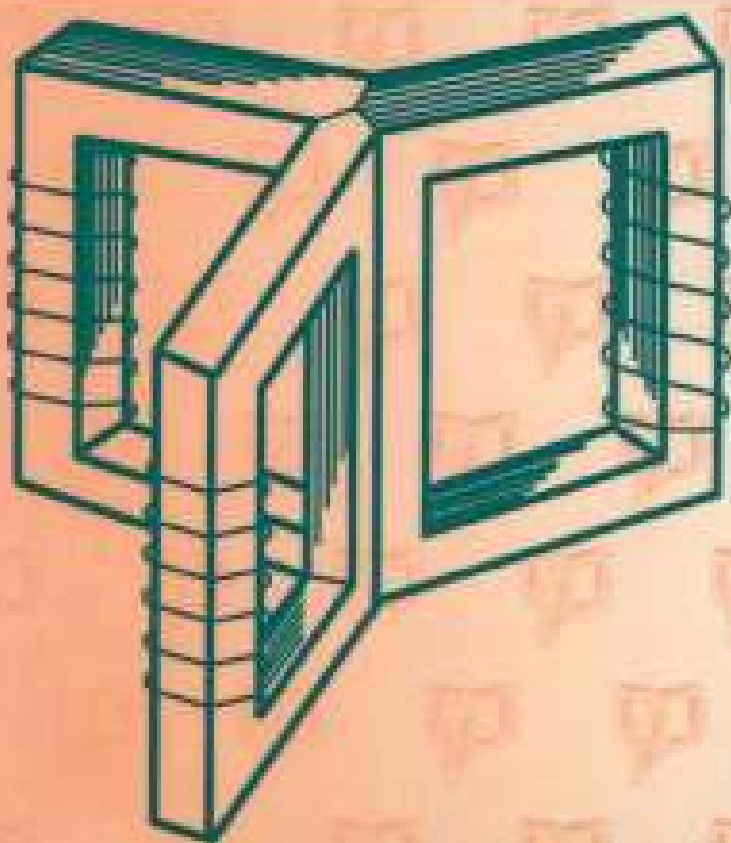


Practical Transformer Handbook



Irving Gottlieb



Practical Transformer Handbook

Practical Transformer Handbook

Irving M. Gottlieb P.E.



Newnes

OXFORD BOSTON JOHANNESBURG MELBOURNE NEW DELHI SINGAPORE

Newnes
An Imprint of Butterworth-Heinemann
Linacre House, Jordan Hill, Oxford OX2 8DP
225 Wildwood Avenue, Woburn, MA 01801-2041

A division of Reed Educational and Professional Publishing Ltd

 A member of the Reed Elsevier plc group

First published 1998
Transferred to digital printing 2004
© Irving M. Gottlieb 1998

All rights reserved. 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 W1P 9HE. Applications for the copyright holder's written permission to reproduce any part of this publication should be addressed to the publishers

British Library Cataloguing in Publication Data

A catalogue record for this book is available from the British Library

ISBN 0 7506 3992 X

Library of Congress Cataloguing in Publication Data

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



Typeset by Jayvee, Trivandrum, India

Contents

<i>Preface</i>	ix
<i>Introduction</i>	xi
1 An overview of transformers in electrical technology	1
Amber, lodestones, galvanic cells but no transformers	1
The ideal transformer – an ethereal but practical entity	3
A practical question – why use an iron core in transformers?	4
Why does a transformer transform?	6
Core or shell-type transformer construction	9
Transformers utilizing toroids and pot cores	13
Bells, whistles and comments	15
Immortality via over-design, sheltered operation and a bit of luck	17
Direct current ambient temperature resistance – just the beginning of copper losses	19
More can be done after ringing door-bells, lighting lamps and running toy trains	22
2 Specialized transformer devices	25
Saturable converters – another way to use transformers	25
Energy transfer without flux linkages – the parametric converter	31
The Lorain subcyclor – an erstwhile parametric converter	36
The constant current transformer	38
The constant voltage transformer	40
Saturable reactor control of a.c. power	41
Transformer action via magnetostriction	45
A double core transformer for saturable core inverters	47
3 Operational features of transformers	49
Operation of transformers at other than their intended frequencies	50

Next to magic, try the auto-transformer for power capability	52
Dial a voltage from the adjustable auto-transformer	55
Transformer phasing in d.c. to d.c. converters	57
Transformer behaviour with pulse-width modulated waveforms	58
Current inrush in suddenly-excited transformers	62
Transformer temperature rise – a possible cooldown	64
High efficiency and high power factors are fine, but don't forget the utilization factor	66
Transformers in three-phase formats – all's well that's phased well	72
Recapturing the lost function of the power transformer	73
Bandpass responses from resonated transformers	74
4 Interesting applications of transformers	78
Remote controlling big power with a small transformer	78
The use of a transformer to magnify capacitance	78
A novel transformer application in d.c.-regulated supplies	81
Polyphase conversions with transformers	83
The hybrid coil – a transformer gimmick for two-way telephony	84
Transformer schemes for practical benefits	85
Transformers in magnetic core memory systems	87
The transmission line transformer – a different breed	88
Transformers as magnetic amplifiers	97
Transformer coupling battery charging current to electric vehicles	99
5 High-voltage transformers	104
Emphasis on high voltage	104
Stepped-up high voltage from not so high turns ratio	105
A brute force approach to tesla coil action	107
Transformer technique for skirting high voltage problems	108
A novel winding pattern	111
Other techniques for high voltage transformation	112
The current transformer	115
Stepping-up to high voltage and new transformer problems	116
6 Miscellaneous transformer topics	119
Transformer saturation from geomagnetic storms	119
The strange 4.44 factor in the transformer equation	122
Link coupling – transformer action over a distance	124
Homing in on elusive secondary voltage	126
An apparent paradox of electromagnetic induction	128
Transformer technique with 'shorts' – the induction regulator	130
Constant current transformers incorporating physical motion	131
A novel way to null transformer action between adjacent tuned circuits	133
Transformers only do what comes naturally	138

The flux gate magnetometer	139
Transformer balancing acts – the common-mode choke	140
Mutual inductance, coefficient of coupling and leakage inductance	143
Short-circuit protection of power lines via a unique transformer	145
Appendix – Useful information	149
<i>Index</i>	169

This Page Intentionally Left Blank

Preface

Since the dawn of the electrical age, many excellent treatises dealing with transformers have been available, one differing from another in the depth of the mathematical rigour involved. However, the salient feature common to most has been their preoccupation with the transformers used in the 50/60Hz utility industry.

This book takes a somewhat different tack; it deals largely with transformers more relevant to electronic technology, control techniques, instrumentation, and to unusual implementations of transformers and transformer-like devices. In so doing, the author feels that the highest usefulness will ensue from emphasis on the *practical* aspects of such transformers and their unique applications. For, if properly done, those readers wishing to probe further will have been guided along appropriate paths to extended investigation. The underlying objective is to stimulate the creativity of engineers, hobbyists, experimenters and inventors, rather than to provide a conventional classroom-like text.

It is to be hoped that a readable and interesting exposition of the topic will help dispel much of the prevalent notion that transformers are mundane devices of a mature technology with little prospect of further evolutionary progress. To this end, it should be easy to show that the traditional treatment of transformers, although providing a solid academic-foundation, tends to confine the modern practitioner to a bygone era.

Irving M. Gottlieb P.E.

This Page Intentionally Left Blank

Introduction

Many practical aspects of transformers are ‘old-hat’ to a large body of practitioners in electrical and electronic technology. It is generally taken for granted, for example, that transformers conveniently serve the purposes of stepping up and stepping down a.c. voltages and currents. And it is widely appreciated that transformers are relatively-reliable devices with a good measure of immunity to abusive treatment. Perhaps best of all, the deployment of transformers in many projects does not require the professionalism of the narrow specialist. Indeed, in many situations, transformers can yield desired results under operating conditions quite different from the ratings on the nameplate.

Having said all this, it is disconcerting and unfortunate that hobbyists and engineers, alike, can be observed to be oblivious to many interesting and useful implementations of transformers and transformer-like devices; this is because many practical insights are not easily extracted from traditional transformer literature. This calls for a focus on various cause-and-effect relationships glossed over, smothered in ‘hairy’ mathematics, or simply omitted in texts and handbooks. It is far from the intent of the author to degrade the very fundamental value of the rigorous engineering approach to this or any other topic. However, it is often the case that added values are found to obtain from a more informal approach to circuits, systems, and devices.

Our mission, accordingly, will be to explore what can be done with other than the ‘garden-variety’ techniques of manipulating voltage and current in conventional ways. Along the way, we shall encounter transformation schemes based upon *different* principles than those of conventional-transformer operation. These should be of great interest to experimenters, innovators, and practical engineers seeking to extend the control domain now possible via unique applications of ICs and other solid-state components. Thus, the central aim of the ensuing discussions will be to drive home practical insights which can be put to use in new and rewarding ways. And hopefully, we can even make lemonade from the ‘lemons’ of transformer shortcomings.

This Page Intentionally Left Blank

1 An overview of transformers in electrical technology

This chapter will be concerned with the broad generalizations of transformers in order that the interesting aspects of specific types can more easily be dealt with in the ensuing discussions. It will be seen that this truly ubiquitous device plays an important role in an indefinitely wide spectrum of systems applications and circuit techniques. On the one hand, the transformer is about as simple an electrical component as one could envisage – a crude version could be made out of various scraps and artifacts found in one's living-quarters. On the other hand, it is interesting to observe that the demonstration of transformer action did not occur until a fortunate blend of theoretical comprehension and practical insight had evolved.

From a historical viewpoint, the beauty of the transformer is that it finally brought together the apparently isolated phenomena of electricity and magnetism. The other electromagnetic devices we now blandly take for granted (electromagnets, relays, electric motors etc.) were the natural outgrowth of this important recognition. Today, it can be said either that the transformer gave great impetus to the widespread use of a.c., or that the use of transformers was largely due to the development of a.c. systems.

Amber, lodestones, galvanic cells but no transformers

The reader is undoubtedly familiar with the basic phenomenon of transformer action. A couple of coils of wire in close proximity and an a.c. source suffice for the transfer of electrical energy via mutual induction. It is, however, only too easy to trivialize the rather simple device known as the transformer. Although a.c. sources weren't around in earlier times, the principle

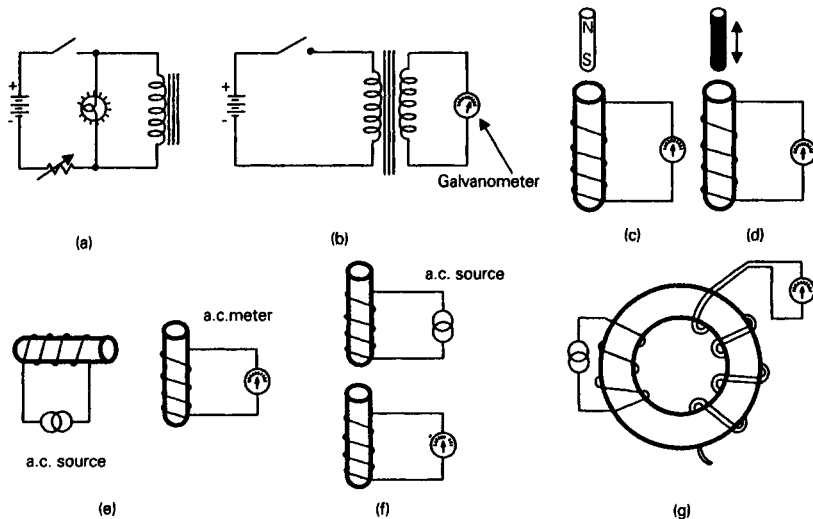


Figure 1.1 'Missing the boat' in the quest for the electricity-magnetism relationship. (a) When the switch is opened, the previously dimly-lit lamp flares up brilliantly for a moment before extinguishing. This demonstrates energy storage in the magnetic field, counter-EMF and self-inductance. What is lacking for transformer action is the proximity of a secondary coil to participate in these events. (b) The galvanometer in this true transformer circuit remains at zero as long as there is constant battery current. However, at about the instant of opening or closing of the switch, there is a momentary deflection. This can easily be missed or misinterpreted. (c) We may suppose the permanent magnet is casually brought near the coil and then allowed to remain stationary. There is no evidence of electric current in the coil. The observer, moreover, has a 'blind spot' to the effect of withdrawing the magnet. (d) The experimenter intuitively imparts motion to the rod. Unfortunately, the rod is made of non-magnetic material. (e) The 'primary coil' is excited from an a.c. source. However, there is no mutual inductance and no transformer action because of the perpendicular relationship between the two coils. (f) This is an appropriate setup for observing transformer action if an a.c. meter is used. However, if a d.c.-responding meter is used, its pointer will remain at zero for ordinary a.c. frequencies. (g) The bifilar secondary windings are inadvertently connected in the series-bucking format. No net induced voltage is detected. The same situation can arise with separate secondaries if they have equal turns.

of the transformer could have, and one can argue, should have, been recognized at a very early date. This is because the erstwhile experimenters long suspected an elusive relationship between electricity and magnetism. To stumble upon the techniques of freely converting one into the other would have ushered in the electrical age at an earlier date. Yet, even after the evolution of the electromagnet, the very next step – the recognition of transformer action, did not occur immediately. Michael Faraday in England

and Joseph Henry in America finally postulated the laws of electromagnetic induction and performed the laboratory experiments that demonstrated transformer action. Others, of course were involved in the complex chain of scientific and empirical events dealing with the behaviour of electricity and magnetism, but Faraday and Henry deserve credit for finally bringing the device to fruition.

The long-elusive essence of transformer action turned out to be embodied in the word *change*. Specifically, a magnetic field linking the turns of an 'input' and an 'output' coil had to undergo *change* by one way or another. It was not enough for a stationary permanent magnet to lie next to such coils. Nor was it sufficient for the d.c. from a battery to circulate through one of the coils. One can sadly speculate that the inductive 'kick' of a galvanometer connected to one coil when the battery circuit to the other coil was made or broken must have elicited a 'so what' response from some of the earlier experimenters. Would you have thought otherwise? The elusiveness of this chance discovery is illustrated in Fig. 1.1.

The ideal transformer – an ethereal but practical entity

A good practical way of dealing with transformers is to consider them *ideal* devices, except where one or more of their non-ideal characteristics merits focused discussion. The validity of this approach lies in the fact that even the narrow specialist finds it expedient to make assumptions and approximations in the otherwise rigorous relationships he uses. The features of the ideal transformer are as follows:

1. The ideal transformer has windings of perfect conductivity – there are *no* I^2R losses in either the primary or secondary coils.
2. When endowed with a magnetic core, the ideal transformer exhibits no core losses. That is, energy dissipation from hysteresis and eddy currents in the core is *zero*.
3. Perfect electromagnetic coupling exists between the windings. In other words, there is *zero* leakage inductance.
4. *No* exciting or magnetizing current is needed to set up the magnetic flux.
5. *No* capacitances are incorporated in the windings.
6. The core exhibits *linear* magnetic characteristics (constant permeability).

These are the basic 'personality traits' of such a 'sinless' device. One could go into greater detail by stating that eddy currents, proximity effect and skin effect should be *absent* in the windings. There should be *no* electromagnetic radiation. There should be immunity to effects of temperature on

operation. Most often, our ideal transformer will operate from a sinusoidal source of a.c. current. Practicality is fortunately well-served by the initial assumptions of *losslessness* and *linearity*. Deviations from such perfection can then be handled as the needs of the particular situation dictate (Fig. 1.2 illustrates these concepts).

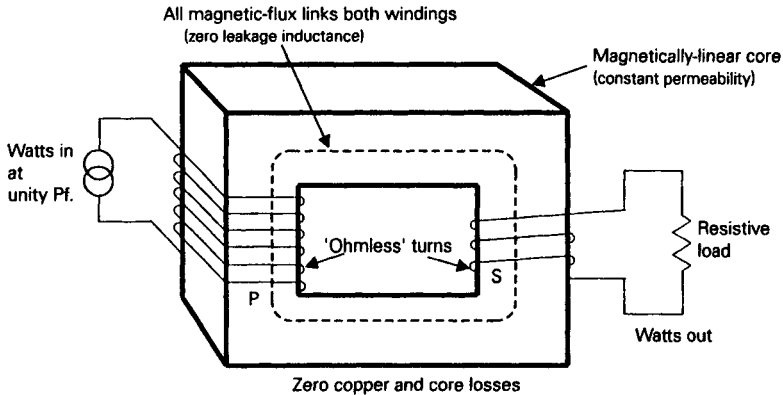


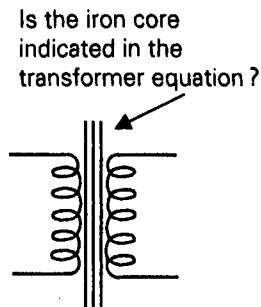
Figure 1.2 The basic features of the ideal transformer. This hypothetical creature, like perpetual motion, free energy or fulfillment of a politician's promises, is not found in nature. However, the concept tends to be useful in dealing with many practical aspects of real-life transformers. Among other things, the ideal transformer is 100% efficient and exhibits zero (perfect) voltage regulation.

A practical question – why use an iron core in transformers?

Students, neophytes and even those fairly well-versed in electrical technology often ponder, but feel embarrassed to ask, why 50/60 Hz transformers have to have iron cores. Why not just do away with the core and save on weight, cost and manufacturing problems? Such speculation is not to be thought of as foolish and is not on the level of those obsessed with perpetual motion or free energy.

It so happens that there is nothing in the transformer equation used by physicists or electrical engineers that explicitly indicates the need for iron or other magnetic materials (see Fig. 1.3). Such equations simply tell us that so much voltage is forthcoming from a certain magnetic flux density undergoing a certain rate of change. If one did not already know the outcome, it would be perfectly natural to suppose that practical power transformers could be designed around air cores. In attempting to build such a transformer, one would stumble upon a disconcerting practical fact of life: in order to obtain the flux density required for handling volts and amperes

rather than, say, microvolts and microamperes, the windings would be massive coils comprising perhaps miles of wire and would have impractically high resistances. The very geometry of such windings would self-defeat the objective of developing high flux *density*. Put another way, high *inductance* is required in the primary in order to keep the transformer exciting current at a small fraction of the allowable full-load current. Here again, such high inductance is not practical with a conductor that has sufficient cross-sectional area to carry any practical current. The mathematics is not incorrect – it is conceivable that cryogenically cooled windings with zero resistance could dispense with the conventional magnetic core.



$$E = \frac{kNf(10^{-8})}{B_m A}$$

where E is the number of volts induced in N turns.

k is 4.44 for sine waves, or

4.00 for square waves.

N is the number of turns on a winding.

B_m is the peak flux density.

A is the cross-sectional area of the core.

(B_m and A should be expressed in the same units).

f is the frequency in Hz.

Figure 1.3 The basic equation for transformer windings. Interestingly, this equation says nothing about using iron or other magnetic cores in transformers. Without fore-knowledge of practical ramifications, it would be natural enough to suppose a 50/60 Hz transformer could be designed with an air core.

It is the *high magnetic permeability* of the iron or other core which saves the day by making possible high magnetic flux density with minimal magnetomotive force. That is, just a few ampere turns can produce much greater flux density than is attainable with air as the 'core'. Stated yet another way, we get high inductance from *practical* coils of wire.

Why does a transformer transform?

Once one gains insight into the need for an iron core or other magnetic material in low frequency transformers, the nature of the energy transfer from the primary to the secondary winding often poses a dilemma. Although it should be quite clear that the source of the electrical energy delivered to the loaded secondary must be the a.c. power-line, why, indeed, should the primary obligingly draw and then transfer the required load power to the secondary? Inasmuch as the unloaded transformer is a very miserly consumer of line current, some *change* must occur within the transformer when an appreciable load is imposed on the secondary.

In order to get to the bottom of this situation, we must first consider the reason the unloaded transformer draws only a very small primary current from the line. Although primary windings can have effective resistances of a fraction of an ohm to several ohms in a wide variety of ordinary transformers and line voltages of tens and hundreds of volts are routinely encountered, there appears to be opportunity for a heavy current flow even with an open-circuited secondary. The answer is that a transformer which is properly designed will exhibit a sufficiently high *inductance* to keep the primary current very low when there is no secondary load demand.

It is relevant at this point to focus on the nature of inductive reactance; it exists because of a counter-EMF which opposes the impressed voltage. Thus, if 100 V are applied to the primary from the a.c. line, it is conceivable that 99 V will be induced as the counter-EMF, effectively leaving only 1 V to force current in the primary. This surely is old hat to many readers, but let's see what happens when a load is connected to the secondary. Because of the mutual magnetic flux threading through both windings, the secondary develops an induced voltage which now drives current through itself and the load. These secondary ampere turns neutralize just the right amount of the primary ampere turns, or field intensity, to *reduce* the primary inductance by the amount needed to draw the required line current. Thus, we may imagine that, in response to the secondary load, the primary counter-EMF drops to 98 V, thereby allowing twice the line current to be drawn by the primary.

It is as if feedback information of increased secondary load causes the opening of a valve to admit more line current to the primary. As this takes place, the relative primary and secondary ampere turns change in such a way as to keep the *mutual flux* fairly constant. Thus, from no load to full load, the primary counter-EMF decreases to allow needed line current to supply the secondary load circuit. Viewed from a slightly different viewpoint, one can construe that increasing secondary load owes its existence to diminishing primary inductance. What might initially appear as a dull and passive device actually has an interesting and dynamic inner life.

Let us look a bit more closely at the small primary current drawn by the

unloaded transformer. Some of it simply supplies copper and core losses. For many practical purposes, however, it is permissible to deal with a so-called excitation current. This is approximately the magnetization current which sets up the mutual magnetic flux in the core. As already mentioned, this mutual flux is the energy transfer agent between the primary and secondary. Although primary and secondary ampere turns undergo change with varying load conditions, these changes are such as to keep the mutual flux nearly constant. The waveshapes and phase relationships of the exciting current, the mutual flux and the induced secondary voltage are shown in Fig. 1.4. The peaked excitation current shape betrays its largely *third harmonic* content.

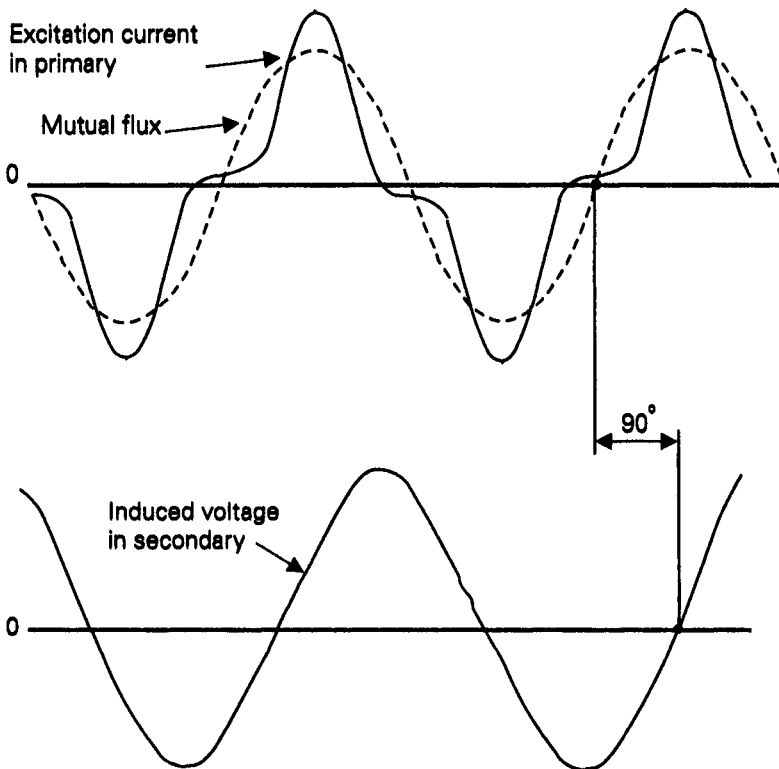


Figure 1.4 Waveforms in the iron core transformer. The non-sinusoidal excitation current is driven from a sine wave source but is distorted due to the effect of the non-linear iron core. Nevertheless, it produces the sinusoidally-shaped mutual flux. In response to the mutual flux, the induced voltage in the secondary is also a sine wave. The mutual flux leads the induced voltage by 90° . Note that the wave-shaping property of the core simultaneously causes and removes distortion. In this manner, the sinusoidal line voltage impressed on the primary again emerges as a sine wave of induced secondary voltage.

Interestingly, the small excitation current is *non-sinusoidal* despite the fact that the transformer is driven from a sine wave source. However, the excitation current brings into being sinusoidal mutual flux, which then induces a *sine wave voltage* in the secondary. Lest this prove confusing, it must be pointed out that these cause and effect relationships stem from the non-linear magnetic characteristics of the iron core. Being small, the peaked excitation current can often be neglected in practical transformer work. It is nevertheless very important in transformer operation; in some systems where the third harmonic current is blocked, the secondary voltage itself, then becomes distorted often leading to various system disturbances and interference with any communications circuits which might be in the vicinity.

The fact that a mutual flux links both the primary and the secondary turns means that induced voltages will be developed in these windings in the same way. Stated another way, the same number of volts per turn or turns per volt will characterize both the primary and secondary windings. Thus, a secondary winding with the same number of turns as the primary will exhibit the same induced voltage as the primary. If the secondary has twice as many turns as the primary, the induced secondary voltage will be twice that of the primary; it is likewise true that a secondary with one-third the number of primary turns will be induced with one-third of the induced voltage in the primary. This is one of the most important uses of transformers – stepping voltages up and down. Keep in mind, however, that available secondary current is necessarily the *inverse* of the voltage change. If the voltage is doubled, the available current is halved; if the voltage is stepped-down to one-third of the primary voltage, the available secondary current is then three times the current consumption of the primary. These rules apply to both ideal and practical transformers designed to approach ideal characteristics. If it were otherwise, the transformer would be in the same league as perpetual motion machines. Historically, transformers decided the choice of a.c. over d.c. utility power.

A second, extremely useful, feature of transformers is *impedance matching*. This characteristic is the natural consequence of the voltage and currents in the primary and secondary windings. Simply stated, an impedance connected to the primary appears at the secondary multiplied by the square of the voltage gain in the secondary. Thus, if the secondary to primary voltage *step-up* is three, nine times the primary circuit impedance is 'seen' at the secondary. Similarly, if the voltage *step-down* ratio going from primary to secondary is one-half, the secondary will reflect one-quarter of the impedance associated with the primary. For most practical purposes, it is equally valid to say that the impedance transformation occurs as the square of either the secondary to primary voltage ratio *or* their turns ratio. This is shown in Fig. 1.5.

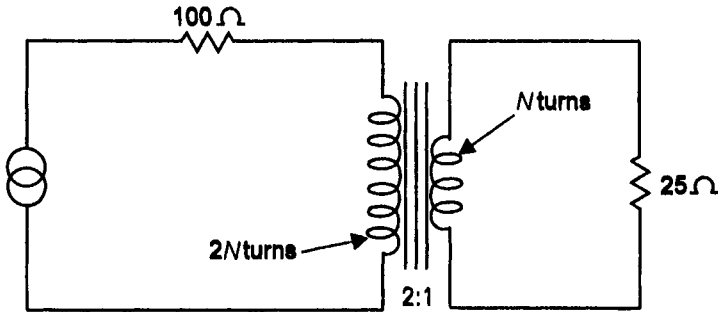


Figure 1.5 Typical example of impedance matching with a transformer. The internal resistance of the a.c. generator is $100\ \Omega$. Maximum power transfer into a $25\ \Omega$ resistive load is achieved with a transformer having a $\frac{1}{2} : 1$ voltage step-down ratio because the impedance reduction is then $(\frac{1}{2})^2$ or $\frac{1}{4}$. The terms 'resistance' and 'impedance' are often used interchangeably in this technique. (However, maximum power transfer can occur only if the resistive component of impedance is numerically matched and the reactive component is resonated or tuned out.)

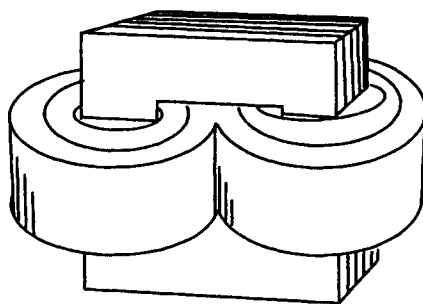
Core or shell-type transformer construction

Imagine yourself to be a creative type of scientific experimenter back in the early part of the nineteenth-century who had just heard of a physicist's claim that electromagnetic coupling (transformer action) could be demonstrated by causing a time-varying magnetic circuit to encircle conductors so that magnetic flux and electric current were interlinked in a chain-like fashion. The implication here is that the electric current would become available because of the physical relationship of the magnetic and electric elements in such a setup.

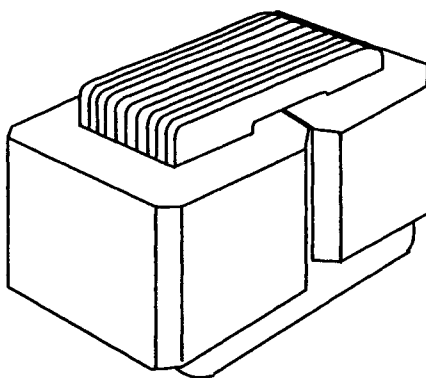
Perhaps you might suspect that the converse situation might also be valid; that is, the conductors might be envisioned to encircle a magnetic circuit in such a way that the overall situation would still be the interlinking of magnetic flux and electric conductors. One might go so far as to ponder whether a time-varying current might serve to produce mutual induction (again, transformer action). With some experimentation to cast light on these cause and effect relationships, you would likely discover the *two* basic practical approaches in constructing transformers. Either wind the conductors around a magnetic core or wind the core around conductors. Although fabrication of modern transformers may not proceed in exactly this way, the net effect remains as described. That is why we have so-called core-type and shell-type transformers. These constitute the two main methods of converting transformer theory into practical hardware. It is not feasible to say which is better – there are too many trade-offs in cost,

efficiency, ease of manufacturing, thermal considerations, insulation matters etc. to judge supremacy for either type.

Single-phase core and shell-type transformers are shown in Fig. 1.6. They have in common the fact that their cores are built up of thin laminations to discourage the flow of eddy currents. For 50/60 Hz transformers these laminations tend to be about 0.3 mm in thickness and are insulated either by natural surface oxidation or by a special coating of varnish or shellac. In audio frequency transformers, much thinner laminations are needed to keep eddy current losses low.



(a)



(b)

Figure 1.6 The two basic physical configurations of iron core transformers. (a) The core-type construction; the magnetic core is surrounded by the primary and secondary windings. (b) The shell-type construction; the windings are surrounded by the core. Note that in both types, portions of the primary and secondary are so distributed as to provide close physical separation. This minimizes leakage inductance. In drawings, however, it is customary to find the primary on one core leg and the secondary on an opposite core leg.

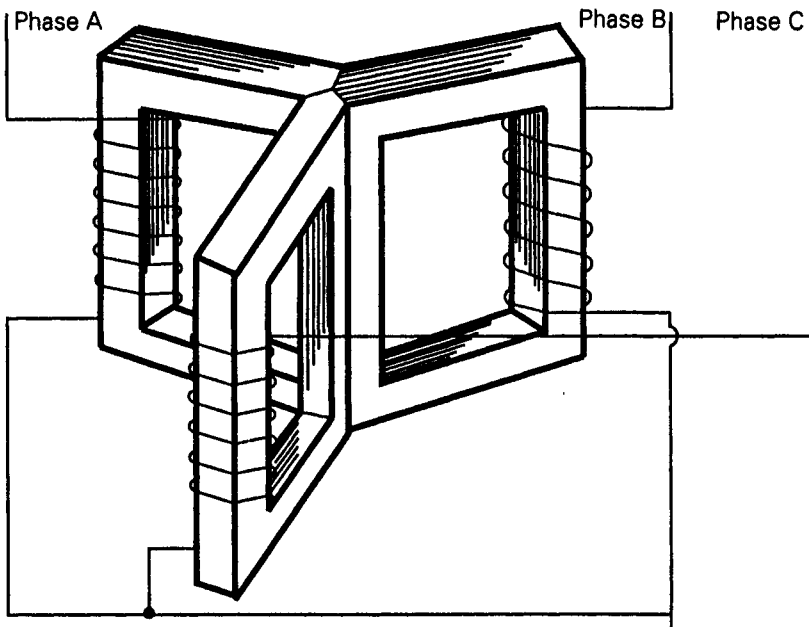


Figure 1.7 Textbook example of a core-type three-phase transformer. This theoretically-valid configuration will not be encountered in practice because of its awkward construction; a re-arrangement of the core elements leads to a format which can be more conveniently manufactured (see Fig. 1.8). (For the sake of simplicity, only the primary windings are depicted in the above figure. The secondaries can be assumed to be wound over the primaries.)

A natural way to construct a three-phase core-type transformer is shown in Fig. 1.7. Such an arrangement provides instructive insight into the nature of three-phase systems but is obviously an unwieldy format for manufacturing and installation purposes. Operationally, it turns out that no net magnetic flux exists in the central three-legged junction. Except for mechanical reasons, this implies that the joined core elements could be eliminated. One can imagine cutting out the vertical portion of the cores forming this junction and butting together the horizontal portions at 120° relative to one another. However, nothing much would be gained in the quest for a simpler device.

The solution is an overall structure such as the three-phase core-type transformer shown in Fig. 1.8. This in-line format incorporates some loss of symmetry in the magnetic circuits of the three phases. However, operational disturbance is minimal and the scheme has proved to be an elegant *practical* solution to the three-phase transformer problem.

Following similar logic, the in-line format provides us with the structure for the shell-type three-phase transformer, as illustrated in Fig. 1.9.

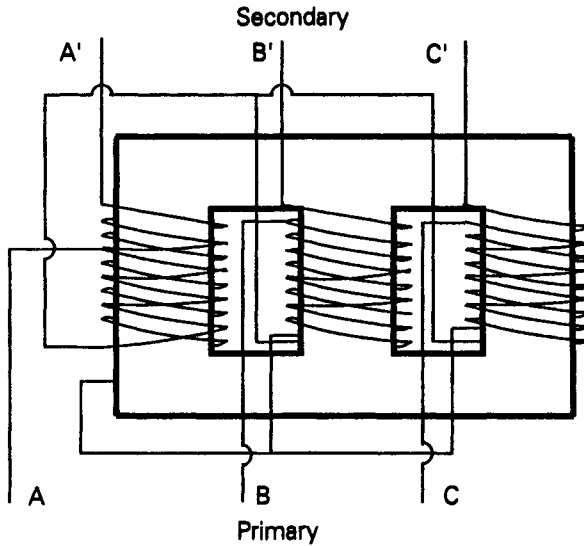


Figure 1.8 *The practical configuration of a core type three-phase transformer. The easily manufactured in-line format is not perfectly balanced for the three excitation currents, but operational disturbance is minimal.*

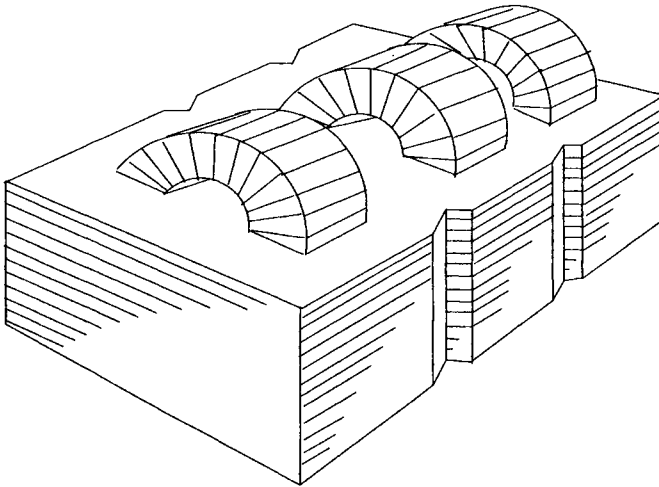


Figure 1.9 *Construction of the shell type three-phase transformer. The in-line format is evident, as is the enclosure of the windings by the magnetic circuit.*

Such three-phase transformers manifest savings of iron and copper when compared to equivalent systems using three single-phase transformers. This

results in a higher overall efficiency and better regulation, as well as lower initial installation costs. However, in the event of failure, maintenance is more difficult compared with a single-phase transformer. In electronic control systems, a single three-phase transformer is usually justified by the savings in space. In any event, the advantages of three-phase technology in power transmission, motor operation, electric vehicles and electronic power control merit interest in three-phase transformers.

Transformers utilizing toroids and pot cores

The core and shell constructions alluded to fit the bill for utility frequency transformers and also work well for audio frequencies. By the time we get to the middle and upper audio frequencies, difficulties begin to set in because of the thin laminations required. Also, the core losses of 50/60 Hz silicon transformer steel tend to become excessive. Transformers working in the 10 kHz region and higher are often designed around toroidal cores. These are made of various magnetic substances such as powdered iron, ferrites and alloy formulations. A typical toroidal transformer is shown in Fig. 1.10. Tape-wound toroidal cores using various magnetically permeable alloys are also found in specialized applications.

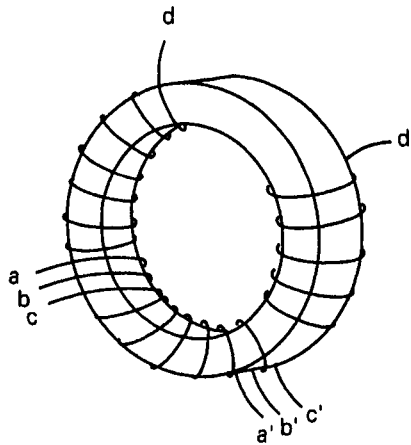


Figure 1.10 The toroidal core has many non-utility transformer applications. These core structures are made from ferrite, powdered metal and alloys. Tape cores are also available in the toroidal shape. For the most part, such transformers find applications at audio, radio and near-microwave frequency. Various formulations yield optimum behaviour at specified frequency ranges, permeability and loss characteristics. Tape cores, for example, work best in the audio frequency band. Both powdered iron and ferrite cores are used at radio frequencies.

Although the many different types of core materials used in toroids may at first be confusing, it is actually a satisfactory situation. This is because it becomes convenient to select a material likely to produce optimum results for the purpose at hand. The manufacturers have rigorous control over specifications, tolerance and reproducibility. Many experimenters like to work with toroids because one is freed from having to stack-up laminations. In terms of performance, toroidal transformers are known for their restricted external fields – such units can be mounted quite close together without intercoupling problems.

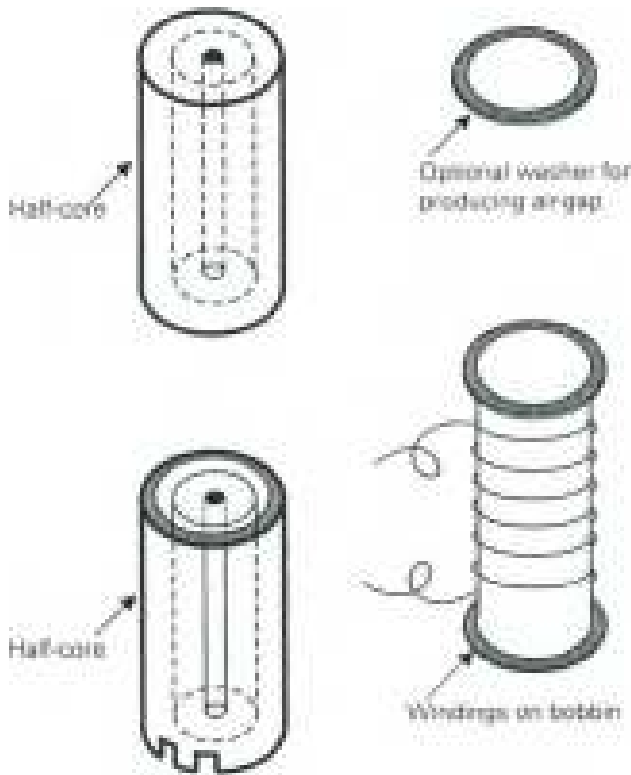


Figure 1.11 *Salient features of construction of the pot core transformer. The cylinder of magnetic material disassembles near its centre, exposing the removable and easy to wind bobbin. A nice feature for experimenters is that an effective air gap can be inserted in the magnetic circuit by putting washers between the cylinder halves. This prevents unwanted core saturation in applications requiring a linear transformer. Power handling capability tends to be less than for toroidal cores.*

The other side of the coin involves manufacturing problems. Although there are no laminations to stack, hand-winding can be quite tedious. Of course, in industry there are toroidal winding machines. Anyone who has seen one of these in operation for the first time must come away surprised at the speed at which a multi-layer winding of fine wire can be applied. Fortunately, many high frequency toroidal transformers only require a few turns of easily manageable wire.

There is considerable overlap in the magnetic parameters of core materials. The experimenter can often obtain satisfactory results from a core material which a specialist might reject because of less than optimum characteristics.

The pot core, also known as the cup core, is a closed cylinder containing an upright round member on which a nylon bobbin is situated. Although closed when finished, the cylinder comes apart near its mid-point to allow access to the bobbin. Usually this core shape is made of powdered iron or one of the ferrite materials. Such materials lend themselves well to moulding techniques during manufacture. It is more convenient to wind a bobbin than it is to hand wind a toroid. Good electromagnetic coupling is readily attainable with this core configuration and it can be said to be self-shielded, in some instances even more so than the toroid. Mounting on PC boards tends to be more straightforward than toroids. One does not, however, often encounter pot cores with as high power-handling capability as is readily forthcoming from toroids. Mechanical problems enter into this discrepancy and it should also be pointed out that the bobbin windings have limited exposure to air circulation. The basic physical features of the pot core transformer are shown in Fig. 1.11.

Interestingly, one can compare the toroidal transformer to the core-type construction discussed previously. This is because the windings surround the core in the toroid. Conversely, the core surrounds the windings in the pot core transformer. Therefore, it can be likened to the shell core transformer discussed earlier.

Somewhere between core-type constructions are various shapes of ferrite, iron core, molybdenum permalloy and other powdered cores. These may be rectangular, square, E's, C's, or may have a shape custom-made for specialized applications. There are also open cores consisting of rods or strips. The ferri-loopstick antennas of portable radios are actually open core transformers. Another example of an open core transformer is the ignition coil used in cars which utilizes a bundle of soft iron wires for its core.

Bells, whistles and comments

A basic transformer may consist of a simple arrangement of primary and secondary windings and a magnetically permeable core. If operation is

intended to be at high or radio frequencies, an air core may be entirely suitable for the required power transfer. However, transformers may be endowed with various embellishments making some models appear as de luxe versions. The windings can incorporate *taps* so as to provide flexibility in the selection of transformation ratio. Litz wire may be used in high frequency windings to reduce skin effect. Low frequency transformers may make use of square wire, or even foil, in order to increase the packing factor of the windings. In some designs, air gaps are introduced in the core in order to linearize the inductance.

There is no end of *ventilation* and *heat removal* techniques. These involve forced air cooling, circulating fluids, finned radiators and sophisticated approaches such as heat pipes and the use of exotic materials such as beryllium oxide (beryllia), or diamond. In large transformers, the objective is to hold down the I^2R losses in the temperature-dependent resistance of the copper windings. This remains true in smaller transformers, but often the main idea is to obtain a lot of power capability in a restricted space. Although aluminium conductors have only about 60% of the electrical conductivity of copper, sometimes a weight-saving can be obtained. Various *insulating* materials are used, but nearly all benefit in reliability and longevity by holding temperatures down.

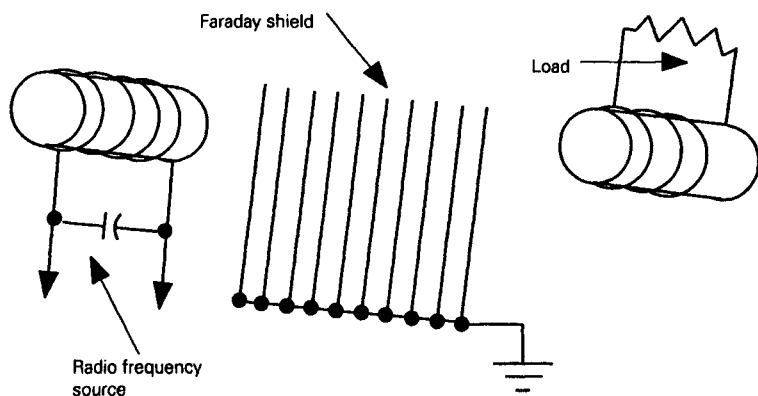


Figure 1.12 Transformer coupling of RF through a Faraday shield. The Faraday shield is a conductive material screen without any closed loops. It allows electromagnetic linkage (transformer action) but blocks capacitive coupling.

An *electrostatic shield* is often desirable between the primary and secondary. Such a shield may serve various purposes. In electronic systems, it prevents both line transients from capacity coupling into the load and helps

block circuit-produced noise from polluting the power line. Such shields can be foil or screen-type conductive material. A main requirement is that the electrostatic shield must *not* behave as a shorted turn. Therefore, it *almost* closes upon itself, but contains a gap. To be effective, the shield usually needs to be grounded and a terminal is often provided for this purpose. In some transformers, the shield is internally connected to the core. At radio frequencies, the *Faraday shield* shown in Fig. 1.12 allows transformer action, but shields against unwanted capacitive coupling.

In high voltage transformers, the internal shield serves another very important purpose. When appropriately shaped and contoured, it controls the *distribution* of electric field flux so as to reduce the likelihood of arcing – an extremely destructive phenomenon. The already high peak voltages can have relatively low energy transients imposed on the crests of the sine waves which then have the potential to break down air or puncture insulation.

At radio frequencies, both electrostatic and electromagnetic shielding of a resonant circuit or a transformer can be achieved by enclosure in a copper or other metallic box. It is only necessary that internal spacing should be at least one coil radius in order to prevent excessive eddy current losses and reduction of Q . On the other hand, at audio and especially at lower frequencies, it becomes difficult to confine the magnetic fields in an enclosure of copper sheet metal. Here, a magnetic material is needed – preferably one having high permeability at low flux densities.

In general, special techniques are often used in order to optimize special performance needs. For example, a transformer handling high power levels in the audio frequency range may be potted with a sound deadening material in order to reduce the intensity of audible noise from the vibrating laminations. (The windings tend to add to the sound level also.) In three-phase utility systems, one encounters an interesting add-on to Y-connected transformers. This is in the form of *delta-connected tertiary windings* which are not called upon to deliver load current. Rather, these auxiliary windings provide a path for excitation current to flow. In the absence of such an ‘artificial’ path, the otherwise sinusoidal input and output voltages become ‘polluted’ with third harmonic energy from the non-linearity of the core. This can cause a more peaked waveform, endangering insulation. It also tends to cause interference with communication services from the third harmonic fields present in the lines.

Immortality via over-design, sheltered operation and a bit of luck

A transformer, being a passive device made up of essentially inert materials should last forever. Indeed, from a practical viewpoint, many, if not most, appear destined for such a lifetime. Nevertheless, as with all man’s

inventions, there is an on-going failure rate. Even under benign operating conditions, transformers develop short- and open-circuits, intermittents, electrical leakage and isolation problems, excessive audible noise, peculiar odours, and with some core materials, substantially changed characteristics. In real life, the terms 'passive', 'inert' and 'benign' tend to be more relative than absolute.

Transformers can be exposed to moisture, ozone and air pollution. Unfortunately, the processes of rust, erosion and corrosion are assisted by electrical potentials and currents. The unseen mechanical forces from electromagnetic reaction (Lenz's law) are constantly at work loosening the attachments of the windings. The previously mentioned 'benign operation' caters more to idealism than to real life exposure; most transformers are called upon to endure and survive their share of overloads, transients and excessively hot or cold environments. When things are restored to normal, small residual impairments to insulation and terminal connection linger on; over the long-term, these accumulate and can lead to a major disruption of performance.

Despite near-perfect quality control, a small percentage of transformers shipped from the factory will incorporate some kind of a defect not easily detected in conventional testing. This is simply a statistical fact of life. Particularly vulnerable are transformers with windings of very small gauge wire, as well as those in which there is heavy reliance upon the quality of insulation. With high voltages and/or high frequencies, inconsequential leakage paths are not likely to remain, so instead of 'clearing up' they tend to carbonize and become even more conductive.

Whether it is easy or economical to repair a defective transformer is highly dependent upon individual circumstances. A good inspection is always in order, for quite often the day may be saved by just re-soldering or repairing the terminal connections.

Shorted turns are a common mode of failure for transformers. An analogue ohmmeter is not very useful in detecting this defect because of the fuzzy precision of such meters. Digital meters may not be much better because of the need correctly to interpret very small percentages of change. Moreover, such defects are often temperature-, voltage-, and current-dependent or may be intermittent for apparently no reason. If the transformer runs abnormally hot and/or develops less than its normal output voltage, one can safely assume the possibility of shorted turns. Hardest to find are one or several shorted turns in a winding consisting of many turns. With the load disconnected from the secondary, the consumption of appreciable primary current is usually a dead give-away of an 'internal load', i.e. *shorted turns*. These can be in *either* the primary or the secondary.

Open circuited windings are, fortunately, often traced to bad terminal connections. A 'cold' or otherwise defective solder joint is a common cause. A more elusive open circuit is sometimes discovered in high voltage trans-

formers utilizing very fine wire in the secondary winding. Such a winding may be physically intact, but electrically discontinuous – the actual copper conductor parts but the enamel or plastic insulation keeps the ‘wire’ together. This headache is generally confined to the factory and is remedied by relaxing the tension imposed by the coil winding machine.

Shorts to the core are also common and make themselves known because cores are often grounded for safety reasons. If feasible, the core may be ungrounded for emergency operation and in some cases may be left like this. In electronic systems, grounding the core through a capacitor often provides a partial shielding effect for components close to each other. Usually, cores of small transformers are left floating.

In radio frequency work, it is often found that a ferrite core transformer no longer works as well as it once did. Disturbances may be observed in the system VSWR and there may be evidence of harmonic generation well above the norm. Here, one can reasonably infer that due to abusive operation and resulting overheating, the permeability and/or other magnetic parameters of the core have suffered permanent change.

Direct current ambient temperature resistance – just the beginning of copper losses

As resistivity in metals, such as copper, increases with temperature, a transformer with narrow-gauge wire in its windings can run hot. If forced to deliver the desired load current, the temperature-dependent resistance will lead to yet higher I^2R loss. If such a transformer survives, its operating efficiency can be seriously reduced. Also, we would have reason to be concerned about the insulation charring and ultimately failing. The magnetic properties of the core can be adversely affected at high temperatures, particularly ferrite material.

Apart from initially selecting the proper conductor and core sizes, some sort of heat-sinking and cooling/ventilating technique is generally required with large utility transformers. This may take the form of oil immersion or may involve a circulating liquid system with pumps, fans and radiators.

Selecting the conductor cross-sectional area of the windings turns out to be both an art and a science, but the most valuable aid is experience. This is because compromises must be made with cold logic – after all, inevitably there is the conflict posed by the practical necessity of fitting the windings into the available space. Using a larger core may prove self-defeating because of the greater length (and, therefore, resistance) of the new winding. An experienced transformer designer finds it necessary to rely a great deal on stacking and packing factors, to say nothing about ‘fudging’ factors. What usually saves the day is that old standby, intuition.

There is even more to so-called ‘copper losses’ of transformer windings. The phenomenon of *skin effect* causes a.c. resistance to be higher than the d.c. resistance of conductors; as the frequency increases, more and more of the current gets concentrated closer to the surface, effectively decreasing the cross-sectional area of the conductor (see Fig. 1.13). The effect is relatively small at 50/60 Hz but picks up at 400 Hz and should not be overlooked at aircraft frequencies of 1200 Hz and higher. When a designer endeavours to attain 98–99% efficiency, every bit of energy dissipation assumes significance. For the user, it simply makes good sense to recognize the practical limits of power transformation; intelligent trade-offs in circuit and system performance are bound to follow.

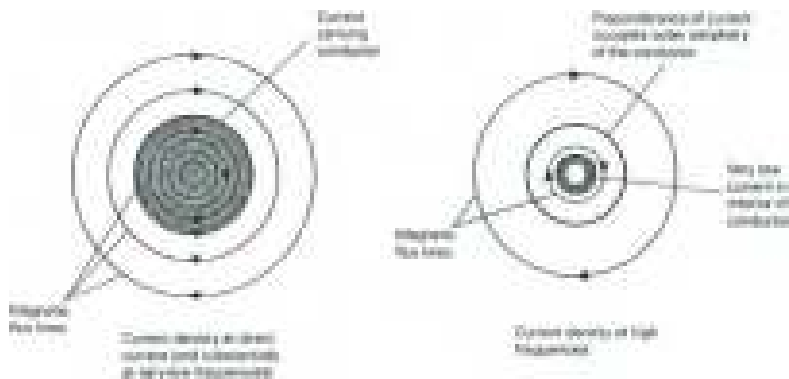


Figure 1.13 The skin effect in current-carrying conductors. In actuality the current density does not abruptly become zero at the boundary lines shown in the drawings. Rather, the boundary lines indicate where current density is $1/e$ or 37% of the value at the surface of the conductor. The basic effect on transformer windings is an increase of the resistance and reduction of the current carrying capacity. A practical explanation of skin effect postulates that the inner portions of a conductor, being encircled by a greater number of flux lines, offer higher inductive reactance than is encountered in the outer portions.

Surprisingly, designers have to account for skin effect at 50/60 Hz too, but only where heavy currents and large cross-sectional area conductors are involved. The ordinary practitioner in electricity and electronics need not be concerned with skin effect at utility frequencies. The following simplified formulae can provide practical insight into the importance of skin effect as a function of frequency. If the skin effect penetration d is sufficiently small, one must recognize that much of the geometrical cross-sectional area of the conductor is not available to carry the current density it would at d.c. The practical manifestation of this current-crowding near the surface is an increased effective resistance of the conductor.

$$d = \frac{2.60}{\sqrt{f}} \quad (\text{in inches})$$

$$d = \frac{6.60}{\sqrt{f}} \quad (\text{in centimetres})$$

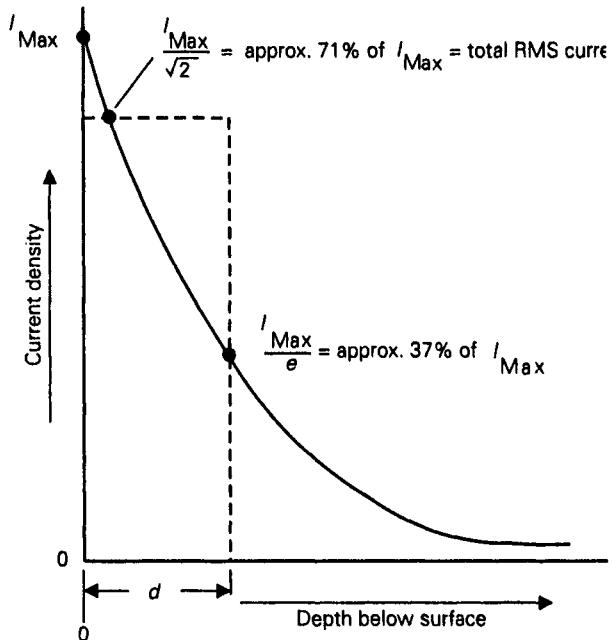


Figure 1.14 Graphical representation of the skin effect. Less utilization of the cross-sectional area of a conductor makes its effective resistance at high frequencies greater than its d.c. resistance. From the graph, it can be seen that the area represented by the dashed rectangle can be thought of as a fictitious conductor with uniform current density. The resistance of this fictitious conductor would be the same as that of the actual conductor under the influence of the skin effect. The smaller the penetration depth d , the less cross-sectional area is available for carrying current.

Figure 1.14 shows the nature of the skin effect relationship and the meaning of the penetration depth d . It is to be noted that d does not designate a circular line dividing all of the current density on one side from zero current density on the other. Nevertheless, at high frequencies, *most* of the current in the conductor will be crowded near the surface. That is why we see hollow tubing being used in radio frequency work.

Litz wire is an effective way to reduce the influence of the skin effect. Litz wire is a twisted and stranded conductor format with each wire *insulated*

from one another. Paralleling is accomplished by conductively joining the *ends* of the individual wires. High frequency resistance is relatively low because of the large surface area of the paralleled wires. The transformers used in modern switching power supplies, inverters and converters operate at sufficiently high frequencies to make effective use of Litz wire windings.

Heavy currents in adjacent turns of large utility transformers force non-uniform current density in the windings. This *proximity effect* acts in a similar manner to the skin effect, increasing the effective resistance of the conductors. Proximity effect also asserts itself in the high frequency transformers used in regulated power supplies, inverters and converters.

There can be yet another loss of current carrying capability. Eddy current generation in the winding conductors themselves dissipates power and produces an additional rise in temperature in the windings. Eddy current induction is perpendicular to the desired direction of winding current. In large transformers, massive windings are sometimes laminated to break the path of these eddy currents. The eddy currents induced in the winding conductors behave as a multitude of short-circuited secondaries and indirectly increase the effective resistance of the windings via the rise in temperature they cause.

At radio frequencies, we can identify yet another culprit contributing to copper losses, namely *electromagnetic radiation*. Escape of energy in this manner manifests itself essentially as an increased loading of the transformer. It is as if an unwanted resistance is connected across the secondary terminals. At power line frequencies, radiation does not merit consideration as a mechanism of dissipation. It slowly asserts its presence at several tens of kHz and can assume significance with transformers operating in the several MHz region.

Practical transformers are surrounded by leakage flux – magnetic force lines that complete their paths through the air instead of through the core. This often calls for care in the mounting and positioning of transformers as well as electronic components. Not only can such leakage flux cause inter-coupling problems but power dissipation is increased when leakage flux induces eddy currents in conductive material. During geomagnetic storms, large utility transformers can experience magnetic saturation from earth currents whereupon leakage flux becomes sufficiently strong to react violently with external metallic objects. In electronic systems, toroidal cores and cup cores are often favoured because of their relatively low leakage flux.

More can be done after ringing door-bells, lighting lamps and running toy trains

Transformer principles can be ingeniously applied to a variety of uses besides the obvious manipulation of a.c. voltages and currents. Two examples

are shown in Fig. 1.15(a) and (b). Tuning of the radio frequency stages in radio receivers can be partially accomplished by means of a copper slug that can be withdrawn or inserted (or rotated within) into a resonant inductor. In such a scheme, the copper slug has eddy currents induced in it and behaves as a short-circuited secondary winding. As such, it is able to *reduce* the effective inductance of the primary, thereby shifting resonance to *higher* frequencies. Instead of, or in conjunction, with the copper tuning-slug, a slug of magnetic material can also be used. Because its permeability is higher than that of air, the magnetic slug *increases* the inductance of the resonated primary in becoming more intimately linked with the primary flux. The resonant frequency is thereby lowered. (Q can be enhanced by permeability tuning but tends to be degraded by the copper slug.)

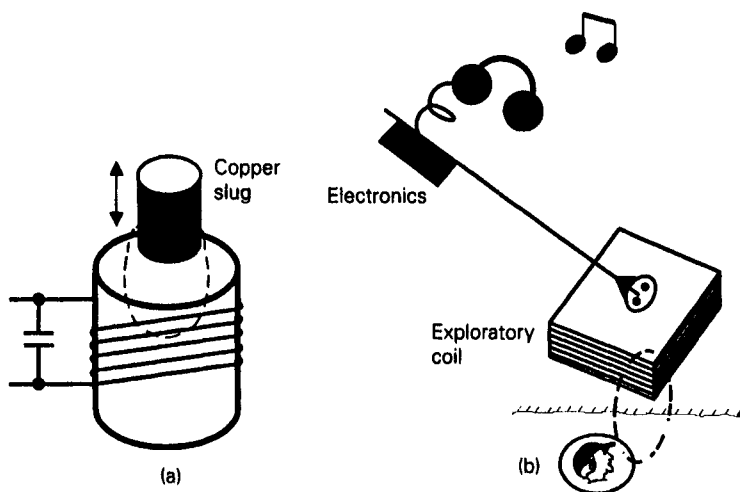


Figure 1.15 Interesting examples of transformer action. (a) In slug tuning, the copper slug acts as a short-circuited secondary and reduces the inductance of the resonant primary winding. (So-called 'permeability' tuning makes use of magnetic material which increases the inductance of the winding.) (b) The exploratory coil of a metal detector constitutes the resonated primary winding of a 'transformer'. The secondary is the metallic object buried in the earth in which eddy currents are induced. Detection takes place because of the reduction of inductance and resultant increase in the resonated frequency. (Ferro-magnetic objects tend to increase inductance and lower the frequency.)

Whoever first devised a *metal detector* must have been mindful of the transformer action mentioned above. The 'primary' of a simple metal detector is

a resonant winding generally several inches to a foot or so in diameter. (Square or rectangular loops can also be used.) This air core tank is associated with an oscillator. Frequencies of several hundred kHz to several MHz often provide a practical compromise between such factors as sensitivity, stability and ground penetration. Loosely coupled to the output of this oscillator is another oscillator with a very stable frequency, shielded from exposure to metal or objects in the environment.

With no metallic objects in the field of the exploratory 'primary' winding, the operator's headset produces a steady audible tone. When metal is detected, the pitch of the tone abruptly shifts in response to the changed resonance of the exploratory winding. (The audio tone is the beat note between the two oscillator frequencies.) More sophisticated models distinguish between ferrous and non-ferrous metals due to the fact that the beat note either increases or decreases in pitch. Some metal detectors used in detecting nails in tyres and in wall-studs incorporate similar *transformer* principles.

2 Specialized transformer devices

Admittedly, it is easy enough to view transformer technology as a boring arithmetical exercise for accomplishing voltage and/or current changes suitable to the particular task at hand. It might appear that it is 'old hat', regardless of whether we wish to ring a door-bell, run a toy train or operate the heater of a travelling wave tube. Some contemplation together with a fair amount of investigative effort reveals such an initial appraisal to be the tip of the iceberg; the topic of transformers embraces much more than the mundane situations that first come to mind.

Indeed, once we embark on the intellectual journey away from the 'basic' voltage-changing transformer, the topic becomes more interesting. Practitioners in electrical and electronic technology will likely perceive many of new ways of solving the ever-nagging technical problems encountered with long-accepted techniques. In particular, the association of these 'off the beaten track' transformer devices with modern electronic circuitry is bound to yield rewarding results.

Saturable core inverters – another way to use transformers

The use of special transformers in d.c./a.c. inverters is widely known. A number of such circuits utilize a pair of switching transistors working in tandem with a transformer that alternately saturates in one direction then the other. Because of the amplification provided by the transistors and the flux-switching characteristic of the associated transformer, the d.c. source is converted to a.c. power; in other words, we have an oscillator, but not of the LC, LR or RC type.

This may appear straightforward enough, but at one time such a scheme would have been rejected by electrical engineers. For one thing, very large

losses would have been anticipated if the principle was put to work with electron tubes. Mainly, however, the notion of deliberately driving a transformer core hard into its magnetic saturation regions was exactly contrary to the design philosophy expounded in the textbooks. Already lamented was the hysteresis loss resulting from operation of transformers in their 'linear' region; extending the flux excursions into the saturation regions (past the 'knee' of the magnetization curve) was sure to increase astronomically core loss. This was expected to happen as a consequence of both hysteresis and eddy current losses. Despite this, modern saturable core inverters often operate at over 90% efficiency *overall* – that is, including transistor losses. How was this accomplished? The salient factors are as follows:

1. Magnetic core materials (such as ferrites) were developed exhibiting *narrow* rectangular hysteresis loops. Because hysteresis dissipation is proportional to the area enclosed by the hysteresis loop, the narrow loops greatly reduced hysteresis loss. Also, the abrupt transition into saturation reduced the time spent in loss operational areas.
2. The new magnetic materials (again, mainly ferrites) feature high-volume resistivity. This greatly discourages eddy current loss. At the same time, it does away with the need for core laminations which tend to prevent full utilization of transformer materials.
3. Rugged high-speed transistors were developed that could be hard-driven so as to produce minimal voltage drop in the conducting state. When in the non-conducting state, these transistors have been found to act very much as open switches.
4. The evolution of power transistors from germanium to silicon devices allowed greater tolerance to operating temperature. Moreover, the practical aspects of inverter technology lends itself well to heat removal techniques; dedicated heat-sink and thermal hardware have become readily available. Maximum cooling effect can be realized from conductive, convective and radiative heat removal.

A typical saturable core inverter is shown in Fig. 2.1. This particular circuit allows the transistors to be mounted on a metal chassis or on heat-sinks, without insulating washers or gaskets. This is both thermally and electrically desirable and is the compelling feature of this grounded collector circuit. Circuits using the grounded emitter or grounded base configurations are also used but are at a thermal disadvantage because of the thermal resistance of insulating hardware.

The transformer shown in Fig. 2.1 is not true to life with regard to the physical separation of the windings. Both windings should be very close together in order to minimize leakage inductance. By the same token, it is generally best for both windings to occupy the major portion of the toroidal core. Leakage inductance manifests itself as 'spikes' on the voltage wave-

form. Not only do such spikes represent wasted energy, but they tend to be a major cause of transistor failure. This is because the peak voltages of the spikes penetrate the safe operating areas specified for the transistors. Leakage inductance can also cause damage due to spurious oscillation.

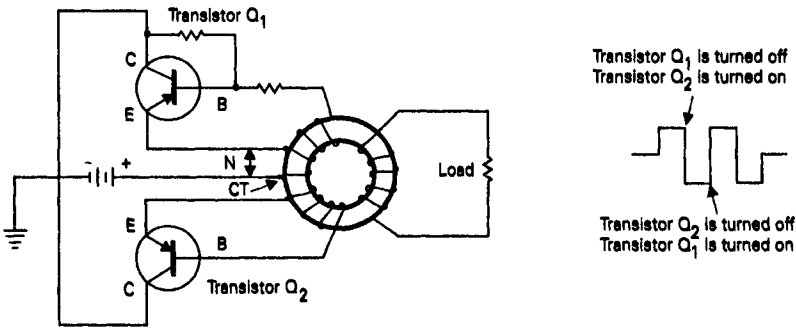


Figure 2.1 A typical inverter based on a saturating transformer. In this circuit, the auto-transformer type winding provides positive feedback to the bases of the transistors. The onset of saturation of the core alternately deprives one transistor then the other of such feedback current, resulting in sustained oscillation.

Various snubbing and energy absorption techniques are employed to attenuate the energy of these spikes, but the less energy so absorbed, the less inroad is made on operating efficiency. It is much better to start out with a tightly-coupled transformer incorporating a good magnetic circuit.

The hysteresis loop of the core is shown in Fig. 2.2, which pin-points the important switching events of the inverter. In a sense, it appears that the transformer is the active device responsible for generating the cyclic circuit action. Although the transformer is uniquely used, it must be properly recognized as a *passive* device. As such, its role is to determine the oscillation frequency in similar way that LC resonant tanks do so in other circuits. The required power amplification comes from the *transistors*.

Let us now follow the interactions between the switching transistors and the transformer core for one cycle of oscillation: Assume that transistor Q_1 has been turned *on*. As a result, there is a rapidly-increasing current between CT of the main winding and the emitter of Q_1 . This increasing current induces two important voltages. One of these voltages forward-biases Q_1 into heavy conduction, reinforcing its *on* state. The other induced voltage reverse-biases Q_2 , reinforcing its *off* state.

The portion of the winding carrying the current demanded by Q_1 ultimately causes abrupt core saturation. The electromagnetic induction previously mentioned can no longer take place. However, it does not immediately cease because of the energy stored in the core and the conductive

state of the transistors persists for a time. Ultimately, however, the collapse of the magnetic field gains speed and opposite biases are applied to the transistors, reversing their conductive states. Q_1 has now turned *off* and Q_2 has turned *on*.

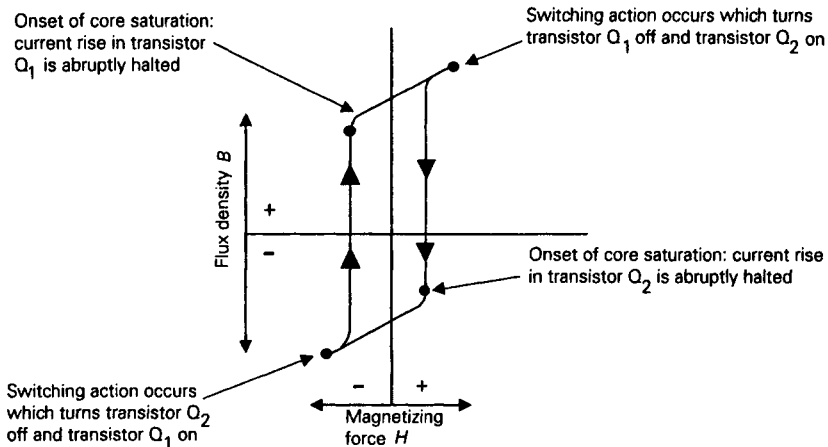


Figure 2.2 Sequence of magnetic events in the saturating core transformer. The conduction of the transistors results from forward bias developed in appropriate feedback windings. Whenever the core goes into saturation, the rate of flux change drops to a low value and one of the transistors will lose its forward bias. Because of the energy content in the core, such a transistor cannot turn off immediately but as soon as its forward bias is depleted, a regenerative switching action takes place which turns the other transistor on. The cycle is repetitive.

A similar sequence of events takes place for the remaining half of the oscillation cycle. Opposite transistors and oppositely polarized ampere turns on the transformer again bringing the core into saturation, but in the opposite magnetic sense as before. Thus, sustained oscillation takes place with the transistors continually alternating their conductive states. Note that switching transitions are *regenerative*, with the changes reinforcing themselves. Because of this, square waves can even be obtained from 50/60 Hz transformer steel with its sloppy hysteresis loop. (Efficiency is greater, however, with 'square loop' magnetic material.)

The feedback windings on the inverter transformers generally require some experimental optimization. A ballpark procedure is to design them for about 5V. This allows a reasonable value of current-limiting resistance to be used. It is necessary to hard-drive the switching transistors so that they are well-saturated while conducting. At the same time, too much base current agitates spike production and adds to circuit I^2R losses. Obviously, there is some interplay between the number of feedback turns and the value

of the current limiting resistance. If too many feedback turns are used, the power loss in the higher current-limiting resistance will make itself felt by a drop in operating efficiency.

In order to ensure reliable *starting*, it is common practice to provide a tiny bit of forward bias to just *one* of the transistors. An experimental high resistance will usually do the trick. Too much forward bias is not desirable as the transistors will run cool if oscillation is stopped by a short in the load circuitry. This biasing resistance can be seen associated with the upper transistors in the two inverter circuits.

The saturation flux densities of a few metallic magnetic materials are listed in Table 2.1. Included is silicon steel of the kind commonly used in utility frequency transformers. Reasonably good results can be obtained well into the lower audio frequencies. At higher frequencies, both hysteresis and eddy current core losses seriously erode into operating efficiency. The compelling feature of using this magnetic material is, of course, its widespread availability. (An interesting reflection on the nature of the saturable core inverter is that it would be futile to think of an air core version for operation at radio frequencies – magnetic saturation of air has not yet been observed.)

Table 2.1 Saturation flux densities for various core materials

<i>Core material</i>	<i>Saturation flux density in kilogauss (thousands of lines/sq. cm*)</i>
60-Hz power transformer steel	16–20
Hipersil, Silecron, Corosil, Tranco	19.6
Deltamax, Orthonol, Permenorm	15.5
Permalloy	13.7
Mollypermalloy	8.7
Mumetal	6.6

*When using these B_s values, make certain that the core area A , is expressed in square centimetres.

Maxwells (lines per square inch) can be obtained by multiplying the B_s values by the conversion factor, 6.25. Not included in the table are ferrite family materials. These generally saturate at lower flux densities. Such ferrite cores merit consideration at several kHz and higher frequencies.

The algebraic permutations of the design equation for the saturable-core inverter transformers are as follows:

$$N = \frac{E \times 10^8}{4fB_sA}$$

$$\frac{N}{E} = \frac{10^8}{4fB_sA}$$

$$A = \frac{E \times 10^8}{4fB_sN}$$

$$f = \frac{E \times 10^8}{4B_sNA}$$

$$B_s = \frac{E \times 10^8}{4fNA}$$

where N is the number of turns from centre tap (CT) to the first tap on the main winding. (The 'main winding' is the *primary* in the single transformer inverter circuits; it is the *secondary* of the small saturating transformer in two-transformer inverters.)

N/E represents turns per volt.

E is the voltage applied to the N -turn portion of the main winding; in practice, it is generally acceptable to use the d.c. supply voltage for E .

f is the oscillation frequency in Hz.

B_s is the saturation flux density of the core material. Some guesswork may be involved here, particularly with 50/60 Hz transformer steel.

A is the cross-sectional area of the core. (A is expressed in square centimetres if B_s is expressed in Gauss. A is expressed in square inches if B_s is expressed in Maxwells.)

The best saturable core inverters employ *two* transformers. Yet the added cost and complexity is very slight because one of the transformers is small and operates at a relatively low power level. A typical circuit of such a two-transformer inverter is shown in Fig. 2.3. Only one of the transformers saturates – the small one which drives the transistor bases. The other transformer delivers the square wave output power but is 'conventional' in that its operation is confined to its linear magnetic characteristics. The principle of operation remains essentially that of the single-transformer inverter with timing of the switching action now governed by the onset of magnetic saturation of the *small* transformer.

This type of inverter circuit tends to be more efficient, better behaved and easier to experiment with than the single-transformer configurations. Spike suppression is likely to be more straightforward because these transients now originate in a lower power section of the circuit. If an experimental circuit based on the scheme depicted in Fig. 2.3 fails to produce the expected oscillations, it would be reasonable to suspect inadvertent negative, rather than the required positive, feedback. In such a situation, the leads to either the feedback winding on the linear transformer or the primary winding of the saturable core transformer should be transposed.

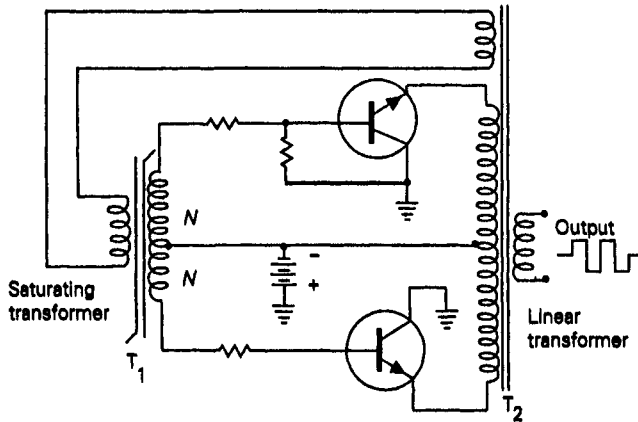


Figure 2.3 The two-transformer saturating core inverter. This is a more efficient oscillator than the single transformer circuit. In this arrangement, the saturating transformer (T_1), is physically small and operates at a relatively low power level. Therefore, it contributes minimally to overall losses. The linear transformer (T_2) is designed to avoid core saturation and so incurs low hysteresis and eddy current loss. Because of these facts, both transformers can be made from E-I laminations of 50/60 Hz transformer steel and still operate efficiently.

The design of a satisfactory saturable core inverter is an interesting challenge to both theoretical and empirical skills. Some of the desired performance parameters tend to be contradictory in nature. The inverter should be reliably self-starting at full load under a wide temperature range. Short-circuiting the output should halt oscillation, allowing the transistors to operate indefinitely at near quiescent current consumption. Nothing in normal use should provoke either high or low frequency parasitic oscillations. A good voltage-regulated d.c. power supply should be used in the interest of frequency stability.

Energy transfer without flux linkages – the parametric converter

The parametric converter ‘transforms’ energy from a primary to a secondary winding yet cannot be classified as a *transformer* in the accepted sense of the word. This is because the two windings do not share mutual flux linkages. In other words, energy transfer does not occur via electromagnetic induction. Although the device bears physical resemblance to ordinary iron core transformers, there is one notable difference in the basic configuration: the primary and secondary windings are positioned so as to be *decoupled* from one another insofar as sharing any common magnetic flux

is concerned. In transformer parlance, one would say that there is 100% leakage flux. Thus, no *ordinary* transformer action can account for transference of energy from primary to secondary.

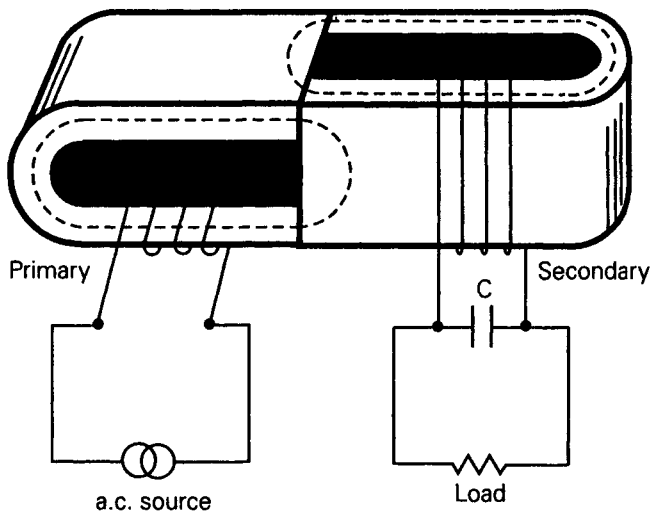


Figure 2.4 Topography of one type of parametric converter. The orthogonal relationship of the two 'C' cores is the important physical feature. Note that unlike the situation in transformers, the magnetic flux from the primary winding does not link the turns of the secondary. Also, the magnetic flux from the secondary winding does not link the turns of the primary. The device is operative in ordinary use when the secondary winding is resonated to the frequency of the a.c. applied to the primary.

One type of parametric converter is shown in Fig. 2.4. Note the orthogonal arrangement of the core sections. Although primary and secondary windings 'see' mutually shared core material, the magnetic flux from the primary does *not* thread through the secondary turns. It is obvious that this peculiar design is implemented with the deliberate intention of *defeating* electromagnetically-induced voltage in the secondary winding. Note the unique phase relationship in Fig. 2.5.

The simpler parametric converter shown in Fig. 2.6 makes use of a single ferromagnetic toroid with the windings arranged to preclude flux-linkage between them. Note that both types incorporate resonating capacitors. These are extremely important adjuncts and are not added optionally to smooth noise pulses or to bypass RF, for example.

In both types of parametric converters, the permeability of the magnetic circuit of the secondary winding is *modulated* by the primary flux even though the primary flux does not link the turns of the secondary winding.

This occurs because of the non-linear relationship between magnetizing force and flux density in ferromagnetic materials. This is tantamount to saying that the parameter, *inductance*, of the secondary is varied at the frequency of the primary excitation.

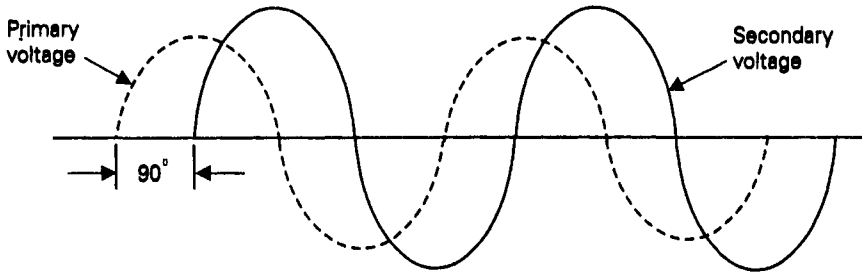


Figure 2.5 Phase relationship of primary and secondary voltages in parametric converters. Note that, unlike the situation in transformers, the input and output voltages are displaced by 90° instead of zero or 180° .

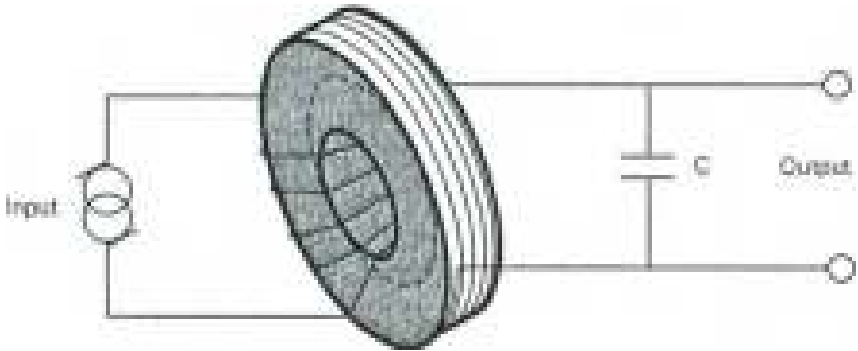


Figure 2.6 A simpler arrangement for parametric conversion. The physical relationship of the two toroidal windings meets the requirements for parametric conversion. Interesting experiments can be conducted with a two-inch OD powdered iron core and several dozen turns on the two windings. Operation in the 50 kHz region is easily achieved with excitation from a test oscillator and a $0.1 \mu\text{F}$ resonating capacitor for C .

This device is, in essence, an *oscillator*. This stems from the fact that if one of the reactive elements of an LC 'tank' circuit is varied at the LC resonant frequency, sustained oscillations will be developed in the tank. The basic reason the LC tank is not self-oscillatory is because of its inherent dissipative losses. Energy 'pumped' into the LC tank via the modulated inductance effectively overcomes these losses and sustained oscillation is the result. This

is not easily grasped intuitively, but it can be mathematically shown that *all* oscillators owe their operation to negative resistance or the cancellation of dissipative losses. What distinguishes different types of oscillator is the mechanism whereby this is done. Thus, it can be done with amplification and feedback, with shocked excitation by pulses, by depressing temperature to near absolute zero, or by *parametric modulation*.

Although it may not be easy to view this transformer-like device as an oscillator, the LC resonant tank comprised of the secondary winding and its associated capacitor undergo oscillation for the same reason that underlies operation of more conventional oscillators – *energy is supplied which replenishes that lost by dissipation* in the L and C elements. From a practical point of view, we can see that energy is transferred from the primary to the secondary, but in an *entirely different* way to the energy-transfer in ordinary transformers. This being the case, one can reasonably expect some unique operational characteristics from the parametric converter. Let us look at some of its salient features.

Suppose the device is in operation with the secondary either open or delivering power to a load. What would happen if a dead-short were placed across the secondary? Unlike ordinary transformers, there would be no tremendous rush of line current into the primary. Quite the contrary – oscillation would immediately *cease* and the device would be out of operation. Interestingly, primary line current remains quite high all the time and is not greatly influenced by secondary conditions. The parametric converter tends to be inefficient at low loads – the line must continue to supply energy to modulate the core permeability.

The parametric converter is an inherent line voltage regulator. That is, the load voltage delivered by the secondary remains fairly constant over a large range of line voltage variation. For example, such a device designed for normal operation from a 115 V, 60 Hz line might provide fairly constant voltage to a rated secondary load even though the line voltage varied from, say, 85 to 150 V. Below 85 V, there simply would be no oscillation because of insufficient permeability modulation of the core. Above 150 V, hysteresis loss from excessive core saturation would prevent oscillation. These are not hard and fast figures; in general, one might reasonably expect such line voltage regulation to be within $\pm 1\%$ over a wider than required line voltage range.

Another interesting and useful feature is high immunity of the sine wave developed in the secondary to both noise and harmonics riding on the waveform of the impressed line voltage. Even if the primary is excited by a square wave, the secondary wave will *still* be sinusoidal. By the same token, a load-generated transient will exert negligible effect on the primary circuit. Thus, one can think of the parametric converter as a high-Q bandpass filter in addition to its operation as an oscillator. The isolation attained between input and output circuits is also suggestive of performance ordinarily associated with vacuum tubes and FETs.

This isolation also manifests itself in load voltage regulation. Much of the load voltage regulation that does occur is due to the resistance of the secondary. Accordingly, large gauge secondary wire can make load voltage nearly independent of load current.

Interesting switchmode power supplies can be implemented in systems in which the parametric converter operates at several tens of kHz, rather than at the power line frequency. Such designs dramatically reduce size and weight, and lend themselves conveniently to empirical optimization.

It is not difficult to manufacture and operate a parametric converter if you can live with random frequency and output voltage, and with less than optimum efficiency of energy transfer. However, contradictory design requirements assert themselves if attempts are made to pin-point more than one performance characteristic. For example, the right number of secondary turns for a certain output voltage may result in a resonant- Q too low for satisfactory output power. To a considerable extent, a high C/L ratio improves the output power capability, but it is possible to go too far in this direction, whereupon output power will decline. If one is committed to a certain frequency, optimization of the performance becomes all the more complicated. Usually, trade-offs have to be accepted resulting in reduced performance.

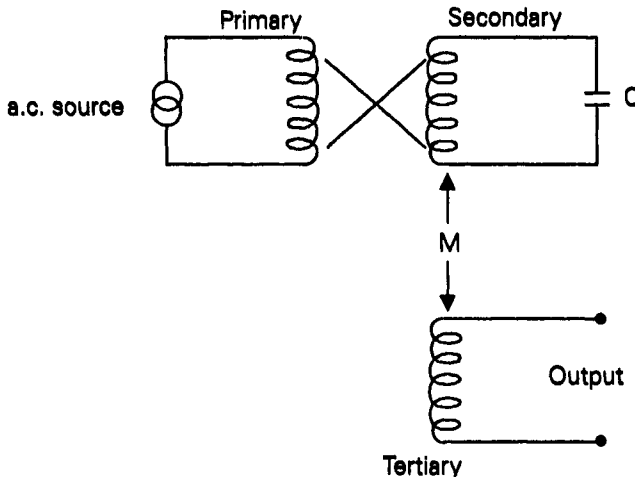


Figure 2.7 Parametric converter with tertiary output winding. By supplying load power from the tertiary winding, independent adjustment of the secondary inductance and the output voltage becomes feasible. Note that the coupling between the secondary and tertiary windings is the same as in conventional transformers due to the secondary and tertiary windings being physically adjacent. The resonant secondary, on the other hand, is parametrically coupled to the primary and produces the stimulated oscillation.

There is an easy way to circumvent these operational conflicts and to reduce the amount of empirical effort needed to improve performance. This entails the addition of a *tertiary winding*, electromagnetically coupled to the resonant secondary winding which allows the empirical adjustments of optimum resonant- Q and desired output voltage to be made independently. One first experiments with various L and C combinations to find which combination is capable of delivering maximum power to a load resistance. This is done by recording the power delivered to various sized resistances connected directly across the secondary winding. Be sure to use appropriate capacitance values to maintain resonance during the process (see Fig. 2.7).

Once the optimum secondary inductance has been determined using the above procedure, the tertiary winding can be wound with the required number of turns to produce the desired output voltage. Note that energy transfer from the secondary to the tertiary winding takes place by *ordinary transformer action*. The tertiary winding, being non-resonant, does not parametrically couple to the primary; it simply performs a voltage step-up or step-down function. An auto-transformer arrangement could also be used.

The Lorain subcycler – an erstwhile parametric converter

A device related to the parametric converter was the once extensively used *Lorain subcycler* which produced the 20 Hz ringing tone in telephones. This was parametrically divided down from the 60 Hz utility line in a scheme best described as a stimulated oscillator; that is, the 20 Hz output was oscillatory, but required the 60 Hz input excitation for its operation. (Those familiar with the *regenerative modulator*, a circuit utilizing amplifying elements will also recognize related phenomena.) The unique aspect of the subcycle ringer was that only non-linear elements were involved. Because of the lack of amplification, the stimulated oscillation, although drawing its energy from the 60 Hz input, was not self-starting. In order for the 20 Hz oscillation to commence, the arrangement had to be shock-excited by an electromechanical relay.

The basic subcycle ringing circuit is shown in Fig. 2.8. The actual device incorporated additional refinements, including a starting relay which momentarily shorted the non-linear inductance. With modern magnetic materials and components, this 'passive' frequency divider could be put to use by a creative experimenter. This is especially true since operation at tens and hundreds of kHz can be realized with powdered iron and molybdenum permalloy cores. Moreover, division can be obtained by both odd and even submultiples of the input frequency. In any event, there is more to the operating principle than meets the eye.

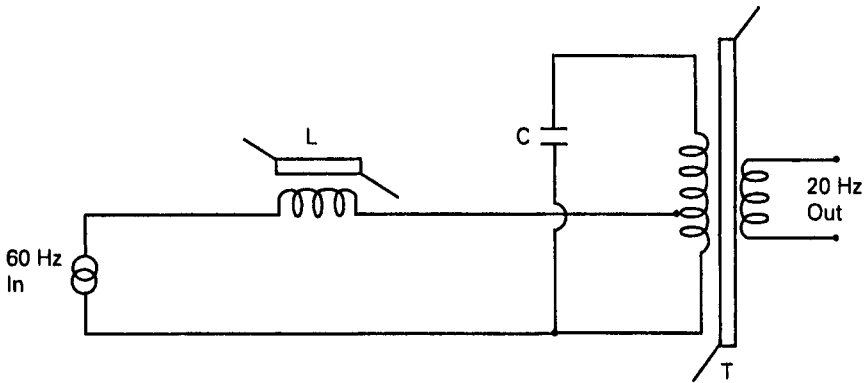


Figure 2.8 Basic circuit of the Lorain subcyclor. Once widely used for telephone ringing, the arrangement provides sustained 20 Hz oscillation, this being the resonant frequency of the tuned transformer (T). Both T and the series inductor (L) are non-linear. Because of these nonlinearities, a 40 Hz 'internal' signal is developed across L which then acts as a modulating element to deliver 20 Hz to the oscillating transformer. Thus, the process becomes self-sustaining. Starting is brought about by momentarily shorting L.

Let us examine the events occurring to enable a 20 Hz oscillation from a 60 Hz input source. This is most conveniently done by supposing the circuit to be already in operation. From our previous discussion of the parametric converter, it is easy to see that the energy needed to sustain a 20 Hz oscillation is provided by the 60 Hz source. The question, of course, is how 60 Hz can become divided down to the lower output frequency.

Note that the symbols used with the two core components are often used to depict *hard* saturation in such circuits as saturable core oscillators. However, hard saturation is not required for the subcyclor; it is only necessary that considerable non-linearity sets in for both the 60 Hz and 20 Hz voltages. As we have assumed the circuit *already* to be in operation, it can be expected that a fairly strong second harmonic of 20 Hz will be developed by transformer T because of its non-linearity. This 40 Hz second harmonic will appear across non-linear inductor L where it will mix and heterodyne with 60 Hz to again produce 20 Hz to reinforce the 20 Hz oscillation in the tuned transformer. It can be seen that the process 'feeds itself' deriving its basic energy from the 60 Hz input.

The basic configuration, as shown in Fig. 2.8, is not self-starting because no amplifying device is incorporated. Momentarily shorting the series inductor will, however, generate the needed transient to initiate sustained oscillation. As with more conventional oscillators, excessive loading will cause oscillation to cease. If this happens, the circuit must be restarted by again momentarily shorting L.

Proprietary information regarding the original Lorain subcycler is not readily forthcoming and it is quite likely that the experimenter will first find the arrangement more inclined to be passive than active. Much depends upon the core materials, the amount of excitation and the location of the transformer tap. Also, it can be helpful to connect a capacitor across L in order to resonate it to the second harmonic of the oscillation frequency. Experimentation is facilitated by also connecting a push-button switch across L so that it can be shorted every few seconds while making modifications. Non-linearity in the two cores can be enhanced with strong magnets or by applying d.c. bias through a high resistance. The resonant- Q of the tuned transformer has an optimum range – too high or too low precludes satisfactory operation. Initial experimentation should be carried out with no output load.

The constant current transformer

There are many situations in technology in which normally detrimental behaviour is exploited in appropriately designed devices. Consider, for example, *leakage inductance* in transformers. For the most part, much time and effort is invested in designs and implementations calculated to reduce leakage inductance to the very smallest value consistent with cost and manufacturing practicalities. This maximizes primary to secondary coupling, operating efficiency and power factor. Nonetheless, situations exist in which good use can be made of the inevitable leakage inductance and it may even prove beneficial to deliberately make it high.

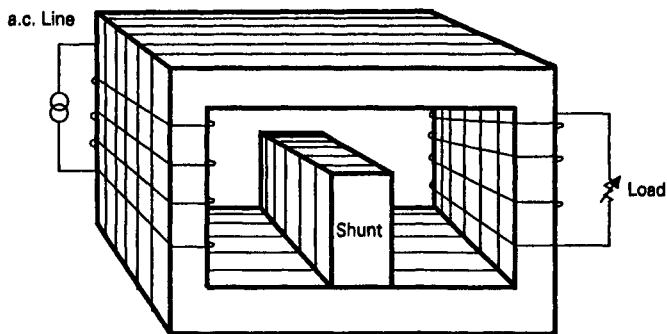


Figure 2.9 One type of constant current transformer. A core type transformer with windings on the most widely separated legs tends to have high leakage inductance. This can result in a near constant current characteristic. The magnetic shunt is optional but it can improve the current regulation. At the same time, its size, placement and possible air gap enables selection of the load current which will be stabilized against variations in load.

The ordinary looking transformer shown in Fig. 2.9 is a less than optimum arrangement from the standpoint of minimizing leakage inductance. It turns out, however, that one can benefit from what appears to be poor performance due to excessive leakage inductance, as can be seen in Fig. 2.10. Considering the operational range between points a and b, it is obvious that *current is nearly constant* over a wide range of voltages. Indeed, what we have is a type of constant current transformer.

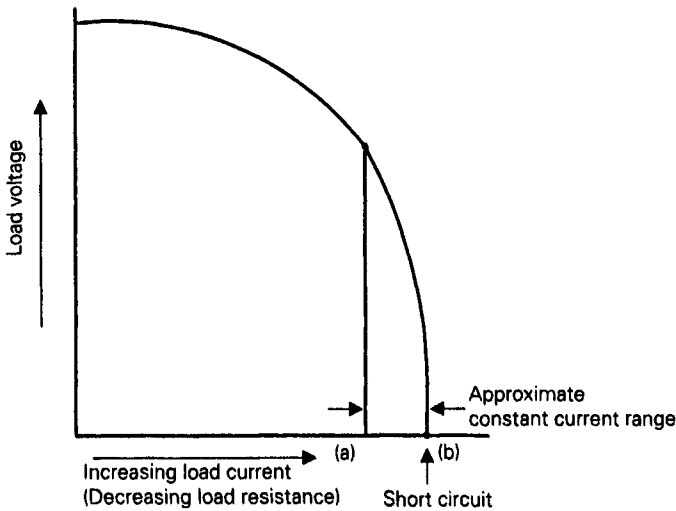


Figure 2.10 Load current stabilization by transformer with high magnetic leakage. Current constancy is not as tight as might be obtained with electronic feedback control. However, certain applications are well served with this characteristic. One is series street lighting systems. Another is the high voltage supply for neon lamps in advertising signs.

Neon and X-ray transformers are often manufactured in this fashion. Such high-leakage transformers have an immunity to catastrophic destruction from short-circuits. A magnetic shunt can be placed to divert flux that would otherwise tend to link the primary and secondary windings. This is tantamount to *increasing* the leakage flux even though the leakage takes place in the interior of the core structure, rather than external to the core. The net result of the magnetic shunt is to further decouple the two windings. By experimenting with magnetic shunts, one can select the near-constant current most suitable for the intended application. Note that no attempt to saturate any portion of the core is usually involved in the design of such transformers.

The constant voltage transformer

It may initially appear a paradox, but the constant voltage transformer basically depends on the same phenomenon at work in the constant current transformer. What is needed in *both* of these devices is relatively loose coupling between the primary and secondary windings. Once this is achieved, one gains considerable freedom in the control of current and/or voltage developed in the secondary. In other words, the secondary is no longer completely dominated by flux linkage emanating from the primary. Such a transformer exhibits a lot of *leakage inductance*. Let us see how this leads to both load and line regulation.

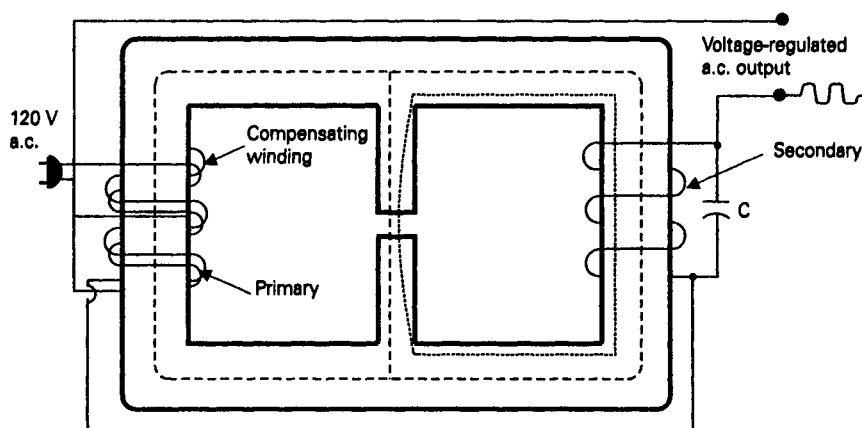


Figure 2.11 *The constant voltage transformer. Also known as the ferro-resonant transformer, no actual resonance exists between capacitor C and the secondary winding. Because of the magnetic shunt, a portion of the primary flux does not link the secondary. Also, due to the reactive capacitive current, the right-hand (secondary) section of the core operates in its saturation region. Stabilizing feedback is obtained via the action of the compensating winding, which carries load current in order to regulate the primary flux.*

The configuration of the constant voltage transformer is shown in Fig. 2.11. The resemblance to the constant current transformer described previously is quite evident. Specifically, both devices have core designs that promote high leakage inductance. In the constant voltage transformer, several additional implementations are operative. It should be realized that the alternate name, 'ferro-resonant transformer', tends to be confusing. Despite the 'resonating capacitor' connected across the secondary winding, this circuit is *not* tuned to parallel resonance. The objective here is to operate a portion of the core in the secondary circuit in its *saturation region*. Because of the additional reactive current drawn by the capacitor, such core saturation is

more readily achieved. With saturation in the magnetic circuit of the secondary, its induced voltage can no longer be entirely governed by either flux linkage or by load conditions. Such voltage stabilization is characterized by a near-square wave output. This waveshape is actually easier on rectifiers than is the usual sine wave.

This transformer also has its internal electromagnetic feedback provision in the form of the compensating winding which carries load current to *oppose* the primary magnetic flux. This promotes increased decoupling between primary and secondary. The result is better voltage regulation than would otherwise be achieved. The penalties one must pay in the use of this device are that they are much *larger* and *heavier* than conventional power transformers. They are also more costly, frequency sensitive and operate at a lower power factor. Balancing these drawbacks is the excellent *reliability* inherent in the constant voltage transformer.

Saturable reactor control of a.c. power

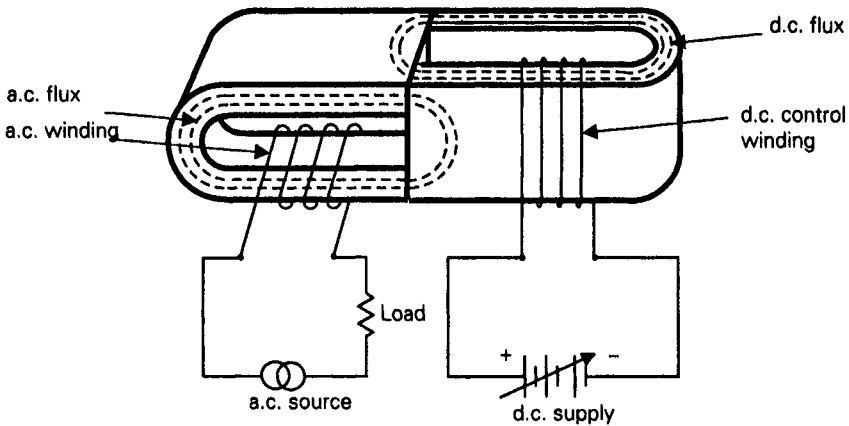


Figure 2.12 Saturable reactor with the structure of parametric devices. Note the unconventional relationship of the two 'C' cores. As a consequence, there are no mutual flux linkages to couple the two windings in transformer fashion. a.c. load current control takes place because of variance in the inductance of the a.c. winding. This, in turn, occurs via permeability control by the ampere turns in the d.c. control winding.

The saturable reactor depicted in Fig. 2.12 bears a physical resemblance to the parametric converter shown in Fig. 2.4. This is why it is convenient to associate the term 'parametric' with it. However, *no* parametric phenomena actually takes place in this control device. To begin with, there is no

resonance or oscillation involved in its operation. Indeed, it controls a.c. load current in much the same way as in the previously discussed saturable core devices, i.e. by *varying the inductance* of the a.c. winding via adjustment of the d.c. in the control winding.

As in the parametric converter, there is no conventional transformer action between the two windings. Mutual flux linkage between these windings does not exist because of the physical arrangement of the two 'C' cores. Nevertheless, a.c. and d.c. magnetic fluxes do share a common path in the central region of the device. Because of this, the d.c. ampere turns in the control winding are able to control this mutual core permeability, and therefore the inductance of the a.c. load winding.

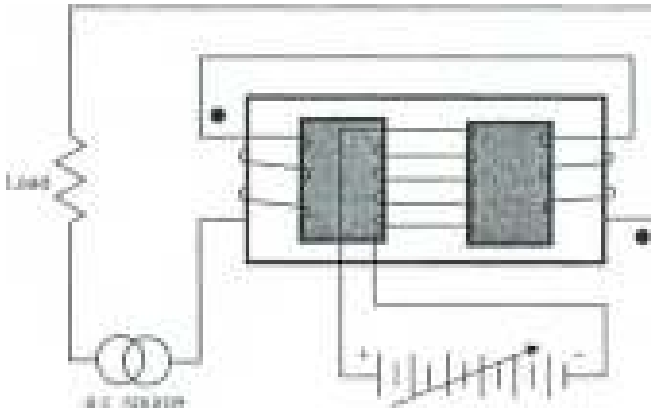


Figure 2.13 *Saturable reactor with isolated control winding. In this arrangement, operation is still as a magnetic amplifier, but no net a.c. voltage is induced in the d.c. control winding. Therefore, there is no need for a resistance or an inductance in series with the d.c. supply. Although usually found in 50/60 Hz applications, modern magnetic materials allow satisfactory designs through the audio frequency range and higher. At these higher frequencies, a pair of toroids can be substituted for the E or E/I laminations. (Some patience would be needed for winding the d.c. control coil, however.)*

The salient feature of this type of saturable reactor is that *no* a.c. is induced in the d.c. control winding. No isolation techniques are needed to keep a.c. out of the d.c. supply. By the same token, there is no tendency for problems in the a.c. winding from inadvertent loading effects. The author has carried out some interesting experiments with these devices in which both series and parallel resonance of the a.c. winding was investigated. Although this enhances the resemblance to a parametric converter, the

basic function of the saturable reactor was retained insofar as variable d.c. was still applied to the control winding. It was found that the influence of resonance in the load circuit tended to increase the sensitivity of the device – a heavier a.c. load could be controlled with less d.c. power. Experimentation is facilitated by working at low audio frequencies rather than at 50/60 Hz. (Select core material and laminations so as to avoid excessive hysteresis and eddy current loss.)

A more elegant type of saturable reactor can be made from E or E-I laminations and takes the form shown in Fig. 2.13. The central winding usually consists of many turns of small gauge wire. This is the d.c. control winding; no series resistance or inductance is needed as no net a.c. voltage is induced in this winding. The two outer windings are associated with the a.c. load and are wound with wire as heavy gauge as is consistent with the requirements of current-carrying capability and with physical constraints. When used with welders, for example, just a few turns of large wire suffice. Saturation from load current is avoided with adequate core volume.

The *phasing* of these load-supplying windings is of utmost importance. They must be connected in series so that their instantaneous magnetic fluxes cancel out in the centre bar of the 'E'. It is because of this flux cancellation that no net induced a.c. voltage is developed in the d.c. control winding. Fortunately, to attain this objective the proper phasing also corresponds to the *series-aiding* combination of the two a.c. voltages (otherwise, there would be no net inductance to be controlled).

This is also a magnetic amplifier in the sense that relatively little d.c. power is needed to control a large amount of a.c. power to the load. This control occurs as a consequence of varying the *permeability* of the core and therefore the *inductance* of the two a.c. windings. Despite the modern tendency to utilize Triacs, IGBTs, GTOs and other solid-state devices for controlling heavy currents and high power levels, an electromagnetic device of this kind remains advantageous for certain applications. It has a higher immunity to damage from transients, is more tolerant to temperature swings and is more likely to survive abusive operation. Another feature that deserves consideration is that it can be used in conjunction with either analogue or digital semiconductor circuits to achieve greater overall system versatility.

A two-winding transformer of conventional design can make a good current control device for a.c. circuits. One winding is connected as a series inductor in the a.c. circuit while an adjustable d.c. is applied to the other winding. The transformer thus becomes a magnetic amplifier in which the permeability of the core is controlled by the d.c. in the non-load winding. This is tantamount to control of the *inductance* of the winding associated with the a.c. load. This technique enables almost dissipationless control of the a.c. load current. The basic setup for utilizing a transformer in this manner is shown in Fig. 2.14.

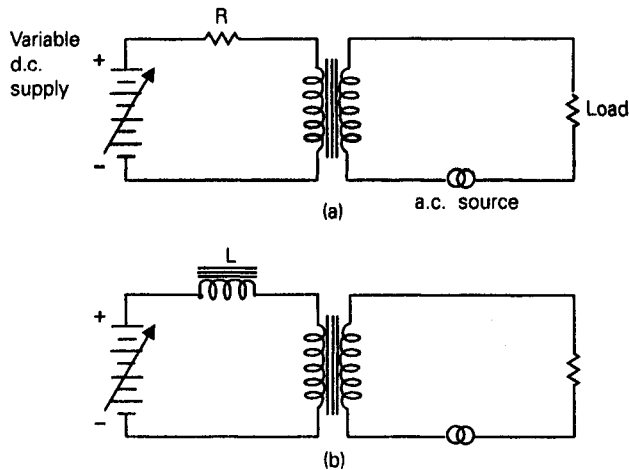


Figure 2.14 Using a transformer as a magnetic amplifier to control power in a load. A small d.c. varies the permeability of the transformer core. This, in turn, varies the inductance of the winding associated with the a.c. load current. A small change in the power level of the d.c. control circuit produces a large change in a.c. power delivered to the load. (a) A resistance is inserted in series with the d.c. supply in order to prevent loading of the a.c. voltage in the control winding. (b) A better method of isolating a.c. and d.c. in the control winding is with a series inductance. Various turn ratios can be used. d.c. polarity and transformer phasing are not important.

To make the scheme practical, some means of preventing shorting or loading of the a.c. voltage is needed, naturally induced in the non-load or control winding via ordinary transformer action. A simple way to do this is to insert a sufficiently high resistance in series with the d.c. source. The shortcoming of this method is that a considerably higher d.c. voltage may then be required to obtain a satisfactory control range. Also, the power loss in the resistance may prove objectionable.

A better approach makes use of an inductor connected in series with the control winding. This is shown in Fig. 2.14(b). With this arrangement, the d.c. control circuit is virtually unaffected, whereas the inductive reactance offered to the induced a.c. prevents any appreciable a.c. loading of the control winding.

Although many step-up and step-down transformer connections can be made to work, certain practical aspects tend to make particular implementations desirable. For example, the a.c. load winding must not have too much resistance. Ideally, control of the a.c. load current is to be accomplished via variation of a 'pure' inductance. Resistance is not so undesirable

in the control winding where a high number of turns results in less d.c. being needed for a given change of core permeability.

Usually, a few empirical trials with several transformers and with different ways of using their windings suffice to produce a workable arrangement. As a start, at least, it is best to aim for a desired control range of the a.c. load and to provide whatever d.c. voltage and current is needed.

Transformer action via magnetostriction

Figure 2.15 appears to represent an ordinary transformer demonstrating electromagnetic coupling. There may, indeed, be some conventional transformer action displayed by a primary and secondary coil wound on a rod of magnetic material. The presence of a d.c. biasing winding is not uncommon; in many circuits such magnetic polarization just 'comes along for the ride'.

There are, however, ways in which this arrangement can exhibit behaviour not observed in conventional transformer circuits:

1. The rod of magnetic material is a nickel alloy with the pronounced property of *magnetostriction*. This means that the rod changes its physical length in response to a varying magnetic field. The inverse phenomenon also prevails – the magnetic properties of the rod changes with mechanical deformation – the flux density is 'modulated'.
2. The rod undergoes a resonant response governed by its physical length. When excited at its resonant frequency, the magnetostrictive effect becomes very intense and a resonant voltage is developed in the secondary coil. At off-resonant frequencies, the effective coupling between the primary and secondary coils is very loose and only a very small conventionally induced voltage will be developed in the secondary. At resonance, it is as if much tighter coupling existed.
3. The phase of the resonant voltage is *opposite* to that associated with 'ordinary' transformer action. For example, the magnetostriction oscillator shown in Fig. 2.15 resembles a familiar Hartley circuit. However, there is no positive feedback via auto-transformer action as in an intentional Hartley oscillator. (Even if the coils were made into LC tanks with capacitors, the oscillations would not be of the Hartley type.) This is no LC oscillator but operates at a resonant frequency determined by the *magnetostrictive action* of the unique magnetic core.
4. Actually, the natural resonating frequency of the magnetostrictive rod is *twice* the exciting frequency. This comes about because the effect is 'blind' to magnetic polarity; it will lengthen or shorten once for each half-cycle of excitation, making it a frequency-doubler. Although such behaviour could be exploited for certain purposes, most applications

benefit from having the resonance occur at the exciting frequency. This is readily brought about by imparting a steady magnetic field to the rod. An extra d.c. bias winding can be used but is not always needed. The oscillator shown in Fig. 2.15 dispenses with such a tertiary winding because the biasing ampere turns are provided by the d.c. drawn by the amplifying device through the primary winding. Some magnetostrictive alloys expand, others contract with magnetization. Permanent magnet biasing is also used.

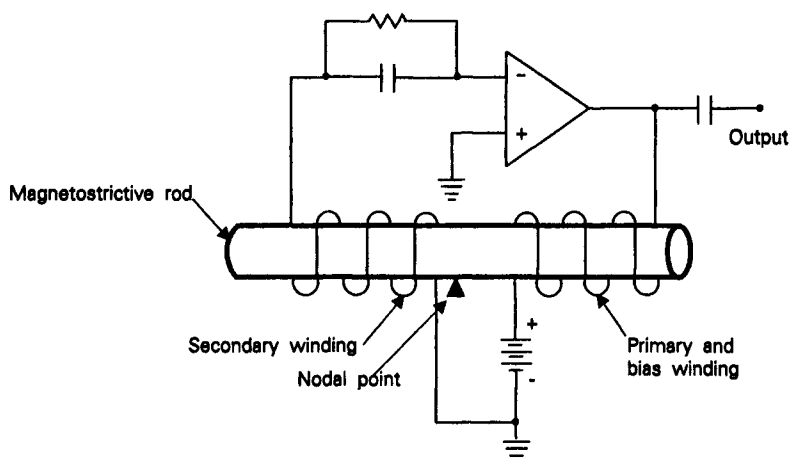


Figure 2.15 *The basic magnetostriction oscillator. The inductive circuit is phased to provide negative, rather than positive, feedback. Despite its topography, this is not a Hartley LC oscillator; instead, mechanical resonance of the magnetostrictive rod determines the frequency. The vibrating rod undergoes changes in its magnetization which then enables positive feedback from the transformer action of the windings. The d.c. consumed by the amplifying device provides the needed magnetic bias for the rod.*

5. Unlike the conventional transformer, the magnetostriction device cannot be physically mounted in just any way. Rather, the magnetostrictive rod must be supported at its nodal point, that is at a point along its length where vibrational amplitude is at a minimum, otherwise the effective Q will be damped. Such a 'loaded' rod will not be amenable as a reliable oscillating element. In contrast to frictional loading, the appropriate loading of the rod with metal disks raises Q and provides considerable control of the frequency response.

The resonant frequency with bias is simply one-half of the speed of sound in the rod divided by its length. The speed of sound in these nickel alloy rods is of the order of 4800 m/s. A 1m rod would therefore resonate at 2400 Hz. A

3cm rod would be expected to resonate at about 50 kHz. Between such extremes, one gets some insight into the practical limits of the frequency range of such rods. However, mechanical 'tricks' can extend the range considerably.

Excellent 60–600 kHz intermediate-frequency transformers were at one time used in high-quality communications receivers; these IF filters were based on magnetostrictive resonance. Unfortunately, they were very costly to produce and implement. However, their 60 to 6 db shape-factor of 1.2 to 1 remains difficult to surpass with crystal and ceramic resonators. The mechanical filter is shown in Fig. 2.16.

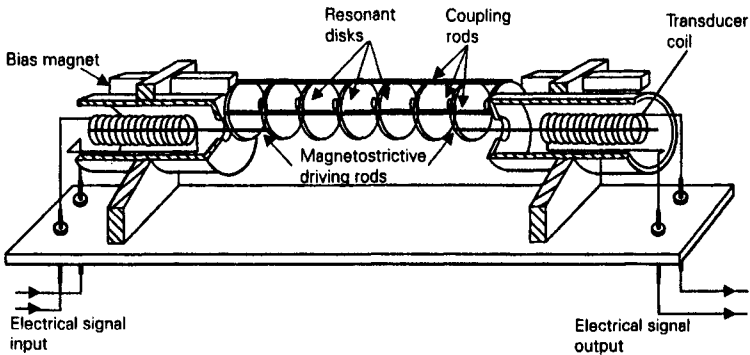


Figure 2.16 Mechanical filter – selective transformer action via magnetostriction. The nickel alloy driving rods are 'loaded' by the mass and compliance of the resonant metal disks. This is tantamount to adding inductance and capacitance and lowers the self-resonant frequency of the driving rods. The overall result is a nearly-rectangular frequency-response curve in the intermediate-frequency band of superheterodyne radio receivers. To reduce insertion loss, the transducer coils are resonated at mid-band by added capacitors. The device is symmetrical in that the input and output can be reversed.

A double core transformer for saturable core inverters

Many solutions have been proposed and implemented to tame the switching spikes produced by saturable core oscillators. These circuits are basically d.c. to a.c. inverters and have a myriad of uses. The switching transients, however, not only tend to damage or destroy the transistors but create a great deal of EMI and RFI. These, together with remedial circuitry, also can make serious inroads in the efficiency and tend to be detrimental to other performance parameters as well. It is always best to have as little spike generation as possible before adding snubber circuits and energy absorption

networks. Other things being equal, it is first necessary to design the saturating transformer so that it exhibits minimal leakage inductance. In practice, one can only go so far in this direction due to the contradictory requirements of the various performance parameters. Also, it is not easy to anticipate the effect of saturation on leakage inductance.

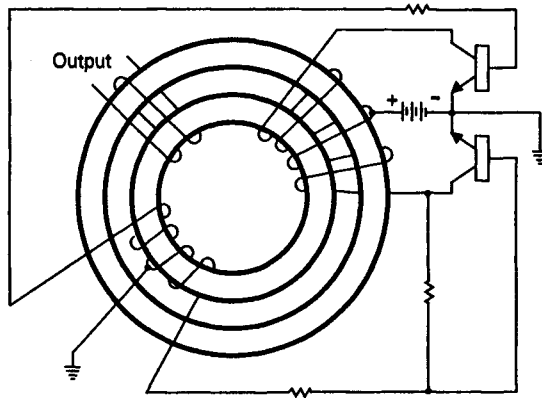


Figure 2.17 Double core transformer for saturable core oscillator. Conventional circuitry is employed except that the base-drive is supplied by the inner-core of the arrangement. If the inner-core is of the same magnetic material as the outer-core, saturation will occur first in the inner-core. This causes transistor switching to take place while the collector winding still has appreciable inductance, not yet being in its saturation region. The collector current spike is thereby greatly reduced. (Transpose base leads if the circuit does not oscillate.)

An interesting scheme for reducing the current spikes is shown in Fig. 2.17. Making use of two cores, this transformer is by its nature experimental, but it does point in the right direction. The two cores are made of the same magnetic material. The inner core will tend to saturate *first* because of its shorter circumference – it will ‘see’ more magnetomotive force per unit length than the outer core. This being the case, the transistors will *already* be switched by the time the outer core saturates. The collector switching spike will be considerably attenuated by the linear inductance of the outer core. This contrasts to the near-zero impedance offered to collector current in conventional single transformer circuits.

Note that the collector and the output windings embrace *both* cores. Also, it should be understood that, unlike Fig. 2.17, the windings should occupy a major portion of the cores in order to keep leakage inductance low. In order to keep winding resistance low, a minimal air gap should exist between the two cores. This approach should prove an experimenter’s delight – after a bit of experience with the scheme, further optimization of performance may be forthcoming with cores of different magnetic materials.

3 Operational features of transformers

It would be nice if the operation of a transformer was just a matter of applying voltage to the primary and meeting the load's requirements from the secondary. One tends to think of the transformer as a passive device with little capability of any gross malperformance. However, experience teaches that the passivity can only too easily translate into active malperformance. Somewhat surprisingly, this pertains to a 'good' transformer, not necessarily one with shorted or open windings (although these defects often enough develop from the malperformance).

In the first place, our perfectly 'good' transformer is, in practice, always a considerable departure from the role model – the *ideal* transformer. The practical transformer will be endowed with such blemishes as winding resistance, core losses, leakage inductance, exciting current, non-linearity, and an ineffective thermal situation. The core can have residual magnetism and some core materials will suffer irreversible damage from an excessive temperature rise.

As a consequence of these real-life defects, transformers can display undesirable characteristics in a circuit or system. These include audible noise, heat development, transient generation, harmonic production, oscillation and ringing, frequency attenuation, self-resonance, waveform displacements from d.c. components and current in-rush. Also, it would not be far-fetched to add expense, bulk and weight. Lest excessive cynicism sets in, it should be mentioned that many *desirable* modes of operation can also be realized by paying heed to the transformer's practical behaviour. In this manner, shortcomings can often be exploited profitably as *useful* responses. As will be seen, much depends on *how* the transformer is operated.

Operation of transformers at other than their intended frequencies

Hobbyists and experimenters often consider 50 and 60 Hz transformers to be interchangeable. Although no immediate consequences may be suffered from this assumption, it is not good engineering practice to ignore the different operating conditions. To get away with disregarding the frequency-dependent parameters of the transformer, one must usually rely on over-design, thermal inertia or just plain luck. With regard to the latter, one sometimes stumbles upon a transformer apparently designed for 55 Hz which is equally at home on 50 and 60 Hz power lines. However, in the majority of cases, it will prove best to comply with the following. An important parameter which generally should not be allowed to change when operating a transformer at other than its rated frequency, is the maximum allowable flux density, B_m . This is straightforward enough. The voltage applied to the primary and the new frequency must be changed in the *same* direction and by the *same* percentage. For example, to properly and safely operate a 120 V, 60 Hz transformer from a 50 Hz power line, the primary voltage must be *reduced* by the fraction, 50/60, yielding 100 V. Although this will, indeed, make the transformer 'happy', there are two side effects: the systems planner must now figure out how to conveniently and efficiently reduce the line voltage and, once having done this, how best to accommodate the new secondary voltage (after all, the turns ratio of the transformer has remained unchanged).

It should be realized, too, that the kVA capability is now 5/6 of its 60 Hz rating. Somewhat paradoxically, the operating efficiency of the transformer may show an *improvement* when operating from 50 Hz because of the lower core losses (see Fig. 3.1). It stands to reason that the *converse* situations apply when operating a 50 Hz transformer from 60 Hz. Here, the kVA rating may now exceed that pertaining to the original 50 Hz rating. The gain, however, will tend to be less than $\frac{1}{5}$ because of the increased core losses. The primary can now use $\frac{6}{5} \times 120$, or 144 V.

In making the adaptations described above to a new operating frequency, care is needed to keep the secondary *load current* within original limits; a load change may be required.

The case of operating the 50 Hz transformer from a 60 Hz line may, in some situations be easier to implement if one can accommodate *less* than maximum kVA capability. Here, it may be acceptable to merely substitute the new power line, not making any changes in either the applied primary voltage or in the load. We will have violated the postulate to keep the maximum allowable flux density at the designed level, but the departure is on the safe side. The 50 Hz transformer will simply operate as if it had been designed with a greater than needed safety factor. Operation will be cool and linearity will be good, but the kVA rating will be about $\frac{120}{140}$ or 83% of what it *could* be at elevated line voltage.

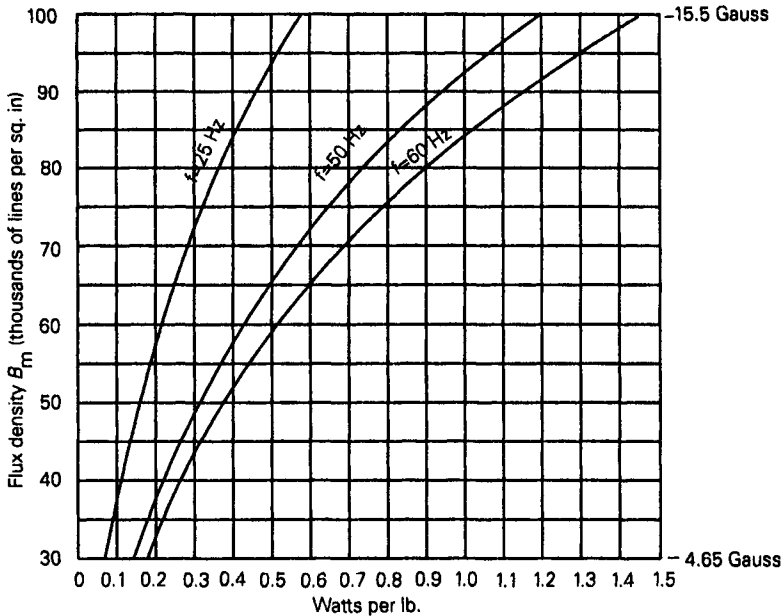


Figure 3.1 Core losses in 0.014 inch, 4% Si steel transformer laminations. These design curves sum up hysteresis and eddy current losses in the much-used core material for utility frequency transformers. (Multiply lines per square inch by 0.155 in order to represent the flux density in Gauss, or lines per square cm.)

A nice thing about the transformer substitution described above is that one can legitimately plan for a smaller, less massive transformer for a final prototype of the system. (It is easy enough to observe that appropriately designed high-frequency transformers tend to be physically smaller than their lower-frequency counterparts.)

The preceding discussion pertains to power line frequencies of sinusoidal waveform. With square waves and audio frequencies, various factors come into play tending to make transformer comparisons more nebulous. A few of these can be listed as follows: type of magnetic core material, core fabrication (butt joint, lap joint, tape-wound, powdered material, air gap etc.), primary inductance, distributed and stray capacitance, lamination thickness, insulation quality of laminate surfaces, skin-effect, leakage inductance, harmonic response, and others. Empirical investigation is generally the best practical procedure with such operation.

A special situation arises when dealing with transformers in a variable duty cycle circuit, such as in pulse-width modulation (PWM) switchmode power supplies. Unique in such operation is the presence of a d.c. component in the waveform. This situation is dealt with elsewhere.

Next to magic, try the auto-transformer for power capability

An amazing feature of the auto-transformer is its possible power-handling capability. Compared to an isolation-transformer of the same size, an auto-transformer can sometimes exhibit load-supplying ratings many times higher. This is best illustrated by considering a conventional isolation-transformer reconnected as an auto-transformer. In demonstrating such a comparison, round numbers for ease of mental arithmetic will be used. Also, it will be assumed that such minor factors such as exciting current, leakage inductance etc. can be safely disregarded in order to exemplify the basic idea of relative kVA ratings.

Let us deal with a 50 kVA two-winding transformer. In normal use the primary is impressed with 10 kV and the secondary delivered 2.5 kV. Such an isolation transformer is shown in Fig. 3.2. If this isolation transformer is reconnected as an additive auto-transformer, we have the situation depicted in Fig. 3.3(a). At first glance, the operating parameters do not appear to be radically changed aside from loss of the isolation feature. A little work with the numbers reveals otherwise, however.

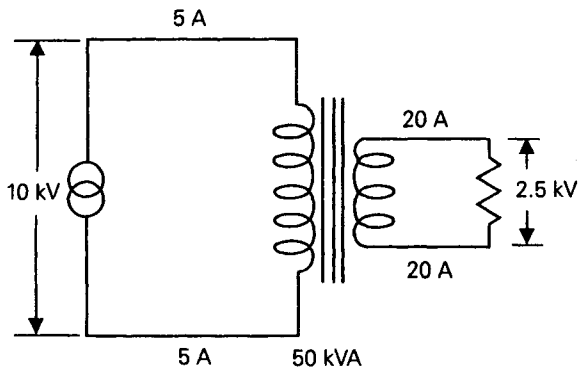


Figure 3.2 Example 50 kVA isolation-transformer to be reconnected as an auto-transformer. If the loss of isolation between the primary and secondary circuits can be tolerated, the auto-transformer connections can provide the experimenter with a number of different voltage transformations.

The allowable current in the 2.5 kV winding of the auto-transformer is $\frac{50 \text{ kVA}}{2.5 \text{ kV}} = 20 \text{ A}$

The allowable current in the 10 kV winding of the auto-transformer is $\frac{50}{10} = 5 \text{ A}$.

These currents combine additively to provide the new load current of 25 A. Then the load is supplied with $10 \text{ kV} \times 25 \text{ A} = 250 \text{ kVA}$. Note that this is *five-times* the kVA rating when the isolation-transformer connections

were used. Where did the additional power-handling capability come from? The answer is that the auto-transformer supplies some of the load by electromagnetic induction, but some by *conduction*. As the transformation ratio approaches unity, the *preponderance* of power transfer occurs by conduction.

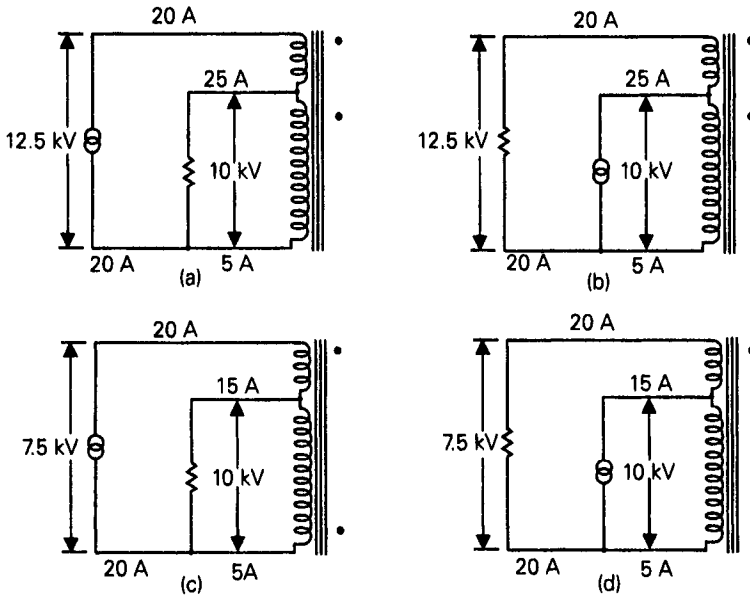


Figure 3.3 Auto-transformer connections of the 50 kVA two-winding isolation-transformer. (a) Additive-phasing with 10 kV winding as output; (b) additive-phasing with 10 kV winding as input; (c) subtractive-phasing with 10 kV winding as output; (d) subtractive phasing with 10 kV winding as input. Note the 250 kVA power-handling capability in circuits (a) and (b).

The extension of power-handling capability exemplified in our numerically illustrated auto-transformer can apply in a similar fashion for subtractively-polarized connections, and for either step-up or step-down of applied voltage. In all cases, power-handling capability becomes very large for near-unity ratios of voltage transformation. Conversely, when this ratio exceeds 10 or so, not much is to be gained over the power-handling capability of an equivalent two-winding isolation transformer. Note, however, there is little gain in Figs. 3.3(e) and (f), and an actual *reduction* in power-handling ability in the circuits of Figs. 3.3(g) and (h).

The situation prevailing for unity transformation ratio is virtually that of the a.c. line itself. This is easy enough to visualize, for all we have is a high inductance connected across the line. *All* power transfer then occurs via

conduction and *none* by electromagnetic induction. A small exciting current flows in this shunt-inductance, but is of negligible consequence in the scheme of things.

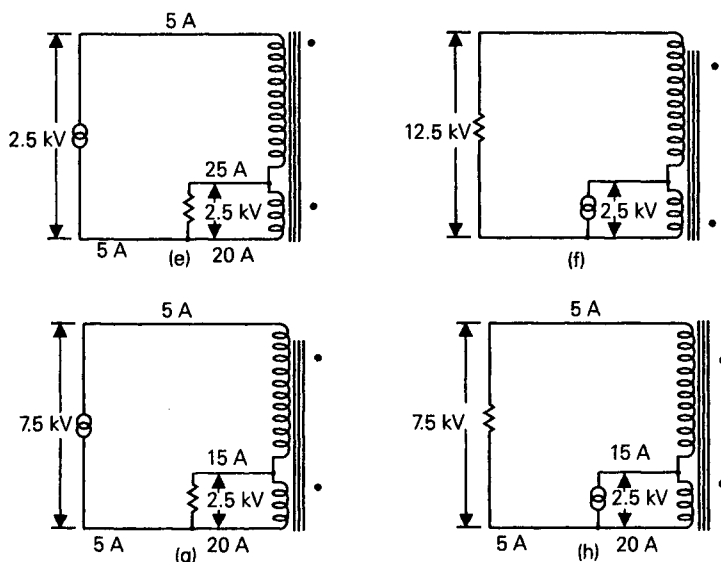


Figure 3.3 (continued) (e) Additive-phasing with 2.5 kV winding as output; (f) additive-phasing with 2.5 kV winding as input; (g) subtractive-phasing with 2.5 kV winding as output; (h) subtractive-phasing with 2.5 kV winding as input. Note the reduction in power handling capability of circuits (g) and (h).

From what has been said, one can reasonably infer that the operating efficiency of an auto-transformer is likely to be even higher than that of an equivalent rating isolation-transformer. The auto-transformer can get by with a smaller core and the tapped single winding will have a shorter length than the sum of the two windings in the isolation-transformer. Translated into the parlance of losses, these features lead to lower core and copper losses. Efficiencies exceeding 99% are possible.

Before leaving the topic of auto-transformers, it is appropriate to point out a *danger* associated with this transformer's format. In Fig. 3.4(a) an auto-transformer is operating between high and lower-voltage circuits. In the situation depicted in Fig. 3.4(b), an open circuit has developed as indicated. Because of this not uncommon defect, the 'output' side of the auto-transformer is now at the high voltage of the input side. The hazard to personnel and equipment is obvious. Unlike the utility industry, the safety factor in electronics work is less compelling, however.

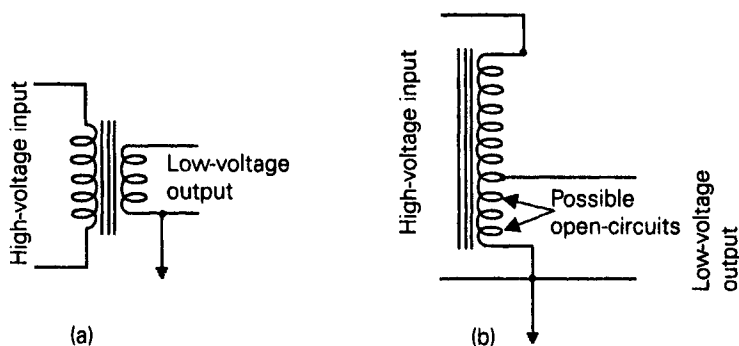


Figure 3.4 Possible hazard from an auto-transformer application. (a) A conventional isolation-transformer interfacing high- and low-voltage circuits. An open-circuit in either primary or secondary simply interrupts the low-voltage output. (b) An auto-transformer used for the same purpose. Note that if an open-circuit were to occur at the end-terminals of the low-voltage portion of the winding, high voltage would appear at the low-voltage output line. Such conjecture is not a fanciful fear as the low-voltage portion of the auto-transformer winding already operates at high current, sustained overload can develop 'hot-spots' at the winding terminations which can ultimately result in an open-circuit.

As if the possibility of a fault placing high voltage on the low-voltage side of a utility system were not a serious enough shortcoming of the auto-transformer, another of its characteristics serves to limit a more extensive use. Other things being equal, the auto-transformer tends to deliver a short-circuit current roughly greater than that of an isolation-transformer by the same ratio as its kVA advantage. The very operating parameters that endow the auto-transformer with superior power-handling capability, regulation and efficiency unfortunately reduce its internal impedance to heavy overload to the extent that there is minimal self-protection.

Of course, a classical argument on behalf of the auto-transformer points to the use of external protective devices and circuits. When we look at other application areas other than the utility industry, it is quite clear that the economy, convenience, efficiency and flexibility of auto-transformers are generally welcome wherever the lack of isolation poses no problem. In any event, it should prove rewarding for the experimenter to keep in mind that useful transformation ratios can often be realized via various auto-transformer connections of the windings on ordinary isolation-transformers.

Dial a voltage from the adjustable auto-transformer

The adjustable auto-transformer such as depicted in Fig. 3.5, is one of the most useful devices in the laboratory. It allows a smooth variation of a.c.

voltage from zero to about 110% of line voltage. Although it appears to be a simple enough device, there could be formidable arguments against its successful implementation if it did not already exist. It would clearly be seen that bringing a tap out from every turn would be a tap-switching nightmare. It would also be realized that the notion of a brush contacting only one exposed turn at a time must be disregarded as being both electrically and mechanically impractical. Seemingly, the allowable compromise of tap-switching just six or so voltages defeats the dream of providing *continuous* voltage variation.

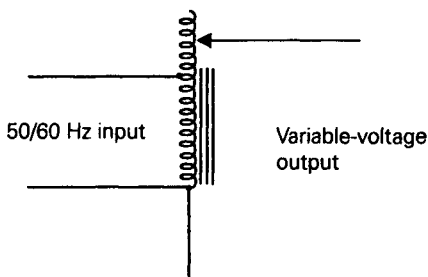


Figure 3.5 *The adjustable voltage auto-transformer. A continuously variable output voltage from zero to about 110% of the input line voltage is available. In auto-transformers, when the transformation ratio is not far from unity, a large part of power transfer occurs conductively. This enables a device of this kind to have a respectable kVA rating despite its modest size. Care must be observed with regard to the lack of electrical isolation between input and output circuits.*

Fortunately, there is a practical solution to our dilemma. Although abhorrent to the purist, we can have the brush contact several turns at a time. The short-circuit currents thereby produced can be relatively benign if our design incorporates two stratagems:

1. The core must be selected to provide a low volts per turn ratio.
2. The brush resistance must be as high as is consistent with reasonable efficiency and voltage regulation of the transformer. Brushes for the larger auto-transformers can be laminated structures with higher lateral rather than longitudinal resistance. In any event, the heat removal provided by convection helps make the scheme practical. Additional heat removal is provided conductively by the aluminium rotor which is made fairly massive for this purpose.

Commonly referred to as 'Variacs' from the original name bestowed by The General Radio Company, these adjustable auto-transformers usually have 50/60 Hz ratings. They are available also in triple-ganged assemblies for use in three-phase systems.

Transformer phasing in d.c. to d.c. converters

A not-uncommon transformer-phasing quandary pops up in engineering laboratories experimenting with d.c. to d.c. converters. The simplified circuits of two widely-used topographies are shown in Fig. 3.6. Circuit (a) is the *forward* converter. Circuit (b) is the *flyback* converter. Despite the similarity of the two converter circuits, both the operational mode and the application areas are different for these circuits. The forward converter has the higher power-handling capability and can more readily deliver tightly-regulated and well-filtered d.c. The flyback converter is uniquely adapted to produce a high-voltage, low-current format.

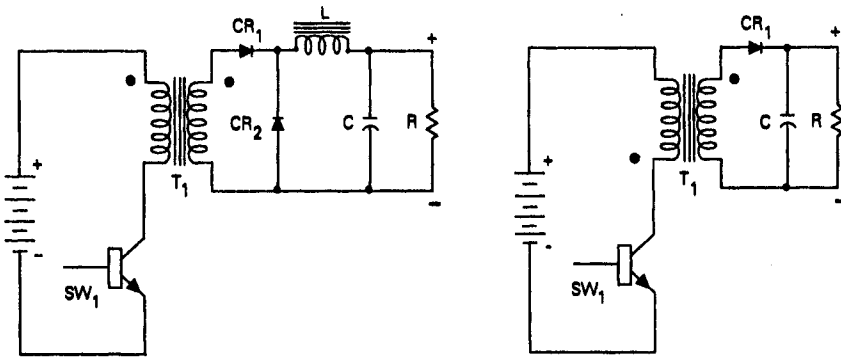


Figure 3.6 Transformer phasing requirements in d.c. to d.c. converters. (a) Basic forward converter: closure of the solid-state switch causes current to be sent through diode CR_1 , through inductance L , and to the load. When the switch opens, load current is maintained from the energy stored in inductance L , the path then provided by free-wheeling diode CR_2 . The secondary voltage induced by the opening of the switch is blocked by then reverse-biased CR_1 . (b) Basic flyback converter: while the switch remains closed, no load current is passed by reversed-biased CR_1 , but energy is stored in the transformer primary. Opening of the switch enables the resultant pulse in the secondary to charge capacitor C to near peak voltage. Diode CR_1 then prevents discharge of this capacitor by the secondary. Transformer phasing is different for the two circuits.

The significant difference between the two converters is the *phasing* of the transformers. If one is alert, this difference should be immediately evident from the phasing dots which accompany the circuit diagrams. Often, these 'go into one eye and out of the other'. They may also be misplaced or altogether absent. As either circuit can give *some* measure of performance with an incorrectly-phased transformer, a rather elusive troubleshooting task can ensue.

In the *forward converter*, closure of the solid-state switch sends current through rectifying diode CR_1 , and through inductance L to the load. When

the switch opens, the oppositely-polarized secondary voltage reverse biases CR_1 , effectively taking it out of the picture. However, the energy storage in the inductance L gives rise to *continuation* of load current, now completing its path through free-wheeling diode, CR_2 . This is a fortunate situation, for the load is supplied with nearly-constant current. A bit of filtering readily leads to a very smooth output. This type of operation cannot occur if the primary leads *or* the secondary leads of the transformer are transposed.

In the *flyback converter*, no current flows from the transformer secondary as long as the solid-state switch *remains* closed. However, during such time, energy is being stored in the primary inductance of the transformer. When the switch *opens*, the abrupt collapse of this stored energy induces a high-voltage pulse in the secondary. Diode CR_1 then enables capacitor C almost to charge to this peak voltage.

Transformer behaviour with pulse-width modulated waveforms

On several occasions during involvement with the design of regulated power supplies, the author has witnessed system malperformance from a straightforward, but often unappreciated, aspect of transformer operation. Engineering textbooks deal in great detail with transformer operation from a sine wave power source. Somewhat less mention is accorded to square wave operation, and harmonic-laden waveshapes are discussed primarily as they stem from the non-linear characteristics of the core material. One does not, however, readily find treatment of transformer response to pulse-width modulated waves such as are commonly used in switching power supplies. It turns out that the discourses on square or rectangular waves usually do not suffice to alert the designer to the peculiar behaviour of transformers to variable duty-cycle, or PWM, waves.

Consider, for example, a power MOSFET switching element in a regulated d.c. supply. The amplitude of the gate-drive voltage required for device saturation and efficient switching is often about 15 V. The basic mechanism of output voltage regulation of these supplies entails modulation of the on-time of the switching cycle. In order to accomplish this, the gate drive voltage should maintain approximately its 15 V amplitude, but should undergo a variable duty-cycle, as controlled by the feedback loop of the supply.

At first sight, a drive transformer tested for 15 V peak square wave amplitude at the operating frequency of the power supply should do the job. 'Common sense' might well suggest everything should go well provided the frequency capability of the transformer did not impose any restrictions. This, unfortunately would *not* be the case even for a transformer with overall ideal parameters. Although the explanation of impending trouble tends to be elusive, it is not difficult to grasp.

Consider the situations depicted in Fig. 3.7(a)–(c). In Fig. 3.7(a), a 30 V unidirectional square wave is applied to the primary of a MOSFET drive-transformer. The secondary waveform has positive and negative excursions of 15 V and preserves the 50% duty-cycle of the primary voltage waveform. Let us now see what happens for duty-cycles *other* than 50%:

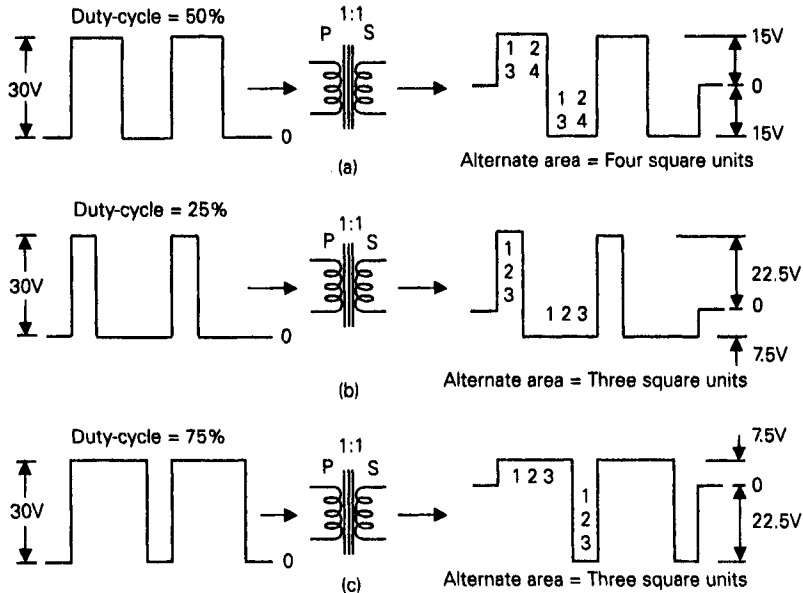


Figure 3.7 The effect of pulse-width modulated waveforms on transformer behaviour. The peak-to-peak amplitude of the positive plus negative excursions in the secondary winding remain constant. However, the peak amplitudes of these excursions vary with the duty-cycle. The basic rule is that the areas above and below the zero-axis of the secondary waveforms remain equal to one another regardless of duty-cycle.

In Fig. 3.7(b) the amplitude of the peak-to-peak voltage on the primary of the transformer remains at 30 V. The duty-cycle is now 25%, meaning that on-time (the positive excursion) is now one-quarter of the cyclic period. Note that the peak amplitude of the positive excursion of the secondary waveform is now 22.5 V and that of the negative excursion is 7.5 V. The gate of the MOSFET could now be endangered, particularly at higher than normal a.c. line voltage. Before proceeding with our investigation, let us see where these numbers come from.

The induced voltage in the secondary of a transformer must show a steady-state equality of volt-seconds for the positive and the negative excursions. If this were not so, there could be a steady d.c. component in the secondary. However, we know that d.c. in the primary of a transformer cannot induce sustained d.c. in the secondary. Accordingly, the transformer automatically adjusts its secondary waveform to preclude possibility of a steady d.c. component. Graphically, this statement translates itself into stating that the positive and negative *areas* of the secondary waveform must be equal. This rule has been applicable all the time, but we did not have to deal with its consequences for 50% duty-cycle waves, whether sinusoidal or square.

Referring to Fig. 3.7(c), we see an even more drastic consequence of the constant volt-second rule. For the 75% duty cycle primary wave, the positive excursion of the secondary voltage is only 7.5 V – probably too low to drive the MOSFET switch into saturation. At the same time, the negative excursion is 22.5 V, again endangering the gate of the MOSFET.

In summary, transformers faithfully reproduce *peak-to-peak* voltages in accordance with their transformation ratio. However, *peak voltages* are at the mercy of the duty-cycle.

Now that we are aware of and understand the ‘strange’ behaviour of transformers operating with pulse-width modulated waveforms, mention of a practical wrinkle may be in order. In many designs of switchmode power supplies, operation tends to be restricted to a relatively narrow range of duty-cycle modulation. For example, the change in pulse width is likely to be much less than the 25–75% duty-cycle depicted in Fig. 3.7(a)–(c). This being the case, it will often prove profitable to *transpose* the connections to either the primary or the secondary winding of the transformer. By so-doing, an acceptable amplitude range of the gate-drive voltage may be obtained to replace amplitude levels either too high or too low. For example, if one initially finds the positive drive voltage from the transformer secondary varies between 12 and 7 V, transposing transformer leads could project the amplitude range to 13–18 V, a possibly more workable region for many power MOSFETS.

The PWM actuated transformer just discussed pertains to the driving requirements of power MOSFETS. When such a transformer is used to drive bipolar power transistors, a somewhat different strategy comes into play. Suppose, for example, it is required to drive NPN-type power transistors. Similar watt-seconds adjustment of the secondary voltage would still take place and the NPN transistor would, like the N-type power MOSFET, be made conductive by the positive-going excursion of the secondary voltage. The simplified bipolar-transistor switching circuit shown in Fig. 3.8 is quite similar to its power-MOSFET counterpart.

In this case, however, good use is also made of the *negative* excursion of the driving voltage. Specifically, the negative-going pulse serves to ‘sweep out’

the stored charge from the base region of the transistor. This greatly speeds up the turn-off and allows the use of a higher switching rate if this is desired. Both the forward bias current and the current needed to neutralize the stored charge exceed the driving current of the power MOSFET. Therefore, the drive transformer must be somewhat larger for bipolar transistors than for power MOSFETS.

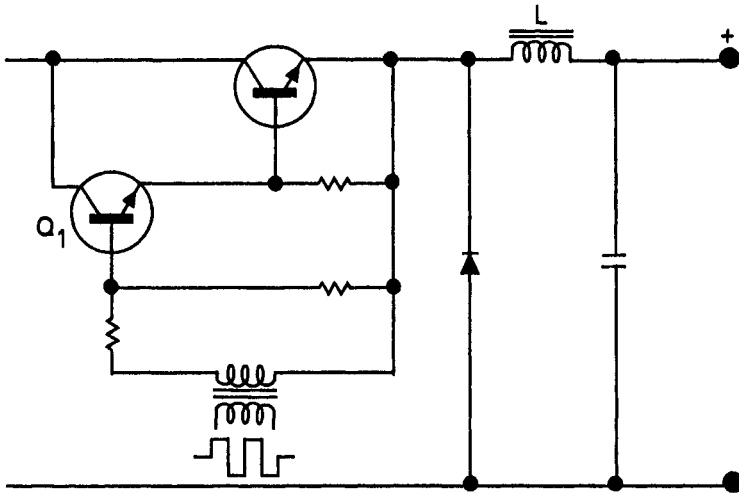


Figure 3.8 Forward and reverse-bias from a PWM drive-transformer. There is more than meets the eye in this simplified bipolar transistor switcher. After each variable duration conductive state, the stored charge in the base region of Q_1 is swept out by the negative excursions of the drive pulses. The drive transformer must be capable of supplying both forward and reverse conduction in the base-emitter junction of Q_1 .

The implementation of the bipolar transistor drive transformer is facilitated by having a *surplus* of positive voltage and a base current limiting resistance in the base circuit. A bit of empirical work will usually be necessary in order to obtain a reasonable compromise between positive and negative secondary voltages in the face of the varying duty-cycle of the PWM waveform. The bipolar transistor must be driven into saturation with an ample safety-margin. However, an increase in base current beyond this criterion introduces more charge storage and becomes counter-productive. In general, it is best to avoid zener breakdown of the base-emitter diode, although there has been some controversy over such operation. One school of thought contends that any further speed-up thus obtained has diminished life-span of the transistor as a side-effect.

Current inrush in suddenly-excited transformers

It is common to find detailed exposition of transformer theory, supported by rigorous mathematics, yet lacking in some important practical aspects of operation. An example involves the initial energization of the primary winding. Experience probably makes this a triviality – this is easily accomplished via a switch provided for the purpose, or by simply inserting the line-plug into the wall socket. Although this common-sense approach is valid for small transformers, it tends not to be for larger ones.

What we have to contend with is current inrush and statistics. With a bit of luck, the ‘natural impedance’ of the primary will be seen by the power line at the moment of energization and current demand will approximate what the transformer designer intended. This could be true whether or not the secondary is open, or is connected to its rated load.

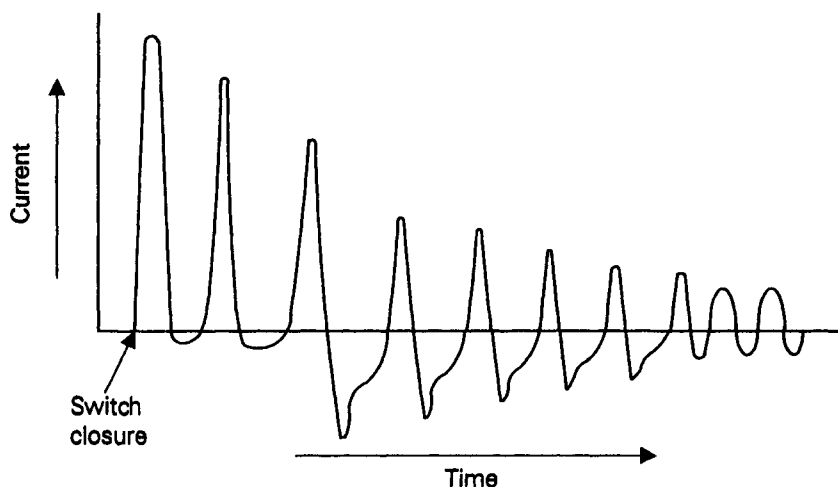


Figure 3.9 Possible inrush current when a transformer is switched on. The magnitude of the disturbance depends upon the residual magnetism in the core and on the phase of the voltage cycle at the moment of switch closure. There is only one set of circumstances for which transient phenomena can be avoided – the residual magnetism must be zero and the switch in the primary circuit must be closed at the instant of maximum voltage.

We are not likely, however, to be this fortunate because of two variables. First, we have no control of the *residual magnetization* left in the core from its most recent operation. Second, we do not know what *portion of the cyclic voltage wave* will first be impressed on the primary winding. Best bet will be such a response as shown in Fig. 3.9.

If the residual flux in the core is *zero* and line voltage is applied to the

primary at the very *crest* of the voltage cycle, the disturbance will be absent and steady-state line current will immediately ensue. Note that this appears anti-intuitive; it is easy to suppose that proper timing for avoiding transient phenomena would be energization at the zero-crossing of the voltage wave. This, however, results in the *greatest* inrush current, the actual magnitude depending upon the polarity of the residual magnetism in the core. If this residual flux is favourably polarized for one voltage zero-crossing, it will be unfavourably polarized for the next zero-crossing. So, in the presence of residual core-magnetism, the practical effect will be a violent current inrush if the switching is done at or near a voltage minimum. Even with zero residual magnetism, there can be considerable inrush current if the primary voltage is applied too far from its maximum value.

The amplitude of the initial current inrush can easily be several times full load current under steady-state operation. This is largely due to *saturation* of the core during the most adverse turn-on circumstances. It is neither practical nor economic to reduce this current transient by much over-design of core volume. Also, delaying the application of the load to the secondary winding, although sometimes helpful, fails to deal with the root of the problem. Undesirable inrush current can occur with the secondary open-circuited insofar as the primary can appear as an inordinately low impedance during the first few cycles of excitation.

In electronic systems, particularly in power supply and inverter circuitry, inrush current stems from the combined action of transformer behaviour and the sudden charge replenishment of discharged capacitors. The inrush current is undesirable because it tends to violate the *SQA* ratings of the semiconductor switches. These devices are not very forgiving; not only are they susceptible to destruction, but when they fail, they generally take out a chain of other solid-state components. Although the inrush current can be softened by electronic methods, greatest reliance for this function is usually delegated to *thermistors* inserted in the a.c. line. A typical example is shown in Fig. 3.10.

Thermistors are temperature-dependent resistors and the type shown in Fig. 3.10 exhibit a strong negative temperature coefficient. This implies a drastic reduction of resistance with a slight increase in operating temperature. When cold, the resistance is high, but the I^2R heat developed then reduces the resistance. Of course, an equilibrium must ultimately set in such that the thermistor maintains just enough resistance to support its operating temperature. In practice, this equilibrium occurs at a sufficiently-low resistance so that there is negligible adverse effect on the protected circuitry. The important point is that several seconds may be required for the thermistor to attain its low operating-resistance. Inrush current is thus avoided.

Significantly, it is of no consequence at what part of the a.c. cycle the circuitry is switched on when protected by thermistors. Also, the thermistor must necessarily operate *above* ambient temperature. The technique

becomes more difficult to apply at high power levels because the thermal time-constant of a physically large thermistor then imposes longer delay than is often desired. Where utmost efficiency is paramount, the thermistor may be shorted out by relay or contactor contacts after it has performed its inrush-softening function.

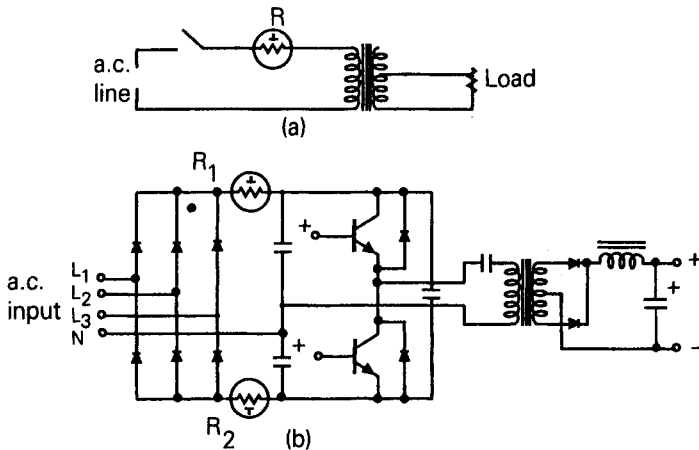


Figure 3.10 *Thermistor implementations for reducing or eliminating inrush current. The thermistor offers high resistance to current when a circuit is first turned on. Thereafter, the thermistor temperature slowly rises and its resistance decreases to the extent that the protected equipment can attain normal operation. No more than a few seconds of such delay can confer the needed protection. (a) The simplest application of thermistor protection against current inrush; (b) Protection in more complex circuit which guards against current inrush from all storage elements.*

Transformer temperature rise – a possible cooldown

It will sometimes be observed that the output transformer of a commercially-produced PWM switching regulator runs quite hot. What generally has happened is that the designer, in trying to place all components into a small space, has used rather marginal core dimensions and skimpy wire gauges for the windings. Justification for such choices can all too easily be rationalized by the notion of intermittent or non-continuous operation. The trouble is that such descriptions of operational time tend to be nebulous. It just isn't easy to pin down the thermal time-constant of a collection of devices in a nearly-sealed enclosure. Sometimes a simple expedient can hold down the temperature rise of the output transformer. For example, in the

three simplified circuits shown in Fig. 3.11, load current continuity is provided by 'free-wheeling' diodes. These diodes allow recapture of the energy stored in the inductor L during off-times of the switching process. Thus, in Fig. 3.11 (a), when the transistor switch goes into its non-conductive state, the sudden collapse of the magnetic field in the inductor induces a back-EMF which is so polarized as to *continue* the current circulating through the load; in order to do so, the current must complete its path through the free-wheeling diode.

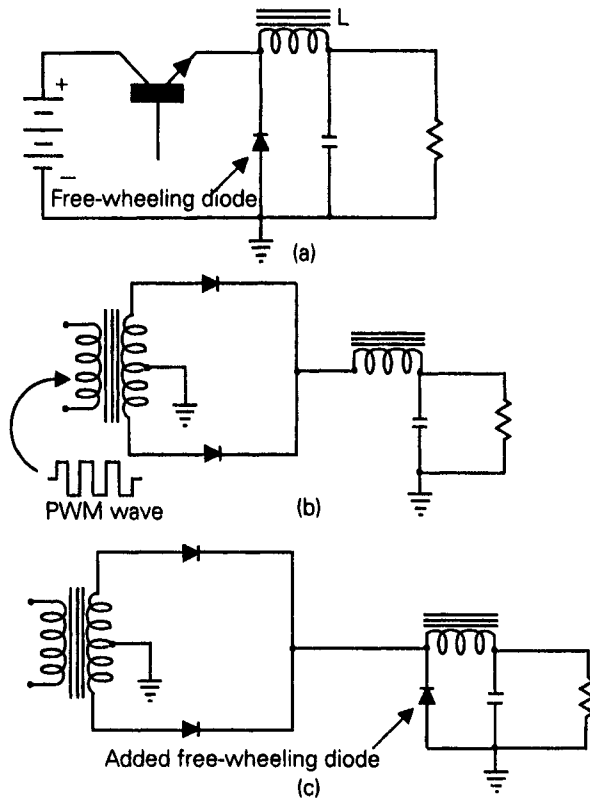


Figure 3.11 Effect of free-wheeling diodes in PWM switching circuits. (a) In the single-ended forward-converter, the free-wheeling diode continues the load current from energy stored in inductor L when the transistor switch is off. (b) The rectifying diodes in this push-pull switching circuit also function as free-wheeling diodes between positive-going pulses in the PWM waveform carried by the transformer. However, the path of the continuing load current is through the secondary winding of the transformer. This contributes to its temperature rise. (c) The addition of an actual free-wheeling diode diverts the current path so that the continuing load current no longer passes through the transformer secondary.

Interestingly, we do not find the free-wheeling diode in circuits having the topography shown in Fig. 3.11(b). This is because the rectifying diodes themselves also perform this function. It is clear, however, that the completed current path now imposes a further burden on the secondary winding of the transformer. This can be 'the straw that breaks the camel's back' from a thermal standpoint.

The transformer's current burden can be reduced by *adding* a free-wheeling diode as shown Fig. 3.11(c). A Schottky diode may be preferable for this new current path in order to discourage competition from the rectifying diodes. Fortunately, however, the resistance and inductance of the secondary winding is often enough to delay sufficiently or prevent this now unwanted current path through the rectifying diodes.

High efficiency and high power factors are fine but don't forget the utilization factor

The power-handling capability of transformers may be gleaned from their primary and secondary voltage and current ratings. However, the most straightforward parameter is the VA or kVA rating. This is the kilovolt-amperes consumed from the a.c. power line with an appropriate load resistance connected across the secondary terminals. One reason that kVA and not kW are generally used is that transformers tend to operate at slightly less than unity power factor. This is due to the small excitation current drawn and also because of inevitable leakage inductance. Despite this formality, many practitioners habitually specify transformers in terms of their approximate wattage or kilowattage ratings. Usually, no harm is done, especially if one is also provided with current ratings. However, for supplying loads containing both resistance and reactance, such as motors, kVA ratings must be used.

There are some pitfalls remaining. These generally assert themselves when the current waveshapes in either the primary or secondary circuits are non-sinusoidal. What happens is that a transformer operating under such conditions will develop an abnormal temperature rise even though the *power* drawn by the load is within the transformer's rating. Non-sinusoidal current waves do several things: they bring harmonic energy which increases both core and copper losses; they lower the power factor even though they may appear to be 'in phase' with their driving voltages; they may involve highly-peaked shapes in order to yield a moderate RMS, or effective current level. The high amplitude portions of such waves then develop inordinately high I^2R copper losses, as well as increased core losses.

The practical ramification of these matters is commonly found in rectifier circuits, where one encounters the concept of *utilization factor*. The utilization factor of a transformer used in a half-wave rectifier is very low, of the

order of 45%. It can even be lower if excessive core saturation sets in from the d.c. component in the secondary. The full-wave centre-tapped rectifier circuit also suffers from incomplete usage of its secondary windings. Its transformer secondary utilization factor can be about 0.64 providing that a sufficiently high inductance is used in the input series arm of the filter. With capacitor input, or with a single capacitor, it will be even lower.

The basic single-phase rectifier circuits are shown in Fig. 3.12. The half-wave rectifier is very hard on the transformer. Besides the interrupted and unipolar magnetization of the core, the saturation effect of the d.c. component causes the utilization factor in both windings to be very low. This is one reason why this primitive circuit is not commonly encountered, at least in sine wave systems. Another shortcoming is that the transformer utilization factor is not much improved by means of an appropriate filter choke as with single-phase, full-wave, rectifier circuits.

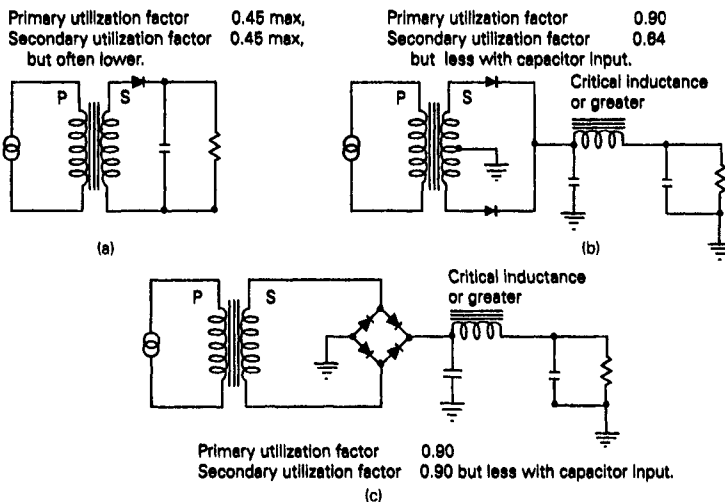


Figure 3.12 Transformer utilization factor in rectifier applications. Because of non-sinusoidal current, a transformer winding with a utilization factor of 0.5 is only able to handle about half of the current it would safely provide with sine wave operation and with a resistive load. Critical inductance in rectifier filters is the minimum inductance which will ensure non-interruptable d.c. For half-wave rectifiers, an infinite inductance would be needed. However, practical inductance values can still provide some smoothing action. (a) Half-wave; (b) full-wave centre-tap; (c) full-wave bridge.

In the helter-skelter of quick design, it is often assumed that the two full-wave circuits of Fig. 3.12 are equivalent. To be sure, there are common denominators – both provide similar full-wave rectification with nearly the

same ripple-voltage spectrum. Also, making allowances for secondary voltages and for the voltage drops in the differing rectifier-arrangements, both rectification schemes can serve a wide combination of d.c. voltage and current needs. The thing that may easily be overlooked is the better transformer utilization factor of the bridge rectifier circuit. Additionally, the construction of the transformer is not burdened with essentially *two* secondary windings as required in the centre-tap circuit. Counterbalancing these features is the fact that the bridge rectifying scheme involves *two* forward conduction voltage drops as opposed to only *one* in the diodes of the centre-tap circuit.

With both single-phase full-wave circuits, better transformer utilization can be realized when the primaries are driven by a square-wave source. This is because of the greatly reduced inactive interval between rectified pulses. A very small filter choke inductance suffices, or even none at all. Even the half-wave rectifier circuit tends to give better utilization of the transformer with square-wave excitation than with sine wave drive. Transformers used in rectifier circuits operating from a square-wave source should respond well to about 10 of the higher odd harmonics of the basic repetition-rate and should have low core losses. (The nice thing about the 50% duty-cycle square-wave is that its peak, effective, and average values are all one and the same.)

The utilization factor is often said to be the ratio of the d.c. output power from a rectifier filter circuit to the a.c. power the transformer secondary *could* deliver to a resistance load, with transformer copper losses being the same in both cases. A practical manifestation of this statement is that the transformer undergoes the same temperature rise whether delivering a.c. power to a resistance load, or causing d.c. power to be developed as the result of rectification and filtering.

Let us compare transformer utilization in the two single-phase, full-wave rectifier circuits shown in Fig. 3.13. Assume that the transformers are matched in the sense that both can deliver the same a.c. power from their secondaries feeding resistive loads. Except for the different formats in their secondary windings, the two transformers have identical cores and identical primary windings. Both transformers would develop the same temperature rise if their secondaries delivered the same a.c. power into appropriate resistive loads.

Having now set up a basis for comparison, it can be surprising to find that the bridge circuit will provide approximately 67% more d.c. output power than the centre-tap rectifier. Moreover, this is despite the higher d.c. voltage capability of the centre-tap circuit with its capacitor filter. The main reason for this large disparity is the *higher* transformer utilization factor in the bridge rectifier circuit. Bear in mind, however, that the *choke* in the bridge circuit must have a certain minimum inductance known as the *critical* inductance L_c . (In most practical situations, it is best, for a variety of reasons, to aim for $2L_c$.)

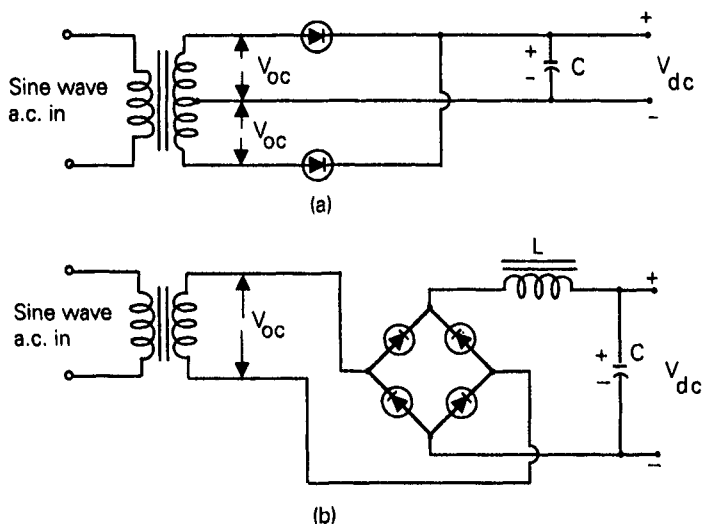


Figure 3.13 Transformer utilization factor in two rectifier circuits. (a) Full-wave, centre-tap, capacitor input; (b) Full-wave bridge, choke-input filter. The two transformers are identical in the sense that both could deliver the same a.c. power into resistors connected across their secondaries. Even though the centre-tap circuit produces a higher d.c. output voltage, the bridge rectifier with the choke input can deliver about 67% more d.c. power than the centre-tap circuit. This is due to the better transformer utilization in the bridge choke circuit.

The secondaries of this centre-tap circuit are exposed to highly-peaked, short-duration current-waves. The secondary of this bridge circuit 'sees' nearly 50% duty-cycle square-waves. Neither are ideal, but the transformer is better able to utilize the square current waves. Choke input in the centre-tap circuit would improve its performance. Capacitor input in the bridge circuit would degrade its performance. Nonetheless, the bridge circuit provides better transformer utilization *regardless* of the type of filter used. It tends to be about 40% higher for *like* filters.

A few words are appropriate regarding the *calculation* of the critical inductance that bears such relevance to transformer utilization in rectifier circuits. Slightly different versions of the equation for critical inductance are found in the technical literature. There is universal agreement, however, that this term identifies the minimum inductance sufficient to just prevent the d.c. output current from becoming zero. In other words, critical inductance is that inductance just capable of *maintaining* d.c. load current. Ideally, with critical inductance there should be no time-lapse when current switches from one rectifier or rectifier pair to the other. A further corollary is that all of the rectifiers in a full-wave circuit will conduct for 180° of the a.c. cycle.

For 60 Hz full-wave rectifier circuits, one often encounters the relationship, $L_c = R/1130$ where L_c is the critical inductance in henries and R is the load resistance in ohms. With a little algebra we obtain $1130(L_c) = R$, which is strange indeed! The 'strangeness' stems from the fact that the number 1130 just happens to be $2\pi(180)$ where 180 represents the *third harmonic* of 60 Hz. We see that the equation for L_c is telling us that a certain inductance is called for in which the inductive reactance at the third harmonic of 60 Hz is equal to (or greater than) the load resistance R . What is wrong with this? The trouble is that the mathematics is at odds with the facts of full-wave rectification – there is *not* supposed to be any third harmonic, or other odd harmonics, in full-wave rectification. (For 50 Hz operation, the critical inductance equation becomes $L_c = R/942$.)

In the technical literature, one may even encounter this relationship expressed $L_c = R/3\omega$ in which no bones are made about the 3ω being equal to $3 \times 2\pi(60)$ which, of course, pertains to the third harmonic of 60 Hz. Although such calculation of the critical inductance is simple enough, confusion is very likely unless we can reconcile the equation with the actual process of rectification. How, indeed, does the *third harmonic* come into the picture?

The mathematics of the situation are quite complicated because the terms in the Fourier series are not simple ones. However, it is easy to see that there are no odd harmonics, such as the third, in the ripple voltage output of single-phase, full-wave rectifiers. It is also obvious that the second harmonic has by far the greatest amplitude of all of the numerous even harmonics. These statements pertain to operation of the transformer from a sine-wave source.

So far one perceives no help from the mathematical outlook. However in order to absolve mathematics from being a liar, consider the following: although the second harmonic is the strongest term in the ripple voltage series, it is not by itself. Rather, the composite wave of the ripple voltages is comprised of the second harmonic *together* with all of the other even harmonics. Because of this, it would necessarily be an over-simplification to assume that only the second harmonic has to be dealt with in determination of critical inductance. In other words, it is *as if* a higher harmonic is the relevant one. Of course, the next-higher harmonic happens to be the *third*. That is, the summation of *all* even harmonics behave as a fictitious or 'phantom' third harmonic.

Admittedly, this is a *contrived* justification. For those who do not find it intellectually digestible, the involvement of the third harmonic in the calculation of critical inductance will simply have to be viewed as a *coincidence*. In any event, once this calculation is made, it is usually best to double it in actual practice. This will compensate for various omitted factors such as the resistances of the transformer secondary, the choke itself, the rectifiers etc. It will also allow wider variation in the load. Additionally, greater inductance generally provides improved filtering.

For those endowed with a mathematical bent, an all-inclusive formula has been derived which enables calculation of critical inductance for all half-wave and full-wave rectification systems operated from either single or polyphase sinusoidal voltage sources. Moreover, it can be seen that this universal formula embraces the equation already given for single-phase, full-wave rectifier circuits. The formula is as follows:

$$L_c = \frac{2R}{p(p^2 - 1)\omega_s}$$

where L_c is the critical inductance in henries.

R is the d.c. load resistance in ohms.

p is the number of pulses per rectification cycle.

ω_s is the radian velocity of the supply frequency, that is $2\pi f_s$.

Putting this interesting formula to the test, it is quickly seen that the critical inductance for a single-phase, half-wave rectifier is infinite. And, it is intellectually rewarding to find that the formula 'degenerates' into $R/3\omega_s$ for single-phase, full-wave rectifiers in agreement with our equation. This is readily confirmed by substituting the number 2 for p . (With the half-wave rectifier, $p = 1$.)

A three-phase, full-wave rectifier has excellent transformer utilization factors, 0.955 for both primary and secondary. As such a circuit develops six pulses per cycle, a minimal inductance suffices for the series filter choke. Moreover, the relatively high ripple frequency is easy to filter for production of nearly 'pure' d.c. The basic circuit is shown in Fig. 3.14.

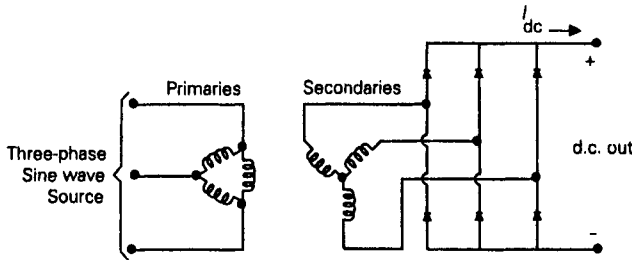


Figure 3.14 Basic three-phase, full-wave rectifier circuit. With minimal filtering, this rectification scheme provides nearly-ideal utilization of the transformers. Long used with tube rectifiers, there was a cost and weight penalty because of the required filament transformers. With solid-state rectifying diodes, these disadvantages have been removed. (Note that the secondary circuit is the same as is used in automobile alternators.) Three single-phase transformers, or one three-phase transformer can be used.

It should be pointed out that the concept of 'critical inductance' loses some of its meaning in polyphase rectifiers because the unfiltered output never goes to zero (see Fig. 3.15). However, designers find that a small choke nevertheless benefits filtering, regulation and tames inrush current. Also, because of non-linearities in the rectifying diodes, such an input filter choke tends to improve further the already good power factor and transformer utilization factor.

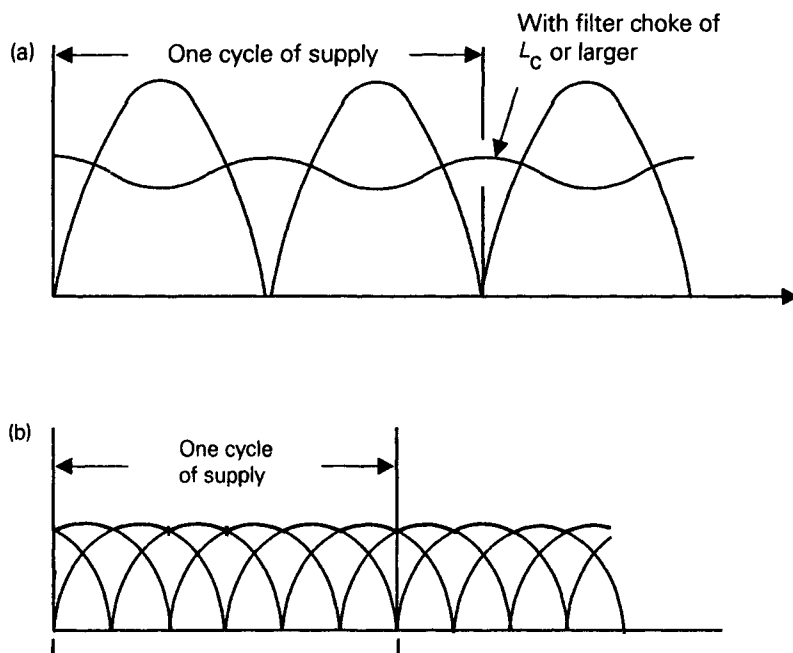


Figure 3.15 Output current in single and three-phase full-wave rectifier circuits. (a) Single-phase centre-tapped and bridge circuits; (b) three-phase, full-wave circuit. Note that the 'natural' output of the single-phase circuits go to zero. The three-phase output never zeros, even without a so-called 'critical' inductance.

These discussions naturally get us back to transformer ratings. As mentioned, VA and kVA are used rather than the power transfer capability in W or kW. This falls into place nicely because the combined effects of low power factor and low utilization factor can produce high currents in the windings despite low power transfer to the load. On the other hand, the VA rating system ensures safe operation for any power factor from zero to unity. Of course, the devil's advocate can cook up hypothetical situations involving drastic changes in applied voltage and operating frequency in which the

transformer could be over-heated or damaged while subjected to its kVA rating. Nonetheless, the kVA rating reliably establishes operating constraints in situations most generally found in practical systems.

Transformers in three-phase formats – all's well that's phased well

Paralleling the outputs from three-phase transformer networks can present interesting dilemmas. At first, it would appear that one only need pay heed to the direction of phase rotation and voltage. Unlike the situation with rotating machinery, there should be no need for synchronism – after all, the branches to be paralleled derive their currents from the same basic source. The *frequencies*, therefore, are the same.

It happens, however, that 'the same basic source' may have been subjected to transformations involving *different combinations* of delta and wye formats in the branches hopefully to be paralleled. This can result in a 30° phase difference between the branches to be paralleled. In such situations, there is no easy way to attain the condition of phase-opposition required for paralleling. What this boils down to in practice is that a bank of delta-wye or wye-delta connected transformers *cannot* be paralleled with a wye-wye or delta-delta bank.

A straightforward phase-correction technique would be the appropriate insertion of a wye-delta or delta-wye transformer bank in to one of the branches. The primary and secondary windings would have to be proportioned to yield an overall unity transformation ratio. However, this would be a costly remedy and would probably be considered self-defeating by the utility industry. It might be worthy of consideration in an electronic system.

Another phasing situation commonly encountered in three-phase circuits involves the proper connection of the third secondary winding in a *delta* format. The first two secondaries can be connected without regard to which terminals tie together. However, definite phasing observance applies to completion of the delta by the *third* secondary. The *wrong* connection will result in a heavy short-circuit. Because of the mainly third harmonic magnetizing current, a *small* circulating current can be detected for the right connection, but should not be grounds for concern. (The third harmonic voltages, unlike the fundamentals, do not cancel in the delta network.)

Recapturing the lost function of the power transformer

At one time it was unthinkable to market a consumer product such as a radio receiver or a television set without a so-called 'power transformer'. Although this component conveniently provided the various operating

voltages, its more important function was to conductively isolate the apparatus from the utility power line. This was construed not only to be protection against electric shock but was felt to confer protection from fire as well. Times have, alas, changed! The recognition that costs could be reduced via elimination of the power transformer proved overwhelming. The increased hazard to the public was explained away with such arguments that one wasn't supposed to operate a radio while immersed in the bath-tub anyway.

Servicing these transformerless products poses a real problem. One must not only avoid shock from a 'hot' chassis, but it is only too easy to cause a nasty short-circuit because of a ground conflict with test and measuring instruments. The straightforward remedy to these hazards is to insert a 1:1 isolation transformer between the a.c. line and the equipment being serviced. For some reason these isolation transformers seems to be unavailable when most needed. Fortunately, a simple expedient can often be resorted to when such is the case.

The two-transformer back-to-back arrangement shown in Fig. 3.16 can provide the proper operating conditions together with the all-important *isolation*. The basic idea is to find two identical step-down transformers and

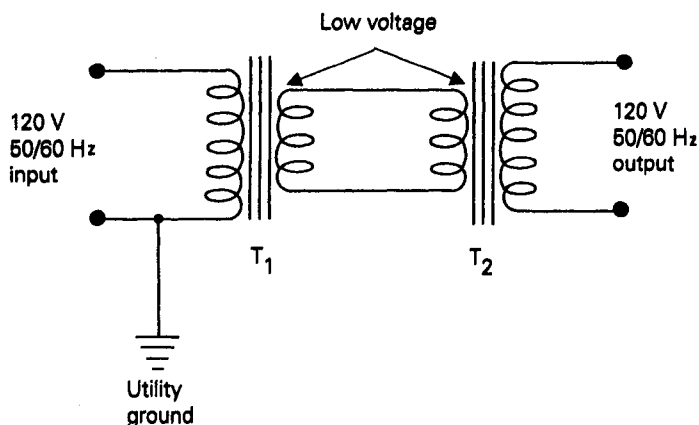


Figure 3.16 Isolation technique via pair of identical step-down transformers. Original purpose of the two transformers was to step-down the 120 V 50/60 Hz a.c. line to a relatively low voltage. With the above connections, line voltage safely isolated from ground is available from transformer T_2 . Phasing of the low voltage windings is not important for power transfer, but one connection may prove favourable in coupling in less noise from the line.

connect their secondary windings together. Conductively isolated line voltage will then be available from what was normally the primary winding of one of these transformers. The pragmatic philosophy here is that one is, indeed, likely to dig up a pair of such transformers in the service shop or

the engineering laboratory. There is nothing sacred about the *equal* secondary voltages – one can think of bell transformers, filament transformers or transformers originally intended for solid-state power supplies. Of course, the kVA ratings should be compatible with their new use. One can take some liberties here because many test and measurement procedures do not require much time.

Bandpass response from resonated transformers

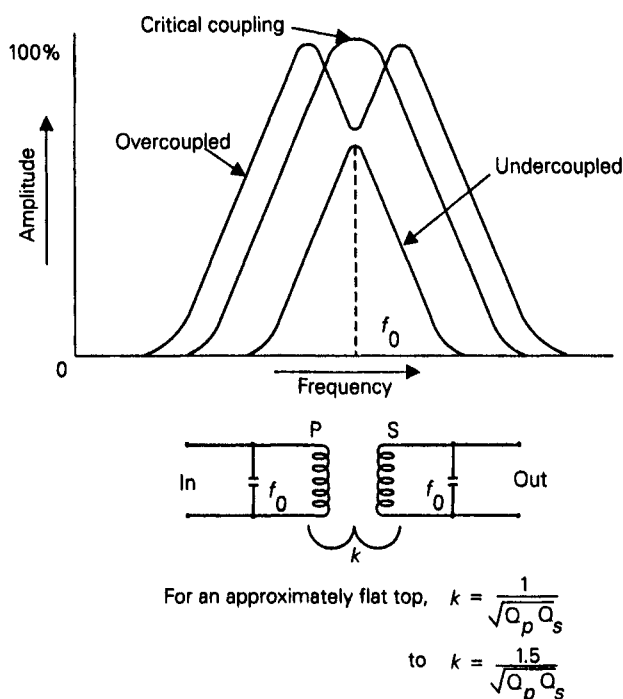


Figure 3.17 Frequency response of the doubly-tuned transformer. Near-identical primary and secondary windings are used and both are resonated at the same mid-band frequency. The windings are loosely coupled at, or near, the so-called 'critical' coefficient of coupling which corresponds to a flat top in the frequency-response curve. Driving and load impedances should be much higher than the characteristic impedances of the resonated circuits.

Doubly-tuned transformers have had a long history of applications in radio, television and other communications systems. The usual objective has been to obtain a band-pass characteristic with reasonably flat response over the pass-region. The most commonly encountered use has been in the intermediate frequency (IF) amplifier of superheterodyne circuits. Although

various other bandpass devices, such as those relying upon crystal, ceramic or mechanical resonators have been able to provide superior skirt-selectivity and deeper rejection in the stop-band, the doubly-tuned IF transformer has enjoyed a well-deserved reputation as a cost-effective workhorse in numerous practical situations. And where narrower and steeper responses have been needed, it has proven easy enough to cascade two or even three such transformer-coupled amplifier stages.

Unlike the need to use tight coupling between primary and secondary in *power* work, the tuned windings of the IF transformer are loosely-coupled. Air core design has commonly been used, although certain performance enhancements can be realized with the proper use of magnetic material. One of the most popular of the IF passbands centres at 455 kHz. Proper performance hinges very much on the 'critical' coefficient of coupling k . A flat passband is first attained with k ; in practice, it is found that substantially flat passbands prevail over the coupling region of k to $1.5k$. The detrimental effect of looser or tighter coupling is shown in Fig. 3.17. Note that k is defined as 'critical' for the condition of:

$$\frac{1}{\sqrt{Q_p Q_s}}$$

The coefficient of coupling is established by two factors, the operating Q values of the resonant circuits and the geometric spacing of the two windings. Experimentation is conveniently brought about by changing the physical separation between the coils. The characteristic impedance Z_0 of both primary and secondary is given by $\sqrt{L/C}$, where L and C are the resonating values. Both the generator and the amplifier input impedances should be at least 10 times Z_0 in order not to substantially degrade the Q values.

4 Interesting applications of transformers

It comes as no surprise to encounter interesting applications of 'out of the ordinary' transformers. It also happens, however, that there are interesting applications of transformers described as mundane or 'garden variety'. With this in mind, some intriguing implementations of ordinary transformers will be described in this chapter. Along the way, we will discuss one or two interesting uses of perhaps not so ordinary transformers as well.

Remote-controlling big power with a small transformer

An interesting use of an ordinary transformer is as an isolation device to facilitate remote control. Admittedly, in this regard, the transformer can be perceived as being low-technology, one that most often is the task of more sophisticated devices, such as opto-isolators, or one of the several three-terminal active semiconductor components. Generally, there is no argument here, but in some instances the garden variety transformer will be found to possess some advantageous features.

Consider the power control circuit shown in Fig. 4.1. Here the gate of a Triac cannot receive sufficient voltage for triggering the Triac unless the primary of the transformer is short-circuited by a switch. With the switch open, the inductance of the secondary is high. Due to the potentiometric action of this inductance in association with resistance R_1 , only a sub-trigger voltage is available for the gate of the Triac. However, when the switch in the primary circuit is closed, this short-circuit is reflected in the secondary, causing the secondary inductance to drop almost to zero. The gate trigger level is then easily exceeded and the Triac will fire on both halves of the a.c. cycle.

For many applications of this kind, a small filament or bell-ringing transformer can readily be adapted for a wide variety of Triacs by experimenting with the value of R_1 . Remote control is achieved via an appropriate length

of twisted wire or coaxial line. Even though this may preclude attainment of a true short-circuit, the change in secondary inductance will nevertheless be quite large. (Residual secondary inductance can cause some phase shift, but generally not enough adversely to affect load power. In any event, phase can be 'straightened out' with a large capacitor across R_1 .)

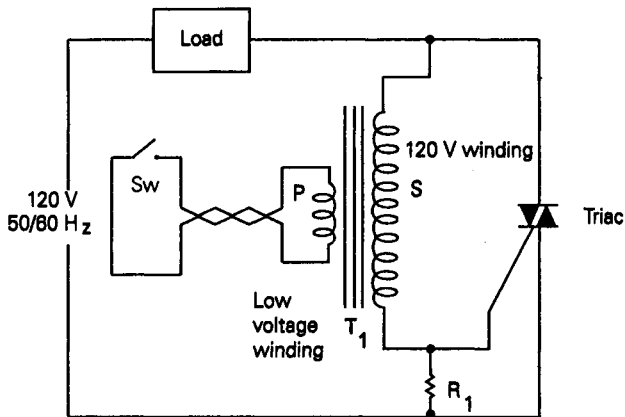


Figure 4.1 Use of a transformer as a passive remote control device. Closure of switch *Sw* nearly reduces the inductive reactance of the transformer secondary to zero. The potentiometric voltage division of the secondary resistance and R_1 then rises to the triggering level of the Triac gate. By experimenting with R_1 , a long length of twisted cord can be accommodated for remote actuation. Grounding in the primary circuit can help eliminate RFI in some cases. T_1 can be a small filament or bell-ringing transformer.

One of the features of such implementation is the *reliable* isolation provided by the transformer. Another is the high *immunity to abusive operation*. Finally, the *low impedance* of the primary circuit is unlikely to be vulnerable to electromagnetic disturbances. A tiny transformer suffices to control a large Triac because of the relatively low gate current involved. It is conceivable, therefore, that *low cost* may also influence selection of this transformer function.

The use of a transformer to magnify capacitance

Electronics practitioners versed in both audio and radio frequency techniques are familiar with impedance-matching applications of transformers. Generally, the objective is to convert a resistance to a higher value in a step-up transformer or to convert a resistance to a lower value in a step-down transformer. In the first case, the converted value of resistance will

be greater than that associated with the primary by the square of the secondary to primary turns ratio. In the second case the converted resistance will be less than that associated with the primary by the ratio of the secondary to primary turns, squared.

More generally, impedances are also convertible in this manner with transformers. Moreover, the conversion can be accomplished with auto-transformers as well as with conventional two winding transformers. A familiar impedance transformation takes place in the output transformer used to match the output impedance of an audio output stage to the inordinately low voice-coil impedance of a dynamic speaker via a step-down transformer.

An interesting application of impedance transformation is encountered with single-phase a.c. induction motors in which a large capacitor is needed to split the phase of the applied line voltage so that the motor starts essentially as a two-phase motor. Unless something of this nature is done, the motor will not develop starting torque. Once under acceleration, the capacitor is cut out of the circuit by a centrifugal switch. A practical problem arises because of the large capacitance required; because of cost and physical size, such a starting capacitor usually has to be an electrolytic type. These can be marginally satisfactory, but tend to have adverse ageing characteristics, high leakage current, sloppy capacitance values, limited temperature tolerance and other manifestations of unreliability. Avoiding the need for the electrolytic starting capacitor is a worthwhile design objective.

It turns out that oil-impregnated paper capacitors as well as capacitors with various plastic film dielectrics have excellent electrical characteristics, are readily available with high voltage ratings and are capable of stable and long-lived operation, even at fairly elevated temperatures. However, such capacitors would be impractically massive in physical size for the requisite capacitance needed for motor starting. This is where the *transformer* comes to the rescue. With the appropriate turns ratio, a step-down transformer can make a relatively low capacitance oil or plastic type capacitor appear to the motor as a very high capacitance. Auto-transformers are generally used for this purpose. Interestingly, many tens of amperes of near-quadrature current can be applied to the starting winding of the motor; at the same time, the oil or film type capacitor can withstand many hundreds of volts during starting. (It 'sees' the transformer as a step-up type with respect to the a.c. line voltage.)

A typical arrangement of the kind described is shown in Fig. 4.2. One can consider the motor a two-phase type inasmuch as the starting winding is displaced 90° on the stator from the running winding. The essential purpose of a starting capacitor (whether 'physical' or impedance transformed) is to make the single-phase supply appear as a two-phase source to the motor, for only then will it develop starting torque. Note that the high starting current of the motor is mainly carried by the low voltage turns of the

auto-transformer. This is also a great advantage of this technique inasmuch as high current in an electrolytic capacitor quickly produces a rise in temperature and is detrimental to longevity. In contrast, the high-voltage applied to the oil or film capacitor has little adverse effect and it is not called upon to pass high current.

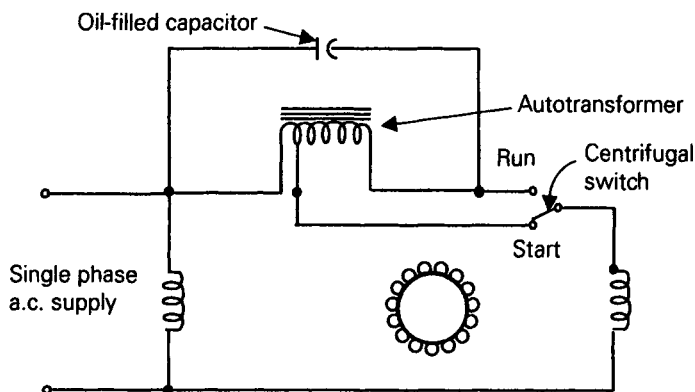


Figure 4.2 *Magnification of capacitance by transformer action. With a special starting arrangement, a two-phase induction motor can be operated from a single-phase a.c. line. What is needed is a momentary association of a large capacitor with one of the motor windings. This is a phase-splitting technique which creates a virtual two-phase a.c. source. The electrolytic capacitors used in large capacitance values are often unreliable and become a maintenance problem. In the above scheme, the autotransformer 'steps-up' the small capacitance of the reliable oil-type starting capacitor.*

A typical example of such transformation of capacitance might be an auto-transformer with, say, 140 total turns and a tap at 20 turns. Assume such an auto-transformer is associated with an $8\ \mu\text{F}$ oil-type capacitor as in Fig. 4.2. The capacitance 'seen' by the motor when the centrifugal switch is in its start position is

$$(140/20)^2 \times 8 = 392\ \mu\text{F}.$$

That is, the motor behaves just as it would if a $392\ \mu\text{F}$ 'physical' capacitor were used during its starting phase. As pointed out, such a large capacitance would have to be obtained from an electrolytic capacitor with its several shortcomings. The $8\ \mu\text{F}$ oil-type capacitor has to be rated for high voltage, specifically to withstand $140/20$ times the line voltage. Assuming a 120 V line, this would be 840 V. A 1000 or 1200 V capacitor could be selected with very little cost penalty.

To avoid the confusion which sometimes arises from the unfamiliarity with this technique, it should be noted that the parameter which really undergoes transformation is *capacitive reactance*. Keeping this in mind, capacity bears a reciprocal relationship to its reactance, i.e. the higher the capacity, the lower the capacitive reactance and vice versa.

Although the implementation of this technique is quite straightforward, trouble can be encountered with skimpy design of the auto-transformer. If there is excessive leakage inductance, or if saturation sets in, either a smaller than predicted transformation of capacitance will occur, or a resistive component will appear. The net result will be a *reduced* starting-torque developed by the motor. Otherwise, the auto-transformer can easily be the most reliable item in the system.

A novel transformer application in d.c. regulated supplies

Figure 4.3 shows a *phase-modulation* regulated power supply. This type of d.c. supply relies on duty-cycle variation to accomplish voltage regulation. However, the principle of operation is different from the commonly encountered pulse-width modulation technique. In particular, the phase-modulation supply incorporates a unique application of a transformer.

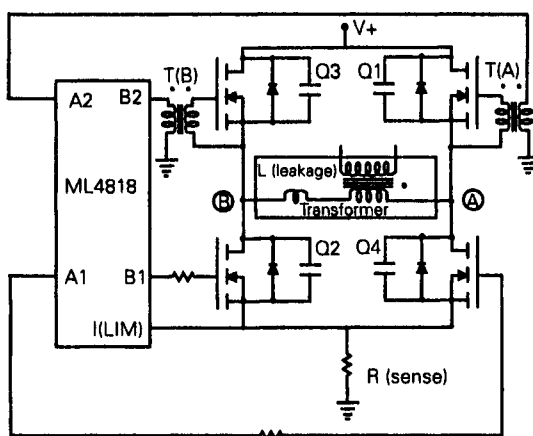


Figure 4.3 Partial circuit of regulated supply utilizing transformer leakage inductance. In this arrangement, the normally tolerated leakage inductance of a transformer is deliberately used in the operating scheme. High efficiency results because the power MOSFETs undergo their switching transitions at zero drain emitter voltage.

It should be pointed out that the four power MOSFETs do not constitute a bridge rectifier; rather their configuration is that of an H-bridge. Whereas the bridge rectifier converts a.c. to d.c., the H-bridge converts d.c. to a.c. The electromechanical analogue of the electronic H-bridge is shown in Fig. 4.4. It can be seen that a single-poled d.c. source can be made into symmetrical a.c. for operation of the transformer primary. It can be appreciated that appropriate switching logic can be implemented for various control techniques. The logic and overall operation of the regulated power supply will not be discussed; rather, we will focus on the way the *transformer* is used.

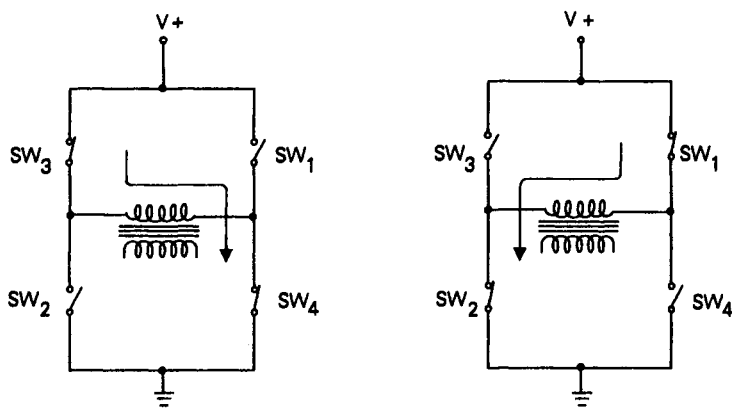


Figure 4.4 Mechanical analogue of the H-bridge. Whereas the bridge rectifier converts a.c. to d.c., the H-bridge converts d.c. to a.c. In this way, a transformer can be operated from a single d.c. power source. Efficient control can be achieved with a logic-controlled H-bridge.

The transformer, it should be noted, is not tuned as in so-called resonant supplies. In a sense, resonance phenomena exist, but neither the primary nor secondary of the transformer is made oscillatory by either physical or stray capacitance. Rather, there is an ongoing energy exchange between the output capacitance of the power MOSFETs and the *leakage inductance* of the transformer. Sustained oscillation is prevented by the switching action and regulation is achieved by varying the delay of the switching process. It would be only natural to ponder how a transformer designer could reliably target a desired leakage inductance. The answer is that the actual leakage inductance is not critical – a wide range of values can be accommodated. A resistance in the logic circuit can be conveniently adjusted to optimize operation for a wide range of leakage inductance. To put the situation another way, it would be most difficult to make a practical transformer with *no* leakage inductance.

Polyphase conversions with transformers

It is a common fallacy that very specialized transformers would be needed to change the number of phases in a power system. Usually it is thought that some kind of non-linearity must be involved in order to produce harmonics and sub-harmonics which can be combined in new ways. This is not so; it turns out that very ordinary linear transformers can be connected to convert efficiently from one polyphase format to another. The single exception is that such transformers *alone* cannot convert a single-phase system to any polyphase system. (We are discussing the transfer of large blocks of energy, not the phase-shifting manipulations resulting from RC circuits for low-level uses in electronic systems.)

The Scott connection makes use of two more or less conventional transformers which enable efficient conversion from a three-phase line to two-phase output. It should go without saying that the conversion can be *reversed* in which case one would undergo change from an incoming two-phase power line to a three-phase format.

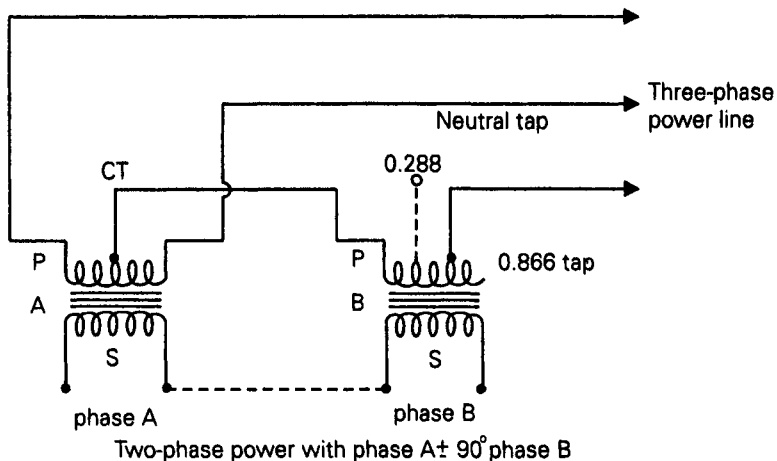


Figure 4.5 Scott connection for three-to-two or two-to-three-phase transformation. The only difference in the two similar transformers is the position of the taps. A three wire, two-phase format can be had by connecting the secondary windings in series. The manner in which this is done then determines the quadrature lead or lag situation in the converted two-phase system. The unused tap can provide an exact neutral for the three-phase side. (The CT tap on the main transformer cannot be used as a neutral point – its use for this purpose would unbalance the three-phase voltages.)

The basic Scott connection is shown in Fig. 4.5. Except for the taps, assume the two transformers to be identical. For simplicity, consider that

both transformers would also show a one-to-one turns ratio if used in a simple single-phase circuit with no connections made to any of the taps. Transformer A is known as the *main transformer* and transformer B is referred to as the *teaser transformer* in the technical literature. However, both play equally important roles in actual operation.

Note the unused tap on the main transformer. Its use is as follows. Assuming the Scott connection is being used to transform from a two-phase power line to a three-phase output format, then this previously unused tap will provide an exact neutral, yielding a balanced three-phase, four-wire system. The way of locating this tap is to state that it is at one-third of the turns occupied by the 0.866 tapped winding. This establishes the neutral tap at 0.288 of the total turns from the end of the winding which connects to the centre-tap of the main transformer. In many practical situations, one can dispense with the neutral connection.

The hybrid coil – a transformer gimmick for two-way telephony

An interesting transformer application, long used in telephone technology, was the so-called hybrid coil. In its simplest form, this was a unity-ratio transformer with centre-tapped primary and secondary. The laboratory set-up shown in Fig. 4.6 is useful in demonstrating the principle involved. Assume a signal is impressed at terminals *ab*. It turns out that no induced signal will be present across impedance Z_3 if impedances Z_1 and Z_2 are identical. Under this condition, the overall arrangement is essentially a balanced bridge. Let an amplifier be associated with the transformer in such a way that the input of the amplifier samples the voltage (if any) across Z_3 ; and let the output of the amplifier be developed across terminals *ab*.

From what has been postulated so far, a significant fact emerges. The amplifier will not 'sing', or self-oscillate due to the bridge-balanced isolation between its output and input. Let us now see what other features are manifested.

Dispensing with the a.c. signal source used for our demonstration of bridge balance, assume a signal to be introduced across impedance Z_1 . Tracing the path of this to Z_3 signal, it is amplified and then via ordinary transformer coupling it is passed to impedance Z_2 . In our mind's eye, we can see a weakened telephone signal appearing across Z_3 and ultimately arriving at the receiver Z_2 , boosted in power level. Interestingly, Z_2 can be construed to be the transmitter, whereupon the amplified signal arrives at 'receiver' Z_1 . We can see that the scheme works in *both* directions. Of course, in the real world the attenuation of the telephone signals is primarily caused by the long transmission lines. Thus, 'repeaters' comprising a hybrid coil and an amplifier are required at appropriately-spaced intervals. (Also,

telephone systems generally used hybrid coils at Z_1 and Z_2 as well as at Z_3 . Many versions of the arrangement shown in Fig. 4.6 are possible without altering the basic operating principle. Either centre-tapped or separate windings can be used. Depending upon frequency, either iron or air core transformers can serve the purpose. For the sake of symmetry and balance, two secondaries were often used, one inserted in each line. In all cases, however, Z_1 must equal Z_2 to attain nulling at Z_3 .

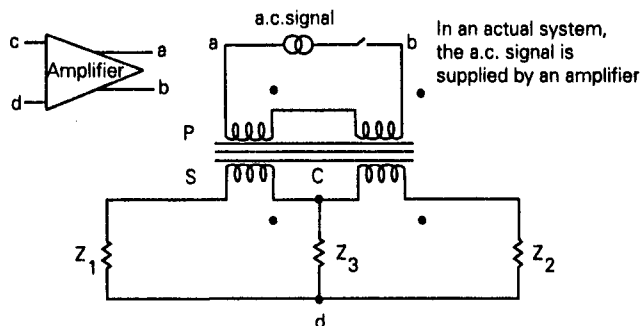


Figure 4.6 Laboratory set-up for demonstrating bridge behaviour of a hybrid coil. Despite its common name, the hybrid coil is actually a transformer. The objective in telephone practice is to associate an amplifier with the circuit so that a weak signal from Z_1 can be boosted in level and delivered to Z_2 (or the converse direction, from Z_2 to Z_1). Without the hybrid coil, the amplifier could not be used – it would ‘sing’ because of lack of isolation between its input and output. In the above scheme the input of the amplifier would sample signals across Z_3 ; the output of the amplifier would take the place of the a.c. signal generator.

Transformer schemes for practical benefits

The practical ‘tricks’ about to be described stem from transformer techniques long used by the utility industry. However, similar schemes might suggest themselves for useful applications in electronic systems. In Fig. 4.7(a), we see a delta–delta connection of transformers in a three-phase network. Assuming that the phase windings in the transformers are all the same, the overall voltage transformation ratio is unity – the output voltages are the same as the input voltages.

Utilizing the *same* transformers, let us now deal with a delta–wye arrangement as shown in Fig. 4.7(b). The output voltages will now be $\sqrt{3}$ or 1.73 times the input voltages. Thus, we have obtained an appreciable step-up of voltage without altering the phase windings on any transformer. At first glance, this is hardly a newsworthy achievement. It has, however a very important feature in high voltage engineering. Suppose the delta–wye

configuration is employed to step a three-phase 58 kV line up to 100 kVs by taking advantage of the 1.73 multiplying factor of the delta-*Wye* connection. A little contemplation reveals that the phase windings of the *Wye*-connected secondaries only have to be insulated for 58 kV, *not* for 100 kV.

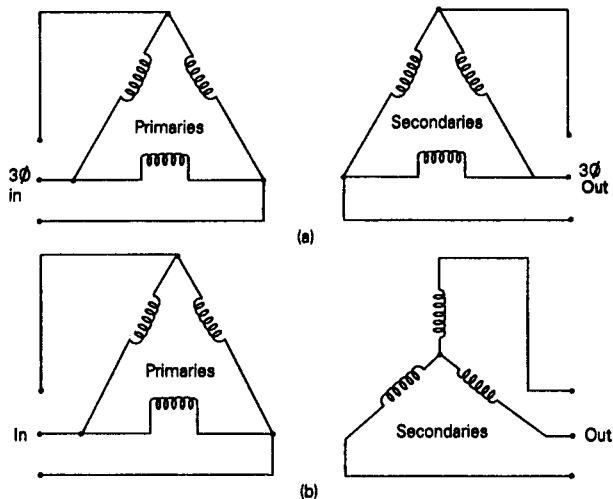


Figure 4.7 Use of *Y*-connected secondaries for voltage step-up. (a) Unity voltage transformation in delta-delta format. It is assumed that all phase windings have the same turns. (b) 1.73 voltage step-up in delta-*Wye* format. The salient feature of this arrangement is that the secondary windings only have to be insulated for 58% of the output voltages.

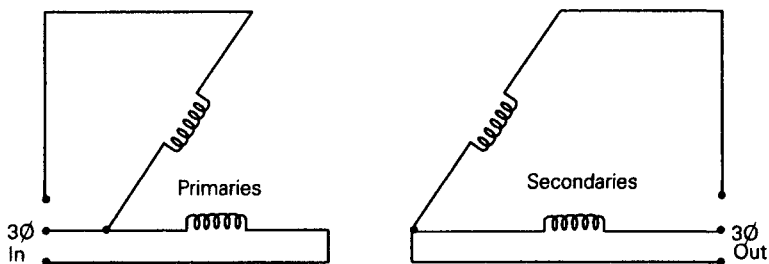


Figure 4.8 Open-delta or 'V' three-phase transformer format. Except for the omission of one transformer, this arrangement remains similar to the delta-delta format of Fig. 4.7(a). Three-phase operational requirements continue to be met but at reduced power-handling capability. In systems where the open-delta is adequate, it becomes convenient to add a third transformer for increased load demand at a later time. Conversely, a delta-delta system can continue to provide emergency service by removing a defective transformer.

Not to be outdone, delta–delta connected transformers also have an interesting application. Suppose that one of the three transformers were *omitted*. What would be left is known as an open-delta, or ‘V’ arrangement. This is shown in Fig. 4.8. This connection scheme continues to fulfil the requirements of three-phase transformation. As would be expected, the power-handling capacity is reduced. Actually, it is reduced *more* than the anticipated two-thirds of the three-transformer delta–delta format. That is because the two-transformer format operates at a power factor of 86.6% even with resistive loads. The net result is that only 58% of the three-transformer capability is available. There will be a cost saving if this is acceptable.

Suppose now that an open-delta system is initially installed. At a *later* time, when increased load demand dictates addition of a third transformer for delta–delta operation; a 50% incremental investment then provides a 73% increase in power rating.

Transformers in magnetic core memory systems

Computer memory systems have had an extensive history of utilizing special arrays of saturable core toroidal transformers acting as bistable logic elements. Although these have been largely displaced by semiconductor ‘chips’, many useful applications remain where speed and cost assume less importance than non-volatile operation, that is, the retentivity of stored information despite shut-off of power. Other features, too, can be readily provided by such so-called core memories. They can yield reliable performance in fairly hostile environments and are relatively immune to damage or upset from transients, especially the kind commonly delivered by power lines.

Fig. 4.9 shows such an array of core memory devices in its simplest form. Each element is actually a three-winding isolation transformer in which each one turn winding consists of a conductor passing through the aperture. The ferrite-core material is generally selected primarily for its near-ideal rectangular hysteresis loop. The basic idea in operation is that a given core can be switched to one or the other of its flux remanence states by the *mutual* currents at the intersection of a row and column conductor. *Other* cores, although also exposed to a current pulse, will not undergo magnetic switching because only *half* of the requisite magnetomotive force is experienced by such non-selected cores.

Memory readout is accomplished by probing the array with current pulses injected at the readout terminals. Only an appropriately magnetized core will thereby produce the identifying counter-EMF. Thus, the *address* of a selected logic bit is obtained. However, the readout process is destructive in the sense that it wipes out the stored information. To prevent this, a rewrite current pulse follows the readout pulse in order to re-establish the memory status of the array.

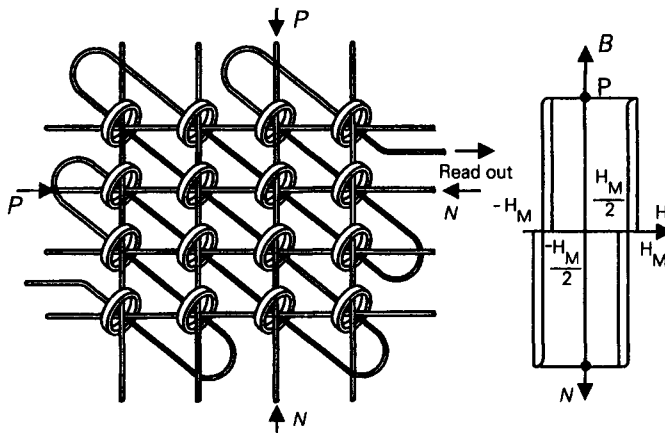


Figure 4.9 Basic core logic memory system. These are arrays of three-winding saturable core transformers. Each winding comprises the through-passage of a conductor and is the equivalent of one turn. Logic address is to columns and rows, so that any single core can be selectively addressed. The two saturable states of the cores represent logic zero and one. At the readout terminal, the logic state of any core is indicated by the polarity of the counter-EMF in response to a current pulse.

In order to speed up switching time and to reduce physical volume, much emphasis was once placed on core memory technology. Complexity and sophistication well beyond the basic arrangement of Fig. 4.9 became the order of the day.

The transmission line transformer – a different breed

Those involved with the technology of handling radio frequency power have attained a respectable amount of progress in adapting conventional transformers to the job. However, they work against great difficulties in attempts to optimize more than any single performance parameter. The simple reason is that irreconcilable design conflicts inevitably assert themselves. One might like to obtain broadband response, minimal leakage inductance, low losses and high efficiency, linear operation, high power capability, easy reproducibility and low cost all in the same device. Via a unique blend of science and art, an acceptable compromise can often be made. However, a better balance between such desirable features is likely to come in a *different* device, the so-called *transmission line transformer*.

To the first-time viewer, the transmission line transformer bears close resemblance to the conventional transformer (see Fig. 4.10). Indeed, a

bifilar-wound transformer might require close scrutiny to determine whether it is a conventional transformer dependent upon mutual flux linkage of its primary and secondary windings, or a transmission line transformer which operates on a different principle. A couple of examples of the operational divergence of the two devices quickly reveals that they are of different species; the ferromagnetic core of the conventional transformer is intimately involved in energy transfer, but this is *not* the case with the transmission line transformer. (Although a school of thought claims this may not be entirely true at the low frequency end of the response characteristic, it is nonetheless acknowledged that core losses and core non-linearity remain much less of a problem than in a conventional transformer of equal power-rating.)

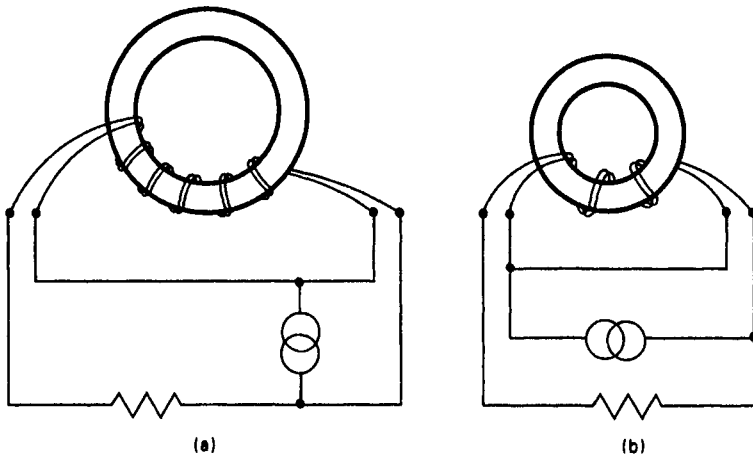


Figure 4.10 Conventional and transmission line transformers can bear a close resemblance. With similar ratings, the transmission line transformer will tend to have fewer turns and a smaller core than the conventional transformer. (a) Conventional transformer connected for 2:1 voltage step-up. This corresponds to 4:1 impedance transformation. (b) Transmission line transformer connected to provide a 4:1 transformation of impedance.

Whereas distributed capacitance associated with the windings greatly limits high frequency response in conventional transformers, most of such capacitance is constructively absorbed in the characteristic impedance of the transmission line transformer.

It might sound like a riddle to ask what a *travelling wave tube (TWT)*, a *Beverage wave antenna*, and a *transmission line transformer* have in common. Suffice it to say that expertise in the theory and applications of one of these implementations could provide the intellectual springboard for inventing the other two. More to the point, the shared operational feature is *energy transfer via an excited transmission line*. The excitation can assume different

forms; interestingly, however, it need not occur by flux linked coils as in a conventional transformer. Excitation introduced at one end of a transmission line propagates energy to the far end via travelling wave action.

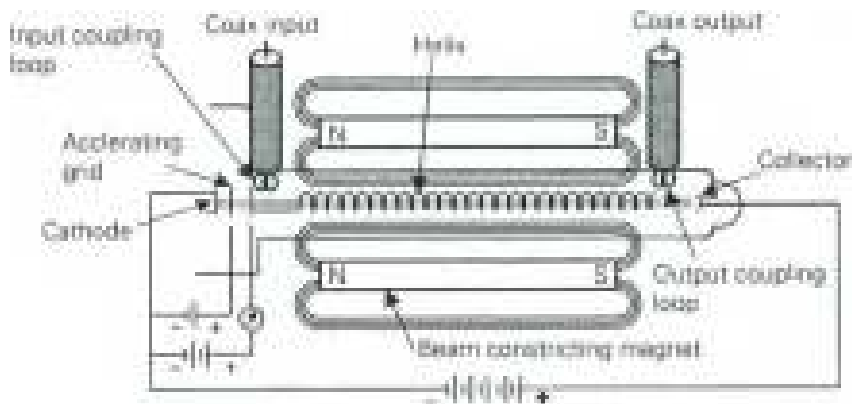


Figure 4.11 The travelling wave tube. The helix is actually a coiled single conductor transmission line. As such, a signal impressed at its input is guided to its output in the form of travelling waves. Along the way, interaction with the electron beam imparts energy so that the output signal is at a higher power level than the input signal. (At UHF and microwave frequencies single wire lines require no image line or 'return' path.)

The TWT (Fig. 4.11) represents the most sophisticated utilization of the principle. In this device, the helix is in reality a coiled transmission line. Travelling wave energy transport in the helix is slowed so it can take energy from the electron beam. This transmission line helix is excited by, and continually supplied with, energy from the electron beam. The energy extracted from the electron beam not only gives rise to a *travelling wave* in the helix, but progressively increases its amplitude. Hence, this device is by its nature an *amplifier*. A weak signal coupled to the input end of the helix structure emerges greatly strengthened at the output end. Also, as is commonly attainable from transmission lines, broadband operation just 'comes along for the ride'.

The Beverage wave antenna (Fig. 4.12), at one time popular for low frequency radio communication links, in its simplest form is a long wire several tens of feet above the earth. This wire, in conjunction with its 'image' at some depth depending upon soil conductivity, forms an elementary transmission line. As such, it is terminated at one end by a resistance equal to the line's characteristic impedance. The other end feeds a radio receiver, also on an impedance match basis. A radio wave incident at the 'dummy' resistance end of the line excites travelling wave energy which arrives at the receiver.

Radio waves coming from the *opposite* direction cause travelling wave energy to propagate to the dummy terminating resistance and undergo no reflection. Thus the receiver does not 'see' such signals. Of salient importance is the fact the velocity of the radio wave governs its 'coupling' to this transmission line antenna. That is, energy is extracted in the mode of a travelling wave.

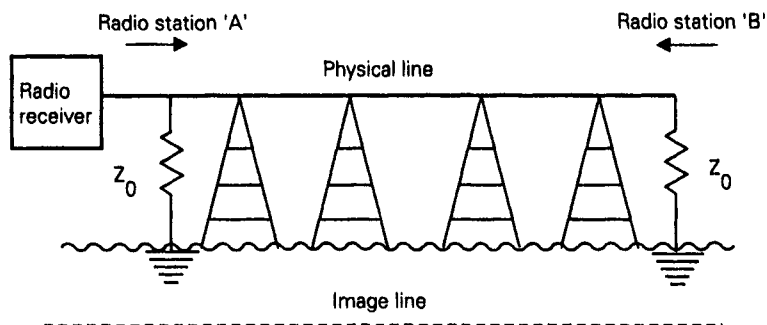


Figure 4.12 *The Beverage wave antenna. This antenna system operates as a transmission line. Often several wavelengths long at low radio frequencies, it guides travelling waves 'captured' from a passing radio signal to the receiver. It is inherently directional, providing maximum energy from radio station 'B' as shown above. Although station 'A' sets up travelling waves proceeding to the right, these never get reflected back to the receiver, being dissipated in far-end termination Z_0 .*

A rigorous analysis of the properties of transmission lines can be a mathematician's delight yet fail to focus attention on behaviour that is relevant to the operation of the transmission line transformer. For example, long and short transmission lines are ideally equivalent if their characteristic impedances are matched by both generator and load. However, since practical transmission lines cannot be lossless, energy will be transferred more efficiently by the short line. Moreover, if the load is anything other than a resistance equal to the characteristic impedance of the line, standing waves (or a portion thereof in the short line) will be formed with various adverse effects on losses and on frequency response. It will be seen to be of practical significance that a transmission line *can* be but a small fraction of a wavelength long without alteration of its basic operating principles.

Ideally, a transmission line can be surrounded by either magnetic or dielectric material with no effect on any of its propagation characteristics. With practical lines, there may be negligible modification of the line's characteristic impedance. This may be counter-intuitive in light of the fact that the characteristic impedance $\sqrt{L/C}$ is clearly a function of the per unit length inductance and capacitance. It would seem that these parameters

surely must be directly affected by adjacent magnetic or dielectric material. The resolution to this dilemma is indicated in Fig. 4.13. It is seen that the magnetic or dielectric material must be situated *between* the conductors of the line in order to exert effect. (Similarly, only materials situated *between* the central conductor and the outer conductor of coaxial lines influence the line's electrical characteristics.)

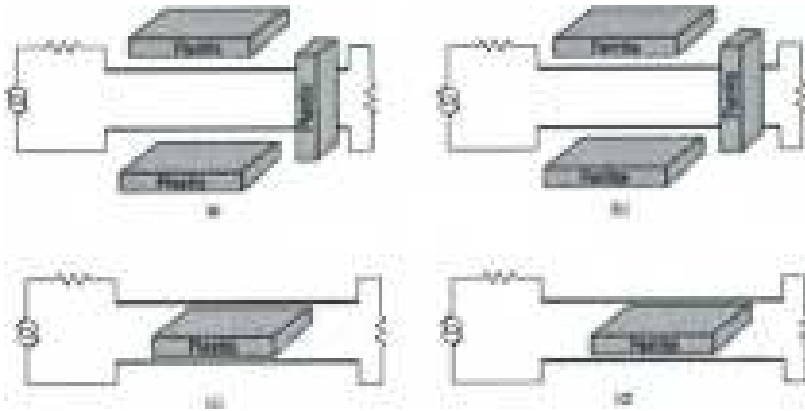


Figure 4.13 *Transmission lines with adjacent magnetic and dielectric substances. Ideally, the situations depicted in (a) and (b) have no effect on the characteristics or operation of the transmission line. This is why frequency response is nearly immune to stray capacitance and core loss is negligible in the transmission line transformer. Situations (c) and (d) strongly modify the characteristic impedance Z_0 of the line – dielectric material reduces Z_0 , whereas magnetic material increases Z_0 .*

Another behaviour of transmission lines should be kept in mind before investigating the transmission line transformer. Ideally, a transmission line need not be straight and unidirectional. Practical lines can be *coiled up* in various ways with negligible adverse effects. Thus, what may appear as an ordinary winding of some sort, may actually be used for its transmission line properties. Indeed, we shall see that such a winding may simultaneously serve as an inductance coil *and* as a transmission line.

The transfer of energy in transmission lines can occur in different *modes*; this is an extensive study in itself. In order to obtain a practical 'feel' for these manifestations, brief references to diverse implementations should help remove the aura of mystery often associated with transmission line transformers. For example, mention of the Beverage wave antenna suggests another single wire transmission line which, because of a different operational mode, can transfer energy without need of a 'return' line. This is the G-line, which has been useful in conveying television signals from an antenna to a remotely located television receiver.

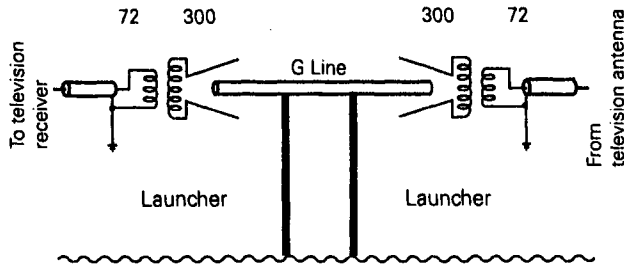


Figure 4.14 Energy transfer with a guide wire transmission line. A particular mode of electric and magnetic fields is required for travelling wave propagation along the single-wire line. This is accomplished by the tapered horn launchers. The insulation on the wire plays a significant role – it causes bending of the electric field to the extent that radiation is prevented. A transmission line operating in this propagation mode is also known as a surface waveguide. Practical lengths are in the order of several hundreds to several thousands of feet inasmuch as losses are lower than in conventional lines.

The unique feature of the G-line shown in Fig. 4.14 is the means of injecting energy into it and extraction of energy from it. It will be seen that tapered horn structures are used for these purposes. The travelling waves thereby set up in the single line derive from different electric field and magnetic field configurations than those associated with the more common situations in twin wire or coaxial lines. Nevertheless, energy is invested in travelling waves which *remain guided* along the wire as they proceed from the sending to the receiving end. Significantly, loss by radiation is not a major problem with the travelling waves excited in this mode. It must be pointed out, however, that such a wire must be enclosed by high-quality insulation in order to prevent radiation. The adjacent dielectric material causes the electric field to *bend* so that moving charges terminate on the wire itself rather than becoming detached as a radiation field.

Polyethylene-jacketed No. 6 aluminum wire is electrically and mechanically suitable for such a line. The line may be about 10 feet above ground level, but otherwise it should be at least a half wavelength clear of objects at the lowest frequency of intended use. Such lines tend to have impedances in the neighbourhood of $300\ \Omega$. For convenience, this can be reduced to $72\ \Omega$ to accommodate coaxial cable. Such impedance transformers are readily available.

Having looked at examples of energy transfer by travelling wave systems, let us return to our transmission line transformer. Its operating principle should be easier to grasp.

Our contemplation of various travelling wave implementations and transmission line devices has not been an idle diversion. Rather, it has been with the intention of showing that energy transfer can be accomplished

with electromagnetic fields in a manner quite different from that in conventional transformers. This being the case, it should come as no surprise that different design philosophies are used in designing transmission line and ordinary transformers. Indeed, their similar features such as windings on magnetic cores can almost be construed as a coincidence.

Let us first deal with the conventional transformer. From fundamental considerations of the physics of electromagnetic induction, a basic equation emerges which relates the voltage developed in a winding, the number of turns on the winding, the frequency, and the peak value of mutual flux, with the idea being to limit operation to the essentially linear portion of the core's magnetization curve. This equation gets us started – thereafter, the transformer designer must also deal with various practical aspects such as wire size, core loss, temperature rise, leakage inductance, insulation, manufacturability and economics. Inasmuch as these matters are inter-related, an effective blend of trade-offs is called for.

The design equation of the conventional transformer is:

$$E = 4.44fN\Phi_p \times 10^{-8}$$

where E is the impressed or developed voltage, f is the frequency in Hz, N is the number of turns, and Φ_p is the peak magnetic flux in the core.

The fourth parameter can be determined if we know or specify three of the parameters of the equation. In most practical situations, this enables determination of the number of turns N in a winding to be determined. However, a comment is in order regarding the *sizing* of the core. Φ_p in the basic equation provides the clue here:

Let $B_{\max}$ represent the maximum permissible *flux density* of the core. This is readily determined from charts and graphs pertaining to the selected core material. The basic goal of the selection is to retain a reasonable amount of linearity and to stay below flux density values which would incur excessive hysteresis losses. The area of the core will then be

$$\frac{\Phi_p}{B_{\max}}$$

If, for any reason we find the calculated area of the core unacceptable, it will be necessary to go back and assume a different value for N in the basic transformer equation.

When we consider the basic design approach to the transmission line transformer, an entirely different philosophy prevails. The number of turns, for example, may figure in the low frequency response of the transformer, but is not otherwise involved in the energy transfer. It is a similar situation with the core; it contributes to the low frequency response but is also not involved with basic energy transfer. This immediately implies that we can *dispend* with the core if wide-band performance is not of interest.

Rather than be concerned with the number of turns, our concern is with the *linear length* of windings. This should not exceed one-eighth of a wavelength for the highest frequency of intended operation. A second requirement is that the characteristic impedance of the winding should be the square root of the product of the input and output impedances seen by the transformer. Thus, if such a transmission line transformer is to work between a $12.5\ \Omega$ source and a $50\ \Omega$ load, the bifilar, twisted wire or coaxial winding should have a characteristic impedance of $\sqrt{12.5 \times 50}$ or $25\ \Omega$. In usual practice, it suffices to reasonably approximate these impedance requirements. The constructional format of a cored transmission line transformer is shown in Fig. 4.15.

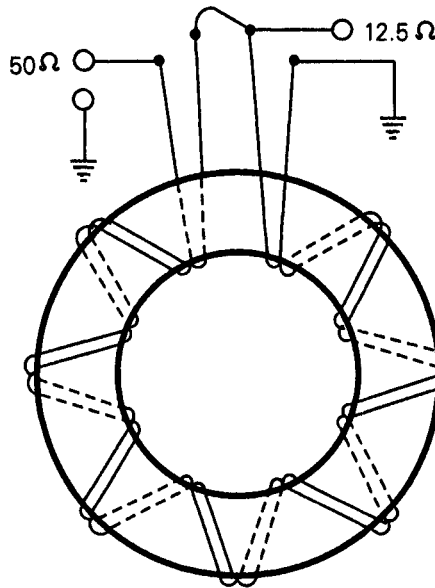


Figure 4.15 Practical construction of a transmission line transformer. The winding is spread out so as to make use of about 330° of the toroidal core. This favours short leads at the terminations. If need be, each of the sketched windings can actually be a pair of windings in parallel. That would make it easier to attain the low characteristic impedance ($25\ \Omega$) required in this transmission line transformer.

Considerable discussion has been allotted to the energy transfer feature of travelling waves on transmission lines. In so doing, it has been pointed out that the coiled configuration of the lines, as well as their associated magnetic cores, are not directly involved in the basic operation. Nonetheless, these windings, together with their cores contribute to the performance in an important way.

Even though the transmission line transformer handles electromagnetic energy in a certain mode, it still exhibits ordinary circuit effects. Thus, there is a *shunting* effect from any line that connects directly across the source or the load. This shunting effect is particularly harmful at *lower frequencies* because the inductive reactance is then lower. This immediately tells us that higher inductance is called for in order to broaden the frequency response to the lower frequencies. Keep in mind that we can accomplish this without affecting the characteristic impedance or the electrical length of the lines. First, the coiled configuration naturally increases this 'external' inductance. Second, a magnetic core increases the inductance still more.

What has been done is to get the windings to behave *simultaneously* as transmission lines and as RF chokes. This enables the transmission line behaviour to be extended down to much lower frequencies than would otherwise be the case. The high inductive reactance of these simulated RF chokes now free the source and/or the load from the ordinary circuit shunting effects. All the while, the behaviour of the lines to travelling waves has not been affected. Unfortunately, the symbols used with transmission line transformers can give the illusion that the windings are flux link coupled as in conventional transformers; once one learns to dispense with this notion, it becomes much easier to grasp the profound differences between the two devices. The extension of low frequency response due to the inductance-increasing effect of a magnetic core in the transmission line transformer is shown in Fig. 4.16.

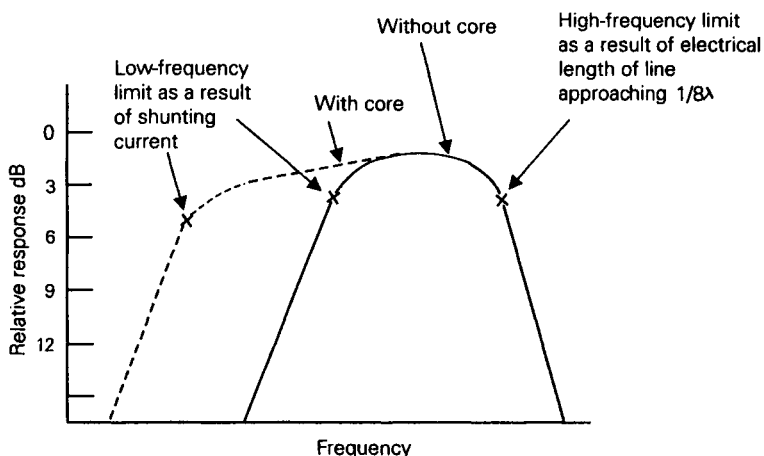


Figure 4.16 The use of a magnetic core to extend the wideband response. The transmission line windings must also act as RF chokes so that inductive reactance will be provided to soften the circuit shunting effect of the lines at lower frequencies. Although the coiled format of the lines partially meets this objective, low frequency response is greatly enhanced with the addition of the core.

Transformers as magnetic amplifiers

The so-called saturable reactor in which a control winding is impressed with d.c. in order to vary the inductance of an a.c. carrying winding can be considered a primitive magnetic amplifier. Except for the way in which it is used, the device resembles a transformer. Indeed, a conventional transformer can be operated in this fashion – the basic idea is to control core saturation so that the second winding changes its inductance, thereby controlling the a.c. in the load.

Usually, however, the term ‘magnetic amplifier’ is reserved for a very similar device which also exerts its control through core saturation. Known also as a *magamp*, this device does not rely upon the d.c. control winding to actually saturate the core, but only to govern the amplitude level (and therefore, the *time*) at which core saturation occurs during the cyclic excursions of the a.c. in the load winding. The onset of core saturation switches the inductance of the load winding from a relatively high to a very low value. Therefore, the *duration* of load current within a half-cycle is subject to control.

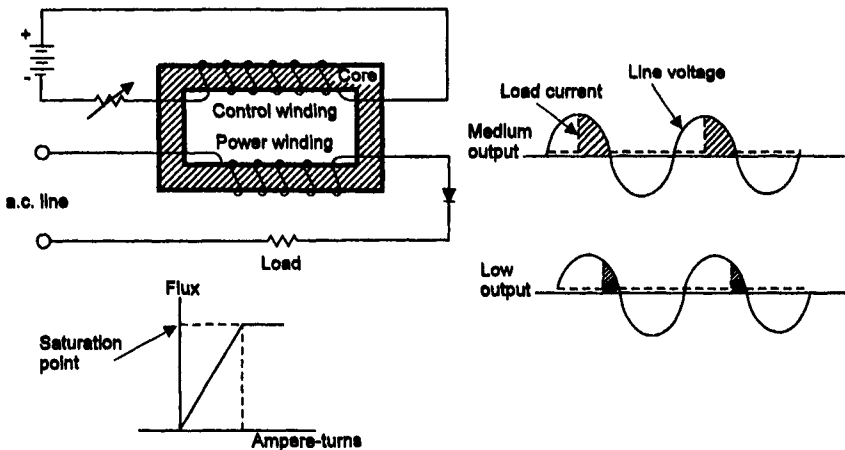


Figure 4.17 Operating principle of the magamp. The inclusion of the rectifier diode in the load circuit changes the way in which load current is controlled. Because of the diode, the core flux is due not only to the ampere-turns contributed by the control winding, but also to the d.c. component of the load current. The control winding current then determines when in the a.c. cycle the core abruptly goes into magnetic saturation. Upon saturation, the inductance of the power winding drops dramatically, allowing high load current for the remainder of the a.c. cycle. (Note the resemblance to SCR or thyatron operation where one can almost substitute ‘ionization’ for magnetic saturation.)

As the d.c. in the control winding is not employed to saturate the core but only to control the *timing* of core saturation, we must look to another means of projecting the core into magnetic saturation. This is accomplished by inserting a rectifying diode in the load circuit. The resultant half-wave rectification provides the d.c. component in the load winding which saturates the core. Operated in this mode, the device is said to be self-saturating. As stated earlier, control of the load current is determined by the current in the control winding—core saturation can be made to occur *early* in the a.c. cycle for high load-current, or *late* for low load current (see Fig. 4.17). Such control is continuous and smooth. It is also characterized by *high amplification*—small current in the control winding can control heavy load current. The operation is analogous to that of an SCR in a phase control circuit. Indeed, the load current waveforms of the two devices are substantially the same.

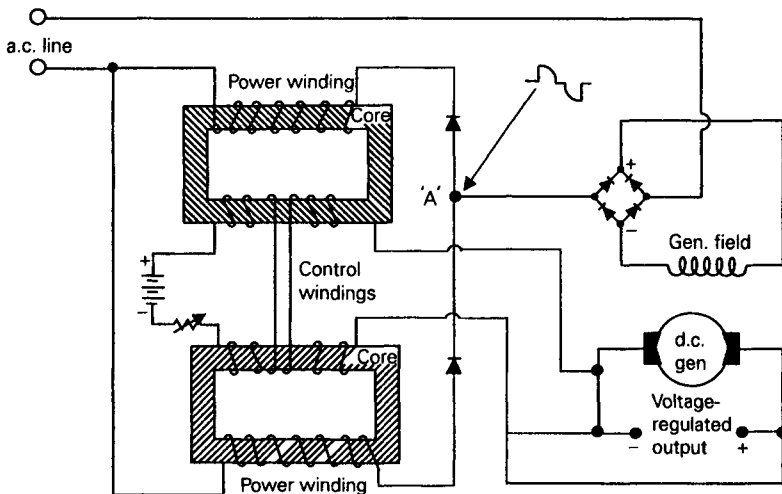


Figure 4.18 *Magamps in a full-wave circuit as a voltage regulator. Each of the two magamps operates as a half-wave phase control device, essentially the same as the magamp of Fig. 4.17. Together, full-wave control is realized. Note that the load current waveform is the same as that delivered by a Triac. Each of these magamps has two d.c. control windings. One provides ampere turns from an adjustable current supply. The other senses voltage or current which is to be regulated by feedback.*

An interesting magamp application is shown in Fig. 4.18 where the d.c. output voltage of a shunt field generator is regulated. The two half wave magamps are virtually used in a push-pull format to provide full-wave control of a.c. output current. The output current waveform from the magamps at circuit junction 'A' in Fig. 4.18 is the same as one obtains from a Triac phase-control circuit. In this application, the objective is to automatically

vary the d.c. field current of the generator so as to maintain constant d.c. output voltage. It will be seen that the output voltage of the generator is sensed by the 'extra' control windings on the magamps. Overall the stabilization technique is analogous to schemes commonly employed in solid state regulated power supplies. The reliability of the transformer-like devices tends to be better than would likely be achieved with Triacs, SCRs and other semiconductor components; this is especially true in hostile environments.

The manner of connecting the control windings to one another is important for optimum performance. With appropriate phasing, the effect of any a.c. voltages induced in the control windings can be cancelled. Thus, a.c. can be kept out of the adjustable d.c. current source. The control windings are often designated by other names, according to their intended circuit function. Thus, one encounters such terms as bias or gate windings. In the circuit shown in Fig. 4.18, the control windings associated with the adjustable current source can be said to be *reference* windings. And those control windings which sample the output of the system are sometimes termed *sensing* windings. Whatever their nomenclature, such windings exert control by establishing and/or varying the magnetic flux density of the core. Unlike the a.c. in the power windings, the control windings require relatively little power to exert their influence.

High permeability together with a narrow rectangular hysteresis loop are the prime targets for core selection in magamps. Other things being equal, the criterion of fast response is best met at higher frequencies.

Transformer coupling battery charging current to electric vehicles

One of the procedures necessary for owning and operating an electric vehicle is to subject it to recharge during the night hours. The straightforward method would be simply to mate a plug from the home-based charger to an accessible socket on the vehicle and then switch on the power. Direct current could either be delivered from the charger or the rectifier could be in the vehicle. Of course, various sophistications might be invoked – timing, regulation and charge-status sensing could be implemented. There could even be a faster than normal charging rate to take care of special circumstances. It is not easy to play with numbers inasmuch as there are a multiplicity of trade-offs and contradictions. These essentially pertain to the nature of the battery system, the allowable current drain at the residence, the battery life-span in terms of charge cycles and economic aspects involving the hours of energy delivery from the utility. For the sake of this discussion, let us just assume that an 8- or 10-hour overnight charge would suffice for energy replenishment for many electric vehicle owners.

A shortcoming of the above scenario is the problem with reliability and safety of the plug and receptacle connection. Regardless of safeguards, there could ultimately be shorts, open circuits and hazardous electric shocks. A few of the causative factors one can mention are abuse, ageing, chemical corrosion, mechanical failure, water on the floor, etc. Unlike other several hundred volt systems using plug and socket connection, *this* connecting link is likely to be made and parted several or more times per week. How much better the situation could be if the charging current was delivered to the vehicle via inductive coupling, that is, by *transformer action*. This would be true whether the vehicle was parked over an energized 'primary' winding, or whether the operator inserted a lightweight paddle containing such a winding in a vehicle-slot provided for such purpose. These methods are illustrated in Fig. 4.19. Both schemes make use of vehicle-mounted pickup loops and rectifiers.

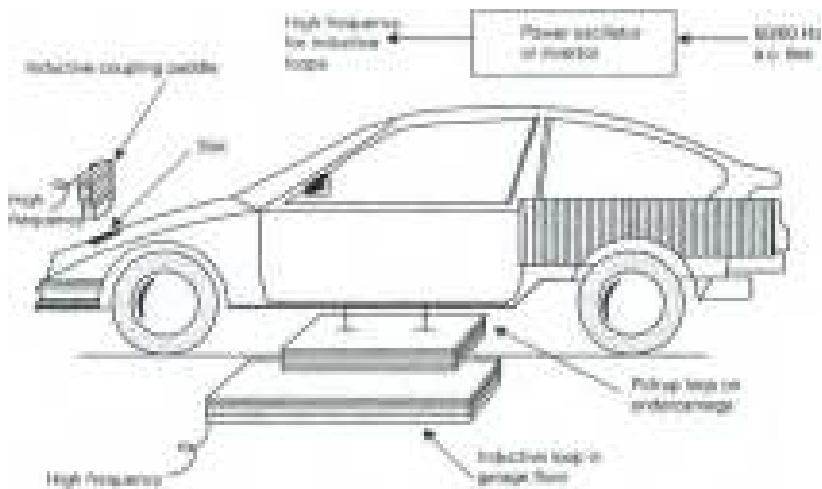


Figure 4.19 *Transformer-coupled formats for charging electric vehicle batteries. One method makes use of a large 'primary' winding in the garage floor. Energy pickup by induction would take place in the 'secondary' winding on the underside of the vehicle. The rectifier is inside the vehicle. The other, basically similar, technique utilizes a small paddle which the operator inserts into a slot in the vehicle provided for the purpose. Unless radio frequencies were used, both schemes would be likely to require physical contact of magnetic core material to achieve good transformer action.*

At the time of writing, inductively-coupled charging has actually been implemented by both large and small manufacturing firms. Several versions appear to have proven workable, but little has been achieved in the way of standardization and consensus that this is, indeed, the way to go. A

formidable barrier to further progress is the contention that the plug and socket link used in conduction-connected charging *can* be made reliable, foolproof and safe. If this school of thought wins out, it will enjoy the benefit of being a less costly technique.

Those with an experimental turn of mind generally feel that the final appraisal of inductively coupled charging should not yet be made. More investigation is needed, they contend, of coupling architecture, optimum frequencies, inverter devices and circuit techniques. For example, the introduction and development of IGBT solid-state switching devices has given new impetus to the search for an overall satisfactory inductively-coupled charging system. The IGBT is electrically rugged, inexpensive, and has a good record of efficient inverter operation in the 5–50 kHz region. Another device, the GTO, has demonstrated its switching capability in rail and traction vehicles and is also a good candidate for the design of the inverter which appears necessary for the induction-charging technique – the basic requisite being a source of power much higher in frequency than the 50/60 Hz utility supply.

It turns out that the choice of frequency is of prime importance. At one extreme, looms the difficulty of transferring 50/60 Hz power across an air gap or without a low reluctance closed magnetic core. At the other extreme, that is, in the several-hundred kHz radio frequency region, power transfer is practical, but it would bode ill for communications to have thousands of rather powerful transmitters in the environment. The problems associated with radiation and harmonics would bound to be overwhelming. Energy coupling at 5 kHz might marginally suffice in certain arrangements. 50 kHz could be expected to behave as a low radio frequency and provide a high coupling coefficient with minimal magnetic material as the transformer 'core'.

50 kHz operation is intriguing, but at this frequency it is not so easy to obtain high switching efficiency from the active elements in the inverter. RFI might also loom up as a problem, particularly with square waveforms. The audio frequency region between 5 kHz and, say, 18 kHz merit consideration if we are satisfied that energy transfer at 5 kHz could be marginally practical. However, it would be wise to avoid the whole audio frequency region – there are bound to be audible noise problems. From the standpoint of consumer acceptance, it is best to provide silent electrical equipment. Innovation is needed and one shouldn't exclude even a 'far-out' scheme such as microwave linkage using proven magnetrons from oven technology.

As a consequence of our off-the-cuff deliberations, it appears that a reasonable frequency range with which to commence experimentation and evaluation would be in the 20–25 kHz region. Our efforts would be aided by the considerable accumulation of technical information pertaining to inverter operation at such frequencies. Most certainly, the gathering together of components needed to build experimental apparatus would

pose few problems. Later, one could scale the operating frequency higher, say to the 100 kHz region. Here, good transformer action with a lightweight coupling 'paddle' should present no problem. Power levels of 2–20 kW can be envisaged, depending upon the desired charging rate. Commercial practicality would very much depend upon success in suppression and shielding of radiation.

The experimenter should strive for maximum area of the primary and secondary windings, and for a minimal air gap between the two. The effect on power transfer of resonating the windings should be investigated and attempts should be made at optimizing the Q values. Much needs to be learned about the use of magnetic material in and around the windings. It is conceivable that U- or C-shaped cores could make physical contact so that a fairly good transformer format is realized. With such a complete magnetic circuit, transferring power across an air gap would no longer be the major problem, although minimizing leakage inductance would still tax the experimenter's ingenuity.

Despite the imperfections of the transformer format for delivery of charging current, it is conceivable that the overall scheme could eventually prove practical, reasonably efficient and economically justifiable.

5 High-voltage transformers

Although it seems that high voltage transformers are just simple extensions of those operating at more moderate voltages, there are good practical reasons for considering them as a separate group. This is because they are particularly prone to certain problems that tend to be benign or even absent in ordinary voltage transformers. Among such problems are arc-over, aggravated insulation stress, and ozone generation. The severity of these malfunctions depend very much on temperature, humidity and time. What might look passable in a quick laboratory test might, nonetheless, lead to early failure in the actual operational environment.

Troubles of a more subtle nature accrue from simply increasing secondary turns in order to gain a higher step-up ratio. Leakage inductance is likely to increase faster than anticipated because the added turns now get spaced too far from the core and the primary winding. The self-resonant frequency may decrease to a region where interference with proper circuit action can be expected. (Not only does inductance go up as the *square* of the number of turns but additional distributed capacitance is introduced as well.)

One can easily see practical consequence of these matters. A quick example is the flyback transformer in television sets. Instead of attempting to design these high voltage transformers for the needed 25 kV or so, it is common to settle for 8 kV flyback transformers used in conjunction with d.c. voltage multipliers such as triplers. Further ways to avoid the undesirable side-effects of high voltage in transformers will be discussed in the ensuing text.

Emphasis on high voltage

It is generally understood that devices such as the Tesla coil, the Ruhmkorff induction coil, the Ford Model-T spark coil, ordinary ignition coils, and television flyback transformers are all basically step-up high voltage

transformers. These devices differ in operation from an ordinary step-up transformer in which the primary would be impressed with 50/60 Hz, and in which the high voltage induced in the secondary would be expected to more or less reproduce the sinusoidal wave form of the primary excitation. By contrast, the family of devices mentioned depend upon shock excitation of their primary windings by steeply rising and falling wave-forms.

The consequence of such operation is that the induced secondary voltage is oscillatory at a low radio frequency largely determined by the inductance and the distributed capacity of the secondary winding. As a practical observation, this frequency tends to be in the 100–300 kHz region for a large number of these high voltage devices. Usually, there is considerable leakage inductance, tantamount to saying that fairly loose coupling exists between the windings. This enables the resonance in the secondary to boost the peak voltage to a higher level than would be dictated by the turns-ratio alone. Although energy considerations are often important (such as in an automobile ignition coil), the primary function of these devices is the production of high voltage. The design philosophy of conventional power transformers applies only minimally here.

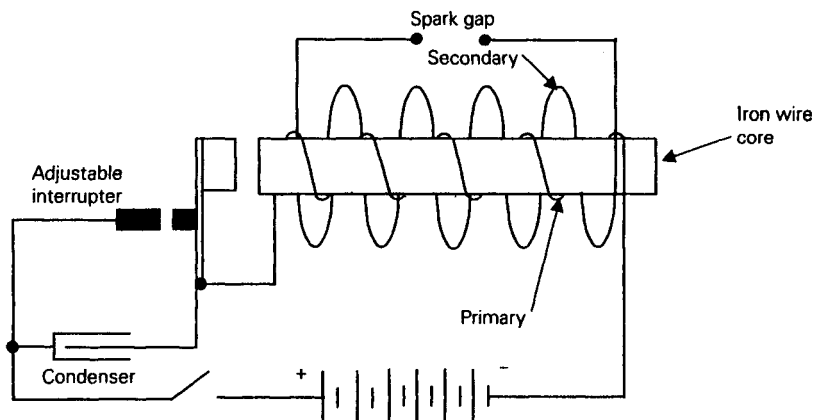


Figure 5.1 *The Ford Model-T spark coil. This classic example of an induction coil is actually a high ratio step-up transformer. The current interruption rate is adjustable, extending from low to medium audio frequencies. (In auto-repair, it remains common to refer to the capacitor as a 'condenser'.) The automatically-vibrating interrupter is actuated from magnetic flux from the core and is essentially a buzzer. 6–10 V and 2–3 A suffice to produce an energetic spark discharge.*

The Ford Model-T spark coil shown in Fig. 5.1 is representative of many of these high voltage devices, regardless of name. The capacitor across the interrupter contacts prevents arcing. In so doing, the longevity of the

contacts is greatly extended. Also, the near instantaneous interruption of the primary current increases the high voltage developed in the secondary. Current interruption occurs at an audio frequency rate, giving rise to bursts of damped RF voltage in the secondary. Modern ignition coils generally have three terminals, being internally connected as auto transformers and the non-vibratory current interruptor is physically separate from the coil itself.

Stepped-up high voltage from not so high turns ratio

An experimenter is interested in determining the approximate turns ratio of an automobile ignition coil. He impresses a small audio frequency sine wave on the primary terminals and uses an oscilloscope to measure the induced voltage in the secondary. As a result of this measurement technique, it is found that the voltage step-up is approximately 100. This does not appear reasonable because 100 times the 12 V of the automobile battery falls far short of the 15 000–30 000 V needed for firing spark plugs. What is the nature of the discrepancy?

Although not commonly referred to as a 'flyback' transformer, the ignition coil develops its high secondary voltage in a similar manner to the fly-back transformer in a television set. In both instances, the primary winding requires a waveform with a very high rate of voltage or current change. Such a waveform induces a high voltage counter EMF in the primary and it is *this* induced voltage which is stepped-up further in the secondary. Thus in an automobile ignition system, the abrupt cut-off of the applied 12 V induces a counter EMF in the primary with a peak amplitude of about 250–300 V. When this purposely-produced 'transient' is multiplied by a 100 to 1 step-up turns ratio, one obtains the 25 kV or so needed for reliable firing of the plugs. The basic breaker point ignition system is shown in Fig. 5.2.

This situation provides an interesting insight into the operation of all transformers regardless of impressed waveshape. The *true* transformation ratio of any transformer in operation is the ratio of the induced voltage in the secondary to the induced voltage in the primary and *not* the ratio of secondary to primary terminal voltages. However, for *sine wave* operation, the terminal voltage relationship provides a good approximation in ordinary practice. Note that although secondary induced voltage will be slightly higher than secondary terminal voltage, the primary induced voltage will be slightly less than primary terminal voltage. Consideration of these facts leads to a true secondary to primary transformation ratio slightly higher than the value commonly alluded to.

In Fig. 5.2 the capacitor (still known as a 'condenser' in automotive circles) generally runs between 0.20 and 0.25 μF . Its value can be somewhat critical for best results. At its optimum value, the burning of the two breaker points will be minimal and about equal. The most common failure mode of

this capacitor is the onset of high leakage. This reduces the peak amplitude of the counter EMF when the breaker points opens, causing a weak and cold spark to occur in the spark plug. As normally tested, an old capacitor may appear reasonably good, but under actual operating conditions, a weak or intermittent ignition system may be in evidence. Breaker point opening produces *single* bursts of high voltage transients in the secondary; this contrasts to the *sustained* high voltage pulse train produced in the Model-T coil because of its vibrator.

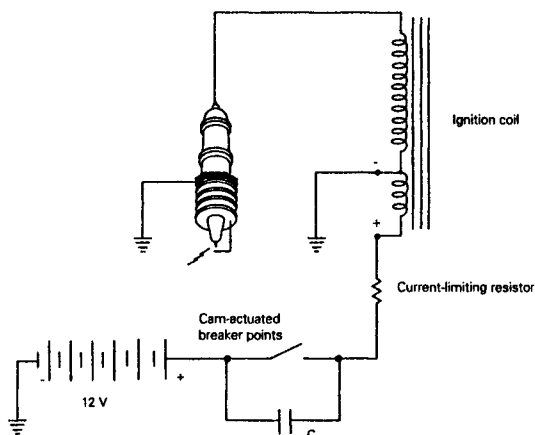


Figure 5.2 High voltage generation in automobile ignition system. Although the old-style mechanical breaker scheme is shown, the principle remains valid for modern electronic systems where similar coils (autotransformers) continue to be used. Although operating from a 12 V battery, 20 kV or more can be developed for the spark plug. When the breaker points are opened, very rapid collapse of the energy stored in the magnetic field of the coil takes place. This generates a pulse voltage of about 300 V which is then stepped-up by auto transformer action. The capacitor prevents arcing when the breaker points open.

It will be noted that the primary terminals on the ignition coil have polarity markings. At first thought, this should not be necessary for one would not expect much difference in the stepped-up voltage either way. Surprisingly, it turns out that the spark plug has a polarity preference for easy firing. This is because the central electrode runs *hotter* than the grounded one. The basic phenomenon of voltage breakdown of the compressed fuel-charge is assisted by thermal emission of electrons. The spark discharge path is then an ionization event. Also probably contributing to polarity preference is the fact that the spark plug gap is not comprised of identical surface geometries as would be the case for two spherical 'points' such as seen in high-school physics laboratories for demonstrating high voltage discharges.

The core of the ignition coil is a bundle of soft iron wires. This is ample for

the purpose as a closed, low reluctance magnetic circuit is not needed for energy transfer of short duration pulses. Harmonics, shock excitation, leakage inductance and winding capacitance can combine to produce a wide spectrum of high frequency energy which radiates and tends to interfere with radio, television and other communications. This has been brought under quite good control by means of resistive spark plug cables and by incorporating resistors within the spark plugs themselves. Sufficient dampening of radio frequency oscillations can be realized without excessive depletion of sparking intensity.

A brute force approach to Tesla coil action

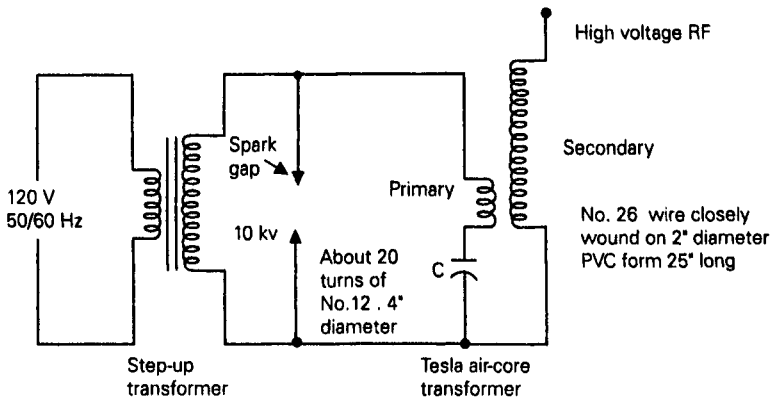


Figure 5.3 Simple Tesla coil system for producing high voltage at radio frequency. The line voltage step-up transformer will often have a 25–40 mA rating. A wide-gapped automobile spark plug can be used. The high voltage capacitor can initially be about 500 pF. Optimum operation roughly coincides with this capacitor resonating the primary winding to the self-resonant frequency of the Tesla coil secondary. It is obviously easier to experiment with the Tesla primary winding than with the capacitor. Caution – lethal voltages are present.

The high voltage system shown in Fig. 5.3 is a tried and tested method of generating high voltage in the low radio frequency region. It has had considerable use in welding to initiate ionization of the air so that the main welding arc can take place. Otherwise, the initial radio frequency arc has little welding influence inasmuch as it is not backed up by high current. At the same time, the radio frequency voltage poses no threat to the operator of the welding apparatus. The nervous system does not perceive radio frequencies as an electrical shock because such currents travel on the very outer surface of objects or people (the so-called ‘skin-effect’). Moreover, radio

frequencies at a sufficiently high voltage exhibit a greater tendency for arcing and for ionizing the air than power line frequency.

As can be seen, a step-up transformer is used to bring the 50/60 Hz a.c. up to sufficient potential to produce continuous actuation of a spark gap. Such a spark gap acts as a rapid-fire switch, interrupting the current twice per cycle. These current waveforms are very steep and are able to shock-excite the resonant circuit at a low radio frequency. The nice thing about this technique is that the hazardous low frequency a.c. is completely isolated from the output. An ordinary automobile spark plug suffices well for experimental purposes. Also, a Model-T type induction coil can be used to operate the spark plug. In any event, optimum operation depends upon energy back-up and on satisfying the resonance of the secondary output winding. This may call for a bit of 'cut and try' regarding both the high voltage capacitor and the design of the radio frequency output transformer.

Transformer technique for skirting high voltage problems

An interesting and useful transformer technique derives from radar technology. In many systems, it is desirable to deliver high voltage pulses to a magnetron which then becomes the source of the microwave radiation. The pulses are formed by a high power modulator at a selected pulse repetition rate and applied to the magnetron through a pulse transformer. As the magnetron has the format of a simple diode vacuum tube, the circuit implementation appears straightforward enough. Some practical problems loom up, however.

A more or less garden variety pulse transformer with an appropriate step-up ratio might be considered to supply the 20 kV or so pulses between the cathode and anode structure of the magnetron. Because of the magnetron's construction, it is common practice to apply negative-going pulses to the cathode filament terminal thereby enabling the anode structure together with the metal outer portion of the tube to be at ground potential. This is in the interests of safety and also obviates any need to insulate the magnetron from chassis metal. Insofar as concerns the operation of the tube, it 'sees' the required positive anode with respect to cathode polarity condition.

It can be seen by referring to Fig. 5.4 that the use of a two-winding pulse transformer would result in the filament being at a very high voltage with respect to ground. This, of course, would call for a large and costly filament transformer with a high voltage insulated secondary winding. Even if this 'brute force' technique were pursued, it is likely that there would be problems because of degradation of the pulse shape. Fortunately, a clever deployment of a three-winding pulse transformer can allow use of an ordinary filament transformer and at the same time enable the magnetron to operate with its anode structure at ground potential.

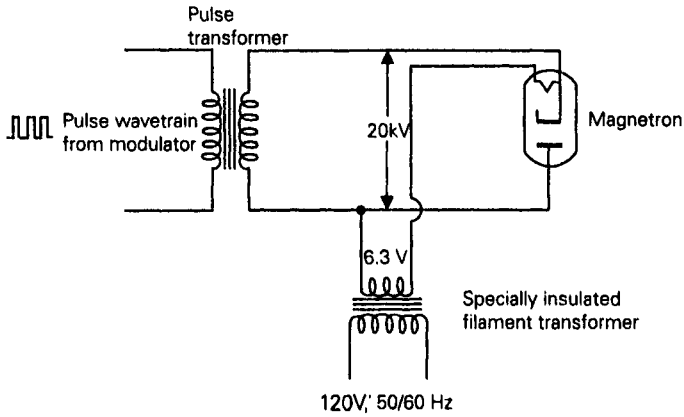


Figure 5.4 'Brute-force' circuitry for the operation of a magnetron. Although conceptually simple, this arrangement poses size, weight and cost problems; moreover, it would not be easy to preserve pulshape fidelity. The prime difficulty centres around the required high voltage insulation for the secondary winding of the filament transformer. Note that the pulse transformer secondary must be wound with wire capable of carrying the filament current of the magnetron.

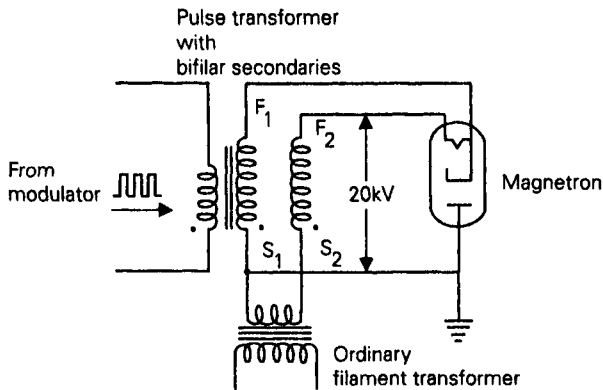


Figure 5.5 Magnetron drive scheme using three-winding pulse transformer. The bifilar-wound secondaries remove the high voltage insulation requirement from the filament transformer. At the same time, the anode structure of the magnetron operates at ground potential.

A basic circuit of a three-winding pulse transformer used for driving a microwave magnetron is shown in Fig. 5.5. The two secondary windings are bifilar-wound with conductor sizes capable of carrying the magnetron filament current. Because these secondaries are identical, there is zero

potential between winding ends S_1 and S_2 . Similarly, there is zero potential between winding ends F_1 and F_2 . Note that the filament transformer 'sees' only ground or near ground potential. Therefore, no high voltage insulation is needed in the construction of the filament transformer. The anode structure of the magnetron is also at ground potential, but negatively polarized high voltage is applied to its cathode. This provides the same operation one would expect from a grounded cathode circuit with positive anode.

Overall, this transformer scheme leads to a system featuring safety, reliability, freedom from flash-overs and savings in cost and weight. This setup also proves helpful in attaining pulse integrity of the microwave bursts. The phasing-dots depicted for the three windings are proper for a modulator delivering positive-going pulses to the primary. For negatively-going modulation pulses, the phasing of the primary winding would be reversed. (Or, alternatively, the phasing of the two secondaries could be reversed *as a pair*.)

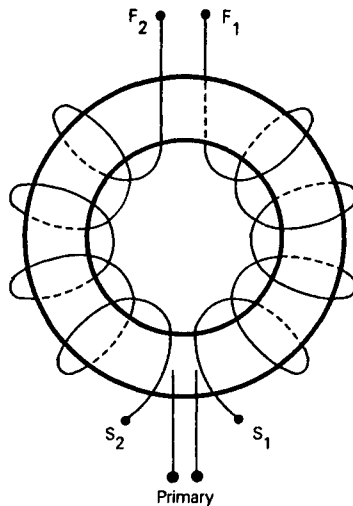


Figure 5.6 An alternative winding format for a pulse transformer. Note that one secondary is wound in a clockwise direction, while its mate is wound in a counter clockwise direction. In operation, there is zero potential between F_2 and F_1 ; likewise, there is zero potential between S_2 and S_1 . This transformer provides the same features as the bifilar-wound type. The construction utilizes insulation between the core and the primary and then again between the primary and the secondaries.

An alternative winding format for a high voltage pulse transformer is shown in Fig. 5.6. The circuit connections remain the same as the transformer shown in Fig. 5.5 except that bifilar windings are not used for the secondaries. Instead, the secondaries comprise two separate windings.

These windings are wound in *opposite* directions. In Fig. 5.6 the primary winding is not seen because it is beneath the insulation that separates the primary from the secondaries. As with the bifilar type, this pulse transformer enables the anode structure of a magnetron to be at ground potential, with negatively polarized pulses applied to the cathode. This pulse transformer also frees the magnetron filament transformer from requiring high voltage insulation.

Both of these pulse transformers have been commonly made with step-up ratios of 4 or 5 to 1. The 20 kV, or so, developed in the secondaries has been found convenient for popular magnetrons. At the same time, 4 or 5 kV has been readily forthcoming from modulators feeding the primary. The experimenter may wish to incorporate these winding techniques in other high voltage applications such as photo-flash, laser or X-ray apparatus.

A novel winding pattern

At high radio frequencies, certain departures from conventional practice often pay dividends in performance. Consider, for example, the two toroidal transformers shown in Fig. 5.7. Both have *bifilar* windings and both have one to one transformation ratios. The transformer at A may be said to be conventionally wound. It is good in many respects and will be found to have tight coupling and minimal leakage inductance. However, at high frequencies and high power levels, the capacitance between the physically-close leads can prove troublesome. Because the potential between the start and finish ends of the windings can be quite high, a very small capacitance is sufficient to produce an undesired resonance or to limit the tuning range. Even worse, there is the possibility of *voltage breakdown* and subsequent arcing.

Transformer B retains the good features of transformer A but the start and finish leads of the windings are physically far apart. The capacitive effect prevalent in transformer A is practically absent. Also, the likelihood of *arcing* is greatly reduced. As in transformer A, transformer B utilizes almost the entire core. The electromagnetic operation of the two transformers is practically identical, despite their different winding patterns.

Note what has been done in the winding of transformer B. Half of the winding is in one direction and half is in the opposite direction. Contrary to one's first thought, this does not bring about cancellation of the inductance. Indeed, overall inductance is about N^2 or four times the inductance of each half-winding. This is due to the manner in which the series connection of the two half-windings is made. One might say that the phasing of the two half-windings is additive. If it is practically feasible, it is a good idea to have the interconnecting coil segment go right across the central region of the toroid as depicted in transformer B. Otherwise, there is nothing remarkable about this winding technique. At lower frequencies and/or power levels

it might show no advantages. Also, either type may prove best suited to a particular PC board layout.

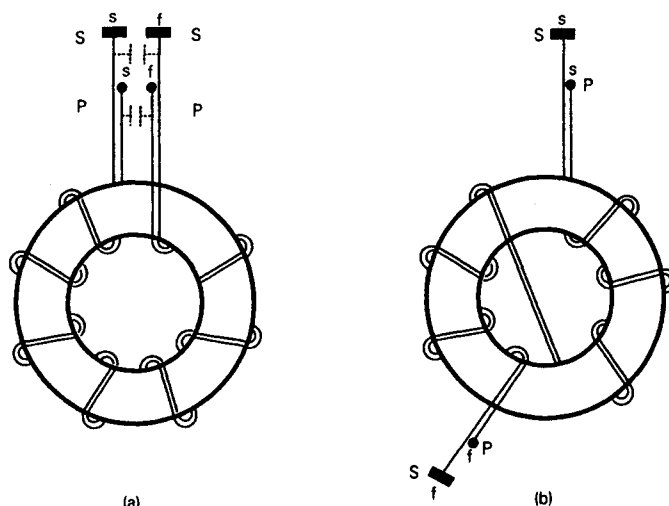


Figure 5.7 Winding techniques for high frequency toroidal transformers. Both are 1:1 bifilar-wound types, but with different winding patterns. (a) This is a conventional winding. Note the physical proximity of the start and finish leads in both primary and secondary windings. At high frequencies, these lead capacitances can exert strong effects on tuning, broadbanding and Q values. In transmitters, arc-overs can occur because the highest RF potentials exist between these leads. (b) With this super toroid winding pattern, the start and finish leads are far apart. End lead capacitance is reduced to a negligible level and voltage arc-over is no longer likely.

Other techniques for high voltage transformation

Modern versions of the Tesla coil can make use of solid-state devices to generate the harmonic-rich pulse to excite the primary of the Tesla coil. (Of course, the so-called Tesla coil is actually an RF transformer.) Perhaps, however, most of these implementations more closely resemble the flyback circuits used to generate high voltage in television sets. It is not easy to replace a spark gap or a mechanical interruptor with a solid-state device when truly high voltage is required from the Tesla coil arrangement and attempts to do so remain on an experimental basis. The problems are with the rise and fall times of such devices as Triacs, SCRs, GTOs and IGBTs. Power MOSFETs are faster and lend themselves to easy paralleling.

A compromise can be made by using any of the above devices in conjunction with a step-up transformer to operate a spark gap. The only

shortcoming of this approach is that one cannot say that it is completely solid-state. There is, however, an interesting design technique for achieving Tesla coil performance with all solid-state circuitry.

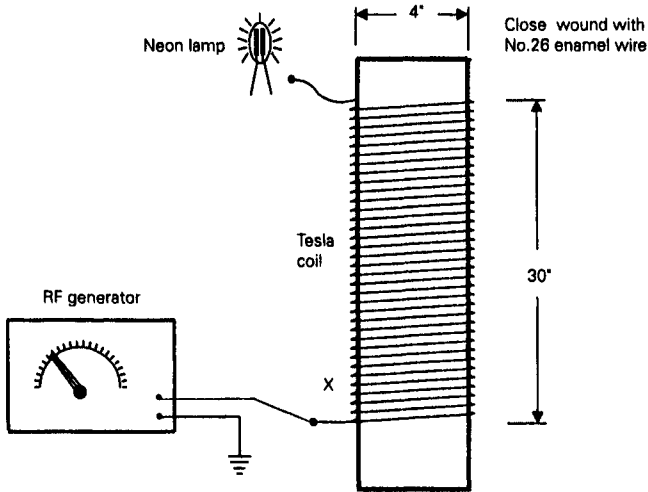


Figure 5.8 Test setup for ballparking quarter-wave resonance of Tesla coil. For this operational mode, no primary winding is used. Resonance and voltage step-up is determined more by the length of the wire in the winding rather than by the inductance of the coil and its distributed capacitance. This Tesla 'coil' is actually a helix quarter-wave transmission line.

Let us look at the high voltage output winding as a *quarter wave transmission line* rather than as the secondary of a transformer. After all, it has the geometric configuration of a long cylinder and we can surmise that it will behave as a quarter wave line at *some* frequency. With this philosophy we are less concerned with the inductance or with the distributed capacity which figured predominantly for behaviour as a parallel-resonant tank circuit. We are now more concerned with the actual *physical length* of the winding. Although in the form of a helix, transmission line properties can be elicited from such a format. In other words, the idea is to operate our 'coil' in a different mode than conventional LC resonance. In so doing, the free far-end of this quarter wave 'helix-line' will develop high impedance together with high voltage. A simple test setup for investigating the quarter-wave resonance behaviour of a helix transmission line is shown in Fig. 5.8.

There is a profound difference in the way the quarter-wave helix is driven from the more conventional Tesla-type winding. The 'cold' end of the helix is not excited by a several-turn primary coil. Instead, direct connection is made to the helix from a low impedance radio frequency source. And instead of shock excitation, the radio frequency source is tuned to the

quarter-wave resonance established by the helix. For practical operation, a power driver is shown in Fig. 5.9. Experimentation with gate biasing will probably be required.

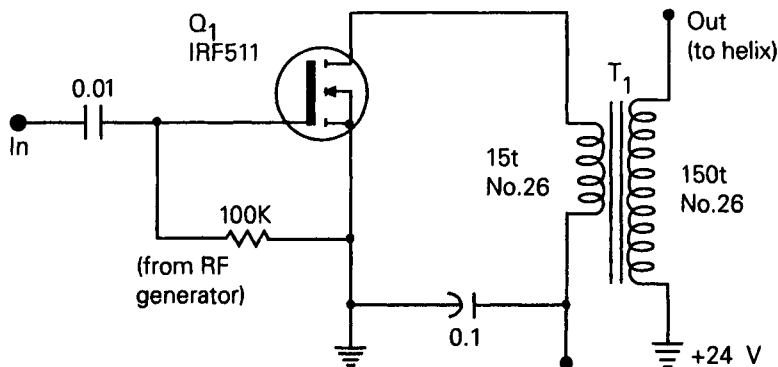


Figure 5.9 Driver circuit for the quarter-wave high voltage helix. For practical operation of the helix, this driver can be inserted at point 'X' of Fig. 5.8. Do not attempt to use utility frequency silicon steel for the core of transformer T_1 . A ferrite core such as the Amidon EA-77-250 'E' core (or equivalent) will handle the radio frequency energy with negligible core losses.

Despite the use of the helix winding as a quarter-wave resonator, it stands to reason that the ordinary LC behaviour must also be somehow involved. Usually it will be found that the natural LC resonance is at a lower frequency of than that the transmission line resonance. As the impedance of such a parallel resonant tank is inductive, it is as if the quarter-wave line were shunted by a high reactance RF choke. This is fine but there is obviously much room for optimization by experimenting with the length, diameter and winding pitch of the helix. While doing so, one can make use of a temporary primary coil in order to determine the parallel resonant frequency of the helix. Other things being equal, it probably should be just slightly lower than that of the quarter-wave transmission line resonance. This does not, however, appear to be a critical criterion of performance.

It would be only natural to ponder the feasibility of employing this voltage transformation technique at a much higher frequency, thus simplifying the quarter-wave line. It happens, however, that for frequencies in the neighbourhood of a couple of hundred kHz, the metre-long helix is a relatively poor antenna. At higher frequencies, there would be greater radiation efficiency. This would not only have the effect of loading down the high voltage, but would be much more involved in interfering with various radio, television and communications services.

The current transformer

The current transformer traces its origin to high voltage instrumentation. Specifically, it provided a convenient means of safely monitoring the current in a high voltage a.c. line. Obviously, cutting the line and inserting an ammeter involved service interruption and danger. In its basic form, the current transformer has a primary of one or several turns. The secondary then has a multiplicity of turns and is intended to be short-circuited through the indicating ammeter. Most often the design is such as to provide 5 A of secondary current as corresponding to a full-scale reading. An approximate inverse turns ratio exists between the currents in the two windings. The ratio is not exact because of the effect of the magnetizing current. At low loads, the error can be appreciable. This is why calibration is generally for a specific ammeter. Current transformer types and applications are shown in Fig. 5.10.

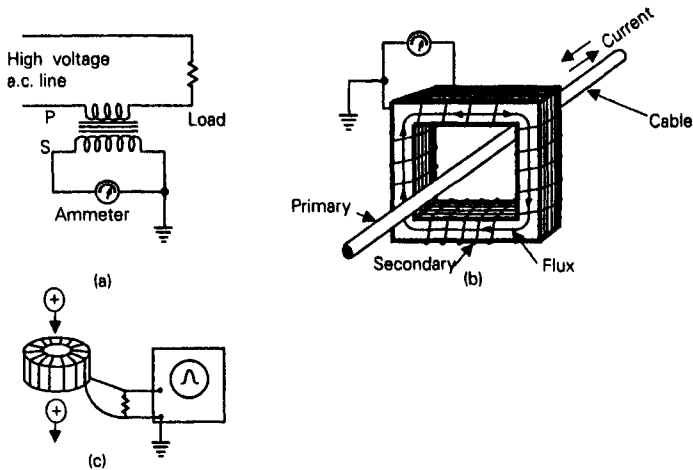


Figure 5.10 Basic current transformers. (a) Current measurement technique in high voltage utility line. Isolation, safety and convenience are provided by this classic use of the current transformer. (b) Physical arrangement of a through-type current transformer. This format makes high voltage insulation easy to implement. The current-carrying conductor acts as a single turn. (c) Adaptation of a current transformer for research projects. Charged particles in transit through the aperture of the solenoid behave as electric current in a single turn 'primary'. The load resistance helps attain flat frequency response over a wide band.

Some precautions are in order when using the current transformer in high voltage systems. One secondary lead should be well grounded. For extra safety, it is probably a good idea to ground the core and container too. The current transformer should never be open-circuited; if the

ammeter is to be removed, a short-circuit should *first* be placed across the secondary. There are several reasons for this precaution. When the current transformer is open-circuited, it operates as a voltage step-up transformer. The stepped-up voltage can puncture its insulation. It also becomes a hazard to personnel. Also, the mutual magnetic flux is much greater for open-circuit than for short-circuit operation. This can result in strong residual magnetism left in the core, thereby increasing magnetization current and introducing error in the transformation ratio. (As a corollary, where one strives for measurement accuracy, it is wise to demagnetize the core before placing the current transformer in service.)

Many specialized current transformers are now available. A popular type is associated with clamp-on meters. These open and then close around a current carrying conductor, providing quick measurement and maximum convenience. They find wide use in high current, moderate voltage systems and have been made sufficiently sensitive to be very useful in electronic work. By means of various techniques, such as disturbance of resonance, the sensing of d.c. has been made possible. However, many 'd.c. current transformers' actually base their operation on Hall effect elements.

Physicists often use a current transformer in the form of a magnetic material toroid with only a 'secondary winding'. Although there is no physical primary, the passage of *charged particles* through the aperture of the toroid yields a response just as if a current-carrying conductor passed through the toroid. Although simple, these devices are carefully designed and constructed to exhibit wide-band frequency response so that a variety of charged particles with differing velocities can be accurately evaluated. To facilitate installation, the device is available in split-halves; this enables quick clamp-on to the pipe carrying the particle beam.

An important attribute of the current transformer is that it exerts minimal effect on the line or circuit undergoing measurement. This stems from the low inductance of the one- or several-turn primary and from the near short-circuit reflected by the secondary and its meter. In addition to measurement instrumentation, current transformers are widely encountered in electronic systems. In regulated power supplies, for example, current transformers may be involved in the feedback loop in order to stabilize load current; they may also be used to sense development of current overshoots and spikes which tend to damage transistors. (Once sensed, the offending transients can be limited in magnitude.)

Stepping-up to high voltage and new transformer problems

The design and construction of the high voltage transformer is an art beyond that pertaining to its low voltage counterpart. Problems unique to

high voltages can commence at several kV and are definitely manifest at 10 kV. Contradictions arise between the space requirements of insulation and the desire to keep leakage inductance within bounds. The relatively small gauge secondary winding can pose new mechanical and electrical problems. Core insulation and inter-layer insulation of the windings place greater priority on window space than with a lower voltage transformer of similar kVA rating. At the same time, one must continue to be mindful of saturation, temperature rise and exciting current.

The important new factor involved in the operation of the high voltage transformer is *corona*. This phenomenon is the ionization of the air molecules. As such, it extracts energy from the transformer and is tantamount to a load resistance placed across the secondary. However, the reduction in efficiency due to this unwanted power drain is not the chief complaint. Corona exerts several more deleterious effects than mere loading. Ionized air, comprising charged atoms and freed electrons no longer exhibits its normal dielectric property; indeed, it behaves more as a conductor than an insulator. As such, it readily becomes the precursor to sparking or to damaging arcing. While mild intermittent spark discharges may be self-healing, once the cumulative ionization of the arc develops, service interruption, short circuit and possibly fire tend to follow. The simple setup shown in Fig. 5.11 serves to demonstrate corona. Observe *caution* – stay several feet away from the high voltage portion of the setup.

Often, corona is evidenced by a hissing sound, a faint bluish light and by the odour of ozone ('electrified' oxygen, O_3). All may seem serene and benign in the sense that no crackling of spark discharges is heard or seen. Most often, however, such a 'normal' high voltage manifestation is insidious, portending serious trouble. For one thing, sparking and arcing may just be a matter of a change in humidity or in dust concentration in the air.

The hidden culprit of corona is the *ozone*; it is a chemically active gas with particular affinity for organic substances. It either decomposes or otherwise destroys rubber and many of the commonly-used insulating materials. In so doing, the ozone reverts to oxygen. Also helping to damage the dielectric strength of transformer potting and insulating materials is the ultraviolet radiation which accompanies the visible bluish glow of the corona discharge.

The practical way to reduce corona is to avoid points, sharp edges, abrupt bends and, as much as possible, small diameter wires. At such geometrical shapes there is dense concentration of electric charges which more readily stresses the adjacent air to its ionization level. So-called 'corona balls' are generally brass spheres which, acting oppositely to sharp or pointed conductors, tend to reduce the electrostatic stress on the air so that very little ionization current flows.

Even more care is called for when the transformer handles both high voltage and high frequency. Radio frequencies 'like' to jump and establish

arcs to nearby surfaces. The surfaces need not be metallic; many otherwise insulating materials have sufficient conductivity to distort or interrupt the electric field. Air molecules ionized in a high frequency electric field are imparted with more energy than in a low frequency field. And finally, capacitances that are negligible at low frequencies offer little reactance to the higher frequencies. A favourite 'test' procedure of experimentally inclined radio amateurs has been to bring the sharpened end of a pencil close to the high impedance end of a resonant tank circuit. The resultant arc would then be evaluated in terms of its length, breadth and ferocity. Such a procedure is most certainly not advocated but is mentioned as illustrative of the jumping tendency of the high voltage, high frequency arc.

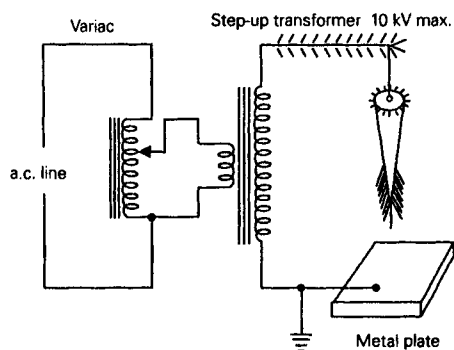


Figure 5.11 Setup for demonstrating conditions favourable to corona. Observation is best made in the dark. Emission of a pale bluish glow is indicative of ionization of air molecules known as corona. Corona consumes electrical power and produces ozone and ultraviolet radiation. It also causes electrical noise which can interfere with communications and electronic equipment. Ionized air, being conductive, is often the prelude to electrical breakdown culminating in sparking and arcing. Note that corona first appears at small diameter wires, sharp edges, abrupt bends and at points as the voltage is raised. (Caution – remain well away from the high voltage.)

Where cost and bulk have low priority, high voltage transformers are best packaged so as to exclude atmospheric air. One technique is oil immersion. Another is the use of a pressurized gas with high dielectric strength. A good vacuum would be suitable but is not generally considered practical in production.

6 Miscellaneous transformer topics

In the process of viewing the topic of transformers from several different angles, one inevitably encounters interesting material which, although difficult to classify, nonetheless bears relevance to the basic theme. The easy way out, admittedly, is to present such material under the 'miscellaneous' heading. Although the discussions in this chapter did not really qualify for inclusion in previous chapters, these miscellaneous discussions will add enlightenment and insights to the overall topic of transformers.

Transformer saturation from geomagnetic storms

Several instances have been discussed in which controlled or intended core saturation of transformers and transformer-like devices have yielded beneficial results. Of perhaps greater concern to not only practitioners of electric/electronic technology but also to the general public, is inadvertent saturation of transformer cores from *geomagnetic storms*, for the possible result of this phenomenon is power outage for protracted periods and over hundreds or thousands of square miles of normally reliable electrification.

The origin of these storms appears to be deep within the interior of the sun, the outward manifestation being the approximately 11-year sun spot cycle. Among other things, the sun spots signify the ejection of the 'solar wind', a heated plasma of charged particles – mostly protons and electrons. An extremely large magnetohydrodynamic generator develops when the solar wind interacts with the earth's magnetic field. Indications of this interaction are the northern and southern lights, or in more technical parlance, the aurora borealis and aurora australis, respectively. The ionized gases of the upper atmosphere glow with the diverse colours of the rainbow and with ever-changing shapes and patterns.

Less beautiful, however, are the effects of induced earth currents comprising up to millions of amperes. The frequency of these currents are very

low, perhaps in the order of 0.003 Hz. This frequency is so low compared to the 50/60 Hz frequency of utility power systems that no discernible error results from viewing the induced earth current as d.c. As such, it has the habit of finding its way into the windings of giant power station transformers where it produces violent half-cycle saturation effects because of magnetic biasing of the core. (It makes no predictable difference which half-cycle is affected, this being determined by the direction of the induced current.) This phenomenon is illustrated in Fig. 6.1.

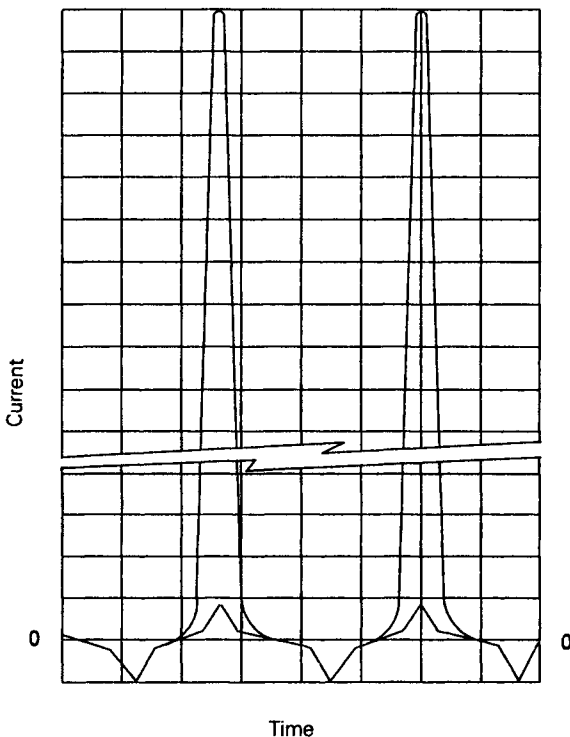


Figure 6.1 Half-cycle core saturation from geomagnetically-induced current. Normally a transformer has a low amplitude exciting current, as shown by the a.c. wave centred on the zero axis. This represents a relatively linear use of the core's magnetic characteristics. However, with injection of a biasing current from a geomagnetic storm, the transformer then becomes extremely non-linear every half-cycle with very large peak currents. The 'broken' portion of the graph drives home the disparity between normal and abnormal operation.

Keep in mind that we are alluding to very large energy systems – utilities involved with the generation and distribution of hundreds of MW. It may not be easy to mentally scale-up the consequences of such core saturation

from the relatively benign malfunction of a small transformer we might be dealing with on an experimental breadboard.

When a giant power transformer is subjected to half-cycle saturation of such violence as depicted in Fig. 6.1, both copper and core losses can produce severe overheating, with secondary damage resulting from stray flux leakage effects. Whether or not the transformer survives this abuse, its electrical malfunctioning can wreak havoc with the utility transmission system. For example, the extreme non-linearity of the half cycle saturation injects powerful odd and even harmonics into the network. These, in turn, cause malfunction of protective relays and disturb the functioning of capacitor banks. At the same time, the affected transformer consumes inductive reactive power in such quantity as to reinforce the other malfunctions.

It would be only natural to ponder why sufficient over-design is not incorporated to enable such transformers to better handle these geomagnetically-induced currents. It turns out that the volume of the steel core needed for such an objective would be neither technically nor economically feasible. It is difficult enough to achieve near-linear operation under normal conditions.

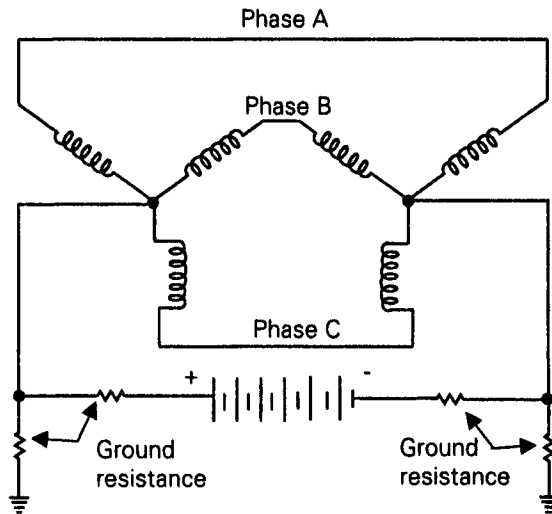


Figure 6.2 The geomagnetic-induced voltage in a three-phase transmission system. For simplicity, only the high voltage windings of the transformers are shown. The battery simulates the almost d.c. voltage gradient induced in the earth from the geomagnetic storm. The resultant biasing current through the transformer windings produces violent half-cycle saturation of the cores. This endangers the transformers and causes malfunctioning of the power network because of abnormally high harmonics and reactive energy.

The detrimental effects of half-cycle saturation are probably even worse than has been described because major transmission networks are involved with three-phase power. It is interesting to contemplate one way the near-d.c. geomagnetic half-cycle saturation can happen in a three-phase power system. Figure 6.2 depicts, in a simplified manner, the situation on an 'as if' basis – it is *as if* a battery were inserted between 'grounded' neutrals of transformers distanced hundreds of miles apart. The two grounds, although represented by the same symbol, may be 'ohmically' far apart, especially where igneous rock terrain is involved. Moreover, the geomagnetic ground current is driven by a voltage gradient which increases with geographic distance. The net result is that there is a more than adequate voltage (represented by the battery) to force biasing current through the transformer windings. In turn, the normally benign exciting current becomes extremely large because of core saturation.

The strange 4.44 factor in the transformer equation

In the design of a conventional transformer, it is usually good practice to make an early determination of the number of volts per turn that will prevail in the tentative design. This makes sense because one usually knows the operational frequency and has made trial assumptions about the core material, core area and the maximum allowable magnetic flux density. Whether these initial assumptions can materialize into a practical transformer will depend upon whether the needed wire size and insulation can be neatly fitted within the window space of available and reasonable core sizes corresponding to the initially-chosen core area. Obviously, this is an art and a science in which the best path to expertise is lots of previous experience.

The equation used in such circumstances is:

$$E/N = 4.44fB_{\max}A \times 10^{-8}$$

where E/N is the number of volts per turn, f is the frequency in Hz, $B_{\max}$ is the maximum flux density and A is the area of the core.

Providing one maintains consistency of units, this provides a straightforward way of starting the design of a transformer. It is only natural to ponder the *origin* of the strange factor, 4.44. The practitioner will no doubt have observed that this factor is simply 4 when the transformer is to be operated in the saturating mode, such as in some saturable core oscillators. Indeed, one can find the implication in the technical literature that the factor is 4.44 in linear transformers, but 4 in saturating transformers. This is *not* altogether true, as will be demonstrated by a look at the basic principle of electromagnetic induction and also by a practical circuit using both saturating and linear transformers. (Incidentally, B_{sat} is generally used to represent $B_{\max}$ in the saturating core transformer.)

A basic equation pertaining to electromagnetic induction is simply:

$$\frac{E_{av}}{N} = \frac{B_{max}A}{t} \times 10^{-8}$$

where E_{av} is *average* induced-voltage per turn, B_{max} is the *peak* magnetic flux-density, A is the core cross-sectional *area*, t is the *time* required for the flux-density to change from zero to B_{max} .

Suppose we have a two-winding transformer and that all of the magnetic flux links both primary and secondary. As is commonly the case, assume a sine wave of frequency f responsible for the flux variation in this transformer. That being the case, t will necessarily be *one-quarter* of a cycle. If we now substitute $\frac{1}{4}f$ for t , the equation takes on the form:

$$\frac{E_{av}}{N} = \frac{B_{max}A}{\frac{1}{4}f} \times 10^{-8}$$

Rearranging this becomes

$$\frac{E_{av}}{N} = 4fB_{max}A \times 10^{-8}$$

which is valid, but not practical because we are accustomed to dealing with the *effective* value of sine-waves, rather than the average value. This requires multiplication of the average value by the conversion factor, 1.11. That is how we end up with:

$$\frac{E_{eff}}{N} = 4.44fB_{max}A \times 10^{-8}$$

but for square waves, average, effective and peak values are the *same* quantity. Therefore, the multiplying factor simply remains 4 and we have:

$$\frac{E_{eff}}{N} = 4fB_{max}A \times 10^{-8}$$

We are now in a good position to make a practical interpretation of our brief mathematical exercise. It is simply that the use of the 4.44 or 4 factor is *not necessarily* determined by whether or not the transformer is a saturating or a linear type but rather by whether the transformer handles a *square wave* or a *sine wave*. This is clearly demonstrated by consideration of transformers T_1 and T_2 in the two-transformer saturable core oscillator shown in Fig. 6.3. Transformer T_1 undergoes hard saturation. In so doing, it initiates the alternate switching action of the transistors and thereby determines the oscillation frequency. Because the switching transistors are alternately on and off, we are dealing with a *square wave*. Thus, the design of transformer T_1 involves the factor 4 in the electromagnetic induction equation. Transformer T_2 ,

however, does *not* saturate. Indeed, one of the salient features of this circuit is that high efficiency is attained by *avoiding* the core losses inevitable in a saturating output transformer.

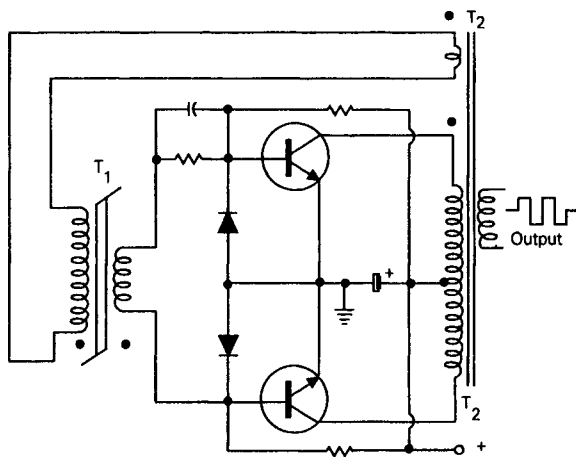


Figure 6.3 *Saturating core oscillator with a linear output transformer. A square voltage waveform is associated with T_2 , the 'linear' output transformer. For that reason, the volts per turn design equation makes use of the factor 4. (The commonly-used factor, 4.44, used in a.c. power work involves sinusoidal voltage waveforms.) Note that it is not saturation per se that governs choice of this factor in the design equation.*

Even though transformer T_2 is a linear type, its design equation would *not* make use of the 4.44 factor commonly used when transformer operation is confined to its relatively linear region. Because T_2 in this circuit operates with a *square wave*, its design would be calculated with the factor 4 in the equation for determining the volts per turn. In the technical literature where the 4.44 factor is associated with linear transformers, the *assumption*, although not always stated, is that *sine wave* operation is involved.

It might seem natural to ponder what the situation would be with a *saturating* transformer hard-driven from a sine wave source. With such a setup, the originally sinusoidal wave would no longer remain so 'inside' the transformer as it is alternately saturated with the changing polarity of the excitation. This would qualify as a square wave situation best served by the factor 4 in the volts per turn equation.

Link coupling – transformer action over a distance

Link coupling is an interesting transformer scheme used in radio frequency systems; in essence, it enables physical separation of two resonated LC 'tank'

circuits. Ordinary electromagnetic induction takes place *twice* in this arrangement. The general setup is shown in Fig. 6.4. It allows transport of high frequency energy at a convenient low impedance level. The inter-connecting line can be a twisted pair, but more advantage is achieved if shielded coaxial cable is used. Quite often the links comprise just a single, or several turns. Ideally, the impedance represented by these links should match the characteristic impedance of the line. In actual practice, one strives for as low an impedance as is consistent with the energy that must be passed from 'primary' to 'secondary'. Performance is not critical, but there are a few rules that should be followed.

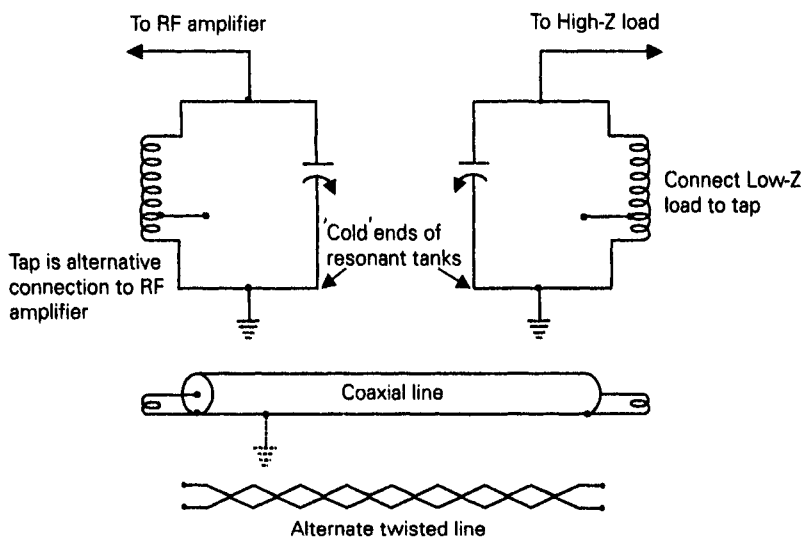


Figure 6.4 The basic link-coupled transformation system. This setup provides transformer action over a distance. One or several turns of relatively heavy conductor generally suffices for the links. The tightness of coupling is essentially governed by the physical separation between the tank circuits and the links and is best optimized experimentally. Sometimes a variable capacitor associated with the links is used to compensate reactance from non-ideal conditions such as impedance mismatches. Links can be wound in either direction.

The physical location of the links should be at the RF 'cold' end or portion of the resonant tanks. This minimizes capacitive coupling and detuning effects. Grounding of the coaxial shield or of one of the twisted wire conductors should be investigated empirically. In most cases, such grounding tends to reduce radiation from the line but sometimes the line is better left ungrounded. There is also the question of where to introduce the ground(s). Most often, the most effective connection point for the ground is close to the

transmitter end of the system. The degree of coupling is a function of the number of turns on the links as well as their physical proximity to the resonant circuits. With a little experimentation, one can obtain both high energy transfer and high Q resonances. By transposing connections to one of the links, a more favourable phase situation can be obtained for neutralization purposes, or for such loads as antennas.

This scheme works very well for air core inductors, but is also implementable with toroids or with other inductors using magnetic cores. Increasing the diameter of the links is a practical way of reducing coupling if this is needed. The easiest link-coupled system to get working properly is one that is symmetrical with the same resonant circuits and link geometries at both ends. This is not a requisite, however.

Homing in on elusive secondary voltage

After the need for the magnetic core in all but high frequency transformers is understood, another enigma confronts students and often practitioners reasonably versed in electrical technology. This has to do with the matter of fractional turns and is best illustrated with a typical example. Suppose the value of two turns per volt applies to the design of a simple power transformer in which the intent is to step down 120 V to 5.75 V. It is noted that a 12-turn secondary winding will provide 6 V and that 11 turns will produce 5.5 V. From a purely mathematical viewpoint, our course appears straightforward enough – an $11\frac{1}{2}$ -turn secondary should meet our goal for the desired output of 5.75 V.

Actually, this will not work. Only increments of whole numbers of turns are permissible. Fractional turns cannot be implemented in practice to change induced voltage. Some insight into the nature of this paradox may be gleaned from the simple toroidal transformers shown in Fig. 6.5. Note that ‘one turn’s worth’ of voltage is obtained from *both* secondary arrangements. This shows that only the segment of the secondary which passes *through* the toroid is active in the electromagnetic induction of voltage. The remainder of the turn ‘just comes along for the ride’ – even worse, it contributes to the I^2R loss in the winding. This same situation prevails in core and shell type transformers built up of laminations.

Fortunately, it is possible to closely approach such a desired secondary voltage by modifying the turns on the *primary* winding. In the above example, we can work with a 12-turn secondary. In order to obtain 5.75 V, rather than 6 V from this secondary, the step-down ratio of the transformer must be *increased* by the factor $6/5.75$ or 1.043. To accomplish this, the 2×120 , or 240-turn primary must be increased by 1.043, yielding 250.3 turns. Dispensing with the fractional turn, the new primary must have 250 turns. The new step-down turns ratio is $250/12$, or 20.83. Then $120/20.83$ yields the new

secondary output of 5.76 V. (The fact that our modified transformer is no longer based on exactly two turns per volt is of negligible practical consequence.)

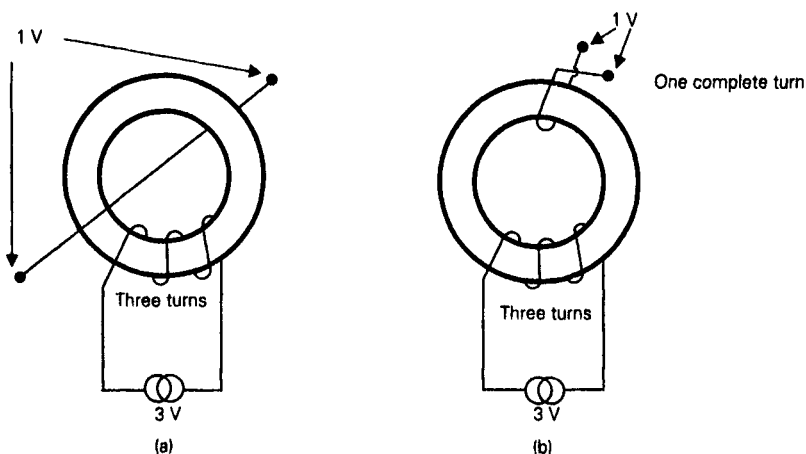


Figure 6.5 Demonstration that voltage adjustment cannot be made with fractional turns. (a) The secondary is a single conductor that passes through the toroidal core. (b) The secondary is a complete turn. The induced voltage in these two secondary 'windings' is the same. This shows that the added turn segments in (b) do not take part in transformer action. Therefore, such fractions of a turn are not useful for manipulating the secondary voltage.

An alternative approach naturally suggests itself. Let us try to develop the desired 5.75 V by starting with an 11-turn secondary winding which provides only 5.5 V. This time the number of turns on the primary must be *decreased* by the factor $5.75/6$, or 0.958. Thus, $240 \times 0.958 = 230$ 'practical' turns on the new primary. The new step-down transformation ratio is $230/11$, or 20.91. Finally, the new secondary voltage is $120/20.91$, or 5.74 V.

A possible advantage of the second approach is that no splices need be made in the primary winding. It may be argued on the other hand that it tends to be better to have a surplus rather than a deficit of turns on the primary. Because of the small percentage change in either case, the technique selected may be said to be a toss-up. Both methods suffer in absolute accuracy because of the necessary rounding-off of the calculated turns. Transformers with different turns ratios and other design parameters can likewise be made to home-in on a desired secondary voltage. Whenever this cannot be accomplished directly by altering the secondary turns, the basic approach is to set up a more favourable transformation ratio by modifying the *primary* winding. Sometimes a combination of changes in *both* windings yields the best results.

Keep in mind that voltage is induced in the segments of turns that are encircled by magnetic flux. Therefore, the accurate way to count turns is by summing up the *inner* conductor segments. Outer portions of turns are physically necessary but they are not, however, active participants in electromagnetic induction. (If you can successfully eliminate these outer segments, rush immediately to the patent office.) Interestingly, the commonly encountered current transformer represented by Fig. 6.5(a) does indeed operate without any needless turn segment encumbering the straight-through secondary conductor. Here, we have the exception that proves the rule, although the purist may argue that the turn is completed by the external circuit.

The example dealt with involved just 0.25 V one way or the other with regard to the choice of secondary turns. Some might say that the sought-after precision is not generally worth the trouble in practical situations. The transformer in the example was a relatively small one, of the order of a couple kVA. Let us take a look at a larger transformer, with a rating of about 4 kVA. Here, we would find that a typical design would call for two volts per turn, rather than the two turns per volt pertaining to the smaller transformer. This time, our prospects of obtaining the proper number of secondary turns for a desired output voltage are four times worse than with the smaller transformer.

With even larger transformers, the volts per turn needed continues to increase. Thus, it continues to become increasingly difficult to come up with a satisfactory design based upon a selected number of secondary turns. It then becomes increasingly necessary to experiment with the *combination* of primary and secondary turns which best approaches the secondary output voltage required. Such experimentation may be mental, physical or a combination of efforts. In any event, there is usually enough leeway allowable to modify an initially chosen turns per volt or volts per turn ratio in order to arrive at a more suitable transformation ratio.

The bottom line is that it is unfortunate that the secondary voltage is not amenable to trimming by means of fractional secondary turns. Were it possible, that would certainly be the quick and easy approach. However, by using the techniques discussed, it is very nearly possible to obtain any desired secondary voltage.

An apparent paradox of electromagnetic induction

Valuable insights often result from the resolution of paradoxes. In technology, one sometimes encounters useful deployment of cause and effect relationships that seemingly violate our notion of the possible. Consider, for example, a salient feature associated with toroidal transformers and inductors. This pertains to the negligible external electromagnetic fields

surrounding these components. Because of this, practical circuit boards can be made with a number of such devices in close physical proximity to one another with little danger of undesirable transformer action or mutual coupling. This is especially important in electronic circuits where such stray couplings can lead to unwanted oscillation or other operational disturbances. The same is true of electric-wave filters where the toroidal transformers and inductors can be tightly packed together with little alteration of frequency response.

The confined fields of these toroidal components is not surprising in the light of our previous discussion in which it was pointed out that only the *inner* segments of the windings actively participate in electromagnetic induction. The outer segments of the turns justify their existence simply as connecting links to the successive inner segments. It should additionally be pointed out that the magnetic core of the toroid offers much less reluctance to the magnetic flux – so why should it occupy much of the air *external* to the toroidal core?

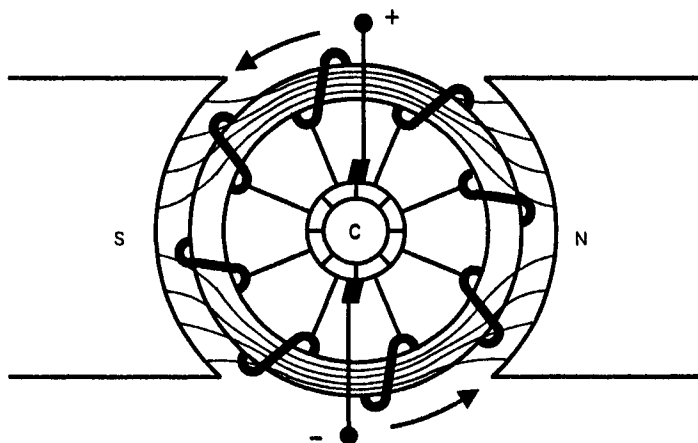


Figure 6.6 Electromagnetically active outer conductors in toroidal winding. In the Gramme ring armature, a voltage is generated due to the cutting of magnetic flux by the outer wires of the rotating ring. This is contrary to the behaviour of toroidal inductors and transformers where it is the inner segments of the turns that are electromagnetically active. The apparent paradox stems from the mechanical motion in the dynamo. Note that the inner wires of the winding are magnetically shielded in the dynamo.

This is where we find an apparant paradox: the now-obsolete Gramme ring electric dynamos used a rotating toroidal winding with an iron core. When used as a generator, this machine depended for its induced voltage on the *outer* segments of the toroidal winding cutting magnetic flux from

pole pieces, as illustrated in Fig. 6.6. Indeed, writers of the time lament the fact that the *inner* segments of the toroidal windings, being shielded from the magnetic flux, do *not* participate in electromagnetic induction. This scenario appears contradictory to what was previously mentioned regarding the confined fields of toroidal transformers and inductors. Happily, the apparent paradox is resolved by pointing out that the induction in the Gramme ring machine is *forced* because of the *mechanical motion* involved.

Transformer technique with 'shorts' – the induction regulator

A clever implementation of transformer action is encountered in the *induction regulator*. This device automatically improves the voltage regulation of single-phase feeder lines. Actually, it is comprised of a built-in system with a voltage sensor and an electric motor for turning the shaft of the induction generator proper. If one did not know otherwise, the rotatable induction generator itself might be mistaken for a single-phase induction motor insofar as concerns its construction. (Here, it is worth mentioning that no tendency for motor action exists – it will be recalled that an actual single-phase induction motor develops no starting-torque on its own; rather some 'artificial' means is needed to get such a motor going.)

The schematic diagrams of the induction generator are shown in Fig. 6.7. Note the connections to the a.c. line. The operating principle is simple enough; the rotatable primary can be turned so that the induced secondary voltage can either series-aid or series-buck the line voltage. Any degree of aiding or bucking the line voltage over a wide range can be obtained smoothly by positioning the primary appropriately. In the overall system, a line voltage sensor, such as a tapped voltmeter, signals the rotational direction of a motor needed to regulate the line voltage. The motor is coupled to the induction regulator shaft in such a way as to turn it one way or the other through a 180° arc.

This is straightforward enough, but the 90° position of the primary cannot correspond to a neutral position in the above discussion. This is because of the inductive reactance of the secondary winding which is in series with the a.c. line. However, if a *tertiary* short-circuited winding is quadrature-mounted on the rotatable primary, the actual secondary will 'see' *this* winding as a shorted 'secondary' when the primary is in its geometrically-neutral, or 90° position. This will, then, reduce the inductive reactance of the actual secondary to more or less zero. At other angular positions of the rotatable primary, the coupling of the shorted tertiary to the actual secondary winding will progressively diminish to zero at the maximum aid or buck positions.

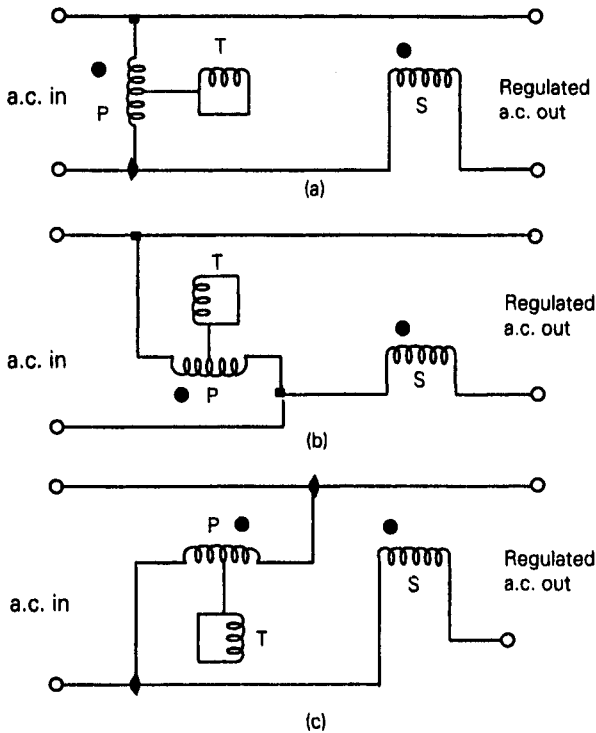


Figure 6.7 Transformer action of the induction regulator. (a) The rotatable primary winding P is positioned perpendicular to the series-connected secondary S. Coupling between these two windings is zero. However, the shorted tertiary winding T, which is mechanically linked to the primary, couples directly to the secondary. This very nearly reduces the secondary inductive reactance to zero. (b) The primary has been rotated 90° counter-clockwise. Mutual inductance between primary and secondary is additive, thereby boosting the line voltage. The tertiary winding is physically and electromagnetically out of the way. (c) The primary has been rotated 90° clockwise from its neutral position at (a). Mutual inductance between primary and secondary is now subtractive, thereby bucking the line voltage. Again, the shorted tertiary is decoupled from the secondary.

Constant current transformers incorporating physical motion

Another type of current-regulating transformer is shown in Fig. 6.8. Here, one winding, often the secondary, is free to move up and down. As a manifestation of Lenz's law, there is a force of repulsion between the windings

that increases with load current. The actual physical displacement of the moving winding is a balance between load current and weight. The greater the gap between primary and secondary, the greater is the *leakage inductance* associated with these windings. Current regulation takes place because higher leakage inductance results in lower current availability from the secondary. A dashpot dampens the tendency for undesirable mechanical oscillation. A nice feature of this scheme is that various levels of constant current can be conveniently selected by changing the *weight*. Load current is directly stabilized against load resistance and indirectly stabilized against line voltage. Good constancy of load current results.

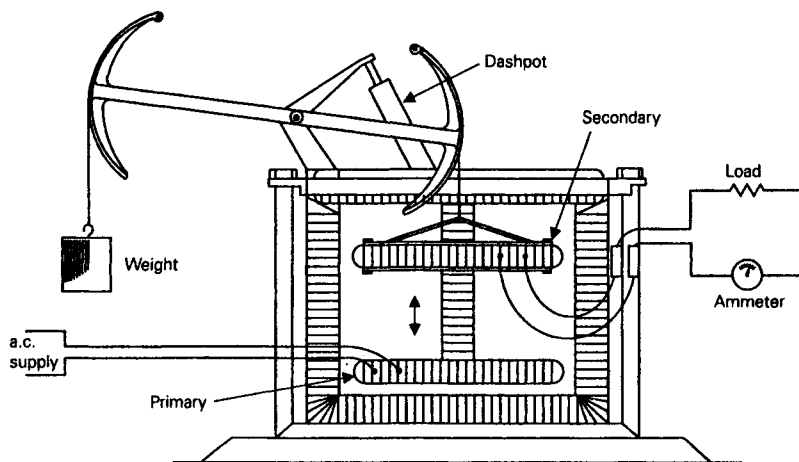


Figure 6.8 *Dynamic current transformer using physical motion for regulation. The electromagnetic repulsion between the primary and secondary windings increases with current but is balanced by the weight. The net effect is automatic regulation of load current.*

This type of current-regulating transformer has found its greatest use with street lighting systems. When the lamps used in such networks require d.c., a rectifier can be used without affecting the basic operation of the current-regulating transformer. Although this constant current technique is straightforward and reliable, it has an inherent shortcoming. Because of the purposefully high leakage inductance, the power factor is inordinately low. In large systems, this leads to high installation and operating costs; it can also upset line voltages.

An interesting aspect of this motion dependent transformer is the simple fact that the same mechanical forces exist in *all* transformers. The small transformers used in electronic technology develop forces of little consequence. In larger transformers used in utility systems, dedicated design effort is needed to counteract these forces. Although entire windings may

have little motional freedom, individual conductors are subject to displacement and can produce maintenance problems from short-circuits and arcing. Aggravating the situation is the fact that a line short-circuit current may be 50 times normal full-load current.

A novel way to null transformer action between adjacent tuned circuits

An interesting transformer technique, one of potential use in today's technology, relates to the nulling principle employed in the *neutrodyne* radio receivers of the mid-1920s. The objective was to discourage, rather than enhance, electromagnetic coupling between physically close inductances. The manner in which this was accomplished was simple in implementation, but not intuitively obvious. Indeed, one cannot help but see a paradox in the fact that the scheme had for years been overlooked by engineers, radio amateurs and experimenters; even now, it is not easy to find an engineer with a quick grasp of the true nature of the phenomenon.

The neutrodyne set overcame one of the irritating bugs that plagued the otherwise satisfactory radio circuits of the time – two stages of tuned radio – frequency amplification, a tuned detector and one or two stages of audio amplification. The basic problem was *oscillation* in the radio frequency stages. Neutralization of the tube capacitances was only partially effective because of the physical proximity of the resonant circuits. Put simply, the traditional neutralizing techniques could not cancel the electromagnetic coupling between the coils. Most attempted remedies were frequency sensitive, making it difficult for the non-technical enthusiast to properly tune in a station.

To circumvent this *unwanted* transformer action, some manufacturers resorted to shielding. This, however, tended to detune or degrade the resonant circuits. Also, the copper boxes were costly to fabricate and install. Another scheme involved mounting the three coils involved mutually perpendicular to one another. This worked to an extent, but proved too difficult for mass production. The culprit was probably stray coupling from capacitance and radiation. What ultimately proved practical and 'idiot-proof' was traditional neutralization of tube capacitance in conjunction with the neutrodyne technique of nulling transformer action between adjacent coils. Let us see how this was done.

The secret of the neutrodyne radio was the *positioning* of the three coils in the radio frequency part of the set. These were mounted at a 'magic' angle, but not in the older pattern of the mutually perpendicular layout. Rather, the coils in the neutrodyne set all had the same angular inclination with respect to the horizontal plane of the breadboard. Explanation is simplified if we consider just two of the coils. In Fig. 6.9(a), two such coils are depicted.

Note, also, the alluded 'magic' angle X . Actually, this angle was 54.7° . Although it is intriguing that the tangent of this angle is the square root of 2, the mathematical proof of the zero coupling behaviour is quite complex; fortunately, it is not necessary. A simple strategem allows a satisfactory qualitative explanation for practical purposes.

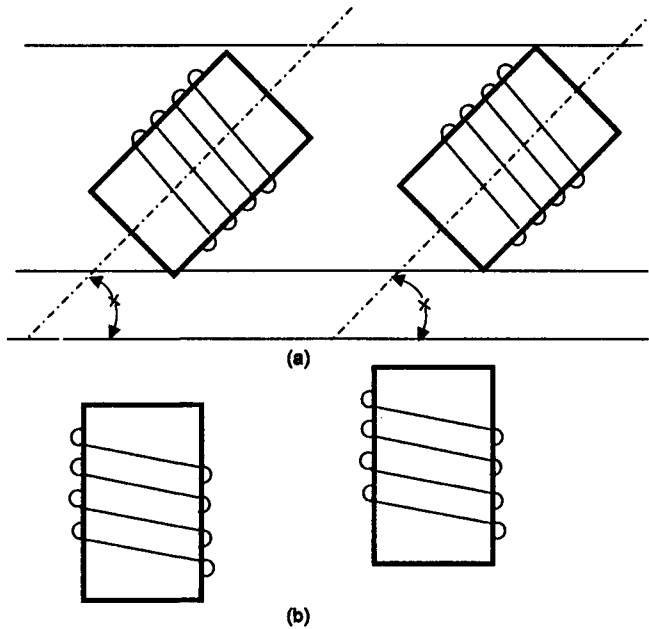


Figure 6.9 Neutrodyne coil orientation and an equivalent layout. The situation shown in (b) is obtained by rotating the coils depicted in (a) counter-clockwise by the angle ' X '. In so doing, the electrical relationship remains undisturbed and it is no longer necessary to deal with angle ' X '. The objective of this strategem is to simplify both experimentation and analysis. This will be made evident by the fact that the zero-coupling null can be 'detuned' by sliding either coil in (b) slightly one way or the other along its longitudinal axis.

In Fig. 6.9(b), the two coils are given a graphical twist by the angle X so that they are not slanted to one's view. Note that the two coils retain their geometrical relationship to one another. We have the right to expect zero transformer action between paired coils in both Fig. 6.9(a) and (b). The nice thing about the situation as shown in Fig. 6.9(b) is that we can go in and out of the nulling phenomenon by merely sliding one of the coils back and forth along its longitudinal axis. This enable us to dispense with consideration of angle X .

The 'common-sense' approach to the situation in Fig. 6.9(b) is that there

must be some opportunity for electromagnetic transfer of energy between the coils. One would, of course, suspect less than maximum coupling, but one doesn't readily see this as a condition of *zero* mutual inductance. In the bygone days of 'wireless', *vario-couplers* of this nature were widely used and energy transfer went from maximum to minimum as physical separation of the coils increased. To have a nulling action take place at some *intermediate* position of separation taxes one's sense of logic, to be sure. The practical way to resolve this dilemma is to set up a very simple experiment with two resonant circuits, a dual trace oscilloscope and an RF oscillator.

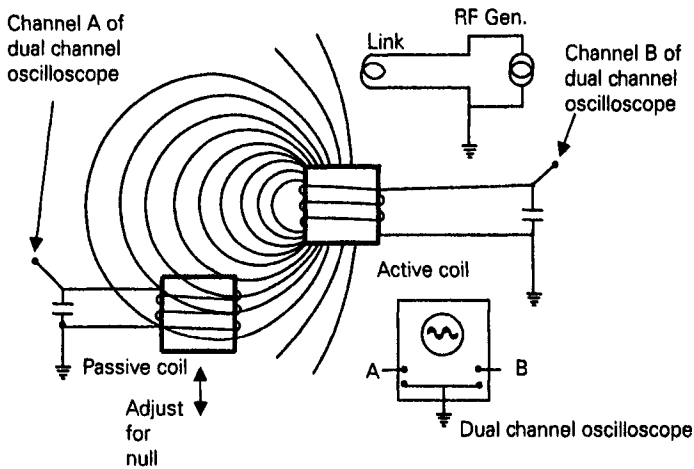


Figure 6.10 Experimental setup for demonstrating nulling of transformer coupling. This arrangement duplicates the coil positional scheme shown in Fig. 6.9(b). For practical purposes, it is electrically equivalent to the slant-mounted coils originally used in the neutrodyne radio. RF energy is imparted to the 'active' coil via the single or several-turn link. This link is located at the far-end of the active coil so as to have negligible coupling to the passive coil. Both coils are resonated to the same frequency.

The demonstration setup shown in Fig. 6.10 enables observation of the neutrodyne nulling phenomenon. All that is needed in the way of instrumentation are 'garden varieties' of an RF generator and a dual channel oscilloscope. Otherwise, the setup is just an extension of Fig. 6.9(b). Note that the coils are resonated to the same frequency. This doesn't have to be precise – assuming identical coils, simply shunt them with the same value of 10% capacitors. If the coils are from the input section of an AM radio, capacitors can be in the $0.001\ \mu\text{F}$ region. The coils will then be tuned below the AM broadcast band and powerful local AM signals will not be likely to sneak in and interfere with the oscilloscopic waveforms.

Once the link-coupled RF signal is tuned to the vicinity of the resonant coils, it will be found that the induced voltage in the passive coil can be zeroed by carefully sliding this coil along its longitudinal axis. At null, the physical relationship of the two coils will resemble the situation shown in Fig. 6.9. If one wishes, the two coils can then be rotated to duplicate the slanted-layout shown in Fig. 6.9(b). Measurement of angle X will then be found to be close to 55° .

The best results will be obtained with 'square' coils in which the diameter and the winding length are the same. This happens to be the geometry corresponding to optimum Q, but such coils are more likely to come out of older rather than modern, radios. Optimum spacing of the coils is from one radius to one diameter. Spacing too closely adversely affects frequency independence because of capacitive coupling. Too much separation makes the null adjustment too critical and the demonstration may be impaired by noise and RFI.

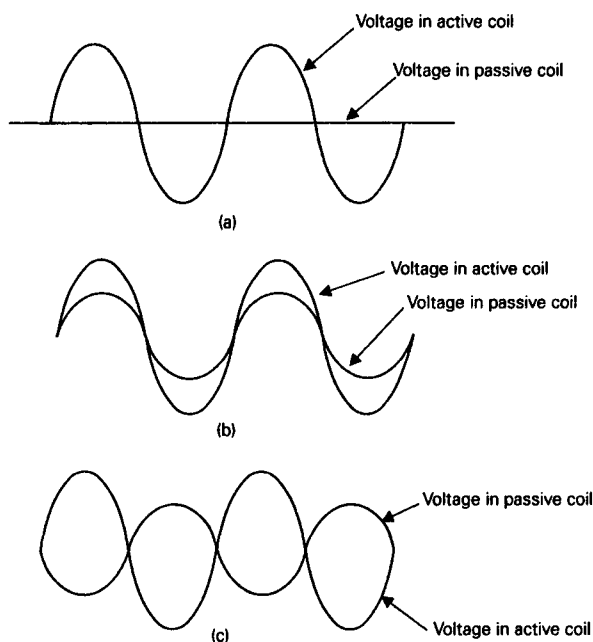


Figure 6.11 The electromagnetic nulling demonstrated in the setup of Fig. 6.10. (a) Critical position of the coils – no transformer action is evident as zero voltage is monitored in the passive coil. (b) and/or (c) The nulling action displayed in (a) has been upset by moving the passive coil one way or the other along its longitudinal axis.

The nulling phenomenon that will be observed is shown in Fig. 6.11. Note the 180° phase displacement of the induced signal on the two sides of null.

The null results from cancellation of oppositely polarized voltages induced in the passive coil, resulting in the *net* zero voltage as seen in Fig. 6.11 (a).

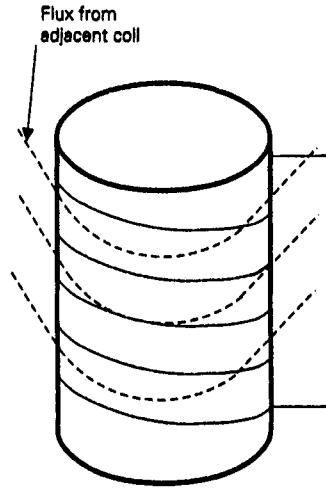


Figure 6.12 *Flux-cutting in opposite directions induces no net voltage in coil. This magnified version of the passive coil in Fig. 6.10 shows the situation at null. Both downward and upward inclined flux cuts the coil conductors, thereby inducing oppositely polarized EMFs which cancel one another. For practical purposes, this manifests itself as zero electromagnetic coupling.*

The nulling process is further illuminated in Fig. 6.12. This is essentially a magnified view of the passive coil in the experimental setup shown in Fig. 6.10. It can be clearly seen that the coil conductors are cut by both downward and upward sloping magnetic flux lines. Thus, at null there is cancellation of the two oppositely polarized induced EMFs in the coil and there is no *net* transformer action. Although the purist might object to the statement that no mutual inductance exists at null, most practical applications of the phenomenon would 'see' a condition of electromagnetic isolation between pairs of coils deployed in this manner.

Despite the RF amplification in the neutrodyne radio set, the cancellation of transformer coupling between the two end-coils, as well as between adjacent coils probably also contributed to the set's stability over the entire broadcast band. The experimental setup shown in Fig. 6.10 can be used to confirm the null's independence of frequency by changing the frequency of the RF generator. It will be seen that off-resonant operation affects amplitude much more than phase, providing that the two coils themselves have the same resonant frequency. Even when this is not the case (as in adjustment of the three tuning dials on the radio), transformer coupling between coils remains greatly reduced.

Transformers only do what comes naturally

An elusive failure mode once plagued driven inverters usually culminating in the destruction of one or both of the power transistors. It was realized at an early stage that the linear output transformer was somehow involved. However, over-design and greater safety factors did not appear to remedy the situation. After costly and embarrassing experiences in the field, the destruct phenomenon was diagnosed as the tendency for the transformer to 'walk' up its hysteresis loop. Devoid of hairy mathematics, what was happening was that the flux-density operating point of the transformer would depart from its symmetrical position between the core saturation regions and progressively move towards one or the other of these saturation regions. As this went on, one of the transistors was subjected to ever increasing current until catastrophic damage set in. The whole scenario tended to be regenerative – once under way, the transformer core would be 'walked' into one of its saturation regions, as shown in Fig. 6.13.

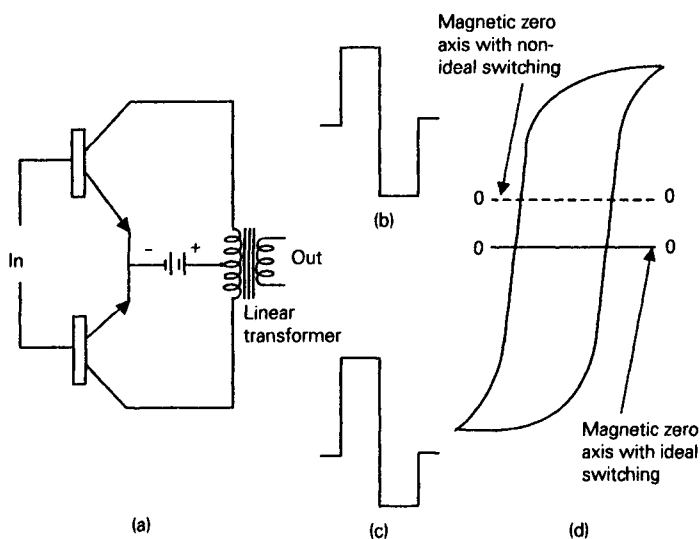


Figure 6.13 Transformers seek volt-second equilibrium and may 'walk' in pursuit thereof. (a) Simplified circuit of a driven inverter. The output transformer is supposed to operate in its linear region. Collector currents are moderate. (b) Ideal switching waveform – alternate half-cycles have equal volt-seconds. (c) Non-ideal switching – inequality in the volt-second areas of alternate half-cycles. (Usually from unbalanced transistors.) (d) The zero axis of the transformer core's hysteresis loop 'walks' in quest of volt-second equilibrium. In so doing, the magnetic saturation region may be approached or reached. This results in abnormally high collector current, endangering the transistor(s).

We know that transformers do not transfer sustained direct currents. The transformer preserves such behaviour by equalizing the volt-second product of successive half-cycles. Normally, we take this for granted. However, in a driven inverter with non-ideal transistor switches, the transformer, although accurately centre-tapped, may not experience equal volt-second half-cycles. Its attempt at correction of this 'forbidden' operating mode is to readjust the 'centre of gravity' of its flux excursions. In doing so, it 'walks' up its hysteresis curve. Soon it gets into one of its saturation regions and we have the 'mysterious' failure mode. (Keep in mind that although the output transformer of a driven inverter handles square waves, it is supposed to be operated well within the near-linear magnetic region of its core.)

The true culprits of such malperformance are generally transistors with excessively sloppy switching characteristics. Unless they are well-matched, they will present the transformer with unequal volt-second pulses during successive half-cycles. Tolerance to such non-ideal switching can be increased by an air gap in the core.

The flux gate magnetometer

The flux gate magnetometer is an interesting transformer application in the field of electronic instrumentation. It can serve as a compass or as a very sensitive detector of small magnetic anomalies. It finds extensive use in oil exploration and in archaeology. The basic transformer nature of this device is evident from Fig. 6.14. Here, we have a primary winding, the drive coil, on a high-permeability toroidal core. There are two pairs of secondary windings orthogonally-spaced around the core. The operation stems from the unique way in which this transformer scheme is used.

The excitation of the drive coil is intense enough to *symmetrically* saturate the core in both magnetic polarities. Under this condition, the induced secondary voltages are rich in *odd* harmonics. However, no *even* harmonics are developed. To understand how an external magnetic field can upset this situation, we can focus attention on just one pair of secondary windings. The basic idea is that an external magnetic field will impose a slight magnetic bias on the core and will thereby upset the symmetry of the hysteresis loop. This is tantamount to a vertical shifting of the zero axis of the hysteresis loop and will manifest itself by induction of even harmonics in the secondaries. The strongest of these even harmonics will be the second harmonic of the drive coil frequency.

Other things being equal, a high operating frequency contributes to sensitivity; several hundred hertz can suffice in a practical instrument. A good sine wave is preferable for the drive, even though the core saturation will distort the resultant flux. The possible drawback of applying a square wave is the imposition of an extra filtering burden on the logic circuits that work

from the secondary outputs. The reason for the two pairs of secondary windings is provision of unambiguous directional data when the device is used as a compass. Because the sensed field has a sine and a cosine component associated with its direction, a suitable logic system can resolve such information for analogue display on a 360° chart, as in an ordinary compass. A digital readout can also be implemented from this data.

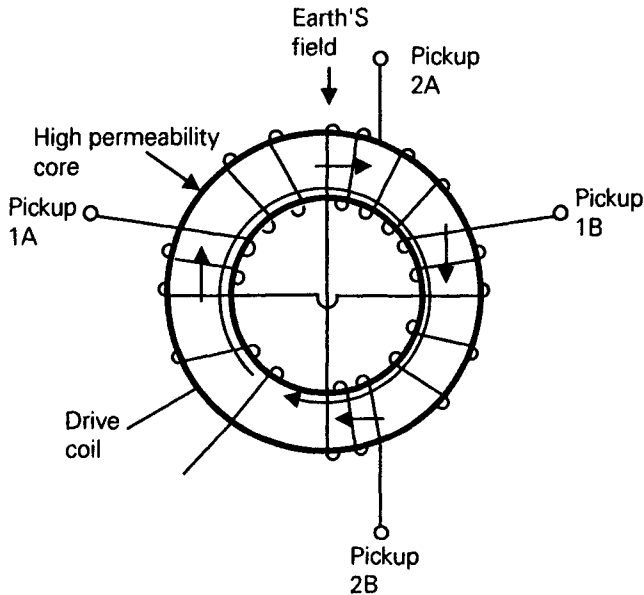


Figure 6.14 *The flux gate magnetometer. The a.c. voltage applied to the drive coil saturates the high permeability core equally for both magnetic polarities. No even harmonics are induced in the paired pickup coils; the symmetrical non-linearity develops only odd harmonics. The introduction of an external magnetic field then alters the symmetry of the magnetization (hysteresis) curve of the toroid. This causes second harmonic signals to be available from the pickup coil's terminals.*

Transformer balancing acts – the common-mode choke

Common-mode chokes are 1:1 transformers which can serve several useful purposes when properly phased and deployed. In a grounded-grid RF amplifier, the use of such a device enables convenient injection of the drive power, keeps RF from the filament transformer and prevents RF current from circulating through the filament where it could raise the temperature above safe limits. A simplified circuit of a grounded-grid stage is shown in

Fig. 6.15(a). Note the phasing dots alongside the common-mode choke. The windings are of bifilar format, or at least are physically very close to one another. Because of the phasing, the 50/60 Hz filament current experiences no drop from inductive reactance. Indeed, mutual flux is cancelled and the core is not subject to magnetic saturation. From a practical viewpoint, virtually no change in filament current would occur if the magnetic core of the common-mode choke were withdrawn or if the two windings were short-circuited.

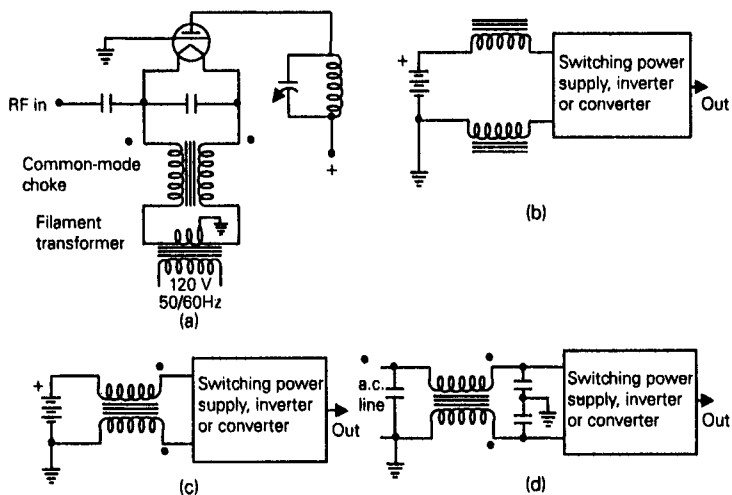


Figure 6.15 Common-mode choke applications. This transformer-like device provides unique response to differential and common-mode currents. (a) Common-mode choke in grounded-grid RF amplifier. Filament current 'sees' negligible inductive reactance and causes no core-saturation. Incoming RF is isolated from the filament transformer by choke action. (b) Brute-force method of keeping switching transients out of the d.c. supply. The individual inductors must be physically large to avoid saturation. (c) Using a common-mode choke to attain the goal set forth in (b). There is no core-saturation and a relatively-small device can be used. (d) Similar setup as (c) operates from an a.c. source. Such common-mode choke circuits help prevent noise contamination of the a.c. power line.

With regard to the incoming RF, a different situation prevails. Here, a high inductive reactance between the filament and the filament transformer is needed. Such a requirement is met because the RF voltage 'sees' the choke in a different manner from that of the filament voltage. The two windings are essentially in parallel for RF current and the core helps develop a high inductive reactance at these high frequencies. Stated another way, the

transformer device does not participate in common-mode operation for the RF drive – there is no cancellation of mutual flux in the core; rather, the RF is isolated from the filament transformer by the full inductance presented by these effectively parallel windings.

The simple open core construction of this type of common-mode choke is explained by the fact that its magnetic circuit bears little relevance to its operation at 50/60 Hz. On the contrary, at radio frequencies where high inductive reactance is needed, the open core is entirely adequate. Note the capacitor across the filament-end of the windings; this ensures that the same RF voltage is applied at these circuit junctions.

Many electronic devices have a tendency to contaminate the a.c. power line with switching transients, noise spikes and various high frequency disturbances. Although the 50/60 Hz power has always arrived as a more or less distorted sine wave with all manner of ‘piggy back’ impulses, the situation rapidly began to deteriorate once solid-state power devices became popular. It has been widely recognized that the remedy must begin with *prevention* of the disturbances from getting into the a.c. line. To this end, an extensive technology based on isolating, filtering and shielding the EMI and RFI has developed. Many situations require more than a bit of empirical work because of various combinations of conduction, ground currents, electrical and magnetic fields, and radiation. In any event, it is often found that the common-mode choke plays a significant role.

The simplified circuit shown in Fig. 6.15(b) depicts the use of individual inductors inserted to attenuate noise transients that would otherwise find their way into the power source – a d.c. supply in this case. A shortcoming of this ‘brute force’ technique is that the inductors have to be physically large in order to avoid magnetic saturation of their cores. A more elegant approach is realized through the use of a common-mode choke, as shown in Fig. 6.15(c). With the phasing as indicated, ampere turns of the two windings balance out so that there is no net mutual flux from the d.c. At the same time, *unbalanced* currents are attenuated by the inductive reactance of the windings. In this way, ground-path a.c. currents are stopped in their tracks.

Essentially the same scheme operates similarly when working from the a.c. line. This is shown in Fig. 6.15(d). Enhanced suppression of the noise voltages results from the use of capacitors as shown. Similar techniques can be used at the *output* of noise-generating equipment. Excessive inductance of the windings can prove counter-productive at high frequencies because of series resonances and bypass action from the distributed capacitance which inevitably accompanies added turns. The desired choke action must be available in the region of the frequencies to be suppressed.

Mutual inductance, coefficient of coupling and leakage inductance

Mutual inductance M can be determined from inductance measurements of the primary and secondary windings L_1 and L_2 . Make a second pair of inductance measurements with L_1 and L_2 first series-connected one way, then series-connected with the opposite phase relationship. Let the larger of these measurements be L_3 and let the smaller of these measurements be L_4 . With these easily-measured data, either of two formulae may be used to calculate M ;

$$M = \frac{L_3 - L_2 - L_1}{2}$$

or

$$M = \frac{L_2 + L_1 - L_4}{2}$$

As a check, both equations should yield the same value of M (refer to Fig. 6.16(a)). Also, for conventional iron core power transformers, M will be found to be very-close to its maximum possible value

$$\sqrt{L_1 L_2}$$

Once M has been determined, the coefficient of coupling k is readily calculated from the expression

$$k = \frac{M}{\sqrt{L_1 L_2}}$$

For iron core power transformers, k tends to be close to unity. For air core transformers k can be a small fraction. Here, k can be understood to represent the *ratio* of the actual mutual inductance to the maximum possible value which would exist for very tight coupling. k is a useful parameter in high frequency air core transformers where the objective often is *not* maximum energy transfer.

Note that the *leakage inductance* of a transformer is zero for $k = 1$ but becomes progressively greater for smaller values of k .

To the designer of power transformers, *leakage inductance* should be as low as possible because of its adverse effect on the load voltage regulation. Because of the near-success in attaining this objective, it turns out that accurate measurements of leakage inductance are difficult to make. Fortunately, this is not often a problem for those who work with already-designed transformers. Thus, for many practical purposes, the assumption of zero leakage inductance is not likely to lead one astray. In any event, it is unlikely that much can be done about the residual leakage inductance that may assert itself in a manufactured transformer.

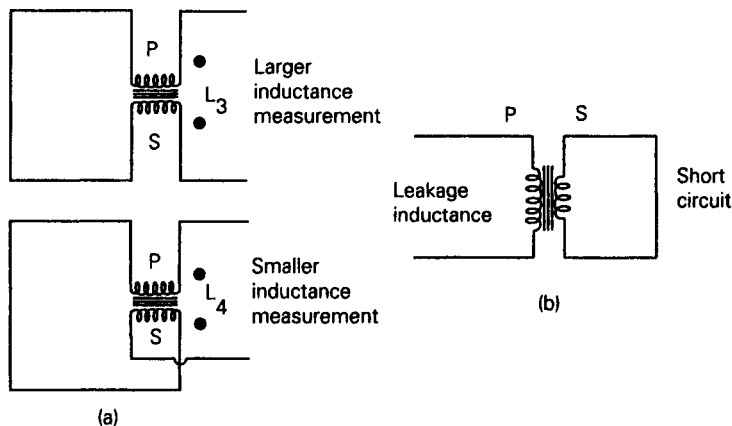


Figure 6.16 Simple test procedures for transformer parameters. Although magnetic cores are shown, the illustrated test setups tend to yield more meaningful results with air core transformers. This is particularly true with iron core power transformers where leakage inductance can be negligibly low and the coefficient of coupling is very close to one. In contrast, air core types often use loose coupling as do some specialized iron core transformers. (a) Transformer connections for obtaining L_3 and/or L_4 . This is an intermediate procedure for the determination of mutual inductance M , as explained in the text. (b) Setup for obtaining approximate leakage inductance. This technique is very useful on a relative basis in which the inductance measurement is repeated as the transformer design is changed.

In inverters and saturable core oscillators, leakage inductance causes generation of spikes which can destroy the transistors. Here, we have a different situation in the pursuit of low leakage inductance because practitioners other than transformer specialists get involved in the empirical work required. Usually, considerable experimentation is needed with regard to core material, core size and shape, winding format and wire size; at the same time, various circuit and snubber arrangements are likely to enter the picture. From the practical standpoint, what is needed is a simple test procedure for monitoring the *relative* levels of leakage inductance as experimentation progresses. Fig. 6.16(b), together with the following discussion, depicts such a test. This may be brought about by measurements of the inductance of the primary winding with the secondary short-circuited. Despite several assumptions and approximations which enter into the determination of leakage inductance in this manner, the procedure enables one to 'home in' on an optimum design format. As helpful as this method tends to be, the not easily predicted side-effects of core saturation during actual operation can sometimes still pose problems with spike generation. It is at this stage that energy-absorbing networks (snubbing circuits) become a

better investment of effort. In summary, the less inductance that can be 'seen' in one winding with the other shorted, the less leakage inductance can be assumed to reside in the transformer.

Short-circuit protection of power lines via a unique transformer

Transformers also have their applications in utility power systems. It is well known that a malperformance greatly feared is a sudden short circuit; much damage and disturbance tends to result from the tremendous surge of over current following such a fault. Conventional protection has generally incorporated a sensing circuit working in conjunction with a mechanically-operated method of inserting current-limiting impedance in the line. The implementation of such protection has been both a science and an art, with the never-quite-attainable goal of fast action, reliability and convenient restoration to normal service.

A possibly all-around better protective scheme has evolved from an appropriate blend of classical transformer principles and high technology, specifically the use of high-temperature superconducting materials. What is to be described has been technically feasible with helium-temperature superconductors ($<23\text{K}$), but the cryogenic cooling system imposed an unacceptable financial burden. With the newer superconducting materials, the transition temperature being in the vicinity of 100K , it becomes practical to use relatively inexpensive liquid nitrogen as the refrigerant.

The basic current limiter is a transformer configuration with a short-circuited secondary as depicted in Fig. 6.17. As may be suspected, there is more than meets the eye, the very simplicity of the device being somewhat deceptive. For one thing, the shorted secondary, because of its super conductivity, has zero ohmic resistance during normal operation of the power system. Accordingly, no resistance is reflected into the series-connected primary winding during normal loading of the power system. This is fine and dandy, for it eliminates what could be a major source of power dissipation. It should be kept in mind, too, that the design of this transformer must be such that the normal inductive reactance of the primary is minimal. This is brought about with as few primary turns as is consistent with other objectives, and by the use of an *open* rather than a closed core. All in all, there is negligible upset in the normal operation of the power line from the insertion of this device.

We come now to a major departure in operating principle from conventional transformers: when this transformer experiences operation under normal power-line loads, its core is not exposed to magnetic fields. This is because of shielding action provided by the superconducting secondary

which is positioned between the primary winding and the core. One's natural reaction to this statement is that there is nothing unusual about such transformer construction. It is, however, the nature of superconductivity to *repel* penetration of magnetic fields. Indeed, this device is known as a shielded-core superconductor fault-current limiter. Under normal loading, the power line 'sees' the approximate equivalent of an air-core inductor with very low inductance. There is therefore miniscule voltage drop and negligible operational disturbance to the power line.

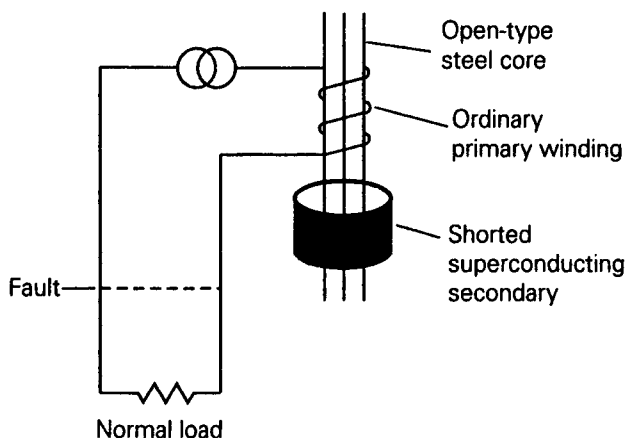


Figure 6.17 *Current limiting by transformer-coupled superconductor behaviour. The operation of this device is explained in the text. The salient features of its constructional format are as follows. (a) The superconducting secondary is closest to the core. Only the secondary is refrigerated. (b) The primary is the outermost winding; its magnetic field can get to the core only by going through the secondary material. (This it cannot do as long as the secondary material remains in its superconducting state.) In normal operation the core is shielded from magnetic fields. In fault-current operation, the core participates in ordinary transformer-action and the power-line 'sees' current-limiting impedance.*

Suppose, now, that a short circuit develops in the power line. Ordinary transformer action couples the higher-than-normal primary current into the superconducting secondary. The secondary current-density quickly exceeds the superconducting transition-level, whereupon ordinary ohmic resistance is manifested. This secondary resistance reflected into the primary winding appears as a current-limiting resistance in series with the power line and the short-circuiting fault. Nor is this all that happens; inasmuch as the magnetic shielding effect of superconductivity has vanished, the steel core now participates in the transformer action so that the power-line fault current is further limited by higher inductive reactance of the pri-

mary. Moreover, self-heating of the now-resistive secondary reflects even more resistance into the primary.

It should be understood that this quick-acting scheme for limiting fault current does not dispense with mechanically actuated circuit breakers. Rather, it confers protection during the inherent delay the breakers require to respond. The commercialization of this technique stems from the discovery of the high-temperature superconducting materials because liquid nitrogen is a more practical and less costly refrigerant than liquid helium. Another favourable cost-factor is the need for only a simple open-type steel core. Best of all, this sophisticated transformer device neatly overcomes the drawbacks of mechanized motion.

This Page Intentionally Left Blank

Appendix – useful information

Technical data is generally presented in either tabular or graphical form. With regard to the latter, it is often more convenient to depict such information on semi-log or on log-log paper rather than on conventional cartesian graph paper. There are several advantages for doing this. A much greater range of variation can be accommodated before visual difficulties set in. Interpretation of coordinate points can be quickly grasped – less intellectual effort is needed. The plotting of such ‘curves’ can often be accomplished with far fewer coordinate points. To the knowledgeable reader, there is important ‘hidden’ information. For example, the slope of the straight line log-log plot reveals the power of the exponent in the relationship. If eddy current loss is plotted as a function of frequency, a slope of 2 indicates that eddy current dissipation increases as the square of frequency. Similarly, a slope of 1.6 for the log-log plot of hysteresis loss *versus* flux density yields the Steinmetz relationship which tells us that hysteresis loss goes up as the 1.6 power of maximum flux density (ideally).

As a corollary of the above, a *departure* from a straight line relationship in a log-log plot tells us that a change in the relationship has occurred. In transformer technology, it is likely to be core saturation. This would not be as detectable in a plot on cartesian graph paper for what we would see is a mere change in curvature of the already curved plot.

Although the differently-named magnetic units and the different measuring systems tend to be initially confusing, one is not likely to go wrong if strict consistency is always observed. For example, lines per square *inch*, area in square *inches* and ampere-turns per *inch* always go together. Similarly, metric measure expressed in centimetres, square centimetres etc. can be used. In order to avoid the forbidden *mixture* of units, the table of *conversions* will serve to ‘clear the deck’.

Table A.1 Systems used in magnetics

<i>Unit</i>	<i>CGS system</i>	<i>MKS system</i>	<i>English system</i>
ϕ Flux	Maxwell or line	Weber	Maxwell or line
B Flux density	Gauss or lines/cm ²	Webers/m ² (Tesla)	Lines/in ²
F MMF	Gilbert	Ampere turns	Ampere turns(NI)
H Magnetizing force	Oersted or Gilbert/cm	NI/m	NI/in
A Area	cm ²	m ²	in ²
I Length	cm	m	in
V Volume	cm ³	m ³	in ³
μ Permeability	1	$0.4\pi \times 10^6$	3.192

In the technical literature one sometimes finds combinations used. It isn't as bewildering as one might suspect; the different terms for the same entity differ from one another only by a conversion factor, not in concept. Many practitioners favour the basic CGS system because *permeability* then has the convenient reference value of unity.

Table A.2 Magnetic units conversion factors and wire area

<i>To convert</i>	<i>Into</i>	<i>Multiply by</i>	<i>Conversely, multiply by</i>
Ampere turns	Gilberts	0.4π	0.796
Ampere turns/in	Oersteds	0.495	2.02
Ampere turns/m	Oersteds	0.004π	79.6
Lines/in ²	Gauss	0.155	6.45
Webers	Lines	10^8	10^{-8}
Webers/m ² (Tesla)	Gauss	10^4	10^{-4}
in ²	Circular mils	1.273×10^6	0.785×10^{-6}
cm ²	Circular mils	0.197×10^6	5.063×10^6

Although initially confusing, calculations are easily handled by maintaining strict consistency – stick to *either* English or metric units and pay heed to the *magnitude* of the parameters selected. For example, work with *either* metres or centimetres.

Table A.3 Characteristics of magnetic materials

<i>Trade names</i>				<i>Trade marks</i>			
A	Hypersil, Microsil, Silectron, Orthosil, Magnesil			E	Supermendur, Permendur, Hyperco		
B	Deltamax, Orthonol, Orthonik, Hypernik V, Square mu-49			F	Molybdenum Permalloy Powder, MPP		
C	Permalloy 45, 48 Alloy, Carpenter 49, 4750, Hypernik			G	Carbonyl, Iron Oxides, Crolite, Polyrion		
D	Supermalloy, Square mu-79, Hymu 80, Square Permalloy 80, Mumetal, Hy Ra 80			H	Ferrite, Ferramic, Ceramag, Sifertrit, Ferroxcube		
<i>Material</i>	<i>% Alloy</i>	<i>Characteristic property</i>	<i>Typical application</i>	<i>Initial permeability</i>	<i>B_m, kG</i>	<i>B, kG</i>	<i>Typical frequency</i>
Silicon – iron ANSI	4% Si	Low cost	Transformers and reactors	400 typ.	19	12	To 100 Hz
Silicon – iron A	3% Si	Low cost, grain oriented	Low-loss transformers and reactors	1500 typ. 100 – 1000 Hz	18	13.6	To 100 Hz 100 – 1000 Hz 2 – 5 kHz
Nickel – iron B	45 – 50% Ni	Rectangular hysteresis loop	Saturating transformers and reactors	1000 typ.	15 – 16	12 – 15	100 – 1000 Hz 1 – 8 kHz 10 – 50 kHz
Nickel – iron C	45 – 50% Ni	Semioriented, soft	Combined permeability and flux density characteristics	7500 typ.	12 – 15	6.2 – 8	
Nickel – iron D	70 – 80% Ni, 3 – 5% Mo	High permeability Low core loss high stability	High-frequency reactors	20 000 – 120 000	7 – 8	5 – 6	To 500 Hz 2 – 20 kHz 40 – 100 kHz
Cobalt – iron E	3 – 50% Co, (2% V)	High flux density, high cost	Small size and weight Saturating transformers	800 typ.	22	21	1 – 5 kHz
Powder F permalloy		High stability, selectable permeability	High Q filters	14 – 550	3 – 6		10 kHz – 1 MHz
Powdered G iron		Stable Q and permeability	RF transformers, wave filters	10 – 100	9 – 10		10 kHz – microwave
Ferrites H	Mn Zi, Ni Zi	Wide variety of shapes and sizes, lower cost than nickel – iron, high resistivity	High-frequency pulse and power transformers and reactors	250 – 5000	2 – 4	1 – 2	10 kHz – microwave

The magnetic materials used in transformer cores are alloys, powdered or ferrites. At radio frequencies, the air core is also extensively used. The peculiarity of air as a magnetic 'material' is that it exhibits no known saturation level.

Table A.4 American wire gauge, standard annealed copper at 20 °C

<i>Gauge</i>	<i>Diameter in mm</i>	<i>Cross section</i>		<i>Ohms per 1000 ft</i>	<i>Feet per ohm</i>	<i>Pounds per 1000 ft</i>
		<i>Circular mils</i>	<i>Square inch</i>			
0000	460.0	211 600	0.1662	0.049 01	20 400	640.5
000	400.6	167 800	0.1318	0.061 82	16 180	507.8
00	364.8	133 100	0.1045	0.077 93	12 830	402.8
0	324.9	105 600	0.08291	0.098 25	10 180	319.5
1	289.3	83 690	0.065 73	0.1239	8070	253.3
2	257.6	66 360	0.052 12	0.1563	6398	200.9
3	229.4	52 620	0.041 33	0.1971	5074	159.3
4	204.3	41 740	0.032 78	0.2485	4024	126.3
5	181.9	33 090	0.025 99	0.3134	3190	100.2
6	162.0	26 240	0.020 61	0.3952	2530	79.44
7	144.3	20 820	0.016 35	0.4981	2008	63.03
8	128.5	16 510	0.012 97	0.6281	1592	49.98
9	114.4	13 090	0.010 28	0.7925	1262	39.62
10	101.9	10 380	0.008 155	0.9988	1001	31.43
11	90.7	8230	0.006 46	1.26	793	24.9
12	80.8	6530	0.005 13	1.59	629	19.8
13	72.0	5180	0.004 07	2.00	500	15.7
14	64.1	4110	0.003 23	2.52	396	12.4
15	57.1	3260	0.002 56	3.18	314	9.87
16	50.8	2580	0.002 03	4.02	249	7.81
17	45.3	2050	0.001 61	5.05	198	6.21
18	40.3	1620	0.001 28	6.39	157	4.92
19	35.9	1290	0.001 01	8.05	124	3.90
20	32.0	1020	0.000 804	10.1	98.7	3.10
21	28.5	812	0.000 638	12.8	78.3	2.46
22	25.3	640	0.000 503	16.2	61.7	1.94
23	22.6	511	0.000 401	20.3	49.2	1.55
24	20.1	404	0.000 317	25.7	39.0	1.22
25	17.9	320	0.000 252	32.4	30.9	0.970
26	15.9	253	0.000 199	41.0	24.4	0.765
27	14.2	202	0.000 158	51.4	19.4	0.610
28	12.6	159	0.000 125	65.3	15.3	0.481
29	11.3	128	0.000 100	81.2	12.3	0.387
30	10.0	100	0.000 078 5	104	9.64	0.303
31	8.9	79.2	0.000 062 2	131	7.64	0.240
32	8.0	64.0	0.000 050 3	162	6.17	0.194
33	7.1	50.4	0.000 039 6	206	4.86	0.153
34	6.3	39.7	0.000 031 2	261	3.83	0.120
35	5.6	31.4	0.000 024 6	331	3.02	0.0949
36	5.0	25.0	0.000 019 6	415	2.41	0.0757
37	4.5	20.2	0.000 051 9	512	1.95	0.0613
38	4.0	16.0	0.000 012 6	648	1.54	0.0484
39	3.5	12.2	0.000 009 62	847	1.18	0.0371
40	3.1	9.61	0.000 007 55	1080	0.927	0.0291
41	2.8	7.84	0.000 006 16	1320	0.756	0.0237
42	2.5	6.25	0.000 004 91	1660	0.603	0.0189
43	2.2	4.84	0.000 003 80	2140	0.467	0.0147
44	2.0	4.00	0.000 003 14	2590	0.386	0.0121
45	1.76	3.10	0.000 002 43	3350	0.299	0.00938

Table A.4 – continued

<i>Gauge</i>	<i>Diameter in mm</i>	<i>Cross section</i>		<i>Ohms per 1000 ft</i>	<i>Feet per ohm</i>	<i>Pounds per 1000 ft</i>
		<i>Circular mils</i>	<i>Square inch</i>			
46	1.57	2.46	0.000 001 94	4210	0.238	0.00746
47	1.40	1.96	0.000 001 54	5290	0.189	0.00593
48	1.24	1.54	0.000 001 21	6750	0.148	0.00465
49	1.11	1.23	0.000 000 968	8420	0.119	0.00373
50	0.99	0.980	0.000 000 770	10 600	0.0945	0.00297
51	0.88	0.774	0.000 000 608	13 400	0.0747	0.00234
52	0.78	0.608	0.000 000 478	17 000	0.0587	0.00184
53	0.70	0.490	0.000 000 385	21 200	0.0472	0.00148
54	0.62	0.384	0.000 000 302	27 000	0.0371	0.00116
55	0.55	0.302	0.000 000 238	34 300	0.0292	0.000916
56	0.49	0.240	0.000 000 189	43 200	0.0232	0.000727

Listed resistance values are for direct current.

The resistance of hard-drawn copper is about 2.5% higher than listed values. Copper temperature coefficient is approximately $0.00393 \Omega/^{\circ}\text{C}$.

Table A.5 Partial list of Litz wire formats (MWS Wire Industries)

Number of strands	Size (AWG)	Circular mils nominal	Nearest AWG equiv. (cir. mils)	Resistance (ohms per 1000 ft. at 20°C) nominal	Single polyurethane		
					Mean o. d. (in.) nominal	ft per pound	pounds per 1000 ft
3	30	300.00	25½	34.57	0.022	1068	0.936
4	30	400.00	24	25.93	0.025	801	1.25
5	30	500.00	23	20.74	0.028	641	1.56
6	30	600.00	22½	17.29	0.031	534	1.87
7	30	700.00	21½	14.82	0.033	458	2.18
8	30	800.00	21	12.96	0.036	401	2.49
9	30	900.00	20½	11.52	0.038	356	2.81
10	30	1000.00	20	10.37	0.040	321	3.12
15	30	1500.00	18½	6.91	0.049	214	4.67
20	30	2000.00	17	5.19	0.056	160	6.25
3	32	192.00	27	54.00	0.018	1695	0.590
4	32	256.00	26	40.50	0.020	1272	0.786
5	32	320.00	25	32.40	0.023	1017	0.983
6	32	384.00	24½	27.00	0.025	848	1.18
7	32	448.00	23½	23.14	0.027	727	1.38
8	32	512.00	23	20.25	0.029	636	1.57
9	32	576.00	22½	18.00	0.031	565	1.77
10	32	640.00	22	16.20	0.032	509	1.97
15	32	960.00	20½	10.80	0.039	339	2.95
20	32	1280.00	19	8.10	0.046	254	3.94
3	34	119.07	29½	87.10	0.014	2680	0.373
4	34	158.76	28	65.33	0.016	2010	0.498
5	34	198.45	27	52.26	0.018	1608	0.622
6	34	238.14	26½	43.55	0.020	1340	0.746
7	34	277.83	25½	37.33	0.021	1148	0.871
8	34	317.52	25	32.66	0.023	1005	0.995
9	34	357.21	24½	29.03	0.024	893	1.12
10	34	396.90	24	26.13	0.026	804	1.24
15	34	595.35	22½	17.42	0.031	536	1.87
20	34	793.80	21	13.07	0.036	402	2.49
3	36	75.00	31	138.27	0.011	4230	0.236
4	36	100.00	30	103.70	0.013	3173	0.315
5	36	125.00	29	82.96	0.015	2538	0.394
6	36	150.00	28	69.13	0.016	2115	0.473
7	36	175.00	27½	59.26	0.017	1813	0.552
8	36	200.00	27	51.85	0.018	1586	0.631
9	36	225.00	26½	46.09	0.019	1410	0.709
10	36	250.00	26	41.48	0.021	1269	0.788
15	36	375.00	24½	27.65	0.025	846	1.18
20	36	500.00	23	20.74	0.029	635	1.58
25	36	625.00	22	16.59	0.032	508	1.97

Table A.5 – *continued*

<i>Number of strands</i>	<i>Size (AWG)</i>	<i>Circular mils nominal</i>	<i>Nearest AWG equiv. (cir. mils)</i>	<i>Resistance (ohms per 1000 ft. at 20°C) nominal</i>	<i>Single polyurethane</i>		
					<i>Mean o. d. (in.) nominal</i>	<i>ft per pound</i>	<i>pounds per 1000 ft</i>
30	36	750.00	21½	13.83	0.035	423	2.36
40	36	1000.00	20	10.37	0.041	317	3.16
50	36	1250.00	19	8.30	0.046	254	3.94
60	36	1500.00	18½	6.91	0.050	212	4.72

Transformers used in switching power supplies, inverters and converters can operate from 20 to 200 kHz and higher. The use of Litz wire in high frequency windings can appreciably reduce copper losses from skin-effect.

Table A.6 Quick design data for 60 Hz power transformers

<i>Watts</i>	<i>Section of core (in.)</i>	<i>Area of core (square in.)</i>	<i>Primary turns</i>	<i>Primary wire size</i>	<i>Turns per volt</i>
10	$\frac{1}{2} \times \frac{1}{2}$	0.25	3500	31	32
10	$\frac{1}{2} \times \frac{5}{8}$	0.31	2800	31	24.2
12	$\frac{1}{2} \times \frac{3}{4}$	0.37	2300	30	20.0
12	$\frac{5}{8} \times \frac{5}{8}$	0.38	2280	30	19.6
15	$\frac{5}{8} \times \frac{3}{4}$	0.46	1875	29	16.1
22	$\frac{5}{8} \times 1$	0.62	1400	28	12.2
20	$\frac{3}{4} \times \frac{3}{4}$	0.55	1570	28	13.6
25	$\frac{3}{4} \times 1$	0.75	1150	27	10.0
30	$\frac{3}{4} \times 1\frac{1}{4}$	0.93	930	26	8.1
50	$\frac{3}{4} \times 1\frac{1}{2}$	1.12	770	24	6.7
50	1×1	1.0	860	24	7.5
60	$1 \times 1\frac{1}{4}$	1.25	690	23	6.0
65	$1 \times 1\frac{1}{2}$	1.50	575	23	5.0
75	$1 \times 1\frac{3}{4}$	1.75	490	22	4.2
110	1×2	2.0	430	21	3.7
105	$1\frac{1}{4} \times 1\frac{1}{4}$	1.56	550	21	4.8
100	$1\frac{1}{4} \times 1\frac{1}{2}$	1.87	460	21	3.8
120	$1\frac{1}{4} \times 1\frac{3}{4}$	2.18	400	20	3.5
140	$1\frac{1}{4} \times 2$	2.5	350	19	3.2
125	$1\frac{1}{2} \times 1\frac{1}{2}$	2.25	380	20	3.3
150	$1\frac{1}{2} \times 1\frac{3}{4}$	2.64	330	18	2.9
200	$1\frac{1}{2} \times 2$	3.0	290	17	2.42
300	2×2	4.0	215	15	1.87
400	$2 \times 2\frac{1}{2}$	5.0	175	14	1.52
500	2×3	6.0	145	13	1.26

Use 29-gauge E-I silicon steel laminations.

For 50 Hz operation, select a core area 120% greater than indicated for the same wattage 60 Hz transformer. If this is not feasible, get as close as possible to the 120% factor; then change the number of primary turns so that the *product* of turns and area are again the same as for 60 Hz operation. Record the new number of primary turns so that the new turns per volt becomes calculable.

Table A.7 Voltage or current ratio versus power ratio and decibels

<i>Voltage or current ratio</i>	<i>Power ratio</i>	<i>−dB +</i>	<i>Voltage or current ratio</i>	<i>Power ratio</i>
1.000	1.000	0.0	1.0000	1.0000
0.989	0.977	0.1	1.0116	1.0233
0.977	0.955	0.2	1.0233	1.0471
0.966	0.933	0.3	1.0351	1.0715
0.955	0.912	0.4	1.0471	1.0965
0.944	0.891	0.5	1.0593	1.1220
0.933	0.871	0.6	1.0715	1.1482
0.923	0.851	0.7	1.0839	1.1749
0.912	0.832	0.8	1.0965	1.2023
0.902	0.813	0.9	1.1092	1.2303
0.891	0.794	1.0	1.1220	1.2589
0.881	0.776	1.1	1.135	1.288
0.871	0.759	1.2	1.1482	1.3183
0.861	0.741	1.3	1.161	1.349
0.851	0.724	1.4	1.175	1.380
0.841	0.708	1.5	1.189	1.413
0.832	0.692	1.6	1.202	1.445
0.822	0.676	1.7	1.216	1.479
0.813	0.661	1.8	1.230	1.514
0.803	0.646	1.9	1.245	1.549
0.749	0.631	2.0	1.2589	1.5849
0.776	0.603	2.2	1.288	1.660
0.759	0.575	2.4	1.318	1.738
0.750	0.562	2.5	1.334	1.778
0.724	0.525	2.8	1.380	1.905
0.708	0.501	3.0	1.4125	1.9953
0.692	0.479	3.2	1.445	2.089
0.676	0.457	3.4	1.479	2.188
0.668	0.447	3.4	1.4962	2.2387
0.661	0.436	3.6	1.514	2.291
0.646	0.417	3.8	1.549	2.399
0.631	0.398	4.0	1.5849	2.5119
0.596	0.355	4.5	1.6788	2.8184
0.562	0.316	5.0	1.7783	3.1623
0.531	0.282	5.5	1.8836	3.5481
0.501	0.251	6.0	1.9953	3.9811
0.473	0.224	6.5	2.113	4.467

Table A.7 – *continued*

<i>Voltage or current ratio</i>	<i>Power ratio</i>	<i>–dB +</i>	<i>Voltage or current ratio</i>	<i>Power ratio</i>
0.447	0.200	7.0	2.239	5.012
0.422	0.178	7.5	2.371	5.623
0.398	0.159	8.0	2.512	6.310
0.376	0.141	8.5	2.661	7.079
0.355	0.126	9.0	2.818	7.943
0.335	0.112	9.5	2.985	8.913
0.316	0.100	10	3.162	10.00
0.282	0.0794	11	3.55	12.6
0.251	0.0631	12	3.98	15.9
0.224	0.0501	13	4.47	20.0
0.200	0.0398	14	5.01	25.1
0.178	0.0316	15	5.62	31.6
0.159	0.0251	16	6.31	39.8
0.141	0.0200	17	7.08	50.1
0.126	0.0159	18	7.94	63.1
0.112	0.0126	19	8.91	79.4
0.10000	0.0100	20	10.00	100.0
0.08913	0.0079	21	11.22	125.9
0.07943	0.0063	22	12.59	158.5
0.70709	0.0050	23	14.13	199.5
0.06310	0.00398	24	15.85	251.2
0.05623	0.03162	25	17.78	316.2
0.05012	0.002512	26	19.95	398.1
0.04467	0.001995	27	22.39	501.2
0.03981	0.001585	28	25.12	631.0
0.03548	0.001259	29	28.18	794.3
0.03162	0.001000	30	31.62	1000
0.02818	0.000794	31	34.48	1259
0.02512	0.000631	32	39.81	1585
0.02239	0.000501	33	44.67	1995
0.01995	0.000398	34	50.12	2512
0.01778	0.000316	35	56.23	3162
0.01585	0.000251	36	63.10	3981
0.01413	0.000199	37	70.79	5012
0.01259	0.000158	38	79.43	6310
0.01122	0.000126	39	89.13	7943
0.01000	0.000100	40	100.00	10000
0.00891	0.000079	41	112.2	12590

Table A.7 – continued

<i>Voltage or current ratio</i>	<i>Power ratio</i>	<i>–dB +</i>	<i>Voltage or current ratio</i>	<i>Power ratio</i>
0.00794	0.000063	42	125.9	15850
0.00708	0.000050	43	141.3	19950
0.00631	0.000040	44	158.5	25120
0.00562	0.000032	45	177.8	31620
0.00501	0.000025	46	199.5	39810
0.00447	0.000020	47	223.9	50120
0.00398	0.000016	48	251.2	63100
0.00355	0.000013	49	281.8	79430
0.00316	0.000010	50	316.2	100000

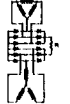
Decibels are abbreviated dB. Decibels express the power ratio between two power levels, such as P_1 and P_2 . Expressing the power ratio P_2/P_1 in this way, $\text{db} = 10\log_{10}(P_2/P_1)$. Also, if dB is known, $P_2/P_1 = \text{antilog}(\text{dB}/10)$. Moreover, $\text{dB} = 20\log_{10}(V_2/V_1)$ or $20\log_{10}(I_2/I_1)$ if the two voltages or the two currents have been measured in equal impedance circuits. For example, the input and output impedances of an amplifier would have to be the same.

Table A.8 Characteristics of single-phase rectifier circuits with resistive loads

Rectifier circuit connection	Half-wave	Full-wave centre-tap	Full-wave bridge
			
Load voltage and current waveshape			
Characteristic			
Diode average current			
$I_{F(AV)}/I_{L(DC)}$	1.00	0.50	0.50
Diode peak current			
$I_{FM}/I_{F(AV)}$	3.14	3.14	3.14
Form factor of diode			
$I_{F(RMS)}/I_{DC}$	1.57	1.57	1.57
Diode RMS current			
$I_{F(RMS)}/I_{L(DC)}$	1.57	0.785	0.785
RMS input voltage per transformer leg			
$V_i/V_{L(DC)}$	2.22	1.11	1.11
Peak inverse voltage			
$V_{RRM}/V_{L(DC)}$	3.14	3.14	1.57
Transformer primary rating $VA/P_{DC} \times$	3.49	1.23	1.23
Transformer secondary $\times$ rating VA/P_{DC}	3.49	1.75	1.23
Total RMS ripple (%)	121	48.2	48.2
Lowest ripple frequency (f_r/f_i)	1	2	2
Rectification ratio (conversion efficiency) (%)	40.6	81.2	81.2

The ratios designated by $\times$ are the *reciprocals* of the utilization factor.

Table A.9 Characteristics of three-phase rectifier circuits with resistive loads

Rectifier circuit connection	Half-wave star	Bridge	Double wye with interphase transformer	Full-wave star	Wye–delta connections	
						
Load voltage and current wave shape						
Characteristic						
Average current through diode $I_{F(AV)}/I_{L(DC)}$	0.333	0.333	0.167	0.167	0.167	0.333
Peak current through diode $I_{FM}/I_{F(AV)}$	3.63	3.14	3.15	6.30	6.30	6.30
Form factor of current through diode $I_{F(RMS)}/I_{F(AV)}$	1.76	1.74	1.76	2.46	2.46	2.46
RMS current through diode $I_{F(RMS)}/I_{L(DC)}$	0.587	0.579	0.293	0.409	0.409	0.818
RMS input voltage per transformer leg $V_i/V_{L(DC)}$	0.855	0.428	0.855	0.741	0.715	0.37
Diode peak inverse voltage (PIV), $V_{RRM}/V_{L(DC)}$	2.09	1.05	2.42	2.09	1.05	1.05
Transformer primary rating VA/P_{DC} $\times$	1.23	1.05	1.06	1.28	1.01	1.01
Transformer secondary rating VA/P_{DC} $\times$	1.50	1.05	1.49	1.81	1.05	1.05
Total RMS ripple (%)	18.2	4.2	4.2	4.2	1.0	1.0
Lowest ripple frequency f_r/f_i	3	6	6	6	12	12
Rectification ratio (conversion efficiency) (%)	96.8	99.8	99.8	99.8	100	100

The ratios designated by $\times$ are *reciprocals* of the utilization factor.

Table A.10 Single-phase rectifier circuits with critical inductance loads

Rectifier circuit connection	Single-phase full-wave centre-tap	Single-phase full-wave bridge	Three-phase half-wave star	Three-phase full-wave bridge
				
Load voltage waveshape				
Characteristic				
Average current through diode $I_{F(AV)}/I_{L(DC)}$	0.500	0.500	0.333	0.333
Peak current through diode $I_{FM}/I_{F(AV)}$	2.00	2.00	3.00	3.00
Form factor of current through diode $I_{F(RMS)}/I_{F(AV)}$	1.41	1.41	1.73	1.73
RMS input voltage per transformer leg $V_i/V_{L(DC)}$	1.11*	1.11	0.855	0.428
Diode peak inverse voltage (PIV) $V_{RRM}/V_{L(DC)}$	3.14	1.57	2.09	1.05
Transformer primary rating VA/P_{DC} X	1.11	1.11	1.21	1.05
Transformer secondary rating VA/P_{DC} X	1.57	1.11	1.48	1.05
Ripple $(V_r/V_{L(DC)})$				
lowest frequency in rectifier output (f_r/f_i)	2	2	3	6
Peak value of ripple components:				
ripple frequency (fundamental)	0.667	0.667	0.250	0.057
second harmonic	0.133	0.133	0.057	0.014
third harmonic	0.057	0.057	0.025	0.006
Ripple peaks with reference to d.c. axis:				
positive peak	0.363	0.363	0.209	0.0472
negative peak	0.637	0.637	0.395	0.0930

Usually, the immediate load will be a filter-choke. The choke inductance can exceed the critical value, the practical limit being imposed by its winding resistance which degrades voltage regulation. The ratios designated by x are reciprocals of the utilization factor. (Critical inductance for a half-wave circuit would be infinite.)

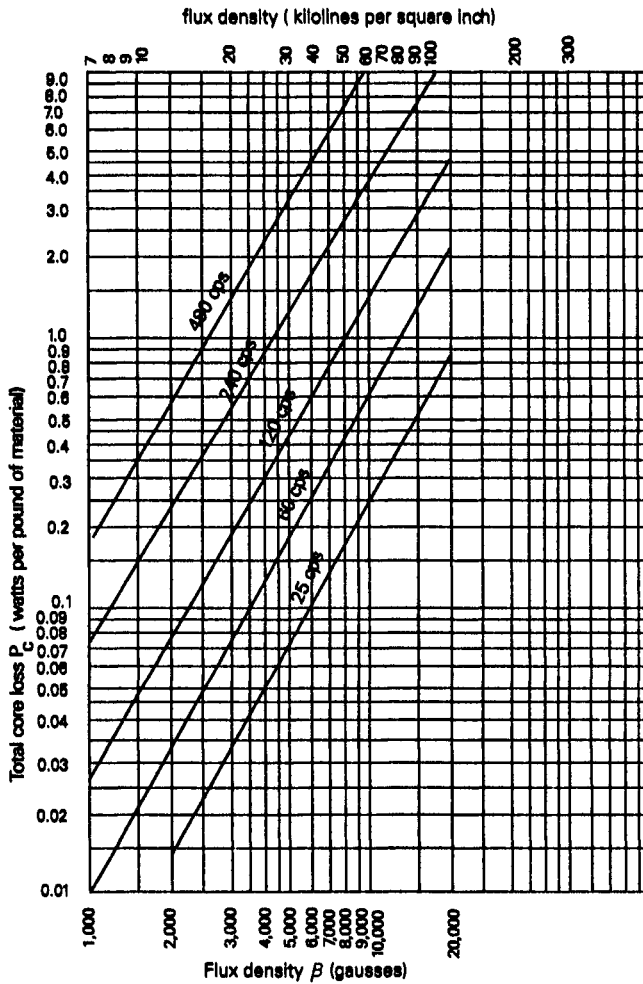


Figure A.1 Total core loss for 29 gauge, 4.2% silicon steel laminations. These curves summarize the hysteresis and eddy-current losses. (Courtesy of Allegheny Steel Co.)

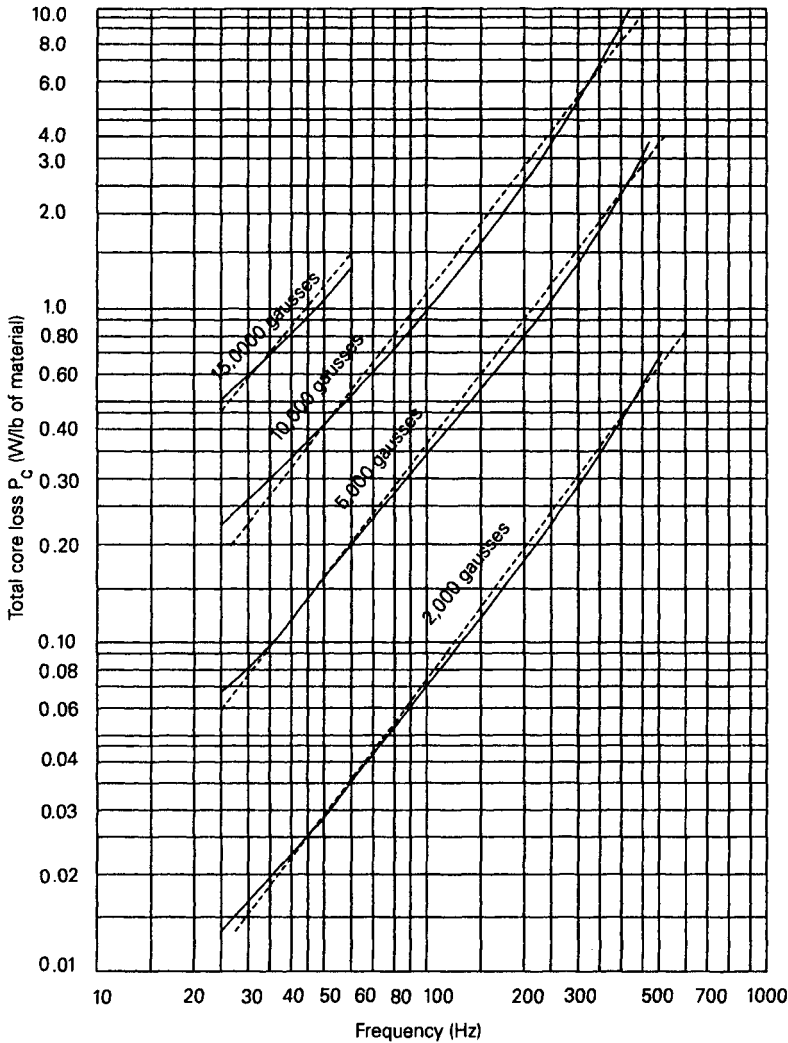


Figure A.2 Alternate core loss curves for 29 gauge, 4.2% silicon steel laminations. The solid lines derive from laboratory measurement data; the dashed lines are idealized. The differences are negligible for most practical purposes. (Courtesy of Allegheny Steel Co.)

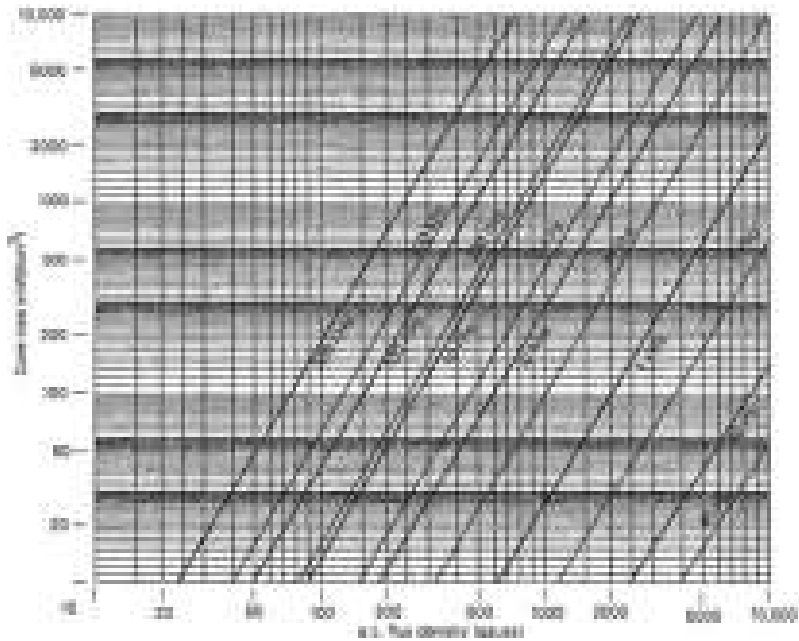


Figure A.3 Core loss characteristics of micrometals No. 26 iron powder cores. Iron powder material has a high temperature Curie point and has long been used at high frequencies because of its low hysteresis and eddy-current losses. Mix 26, however, merits consideration for power line frequencies as well. Powdered cores have what may be considered a distributed air gap. Very good linearity and permeability constancy is attained.

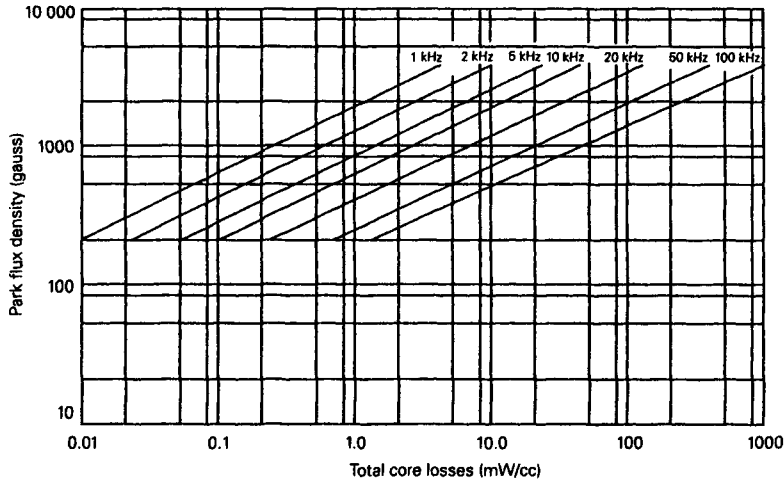


Figure A.4 Total core losses in Ferroxcube 3C8 ferrite material. This is a popular transformer core material in regulated power supply technology. Although abscissa and ordinates in the above graph are inverted from the usual presentation, the 'curves' nonetheless remain straight lines and are easy to deal with.

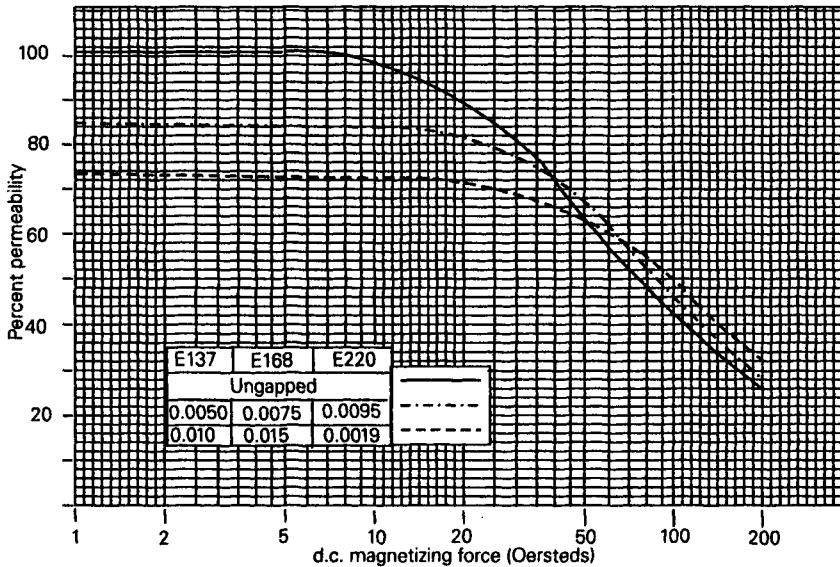


Figure A.5 The effect of d.c. in a core winding. The relatively gradual magnetic saturation in the No. 26 micrometals powdered-iron core is due to the 'distributed air gap' inherent in such core material.

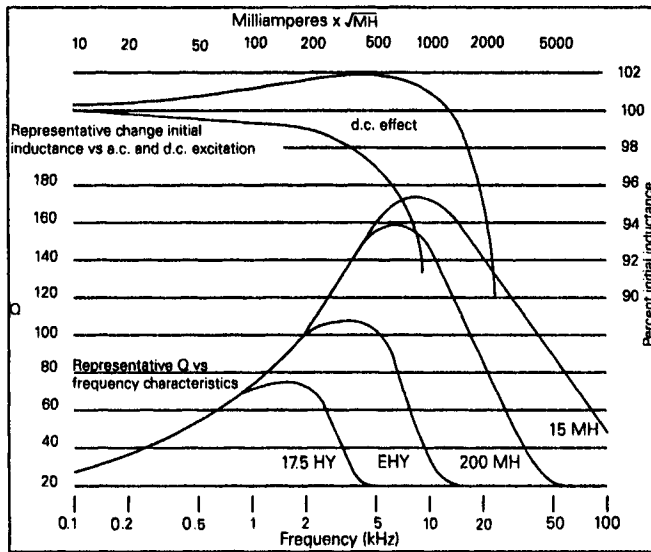


Figure A.6 'Q' curves for magnetics Inc. No. 930 moly-permalloy core material. These types of data are intended for inductors and resonant circuits. However, it is possible to gain useful insights for transformer applications also. For instance, high-Q regions denote low core loss, one of the important parameters in transformer operation.

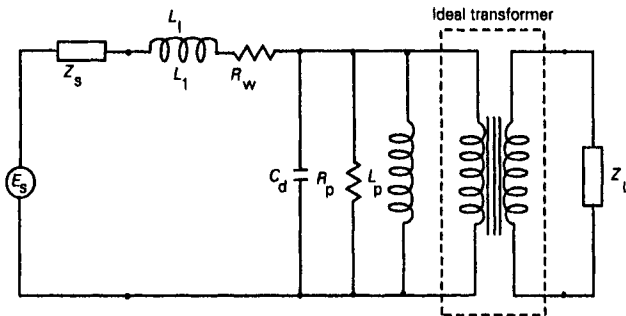


Figure A.7 Transformer equivalent circuit for assessing frequency response. This type of equivalent circuit is particularly useful for gaining insight into the behaviour of wide-band transformers. At high frequencies, it can be seen that the predominant factors are the equivalent winding resistance, the leakage inductance L_1 , core loss R_p , distributed capacitance and the output impedance of the source. At low frequencies, the primary magnetising inductance, being a shunt-arm, progressively increases attenuation as its reactance decreases. In this circuit, R_w represents the combined resistive-loss in both primary and secondary. R_p , however, simulates core loss.

This Page Intentionally Left Blank

Index

- a.c. circuit control, with saturable reactor, 41–5
- a.c. power control, 41–5
- American wire gauge table, 152
- Auto-transformers:
 - adjustable voltage, 55–6
 - hazard factor, 54
 - limitations, 55
 - for motor starting, 79–81
 - operating efficiency, 54
 - power capability, 52–3
- Automobile ignition systems, 105–107, 108
- Bandpass response, 74–5
- Battery charging, electric vehicles, 99–102
- Beverage wave antenna, 90–91, 92
- Biasing, 46
 - starting resistance, 29
- Bifilar windings, 111
- Bipolar transistor drive transformer, 61
- Bridge rectifier circuit, 68–9
- Capacitance magnification, by transformer, 78–81
- Capacitive reactance, 81
- Capacitor:
 - characteristics, 79
 - failure modes, 105–106
- Charging batteries, inductively-coupled, 100–102
- Coefficient of coupling, 76, 143
- Common-mode choke, 140–42
- Communications receivers, using IF filters, 47
- Computer memory systems *see* Magnetic core memory systems
- Condenser *see* Capacitor
- Constant current transformer, 38–9
- Constant voltage transformer, 40–41
- Conventional transformers, design equation, 94
- Copper losses, in windings, 20, 66
- Core construction:
 - alternative materials table, 29
 - double, 47–8
 - flux density, 94
 - types, 13–15
 - use of iron, 4–5, 10
- Core saturation, 40–41, 147
 - from geomagnetic storms, 119–22
 - half-cycle effects, 120–22
- Core-type transformers, 9–11
 - three-phase, 11
- Corona phenomenon, 117, 118
- Critical inductance, 68–72
- Cup cores *see* Pot cores
- Current inrush *see* Inrush current disturbance
- Current regulation, 131–2
- Current transformers, 115
 - for clamp-on meters, 116
 - dynamic, 131–2
 - for electronic usage, 116
 - precautions, 116

- in regulated power supplies, 116, 132
 - specialized, 116
- d.c. regulated power supply, 81–2
- d.c. to a.c. inverters, 25–6
- d.c. to d.c. converters, 57–8
- Decibels table, 157–9
- Delta connections, 73
 - in tertiary windings, 17
- Delta–delta connections, 85
- Delta–wye connections, 73, 85–6
- Design 4.44 factor, 122–4
- Double core, 47–8
 - for saturable core oscillator, 48
- Doubly-tuned intermediate frequency transformers, 74–5
- Dynamic current transformers, 131–2
- Eddy current induction, 22
 - losses, 26
- Electric vehicles, battery charging, 99–102
- Electric-wave filters, electromagnetic induction, 129
- Electricity/magnetism relationship, 2–3
- Electromagnetic coupling, 45
- Electromagnetic induction:
 - laws formulated, 3
 - paradox, 128–30
- Electromagnetic radiation, 22
- Electronic circuits:
 - electromagnetic induction, 129
 - transformer applications, 85–7
- Electrostatic shield, 16–17
- Energy transfer, travelling waves, 94–5
- Excitation, in transmission line, 89–90
- Excitation current, 7
- Failures:
 - of driven inverters, 138–9
 - hysteresis ‘walk-up’, 138, 139
- Faraday, Michael, 2
- Faraday shield, 17
- Ferro-resonant transformers, 40
- Flux density, 5
 - saturation table, 29
- Flux gate magnetometer, 139–40
- Flyback converter, 57, 58
- Flyback transformer, 103, 105
- Ford Model T spark coil, 104–105
- Forward converter, 57–8
- Fractional turns, 126–8
- Free-wheeling diodes, for current continuity, 65–6
- Frequencies, 50/60 Hz operational differences, 50–51
- Full-wave rectifier circuits:
 - bridge, 67–8
 - centre-tap, 67–9
 - critical inductance, 69–72
 - design formula, 71
 - presence of harmonics, 70
- G-line, 92–3
- Geomagnetic storms, 119
- Gramme ring dynamos, 129–30
- Grounding:
 - of cores, 19
 - of link-coupled system, 125–6
- GTO (switching device), 101
- H-bridge configuration, 82
- Half-cycle saturation effects, from geomagnetic storms, 120–22
- Hard saturation, 37
- Harmonics, critical inductance, 70
- Hartley circuit, 45
- Heat removal techniques *see* Ventilation techniques
- Helix winding, 113–14
- Henry, Joseph, 3
- High voltage transformers:
 - air exclusion, 118
 - alternative techniques, 112–14
 - basics, 103–105
 - design and construction problems, 116–17
 - pulse formation, 108–11
- High-leakage transformers, 39
- High-Q bandpass filter *see* Parametric converters
- High-voltage transformers, basics, 103–104
- Hybrid coil, 84–5
- Hysteresis losses, 26, 138–9
- Ideal transformer concept, 3–4, 49
- IF amplifiers, 75–6
 - doubly-tuned, 75

- popular passbands, 75
- IF transformers, 47
- Ignition coils, 105
 - turns ratio, 105
- IGTB (switching device), 101
- Impedance matching, 8
- Impedance transformation, 79
- Inductance, 5, 6
- Induction motor, starting, 79
- Induction regulator, 130–31
- Inductive reactance, 6
- Inductively-coupled battery charging, 100–102
- Inrush current disturbance, 62–4
- Inverter circuits, 25–31
- Inverter transformers *see* Saturable core inverters
- Iron core, 4–5, 10
- Isolation devices, 77–8
- Isolation transformers, 73–4

- Kilovolt-ampere (kVA) rating, 66, 72
- Kilowattage ratings, 66, 72

- Laser apparatus, winding techniques, 111
- Leakage flux, 22
- Leakage inductance, 38, 82, 143–5
- Lenz's law, 18
- Line voltage regulator, 34
- Linear output transformer, walking process, 138
- Linearity assumption, 4
- Link coupling transformation system, 124–6
- Litz wire:
 - for skin effect reduction, 16, 21–2, 153
 - table of formats, 152–3
- Load voltage regulator, 34
- Log-log plots, 149
- Lorain subcyclor, 36–8
- Losslessness assumption, 4
- Low frequency response enhancement, 96

- Magnetic amplifier (magamp), 43, 44
 - application, 98–9
 - principles, 97–8
- Magnetic core memory systems, transformer applications, 87
- Magnetic materials, table of properties, 151
- Magnetic permeability, 5
- Magnetics:
 - conversion factors, 150
 - systems of units, 150
- Magnetometer, flux gate *see* Flux gate magnetometer
- Magnetostriction properties, 45–6
- Magnetron drive operation, 108–109
- Mechanical filters, 47
- Metal detectors, 23–4
- MOSFET switching element, 58–61
- Mutual conductance, 143
- Mutual flux, 6–8

- Neon transformers, 39
- Neutralization of tube capacitances, 133
- Neutrodyne radio receivers, 133
- Noise spikes, avoidance, 142
- Non-sinusoidal waveforms, 7–8, 123–4
- Nulling process, 133–5
 - demonstration setup, 135–8

- Oil immersion, high voltage transformers, 118
- Open core transformers, 15
- Open-delta connections (Varrangement), 87
- Operating principles, 6–8
- Oscillator, 33–4
 - magnetostrictive, 46
- Ozone factor *see* Corona phenomenon

- Parametric converters, 31–6, 41
 - see also* Lorain subcyclor
- Parametric modulation, 34
- Phase-modulation, for d.c. regulated supplies, 81–2
- Phasing transformer, 57–8
- Photo-flash apparatus, winding techniques, 111
- Polarity preference, automobile ignition system, 106
- Polyphase conversion, 83–4
- Pot core transformers, 14–15
- Pot cores, 15
- Power control, remote, 77–8
- Power lines, short-circuit protection, 147
- Power MOSFETS, 58–61

- Power ratio table, 157–9
- Power transformers, 73
 - design data table, 156
- Power-handling capabilities, 66
- Pressurized gas, high voltage transformers, 118
- Proximity effect, 22
- Pulse transformers, 108
 - three-winding, 108–11
 - two-winding, 108
- Pulse-width modulation:
 - switchmode power supplies, 51
 - transformer behaviour, 58–61
- Quarter wave transmission line, helix, 113–14
- Radio frequencies:
 - dangers from, 107
 - high voltage jumping tendency, 117–18
 - stage tuning, 23
 - system link coupling, 124–6
- Radio interference, automobile ignition system, 107
- Rectifier circuits, 66–9
 - 60 Hz full-wave, 70
 - single-phase characteristics, 67, 160, 162
 - three-phase characteristics, 71–2, 161
- Reference windings, 99
- Regenerative modulator, 36
- Remote control devices, use of transformers, 77–8
- Repair techniques, 18
- Repeaters, telephone, 84
- RF amplifier, use of common-mode choke, 140–42
- RF chokes, simulated, 96
- Ringing circuit, 36
- Safety hazards, avoiding, 73–4
- Saturable core inverters, 25–31, 47–8
 - design equation, 29–30
- Saturable reactors:
 - for a.c. circuit control, 41–5
 - with isolated control winding, 42, 43
 - parametric device structure, 41
- Saturation flux densities, table, 29
- Schottky diode, for current continuity, 66
- Scott connection (for poly-phase conversion), 83–4
- Secondary voltage, 126–8
- Sensing windings, 99
- Shell-type transformers, three-phase, 11–13
- Shields:
 - electrostatic, 16–17
 - radio frequency, 17
- Short-circuit protection of power lines, 145
- Shorted turns failure mode, 18–19
- Shunting effect, 96
- Sine wave voltage, 8
- Skin effect, in windings, 20–22
- Slug tuning, 23
- Snubbing circuits, 144
- Spikes, 26
 - suppression/reduction, 27, 30, 48, 142
- Square waveform effects, 123–4
- Starting capacitors, 79
- Steinmetz relationship, 149
- Step-down ratio, 8
- Step-up ratio, 8
- Street lighting systems, current-regulating transformers, 132
- Subcycle ringer, 36
- Switch mode power supplies, 60
- Switching transients, avoidance, 142
- Synchronism, 72–3
- Taps, for ratio selection, 16
- Teaser transformer, 84
- Telephony, transformer applications, 84–5
- Temperature control, output transformer, 64–5
- Temperature effects, on transformer operation, 19
- Tertiary winding addition, 36
- Tesla coils:
 - modern versions, 112
 - solid-state circuitry, 113
 - systems, 107–108
- Thermistors, inrush current protection, 63–4
- Third harmonic, in current shape, 7, 70
- Three-phase transformer network circuits, 72–3
- Three-winding pulse transformer, 109–10
- Toroidal cores, 13–14

- Toroidal transformers, 13–14
 - high frequency, 111–12
 - winding, 129–30
- Transformation ratio, 105
- Transformer action
 - basics, 6–8
 - discovery, 2–3
 - failure modes, 18–19
 - field change, 3
 - sudden excitation, 62–4
- Transformer applications, 22–4, 39, 43, 44
 - see also under specific uses*
- Transformer construction
 - core-type, 9–10
 - embellishments, 16–17
 - open core, 15
 - pot cores, 15
 - shell-type, 9–11
 - for special needs, 17
 - three-phase, 11–13
 - toroidal cores, 13–14
- Transformer design
 - 4.44 factor, 122–4
 - conventional equation, 94
 - involving physical motion, 132
- Transistors:
 - damage from transformer behaviour, 138–9
 - development, 26
 - power amplification, 27
- Transmission lines:
 - behaviour, 91–3
 - transformer design philosophy, 94–6
 - transformer operating principle, 93–4
 - transformers for, 88–90
- Travelling wave systems, 93
- Travelling wave tube, 90
- Two-winding transformers, 43
- Utilization factor, 66–8
- Variable duty cycle circuits, 51
- Variac (adjustable auto-transformer), 56
- Vario-couplers, 135
- Ventilation techniques, 16, 19
- Volt-second rule, 60, 138, 139
- Voltage or current ratio table, 157–9
- Voltage variation, by auto-transformer, 55–6
- Waveform spikes, leakage inductance, 26–7
- Waveforms:
 - in iron core transformer, 7–8
 - pulse-width modulated, 58–61
 - square, 123–4
 - third harmonic, 7, 70
- Welding, use of Tesla coil, 107
- Windings:
 - basic equation, 5
 - cross-section area selection, 19
 - patterns in high frequency toroidal transformers, 111–12
- Wire areas, conversion table, 150
- Wire gauge table, American, 152
- Wye format, 73
- Wye-delta connections, 73
- X-ray transformers, 39
 - winding techniques, 111

Printed in the United Kingdom
by Lightning Source UK Ltd.
118378UK00001B/228

