oscilloscope uses a CRT to draw the graph by repeatedly tracing a beam of electrons onto the fluorescent phosphor screen to form a visible trace. The CRT is a large valve with focussing anodes (collectively known as an electron lens) that shape the electron stream from cathode to final anode into a beam that is brought into focus on the fluorescent phosphor coating on the inside of the tube face to form a sharply defined spot. The phosphor coating is not conductive, but current flows because the electrons in the beam strike the phosphor so hard that one fast-striking electron releases a slow-moving secondary electron in addition to a photon of light. The slow-moving electron is easily attracted to the concentric final anode, and thus a circuit is formed. The spot is scanned repeatedly across the tube face by applying a ramp waveform to the horizontal deflection anodes, and the persistence of the phosphor coating coupled to the persistence of the eye produces a continuous display.

In order to produce a stationary trace, **all** oscilloscopes have three main blocks which have to be adjusted correctly (see [Figure 4.13](#)).

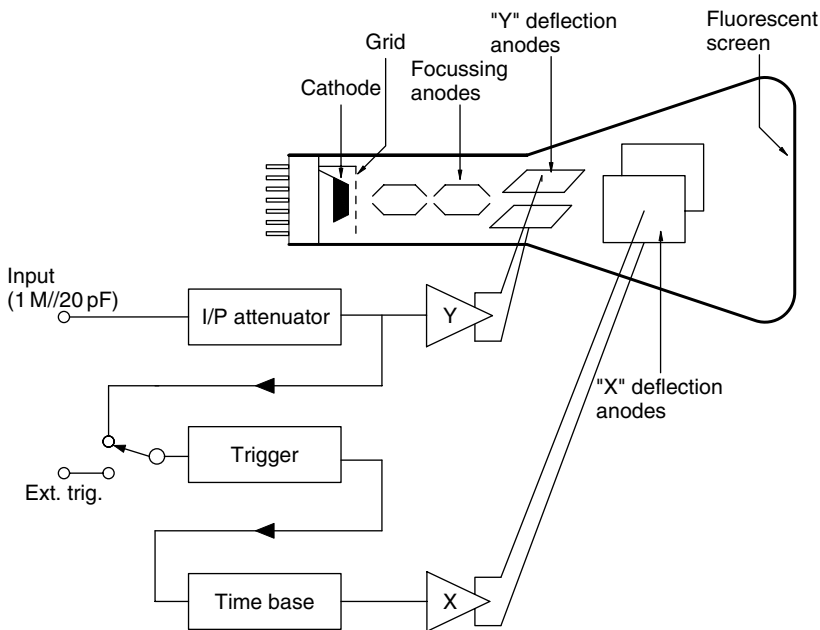


Figure 4.13
Analogue oscilloscope block diagram

The “Y” amplifier

The “Y” **amplifier** is responsible for controlling the deflection and positioning of the beam in the vertical direction of the display tube. Adjustment of the relative DC potentials on the “Y” deflection anodes **shifts** the beam. Adjustment of the input attenuator controls the sensitivity of the deflection – commonly labelled in **volts/div**. This control uses a 1, 2, 5 sequence because it gives equal spacing on a logarithmic graph and allows the oscilloscope to adapt smoothly to the natural world. (Currency uses this logarithmic sequence for the same reason; 1p, 2p, 5p, 10p, 20p, 50p, £1, £2, £5, £10, £20, £50.)

There is often a **fine** control (sometimes called **variable**, or **vernier**) that allows the vertical scaling to be finely adjusted to make the waveform conveniently fit on the screen. You can't measure the voltage anymore, but most of the time, we use an oscilloscope because we are more concerned about the shape of the waveform.

All oscilloscopes also have a choice of input **coupling**, marked **AC**, **DC**, or **GND** (see Figure 4.14).

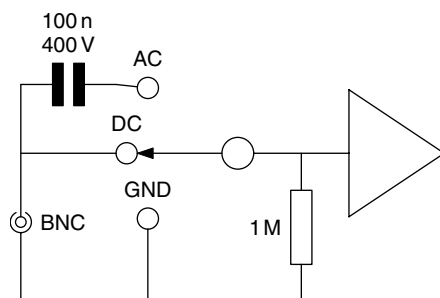


Figure 4.14
Input coupling to the “Y” amplifier

- **DC:** All signals, including DC, are passed to the deflection anodes. This is the preferred mode of operation as it does not attenuate low-frequency signals, and therefore does not distort square waves, but it may not be practical when investigating small audio signals in a valve circuit.
- **AC:** DC is blocked by a capacitor, this is useful for investigating the AC conditions of a circuit separately from the DC conditions – such as faultfinding.
- **GND:** This connects the input of the oscilloscope to ground/earth. If the trace is then moved to a convenient reference point, when returned to DC coupling, the movement of the trace from the reference allows the applied DC to be measured easily. (Using an oscilloscope to measure DC might seem like a sledgehammer to crack a nut, but it responds faster than even an analogue meter, so it can be very useful when faultfinding.)

The time base and “X” amplifier

The **time base** provides the horizontal, or “X”, sweep, in a 1, 2, 5 sequence typically ranging from milliseconds per division to nanoseconds per division. As with the “Y” amplifier, there is often a variable control that finely adjusts the horizontal scaling to make the waveform conveniently fit on the screen, but the horizontal sweep is now uncalibrated, so absolute time measurements can no longer be made.

The time base control adjusts the frequency of a precision ramp generator. Because the oscilloscope display tube uses electrostatic deflection anodes to deflect the beam by an amount proportional to applied voltage, the combination of the two produces a sweep where horizontal deflection is proportional to time.

Depending on the tube type and final EHT voltage (higher HT makes the electrons travel faster, making them more difficult to deflect), the deflection anodes require $\approx 200 V_{pk-pk}$ to deflect the beam from one side of the tube face to the other. A dedicated high-voltage “X” amplifier very similar to the “Y” amplifier is therefore required to deliver the sweep voltages required by the tube.

If the sensitivity at the input of the “X” amplifier is made the same as that at the input of the “Y” amplifier, the “X” amplifier’s input can be switched to accept either the ramp from the time base generator, or Ch2 (leaving Ch1 to go to the “Y” amplifier). This is known as **“XY” mode**, and because it is so cheap to provide, all oscilloscopes offer it. “XY” mode allows the display of **Lissajous** figures – these are virtually useless, but are loved by producers of science fiction programmes and “James Bond” films (see [Figure 4.15](#)).

The Lissajous figure in the diagram compares the voltage waveform across a choke with the current waveform through it. Because the phase between current and voltage is zero at

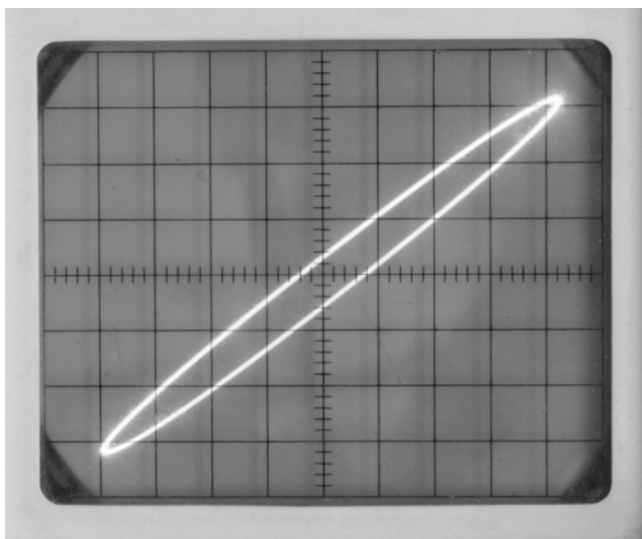


Figure 4.15

A Lissajous figure showing a perfectly straight line (rather than a slight ellipse) would indicate zero phase error

resonance, and changes very sharply around it, testing for phase is a very sensitive way of determining resonant frequency. When the two waveforms are in phase, a perfect straight line results, rather than an ellipse.

Triggering

If we were to allow the time base to sweep randomly with respect to the input signal, we would see a mess of moving patterns on the screen. We need to sweep the electron beam in such a way that it draws each trace perfectly overlaid onto the previous trace, thus producing a single trace of the input waveform. If instead of letting the time base sweep continuously, we only allowed it to begin a sweep at the instant that a particular feature of the input waveform had been identified, this would **synchronise** the time base to the input signal and force the traces to overlay.

The **trigger** identifies the significant part of the input signal's waveform both by its absolute level and whether that voltage is rising or falling.

The simplest form of triggering is **AC**, which triggers the time base each time the input waveform passes through zero. **DC** or **normal** triggering allows the trigger voltage to be determined by the user adjusting the **trigger level** control. Whether in AC or DC mode, trigger **slope** can be chosen to be + (rising edge) or – (falling edge) (see Figure 4.16).

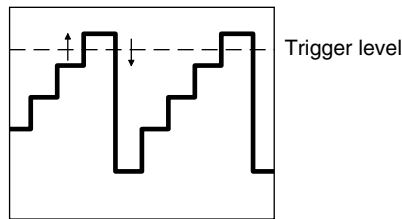


Figure 4.16

The trigger identifies a unique point on the waveform using a combination of voltage and slope

In this example, the trigger level has been set so that the oscilloscope trigger fires on the top step of the staircase waveform. But the waveform goes **up** the step, then **down**, so if the trigger fired on both edges, we would see two traces, one shifted horizontally from the other. To avoid this problem, trigger slope enables us to select whether the trigger fires on the rising or falling edge.

Note that most analogue oscilloscopes produce a sweep that is slightly wider than the display, so the start of the sweep at the trigger point and the end of the sweep are not seen unless deliberately shifted onto the screen.

Because in this example the time base has been set to show more than one cycle of the waveform, the oscilloscope would

only manage to sweep half way across the screen before being triggered again, starting a new sweep, so it would never manage to sweep all the way across the screen. To prevent the problem of truncated sweeps, once the trigger has fired, it is **guarded** until the time base has completed the sweep and is ready to begin the next.

If the time base is only allowed to sweep after the trigger has been fired, it is known as **normal** triggering. However, the disadvantage of normal triggering is that if you haven't been able to trigger the oscilloscope, you don't have a display, so you don't know why you have failed to trigger. To circumvent this problem, oscilloscopes have **AC auto** triggering (occasionally known as **bright line**) which automatically triggers the time base after each sweep (**free run**), or is triggered when the trigger waveform passes through 0 V. Although this mode of triggering can cause an unlocked trace, it makes it easier to determine which controls need to be adjusted to produce the required display.

The signal may contain noise that we want the trigger to ignore. If we were looking at a small signal in a pre-amplifier, it might well have high-frequency noise on it, and if we were to allow that noise to reach the trigger, it would make the display of our wanted signal unstable. (Remember that noise is proportional to the square root of bandwidth, so a 20 MHz oscilloscope sees 30 dB more noise than we can hear in the 20 kHz audio bandwidth.) For this reason, manufacturers include **HF reject**, which is typically a ≈ 30 kHz low-pass filter at the input to the trigger. Conversely, our small signal might be riding on some mains hum, and we might be able to reject this using **LF reject** (high-pass filter at ≈ 80 kHz), but we will see a much better way in a moment.

The trigger needs to be provided with an input signal from which to trigger. **Internal** triggering picks off one of the signal channels after their attenuators, and is typically marked Ch1 or

Ch2. **External** triggering may come from a front panel BNC or from **AC line**, which is the mains supply.

Suppose that you are investigating an amplifier, and have applied a 1 kHz sine wave from an oscillator to its input, by definition, signal levels within the amplifier change as you move from test point to test point, probably requiring adjustment of trigger level if internal trigger is selected. However, if you also connect the oscillator to the external trigger input, and trigger from this point, the display will always be correctly triggered, no matter where you probe, and no matter how much noise or equalisation has been added to your signal. External triggering is almost always better than internal, so it is a good idea to get into the habit of using it.

Similarly, when you are investigating mains hum, if you select AC line, the trace will always be locked. Incidentally, if you suspect that you are seeing mains hum, but are not sure, if the trace becomes stationary when you flick the trigger to AC line, mains hum is confirmed.

Bandwidth

The most important single specification for a motorcycle is its engine capacity and the counterpart for an analogue oscilloscope is bandwidth. And just like the motorcycle, an oscilloscope's bandwidth is invariably prominently displayed next to its maker's name.

Strictly, bandwidth is defined as the difference in frequency between the HF and LF roll-offs where the response has fallen by 3 dB. Thus, an FM tuner might define the bandwidth of its 10.7 MHz intermediate frequency amplifier as being $10.85 \text{ MHz} - 10.55 \text{ MHz} = 300 \text{ kHz}$. Because all oscilloscopes respond to DC (which is 0 Hz), it is sufficient to specify an oscilloscope's bandwidth simply by quoting the HF roll-off.

A 20 MHz oscilloscope is probably the slowest new oscilloscope that you could buy today, and this is just fast enough for audio. But if the ear can only hear to 20 kHz, why do we need one thousand times as much bandwidth? The first clue is in the definition of bandwidth. A 20 MHz oscilloscope will display a 20 MHz sine wave, not at its correct amplitude, but attenuated by 3 dB. A well-focussed trace is just capable of displaying a 1% amplitude error, so we would ideally like the limited bandwidth of our oscilloscope to contribute a smaller error, perhaps $\frac{1}{2}\%$ error, at the highest frequency for which we want to be able to measure amplitude correctly. If the oscilloscope's transient response has been optimised, its HF response corresponds to that of a single CR network, and we find that for $\frac{1}{2}\%$ error we need the -3 dB frequency to be ten times higher. In other words, our 20 MHz oscilloscope introduces its own amplitude errors above 2 MHz. (It is precisely this 10:1 ratio that justified the existence of 60 MHz oscilloscopes – traditional analogue television signals had a bandwidth of 5.5 MHz.)

Having established that our 20 MHz oscilloscope is only accurate to 2 MHz, this is still way beyond 20 kHz, so why do we need this bandwidth? Unfortunately, prototype audio power amplifiers do not always behave as they should, and it is not uncommon to find them oscillating between 1 and 2 MHz, so our oscilloscope needs to be able to display any oscillation perfectly in order that we can see it, and do something about it. There's an old saying that you can be neither too rich nor too thin. As far as oscilloscopes are concerned, you can't have too much bandwidth...

The other parameter that goes hand in hand with bandwidth is the fastest time base speed. The ideal display contains one or two cycles of the waveform. Taking 20 MHz as an example, the **period** (duration of one cycle) is:

$$T = \frac{1}{f} = \frac{1}{20 \times 10^6} = 50 \text{ ns}$$

Oscilloscope screens generally have ten horizontal divisions, so to display two cycles of 20 MHz, we need a time base that sweeps five divisions in 50 ns. In other words, it needs to be able to sweep at 10 ns/div. In practice, very few 20 MHz oscilloscopes are able to sweep as fast as this, and the author's (cheap and cheerful) Gould OS245A fastest sweep is 1 μ s/div, so it can only adequately display 200 kHz, despite being described as a 10 MHz oscilloscope. This isn't really fast enough for faultfinding audio, because the easiest route to curing unwanted oscillations is to find a circuit change that alters the frequency of oscillation. Once that has been done, curing the oscillation is usually easy.

For faultfinding power amplifiers, we would ideally like a time base that can sweep at 100 ns/div, or better. The next step up from 20 MHz is a 100 MHz oscilloscope such as the 1970s vintage Tektronix 465, which can sweep at 20 ns/div and has lots of extra features which we now need to investigate.

Bells and whistles

Bells and whistles allow detailed examination of more complex signals, but these refinements still fall into the basic three blocks of "Y" amplifier, time base, and trigger.

The "Y" amplifier

Frequently, we need to investigate more than one signal at a time and compare relative timings or voltages. To do this, we add extra input attenuators and amplifiers, and switch sequentially between them at the input to the final "Y" deflection amplifier (see [Figure 4.17](#)).

Older oscilloscopes require you to choose between **alternate** or **chop** modes to switch between the channels. As the names suggest, alternate mode sweeps alternately between Ch1 and

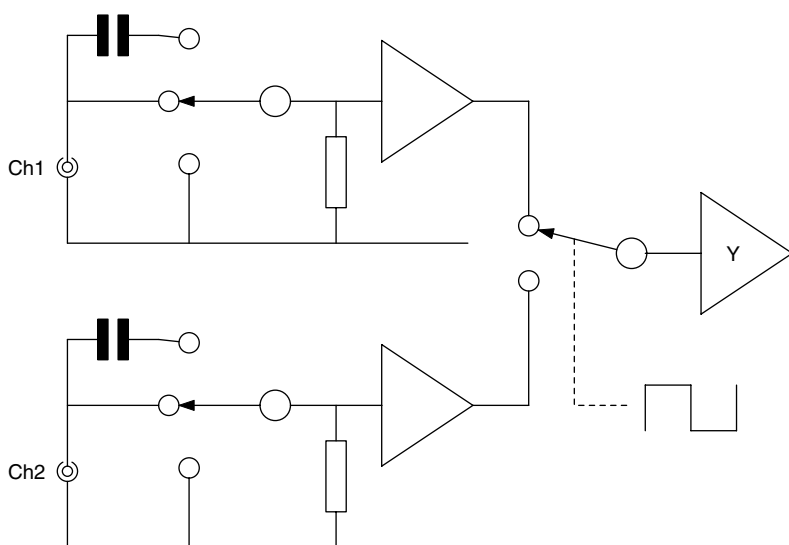


Figure 4.17

Adding channels to an oscilloscope simply requires additional attenuators, input amplifiers, and an electronic switch

Ch2, and this is best for time base settings $>2\text{ms/div}$, but below that, chop mode is better because chopping between the channels as they are swept avoids flicker at slow time base speeds.

Although most oscilloscopes have two channels, some modern oscilloscopes can display up to four input channels at once, but the screen becomes rather cluttered, and the display dims. (Brightness at a given point is proportional to the number of electrons striking that point, so if the electron beam has to produce four traces instead of one, each trace must be reduced to one quarter of the brightness of a single trace.) Surprisingly, the more significant limitation is that four channel oscilloscopes often do not have a separate triggering channel, so the fourth channel ends up being used as the external trigger, and your extra money actually bought a three-channel oscilloscope (possibly with limited “Y” attenuators on Ch3 and Ch4), so check this very carefully.

Cursors are internally generated vertical or horizontal moveable lines that are added at the input to the “Y” amplifier and are the modern development of crystal calibrator **pips** that were sometimes included in early oscilloscopes and radar. Because cursors are digitally generated, the count that produces their position can be displayed as a number, and this number can be modified automatically by “Y” amplifier setting or time base setting to give a read-out directly in terms of time or voltage. Cursors have three valuable advantages:

- Without cursors, you measure voltage (or time) by counting squares and multiplying by the volts or time per division – which is tedious and prone to error.
- Because the phosphor trace is behind the (thick) glass face of the CRT and the ruled **graticule** (scale) is on the outside, the absolute position of your eye changes their relative position. On superior oscilloscopes this **parallax** error was eliminated by painting the graticule lines on the **inside** of the CRT face (**internal graticule**).
- Measurement referred to a ruled scale assumes that CRT deflection is linear, that deflection amplifiers are linear, and that the time base ramp is linear. In practice, none are perfectly linear, but cursor position is distorted by the same amount as the waveform’s features, so all the errors cancel.

The addition of cursors converts the oscilloscope from being a display device into a measuring instrument (see [Figure 4.18](#)).

The time base

Although we argued earlier that we want as fast a time base as possible, a fast time base is rather like looking through a telescope – you can see in great detail, but you don’t know where you are looking. High-powered astronomical telescopes solve the problem by fitting a low-power sighting telescope along the barrel of the main instrument. Once the sighting

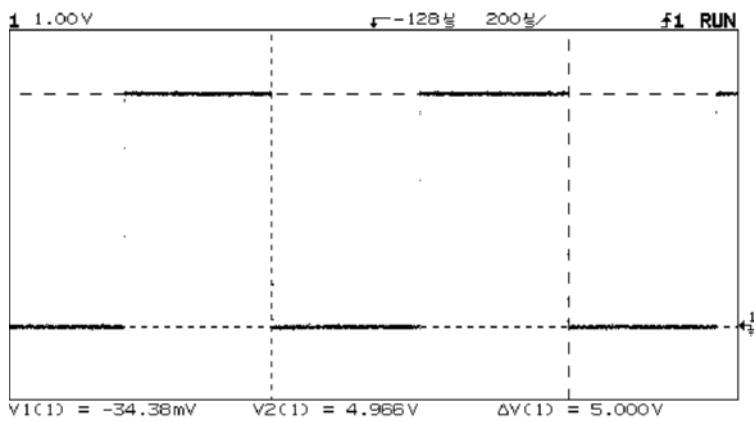


Figure 4.18
Cursors enable quick and accurate measurement of voltage or time

telescope has located the region to be explored, the observer switches to the main instrument. Similarly, fast oscilloscopes use their **main** or “A” time base to find the region on the waveform to be investigated, then the **delayed** or “B” time base is engaged to give the detailed view. Since position across a sweep is time, adjusting the **delay** before the “B” time base fires determines which part of the waveform is to be investigated in detail (see [Figure 4.19](#)).

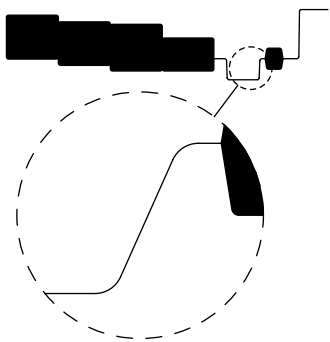


Figure 4.19
Delayed time base enables detailed examination of any part of this complex waveform

To help the user, typical oscilloscopes have two intermediate modes between the “A” and “B” time bases. **Intensified** mode adds a **bright-up** to the main display which highlights where the “B” time base will sweep. Once the correct region has been chosen by adjusting delay, **mixed** mode can be engaged, which alternates between “A” time base and “B” time base, allowing fine delay adjustment with confidence. Finally, the “A” display can be switched off by selecting “B” time base only, resulting in an uncluttered display of the precise area of interest.

Delayed time base is intended for investigating high-frequency detail buried in a low-frequency waveform, so it is particularly suitable for investigating:

- Ringing in HT supplies caused by diode switching
- Ringing at the leading edge of square waves caused by incorrectly terminated audio transformers
- Heater noise spikes caused by diode switching.

Using the delayed time base tends to dim the display, so the fastest oscilloscopes use CRT image intensifier techniques (as used in night sights) to make the display brighter, but at increased cost.

Triggering

In order to use the improved time base system effectively, the triggering must become more sophisticated. Early delayed time base oscilloscopes offered a separate trigger for the “B” time base, but this isn’t particularly useful, so most modern oscilloscopes simply start the “B” time base immediately after the delay. (This must be one of the few whistles that has been discarded!) Other features include triggering on numbered television lines or pattern triggering where the trigger recognises a pattern such as a particular sequence of logic pulses.

Holdoff is a very useful facility that allows the trigger to be guarded by an adjustable time, so that only the intended

transition triggers the oscilloscope. This facility is particularly useful for persuading an oscilloscope to synchronise to the data words in the serial digital data stream leaving the S/PDIF output of a CD player.

Digital oscilloscopes

Analogue oscilloscopes apply an amplified input waveform directly to the CRT, so it must have the same bandwidth as the “Y” amplifiers. 20 MHz CRTs are easily and cheaply made, and even 100 MHz is not too much of a problem, but >400 MHz is much more expensive, so the 1 GHz Tektronix 7104 was a very rare beast.

Digital oscilloscopes solve the CRT bandwidth problem by divorcing the wide bandwidth of the input waveform from the display. They do this by converting the analogue waveform into digits, storing them in memory, and reading them out at a (much slower) rate chosen to suit the display. Thus, when digital oscilloscope manufacturers refer to bandwidth, they mean the **analogue** bandwidth of the attenuator and input amplifier system up to the input of the ADC. Not only does storing and displaying the waveform at a slower rate mean that an extremely cheap display is perfectly suitable for a 5 GHz oscilloscope, but it also means that a faster time base is easily achieved. As an example, the (now obsolete) 100 MHz HP54600B can sweep at the ideal 2 ns/div, enabling two cycles of 100 MHz to fill the screen.

Because digital oscilloscopes are so much cheaper to make, the major manufacturers no longer offer analogue oscilloscopes. Nevertheless, despite all the fuss made by marketing departments, a good analogue oscilloscope such as the 1970s vintage 350 MHz Tektronix 485 takes a lot of beating. To make an informed choice between buying a new digital or an old analogue oscilloscope, we will need to investigate some of the

murkier areas of digital principles about which some oscilloscope manufacturers are distinctly coy...

The ADC (Analogue to Digital Convertor)

Only digital oscilloscopes have an ADC, commonly known as the **acquire** block, and its principles of operation need to be understood to obtain good results.

Sampling and quantising

An analogue signal can change continuously in both time and voltage. Digital oscilloscopes take the analogue signal and break it up in time (**sampling**) and in voltage (**quantising**). Having broken the signal, they convert it to a stream of binary digital numbers that describe that signal.

When an analogue signal is quantised, there are always errors because the analogue voltage cannot be perfectly described by a finite number of quantising levels (see [Figure 4.20](#)).

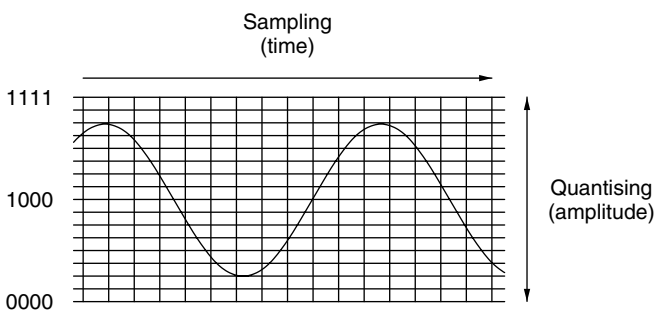


Figure 4.20

Quantising breaks voltage into steps; sampling breaks time into steps. The two operations are separate and can be applied in either order

As an example, the 100 MHz Tektronix TDS3012 has 2^9 (512) quantising levels, which means that if a voltage falls between two levels, there will be a maximum error of $\pm 0.1\%$. Given that

$\pm 0.1\%$ quantisation error is trivial compared to the $\pm 2\%$ analogue errors in the input attenuator and amplifier section, you might feel that using nine bits is extravagant in terms of ADC and memory cost. There are two justifications for this bit depth:

- 512 quantising levels map ideally (with a small overlap) onto a standard LCD screen having 480 vertical pixels, so this is a worthwhile advantage.
- Because the quantising errors are so small, a trace can be expanded by a factor of ten **after** capture, allowing investigation of details without quantising errors becoming too noticeable.

We have seen that quantisation (voltage) errors are almost insignificant. Unfortunately, sampling (time) errors are far more of a problem...

Nyquist and equivalent sample rate

The Nyquist criterion states that the original waveform can be correctly reconstructed provided that the sampling frequency is at least double the highest frequency to be sampled. Practical considerations increase this frequency slightly, so CD uses 44.1 kHz sampling frequency even though its audio bandwidth is limited to 20 kHz. This simple theory implies that a 100 MHz oscilloscope requires an ADC that samples at 200 MS/s (megasamples per second), yet the HP54600B can only sample at 20 MS/s, whereas the TDS3012 can sample at 1.25 GS/s. Why do these two 100 MHz oscilloscopes have such wildly different sample rates?

Equivalent sample rate

The HP uses a technique known as **equivalent sample rate** (called **repetitive sample rate** by LeCroy), and it works like this: We assume that the input signal is repetitive, perhaps a 100 MHz sine wave. Sampling at 20 MS/s means that we

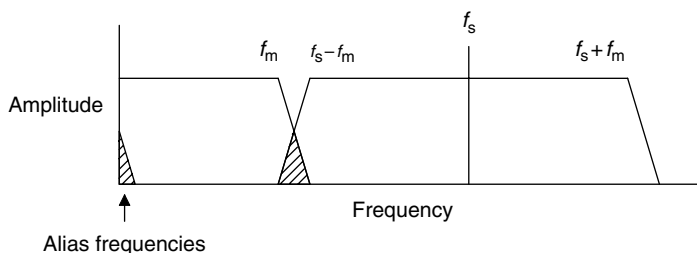
sample at 50 ns intervals, but an entire cycle of 100 MHz only lasts for 10 ns, so we miss four entire cycles completely and sample a single point on the fifth cycle. This doesn't sound very useful, but if we were to delay sampling after the next trigger point by a time equivalent to one hundredth of the screen width, our next sample would be at a slightly later point on the waveform, and would therefore plot a different voltage. If we keep on incrementing the sampling time delay after trigger by one-hundredth screen width intervals, we eventually build up an entire waveform that appears to have been sampled at a much higher rate. Thus, if we had set our time base to 2 ns/div (to display two cycles of 100 MHz), the sweep would display 20 ns, and if we had set our delay increments to give one hundred points, we would appear to be sampling every 0.2 ns, giving an equivalent sample rate of 5 GS/s. Very clever, yes?

The problem with equivalent sample rate is in the assumption that the input waveform is repetitive. A general rule of thumb for oscilloscopes is that you need at least ten points in a cycle to give a reasonable approximation to the shape of that waveform. But if your waveform is a one-off pulse, the ADC needs to take a genuine ten samples in one cycle, so a sample rate of 1 GS/s is needed to give a reasonable approximation of a 100 MHz sine wave in one-shot mode. Thus, the TDS3012 **does** have a one-shot bandwidth of 100 MHz, but the 20 MS/s sample rate of the (much earlier) HP54600B means that it has a one-shot bandwidth of only 2 MHz.

Contravening Nyquist and aliasing

If we apply a frequency (f_m) that is higher than half the sample frequency (f_s) to an ADC, the frequency of the input signal will be misinterpreted as a low frequency, and this phenomenon is known as **aliasing** (see [Figure 4.21](#)).

When digital audio is recorded, the ADC is preceded by a low-pass **anti-aliasing filter** set to slightly less than half the sampling

**Figure 4.21**

Once $f_m > f_s - f_m$, alias frequencies begin to crawl up from 0 Hz

frequency. In this way, aliasing is eliminated, but as we will see in a moment, digital oscilloscopes **cannot** incorporate an anti-aliasing filter, so they are susceptible to this problem.

Sample rate, record length, and time base speed

Remember that a digital oscilloscope samples the incoming signal and writes the results into memory. Later, the memory is read to the display at a convenient rate. Suppose that we want to display two cycles of 100 MHz, and we have chosen an ADC that can run at 1 GS/s. Each cycle generates ten points and we have two of them, so we generate twenty points across the screen. Now suppose that we want to display two cycles of 50 Hz. Each cycle lasts 20 ms, so the total sweep lasts for 40 ms. In that time, sampling at 1 GS/s, we generate 40 000 000 points, requiring a lot of memory. Worse, we write to that memory at 1 GHz. Unlike computer memory, we don't need to be able to access data points in this memory in a totally random fashion, so there are various fudges, collectively known as **demultiplexing**, that allow the use of slower memory. Nevertheless, oscilloscope memory is expensive, and in 2003, upgrading a LeCroy 8600A 6 GHz 10 GS/s 4ch oscilloscope from its standard 2 Mpt memory to 50 Mpt increased the price of the oscilloscope by 57%.

We ideally need an entire screen of memory either side of the screen, which would enable a feature at one extreme side of the

screen to be moved to the other extreme without bringing a blank space onto the screen. In addition, it would be nice to have enough samples in memory to allow the trace to be expanded without generating visible gaps on the display. If we had 10 000 samples allocated across the screen, a basic computer display 640 pixels wide would allow us to expand by a factor of ten without producing visible gaps. To allow a screen either side, we would need to store a total **record length** of 30 000 samples, of which 10 000 would occupy the screen at any one time (see Figure 4.22).

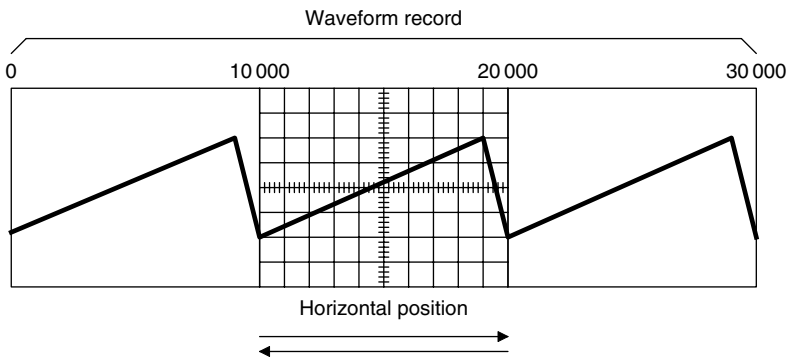


Figure 4.22

The horizontal position control moves the display window back and forth across the waveform record

The screen has ten divisions, so we can find the maximum permissible sample rate using:

$$\text{Maximum permissible sample rate} = \frac{\text{visible record length}}{10 \times \text{time/div}}$$

Note that sample rate is now **totally dependent** on the amount of memory, so if we set the time base to $10 \mu\text{s/div}$, our example oscilloscope having a visible record length of 10 000 points would be forced to drop its sample rate from 1 GS/s to 100 MS/s . This isn't a problem, because even if we required

ten points per cycle, this sample rate and time base setting could correctly capture one thousand cycles at 10 MHz, and we would never be able to distinguish this number of cycles across the screen, so the display would be **identical to an analogue oscilloscope**. In practice, the argument is complicated by the fact that we only display 640 of the 10 000 points in memory, but because we have correct data in memory, tricks can be applied to greatly reduce the inevitable aliasing caused by the re-sampling needed to accommodate the display.

It might seem odd to choose 10 000 points across the screen when we know that we will have to re-sample to 640 pixels, but 10 000 produces nice round numbers for the sample rates generated by the 1, 2, 5 sequence of the time base.

Because sample rate varies with time base setting, we cannot precede the ADC with an anti-alias filter, and this means that aliasing is always a possibility. As an example, analogue video contains high and low frequencies simultaneously, and an oscilloscope set to 10 $\mu\text{s}/\text{div}$ will display almost two television lines which contain significant energy at 4.43 MHz (3.58 MHz USA), and although the TDS3012 can display this with insignificant aliasing, the short record length of the HP54600B causes sufficient aliasing to render the display almost unusable.

Summarising; at slow time base speeds, it is record length that determines sample rate and you can't have too much memory. Be warned that because oscilloscopes access their memory at such high speeds, their data busses are transmission lines and extra memory cannot simply be added in the manner of a PC, so the choice has to be right at the moment of purchase.

Features unique to digital oscilloscopes

We have found that a high sample rate is crucial, but that it can only be maintained at slow sweep speeds by having

sufficient record length, and that analogue bandwidth becomes a secondary consideration. Worse, because digital oscilloscopes are really computers with knobs on, they obey Moore's law (speed doubles every 18 months), and depreciate almost as fast as a PC, so any oscilloscope with a maximum sample rate of less than 100 MS/s (appropriate for 20 MHz analogue bandwidth) is effectively junk. Despite these caveats, digital oscilloscopes have many features that simply **cannot** be provided in an analogue oscilloscope, so it is time to look at their advantages.

Negative time

Because digital oscilloscopes acquire, store, and display the input signal continuously, they can show events **before** the trigger point, allowing them to look backwards in time. By definition, this is impossible for an analogue oscilloscope.

Usable slow time base settings

At very slow sweep speeds, an analogue CRT no longer produces a graph, but a decaying swept spot, making the display very difficult to interpret. At these slow sweep speeds, digital oscilloscopes stop triggering, and free-run (known as **roll mode**), which produces a non-decaying display akin to an analogue chart recorder, allowing sweep speeds of 10 s/div, which is very useful for observing events like heater warm-up times.

Capturing and displaying infrequent events

An analogue oscilloscope produces trace brightness proportional to the time spent by the beam sweeping overlaid traces per second. If the oscilloscope spends most of its time waiting to sweep, the trace dims. By contrast, even a single trace written into digital memory can be written to the display repetitively to give a bright trace.

Colour

The colour of the trace on an analogue oscilloscope is determined by the required persistence of the phosphor screen coating, so the trace produced by Ch1 is the same colour as that produced by Ch2. Digital oscilloscopes can use television-type displays (CRT or LCD), allowing the two traces to be different colours, which means that instead of using the top half of the screen for Ch1 and the lower for Ch2, both traces can occupy the full height of the screen because they are easily identified.

The previous point is more important than at first appears. The screen's "Y" axis corresponds directly to the codes ranging from 0 to 512 (9 bit ADC) leaving the ADC. If we only use half the screen for one channel (0 to 256), we have effectively thrown away one bit of ADC resolution, so we must always use the full height of the screen to maximise accuracy. Of course, the same argument could be levelled at an analogue oscilloscope, but the significance is that reduced trace height degrades the accuracy of automated measurements. If we have two waveforms occupying the full height of the screen, they can only be easily distinguished if they are of different colours, so for a multi-channel digital oscilloscope, a colour display is a necessity rather than a luxury.

Storing and exporting traces

The digital advantages mentioned previously pale into insignificance compared with the ability to store traces for later comparison. When fettling an amplifier and testing the effect of changing a component, it is extremely useful to be able to compare "before" and "after". Most digital oscilloscopes can export traces, but some require an add-on module, plus proprietary software on the PC, and perhaps a specific card in the PC. This all costs money, and data transfer requires the oscilloscope to be connected to the PC, which is a nuisance. The ideal oscilloscope has a 3.5" floppy drive using PC

formatted diskettes so that you can simply store a trace and take the diskette to any PC at your leisure. This requirement greatly reduces the value of older digital oscilloscopes, so if you are considering buying one, check very carefully how it exports traces.

A very powerful oscilloscope may encounter a somewhat different problem in that it has such a large record length that its trace files are simply too large to fit on a floppy disk. These oscilloscopes must transfer their data via a network port. At the time of writing, amateurs are unlikely to encounter this problem because such oscilloscopes are still quite expensive, but prices are falling all the time.

Peak detect/data compaction and glitch detection

At slow time base speeds, finite record length forces the oscilloscope to lower the rate at which it writes to memory, so it may miss an important event (perhaps a glitch) that occurs between samples. **Peak detect** mode captures such transients by running the ADC at its maximum sample rate, but only records the highest (or lowest) voltage and drops this result into the nearest available memory location (see [Figure 4.23](#)).

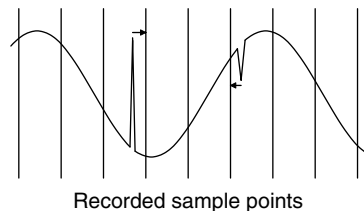


Figure 4.23

When sample rate falls, a fast event can fall **between** recorded samples, so peak detect detects such events and records them at the nearest memory location

The fact that such events are recorded with a slight time error and gross distortion is not nearly as serious as missing the event altogether. Peak detect alerts the user that a more detailed investigation is necessary and is invaluable for finding the cause

of otherwise insoluble faults, such as those caused by random spikes on the output of a faulty switched mode power supply.

A slightly more sophisticated alternative called **data compaction** by LeCroy is to measure the highest **and** lowest voltage, and record both at the nearest available memory location. Since two items of data are recorded, rather than one, this approach immediately doubles the amount of memory required for a specified record length.

Averaging and noise

Because a digital oscilloscope effectively stores the data resulting from each sweep as a database, it can perform mathematical calculations such as **averaging**. The oscilloscope displays a repetitive waveform, so each sweep across the screen draws a trace identical to the one before it. If we average across traces, the repetitive element stays the same, but any random noise on the traces averages to zero. Averaging is a very powerful method of extracting signals seemingly buried in noise. Typically, averaging can be adjusted in binary logarithmic steps (2, 4, 8, etc.) from 2 to 512 traces. The reduction in noise is proportional to $\sqrt{n}$ so averaging over 512 traces reduces noise by 27 dB, but at the expense of drastically slowing the time taken for the display to stabilise. If the oscilloscope is triggered from the source of crosstalk, averaging is very useful for detecting crosstalk from one channel into the power supply or to another audio channel.

Envelope

Envelope mode never deletes a sweep, so successive sweeps are displayed over one another in the manner of a multiple exposure photograph. Although the display may become completely distorted, making it impossible to discern a recognisable waveform, it enables maximum voltage excursions to be measured easily, and this is particularly useful for determining how much overload capability is required in a given stage. We simply apply music for as long as we like and measure the maximum

vertical excursion. Conversely, data communications engineers use envelope for measuring time jitter (horizontal excursion) on their data waveforms.

Displaying a sweep indefinitely is not always ideal, so envelope mode is often adjustable in 2, 4, 8 steps to limit the number of sweeps for which the results of a given sweep are stored and displayed, and this enables the appearance or disappearance of infrequent pulses to be monitored.

Adjustable persistence

All displays (and the human eye) have a property known as persistence, whereby a gradually decaying visible image remains even though its cause has gone. This might seem to be a defect, but it can be quite useful because it assists in seeing momentary events. Persistence in an analogue oscilloscope is pre-determined by the choice of display phosphors and is unchangeable.

Now that memory has fallen in price, digital oscilloscopes can mimic persistence electronically and controllably. The way that this is done is to treat the array of display pixels as a 3D database. Early digital oscilloscopes simply placed a “1” or a “0” in each pixel on each trace. This meant that infrequent events were displayed with exactly the same brightness as frequent events, leading to a confusingly cluttered display. However, if a 4-bit number is stored at each pixel, it can be incremented by one each time the waveform hits that pixel (perhaps from 0100 to 0101). If a calculation is applied simultaneously that gently decrements that number over time, pixel brightness can be made proportional to that number, and we have recovered the analogue advantage of persistence. The technique is known as **digital phosphors** and allows a digital oscilloscope with sufficient sample rate to produce a display almost indistinguishable from a top-quality analogue oscilloscope, but with all the digital advantages. This truly is the best of both worlds, and was the deciding factor in the author’s purchase of a TDS3032.

The key difference between persistence and envelope is that envelope switches pixels off abruptly whereas persistence exponentially reduces the brightness of older pixels (see [Figure 4.24](#)).

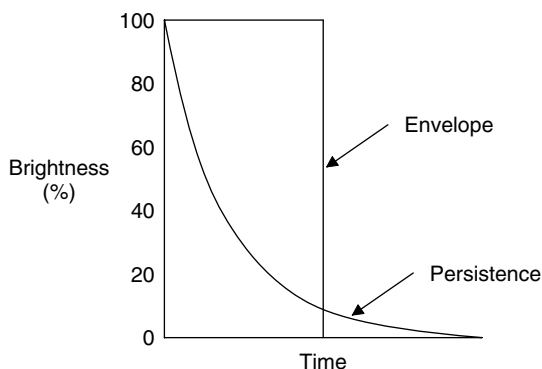


Figure 4.24

Persistence (exponential decay) versus envelope (abrupt switch off) brightness

Dots and vectors

Because the oscilloscope samples and quantises, it plots dots on a screen of graph paper. It is conventional to join the dots on graph paper, and oscilloscopes can also perform this function. Some oscilloscopes offer the user a choice of the way in which the journey is made (straight line, or $\sin(x)/x$ curve) from one point to the next; so this feature is usually known as **vectors**. However, there is an implicit assumption in joining the dots that the curve really did move smoothly from one dot to the next. Finite sample rate means that this assumption may not be true, particularly at slow time bases, so most oscilloscopes can be switched to dots only.

Automated measurements

A display on an oscilloscope is likely to contain one or two cycles of a waveform. Each cycle is thus a histogram composed of many vertical bars, and measurements or calculations that involve maxima, minima, or areas under one cycle of the graph can easily be made. Provided that the oscilloscope can display the waveform

without distortion, a digital oscilloscope can calculate V_{RMS} or V_{mean} , etc. for **any** waveform, and accuracy is limited only by the number of samples taken per cycle and the number of quantising bits used. Normally, as many cycles as are available in the waveform record are used for calculation, but cycle RMS and cycle mean are found by considering only the first cycle after the trigger. The great advantage of automated measurements is that they are live and update as the waveform changes.

Fast Fourier Transform (FFT)

The FFT is a mathematical tool that allows data captured in the time domain to be displayed in the frequency domain. Put simply, although the vertical axis is still voltage, it is now plotted against frequency, rather than time, and the oscilloscope has been converted into a spectrum analyser. This is an extremely useful facility for investigating distortion harmonics (see [Figure 4.25](#)).

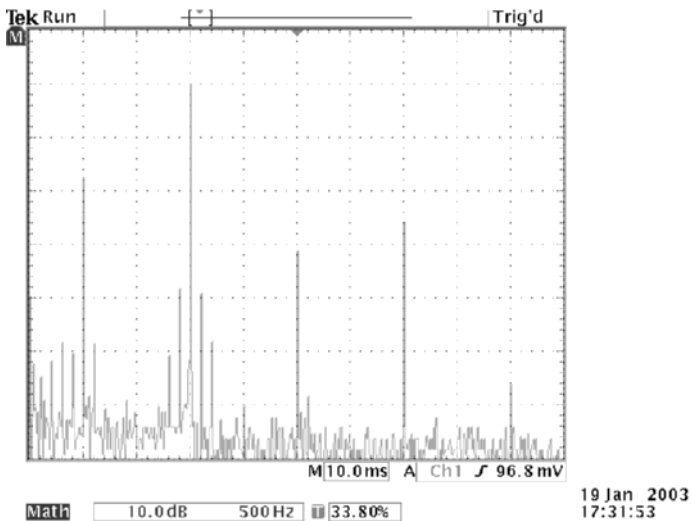


Figure 4.25

This FFT shows the amplitudes of the 2nd to 6th harmonics of a push–pull output stage. As expected, the odd harmonics are higher amplitude than even, but more interestingly, there are 50 Hz spikes around the 2nd and 3rd harmonics caused by intermodulation with a poor power supply

The FFT is an immensely powerful tool, but it has limitations...

In converting from time to the frequency domain, the mathematics of the FFT make the assumption that the waveform to be analysed repeats itself periodically. This assumption may seem trivial, but has **major** repercussions.

If we captured a **single** cycle of the waveform, and drew it around a circular drum (like a seismograph) so that the end of the cycle just met the beginning, then by rotating the drum we could replay the waveform *ad finitum* and reproduce our original signal. Unfortunately, any uncertainty as to the precise timing of the end of the cycle causes a step in level when we attempt to loop the recorded cycle back to itself on replay. However, if we capture more cycles on our drum, this glitch occurs proportionately less frequently and causes less of an error. Thus, capturing a thousand cycles reduces the error by a factor of a thousand, at the cost of needing one thousand times the record length, making an oscilloscope's record length even more important if FFT is to be contemplated.

The other way of reducing the glitch is to force periodicity by applying a **window** to the waveform record. In this context, a window is a variable weighting factor that multiplies the values of the samples at the ends of the waveform record by zero, but applies a greater weighting (≤ 1) to samples towards the middle. Since any number multiplied by zero is zero this forces the end samples to zero, and allows the waveform record to be repeated without glitches (see [Figure 4.26](#)).

Because windowing distorts the waveform record, it must distort the results of the FFT calculated from that record. Windowing either spills energy from high-amplitude bins into adjacent bins, producing skirts around frequencies having high amplitude, or it changes bin amplitudes. (Because the process

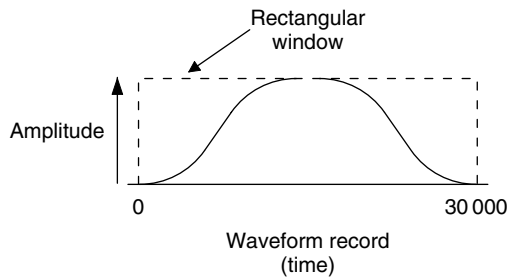


Figure 4.26

Shaping the window reduces the effects of periodicity violation compared to a rectangular window

of sampling broke time into discrete slices, the results of an FFT must produce frequencies in discrete slices, and these are known as **bins**.) All windows are therefore a compromise between frequency and amplitude resolution.

A window that does not modify sample values is known as a rectangular window (because it multiplies by a constant value of 1 over the entire waveform record). Because the rectangular window does not modify sample values, it does not cause spreading between bins, and it offers the best frequency resolution. Unfortunately, amplitudes are likely to be in error because of periodicity violation. Conversely, the Blackman-Harris window modifies the ends of the waveform record to avoid periodicity violation, which causes spreading between bins, but improves amplitude resolution. In the same way that averaging reduced noise on a displayed waveform, averaging can significantly improve the signal to noise of an FFT display, but with the same cost of reduced measurement speed.

Connecting to the oscilloscope

Having acquired our oscilloscope, we need a means of connecting it to the circuit we want to test.

Input capacitance and voltage dividing probes

The most sensitive setting on a typical oscilloscope is 2 mV/div and the standard input resistance is 1 M Ω . This is almost sensitive enough for a microphone, so we must use a screened lead to avoid picking up hum from all the (nearby) mains wiring. Typical coaxial screened lead has a capacitance of ≈ 100 pF/m; perhaps we have bought a 20 MHz oscilloscope, and use a 1 m lead having 100 pF of capacitance to connect the input of the oscilloscope to the output of a common cathode gain stage using a high μ triode such as an ECC83 or 6SL7. Under typical operating conditions, this stage is likely to have an output resistance of 60 k Ω . Unfortunately, the capacitance of the lead and output resistance of the amplifier forms a low-pass filter having a -3 dB frequency of:

$$f = \frac{1}{2\pi CR} = \frac{1}{2 \times 3.14 \times 100 \times 10^{-12} \times 60 \times 10^3} \approx 26.5 \text{ kHz}$$

We cannot make useful audio measurements with this filter in the way. We cannot change the output resistance of the amplifier, so we must reduce the capacitance.

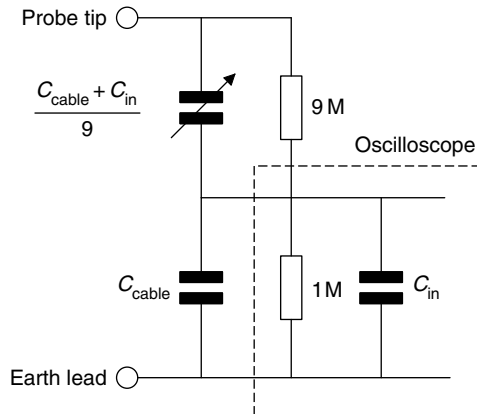
One way of reducing the capacitance would be to cut the lead shorter, perhaps to 10 cm, which would reduce the capacitance by a factor of ten to 10 pF, and would move the filter frequency to 265 kHz, which is safely out of the audio range. The width of one line of text in this book is approximately 10 cm, so it would be quite difficult to work with a lead this short.

Even worse, the oscilloscope itself has input capacitance, typically 10–30 pF, and the precise value is usually stated on the front of the oscilloscope next to the input BNC (see [Figure 4.27](#)).

**Figure 4.27**

Most oscilloscopes state their input loading next to the BNC input socket

The solution to the capacitance problem is to make a paralleled pair of potential dividers, one resistive, one capacitive (see [Figure 4.28](#)).

**Figure 4.28**

An oscilloscope probe is essentially two potential dividers in parallel. The complete circuit is contributed partly by the oscilloscope, partly by the “probe”

The $1\text{ M}\Omega$ resistor is the standard oscilloscope input resistance, and this is in parallel with the oscilloscope’s input capacitance and the lead capacitance. Together with the

1 M Ω resistor, the additional 9 M Ω series resistor forms a 10:1 potential divider. The cunning bit is the addition of the capacitor across the 9 M Ω resistor. There are two ways of looking at this capacitor:

- Potential dividers: The 9 M Ω /1 M Ω resistive potential divider has a loss of 10:1. We can also make potential dividers from capacitors, but because a capacitor's reactance is inversely proportional to its capacitance, a 10:1 capacitive divider requires the upper capacitor to be one-ninth the value of the lower capacitance.
- Time constants: We can view the upper resistor/capacitor combination as a high-pass time constant, and the lower resistor/capacitance combination as a low-pass time constant. If the time constants are equal, their filtering effects cancel out and the network has a flat frequency response.

Looking into the input of the circuit, at DC the capacitors are an open circuit, so we see the two resistors in series, giving an input resistance of 10 M Ω . At high frequencies, the capacitors are short circuits, so we can neglect the effect of the resistors, and we see a pair of capacitors in series:

$$C_{\text{series}} = \frac{C_1 C_2}{C_1 + C_2} = \frac{13.33 \text{ pF} \times 120 \text{ pF}}{13.33 \text{ pF} + 120 \text{ pF}} = 12 \text{ pF}$$

At the expense of oscilloscope sensitivity, and by adding a resistor and capacitor at the far end of our lead, we have reduced the input capacitance by a factor of ten from 120 to 12 pF. **Oscilloscope probes** always state their input capacitance, either in the data sheet that came with the probe, or on their connector (see [Figure 4.29](#)).

Because different oscilloscopes have different input capacitances, we need to be able to adjust the upper capacitor to suit the oscilloscope. This adjustment is often a small screw in the side of the probe (see [Figure 4.30](#)).



Figure 4.29

Some probes state their essential data on the connector housing



Figure 4.30

Probes are often equalised by a small screw in the body

The adjustment is set by connecting the probe to the 1 kHz square wave provided by all oscilloscopes on their front panel, and often called **Cal**. The oscilloscope is set to display one cycle of the square wave, and the capacitor adjusted to give the flattest, squarest, leading edge (see [Figure 4.31](#)).

Other probes

Although the **passive** voltage-dividing probe reduced input capacitance from 120 pF to 12 pF, it is possible to do better. **Active** voltage probes incorporate an amplifier at the probe tip and can reduce their input capacitance to ≈ 2 pF, and because they incorporate an amplifier, no sensitivity is lost.

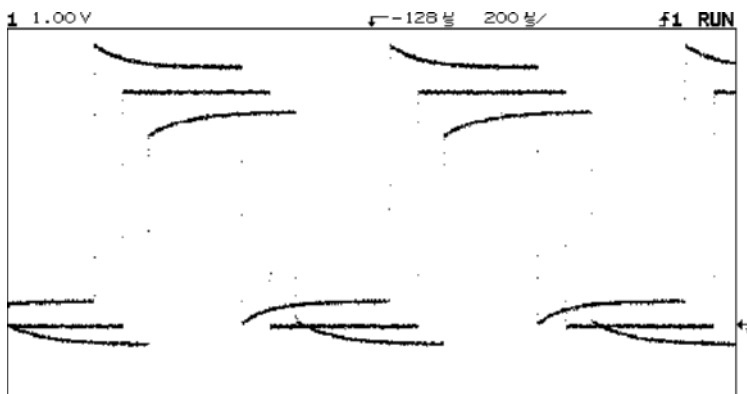


Figure 4.31

Correct probe equalisation produces a square leading edge (middle trace), rather than overshoot or rounding

Current probes measure the current in a wire indirectly by measuring the magnetic field it produces, allowing current to be determined without breaking the circuit. They achieve this by directing the magnetic field through a coil of wire forming the secondary of a transformer, or they sense the field with a semiconductor **Hall effect** device. The passive transformer type cannot sense DC, but the active Hall effect type can sense DC at the cost of an amplifier and power.

Universal passive $\times 10$ probes are cheaply available for oscilloscopes with <250 MHz bandwidth. Faster oscilloscopes need passive probes specifically designed for them if their bandwidth is to be maintained to the probe tip, and such probes are typically ten times the price of a universal probe. Even worse, active probes are typically fifty times the price of a universal passive probe, and can be as much as two hundred and fifty times the price! Moral: Look after your probes – the author keeps his probes in a plastic lunch box when they're not in use.

Oscilloscopes having automated measurement systems need to know that the signal has been applied via a probe. More

sophisticated oscilloscopes have connectors/pins near the signal connector to set the appropriate scaling automatically, but some need to be told manually that a probe has been connected. Really old oscilloscopes require **you** to do the thinking and to remember to multiply the volts/div appropriately.

Transmission lines and terminations

The final way of connecting to an oscilloscope is via a $50\ \Omega$ RF **transmission line**, terminated by $50\ \Omega$ at either end, and this is why faster oscilloscopes have a $50\ \Omega$ switch on their input coupling. If you accidentally operate this switch with a $\times 10$ probe connected, your signal will disappear. The other hazard with the $50\ \Omega$ **termination** is that the resistor has a very low power rating, so it can easily be burnt out if a large signal is accidentally applied directly to it.

Fortunately, transmission line effects do not become apparent until the cable length is a significant proportion of the length of one wavelength of that frequency travelling down the cable. Since the velocity of propagation down most cables is $\frac{2}{3}$ the speed of light ($c \approx 3 \times 10^8$ m/s), this means that one wavelength at 20 kHz occupies 10 km of cable. We can completely ignore transmission line effects in analogue audio.

Oscillators and dedicated audio test sets

If you have an oscilloscope, you need an oscillator, but an oscillator is also **very** useful for supplying an external source of AC to a component bridge so that components can be tested at different frequencies. Air cored inductors are more easily measured at 20 kHz than at 1 kHz ($X_L = 2\pi fL$), whereas the primary inductance of an output transformer should be measured at 20 Hz.

Traditional (Wien) oscillators

All analogue oscillators contain three essential blocks:

- An amplifier with positive feedback
- A frequency-selective network
- An amplitude stabilisation circuit.

The amplifier could use valves, transistors, or integrated circuits. The frequency selective network needs to be tuneable, and since no network can be continuously tuned over the entire 20 Hz–20 kHz range, it is conventional to break that range into switched decades (20–200 Hz, 200 Hz–2 kHz, 2–20 kHz), and continuously tune over the resulting 10:1 ranges.

Traditionally, the most popular frequency selective network was the Wien network (see [Figure 4.32](#)).

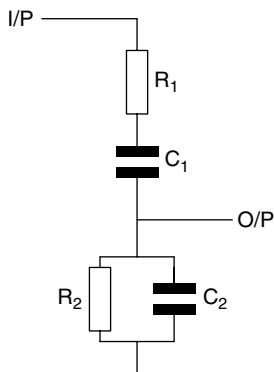


Figure 4.32

The Wien frequency selective network

The frequency of oscillation is given by:

$$f = \frac{1}{2\pi\sqrt{R_1 R_2 C_1 C_2}}$$

In order to change frequency, we only need to change the value of one component. The loss of the network at resonance is:

$$\frac{V_{\text{out}}}{V_{\text{in}}} = \frac{1}{1 + \frac{C_2}{C_1} + \frac{R_1}{R_2}}$$

The significance of this equation is that for the oscillator to maintain constant amplitude as we tune from one end of the range to the other, we must actually change a pair of values simultaneously, either the resistors must change together, or the capacitors must change together. Further, when we gang a pair of variable resistors or capacitors together, any tracking error causes the level to change with frequency. This deviation from perfection is known as **range flatness**. In general, it is easier to make tracking variable air-spaced capacitors than tracked variable resistors, so Wien oscillators usually use a dual-gang variable capacitor, but range flatness $< \pm 0.1$ dB is difficult to achieve. Variable capacitors can only have 180° of rotation, so if you see an oscillator having a 180° scale (rather than the 270° achievable by a variable resistor), you can be pretty certain that you are looking at a Wien oscillator.

Although the Wien network can be connected around an amplifier, sustained oscillation requires that the amplifier's gain is only **just** sufficient to overcome the losses in the frequency selective network. If the Wien network uses paired components, it has a loss of $\frac{1}{3}$, so the amplifier must have a gain of precisely 3. Although it is possible to trim the gain of the amplifier in the laboratory to produce oscillation, it soon drifts and oscillation stops. What is needed is a means of stabilising the gain. The seminal Hewlett-Packard HP200 oscillator showed the way by using the change of resistance with temperature (and therefore applied voltage) of a fine tungsten filament isolated from external influences by placing it in a hard vacuum. In other words, it is used as a lightbulb to stabilise

amplitude. At low frequencies, the temperature of the filament begins to track the waveform, increasing distortion, so thermally stabilised Wien oscillators rarely produce frequencies lower than 20 Hz.

Summing up, a traditional Wien oscillator is likely to produce sine waves from 20 Hz–20 kHz in three ranges with a typical range flatness of ± 0.2 dB. Distortion is mainly dependent on the quality of the amplifier and can be made to be very low, although typical bench oscillators produce distortion ranging from 0.5 to 0.05% at 1 kHz, and rising at low frequencies. Modern low-distortion oscillators tend to be based on the state variable filter.

Function generators

The Wien oscillator can produce a very low distortion sine wave, but an electronics laboratory often needs square waves or pulses. An alternative way of producing oscillations is to charge a capacitor from a constant current source (producing a rising ramp) until it reaches an upper threshold voltage, then immediately start discharging it with an equal and opposite constant current source (producing a falling ramp). When the falling ramp reaches a lower threshold voltage, we start the cycle again. The oscillator thus produces a triangular wave.

The oscillator is very versatile because not only does it produce the triangular wave, but it also produces a square wave and can produce a sine wave.

Because the sine wave is derived from a triangle wave, it always contains significant distortion, and 1% THD is typical. The distortion can often be seen at the tip of the sine wave, where a Norman arch appears rather than a soft arch (see [Figure 4.33](#)).

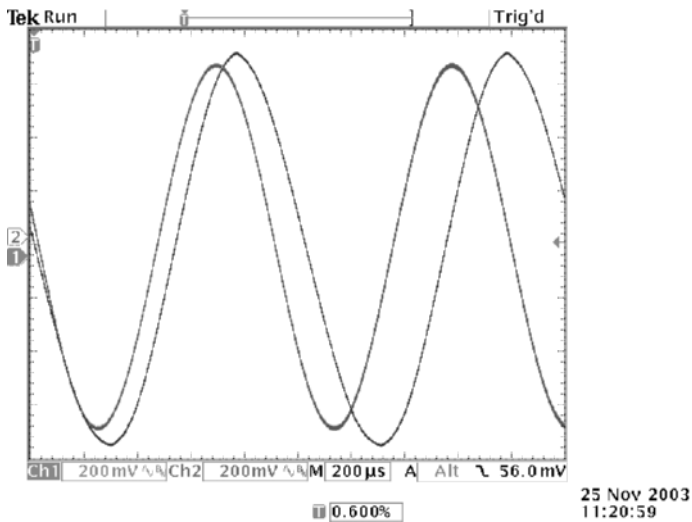


Figure 4.33

Comparison of low distortion sine wave (lower amplitude) with sine wave generated by function generator (higher amplitude). Note the Norman arch produced by the function generator

You might wonder why we should waste any time considering function generators for audio when they have such poor distortion, but even the very cheapest function generator has two very valuable qualities for audio:

- They produce good quality square waves (sharp edges and negligible ringing) that are good for testing amplifier stability.
- They do not require ganged controls to change frequency, so their range flatness is intrinsically perfect, making them excellent for critically measuring amplitude against frequency response of amplifiers.

As an example of the second point, the author was unable to detect any range flatness errors on his cheap function generator even though his dedicated audio test set can indicate 0.02 dB errors clearly.

Dedicated audio test sets

Neither a function generator nor a typical bench oscillator is ideal for audio testing. Fortunately, dedicated audio test sets are available second-hand.

An audio test set should contain at the very least:

- A fairly low distortion ($<0.05\%$) sine wave oscillator (typically 20 Hz–20 kHz)
- A wide-band meter calibrated for reading sine waves
- A simple form of total harmonic distortion (THD) measurement
- A peak programme meter (**PPM**) for measuring noise.

Better test sets may include noise and/or wow and flutter (W&F) measurements. However, the main advantage of audio test sets is that they have meters scaled directly in dB, and can therefore measure audio frequency responses quickly and precisely.

When we looked at AC measurements, we saw that there were various options for specifying an AC waveform:

- Rectify the waveform, apply the resulting DC to a mean reading meter, and calibrate it to V_{RMS} of sine wave. This rectifier is ideal for driving an expanded scale having a range of only 1 dB because it averages noise to zero, improving accuracy, but it is only suitable for measuring sine waves.
- Use a rectifier that produces an output that genuinely is proportional to V_{RMS} and apply it to the meter. This rectifier is appropriate for power and distortion measurements.
- Use a rectifier that captures V_{pk} (both positive and negative), and uses a meter with **ballistics** that not only display short peaks accurately (fast **attack**), but allows them to be read (slow **decay**). This approach is exemplified by the PPM and is appropriate for music and noise.

Because different measurements require different rectifiers, dedicated audio test sets incorporate all three rectifier types and they are selected as appropriate. Thus, a test set might have a moving coil meter having a sufficiently fast response time to meet the PPM attack specification (the slow decay is produced electronically) and driven by a peak-reading rectifier. To measure sine wave amplitude, a mean reading rectifier would be substituted, the meter ballistics would be slugged electronically, and the scale might be expanded electronically to make measurement easier. Finally, to measure distortion, the expanded scale and slugged ballistics would be retained, but an RMS rectifier would be substituted because this sums harmonic powers correctly.

For simple amplitude measurements, the meter section is simply a meter/rectifier preceded by a calibrated amplifier. To measure harmonic distortion, a filter must be added to reject the fundamental. We want to measure the level of the harmonics, but must reject the fundamental without affecting the level of the 2nd harmonic, which is one octave higher in frequency than the fundamental. This can be done in one of two ways:

- Use a high-pass filter. Since the 2nd harmonic is an octave higher than the fundamental, if we want to measure THD to 0.1% (−60 dB), we need a filter with a slope of 60 dB/oct, or more. This is achievable, but not easy, and such a filter is unlikely to be tunable for different frequencies, so this approach leads to a test set that can only measure distortion at one or two fixed frequencies.
- Use a notch filter. Active notch filters can easily achieve rejection of 80 dB, or more, at their notch frequency, but need precise tuning to achieve maximum rejection. Once auto-tuning has been added to maximise rejection, it is a small step to make the filter tuneable over a wide range, and this more modern approach leads to a test set that can measure distortion at any frequency, probably to better than 0.01% (−80 dB).

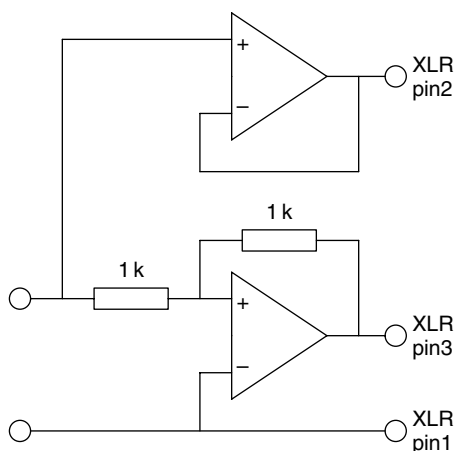
In the UK, there is a variety of test sets (mostly ex-BBC) available at prices attractive to the amateur, but it should be remembered that there is almost always a reason for test equipment being cheap:

- BBC ATM1 plus TS10: Valve based, and both are usually lethally packaged. Don't even think about them. The Wien oscillator Tone Source (TS10) has poor distortion, The Audio Test Meter (ATM1) has superb attenuators and the (separate) mean/flat meter has very low stiction, but the amplifiers are noisy, and it can't even measure THD without external assistance. The PPM movement has incorrect ballistics, and PPM1 to PPM2 is 6 dB, rather than 4 dB.
- BBC EP14/1: This was the first IC-based BBC test set, and includes a PPM for noise measurement. It is a true piece of laboratory equipment; using the expanded scale enables repeatable measurements to an accuracy of 0.05 dB. The THD measurements down to $\approx 0.1\%$ can be made at 100 Hz and 1 kHz but a mean rectifier is used, rather than RMS. The Wien oscillator has somewhat a poorer distortion than the meter.
- Ferrograph ATS1: An idiosyncratic, but very versatile, piece of test gear. Designed (predictably, as Ferrograph made some quite nice tape machines) for comprehensive testing of tape machines, it also includes W&F measurement, but not a PPM. Oddly, it tends to be quite a lot more expensive than the BBC alternatives.
- BBC ME2/5: A "cooking" piece of test gear designed to replace the EP14/1 in less critical usage. The oscillator is digitally synthesised and can sweep frequency (intended, but never used, for automated music line testing). The meter section is only accurate to 0.1 dB but contains a primitive digital frequency meter. The unit is newer, smaller, lighter, more expensive, less accurate, and less reliable than the EP14/1...
- Technical Projects MJS401D and Neutrik derivatives: A splendid piece of equipment with a wide bandwidth meter section far better than the EP14/1, including THD measurement at **any** frequency, accurate frequency meter, and

comprehensive filters and rectifiers. Check for 1 kHz distortion at +20 dBu; with the 20 Hz–22 kHz filter selected, it should typically better 0.002%. Commendably, Neutrik willingly repairs these test sets, even though they may be considerably more than ten years old. Options may add W&F, IMD. Later versions include a phase meter (very useful), so price may be variable; expect it to cost significantly more than an EP14/1, but perhaps need attention. Thoroughly recommended, but balance is a slightly questionable owing to electronic rather than transformer balancing.

- Lindos LA100: Not (yet) available at amateur prices, this is an excellent piece of semi-automated portable equipment with an LCD display rather than a moving coil meter. It indicates to 0.01 dB, and under optimum conditions can measure distortion (at only five frequencies) to $\approx 0.005\%$. Excellent for fast routine testing of tape machines. Includes everything previously described, and more (except real mechanical PPM and transformer balancing). Can talk to printers and PCs. Super. (Rechargeable batteries are included, but they typically die after five years. When they die, they crash the firmware – so just remove them.)
- Audio Precision System One: An amazing piece of test equipment best used for production testing. Outstanding measurement ability combined with stunningly poor user-friendliness when used for one-off measurements. Needs a PC to drive it, but all its results are therefore available as files. Appalling ergonomics aside, the author would love one.

Professional audio equipment is inevitably balanced, so test equipment is designed to interface with balanced audio. Traditionally, transformers were used, but electronic balancing is now common, which can cause problems when connected to unbalanced equipment. The problem usually occurs at the output of the oscillator. The balanced output is often provided by a pair of unbalanced amplifiers producing signals referenced to earth, one producing one phase, and the other the inverted phase (see [Figure 4.34](#)).

**Figure 4.34**

An electronically balanced output is often simply a pair of op-amps, one inverting, one non-inverting

Because the output is referenced to earth, hum loops can occur. This isn't a problem in a balanced system, but when an unbalanced output is taken, it can cause problems. The scaling on an oscillator's attenuators refers to the voltage between the two balanced conductors. If we only use one output, the output is half the amplitude (-6 dB). It is easy to forget this and calculate gains incorrectly. Equipment using transformer balancing has neither of these problems – if we want an unbalanced source, we simply connect one phase to the earthy input of our amplifier under test.

Most audio test sets have a meter section monitoring output typically called “oscilloscope” or “listen”. This output is extremely useful when measuring noise or distortion. Listening to the character of the noise on headphones can often give clues as to its origin (using a loudspeaker often provokes acoustic feedback).

When measuring distortion, the signal at the monitoring output is the distortion waveform that can be taken to an oscilloscope.

Incorrect bias in a Class AB amplifier stage causes sharp cross-over distortion spikes, so monitoring the distortion waveform can be a very quick way for setting bias correctly.

Even better, if the monitoring output is taken to an oscilloscope or computer soundcard that can perform an FFT, the spectrum of the distortion waveform can be investigated. The author finds the ability to analyse distortion invaluable during the early stages of audio design. As an example, the type 76 triode has a very good reputation for low distortion, so a batch was tested. Pleasingly, they produced ≈ 6 dB less distortion than a typical 6J5 under the same conditions. However, the FFT quickly revealed that unlike the 6J5, the 76's harmonics did not decay quickly (see Figure 4.35).

A cheap, but slower alternative to using the FFT to analyse the distortion waveform is to take the monitoring output to a wave

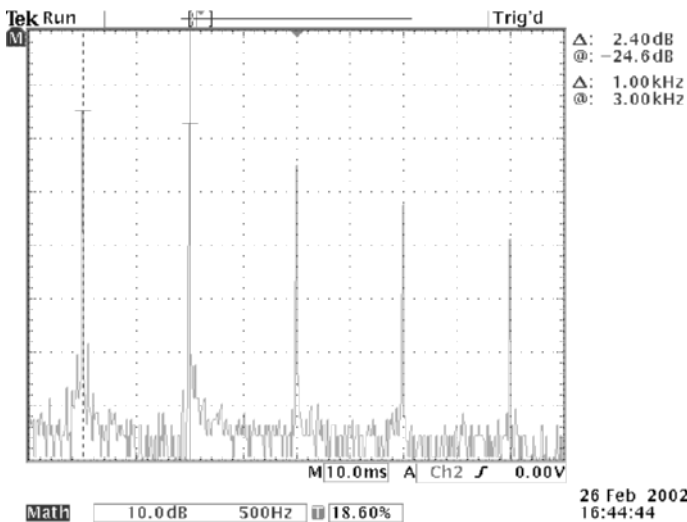


Figure 4.35

This FFT shows the distortion spectrum from 2nd to 6th harmonic of a mu-follower using a 76 as the lower value and triode-strapped D3a as the upper. Note that the levels of the 2nd and 3rd harmonics are comparable

analyser. A wave analyser is effectively a meter preceded by a radio that can be tuned to audio frequencies. Thus, it can be tuned through the distortion harmonics and measure their individual levels. Wave analysers are available second-hand very cheaply, but tend to be quite bulky.

Other test equipment

Most electronics engineers would regard a multimeter, and oscilloscope, and an oscillator to be the absolute minimum required test equipment, but the following is a small selection of other items useful when building and testing valve amplifiers.

Component bridge

A component bridge allows you to measure capacitors and inductors accurately. This is particularly useful when building filters or equalisation networks, and allows you to remove initial component value as a source of error. Component bridges also measure resistance, and are often more accurate at measuring low resistances than even an expensive digital multimeter.

The Marconi TF2700 is an excellent instrument, and you will see it advertised in the electronics magazines for a very reasonable price second-hand. It uses a single PP9 battery, and current consumption is so low that it is not worth the bother of making a mains adapter. It will measure capacitance (0.5 pF–1100 μ F), inductance (0.2 μ H–110 H), and resistance (10 m Ω –11 M Ω) to a basic accuracy of $\pm 1\%$. It can indicate loss factor of capacitors (very useful), and can measure air-cored inductors – digital bridges often can't. The circuit is very simple so it can be fixed easily if it develops a fault.

Even better, the accuracy of the TF2700 can easily be improved. The main range switch uses 0.5% tolerance resistors, most of which can be replaced by 0.1% metal film resistors, but the 10 Ω resistor must be replaced by a 0.1% **non-inductive** wire-wound resistor. The coarse balance switch already uses 0.1% resistors, so no changes are necessary.

The final error comes from the position of the dial on the fine balance variable resistor, and reducing this error takes a little more time. You need four 200 k 0.1% resistors and a 100 k 0.1% resistor. If these are all wired in series, with the 100 k at one end, you can pick off values from 100 to 900 k in 100 k steps. Measure each of these values with the **main range switch set to 10 M**. This forces the coarse balance control to be set to zero, and puts the onus of measurement on the fine balance control. If you now plot a graph of measured value against known value, and draw a line of best fit through it, you will quickly see how much rotation the dial requires to minimise errors.

Assuming that the internal 100 n 0.1% standard capacitor is not in error, the combination of replacing the main range switch resistors and fine dial adjustment usually reduces errors from ± 1 to $\pm 0.25\%$.

Voltstick

These have various brand names and are incredibly useful. They usually look like a fat pen with a white tip. If the tip is near to mains, an internal LED lights (runs off 2 $\times$ AAA batteries for years). They require no contact and are a lifesaver. Always use them before cutting any cable that, potentially, could carry mains. They're very useful for checking the fuse in a mains plug because they can just be waved near the cable. If the voltstick detects mains, the fuse must be intact – and it has been checked without needing a screwdriver to open up the plug and check for continuity.

Continuously variable transformer (variac)

A transformer does not necessarily need a secondary winding. It can perfectly well have a single winding that is tapped part way to produce the lower voltage, this is known as an **autotransformer**. Although an autotransformer is cheap, it does not provide isolation between primary and secondary. If an autotransformer were used to step 240 V to 12 V, but the neutral became disconnected between the mains outlet and the autotransformer, the full 240 V would appear on the 12 V circuit, with possibly fatal results. For this reason, the most common use of an autotransformer is the variable autotransformer, often referred to as a **variac**. A variac is a toroidal autotransformer with the tapping obtained by a rotating wiper that can move from one end of the winding to the other (see [Figure 4.36](#)).

Variacs are commonly used for applying power to equipment gently or for testing tolerance to mains voltage changes. There are various reasons for gently applying power to a device under test (DUT):

- The DUT is newly built and not known to work. Bringing power up gently minimises the smoke in the event of a wiring (or design) error.
- The DUT has not been powered for years, and although it worked once, there is a suspicion that it may have developed a fault.
- The DUT contains old electrolytic capacitors, that if gently re-formed by ramping power up over the course of 30 minutes, will subsequently be fine, but could fail with suddenly applied power.

Because mains voltage is not guaranteed to be exact, variacs are used to test that at the lowest expected voltage:

- Regulators do not drop out of regulation
- Power amplifiers deliver the specified power.



Figure 4.36

A 10 A variac is useful, if rather heavy

And at the highest expected voltage:

- Devices dissipating significant power remain at an acceptable temperature
- Devices with voltage limits such as capacitors or output valves are still within their voltage limits.

In order to provide +10% output voltage, many variacs have a 90% tapping to which the incoming mains is connected. When the wiper reaches the 90% point, full mains is delivered, but as

it sweeps past, the autotransformer steps the voltage up to a maximum of +10% (see [Figure 4.37](#)).

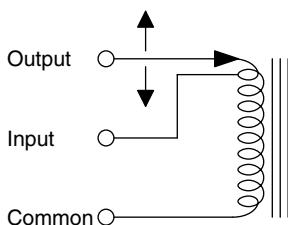


Figure 4.37

Applying the input to a tap part way down the winding allows a variac to develop 0–110% output voltage

A 50 VA transformer could provide 50 V @ 1 A, 25 V @ 2 A, or 10 V @ 5 A, so the gauge of secondary wire is chosen to suit the expected current. Conversely, because a variac has a single winding of wire having constant diameter, the maximum load current is constant, so variacs are rated by their load current, rather than VA.

Variacs are commonly available second-hand, but they are frequently naked, requiring a case to make them safe. New variacs are available neatly cased with a mains outlet, complete with current and voltage monitoring meters, which can be very useful. If you have the choice, a 10 A variac is much more useful (and very little more expensive) than a 2 A version.

Valve tester

Valve testers are useful if you have a very large stock of valves or are a keen designer and want to be able to make measurements enabling you to plot your own curves. However, it should be borne in mind that a universal valve tester inevitably

exposes lethal voltages, so they are intrinsically dangerous, even by the relaxed standards of their day, and should be only approached with great caution. By definition, it is not possible to make a “safe” valve tester – the assumption was always made that they would be operated by people who were aware of all the dangers and were competent to deal with them.

Typical receiving valves have heaters varying from 2.5 V @ 2.5 A (2A3) to 40 V @ 300 mA (PL519), with plenty of variation in between, so a valve tester needs a heater/filament supply to cope with this. One possible solution might use a variable regulated DC supply, but a supply capable of providing 40 V would have to dissipate at least 94 W when supplying 2.5 V @ 2.5 A, so the traditional solution was a tapped transformer. Inevitably, there is considerable wiring between the transformer and heater/filament, and this can cause a significant voltage drop at high currents, so it is well worthwhile to check the actual voltage on the valve’s pins.

Ideally, the supplies required by anodes and screen grids should not change their voltage under load. Typical valve testers can supply up to 400 V at up to 100 mA, and this is just within the bounds of a regulator, but a cunning solution was to use a tapped transformer to apply AC directly to the valve electrodes, and use the valve’s self-rectifying action. The reason for not rectifying the AC directly was that rectification and smoothing inevitably worsen the regulation of the supply, whilst simultaneously adding expense. The current through the valve therefore consists of half-wave rectified pulses (see [Figure 4.38](#)).

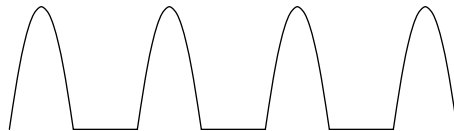


Figure 4.38

Valve anode current in scaled AC valve tester

We saw earlier that the inertia of a moving coil meter causes it to respond only to the DC component of a current. A full-wave rectified sine wave is composed of:

$$V = V_{AC(RMS)}[0.90 + 0.6(2f) - 0.12(4f) + 0.05(6f) - 0.03(8f) \dots]$$

Thus the DC component corresponds to $0.90V_{AC(RMS)}$, or, to put it another way, we require $V_{AC(RMS)} = 1.11V_{DC}$. This is not a problem – we simply scale our anode and screen grid tappings appropriately, so that when we select $V_a = 400\text{ V}$, the tester actually applies $444V_{RMS}$ to the anode.

We can use almost the same trick at the grid, but we cannot allow the grid to go positive as the resulting grid current could damage the valve, so we half-wave rectify the grid bias voltage. (Full-wave rectification is not necessary because the valve cannot conduct on the missing grid half-cycle when the anode and screen grid voltages are negative.) Half-wave rectification halves the mean voltage compared to full-wave rectification, so when read by a moving coil meter, the control grid voltage should correspond to 0.555 of the claimed equivalent DC voltage. In practice, the calibration procedure for both the portable CT160 and the laboratory VCM163 specified a ratio of 0.52.

Grid voltage discrepancies aside, this measurement technique means that although the valve passes appropriate currents when conducting, it only does so for half the time, so the meter reading anode and screen grid currents needs to be calibrated to indicate double the current that would be indicated by an external meter.

Despite the previous caveats, the scaled AC technique makes it possible to take reasonably accurate measurements over a wide range at far less cost than a true laboratory test rig using pure DC supplies.

Although one use of a valve tester is to make sufficient measurements to be able to plot curves, the more common requirement is to investigate performance at a single point, perhaps to match valves in push–pull output stages. When matching valves, we must not only match anode currents (which we can already measure), but also match mutual conductance. Mutual conductance is a moving target, and because it changes with anode current, it is a small-signal parameter, so it should be measured by a low-amplitude AC voltage applied to the valve’s control grid. A laboratory DC test rig could use any convenient frequency, but a scaled AC tester applies mains frequency plus a spray of harmonics to the grid, so a higher frequency is required, but not so high that Miller capacitance can cause a problem. In the VCM163, AVO chose to apply 15 kHz at $\approx 270 \text{ mV}_{\text{pk-pk}}$ to the control grid, sense the anode (or screen grid) current with a 10Ω resistor, then amplify the resulting (rather small) voltage with a tuned amplifier in order to reject the (rather larger) mains related signals. Because this enabled accurate determination of mutual conductance at any operating point, it became known as a **dynamic** measurement. By contrast, earlier AVO testers made a **static** measurement of mutual conductance. They biased the control grid to $-\frac{1}{2} \text{ V}$, measured the resulting anode current, and the meter current was then manually nulled using a “backing off” control. Finally, grid voltage was increased to $+\frac{1}{2} \text{ V}$ and the increase in anode current could be directly read as mutual conductance in mA/V . Interestingly, the complete patent specification [1] mentions that using $-\frac{1}{2} \text{ V}$ to $+\frac{1}{2} \text{ V}$ doesn’t give the right answer, but that if the voltage change is doubled from -1 V to $+1 \text{ V}$, “*the mutual conductance figure is substantially correct*”.

A scaled AC valve tester produces test voltages that are directly proportional to mains voltage, so not only does the mains transformer primary need coarse tappings to set the tester to the nominal mains voltage, it also needs fine tappings that can be adjusted from the front panel to compensate for hourly variations. We can now draw a simplified diagram of a final generation scaled AC valve tester (see [Figure 4.39](#)).

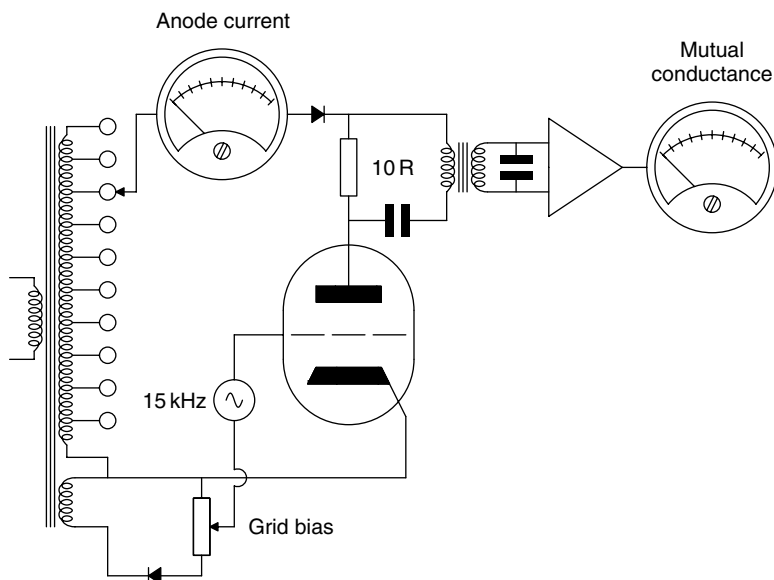


Figure 4.39
Simplified diagram of scaled AC valve tester

In practice, although the anode and screen grid are able to self-rectify, AVO found that adding a series silicon diode prevented parasitic oscillations, so the voltage at these electrodes is already half-wave rectified.

When scaled AC testers were designed, mains demand was far less variable than it is now, and sudden current demands caused by half the population switching on a 2 kW electric kettle at the instant of the adverts part-way through a popular TV programme just didn't occur. In short, scaled AC testers need their mains to be regulated if they are to be used for generating sets of curves.

Synthesised mains supplies with excellent regulation are available, but they are expensive (as much, if not more, than the valve tester). Constant voltage transformers can't be used because although controlled saturation of their iron core holds V_{RMS} constant, it distorts the waveform and therefore invalidates

the AC/DC relationships on which scaled AC testers are based. The cheapest solution is a second-hand traditional AC voltage regulator based on a servo motor-driven variac.

Incidentally, if you have an AVO tester with a dirty electrode selection switch, they can be removed, stripped down, thoroughly cleaned, reassembled, and refitted. The job takes an entire afternoon, but if you opt to do it, take the opportunity to check and if necessary, correct an early design flaw in the rotor's phosphor bronze leaf springs, which should sit flush in a moulded cavity, but some did not have cropped corners, so they sat proud and wore the adjacent rotor. As you reassemble the rotors onto the shafts, check that each one rotates freely in both directions and that the leaf springs don't catch.

AVO testers had double pole mains switches that are likely to be worn or have dirty contacts, possibly both, which can add significant (and variable) resistance in series with the tester, and significantly distort results. It's far cheaper and simpler to replace it immediately than have to diagnose it later as being the fault...

Since the modern use of valves is primarily for audio, the voltages required are rather more restricted, so the scaled AC technique is no longer necessary. Rather than having many sockets wired in parallel with electrode connections selected by a large (and expensive) multi-pole switch, modern testers tend to have dedicated plug-in boards for each socket and removable wire links to select electrode connections.

Curve tracer

One of the uses of the valve tester was to generate data from which curves could be plotted, another was to match valves at a particular operating point. If we had a machine that could plot curves directly, we could not only save time, but we could

match valves over an entire set of curves rather than at a single point. A curve tracer is effectively a low-bandwidth oscilloscope with fixed time base and a step generator instead of the trigger block, so comparison between one set of curves and another either requires a curve tracer with digital memory, or an analogue tracer and a camera. (Traditionally, a Polaroid film camera would have been used, but a modern digital camera can do the job far more conveniently.)

The most useful curves for a valve are the mutual characteristics (I_a against V_g), or the anode characteristics (I_a against V_a). In each case, anode current is monitored and applied to the “Y” anodes, and a swept voltage is applied to the “X” anodes to produce a curve (see [Figure 4.40](#)).

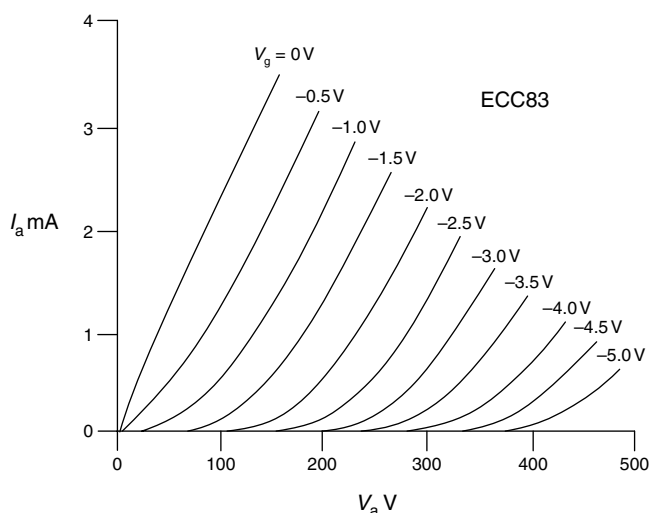


Figure 4.40
Valve anode characteristics

Since the swept voltage is applied to the “X” anodes and the valve simultaneously, it can have any waveform, and the cheapest waveform is a half-wave rectified sine wave derived from a mains transformer. Unfortunately, the half-wave rectified sine wave slows towards the end of the sweep causing

uneven brightness when displayed on an analogue oscilloscope, and other errors can cause the retrace (as the voltage returns to 0 V) not to overlay perfectly. A ramp waveform plus blanking solves both these problems but requires a generator and power amplifier to feed the valve.

Ideally, we would like to produce a family of curves simultaneously, perhaps a complete set of anode characteristics for a 6080 power triode. To do this, we would need to sweep V_a from 0 V to perhaps 200 V, supplying a current of up to 65 mA (to limit P_a to the maximum allowable dissipation of 13 W), requiring a power amplifier based on a HT regulator. In addition, we would need a step generator and high-voltage amplifier for V_g that quickly and repetitively sequences from 0 to -80 V in 10 V steps to generate the individual grid curves (see Figure 4.41).

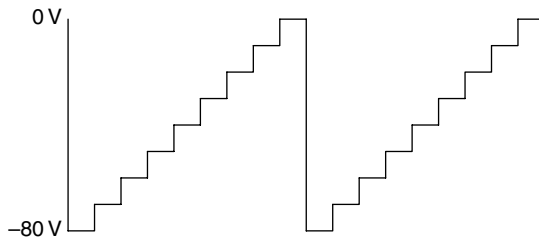


Figure 4.41

A repetitive stepped grid waveform enables a family of anode curves

The previous arrangement would display triode anode characteristics, but displaying pentode anode characteristics is somewhat harder. Measuring pentode anode current forces the current sense resistor to be in the anode circuit, but as the absolute voltage of this resistor is being swept from 0 V to perhaps 400 V, a differential amplifier with very good common mode rejection is required, whereas we could declare that $I_a = I_k$ for the triode, and measure cathode current without worrying about a superimposed sweep voltage. The other problem with measuring a pentode is that when $V_a = 0$ V, $I_a = 0$, so the screen grid becomes an anode and passes the entire cathode

current, yet its maximum dissipation is strictly limited. As a consequence, we cannot simply apply a constant voltage to the screen grid, it must be switched on only when the anode voltage is being swept.

Another problem is that as the anode voltage approaches its maximum, so does anode current, causing anode dissipation to approach its maximum, but a hot anode changes the valve's characteristics slightly. (This is one cause of retrace error in the half-wave rectified sine wave sweep.) To avoid this error, the sweep can be modified to consist of a series of short pulses having amplitudes that follow the path of the original ramp waveform. As an example, if the anode voltage pulses caused anode current to be switched on for only 10% of the time, they would reduce anode dissipation to one-tenth of the continuous ramp.

The preceding thoughts only consider an analogue curve tracer, yet they show that a curve tracer must be quite complex. Very few laboratories needed to buy a curve tracer, so the development cost had to be recovered over a small number of sales, and prices were stratospheric.

The question is, "Is it worth the price?" The answer to this sort of question is always personal, but the author's oscilloscope is switched on ten times as often as his digital transistor curve tracer, so the high prices commanded by geriatric valve-based analogue curve tracers are staggering.

Despite knocking the amateur utility of a professional curve tracer, there is one instance where a valve curve could be useful, yet easily obtained. Suppose that you had a stereo push-pull triode amplifier that needed new output valves, and you found a dozen new old stock (NOS) valves in a dusty box at a radio fair for a song. From these valves, you would like to find pairs, but you don't have a valve tester, or a curve tracer, yet you do have an oscilloscope and a junk box of transformers. A quick test set-up conveniently enables matching [2] (see [Figure 4.42](#)).

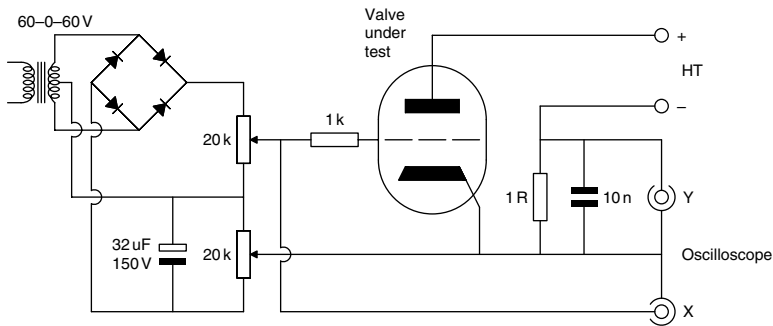


Figure 4.42

This test set-up by Alan Douglas [2] enables a single mutual characteristic to be plotted by an oscilloscope

The set-up enables the entire mutual characteristic (I_a against V_g) of a valve to be determined for one value of anode voltage. Since a power amplifier output stage operates with a fixed HT voltage, this sweep tells us everything we need to know to be able to pick pairs of valves (see [Figure 4.43](#)).

Semiconductor component analyser

This is a tiny handheld device with three leads and a dot matrix LCD screen that scrolls to give comprehensive information. It is incredibly useful for quickly determining the pin-out of a known component or identifying an unknown one, and can also be used for faultfinding if you don't mind removing components from their surrounding circuitry. Unfortunately, they are still quite expensive, so you will need to think carefully before buying one (see [Figure 4.44](#)).

A handy constant current regulator

The author often needs a constant current regulator during testing and prototyping. The device is simply placed in series with

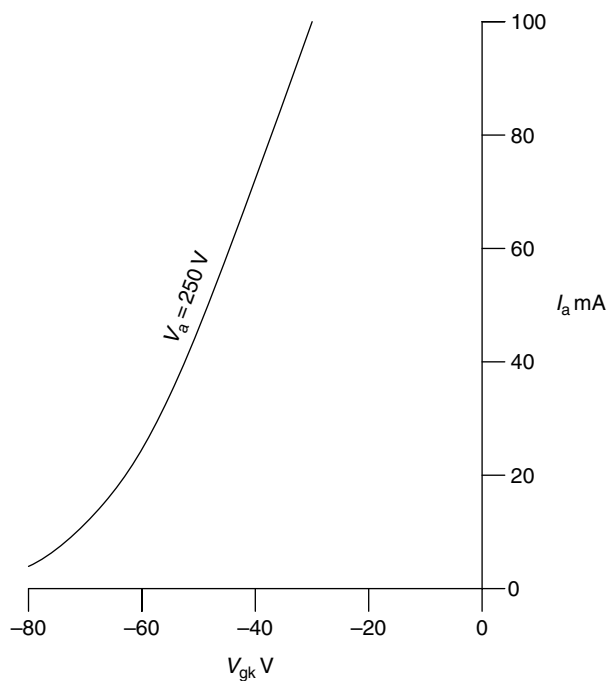


Figure 4.43

Plotting the mutual characteristic of I_a against V_g at the chosen anode voltage enables easy matching of power valves



Figure 4.44

Semiconductor component analyzer

any power supply and turns it into a programmable constant current source, so it only needs two terminals (see [Figure 4.45](#)).



Figure 4.45

This device converts any voltage supply into a regulated constant current supply

The circuit is a modification of standard constant current heater regulator circuits, but with the addition of a multi-position switch and variable resistor it can very conveniently set any current between $\approx 3\text{ mA}$ and 1.5 A , so it can be used for anything from plotting diode curves to powering “P” series valve heaters (see [Figure 4.46](#)).

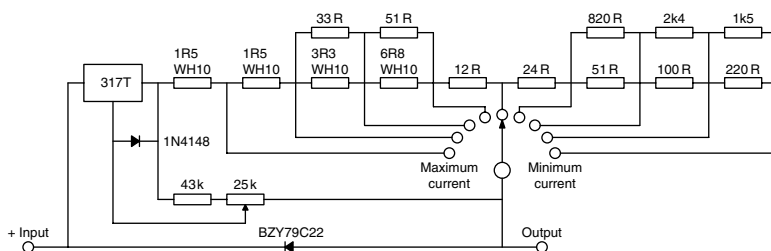


Figure 4.46

Circuit diagram of 3 mA to 1.5 A adjustable constant current sink

The BZY79C22 diode protects the 317 from over-voltage and stored charge in the load. The 317 must be bolted to a heatsink, perhaps the aluminium chassis. The logarithmic range switch doubles current at each click, whilst the 25 k Ω variable resistor and associated 43 k Ω resistor ensure that each range just overlaps the next. The entire circuitry can be very conveniently hardwired on the lid of a 6" $\times$ 4" $\times$ 2" diecast aluminium box, making it very easy to make or repair (see [Figure 4.47](#)).

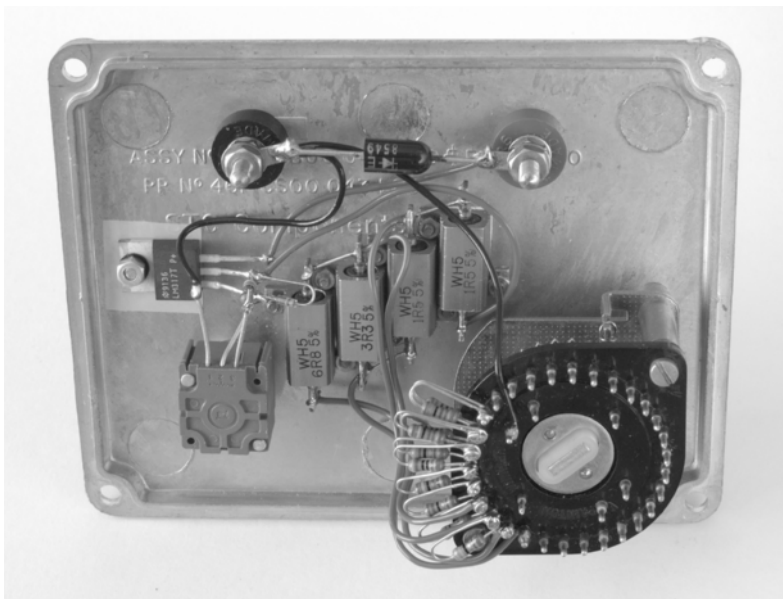


Figure 4.47

The components are hardwired onto the lid of a diecast box

Audio phase meter

It is often useful to be able to measure the relative phase of two related signals. Traditionally, this was done by counting squares on an oscilloscope and scribbling numbers on the

back of an envelope, but this is very inaccurate. A much handier method would be to have a circuit that produces a voltage directly proportional to phase and with reversible polarity to indicate lead or lag. Such a device would convert an ordinary DVM into a digital phase meter (see [Figure 4.48](#)).

The circuit is actually very simple, and is essentially an updated version of the BBC phase meter that was used for testing the equalisation of stereo analogue music circuits used for Radio 3 concerts. The 072 FET input op-amps allow a 1 M input resistance suitable either for oscilloscope probes or for direct connection. The back-to-back Zeners protect the op-amps from over-voltage. The sine waves from the op-amps are passed to a pair of 311 comparators which convert them into square waves.

An XOR gate only produces a logic “1” when its inputs are different, so when the two square waves are 180° out of phase, it produces a “1”, and when they are perfectly in phase, it produces a “0”. Because the change from 0 to 180° changes the width of the pulse leaving the XOR gate, its mean level is directly proportional to phase. Mean level is found by dumping charge onto the 470 n plastic capacitor. The voltage can now be read by a DVM, and is scaled by the potential divider to produce 180 mV when the two inputs are exactly out of phase (180°).

The D-type determines whether Ch2 lags or leads Ch1. If Ch2 lags, there will already be a “1” present when the positive-going edge of Ch1 triggers the D-type, so this “1” will be loaded to the Q output and remains there until the D-type is next triggered. The “1” (5 V) drives a limited current into the base of the BC549, turning it on, and activating the relay which reverses the polarity of the 0–180 mV phase voltage fed to the DVM, thus activating the DVM’s +/– sign.

All ICs to have power supply pins decoupled to ground via 100 n ceramic

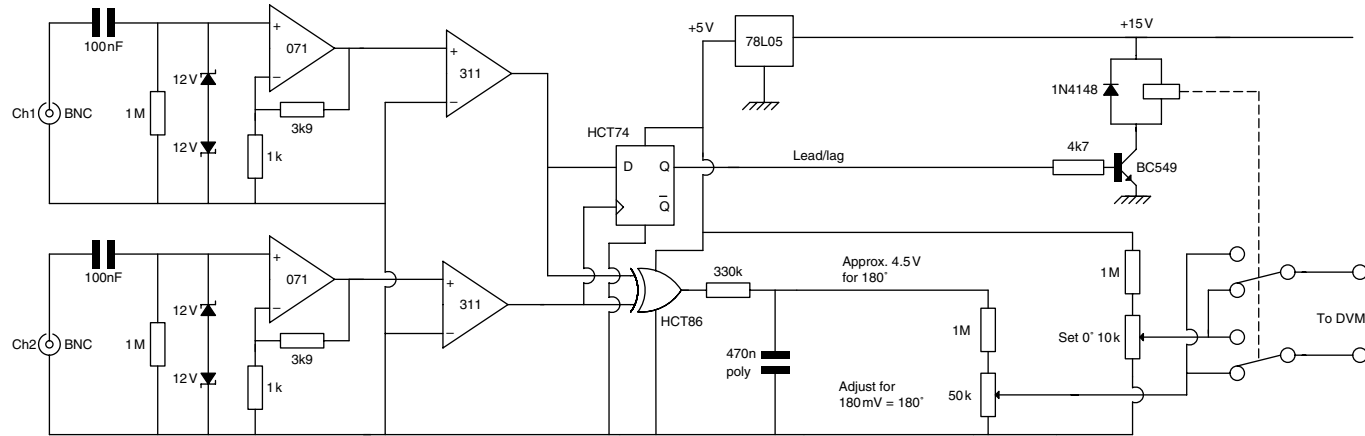


Figure 4.48
This circuit converts a DVM into a precision audio phase meter

References

1. “An Improved Method and Apparatus for Testing Radio Valves”
Sydney Rutherford Wilkins & The Automatic Coil Winder &
Electrical Equipment Company Ltd. UK Patent 480752 (1938).
2. “Tube Testers and Classic Electronic Test Gear” Alan Douglas.
Sonoran Publishing (2000).

Further reading

- “Audio Measurement Handbook” Bob Metzler. Audio Precision.
- “Electronic Measurements” 2nd edn. F E Terman and J M Pettit.
McGraw-Hill (1952).

CHAPTER 5

FAULTFINDING TO FETTLING

In the previous chapter, we looked inside test equipment to see how it worked and to understand its design limitations. We are now in a position to move on, so this chapter is concerned with using test equipment to bring a device to a usable state.

There are three stages of testing:

- Safety testing: Will the device under test (DUT) endanger the user?
- Functionality testing: Does the DUT work as it should?
- Faultfinding (hopefully optional): Having established that the DUT doesn't work correctly, what needs to be fixed?

Rather than scribbling results on scraps of paper and losing them, you will find it extremely useful to keep a logbook of your tests. The logbook should include:

- Circuit diagrams of the circuit tested
- Diagram or description of how the test was undertaken
- Results of the test
- Subjective comments on the test (Were the results expected?)
- Date of the test, and test equipment used
- If any of the test data has been saved to computer (spreadsheet, photograph, etc.), include the file names.

The purpose of the logbook is to enable you to look back and **know** the results of a test, rather than scratching your head and saying, “I’m sure I’ve tested this before . . .”

Safety

Before we go any further, we should understand that there is no such thing as perfect safety. Everything we do, from crossing the road to eating a chicken sandwich carries an element of risk. Lawyers currently enjoy a climate where ludicrous damages are awarded to clients who failed to take any responsibility for their own safety – of course a cup of hot coffee will scald you if you spill it over yourself! The dangers of hot liquids are well known, but the dangers of valve electronics are not as well known as they used to be, and even professional engineers can become complacent about putting their hands inside live equipment.

The aim of this section is to point out some of the more common dangers and show simple ways by which risk may be reduced. Nevertheless, it is impossible to cover all possible situations, and this is why electronics workshops forbid work on live equipment unless two people are present. In that way, one person is ready to rescue the other in the event of an accident. If you have a more experienced friend, it would be a good idea to emulate commercial practice by asking them to attend when live working is necessary and to check your work beforehand. Ensure that they know how to switch off the power quickly.

An understanding of the physiological effects of electric shock, whilst somewhat ghoulish, serves to underline why safety is so important.

How an electric shock can be received

The best way of avoiding electric shock is to understand how it can be received.

The electricity supply leaving the wall socket in your home is a very good approximation to a pure Thévenin source of zero resistance, with one side of the source connected to earth. In this instance, we do not mean earth in its purely technical sense, the supply really is connected to the planet Earth. You and I perform most of our activities on the surface of the earth, and we are therefore electrically connected to it, albeit usually by a high-resistance path. We can improve our electrical contact to earth by standing barefoot on a damp floor, or by firmly gripping something that is electrically bonded to earth.

Manual activities like hobbies are precisely that, **manual**. They involve our hands going inside objects, and touching them. Humans generally possess two hands, and so what is more natural than to put **both** hands inside a piece of equipment?

The scene has now been set for one hand to be holding the (earthed) chassis of a piece of equipment, whilst the other is moving around and accidentally comes into contact with live mains.

Look at your hands. They are on the ends of your arms, and your arms are joined to your torso. Trace a line from the fingers of one hand to the fingers of your other hand, without the line leaving your body. Note the path.

The easiest path for an electric current to flow from one hand to the other crosses near the heart. Similarly, a shock from one hand to earth passes near to the heart.

The effects of electric shock

The main danger from electric shock is fibrillation of the heart. The heart normally beats at a slow pace regulated by electrical impulses from the brain. If we apply 230 V 50 Hz to the heart, it pulses quickly, the flow regulating valves do not operate correctly, and no blood is pumped, leaving the brain to die of

oxygen starvation in about 10 minutes. A sustained current of 20 mA through the heart is sufficient to cause fibrillation, resulting in the adages, “20 mills kills”, and, “it’s the volts that jolts, but the mills that kills”.

A sustained current might not kill, but could cause irreversible injury due to the heating effect of the current. RF burns are notorious for this, and can result in limbs having to be amputated to prevent the spread of gangrene. A lower sustained current may cause injury from which the victim does recover. Eventually. Skin grafts may be necessary.

A current of 20 mA results from a 230 V supply connected across a resistance of $\approx 11\text{ k}\Omega$. Very few power supplies have a source resistance as high as this, so shock current is determined primarily by skin resistance, and a shock from a puny transformer delivering 230 V is just as dangerous as the shock received directly from a 100 A mains feeder. Skin resistance is reduced by damp hands, and standing in the rain in a puddle of water lowers resistance further.

All of the above considerations refer to the direct consequences of electric shock, but do not consider secondary effects. A minor shock that causes the victim to lose their footing and fall over could be fatal if they happen to be standing on a ladder 30 ft above concrete. Another possibility is that the reflex muscle jerk in reaction to the shock could cause the victim to throw themselves through a glass window and bleed to death.

Even after a minor electric shock, the victim will be confused and disoriented, and shock in its full medical sense is a possibility. Shock kills.

Burns

Although the primary hazard of electrical equipment is shock from high voltages, it should be realised that low voltages can

be just as hazardous. A low-voltage/high-current DC supply will have a large, low ESR, reservoir capacitor capable of delivering many amps of current into a short circuit.

Rechargeable batteries are even more dangerous because they are capable of sourcing substantial current for a significant time. (Think about it, a car starter motor requires 100 A, or more, for several seconds.) These batteries are not merely capable of burning, they can vapourise metal bracelets, watch-straps, and tools.

Do not wear jewelry whilst working on live equipment.

Avoiding shock and burns

It should now be obvious that electric shock is potentially lethal, burns can be serious, and that both must be avoided at all costs.

Provided that you have made, or modified, your equipment carefully, there will be no exposed voltages, and all metalwork will be earthed, resulting in a very low risk of shock. The danger arises when you **deliberately** remove the safety covers, and start testing the equipment with power applied.

Some authorities suggest that you should always work on live equipment with your left arm behind your back, so that any shock received will not pass from arm to arm across the heart. Whilst it is true that this will reduce the severity of the shock, it tends to increase the risk of receiving a shock.

The best way of improving safety is to think about safety, and to **think about what you are doing**. It might seem obvious to think about what you are doing, but for most of our lives we think about many things at once. For instance, when driving, are you thinking **only** about driving,

or are you actually thinking about what you are going to say to your boss when you arrive late for work, and when is that idiot in front of you going to turn into the junction, and isn't that a rather attractive male/female/alien over there by the bus stop?

Thinking about what you are doing means not working late. Do not attempt to test a newly completed project at 11.30 at night; you will not be alert and could damage the project and/or yourself.

Electrical safety testing

Consumer electronics falls into two safety categories. Class I has an earthed conductive barrier enclosing the high voltages, whereas Class II has two independent insulating barriers enclosing the high voltages. Many appliances are actually a combination of the two categories because although the appliance could be Class I, the mains lead is invariably Class II. Class I relies on **quickly** blowing the mains fuse, whereas Class II relies on undamaged barriers. Commercial safety testing therefore consists of a careful inspection by a competent person to check barrier integrity, plug wiring, and that the correct fuse has been fitted. These visual tests are backed up by electrical tests.

A safe piece of Class I equipment has an earthed conductive barrier with no conductive path to any of the enclosed high voltages and a sufficiently low-resistance path to earth that any contact to the mains line causes such a large current to flow that the mains fuse ruptures quickly. Electrical safety can therefore be tested by measuring leakage current flowing from the barrier and by measuring the resistance from the barrier to the earth pin of the mains plug. Since many pieces of equipment have to be tested, it makes sense to have a dedicated instrument to make the tests.

Portable appliance testers

Portable appliance testers (PAT) are simply specialised resistance and leakage testers with appropriate connectors for **quickly** testing mains portable appliances. The emphasis on the word “quickly” is important because even a small business could have hundreds of appliances that must be tested for safety each year. In order to cause minimum disruption each appliance must be tested quickly and unambiguously, and a record must be made of the test. Older PAT testers simply make the required pass/fail electrical tests, whereas newer testers guide the operator through the test procedure, make the measurements, give pass/fail status, and log the data together with the appliance’s identification (often a bar code) for later download to a computer. Because the newer testers make it much easier to comply with the various safety regulations, there are plenty of older, manual testers available in perfectly good working order, and you might want to acquire one (see [Figure 5.1](#)).

As part of a Class I test, a PAT tester measures earth loop resistance (resistance from the chassis to the earth pin of the plug) by passing a considerable current, sometimes as much as 25 A. The purpose of such a large test current is to detect frayed earth connections. A single strand would still have low resistance because it would be very short, even though its cross-sectional area would be small, so a simple resistance measurement could not detect the problem. However, the PAT tester deliberately seeks to rupture such frayed connections, which subsequently fail the resistance test. At the same time, the tester applies a high voltage ($>500 V_{\text{RMS}}$) simultaneously to the line and neutral terminals and monitors leakage current returning through the earth circuit. Do not touch the appliance whilst the PAT tester applies this dangerous voltage.

Be aware that PAT testers are not invincible, and can develop faults just like any other piece of equipment. If you rely on a tester to determine safety of other equipment you have to be



Figure 5.1

Older PAT testers are still perfectly capable of making important safety tests

certain that the tester works correctly and that you are operating it correctly. Most testers include operating instructions and a selection of “faulty” test jigs that should be regularly tested to verify correct operation of the tester. You might question the validity of using a second-hand tester, particularly since it will almost certainly be sold without any warranty of fitness for purpose. Provided that you are **not** using it to assist in selling goods or services, the question is not whether the tester perfectly

meets all the latest safety regulations, but whether it reduces risk compared to not being able to test at all . . .

Functionality testing

The word “testing” implies a degree of ambiguity about the results of the test; if we **knew** that our new amplifier was going to work perfectly from the moment that it was completed, we would not need to test it. However, we know that wiring mistakes can be made, and that components could be faulty, so we test our amplifier carefully.

Second-hand equipment versus freshly constructed new equipment

Both types of equipment should be treated with a great deal of suspicion and apprehension. The only sensible state of mind when first switching on is controlled fear.

Quite clearly, a piece of newly constructed equipment **will** be switched on at some point, but some old equipment may be so dangerous, or riddled with faults, that it should never be energised, other than applying kinetic energy to throw it into a skip. The state of old equipment can easily be determined by looking at the components.

Things to avoid are:

- Wire insulated with rubber and covered with cotton
- Enormous resistors marked with tip, body, and spot colour codes
- Electrolytic capacitors with bulges in the rubber surface supporting the tags
- Previous evidence of fire
- Insulating tape anywhere!

Vintage radios may incorporate any or all of these features but may still be of value to someone, so try to check with someone else before destroying them [1].

For a professional, the worst possible sign is previous modification by an amateur. The professional then has to decide whether or not the amateur knew what they were doing. What is the effect of their modification, and was it done competently and safely? For this reason, modified equipment is usually worth **less** than unmodified equipment, so bear this in mind before you embark on modifications.

Second-hand equipment can often be dated by the date on the electrolytic capacitors. If it is over 30 years old, it is likely to need major refurbishment even to make it work, so this should be taken into account if you are considering purchase.

The first application of power

Before applying power for the first time, the following “Ten Commandments” should be observed:

- Inspect the earth bonding. Does the chassis appear to be properly earthed? (Does an undamaged earth wire make a good connection to the chassis? If a tag is used, is it tightly bonded with a star washer between it and the chassis?)
- When measured, is the resistance from the earth pin of the mains plug to the chassis of the equipment significantly less than $0.5\ \Omega$? (Preferably use a PAT tester.)
- With the power switch on the equipment (if fitted) switched on, is the resistance from the other pins of the mains plug to the earth pin infinite? (Preferably use a PAT tester.)
- Does the mains cable look safe? i.e. not frayed, perished, cut, or melted by a soldering iron.
- Is the mains plug wired correctly?
- Does the plug grip the sheath of the cable correctly?

- Does the plug look safe? i.e. no cracks, chips, dirty pins, etc.
- Is a fuse of appropriate rating fitted? (It is unlikely that the fuse rating should be greater than 3 A.)
- Does all the internal wiring of the chassis look secure?
- Is the chassis clear of swarf and odd off-cuts of wire? Turn it so that bits can fall out, and give it a really good shake, whilst blowing vigorously into the chassis to free small parts. Alternatively, if it is too heavy to lift and shake, use a $\frac{1}{2}$ " paintbrush and a powerful vacuum cleaner to remove debris.

If all appears to be well, you can move on to the next stage.

The author **always** assumes that when power is applied, the amplifier will explode, or at the very least, catch fire. It does not, therefore, make sense to stand with your face over the amplifier, or to place it in the middle of a pile of inflammable debris. Power amplifiers should have dummy loads, or “disposable” (cheap) loudspeakers should be connected across their outputs.

The safest way to test a valve amplifier is in stages. Many amplifiers use a valve HT rectifier, so if this is removed, the mains transformer and heaters can be tested before applying HT. Silicon HT rectifiers make disabling the HT a little harder, and require the AC to the rectifiers to be removed. If this is a new amplifier, you will have planned ahead by testing the heaters **before** doing any other wiring (that way, the heater wiring is easily accessible should a fault surface).

The heaters in indirectly heated valves take a moment before they begin to glow, so connect a meter across the heater supply to give an instant indication of whether the heater supply is present, and leave the meter in a clearly visible position. If you have a variac, you can apply power gently, and if the meter monitoring the heater supply doesn't immediately respond when you advance the variac from zero, back it off and investigate. The advantage of using a variac is that even a short circuit across

the heater supply would be unlikely to cause damage because you would spot it before applying significant voltage.

If you don't have a variac, retire to a safe distance and apply the power in silence. This way, you will hear any unusual noises, such as the crackles or pops that presage destruction. If the meter monitoring the heater supply responds appropriately and nothing untoward happens, move a little closer and sniff the air. Can you smell burning? Are there any little wisps of smoke leaving the chassis? If all still seems to be well, look closely at the heaters – they should be glowing, but should not have hotspots (see [Figure 5.2](#)).



Figure 5.2

This 13E1 had a nasty cathode hotspot, so it was rejected without HT ever being applied to its anode

Leave the valves to warm for a while. Ideally, at first switch-on, unused valves should warm their heaters for 30 minutes before HT is applied. However, even 10 minutes gives time for your nerves to calm down and for you to consider your next step. Having left the heaters glowing for a while, check the temperature of the mains transformer, which should be cool. Switch off power,

and **unplug** the mains lead, leaving it in plain sight. If the amplifier has heater regulators, carefully inspect them for signs of damage, but be careful near the valves, which will still be hot.

If nothing has been damaged by warming the heaters, it is time to test the HT as well as the heater supplies. Depending on the complexity or output power of the amplifier, this might be done in a number of ways:

- Classic amplifiers simply need their HT rectifier to be inserted and the amplifier to be switched on.
- Modern amplifiers often use silicon HT rectifiers. Suddenly applying peak HT to an electrolytic that has sat quietly on a shelf for several months is unkind. Use a variac to bring the voltage up gently over 20 seconds, or more.
- If the amplifier has separate HT and heater transformers, it is worth powering the HT transformer from a variac and testing the HT before completing the final wiring that connects the transformer to the amplifier's internal mains distribution.

Individual circumstances will determine how you test the HT for the first time, but if it's possible to do it gently, then do so – it minimises the amount of smoke.

Having chosen your method of testing the HT, apply power. Listen. Are the loudspeakers making any unusual noises? In this instance, silence really is golden. Is the HT voltage correct? If the HT voltage is correct, then it is highly likely that the circuit is working as it should, and you may breathe a quiet sigh of relief.

If you are testing a newly built power amplifier, there is a 50/50 chance that global negative feedback taken from the loudspeaker output will turn out to be **positive** feedback, and the amplifier will become a power oscillator. Often, before oscillation starts, a quickly increasing hum will be heard from the loudspeaker; if the amplifier is switched off at this point, only a brief shriek of oscillation will be suffered.

If, at any point, something untoward happens, switch off immediately at the mains outlet and unplug the mains plug.

Usually, if anything is wrong in an electronic circuit, heat is generated and components are burnt, so look for charred resistors or wires. Once the burnt parts are found, the fault is usually blindingly obvious.

If you switched off hurriedly because of a burning smell, what sort of a smell was it? Old equipment will often be dusty, so a slight burnt dust smell is normal. Bacon smells are sometimes produced by burning mains transformers, whereas burning PCBs often smell like underground railway stations with a hint of charcoal, but burning wiring gives off an acrid smell.

If the amplifier appeared to be satisfactory, leave it switched on for a minute or two longer, whilst keeping an eagle eye on everything; particularly output valves, which should not have glowing anodes, or purple and white flashes. Switch off and sniff the internals closely for unusual smells. Some engineers go one step further and touch components with their finger to check temperature, but this is not recommended as some HT may still be present. If all seems well, the amplifier can be switched on again, and all the DC voltages carefully checked, if it still looks good, then it probably **is** good.

If you have a variac, now is the time to check how well the amplifier responds to mains voltage variations. UK mains voltage is **specified** as being $230\text{ V} +10\% -6\%$, but this is simply a paperwork ruse to enable harmonisation with the rest of the European Union. In practice, the voltage at your wall socket is the same voltage it always was. Nevertheless, although mains voltage variations are smaller than they used to be, the voltage does vary, and could cause problems. Use the variac to check that capacitor voltage limits are not exceeded when 253 V is applied, and that regulators don't become excessively hot. Similarly, drop the mains voltage to 216 V and check that

regulators do not drop out. These two tests are quite severe, and if you **know** that your mains voltage is more stable, you might decide to adopt a less stringent test.

For the first few weeks of service, a new amplifier should be watched like a hawk for signs of incipient self-immolation and should not be left unattended when switched on.

Faultfinding

Unfortunately, we all have to do some faultfinding at some time or another, either because we made a design or construction mistake, or because we need to repair an amplifier that has failed due to old age. Before we dive into details, we should ask one very important question, “Did it work once?”

If it worked once, you are looking for a faulty joint or component, and the resistance range of your multimeter will prove invaluable for finding carbon resistors that have “gone high” in value from age, or capacitors that have become leaky. If the circuit is freshly built, then you are probably looking for a wiring error.

Individually test each component

One way of faultfinding an amplifier might be to remove each component individually and test it:

- Resistors: Check claimed value against measured value using the resistance range of a DVM.
- Capacitors: Use DVM on resistance range to check for leakage, then check claimed value of capacitance against measured value on the capacitance range of the component bridge.
- Inductors: Use DVM on resistance range to check continuity, then measure inductance on the component bridge.

- Valves: Use valve tester to compare with test data on the valve tester.
- Semiconductor diodes: Use diode check range on DVM to check for open circuit when reverse biased, and correct forward drop when forward biased.
- Transistors: Use a semiconductor analyser to check functionality and compare measured h_{FE} with manufacturer's data. Even better, use a curve tracer to plot its full characteristics.
- Transformers: Connect to an oscillator and check each output with an oscilloscope to verify functionality and turns ratio.

You will notice that the previous tests assume that you have a great deal of expensive test equipment and full manufacturer's data to check your measurements against. We only remove and test each component fully when:

- We know no better
- All else has failed.

The very best piece of test gear is your brain. It's readily available, so it seems a shame not to use it...

DC conditions

The most common fault is a lack of signal, or reduced signal accompanied by gross distortion.

Most faults can be found very quickly by measuring the **DC conditions** of the circuit. For a well-designed circuit to be observably faulty, the DC voltages usually need to be very wrong. Consequently, checking the measured voltages against the design voltages quickly pinpoints the fault – it's easy to confuse resistor multiplier band colours. Mark the measured voltages (lightly, in pencil) on the circuit diagram. This usually has the effect of making the fault appear blindingly obvious.

Sometimes you will not have the circuit diagram of the amplifier, let alone its design voltages. No matter, there were very few variations in classic valve circuitry, and their circuits were so simple that it is not difficult to produce a block diagram of the amplifier. At this point, consider how you would design an amplifier using those valves, and look for similarities in the actual amplifier. It should now be possible to obtain a rough idea of what sensible voltages might be, and these can be checked against the faulty circuit. Modern amplifiers are likely to be more complex and contain transistors, so they may require careful circuit tracing. Fortunately, there are now so many websites carrying information on valve audio that it is highly likely that the information you need is somewhere on the web – you just have to find it. If you can't find a complete circuit diagram, valve data sheets are a good second best because they give maximum ratings, typical applications, and pin connections.

A calculator is invaluable for calculating currents through resistors, and generally deciding whether measured voltages make sense.

Do not implicitly believe what your digital voltmeter tells you. Even the standard $10\text{ M}\Omega$ input impedance of a digital voltmeter loads some circuits, particularly the grid circuit of cathode followers or circuits with grid battery bias. The author was once convinced that audible distortion was due to the DC conditions within a valve active crossover, and a digital voltmeter appeared to confirm the theory, but a valve voltmeter with $90\text{ M}\Omega$ input resistance measured a more correct value, and the distortion finally turned out to be due to an intermittently scraping voice coil in a loudspeaker.

Blocks and attitudes

Imagine that you have just been told that the Hi-Fi isn't working. If you have simultaneously been plunged into darkness, you assume that the power has failed. Alternatively, if you look at the

CD player, and see that its display isn't counting, despite you having pressed "play", you conclude that either the player or the CD is faulty, and you try another CD. In each instance you are breaking the system down into blocks and checking each block. Exactly the same technique can be applied to internal electronic faultfinding.

A power amplifier example

Imagine that you switch the amplifier on, play a CD, yet nothing comes out of the loudspeakers. You look at the amplifier and observe that the heaters are glowing. You have just eliminated the mains lead and associated fuse. Next, it could be that the amplifier simply isn't receiving any signal, or that the loudspeakers have become disconnected. Of course, the loudspeakers could be faulty, but most loudspeakers contain two drive units, and possibly more, so for both loudspeakers to be faulty, you require four simultaneous failures. That just isn't likely, even if you had a **very** good party the night before. Similarly, although the leads to both loudspeakers could both be faulty, it's unlikely. It is far more likely that we are looking for a single fault that affects both channels.

We ought to check that a signal is reaching the amplifier. We could try a different source, or if the CD player has an integral volume control (and some do), we could plug it directly into the power amplifier and increase the volume gently from zero.

Assuming that there is still no sound, we know that we are looking for something that is common to both channels and that's usually the power supply. We know that the heaters work, so that narrows the investigation to the HT supply. There haven't been any nasty smells, explosions, or fizzing noises, so something has died quietly and completely. If a component was merely poorly, we would have had low HT, and the amplifier would have produced some sound, although probably very distorted.

All HT rectification is full-wave, so it is unlikely that both diodes have failed, unless it's a valve rectifier, where the

common cause for two diode failures would be the heater, but we checked that all the heaters were glowing. We can lose HT either because something in series has failed open circuit, or because something has failed short circuit down to 0 V. If an HT capacitor had failed short circuit, it would have announced its failure with some noise (explosion, fizzing, distorted sound plus hum from the loudspeakers). It is far more likely that a series resistor or choke has failed.

Now that we know what we are looking for, we can switch off, unplug the amplifier from the mains (leaving the plug in plain view), take the covers off and (carefully) investigate. There could still be charged HT capacitors, so it's worth using the DC voltage range of your DVM to check that they are discharged, but be careful when probing inside even a notionally unpowered amplifier, more than one piece of equipment has been destroyed by the slip of a probe. If capacitors still retain charge, a 10 k Ω wirewound resistor is a handy way to discharge them safely (see [Figure 5.3](#)).



Figure 5.3

A 10 k 6 W resistor with leads and insulated crocodile clips is very handy for safely discharging capacitors

Quite apart from the fact that a capacitor's residual charge could give you a nasty surprise, it would interfere with any attempt at resistance measurement by your meter. Bear in mind that electrolytic capacitors suffer badly from dielectric absorption, so leave your meter monitoring their voltage whilst discharging. You must pull the capacitor's voltage to well below 1 V, and this can easily take 20 seconds. Assuming that all the capacitors are discharged, we can use the resistance range to check the HT choke and series HT resistors, and find the problem quickly.

A pre-amplifier example

When the RIAA pre-amplifier is selected, one channel is very low level and distorted. Bear in mind that the cartridge or associated pick-up arm wiring could have failed. One obvious test is to set your DVM to its resistance range and check continuity from tip to sleeve of each phono plug, but this is **not** a good idea. Momentarily passing DC through a moving magnet cartridge is likely to magnetise its core permanently and increase distortion. A far better solution is simply to swap the phono plugs over and see if the fault swaps channels. If it does, the fault is in the arm or cartridge, if it doesn't, the fault is in the pre-amplifier.

If, like the author, you have chosen to use a DIN plug as your connection from pick-up arm to pre-amplifier, you can't easily do this swap (although you could make up a short adaptor lead for swapping channels). There are other ways of testing each channel of the pre-amplifier:

- Tap the input valve of each channel gently and listen to each loudspeaker. Since valves are invariably slightly microphonic, the thump/ting should be equally loud from each loudspeaker.
- Turn up the volume fully, and listen for hiss on each channel. If one channel is significantly quieter than the other, it suggests that the signal from the first stage is not being amplified. The noise should be a clean hiss. Uneven noise suggests a faulty connection.

- With the volume turned fairly well down, put your finger on each input of the amplifier. This should cause a loud hum. (If you do this test at the cartridge pins, beware that the loud hum doesn't make you jump and hit the stylus.)

Having established that the pre-amplifier is genuinely at fault, we know that it is very unlikely to be the LT or HT power supplies (unless built as dual mono) because one channel works, so it's time to look at a circuit diagram (see [Figure 5.4](#)).

Looking at the circuit, we see that there is a lot of fragile silicon (is there any other kind?). We ought to first work out what the circuit is doing. The first stage is pretty conventional, although the LED bias in the cathode is a useful indicator that shows whether the stage is passing HT current. The first stage is followed by conventional 75 μ s passive equalisation, although the 12 k resistor in series with the 270 pF capacitor indicates that 3.18 μ s has also been implemented. The second stage also has LED bias, and has a cascode constant current load (to minimise distortion), and this is direct coupled to a cathode follower. The cathode follower has a simple constant current load and provides a low (and unchanging) output resistance to drive the 3180/318 μ s passive equalisation which is direct coupled to the output cathode follower.

The quickest way to find the fault in this circuit would be to start at the input valve, and measure the DC voltage on the output of each valve with the circuit fully powered. We should be careful not to slip with the probes and create extra faults!

The input valve is a conventional common cathode, so even if we didn't know that the design voltage is 126 V, we could check to see that it is somewhere between $\frac{1}{3}$ HT and $\frac{2}{3}$ HT. When faultfinding, we really don't quibble about precise voltages, we look for things to be "roughly right" or clearly wrong.

Although the anode circuit of the second stage looks complex, we can ignore the silicon for the moment, and just check the

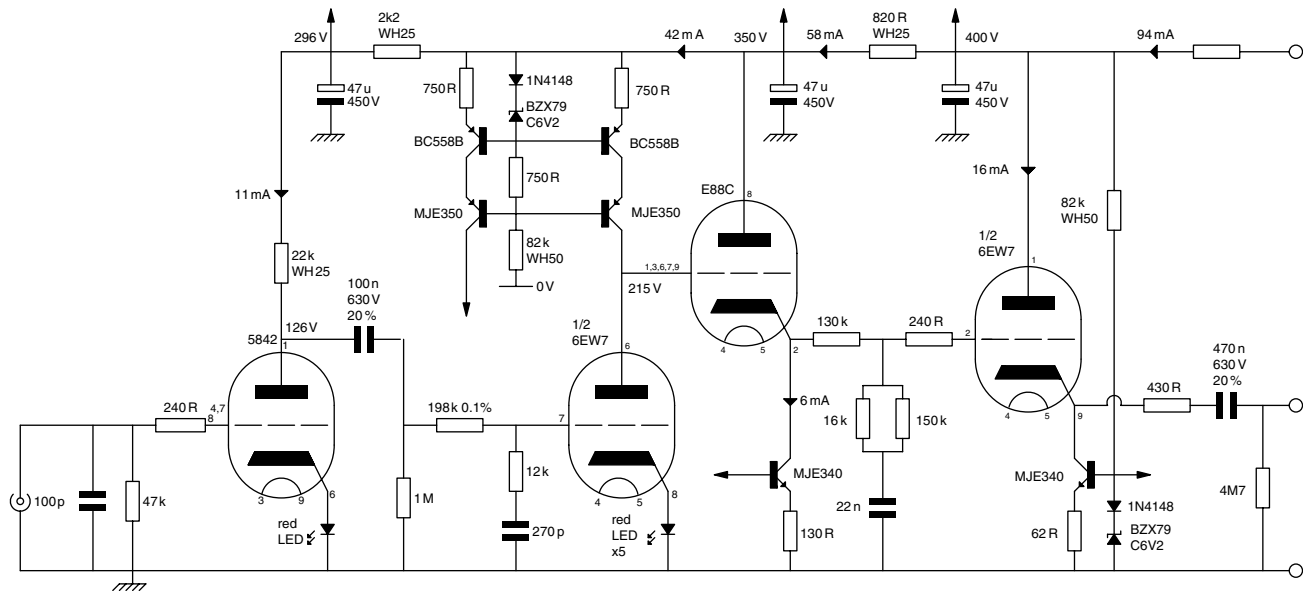


Figure 5.4
 RIAA stage with active loads

anode voltage using the same $\frac{1}{3}$ HT to $\frac{2}{3}$ HT criterion. Because it's so very easy to slip with a DVM probe, it's safer to measure the anode voltage on the valve socket rather than on the collector of the MJE350 (where you could slip and short to either of the adjacent pins). Because the third stage is a cathode follower, its output is on its cathode, and for any cathode follower DC coupled from the preceding anode, its cathode should be a few volts higher than that anode voltage. The fourth stage is also a DC-coupled cathode follower, so its voltage should be a few volts higher than the voltage at the cathode of the preceding cathode follower.

Perhaps when we measure, we find that the anode of the second valve is at 30 V. This is too low and indicates a fault. We also notice that the LEDs in the cathode circuit aren't glowing, and this suggests that the valve is not drawing any significant current. It looks as though we need to investigate the silicon. If the valve is not drawing current, the fault is far more likely to be in the constant current load, so we need to check that the constant current load is being told to do the right thing.

When working with transistors, we generally make the sweeping assumption that the base draws zero current. This assumption hugely simplifies faultfinding because it means that we can predict the voltages looking down the bias chain formed by the 1N4148 diode, BZX79 C6V2 Zener diode, $750\ \Omega$ resistor and $82\ \text{k}\Omega$ wirewound resistor. Because the wirewound resistor is nice and big, it is easy to touch with the DVM's probe. Together, we would expect the two diodes to drop 6.9 V, and in comparison with the $82\ \text{k}\Omega$, the $750\ \Omega$ won't drop much, so we should expect the voltage across the $82\ \text{k}\Omega$ resistor to be about 10 V less than the HT.

Having measured the drop across the $82\ \text{k}\Omega$ resistor, and found that it looks roughly correct, that suggests that the bias chain is correct. We now need to check the transistors. The easiest way to check a transistor is to check the voltage drop across its base-emitter junction, which should be $\approx 0.7\ \text{V}$. It's very tricky

to connect two probes directly to a transistor safely, so we find other, larger, points that are connected to the base and emitter to check the transistors. They turn out to be correct, so we deduce that the $750\ \Omega$ current programming resistor has failed open circuit. Just to be certain, we switch the power off, leaving a meter monitoring the HT, and use a DVM that is guaranteed not to switch diodes on, to measure the $750\ \Omega$ resistor before removing and replacing it. Unfortunately, when we measure, we find that the $750\ \Omega$ resistor is innocent.

Despite having confidently predicted that the silicon circuitry would be at fault, it seems that it is not faulty. If it isn't faulty, then it must be sourcing a current down to $0\ \text{V}$, and if the current isn't going through the valve, it must be going somewhere else. The next check we could make is to measure the voltage across the $750\ \Omega$ current programming resistor. Between them, the Zener and 1N4148 diodes drop $6.9\ \text{V}$ and this voltage is across the base-emitter junction of the transistor plus $750\ \Omega$ resistor. The base-emitter junction drops $0.7\ \text{V}$, so we should expect to see $6.2\ \text{V}$ across the $750\ \Omega$ resistor, exactly the same as the Zener voltage. We measure the drop and it is correct. There's no longer any doubt about it, the constant current load is working correctly, and is delivering $6.2\ \text{V}/750\ \Omega = 8.27\ \text{mA}$; it's just not going where it should.

We now need to look very carefully, perhaps with a magnifying glass and very bright torch to see if there are any whiskers of wire lurking on the second valve's socket, and because it is DC coupled to the third valve, we need to check that too. Cleaning the bases with a stiff brush whilst vacuuming is often a good idea. Unfortunately, even this doesn't clear the fault.

Power transistors have their collectors connected to their case, so when we bolt them to a heatsink, we have to use an insulating kit. The MJE350 dissipates $1.1\ \text{W}$, which is a little more than it can comfortably dissipate without a heatsink, so it was screwed to the chassis. (This is not ideal because it adds $\approx 6\ \text{pF}$ to the

output capacitance of the constant current source.) We unscrew the transistor, lift it clear, and apply power for just long enough to see the five LEDs light up and the anode voltage of the second stage rise to 180 V. The insulating washer was faulty!

Although it is satisfying to find the fault, we must think a little further. What caused the fault? Did it fall or was it pushed? There's no point in replacing a component only to have it fail again three months later. We need to know what caused the failure. The insulating washer **might** have been of faulty manufacture, but the far more likely reason is that it has been pulled against a hole that was not deburred fully. Check the hole on the chassis and the transistor very carefully. Do not re-use the fixing screw – it might have had a sharp burr that damaged the washer.

Unusually, because of the capacitance problem, the best repair would be to fit a small heatsink to the transistor and lose the heat to the air, rather than to the chassis.

These two examples demonstrated that a little thought at the scene of the crime can take you to the guilty component, and that no special test equipment is needed other than a willingness to observe clues and think about what they are telling you. The second example, in particular, demonstrates that we need to think about currents and where they flow. Current doesn't just flow into somewhere and disappear – if it did, there would be untidy heaps of electrons everywhere.

AC conditions

Occasionally, testing DC conditions will not reveal the fault. An oscilloscope to probe around and trace the signal is the obvious approach, but even if you don't have an oscilloscope, the inherent microphony in valves can be useful. Valves can be thumped gently with a screwdriver handle whilst listening to a “disposable” loudspeaker on the output. When the “ting” stops

dead, the area of the fault has been found and components can be replaced or tested until the fault is found.

Diagnosing and eliminating hum

There are various ways that hum can find its way into an amplifier:

- Directly injected from the HT supply
- From heater wiring
- Electrostatic pick-up
- Electromagnetic pick-up
- Hum loop.

Assuming that hum has been discovered, there are quick checks that can be made before dragging out the oscilloscope. Does the hum disappear the instant that you switch the amplifier off, or does it gently fade away? If it disappears instantly, it could be electromagnetic pick-up from the amplifier's mains transformer, or power supply hum – either from the heaters or the HT supply. If it gently fades away, the hum is at the input of the amplifier and could be electrostatic pick-up (poor screening), a hum loop, or it could already be present on the signal entering the amplifier.

Does the volume control affect how loud the hum is? If it does, then the hum is before the volume control. Is the hum only present when one particular input is selected? If so, unplug that input and if the hum goes away, you are looking for hum on that particular source equipment. Does the hum change when you touch the chassis or a lead? If so, you have electrostatic pick-up due to a failed earth connection. Does the hum change when you move a lead? Leads might be screened, but if you have a moving coil cartridge, and trail pick-up arm leads across a mains transformer, you can expect electromagnetic pick-up. Is the hum dependent on which pieces of equipment are plugged into the mains? If so, this suggests a hum loop. Unplug all the audio and mains leads, and for each piece of equipment, check continuity ($<0.5\ \Omega$) between

the earth pin of the mains plug and the body of its phono sockets. In the unlikely event that two pieces of equipment have continuity between mains earth and signal earth, you have a hum loop, and it may be necessary to break the bond between the chassis and 0 V signal earth within a piece of equipment.

Assuming that you have pinned the hum problem down to one piece of equipment, it is now time to switch on the oscilloscope. Trigger the oscilloscope from “line” so that the oscilloscope is always triggered when you poke around looking for hum, then touch the probe tip with your finger to check that the oscilloscope genuinely is triggered, and that it is ready and able to detect hum (see Figure 5.5).

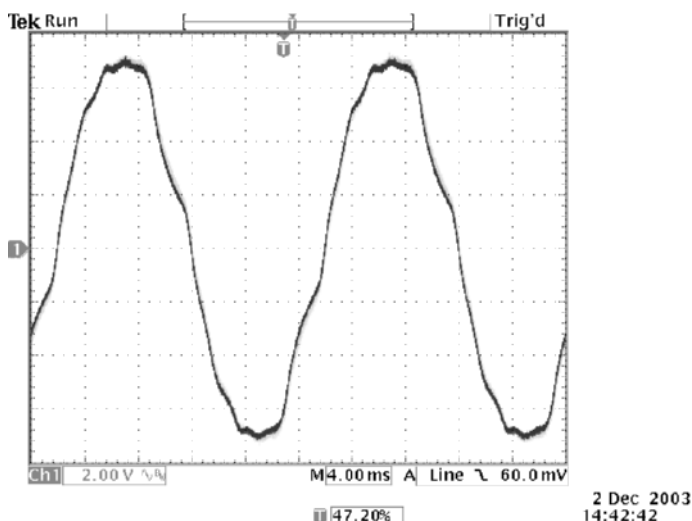


Figure 5.5

Typical hum waveform obtained by touching probe tip with finger

A power amplifier example

Since the patient is a power amplifier, the oscilloscope may be able to show the hum at the output, so it's worth looking, because the type of hum gives a clue as to its origin. Bear in mind whilst investigating that you are unlikely to find squeaky

clean waveforms, they will always be messy – either covered in noise or distorted. If you find something approximating to a sawtooth waveform at the output of the amplifier, you have HT supply hum from somewhere near the reservoir capacitor, and it's likely that the hum is being injected directly into the output stage. The next step is to look for hum on the HT supply, but this needs to be done very carefully.

Switching the input coupling of the oscilloscope to AC will reject the DC, allowing the sensitivity of the oscilloscope to be increased until the hum is clearly visible, but the full HT is still being applied to the probe. Typical $\times 10$ probes can only withstand 200 V, and the HT at a power amplifier output stage is likely to be >300 V. In addition to your $\times 10$ probes, you also need a $\times 100$ probe that is rated for high voltages. Bought new, high-voltage oscilloscope probes are expensive, but you may well be able to find a second-hand probe for a much more reasonable price. If you're looking through a box of old probes, remember that high-voltage probes tend to be rather bulkier than normal probes (see [Figure 5.6](#)).

Assuming that you have a safe means of looking at the HT feeding the output transformer, $<1\text{ V}_{\text{pk-pk}}$ is an acceptable

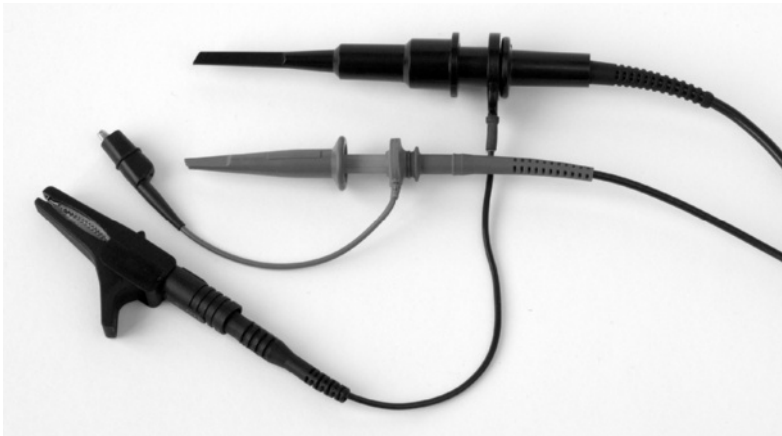


Figure 5.6

High-voltage probes tend to be bulkier than normal ones

ripple amplitude for a push–pull amplifier, but a single-ended output stage is unable to reject power supply ripple, so $<30\text{ mV}_{\text{pk-pk}}$ ripple is preferable. These rough guides are appropriate for conventional loudspeakers, but high-efficiency loudspeakers such as horns are less tolerant, so $<100\text{ mV}_{\text{pk-pk}}$ (PP) and $<3\text{ mV}_{\text{pk-pk}}$ (SE) might be more appropriate. The Quad II has a trap for the unwary in that it has substantial ripple (typically $>60\text{ V}_{\text{pk-pk}}$) at the HT feeding the output transformer, but smooths the HT to the screen grids and relies on tetrode action to reject ripple at the anodes.

Assuming that there is excessive ripple on the HT feeding the output stage, something must be done about it. If it is a new fault on an old amplifier, then an electrolytic capacitor has probably dried out, and it must be replaced by another of similar value and the same, or higher, voltage rating. If it is a new amplifier, a component fault is unlikely, and the design needs to be changed, perhaps by adding a stage of LC smoothing.

A valve microphone example

It's not unusual for hum to come from a variety of sources. Condenser microphones require an amplifier just behind the capsule that not only has high input impedance, but amplifies a very small signal, so eliminating hum is quite a problem. A forty-year-old valve microphone had hum, and although replacing connectors/cables and attending to earth bonds substantially reduced the hum, it did not eliminate it, so attention turned to the HT supply (see [Figure 5.7](#)).

The supply is a conventional bridge rectifier feeding a reservoir capacitor followed by resistor/capacitor smoothing and a neon regulator valve. When the hum at the reservoir capacitor was investigated, instead of it being an even sawtooth, it had alternating large and small teeth, suggesting that one path of the bridge rectifier was less able to charge the capacitor than the other (see [Figure 5.8](#)).

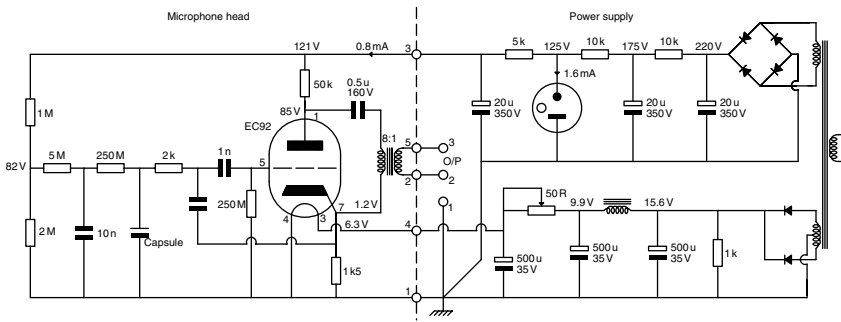


Figure 5.7

The power supply for this valve microphone employs extensive filtering

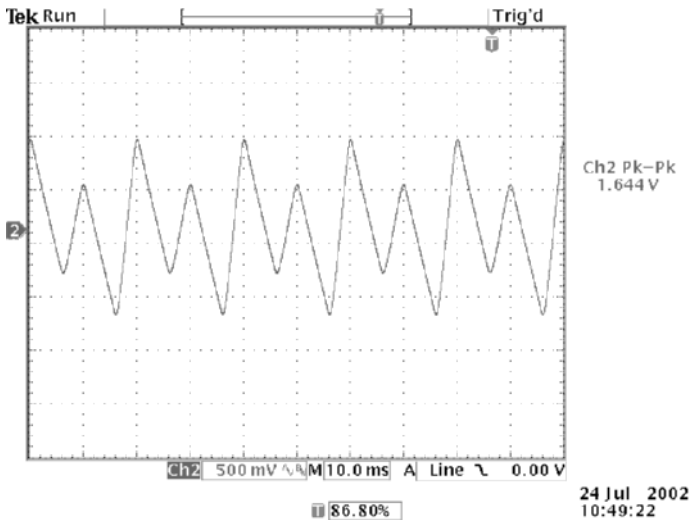
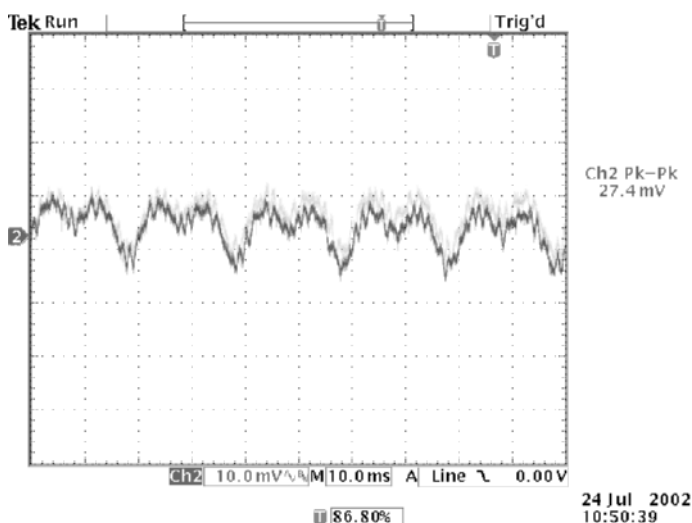


Figure 5.8

The unequal sized teeth on this reservoir capacitor ripple waveform were caused by faulty rectification

Replacing the bridge rectifier would restore an even sawtooth, but would only make a very minor difference to the hum. The hum on the next capacitor down the chain was investigated. The hum on this capacitor should have been an almost pure 100 Hz sine wave, yet noise and spikes were present. The capacitor was simply not doing its job (see [Figure 5.9](#)).

**Figure 5.9**

The hum on this capacitor **ought** to be almost pure 100 Hz, without noise spikes

Given that the power supply was forty years old, it seemed likely that if one electrolytic capacitor was faulty, then they all were, and if they weren't, then they soon would be. Replacing the bridge rectifier and **all** the capacitors made sense...

However, the nastiest fault came to light as the power supply was being gutted of its (chassis mounting) capacitors. Instead of having positive and negative solder tags, each capacitor had a single positive solder tag, and the negative connection was made simply by pressing the aluminium can onto the chassis. Because there wasn't a star washer between the can and the chassis, the mechanical joint was not gas-tight and gradually increased in resistance over the years, so part of the hum was due to this increased resistance. This is a very poor construction technique, and should be replaced on sight (see [Figure 5.10](#)).

Modern capacitors are so much smaller that the entire replacement circuit was built on a small piece of strip board, and this cured the



Figure 5.10

This capacitor could not make a durable low-resistance negative connection simply by contacting the chassis!

hum. Moral: Sometimes there are so many small faults that nothing less than a comprehensive rebuild can effect a full cure.

Oscillation

There are various causes of oscillation, and each tends to produce characteristic frequencies:

- Incorrect polarity of global feedback in power amplifiers tends to produce a loud shriek of oscillation.
- Excessive global feedback around a power amplifier with poor compensation or poor output transformers tends to occur between 30 and 300 kHz.

- Oscillation in individual stages tends to be at radio frequencies, and could be anywhere between 100 kHz and 100 MHz.
- Motorboating is due to feedback around a number of stages via the common power supply, and tends to occur at 1 Hz.

Global feedback and power oscillators

The simplest fault is that a new amplifier has been built, but when the global negative feedback loop is connected, the amplifier turns into an oscillator. If, as is usual, the negative loudspeaker terminal of the amplifier is connected to 0 V, then an amplifier with series applied negative feedback forces the other terminal to become the non-inverted output. This means that in order to maintain correct polarity, oscillation in a push–pull amplifier is most easily cured by swapping the signals to the output valves at the output from the driver stage (see [Figure 5.11](#)).

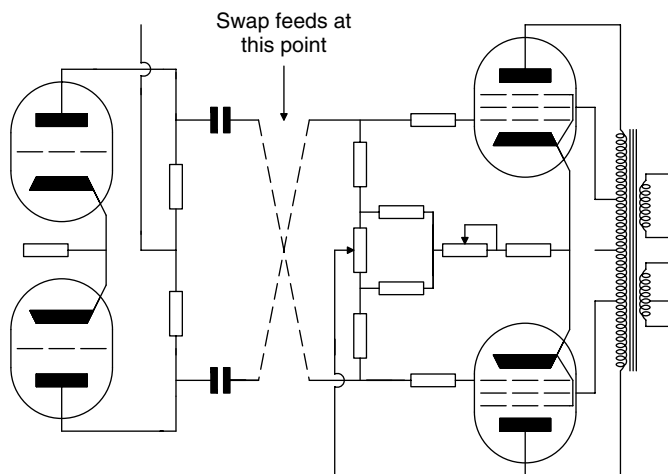


Figure 5.11

Correcting polarity in a push–pull amplifier is most easily done by swapping over feeds at the output of the driver

Oscillation caused by incorrect polarity of global feedback is slightly more awkward to cure in single-ended amplifiers, so it's fortunate that they don't often have it. Because the output

transformer primary end are usually designated anode or HT, the only cure is to swap over the secondary leads (see Figure 5.12).

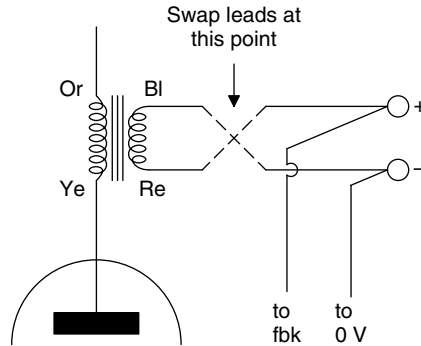


Figure 5.12

Correcting polarity in a single-ended amplifier can only be done by swapping over the output transformer secondary

In theory, we always know the transformer phasing, and we know which way they should be connected, so the previous problems never occur. In practice, the author finds it quicker to wire one channel of a stereo amplifier to one configuration and the other channel to the opposite polarity. This ensures that one channel is wrong and one right. The incorrect channel is thus easily identified and corrected.

Some amplifiers include an output transformer secondary in the cathode circuit of their output valves. Practice shows that it can be quicker simply to wire the cathode feedback winding one way round, apply a signal from an oscillator, and monitor the output on an oscilloscope. Then, without changing any settings, swap the connection of the cathode feedback winding. The connection that produces the lowest output voltage (and no HF oscillation!) is the correct one.

Global feedback and compensation

The amount of global feedback that an amplifier will tolerate is governed primarily by its output transformer, so a new

transformer requires new compensation. For various reasons [1] it is the interaction between the input stage and output transformer that determines the stability of an amplifier when global feedback is applied, so most amplifiers conform to a pattern:

- The anode load R_L is shunted by a capacitor C_1 which may be in series with a resistor R_1 .
- The global feedback resistor (R_{fbk}) is bypassed with a capacitor C_2 which may be in series with a resistor R_2 .

We can therefore draw a generic diagram that shows possibilities for adjusting compensation once R_{fbk} has been set to give the required gain (see Figure 5.13).

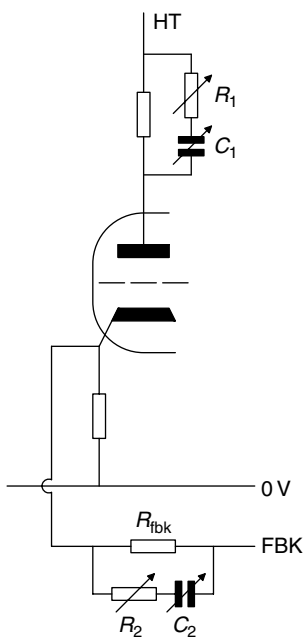


Figure 5.13

Global negative feedback often requires additional compensation components

The way to determine the optimum values is to fit variable resistors and capacitors in the appropriate positions, apply a 10 kHz square wave to the input of the amplifier, monitor the

output across a dummy load made from power metal film resistors – **not** wirewound resistors (which are noticeably inductive at such low values), and adjust the compensation components until the cleanest, sharpest square wave results. Two variable capacitors are needed, and these are best obtained from radio fairs. Any dual-gang variable capacitor that could tune a medium wave radio will do, because they are typically 50–500 pF (best) or 30–365 pF (smaller and slightly newer). For each variable capacitor, connect the two sections in parallel to produce an $\approx 80\text{--}800$ pF variable capacitor (see Figure 5.14).

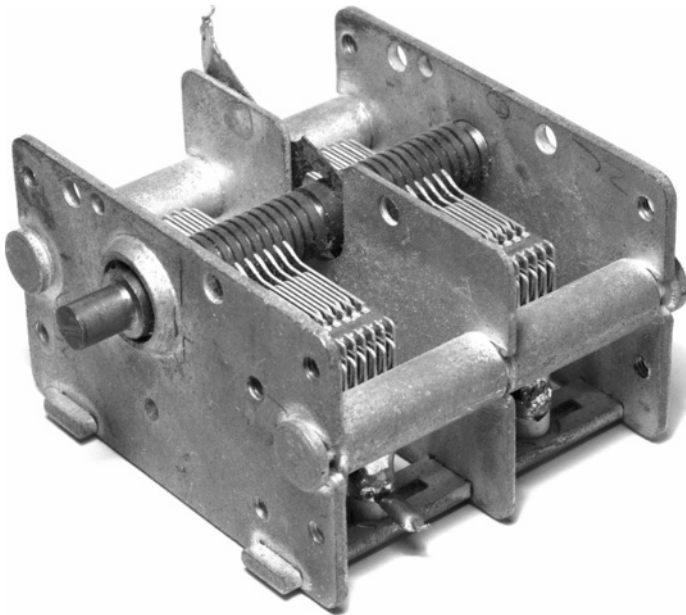


Figure 5.14

A salvaged radio capacitor is ideal for determining optimum feedback compensation

Bear in mind that C_1 and R_1 are connected to high voltages, so they must have insulated knobs and be approached with extreme care.

- The value of C_2 is critical. Too little causes the amplifier to oscillate, too much doesn't significantly lower the amplitude

of the ringing, but rounds the leading edge about which the ringing occurs.

- R_2 may need to be zero to prevent oscillation. A good starting point is $R_{fbk}/10$.
- The value of R_1 is critical for damping the ringing – too little causes overshoot. A good starting point is $R_L/10$.
- The value of C_1 is not critical, but too high a value rounds the square wave, and too little causes overshoot at the leading edge. Don't touch the body of the capacitor while the amplifier is powered.

The quickest way to find the optimum settings is to start without C_1 and R_1 connected. Adjust C_2 to give minimum ringing concomitant with the ringing decaying exponentially and being superimposed on a sharp leading edge, rather than a curving edge. (You will be relieved to learn that this condition takes less time to identify than to describe!)

Switch off the amplifier, and unplug it from the mains, leaving the plug in clear sight. Connect C_1 and R_1 . Apply power to the amplifier, and adjust C_1 and R_1 simultaneously for minimum exponentially decaying ringing on a sharp leading edge (be careful, these components are connected to high voltages). A minor adjustment of C_2 may be necessary. Strive for an optimum setting of all components with a minimum setting of C_2 .

When the input stage is single ended, symmetry between the positive and negative edges can never be achieved because the valve's non-linear anode characteristics mean that it can source current slightly better than it can sink current. However, if the input stage is a differential pair, symmetry **is** possible, provided that the compensation network is connected between the anodes (see [Figure 5.15](#)).

For a differential network, start with $R_1 = R_L/5$, and use only one section of the variable capacitor.

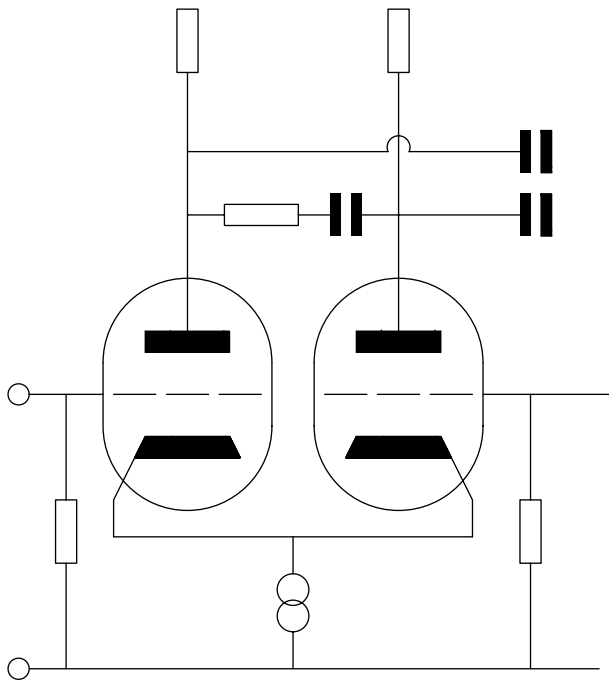


Figure 5.15

A differential pair input stage should connect the compensation components between the anodes

Having carefully set the optimum values, dab a 220 nF (or similar) capacitor across the output of the amplifier whilst observing the 10 kHz square wave. Note that adding the capacitor ruins your carefully optimised square wave response – this implies two things:

- There's no point in struggling to fit precise values of resistors and capacitors in the compensation networks. They are a compromise. This is a **good** thing, because it means that the nearest standard values in your stock will probably do.
- It would be a very good idea if you knew precisely what sort of a load the amplifier was actually going to drive. And it would be even better if the load could be made resistive at high frequencies. We will investigate how this can be done later...

Once you have determined the optimum settings, switch the amplifier off, and having checked that all the capacitors are discharged, use a DVM to measure the value of the resistors, and a component bridge to measure the capacitors. Substitute fixed components and confirm that the amplifier still works as expected. Optimising the compensation components doesn't take long, but the results are worthwhile (see Figure 5.16).

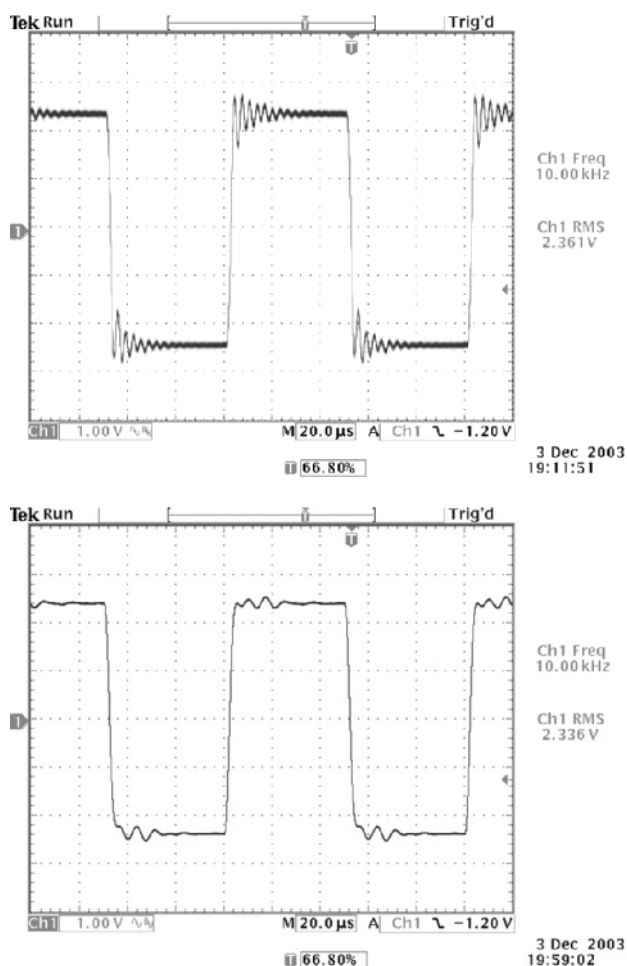


Figure 5.16

Before and after. Before shows significant ringing. Note the slight asymmetry in ringing between positive and negative halves of the square wave.

As an unusual commercial example, the Rogers Cadet III doesn't directly shunt its anode load, but shunts V_{gk} of the second stage, which amounts to the same thing, but avoids injecting noise from the power supply. It bypasses its global feedback resistor with a capacitor, but sets both compensating resistors to 0Ω . In addition, a small amount of neutralisation (positive feedback) has been applied from one output valve's anode to the grid of its counterpart. If the output transformer were perfectly balanced (split bobbin winding), and the phase splitter were perfectly balanced, we would expect to see a similar capacitor from the other anode (see Figure 5.17).

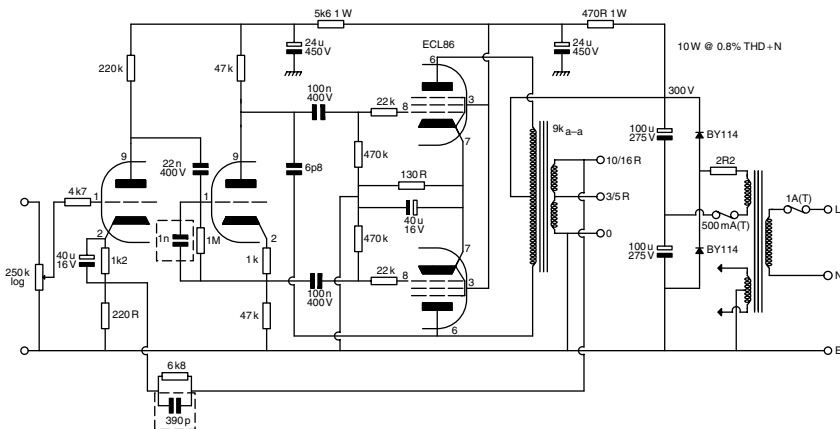


Figure 5.17
Rogers Cadet III power amplifier

Oscillation in individual stages

Cathode followers are particularly likely to oscillate, and when they do so, it can be at a very high frequency, possibly as high as 100 MHz, requiring a good oscilloscope even to show the problem. Fortunately, cathode followers can generally

be tamed quite easily by adding a grid-stopper resistor and possibly a cathode-stopper. The grid-stopper must be a non-inductive resistor (carbon film is ideal) soldered as close the grid pin as possible and between any other circuitry. The value of the grid-stopper resistor must be found by experiment, but typical values range from 1 to 22 k Ω . If needed, the purpose of a cathode-stopper resistor is to buffer a capacitive load from the (slightly inductive) output impedance of the cathode follower. One of the cathode follower's virtues is its low output impedance, so it seems a shame to raise it by adding series resistance. Fortunately, cathode-stopper resistors can usually be quite low value, typically 47–470 Ω . Another cause of local oscillation is poor HT decoupling. Adding a 100 nF capacitor between the top of the anode load and the bottom of the cathode bias resistor often helps (see [Figure 5.18](#)).

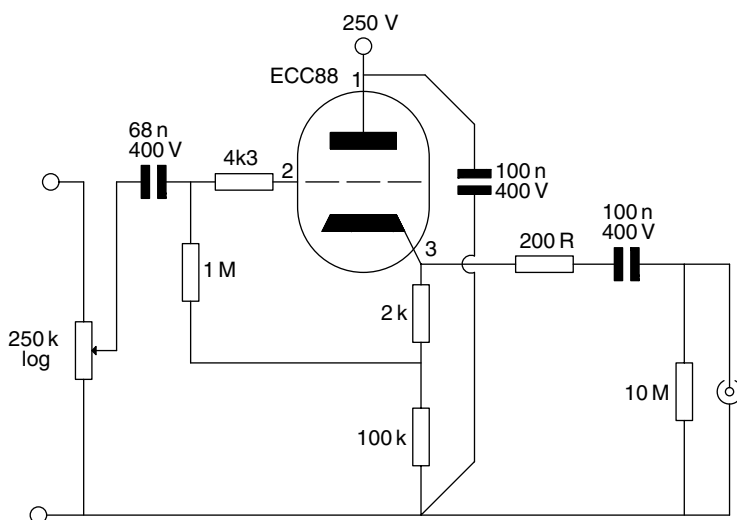


Figure 5.18

Cathode followers are particularly susceptible to RF oscillation, so this example has a grid-stopper, cathode-stopper, and local power supply decoupling

Common cathode stages can also oscillate, and grid-stoppers are an almost universal panacea, but local HT decoupling might also be needed.

Stages can sometimes couple together and oscillate at RF via the heater path. Ideally, each valve should have each heater pin decoupled to **chassis** (not 0 V signal earth) by a 10 nF capacitor having short leads.

Motorboating

Because motorboating is a low-frequency (1 Hz) oscillation due to unwanted power supply coupling, experiment to see if increasing smoothing capacitance at one stage changes the frequency. If you can change something, then you must be near to the source of the fault. The best cure is to reduce the HT source resistance. A regulator is ideal, but reducing an HT series resistor might work. Alternatively, reducing smoothing capacitance might be acceptable, depending on hum. If all else fails, you could resort to the traditional cure of reducing the value of audio coupling capacitors, but this is not recommended – why destroy the audio design because of a power supply problem?

Noise and crackles

Noise and crackles are due to an intermittent conduction path. The most likely cause is a dry soldered joint. Because the fault is caused by intermittent contact, tapping joints with a plastic (insulating) pen can be effective, an even better alternative is to poke the joint with a sharp (insulated) probe in the manner of an unsympathetic dentist (see [Figure 5.19](#)).



Figure 5.19

A sharp (insulated) probe can be useful for finding dry joints

If the amplifier is old, the fingers in the valve socket may be making poor contact with the valve pins, so try moving each

valve gently in its socket. The real cure is to replace the socket, but removing the valve and squirting contact cleaner into the socket (with power off) may effect a temporary cure. Alternatively, using a small probe to tighten the fingers may be a longer-lasting solution.

Although DIN plugs and sockets have a wiping action when the plug is inserted (unlike phonos), the silver plating on the better quality connectors can corrode when a plug or socket is left unmated. The pins in a plug can easily be recovered by sliding the body back to expose the pins and dipping them in Goddard's Silver Dip until they are bright and shiny. Sockets are rather harder, and it is easier simply to replace them.

Luck also plays a part. The author spotted a bulge on the base of an electrolytic HT capacitor and immediately knew that it should be replaced. Unfortunately, one of the wires just wouldn't desolder and a closer inspection revealed that the joint had **never** been soldered! When bought at a market, the amplifier was described as "works well" – but with a production fault like that it could never have worked well in its life, and the unsoldered joint explained the odd crackles heard whilst testing (see [Figure 5.20](#)).

Electrolytic capacitors fail because their electrolyte evaporates, causing poor and possibly intermittent contact to one of the capacitor plates. Old electrolytic capacitors are best replaced en masse, rather than searching to find the noisy offender from a bunch of equally likely suspects.

Intermittent faults

These are the worst to diagnose. They are usually mechanical, so poking around with a sharp insulated probe can be useful. A traditional test for a **semiconductor** was to heat it with a hair



Figure 5.20

Not only is this electrolytic faulty (note the bulge and cracks adjacent to the tag) but its earth wire was never soldered!

dryer to produce the fault, then selectively cool parts with aerosol freezer. This is expensive in environmentally unfriendly freezer, so this technique should only be used as a last resort after probing or thumping with a screwdriver handle. Squirting freezer spray near a valve is likely to crack the envelope.

Classic amplifiers: comments

The following remarks relate to the author's personal experience of a few samples of each amplifier, but the comments are included because some guidance is better than none. Various amplifiers, such as Radford, are not included, not because the author has any bias against them, but simply because he has not owned one.

As a very broad generalisation, classic amplifiers using more expensive output valves are likely to be better. Amplifiers using KT66 may be better than EL34, which will be more powerful than EL84, and ECL82 or ECL86 are at the bottom of the

heap. Curiously, amplifiers using KT88 **may** be worse than any of these, because they may have been designed as public address amplifiers, purely for their high output power.

The quality of output transformers is crucial. Poor output transformers will be small for their rating, although C-core transformers may be an exception to this rule, and are almost always a guarantee of good quality.

Quad II

There are still an awful lot of these about, and they generally require very little work to restore them to their original performance. Typically, the $180\ \Omega$ 3 W cathode resistor and its associated $25\ \mu\text{F}$ capacitor need to be replaced. Quad IIs are popular with tweekers, so there are various modifications. Mostly the modifications replace the GZ34 with silicon to increase output power, and others replace the GEC KT66 with EL34 because NOS GEC KT66 valves are fearfully expensive, although current production KT66 are said to be satisfactory.

Williamson

These were most frequently made by amateurs, so finding a matching pair is tricky, and build quality may be less than wonderful. However, at the very least, the output transformers are well worth salvaging. In the UK, Williamsons used KT66, but some of the American and Australian variants used the somewhat less linear (but much cheaper) 807.

Leak TL12 and BBC LSM/8 derivative

This amplifier has an output stage very similar to the Williamson (triode-strapped KT66 again), and has become

very fashionable (read expensive), but they are likely to be in quite poor condition because of their age.

The BBC LSM/8 lived in a compartment at the bottom of the LSU/10 loudspeaker (where it became very hot). Because it was intended for studio monitoring, the amplifier has a transformer-balanced input and a volume control. In common with many BBC loudspeaker amplifiers, some versions included a bass equaliser. Once these input modifications are removed, the two amplifiers are identical.

BBC amplifiers

Unfortunately, many BBC amplifiers were designed for $25\ \Omega$ loudspeakers and cannot be modified for $8\ \Omega$ without replacing the output transformer. It is also well worth asking why the BBC disposed of the amplifier, particularly if it is thought to have come from an impoverished local radio station. Under the latter circumstances, it is most unlikely that it spent its time cherished in a protective box in a dry cupboard.

Leak TL12+

This is a completely different beast from the TL12, and is very similar to a Mullard 5-10 using EL84 output valves. They are comparatively recent (perhaps only 35 years old), but by modern standards they are noisy, owing to high sensitivity and the EF86 input pentode.

Leak Stereo 20

Far more common than the TL12+, this is almost a pair of TL12+ on one chassis but sharing a slightly under-rated mains

transformer, and with a few other corners cut. Input valve is the dual triode ECC83, but they are still noisy. A pair of TL12+ is preferable unless you are only buying the amplifier for the chassis and transformers. Thanks to the continuous demand by musicians (guitar amplifiers), modern EL84s are cheap and plentiful, but NOS Mullard and Siemens EL84s are quite rare and therefore expensive.

Leak TL10

Similar in design to the Mullard 5-20, but uses a 6SN7 phase splitter. Quite a nice conservatively rated amplifier designed for KT61 (tricky to find), but can be modified for EL34 or 6L6. Although many TL10s were made, fewer seem to have survived than the Stereo 20.

All Leak amplifiers are ridiculously over-sensitive. 110 mV for full power (Stereo 20) is far too sensitive and causes all sorts of hum and noise problems.

Rogers Cadet III

10 W using ECL86 output valves. There were two versions, an integrated version and a separate chassis version. These are ideal for beginners to cut their teeth on, but not of great intrinsic value. Beware that some rather nasty cost-cutting means that cathode bias resistors and decoupling capacitors were shared between the stereo channels, so poor stereo cross-talk is probably due to a failed electrolytic capacitor! Bizarrely, the power supply uses a voltage doubler, and because this imposes high ripple currents, the two doubler capacitors are highly likely to need replacement. Oddly, the disc input stage was very good for its time, but the amplifier needs three ECC807 ($\mu = 140$, and irreplaceable), so if you are considering

buying one, you **must** check that the amplifier has the full complement of undamaged ECC807.

Reference

1. "Valve Amplifiers" 3rd edn. Morgan Jones. Newnes (2003) 0-7506-5694-8.

CHAPTER 6

PERFORMANCE TESTING

The previous chapter enabled us to bring a piece of audio equipment to a state where it could be used safely. In this final chapter, we will assume that we have a working piece of equipment, but that we want to measure precisely how well it works – we might do this for a variety of reasons:

- We hope to verify that it works as designed
- We suspect that it can be improved, so we measure with a view to correction
- We **know** that it works well and want some numbers to boast about.

Linear distortions

We should always measure linear distortion first because an unexpected result could force a minor circuit change affect non-linear distortion. By contrast, it is unusual for any change such as a bias adjustment required as a result of non-linear distortion to affect a linear distortion such as amplitude against frequency response.

Gain

The first measurement is likely to be a gain measurement. We apply a signal of known amplitude at the input of the amplifier and measure the amplitude at the output. The ratio of the two is the **gain** (G), also known as **amplification** (A):

$$\text{Gain} = \frac{V_{\text{out}}}{V_{\text{in}}}$$

The audio band is traditionally considered to range from 20 Hz to 20 kHz, so to avoid errors caused by HF and LF roll-offs, gain is measured in the middle of this band using a sine wave having a frequency of 1 kHz.

If we apply too much signal to the amplifier, the signal at the output will be distorted, making the gain measurement inaccurate. To avoid distortion, and maximise accuracy, we apply a sufficiently small signal to the input of the amplifier to ensure that the output is undistorted. For this reason, gain is known as a **small-signal** measurement.

The whole point of making technical measurements is to provide a prediction of how well the equipment will work when used for its intended purpose, so our test level is important. Professional audio equipment has tightly defined levels, but even domestic audio has reasonably well-defined levels. For example, the original CD standard defined that a CD player reproducing an undistorted sine wave at the maximum level permitted by the digital code (known as 0dBFS) should deliver an analogue amplitude of $2 V_{\text{RMS}}$. If a CD player is capable of delivering $2 V_{\text{RMS}}$, it is not unreasonable to expect a pre-amplifier to be able to cope with this amplitude without distortion, so we could use this as our test level.

Although it is traditional to specify audio voltages in terms of V_{RMS} , and meters are calibrated in terms of V_{RMS} , gain does not

change with amplitude, so there is no reason why we should not modify our test level slightly, to suit our particular test equipment ideally. Thus, if we were using an oscilloscope to measure gain, we would find it more convenient to use a sine wave having an amplitude of $6 V_{\text{pk-pk}}$ (equivalent to $2.12 V_{\text{RMS}}$) because this would occupy exactly six vertical divisions at 1 V/div . (Although most oscilloscopes have eight vertical divisions, an analogue display tube's linearity tends to deteriorate at extreme deflections.)

It doesn't matter which form of measurement we use, just so long as we measure the input and output in the same way:

$$\frac{V_{\text{out}}}{V_{\text{in}}} = \frac{V_{\text{out(pk-pk)}}}{V_{\text{in(pk-pk)}}} = \text{Gain}$$

Dedicated audio test sets

Dedicated audio test sets invariably express their results in terms of **decibels** or **dB**:

$$\text{Gain}_{(\text{dB})} = 20 \log \left[\frac{V_{\text{out(RMS)}}}{V_{\text{in(RMS)}}} \right]$$

It's probably some time since you covered logarithms at school, so at the risk of offending those with better memories, the equation is used by first calculating the voltage ratio, then taking the common logarithm of the result (the "log" button on your calculator), and finally multiplying the result by twenty. Remember: Terms inside brackets first, then functions (logs, trig, etc.), then powers, then multiplication and division, and finally addition or subtraction.

One justification for this apparent complication is that if we have a chain of amplifiers with their gains specified in dB, we can find the total gain simply by adding all the gains (in dB) together. If we want to find the total gain when individual gains

are expressed as ratios, we have to multiply all the gains together, which is somewhat harder.

Note that gain expressed in dB is still a **ratio**, albeit a logarithmic ratio. If we want to use dBs to express an absolute voltage, we must choose a reference voltage and compare our measured voltage to that reference. Many years ago, telecommunications companies sent analogue audio over 600 Ω telephone lines, and in order to express levels in terms of dBs, they defined a reference level of 1 mW dissipated in 600 Ω , and called this level **0 dBm** (m for milliwatt). Power is quite difficult to measure, so their meters actually measured the voltage across a 600 Ω resistor and the meter scales were calibrated in terms of the power dissipated in that resistor. **Nobody** uses 600 Ω any more, yet audio test sets are still saddled with this legacy, even though they usually measure the voltage across entirely different impedances. Because of this, the modern parlance is **dBu**, which corresponds to the same **voltages** as dBm, but ignores the 600 Ω impedance requirement.

We can find the voltage required to dissipate 1 mW in 600 Ω by rearranging $P = V^2/R$ to give:

$$V = \sqrt{PR} = \sqrt{0.001 \times 600} = \sqrt{0.6} \approx 0.775 V_{\text{RMS}}$$

Because this voltage is derived from power, it must be specified in terms of V_{RMS} , and this is why V_{RMS} is popular for audio.

Now we know that 0.775 V_{RMS} would dissipate 1 mW in 600 Ω , we can deem this to be our reference level, and call it 0 dBm if we truly are using 600 Ω impedances, or 0 dBu if we are ignoring impedances.

20 V_{RMS} in dBu:

$$\begin{aligned} \text{dBu} &= 20 \log \left[\frac{V_{\text{RMS}}}{0.775 V_{\text{RMS}}} \right] = 20 \log \left[\frac{20}{0.775 \text{ V}} \right] \\ &= +28.2 \text{ dBu} \end{aligned}$$

Note that when using dBu, it is conventional to state explicitly that the number is positive. Smaller voltages can easily produce negative values...

3.54 mV_{RMS} in dBu:

$$\begin{aligned}\text{dBu} &= 20 \log \left[\frac{V_{\text{RMS}}}{0.775 V_{\text{RMS}}} \right] = 20 \log \left[\frac{0.00354}{0.775 \text{ V}} \right] \\ &= -46.8 \text{ dBu}\end{aligned}$$

It is also conventional to specify AC measurements in dBs to only one decimal place unless you have stunningly accurate test equipment, because it is quite difficult to measure AC more accurately than 1%, and this corresponds to ≈ 0.1 dB.

Conversely, you might have made a measurement in dBs but need to convert it back into a voltage ratio. In which case:

$$\frac{V_1}{V_2} = 10^{\left(\frac{\text{dB}}{20}\right)}$$

16 dB expressed as a voltage ratio:

$$\frac{V_1}{V_2} = 10^{\left(\frac{\text{dB}}{20}\right)} = 10^{\left(\frac{16}{20}\right)} = 6.31$$

-16 dB expressed as a voltage ratio:

$$\frac{V_1}{V_2} = 10^{\left(\frac{\text{dB}}{20}\right)} = 10^{\left(\frac{-16}{20}\right)} = 0.158$$

+16 dBu expressed as an absolute voltage:

$$\begin{aligned}V_{\text{RMS}} &= 0.775 V_{\text{RMS}} \times 10^{\left(\frac{\text{dB}}{20}\right)} = 0.775 V_{\text{RMS}} \times 10^{\left(\frac{16}{20}\right)} \\ &= 4.89 V_{\text{RMS}}\end{aligned}$$

–16 dBu expressed as an absolute voltage:

$$V_{\text{RMS}} = 0.775 V_{\text{RMS}} \times 10^{\left(\frac{-16}{20}\right)} = 123 \text{ mV}_{\text{RMS}}$$

Note that the sign is absolutely critical, and that it is very important to be clear about whether you are using dB as a ratio or dBu as an absolute level. In addition to dBu, other audio references are in common use:

<i>Nomenclature</i>	<i>Reference level</i>
dBm	1 mW into 600 Ω
dBu	0.775 V_{RMS}
dBV	1 V_{RMS}
dBW	1 W into 8 Ω
dBFS	Maximum undistorted sine wave amplitude permitted by digital code (Full Scale)

Peak programme meter scales

Peak programme meter (PPM) scales were briefly mentioned in Chapter 4, but we now need to look at the scale and understand the logic behind it. The PPM was originated by the BBC for measuring the level of live programme material and there were various requirements:

- The programme material would subsequently be presented to analogue tape machines or transmitters employing amplitude modulation, neither of which forgive overload.
- The meter would be read and interpreted by operators who were primarily concerned with making artistic adjustments to achieve the best sound balance.
- The operators would be working under pressure, possibly with poor lighting.

Taken together, these requirements mean that the meter must be easy to read, and have a clear, uncluttered scale. An equal

increment logarithmic scale of 4 dB/div means that meter deflection is proportional to loudness (see Figure 6.1).



Figure 6.1

The PPM scale is 4 dB/div, with 0 dBu = PPM4 (This is a stereo PPM with a pair of back-to-back movements, hence the pair of pointers.)

White lettering on a black background makes the scale easy to read. 0 dBu = PPM4, and since there are 4 dB/div, PPM1 = -12 dBu and PPM7 = +12 dBu. Normal practice is to use 0 dBu as line-up level, and to mix programme material so that peaks do not exceed PPM6, or +8 dBu. Although UK broadcasters use the 1–7 scale, European broadcasters mark the scale -12 to +12. The meters are identical in all other respects [1].

When a PPM is used for engineering, all controls are referenced to PPM4. Thus, the level of the incoming signal can be read directly from the controls provided that they have been adjusted to set the PPM to read PPM4. PPMs intended for engineering use generally also include an expanded scale in addition to the PPM scale, and may include other scales, making them more cluttered, but they're not being used to mix live programmes (see Figure 6.2).

The unfortunate consequence of the 600 Ω legacy

The reason that the telecommunications companies used 600 Ω is that their equipment drove cables that were so long that they really were transmission lines at audio frequencies. As an

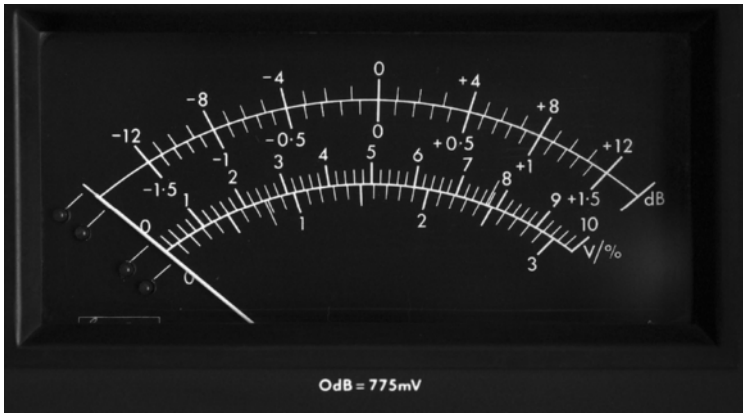


Figure 6.2

Engineering PPMs always include an expanded scale in addition to the PPM scale, and may include other scales

extreme example, there used to be a cable carrying programme from Bush House in London to the transmitter at Daventry (67 miles away) without any intervening amplification. Because the cables were transmission lines, impedance matching was important, so the cables' **characteristic impedance** of $600\ \Omega$ had to be driven from a $600\ \Omega$ source and be **terminated** by a $600\ \Omega$ load. As a consequence, exchange (telecomms) and studio (broadcast) plant (electronic equipment) was all designed for $600\ \Omega$ impedance matching. The significance of this is that the input resistance of the destination formed a potential divider in conjunction with the source resistance driving it (see Figure 6.3).

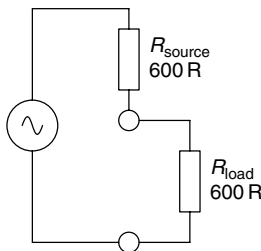


Figure 6.3

The legacy of $600\ \Omega$ analogue telecommunications is a 6 dB potential divider to trap the unwary

Using the potential divider equation:

$$\frac{V_{\text{out}}}{V_{\text{in}}} = \frac{R_{\text{lower}}}{R_{\text{upper}} + R_{\text{lower}}} = \frac{6000\ \Omega}{600\ \Omega + 600\ \Omega} = 0.5 = -6\ \text{dB}$$

The modern technique is to make all devices have a high input impedance and a low output impedance. Thus:

$$\frac{V_{\text{out}}}{V_{\text{in}}} = \frac{R_{\text{lower}}}{R_{\text{upper}} + R_{\text{lower}}} = \frac{00\ \Omega}{0\ \Omega + 00\ \Omega} = 1 = 0\ \text{dB}$$

We can now see that the modern technique avoids a wasteful loss of 6 dB. The consequence is that if you set the attenuators on a 600 Ω audio oscillator to produce a specific output level, its output will be 6 dB high unless it is loaded or **terminated** by 600 Ω . Unless you know your oscillator, it is wisest to **measure** the signal at the input of the amplifier, rather than rely on the attenuators.

Source and load impedances for the Device Under Test (DUT)

Having ensured that the oscillator is correctly terminated, we must also ensure that we terminate our DUT correctly.

If we are measuring a simple valve pre-amplifier, it might have an output impedance of $\approx 6\ \text{k}\Omega$ and expect to see a 1 M Ω load. A typical audio test set has an input impedance of 100 k Ω on its “high” setting, so this would cause an additional loss of 0.53 dB compared to a measurement with the correct 1 M Ω loading. Conversely, measuring with a 10:1 oscilloscope probe (10 M Ω input resistance), would theoretically give a reading 0.054 dB high, but since this corresponds to +0.6%, it is comparable with oscilloscope error, so we would probably never notice it.

If we measure a power amplifier, we should terminate it with a load resistance appropriate to the secondary setting. Thus, if we have set the secondary to match a 4 Ω load, we need a 4 Ω resistor, often known as a **dummy load**. Since we will almost certainly attempt to determine maximum output power, the resistor must be capable of withstanding this power without damage.

Unfortunately, wirewound resistors are **not** suitable as dummy loads because values $< 1\text{ k}\Omega$ have considerable inductance compared to their resistance. Fortunately, 50 W non-inductive thick film resistors are available, and multiples of these can be screwed to a large heatsink to make a perfectly satisfactory dummy load (see [Figure 6.4](#)).

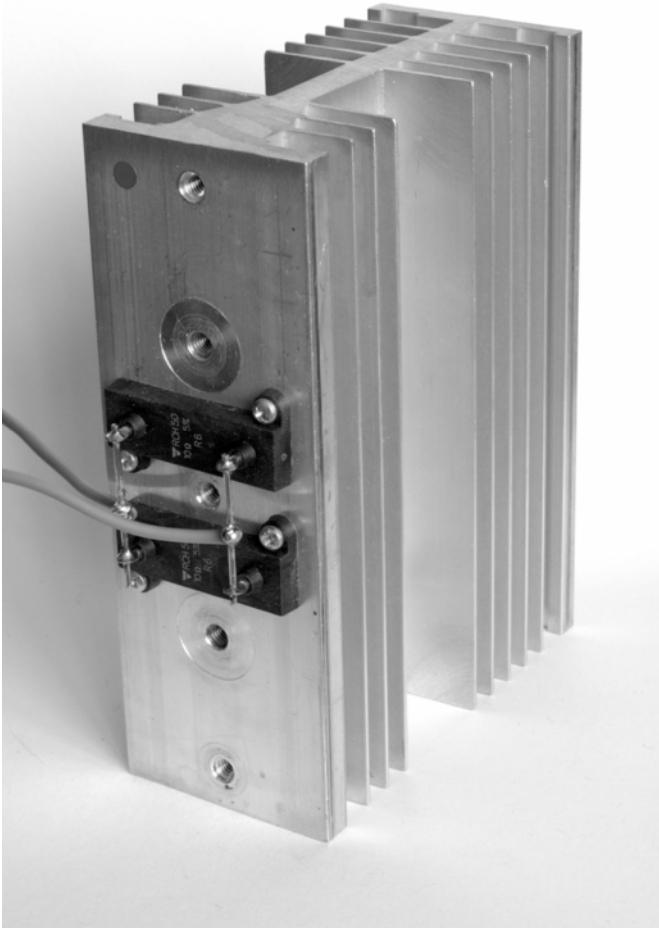


Figure 6.4

A dummy load for testing power amplifiers can be made from thick-film resistors screwed to a large heat sink

Whether the DUT is a power amplifier or a pre-amplifier, it should be driven from an appropriate source resistance. Modern audio electronic equipment such as a CD player typically has an output resistance of $\leq 600\ \Omega$. Some audio oscillators have selectable output resistance, whereas others are fixed, but the highest output resistance is generally $600\ \Omega$, so any of these settings is appropriate. However, if you know that you will drive your DUT from a higher source resistance, perhaps from the output of a simple $100\ \text{k}\Omega$ logarithmic volume control, a $25\ \text{k}\Omega$ series resistor should be added to emulate the maximum expected source resistance.

Measuring gain at different frequencies

Once we have taken the trouble to measure the gain correctly at 1 kHz, it is easy to leave all other controls alone, then change frequency, and take additional measurements to produce a graph of amplitude against frequency, which is habitually called “frequency response”.

When we measured gain, it was equally valid to present the result as a pure voltage ratio (14) or in dBs (22.9 dB), but the purpose of measuring the amplitude against frequency response of an amplifier is to see if the amplitude **deviates** significantly from the 1 kHz level. This observation has two important implications:

- Since we are worried about the audibility of amplitude deviations, we should use a logarithmic measurement because the ear/brain responds logarithmically. We should either measure directly in dBs or convert our measurement into dBs.
- Absolute gain (or loss) is irrelevant, only deviations from the 1 kHz reference level matter. If we measured 1 kHz gain on a cathode follower pre-amplifier, we might have applied 0 dBu to the input of the DUT and measured $-0.6\ \text{dBu}$ at the output. Rather than do lots of arithmetic, it is far easier to

adjust the oscillator output level to produce precisely 0 dBu at 1 kHz at the **output** of the DUT, then (leaving oscillator level unchanged) measure output level at different frequencies and obtain the amplitude against frequency response directly.

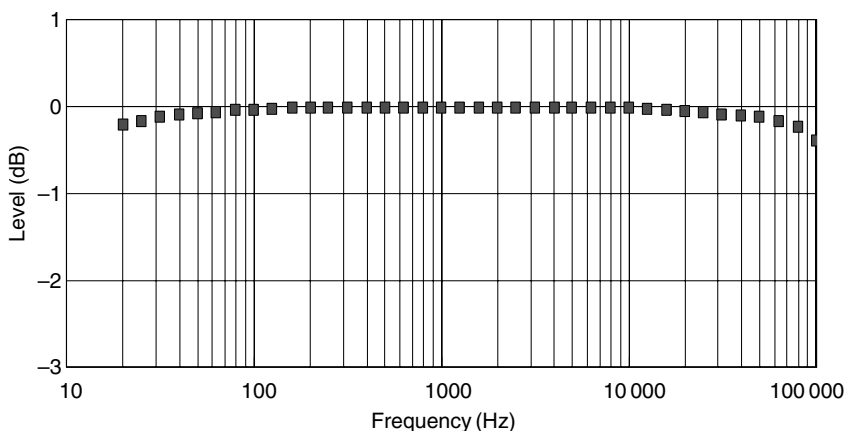
Which frequencies to use?

Because the ear/brain responds logarithmically, audio amplitude against frequency graphs are plotted on paper having a linear vertical scale (remember that dBs are already logarithmic) and a logarithmic horizontal scale. The most useful graph has equally spaced points, rather than a desert punctuated by oases of closely clustered points. Thus, we should choose measurement frequencies that produce equally spaced points on logarithmic paper. We could use the oscilloscope attenuator logarithmic sequence of 1, 2, 5 to give equally spaced points, resulting in test frequencies of; 20 Hz, 50 Hz, 100 Hz, 200 Hz, 500 Hz, 1 kHz, 2 kHz, 5 kHz, 10 kHz, and 20 kHz. Unfortunately, the points are rather sparse, so a better choice is to use International Standards Organisation (ISO) recommendation $266\frac{1}{3}$ octave frequencies:

20 Hz, 25 Hz, 31.5 Hz, 40 Hz, 50 Hz, 63 Hz, 80 Hz, 100 Hz, 125 Hz, 160 Hz, 200 Hz, 250 Hz, 315 Hz, 400 Hz, 500 Hz, 1 kHz, 1.25 kHz, 1.6 kHz, 2 kHz, 2.5 kHz, 3.15 kHz, 4 kHz, 5 kHz, 6.3 kHz, 8 kHz, 10 kHz, 12.5 kHz, 16 kHz, 20 kHz.

Twenty-nine points are thus needed to cover the range from 20 Hz to 20 kHz, and these frequencies would be entirely appropriate if we were measuring an RIAA pre-amplifier because their errors are perfectly capable of producing deviations at almost any point across the audio band. (Incidentally, fuse values follow the same logarithmic sequence; 2 A, 2.5 A, 3.15 A, etc.)

However, if we were measuring a typical power amplifier, this would be a wasted effort because it has a low frequency roll-off, a high frequency roll-off, and is flat in between (see [Figure 6.5](#)).

**Figure 6.5**

Crystal Palace amplifier amplitude against frequency response taken at an output power of 1 W

ISO $\frac{1}{3}$ octave frequencies (and their multiples beyond 20 kHz) were used, resulting in a graph with evenly spaced points. As can be seen, the amplitude response against frequency is very nearly flat, and taking the measurements was tedious. When measuring a valve power amplifier, it is better to sweep low frequencies slowly up to 50 Hz, to find any low-frequency bumps, then sweep from 10 to 500 kHz to spot undamped HF resonances due to the output transformer, and if no unpleasant resonances are found, simply determine the LF and HF frequencies for which the response is either 1 dB or 3 dB down compared to the 1 kHz reference level.

An even faster method than sweeping is to apply a 1 kHz square wave and check for negligible ringing at the leading edge, then apply a 100 Hz square wave and select DC coupling at the input of the oscilloscope before checking that the top and bottom horizontal lines are truly **horizontal**, rather than sagging (sag indicates LF loss). Assuming that the square wave response is satisfactory, the -3 dB points can be determined using sine waves. If it isn't, there's not a lot of point in plotting

a pretty graph of a HF or LF aberration, the problem needs to be corrected!

Plotting the graph

Having recorded our results neatly, we now need to plot them. Considering the horizontal scale first, logarithmic graph paper is specified in terms of cycles and starts at 1. Thus, three-cycle logarithmic paper can cover 10 Hz to 10 kHz, but we want to start at 20 Hz and finish at 20 kHz, necessitating four-cycle paper.

The vertical scale needs to be linear, and typical paper has minor divisions every millimetre and major divisions every centimetre. You were probably taught at school to choose a graph scale that used all of the graph paper, so you might find that a total range of ± 2 dB to be sufficient to encompass your measurements. On typical paper, this approach would give a scale of 0.25 dB/cm, yet the measurement is probably only accurate to ± 0.1 dB, so the very large graph would imply false accuracy. Worse, it could make a perfectly reasonable response look poor – an audibly flat response ought to **look** flat on paper. A sensible vertical scale is 2 dB/cm on typical office paper. This scale produces a graph that correlates well to the audible effects and does not display false accuracy.

Lin/log paper is expensive, and paper ideally suited to audio is unobtainable, so the author generated his own graph paper using a CAD package (see [Figure 6.6](#)).

Measuring RIAA equalisation using an inverse RIAA network

It is difficult to make a practical RIAA stage having perfect equalisation because valves vary from sample to sample and their characteristics change as they wear. It is therefore useful

to have an inverse RIAA network that can be connected between an oscillator and the input of an RIAA stage so that the combination theoretically yields a flat amplitude response against frequency. If the inverse RIAA network is perfect, any deviation from flatness is due to the RIAA stage, so adjustments can be made until flatness is achieved.

Unfortunately, designing a “perfect” inverse RIAA network is not trivial. Factors that must be taken into account are:

- The network is sensitive to source and load resistances, so it needs to be matchable to common oscillator output resistances.
- The output resistance of the network should be similar to that of a typical cartridge, so that the (perhaps varying?) input impedance of the RIAA stage loads it correctly.
- Unfortunately, the typical technique of designing an ideal RIAA equalising circuit, preceding it with the inverse RIAA, and emulating it in an electronic analysis package tends to generate calculation errors. This is because the network analysis equations that can cope with general problems are more complex and more numerous than the optimised equations for a specific problem.
- Capacitors are typically only available in 1%, but resistors are available in 0.1%, so errors should be limited to the capacitors. Further, it would be convenient if standard capacitor values could be used.

With these considerations in mind, the author started with the fundamental RIAA equation, and added the $3.18\mu\text{s}$ time constant that has to be implemented at the time of cutting to protect the (probably Neumann) cutting head. The RIAA replay equation is therefore:

$$G_s = \frac{(1 + 318 \times 10^{-6} \times s)(1 + 3.18 \times 10^{-6} \times s)}{(1 + 318 \times 10^{-6} \times s)(1 + 75 \times 10^{-6} \times s)}$$

Where $s = j\omega$, and $\omega = 2\pi f$.

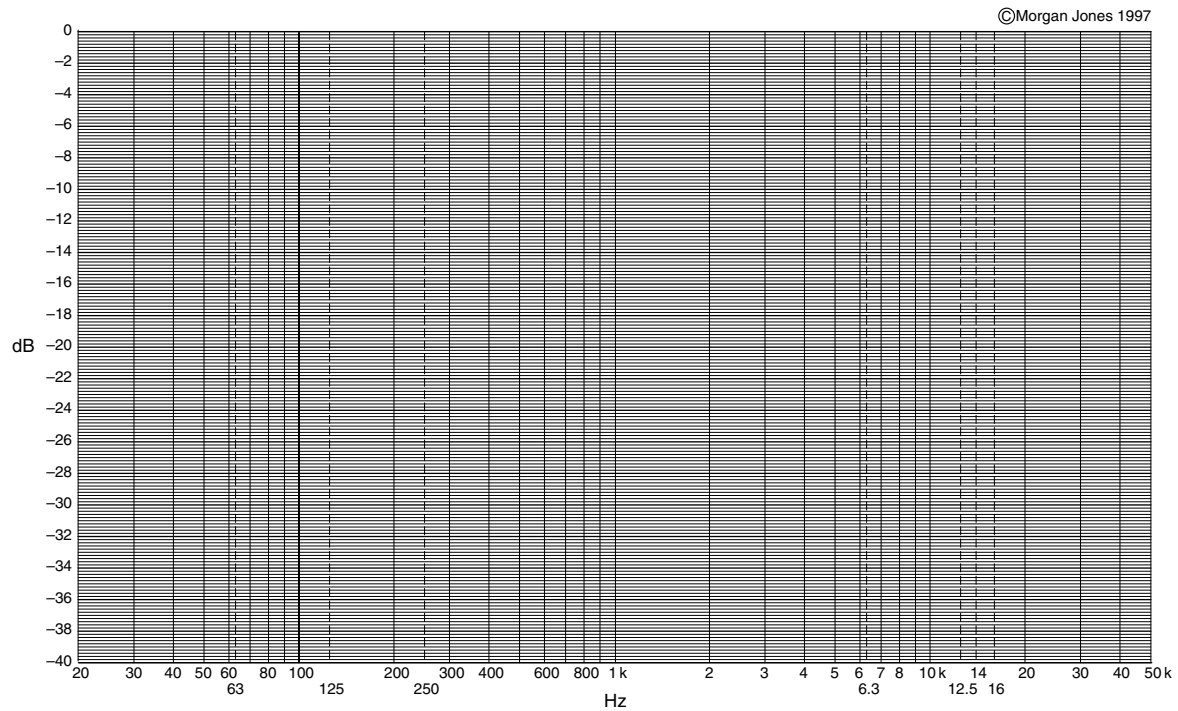


Figure 6.6

Lin-log graph paper is difficult to obtain, so photocopy this “audio” paper

This Page is Intentionally Left Blank

Rather than re-invent the wheel, an existing inverse RIAA network developed by Jim Hagerman [2] from Stanley Lipshitz's [3] seminal AES paper was investigated, and the equation for its attenuation was written from first principles and entered into a spreadsheet. The fundamental RIAA replay equation was also entered into the spreadsheet, enabling the two equations to be multiplied together to generate a graph which should ideally show zero error (see [Figure 6.7](#)).

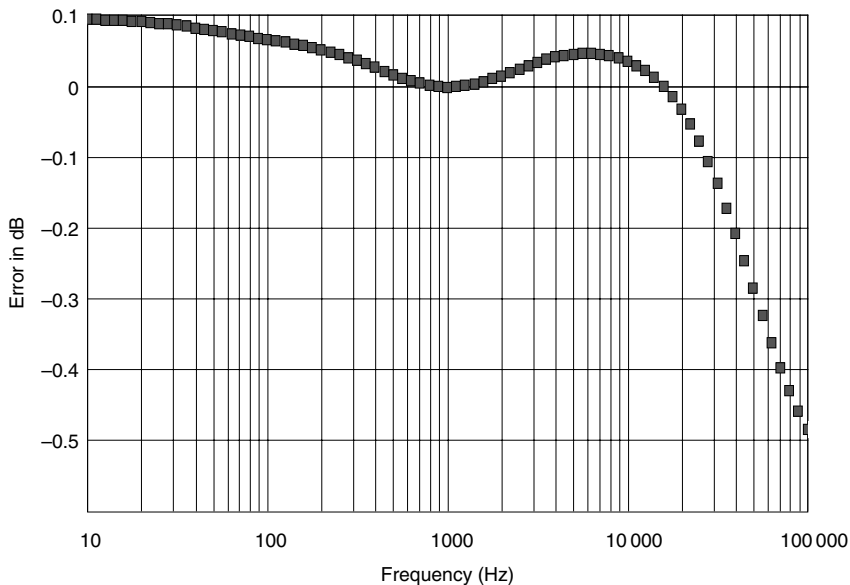


Figure 6.7

The error curve for this commercial RIAA inverse network is quite respectable despite being restricted to E24 component values

Despite being restricted by E24 component values, the response is surprisingly good over the audio band. Nevertheless, every BBC audio test set from the (valve) ATM/1 onwards is capable of measuring to better than 0.1 dB, so it seems worthwhile to improve matters.

The error worsens towards 100 kHz because the original Hagerman article used $3.5\ \mu\text{s}$, rather than $3.18\ \mu\text{s}$. Although the less popular

Ortofon cutting heads used $3.5\text{ }\mu\text{s}$ rather than the Neumann's $3.18\text{ }\mu\text{s}$, it is the author's experience that it is always preferable to use too little equalisation than too much. Changing component values to set $3.18\text{ }\mu\text{s}$, and making the fine adjustments possible by choosing 0.1% tolerance resistors available in E96 values significantly improves the design errors (see [Figure 6.8](#)).

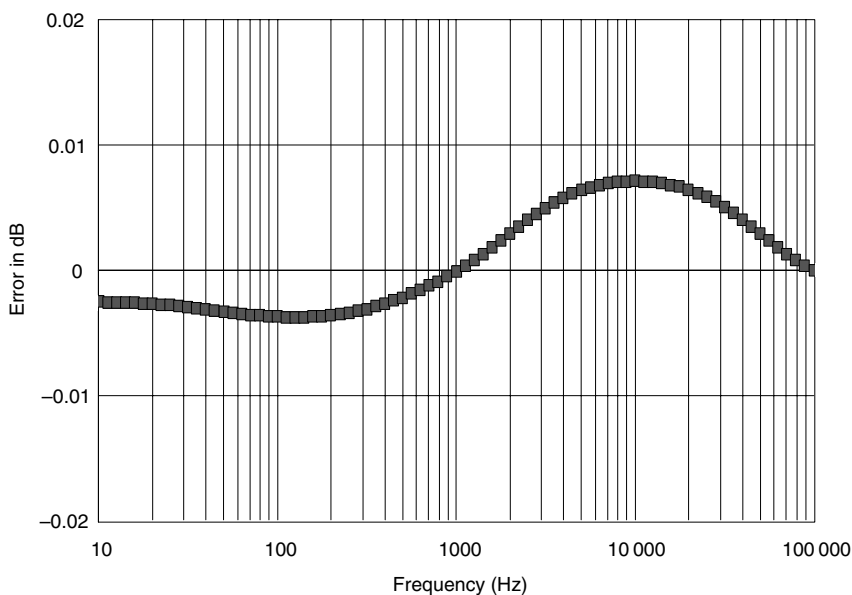


Figure 6.8

Using a mix of E96 and E24 resistor values allows the design error to be reduced greatly

The design errors using single E24 capacitors and combinations of E96 and E24 resistors are now comfortably within $\pm 0.01\text{ dB}$, so the effect of practical components can be investigated. Since the capacitors are only available in 1% tolerance, the worst possible combination of capacitor tolerance extremes was investigated (see [Figure 6.9](#)).

Even the worst possible combination of 1% capacitor tolerances leaves the predicted error better than $\pm 0.1\text{ dB}$. Because even a cheap function generator intrinsically has superb range

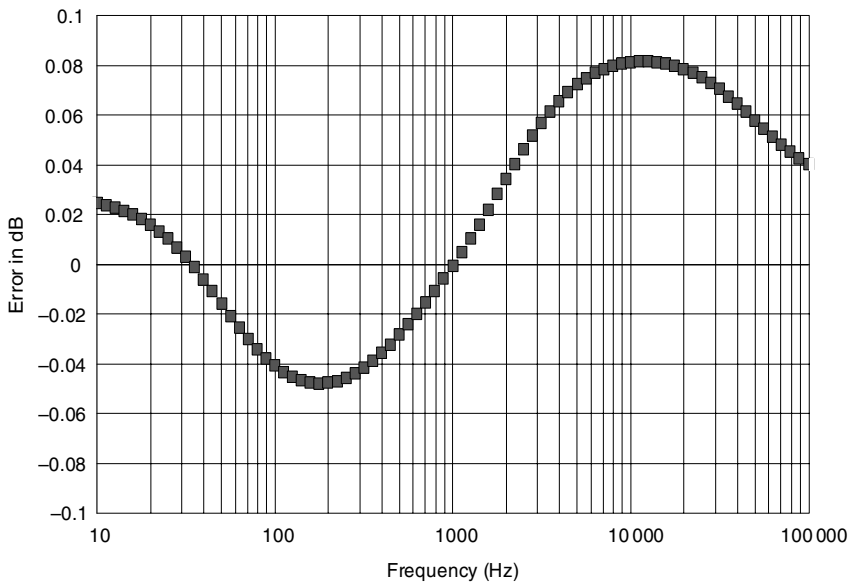


Figure 6.9

The worst possible combination of capacitors errors produces this frequency response error

flatness and constant output resistance with frequency, network values were optimised to suit the output resistance of a typical function generator (see [Figure 6.10a](#)).

Whether your function generator has an output resistance of $50\ \Omega$ or $75\ \Omega$ will determine whether you need a $47\ \Omega$ or $22\ \Omega$ resistor respectively. The further advantage of matching the network to a function generator is that it allows square wave testing at the flick of a switch.

If the wires between the output of the oscillator and the network are more than a few inches long, and the function generator is flat to 10 MHz, transmission line effects cause an overshoot on the leading edge of the square wave that can be misleading. If a longer cable cannot be avoided, use $50\ \Omega$ co-axial cable, and modify the input of the network so that it (almost) correctly terminates the cable (see [Figure 6.10b](#)).

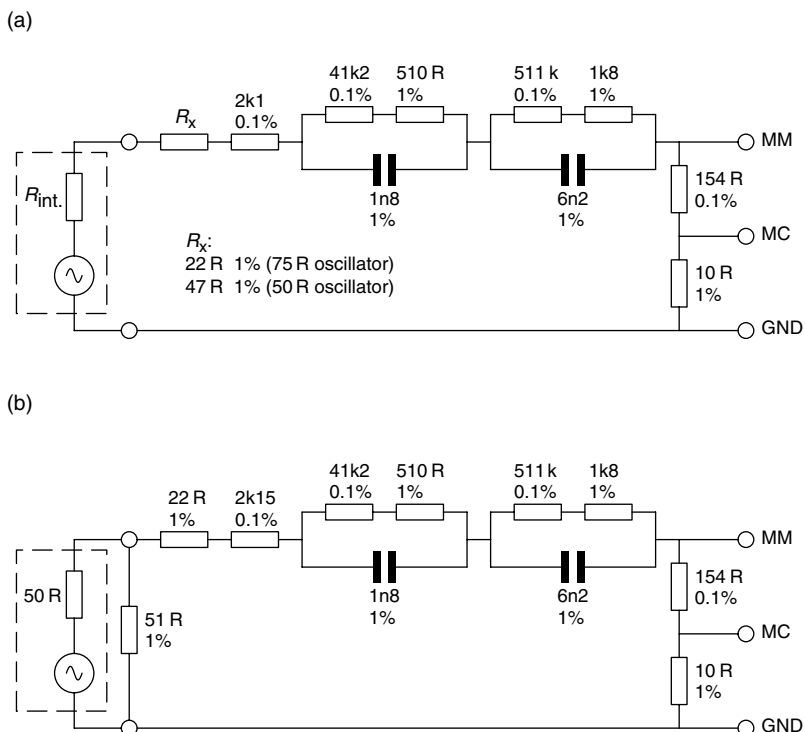


Figure 6.10
Final circuit of inverse RIAA network

This modification eliminates the overshoot at the expense of adding a slight tilt at the leading edge of a 10 kHz square wave. Normal audio leads can be used to connect the output of either version to the pre-amplifier under test.

Impedance against frequency

It can be extremely useful to measure **impedance** against frequency. Power amplifiers are carefully designed by plotting loadlines that cunningly maximise output power and minimise distortion, yet we connect them to real loudspeakers having wildly varying impedance. Thus, we might wish to determine

the actual impedance of our chosen loudspeaker in order to optimise its matching to the amplifier, or to see if it has any unusual features that might cause problems. (A peak in impedance at HF could cause instability in an amplifier.)

Alternatively, we might want to determine the output impedance of a power amplifier to verify that it is sufficiently low not to disturb the Thiele-Small [4] parameters of an associated loudspeaker, causing an unwanted bump at low frequencies.

Measuring the impedance of a loudspeaker

A single magnetic loudspeaker has a non-flat impedance curve for two reasons:

- The moving mass (usually a cone) and its suspension (usually a spider) combine to produce a low-frequency mechanical resonance which is transformed through the motor into a peak in impedance. A typical moving coil bass driver has a resonant frequency between 20 and 50 Hz in free air. Tweeters resonate between 500 Hz and 1.5 kHz, but the resonance may be so heavily damped (perhaps by Ferrofluid®) as to be unobservable.
- As implied by the name, moving coil drivers have a coil which possesses inductance, causing impedance to rise with frequency. Although the moving element of a ribbon tweeter is very nearly resistive, the necessary impedance-matching transformer has leakage inductance, producing an impedance curve very similar to a moving coil tweeter.

Although it is theoretically possible to compensate the impedance curve for the low-frequency mechanical resonance, it requires inconveniently large component values. Fortunately, compensating for voice coil inductance **is** practical, requiring

a simple CR Zobel network across the loudspeaker terminals (see Figure 6.11).

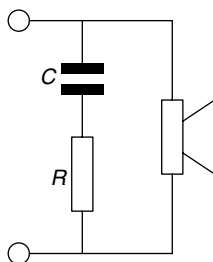


Figure 6.11

A simple Zobel network can compensate for a loudspeaker's voice coil inductance

In theory, the required resistance is equal to voice coil DC resistance, and the required capacitor is found using:

$$C_{\text{Zobel}} = \frac{L_{\text{voice coil}}}{R_{\text{DC}}^2}$$

Most manufacturers specify voice coil inductance, enabling the required capacitance to be calculated easily. In practice, the value is only a guide, so it is better to connect the calculated Zobel network across the driver, then trim its values until a flat impedance results. A variable capacitance box makes life much easier, and they are surprisingly cheap second-hand (see Figure 6.12).



Figure 6.12

Capacitance (and resistance) boxes are surprisingly cheaply available second-hand, and are very useful

The easiest way to measure loudspeaker impedance is to make a potential divider from your oscillator's output resistance and the loudspeaker's impedance, then measure its attenuation (see Figure 6.13).

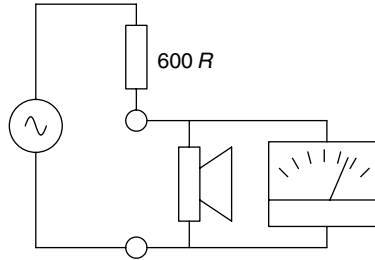


Figure 6.13
Measuring loudspeaker impedance

We make the assumption that the oscillator has a constant output resistance with frequency, then use the potential divider equation in reverse to find the lower resistance:

$$\frac{V_{\text{out}}}{V_{\text{in}}} = \frac{R_{\text{lower}}}{R_{\text{upper}} + R_{\text{lower}}}$$

Rearranging:

$$Z_{\text{loudspeaker}} = \frac{R_{\text{oscillator}}}{\left(\frac{V_{\text{loudspeaker}}}{V_{\text{oscillator (open circuit)}}} - 1 \right)}$$

Note that the attenuation term has been inverted, and that we must measure the output voltage of the oscillator without any loading, **before** connecting the loudspeaker across it. Since an awkward calculation is required to derive the impedance, and it is quite likely that V_{in} and V_{out} were actually measured in dBu but must be converted back into absolute voltages, a spreadsheet is ideal both for the calculations and for plotting the graph. As an example, a modern ribbon tweeter was measured before and after a Zobel network was added (see Figure 6.14).

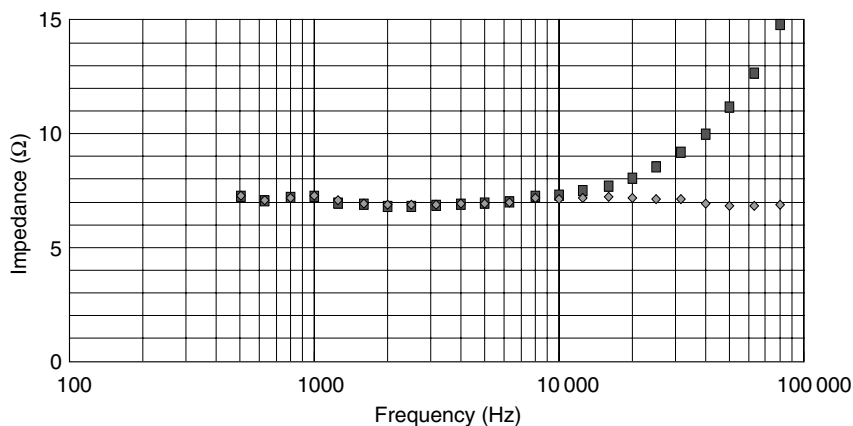


Figure 6.14

Impedance against frequency of a ribbon tweeter before and after compensation with a Zobel network

The graph shows two plots: one naked, and one showing the effect of a parallel Zobel network consisting of an 8R2 resistor and 500 nF capacitor that converted a rising impedance into a very nearly resistive 7 Ω load.

Purely resistive loudspeaker loads are desirable because:

- Non-resistive loads demand current that is out of phase with the applied voltage. Deviation from a perfectly straight load-line in the amplifier is likely to cause distortion.
- Reactive loads can provoke HF instability in some amplifiers.
- Unless the amplifier has **zero** output impedance, a varying load impedance with frequency creates a potential divider having attenuation varying with frequency. This results in an amplifier/loudspeaker combination that has an amplitude response that varies with frequency (because the impedance of the loudspeaker varies with frequency). Even if the errors are quite low amplitude (<1 dB), broadband errors are very noticeable, so they must be avoided.

Measuring the output resistance of a pre-amplifier

The output resistance of a pre-amplifier can be measured very easily by setting a convenient open-circuit output voltage, then loading it with a variable resistance, and adjusting that resistance until the output voltage halves. At this point, the load resistance is equal to the source resistance, so the load resistance is removed and measured with a DVM, or if a resistance box is used as the load, the output resistance can be read directly (see Figure 6.15).

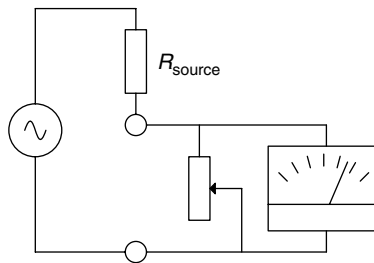


Figure 6.15

A variable resistor across the output of a pre-amplifier allows output resistance to be determined quickly and easily

When used with $\times 10$ probe, an oscilloscope has an input resistance of $10\text{ M}\Omega$, which is certainly high enough to be able to measure the open-circuit voltage of any practical pre-amplifier. So if we were to adjust the signal at the input of the pre-amplifier to produce a $6V_{\text{pk-pk}}$ sine wave on the oscilloscope, it would have a vertical sensitivity of 1 V/div . If the sensitivity were changed to 0.5 V/div , and the variable resistor adjusted to restore the original deflection of six vertical divisions, then the output voltage of the pre-amplifier would have been halved. The variable resistor could now be removed and its measured value would be equal to the output resistance of the pre-amplifier.

Measuring the output impedance of a power amplifier

Power amplifiers are designed to have very low output impedances (ideally a fraction of an Ohm). If the previous method were used to measure the output impedance of a transistor power amplifier, the amplifier would be destroyed or would shut down. We need a different method. The way to measure a power amplifier's output impedance is to use a variation of the method used for determining loudspeaker impedance (see [Figure 6.16](#)).

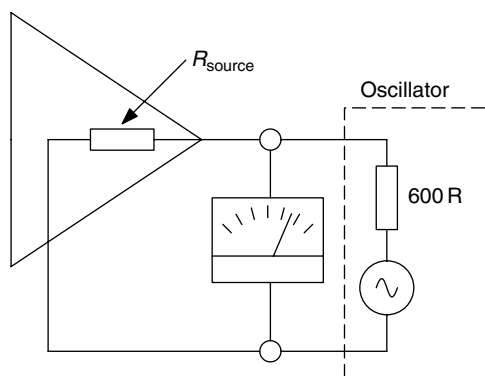


Figure 6.16

To measure a power amplifier's output resistance, we connect it across an oscillator and see how well it can attenuate the oscillator

We short-circuit the **input** of the amplifier to ensure that it cannot produce a signal. We then connect a meter directly across the output terminals and also connect our oscillator across the output terminals. Because the amplifier is designed to have almost zero output impedance, it should be able almost completely to attenuate the output of the oscillator. We use the same equation for determining amplifier output resistance as we used for the loudspeaker:

$$Z_{\text{amplifier output impedance}} = \frac{R_{\text{oscillator}}}{\left[\frac{V_{\text{amplifier}}}{V_{\text{oscillator (open circuit)}}} - 1 \right]}$$

Disappointing results tend to fall out from this measurement (see Figure 6.17).

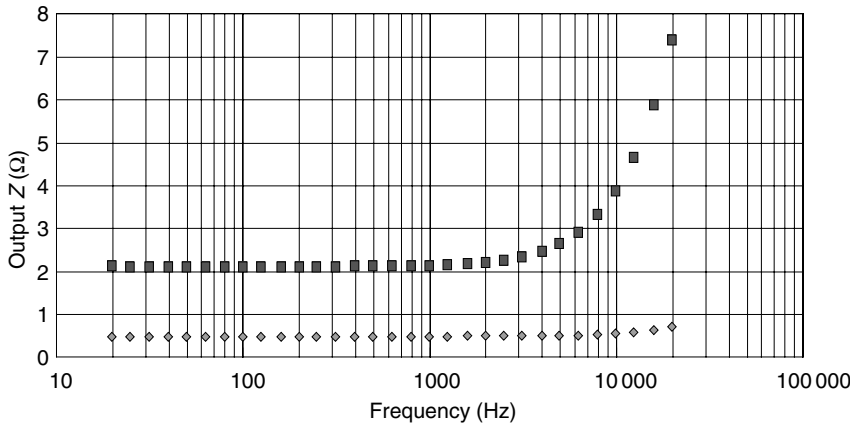


Figure 6.17

Output impedance against frequency for two power amplifiers. Lower: Crystal Palace push–pull amplifier with global feedback. Upper: Scrapbox Challenge single-ended amplifier with zero global feedback. Both amplifiers were configured to drive an 8Ω load

The output impedance is invariably higher and changes more with frequency than would be expected from a simple consideration of output valve r_a transformed by the square of the output transformer turns ratio. The reasons for this are:

- Valves in practical circuits rarely achieve the stunningly low values of r_a quoted by manufacturers on their data sheets.
- The output transformer secondary resistance and reflected primary resistance are in series with r_a .
- The leakage inductance of the output transformer is in series with r_a .
- Feedback amplifiers must reduce their feedback at high frequencies to maintain stability. Feedback reduces output impedance, so if feedback reduces with frequency, its ability to reduce output impedance is reduced, and output impedance rises with frequency.

Precautions must be observed when using this method of measuring output impedance.

The test set-up must not be plugged or unplugged with the power amplifier switched on, or there is a danger of short-circuiting the output of the amplifier (which can destroy a transistor amplifier). The oscillator must be set to produce a high output voltage, but the power amplifier attenuates the oscillator so effectively that the meter/oscilloscope must be set to a high sensitivity. If the amplifier is switched off without first decreasing the sensitivity of the meter/oscilloscope, the resulting gross overload could damage the meter/oscilloscope.

Another possible problem is that the oscillator may object to driving a high level into a short circuit, and its current limit may operate. Sweep the entire audio frequency band to ensure that it doesn't current limit, and if it does, back the level off.

The test makes the assumption that the output resistance of the oscillator is a constant resistance. Some oscillators may not conform to this ideal, so select the lowest source resistance available and add an external series resistor. The value of the resistor depends on the individual design of the oscillator:

- Dedicated audio test sets were designed to be able to drive $600\ \Omega$ music lines, and a constant output resistance with frequency was part of their design requirement. They do not need an external resistor.
- Modern bench audio oscillators tend to base their design on operational amplifiers, so they have an internal output resistance of zero which is built out by a series resistor to give the specified output resistance. They are unlikely to need an external resistor.
- Old bench audio oscillators are likely to use a transformer to couple to their output terminals. Their output resistance is not necessarily constant. If in doubt, select the lowest output

resistance, and add a series resistor of ten times that value to swamp any variations. (The series resistor should be precise and non-inductive, so a metal film type is ideal.)

Non-linear distortions

Non-linear distortion implies that the output is not simply a scaled replica of the input, but that additional frequencies have been added as a result of non-linearity within the DUT.

Maximum output power and distortion versus level

One of the most important specifications of a power amplifier is the amount of power that it can deliver into a specified load impedance. There is plenty of scope for argument in defining what constitutes “maximum power”. Possibilities for defining maximum power are:

- The power developed at, or just before, the point when clipping of the output waveform can be observed on an oscilloscope.
- The power developed when distortion reaches an arbitrary value. Perhaps 10%, perhaps 1%, perhaps even 0.1%.
- The power developed just before the point when distortion begins to rise sharply.

The first definition is quite useful for amplifiers that employ plenty of global feedback, as distortion tends to be very low until clipping occurs, whereupon it rises catastrophically. Even so, the definition is a little fluffy.

The second definition is popular when measuring zero-feedback single-ended amplifiers, although adjusting the criterion from 3 to 10% is a popular ploy for increasing the perceived value of an amplifier by substantially increasing its measured power

output. At the opposite extreme, Leak specified amplifiers by the power they could deliver with 0.1% distortion.

The final definition requires the plotting of a graph of distortion against frequency. This is easy with a modern automated audio test set, but harder with an older, manual test set.

The amplifier is loaded with a dummy load, and the resulting distortion is measured by the audio test set. In addition, the monitoring output of the test set feeds the distortion waveform to an oscilloscope (see [Figure 6.18](#)).

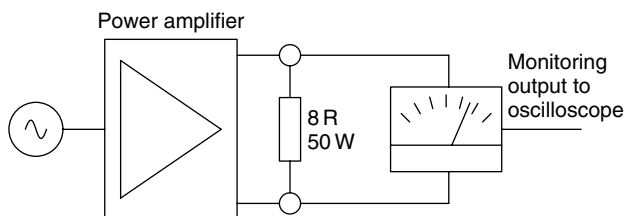


Figure 6.18
Determining maximum output power

Fortunately, when testing valve amplifiers, there is a very distinct point that is very sharply defined and therefore readily identifiable, so it can be used to define maximum output power. We are not looking for a particular distortion figure (after all, that can be adjusted by negative feedback), rather, we are looking for an abrupt change in the distortion waveform that signifies grid current in the output stage (see [Figure 6.19](#)).

Comparing the two distortion waveforms, the second waveform has an additional valley at one of its peaks. This valley appears and disappears very sharply with level, so it is ideal for defining maximum power. In this particular test of the author's Crystal Palace amplifier, it defined maximum output power at 1 kHz to be 47 W with 0.56% THD. The amplifier delivered 52 W with 1.5% THD, just below observable clipping on an oscilloscope.

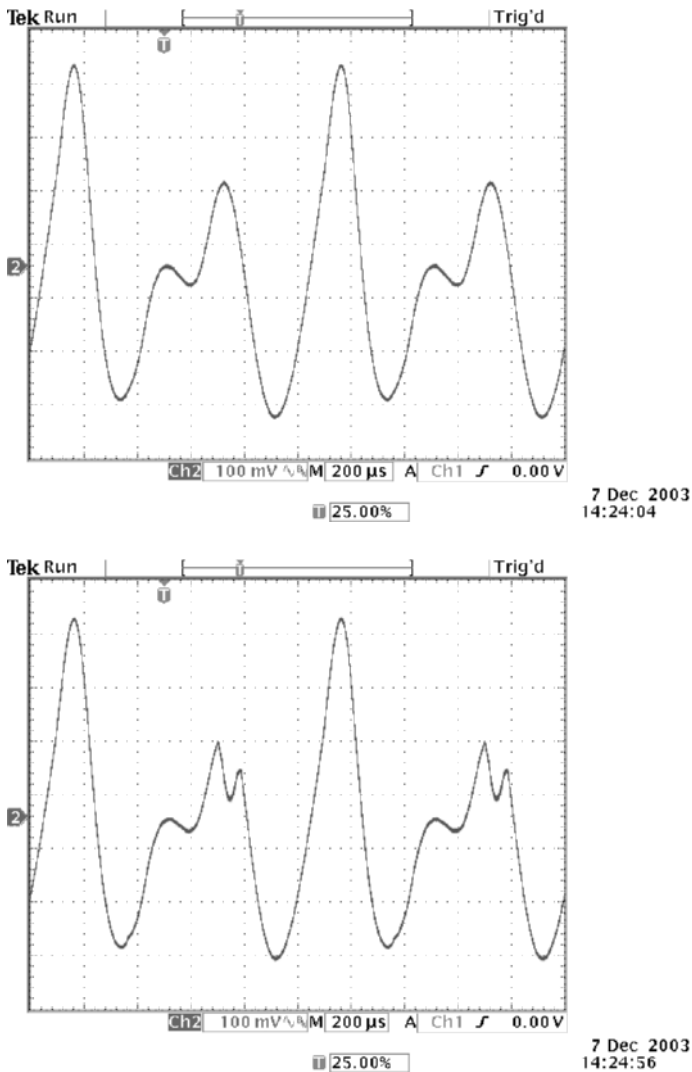


Figure 6.19

A change in output power of only 0.2 dB heralded grid current as evidenced by these distortion waveforms. No grid current (upper), entering grid current (lower)

Measuring distortion against output power is fiddly with a manual test set, but manageable with care. Calculate the level in dBu for 1 W into the test load resistance. Adjust level until the power amplifier develops 1 W. Then, increase level in 1 dB

steps, recording the amount of distortion at each point until gross distortion results. You can then use a spreadsheet to convert the levels in dBu into power, and plot a graph of distortion against output power (see Figure 6.20).

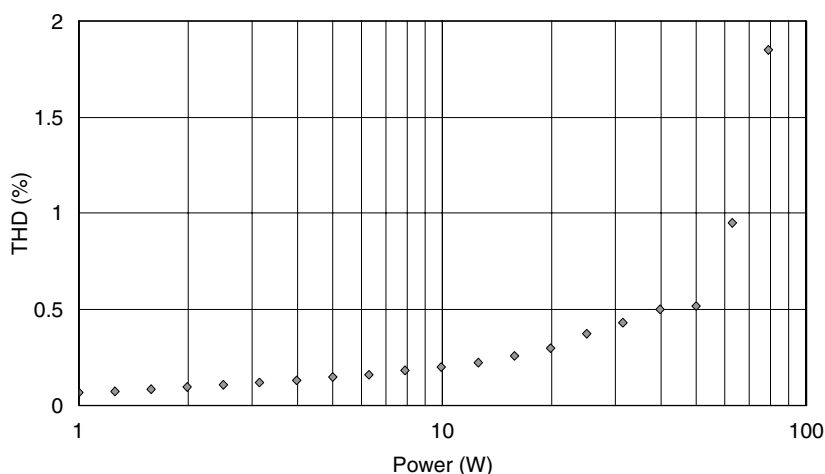


Figure 6.20

Distortion against output power at 1 kHz

Once again, the reason for using a logarithmic scale for output power is that the ear perceives loudness logarithmically. Looking at the graph, we see a smooth curve with rising distortion and what appears to be a discrepancy at 50 W. Considerable effort and number-crunching was required to produce this graph, yet it only tells us what we already knew, and does it less precisely. It seems that plotting distortion against output power is not very useful for finding maximum output power.

Another justification for plotting distortion against output power is that if the distortion rises at low levels, this can be an indicator of crossover distortion. Unfortunately, as all amplifiers produce noise, practical measurements of

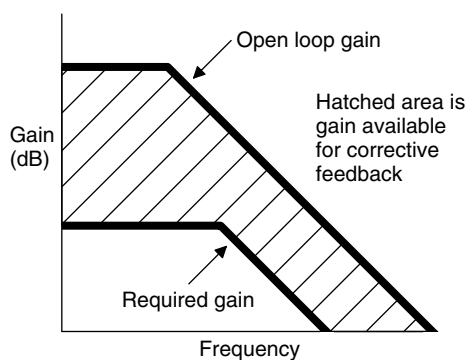
distortion are not so much THD as $\text{THD} + \text{N}$ (THD plus noise). As output signal level falls, the noise becomes more significant, so all amplifiers suffer deteriorating $\text{THD} + \text{N}$ at low levels.

A sudden change in the **distortion waveform** with level is a powerful indicator of a problem. At high levels, the sudden appearance of a dip indicates grid current (valid for power amplifiers and small-signal amplifiers). Conversely, at low levels, the gradual appearance of triangular spikes as level falls indicates crossover distortion. If you designed the circuit you are testing, you already know which features to look for. If you're testing a single-ended amplifier, there is no need to look for crossover distortion because it is an exclusive failing of push-pull Class AB amplifiers.

Distortion against frequency

A popular test with feedback amplifiers is to test distortion against frequency. The reason for this is that if the amplifier relies on global negative feedback to reduce distortion, the gain before feedback is applied must be maximised, but to maintain stability, this gain must fall with frequency. But feedback can only correct by the ratio between the open loop gain and the required closed loop gain. This means that if the open loop gain is 80 dB and the required closed loop gain is 20 dB, 60 dB of feedback is available for correction. But if only 40 dB of open loop gain is available, and the required closed loop gain is still 20 dB, only 20 dB of correction is available, so distortion must rise. Summarising, all feedback amplifiers must have distortion that rises with frequency (see [Figure 6.21](#)).

Transistor power amplifiers often rely on global feedback to lower distortion, so plotting distortion against output power can be quite revealing. Valve power amplifiers almost invariably

**Figure 6.21**

Open loop gain and available corrective feedback

incorporate an output transformer, so they are typically limited to ≈ 30 dB of feedback before stability becomes distinctly questionable.

Measuring distortion against frequency has to be done very carefully if it is to bear any relation to audible effects. If we measure distortion at 1 kHz, then even the 20th harmonic (20 kHz) is theoretically within the audio band, but if we measure distortion at 10 kHz, only the 2nd harmonic is within the audio band. If the measurement is intended to correlate with audible effects, a 20 kHz low-pass filter must be used, but this would have the effect of making an amplifier producing predominately 3rd harmonic distortion to have falling distortion with frequency after 6 kHz.

Conversely, if we were measuring an amplifier with a view to possible improvement, we would not use the 20 kHz low-pass filter, as this colours the results. As an example, a 10 W valve amplifier employing substantial global feedback was tested at 10 W (see [Figure 6.22](#)).

As should be expected from a high feedback amplifier, the distortion is very low at 1 kHz, but rises with frequency in accordance

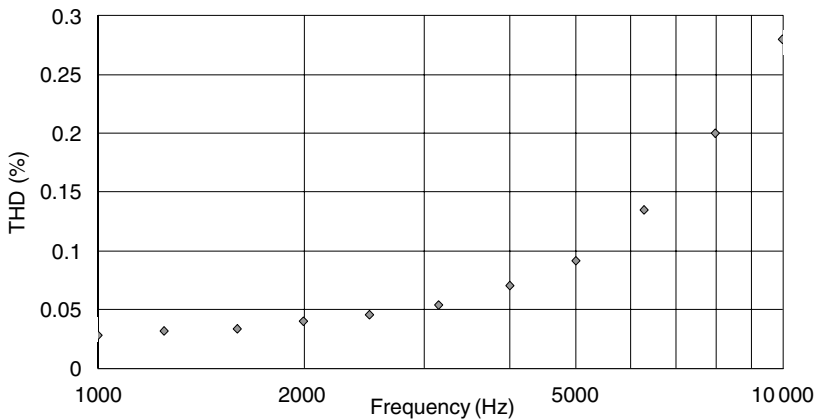


Figure 6.22

This 10 W high feedback amplifier exhibits textbook rising THD with frequency

with the falling open loop gain required to maintain HF stability. In this example, the measured distortion at 10 kHz was **exactly** ten times that at 1 kHz, so the distortion was 20 dB worse at a frequency one decade higher. In other words, the distortion rises at 20 dB/decade, which is the same as 6 dB/octave, and is **exactly** the slope produced by the CR compensation network required to maintain stability. The graph of distortion against frequency mirrors the graph of open loop gain against frequency.

It seems that plotting distortion against frequency doesn't really tell us very much that we didn't already know. However, it does emphasise that over-zealous compensation increases high-frequency distortion. What the graph doesn't tell us very clearly is whether the amplifier might suffer from slewing distortion.

Slewing distortion

A capacitor can change its voltage instantaneously, but it needs an infinite current to do so. At a more practical level,

the faster we want to change the voltage across a capacitor, the more current is required. All amplifier stages have input capacitance, so when the preceding stage attempts to change the voltage at the input of the next stage, it must change the voltage across this shunt capacitance. In small-signal terms, the shunt capacitance forms a low-pass filter in conjunction with the output resistance of the preceding stage (see [Figure 6.23](#)).

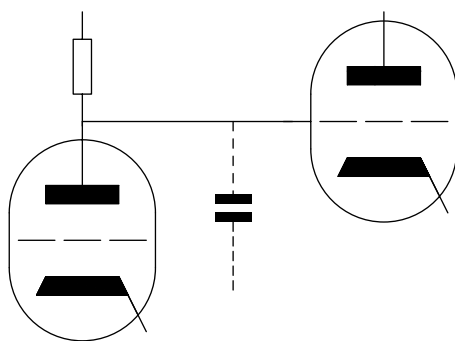


Figure 6.23

Shunt capacitance driven by inadequate current causes slewing distortion

When we attempt to impose a large voltage change quickly, the amount of current that the capacitor can draw may be limited by the quiescent current in the preceding stage. Under this condition, the capacitor charges slower than the input waveform changes voltage, and this distortion is known as **slewing distortion**. Slewing distortion became particularly noticeable when early operational amplifiers were used for audio, but valve amplifiers are by no means immune to the problem. Fortunately, the problem is easily detected, if not quite so easily cured.

To test for slewing distortion, drive the amplifier to full power at 1 kHz, then increase frequency to 10 kHz and beyond. At some

point, the sine wave will begin to look more like a triangular wave, and this is evidence of slewing distortion (see Figure 6.24).

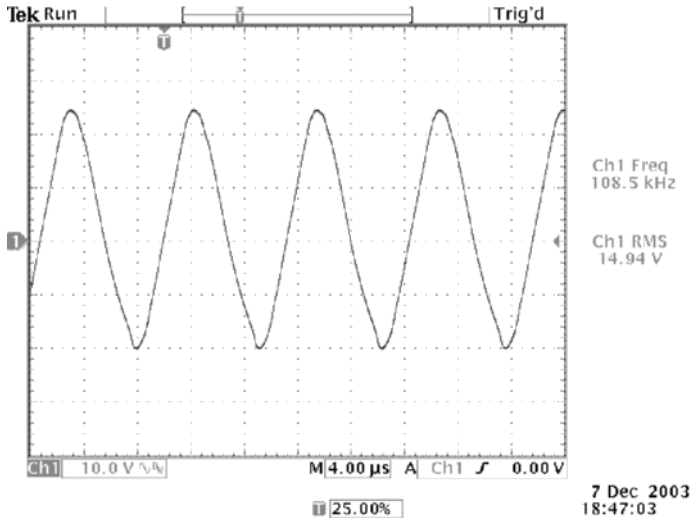


Figure 6.24

Slewing distortion turns a sine wave into a triangular wave

Providing that the effect is not visible at 10 kHz, there probably isn't a problem, but if it occurs at a lower frequency, then the offending combination of stages should be identified and corrected.

The definition of whether slewing distortion is a problem is somewhat woolly, and this is because it all depends on what sort of music you listen to and at what level. Music played on acoustical instruments tends not to generate high amplitudes at 10 kHz, so it rarely provokes slewing distortion, whereas electronically generated music is capable of provoking slewing distortion.

If the patient is a feedback amplifier, the global feedback should be removed and the input level adjusted to restore full power at 1 kHz. Then, change the frequency to 10 kHz and

probe inside the amplifier to find the offending combination of stages. Slewing distortion can be tackled by:

- Substantially increasing the quiescent current in the preceding stage (probably by a factor of five) forces a redesign, and a valve having similar μ but higher mutual conductance will probably be required.
- Reducing the shunt capacitance. Usually the offending combination is the driver stage and output stage, because this is where the largest voltage swings occur, and output stages invariably have high input capacitance. Pentodes have far lower input capacitance than triodes, so a triode-strapped output stage could cause slewing distortion! One way of reducing the input capacitance is to insert a cathode follower, but cathode followers are not immune to slewing distortion, so it will probably need to pass five times the current of the preceding stage. Another possibility for push–pull Class A amplifiers is partly to neutralise the output stage by connecting a small capacitor from each anode to the other valve's control grid. This is positive feedback and needs to be used with great care...

Power bandwidth and transformer saturation

In addition to checking for slewing distortion, it can be useful to check the power bandwidth. This is essentially amplitude against frequency response at full power. The high-frequency -3 dB point at full power is generally limited by slewing distortion and so long as it is beyond 20 kHz, it probably isn't a problem, but the low-frequency power response is more significant.

The low-frequency power response is almost totally governed by transformer core saturation. Rather than looking for the frequency at which output power halves (-3 dB), which would imply gross distortion, it is kinder to the output valves to measure the frequency at which distortion begins to rise (see [Figure 6.25](#)).

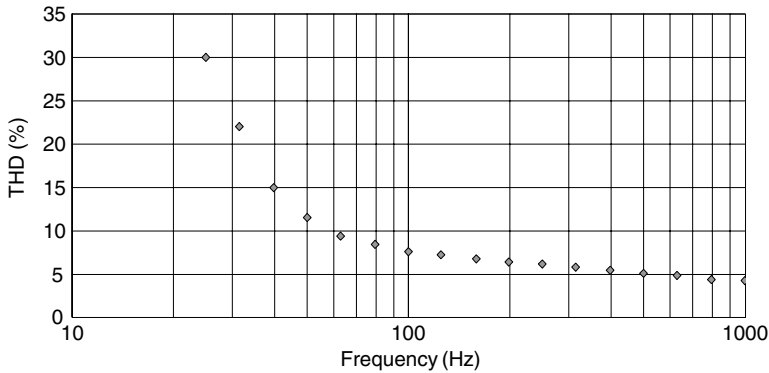


Figure 6.25

Transformer core saturation causes increased distortion at low frequencies (This ill-suited transformer was deliberately chosen to ease measurement.)

The graph shows the distortion of a zero feedback single-ended amplifier with an ill-suited output transformer having insufficient primary inductance. As can be seen, once frequency falls below 60 Hz, distortion rises dramatically.

Typical classic push-pull amplifiers could generally support full power down to 50 Hz before transformer saturation caused distortion that was visible on an oscilloscope. On test, the Crystal Palace amplifier produced 1% THD at full power at 22 Hz, and reducing frequency until transformer saturation became visible on the oscilloscope required 11 Hz, but produced pulsing orange flashes from within the output valves, so the test was hurriedly abandoned.

Investigating the distortion spectrum

If you are able to perform an FFT on the distortion waveform, or have a wave analyser allowing you to measure individual distortion amplitudes, this can be very useful.

To find individual distortion amplitudes, we need to make a sweeping assumption. Fortunately, the assumption is usually

valid. The assumption is that **one** of the distortion harmonics is much higher level than the rest, >6 dB higher will do, but >10 dB is better, it doesn't matter which harmonic it is. Provided that this assumption is true, then the THD measured on the audio test set is equal to the amplitude of that individual harmonic.

As an example, the author tested one section of a 7N7 as the lower valve in a mu-follower amplifier:

<i>Levels measured by FFT or wave analyser</i>						
THD (%)	THD	2nd	3rd	4th	5th	6th
0.175	−55	0	−33	−42	−50	−58
absolute levels		−55	−88	−97	−105	−113

The absolute level at the monitoring output of the audio test set is unimportant. What is important is the relative amplitudes between the 2nd harmonic and the higher harmonics. Because the 3rd harmonic is −33 dB compared to the 2nd, the sweeping assumption that the level of the highest harmonic is equal to the level of the THD is justified, so (having converted THD into dB) we can immediately deem the result to be the level of the 2nd harmonic. The higher harmonics were all measured relative to the level of the 2nd harmonic, so their absolute levels can easily be found. In this example, the technique has produced some remarkably low levels for the 4th and higher harmonics, so it is best to check the distortion residual of the test equipment very carefully before inserting a DUT and believing impressive figures. The author's test equipment is generally reliable to ≈ -95 dB.

Investigating the distortion spectrum of an amplifier is a very powerful technique for choosing between different valves and bias points. In theory, it ought to be good for detecting power amplifier problems, but in practice, observing the distortion waveform on an oscilloscope is faster and more sensitive.

There are many packages that allow computers to perform the FFT. Beware that when you use these, you are relying on the linearity of your sound card. The sound cards that are supplied with computers are not usually very good, and you will need to install a recording-quality sound card, preferably 24 bit, 96 kHz sampling frequency, or better.

References

1. IEC standard 268 (Sound System Equipment). Part 10: Peak programme level meters.
2. “On Reference RIAA Networks” Jim Hagerman. “Audio Electronics” available at: <http://www.hagtech.com/iraa.html>.
3. “On RIAA equalization Networks” Stanley P Lipshitz. *Journal of the Audio Engineering Society*, 1979 June, Volume 27, Number 6, 458–481.(1).
4. A series of papers by A N Thiele and Richard H Small in the *Journal of the Audio Engineering Society* starting from May 1971, Volume 19, No. 5 to January 1972, Volume 20, No. 1.

This Page is Intentionally Left Blank

INDEX

- Active voltage probe, 225
- Aliasing, 209
- Allen (Hex) keys and drivers, 112–13
- Alternating current (AC), 194
- Analogue oscilloscopes, 191–3
- Analogue to digital converter (ADC),
 - 190, 207
 - contravening Nyquist and aliasing, 209–10
 - equivalent sample rate, 208–9
 - Nyquist and equivalent sample rate, 208
 - sample rate, record length, and time base speed, 210–12
 - sampling and quantising, 207–8
- Anode dissipation, 250
- Appliance testers, portable, 264–6
- Audio phase meter, 254–5
- Audio Precision System One, 235
- Audio test sets, 232–8
 - and oscillators, 237
- Autotransformer, 240
- Averaging, 216

- BBC amplifiers, 303
- BBC ATM1 plus TS10, 234
- BBC EP14/1, 234
- BBC ME2/5, 234
- Bells and whistles, 201–6
- Bins (FFT), 221
- Blackman-Harris window, 221
- Bonded acetate fibre (BAF), 37
- Burns, 261–2
 - see also* Safety

- Cal, 225
- Cathode ray tube, 17
- C-core transformers, 74, 302
- Chassis layout:
 - acoustical, 35–7
 - aesthetic, 37–8
 - electromagnetic induction, 5
 - beam valves and mains transformers, 9–10
 - coupling between wound components, 5–7
 - transformers and the chassis, 8–9
- electrostatic induction, 24–5
- electrostatic screening, 27–9
- input sockets, 29–31
- valve sockets, 25–7
- heat, 10–11
 - cold valve heater surge current, 20–1
 - modes of cooling, 11–13
 - using the chassis as a heatsink, 16–20
 - valve cooling and positioning, 13–16
 - wire ratings, 21–2
- mechanical/safety, 32–5
- output-transformerless (OTL)
 - amplifiers, 45–7
- power amplifiers, 38–42
 - meters and monitoring points, 42–3
 - sympathetic power amplifier recycling, 43–5
- power supplies, 52–5

- fitting mains outlets, 53–4
- fitting TO-3 regulators and transistors, 54–5
- regulator positioning, 52–3
- pre-amplifiers, 47–50
- personality, 51–2
- umbilical leads, 50–1
- unwanted voltage drops, 22–4
- Class A amplifiers, 127
- Class AB amplifiers, 127, 237, 339
- Class B amplifiers, 23
- Component bridge, 238–9
- Computer-aided design (CAD), 60
- Conduction, 11–12
- Constant current regulator, handy, 251–3
- Continuously variable transformer, 240–2
- Convection, 12–13
- Cooling, modes of, 11–13
- Corner pillars, 87–8
- Countersinking, 68–9
- Crest factor, 187
- Crystal Palace amplifier, 336, 345
- Current probes, 226
- Cursors, 203
- Curve tracer, 247–51
- Cutters, 106–7

- Data compaction, 216
- Deburring, 65–7
- Decibel (dB and dBu), 308–11
- Dedicated audio test sets, 232–8
- Demultiplexing, 210
- Desoldering, 104–5
- Device under test (DUT), 240, 314
- Digital cameras, 116
- Digital oscilloscopes, 206–12
 - features unique to:
 - adjustable persistence, 217–18
 - automated measurements, 218–19
 - averaging and noise, 216
 - capturing and displaying infrequent events, 213
 - colour, 214
 - dots and vectors, 218
 - envelope, 216–17
 - Fast Fourier Transform (FFT), 219–21
 - negative time, 213
 - peak detect/data compaction and glitch detection, 215–16
 - storing and exporting traces, 214–15
 - usable slow time base settings, 213
- Digital phosphors, 217
- Digital voltmeters (DVMs), 189–90
 - choosing, 191
- Direct current (DC), 194
- Distortion waveform, 339
- Drill speed, 64–5
- Drilling, 157
- Drilling block, 62
- Dummy load, 314

- Earth loop resistance, 129
- Electric shock, 259–61
 - see also* Safety
- Electrical safety testing, 263
 - portable appliance testers, 264–6
- Electromagnet, 172
- Electromagnetic induction, 5
 - beam valves and mains transformers, 9–10
 - coupling between wound components, 5–7
 - transformers and the chassis, 8–9
- Electrostatic induction, 24–5
 - electrostatic screening, 27–9
 - input sockets, 29–31
 - valve sockets, 25–7
- Electrostatic screening, 27–9
- Envelope mode, 216
- Equivalent sample rate, 208–9
 - and Nyquist, 208
- Etching, 155–7

Fast Fourier Transform (FFT), 219

Faultfinding:

AC conditions, 282–9

DC conditions, 273–82

individually test each component,
272–3

intermittent faults, 300–1

noise and crackles, 299–300

oscillation, 289–99

Ferrograph ATS 1, 234

Flat-bladed screwdrivers, 110

Four-wire connection, 53

Full scale deflection (FSD), 173

Function generators, 230–1

Functionality testing:

application of power, 267–72

second-hand equipment versus
freshly constructed new
equipment, 266–7

Gas irons, 98–9

Glass reinforced plastic (GRP),
152

GND, 194

Hard-wired, 33

Hardwiring modules, 161–2

High speed steel (HSS), 157

High-value resistors, 180–1

Holdoff, 205

Hot air gun, 115–16

Humdinger, 123

Identifying the outer foil of a
capacitor, 163–4

Input sockets, 29–31

Leak Stereo 20, 303–4

Leak TL10, 304

Leak TL12 and BBC LSM/8
derivative, 302–3

Leak TL12+, 303

Lindos LA100, 235

Linear distortions:

frequency, impedance against, 326–7

measuring impedance of a

loudspeaker, 327–30

measuring output resistance of a
pre-amplifier, 331–3

source and load impedances for
DUT, 314–16

gain, 307–8

dedicated audio test sets, 308–11

frequencies to use, 317–19

measuring at different

frequencies, 316–17

measuring RIAA equalisation

using inverse RIAA

network, 319–26

peak programme meter scales,
311–12

plotting the graph, 319

unfortunate consequence of 600Ω
legacy, 312–14

Loadlines, 182, 326

Low-value resistors, 179

Lubrication, 63–4

Making a chassis using extruded

aluminium channel and sheet, 85

corner pillars, 87–8

finishing, 92–3

fitting carpet-piercing spikes, 91–2

fitting the pillars, 88–91

fitting the top plate, 91

Marker pens, 116

Marking out:

CAD solution, 60

checking and punching, 60–1

correct drill speed, 64–5

countersinking, 68–9

deburring, 65–7

lubrication, 63–4

making small holes in thin sheet, 74–5

manually, 59–60

measurements, 57–9

- preventing relative movement, 62–3
- safety and drilling, 61–2
- sheet metal punches, 72
- tapping holes, 69–72
- Matching:
 - high-value resistors, 180–1
 - low-value resistors, 179
- Mean level, 182–3
- Meters and AC measurements:
 - crest factor, 187
 - mean level, 182–3
 - peak to mean ratio, 186
 - peak voltage, 181–2
 - power and RMS, 184–6
 - speed of measurement, 187–8
- Miller capacitance, 245
- Moving coil meters and DC measurements:
 - how a moving coil meter works, 172–3
 - if meter unspecified, 175–6
 - measuring larger currents, 174–5
 - measuring output stage cathode current, 178–9
 - measuring voltages, 177–8
- Non-linear distortions:
 - distortion against frequency, 339–41
 - investigating the distortion spectrum, 345–7
 - maximum output power and distortion versus level, 335–9
 - power bandwidth and transformer saturation, 344–5
 - slewing distortion, 341–4
- Nuts and associated tools, 113–14
- Nyquist and equivalent sample rate, 208
- Oscillation:
 - global feedback and compensation, 291–7
 - global feedback and power oscillators, 290–1
 - in individual stages, 297–9
 - motorboating, 299
- Oscilloscopes:
 - analogue oscilloscopes, 191–3
 - connecting to the oscilloscope:
 - input capacitance and voltage dividing probes, 222–5
 - other probes, 225–7
 - transmission lines and terminations, 227
 - digital oscilloscopes, 206–12
- Oscilloscope probes, 224
- Oscillators and dedicated audio test sets, 227
- Output stage cathode current, 178–9
- Output-transformerless (OTL) amplifiers, 45–7
- Passive voltage-dividing probe, 225
- Peak detect mode, 215
- Peak programme meter (PPM), 232, 311
- Peak to mean ratio, 186
- Peak voltage, 181–2
- Pliers, 107–8
- Portable appliance testers (PAT), 264–6
- Power, application of, 267–72
- Power amplifier recycling, 43–5
- Power amplifiers, 38–42
- Power and RMS, 184–6
- Pozidriv and Phillips screwdrivers, 110–11
- Pre-amplifiers, 47–50
- Printed circuit boards, 33, 148
- Quad II, 302
- Radiation, 13
- Range flatness, 229
- Remote sensing *see* Four-wire connection
- Repetitive sample rate *see* Equivalent sample rate
- Rogers Cadet III, 304–5
- Roll mode, 213

-
- Safety:
 - avoiding shock and burns, 262–3
 - burns, 261–2
 - electric shock, 259–61
 - Sawing metal, 75–6
 - cutting irregular holes, 78–81
 - cutting perforated sheet to size, 84
 - making holes in perforated sheet, 83–4
 - making round holes without a punch, 82–3
 - sheet metal, 77
 - Scalpel, 115
 - Scrapbox Challenge amplifier, 333
 - Screwdrivers, 110
 - Semiconductor component analyser, 251
 - Slewing distortion, 341–4
 - Slow time base settings, 213
 - Small-signal measurement, 307
 - Solder tags, 103–4
 - Solder, 99–100
 - Soldering and shrink back, 100–3
 - Soldering irons, 94–6
 - Speed of measurement, 187–8
 - Stefan’s law, 13
 - Synthetic resin bonded paper (SRBP), 152
 - Technical Projects MJS401D and Neutrik derivatives, 234–5
 - Tips, 97–8
 - Toolbox, 117
 - Total harmonic distortion (THD), 232
 - Traditional chassis, 84–5
 - Triggering, 196–9
 - Valve amplifier, planning mechanical layout of *see* Chassis layout
 - Valve cooling and positioning, 13–16
 - Valve sockets, 25–7
 - Valve tester, 242–7
 - Variac, 240
 - Vectors, 218
 - Voltstick, 239
 - Wien oscillators, 228–30
 - Williamson, 302
 - Wire strippers, 109
 - Wiring hand tools:
 - Allen (Hex) keys and drivers, 112–13
 - cutters, 106–7
 - flat-bladed screwdrivers, 110
 - hot air gun, 115–16
 - lighting, 117
 - marker pens and digital cameras, 116
 - nuts and associated tools, 113–14
 - pliers, 107–8
 - Pozidriv and Phillips screwdrivers, 110–11
 - scalpel, 115
 - screwdrivers, 110
 - toolbox, 117
 - wire strippers, 109
 - Wiring techniques, 117–18
 - Class I and Class II equipment, 128
 - earthing, 129
 - 0 V system earthing, 132–3
 - breaking the 0 V signal earth to chassis bond, 133–4
 - earth safety bonding, 129–32
 - ideally positioning the 0 V signal earth to chassis bond, 134–5
 - interconnects, screens and their earths, 135–6
 - internal earth wiring of amplifiers, 136–40
 - rectifiers and high-current circuitry, 140–4
 - electromagnetic fields and heater wiring, 118–21
 - electrostatic fields and heater wiring, 121–4

- enlarging the hole for the wire in a solder tag, 166–7
- fuses, 126–8
- high-current ($> 2\text{A}$) heater
 - regulators that shut-down at switch-on, 164–5
- layout of components, 144–8
- mains switching, 125–6
- mains wiring, 124
- making PCBs, 148
 - drilling, 157
 - etching, 155–7
 - hardwiring modules, 161–2
 - modifying PCBs, 160–1
 - PCB layout, 148–52
 - stuffing and soldering, 157–60
 - transferring the layout to the PCB, 152–5
- NOS valve sockets that won't solder, 166
- tarnished silver-plated stand-offs that won't solder, 165–6
- Wiring tools:
 - desoldering, 104–5
 - gas irons, 98–9
 - solder tags, 103–4
 - solder, 99–100
 - soldering and shrink back, 100–3
 - soldering irons, 94–6
 - tips, 97–8
- “X” amplifier and time base, 195–6
- “Y” amplifier
 - bandwidth, 199
 - triggering, 196–9